

Dizel Motorlarda Silindir İçi Basıncı Tahmini için Veri Odaklı Makine Öğrenimi Modellerinin Performans Analizi

Ahmet KARAOĞLU¹ , Hüseyin SÖYLER^{*2} 

¹ Sinop Üniversitesi, Mühendislik Fakültesi, Bilgisayar Mühendisliği Bölümü, 57000, Sinop, Türkiye

² Sinop Üniversitesi, Gerze Meslek Yüksek okulu, Ulaştırma Hizmetleri, 57600, Sinop, Türkiye

(Alınış / Received: 20.01.2025, Kabul / Accepted: 16.06.2025, Online Yayınlanma / Published Online: 25.08.2025)

Anahtar Kelimeler

Dizel motor silindir
içi basınç tahmini,
XGBoost,
Karar Ağacı,
Random Forest,
Makine öğrenmesi

Öz: Bu çalışma, dizel motorlarda silindir içi basınç tahmini için veri odaklı makine öğrenimi yaklaşımlarının performansını karşılaştırmayı amaçlamaktadır. Krank açısı ve yük değişkenlerine dayalı bir veri seti kullanılarak, Random Forest, Karar Ağacı ve XGBoost algoritmaları değerlendirilmiştir. Modellerin doğruluk oranları, işlem süreleri ve hata metrikleri detaylı bir şekilde analiz edilmiştir. Sonuçlar, Random Forest ($R^2 = 0.9999$) modelinin genelleme başarısı ve düşük hata oranları ile en iyi performansı sergilediğini göstermiştir. Karar Ağacı modeli de benzer şekilde yüksek doğruluk ($R^2 = 0.9999$) sunmasına rağmen genelleme yeteneği sınırlı kalmıştır. XGBoost modeli ise hızlı tahmin yeteneği ile dikkat çekerken, $R^2 = 0.9990$ ile hafif doğruluk kayıpları sergilemiştir. Bu modeller, emisyon kontrolü ve motor verimliliğini artırmaya yönelik çalışmalarda kullanılabilir güçlü tahmin araçları sunmaktadır. Elde edilen bulgular, veri odaklı yaklaşımların, dizel motorların performans ve emisyon analizinde güvenilir ve etkili bir alternatif sunduğunu ortaya koymaktadır.

Performance Analysis of Data-Driven Machine Learning Models for In-Cylinder Pressure Prediction in Diesel Engines

Keywords

Diesel engine in-cylinder
pressure prediction,
Artificial neural networks,
XGBoost,
Random Forest,
Machine learning

Abstract: This study aims to compare the performance of data-driven machine learning approaches for in-cylinder pressure prediction in diesel engines. Using a dataset based on crank angle and load variables, Random Forest, Decision Tree and XGBoost algorithms are evaluated. Accuracy rates, processing times and error metrics of the models are analysed in detail. The results showed that the Random Forest model ($R^2 = 0.9999$) performed the best with generalisation success and low error rates. Similarly, the Decision Tree model also showed high accuracy ($R^2 = 0.9999$), but its generalisation ability was limited. The XGBoost model, on the other hand, was notable for its fast prediction capability, but exhibited slight accuracy losses with $R^2 = 0.9990$. These models provide powerful predictive tools that can be used in emission control and engine efficiency improvement studies. The findings suggest that data-driven approaches offer a reliable and effective alternative for performance and emission analysis of diesel engines.

1. Giriş

Dizel motorlar, günümüzde özellikle kara ve deniz taşımacılığında yaygın olarak kullanılmakta olup, enerji verimliliği ve dayanıklılık açısından önemli bir rol oynamaktadır. Bu motorların performans analizi, yanma odasında oluşan silindir basıncı verilerinin detaylı incelenmesiyle sağlanmaktadır [1]. Silindir

basıncı, motorun verimliliğini, yanma etkinliğini ve emisyon özelliklerini doğrudan etkileyen kritik bir parametredir. Bu bağlamda, krank açısına bağlı olarak silindir içinde oluşan basıncı tahmin etmek, yanma sürecinin dinamiklerinin anlaşılmasında önemli bir adımdır.

Silindir içi basıncın doğru tahmini, motor performansını optimize etmek, emisyonları azaltmak ve dizel motorlarda gelişmiş yanma kontrol stratejileri uygulamak için çok önemlidir. Silindir içi basınç doğrudan sensörlerle ölçülebilse de bu ölçüm yöntemi çeşitli sınırlamalara sahiptir. Yüksek çözünürlüklü basınç sensörleri hem yüksek maliyetlidir hem de motorun çalışması esnasında fiziksel olarak zor koşullarda kalmaları nedeniyle zamanla bozulabilmektedir. Ayrıca, her silindire ayrı bir sensör yerleştirmek çoğu durumda ekonomik ve yapısal olarak mümkün değildir. Özellikle çok sayıda motorun aynı anda izlenmesi gereken test ortamlarında bu yaklaşım pratikliğini yitirmektedir. Bu bağlamda, krank açısı ve yük gibi doğrudan ölçülebilir parametreler kullanılarak silindir basıncının makine öğrenmesi ile tahmin edilmesi, hem gerçek zamanlı izleme hem de maliyet etkinliği açısından güçlü bir alternatif sunmaktadır [2].

Araştırmacılar ek sensörler olmadan silindir içi basıncı tahmin etmek için çeşitli dolaylı yöntemler geliştirmişlerdir [3]. Bu yaklaşımlar arasında krank mili kinematikini kullanan model tabanlı teknikler, ivmeölçer ölçümleri ve yarı ampirik modeller yer almaktadır. Krank mili tabanlı yöntemler, ölçülen krank açısı ve tek bir basınç sensöründen silindir basınçlarını yeniden yapılandırmak için esnek krank mili modellerini ve adaptasyon algoritmalarını kullanır [2], [3], [4], [5]. İvmeölçer tabanlı teknikler, basıncı tahmin etmek için ivme sinyalinin yanma ve yanma dışı bileşenlerini ayırır. Yarı ampirik modeller, basınç izlerini tahmin etmek için politropik sıkıştırma varsayımlarını ısı salınım fonksiyonları ile birleştirir.

Kademeli enjeksiyon, yakıtın tek seferde değil, farklı zaman dilimlerinde püskürtülmesiyle yanma sürecini kontrol altına almayı amaçlayan bir tekniktir. Ancak bu strateji, tahmin modellerinin doğruluğunu zorlaştırmaktadır. Emisyon yönetmelikleri daha katı hale geldikçe, sağlam silindir içi basınç tahmini dizel motor performansını ve verimliliğini optimize etmede giderek daha önemli bir rol oynayacaktır [6].

Son zamanlarda yapılan araştırmalar, kalibrasyon ve hesaplama gereksinimleriyle ilgili zorlukları ele alarak motorlarda silindir içi basıncı tahmin etmek için geleneksel termodinamik modellere alternatifler araştırmıştır. Veri odaklı yaklaşımlar, özellikle de yapay sinir ağları (YSA'lar), silindir basıncını doğru bir şekilde tahmin etme konusunda umut vaat etmektedir. Bu modeller, belirli parametreler için doğruluk ve hesaplama süresi açısından termodinamik simülasyonlardan daha iyi performans gösterebilir [7], [8]. Bununla birlikte, termodinamik tabanlı modeller, bazılarının çeşitli motor parametrelerine karşı hassasiyetler içermesiyle birlikte geçerliliğini korumaktadır [9]. Emisyon tahmini için silindir içi basınç ölçümlerini kullanan gerçek zamanlı uygulanabilir modeller geliştirilmiştir [10]. Araştırmacılar ayrıca simülasyonların otomatik kalibrasyonu için yöntemler ve parametre tahmin

stratejileri önermişlerdir [11], [12]. Çevrimiçi pilot kütle tahmini için sanal sensörler ve basınç tahmini için derin öğrenme yaklaşımları gibi gelişmiş teknikler bu alanda devam eden gelişimi göstermektedir [13], [14].

Son çalışmalar, motorlarda silindir basıncını tahmin etmek için doğruluk ve hesaplama verimliliğinde potansiyel iyileştirmeler sunan veri odaklı yaklaşımları araştırmıştır. Regresyon modelleri, sinir ağları ve zaman serisi analizleri dahil olmak üzere çeşitli makine öğrenimi teknikleri uygulanmıştır. Yapay Sinir Ağları (YSA), silindir basıncını yüksek doğrulukla tahmin etmede umut verici sonuçlar göstermiştir [15], [16]. Ayrıca makine öğrenimine dayalı regresyon modelleri, mühendislik alanlarının yanı sıra, tıp gibi farklı disiplinlerde de başarılı sonuçlar vermektedir [17], [18]. Uzun-Kısa Süreli Bellek (LSTM) modelleri, silindir basıncının zaman serisi tahmininde etkili olduğunu göstermiştir [19]. Frekans Tepki Fonksiyonları ve Evrimsel Sinir Ağları (CNN) sırasıyla krank mili hızı dalgalanmalarından ve basınç ölçümlerinden silindir basıncını tahmin etmek için kullanılmıştır [20], [21]. Zaman gecikmeli sinir ağları, krank kinematığı veya blok titreşim ölçümleri kullanılarak silindir basıncını yeniden yapılandırmak için uygulanmıştır. Bu veri odaklı yaklaşımlar, dizel motorları, HCCI motorları ve doğal gaz motorları dahil olmak üzere çeşitli motor türlerinde potansiyel göstermiştir [22], [23], [24].

Yapılan literatür incelemeleri, dizel motorlarda silindir içi basınç tahmininin motor performansı ve emisyon kontrolü açısından önemini ortaya koymuştur. Geleneksel termodinamik modeller, basınç tahmininde sıkça kullanılmakla birlikte, yüksek hesaplama maliyetleri ve kalibrasyon zorlukları nedeniyle veri odaklı yaklaşımlar alternatif bir çözüm sunmaktadır. Bu bağlamda, makine öğrenmesi yöntemleri, silindir basıncını doğrudan ölçüm yapmadan yüksek doğrulukla tahmin etme potansiyeline sahiptir. Literatürde Random Forest, Karar Ağacı ve XGBoost gibi makine öğrenimi algoritmaları, motor performans analizi ve emisyon kontrolü için umut verici sonuçlar vermiştir [24], [25], [26], [27].

Bu çalışma, söz konusu tahminlerin doğruluk düzeyini ortaya koymak ve farklı modellerin bu amaçla kullanılabilirliğini değerlendirmek üzere gerçekleştirilmiştir. Dizel motorlarda silindir içi basıncı tahmin etmek için farklı makine öğrenmesi modellerinin doğruluğunu karşılaştırarak bu modellerin avantajlarını ve kısıtlılıklarını belirlemektir. Bu doğrultuda, farklı deneysel koşullarda elde edilen silindir basıncı verileri kullanılarak Random Forest, Karar Ağacı ve XGBoost algoritmalarının performansı doğruluk oranları, işlem süreleri ve hata analizleri üzerinden değerlendirilmiştir. Veri toplama süreci, model eğitiminde kullanılan parametreler ve performans değerlendirme kriterleri titizlikle belirlenmiş ve

deney sonuçları karşılaştırmalı analizlerle incelenmiştir. İzlenen yöntem ve analiz süreçleri aşağıdaki materyal ve metot bölümünde detaylı olarak açıklanmaktadır.

