



**NORMAL DAĞILIM İÇİN MAKİNE
ÖĞRENMEİ ALGORİTMASINA DAYALI
GELİŞTİRİLEN UYUM İYİLİĞİ PROSEDÜRÜ**

Yüksek Lisans Tezi

Zahir HAJIZADA

Eskişehir 2024

**NORMAL DAĞILIM İÇİN MAKİNE ÖĞRENMEİ ALGORİTMASINA
DAYALI GELİŞTİRİLEN UYUM İYİLİĞİ PROSEDÜRÜ**

Zahir HAJIZADA

Yüksek Lisans Tezi

İstatistik Anabilim Dalı

İstatistik Bilim Dalı

Danışman: Prof. Dr İlhan USTA

Eskişehir

Eskişehir Teknik Üniversitesi

Lisansüstü Eğitim Enstitüsü

Ocak 2024

JÜRİ VE ENSTİTÜ ONAYI

Zahir HAJIZADA'nın NORMAL DAĞILIM İÇİN MAKİNE ÖĞRENMESİ ALGORİTMASINA DAYALI GELİŞTİRİLEN UYUM İYİLİĞİ PROSEDÜRÜ başlıklı çalışması 18/01/2024 tarihinde aşağıdaki jüri tarafından değerlendirilerek "Eskişehir Teknik Üniversitesi Lisansüstü Eğitim-Öğretim ve Sınav Yönetmeliği"nin ilgili maddeleri uyarınca, İstatistik Anabilim dalında Yüksek Lisans Tezi olarak kabul edilmiştir.

Jüri Üyeleri

Unvan Adı Soyadı

İmza

Üye

: Prof. Dr. İlhan USTA

Üye

: Prof. Dr. Yeliz MERT KANTAR

Üye

: Prof. Dr. Arzu ALTIN YAVUZ

Prof. Dr. Semra KURAMA

Lisansüstü Eğitim Enstitüsü Müdürü

18/01/2024

DANIřMAN ONAYI

Daniřmanlıęını yrttęm Yksek Lisans ęrencisi Zahir HAJIZADA, NORMAL DAęILIM İİN MAKİNE ęRENMEēİ ALGORİTMASINA DAYALI GELİēTİRİLEN UYUM İYİLİęİ PROSEDR baēlıklı tez alıēmasını tamamlamıētır. Hazırlamıē olduęu tez tarafımca incelenmiē ve ęrencinin tez savunma sınavına alınması bilimsel ve etik aıdan uygun grlmētır.

Tez Daniřmanı

Prof. Dr İlhan USTA

ÖZET

NORMAL DAĞILIM İÇİN MAKİNE ÖĞRENMESİ ALGORİTMASINA DAYALI GELİŞTİRİLEN UYUM İYİLİĞİ PROSEDÜRÜ

Zahir HAJIZADA

İstatistik Anabilim Dalı

İstatistik Bilim Dalı

Eskişehir Teknik Üniversitesi, Lisansüstü Eğitim Enstitüsü, Ocak 2024

Danışman: Prof. Dr İlhan USTA

İstatistiksel analizlerin doğruluğu ve güvenilirliği açısından normallik varsayımının sağlanıp sağlanmadığının kontrol edilmesi oldukça önemlidir. İstatistiksel analizlerde kullanılan bir veri setinin normal dağılıma sahip olup olmadığı temel olarak iki farklı yöntemle belirlenir. Bunlar, grafiksel yöntemler ve normal dağılım için uyum iyiliği testleridir. Ancak bu yöntemler tek başına kullanıldığı zaman tüm durumlar için normalliği belirlemek konusunda birbiriyle çelişen sonuçlar verebilmektedir. Bu çalışmada literatürde var olan yöntemlere alternatif olarak makine öğrenmesi algoritmasına dayalı bir veri setinin normal dağılıma uygun olup olmadığını belirleyen bir prosedür geliştirilmiştir. Makine öğrenmesi algoritmasına dayalı olarak geliştirilen bu prosedürün I. Tip hata oranı ve gücü literatürde kabul görmüş ve yaygın olarak kullanılan normal dağılım için uyum iyiliği testleri ile simetrik ve asimetrik dağılımlar altında kapsamlı bir simülasyon çalışmasıyla karşılaştırılmıştır. Ayrıca geliştirilen prosedürün performansı doğruluk, kesinlik, duyarlılık ve F1 skoru kriterlerine göre de incelenmiştir. Yapılan kapsamlı karşılaştırmalar sonucunda, literatüre ilk kez sunulan normal dağılım için makine öğrenmesi algoritmasına dayalı uyum iyiliği prosedürünün bu çalışmada ele alınan normallik testlerinden hemen hemen tüm durumlar için daha güçlü kararlar verdiği gözlemlenmiştir.

Anahtar Sözcükler: Normallik testleri, I. Tip hata, Testin gücü, Makine öğrenmesi algoritmaları.

ABSTRACT

GOODNESS-OF-FIT PROCEDURE DEVELOPED BASED ON MACHINE LEARNING ALGORITHM FOR NORMAL DISTRIBUTION

Zahir HAJIZADA

Department of Statistics

Programme in Statistics

Eskişehir Technical University, Institute of Graduate Programs, January 2024

Supervisor: Prof. Dr İlhan USTA

It is very important to check the normality assumption for the accuracy and reliability of statistical analyses. Whether a data set used in statistical analysis has a normal distribution is basically determined by two different methods. These are graphical methods and goodness-of-fit tests for normal distribution. However, when these methods are used alone, they may provide conflicting results in determining normality for all cases. In this study, as an alternative to the existing methods in the literature, a procedure based on a machine learning algorithm is developed to determine whether a data set conforms to the normal distribution. The Type I error rate and power of this procedure, which is based on a machine learning algorithm, is compared with the well-accepted and widely used goodness-of-fit tests for the normal distribution under symmetric and asymmetric distributions in a comprehensive simulation study. The performance of the developed procedure was also analyzed in terms of accuracy, precision, sensitivity and F1 score. As a result of extensive comparisons, it is observed that the goodness-of-fit procedure based on the machine learning algorithm for the normal distribution, which is presented for the first time in the literature, outperforms the normality tests considered in this study for almost all cases.

Keywords: Goodness of fit test for Normal distribution, Type I error, Power of test, Machine learning algorithms.

TEŞEKKÜR

Yüksek lisans eğitimim boyunca bana yol gösteren, desteğini hiçbir zaman esirgemeyen, maddi manevi her zaman yanımda olan çok değerleri danışmanım sayın Prof. Dr. İlhan USTA'ya, tez çalışmam boyunca yardımlarını esirgemeyen sayın Doç. Dr. Şükrü ACITAŞ'a ve Doç. Dr. Emre ÇİMEN'e teşekkür etmeyi kendime borç bilirim.

Zahir HAJIZADA



ETİK İLKE VE KURALLARA UYGUNLUK BEYANNAMESİ

Bu tezin bana ait, özgün bir çalışma olduğunu; çalışmamın hazırlık, veri toplama, analiz ve bilgilerin sunumu olmak üzere tüm aşamalarında bilimsel etik ve kurallara uygun davrandığımı; bu çalışma kapsamında elde edilen tüm veri ve bilgiler için kaynak gösterdiğimi ve bu kaynaklara kaynakçada yer verdiğimi; bu çalışmanın Eskişehir Teknik Üniversitesi tarafından kullanılan “bilimsel intihal tespit programı”yla tarandığını ve hiçbir şekilde “intihal içermediğini” beyan ederim. Herhangi bir zamanda, çalışmamla ilgili yaptığım bu beyana aykırı bir durumun saptanması durumunda, ortaya çıkacak tüm ahlaki ve hukuki sonuçları kabul ettiğimi bildiririm.

Zahir HAJIZADA

İÇİNDEKİLER

Sayfa

BAŞLIK SAYFASI	I
JÜRİ VE ENSTİTÜ ONAYI.....	II
DANIŞMAN ONAYI	III
ÖZET	IV
ABSTRACT.....	V
TEŞEKKÜR	VI
ETİK İLKE VE KURALLARA UYGUNLUK BEYANNAMESİ.....	VII
İÇİNDEKİLER.....	VIII
TABLolar DİZİNİ	XI
ŞEKİLLER DİZİNİ.....	XII
SİMGELER VE KISALTMALAR DİZİNİ.....	XIV
1. GİRİŞ	1
2. VERİ SETİNİN NORMAL DAĞILIMA UYGUNLUĞUNU BELİRLEYEN YÖNTEMLER	6
2.1. Grafiksel Teknikler	6
2.1.1. Kutu-çizgi grafiği.....	6
2.1.2. Histogram.....	7
2.1.3. Gövde-yaprak grafiği	8
2.1.4. Q-Q grafikleri	9
2.2. Normal Dağılım İçin Uyum İyiliği Testleri.....	11
2.2.1. Kolmogorov-Smirnov testi.....	14
2.2.2. Lilliefors testi	14
2.2.3. Cramer-von Mises testi	16
2.2.4. Anderson-Darling testi.....	16
2.2.5. D'Agostino-Pearson testi	17
2.2.6. Jarque-Bera testi	18
2.2.7. Shapiro-Wilk testi.....	18

3. NORMAL DAĞILIM İÇİN MAKİNE ÖĞRENMESİ ALGORİTMASINA DAYALI GELİŞTİRİLEN UYUM İYİLİĞİ PROSEDÜRÜ.....	20
3.1. Makine Öğrenmesinde Sınıflandırma Problemi	21
3.2. Normal Dağılım İçin Makine Öğrenmesi Algoritmasına Dayalı Geliştirilen Uyum İyiliği Prosedürünün Oluşturulmasında İzlenen Aşamalar	22
3.2.1. Normal dağılım için makine öğrenmesi algoritmasına dayalı geliştirilen uyum iyiliği prosedürü için veri setinin toplanması ve ön işleme.....	22
3.2.2. Normal dağılım için makine öğrenmesi algoritmasına dayalı geliştirilen uyum iyiliği prosedürü için ayırt edici özneliliklerin seçimi.....	23
3.2.2.1. <i>L-momentler</i>	24
3.2.2.2. <i>Kantiller</i>	25
3.2.2.3. <i>Entropi</i>	25
3.2.2.4. <i>Pearson çarpıklık ve basıklık katsayıları</i>	26
3.2.3. Normal dağılım için makine öğrenmesi algoritmasına dayalı geliştirilen uyum iyiliği prosedürü için model seçimi	27
3.2.3.1. <i>Kategoring boosting algoritması</i>	28
3.2.3.2. <i>Rasgele ormanlar algoritması</i>	29
3.2.4. Normal dağılım için makine öğrenmesi algoritmasına dayalı geliştirilen uyum iyiliği prosedüründeki modelin eğitilmesi	30
3.2.5. Modelin değerlendirilmesi	31
4. SİMÜLASYON ÇALIŞMASI.....	35
4.1. I. Tip Hata Değerlerine İlişkin Sonuçlar	38
4.2. Güç Değerlerine İlişkin Sonuçlar	40
4.3. Eğitim Setinde Tanıtılmayan Dağılımlara İlişkin Sonuçlar	61
4.4. Kategoring Boosting ve Rassal Ormanlar Algoritmalarının Karşılaştırılmasına İlişkin Sonuçlar.....	66
5. SONUÇ VE ÖNERİLER.....	71
KAYNAKÇA.....	73

ÖZGEÇMİŞ



TABLÖLAR DİZİNİ

Sayfa

Tablo 2.1. Literatürdeki çalışmalarda yaygın olarak kullanılan normallik testleri.	12
Tablo 2.2. İstatistiksel programlarda bulunan normallik testleri.	13
Tablo 3.1. Dağılımlara ait entropi değerleri.	26
Tablo 4.1. MLPN prosedürü ve Normallik testlerinin I. Tip hata değerleri.	38
Tablo 4.2. MLPN prosedürünün I. Tip hata oranı için sınıflandırma başarısını gösteren kriterler.	40
Tablo 4.3. MLPN prosedürü ve normallik testlerinin $Beta(2, 2)$, $Lojistik(0, 1)$ ve $Uniform(1)$ dağılımları altında güç değerleri.	41
Tablo 4.4. MLPN prosedürü ve Normallik testlerinin $Student t(5)$ ve $Laplace(0, 1)$ dağılımları altında güç değerleri.	46
Tablo 4.5. MLPN prosedürü ve Normallik testlerinin farklı α anlam düzeyleri için $Beta(2, 5)$, $Gama(4, 5)$, $Lognormal(0, 1)$, $Ki - Kare(5)$ ve $Üstel(1)$ dağılımları altında güç değerleri.	50
Tablo 4.6. MLPN prosedürünün güç değerleri için sınıflandırma başarısını gösteren kriterler.	58
Tablo 4.7. Eğitim setinde tanıtılmayan dağılımların $\alpha = 0,05$ anlam düzeyinde güç değerleri.	61
Tablo 4.8. CatBoost ve RF algoritmalarının $\alpha = 0.05$ anlam düzeyinde I. Tip hata ve güç değerleri.	66
Tablo 4.9. CatBoost ve RF algoritmalarının farklı ölçütlere göre performansı.	70

ŞEKİLLER DİZİNİ

Sayfa

Şekil 2.1. Farklı dağılımlara sahip 100 birimlik veri seti için kutu-çizgi grafikleri.	7
Şekil 2.2. Farklı dağılımlara sahip 100 birimlik veri seti histogram grafikleri.	8
Şekil 2.3. Farklı dağılımlara sahip 100 birimlik veri seti gövde-yaprak grafikleri.	9
Şekil 2.4. Farklı dağılımlara sahip 100 birimlik veri seti Q-Q plot grafikleri.	10
Şekil 3.1. Sınıflandırma problemine ait örnek.	21
Şekil 3.2. Catboost örnek şeması.	29
Şekil 3.3. RF örnek şeması.	30
Şekil 3.4. İkili sınıflandırmada karışıklık matrisi.	32
Şekil 4.1. MLPN prosedürü ve normallik testleri için I. Tip hatanın elde edilmesine yönelik akış şeması.	36
Şekil 4.2. MLPN prosedürü ve normallik testleri için güç değerlerinin elde edilmesine yönelik akış şeması.	37
Şekil 4.3. MLPN prosedürü ve Normallik testlerinin I. Tip hata değerlerinin grafik gösterimi.	39
Şekil 4.4. Grup 1 (a), Grup 2 (b), Grup 3'e (c) ait dağılım grafikleri.	41
Şekil 4.5. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için $Beta(2, 2)$ dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.	44
Şekil 4.6. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için $Lojistik(0, 1)$ dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.	45
Şekil 4.7. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için $Uniform(0, 1)$ dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.	46
Şekil 4.8. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için $Student t(5)$ dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.	49
Şekil 4.9. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için $Laplace(0, 1)$ dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.	50

Şekil 4.10. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için <i>Beta</i> (2, 5) dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.	55
Şekil 4.11. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için <i>Gama</i> (4, 5) dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.	56
Şekil 4.12. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için <i>Lognormal</i> (0, 1) dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.....	57
Şekil 4.13. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için <i>Ki – kare</i> (5) ve $\alpha = 0,01$ (d), $\alpha = 0,05$ (e) ve $\alpha = 0,10$ (f) anlam düzeyleri için <i>Üstel</i> (1) dağılımları altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.	58
Şekil 4.14. Eğitim setinde tanıtılmayan dağılımlar altında $\alpha = 0,05$ güven düzeyi için MLPN prosedürü ve normallik testlerine ait güç değerleri.	66
Şekil 4.15. CatBoost ve RF algoritmalarının performans karşılaştırması.	70

SİMGELER VE KISALTMALAR DİZİNİ

A	: Anlam düzeyi
AD	: Anderson-Darling testi
CHI	: Ki-Kare testi
CVM	: Cramer-von Mises testi
DAP	: D'Agostino-Pearson testi
JB	: Jarque-Bera testi
KS	: Kolmogorov-Smirnov testi
LL	: Lilliefors testi
MLPN	: Normal dağılım için makine öğrenmesi algoritmasına dayalı geliştirilen uyum iyiliği prosedürü
SW	: Shapiro-Wilk testi

1. GİRİŞ

İstatistiksel analizin temel testleri arasında yer alan ve parametrik testler olarak adlandırılan t-testi, varyans analizi (ANOVA), Pearson korelasyon anlamlılık testi ve benzeri testlerin temel varsayımı, analizi yapılacak veri setinin normal dağılıma sahip olmasıdır. Bu varsayım kısaca normallik varsayımı olarak adlandırılır. Normallik varsayımı sağlanmadan veya kontrol edilmeden yapılan parametrik testlerin sonucunda verilen kararlar, güvenilirliğini yitirmiş olmasının yanı sıra yanlış verilebileceğinden dolayı maddi kayıplara ve zaman kaybına yol açabilmektedir. Bu nedenle parametrik testlerle ilgili analizlere başlamadan önce analizi yapılacak olan veri setinin dağılımının normal dağılıma sahip olup olmadığının belirlenmesi gerekir. Diğer bir ifade ile normallik varsayımının sağlanıp sağlanmadığının belirlenmesi araştırmacıların önemle üzerinde durması gereken konulardan birisidir.

Bir veri setinin normal dağılıma uygunluğunu kontrol etmek için kullanılan en temel yöntem grafiksel yöntemlerdir. Fakat sadece grafiksel yöntemleri kullanıp dağılımla ilgili karar vermek doğru değildir. Bunun nedeni grafiği yorumlayan araştırmacının gerçekte normal dağılıma sahip olmayan bir veri setinin normal dağılmış gibi veya aksine normal dağılıma sahip bir veri setinin normal dağılmamış gibi değerlendirebileceğinden dolayı grafiksel yöntemler kişiye bağlı subjektif yöntemler olarak kabul edilir. Bundan dolayı bir veri setinin normal dağılıp dağılmadığı ile ilgili kararı grafiksel yöntemlerin yanı sıra normal dağılım için uyum iyiliği testlerinin yapılması araştırmanın güvenilirliğini artıracaktır.

Normallik testleri, örnekleme ait veri setinin normal dağılıma ne kadar uyduğunu veya bu dağılımdan ne kadar saptığını kontrol etmek amacıyla kullanılan uyum iyiliği (goodness-of-fit) testlerine verilen genel addır. Veri setinin normal dağılıma sahip olup olmadığının belirlenmesi yukarıda da vurgulandığı gibi oldukça önemli olduğundan literatürde çok fazla normallik testleri geliştirilmiş ve geliştirilmeye devam etmektedir. Ancak literatürde kullanılan her bir normallik testinin kendine göre bazı varsayımları vardır. Bu durumda farklı örneklem büyüklüklerine veya veri yapılarına sahip veri setleri için farklı normallik testlerinin kullanılması gerekliliği doğmaktadır. Çünkü bir veri setinin normal dağılıp dağılmadığının sınanması için kullanılan farklı normallik testleri birbiriyle çelişkili sonuçlar verebilmektedir. Bundan dolayı hangi durumlarda hangi normallik testinin kullanılması önemli bir problem haline gelmiştir. Bu problemin aşılabilmesi için yeni normallik testlerinin geliştirilmesinin yanı sıra hangi durumda hangi

normallik testinin kullanılacağıının belirlenmesi için kapsamlı simülasyon çalışmaları yapılmaktadır.

Literatürde yaptığımız araştırmalara göre şimdiye kadar kırktan fazla normallik testi geliştirilmiştir. Normallik testlerine yönelik ilk çalışmalar 1900'li yılların başlarından itibaren başlamıştır. İlk olarak Pearson (1900) tarafında ki-kare testi önerilmiştir. Bundan sonra Cramer- von Mises (1928), Kolmogorov ve Smirnov (1933), Anderson-Darling (1952), Kuiper (1962), Lilliefors (1967), Shapiro-Wilk (1965), Ajne (1968), D'Agostino-Pearson (1972), Jarque-Bera (1987), Zhang (2002), Esteban (2007) testleri, geliştirilen normallik testlerinden önemli olanlarıdır.

Fakat bu testler belirli bir koşul ve varsayım altında uygulanabilir. Farklı örneklem büyüklükleri altında testler farklı sonuçlar çıkarabilmektedir. Örneğin, küçük veya orta boyutlu örneklem verildiğinde bazı testler H_0 hipotezini reddederken, bazıları bu hipotezi reddetmekte başarısız olabilir. Bu yüzden birbiri ile çelişkili sonuçlar araştırmacılar için çözülmesi gereken bir sorun haline gelmiştir. Son yıllarda normallik testleri ile ilgili yapılan çalışmaların büyük bir kısmı bu testlerin hangisinin daha güçlü olduğunu bulmaya yöneliktir.

Yazıcı ve Yolaçan (2007), yaptıkları çalışmada on iki tane normallik testini incelemiş ve Monte Carlo (MC) simülasyon çalışmasıyla bu testlerin güçlerini karşılaştırmışlardır. Bu çalışmadan çıkan sonuçlara göre ele alınan normallik testlerine ait güç değerlerinin örneklem büyüklüğü arttıkça arttığı gözlemlenmiştir. Ayrıca küçük örneklem sahip simetrik dağılımlar için Kolmogorov-Smirnov, Modifiye edilmiş Kolmogorov-Smirnov ve Anderson-Darling testleri önerilmiş, asimetrik dağılımlar için ise Shapiro-Wilk testinin daha güçlü olduğu tespit edilmiştir.

Özer (2007), yüksek lisans tez çalışmasında literatürde bulunan dokuz tane normallik testinin performansını I. Tip hata ve testin gücüne göre karşılaştırmıştır. Söz konusu karşılaştırmalar sonucunda ele alınan testlerden I. Tip hata olasılığı bakımından Jarque-Bera testinin, testin gücü bakımından ise Shapiro-Wilk testinin diğer testlerden daha iyi performans gösterdiği gözlemlenmiştir.

Breton ve ark. (2008), çalışmalarında normalliği belirlemek konusunda testlerin gücünü karşılaştırmak için yeni bir yaklaşım önermiştir. Bu çalışmayla Johnson dağılım sistemleri kullanılarak verinin ortalama, varyans, çarpıklık ve basıklık değerleri ile normallik testlerinin gücü tahmin edilmiştir.

Romao ve ark. (2009), araştırmacılar tarafından yaygın olarak kullanılmayan otuz üç tane normallik testini çeşitli örneklem büyüklükleri ve farklı anlam düzeylerini göz önünde bulundurarak karşılaştırmalar yapmıştır.

Noughabi ve Arghami (2010), yedi tane normallik testinin çeşitli örneklem büyüklüklerine göre güçlerinin karşılaştırmasını MC simülasyon çalışmasını kullanarak yapmışlardır. Çalışmada kullanılan alternatif dağılımlar farklı özelliklerine göre dört gruba ayrılmış ve her grup için hangi testin daha güçlü olduğu ortaya koyulmuştur.

Razali ve Wah (2010) çalışmalarında dört tane normallik testini incelemiş ve MC simülasyonu ile bu testlerin güçlerini karşılaştırmışlardır. Çalışmadan çıkan sonuca göre farklı örneklem büyüklüklerinde kullanılan alternatif dağılımlar için en güçlü test Shapiro-Wilk, en az güce sahip olan test ise Kolmogorov-Smirnov testi olmuştur.

Yap ve Sim (2011), simetrik kısa kuyruklu, simetrik uzun kuyruklu ve asimetrik yapılara sahip alternatif dağılımları kullanarak sekiz tane normallik testinin güçlerini MC simülasyonu yardımıyla karşılaştırmıştır. Araştırmanın sonuçlarına göre Shapiro-Wilk testinin diğer testlerden daha güçlü olduğu ortaya çıkmıştır.

Yıldırım ve Gökpınar (2012), bazı normallik testlerini tanıtmış ve bu testlerin güçlerini $(-\infty, +\infty)$ ile $(0, +\infty)$ aralığında tanımlı çeşitli simetrik ve asimetrik dağılımlara göre karşılaştırmıştır. Çalışmada ayrıca hangi testin hangi durumlarda birbirlerine göre daha başarılı olduklarını belirlemişlerdir.

Büyükuysal (2014), yaptığı doktora tez çalışmasında en yaygın olarak kullanılan ve istatistik paket programlarında bulunan beş tane normallik testini seçerek farklı örneklem büyüklüklerinde dağılımlar üretmiş ve MC simülasyonu ile bu testlerin I. Tip hata ve güçlerini karşılaştırmıştır. Simülasyon çalışması sonucunda Shapiro-Wilk testi en iyi sonucu verirken, örneklem büyüklüğü azaldıkça Anderson-Darling testinin de Shapiro-Wilk testi kadar iyi sonuçlar verdiği gözlemlenmiştir.

Lee ve ark. (2016), çalışmalarında farklı kuyruk yapılarına sahip simetrik ve asimetrik dağılımlar altında belirledikleri normallik testlerini I. Tip hata ve güçleri bakımından karşılaştırmıştır. Çalışmanın sonucunda tüm dağılımlar için güçlü bir testin olmadığı, testin gücünün verinin dağılımına bağlı olduğu ortaya çıkarılmıştır.

Bilim (2017), yüksek lisans tez çalışmasında R paket programı ile farklı büyüklük ve özellikte veriler üreterek literatürde yaygın olarak kullanılan dokuz tane normallik testinin karşılaştırmasını yapmıştır. Tez çalışması sonucunda ele alınan testlerin I. Tip

hata oranlarının benzer çıktığı, güç değerlerinde ise Shapiro-Wilk ve Shapiro-Francia testlerinin önde olduğu gözlemlenmiştir.

Ogunleye ve ark. (2018), çalışmalarında literatürde yaygın olarak kabul edilen dört tane normallik testini kullanmış ve MC simülasyon tekniğiyle bu testlerin I. Tip hata ve güçleri bakımından karşılaştırmasını yapmıştır. Seçilen testler arasında Shapiro-Wilk testinin diğer testlerden normalliği belirleme konusunda daha başarılı olduğu kanaatine varılmıştır.

Doulah (2019), çalışmasında yirmi yedi tane normallik testini farklı örnek büyüklüklerine sahip veriler için MC simülasyon tekniğini kullanarak karşılaştırmıştır. Simülasyon çalışmasında kullanılan veriler simetrik (uzun ve kısa kuyruklu) ve asimetrik (uzun ve kısa kuyruklu) alternatif dağılımlardan üretilmiştir. Yapılan çalışma sonucunda hangi senaryo için hangi normallik testinin daha güçlü olduğu gösterilmiştir.

Wijekularathna ve ark. (2019) yaptıkları çalışmada on iki tane normallik testinin güç bakımından karşılaştırmasını simetrik (uzun ve kısa kuyruklu) ve asimetrik (uzun ve kısa kuyruklu) dağılımlar altında yapmışlardır. Ele alınan normallik testlerinin gösterdikleri güç performansları farklı örneklem büyüklükleri için yorumlanmıştır.

Bayoud (2021) yaptığı çalışmada literatüre yeni bir normallik testi sunmuş ve önerilen bu testin güç değerlerini çalışmada ele alınan normallik testleri ile karşılaştırmıştır.

Khatun (2021) yaptığı çalışmada verilerin normalliğini test ederken grafiksel yöntemlerin yetersiz kaldığını, normallik testleriyle bunun desteklenmesi gerektiğini, aynı zamanda normallik testlerinin de küçük örneklem büyüklüklerinde hassas olduğunu, örneklem büyüklüğü arttıkça testlerin gücünün de arttığını ortaya koymuştur. Bu çalışmada ele alınan normallik testleri arasında en güçlü testin Shapiro-Wilk, en az güce sahip testin ise Jarque-Bera testi olduğu gösterilmiştir.

Demir (2022) yaptığı çalışmasında on tane normallik testini farklı örneklem büyüklüklerinde MC simülasyon tekniğiyle karşılaştırmıştır. Çalışmanın sonuçlarına göre, çarpıklık ve basıklık katsayıları sıfıra eşit veya yakın olduğunda normallik testinin performansı örneklem büyüklüğünden etkilenmemektedir.

Literatürden de görüldüğü gibi tüm veri yapıları için güçlü bir testin olmadığı, ayrıca aynı veri yapıları için de yapılan karşılaştırmalarda farklı normallik testlerinin birbiriyle çelişkili sonuçlar verebileceği gözlemlenmektedir. Bu nedenle bu tez çalışmasında tüm veri yapıları için verinin normal olup olmadığına karar verebilecek

makine öğrenmesi algoritmasına dayalı bir prosedür önerilmiştir. Önerilen bu prosedürün performansı I. Tip hata ve testin gücü bakımından literatürde yaygın olarak kullanılan yedi normallik testi ile karşılaştırılmıştır.

Bu tez çalışmasının ikinci bölümünde verinin normal dağılıp dağılmadığını belirlemek için kullanılan grafiksel yöntemler, daha sonra ise literatürde yaygın olarak kullanılan yedi tane normallik testi hakkında genel bilgi verilmiş, her bir testin test istatistiğinin hesaplanmasına dair formüller gösterilmiştir.

Üçüncü bölümde makine öğrenmesi, makine öğrenmesinde sınıflandırma problemi, normal dağılım için makine öğrenmesi algoritmasına dayalı geliştirilen uyum iyiliği prosedürünün (MLPN) oluşturulmasındaki aşamalar anlatılmış, kullanılan algoritmalar ve dağılımların ayırt edici özellikleri hakkında bilgi verilmiştir.

Dördüncü bölümde çalışmanın ikinci bölümünde yer alan yedi tane normallik testi ve üçüncü bölümünde literatüre yeni sunulan MLPN prosedürünün I. Tip hata oranları ve testin gücü bakımında performansları farklı örneklem büyüklükleri ve farklı veri yapıları için MC simülasyon tekniğiyle karşılaştırılmış ve elde edilen sonuçlar yorumlanmıştır.

Son olarak beşinci bölümde ise tez çalışmasına dair sonuçlar verilmiş, önerilen normallik prosedürünün güçlü ve zayıf tarafları, çalışmada ele alınan diğer normallik testlerinden farklı yanları belirtilmiş, gelecekte bu konuyla ilgili neler yapılabileceğine dair önerilerde bulunulmuştur.

2. VERİ SETİNİN NORMAL DAĞILIMA UYGUNLUĞUNU BELİRLEYEN YÖNTEMLER

Bir veri setinin normal dağılıp dağılmadığını belirlemek için temel olarak grafiksel yöntemler ve uyum iyiliği testleri kullanılır. Bu bölümde literatürde yaygın olarak kullanılan grafiksel yöntemler ve normal dağılım için uyum iyiliği testleri hakkında genel bilgiler verilerek tanıtılmıştır.

