



Tekstil üretim süreçlerinde hata tahmini ve önlenmesi: Makine öğrenmesi tabanlı karar destek sistemi yaklaşımı

Error prediction and prevention in textile production processes: Machine learning based decision support system approach

Enes Buğra Özkan¹, Berra Figengöz², Ali Osman Özdemir¹, Fatma Civelek¹, Aslı Aksoy^{1*}, Seval Ene Yalçın¹, Seda Güneri Uslu³

¹Endüstri Mühendisliği Bölümü, Bursa Uludağ Üniversitesi, Bursa, Türkiye.

enes.ozkan@peksavunma.com, al.ozdemir@icloud.com, fatma.civelekdot@gmail.com, asliaksoy@uludag.edu.tr, secalene@uludag.edu.tr

²Tekstil Mühendisliği Bölümü, Bursa Üniversitesi, Bursa Türkiye.

berrafignz@gmail.com

³Akbaşlar Tekstil, Ar-Ge Merkezi, Bursa Türkiye.

seda.guneri@asbaslar.com

Geliş Tarihi/Received: 15.04.2025

Düzeltilme Tarihi/Revision: 26.09.2025

doi: 10.65206/pajes.10594

Kabul Tarihi/Accepted: 06.11.2025

Araştırma Makalesi/Research Article

Öz

Tekstil üretim süreçleri, çok sayıda birbirine bağımlı değişkenin eşzamanlı olarak yönetilmesini gerektiren karmaşık yapılarla karakterize edilir. Bu değişkenler; kumaş kalitesi, renk tutarlılığı, doku, dayanıklılık ve üretim maliyeti gibi temel ürün özelliklerini doğrudan etkiler. Süreçlerdeki yüksek değişkenlik, istenilen kalite seviyesinin tutarlı biçimde sağlanmasını zorlaştırmakta ve üretim hatalarının oluşma olasılığını artırmaktadır. Bu zorluklara çözüm olarak, çalışmada tekstil üretiminde süreç parametrelerini optimize etmeyi ve olası üretim hatalarını önceden tahmin etmeyi amaçlayan makine öğrenmesi tabanlı bir karar destek sistemi önerilmektedir. Sistem, mevcutta bulunan çalışmalardan farklı olarak üretim değişkenlerinin sistematik biçimde yönetilmesini sağlayarak kalite ve verimliliği artırmaya yönelik proaktif müdahalelere olanak tanımaktadır. 2022-2023 yılları arasında bir tekstil işletmesinde gerçekleştirilen kapsamlı analizler sonucunda, rotasyon baskı aşamasında en sık karşılaşılan hata türü belirlenmiş ve bu hatanın temel nedenleri sistem analiz adımlarıyla incelenmiştir. Üretim sonuçlarını etkileyen kritik parametreler tanımlanmış, ilgili veri depolama sistemlerinden elde edilen tarihsel üretim verileri model geliştirme sürecinde kullanılmıştır. Bu çalışmanın temel katkısı, hata nedenlerinin analizini içeren çok yönlü yaklaşımı ile tahmine dayalı bir karar destek modelinin geliştirilmesini bir araya getirmesidir. Model, hem operatör hem de makine performans metriklerini dikkate alan makine öğrenmesi algoritmalarından yararlanarak hatasız üretim senaryolarını öngörmektedir. Böylece yalnızca olası hataları önceden tahmin etmekle kalmayıp, aynı zamanda bu hataların oluşmasını engelleyecek optimum parametre yapılarını da önermektedir. Araştırma, tekstil sektöründe süreç iyileştirme ve hata önleme stratejilerinin geliştirilmesinde veri odaklı yöntemlerin potansiyelini ortaya koymaktadır. Önerilen sistem, üretim güvenilirliğini artıran, atık oranlarını azaltan ve karmaşık üretim ortamlarında gerçek zamanlı karar alma süreçlerini destekleyen ölçeklenebilir ve akıllı bir çerçeveye sunmaktadır.

Anahtar kelimeler: Süreç yönetimi, makine öğrenmesi, hata tahmini, tekstil sektörü, performans iyileştirme.

Abstract

Textile production processes are inherently complex, involving the simultaneous management of numerous interdependent variables. These variables significantly affect key product attributes such as fabric quality, color consistency, texture, durability, and overall production cost. The high degree of variability within these processes presents a substantial challenge to maintaining consistent product quality, thereby increasing the likelihood of production errors. As a solution to these challenges, this study proposes a machine learning-based decision support system aimed at optimizing process parameters in textile production and predicting potential manufacturing defects in advance. Unlike existing approaches, the system enables the systematic management of production variables, allowing for proactive interventions to enhance quality and efficiency. Between 2022 and 2023, a comprehensive analysis was conducted within a textile company to identify and examine the most prevalent defect type occurring during the rotary printing stage. Through structured system analysis, the root causes of these defects were investigated, and the critical parameters influencing production outcomes were defined. Historical production data, retrieved from internal data storage systems, served as the foundation for model development. A key contribution of this study lies in its integrative approach, which combines defect causality analysis with the construction of a predictive decision support model. The model leverages machine learning algorithms that incorporate both operator and machine performance metrics to forecast defect-free production scenarios. By doing so, it not only anticipates potential errors but also recommends optimal parameter configurations to prevent their occurrence. This research underscores the potential of data-driven methodologies in advancing process improvement and error mitigation strategies within the textile industry. The proposed system offers a scalable and intelligent framework for enhancing production reliability, reducing waste, and supporting real-time decision-making in complex manufacturing environments.

Keywords: Process management, machine learning, error prediction, textile industry, performance improvement.

1 Giriş

Tekstil sektörü, sürekli değişen trendler ve artan rekabet koşulları nedeniyle her geçen gün daha fazla yenilik ve gelişim gerektiren bir alan haline gelmiştir. Bu durum, üretim

süreçlerinde daha yüksek verimlilik ve kaliteyi elde etme çabalarını artırmıştır. Tekstil endüstrisinde, üretim hatalarının minimuma indirilmesi ve süreçlerin optimizasyonu, kalite kontrol ve maliyet yönetimi açısından kritik öneme sahiptir. Yapay zeka ve makine öğrenmesi, tekstil endüstrisinde kalite

*Yazışılan yazar/Corresponding author

kontrol süreçlerini iyileştirmek, hataları tespit etmek ve önlemek amacıyla giderek daha fazla kullanılmaktadır. Bu teknolojiler, büyük veri kümelerini analiz etme ve desen tanıma yetenekleri sayesinde, üretim süreçlerinin daha verimli ve etkili hale getirilmesine olanak tanır. Bu çalışma, kullanılan yöntemleri ve elde edilen sonuçları detaylandırarak, tekstil sektöründe yapay zekâ ve makine öğrenmesi uygulamalarının potansiyelini göstermeyi amaçlamaktadır. Bu bağlamda, literatürde yer alan benzer çalışmalar ve elde edilen bulgular ışığında, bir tekstil işletmesinde gerçekleştirilen uygulamalar ve elde edilen sonuçlar karşılaştırmalı olarak ele alınacaktır.