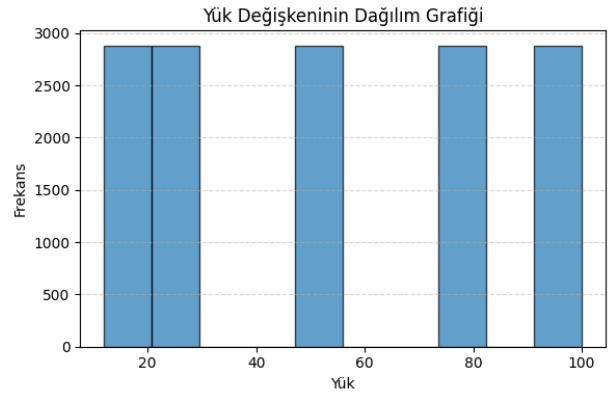
2. Materyal ve Metot

Bu çalışmada, farklı makine öğrenimi algoritmalarının tahmin performanslarını değerlendirmek amacıyla kapsamlı bir analiz gerçekleştirilmiştir. Analiz sürecinde, veri ön işleme, model eğitimi, hiperparametre ayarları ve model değerlendirmesi adımları detaylı bir şekilde ele alınmıştır. Çalışmada kullanılan yöntemler, verilerin yapılandırılmasından model performanslarının görselleştirilmesine kadar disiplinlerarası standartlara uygun bir yaklaşımla uygulanmıştır. Aşağıda, kullanılan algoritmaların özellikleri, performans değerlendirme metrikleri ve sonuçların görsel karşılaştırmalarına ilişkin detaylar sunulmaktadır.

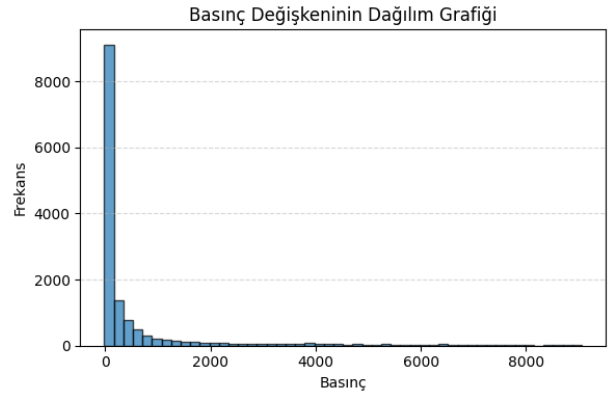
2.1. Veri seti ve ön işleme

Bu çalışmada, [28] tarafından sağlanan, 4 zamanlı common rail enjeksiyon sistemine sahip 6,8 L turboşarjlı 6 silindirli Tier II bir dizel motora ait veriler kullanılmıştır. Veriler, ISO Standardı 8178:4 D2 döngüsüne uygun olarak 1800 rpm'de gerçekleştirilen beş test çalışmasını kapsamaktadır. Bu veri seti, çalışmamızda kullanılan makine öğrenimi modellerinin geliştirilmesi ve değerlendirilmesi için temel oluşturmuştur. Makine öğrenmesi algoritmaları, büyük ve çeşitli veri kümelerini işleyebilme yetenekleri sayesinde motor verisi gibi kompleks yapıları sistemlerde başarılı tahminler sağlamaktadır [29].

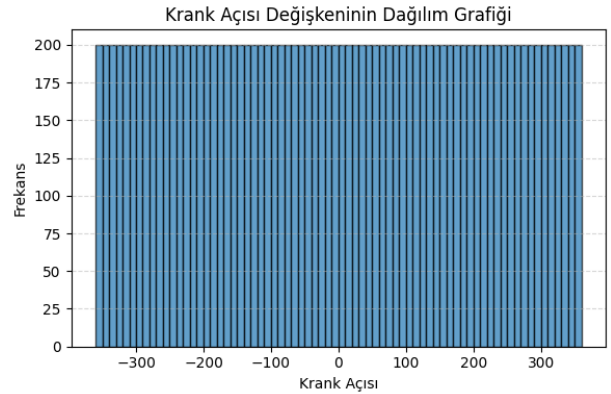
Veri seti, Excel formatında yüklenmiştir ve Tablo 1'de örnek satırları gösterilmiştir. İlk iki sütun giriş değişkenleri (X), üçüncü sütun ise çıkış değişkeni (y) olarak belirlenmiştir. Veri setinin uygunluğu, eksik veri kontrolü ve genel bilgi incelemesi ile doğrulanmıştır. Eksik veri kontrolü sırasında, veri setindeki tüm sütunlar gözden geçirilmiş ve veri bütünlüğü sağlanmıştır. Eksik veya hatalı veriler bulunmadığı doğrulandıktan sonra normalleştirme adımına geçilmiştir. Bu aşamada, giriş ve çıkış değişkenleri MinMaxScaler kullanılarak 0 ile 1 arasında normalize edilmiştir. Bu işlem, farklı ölçeklerdeki değişkenlerin model performansı üzerinde olumsuz bir etkisi olmaması için uygulanmıştır.



(a)



(b)



(c)

Şekil 1. Veri setinin dağılım grafikleri. (a) Yük, (b) Basınç, (c) Krank açısı.

Bu görselleştirme, değişkenlerin model eğitimi sırasında homojen bir şekilde kullanılacağını ve veri kaybı veya tutarsızlık olmadığını vurgulamaktadır. Şekil 1'de görülen yapı, modelleme sürecinde değişkenlerin özellikleri hakkında kapsamlı bir bilgi sunmuştur.

Veri normalleştirme sonrası, giriş (X) ve çıkış (y) değişkenlerinin dağılımları incelenmiştir. Eğitim, doğrulama ve test setlerine ayrılmadan önce verilerin homojen dağılımda olduğu gözlemlenmiştir. Şekil 1'de görüldüğü üzere basınç değişkeni için histogram grafiği incelendiğinde, verinin büyük bir kısmının düşük değerler arasında yoğunlaştığı gözlemlenmiştir. Basınç değişkeni uzun bir sağ kuyruk

sergilemektedir. Bu da yüksek basınç değerlerinin nadiren ölçüldüğünü göstermektedir. Bu durum, model tahmin performansı üzerinde etkili olabilir ve gerekirse veri dönüşüm teknikleri ile ele alınabilir.

Tablo 1. Veri setinden örnek veriler.

	Yük (%)	Krank Açısı	Silindir Basıncı
1	12	-360.00	90.759
2	12	-359.75	89.577
3	12	-359.50	89.746
4	12	-359.25	89.302

Krank açısı değişkeninin dağılım grafiği ise açılar arasında eşit bir dağılım sergilemektedir. Bu durum, krank açısının veri setinde eksiksiz bir şekilde ölçüldüğünü göstermekte ve modelin bu değişkenle sağlıklı bir ilişki kurabileceğine işaret etmektedir. Yük değişkeni için histogramda ise farklı yük seviyelerinin (örneğin 20, 40, 60, 80 ve 100) eşit sıklıkta yer aldığı görülmektedir. Bu durum, veri setinin farklı yük koşullarını kapsadığı ve modelin bu yüklerle iyi bir şekilde eğitilebileceği anlamına gelir.

Tablo 2. Veri setinin istatistiksel verileri.

İstatistik	Yük	Krank açısı (°)	Basınç (Bar)
Gözlem Sayısı (Count)	14400.000	14400.000	14400.000
Ortalama (Mean)	52.400	-0.125	670.614
Standart Sapma (Std)	32.142	207.853	1502.922
En Küçük Değer (Min)	12.000	-360.000	-20.095
1. Çeyrek (Q1, %25)	25.000	-180.062	40.545
Medyan (Q2, %50)	50.000	-0.125	80.972
3. Çeyrek (Q3, %75)	75.000	179.813	409.301
En Büyük Değer (Max)	100.000	359.750	9058.385

Bu dağılım analizleri, veri setindeki dengesizliklerin ve homojenliklerin belirlenmesine yardımcı olmuş ve modelleme sürecinde dikkate alınmıştır. Tablo 2'deki istatistiksel özet, veri setinin boyutlarını ve dağılımlarını göstermektedir.

Tabloda yer alan istatistiksel ölçütler, veri setinin dağılım özelliklerini özetlemektedir. Gözlem Sayısı (Count), her değişkene ait toplam veri sayısını belirtir. Ortalama (Mean), tüm gözlemlerin aritmetik ortalamasını ifade eder. Standart Sapma (Std), verilerin ortalama etrafındaki yayılım derecesini gösterir. Büyük sapma, verinin daha geniş dağıldığını ifade eder. En Küçük Değer (Min) ve En Büyük Değer (Max), sırasıyla veri setindeki minimum ve maksimum

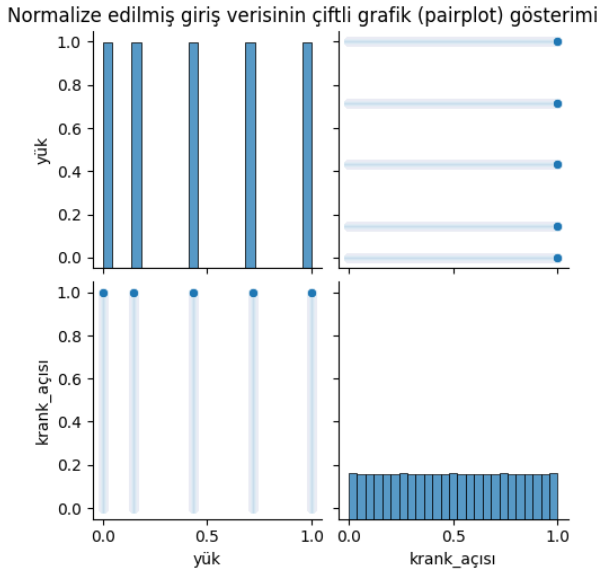
değerleri temsil eder. 1. Çeyrek (%25, Q1), verinin en küçükten büyüğe doğru sıralandığında ilk %25'lik kısmının üst sınırıdır. Medyan (%50, Q2), ortadaki değeri temsil eder ve veriyi iki eşit parçaya böler. 3. Çeyrek (%75, Q3) ise verinin %75'inin altında, %25'inin üstünde kalan değeri belirtir. Bu çeyrek değerler, özellikle veri setindeki asimetrik dağılımların ve aykırı değerlerin tespitinde önemli rol oynar.

Tablo 2 incelendiğinde yük değişkeni 12 ile 100 arasında değerler almakta ve ortalaması yaklaşık 52.4'tür. Krank açısı değişkeni -360° ile 360° arasında eşit dağılımlı bir değişken olup ortalama açısı -0.125° civarındadır. Basınç değişkeni ise -20.1 ile 9058,4 arasında geniş bir aralıkta dağılmıştır ve ortalama basınç değeri 670.6 bar olarak gözlemlenmiştir. Bu sonuçlar, özellikle basınç değerlerinde standart sapmanın yüksek olması nedeniyle veri setinin uç değerlere duyarlı olabileceğini göstermektedir. Verideki minimum ve maksimum basınç değerleri, bazı anomalilerin veya yüksek basınç koşullarının veri setine dahil edildiğini göstermektedir. Bu durum, model eğitiminde ve değerlendirilmesinde dikkat edilmesi gereken bir husus olarak ele alınmıştır.