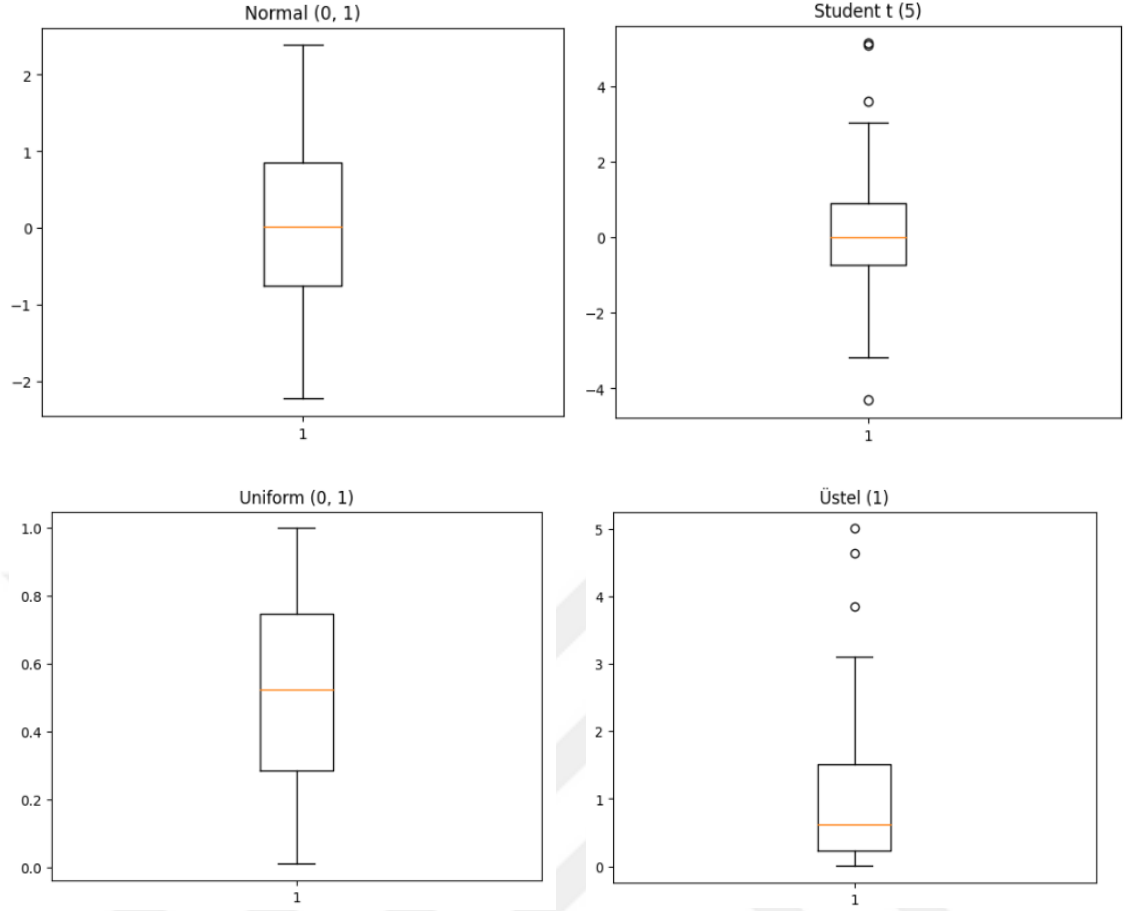
2.1. Grafiksel Teknikler

2.1.1. Kutu-çizgi grafiği

Kutu-çizgi grafiği, veri dağılımlarını görsel olarak incelemek ve normallik varsayımını değerlendirmek için kullanılan bir grafik tekniğidir. Bu grafikte ortalamaya göre büyük ve küçük değerlerin yoğunluklarının simetriden ne kadar uzaklıkta olduğunu, aykırı değerlerin olup olmadığını görmek mümkündür. Ortalama etrafındaki bölgeler simetrik, sağda ve solda bulunan kuyruklar dar ve aykırı değerler yoksa dağılımın normal olabileceği çıkarımı yapılabilir. Kutu-çizgi grafiğini açıklayan temel özellikler aşağıdakilerdir (Akdeniz, 2018):

- **Kutu:** Kutunun alt sınırı verinin alt çeyreğini (Q_1), üst sınırı ise verinin üst çeyreğini (Q_3) temsil eder. Yani kutunun alt sınırından üst sınırına kadar olan kısmı, verinin ortadaki %50'sini içerir.
- **Orta çizgi:** Kutunun içindeki yatay çizgi, verinin ortanca değerini gösterir. Veri seti sıralandığında bu çizgi, verinin yarısını altında ve yarısını da üstünde bırakır.
- **Bıyıklar:** Kutunun dışındaki dikey çizgiler, verinin çeyrek aralığını temsil eder. Üst bıyık, verinin üst çeyrek aralığının dışındaki en yüksek değeri gösterir. Alt bıyık ise verinin alt çeyrek aralığının dışındaki en düşük değeri gösterir.
- **Aykırılar:** Bıyıkların dışındaki, aykırı değerleri gösteren noktalar veya dairelerdir. Aykırı değerler, veri setinin genel dağılımından önemli ölçüde sapan değerleridir.

Simetrik ve asimetrik dağılımların kutu-çizgi grafikleri arasındaki görsel farklar aşağıdaki şekilde gösterilmiştir.



Şekil 2.1. Farklı dağılımlara sahip 100 birimlik veri seti için kutu-çizgi grafikleri.

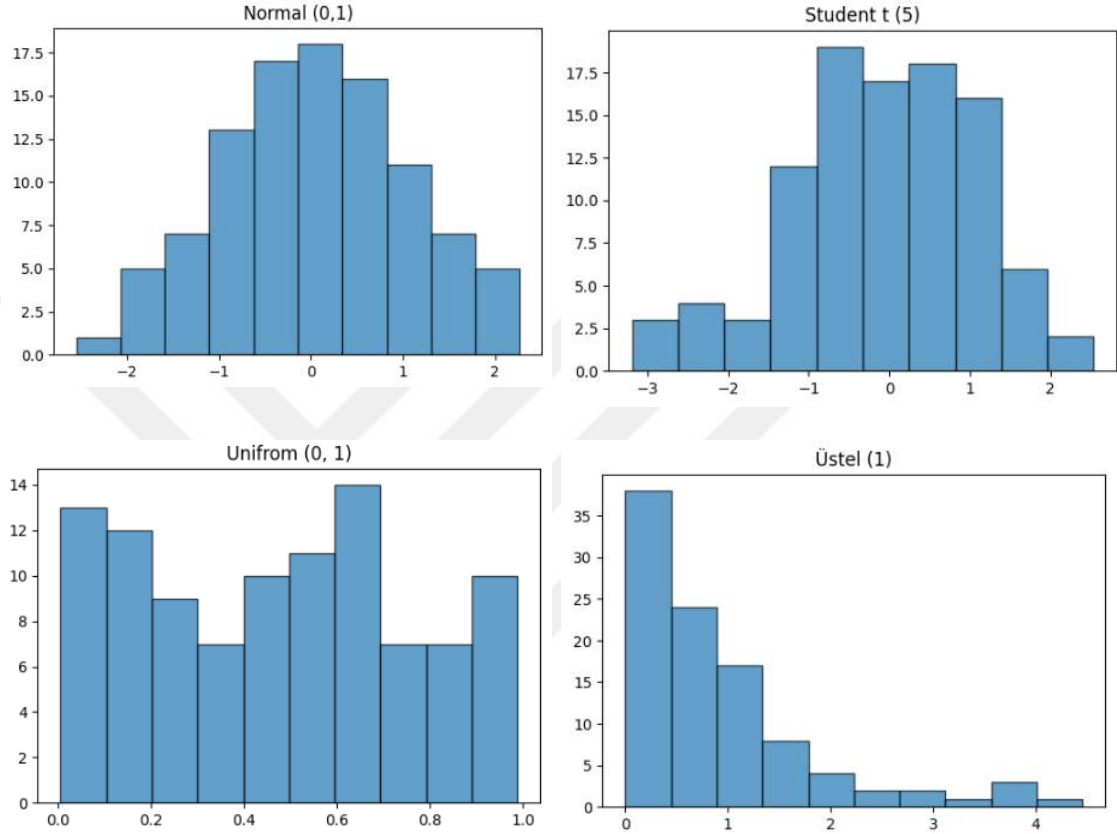
2.1.2. Histogram

Histogram, veri dağılımının görsel olarak incelenmesi ve normallik varsayımının kontrol edilmesi için sıkça kullanılan bir diğer grafik yöntemidir. Bu grafik, veri değerlerinin hangi aralıklarda ne kadar sıklıkla bulunduğunu gösterir. Histogram grafiğini açıklayan temel özellikler aşağıdakilerdir (Akdeniz, 2018):

- **Sütunlar:** Histogram, veri aralıklarını temsil eden dikdörtgen sütunlardan oluşur. Her sütun, belirli bir veri aralığında kaç veri noktasının bulunduğunu gösterir. Sütunların yüksekliği, bu aralıktaki veri noktalarının sıklığını yansıtır.
- **Veri aralıkları:** Veri seti, belirli sayıda aralığa bölünür. Bu aralıklara "bin" veya "class" denir. Her sütun, bir aralığın sınırlarını temsil eder ve bu aralıktaki veri noktalarının sayısını gösterir.
- **X-eksen:** X-ekseni, veri aralıklarını veya sınırlarını temsil eder. Bu eksen, veri setinin değerlerini belirli aralıklara böler.

- **Y-eksen:** Y-ekseni, her bir aralıktaki veri noktalarının sıklığını veya sayısını temsil eder. Y-ekseni, her sütunun yüksekliğini belirler.

Simetrik ve asimetrik dağılımların histogram grafikleri arasındaki görsel farklar aşağıdaki şekilde gösterilmiştir.



Şekil 2.2. Farklı dağılımlara sahip 100 birimlik veri seti histogram grafikleri.

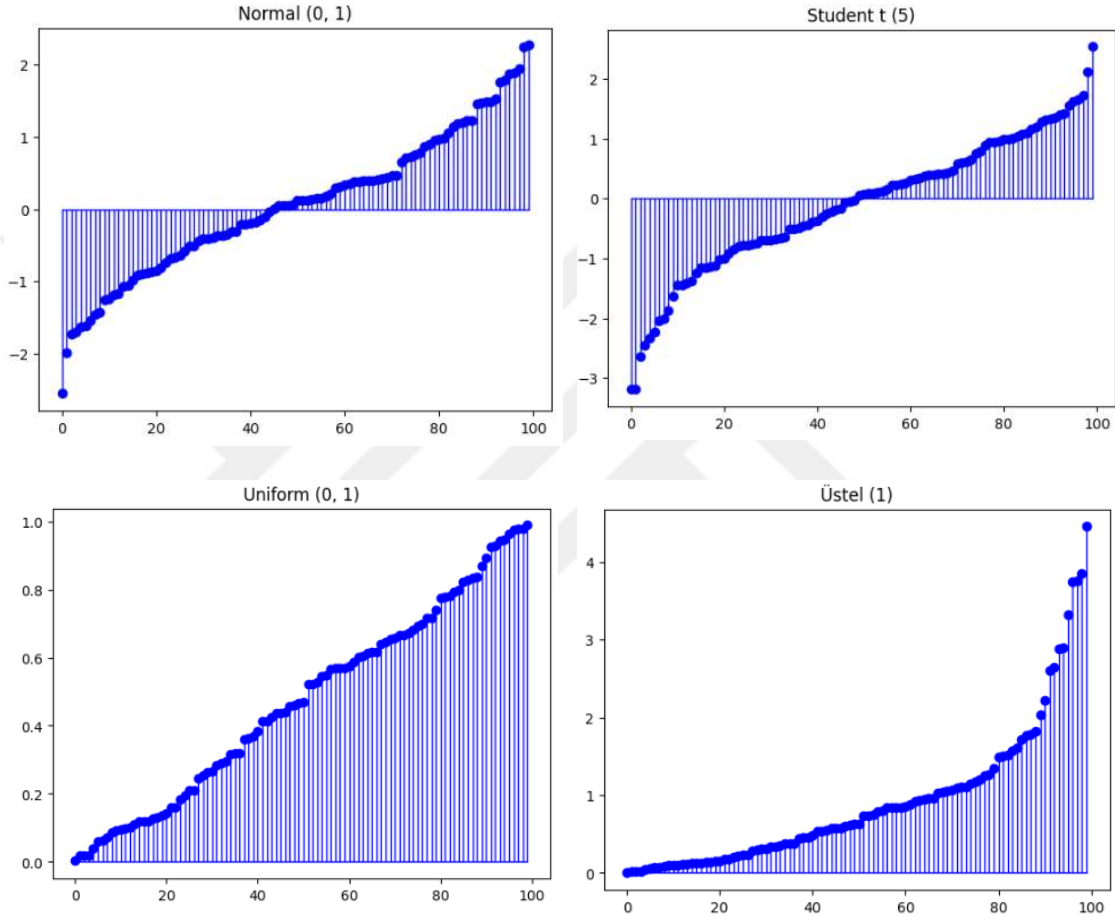
2.1.3. Gövde-yaprak grafiği

Gövde yaprak grafiği, veri setinin dağılımını görsel olarak temsil etmek için kullanılan bir yöntemdir. Bu grafik, her veri noktasının değerini ve dağılımını belirtmek için kullanılır. Gövde yaprak grafiği, özellikle küçük örneklem büyüklüklerinde veya sadece birkaç veri noktasının bulunduğu durumlarda veri dağılımının anlaşılmasına yardımcı olur. Gövde-yaprak grafiğini açıklayan temel özellikler aşağıdakilerdir (Akdeniz, 2018):

- **Gövde:** Grafikte dikey olarak görünen sayılar gövdeyi temsil eder. Gövde, veri değerlerinin yüzler ve onlar basamağını ifade eder. Genellikle solda düşük değerlerden başlayarak yüksek değerlere doğru sıralanır.

- **Yapraklar**, gövdeye bağlı olan yatay olarak görünen sayıları temsil eder. Yapraklar, veri değerlerinin onlar ve birler basamağına karşılık gelir. Her yaprak, bir veri noktasının tam değerini ifade eder.

Simetrik ve asimetrik dağılımların histogram grafikleri arasındaki görsel farklar aşağıdaki şekilde gösterilmiştir.



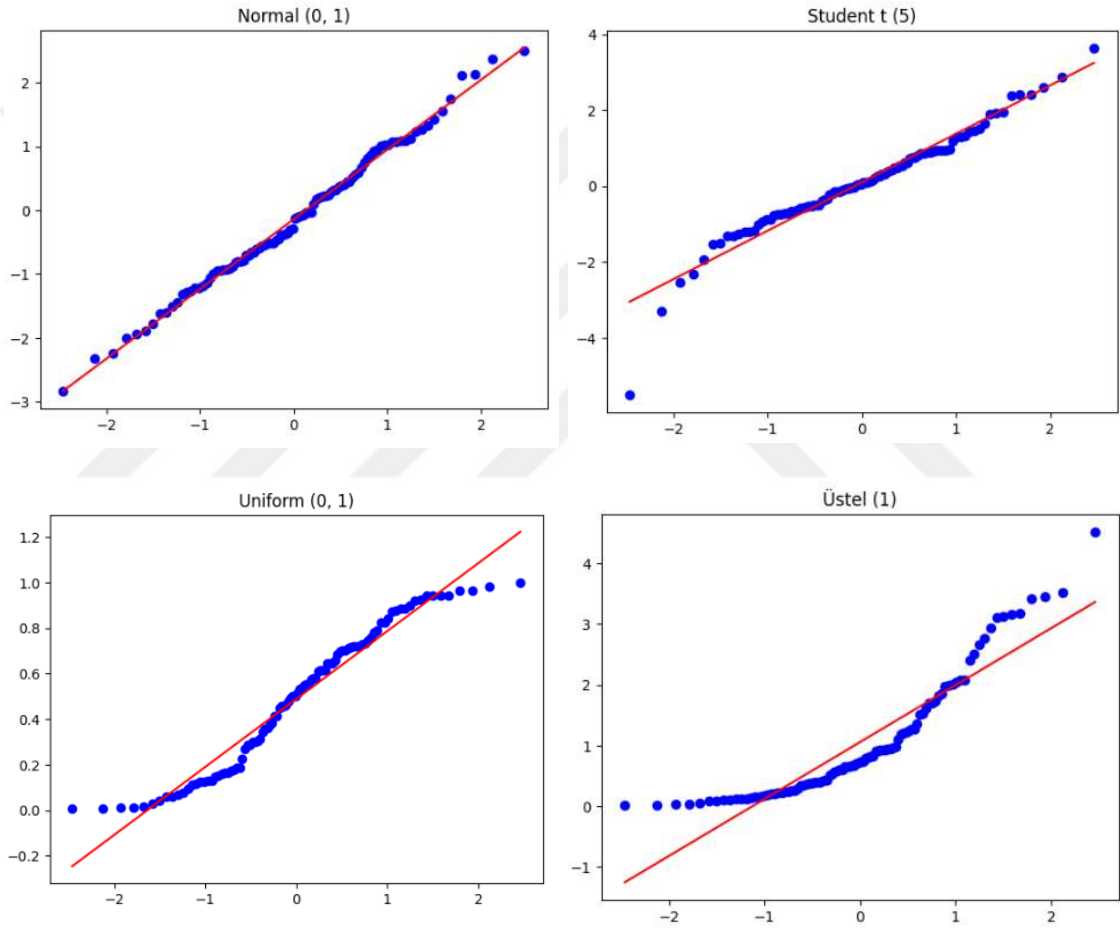
Şekil 2.3. Farklı dağılımlara sahip 100 birimlik veri seti gövde-yaprak grafikleri.

2.1.4. Q-Q grafikleri

Normallik varsayımını kontrol etmek için kullanılan Q-Q (quantile-quantile) grafiği, bir veri setinin normal dağılıma ne kadar uyduğunu değerlendirmek için en yaygın kullanılan grafik yöntemidir. Q-Q grafiği, gözlemlenen veri noktalarının ve teorik bir normal dağılımın kantil değerleri arasındaki ilişkiyi görselleştirir. Eğer veriler normal bir dağılıma uyuyorsa, Q-Q grafiği noktaların yaklaşık olarak bir çizgi boyunca düzgün bir şekilde yerleşmesi beklenir. Q-Q grafiğini açıklayan temel özellikler aşağıdakilerdir (Akdeniz, 2018):

- **Gözlemlenen değerler:** Q-Q grafiğinde x-ekseninde gözlemlenen veri noktaları yer alır. Bu noktalar, veri setinin sıralanmış değerleridir.
- **Kantil değerleri:** Q-Q grafiğinde y-ekseninde teorik bir normal dağılımın kantil değerleri bulunur. Bu değerler, normal dağılımın belirli kesirlerine karşılık gelir.

Simetrik ve asimetrik dağılımların Q-Q grafikleri arasındaki görsel farklar aşağıdaki şekilde gösterilmiştir.



Şekil 2.4. Farklı dağılımlara sahip 100 birimlik veri seti Q-Q plot grafikleri.

Grafiksel yöntemler gözlemsel değerlerin teorik değerlere uygunluğunu sadece görsel bir biçimde sunduğu ve normallikle ilgili kararın araştırmacıdan araştırmacıya değişebileceğinden dolayı normal dağılım varsayımının kontrol edilmesinde objektif yöntemler olmaktan çıkmaktadır. Bundan dolayı bir veri setinin normal dağılım varsayımını sağlayıp sağlamadığını daha objektif ve güçlü yöntemlerle belirlemek için normallik testlerinin kullanılması gereklidir.

2.2. Normal Dağılım İçin Uyum İyiliği Testleri

Anakütleden çekilen bir veri setinin belirli bir dağılıma uyup uymadığını belirlemek için kullanılan istatistiksel testlere uyum iyiliği testleri denir. Uyum iyiliği testleri, çeşitli istatistiksel dağılımlar için kullanılmaktadır: Normal, Uniform, Weibull ve benzer dağılımlar (Gamgam ve Altunkaynak, 2012).

Tez çalışmasının bu bölümünde normal dağılım için uyum iyiliği testleri incelenmiştir. Örneklem verilerinin veya veri setinin normal dağılıma sahip bir anakütleden çekilip çekilmediğini test etmek amacıyla kullanılan istatistiksel testlere normal dağılım için uyum iyiliği testleri veya kısaca normallik testleri denir. Normallik testlerinin uygulanmasında öncelikle hipotezler aşağıda ifade edildiği gibi kurulur:

H_0 (Sıfır hipotezi): Örneklem verisi, normal dağılıma ait bir anakütleden çekilmiştir.

H_1 (Alternatif hipotez): Örneklem verisi, normal dağılıma ait bir anakütleden çekilmemiştir.

Kurulan bu hipotezlerden sonra sıfır hipotezinin reddedilip reddedilmeyeceği hakkında karar vermek amacıyla örneklem verisine göre test istatistiği hesaplanır. Sıfır hipotezinin reddedileceği test istatistiği değerlerinden oluşan küme "ret bölgesi" diye isimlendirilir. Bunun yanı sıra "kabul bölgesi" de H_0 hipotezinin reddedilemeyeceği test istatistiği değerlerinden oluşan kümedir. Ret bölgesi ile kabul bölgesi arasında oluşan sınır test istatistiğinin kritik değeridir. Eğer test istatistiğinin hesaplanan değeri ret bölgesine ait sınırlar içerisine düşerse bu verinin H_0 hipotezinde bahsi geçen normal dağılımdan gelmediği kararına varılır.

Sıfır hipoteziyle ilgili karar verildiği zaman iki tip hata yapılabilir. Bunlar, I. Tip hata ve II. Tip hatalardır. H_0 doğru olduğunda H_0 'ın reddedilmesi olasılığı I. Tip hatayı ifade ederken, H_0 yanlış olduğunda H_0 'ın reddedilememesi ise II. Tip hatayı ifade eder. Testin gücü ise H_0 yanlış olduğunda H_0 'ın reddedilmesi olasılığıdır (Akıllı, 2020).

Tüm uyum iyiliği testlerinde olduğu gibi normallik testlerinde de önceden kabul edilen I. Tip hata (α anlam düzeyi) ve testin gücü bakımından en iyi performansı gösteren testle karar vermek istenilmektedir. Bu nedenle literatürde birçok normal dağılım için uyum iyiliği testi önerilmiştir. Ki-kare, Shapiro-Wilk, Kolmogorov-Smirnov, Cramer-von Mises, Lilliefors, Anderson-Darling, Jarque-Bera testleri literatürde en yaygın kullanılan ve bilinen testlerdir. Bu testlerden bazıları tüm dağılımlar için genel olarak

kullanılırken bazıları da sadece normal dağılım için kullanılmaktadır. Ancak, her normallik testinin kendine özgü varsayımları ve özellikleri vardır. Hangi normallik testinin güçlü bir şekilde kullanılacağı veri setinin yapısına ve araştırmanın amacına bağlı olarak değişebilir.

Bu tez çalışmasında normal dağılım için makine öğrenmesi algoritmasına dayalı uyum iyiliği prosedürü önerilmiş ve literatürdeki normallik testlerine göre I. Tip hata oranı ve testin gücü bakımından performansının değerlendirilmesi amaçlanmıştır. Bu amaç doğrultusunda karşılaştırma için seçilen normallik testleri, literatürdeki çalışmalarda yaygın olarak kullanılan ve istatistiksel paket programlarında mevcut olan normallik testleridir.

Literatürde yaygın olarak kullanılan ve istatistiksel paket programlarda bulunan normallik testleri sırasıyla **Tablo 2.1** ve **Tablo 2.2**'de özetlenmiştir.

Tablo 2.1. Literatürdeki çalışmalarda yaygın olarak kullanılan normallik testleri.

Çalışmalar	KS	LL	CVM	AD	DAP	JB	SW	CHI	SF	KP
Yazici ve Yolaçan (2007)	✓			✓	✓	✓	✓	✓		✓
Özer (2007)	✓	✓		✓	✓	✓	✓	✓		
Breton ve ark. (2008)	✓	✓			✓			✓		
Romao ve ark. (2009)				✓	✓	✓	✓		✓	
Noughabi ve Arghami (2010)	✓		✓	✓		✓	✓			
Razali ve Wah (2010)	✓	✓		✓			✓			
Yap ve Sim (2011)	✓	✓	✓	✓	✓	✓	✓	✓		
Yıldırım ve Gökpınar (2012)	✓			✓		✓				
Büyüküysal (2014)	✓		✓	✓		✓	✓			
Lee ve ark. (2016)	✓	✓	✓	✓		✓	✓		✓	

Tablo 2.1. (Devam) *Literatürdeki çalışmalarda yaygın olarak kullanılan normallik testleri.*

Bilim (2017)	✓	✓	✓	✓	✓	✓	✓		✓	
Ogunleye ve ark. (2018)	✓			✓			✓	✓		
Doulah (2019)		✓	✓	✓	✓	✓	✓	✓	✓	✓
Wijekularathna ve ark. (2019)	✓		✓	✓		✓	✓			
Bayoud (2021)	✓		✓	✓		✓	✓		✓	✓
Khatun (2021)		✓	✓	✓		✓	✓	✓	✓	
Demir (2022)	✓	✓	✓	✓	✓	✓	✓		✓	

Tablo 2.2. *İstatistiksel programlarda bulunan normallik testleri.*

	KS	LL	CVM	AD	DAP	JB	SW	CHI	SF	KP
SPSS	✓	✓		✓			✓	✓	✓	
Minitab	✓			✓	✓		✓			
STATA	✓			✓	✓	✓	✓			
SAS	✓	✓	✓	✓	✓	✓	✓	✓		
MATLAB	✓	✓		✓	✓	✓	✓	✓		
Eviews		✓				✓	✓			
R	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
Python	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓

Normallik testleri ile ilgili yapılan literatür araştırmasında ulaşabildiğimiz on yedi tane çalışmanın yaklaşık olarak %94'ünde Anderson-Darling (AD), %88'inde Shapiro-Wilk (SW), %82'sinde Kolmogorov-Smirnov (KS) ve Jarque-Bera (JB), %64'ünde D'Agostino-Pearson (DAP), %58'inde Cramer-von Mises (CVM), %52'inde Lilliefors (LL), %41'inde Ki-Kare (CHI) ve Shapiro-Francia (SF), %23'ünde ise Kuiper (KP) normallik testleri kullanılmıştır. Benzer sonuçlar istatistiksel paket programlarında bulunan normallik testleri için de geçerlidir. Buna göre bu tez çalışmasında literatürde en az %50 oranında kullanılmış olan KS, LL, CVM, AD, DAP, JB, SW normallik testleri seçilmiş, bu normallik testleri ile ilgili genel bilgiler aşağıda sunulan alt başlıklarda sırasıyla verilmiştir.

2.2.1. Kolmogorov-Smirnov testi

Kolmogorov-Smirnov (KS) testi ilk olarak Kolmogorov tarafından 1933 yılında önerilmiştir. Bundan altı yıl sonra yani 1939 yılında Smirnov da iki bağımlı olmayan örnek için bir normallik testi geliştirmiştir. Uygulamadaki benzerlikleri açısından bu testler literatürde KS testi olarak bilinir (Kolmogorov, 1933; Smirnov, 1939). Küçük örnek büyüklüklerinde ki-kare normallik testinin doğruluğunun şüpheli ve herhangi bir başka örnek büyüklüğü için de genellikle ki-kareden daha güçlü olması sebebinden KS testi daha kullanışlıdır (Lilliefors, 1967).

KS test istatistiği, dağılımın tüm parametreleri bilindiği durumlarda elde edilebilir. Bundan dolayı parametrelerden en az bir tanesinin bilinmediği durumlarda KS test istatistiği için hesaplanan kritik değer tabloları geçersiz hale gelir (Köle, 2014).

KS testi, yokluk hipotezinde gösterilen birikimli dağılım fonksiyonu olan $F(x)$ ile örneklem birikimli dağılım fonksiyonu $F_n(x)$ arasında olan maksimum uzaklığı test istatistiği şeklinde ele alır. KS testinde hipotezler aşağıdaki şekilde kurulur:

$$H_0: F_n(x) = F(x)$$

$$H_1: F_n(x) \neq F(x)$$

Normal dağılım için KS testinde $F(x)$, standart normal dağılımın birikimli dağılım fonksiyonudur. KS testinin test istatistiği olan D_n ise aşağıdaki formülle hesaplanır:

$$D_n = \sup_x \{|F(x) - F_n(x)|\} \quad (2.1)$$

Bu formülle hesaplanan D_n değeri, örneklem büyüklüğü olan n ve α anlam düzeyine karşılık gelen ve KS tablosundan bulunan D_k kritik değeriyle karşılaştırılır. Karşılaştırma sonucunda $D_n \geq D_k$ oluyorsa H_0 hipotezi reddedilir. Aksi olduğu durumda H_0 hipotezi reddedilemez.

2.2.2. Lilliefors testi

Normal dağılıma uygunluğu kontrol etmek için kullanılan testlerden bir diğeri de KS testinin bir uyarlaması olan ve Hubert W. Lilliefors tarafından önerilen Lilliefors (LL) testidir (Lilliefors, 1967). LL testi oldukça küçük örnek büyüklükleri için kullanılabilir ve asimptotik olarak ki-kare testinden güçlü bir testtir (Akıllı, 2020).

Eğer dağılım ve dağılımın parametreleri biliniyorsa KS testi doğru bir şekilde uygulanabilmektedir. Fakat dağılıma ait bazı parametreler örneklem verilerinden tahmin ediliyor ise KS testinin kullanılmasında kritik tablo değerleri doğru bir çıkarsama yapılmasına olanak vermez. Bu durumda KS testi yerine Lilliefors'un önerdiği kritik değerler kullanılmalıdır (Gamgam ve Altunkaynak, 2012). LL testine yönelik hipotezler aşağıdaki gibidir:

H_0 : Örneklem verisi, varyansı ve ortalaması bilinmeyen normal dağılıma sahip bir anakütleden gelmektedir.

H_1 : Örneklem verisi, normal dağılıma sahip bir anakütleden gelmemektedir.