Literatür taraması kapsamında tekstil işletmelerinde gerçekleşen hataların tespiti, tahmini üzerine makaleler incelenmiştir. Hata tahmini üzerine tekstil dışında farklı dallarda yapılan araştırmalar da incelenerek yöntem güçlendirilmeye çalışılmıştır. Şanlıtürk (2018) çalışmasında makine öğrenme algoritmaları kullanarak hatalı ürünleri önceden tahmin etmeyi amaçlamış, rastgele orman (random forest), basit bayes, destek vektör makineleri ve k-en yakın komşu (k nearest neighbour, k-NN) algoritmaları uygulamış ve her bir algoritmanın tahmin performansı karşılaştırmıştır [1]. Demir ve Dinçer (2020), üretimde hata yaratabilecek üretim hattını belirlemek amacıyla matematiksel olarak güçlendirme (boosting) algoritmaları (ADABOOST, XGBOOST, Gradient Boost) kullanmış, yaptıkları simülasyon sonuçlarına göre alt örnekleme metodolojisinin, aşırı örnekleme metodolojisine göre daha iyi performans sergilediği gözlemlenmiştir [2]. Soma ve Pooja (2022), veri türünü tahmin etmek amacıyla gri seviye eşlik matrisi (GLCM) ile 1000 yüksek çözünürlüklü kumaş görüntüsünden özellik çıkarımı yaparak, destek vektör makinesi ve yapay sinir ağıları kullanarak verilerin türünü tahmin etmiş, çalışmada yapay sinir ağıları modelinin daha iyi sonuç verdiği belirtilmiştir[3]. Çıklaçandır ve ark. (2020) sınıflandırma yapmak amacıyla kumaş görüntülerinden özelliklerini çıkarıp K-means kümeleme algoritması ile hatalı veya hatasız olarak sınıflandırmış, özellik çıkarma için gri seviye eş oluşum matrisi ve medyan farkı kullanılarak sonuçları karşılaştırmış, gri seviye eş oluşum matrisi kullanıldığında %97.99'a kadar yüksek bir başarı oranı elde edildiğini göstermiştir[4].

Başka bir çalışmada görüntü işleme tekniklerinden yararlanılarak kumaş üzerinde gerçek zamanlı hata tespiti yapabilecek bir sistem geliştirmiştir. Bu sistem yüksek çözünürlüklü kamera sayesinde anlık kaydedilen dokuma kumaş görüntüleri üzerinde görüntü işleme tekniklerinden Shearlet dönüşümünü kullanarak gerçek zamanlı kumaş hata kontrolünün yapılmasını sağlayan düzenekten oluşmaktadır. %94,25 oranla hataları tespit etmiş ve doğruluğu K katmanlı çapraz doğrulama yöntemi ile doğruluğu test edilmiştir[5]. Sinirsel bulanık çıkarım sistemi (ANFIS) tarafından iplik dayanıklılığı ve düzensizlik tahmin etmek için altı iplik parametresi kullanarak bir model geliştirmiştir. Modeli, iplik dayanıklılığını ve düzensizliğini tahmin etmek üzere eğitilmiş olup tahmin doğruluğu, beş istatistiksel metrik kullanılarak değerlendirilmiştir. Pamuk iplik özellik tahmininde etkili bir şekilde kullanılabilirliği sonucuna ulaşılmıştır[6]. Altunsöğüt (2022) bir tekstil firmasının geçmiş verilerinden faydalanarak siparişler için gerekli kumaş üretiminde en az fireyle kumaşın üretilmesi için makine öğrenmesi algoritmalarını kullanarak tahmin etmeyi amaçlamıştır. Bu amaçla, J48, basit bayes, k-en yakın komşu, destek vektör makineleri (SVM) ve çok katmanlı algılayıcılar (multilayer perceptron) algoritmalarını WEKA uygulamasında test etmiştir. %80.1617 doğruluk oranı 11 ile k-

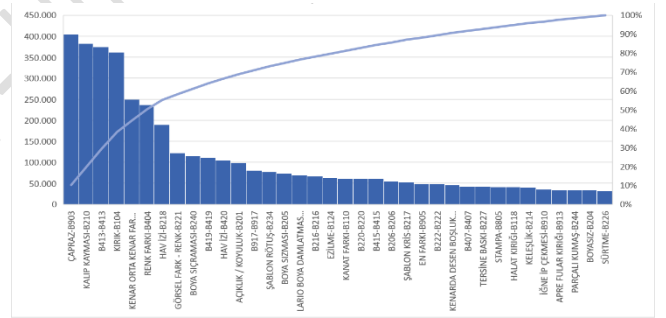
NN algoritması ile elde edilmiştir. Bu belge, kumaşlarda oluşan kusurları otomatik olarak tespit etmek için bir yöntem sunmaktadır[7].

2 Materyal ve yöntem

Bu çalışmada, tekstil üretim sürecinde hataları minimize etmek ve hatasız bir üretim süreci oluşturmak amacıyla makine öğrenmesi algoritmalarından yararlanılmıştır. Çalışmanın başlangıcında, hata tahmini üzerine yapılmış benzer çalışmalar detaylı bir şekilde incelenmiş ve kullanılan yöntem ve araçlar değerlendirilmiştir. Çalışmada kullanılan veri setleri işletme tarafından kullanılan SAP sistemi üzerinde elde edilmiştir. Hata tahminleme çalışmalarında kullanılan farklı yöntemler incelenerek, kalıp kayması gibi hataların kök neden analizleri yapılmıştır. Elde edilen hata ile ilişkili özelliklerin çıkarılması ve bu özelliklerin hatayla olan ilişkilerinin belirlenmesi için korelasyon analizleri gerçekleştirilmiştir.

2.1 Mevcut sistemin analizi

Firmanın 2022 yılına ait üretim verileri incelenerek hataların gerçekleşme miktarları incelenmiş bu bağlamda ilgili hataların gerçekleştiği bölümler ele alınmıştır. Şekil 1. grafikten yola çıkılarak oluşan hataların %51'inin rotasyon baskı bölümüne ait olduğu görülmüştür. Görülme miktarları fazla olan kırık, çaprazlık gibi diğer hataların ise rotasyon baskı işlemlerini etkilediği görülmektedir.