Şekil 2'de, normalleştirilmiş yük ve krank açısı değişkenleri arasındaki dağılım grafiği detaylı bir şekilde sunulmaktadır. Bu çift değişkenli dağılım grafiği, veri setindeki değişkenlerin normalizasyon işleminden sonra hangi aralıklarda yer aldığını ve dağılım düzenlerini açıkça ortaya koymaktadır. Normalizasyon süreci, değişkenler arasındaki ölçek farklılıklarını ortadan kaldırmak ve model eğitimi sırasında her bir değişkenin eşit bir katkıda bulunmasını sağlamak için uygulanmıştır. Grafik, yük değişkeninin belirli seviyelerde kümелendiğini ve farklı motor yük koşullarını temsil ettiğini göstermektedir. Bu kümelenme, veri setinin farklı yük koşullarını kapsadığını ve motor performans analizinde çeşitliliğin sağlandığını işaret etmektedir. Bunun yanı sıra, krank açısı değişkeni krank açısının tam döngüsünü kapsamakta olup, -360° ile +360° arasında eşit bir dağılım sergilemektedir. Bu durum, krank açısına bağlı olarak silindir içi basınç değişimlerinin tam bir döngü boyunca ölçüldüğünü ve analiz için gereken sürekliliğin sağlandığını göstermektedir. Ayrıca, grafikte her bir değişkenin ayrı ayrı normalleştirilmiş değerlerini koruduğu ve veri çeşitliliğinin uygun bir şekilde temsil edildiği görülmektedir.

Bu analiz, veri setinin homojen bir dağılım sergilediğini ve modelleme sürecinde herhangi bir dengesizlik veya eksiklik olmadığını doğrulamaktadır. Grafikte gözlemlenen düzen, modelin eğitim sırasında hem yük hem de krank açısı değişkenlerinden faydalanarak doğru ilişkiler kurabileceğini ve silindir içi basınç tahmini için güçlü bir temel oluşturabileceğini göstermektedir. Böylece, veri setindeki bu düzenli yapı, modelin genelleme

başarısını artırma potansiyeline sahip olduğunu işaret etmektedir.



Şekil 2. Normalleştirilmiş yük ve krank açısı değişkenleri

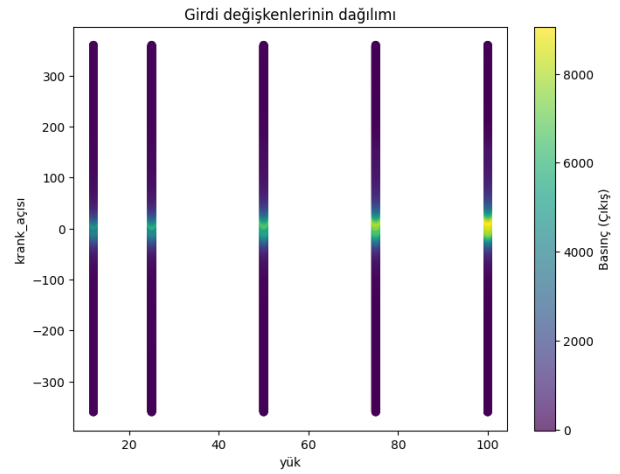
2.2 Veri setinin ayrılması

Normalizasyon işleminin ardından, veri seti %70 eğitim, %15 doğrulama ve %15 test olarak üç alt kümeye bölünmüştür. Veri bölme işlemi sırasında, modelin genellenabilirliğini artırmak amacıyla rastgele veri dağılımı için random_state=42 parametresi kullanılmıştır. Veri setinin eğitim, doğrulama ve test kümelerine ait örnek büyüklükleri aşağıda Tablo 3'te verilmiştir.

Tablo 3. Veri setininne ait örnek büyüklükleri

Veri Kümesi	Örnek Sayısı	Açıklama
Eğitim (X_train, y_train)	9792	Modelin öğrenme sürecinde kullanılan veri kümesi
Doğrulama (X_val, y_val)	2448	Modelin hiperparametre ayarlamasında kullanılan doğrulama kümesi
Test (X_test, y_test)	2160	Modelin performans değerlendirmesi için ayrılmış, daha önce görülmemiş veri kümesi

Bu yapılandırma, veri setinin eğitim sırasında aşırı öğrenme (overfitting) veya yetersiz öğrenme (underfitting) problemleri yaşamaması için dengeli bir şekilde bölünmesini sağlamıştır. Model performansı, doğrulama ve test kümeleri üzerinde hesaplanarak genel sonuçların değerlendirilebilirliği artırılmıştır.



Şekil 3. Girdi değişkenlerinin dağılım grafiği

Şekil 3, çalışma kapsamında kullanılan veri setinin giriş değişkenleri arasındaki ilişkiyi görselleştirmektedir. Girdi değişkenleri, krank açısı (°) ve yük (%) olarak belirlenmiş, çıkış değişkeni ise basınç (Cyl kPa) değeridir. Yatay ekseninde yük değerleri (%), dikey ekseninde krank açıları (°) yer almaktadır. Renk skalası, her bir gözlem için basınç değerlerini temsil etmekte olup, yüksek basınç değerleri daha açık tonlarla gösterilmiştir. Veri seti, farklı yük seviyelerinde motor performans parametrelerinin detaylı şekilde ölçülmesi amacıyla toplanmıştır. Her yük seviyesinde krank açısı -360° ile +360° arasında değişmekte ve yanma döngüsünün tamamını kapsamaktadır. Grafik, veri setindeki ölçüm yoğunluğunun farklı krank açıları için sabit yük seviyeleri için homojen dağılım gösterdiğini ortaya koymaktadır. Bu yapılandırma, basınç değişimlerinin krank açısı ve yük ile ilişkisini analiz edebilmek için veri setinin kapsamlı bir şekilde düzenlendiğini göstermektedir. Böylece, modelleme ve analiz süreçlerinde kullanılan veri seti, farklı motor çalışma koşullarını temsil eden ölçüm değerlerinden oluşmakta ve her bir yük seviyesi için krank açısına bağlı basınç ölçümleri detaylı bir şekilde kayıt altına alınmaktadır.

2.3. 5-katmanlı (5-fold) çapraz doğrulama (cross-validation)

Bu çalışmada, model performansının daha güvenilir ve nesnel bir biçimde değerlendirilmesi amacıyla 5-katmanlı (5-fold) çapraz doğrulama yöntemi kullanılmıştır. Veri seti, rastgele karıştırılarak beş alt kümeye bölünmüş ve her bir fold için modeller ayrı ayrı eğitilip doğrulama verisi üzerinde test edilmiştir. Her bir iterasyonda modeller, eğitim alt kümesinde GridSearchCV yöntemi ile optimize edilmiş hiperparametreler ile eğitilmiş, doğrulama verisinde en yüksek R^2 değerini veren yapılandırma kaydedilmiştir.

Tablo 4. Model Katmanlarının Yapısı

Bileşen	Açıklama
Model Tipi	Decision Tree, Random Forest, XGBoost
Girdi Özellikleri	2 (örneğin: sıcaklık, basınç)
Çıktı	1 (örneğin: verim, emisyon)
Doğrulama Yöntemi	5-katmanlı (5-fold) çapraz doğrulama
Test Verisi	%15 oranında, tamamen ayrılmış ve eğitimde hiç kullanılmamış
Hiperparametre Seçimi	Her fold'da GridSearchCV ile en iyi parametre kombinasyonu seçildi
Performans Metrikleri	MAE, RMSE, R ²

Nihai olarak her model için en başarılı fold'da eğitilen yapı, daha önce eğitimde hiç kullanılmamış %15'lik test veri seti üzerinde değerlendirilmiştir. Model başarımı, ortalama mutlak hata (Mean Absolute Error - MAE), kök ortalama kare hata (Root Mean Squared Error - RMSE) ve determinasyon katsayısı (R²) ile ölçülmüştür. Kullanılan makine öğrenmesi yöntemleri ve değerlendirme yapısı Tablo 4'te sunulmuştur.

2.4. Regresyon modelleri

Bu çalışmada, regresyon problemini çözmek için XGBoost Regressor, Random Forest Regressor, Decision Tree Regressor olarak üç farklı model kullanılmıştır. Bu çalışmada hiperparametre seçimi, her bir makine öğrenmesi algoritmasının yapısal özellikleri ve genelleme kapasitesi dikkate alınarak, sistematik bir tarama yaklaşımı ile gerçekleştirilmiştir. Her fold içerisinde GridSearchCV yöntemi kullanılarak belirli hiperparametre aralıkları üzerinden kapsamlı bir arama yapılmış ve en iyi çapraz doğrulama performansını sağlayan parametre kombinasyonu belirlenmiştir. Hiperparametre aralıkları, ilgili algoritmaların aşırı öğrenmeye yatkınlıklarını dengeleyecek şekilde, özellikle ağaç tabanlı modellerde derinlik, örnek sayısı ve budama gibi yapısal sınırlayıcılar gözetilerek oluşturulmuştur. XGBoost modeli için ise L1 (reg_alpha) ve L2 (reg_lambda) düzenleme katsayıları gibi yerleşik kontrol mekanizmaları da arama kapsamına dâhil edilmiştir.

Parametre seçim süreci sadece doğrulama metriklerine değil; aynı zamanda kalıntı analizi, parite grafikleri ve hata dağılımı gibi görsel destekleyici unsurlarla da değerlendirilmiştir. Nihayetinde, her model için en iyi sonuç veren fold'daki yapılandırma, daha önce eğitimde hiç kullanılmamış %15'lik test veri seti üzerinde test edilerek, modellerin hem öğrenme hem de genelleme kapasiteleri nesnel biçimde analiz edilmiştir. Bu bütüncül yaklaşım, sadece doğruluk metriklerine değil, model davranışlarına da odaklanan çok yönlü bir değerlendirme süreci sunmuştur.

2.4.1. XGBoost

XGBoost (Extreme Gradient Boosting), gradyan artırmalı ağaç (gradient boosting tree) algoritmalarının optimize edilmiş ve düzenleme mekanizmaları entegre edilmiş bir sürümüdür.

Tablo 5. XGBoost Regressor model parametreleri.

Parametre	Değer	Açıklama
n_estimators	100	Toplam zayıf model sayısı
max_depth	6	Her bir ağacın maksimum derinliği
learning_rate	0.1	Öğrenme oranı
subsample	1	Her ağaç için kullanılacak örnek oranı
colsample_bytree	1	Her ağaçta kullanılacak özellik oranı
gamma	0	Bölünme için gereken minimum kayıp azaltımı
reg_alpha	0	L1 düzenleme katsayısı
reg_lambda	2	L2 düzenleme katsayısı
random_state	42	Rastgelelik kontrolü

Model, hata fonksiyonunu minimize etmeyi amaçlayan art arda ağaçlar inşa eder ve her yeni ağaç, önceki ağaçların yapamadığı tahminleri iyileştirmeye çalışır. XGBoost, L1 (lasso) ve L2 (ridge) regularizasyon terimlerini içeren kayıp fonksiyonu sayesinde overfitting riskine karşı oldukça dayanıklıdır. Ayrıca yüksek hız, paralelleştirilebilirlik ve eksik veri toleransı gibi avantajlara sahiptir. XGBoost Regressor algoritması model parametreleri Tablo 5'te gösterilmiştir.