$X_1, X_2, \dots, X_n$ rassal örneklem olmak üzere, normal dağılımın bilinmeyen μ ve σ^2 parametrelerinin en iyi yansız tahmin edicileri sırasıyla aşağıdaki gibi gösterilsin:

$$\bar{X} = \frac{\sum_{i=1}^n X_i}{n} \quad (2.2)$$

ve

$$s^2 = \frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n - 1} \quad (2.3)$$

H_0 hipotezinde ifade edilen normal dağılımın birikimli dağılım fonksiyonu $F(x)$, μ ve σ^2 'nin değerleri yerine onların tahmin değerleri olan $\bar{x}$ ve s^2 değerleri kullanılarak aşağıdaki gibi elde edilir:

$$F(x) = P(X \leq x) = P\left(Z \leq \frac{x - \bar{x}}{s}\right) \quad (2.4)$$

LL testinde n tane veriden oluşan örneklem için dağılım fonksiyonu olan $F_n(x)$ fonksiyonu Z tablo değerleri kullanılarak hesaplanmaktadır. LL testinin test istatistiği aşağıdaki formülün yardımıyla elde edilir:

$$D_h^* = \max_x = |F^*(X) - F_n(X)| \quad (2.5)$$

Bu formülle hesaplanan LL test istatistiği değeri D_h^* , α ve n değerine uygun olarak LL tablosundan bulunan D_k kritik değeriyle karşılaştırılır ve $D_h^* \geq D_k$ oluyorsa H_0 hipotezi reddedilir (Lilliefors, 1967).

2.2.3. Cramer-von Mises testi

Cramer-von Mises (CVM) testi de KS testi yerine kullanılabilen ve KS testine benzer olarak deneysel dağılımı esas alan parametrik olmayan bir normallik testidir. CVM testi, ilk olarak Harald Cramér ve Richard von Mises tarafından geliştirilmiştir. Bu test, bir örneklem verisinin normal, uniform veya diğer olasılık dağılım modellerine uygunluğunu test etmek için kullanılabilir. CVM testi, KS testine göre daha güçlü sonuçlar verir.

CVM normallik testinin test istatistiğinin formülü aşağıda verilmiştir:

$$W^2 = \sum_{i=1}^n \left(F(x_{(i)}) - \frac{2i-1}{2n} \right)^2 + \frac{1}{12n} \quad (2.6)$$

Burada, n örneklem büyüklüğü, $x_{(i)}$ i . sıralı gözlem değeri, $F(.)$ ise örneklem birikimli dağılım fonksiyonudur.

Hesaplanan W^2 test istatistiğinin değeri α anlam düzeyine göre elde edilen kritik değerle karşılaştırılarak, veri setinin normal dağılıma sahip bir anakütleden çekilip çekilmediği kararı KS testinde olduğu gibi verilir. CVM, özellikle küçük ve orta örneklem büyüklükleri için uygun olan bir testtir. Büyük örneklem büyüklükleri için, bu testin hassasiyeti azalabilir (Kutsal, 2016).

2.2.4. Anderson-Darling testi

Anderson Darling (AD) testi, KS testinden uyarlanarak 1952 yılında Theodore Wilbur Anderson ve Donald Allan Darling tarafından geliştirilmiştir (Anderson ve Darling, 1952). AD testi için test istatistiğinin formülü aşağıdaki gibidir:

$$A^2 = -n - \frac{1}{n} \sum_{i=1}^n (2i-1) \left[\ln F(X_{(i)}) + \ln (1 - F(X_{(n+1-i)})) \right] \quad (2.7)$$

Küçük örnek büyüklükleri için A^2 formülüne aşağıda ifade edilen düzeltme uygulanır:

$$A_d^2 = A^2 \left(1 + \frac{0.75}{n} + \frac{2.25}{n^2} \right) \quad (2.8)$$

Burada, A_d^2 ifadesi düzeltilmiş A^2 'yi gösterir (Kutsal, 2016).

Hesaplanan test istatistiğinin değeri α anlam düzeyine göre elde edilen kritik değerle karşılaştırılarak, veri setinin normal dağılıma sahip bir anakütleden çekilip çekilmediği kararı daha önceki alt bölümlerde anlatılan normallik testlerinde olduğu gibi verilir. Parametrelerin tahmin edilebildiği durumlarda AD test istatistiğinin limit dağılımı tahmin edilmiş parametreler açısından değişiklik gösterir (Stephens, 1974).

2.2.5. D'Agostino-Pearson testi

D'Agostino- Pearson (DAP) testi 1971 yılında Ralph D'Agostino ve Albert Pearson tarafından geliştirilmiştir (D'Agostino, 1972). Test, normal dağılımın çarpıklık (skewness) ve basıklık (kurtosis) ölçütlerine dayanır. Normal dağılım için çarpıklığın 0, basıklığın ise 3'e eşit olması beklenir.

Örneklem için çarpıklık ve basıklık ölçütlerinin tahmin edicileri sırasıyla g_1 ve g_2 olarak aşağıdaki eşitliklerde verilmiştir.

$$g_1 = \frac{m_3}{m_2^{3/2}} = \frac{\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^3}{\left(\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \right)^{3/2}} \quad (2.9)$$

$$g_2 = \frac{m_4}{m_2^2} = \frac{\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^4}{\left(\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \right)^2} \quad (2.10)$$

Burada m_2 , m_3 ve m_4 sırasıyla ortalama etrafındaki ikinci, üçüncü ve dördüncü merkezsiz momentlerin tahmin edicileridir.

DAP testinin test istatistiği aşağıdaki formül yardımıyla hesaplanır:

$$K^2 = Z_1(g_1)^2 + Z_2(g_2)^2 \quad (2.11)$$

Burada $Z_1(.)$ ve $Z_2(.)$ fonksiyonları D'Agostino ve ark. (1972) tarafından önerilen dönüşümlerdir. Yapılan çalışmalara göre, K^2 test istatistiği serbestlik derecesi iki olan ki-kare dağılımına sahiptir.

Hesaplanan K^2 test istatistiğinin değeri α anlam düzeyine göre iki serbestlik dereceli ki-kare dağılımından elde edilen kritik değerle karşılaştırılarak, veri setinin normal dağılıma sahip bir anakütleden çekilip çekilmediği kararı daha önceki alt bölümlerde anlatılan normallik testlerinde olduğu gibi verilir.

2.2.6. Jarque-Bera testi

JB testi, 1980 yılında Carlos Jarque ve Anil K. Bera tarafından geliştirilmiştir ve DAP testine benzer olarak örneklem verilerinin çarpıklık ve basıklık ölçütlerine dayanır. JB testinde çarpıklık ve basıklık ölçütlerinin birleşik bir ölçüsü olan JB test istatistiği kullanılır. JB test istatistiği denklem 2.9 ve 2.10'da verilen g_1 ve g_2 tahmin edicileri kullanılarak aşağıdaki gibi hesaplanır.

$$JB = \frac{n}{6} \left(g_1^2 + \frac{1}{4} (g_2 - 3)^2 \right) \quad (2.12)$$

JB test istatistiği, H_0 hipotezi altında iki serbestlik dereceli ki-kare dağılımına asimptotik olarak yakınsar (Jarque ve Bera, 1987). Hesaplanan JB test istatistiğinin değeri α anlam düzeyine göre iki serbestlik dereceli ki-kare dağılımından elde edilen kritik değerle karşılaştırılarak, veri setinin normal dağılıma sahip bir anakütleden çekilip çekilmediği kararı daha önceki alt bölümlerde anlatılan normallik testlerinde olduğu gibi verilir.

2.2.7. Shapiro-Wilk testi

Shapiro- Wilk (SW) testi, 1965 yılında Samuel Shapiro ve Martin Wilk tarafından geliştirilmiş ve normallik testleri içinde en güçlü testlerden birisidir (Shapiro ve Wilk, 1965). Bu test, örnekleme ait olan sıra istatistiklerinin uygun bir lineer bileşeninin karesinin, gözlem değerlerinin ortalamadan olan farklarının kareler toplamına bölünmesiyle elde edilir. SW testi için test istatistiği (W) aşağıdaki gibidir:

$$W = \frac{\left(\sum_{i=1}^n a_i x_{(i)} \right)^2}{\sum_{i=1}^n (x_i - \bar{x})^2} \quad (2.13)$$

Burada,

$$a^T = (a_1, a_2, \dots, a_n) = \frac{m^T V^{-1}}{(m^T V^{-1} V^{-1} m)^{1/2}} \quad (2.14)$$

ve m , standart normal dağılıma ait n tane sıra istatistiğinin beklenen değerlerinin vektörünü, V^{-1} de sıra istatistiklerinin varyans-kovaryans matrisinin tersidir. W test istatistiğinin hesaplanmasında a_i katsayılarının elde edilmesinin zorluğundan dolayı a_i , $i = 2, \dots, 50$ katsayı değerleri tablo şeklinde verilmiştir (Yildirim, 2012).

Hesaplanan W test istatistiğinin değeri α anlam düzeyine göre kritik değerle karşılaştırılarak, veri setinin normal dağılıma sahip bir anakütleden çekilip çekilmediği kararı daha önceki alt bölümlerde anlatılan normallik testlerinde olduğu gibi verilir.



3. NORMAL DAĞILIM İÇİN MAKİNE ÖĞRENMESİ ALGORİTMASINA DAYALI GELİŞTİRİLEN UYUM İYİLİĞİ PROSEDÜRÜ

Makine öğrenmesi, bilgisayar sistemlerine veri analizi yapma yeteneği kazandırmak amacıyla kullanılan bir yapay zeka dalıdır. Makine öğrenmesi, bir bilgisayar programının belirli bir görevi belirli bir performans ölçütüne göre optimize etmesini sağlamak için istatistiksel ve matematiksel teknikler kullanır. Makine öğrenmesinin temel amacı, bir sistem üzerindeki deneyimlerden öğrenme ve bu deneyimlere dayanarak yeni verilere uyum sağlamaktır. Makine öğrenmesi algoritmaları genellikle örüntü tanıma, tahmin, sınıflandırma ve kümeleme gibi çeşitli problemlere çözüm üretmek amacıyla yaygın olarak kullanılır.

Makine öğrenmesi yaklaşımı, ayrıca büyük veri kütlelerini analiz etmek için önemlidir. Geleneksel yöntemlerle bu verilerin analizi çok zor olmanın yanı sıra oldukça zaman alabilir. Ancak makine öğrenmesi algoritmaları, büyük veri kütlelerindeki kalıpları tespit etmek ve verileri daha hızlı ve doğru bir şekilde analiz etmek için tasarlanmıştır (Alpaydın, 2014).

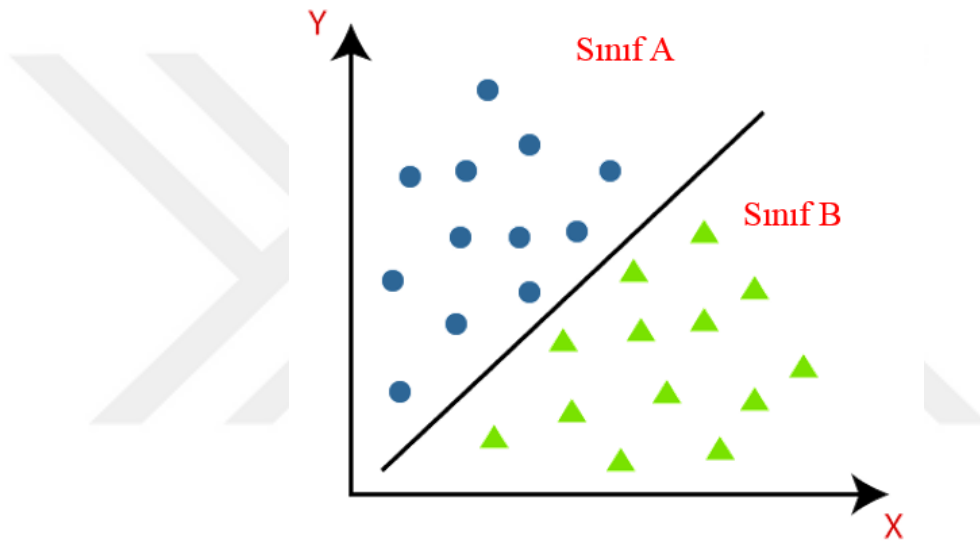
Tüm normallik testlerinin bir birilerine göre avantaj ve dezavantajları vardır. Bazı normallik testleri küçük örneklem büyüklüklerinde daha iyi sonuçlar verirken (SW, JB, AD), bazıları ise yanıltıcı sonuçlar verebilmektedir. Bunun dışında test edilen verinin simetrik uzun kuyruklu veya kısa kuyruklu şeklindeki yapısı da normallik testlerinin sonuçlarını etkilemektedir.

Normallik testlerinin tüm bu dezavantajları göz önünde bulundurulduğunda alternatif bir yönteme ihtiyaç duyulmaktadır. Bu yöntemin oluşturulabilmesi için herhangi bir örneklem büyüklüğünde ve herhangi bir yapıya sahip veri seti, “normal dağılmaktadır” veya “normal dağılmamaktadır” şeklinde bir sınıflandırma problemi olarak düşünülmüştür. Bu düşünce MLPN prosedürünün oluşturulmasında kilit öneme sahiptir.

MLPN prosedüründe farklı veri yapılarında ve farklı örneklem büyüklüklerinde veri setleri oluşturulduktan sonra bu veri setlerine ait ayırt edici özneliklere ilişkin değerler hesaplanıp bu tahmin değerleri ile de ilgili makine öğrenmesi algoritması eğitilmektedir. Bu yapı sayesinde makine öğrenmesi algoritması ona sunulan sınıflandırma problemini öğrendikten sonra yeni bir veri seti geldiğinde bu veri setinin normal dağılıp dağılmadığını tahmin edebilmektedir.

3.1. Makine Öğrenmesinde Sınıflandırma Problemi

Sınıflandırma, bir veri setinde bulunan gözlemleri belirli kategorilere veya sınıflara atama işlemidir. Temel amacı, veri setindeki desenleri anlamak, öğrenmek ve gelecekteki verileri bu desenlere göre sınıflandırmaktır. Makine öğrenmesinde sınıflandırma ise, bir veri kümesinin her bir örneğini belirli bir sınıfa ait olarak etkileme işlemidir. Sınıflandırma problemi, makine öğrenmesinin en önemli uygulama alanlarından biridir. Sınıflandırmanın temel mantığı **Şekil 3.1**'de gösterildiği gibidir.



Şekil 3.1. Sınıflandırma problemine ait örnek.

Makine öğrenmesi algoritmaları yardımıyla sınıflandırma probleminin çözümünde kullanılan aşamalar aşağıda özet şeklinde verilmiştir:

Veri toplama ve ön işleme. Veri toplama, sınıflandırma problemini çözerken dikkate alınması gereken en önemli aşamalardan biridir. Bu aşamada, problemin çözümünde kullanılacak olan veriler elde edilir. Bu veriler, yapay olarak oluşturulabilir veya gerçek dünyadan toplanabilir. Ön işleme aşamasında ise, verilerin temizlenmesi, standartlaştırılması ve eksik verilerin doldurulması gibi işlemler yapılır. Oluşturulan veri seti yeteri kadar büyük değilse, modelin performansı düşük olacaktır. Bir veri setinin büyüklüğü, problemin özelliklerine göre değişir. Genel olarak, veri seti ne kadar büyük olursa modelin performansı da o kadar iyi olur.

Öznitelik seçimi. Sınıflandırma problemini çözerken dikkate alınması gereken bir diğer önemli aşama öznitelik (değişken) seçimidir. Veri setinde bulunan öznitelikler, sınıflandırma probleminin çözümü için gerekli olan bilgileri makine öğrenmesi algoritmasına veya modele sağlamalıdır. Gereksiz öznitelikler, modelin performansını düşürebilir. Örneğin, bir resim kümesinde her bir resmin türüne ait olarak etiketlenmesi için veri kümesinde resmin boyutu, renkleri, şekilleri gibi öznitelikler bulunmalıdır. Ancak, resimlerin alındığı kameranın markası, çekim tarihi gibi öznitelikler gereksizdir.

Model (algoritma) seçimi. Model seçiminde, problemin özellikleri, mevcut veriler ve hedeflenen performans gibi faktörler dikkate alınır. Makine öğrenimi alanında, sınıflandırma için birçok farklı modeller mevcuttur. Bu modeller arasında, doğrusal regresyon, destek vektör makineleri, karar ağaçları ve rastgele ormanlar gibi modeller yaygın olarak kullanılmaktadır. Veri seti küçükse basit bir model kullanmak daha uygun olabilir. Aksine veri seti büyükse daha karmaşık bir algoritma veya model kullanmak gerekebilir.

Model eğitimi. Bu aşamada, seçilen model, mevcut veriler üzerinde eğitilir. Eğitim sonucunda, model, verilerin sınıflarına nasıl ayrılacağını öğrenir. Model eğitimi aşamasında, modelin performansını iyileştirmek için çeşitli teknikler kullanılabilir.

Modelin değerlendirilmesi. Modelin performansı, veri seti üzerinde yapılan testler ile değerlendirilir. Model değerlendirmede, çeşitli performans ölçütleri kullanılabilir. Bunlar doğruluk oranı, hassasiyet, özgünlük ve bu gibi benzeri ölçütlerdir (James, Witten, Hastie and Tibshirani, 2017).

MLPN prosedürünün oluşturulması bu aşamalar doğrultusunda yapılarak aşağıdaki bölümlerde sırasıyla sunulmuştur.

3.2. Normal Dağılım İçin Makine Öğrenmesi Algoritmasına Dayalı Geliştirilen Uyum İyiliği Prosedürünün Oluşturulmasında İzlenen Aşamalar

3.2.1. Normal dağılım için makine öğrenmesi algoritmasına dayalı geliştirilen uyum iyiliği prosedürü için veri setinin toplanması ve ön işleme

MLPN prosedürü için kullanılan veri seti “normal dağılıma sahip olan veri setleri” ve “normal dağılıma sahip olmayan veri setleri” şeklinde iki ayrı grup olarak simülasyon veya benzetim teknikleri kullanılarak farklı örneklem büyüklüklerinde üretilmiştir. “Normal dağılıma sahip olmayan veri setleri” kendi içinde simetrik ve asimetrik olarak

iki grup halinde ele alınmaktadır. Aynı zamanda simetrik dağılımlar da uzun kuyruklu simetrik ve kısa kuyruklu simetrik olarak kendi içinde de iki gruba ayrılmıştır. Veri setinin oluşturulmasına ilişkin aşama, modelin eğitilmesi aşamasında daha ayrıntılı şekilde anlatılmıştır.

3.2.2. Normal dağılım için makine öğrenmesi algoritmasına dayalı geliştirilen uyum iyiliği prosedürü için ayırt edici özniteliklerin seçimi

Belirli bir dağılıma sahip anakütleyi temsil eden tanımlayıcı ve aynı zamanda ayırt edici olan karakteristik değerlere anakütle parametresi veya kısaca parametre denir. Örneklemi temsil eden ilgili değerlere ise tahmin edici veya istatistik denir (Akdeniz, 2018). Bu tanımlayıcı istatistikler örneklem ortalaması, varyansı, standart sapması, çarpıklık katsayısı, basıklık katsayısı, modu, medyanı, kantilleri, minimum ve maksimum değerleri şeklinde sıralanabilir. Fakat belirli bir dağılıma sahip anakütleden çekilmiş örneklemin veya veri setinin dağılımını daha iyi bir şekilde tanımlamak için bu istatistikler yeterli olmayabilir. Örneğin, veri setinin çekildiği anakütlenin dağılımı eğer simetrikse bazı istatistiklerin değerleri birbirine oldukça yakın olduğu için anakütlenin dağılımını belirlemek çok zor ve bazen de imkansız olabilir. Bu nedenle daha farklı karakteristik değerler ve onların tahmin edicileri (istatistikleri) kullanılması gerekmektedir. Literatürde yapılan araştırmalar sonucunda belirli bir dağılıma sahip anakütleyi tanımlamak için L-moment, LH-moment, TL-moment (Shabri, 2002), Bowley çarpıklık katsayısı, Hinkley çarpıklık katsayısı, Moors basıklık katkatasayısı ve Crown basıklık katsayısı (Mala, Sladek ve Habarta, 2021) gibi farklı karakteristik değerlerin ve onların istatistiklerinin kullanılabileceği bulunmuştur.

Yukarıda adı geçen tüm bu anakütle parametreleri ve onların tahmin edicileri MLPN prosedüründe makine öğrenmesi modeli için öznitelik veya değişken olarak kullanılmıştır. Yapılan çeşitli denemeler ve performans testlerinden sonra makine öğrenmesi modelinin performansını olumlu yönde etkileyen parametrelerin tahmin edicileri öznitelik olarak seçilmiştir. MLPN prosedüründe kullanılan ayırt edici öznitelikler alt bölümlerde kısaca tanıtılmıştır.

3.2.2.1. L-momentler

Hosking (1990) tarafından öne sürülen L-moment parametresi sıra istatistiklerinin lineer bir birleşimi şeklinde elde edilmektedir. L-momentin en büyük avantajı veri setlerindeki aykırı değerlerden çok etkilenmemesidir.

$X_1, X_2, \dots, X_n$ olasılık yoğunluk fonksiyonu $f(.; \theta), \theta \in \Theta \subset R^n, n \geq 1$ olan dağılımdan rassal bir örneklem ve $X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(n)}$ bu örnekleme karşılık gelen sıra istatistikleri olmak üzere, anakütle dağılımına ait L-momentler aşağıdaki gibi tanımlanır:

$$\lambda_k = \frac{1}{k} \sum_{j=0}^{k-1} (-1)^j \binom{k-1}{j} E(X_{(k-j)}) \quad k = 1, 2, 3 \dots \quad (3.1)$$

ve örneklem L-momentleri,

$$\begin{aligned} \hat{\lambda}_k &= \frac{1}{k \binom{n}{k}} \sum_{i=1}^n \sum_{j=0}^{k-1} (-1)^j \binom{k-1}{j} \binom{i-1}{k-j-1} \binom{n-i}{j} X_{(i)} \quad k \\ &= 1, 2, 3 \dots \end{aligned} \quad (3.2)$$

şeklinde elde edilir (Hosking ve Wallis, 1997).

$\hat{\lambda}_k$ anakütle dağılımının k 'inci L-momentidir. Özel olarak λ_1 , $E(X)$ 'e eşit olur. L-momentlerin birbirine oranlamasıyla anakütle dağılımının L-katsayıları aşağıdaki şekilde tanımlanmaktadır (Hosking ve Wallis, 1997).

$$L - \text{Varyasyon katsayısı: } \tau = \lambda_r / \lambda_2 \quad (3.3)$$

$$L - \text{Çarpıklık katsayısı: } \tau_3 = \lambda_3 / \lambda_2 \quad (3.4)$$

$$L - \text{Basıklık katsayısı: } \tau_4 = \lambda_4 / \lambda_2 \quad (3.5)$$

$L - \text{Varyasyon katsayısı}$ olan τ 'nın, anakütle dağılımının darlığını veya genişliğini nicel şekilde gösterdiği, $L - \text{Çarpıklık katsayısı}$ τ_3 'ün de dağılımının çarpıklığını ve $L - \text{Basıklık katsayısı}$ τ_4 'ün ise dağılımın mod bölgesinin basıklığını nicel bir şekilde temsil edildiği Hosking ve Wallis (1997) tarafından açıklanmaktadır.

3.2.2.2. Kantiller

Kantiller, bir olasılık dağılımın belirli noktalarını temsil eden istatistiksel ölçümlerdir. Kantiller, yüzdesel dilimleri temsil eder ve dağılımın merkezi eğilimi, yayılımı ve şekli hakkında bilgi sağlar. Özellikle çeyrek kantiller (birinci, ikinci ve üçüncü çeyrek), ondabirlik kantiller (beşinci, on beşinci vb.) ve yüzdebirlik kantiller sıkça kullanılan kantillerdir.

$X_1, X_2, \dots, X_n$ birikimli dağılım fonksiyonu $F(., \theta), \theta \in \Theta \subset R^n, n \geq 1$ olan dağılımdan rassal bir örneklem ve $X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(n)}$ bu örnekleme karşılık gelen sıra istatistikleri olmak üzere, anakütle dağılımına ait kantil aşağıdaki gibi tanımlanır:

$$Q(p) = F^{-}(p) = \inf\{x: F(x) \geq p\}, \quad 0 < p < 1 \quad (3.6)$$

Örnekleme için kantil ise,

$$\hat{Q}(p) = \begin{cases} X_{(np)}, np \text{ tam sayı olduğunda} \\ X_{([np]+1)}, np \text{ tam sayı olmadığında} \end{cases} \quad (3.7)$$

şeklinde tanımlanır.

Örnekleme kantilleri bir veri setini belirli dilimlere ayırmak ve farklı bölgelerindeki değerleri analiz etmek için kullanılır. Örneğin, birinci çeyrek kantili (Q1), veri setinin en küçük %25'ini temsil ederken, üçüncü çeyrek kantili (Q3), veri setinin en büyük %25'ini temsil eder. İkinci çeyrek kantil (Q2) veya medyan, veri setini yarıya böler ve dağılımın ortasındaki değeri temsil eder.

Kantiller, kutu grafiklerinde, istatistiksel dağılımların tanımlanmasında, veri setinin yayılımının analizinde ve aykırı değerlerin tespitinde yaygın olarak kullanılır.

3.2.2.3. Entropi

Entropi, bir rassal deneyin veya bir veri setinin belirsizliğini sayısal olarak ortaya koyan bir ölçüdür. Bir veri setindeki belirsizlik seviyesi arttıkça, entropi değeri artmaktadır. Entropi değeri kazanılan bilgi miktarının bir ölçüsü olarak da düşünülebilir. Bunların yanı sıra, entropi bir olayın ortaya çıkması olasılığı ile ters orantılıdır. Diğer bir ifadeyle, bir olayın ortaya çıkma olasılığı ne kadar çok düşükse, olay gerçekleştiğinde belirsizlik o kadar azalır. Kesikli ve sürekli rassal değişkenlerin belirsizlik ölçüleri için Shannon (1948) tarafından tanımlanan Shannon entropi ölçütü literatürde en çok kullanılan ölçüttür.

Olasılık veya olasılık yoğunluk fonksiyonu $f(.; \theta), \theta \in \Theta \subset R^n, n \geq 1$ olan X rassal değişkeni için Shannon entropi ölçütü aşağıdaki gibi tanımlanır:

$$H(X) = E[-\log(f(X))] = \begin{cases} -\sum_{x \in R} \log(f(x; \theta)) f(x; \theta), & X \text{ kesikli} \\ -\int_{-\infty}^{+\infty} \log f(x; \theta) f(x; \theta) dx, & X \text{ sürekli} \end{cases} \quad (3.8)$$

Bu tez çalışmasında ele alınan dağılımların entropi değerleri aşağıdaki tabloda sunulmuştur.

Tablo 3.1. Dağılımlara ait entropi değerleri.

Dağılım	$H(X)$
$Normal(\mu, \sigma)$	$\ln(\sigma\sqrt{2\pi e})$
$Beta(\alpha, \beta)$	$\ln(B(\alpha, \beta)) - (\alpha - 1)\psi(\alpha) - (\beta - 1)\psi(\beta) + (\alpha + \beta - 2)\psi(\alpha + \beta)$
$Lojistik(\mu, s)$	$\ln(s) + 2$
$Student t(v)$	$\frac{v+1}{2} \left[\psi\left(\frac{1+v}{2}\right) - \psi\left(\frac{v}{2}\right) \right] + \log \left[\sqrt{v} B\left(\frac{v}{2}, \frac{1}{2}\right) \right]$
$Uniform(a, b)$	$\ln(b - a)$
$Laplace(\mu, b)$	$\log_2(2eb)$
$Gama(a, b)$	$a + \ln\Gamma(a) + (1 - a) \cdot \psi(a) + \ln b$
$Lognormal(\mu, \sigma)$	$\log_2 \left(\sigma e^{\mu + \frac{1}{2}} \sqrt{2\pi} \right)$
$Ki - Kare(k)$	$\frac{k}{2} + \ln \left(2\Gamma\left(\frac{k}{2}\right) \right) + \left(1 - \frac{k}{2}\right) \psi\left(\frac{k}{2}\right)$
$\text{Üstel}(\theta)$	$1 - \ln\theta$

3.2.2.4. Pearson çarpıklık ve basıklık katsayıları

Çarpıklık ve basıklık katsayıları, normallik varsayımının kontrol edilmesinde önemli ölçütlerinden biridir. Bu katsayılar, dağılımların şeklini ve yapısını tanımlayan önemli ölçütlerdir. Çarpıklık bir dağılımın simetrisini ölçer, pozitif çarpıklık verinin sağa doğru çarpık olduğunu, negatif çarpıklık ise verinin sola doğru çarpık olduğu gösterir. En sık kullanılan çarpıklık katsayısı Pearson (Pearson, 1900) tarafından tanımlanmış olan çarpıklık katsayısıdır. Bu katsayı üçüncü standartlaştırılmış ortalama etrafındaki moment olarak da bilinir (Mala, Sladek ve Habarta, 2021).

Olasılık veya olasılık yoğunluk fonksiyonu $f(.; \theta), \theta \in \Theta \subset R^n, n \geq 1$ olan X rassal değişkeni veya anakütlesi için Pearson çarpıklık katsayısı aşağıdaki gibi tanımlanır:

$$\gamma_1 = E \left[\left(\frac{X - \mu}{\sigma} \right)^3 \right] = \frac{\mu_3}{\sigma^3} = \frac{E[(X - \mu)^3]}{(E[(X - \mu)^2])^{3/2}} \quad (3.9)$$

Örneklem için Pearson çarpıklık katsayısı ise aşağıdaki gibi elde edilir:

$$\hat{\gamma}_1 = \frac{m_3}{m_2^{3/2}} = \frac{\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^3}{\left(\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \right)^{3/2}} \quad (3.10)$$

Burada, $\bar{X}$ örneklem ortalaması, m_2 ve m_3 ise sırasıyla örneklemin ikinci ve üçüncü merkezi momentidir.

Basıklık ise bir dağılımın kuyruklarında yoğunlaşma şeklinin bir ölçüsüdür. Verinin dördüncü standartlaştırılmış sıralı momenti basıklık katsayısı olarak bilinir: Olasılık veya olasılık yoğunluk fonksiyonu $f(\cdot; \theta), \theta \in \Theta \subset R^n, n \geq 1$ olan X rassal değişkeni veya anakütlesi için Pearson basıklık katsayısı aşağıdaki gibi tanımlanır:

$$\gamma_2 = E \left[\left(\frac{X - \mu}{\sigma} \right)^4 \right] = \frac{\mu_4}{\sigma^4} = \frac{E[(X - \mu)^4]}{(E[(X - \mu)^2])^2} \quad (3.11)$$

Örneklem için hesaplanan dördüncü sıralı moment ise aşağıdaki gibidir:

$$\hat{\gamma}_2 = \frac{m_4}{m_2^2} = \frac{\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^4}{\left(\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \right)^2} \quad (3.12)$$

Burada, m_4 örneklemin dördüncü merkezi momentidir.