Şekil 1. 2022 yılında görülen hatalar

Figure 1. Defects seen in 2022

Şekil 1.' den elde edilen bilgilerden yola çıkılarak en sık karşılaşılan hataların gerçekleştiği rotasyon baskı bölümü pilot bölge olarak seçilmiştir. Rotasyon baskı üretim verimliliğini etkisini incelemek için baskı öncesi ve baskı sonrası aşamalar Tablo 1' de verildiği gibi ayrılmış ve incelenmiştir.

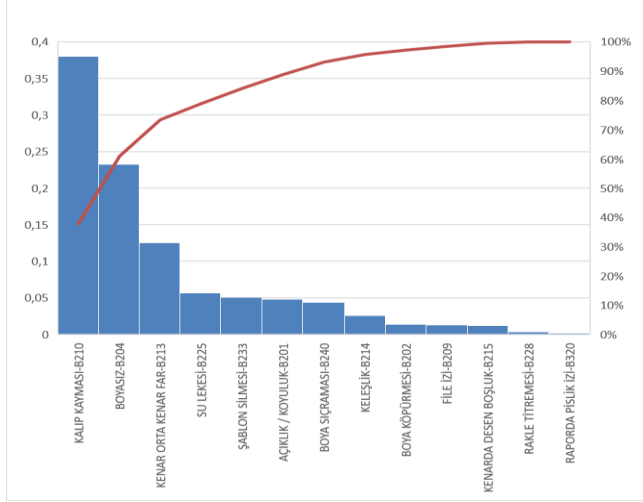
Tablo 1. Baskı Öncesi, Baskı Anı ve Baskı Sonrası İşlemler

Table 1. Pre-printing, During printing and Post-printing Processes

Baskı Öncesi İşlemler	Baskı Anı	Baskı Sonrası İşlemler
Ham kumaş açma		Fıkse
Kumaş ön terbiye	Rotasyon baskı	Yıkama
Ram baskı altı		Ram Apre
Şablon üretimi		

Rotasyon baskı anında gerçekleşen hatalara ait veriler incelendiğinde kalıp kayması, %37 ile en çok görülen hata türüdür. Rotasyon baskı esnasında gerçekleşen hatalara ait pareto grafiği Şekil 2’de verilmiştir.

Kalıp kayması baskı şablonunun olması gereken konumdan sapmasıdır. Kumaş üzerine aktarılmak istenen desen veya yazının olması gereken yerden kaymasına neden olan görsel bir problemdir.



Şekil 2. Rotasyon Baskı Esnasında Gerçekleşen Hatalar
Figure 2. Defects that occur during rotation printing

2.2 Geliştirilen modeller ve çözüm yöntemleri

Bu çalışmada, denetimli ve denetimsiz öğrenme yöntemlerinden faydalanılmıştır. Bu yöntemler, farklı veri kümeleri üzerinde test edilerek en yüksek doğruluk oranına sahip modeller belirlenmiştir.

2.2.1 Denetimli öğrenme

Denetimli öğrenme modelinde problem, sınıflandırma problemi olarak ele alınır ve eğitilmiş sistem test setine yönelik tahmin ve tanıma amacıyla kullanılır[8]. Yaygın olarak kullanılan öğrenme biçimidir. Veri sırasıyla eğitim ve test kümelerine ayrılır. Yapılacak işleme göre veri önceden işaretlenmelidir. Regresyon, tahminleme, sınıflandırma gibi işlemlerde kullanılır. Karar Ağacı, Rassal Orman, XGBoost ve CatBoost modelleri denetimli öğrenme biçimi ile çalışan model örneklerindendir [9].

2.2.1.1 Karar ağacı algoritması

Karar ağacı algoritması, veriyi özniteliklere göre dallara ayırarak sınıflandırma veya regresyon yapan hiyerarşik bir modeldir. Her düğüm, en iyi ayrımı sağlayan kriterle bölünür; yapraklar ise nihai kararları temsil eder. Yorumlanabilirliği yüksek olup, bilgi kazancı veya Gini indeks gibi ölçütlerle çalışır [10][11].

$$H = - \sum p(X) \log p(X) \quad (1)$$

Eşitlik 1’de $p(X)$ belirli bir sınıfa ait grup yüzdesini temsil etmektedir. Karar ağacının entropi değerini en aza

indirgeyecek şekilde bölünme yapması istenmektedir. En iyi bölünmeyi belirlemek için bilgi kazancı Eşitlik 2’de verildiği gibi kullanılır:

$$Gain(S, D) = H(S) - \sum_{V \in D} VSH(V) \quad (2)$$

S ile gösterilen değer orijinal veri kümesini ifade eder D ise kümenin bölünmüş bir parçasıdır. Her V, S’nin bir alt kümesidir. V’nin tümü ayrıktır ve S’yi oluşturmaktadır. bilgi kazancı, bölünmeden önceki orijinal veri setinin entropisi ile her bir özniteliliğin entropi değeri arasındaki fark olarak tanımlanmaktadır [12].

2.2.1.2 Rassal Orman algoritması

Rassal orman algoritması (Random Forest), karar ağaçlarının toplulaştırılmasıyla oluşturulan güçlü ve esnek bir denetimli öğrenme yöntemidir. Her bir ağaç, eğitim verisinin rastgele bir alt kümesiyle ve özniteliklerin rastgele seçimiyle eğitilir; bu çeşitlilik, modelin genelleme yetisini artırır ve aşırı öğrenmeyi azaltır[13][14]. Tahmin sürecinde, sınıflandırma için çoğunluk oyu (majority voting), regresyon için ise ortalama değer kullanılır. Rassal ormanlar, yüksek doğruluk oranı, dirençli hata toleransı ve öznitelik öneminin ölçülebilmesi gibi avantajlarıyla özellikle karmaşık ve yüksek boyutlu veri setlerinde tercih edilmektedir[15].