2.4.2. Karar Ağacı (Decision Tree)

Bu çalışmada bir diğer yöntem olarak, tahminleme sürecinde Karar Ağacı Regressor (Decision Tree Regressor) algoritması kullanılmıştır. Karar ağacı algoritması, veriyi ardışık biçimde bölerek tahmin işlemini gerçekleştiren sezgisel bir öğrenme yöntemidir.

Ağaç yapısı, düğümler aracılığıyla karar kuralları oluşturur ve her yaprak düğümde bir çıktı değeri üretilir. Regresyon amaçlı kullanımında, her bir bölmede varyansı minimize edecek şekilde veri alt kümelerine ayrılır. Karar ağaçları, yorumlanabilirliği yüksek ve hızlı uygulanabilir yapılar sunmasına rağmen, aşırı dallanma durumunda genelleme kabiliyetini kaybederek overfitting'e yatkın hâle gelebilir. Bu nedenle, derinlik ve örnek sayısı gibi hiperparametrelerin dikkatle sınırlandırılması gerekmektedir. Model parametreleri Tablo 6'da gösterilmektedir.

Tablo 6. Karar Ağacı Regressor (Decision Tree Regressor) model parametreleri.

Parametre	Değer	Açıklama
criterion	'squared_error'	Hata fonksiyonu (MSE)
splitter	'best'	En iyi bölmeyi arar
max_depth	None	Maksimum derinlik sınırlaması yok
min_samples_split	2	Bölme için gerekli minimum örnek sayısı
min_samples_leaf	1	Yaprakta bulunması gereken minimum örnek sayısı
max_features	None	Özellik alt kümesi sınırı yok
random_state	42	Sonuçların tekrarlanabilirliğini sağlar
criterion	'squared_error'	Hata fonksiyonu (MSE)

2.4.3. Rastgele Orman (Random Forest)

Bu çalışmada, tahminleme sürecinde bir diğer yöntem olarak Random Forest Regressor algoritması kullanılmıştır.

Tablo 7. Random Forest Regressor model parametreleri.

Parametre	Değer	Açıklama
n_estimators	100	Kullanılan karar ağacı sayısı
criterion	'squared_error'	Hata hesaplama yöntemi (MSE)
max_depth	None	Her ağacın maksimum derinliği
min_samples_split	2	Bir iç düğümün bölünmesi için gereken minimum örnek
min_samples_leaf	1	Bir yaprakta olması gereken minimum örnek sayısı
max_features	'auto'	Her ağaç için kullanılacak maksimum özellik sayısı
bootstrap	True	Örnekleme yöntemi: geri dönüşümlü
random_state	42	Sonuçların tekrarlanabilirliği

Random Forest algoritması, çok sayıda karar ağacından oluşan bir topluluk öğrenme (ensemble learning) yöntemidir. Her bir ağaç, eğitim verisinin rastgele bir alt kümesi ve özelliklerin rastgele bir alt kümesi ile eğitilir. Bu yaklaşım, karar ağaçlarının varyansını azaltır ve modelin genelleme kapasitesini artırır. Rastgele ormanlar, gürültülü verilerle başa çıkabilme kabiliyeti ve aşırı öğrenmeye karşı dirençli

yapısıyla öne çıkar. Ancak, yorumlanabilirliği tekil karar ağaçlarına göre daha düşüktür. Model parametreleri Tablo 7'de gösterilmiştir.

2.5. Performans metrikleri

Model, eğitim kümesindeki verilerle eğitilmiş ve test kümesi kullanılarak değerlendirilmiştir. Performans metrikleri olarak Tablo 8'de gösterildiği gibi determinasyon katsayısı (R^2), ortalama mutlak hata (MAE), ortalama kare hata (MSE) ve kök ortalama kare hata (RMSE) kullanılmıştır.

Model performansını daha bütüncül değerlendirebilmek adına hem Ortalama Mutlak Hata (MAE) hem de Kök Ortalama Kare Hata (RMSE) birlikte kullanılmıştır. MAE, tüm hataların mutlak değerlerinin ortalamasını alarak sapmaların genel büyüklüğünü ölçerken, RMSE hataların karesini alarak büyük sapmalara daha fazla ağırlık verir. Bu nedenle RMSE, uç değerlerin model üzerindeki etkisini daha belirgin şekilde yansıtırken; MAE, geneldeki tahmin hatalarının ortalama düzeyini verir. Her iki ölçüt birlikte değerlendirildiğinde, modelin hem genel doğruluğu hem de uç değerlere karşı hassasiyeti hakkında daha kapsamlı bir değerlendirme yapılabilmektedir. Bu durum, özellikle enerji tüketimi, emisyon tahmini veya maliyet öngörüsü gibi yüksek varyanslı çıktılarda önem arz etmektedir.

Tablo 8. Kullanılan performans metrikleri.

Performans Metrikleri	Açıklama
R^2 Score	Modelin açıklanabilirlik oranını temsil eden determinasyon katsayısı
MAE (Mean Absolute Error)	Tahmin edilen değerler ile gerçek değerler arasındaki ortalama mutlak fark
RMSE (Root Mean Squared Error)	Tahmin edilen değerler ile gerçek değerler arasındaki kök ortalama kare fark

Genelleme kapasitesi, bir modelin yalnızca eğitim verisi üzerinde değil, daha önce görmediği yeni veri üzerinde de benzer düzeyde başarı gösterebilme yeteneğini ifade eder. Bu çalışmada genelleme kapasitesi, eğitim+doğrulama süreci dışında bırakılan bağımsız %15'lik test veri seti üzerindeki performansla değerlendirilmiştir. Eğitim sürecinde elde edilen başarı metrikleri ile test sürecindeki metrikler karşılaştırılarak, modelin yeni veriler üzerindeki tutarlılığı ölçülmüştür. Özellikle eğitim doğruluğu ile test doğruluğu arasında belirgin bir fark olmaması, modelin aşırı öğrenmeden kaçındığını ve genelleme yeteneğinin güçlü olduğunu göstermektedir. Bu yaklaşım, genelleme kapasitesinin sayısal olarak değerlendirilmesini sağlar ve modelin pratik uygulamadaki güvenilirliğini destekler.

2.6. Sonuçların karşılaştırılması ve görselleştirilmesi

Modellerin performans karşılaştırması, normalize edilmiş performans metrikleri kullanılarak radar grafiği ile görselleştirilmiştir. Radar grafiği, her modelin her bir performans metriğinde sergilediği göreceli başarıyı aynı düzlemde incelemeyi mümkün kılmıştır ve tüm metrikler, farklı ölçeklerdeki değerlerin karşılaştırılabilmesi için normalize edilmiştir. Tahmin edilen değerlerin gerçek değerlere göre dağılımını gösteren parite grafikleri, modelin doğruluğunu değerlendirmek amacıyla kullanılmış ve ideal bir model için tüm veri noktalarının 45°'lik doğrultu üzerinde yer alması gerektiği vurgulanmıştır. Kalıntı (residual) dağılım grafikleri, tahmin değerleri ile gerçek değerler arasındaki farkların rastgele dağılımı durumunda modelin uygun bir tahmin performansı sergilediğini ortaya koymaktadır. Hata histogramları, tahmin hatalarının frekans dağılımını inceleyerek normal dağılıma benzer bir yapı olup olmadığını değerlendirmek için kullanılmıştır. Son olarak, radar grafiği, her modelin farklı performans metriklerine göre sergilediği göreceli başarısını bir arada sunarak model karşılaştırmalarının bütüncül bir şekilde incelenmesine olanak tanımıştır.

Modellerin performans metrikleri arasındaki istatistiksel farklılıkları değerlendirebilmek amacıyla, her bir modelin 5 katmanlı çapraz doğrulama sürecinde elde ettiği RMSE, MAE ve R^2 değerlerine tek yönlü varyans analizi (One-Way ANOVA) uygulanmıştır. Bu analizde, her metrik için modeller arası ortalama farkların anlamlı olup olmadığı test edilmiştir. ANOVA testi, özellikle aynı veri seti üzerinde eğitilen modellerin hata dağılımları arasında istatistiksel olarak anlamlı fark bulunup bulunmadığını ortaya koymak açısından önemli bir araç olarak tercih edilmiştir. Analizler, Python ortamında `scipy.stats.f_oneway` fonksiyonu kullanılarak gerçekleştirilmiştir. Elde edilen p-değerlerine göre, istatistiksel olarak anlamlı farklar gözlemlenen metriklerde model tercihinin yönelik güçlü çıkarımlar yapılmıştır.

2.7. Kullanılan yazılım ortamı ve araçlar

Bu çalışma Python programlama dili kullanılarak gerçekleştirilmiştir. Veri işleme ve modelleme için Scikit-learn, XGBoost ve Pandas kütüphaneleri; görselleştirme için Matplotlib ve Seaborn kütüphaneleri kullanılmıştır.

3. Bulgular

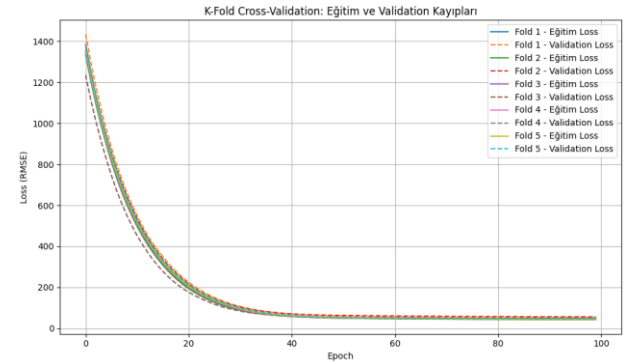
Bu çalışmada, kullanılan modeller, K-Fold çapraz doğrulama ($K=5$) ile değerlendirilmiştir. Eğitim süreci boyunca modelin doğrulama kaybı (validation loss), doğrulama ortalama mutlak hatası (MAE) ve R^2 skorları hesaplanmıştır.

K-Fold yöntemi kullanılarak yapılan değerlendirme sonuçlarında, her bir modelin fold süreçleri için eğitim ve doğrulama kayıp değerleri incelenmiştir. R^2 skorları, modelin doğrulama veri kümesi üzerindeki performansını doğrulamak için hesaplanmış olup, modelin doğruluk düzeyini temsil etmektedir. Modellerin K-Fold çapraz doğrulama sürecinde elde edilen değerleri aşağıdaki Tablo 9'da gösterilmiştir.

Tablo 9. K-Fold çapraz doğrulama değerleri.