Yukarıda anlatılan tanımlayıcı özneliliklerin yanı sıra kullanılacak makine öğrenmesi algoritmasının veri setinin dağılımını daha iyi ayırt edebilmesi için örneklem büyüklüğü (n), minimum ($\min(X_1, X_2, \dots, X_n)$) ve maksimum ($\max(X_1, X_2, \dots, X_n)$) değerleri de tanımlayıcı öznelilik olarak MLPN prosedüründe yer almıştır.

3.2.3. Normal dağılım için makine öğrenmesi algoritmasına dayalı geliştirilen uyum iyiliği prosedürü için model seçimi

Literatürde sınıflandırma problemi için birçok algoritma mevcuttur (Alpaydın, 2014). Bu algoritmalar arasında veri ayırtediciliyi, performansı, dayanıklılığı ve işlem hızı gibi özelliklerinden dolayı Destek Vektör Makineleri (SVM), Sınıflandırma Ağaçları

(CART), Rasgele Ormanlar (RF), Kategorik Boosting (CatBoost) ve Yapay Sinir Ağları (ANN) algoritmaları ön plana çıkmaktadır.

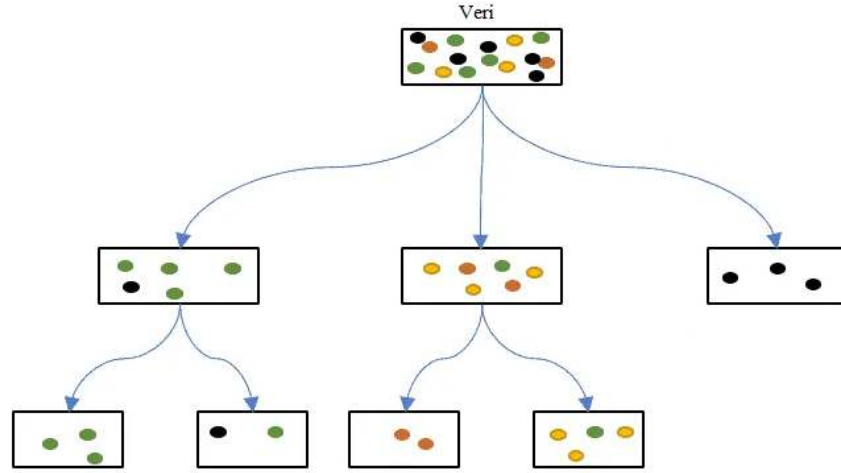
Ön plana çıkan bu algoritmalar MLPN prosedürünün oluşturulabilmesi için ayrı ayrı denenmiş, işlem hızı, dayanıklılığı ve daha az veriyle daha iyi sonuçlar verebildiği için CatBoost algoritmasının kullanılmasına karar verilmiştir. Ayrıca CatBoost algoritmasının diğer ön plana çıkan algoritmalarından çeşitli kriterlere göre daha iyi performans sergilediğini gösterebilmek için RF algoritmasıyla ilgili karşılaştırma sonuçları da tezin ilerleyen bölümlerinde verilmiştir.

3.2.3.1. Kategorik boosting algoritması

Boosting (güçlendirme) 1990'lı yıllarda Yoav Freund ve Robert Schapire tarafından geliştirilmiştir. Literatürde pek çok Boosting algoritması bulunmaktadır ve bunlardan en geniş kullanılanları Adaboost ve Gradient Boosting algoritmalarıdır. Gradient Boosting algoritmasının bir diğer adı ise Categorizing Boosting (CatBoost)'dur ve bu sınıflandırma problemleri için en yaygın uygulanan algoritmalarından biridir (Ceylan, 2018).

CatBoost, Yandex şirketinin geliştirmiş olduğu Gradyan Artırma tabanına sahip açık kaynak kodlu bir makine öğrenmesi algoritması olarak da bilinmektedir. Nisan 2017 tarihinde tanıtılan bu algoritma, Gradyan Artırma algoritmasının çalışma performansını artırmanın yanı sıra XGBoost ile LightGBM algoritmalarına alternatif olarak geliştirilmiştir. CatBoost algoritmasını diğer algoritmalarından ayıran en önemli özellikler; yüksek öğrenme hızı sayesinde hem sayısal hem kategorik hem de metin verileriyle çalışabilmesi, GPU desteği ve pratik görselleştirme seçenekleri sunmasıdır (Tokmak, 2023).

CatBoost'un çalışma mantığı, öncelikle zayıf bir öğrenici oluşturmakla başlar. Daha sonra bu öğrenicinin performansı ölçülür ve hatalar belirlenir. Belirlenen hatalar ağırlıklar ile çarpılır ve bir sonraki öğrenici için veri örneklerine uygulanır. Bu ağırlıklar, daha önce yanlış sınıflandırılan verilerin doğruluğunu artırmak için ayarlanır. Bu adımın bitiminde, bir sonraki zayıf öğrenici oluşturulur ve süreç tekrarlanır. Böylece, her adımda bir önceki öğrenicinin yetersiz kaldığı durumlar düzeltilir ve sonunda güçlü bir öğrenici yaratılır. CatBoost algoritmasının çalışma mantığına ilişkin örnek şema **Şekil 3.2'**de sunulmuştur.



Şekil 3.2. Catboost örnek şeması.

CatBoost algoritması, veri hazırlığı sürecini daha etkin bir şekilde yönettiği için diğer modellere kıyasla daha avantajlı bir konumdadır. Aynı zamanda bu algoritma, eksik verilerle çalışabilir ve kategorik verilere özgün bir kodlama yöntemi uygulayabilir. Bu sayede, kategorik verilerle yüksek performanslı bir şekilde çalışabilir. Dolayısıyla, veri hazırlığı aşamasında ayrı bir kodlama işlemine ihtiyaç duyulmaz ve hatta kodlamanın yapılmaması önerilir. Ayrıca, CatBoost algoritması, simetrik ağaçlar kullanarak yüksek tahmin doğruluğu elde ederken aşırı öğrenme (overfitting) sorununu da engeller. Eğer aşırı öğrenme oluşursa, algoritma hedefe ulaşmadan önce öğrenme prosedürünü durdurur (James, Witten, Hastie ve Tibshirani, 2017).

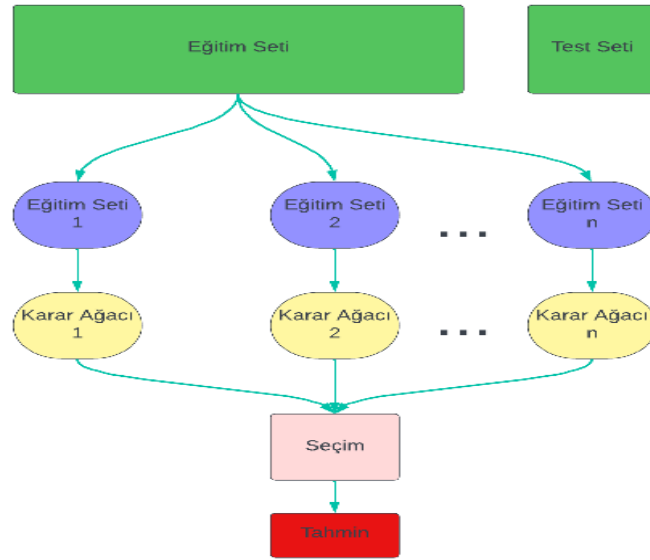
3.2.3.2. Rasgele ormanlar algoritması

Rasgele Ormanlar (RF) algoritması, Leo Breiman tarafından 2001 yılında geliştirilmiştir. Bu algoritma, regresyon ve sınıflandırma problemlerinde yaygın olarak kullanılan bir makine öğrenmesi algoritmasıdır. Ormanı yaratan nesne ağaçlardır ve ağaçların daha çok olması ormanın da sağlam olduğunu ifade eder. RF algoritması da buna benzer bir şekilde veriler üzerinde karar ağaçlarını oluşturur. Oluşturulmuş her bir karar ağacından tahmin değeri alınır ve en sonunda da oylama yaparak optimum çözüm bulunur (Momayez, 2023).

RF algoritması, yüzlerce değişkeni silmeye gerek duymadan onların özelliklerini kullanarak sınıflandırma yapabilir ve değişkenlerin bu sınıflandırmadaki önemini ortaya

çıkartır. RF algoritmasının tek bir karar ağacına kıyasla daha güçlü, dayanaklı olduğu ve tahminleme performansının da diğer makine öğrenmesi algoritmalarına göre daha güçlü olduğu söylenebilir. RF algoritması hem kategorik hem de sürekli değişkenlerle çalışabileceği gibi sürekli ve kategorik değişkenlerin aynı olduğu veri setleri ile çalışabilmektedir. Aynı zamanda sınıflandırma işlemi gerçekleştirdiği zaman veri setinde bulunan kayıp verilere karşı dayanıklı bir yapısı vardır. Çok sayıda değişkenin olduğu veri setleri için de kullanışlı bir algoritmadır. Bu algoritma birden fazla karar ağacından oluştuğu için görselleştirilmesi kolay olmayan bir algoritma olmasıyla birlikte her bir karar ağacını oluştururken oldukça çok bellek kullandığından dolayı düşük belleğe sahip bilgisayarlarda çalıştırılması zor bir algoritmadır (Demirtürk, 2023).

RF algoritmasının çalışma mantığı görsel olarak **Şekil 3.3**'de sunulmuştur.



Şekil 3.3. RF örnek şeması.

3.2.4. Normal dağılım için makine öğrenmesi algoritmasına dayalı geliştirilen uyum iyiliği prosedüründeki modelin eğitilmesi

MLPN prosedürü son zamanlarda geniş kullanım alanına sahip olması ve karmaşık yapıların daha rahat bir şekilde kodlanabilmesinden dolayı Python programlama dili kullanılarak oluşturulmuştur. MLPN prosedürünün oluşturulmasında Python programlama dilinde özel program yazılmış ve ayrıca bu programda Python programlama dilinde tanımlı gerekli kütüphaneler kullanılmıştır.

MLPN prosedüründeki algoritmanın eğitimi için aşağıdaki adımlar izlenerek eğitim veri seti oluşturulmuştur :

Adım 0. n , α , ve tekrar sayısını (N) belirle.

Adım 1. $N(0,1)$ dağılımına sahip anakütleden belirlenen α anlam düzeyi ve tekrar sayısında, örneklem büyüklüğü n olan veriler üret.

Adım 2. Üretilen verileri standartlaştır.

Adım 3. Standartlaştırılmış verilerin minimum ve maksimum değerlerini, çarpıklık ve basıklık katsayılarını, farklı kantil değerlerini, entropi değerini, 1. ve 2. L-moment değerlerini, L-çarpıklık ve L-basıklık katsayılarını tahmin et ve bu tahmin değerlerinden öznitelikler için bir liste hazırla.

Adım 4. Sınıflandırma için bu listeyi "1" kodu ile etiketle.

Adım 5. Öznitelikleri bağımsız değişkenler, etiketlenen değeri ise bağımlı değişken olarak kabul et.

Adım 6. Adım 0'da n örneklem büyüklüğünü sırasıyla $n = 10, 20, 30, 40, 50, 100, 500$ örneklem büyüklükleri için Adım 1-5'i tekrarla.

Adım 7. Adım 0-6'yı normal dağılım yerine sırasıyla beta, lojistik, student t, uniform, laplace, gama, lognormal, ki-kare ve üstel dağılımları için tekrarla. Bu verileri Adım 4'te "0" kodu ile etiketle.

Adım 8. Elde edilen bu verileri eğitim seti olarak kabul et.

Yukarıda ifade edilen adımlar kullanılarak $\alpha = 0,01$, $\alpha = 0.05$ ve $\alpha = 0.10$ anlam düzeyleri için üç ayrı eğitim seti oluşturulmuş ve her bir eğitim seti kullanılarak CatBoost ve RF algoritmaları sınıflandırılma için eğitilmiştir. Ayrıca oluşturulan test seti ile de ilgili modellerin performansının değerlendirilmesi yapılmıştır.

3.2.5. Modelin değerlendirilmesi

MLPN prosedürünün değerlendirilmesi üç aşamada yapılmıştır. Değerlendirmenin ilk aşamasında çalışmada ele alınan normallik testleri ile MLPN prosedürünün performanslarının karşılaştırılması eğitim setinde yer alan yapıdaki veriler üzerinden yapılmıştır. İkinci aşamada eğitim setinde yer almayan farklı yapıdaki veriler için MLPN prosedürünün performansı normallik testlerine göre incelenmiştir. Son olarak üçüncü aşamada ise MLPN prosedürünün geliştirilmesinde kullanılan CatBoost algoritması ile RF algoritmasının hız ve verimlilik açısından karşılaştırması yapılmıştır.

MLPN prodedürünün normallik testlerine göre performansının değerlendirilmesi ile ilgili ilk iki aşamada kullanılan I. Tip hata ve güç değerleri kriterleri aşağıdaki gibi elde edilmiştir.

$$\begin{aligned}
 I. \text{Tip Hata} &= P(H_0 \text{ Red} \mid \text{Gerçekte } H_0 \text{ Doğru}) \\
 &= \frac{1}{TS} \sum_{i=1}^{TS} \mathbf{1}_{(H_0 \text{ Red} \mid H_0 \text{ Doğru})}
 \end{aligned}
 \tag{3.13}$$

$$\text{Testin Gücü} = P(H_0 \text{ Red} \mid \text{Gerçekte } H_0 \text{ Yanlış})$$

$$= \frac{1}{TS} \sum_{i=1}^{TS} \mathbf{1}_{(H_0 \text{ Red} \mid H_0 \text{ Yanlış})}$$

Burada, TS toplam yapılan test sayısını ve $\mathbf{1}_{(A)}$ ise 1 veya 0 değerlerini alan indikatör fonksiyonu ifade etmektedir.

MLPN prosedüründe kullanılan makine öğrenmesi modelinin sınıflandırma başarısını göstermek için I. Tip hata ve güç değerler kriterlerine ek olarak karışıklık matrisi (confusion matris), doğruluk (accuary), kesinlik (precision), duyarlılık (recall) ve F1-skorları da kullanılmıştır.

Karışıklık matrisi, gerçek etiketlerle tahmin edilen etiketler arasındaki uyumu gösterir. İki sınıflı bir sınıflandırma problemi düşünüldüğünde, karışıklık matrisi aşağıdaki gibidir.

		Gerçek	
		Pozitif	Negatif
Tahmin	Pozitif	GP	YP
	Negatif	YN	GN

Şekil 3.4. İkili sınıflandırmada karışıklık matrisi.

Burada;

- Gerçek Pozitif (GP): Kullanılan modelin sayı bakımından kaç adet pozitif sınıfı doğru olarak tahmin ettiğini (True Positives),
- Gerçek Negatif (GN): Kullanılan modelin sayı bakımından kaç adet negatif sınıfı doğru olarak tahmin ettiğini (True Negatives)
- Yanlış Pozitif (YP): Kullanılan modelin sayı bakımından kaç adet negatif sınıfı yanlış olarak tahmin ettiğini (False Positives),
- Yanlış Negatif (YN): Kullanılan modelin sayı bakımından kaç adet pozitif sınıfı yanlış olarak tahmin ettiğini (False Negatives) göstermektedir.

Doğruluk, kesinlik, duyarlılık ve F1-Puanın hesaplanması ise sırasıyla aşağıdaki formüllerin yardımıyla yapılmıştır:

Doğruluk

$$= \frac{\text{Gerçek Pozitif} + \text{Gerçek Negatif}}{\text{Gerçek Pozitif} + \text{Gerçek Negatif} + \text{Yanlış Pozitif} + \text{Yanlış Negatif}} \quad (3.14)$$

$$\text{Kesinlik} = \frac{\text{Gerçek Pozitif}}{\text{Gerçek Pozitif} + \text{Yanlış Pozitif}} \quad (3.15)$$

$$\text{Duyarlılık} = \frac{\text{Gerçek Pozitif}}{\text{Gerçek Pozitif} + \text{Yanlış Negatif}} \quad (3.16)$$

$$F1 = 2 * \frac{1}{\frac{1}{\text{kesinlik}} + \frac{1}{\text{duyarlılık}}} = \frac{2 * \text{kesinlik} \times \text{duyarlılık}}{\text{kesinlik} + \text{duyarlılık}} \quad (3.17)$$

Modelin değerlendirilmesinin üçüncü aşamasında kullanılan algoritmaların hız ve verimlilik açısından karşılaştırılabilmesi için birçok farklı ölçüt bulunmaktadır. Bu çalışmada kullanılan ölçütler aşağıda verilmiştir.

- **Eğitim zamanı** – her iki algoritma aynı veri seti üzerinde eğitilerek eğitim süreleri ölçülebilir. Algoritmaların eğitim süreleri karşılaştırılarak, daha kısa eğitim süresine sahip olan algoritma, hızlı çalışan algoritma olarak değerlendirilir.

- **Model boyutu** – bu ölçütle eğitilmiş modellerin boyutu karşılaştırılabilir. Algoritmalar için eğitim sonucunda oluşan model dosyalarının boyutları kontrol edilir, daha küçük model boyutuna sahip olan algoritma, daha verimli bir algoritma olarak değerlendirilir.
- **Tahmin süresi** – eğitilmiş modeller kullanılarak tahmin süreleri ölçülebilir. Algoritmalar için aynı veri seti üzerinde tahmin yapılarak tahmin süreleri karşılaştırılır ve daha kısa tahmin süresine sahip olan algoritma, daha hızlı çalışan bir algoritma olarak kabul edilir (James, Witten, Hastie ve Tibshirani, 2017).

Tüm bu performans ölçütlerine ilişkin sonuçlar izleyen bölümde sunulmuştur.



4. SİMÜLASYON ÇALIŞMASI

Çalışmanın bu kısmında, ikinci bölümde hakkında kısaca bilgi verilen normallik testlerinin ve üçüncü bölümde tanıtılan MLPN prosedürünün I. Tip hata oranlarını ve güçlerini hesaplamak ve karşılaştırmalarını yapmak için aşağıda ayrıntıları ile verilen kapsamlı bir simülasyon çalışması yapılmıştır. Simülasyon çalışmasındaki hesaplamaların yapılabilmesi için Python programlama dili kullanılarak bir program yazılmıştır.

Simülasyon çalışmasında KS, LL, CVM, AD, DAP, JB ve SW normallik testleri ve MLPN prosedürünün I. Tip hata oranı aşağıdaki adımlar izlenerek hesaplanmıştır:

Adım 1: $n = 10$ örneklem büyüklüğü için $N(0,1)$ dağılımına sahip anakütleden veri seti üret.

Adım 2: Üretilen veri setini standartlaştır ve MLPN prosedürü için gerekli özniteliklerini hesapla.

Adım 3: Çalışmada adı geçen normallik testlerini kullanarak standartlaştırılmış veri seti için $\alpha = 0,01$; $\alpha = 0,05$ ve $\alpha = 0,10$ anlam düzeyinde normal dağılıp dağılmadığına karar ver.

Adım 4: Öznitelik değerlerine göre her bir α anlam düzeyi için oluşturulan MLPN prosedürünü kullanarak tahmin yap (0 veya 1).

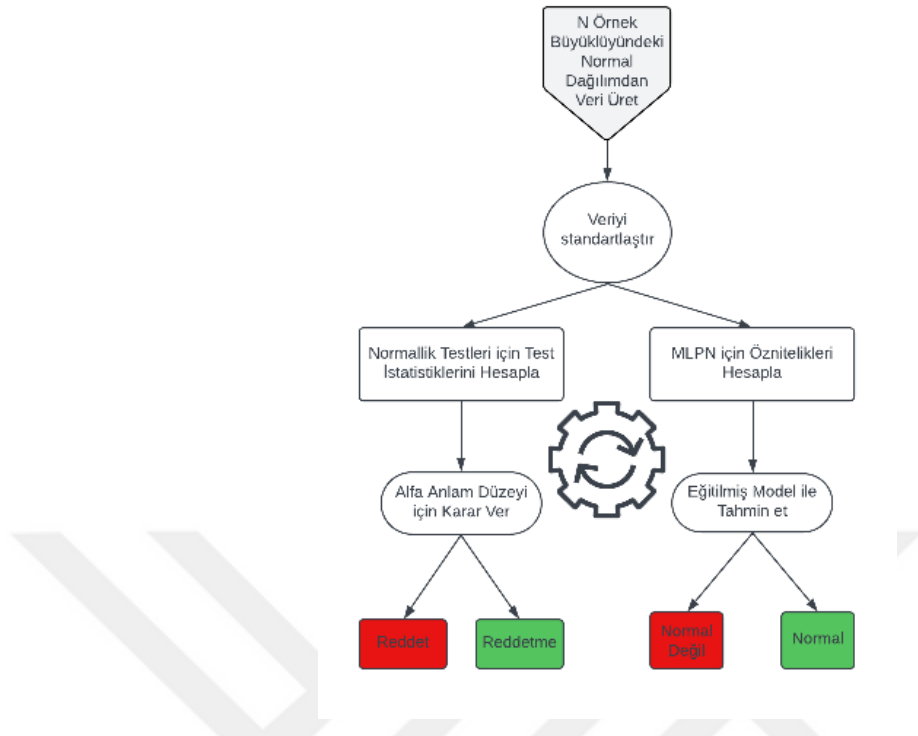
Adım 5: Adım 1-4'ü $TS = 1000$ defa tekrar et. Her bir α anlam düzeyi için yanlış karar ve yanlış tahmin sayılarını belirle.

Adım 6: Her bir α anlam düzeyine göre yanlış karar sayılarını kullanarak normallik testleri için, yanlış tahmin sayılarını kullanarak da MLPN prosedürü için I. Tip hata oranını Eşitlik 3.13'ün yardımıyla hesapla.

Adım 7: Sistemik hatayı minimuma indirmek için Adım 1-6'ı 100 defa tekrarla ve her bir α anlam düzeyi için ortalama I. Tip hata oranını hesapla.

Adım 8: Adım 1'deki n 'i sırasıyla $n = 10, 20, 30, 40, 50, 100, 500$ örneklem büyüklüklerini kullan ve Adım 2-7'yi tekrar et.

I. Tip hatanın oranının elde edilmesine yönelik yukarıda verilen adımlara ilişkin akış şeması **Şekil 4.1**'de görselleştirilmiştir.



Şekil 4.1. MLPN prosedürü ve normallik testleri için I. Tip hatanın elde edilmesine yönelik akış şeması.

Simülasyon çalışmasında KS, LL, CVM, AD, DAP, JB ve SW normallik testleri ve MLPN prosedürünün gücü ise aşağıdaki adımlar izlenerek hesaplanmıştır:

Adım 1: n örneklem büyüklüğü için $Beta(2, 2)$ dağılımına sahip anakütleden veri seti üret.

Adım 2: Üretilen veri setini standartlaştır ve MLPN prosedürü için gerekli özniteliklerini hesapla.

Adım 3: Çalışmada adı geçen normallik testlerini kullanarak standartlaştırılmış veri seti için $\alpha = 0,01$; $\alpha = 0,05$ ve $\alpha = 0,10$ anlam düzeyinde normal dağılıp dağılmadığına karar ver.

Adım 4: Öznitelik değerlerine göre her bir α anlam düzeyi için MLPN prosedürünü kullanarak tahmin yap (0 veya 1).

Adım 5: Adım 1-4'ü $TS = 1000$ defa tekrar et. Her bir α anlam düzeyi için doğru karar ve doğru tahmin sayılarını belirle.

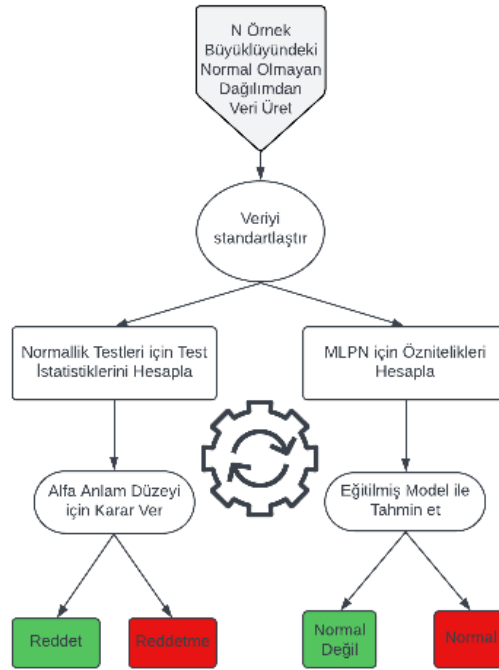
Adım 6: Her bir α anlam düzeyine göre doğru karar sayılarını kullanarak normallik testleri için, doğru tahmin sayılarını kullanarak da MLPN prosedürü için Eşitlik 3.13 yardımıyla gücü hesapla.

Adım 7: Sistematik hatayı minimuma indirmek için Adım 1-6'ı 100 defa tekrarla ve ortalama gücü hesapla.

Adım 8: Adım 1'deki n 'i sırasıyla $n = 10, 20, 30, 40, 50, 100, 500$ örneklem büyüklüklerini kullan ve Adım 2-7'yi tekrar et.

Adım 9: Adım 1-8'i $Beta(2, 2)$ dağılım yerine sırasıyla $Lojistik(0, 1)$, $Student t(5)$, $Uniform(0, 1)$, $Laplace(0, 1)$, $Beta(2, 5)$, $Gama(4, 5)$, $Lognormal(0, 1)$, $Ki - Kare(5)$ ve $Üstel(1)$ dağılımları için tekrar et.

Güç değerlerinin elde edilmesine yönelik algoritma şeması **Şekil 4.2'**deki gibidir:



Şekil 4.2. MLPN prosedürü ve normallik testleri için güç değerlerinin elde edilmesine yönelik akış şeması

Yukarıda ifade edilen adımlar izlenerek MLPN prosedürü ve çalışmada adı geçen normallik testlerine ait deneysel I. Tip hata oranları ve güç değerleri, $\alpha = 0,01$, $\alpha = 0,05$ ve $\alpha = 0,10$ anlamlılık düzeyleri için elde edilmiş ve elde edilen sonuçlar tablo ve görsel şeklinde özetlenmiştir.

4.1. I. Tip Hata Değerlerine İlişkin Sonuçlar

Normal dağılımlı örneklerden veri üretildiğinde elde edilen I. Tip hata değerleri aşağıdaki tabloda verilmiştir:

Tablo 4.1. MLPN prosedürü ve Normallik testlerinin I. Tip hata değerleri.

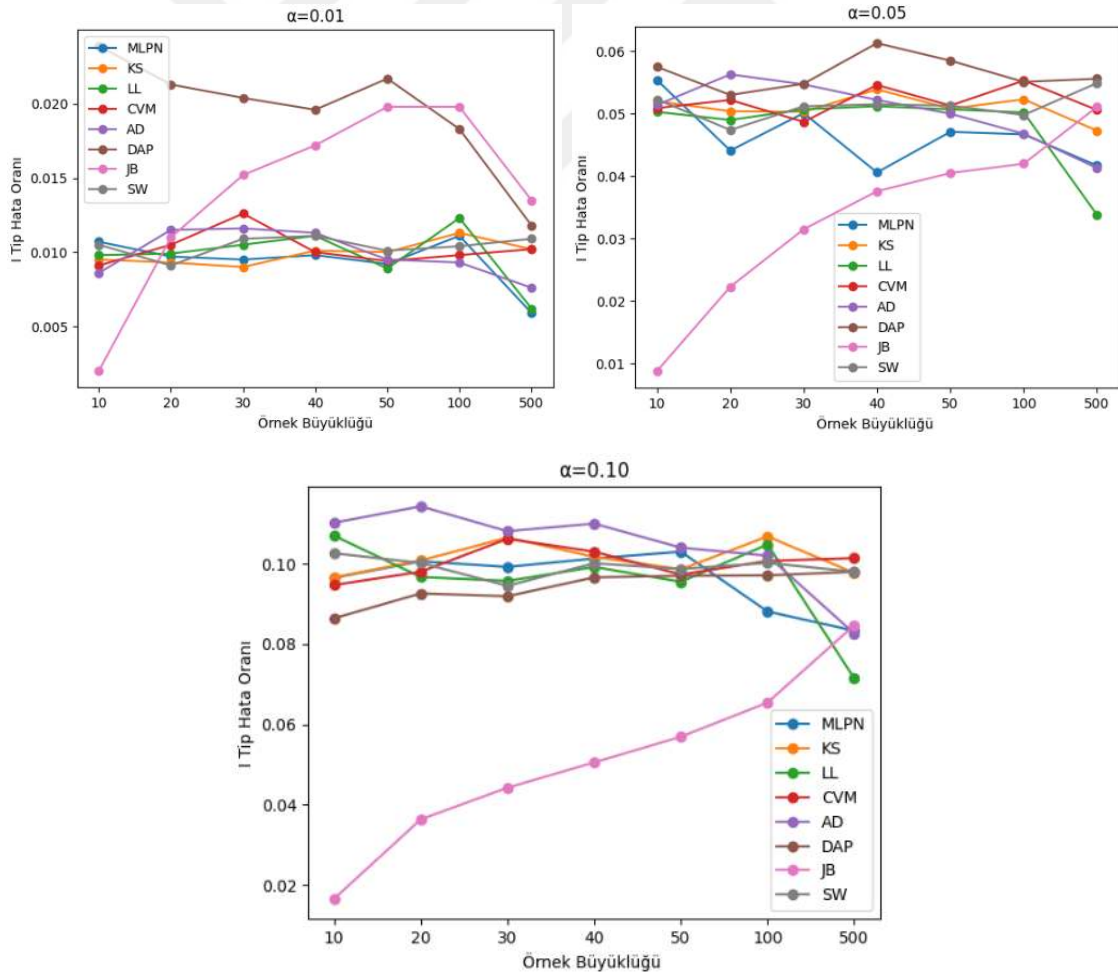
Anlam Düzeyi		Örneklem Büyüklüğü n						
		10	20	30	40	50	100	500
$\alpha = 0,01$	MLPN	0,0107	0,0097	0,0095	0,0098	0,0092	0,0111	0,0059
	KS	0,0095	0,0093	0,009	0,0101	0,01	0,0113	0,0102
	LL	0,0098	0,0099	0,0105	0,0111	0,0089	0,0123	0,0062
	CVM	0,0091	0,0105	0,0126	0,01	0,0094	0,0098	0,0102
	AD	0,0086	0,0115	0,0116	0,0113	0,0095	0,0093	0,0076
	DAP	0,0239	0,0213	0,0204	0,0196	0,0217	0,0183	0,0118
	JB	0,002	0,011	0,0152	0,0172	0,0198	0,0195	0,0135
	SW	0,0105	0,0091	0,0109	0,0111	0,0101	0,0104	0,0109
$\alpha = 0,05$	MLPN	0,0554	0,0441	0,0501	0,0406	0,0471	0,0467	0,0417
	KS	0,052	0,0504	0,0503	0,0539	0,0508	0,0523	0,0473
	LL	0,0503	0,049	0,0507	0,0512	0,0507	0,0502	0,0338
	CVM	0,0509	0,0522	0,0487	0,0546	0,0513	0,0553	0,0506
	AD	0,0514	0,0563	0,0547	0,0522	0,05	0,0468	0,0413
	DAP	0,0575	0,053	0,0548	0,0613	0,0585	0,0551	0,0556
	JB	0,0088	0,0223	0,0315	0,0376	0,0405	0,042	0,0511
	SW	0,0523	0,0474	0,0512	0,0515	0,0513	0,0498	0,0549
$\alpha = 0,1$	MLPN	0,0965	0,1006	0,0992	0,1013	0,103	0,0881	0,0833
	KS	0,0966	0,1008	0,1065	0,1017	0,0986	0,1068	0,0977
	LL	0,107	0,0967	0,0957	0,0992	0,0955	0,1048	0,0714
	CVM	0,0947	0,098	0,1062	0,103	0,0973	0,1007	0,1014
	AD	0,1102	0,1143	0,1081	0,11	0,104	0,102	0,0826
	DAP	0,0864	0,0926	0,0919	0,0966	0,097	0,0971	0,098
	JB	0,0165	0,0363	0,0442	0,0505	0,0569	0,0654	0,0847
	SW	0,1026	0,1002	0,0945	0,1001	0,0987	0,1002	0,098

Tablo 4.1'de yedi normallik testi ve MLPN prosedürüne ait farklı örneklem büyüklükleri için I. Tip hata değerleri verilmiştir. Tablodan görüldüğü gibi $\alpha = 0,01$ anlam düzeyi için örneklem büyüklüğü $n = 10$ olduğunda JB testi en düşük hatayı vermiş, diğer normallik testleri ve MLPN prosedürü α etrafında değerler almıştır.