2.2.1.3 XGBoost algoritması

XGBoost, gradyan yükseltilmiş ağaçlar algoritmasının en popüler ve verimli uygulamalarından biridir. Belirli kayıp fonksiyonlarını optimize ederek ve çeşitli düzenleme tekniklerini uygulayarak işlev yaklaşımına dayalı bir denetimli öğrenme yöntemidir. Aşırı gradyan artırma adı verilen bir yöntemle karar ağacı toplulukları oluşturur. Karar ağacı toplulukları, birden çok karar ağacının tahminlerini birleştirir [16].

$$\hat{y}_i^{(t)} = \sum_{k=1}^t f_k(x_i) = \hat{y}_i^{(t-1)} + f_t(x_i) \quad (3)$$

$$obj^{(t)} = \sum_{i=1}^n l(y_i, \hat{y}_i^{(t-1)} + f_t(x_i)) + \omega(f_t) + \text{sabit} \quad (4)$$

$$l(y_i, \hat{y}_i^{(t-1)} + f_t(x_i)) \approx l(y_i, \hat{y}_i^{(t-1)}) + g_i f_t(x_i) + \frac{1}{2} h_i f_t^2(x_i) \quad (5)$$

$$g_i = \partial_{\hat{y}_i^{(t-1)}} l(y_i, \hat{y}_i^{(t-1)}) \quad (6)$$

$$h_i = \partial_{\hat{y}_i^{(t-1)}}^2 l(y_i, \hat{y}_i^{(t-1)}) \quad (7)$$

Modelde Eşitlik 3 ile y_i gerçek değeri üzerinden $\hat{y}_i^{(t-1)}$ $t-1$ adımıdaki değeri tahmin eder. Her bir karar ağacının tahminine göre bir sonraki karar ağacını oluşturur. Eşitlik 4’ de ki amaç fonksiyonu, $l(y_i, \hat{y}_i)$ kayıp ve $\omega(f_k)$ düzenleme fonksiyonundan oluşmaktadır. $f_t(x_i)$ t adımıdaki yeni öğrenici (learner) fonksiyonudur. Eşitlik 5’de detaylı şekilde verilen kayıp fonksiyonu tahmin edebilirliğini ölçer. Bu fonksiyon için genel olarak ortalama hataların karesi veya lojistik regresyon yöntemlerinden faydalanılır. Eşitlik 6 ve 7 g_i ve h_i , kayıp fonksiyonunun birinci ve ikinci dereceden türevleridir. $\omega(f_k)$

Düzenleştirme terimi, modelin karmaşıklığını ve aşırı uyum sorununun gerçekleşmesinin önüne geçer [17].

2.2.1.4 CatBoost algoritması

CatBoost, gradient boosting temelli, özellikle kategorik verilerle yüksek verimlilikle çalışabilen bir makine öğrenme algoritmasıdır. Kategorik değişkenleri ön işleme gerek duymadan doğrudan işleyebilmesi, veri hazırlık sürecini önemli ölçüde kısaltır. Simetrik ağaç yapısı sayesinde aşırı öğrenme riskini azaltırken, az veriyle dahi yüksek doğrulukta tahminler üretebilir. Sayısal, kategorik ve metin türündeki verilerle etkili biçimde çalışabilen CatBoost, eksik verilerle başa çıkma yeteneği ve hızlı öğrenme kapasitesiyle büyük veri kümeleri gerektiren yöntemlere alternatif olarak öne çıkar. [18] [19].

2.2.2 Denetimsiz öğrenme yöntemi

Denetimsiz öğrenme, etiketlenmemiş veriler üzerinde çalışarak veri kümesindeki gizli örüntüleri ve yapıları keşfetmeyi amaçlayan makine öğrenimi yöntemidir [20]. Denetimsiz öğrenme biçiminde modele eğitim için herhangi bir sınıf ya da sayısal bilgi verilmemektedir. Model veri içindeki ortak noktalar üzerinden sonuç üretmeye çalışır. K-ortalama modeli denetimsiz öğrenme ile çalışan bir algoritmadır. K-Ortalama Algoritması incelenmiştir.

2.2.2.1 K-Ortalama algoritması

K-ortalama algoritması, etiketlenmemiş veriler üzerinde çalışan denetimsiz bir kümeleme yöntemidir. Belirlenen k sayıda merkez etrafında veri noktalarını gruplayarak, aynı küme içindeki benzerliği maksimize, farklı kümeler arasındaki benzerliği minimize etmeyi hedefler. Hata fonksiyonunu en aza indirerek homojen ve ayrık kümeler oluşturur. Büyük veri setlerinde etkili çalışması ve yorumlanabilirliği sayesinde veri bilimi ve makine öğrenmesi uygulamalarında yaygın olarak kullanılmaktadır[21].

3 Bulgular

3.1 Veri seti

Bu çalışmada kullanılan üretim verileri, ilgili tekstil firmasının SAP üretim yönetim sistemi üzerinden elde edilmiştir. 2023 yılına ait toplam 231.194 satır veri incelenmiş, bu verilerin 83.769 satırı rotasyon baskı bölümüne ait kayıtları içermektedir. Veri temizleme süreci kapsamında yanlış girildiği tespit edilen kayıtlar silinmiş ve veri seti 71.388 satıra düşürülmüştür. Ardından, yinelenen satırlar birleştirilerek 3.839 satırlık nihai veri seti oluşturulmuştur.

SAP sisteminde kayıtlı olmayan bazı üretim parametreleri (örneğin; üretim viskozitesi, pres basıncı, mil tipi) firma bünyesindeki boya-baskı kartlarından manuel olarak alınmış ve Excel ortamına aktarılmıştır.

3.2. Veri ön işleme

Modelin doğruluğunu ve genelleme yeteneğini artırmak amacıyla veri ön işleme süreci titizlikle yürütülmüştür. Bu süreçte aşağıdaki adımlar uygulanmıştır:

Eksik ve Hatalı Verilerin Temizlenmesi:

- SAP sisteminden alınan verilerde eksik veya hatalı girişler tespit edilerek veri setinden çıkarılmıştır.

- Manuel girilen üretim parametreleri (örneğin viskozite, pres, mil tipi) boya-baskı kartlarından alınarak Excel ortamında düzenlenmiştir.

Yinelenen Kayıtların Birleştirilmesi:

- Aynı üretim sürecine ait tekrarlayan satırlar birleştirilerek veri seti sadeleştirilmiştir.

- Bu işlem sonucunda nihai veri seti 3.839 satıra düşürülmüştür.

Kategorik Verilerin Kodlanması:

- Modelin çalışabilmesi için kategorik değişkenler sayısal değerlere dönüştürülmüştür.

- Etiket kodlama (label encoding) yöntemi kullanılmıştır.

Korelasyon Analizi:

- Değişkenler arasındaki ilişkiler Pearson korelasyon katsayısı ile değerlendirilmiştir.