Model	Fold	RMSE	MAE	R^2
XGBoost	1	51.803460	16.780338	0.998926
	2	55.567916	18.438903	0.998677
	3	45.995282	15.067042	0.998861
	4	46.904797	15.440563	0.999013
	5	47.080374	15.291440	0.999033
Random Forest	1	20.037200	5.036235	0.999839
	2	14.546379	4.697721	0.999909
	3	13.449586	4.519034	0.999903
	4	13.677057	4.160759	0.999916
	5	13.220504	4.214282	0.999924
Decision Tree	1	23.666829	7.454356	0.999776
	2	19.567332	7.113923	0.999836
	3	17.596448	6.413074	0.999833
	4	16.389079	5.823862	0.999880
	5	15.792187	6.005221	0.999891

XGBoost modelinin her bir fold için eğitim ve doğrulama kayıp değerlerinin epoch sayısına göre değişimi Şekil 4 ile görselleştirilmiştir. Grafiklerde, modelin erken durdurma (early stopping) kullanılarak aşırı öğrenme (overfitting) problemi minimize edildiği gözlemlenmiştir. Eğitim ve doğrulama kayıpları, modelin stabil bir hale ulaştığını göstermektedir.



Şekil 4. XGBoost modeli eğitim ve doğrulama kayıp grafiği.

Modelin farklı doğrulama kümeleri üzerindeki R^2 skorlarının yüksek olması, modelin doğruluk performansının tatmin edici olduğunu ortaya koymuştur.

MAE değerlerinin düşük olması, modelin tahmin ettiği değerlerle gerçek değerler arasındaki sapmanın az olduğunu ifade etmektedir. Model kaybının düşük olması ise modelin doğrulama verileri üzerindeki genelleme kapasitesinin yüksek olduğunu işaret etmektedir. Bu bulgular, önerilen modelin regresyon tabanlı problemler için uygun bir yapılandırma sunduğunu ve K-Fold çapraz doğrulama yaklaşımı ile

güvenilir tahminler sağladığını göstermektedir. Tablo 10'da modellerin doğruluk oranları ve hata değerleri verilmiştir.

Tablo 10. Modellerin doğruluk ve hata değerleri

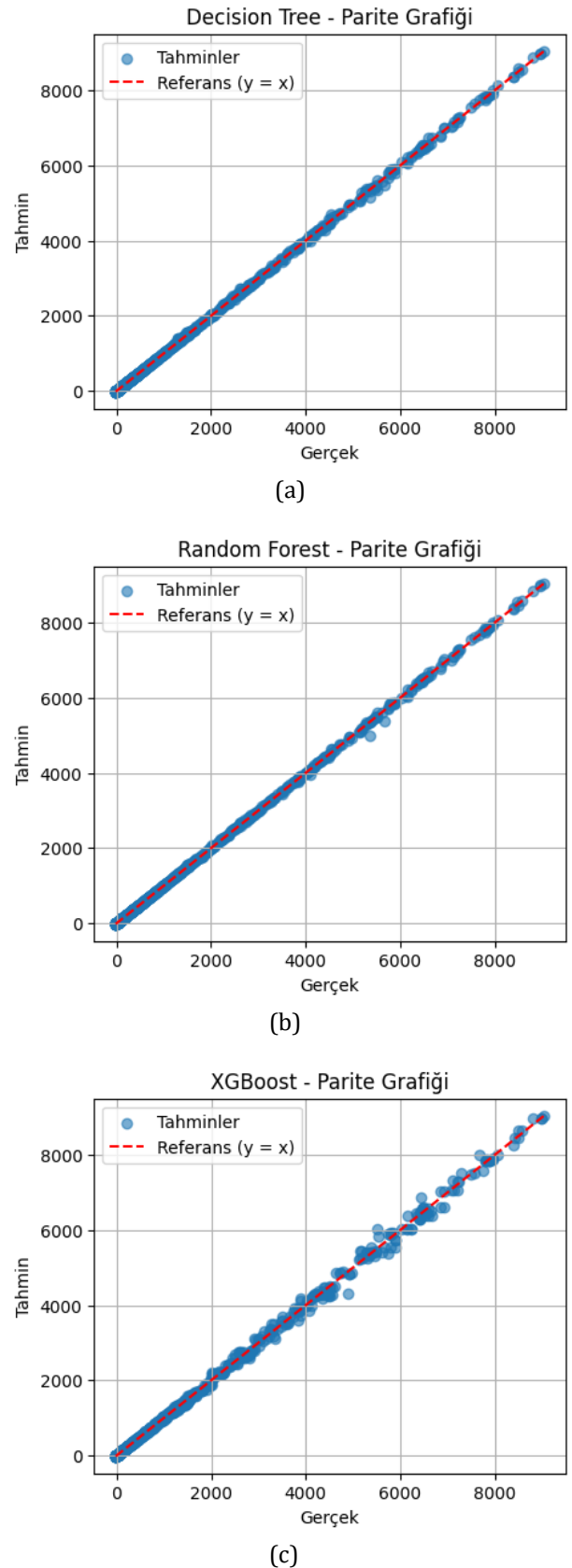
Model	R ² Score	MAE	RMSE
DecisionTreeRegressor	0.9999	6.639	17.785
RandomForestRegressor	0.9999	4.546	15.552
XGBoost	0.9990	16.212	47.394

Decision Tree modeli, test veri kümesi üzerinde neredeyse mükemmel bir performans sergilemiştir. Çok düşük MAE ve RMSE değerleri, modelin tahminlerinin gerçek değerlere oldukça yakın olduğunu göstermektedir. Random Forest modeli, test veri kümesi üzerinde Decision Tree modeline benzer şekilde mükemmel bir doğruluk sergilemiştir. Ancak, MAE ve RMSE değerleri daha düşük olup, tahmin başarısı açısından en iyi model olarak öne çıkmıştır. XGBoost modeli, diğer modellere göre hafif bir doğruluk düşüşü gösterse de hala mükemmel bir performans sunmuştur. MAE ve RMSE değerleri diğer modellere göre bir miktar daha yüksektir, ancak genelleme başarısı yüksektir. Tüm modellerin gerçek değerler ile tahmin edilen değerleri Tablo 11'de gösterilmiştir.

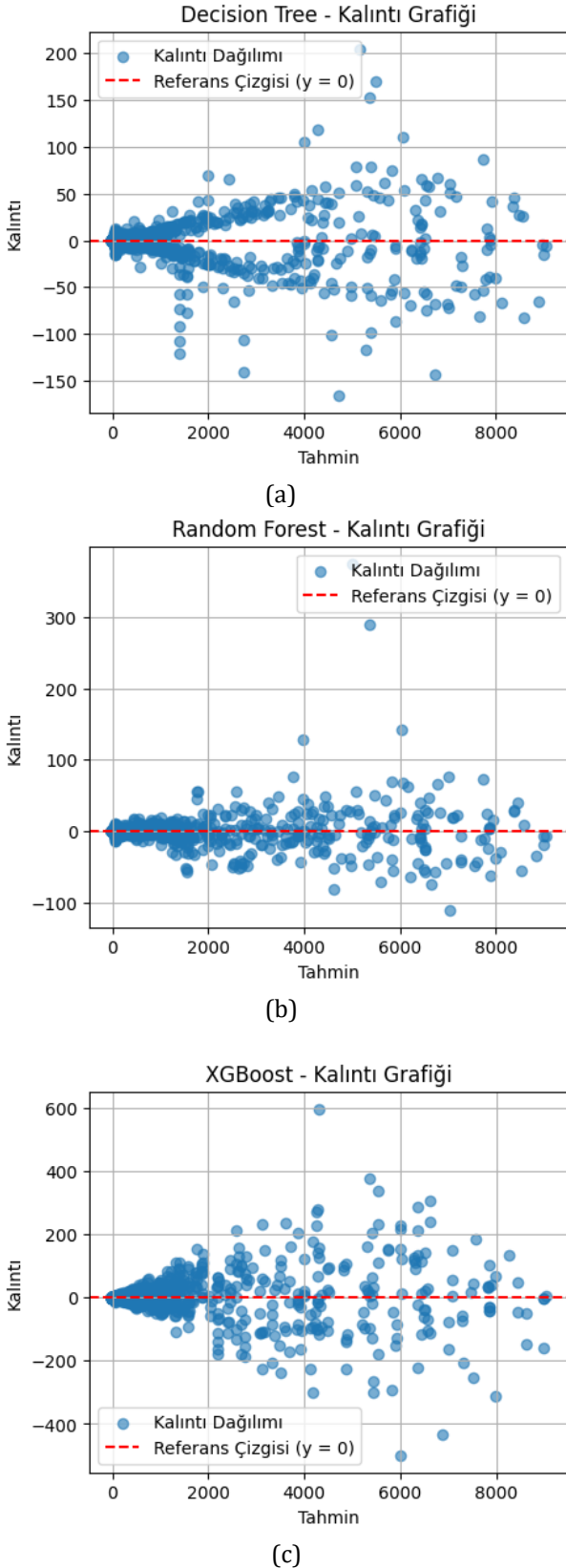
Tablo 11. Tüm modellerin gerçek değerler ile tahmin edilen basınç değerleri.

Gerçek Değer	Tahmin Değeri (Decision Tree)	Tahmin Değeri (Random Forest)	Tahmin Değeri (XGBoost)
49.02229	48.471192	47.970658	48.399296
5293.479	5311.286	5299.47771	5255.908203
60.22967	56.388141	58.397946	58.904957
70.08236	70.533620	70.172157	67.384402
19.59600	19.234800	19.345404	20.927893
-5.324062	-5.466102	-5.359881	-5.633992
2021.785	2039.418193	2036.140771	2016.367310
96.46930	97.209415	98.604179	97.069702
33.42666	43.458730	33.441528	27.610344

Şekil 5, XGBoost, Decision Tree Regressor ve Random Forest Regressor modellerinin test veri kümesi üzerinde tahmin ettikleri değerler ile gerçek değerler arasındaki ilişkiyi göstermektedir. Mavi noktalar, modellerin tahmin değerlerini; kırmızı kesikli çizgi ise ideal tahmin doğruluğunu temsil etmektedir. XGBoost modeli, tahmin değerlerinin büyük çoğunluğunun ideal doğruluk çizgisi üzerinde veya çok yakınında yer aldığı bir performans sergilemiş, ancak bazı veri noktalarında hafif sapmalar gözlemlenmiştir. Decision Tree modeli, tahmin değerlerinin büyük ölçüde kırmızı çizgi ile örtüştüğünü göstermiştir; ancak tek bir ağaç yapısına dayalı olması nedeniyle bazı veri noktalarında genelleme kabiliyeti sınırlı kalmıştır.



Şekil 5. Modellerin parite grafiği. (a) Decision Tree, (b) Random Forest, (c) XGBoost



Şekil 6. Modellerin kalıntı grafiği. . (a) Decision Tree, (b) Random Forest, (c) XGBoost

Random Forest modeli ise tahmin değerlerinin neredeyse tamamen ideal doğruluk çizgisi ile örtüştüğü, sapma ve varyansın dengelendiği bir yapı sunarak en düşük hata değerleriyle en iyi genelleme başarısını elde etmiştir. Genel olarak, tüm modellerin yüksek doğruluk sağladığı görülmüş; ancak Random

Forest modelinin genelleme kapasitesinin diğer modellere kıyasla daha iyi olduğu tespit edilmiştir.