Örneklem büyüklüğü $n = 500$ olduğunda ise en düşük I. Tip hatayı MLPN prosedürü göstermiştir.

$\alpha = 0,05$ anlam düzeyi için örneklem büyüklüğünün $n = 10, 20, 30, 40, 50, 100$ olduğu durumlarda en düşük I. Tip hatayı JB testi göstermiş, diğer normallik testleri ve MLPN prosedürü ise α etrafında değerler almıştır. Örneklem büyüklüğü $n = 500$ olduğunda ise en düşük I. Tip hatayı LL testi göstermiştir.

$\alpha = 0,10$ anlam düzeyi için örneklem büyüklüğünün $n = 10, 20, 30, 40, 50, 100$ olduğu durumlarda en düşük I. Tip hatayı JB testi, $n = 500$ olduğunda ise en düşük I. Tip hatayı LL testi göstermiştir. Örneklem büyüklüğü $n = 10, 20, 30, 40$ olduğunda AD testi α anlam düzeyinin üzerinde olsa da diğer testler ve MLPN prosedürü α anlam düzeyi etrafında değerler almıştır. Elde edilen sonuçlar Şekil 4.3'de görselleştirilerek özetlenmiştir.



Şekil 4.3. MLPN prosedürü ve Normallik testlerinin I. Tip hata değerlerinin grafik gösterimi.

Şekil 4.3’den de görüldüğü gibi her bir α anlam düzeyi için MLPN prosedürünün ve normallik testlerinin I. Tip hata değerlerinin değişimi daha anlaşılır bir şekilde görülebilmektedir.

MLPN prosedüründe kullanılan makine öğrenmesi modelinin sınıflandırma başarısını göstermek için I. Tip hata değerlerine ait karışıklık matrisi, doğruluk, kesinlik, duyarlılık ve F1-skorları Tablo 4.2’de olduğu gibidir.

Tablo 4.2. MLPN prosedürünün I. Tip hata oranı için sınıflandırma başarısını gösteren kriterler.

I.Tip Hata		Gerçekte <i>Normal</i> (0, 1) Dağılımı													
		10		20		30		40		50		100		500	
		<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>
MLPN Tahmini	<i>N</i>	945	0	956	0	950	0	960	0	953	0	954	0	959	0
	<i>ND</i>	55	0	44	0	50	0	40	0	47	0	46	0	41	0
	Doğruluk	0,945		0,956		0,95		0,96		0,953		0,954		0,959	
	Kesinlik	1		1		1		1		1		1		1	
	Duyarlılık	0,945		0,956		0,95		0,96		0,953		0,954		0,959	
	F1 Skoru	0,971		0,977		0,974		0,979		0,975		0,976		0,979	

Burada, *N* “normal dağılım”, *ND* ise “normal dağılım değil” ifadelerinin kısaltmalarını göstermektedir.

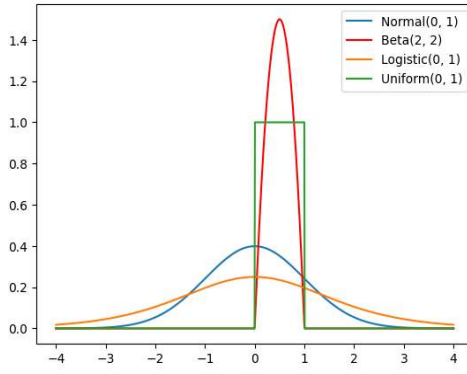
4.2. Güç Değerlerine İlişkin Sonuçlar

MLPN prosedürü ve normallik testlerinin güç değerlerini karşılaştırmak için normal dağılmayan alternatif dağılımlardan (beta, lojistik, student t, uniform, laplace, gama, lognormal, ki-kare ve üstel) farklı örnek büyüklüklerinde veri setleri üretilmiştir. Bundan başka üretilen alternatif dağılımlar üç grup halinde ele alınmaktadır:

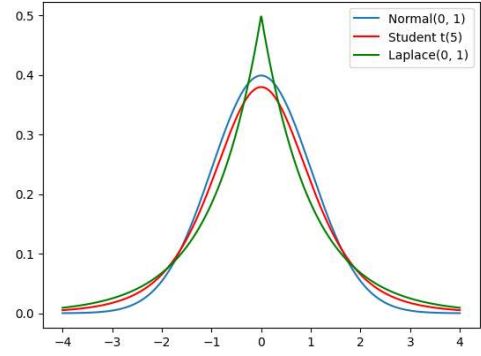
Grup 1 - Kısa kuyruklu simetrik dağılımlar: *Lojistik*(0, 1), *Beta*(2, 2) ve *Uniform*(0, 1).

Grup 2 - Uzun kuyruklu simetrik dağılımlar: *Student t*(5) ve *Laplace*(0, 1).

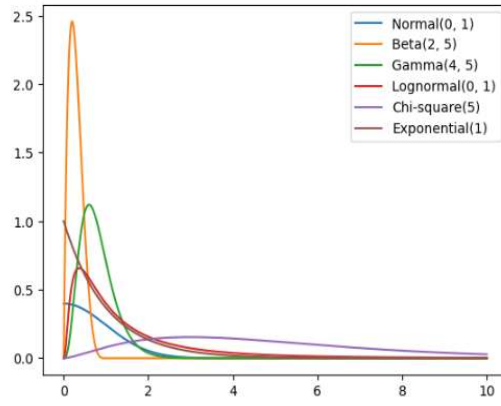
Grup 3 – Asimetrik dağılımlar: *Beta*(2, 5), *Gama*(4, 5), *Lognormal*(0, 1), *Ki – Kare*(5), *Üstel*(1).



(a)



(b)



(c)

Şekil 4.4. Grup 1 (a), Grup 2 (b), Grup 3'e (c) ait dağılım grafikleri.

Grup 1'de yer alan simetrik ve kısa kuyruklu dağılımlar için MLPN prosedürünün ve normallik testlerinin güç değerlerine ilişkin sonuçlar **Tablo 4.3'**de verilmiştir.

Tablo 4.3. MLPN prosedürü ve normallik testlerinin Beta(2,2), Lojistik(0,1) ve Uniform(1) dağılımları altında güç değerleri.

Örneklem Büyüklüğü n									
			10	20	30	40	50	100	500
$Beta(2, 2)$	$\alpha = 0,01$	MLPN	0,2766	0,4827	0,5077	0,5326	0,5786	0,6618	0,9909
		KS	0	0	0	0	0	0	0,0017
		LL	0,0081	0,0092	0,0114	0,0129	0,0162	0,0317	0,4276
		CVM	0	0	0	0	0	0	0,001
		AD	0,0061	0,0116	0,0161	0,0237	0,0287	0,0919	0,9657
		DAP	0,0053	0,0085	0,0213	0,0523	0,0981	0,4083	1

Tablo 4.3. (Devam) MLPN prosedürü ve normallik testlerinin Beta(2, 2), Lojistik(0, 1) ve Uniform(1) dağılımları altında güç değerleri.

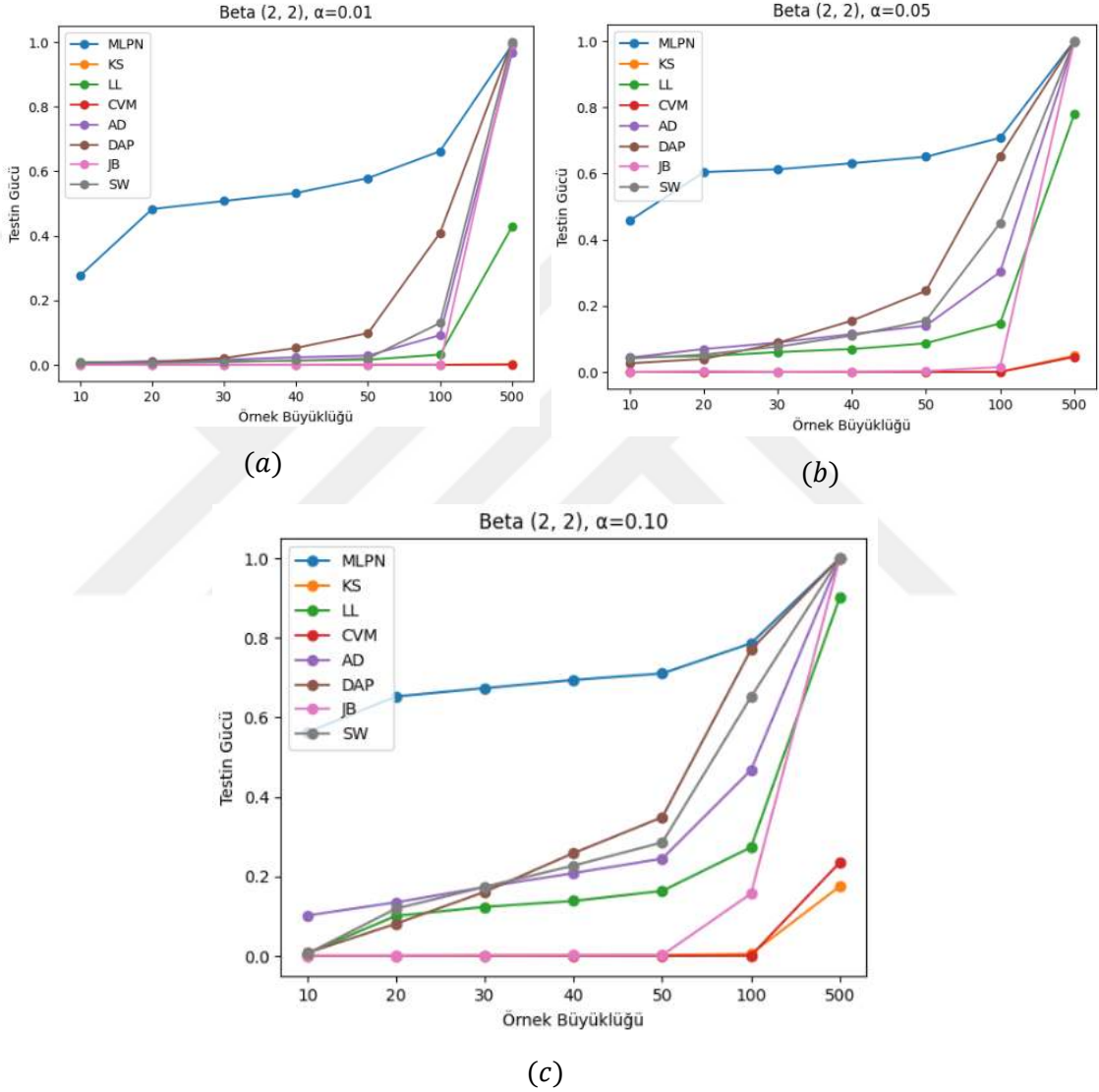
		JB	0	0	0	0	0,001	0,0009	0,9938
		SW	0,0048	0,0061	0,0085	0,0139	0,02	0,1294	0,9999
	$\alpha = 0,05$	MLPN	0,4576	0,6041	0,6125	0,631	0,6507	0,7074	1
		KS	0	0,0001	0,0001	0,0001	0,0003	0,0007	0,0498
		LL	0,0446	0,0478	0,0602	0,0695	0,0869	0,1472	0,7787
		CVM	0	0	0	0	0	0	0,0455
		AD	0,0436	0,0694	0,0892	0,1147	0,1402	0,3022	0,9974
		DAP	0,0263	0,0392	0,0879	0,1553	0,2455	0,6528	1
		JB	0	0,002	0,0003	0,0011	0,0027	0,0152	1
		SW	0,0408	0,0529	0,0766	0,11	0,1559	0,4501	1
	$\alpha = 0,10$	MLPN	0,5621	0,6526	0,6734	0,6939	0,7103	0,7861	1
		KS	0	0,0003	0,0011	0,0017	0,0021	0,0051	0,1749
		LL	0,0073	0,1008	0,1233	0,1384	0,1635	0,2733	0,9011
		CVM	0	0	0,0001	0,0001	0,0002	0,0006	0,235
		AD	0,1019	0,1348	0,1734	0,2083	0,2447	0,4675	0,9993
		DAP	0,0079	0,0811	0,1613	0,2586	0,3488	0,7703	1
		JB	0,0007	0,0008	0,0013	0,0016	0,0022	0,1566	1
		SW	0,0058	0,1191	0,1739	0,2272	0,2854	0,6511	1
			10	20	30	40	50	100	500
Lojistik(0, 1)	$\alpha = 0,01$	MLPN	0,0277	0,0545	0,0834	0,1238	0,1602	0,3126	0,8663
		KS	0	0	0	0	0	0,0001	0,0009
		LL	0,0178	0,0243	0,0256	0,0287	0,0327	0,0473	0,197
		CVM	0	0	0	0	0	0	0,0001
		AD	0,0228	0,0406	0,0465	0,0514	0,0605	0,0933	0,4948
		DAP	0,0556	0,0834	0,1027	0,1168	0,1401	0,2113	0,6979
		JB	0,0084	0,0616	0,0941	0,1226	0,1559	0,2734	0,8029
		SW	0,0267	0,0476	0,0607	0,0745	0,0941	0,1665	0,6894
	$\alpha = 0,05$	MLPN	0,1049	0,1454	0,2034	0,2484	0,2924	0,4802	0,9428
		KS	0	0,0002	0,0005	0,0009	0,0012	0,0013	0,0268
		LL	0,0791	0,0848	0,0963	0,1	0,1153	0,1493	0,449
		CVM	0	0	0	0,0001	0,0003	0,0009	0,0159
		AD	0,0834	0,1146	0,139	0,1409	0,1561	0,2292	0,7286
		DAP	0,1083	0,1491	0,183	0,2037	0,2388	0,3482	0,8525
		JB	0,027	0,0938	0,143	0,178	0,2249	0,3735	0,8901
		SW	0,0815	0,1109	0,1449	0,1652	0,1941	0,3077	0,8404

Tablo 4.3. (Devam) MLPN prosedürü ve normallik testlerinin $Beta(2, 2)$, $Lojistik(0, 1)$ ve $Uniform(1)$ dağılımları altında güç değerleri.

	$\alpha = 0,10$	MLPN	0,1488	0,205	0,2734	0,3434	0,3845	0,5583	0,9619
		KS	0,0014	0,0025	0,0039	0,0042	0,0059	0,0095	0,0795
		LL	0,136	0,1492	0,1669	0,1773	0,1901	0,2515	0,598
		CVM	0,0001	0,0002	0,0005	0,0009	0,0021	0,0039	0,0742
		AD	0,1515	0,1958	0,2151	0,2324	0,2585	0,3322	0,8216
		DAP	0,1583	0,2085	0,2457	0,282	0,3089	0,4287	0,9008
		JB	0,0402	0,1195	0,1802	0,2314	0,2707	0,429	0,9177
		SW	0,1421	0,18	0,2141	0,2421	0,2741	0,3924	0,8877
			10	20	30	40	50	100	500
Uniform(0, 1)	$\alpha = 0,01$	MLPN	0,2154	0,4134	0,4671	0,5204	0,6183	0,9489	1
		KS	0	0	0	0	0	0,0002	0,4509
		LL	0,011	0,0205	0,0348	0,0526	0,0738	0,2532	0,9994
		CVM	0	0	0	0	0	0	0,6555
		AD	0,0116	0,0485	0,0993	0,1755	0,2631	0,768	1
		DAP	0,0058	0,043	0,177	0,3889	0,5902	0,9826	1
		JB	0	0	0	0,0001	0,0004	0,0024	1
		SW	0,0115	0,0298	0,0878	0,2028	0,3563	0,9435	1
	$\alpha = 0,05$	MLPN	0,3484	0,4878	0,5063	0,6963	0,7808	0,9573	1
		KS	0	0,0002	0,0005	0,0011	0,0016	0,0136	0,9329
		LL	0,0649	0,1049	0,1495	0,1992	0,2639	0,5873	1
		CVM	0	0	0	0	0	0,0041	0,9926
		AD	0,08	0,2038	0,3189	0,4556	0,5868	0,9444	1
		DAP	0,0249	0,1537	0,3878	0,6344	0,8087	0,9974	1
		JB	0,0004	0,0005	0,0007	0,0011	0,0018	0,5695	1
		SW	0,0825	0,2035	0,38	0,5852	0,7544	0,9969	1
	$\alpha = 0,10$	MLPN	0,4608	0,6422	0,6725	0,8062	0,8995	0,9816	1
		KS	0,0006	0,0027	0,0038	0,0053	0,0113	0,0559	0,992
		LL	0,1264	0,1911	0,268	0,3377	0,4197	0,7494	1
		CVM	0	0,0003	0,0005	0,0008	0,0037	0,0507	0,9996
		AD	0,1688	0,3357	0,4924	0,6254	0,7405	0,978	1
		DAP	0,0708	0,2586	0,5467	0,7488	0,8824	0,9989	1
		JB	0,0033	0,0011	0,0012	0,0019	0,0333	0,9051	1
		SW	0,1755	0,3658	0,5871	0,7636	0,8784	0,9996	1

Simetrik kısa kuyruklu $Beta(2, 2)$ dağılımı altında MLPN prosedürü ve normallik testlerinin güç değerlerine bakıldığında tüm örneklem büyüklükleri ve α anlam düzeyleri

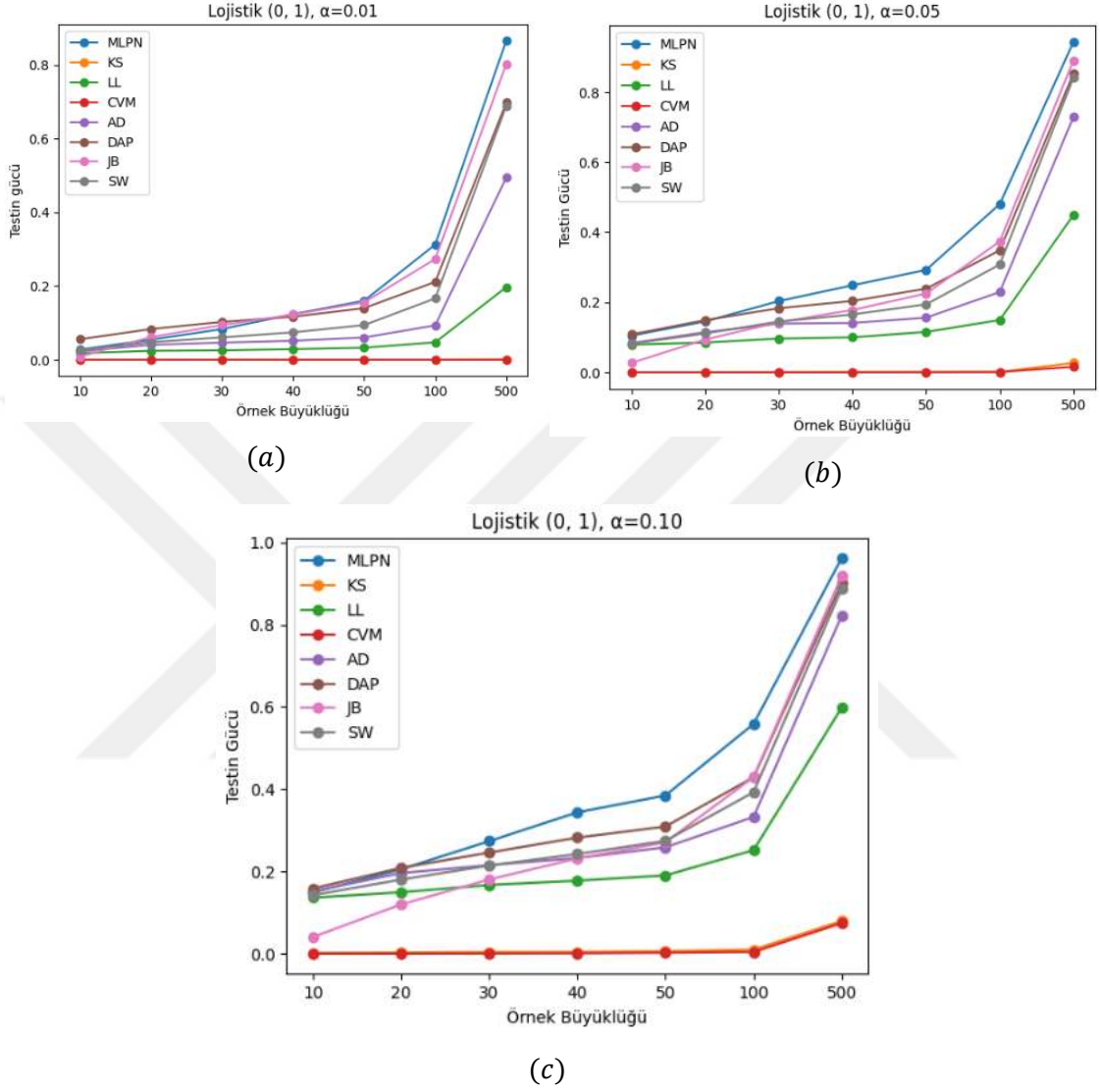
için MLPN prosedürünün gücünün diğer normallik testlerinden daha yüksek olduğu görülmektedir. Genel olarak bu dağılım için ele alınan normallik testleri iyi performans göstermemiş, sadece örneklem büyüklüğü $n = 500$ olduğu durumda AD, DAP, JB ve SW testlerinin gücü 1 değerine ulaşmıştır. Elde edilen sonuçlar Şekil 4.5’de görselleştirilerek özetlenmiştir.



Şekil 4.5. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için Beta(2, 2) dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.

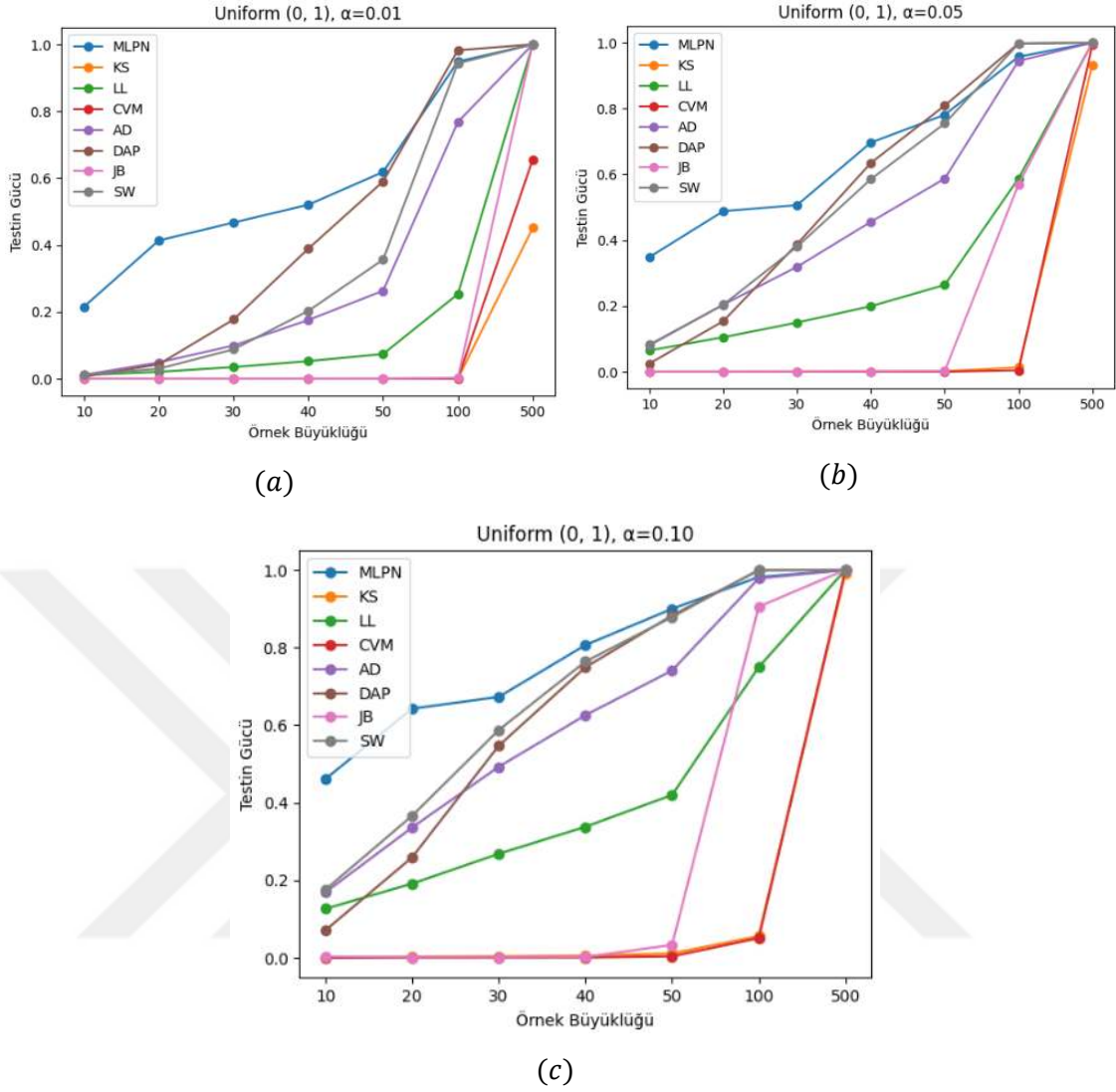
Simetrik kısa kuyruklu $Lojistik(0, 1)$ dağılımı altında $n = 10, 20$ örneklem büyüklükleri için en iyi sonucu DAP testi göstermiş, örneklem büyüklüğünün $n \geq 30$ olduğu durumda ise tüm α anlam düzeyleri için MLPN prosedürünün güç değerleri diğer ele alınan normallik testlerinden daha yüksek olmuştur. Bununla birlikte örneklem

büyüküğünün $n = 500$ olduğunda bile hiçbir test maksimum güç seviyesi olan 1'e ulaşamamıştır. Elde edilen sonuçlar **Şekil 4.6'**da görselleştirilmiştir.



Şekil 4.6. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için Lojistik(0, 1) dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.

Grup 1'in son dağılımı olan $Uniform(0, 1)$ dağılımına ait sonuçlara bakıldığında tüm örneklem büyüklüklerinde ve α anlam düzeyinde MLPN prosedürünün diğer normallik testlerinden çok daha iyi sonuçlar verdiği görülmektedir. Bu dağılım altında ele alınan normallik testleri içinde MLPN prosedüründen sonra testin gücü bakımından ikinci en iyi sonucu SW testi göstermiştir. Elde edilen sonuçlar görsel olarak **Şekil 4.7'**de sunulmuştur.



Şekil 4.7. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için Uniform(0,1) dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.

Grup 2’de yer alan simetrik ve uzun kuyruklu dağılımlar için MLPN prosedürü ve ele alınan normallik testlerinin güç değerlerine ilişkin sonuçlar **Tablo 4.4** ve **Şekil 4.8**’de verilmiştir.

Tablo 4.4. MLPN prosedürü ve Normallik testlerinin Student $t(5)$ ve Laplace(0,1) dağılımları altında güç değerleri.

		Örneklem Büyüklüğü n							
			10	20	30	40	50	100	500
Student $t(5)$	$\alpha = 0,01$	MLPN	0,0466	0,1064	0,1723	0,2644	0,3318	0,5865	0,9928
		KS	0	0,0005	0,0009	0,0024	0,0023	0,005	0,0845
		LL	0,0309	0,0475	0,0621	0,0744	0,0907	0,1574	0,6843

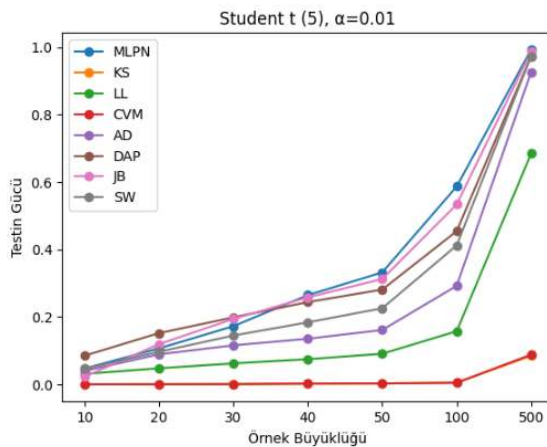
Tablo 4.4. (Devam) MLPN prosedürü ve Normallik testlerinin Student $t(5)$ ve Laplace(0,1) dağılımları altında güç değerleri.