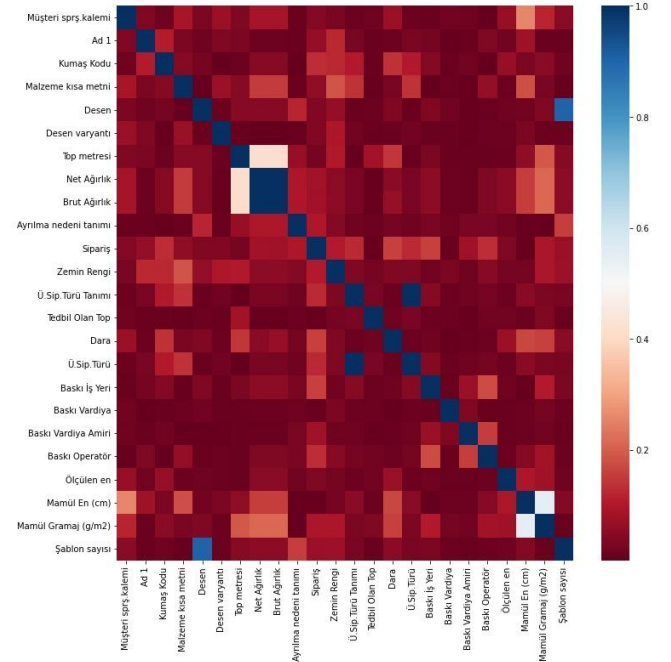
- Yüksek korelasyon gösteren değişkenler bir araya getirilmiş, düşük korelasyon düzeyine sahip olanlar veri setinden çıkarılmıştır.

- Korelasyon analizi ile başlangıçta 24 olan parametre sayısı (Şekil 3'de yer alan korelasyon ısı haritası ürün, üretim süreci, malzeme, çalışan ile ilgili 24 farklı parametreyi içermektedir) Tablo2'de gösterildiği gibi yedi parametreye indirgenmiştir. Tablo 2'de yer alan kumaş kalite göstergesi veri seti için etiket olarak belirlenmiştir.

- Bu sayede modelin daha anlamlı öğrenmesi ve daha etkili sonuçlar üretmesi hedeflenmiştir [22, 23].

3.3 Özellik tanımları

Modelde kullanılan parametrelerin açıklamaları ve kategorik dağılımları Tablo 2.'de verilmiştir.



Şekil 3. Korelasyon ısı haritası

Figure 3. Correlation heat map

Tablo 2. Parametrelerin tanımı

Table 2. Descriptions of parameters

Özellik ismi	Tanım	Özellik Tipi	Miktar
Kumaş Kodu	Üretimi yapılan kumaşın türünü gösteren kod	Kategorik	158
Kumaş Kalite Göstergesi	Hatalı ve hatasız üretimi gösteren özellik	Kategorik	5
Şablon Sayısı	Baskı yapılacak desene uygun şablon sayısı	Kategorik	12
Üretim Sipariş Türü	Baskıya hazır veya ham kumaş	Kategorik	2
Makine	Üretimde kullanılan makine	Kategorik	3
Vardiya	Üretimin gerçekleştiği vardiya	Kategorik	3
Vardiya Amiri	Üretimin takibini gerçekleştiren amir	Kategorik	4
Operatör	Üretimi gerçekleştiren operatör	Kategorik	11

Kök neden analizleri sonucu belirlenen üretim parametrelerinin (üretim viskozite, üretim pres, mil) SAP sistemi üzerinde kaydı tutulmadığından firmanın boya-baskı kartlarından alınan veriler Excele işlenmiştir.

3.4 Deneyler

Materyal ve yöntem bölümünde incelenen denetimli öğrenme yöntemlerinde, hataların sınıflandırılması ve analizi için aynı veri seti kullanılmıştır. Toplam 3839 satırlık veri setinin 3390 satırı hatasız üretim verilerinden oluşmaktadır. Kalıp kayması hatası için oluşturulan dört grubun toplam satır sayısı 449'dur. Etiket gruplarındaki dengesiz dağılımı dengelemek amacıyla üst örnekleme (oversampling) ve alt örnekleme (undersampling) teknikleri uygulanmıştır.

Kullanılan modellerin hiperparametre optimizasyonu, GridSearchCV ve RandomizedSearchCV yöntemleriyle gerçekleştirilmiştir. Eğitim ve test verisi oranını belirlemek için özel bir fonksiyon geliştirilmiş, farklı oranlar deneyerek en yüksek doğruluk değerini sağlayan oran tercih edilmiştir. Optimum hiperparametre değerleri Tablo 3'te sunulmuştur.

Denetimsiz öğrenme yöntemi olan K-ortalama (K-means) algoritması ile en uygun küme sayısını belirlemek amacıyla dirsek yöntemi, Silhouette skoru, Calinski-Harabasz indeksi ve Davies-Bouldin skoru analiz edilmiştir. Bu çoklu değerlendirme sonucunda optimum k değeri belirlenmiştir.

Modelin başarısı, uygulama sonuçlarında kesinlik (precision), doğruluk (accuracy) ve F1 skoru gibi performans metrikleriyle karşılaştırmalı olarak değerlendirilmiştir. Ayrıca, sınıflandırma modellerinin karar sınırlarını daha iyi öğrenebilmesi amacıyla veri uzayında maksimum derinlik (maximum depth) parametresi dikkate alınmış; bu parametre, modelin karmaşık örüntüleri ayırt etme kapasitesini artırmak için optimize edilmiştir.

Tablo 3. Veri seti özellikleri

Table 3. Data set properties

Kategori	Veri Sayısı
Hatasız Üretim	3390
Kalıp Kayması	449
Test/Eğitim Verisi Oranı	0.6/0.4
Maksimum Derinlik	15

3.5 Bulgular ve tartışma

Karar ağacı, xgboost, rassal orman ve catboost algoritmaları kullanılarak geliştirilen modellerin performansını ölçmek için F1, doğruluk ve keskinlik skorları hesaplanmıştır. Karmaşıklık matrisi, sınıflandırma algoritmalarının performansını ölçmek için kullanılan bir araçtır. Bu matris, gerçek ve tahmin edilen sınıflar arasındaki ilişkiyi gösteren dört temel değeri içerir:

- Gerçek pozitif (TP): Gerçekte sınıf 1 olan ve algoritma tarafından doğru olarak sınıf 1 olarak tahmin edilen örnekler.
- Yanlış pozitif (FP): Gerçekte sınıf 0 olan ve algoritma tarafından yanlış olarak sınıf 1 olarak tahmin edilen örnekler.
- Gerçek negatif (TN): Gerçekte sınıf 0 olan ve algoritma tarafından doğru olarak sınıf 0 olarak tahmin edilen örnekler.
- Yanlış negatif (FN): Gerçekte sınıf 1 olan ve algoritma tarafından yanlış olarak sınıf 0 olarak tahmin edilen örnekler.

Keskinlik (Precision): Pozitif bir sınıftaki toplam öngörülen biçimlerden doğru şekilde öngörülen pozitif kalıpları ölçmek için kullanılır [24].