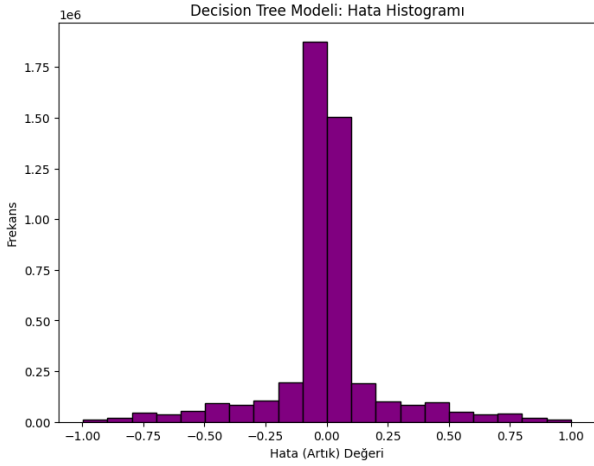
Model performanslarını değerlendirmek amacıyla tahmin hatalarının dağılımı incelenmiş ve sonuçlar her bir model için ayrı ayrı Şekil 6 'da görselleştirilmiştir. Decision Tree, Random Forest ve XGBoost modellerine ait kalıntı (residual) grafiklerine bakıldığında, hataların büyük oranda sıfır çevresinde yoğunlaştığı görülmektedir. Bu durum, modellerin genel olarak yüksek doğrulukla tahmin gerçekleştirdiğini göstermektedir.

Decision Tree modeline ait kalıntılar, sınırlı bir dağılım göstermekte ve tahmin hataları genellikle dar bir aralıkta yoğunlaşmaktadır. Bu da modelin istikrarlı bir şekilde tahmin yaptığına işaret etmektedir. Random Forest modelinin kalıntı dağılımı da benzer şekilde simetrik ve dar bir yapıdadır; bu da modelin genelleme başarısının yüksek olduğunu düşündürmektedir. XGBoost modelinde ise bazı uç değerlerin daha belirgin olduğu görülmektedir. Bununla birlikte, XGBoost'un kalıntılarının büyük çoğunluğu yine sıfır çevresinde toplanmıştır; bu durum, modelin doğruluğunun yüksek olduğunu ancak bazı örneklerde göreceli sapmaların oluşabileceğini göstermektedir.

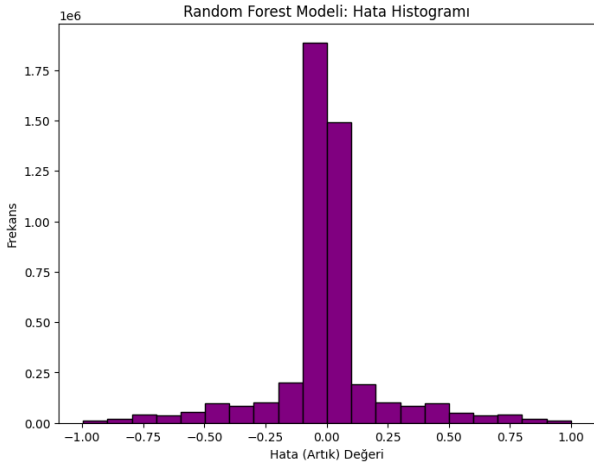
Grafiklerdeki kalıntıların simetrik yapıda ve referans çizgisi etrafında kümelenmesi, sistematik bir hata bulunmadığını ve modellerde aşırı öğrenme (overfitting) sorunu yaşanmadığını desteklemektedir. Genel olarak tüm modeller benzer hata dağılımı profiline sahip olmakla birlikte, Random Forest ve Decision Tree modellerinin kalıntı dağılımları daha dar ve dengeli bir görünüm sergilediğinden, genelleme açısından daha güvenilir sonuçlar verdiği söylenebilir.

Şekil 7'de, tahmin hatalarının dağılımını değerlendirmek amacıyla XGBoost, Decision Tree (Karar Ağacı) ve Random Forest algoritmaları ile oluşturulan modellerin hata histogramları incelenmiştir. Her modelin tahmin hatalarının merkezi eğilimleri ve dağılım özellikleri bu grafikler üzerinden görselleştirilmiştir. Random Forest modelinin hata dağılımı oldukça dar bir bantta yoğunlaşmış olup, hataların büyük kısmı ± 15 aralığında toplanmıştır. Bu durum, modelin hem düşük hem de dengeli hatalarla tahmin yaptığını göstermektedir. Aynı zamanda bu gözlem, modelin en düşük MAE (4.546) ve RMSE (15.552) değerlerine sahip olmasıyla da örtüşmektedir. Decision Tree modeli de sıfır merkezli simetrik bir dağılım göstermekte olup, tahmin hataları genellikle ± 20 civarında dağılmıştır. Ancak uç değerlerde Random Forest'a göre biraz daha fazla sapma gözlenmiştir. XGBoost modelinin hata dağılımı da genel olarak sıfır çevresinde yoğunlaşmaktadır; ancak dağılımın daha yaygın ve uç değerlere daha açık olduğu görülmektedir. Bu modelin MAE (16.212) ve RMSE (47.394) değerleri diğer iki modele kıyasla belirgin şekilde daha yüksek olduğundan, tahmin hatalarının

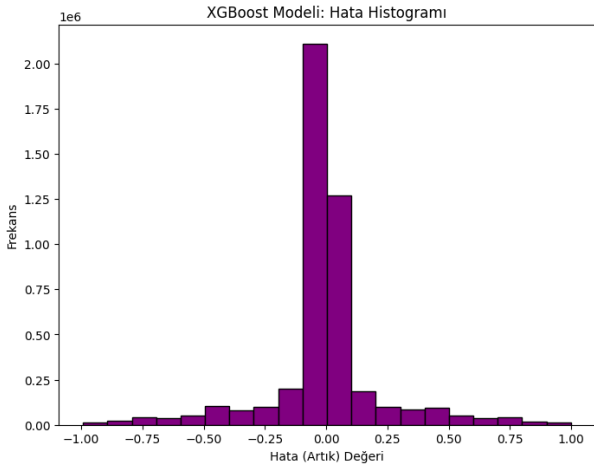
hem daha değişken hem de daha büyük olduğu söylenebilir.



(a)



(b)

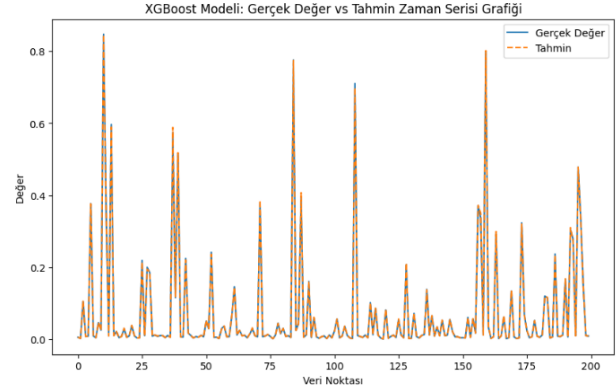


(c)

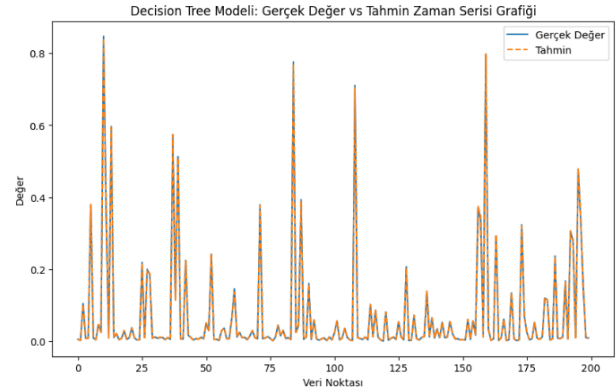
Şekil 7. Modellerin hata histogramları. (a) Decision Tree, (b) Random Forest, (c) XGBoost

Genel olarak, tüm modellerin hata dağılımları normal dağılıma yakın bir eğilim göstermektedir. Ancak, Random Forest modeli, en düşük hata değerleri ve en dar dağılımla en dengeli ve güvenilir tahmin performansını sergileyen model olmuştur. Bu

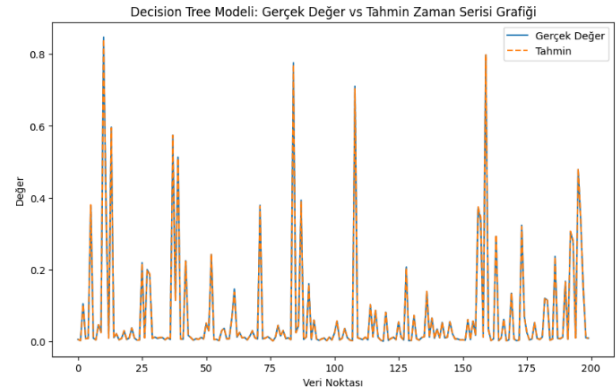
bulgular, özellikle bu çalışmada kullanılan veri seti bağlamında Random Forest algoritmasının en iyi genelleme yeteneğine sahip model olduğunu göstermektedir.



(a)



(b)



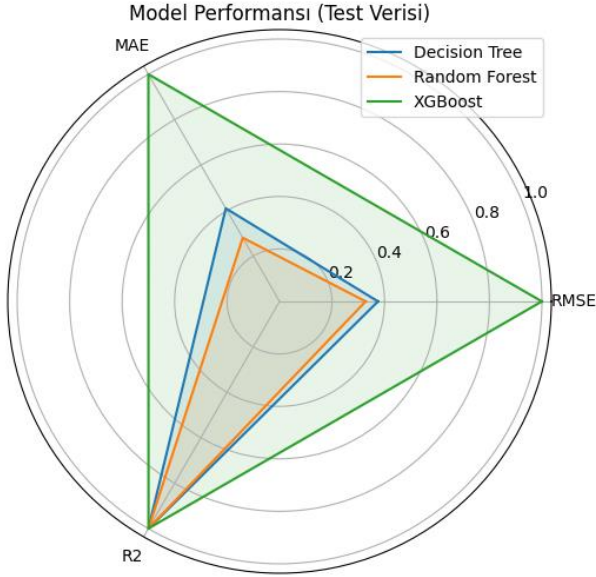
(c)

Şekil 8. Modellerin zaman serisi grafikleri. (a) Decision Tree, (b) Random Forest, (c) XGBoost

Şekil 8’de, XGBoost, Decision Tree ve Random Forest modellerinin tahmin performanslarını değerlendirmek amacıyla oluşturulan zaman serisi grafikleri incelenmiştir. Grafiklerde, modeller tarafından yapılan tahminler ile gerçek değerler karşılaştırılmıştır. Tüm modeller için tahmin değerlerinin zaman serisi boyunca gerçek değerlere oldukça yakın olduğu gözlemlenmiştir.

Şekil 9’da farklı makine öğrenimi modellerinin (Random Forest, Decision Tree ve XGBoost) test verisi

üzerindeki performans metrikleri görselleştirilmiştir. Grafikte, R^2 skoru, MAE ve RMSE gibi temel performans göstergeleri kullanılarak modellerin karşılaştırması yapılmıştır.