		CVM	0	0,0002	0,0006	0,0019	0,0026	0,0047	0,0879
		AD	0,0396	0,0889	0,1156	0,1351	0,1611	0,292	0,926
		DAP	0,0848	0,1518	0,1985	0,2438	0,2809	0,4541	0,9729
		JB	0,0223	0,1191	0,1938	0,2577	0,3124	0,5333	0,9876
		SW	0,0459	0,0965	0,1446	0,1838	0,2254	0,4119	0,973
	$\alpha = 0,05$	MLPN	0,1293	0,2225	0,3358	0,4029	0,4756	0,724	0,9983
		KS	0,0005	0,0034	0,0072	0,0098	0,0132	0,0247	0,296
		LL	0,0937	0,128	0,1567	0,1845	0,2097	0,3223	0,8731
		CVM	0	0,0019	0,0039	0,008	0,0107	0,022	0,3421
		AD	0,11	0,1862	0,2325	0,2637	0,3096	0,4688	0,9773
		DAP	0,1482	0,2312	0,2981	0,3531	0,4035	0,6048	0,9912
		JB	0,0471	0,1668	0,2568	0,326	0,3949	0,6328	0,9946
		SW	0,1101	0,1868	0,2491	0,2977	0,3545	0,5679	0,9913
	$\alpha = 0,10$	MLPN	0,182	0,2937	0,3798	0,4808	0,5658	0,7808	0,9992
		KS	0,0029	0,009	0,0173	0,0223	0,0283	0,0559	0,4795
		LL	0,1559	0,2041	0,2467	0,2776	0,3103	0,4529	0,9272
		CVM	0,0011	0,0047	0,0125	0,0158	0,0217	0,0474	0,5539
		AD	0,178	0,2706	0,3231	0,3661	0,4048	0,5804	0,9871
		DAP	0,1918	0,287	0,3642	0,4276	0,4822	0,6806	0,9955
		JB	0,0635	0,1946	0,3005	0,3817	0,4515	0,6857	0,997
		SW	0,1703	0,2547	0,3222	0,3805	0,4387	0,6486	0,9951
			10	20	30	40	50	100	500
	Laplace(0, 1)	MLPN	0,0675	0,1854	0,2685	0,3554	0,4597	0,8267	1
		KS	0	0,0002	0,0004	0,0011	0,0017	0,0103	0,6963
		LL	0,0486	0,0902	0,1351	0,1753	0,2281	0,4512	0,9974
		CVM	0	0	0	0,0001	0,0004	0,0035	0,7746
		AD	0,0599	0,1429	0,2167	0,2735	0,344	0,6405	1
		DAP	0,1152	0,1915	0,2563	0,2943	0,3574	0,5608	0,998
		JB	0,0279	0,1516	0,2514	0,3239	0,4151	0,6787	0,9995
		SW	0,0679	0,1344	0,2089	0,2707	0,345	0,6349	0,9999
	$\alpha = 0,05$	MLPN	0,1638	0,3321	0,4858	0,5914	0,6795	0,9155	1
		KS	0,0014	0,0047	0,01	0,0165	0,0243	0,0928	0,9517
		LL	0,1388	0,2145	0,3	0,3746	0,4391	0,698	0,9998
		CVM	0,0001	0,0014	0,0035	0,007	0,0118	0,072	0,9795
		AD	0,1621	0,2976	0,3866	0,4776	0,5494	0,8135	1

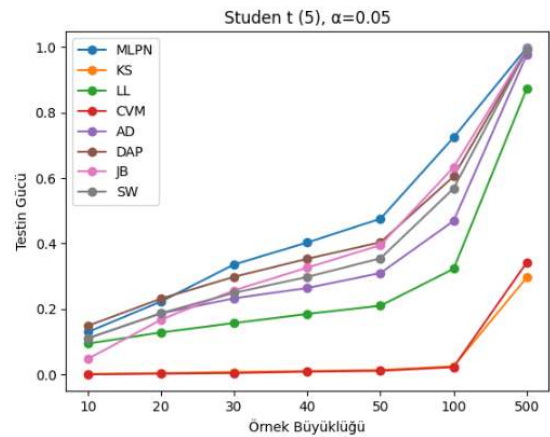
Tablo 4.4. (Devam) MLPN prosedürü ve Normallik testlerinin Student $t(5)$ ve Laplace(0,1) dağılımları altında güç değerleri.

$\alpha = 0,10$	DAP	0,1899	0,2997	0,3782	0,4494	0,5124	0,7293	0,9998
	JB	0,0619	0,2158	0,3308	0,4269	0,513	0,7754	0,9999
	SW	0,1516	0,262	0,3533	0,4435	0,5253	0,7912	1
	MLPN	0,2276	0,4432	0,589	0,6943	0,7694	0,9474	1
	KS	0,005	0,0157	0,0334	0,0522	0,0727	0,2097	0,9877
	LL	0,2199	0,3205	0,4072	0,4851	0,5619	0,8063	1
	CVM	0,0011	0,0069	0,0153	0,0322	0,0499	0,1927	0,9964
	AD	0,2459	0,3917	0,496	0,5829	0,6545	0,8823	1
	DAP	0,2538	0,3734	0,4605	0,5336	0,6035	0,812	1
	JB	0,0905	0,2623	0,3863	0,4871	0,5745	0,8247	1
	SW	0,2256	0,347	0,4466	0,5379	0,6174	0,8598	1

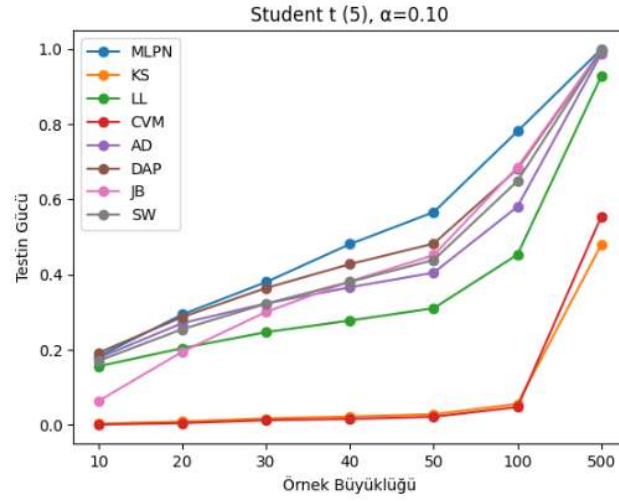
Simetrik uzun kuyruklu 5 serbestlik dereceli *Student t*(5) dağılımı altında örneklem büyüklüğü $n = 10, 20$ olduğu zaman tüm α anlam düzeylerinde ele alınan normallik testleri arasında DAP testi diğer testlerden daha yüksek performans göstermiştir. Örneklem büyüklüğünün $n \geq 30$ olduğu durumlarda ise MLPN prosedürünün performansı ele alınan diğer normallik testlerinden daha iyidir. Örneklem büyüklüğü $n = 500$ olduğunda sadece MLPN, DAP, JB ve SW maksimum güç düzeyi olan 1 değerine yakın değerler almaktadır. Elde edilen sonuçlar Şekil 4.8’de görselleştirilerek özetlenmiştir.



(a)



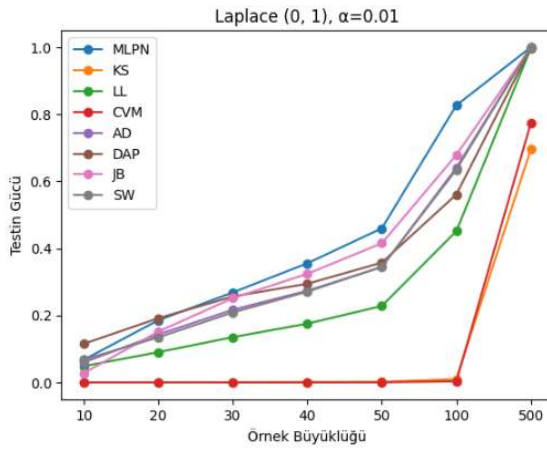
(b)



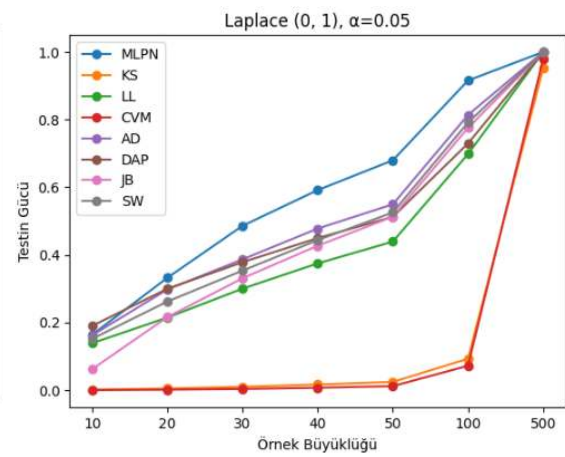
(c)

Şekil 4.8. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için Student $t(5)$ dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.

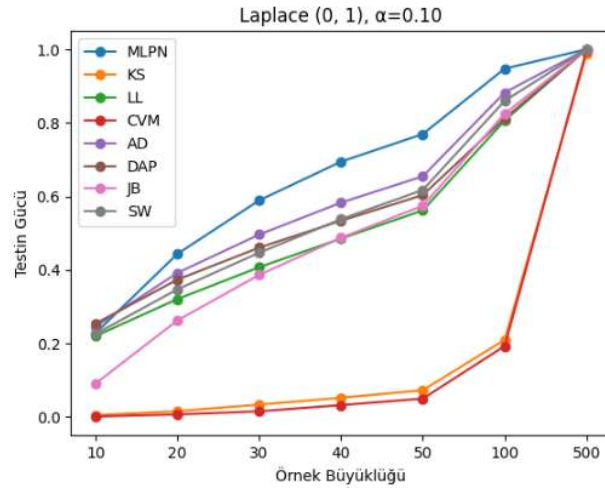
Grup 2'nin ikinci dağılımı olan $Laplace(0, 1)$ dağılımına ait sonuçlarda örneklem büyüklüğü $n = 10$ olduğun zaman DAP testinin gücü diğer testlerden daha yüksek çıkmasına rağmen örneklem büyüklüğünün $n \geq 20$ olduğu durumlarda MLPN prosedürünün güç değerleri diğer normallik testlerinden daha yüksektir. Elde edilen sonuçların görsel sunumu **Şekil 4.9'**da verilmiştir.



(a)



(b)



(c)

Şekil 4.9. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için Laplace(0, 1) dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.

Son olarak Grup 3'deki asimetrik dağılımlara ait MLPN prosedürü ve ele alınan normallik testlerinin güç değerleri **Tablo 4.5** ve **Şekil 4.10**'da verilmiştir.

Tablo 4.5. MLPN prosedürü ve Normallik testlerinin farklı α anlam düzeyleri için Beta(2, 5), Gama(4, 5), Lognormal(0, 1), Ki – Kare(5) ve Üstel(1) dağılımları altında güç değerleri.

Örneklem Büyüklüğü n									
			10	20	30	40	50	100	500
Beta(2, 5)	$\alpha = 0,01$	MLPN	0,2628	0,4252	0,441	0,5296	0,565	0,7648	1
		KS	0	0	0	0	0	0,0007	0,3337
		LL	0,015	0,0284	0,0505	0,0686	0,0907	0,2387	0,982
		CVM	0	0	0	0	0	0	0,24
		AD	0,0186	0,0473	0,0832	0,1232	0,1658	0,4697	1
		DAP	0,0334	0,0403	0,0593	0,0755	0,0873	0,2304	1
		JB	0,0046	0,0211	0,0366	0,0531	0,0641	0,2072	1
		SW	0,021	0,04	0,0843	0,1422	0,2026	0,6598	1
	$\alpha = 0,05$	MLPN	0,4469	0,5735	0,5896	0,6332	0,6846	0,8004	1
		KS	0	0,0005	0,0013	0,0031	0,0043	0,0241	0,7745
		LL	0,0769	0,1197	0,1604	0,2159	0,257	0,5011	0,999
		CVM	0	0	0	0	0,0004	0,0058	0,8329
		AD	0,0897	0,1749	0,2442	0,3254	0,3952	0,7439	1
		DAP	0,0793	0,1144	0,1585	0,2058	0,2466	0,6147	1
		JB	0,0153	0,0502	0,0862	0,121	0,157	0,5141	1

Tablo 4.5. (Devam) *MLPN prosedürü ve Normallik testlerinin farklı α anlam düzeyleri için Beta(2,5), Gama(4,5), Lognormal(0,1), Ki-Kare(5) ve Üstel(1) dağılımları altında güç değerleri.*

		SW	0,0922	0,1749	0,2714	0,3884	0,4988	0,8965	1
	$\alpha = 0,10$	MLPN	0,57	0,7278	0,7601	0,7873	0,8244	0,8575	0,9982
		KS	0,0008	0,0029	0,0079	0,0147	0,0183	0,0738	0,9232
		LL	0,1351	0,1968	0,2526	0,3198	0,3798	0,6481	0,9995
		CVM	0,0001	0,0003	0,0011	0,0031	0,0044	0,0348	0,966
		AD	0,1568	0,2751	0,3591	0,4548	0,5473	0,8557	1
		DAP	0,1201	0,1738	0,2342	0,3127	0,4005	0,8179	1
		JB	0,0244	0,0736	0,1257	0,1851	0,2507	0,7363	1
		SW	0,1603	0,2822	0,4061	0,5404	0,6555	0,955	1
			10	20	30	40	50	100	500
Gamma(4, 5)	$\alpha = 0,01$	MLPN	0,3188	0,56	0,6206	0,6826	0,6984	0,883	0,9988
		KS	0	0	0,0001	0,0004	0,0009	0,0075	0,7384
		LL	0,0325	0,0641	0,0976	0,1434	0,1814	0,444	0,9985
		CVM	0	0	0	0	0,0001	0,0015	0,7903
		AD	0,042	0,1182	0,1888	0,2651	0,3465	0,7194	1
		DAP	0,0799	0,1525	0,2229	0,294	0,3737	0,7051	1
		JB	0,0161	0,1086	0,1857	0,2662	0,3525	0,7123	1
		SW	0,0508	0,1316	0,2316	0,3452	0,4607	0,8674	1
	$\alpha = 0,05$	MLPN	0,5216	0,6845	0,7734	0,7807	0,8291	0,9275	0,9989
		KS	0,0004	0,002	0,0072	0,0102	0,0176	0,0842	0,957
		LL	0,1045	0,1834	0,2593	0,3351	0,3986	0,6997	1
		CVM	0	0,0003	0,0025	0,0031	0,0056	0,0477	0,9876
		AD	0,1245	0,2704	0,3921	0,4907	0,5915	0,8906	1
		DAP	0,1386	0,2459	0,3624	0,4585	0,5504	0,8781	1
		JB	0,0379	0,1581	0,285	0,3936	0,4976	0,8654	1
		SW	0,1319	0,2871	0,4483	0,5784	0,6966	0,9607	1
	$\alpha = 0,10$	MLPN	0,6527	0,7959	0,8412	0,8644	0,9055	0,9742	1
		KS	0,0031	0,0087	0,0243	0,0367	0,0583	0,2015	0,9929
		LL	0,1814	0,2808	0,378	0,4623	0,5389	0,8132	1
		CVM	0,0007	0,0025	0,0105	0,0165	0,0296	0,1665	0,9988
		AD	0,218	0,3896	0,53	0,6219	0,7047	0,9429	1
		DAP	0,1883	0,3305	0,4608	0,5598	0,671	0,9442	1
		JB	0,0666	0,2063	0,3561	0,4777	0,5976	0,9308	1
		SW	0,2189	0,4066	0,5794	0,6969	0,7915	0,9816	1

Tablo 4.5. (Devam) MLPN prosedürü ve Normallik testlerinin farklı α anlam düzeyleri için Beta(2,5), Gama(4,5), Lognormal(0,1), Ki-Kare(5) ve Üstel(1) dağılımları altında güç değerleri.

			10	20	30	40	50	100	500
Lognormal(0,1)	$\alpha = 0,01$	MLPN	0,4276	0,8511	0,9614	0,9833	0,9985	1	1
		KS	0,001	0,0465	0,1528	0,3128	0,492	0,9638	1
		LL	0,2619	0,6025	0,8205	0,9292	0,9765	1	1
		CVM	0	0,0381	0,1537	0,3226	0,5259	0,978	1
		AD	0,3773	0,8126	0,9551	0,9908	0,9983	1	1
		DAP	0,3685	0,6936	0,8668	0,9478	0,9801	1	1
		JB	0,1905	0,6216	0,8401	0,9434	0,9805	1	1
		SW	0,4109	0,8282	0,9692	0,9957	0,9992	1	1
	$\alpha = 0,05$	MLPN	0,6605	0,9148	0,9779	0,9966	0,9989	1	1
		KS	0,0291	0,1956	0,4387	0,6488	0,8028	0,9973	1
		LL	0,4621	0,794	0,9362	0,9797	0,995	1	1
		CVM	0,0142	0,1999	0,4696	0,6955	0,8482	0,9989	1
		AD	0,5788	0,9197	0,9857	0,998	0,999	1	1
		DAP	0,4857	0,8093	0,9345	0,9819	0,9963	1	1
		JB	0,2853	0,7296	0,9122	0,9774	0,9953	1	1
		SW	0,6009	0,9343	0,9914	0,9992	1	1	1
	$\alpha = 0,10$	MLPN	0,7508	0,97	0,9952	0,9998	1	1	1
		KS	0,0027	0,0122	0,0241	0,0329	0,0591	0,1975	0,9925
		LL	0,1847	0,2906	0,3792	0,4566	0,5367	0,8121	1
		CVM	0,0001	0,0027	0,0077	0,0169	0,0309	0,1586	0,9992
		AD	0,2202	0,3992	0,5109	0,6172	0,7067	0,9377	1
		DAP	0,1896	0,3258	0,4473	0,5549	0,6642	0,949	1
		JB	0,0599	0,2112	0,3508	0,4681	0,5906	0,9347	1
		SW	0,2207	0,4118	0,5614	0,6941	0,7894	0,9819	1
			10	20	30	40	50	100	500
Ki – Kare(5)	$\alpha = 0,01$	MLPN	0,2895	0,5434	0,6735	0,6968	0,7665	0,9777	1
		KS	0	0,0001	0,0008	0,0013	0,0036	0,0396	0,9857
		LL	0,0452	0,1148	0,1823	0,2663	0,3438	0,7062	1
		CVM	0	0	0,0001	0,0002	0,0008	0,0146	0,997
		AD	0,0635	0,204	0,3318	0,4867	0,6044	0,9389	1
		DAP	0,1108	0,2214	0,3251	0,45	0,5452	0,8857	1
		JB	0,0247	0,1633	0,2773	0,4187	0,5257	0,8964	1
		SW	0,0782	0,2217	0,397	0,5874	0,7289	0,9835	1

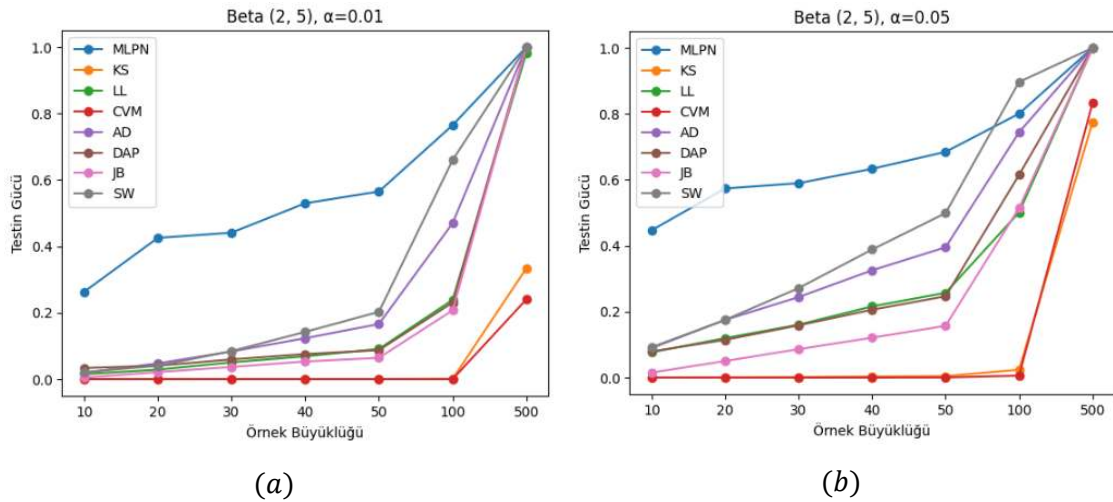
Tablo 4.5. (Devam) *MLPN prosedürü ve Normallik testlerinin farklı α anlam düzeyleri için Beta(2,5), Gama(4,5), Lognormal(0,1), Ki-Kare(5) ve Üstel(1) dağılımları altında güç değerleri.*

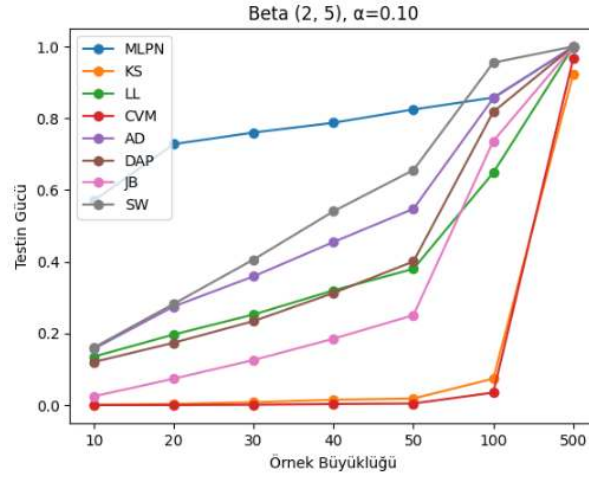
	$\alpha = 0,05$	MLPN	0,4944	0,6969	0,816	0,8321	0,8734	0,9805	1
		KS	0,0003	0,005	0,0151	0,0281	0,0497	0,2421	0,9998
		LL	0,1409	0,2665	0,3862	0,4908	0,5961	0,896	1
		CVM	0	0,001	0,005	0,0134	0,0287	0,2119	1
		AD	0,1806	0,4088	0,5771	0,7095	0,8116	0,9884	1
		DAP	0,1846	0,3462	0,4889	0,6031	0,7159	0,9793	1
		JB	0,0601	0,2372	0,4014	0,5415	0,6754	0,9737	1
		SW	0,1932	0,4369	0,6472	0,7959	0,8899	0,9977	1
	$\alpha = 0,10$	MLPN	0,6204	0,8148	0,8872	0,9123	0,9577	0,9908	1
		KS	0,0854	0,3451	0,6118	0,7967	0,9091	0,9997	1
		LL	0,5724	0,8651	0,9635	0,9922	0,9982	1	1
		CVM	0,0628	0,3668	0,658	0,846	0,9385	0,9998	1
		AD	0,684	0,9527	0,9926	0,9995	0,9999	1	1
		DAP	0,5599	0,8572	0,9644	0,9928	0,9994	1	1
		JB	0,344	0,785	0,9469	0,9888	0,9986	1	1
		SW	0,7007	0,9626	0,9958	0,9997	1	1	1
			10	20	30	40	50	100	500
Üstel(1)	$\alpha = 0,01$	MLPN	0,2681	0,7125	0,889	0,9431	0,9835	1	1
		KS	0	0,0026	0,0101	0,0365	0,0809	0,5997	1
		LL	0,1322	0,3358	0,5471	0,7296	0,8552	0,9979	1
		CVM	0	0,0007	0,0047	0,0259	0,0635	0,6157	1
		AD	0,2075	0,6035	0,8279	0,941	0,9806	1	1
		DAP	0,2234	0,4562	0,6367	0,7845	0,8793	0,999	1
		JB	0,0802	0,3665	0,5886	0,7635	0,8753	0,9994	1
		SW	0,2378	0,6362	0,8802	0,9717	0,994	1	1
	$\alpha = 0,05$	MLPN	0,5231	0,829	0,9425	0,9884	0,9981	1	1
		KS	0,0041	0,0457	0,1165	0,2205	0,3782	0,9295	1
		LL	0,3023	0,5884	0,7754	0,9028	0,9621	0,9999	1
		CVM	0,0005	0,0286	0,1058	0,2263	0,4027	0,951	1
		AD	0,4173	0,8009	0,935	0,9845	0,9977	1	1
		DAP	0,3298	0,6047	0,781	0,9022	0,9613	0,9999	1
		JB	0,1518	0,4826	0,7205	0,8791	0,9531	0,9999	1
		SW	0,4442	0,8367	0,9657	0,995	0,9996	1	1

Tablo 4.5. (Devam) MLPN prosedürü ve Normallik testlerinin farklı α anlam düzeyleri için $Beta(2,5)$, $Gama(4,5)$, $Lognormal(0,1)$, $Ki-Kare(5)$ ve $Üstel(1)$ dağılımları altında güç değerleri.

$\alpha = 0,10$	MLPN	0,6666	0,9216	0,9775	0,9953	0,9994	1	1
	KS	0,005	0,0247	0,0536	0,089	0,1277	0,423	1
	LL	0,2309	0,3886	0,5165	0,6343	0,7221	0,9445	1
	CVM	0,0013	0,0092	0,0278	0,0605	0,0978	0,4438	1
	AD	0,2869	0,5366	0,6929	0,8036	0,8894	0,9942	1
	DAP	0,2464	0,4234	0,5804	0,7148	0,8265	0,9943	1
	JB	0,0887	0,2968	0,4788	0,6442	0,7739	0,9915	1
	SW	0,2938	0,5648	0,7462	0,8712	0,9408	0,9991	1

$Beta(2,5)$ asimetrik dağılımı altında tüm α anlam düzeyleri ve örneklem büyüklüklerinde MLPN prosedürü en iyi performansı göstermiştir. Bu dağılım altında MLPN prosedüründen sonra ele alınan normallik testleri arasında güç sıralmasına göre SW, daha sonra ise AD normallik testi gelmektedir. KS ile CVM testlerinin dışındaki MLPN prosedürü ve ele alınan tüm normallik testleri $n = 500$ için maksimum güç düzeyine olan 1'e ulaşmıştır. **Tablo 4.4'**de $Beta(2,5)$ için sunulan sonuçlar sonuçlar **Şekil 4.10'**da görselleştirilmiştir.

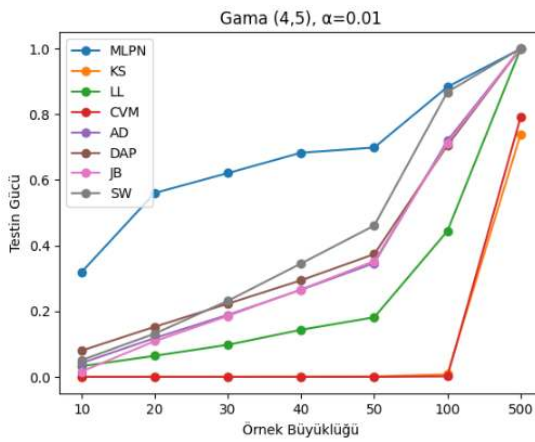




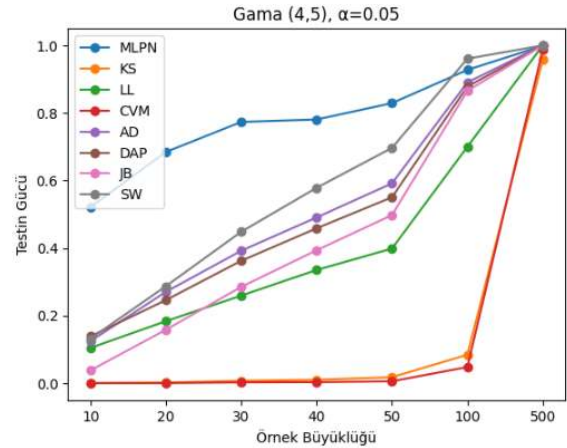
(c)

Şekil 4.10. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için Beta(2,5) dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.

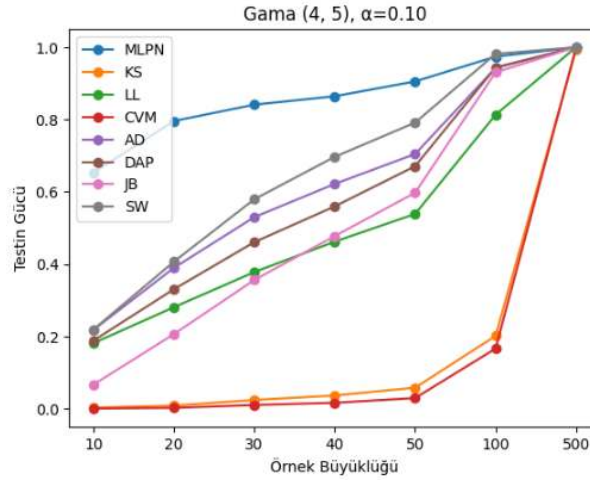
Asimetrik $Gama(4, 5)$ dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerleri hemen hemen $Beta(2, 5)$ dağılımında olduğu gibidir. Burada sıralamasına göre en güçlü üç test, MLPN prosedürü, DAP ve SW normallik testleridir. KS ve CVM testinden başka diğer normallik testleri ve MLPN prosedürü örneklem büyüklüğü $n = 500$ olduğunda maksimum güç düzeyi olan 1'e ulaşamamıştır. Elde edilen sonuçlar Şekil 4.11'de görselleştirilerek özetlenmiştir.



(a)



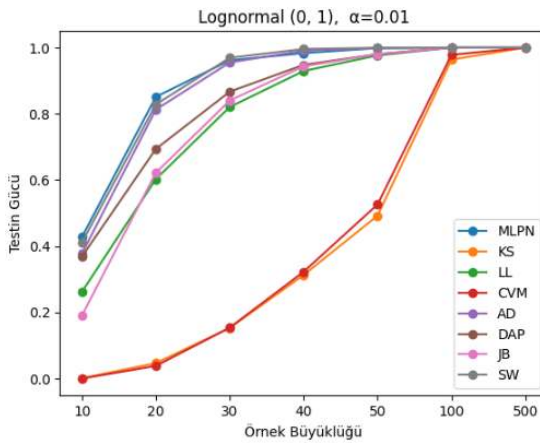
(b)



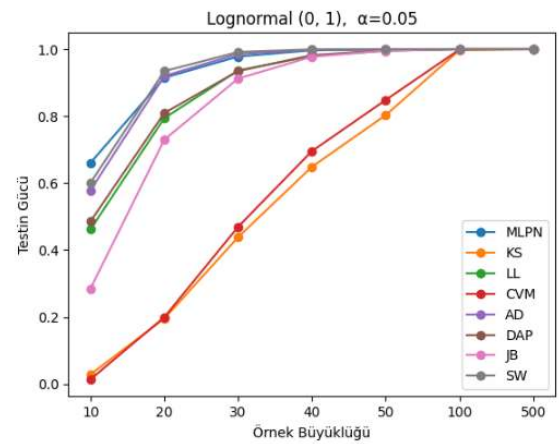
(c)

Şekil 4.11. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için Gama(4,5) dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.