F1-skor: Keskinlik ve hassasiyet kavramları üzerinden hesaplanır.

$$F1skor = 2 * \frac{(Hassasiyet * Duyarluluk)}{(Hassasiyet + Duyarluluk)} \quad (8)$$

Doğruluk: Doğru sınıflandırılan sınıfların tüm sınıflara oranını ifade etmektedir.

$$Doğruluk = \frac{(TP+TN)}{(TP+TN+FP+FN)} \quad (9)$$

Elde edilen veri setindeki etiket değerlerinin dengesiz yapısından dolayı bu skorların düşük olduğu gözlemlenmiştir. Dengesiz veri seti problemi hedef değişkenin etiketinde eşit bir dağılım bulunmamasıdır. Bu problemi çözmek için farklı yöntemler kullanılmaktadır. SMOTE en fazla tercih edilen yöntemdir [25]. Dolayısıyla modellerin performans değerlerini arttırmak ve en iyi performans değerine sahip modele karar vermek için veri setine alt örneklem, SMOTE ve üst örneklem uygulanmıştır. Yapılan bu uygulamalar veri setindeki etiket parametreleri arasındaki sınıf dengesizliğini gidererek modellerin performans değerlerinin artmasını sağlamıştır. Elde edilen tüm çıktılar değerlendirildiğinde en iyi performans metriklerine sahip modelin Tablo 4. ve Tablo 5'de görüldüğü gibi Catboost modeli olduğu belirlenmiştir. Böylelikle çalışmalara Catboost modeli ile devam edilmesine karar verilmiştir.

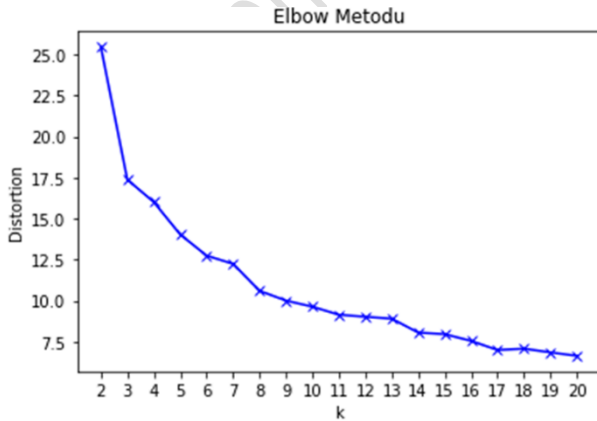
Tablo 4. Yöntemlerin karşılaştırılma tablosu
Table 4. Comparison table of methods

	Karar Ağacı Algoritması			Rassal Orman Algoritması		
	Doğruluk	F1 Skoru	Kesinlik	Doğruluk	F1 Skoru	Kesinlik
Normal	62.44%	62%	63%	80.53%	46%	47%
Alt Örneklem	62.50%	62%	70%	58.61%	58%	60%
SMOTE	72.42%	73%	62%	82.50%	82.72%	82%
Yeni Veri	62.96%	62.91%	63%	80.21%	82%	81%
Aşırı Örneklem	68.30%	66.20%	66%	85.80%	86.20%	87%

Tablo 5. Yöntemlerin Karşılaştırılma Tablosu
Table 5. Comparison table of methods

	XGBoost Algoritması			CatBoost Algoritması		
	Doğruluk	F1 Skoru	Kesinlik	Doğruluk	F1 Skoru	Kesinlik
Normal	88.89%	47%	45%	88.89%	47%	45%
Alt Örneklem	64.38%	64%	65%	64.38%	64%	65%
SMOTE	82.20%	82%	83%	82.20%	82%	83%
Yeni Veri	80.50%	81%	80.52%	84%	85%	84%
Aşırı Örneklem	80.86%	81.20%	81.30%	87.96%	88.12%	88.56%

Denetimsiz öğrenme yöntemlerinde K-ortalamalar modelinin kaç kümeye ayrılacağını belirlemek için Dirsek yöntemi kullanılmıştır. Bu yöntem her bir noktanın oluşan küme merkezlerine olan uzaklığının karelerinin toplamı ile hesaplanmaktadır, yöntemle göre uzaklığın karelerinin toplamındaki değişim miktarının azaldığı nokta dirsek noktasıdır. Bu nokta en iyi k küme sayısını ifade etmektedir [26]. Çalışmadaki dirsek yöntemi örneği Şekil 4'de yer almaktadır.



Şekil 4. Dirsek (Elbow) yöntemi grafiği
Figure 4. Elbow method graph

Modelin kaç kümeye ayrılacağını belirlemek, performanslarını ölçmek için Silhouette, Calinski-Harabasz, Davies-Bouldin farklı k değerleri için skor değerleri incelenmiş ve bu sonuçlara göre optimum k değeri belirlenmiştir.

Silhouette Skoru: Her bir veri noktasının kendi kümesi ile ne kadar iyi eşleştiğini ve diğer kümelerden ne kadar ayrıldığını ölçer. Aşağıdaki eşitlik ile hesaplanmaktadır:

$$s = \frac{(b-a)}{\max(a,b)} \quad (10)$$

a : Bir örnek ile aynı kümedeki diğer tüm noktalar arasındaki ortalama mesafe

b : Bir örnek ile en yakın kümedeki diğer tüm noktalar arasındaki ortalama mesafe

Elde edilen s değeri -1 ile 1 arasında bir sayıdır. 0.5-1 aralığı kümelemenin iyi olduğunu ifade etmektedir [28].

Calinski-Harabasz Skoru: K kümeye ayrılmış bir veri setinin kümeleme doğruluğunu değerlendirmek amacıyla aşağıdaki formül kullanılır:

$$ch = \frac{tr(B_k)x(n_g-k)}{tr(W_k)x(k-1)} \quad (11)$$

$tr(B_k)$: kümeler içi kareler toplamı

$tr(W_k)$: kümeler arası kareler toplamı

k : Toplam küme sayısı

n : Toplam veri noktası sayısı

En yüksek ch değeri en iyi kümeyi ifade eder [28].

Davies-Bouldin Skoru: Küme içerisinde bulunan verilerin merkeze olan uzaklığını minimum, kümeler arası uzaklığı ise maksimum yapmayı hedefleyen bu yöntem ile kümeleme geçerliliğini ölçmek için aşağıdaki eşitlik kullanılır:

$$db = \frac{1}{k} \sum_{i=1}^k \max \left(\frac{S_i + S_j}{d_{ij}} \right) \quad (12)$$

d_{ij} : kümelerde bulunan merkezler arasındaki mesafe

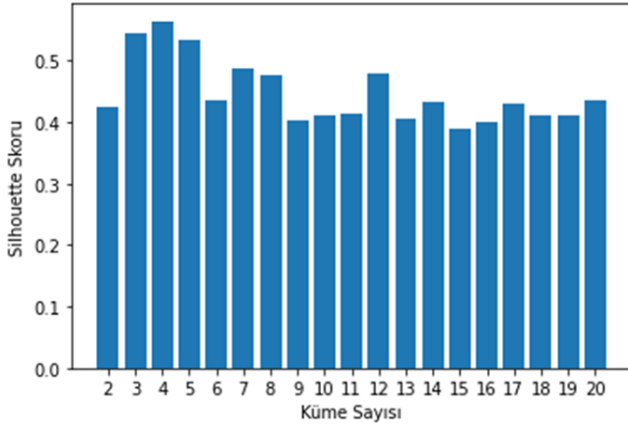
S_i & S_j : küme gözlemlerinin bulunduğu kümenin merkezlerine olan ortalama mesafe

Daha düşük db indeksi değeri, daha iyi kümeleme performansına işaret etmektedir [29].