Şekil 9. Modellerin radar grafiği ile performanslarının gösterimi

Elde edilen sonuçlara göre, Random Forest ve Decision Tree modelleri, $R^2 = 0.9999$ gibi oldukça yüksek bir doğruluk düzeyine ulaşmıştır. Özellikle Random Forest modeli, MAE = 4.546 ve RMSE = 15.552 ile hata metriklerinde en düşük değerlere ulaşarak en başarılı model olmuştur. Decision Tree, benzer bir R^2 değerine sahip olmasına rağmen, hata metriklerinde biraz daha yüksek değerler üretmiştir (MAE = 6.639, RMSE = 17.785).

Öte yandan, XGBoost modeli $R^2 = 0.9990$ ile görece olarak daha düşük bir doğruluk sağlamış ve hata metriklerinde de MAE = 16.212, RMSE = 47.394 gibi daha yüksek değerlere ulaşmıştır. Bu durum, XGBoost'un bu özel veri setinde diğer modellere kıyasla daha düşük bir genel performans sergilediğini göstermektedir.

Tablo 12. Performans Metriklerine Göre ANOVA Sonuçları.

Performans Ölçütü	F Değeri	p Değeri	İstatistiksel Yorum
RMSE	150.5844	0.000	RMSE değerleri arasında anlamlı fark var.
MAE	307.7429	0.000	MAE değerleri arasında anlamlı fark var.
R^2	189.3405	0.000	R^2 değerleri arasında anlamlı fark var.

Sonuç olarak, grafik ve sayısal veriler birlikte değerlendirildiğinde, Random Forest modelinin hem yüksek doğruluk hem de düşük hata oranları ile en başarılı model olduğu; Decision Tree modelinin makul performans gösterdiği, XGBoost modelinin ise

özellikle hata metrikleri açısından daha zayıf kaldığı görülmektedir. Bu da model seçiminde yalnızca R^2 skoru değil, hata ölçütlerinin de birlikte değerlendirilmesinin önemini ortaya koymaktadır.

Modellerin doğruluk ve hata performansları arasındaki istatistiksel farklılıkların belirleyebilmek amacıyla, RMSE, MAE ve R^2 metriklerine yönelik olarak tek yönlü varyans analizi (ANOVA) uygulanmıştır. Her bir modelin 5 katmanlı çapraz doğrulama sürecinden elde edilen verileri üzerinden gerçekleştirilen analiz sonuçları, Tablo 12'de özetlenmiştir. Elde edilen F ve p-değerleri dikkate alındığında, hem hata metrikleri (RMSE ve MAE) hem de doğruluk metriği (R^2) açısından modeller arasında istatistiksel olarak anlamlı farklar olduğu görülmektedir ($p < 0.05$). Bu durum, yöntemlerin performans düzeylerinin sadece ortalama değerlerle değil, aynı zamanda varyans temelli istatistiksel kıyaslamalarla da farklılaştığını ortaya koymaktadır.

4. Tartışma ve Sonuç

Bu çalışmada, dizel motorlarda silindir içi basıncı tahmin etmek için uygulanan veri odaklı makine öğrenimi modellerinin performansları karşılaştırılmıştır. Random Forest, Karar Ağacı ve XGBoost algoritmaları kullanılarak gerçekleştirilen analizler, her bir modelin farklı avantajlarını ve sınırlamalarını ortaya koymuştur.

Elde edilen bulgular, Random Forest modelinin genelleme kapasitesinin üstün olduğunu ve en düşük hata oranlarıyla tahminlerde bulunduğunu göstermiştir. Bu modelin, çoklu karar ağaçlarını kullanarak genelleme kabiliyetini artırması ve hataları azaltması, motor performans analizi gibi yüksek doğruluk gerektiren uygulamalarda kullanımını desteklemektedir. Öte yandan, modelin işlem süresi diğer algoritmalara göre daha uzun olsa da bu süre pratik kullanım için kabul edilebilir düzeydedir.

XGBoost modeli, özellikle hızlı tahmin süresi ve büyük veri setlerinde paralel hesaplama yeteneği ile öne çıkmıştır. Ancak, modelin doğruluk oranının diğer modellere kıyasla hafif bir düşüş göstermesi, genelleme başarısının sınırlı olabileceğini işaret etmektedir. Yine de bu model, gerçek zamanlı uygulamalar ve büyük ölçekli veri işleme gereksinimi olan projeler için uygun bir seçenek olarak değerlendirilebilir.

Karar Ağacı modeli ise, basit yapısı sayesinde nispeten hızlı tahminler yapmış, ancak genelleme yeteneğinin yetersiz olması nedeniyle diğer modellere göre daha düşük performans sergilemiştir. Tek bir karar ağacı ile sınırlı olan bu model, veri setindeki varyans ve uç değerlerin tahmin doğruluğu üzerindeki etkilerine karşı daha duyarlıdır.

Bu çalışmada elde edilen yüksek R^2 skorları, ilk bakışta modellerin aşırı öğrenme (overfitting) yaptığı

izlenimini uyandırabilir. Ancak bu durum, kullanılan modelleme stratejileri ve doğrulama yöntemleri kapsamında sistematik olarak değerlendirilmiş ve kontrol altına alınmıştır. Modeller, yalnızca eğitim verisi üzerinde değil, aynı zamanda beş katmanlı çapraz doğrulama yöntemi ile doğrulama kümelerinde test edilmiş; ardından, bu süreçte kullanılmamış olan bağımsız %15'lik test veri seti üzerinde nihai değerlendirmeye tabi tutulmuştur. Elde edilen yüksek test R^2 değerleri, modellerin yalnızca eğitim verisini ezberlemediğini, aynı zamanda genel veri yapısını başarılı şekilde öğrendiğini göstermektedir.

Buna ek olarak, model kalıntıları (residuals) grafiksel olarak incelenmiş; kalıntıların büyük oranda sıfır etrafında rastgele dağıldığı ve sistematik bir sapmanın bulunmadığı gözlemlenmiştir. Bu, modellerin aşırı öğrenme eğiliminde olmadığını ve genelleme yeteneğinin korunduğunu desteklemektedir. Ayrıca, normalize edilmiş kalıntı analizleri sonucunda büyük tahmin değerlerinde varyans artışı gözlenmiş olsa da bu durum, klasik overfitting tanımına karşılık gelmemektedir; aksine modelin bazı uç değerlerde doğal sapma eğilimi gösterdiğini, ancak genelde güvenilir tahminler sunduğunu göstermektedir.

Sonuç olarak, yüksek R^2 değerleri modelin öğrenme kapasitesini yansıtmaktadır ve overfitting'e işaret eden belirgin bir kanıt bulunmamaktadır. Bununla birlikte, modelin uç değerlerdeki davranışları ve tahmin kararlılığı ilerleyen çalışmalarda farklı test kümeleri ve alternatif düzenleme stratejileri ile daha ayrıntılı olarak incelenebilir.

Literatürdeki çalışmalarla uyumlu olarak, makine öğrenimi tabanlı yaklaşımlar, termodinamik modellere kıyasla hem hesaplama süresi hem de doğruluk açısından önemli avantajlar sunmaktadır. Ancak, veri setindeki değişkenlerin dağılımı, normalizasyon süreçleri ve model hiperparametrelerinin optimizasyonu, elde edilen sonuçların başarısını doğrudan etkileyen kritik faktörlerdir. Özellikle, yüksek standart sapmaya sahip basınç değerleri gibi uç değerlere karşı modellerin duyarlılığı, tahmin performansı üzerinde belirleyici olmuştur.

Tüm modellerin R^2 değerlerinin oldukça yüksek olduğu ve doğrusal ilişkiyi güçlü biçimde yakalayabildiği görülmektedir. Ancak hata metrikleri (MAE ve RMSE) dikkate alındığında, modeller arasında belirgin farklılıklar mevcuttur. Random Forest Regressor modeli, en düşük MAE (4.546) ve RMSE (15.552) değerleri ile en başarılı tahmin performansını sergilemiştir. Bu sonuç, çoklu ağaç yapısının aşırı öğrenmeyi önlemede ve genelleme kapasitesini artırmada etkili olduğunu göstermektedir. Decision Tree modeli de yüksek doğruluk göstermesine rağmen, tekil ağaç yapısı nedeniyle sınırlı genelleme kapasitesine sahiptir. XGBoost modeli ise teorik olarak gelişmiş bir algoritma olmasına rağmen bu çalışmada yüksek MAE

(16.212) ve RMSE (47.394) değerleriyle göreceli zayıf bir performans sergilemiştir. Bu durum, modelin hiperparametre hassasiyetine ve belirli veri dağılımlarında gösterdiği kararsızlığa işaret edebilir. Bu bulgular, literatüre önemli bir karşılaştırmalı katkı sunmakta olup, özellikle regresyon temelli kestirim problemlerinde doğru algoritma seçiminin model başarımı üzerinde belirleyici olduğunu ortaya koymaktadır. Çalışmanın sonuçları, mühendislik uygulamalarında, özellikle sistem performansının tahmini, süreç optimizasyonu ve veri temelli karar destek sistemlerinin geliştirilmesi açısından pratik bir referans oluşturmaktadır.

Çalışmada kullanılan üç farklı makine öğrenmesi modelinin (Decision Tree, Random Forest, XGBoost), 5 katmanlı çapraz doğrulama sürecinde elde ettiği performans metrikleri (RMSE, MAE ve R^2) üzerinde tek yönlü ANOVA (One-Way ANOVA) analizi gerçekleştirilmiştir. Bu analiz ile modellerin hata ve doğruluk ölçütleri açısından istatistiksel olarak anlamlı farklılık gösterip göstermediği test edilmiştir. Elde edilen analiz sonuçları Tablo 12'de özetlenmiştir. RMSE ($F = 150.5844$, $p = 0.0000$), MAE ($F = 307.7429$, $p = 0.0000$) ve R^2 ($F = 189.3405$, $p = 0.0000$) metriklerinin tamamı için p-değerleri 0.05 anlamlılık düzeyinin oldukça altında bulunmuş; bu da modellerin performansları arasında istatistiksel olarak anlamlı farklar olduğunu göstermektedir. Bu bulgu, yalnızca ortalama hata ve doğruluk değerlerine değil, aynı zamanda varyans yapısına dayalı karşılaştırmalarla da modellerin birbirinden ayrıldığını ortaya koymaktadır. Özellikle Random Forest modelinin RMSE ve MAE açısından diğer yöntemlere göre daha düşük hata değerleriyle öne çıktığı ve bu başarımlarının istatistiksel olarak da desteklendiği görülmüştür. R^2 metriğindeki fark da anlamlı bulunmuş olup, modellerin doğruluk düzeylerinde de istatistiksel farklılıklar olduğu sonucuna ulaşılmıştır. Bu analiz, yöntemler arası kıyaslamaların nesnel biçimde yapılmasını sağlamış ve çalışmanın bilimsel güvenilirliğini önemli ölçüde artırmıştır.