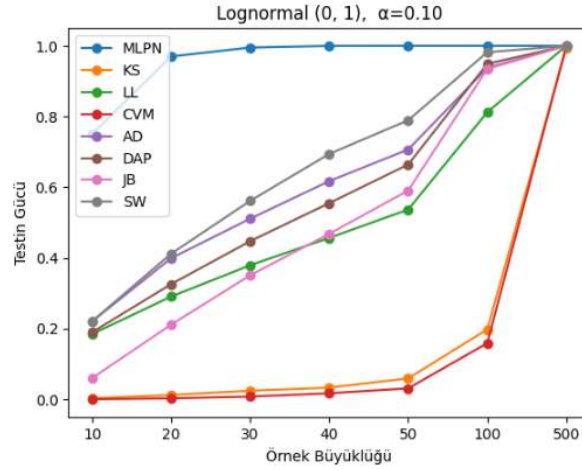
Lognormal(0, 1) dağılımı altında MLPN prosedürü, ele alınan diğer normallik testlerinden daha yüksek değerler almıştır. Bununla birlikte bu dağılım altında diğer testler de küçük örneklem büyüklüklerinde bundan önceki dağılımlara kıyasla daha yüksek güç sergilemiştir. Örneklem büyüklüğü $n = 100$ olduğunda KS ve CVM testlerinden başka MLPN prosedürü ve ele alınan diğer normallik testleri 1 değerini almışlardır. Bu dağılım altında ele alınan normallik testleri arasında MLPN prosedüründen sonra ikinci güçlü test SW testidir. Elde edilen sonuçlar Şekil 4.12’de görselleştirilerek özetlenmiştir.



(a)



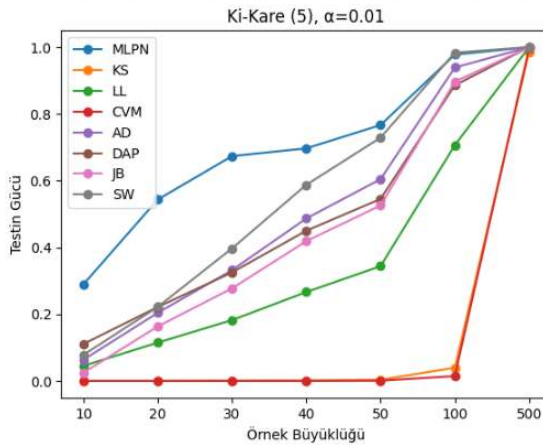
(b)



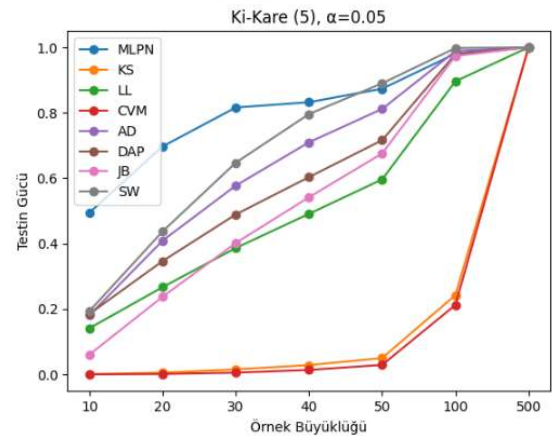
(c)

Şekil 4.12. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için Lognormal(0,1) dağılımı altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.

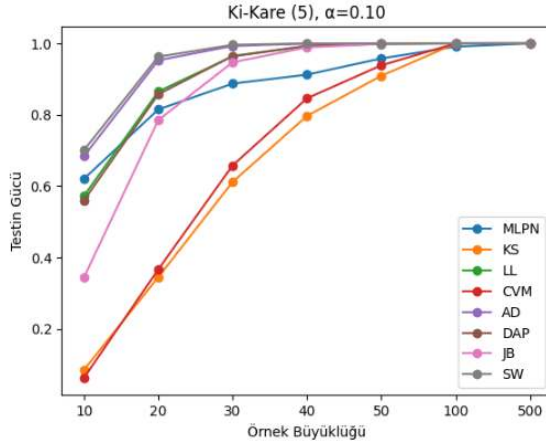
Ki – Kare(5) ve Üstel(1) dağılımları altındaki normallik testleri ve MLPN prosedürünün sonuçları birbirine çok benzerdir. Burada da MLPN prosedüründen sonra en güçlü testin SW, ondan sonra ise DAP testi olduğu görülmektedir. Elde edilen sonuçlar Şekil 4.13’de görselleştirilerek özetlenmiştir.



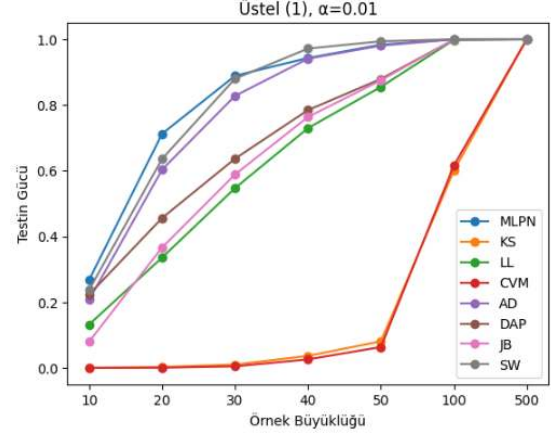
(a)



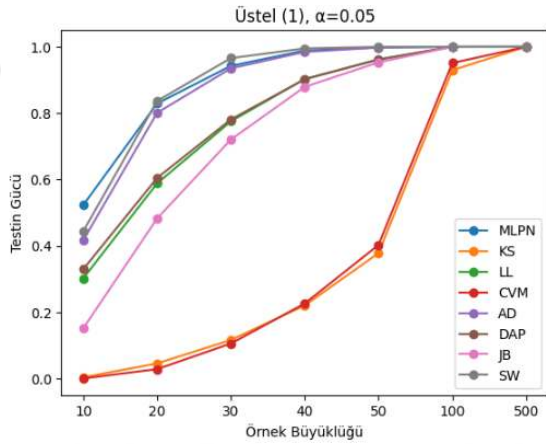
(b)



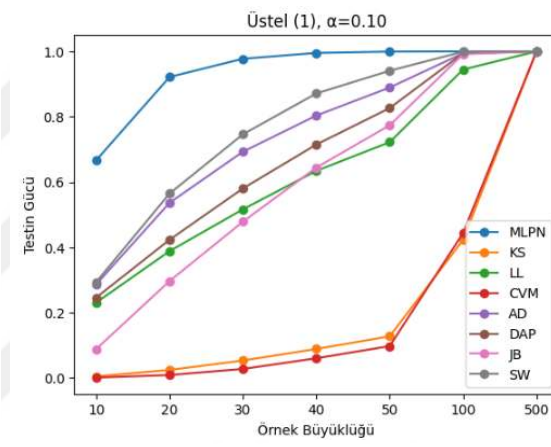
(c)



(d)



(e)



(f)

Şekil 4.13. $\alpha = 0,01$ (a), $\alpha = 0,05$ (b) ve $\alpha = 0,10$ (c) anlam düzeyleri için Ki – kare(5) ve $\alpha = 0,01$ (d), $\alpha = 0,05$ (e) ve $\alpha = 0,10$ (f) anlam düzeyleri için Üstel(1) dağılımları altında MLPN prosedürü ve normallik testlerine ait güç değerlerinin grafiksel gösterimi.

Çalışmada ele alınan normallik testleri için tüm dağılımlar altında en az güce sahip test KS ve CVM testleridir. Aynı zamanda çalışmadaki normallik testlerinin güçleri bu konuyla ilgili literatürde bulunan diğer çalışmaların sonuçları ile örtüşmektedir.

MLPN prosedüründe kullanılan makine öğrenmesi modelinin sınıflandırma başarısını göstermek için güç değerlerine ait karışıklık matrisi, doğruluk, kesinlik, duyarlılık ve F1-skorları Tablo 4.6’da olduğu gibidir.

Tablo 4.6. MLPN prosedürünün güç değerleri için sınıflandırma başarısını gösteren kriterler.

Testin Gücü	Gerçekte Beta(2, 2) dağılımı													
	10		20		30		40		50		100		500	
	N	ND	N	ND	N	ND	N	ND	N	ND	N	ND	N	ND

Tablo 4.6. (Devam) *MLPN prosedürünün güç değerleri için sınıflandırma başarısını gösteren kriterler.*

MLPN Tahmini	<i>N</i>	0	543	0	396	0	388	0	369	0	350	0	293	0	0
	<i>ND</i>	0	457	0	604	0	612	0	631	0	650	0	707	0	1000
	Doğruluk	0,457		0,604		0,612		0,631		0,65		0,707		1	
	Kesinlik	1		1		1		1		1		1		1	
	Duyarlılık	0,457		0,604		0,612		0,631		0,65		0,707		1	
	F1 Skoru	0,627		0,753		0,759		0,773		0,78		0,828		1	

Testin Gücü		Gerçekte <i>Lojistik</i> (0, 1) dağılımı													
		10		20		30		40		50		100		500	
		<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>
MLPN Tahmini	<i>N</i>	0	896	0	855	0	797	0	716	0	708	0	520	0	58
	<i>ND</i>	0	104	0	145	0	203	0	284	0	292	0	480	0	942
	Doğruluk	0,104		0,145		0,203		0,284		0,292		0,48		0,942	
	Kesinlik	1		1		1		1		1		1		1	
	Duyarlılık	0,104		0,145		0,203		0,284		0,292		0,48		0,942	
	F1 Skoru	0,188		0,253		0,337		0,442		0,452		0,648		0,97	

Testin Gücü		Gerçekte <i>Uniform</i> (0, 1) dağılımı													
		10		20		30		40		50		100		500	
		<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>
MLPN Tahmini	<i>N</i>	0	652	0	513	0	494	0	304	0	220	0	43	0	0
	<i>ND</i>	0	348	0	487	0	506	0	696	0	780	0	957	0	1000
	Doğruluk	0,348		0,487		0,506		0,696		0,78		0,957		1	
	Kesinlik	1		1		1		1		1		1		1	
	Duyarlılık	0,348		0,487		0,506		0,696		0,78		0,957		1	
	F1 Skoru	0,516		0,655		0,671		0,82		0,876		0,978		1	

Testin Gücü		Gerçekte <i>Student t</i> (5) dağılımı													
		10		20		30		40		50		100		500	
		<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>
MLPN Tahmini	<i>N</i>	0	871	0	778	0	665	0	598	0	543	0	276	0	2
	<i>ND</i>	0	129	0	222	0	335	0	402	0	457	0	724	0	998
	Doğruluk	0,129		0,222		0,335		0,402		0,457		0,724		0,998	
	Kesinlik	1		1		1		1		1		1		1	
	Duyarlılık	0,129		0,222		0,335		0,402		0,457		0,724		0,998	
	F1 Skoru	0,228		0,363		0,501		0,573		0,627		0,839		0,998	

Testin Gücü		Gerçekte <i>Laplace</i> (0, 1) dağılımı													
		10		20		30		40		50		100		500	
		<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>

Tablo 4.6. (Devam) *MLPN prosedürünün güç değerleri için sınıflandırma başarısını gösteren kriterler.*

MLPN Tahmini	<i>N</i>	0	837	0	668	0	515	0	409	0	321	0	85	0	0
	<i>ND</i>	0	163	0	332	0	485	0	591	0	679	0	915	0	1000
	Doğruluk	0,163		0,332		0,485		0,591		0,679		0,915		1	
	Kesinlik	1		1		1		1		1		1		1	
	Duyarlılık	0,163		0,332		0,485		0,591		0,679		0,915		1	
	F1 Skoru	0,28		0,498		0,653		0,742		0,808		0,955		1	

Testin Gücü		Gerçekte $Beta(2, 5)$ dağılımı													
		10		20		30		40		50		100		500	
		N	ND	N	ND	N	ND	N	ND	N	ND	N	ND	N	ND
MLPN Tahmini	N	0	554	0	427	0	411	0	367	0	316	0	200	0	0
	ND	0	446	0	573	0	589	0	633	0	684	0	800	0	1000
	Doğruluk	0,446		0,573		0,589		0,633		0,684		0,8		1	
	Kesinlik	1		1		1		1		1		1		1	
	Duyarlılık	0,446		0,573		0,589		0,633		0,684		0,8		1	
	F1 Skoru	0,616		0,728		0,741		0,775		0,812		0,888		1	

Testin Gücü		Gerçekte $Gama(4, 5)$ dağılımı													
		10		20		30		40		50		100		500	
		N	ND	N	ND	N	ND	N	ND	N	ND	N	ND	N	ND
MLPN Tahmini	N	0	479	0	316	0	227	0	220	0	171	0	73	0	2
	ND	0	521	0	684	0	773	0	780	0	829	0	927	0	998
	Doğruluk	0,521		0,684		0,773		0,78		0,829		0,927		0,998	
	Kesinlik	1		1		1		1		1		1		1	
	Duyarlılık	0,521		0,684		0,773		0,78		0,829		0,927		0,998	
	F1 Skoru	0,685		0,812		0,871		0,876		0,906		0,962		0,998	

Testin Gücü		Gerçekte <i>Lognormal</i> (0, 1) dağılımı													
		10		20		30		40		50		100		500	
		<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>
MLPN Tahmini	<i>N</i>	0	340	0	86	0	23	0	4	0	2	0	0	0	0
	<i>ND</i>	0	660	0	914	0	977	0	996	0	998	0	1000	0	1000
	Doğruluk	0,66		0,914		0,977		0,996		0,998		1		1	
	Kesinlik	1		1		1		1		1		1		1	
	Duyarlılık	0,66		0,914		0,977		0,996		0,998		1		1	
	F1 Skoru	0,795		0,955		0,988		0,997		0,998		1		1	

Testin Gücü		Gerçekte $Ki - Kare(5)$ dağılımı													
		10		20		30		40		50		100		500	
		<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>

Tablo 4.6. (Devam) MLPN prosedürünün güç değerleri için sınıflandırma başarısını gösteren kriterler.

MLPN Tahmini	<i>N</i>	0	506	0	304	0	184	0	168	0	127	0	20	0	0
	<i>ND</i>	0	494	0	696	0	816	0	832	0	873	0	980	0	1000
	Doğruluk	0,494		0,696		0,816		0,832		0,873		0,98		1	
	Kesinlik	1		1		1		1		1		1		1	
	Duyarlılık	0,494		0,696		0,816		0,832		0,873		0,98		1	
	F1 Skoru	0,661		0,82		0,898		0,908		0,932		0,98		1	

Testin Gücü		Gerçekte Üstel(1) dağılımı													
		10		20		30		40		50		100		500	
		<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>	<i>N</i>	<i>ND</i>
MLPN Tahmini	<i>N</i>	0	477	0	171	0	58	0	12	0	2	0	0	0	0
	<i>ND</i>	0	523	0	829	0	942	0	988	0	998	0	1000	0	1000
	Doğruluk	0,523		0,829		0,942		0,988		0,998		1		1	
	Kesinlik	1		1		1		1		1		1		1	
	Duyarlılık	0,523		0,829		0,942		0,988		0,998		1		1	
	F1 Skoru	0,686		0,906		0,97		0,993		0,998		1		1	

4.3. Eğitim Setinde Tanıtılmayan Dağılımlara İlişkin Sonuçlar

Bu tez çalışmasında önerilen MLPN prosedürü, $n = 10, 20, 30, 40, 50, 100$ ve 500 örneklem büyüklüğündeki *Normal*(0,1), *Beta*(2,2), *Lojistik*(0,1), *Student t*(3), *Uniform*(0,1), *Laplace*(0,1), *Beta*(2,5), *Gama*(4,5), *Lognormal*(0,1), *Ki – Kare*(5) ve *Üstel*(1) parametrelili dağılımlara sahip veri setlerine ilişkin öznelilik değerleri kullanılarak eğitilmiştir. Daha sonra yine aynı parametre değerlerine sahip bu dağılımlardan yeni veri setleri üretilip MLPN prosedürünün performansı sorgulanmıştır. Tezin bu kısmında MLPN prosedürünün eğitim setinde tanıtılmayan farklı parametre değerlerine sahip dağılımlardan üretilen veri setleri için MLPN prosedürünün performansı incelenmiştir. Farklı α anlam düzeyleri arasında belirgin bir fark olmadığından dolayı tablolarda sadece $\alpha = 0,05$ anlam düzeyine ait sonuçlar paylaşılmıştır.

Tablo 4.7. Eğitim setinde tanıtılmayan dağılımların $\alpha = 0,05$ anlam düzeyinde güç değerleri.

		Örneklem Büyüklüğü <i>n</i>						
		10	20	30	40	50	100	500
<i>Beta</i> (3,	MLPN	0,5318	0,6085	0,6547	0,6942	0,7151	0,7495	0,7908
	KS	0	0	0	0,0001	0,0003	0,0005	0,0049

Tablo 4.7. (Devam) Eğitim setinde tanıtılmayan dağılımların $\alpha=0,05$ anlam düzeyinde güç değerleri.

	LL	0,0381	0,0475	0,0515	0,0521	0,059	0,0827	0.3526
	CVM	0	0	0	0	0	0	0.0013
	AD	0,043	0,0587	0,0602	0,0643	0,068	0,1255	0.7805
	DAP	0,0269	0,0312	0,0407	0,0673	0,0871	0,2887	0.9927
	JB	0,0004	0,0007	0,0012	0,0067	0,0104	0,0328	0.9432
	SW	0,0393	0,043	0,0446	0,0523	0,0613	0,1528	0.977
		10	20	30	40	50	100	500
<i>Lojistik(0,2)</i>	MLPN	0,0872	0,1637	0,213	0,2254	0,2528	0,4132	0.924
	KS	0	0,0001	0,0007	0,0009	0,0013	0,0017	0.0247
	LL	0,0726	0,0864	0,0909	0,1048	0,1117	0,1558	0.4625
	CVM	0	0	0	0,0001	0,0007	0,0011	0,0146
	AD	0,084	0,1243	0,1386	0,1513	0,1588	0,2289	0.7362
	DAP	0,1117	0,1562	0,1865	0,2156	0,2381	0,3388	0.8485
	JB	0,0278	0,0967	0,1482	0,1936	0,2241	0,3626	0.8865
	SW	0,0844	0,1205	0,1486	0,1732	0,1939	0,2943	0.839
		10	20	30	40	50	100	500
<i>Student t(3)</i>	MLPN	0,1821	0,4159	0,559	0,6492	0,7176	0,9276	1
	KS	0,0061	0,0299	0,05	0,0746	0,1018	0,2473	0.9718
	LL	0,1621	0,258	0,3474	0,4166	0,4806	0,7307	0.9998
	CVM	0,0023	0,0238	0,0429	0,0673	0,0962	0,2537	0.9863
	AD	0,19	0,3488	0,4521	0,5348	0,5997	0,8463	1
	DAP	0,233	0,3926	0,5069	0,589	0,6549	0,8778	1
	JB	0,1014	0,318	0,4661	0,5709	0,6564	0,8953	1
	SW	0,1887	0,3419	0,4624	0,5555	0,6283	0,8756	1

		10	20	30	40	50	100	500
<i>Uniform(0,2)</i>	MLPN	0,4569	0,6115	0,6907	0,8607	0,8875	0,9838	1
	KS	0	0	0,0002	0,001	0,0021	0,032	0.9337
	LL	0,0677	0,0968	0,1426	0,2075	0,2596	0,5875	1
	CVM	0	0	0	0	0,0002	0,0041	0.9905
	AD	0,0822	0,1914	0,3274	0,4595	0,5829	0,9453	1
	DAP	0,0277	0,1507	0,3917	0,6261	0,7987	0,9967	1
	JB	0,0028	0,001	0,0026	0,0031	0,0064	0,5608	1
	SW	0,0822	0,1932	0,3902	0,586	0,7502	0,9965	1
		10	20	30	40	50	100	500
<i>Lap</i>	MLPN	0,1539	0,3548	0,4792	0,5553	0,6436	0,8918	1

Tablo 4.7. (Devam) Eğitim setinde tanıtılmayan dağılımların $\alpha=0,05$ anlam düzeyinde güç değerleri.

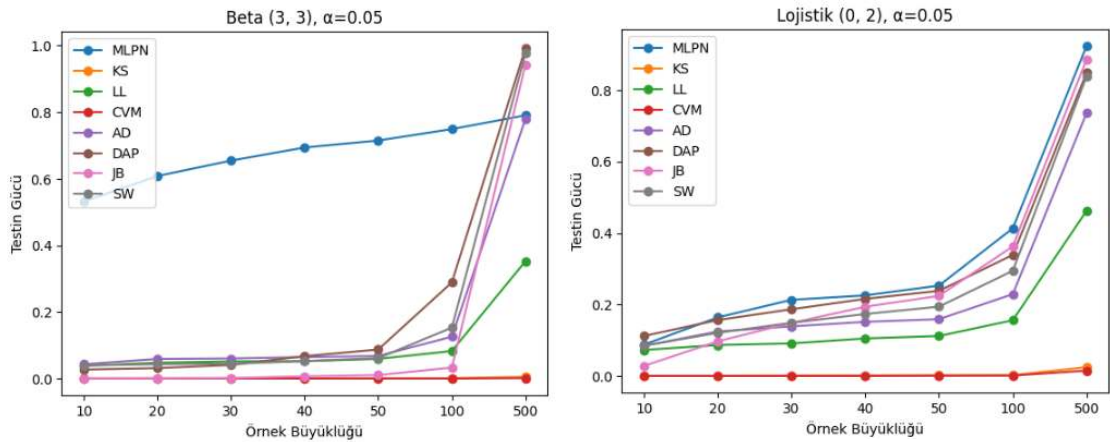
	KS	0,0013	0,0052	0,0114	0,016	0,0282	0,094	0.9526
	LL	0,1375	0,2274	0,2998	0,3692	0,432	0,7064	0.9996
	CVM	0,001	0,0009	0,0036	0,0072	0,0135	0,0693	0.9801
	AD	0,1609	0,2985	0,3921	0,4783	0,5462	0,8238	1
	DAP	0,191	0,3044	0,3824	0,4474	0,4981	0,7341	0.9998
	JB	0,0614	0,2192	0,3342	0,4301	0,5012	0,7829	0.999
	SW	0,1526	0,2674	0,3571	0,4479	0,5159	0,7948	0.9999
		10	20	30	40	50	100	500
<i>Beta(1,3)</i>	MLPN	0,4501	0,6752	0,8278	0,9598	0,968	0,9997	1
	KS	0,0003	0,004	0,0121	0,0251	0,0476	0,2699	1
	LL	0,1397	0,2708	0,4099	0,5276	0,647	0,9546	1
	CVM	0,0001	0,0004	0,0019	0,0072	0,0179	0,2362	1
	AD	0,1913	0,4426	0,6465	0,7885	0,8927	0,998	1
	DAP	0,141	0,2239	0,315	0,4181	0,5516	0,9824	1
	JB	0,036	0,1145	0,2023	0,3154	0,4581	0,9685	1
	SW	0,2065	0,4783	0,7312	0,8816	0,9612	1	1
		10	20	30	40	50	100	500
<i>Gama(2,1)</i>	MLPN	0,5419	0,7609	0,8698	0,951	0,9591	0,9992	1
	KS	0,0008	0,0073	0,0256	0,0486	0,0877	0,3647	1
	LL	0,175	0,3211	0,4722	0,5937	0,7006	0,9524	1
	CVM	0	0,0021	0,0112	0,0264	0,0583	0,3739	1
	AD	0,2274	0,4997	0,6813	0,8061	0,8897	0,9972	1
	DAP	0,2141	0,4013	0,5693	0,6883	0,795	0,9923	1
	JB	0,0792	0,2904	0,4848	0,6323	0,7614	0,9904	1
	SW	0,2432	0,5374	0,7481	0,8797	0,9445	0,9998	1

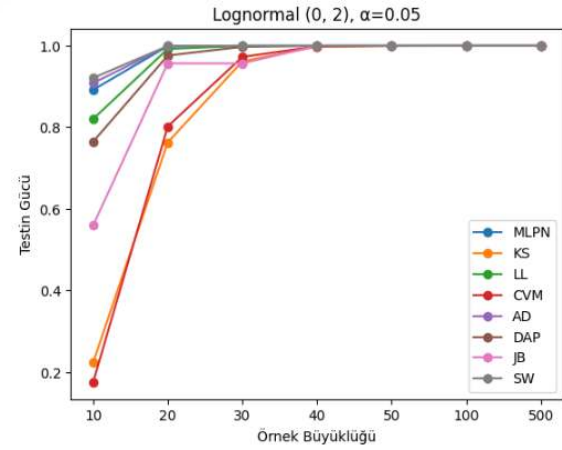
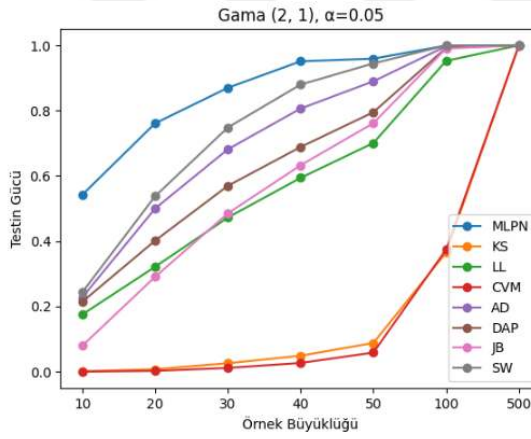
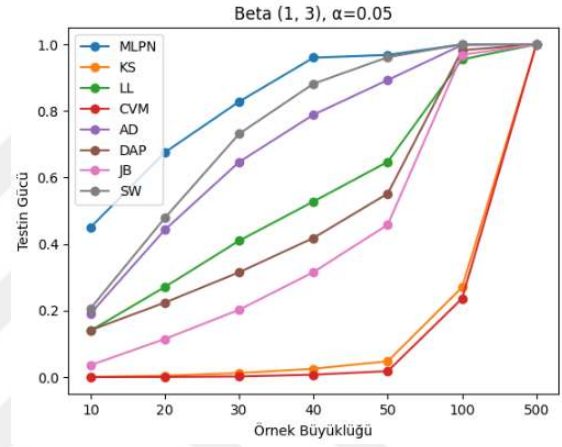
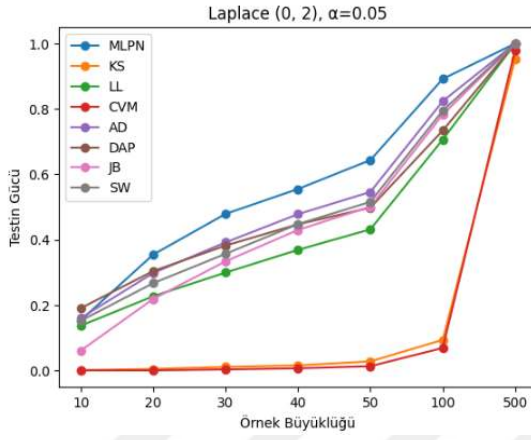
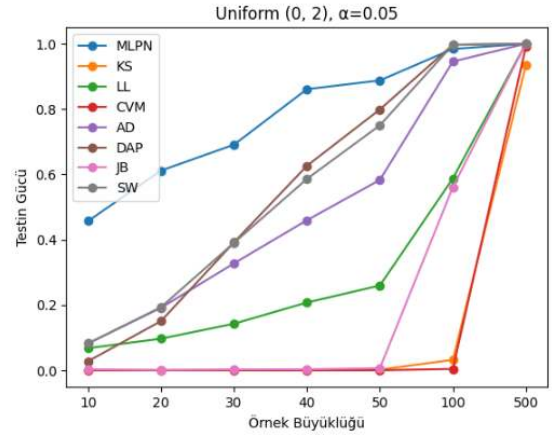
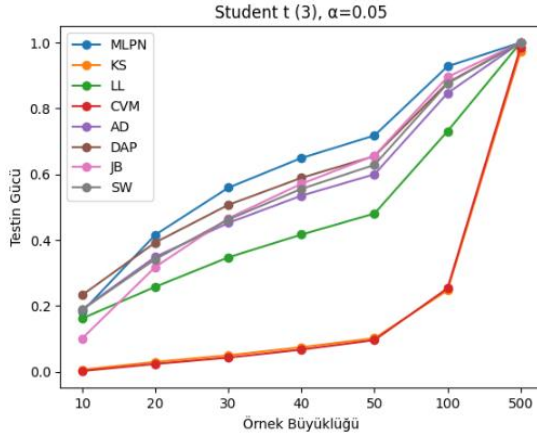
		10	20	30	40	50	100	500
<i>Lognormal(0,2)</i>	MLPN	0,8919	0,999	1	1	1	1	1
	KS	0,2236	0,7631	0,9617	0,9978	0,9999	1	1
	LL	0,8205	0,9922	0,9999	1	1	1	1
	CVM	0,1755	0,8023	0,9725	0,9984	0,9998	1	1
	AD	0,9085	0,999	1	1	1	1	1
	DAP	0,7649	0,976	0,9974	0,9999	1	1	1
	JB	0,5597	0,9568	0,9961	0,9999	1	1	1
	SW	0,921	0,9992	1	1	1	1	1

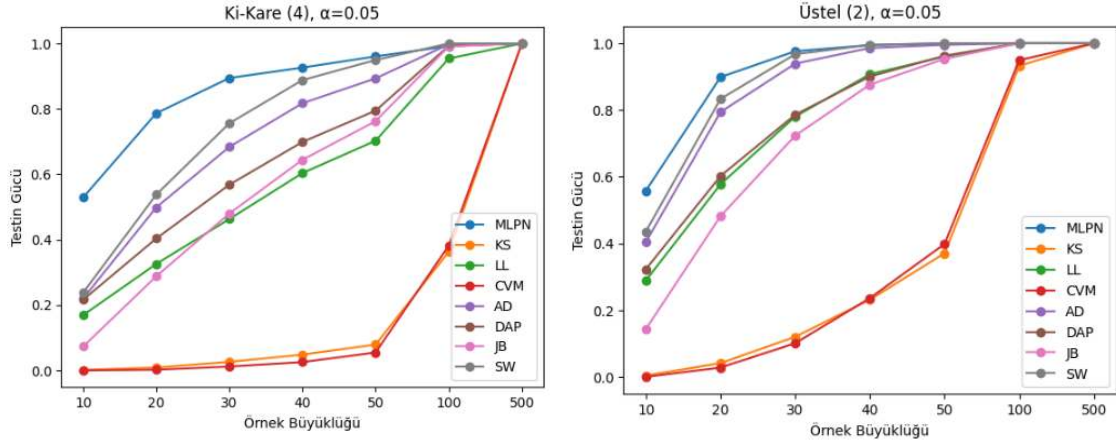
Tablo 4.7. (Devam) Eğitim setinde tanıtılmayan dağılımların $\alpha=0,05$ anlam düzeyinde güç değerleri.