Farklı küme merkezleri için performans metriklerinden faydalanılarak küme değeri 5 olarak belirlenmiş ve skor değerleri Tablo 6'da gösterilmiştir. Şekil 5 ve Şekil 6'daki Silhouette ve Davies-Bouldin skorlarına ait histogramlardan faydalanılarak teyit edilmiştir.

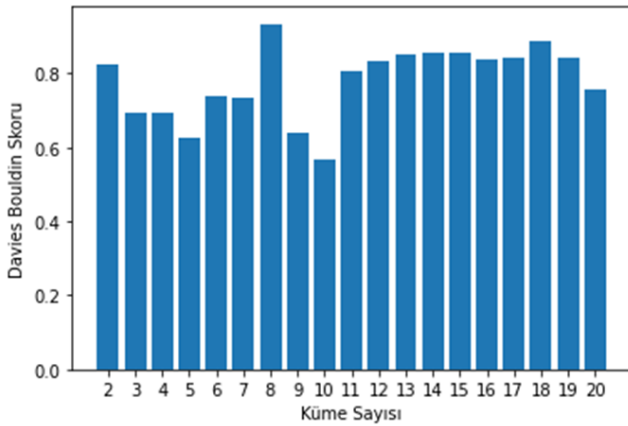
Tablo 6. Kümeleme performans çıktıları
Table 6. Clustering performance outputs

k	Silhouette	Davies-Bouldin	Calinski-Harabasz
4	0.53	0.63	580
5	0.58	0.60	678
6	0.48	0.90	426



Şekil 5. Farklı küme sayıları için Silhouette skoru

Figure 5. Silhouette score for different numbers of clusters



Şekil 6. Farklı küme sayıları için Davies-Bouldin Skoru

Figure 6. Davies-Bouldin Score for different numbers of cluster

3.6 Uygulama

Bu çalışmada geliştirilen tahminleme modeli, Python programlama dili kullanılarak oluşturulmuş ve kullanıcı etkileşimi için Tkinter kütüphanesi aracılığıyla grafiksel bir arayüz tasarlanmıştır. Arayüz, kullanıcıların üretim sürecine ilişkin belirli parametreleri manuel olarak girmesine olanak tanımakta ve bu parametreler doğrultusunda kumaş kalite sınıfı tahmin edilmektedir.

Kullanıcı arayüzü üç temel bileşenden oluşmaktadır:

Giriş ekranı (Şekil 7.): Kullanıcıların modelin ihtiyaç duyduğu üretim parametrelerini (örneğin; kumaş kodu, desen varyantı, metraj, ağırlık vb.) girdiği bölümdür. Bu ekran, kullanıcı dostu bir yapı ile tasarlanmış olup, her parametre için ayrı giriş alanları içermektedir.

Tahminleme Ekranı (Şekil 8): Tablo 2’de belirtilen üretim parametrelerinin eksiksiz şekilde girilmesinin ardından, “Giriş” butonuna basılarak tahminleme modeli çalıştırılır. Sistem, kullanıcıdan alınan verileri makine öğrenmesi algoritması aracılığıyla analiz ederek kumaşın kalite sınıfını öngörür. Tahmin edilen kalite sınıfı, firma bünyesinde tanımlı dört farklı kalite seviyesinden biridir. Bu sınıflar arasında 0 en yüksek kaliteyi, 3 ise en düşük kaliteyi temsil etmektedir.

Eğer tahmin edilen kalite sınıfı 0’dan farklıysa, sistem geçmiş üretim verilerini analiz ederek benzer koşullar altında en yüksek kaliteyi sağlayan optimum parametre kombinasyonlarını belirler ve kullanıcıya önerir. Bu yapı sayesinde, hem üretim sürecinde oluşabilecek kalite düşüşleri önceden tespit edilebilmekte hem de kaliteyi artırmaya yönelik veri destekli iyileştirme önerileri sunulmaktadır.

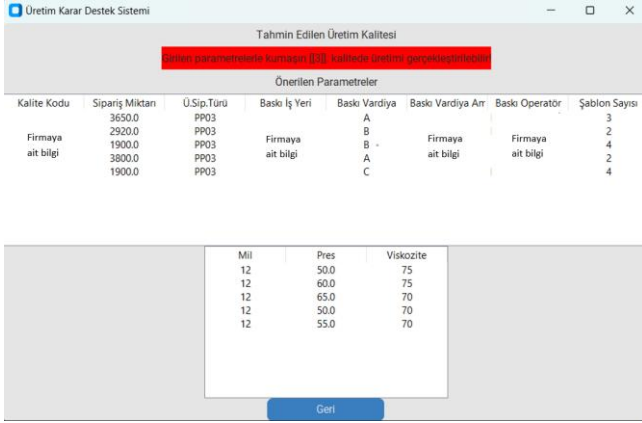
Sonuç ekranı (Şekil 9.): Tahminleme süreci tamamlandıktan sonra, kullanıcıya hem tahmin edilen kalite sınıfı hem de bu kaliteye ulaşmak için önerilen parametre seti sunulur. Bu çıktı, hem metin hem de görsel bileşenlerle desteklenerek kullanıcıya anlaşılır biçimde aktarılır.

Şekil 7. Kullanıcı girişi ekranı

Figure 7. User login screen

Şekil 8. Parametre giriş ekranı

Figure 8. Parameter entry display

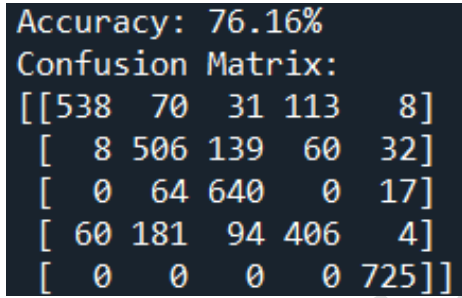


Şekil 9. Sonuç Ekranı

Figure 9. Result Screen

Doğruluk, kesinlik, F1-skoru metrikleri kullanılarak performansları değerlendirilen modeller eğitim ve testte kullanılmayan farklı verilerle test edilerek değerler karşılaştırılmıştır. Şekil 10'da farklı veri setine ait sonuçlar karmaşıklık matrisinde gösterilmiştir.