Bu çalışmada elde edilen sonuçlar, literatürde benzer amaçla gerçekleştirilen çalışmalarla karşılaştırıldığında oldukça başarılı bulunmuştur. Reimer (2024) tarafından Otto tipi içten yanmalı motorlar üzerinde gerçekleştirilen çalışmada çok değişkenli lineer regresyon (MvLR) ve yapay sinir ağları (ANN) yöntemleri kullanılmış; ANN ile 0.9973 R^2 skoru ve 0.2126 Bar² MSE elde edilmiştir. Sontheimer et al. (2023) tarafından HCCI motorları üzerinde yapılan çalışmada ise uzun kısa süreli bellek (LSTM) mimarisi ile 0.20 MSE ve 0.37 MAE değerlerine ulaşılmıştır. Öte yandan Patil ve Theotokatos (2023), deniz dizel motorlarında Fourier serileri üzerinden ANN ile gerçekleştirilen modellemelerde ± 0.2 bar civarında RMSE elde etmişlerdir. Bu çalışmada kullanılan Random Forest ve Decision Tree modelleri ile test verisi üzerinde sırasıyla $R^2 = 0.9999$, MAE = 4.546 ve RMSE = 15.552 gibi oldukça düşük hata değerleri elde edilmiştir.

Tablo 13 Literatürde yapılan güncel çalışmalar.

Makale	Yöntem	Uygulama	MSE	R ²	MAE
[30]	MvLR, ANN	Otto (SI-ICE)	0.2126 Bar ² (ANN)	0.9973 (ANN)	-
[19]	LSTM, DNN	HCCI	0.20 (LSTM)	-	0.37 (LSTM)
[31]	Linear, Elastic, Polynomial, SVM, DT, ANN	Marine Diesel	±0.2 Bar (ANN)	-	-

Bu değerler, sadece doğruluk açısından değil, aynı zamanda genelleme yeteneği bakımından da Random Forest modelinin literatürdeki gelişmiş yapay zeka teknikleriyle kıyaslandığında oldukça rekabetçi ve hatta bazı açılardan üstün olduğunu göstermektedir. Ayrıca, modelin sade yapısı ve hesaplama maliyetinin düşük olması, onu özellikle gerçek zamanlı uygulamalar açısından güçlü bir aday haline getirmektedir.

Bu çalışmanın sonuçları, makine öğrenimi yöntemlerinin dizel motor performansı ve emisyon kontrolü gibi kritik alanlarda güçlü bir alternatif sunduğunu göstermektedir. Ancak, modellerin daha geniş bir veri setiyle ve farklı motor tiplerinde test edilmesi, elde edilen bulguların genellenebilirliğini artıracaktır. Ayrıca, gelecekteki çalışmalar, derin öğrenme yöntemlerini ve daha karmaşık zaman serisi modellerini içerecek şekilde genişletilebilir. Bu sayede, dizel motorların yanma süreçlerini ve emisyon dinamiklerini daha iyi anlamaya yönelik daha kapsamlı tahmin sistemleri geliştirilebilir.

Etik Beyanı

Bu çalışmada, “Yükseköğretim Kurumları Bilimsel Araştırma ve Yayın Etiği Yönergesi” kapsamında uyulması gerekli tüm kurallara uyulduğunu, bahsi geçen yönergenin “Bilimsel Araştırma ve Yayın Etiğine Aykırı Eylemler” başlığı altında belirtilen eylemlerden hiçbirinin gerçekleştirilmediğini taahhüt ederiz.

Kaynakça

[1] E. Alptekin, “Metil ve Etil Ester Kullanılan Bir Common-Rail Dizel Motorda Performans, Yanma ve Enjeksiyon Karakteristiklerinin Karşılaştırılması”, Süleyman Demirel Üniversitesi Fen Bilim. Enstitüsü Derg., c. 21, sy 2, s. 578, Şub. 2017.

[2] J. F. Dunne ve C. Bennett, “A crank-kinematics-based engine cylinder pressure reconstruction model”, Int. J. Engine Res., c. 21, sy 7, ss. 1147-1161, Eyl. 2020.

[3] Y. Lee, S. Lee, ve K. Min, “Semi-empirical estimation model of in-cylinder pressure for compression ignition engines”, Proc. Inst. Mech. Eng. Part J. Automob. Eng., c. 234, sy 12, ss. 2862-2877, Eki. 2020.

[4] W. Jeon vd., “Accelerometer-Based Robust Estimation of In-Cylinder Pressure for Cycle-to-Cycle Combustion Control”, IEEE Trans. Instrum. Meas., c. 72, ss. 1-13, 2023.

[5] S. Kulah, A. Forrai, F. Rentmeester, T. Donkers, ve F. Willems, “Robust cylinder pressure estimation in heavy-duty diesel engines”, Int. J. Engine Res., c. 19, sy 2, ss. 179-188, Şub. 2018.

[6] R. Ristow Hadlich, J. Loprete, ve D. Assanis, “A Deep Learning Approach to Predict In-Cylinder Pressure of a Compression Ignition Engine”, J. Eng. Gas Turbines Power, ss. 1-12, Oca. 2024.

[7] C. Patil ve G. Theotokatos, “Comparative Analysis of Data-Driven Models for Marine Engine In-Cylinder Pressure Prediction”, Machines, c. 11, sy 10, s. 926, Eyl. 2023.

[8] J. Castresana, G. Gabiña, L. Martin, ve Z. Uriondo, “Comparative performance and emissions assessments of a single-cylinder diesel engine using artificial neural network and thermodynamic simulation”, Appl. Therm. Eng., c. 185, s. 116343, Şub. 2021.

[9] B. Lee ve D. Jung, “Thermodynamics-based mean-value engine model with main and pilot injection sensitivity”, Proc. Inst. Mech. Eng. Part J. Automob. Eng., c. 230, sy 13, ss. 1822-1834, Ağu. 2016.

[10] H. Cho, B. Fulton, D. Upadhyay, T. Brewbaker, ve M. van Nieuwstadt, “In-cylinder pressure sensor-based NOx model for real-time application in diesel engines”, Int. J. Engine Res., c. 19, sy 3, ss. 293-307, Nis. 2017.

[11] J. Höller vd., “Parameter Estimation Strategies in Thermodynamics”, ChemEngineering, c. 3, sy 2, s. 56, Haz. 2019.

[12] C. Mährle, S. Held, S. Huber, ve G. Wachtmeister, “A new method to determine the causes of deviation in cylinder pressure curves of motored reciprocating piston engines”, Int. J. Engine Res., c. 23, sy 2, ss. 243-261, Oca. 2021.

[13] C. Jorques Moreno, O. Stenlaas, ve P. Tunestal, “Cylinder Pressure-Based Virtual Sensor for In-Cycle Pilot Mass Estimation”, SAE Int. J. Engines, c. 11, sy 6, ss. 1167-1182, Nis. 2018.

[14] R. Ristow Hadlich, J. Loprete, ve D. Assanis, “A Deep Learning Approach to Predict In-Cylinder Pressure of a Compression Ignition Engine”, J. Eng. Gas Turbines Power, ss. 1-12, Oca. 2024.

[15] C. Patil ve G. Theotokatos, “Comparative Analysis of Data-Driven Models for Marine Engine In-

- Cylinder Pressure Prediction", *Machines*, c. 11, sy 10, s. 926, Eyl. 2023.
- [16] R. Ristow Hadlich, J. Loprete, ve D. Assanis, "A Deep Learning Approach to Predict In-Cylinder Pressure of a Compression Ignition Engine", *J. Eng. Gas Turbines Power*, ss. 1-12, Oca. 2024.
- [17] E. Aslan ve Y. Özüpak, "Comparison of machine learning algorithms for automatic prediction of Alzheimer disease", *J. Chin. Med. Assoc.*, c. 88, sy 2, ss. 98-107, Şub. 2025.
- [18] E. Aslan, "Temperature Prediction and Performance Comparison of Permanent Magnet Synchronous Motors Using Different Machine Learning Techniques for Early Failure Detection", *Eksplot. Niezawodn. – Maint. Reliab.*, c. 27, sy 1, Ağu. 2024.
- [19] M. Sontheimer, A.-K. Singh, P. Verma, S.-Y. Chou, ve Y.-L. Kuo, "LSTM for Modeling of Cylinder Pressure in HCCI Engines at Different In take Temperatures via Time-Series Prediction", *Machines*, c. 11, sy 10, s. 924, Eyl. 2023.
- [20] X. Jin, P. Cheng, W.-L. Chen, ve H. Li, "Prediction model of velocity field around circular cylinder over various Reynolds numbers by fusion convolutional neural networks based on pressure on the cylinder", *Phys. Fluids*, c. 30, sy 4, Nis. 2018.
- [21] S. A. Ali ve S. Saraswati, "Cycle-by-cycle estimation of cylinder pressure and indicated torque waveform using crankshaft speed fluctuations", *Trans. Inst. Meas. Control*, c. 37, sy 6, ss. 813-825, Eyl. 2014.
- [22] S. Trimby, J. F. Dunne, C. Bennett, ve D. Richardson, "Unified approach to engine cylinder pressure reconstruction using time-delay neural networks with crank kinematics or block vibration measurements", *Int. J. Engine Res.*, c. 18, sy 3, ss. 256-272, Tem. 2016.
- [23] Q. Huang, J. Liu, C. Ulishney, ve C. E. Dumitrescu, "On the use of artificial neural networks to model the performance and emissions of a heavy-duty natural gas spark ignition engine", *Int. J. Engine Res.*, c. 23, sy 11, ss. 1879-1898, Tem. 2021.
- [24] F. Zhao ve D. L. S. Hung, "Applications of machine learning to the analysis of engine in-cylinder flow and thermal process: A review and outlook", *Appl. Therm. Eng.*, c. 220, s. 119633, Şub. 2023.
- [25] M. Henningsson, P. Tunestål, ve R. Johansson, "A Virtual Sensor for Predicting Diesel Engine Emissions from Cylinder Pressure Data", *IFAC Proc. Vol.*, c. 45, sy 30, ss. 424-431, 2012.
- [26] D. P. Viana vd., "Diesel Engine Fault Prediction Using Artificial Intelligence Regression Methods", *Machines*, c. 11, sy 5, s. 530, May. 2023.
- [27] E. Tian, G. Lv, ve Z. Li, "Evaluation of emission of the hydrogen-enriched diesel engine through machine learning", *Energy*, c. 307, s. 132303, Eki. 2024.
- [28] W. N. Wan Mansor, S. Abdullah, M. N. K. Jarkoni, J. S. Vaughn, ve D. B. Olsen, "Data on combustion, performance and emissions of a 6.8 L, 6-cylinder, Tier II diesel engine", *Data Brief*, c. 33, s. 106580, Ara. 2020.
- [29] E. Aslan, "Prediction and Comparative Analysis of Emissions from Gas Turbines Using Random Search Optimization and Different Machine Learning Based Algorithms", *Bull. Pol. Acad. Sci. Tech. Sci.*, ss. 151956-151956, Eyl. 2024.
- [30] G. Reimer, "Online in-cylinder pressure and temperature prediction using a modeling approach: Masters Thesis".
- [31] C. Patil ve G. Theotokatos, "Comparative Analysis of Data-Driven Models for Marine Engine In-Cylinder Pressure Prediction", *Machines*, c. 11, sy 10, s. 926, Eyl. 2023.