		10	20	30	40	50	100	500
$Ki - Kare(4)$	MLPN	0,5288	0,7863	0,8942	0,926	0,9602	0,9918	1
	KS	0,0012	0,0088	0,0259	0,0483	0,0789	0,3626	1
	LL	0,1689	0,3255	0,4629	0,6035	0,7024	0,9538	1
	CVM	0	0,0019	0,0118	0,0253	0,0548	0,3815	1
	AD	0,2198	0,4984	0,6836	0,8171	0,893	0,9971	1
	DAP	0,2161	0,4033	0,5679	0,6987	0,7943	0,9924	1
	JB	0,0729	0,2877	0,4796	0,6441	0,7622	0,9908	1
	SW	0,2379	0,5376	0,7558	0,8879	0,9491	0,9999	1
		10	20	30	40	50	100	500
$\dot{U}stel(2)$	MLPN	0,5581	0,8987	0,976	0,9944	0,9995	1	1
	KS	0,0049	0,0419	0,1209	0,2321	0,3709	0,9323	1
	LL	0,2896	0,5776	0,7804	0,9072	0,9602	1	1
	CVM	0,0008	0,0287	0,1015	0,2366	0,3989	0,9494	1
	AD	0,4047	0,7939	0,9388	0,9857	0,9954	1	1
	DAP	0,3232	0,6024	0,7859	0,9002	0,9624	1	1
	JB	0,1439	0,4821	0,7234	0,876	0,9532	1	1
	SW	0,4362	0,8332	0,9671	0,9952	0,9996	1	1

Sonuçlardan da görüldüğü gibi MLPN prosedürü, eğitim setinde tanıtılmayan veri yapılarında da veri setinin normal dağılıma uygun olup olmadığını tahmin etme konusunda ele alınan normallik testlerinden çoğu durumda daha güçlü performans göstermektedir. Elde edilen sonuçlar **Şekil 4.14**'de görselleştirilerek özetlenmiştir.







Şekil 4.14. Eğitim setinde tanıtılmayan dağılımlar altında $\alpha = 0,05$ güven düzeyi için MLPN prosedürü ve normallik testlerine ait güç değerleri.

4.4. Kategoring Boosting ve Rassal Ormanlar Algoritmalarının Karşılaştırılmasına İlişkin Sonuçlar

Literatürde sınıflandırma problemleri için birçok makine öğrenmesi algoritması bulunmaktadır. Bölüm 3’de ifade edilen bu algoritmaların her biri önerilen uyum iyiliği prosedürünün oluşturulmasında kullanılmış ve en iyi performansı CatBoost algoritmasının verdiği gözlemlenmiştir. Ayrıca CatBoost algoritmasından sonra en iyi performans, RF algoritmasıyla elde edilmiştir. CatBoost ve RF algoritmalarının $\alpha = 0.05$ anlam düzeyinde I. Tip hata oranı ve güç değerlerine ilişkin sonuçlar aşağıdaki tabloda derlenmiştir.

Tablo 4.8. CatBoost ve RF algoritmalarının $\alpha = 0.05$ anlam düzeyinde I. Tip hata ve güç değerleri.

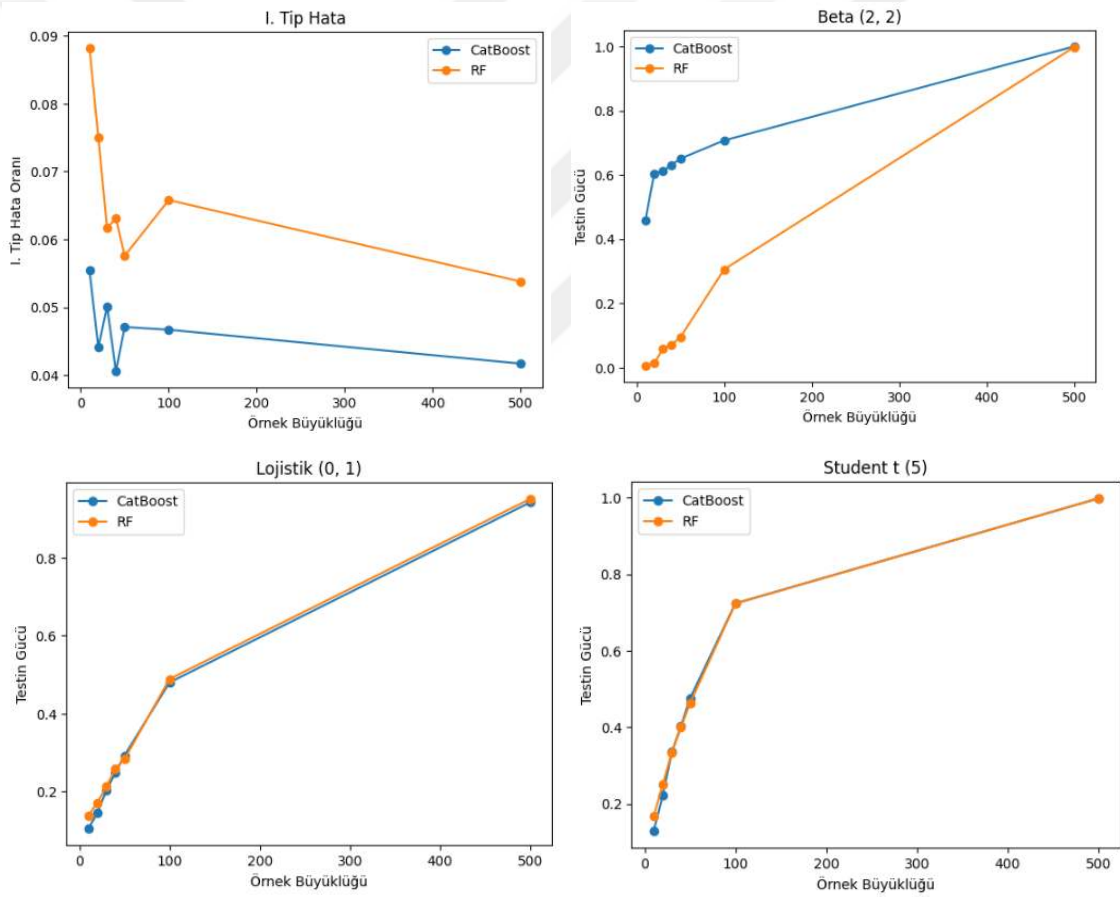
I. Tip hata		Örneklem Büyüklüğü n						
$Normal(0,1)$	Algoritma	10	20	30	40	50	100	500
	CatBoost	0,0554	0,0441	0,0501	0,0406	0,0471	0,0467	0,0417
	RF	0,0882	0,075	0,0617	0,0631	0,0576	0,0658	0,0538
Testin Gücü		Örneklem Büyüklüğü n						
$Beta(2,2)$	Algoritma	10	20	30	40	50	100	500
	CatBoost	0,4576	0,6041	0,6125	0,631	0,6507	0,7074	1
	RF	0,0056	0,0136	0,0575	0,0717	0,0947	0,3059	0,9979
		Örneklem Büyüklüğü n						
$Lojistik(0,1)$	Algoritma	10	20	30	40	50	100	500
	CatBoost	0,1049	0,1454	0,2034	0,2484	0,2924	0,4802	0,9428

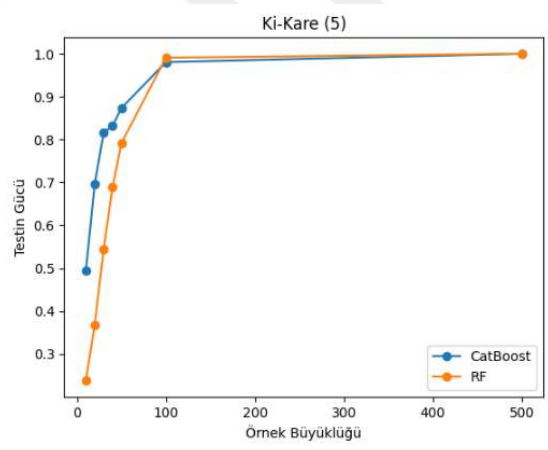
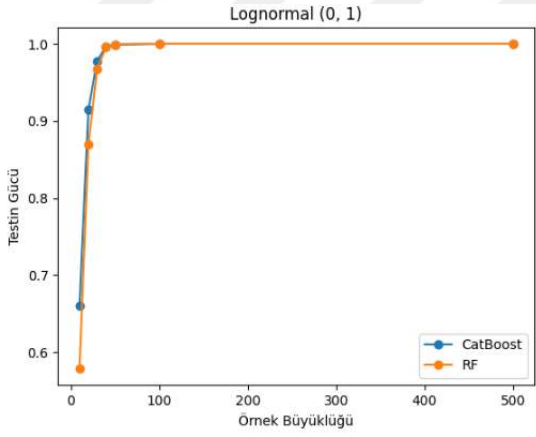
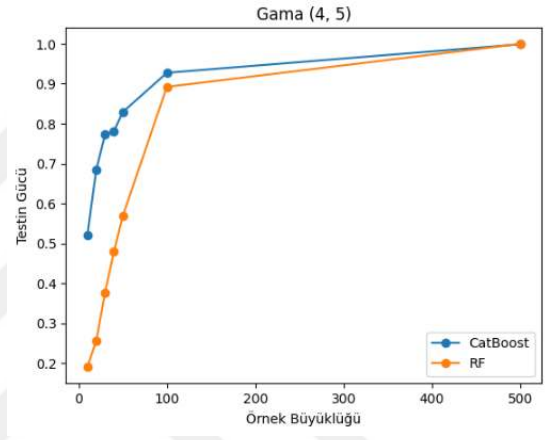
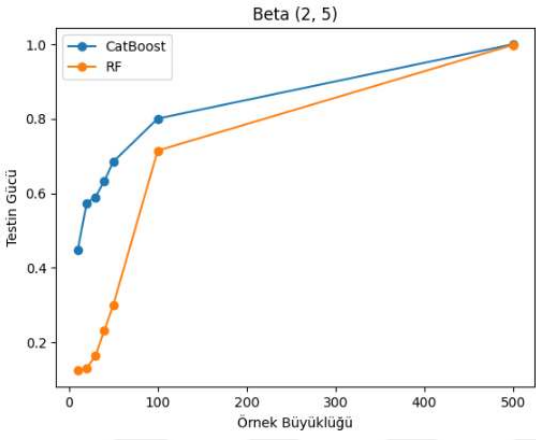
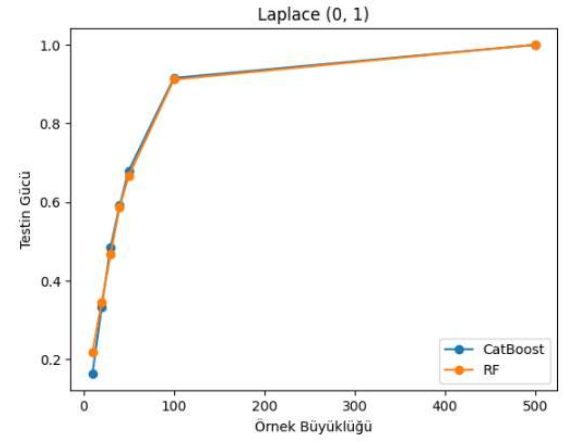
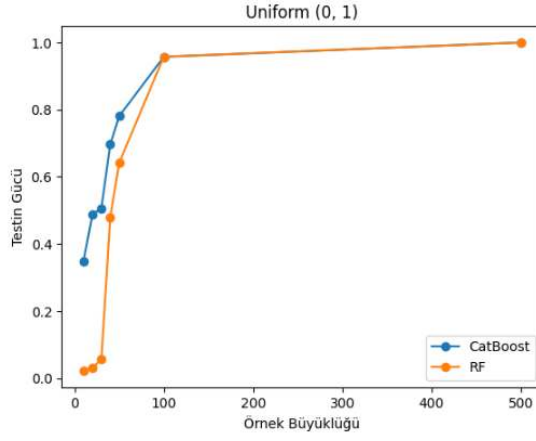
Tablo 4.8. (Devam) CatBoost ve RF algoritmalarının $\alpha = 0.05$ anlam düzeyinde I. Tip hata ve güç değerleri.

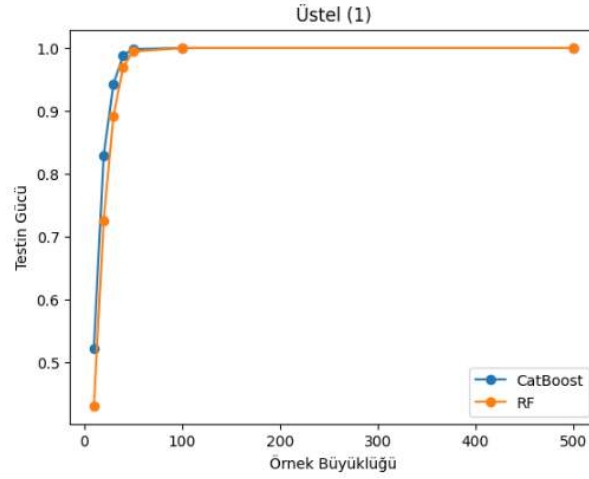
	RF	0.1379	0.1715	0.2126	0.2586	0.2831	0.489	0.9516
<i>Student t(5)</i>	Algoritma	10	20	30	40	50	100	500
	CatBoost	0.1293	0.2225	0.3358	0.4029	0.4756	0.724	0.9983
	RF	0.1678	0.2513	0.3341	0.4002	0.4628	0.7235	0.9986
		Örneklem Büyüklüğü n						
<i>Uniform(0,1)</i>	Algoritma	10	20	30	40	50	100	500
	CatBoost	0.3484	0.4878	0.5063	0.6963	0.7808	0.9573	1
	RF	0.0227	0.0308	0.0571	0.4788	0.6418	0.9577	1
		Örneklem Büyüklüğü n						
<i>Laplace(0,1)</i>	Algoritma	10	20	30	40	50	100	500
	CatBoost	0.1638	0.3321	0.4858	0.5914	0.6795	0.9155	1
	RF	0.2186	0.3442	0.4674	0.5874	0.6663	0.9116	1
		Örneklem Büyüklüğü n						
<i>Beta(2,5)</i>	Algoritma	10	20	30	40	50	100	500
	CatBoost	0.4469	0.5735	0.5896	0.6332	0.6846	0.8004	1
	RF	0.1247	0.1294	0.1639	0.232	0.2997	0.7145	0.9984
		Örneklem Büyüklüğü n						
<i>Gama(4,5)</i>	Algoritma	10	20	30	40	50	100	500
	CatBoost	0.5216	0.6845	0.7734	0.7807	0.8291	0.9275	0.9989
	RF	0.1908	0.2558	0.377	0.4791	0.5684	0.8923	1
		Örneklem Büyüklüğü n						
<i>Lognormal (0,1)</i>	Algoritma	10	20	30	40	50	100	500
	CatBoost	0.6605	0.9148	0.9779	0.9966	0.9989	1	1
	RF	0.579	0.8695	0.9675	0.9959	0.9991	1	1

		Örneklem Büyüklüğü n						
<i>Ki – Kare(5)</i>	Algoritma	10	20	30	40	50	100	500
	CatBoost	0.4944	0.6969	0.816	0.8321	0.8734	0.9805	1
	RF	0.2385	0.3681	0.545	0.6885	0.792	0.9907	1
		Örneklem Büyüklüğü n						
<i>Üstel(1)</i>	Algoritma	10	20	30	40	50	100	500
	CatBoost	0.5231	0.829	0.9425	0.9884	0.9981	1	1
	RF	0.4309	0.7258	0.8921	0.9699	0.9944	1	1

Tablo 4.8'den algoritmaların I. Tip hata oranlarına bakıldığında, CatBoost algoritmasının α anlam düzeyine yakın sonuçlar verdiği, fakat RF algoritmasının α anlam düzeyinden daha yüksek sonuçlar verdiği görülmektedir. *Lojistik(0,1)*, *Student t(5)* ve *Laplace(0,1)* alternatif dağılımları altında $n = 10$ ve $n = 20$ örnek büyüklükleri için RF algoritması CatBoost algoritmasından daha güçlü bir performans gösterdese de diğer örnek büyüklükleri için CatBoost algoritması RF algoritmasından daha güçlü bir performans sergilemiştir. Diğer alternatif dağılımların tüm örnek büyüklüklerinde ise CatBoost algoritması RF algoritmasına kıyasla güç bakımından daha iyi bir performans göstermiştir. Elde edilen sonuçlar **Şekil 4.15'**de görselleştirilerek özetlenmiştir.







Şekil 4.15. CatBoost ve RF algoritmalarının performans karşılaştırması.

CatBoost ve RF algoritmasının eğitim zamanı, model boyutu ve tahmin süresi ölçütlerine göre performans değerlerine ilişkin sonuçlar ise aşağıdaki tabloda sunulmuştur.

Tablo 4.9. CatBoost ve RF algoritmalarının farklı ölçütlere göre performansı.

Ölçütler	CatBoost Algoritması	RF Algoritması
Eğitim Zamanı	0,9905 saniye	13,0126 saniye
Model Boyutu	117060 bayt	39363912 bayt
Tahmin Süresi	0,0082 saniye	0,0291 saniye

Tablo 4.9'dan da görüldüğü gibi eğitim zamanı, model boyutu ve tahmin süresi gibi performans değerlendirme ölçütlerine göre MLPN prosedürünün oluşturulmasında CatBoost algoritması RF algoritmasından daha iyi performans göstermiştir.

5. SONUÇ VE ÖNERİLER

Literatürdeki mevcut testlerin veri yapısına ve örneklem büyüklüğüne bağlı olarak farklı koşullar altında farklı sonuçlar vermesi problemi, bu tez çalışmasının ortaya çıkması için ilham kaynağı olmuştur. Buna bağlı olarak bu tez çalışmasında literatürde yer alan normallik testlerine alternatif olarak normal dağılım için makine öğrenmesi algoritmasına dayalı uyum iyiliği prosedürü (MLPN) önerilmiştir.

MLPN prosedürünün çalışma mantığının en önemli özelliği bir veri setinin normal dağılıma sahip bir anakütleden çekilip çekilmediğini bir sınıflandırma problemi olarak ele almasıdır. Bu mantıktan yola çıkarak, farklı dağılım yapıları ve farklı örnek büyüklüklerinde üretilen veri setlerinden elde edilen öznitelik değerlerine göre sınıflandırma problemi makine öğrenmesi algoritmalarıyla yapılarak MLPN prosedürü oluşturulmuştur.

Bunlara ek olarak literatürde yaygın olarak kullanılan normallik testlerinden yedi tanesi seçilmiş ve önerilen MLPN prosedürünün bu normallik testlerine göre performansı I. Tip hata ve güç değerleri bakımından kapsamlı bir simülasyon çalışması yapılarak karşılaştırılmıştır. Bu tez çalışmasında yürütülen tüm hesaplamalar ve aynı zamanda önerilen MLPN prosedürü Python programlama dilinde kodlamalar yapılarak gerçekleştirilmiştir.

Simülasyon çalışmasında elde edilen sonuçlar, literatüre sunulan MLPN prosedürünün tez çalışmasında ele alınan normallik testlerinden hemen hemen tüm durumlar için daha güçlü kararlar verdiğini göstermiştir. $Beta(2, 2)$ ve $Uniform(0, 1)$ dağılımlarının tüm örneklem büyüklüklerinde, $Lojistik(0, 1)$ dağılımı altında ise $n \geq 30$ örneklem büyüklüklerinde MLPN prosedürüne ait güç performansları çalışmada ele alınan normallik testlerinin performanslarından daha iyidir.

$Student\ t(5)$ ve $Laplace(0, 1)$ dağılımları altında örneklem büyüklüğü $n \geq 30$ olduğu durumlarda MLPN prosedürünün performansları çalışmada ele alınan normallik testlerinin performanslarından güç bakımında daha yüksek çıkmıştır.

Asimetrik dağılımlar altında ise tüm örneklem büyüklükleri için MLPN prosedürü çalışmada ele alınan normallik testlerinden daha iyi güç performansı göstermiştir.

Ayrıca çalışmada ele alınan normallik testleri arasında belirlenen alternatif dağılıma bağlı olarak MLPN prosedüründen sonra en güçlü normallik testlerinin, DAP, AD ve SW olduğu, en zayıf normallik testlerinin ise KS ve CVM olduğu gözlemlenmiştir. Simülasyon çalışması sonucunda örneklem büyüklüğü arttıkça MLPN prosedürü ve

alıřmada ele alınan normallik testlerine ait g deęerlerinin arttıęı ve rneklem byklę 500 olduęunda ise testlerin gcnn genellikle 1'e ulařtıęı grlmřtr.

Bu tez alıřmasında nerilen uyum iyilięi prosedr, normal daęılıma uygunluęu belirlemek iin tasarlanmıřtır. Gelecekte yapılacak alıřmalar iin bu tez alıřması ncl kaynak kabul edilip dięer daęılımlara uygunluęu belirlemek amacıyla yeni bir uyum prosedrleri geliřtirilebilir.



KAYNAKÇA

- Ajne, B. (1968). A simple test for uniformity of a circular distribution. *Biometrika*, 55(2), 343-354.
- Akdeniz, F. (2018). *Olasılık ve İstatistik*. Ankara: Akademisyen Kitabevi.
- Akıllı, K. (2020). *Ampirik Dağılım Fonksiyonlarının Uyum İyiliği Testlerine Etkisi*. Konya: Yüksek Lisans Tezi, Necmeddin Erbakan Üniversitesi.
- Alpaydın, E. (2014). *Introduction to Machine Learning*. Massachusetts Institute of Technology.
- Anderson, T., and Darling, D. (1952). *Asymptotic theory of certain “goodness of fit” criteria based on stochastic processes*. The annals of mathematical statistics, 23(2), 193-212.
- Bayoud, H. A. (2021). Tests of normality: new test and comparative study. *Communications in Statistics - Simulation and Computation*, 4442-4463.
- Bilim, B. (2017). *Büyük Örnekler İçin Normallik Testlerinin Değerlendirilmesi*. Çukurova Üniversitesi. Yüksek Lisans Tezi.
- Breton, M. D., Devore, M. D., and Brown, D. E. (2008). A tool for systematically comparing the power of tests for normality. *Journal of Statistical Computation and Simulation*, 623-638.
- Büyükuysal, M. C. (2014). *Farklı Örneklem Genişliklerinde Normal Dağılım Testlerinin Karşılaştırılması*. Bülent Ecevit Üniversitesi. Doktora Tezi.
- Ceylan, T. (2018). *Perakende Sektöründe Makine Öğrenmesine Dayalı Yaklaşımlar*. İstanbul: Yıldız Teknik Üniversitesi.
- Cramer, H. (1928). *On the composition of elementary errors*. Skand. Aktuar, 11:141-180.
- D'Agostino, R. (1972). *Small sample probability points for the D test of normality (complete samples)*. Biometrika, 59, 219-221.
- Demir, S. (2022). Comparison of Normality Tests in Terms of Sample Sizes under Different Skewness and Kurtosis Coefficients. *International Journal of Assessment Tools in Education*, Vol. 9, No. 2, 397-409.

- Demirtürk, S. (2023). *Makine Öğrenmesinde Sınıflandırma Yöntemleri Kullanılarak Ulaşım Kartı Suistimalinin Tespit Edilmesi*. Samsun: On Dokuz Mayıs Üniversitesi.
- Doulah, S. (2019). A Comparison among Twenty-Seven Normality Tests. *Research & Reviews: Journal of Statistics.*, 8(3): 41-59p.
- Esteban, M., Marhuenda, Y., Morales, D., & Sanchez, A. (2007). *New goodness-of-fit tests based on sample quantiles*. Communications in Statistics-Simulation and Computation, 36(3), 631-642.
- Gamgam, H., ve Altunkaynak, B. (2012). *Parametrik olmayan yöntemler: SPSS uygulamalı*. Seçkin Yayıncılık.
- Hosking, J. (1990). *L-moments analysis and estimation of distributions using linear combinations*. F statistics. J. Roy. Statistics. Soc. B., 52: 105-124.
- Hosking, J., and Wallis, J. (1997). *Regional Frequency Analysis An Approach Based on L-Moments*. London, UK: Cambridge University Press.
- James, G., Witten, D., Hastie, T. ve Tibshirani, R. (2017). *An introduction to statistical learning*. Springer Texts in Statistics.
- Jarque, C., and Bera, A. (1987). *A test for normality of observations and regression residuals*. International Statistical Review, 55, 163–172.
- Khatun, N. (2021). Applications of Normality Test in Statistical Analysis. *Open Journal of Statistics*, 11, 113-122.
- Kolmogorov, A. (1933). *Sulla determinazione empirica di una lgge di distribuzione*. Inst. Ital. Attuari, Giorn., Vol. 4, 83-91.
- Köle, C. (2014). *Üstel dağılım için uyum iyiliği testleri ve bir karşılaştırma*. Ankara: Yüksek Lisans Tezi, Gazi Üniversitesi Fen Bilimleri Enstitüsü.
- Kuiper, N. (1962). *Tests Concerning Random Points on a Circle*. Proc. Koninkl. Neder. Akad. Van. Wetenschappen. 63: 38-47.
- Kutsal, Ö. O. (2016). Entropiye Dayalı Düzgün Dağılıma Uygunluk Testleri. *Gazi Üniversitesi, Yüksek Lisnans Tezi*, 11-12.

- Lee, C., Park, S., and Jeong, J. (2016). Comprehensive comparison of normality tests: Empirical study using many different types of data. *Journal of the Korean Data & Information Science Society*, 1399-1412.
- Lilliefors, H. (1967). *On the Kolmogorov-Smirnov test for normality with mean and variance unknown*. Journal of The American Statistical Association, Vol.64, 534-544.
- Mala, I., Sladek, V. ve Habarta, F. (2012). Comparison of estimates using L and TL moments and other robust characteristics of distributional shape and tail heaviness. *Statistical Journal*, 20(5), 529-546.
- Momayez, D. (2023). *Makine Öğrenmesi Kullanarak İnsan Hareketlerinin Sınıflandırılması: Akıllı Telefon Algılayıcıları Örneği*. Samsun: On Dokuz Mayıs Üniversitesi.
- Noughabi, H., and Arghami, N. (2010). Monte Carlo comparison of seven normality tests. *Journal of Statistical Computation and Simulation* 81 (8):, 965–72. doi:10.1080/00949650903580047.
- Ogunleye, L. I., Oyejola, B. A., and Obisesan, K. O. (2018). Comparison of Some Common Tests for Normality. *International Journal of Probability and Statistics*, 130-137.
- Özer, A. (2007). Normallik Testlerinin Karşılaştırılması. *Anakara Üniversitesi. Yüksek Lisans Tezi*.
- Pearson, K. (1900). On the criterion that a given system of deviations from the probable in the case of a correlated system of variables is such that it can be reasonable supposed to have arisen from random sampling. *Philosophical Magazine* 50, 157.
- Razali, N., and Wah, B. (2010). Power comparisons of some selected normality tests. *ResearchGate*, 126-138.
- Romao, X., Delgado, R., and Costa, A. (2009). An empirical power comparison of univariate goodness-of-fit tests for normality. *Journal of Statistical Computation and Simulation*, 545-591.
- Shabri, A. (2002). Comparisons of the LH Moments and the L Moments. *Research Gate*.

- Shannon, C. E. (1948). A Mathematical Theory of Communication. *The Bell System Technical Journal*, Vol 17, No 3.
- Shapiro, S., and Wilk, M. (1965). "An Analysis of Variance Test for Normality (Complete Samples) ". *Biometrika*, 52, 3/4.591-611.
- Stephens, M. A. (1974). EDF Statistics for Goodness of Fit and Some Comparisons. *Journal of the American Statistical Association*, Vol. 69, 730-737.
- Tokmak, A. (2023). *Mikroservis Ekosisteminde Servis Keşfi Mekanizması*. İstanbul: İstanbul Kültür Üniversitesi.
- Wijekularathna, D. K., Manage, A. B., and Scariano, S. M. (2019). Power analysis of several normality tests: A Monte Carlo simulation study. *Communications in Statistics - Simulation and Computation*, 757-773.
- Yap, B., and Sim, C. (2011). Comparisons of various types of normality tests. *Journal of Statistical Computation and Simulation* 81 (12), 2141–55. doi:10.1080/00949655.2010.520163.
- Yazici, B., and Yolacan, S. (2007). A comparison of various tests of normality. *Journal of Statistical Computation and Simulation* 77 (2), 175–83. doi:10.1080/10629360600678310.
- Yildirim, N. (2012). Bazı Normallik Testlerinin I Tip Hataları ve Güçleri Bakımından Kıyaslanması. *Süleyman Demirel Üniversitesi, Fen Bilimleri Enstitüsü Dergisi*, 109-115.
- Yürekli, A. (2019). Lisansüstü Eğitim Enstitüsü Tez Yazım Kılavuzu. *Eskişehir Teknik Üniversitesi Lisansüstü Eğitim Enstitüsü*, 1-12.
- Zhang, J. (2002). *Powerful goodness-of-fit tests based on the likelihood ratio*. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 64(2), 281-294.

ÖZGEÇMİŞ

ORCID NO : 0009-0002-4976-4813

Ad Soyad : Zahir Hajizada

Yabancı Dil : İngilizce

Eğitim ve Mesleki Geçmişi:

- 2020-2024, Eskişehir Teknik Üniversitesi, Fen Bilimleri Fakültesi, İstatistik Ana Bilim Dalı, İstatistik Bölümü, Tezli Yüksek Lisans.
- 2018-2020, Anadolu Üniversitesi, Sosyal Bilimler Enstitüsü, Para Banka bölümü, Tezsiz Yüksek Lisans.
- 2014-2018, Azerbaycan Devlet İktisat Üniversitesi, SABAH grupları, Muhasebe ve denetim bölümü, Lisans.

Yayınları ve/veya Bilimsel/Sanatsal Faaliyetleri:

- 2022, 3-6 Kasım - 6. Uluslararası Araştırmacılar, İstatistikçiler ve Genç İstatistikçiler Sempozyumu, Bildirimi Sunumu, Antalya.