Karmaşıklık Matrisi (Confusion Matrix), özellikle sınıflandırma problemlerinde bir makine öğrenmesi modelinin ne kadar doğru ve nerede hata yaptığını anlamak için kullanılan temel bir değerlendirme tablosudur.



Şekil 10. Farklı veri seti ile elde edilen karmaşıklık matrisi

Şekil 10. Complexity matrix obtained with different datasets

Modelin belirlenen kumaş kalite göstergesi etiketlerini %76 dolaylarında doğru tahmin ettiği görülmektedir. Test veri setinden yaklaşık %10'luk bir sapma bulunmaktadır.

4 Sonuçlar

Bu çalışmada, tekstil üretim süreçlerindeki hataları önceden tahmin etmek ve önlemek amacıyla makine öğrenmesi tabanlı bir karar destek sistemi geliştirilmiştir. Üretim verileri, bir tekstil işletmesinden alınmış ve hem denetimli hem de denetimsiz öğrenme algoritmalarıyla analiz edilmiştir.

Denetimli öğrenmede en başarılı sonuçlar CatBoost algoritmasıyla elde edilmiş ve sistemin temel tahmin modeli olarak seçilmiştir. Denetimsiz öğrenmede ise K-means algoritması kullanılmış, optimum küme sayısı k=5 olarak belirlenmiştir.

Bu iki yaklaşım entegre edilerek, sistem hem olası üretim hatalarını önceden tahmin edebilmekte hem de riskli durumlarda geçmişte başarılı olmuş parametre kombinasyonlarını önererek kaliteyi artırmaktadır.

Kullanıcı dostu arayüzü sayesinde operatörler üretim parametrelerini kolayca girebilmekte; sistem, kumaş kalite

sınıfını tahmin ederek düşük kalite riski varsa en uygun parametreleri önermektedir. Yapılan testlerde modelin yaklaşık %76 doğruluk ile kalite tahmini yaptığı görülmüştür.

Sistem, özellikle baskı ve boyama gibi hata riski yüksek aşamalarda etkili olup, SAP gibi kurumsal sistemlerle entegre çalışabilmekte ve gerçek zamanlı veri akışıyla anlık geri bildirim sunabilmektedir.

Gelecekte, sensör verileriyle gerçek zamanlı entegrasyon ve operatör geri bildirim mekanizması eklenerek sistemin doğruluğu ve adaptif öğrenme yeteneği daha da geliştirilebilir.

5 Conclusions

In this study, a machine learning-based decision support system was developed to predict and prevent errors in textile production processes. Production data were obtained from a textile company and analyzed using both supervised and unsupervised learning algorithms.

Among the supervised learning models tested, the CatBoost algorithm yielded the highest performance in terms of accuracy, precision, and F1-score, and was therefore selected as the core predictive model of the system. For unsupervised learning, the K-means algorithm was employed, and the optimal number of clusters was determined as k=5 using the Silhouette, Calinski-Harabasz, and Davies-Bouldin indices.

By integrating these two approaches, the system is capable of both forecasting potential production errors and recommending parameter combinations that have previously led to successful outcomes under similar conditions, thereby enhancing overall product quality.

Thanks to its user-friendly interface, operators can easily input production parameters. Based on these inputs, the system predicts the fabric quality class; if the predicted class indicates a risk of low quality, it suggests the most suitable parameter set derived from historical production data. Tests conducted with various datasets demonstrated that the model can predict fabric quality with approximately 76% accuracy.

The system is particularly effective in high-risk stages such as printing and dyeing, and can be integrated with enterprise platforms like SAP, enabling real-time data flow and instant feedback for operators. This contributes to the digitalization of quality control processes and accelerates decision-making on the production line.

In future work, the system's accuracy and adaptive learning capabilities could be further enhanced through real-time integration with sensor data and the implementation of an operator feedback mechanism.

6 Yazar katkı beyanı

Yazar isimleri içerdiğinden değerlendirme sonrası eklenecektir.

7 Etik kurul onayı ve çıkar çatışması beyanı

"Hazırlanan makalede etik kurul izni alınmasına gerek yoktur".
"Hazırlanan makalede herhangi bir kişi/kurum ile çıkar çatışması bulunmamaktadır".

8 Kaynaklar

- [1] Şanlıtürk E. Makine Öğrenme Algoritmalarıyla Hatalı Ürün Tahmini. MSc Thesis, İstanbul Teknik Üniversitesi, İstanbul, Türkiye, 2018.
- [2] Demir E, Dinçer SE. "Determining the Faulty and Refund Products in Manufacturing System: Application on a Textile Firm". Research Journal of Business and Management, 7(3), 178-187, 2020.
- [3] Soma S, Pooja H. "Machine Learning System for Textile Fabric Defect Detection Using GLCM Technique". In: Proceedings of Second International Conference on Advances in Computer Engineering and Communication Systems: ICACECS 2021, 171-182, Singapore: Springer Nature Singapore, 2022.
- [4] Çıklaçandır FGY, Utku S, Özdemir H. "Kumaşlarda Hatayı Yerel Olarak Arayan Denetimsiz Bir Sistem". Tekstil ve Mühendis, 27(120), 252-259, 2020.
- [5] Güvenoğlu E, Bağırman M. "Shearlet Dönüşümü ve Görüntü İşleme Teknikleri Kullanılarak Kot Kumaşlar Üzerinde Gerçek Zamanlı Hata Tespiti". El-Cezeri, 6(3), 491-502, 2019.
- [6] Bhattacharya S, Das PP, Chatterjee P, Chakraborty S. "Prediction of Responses in a Sustainable Dry Turning Operation: A Comparative Analysis". Mathematical Problems in Engineering, 2021(1), 9967970, 2021.
- [7] Altunsöğüt Ö. Tekstil Boyahane İşletmelerinde Makine Öğrenmesi ile Fire Tahmini. PhD Thesis, Trakya Üniversitesi, Edirne, Türkiye, 2022.
- [8] Chao WL. Machine Learning Tutorial. 2011.
- [9] Caruana R, Niculescu-Mizil A. "An Empirical Comparison of Supervised Learning Algorithms". In: Proceedings of the 23rd International Conference on Machine Learning, 161-168, ACM, 2006.
- [10] Daş B, Türkoğlu İ. "DNA Dizilimlerinin Sınıflandırılmasında Karar Ağacı Algoritmalarının Karşılaştırılması". 2014.
- [11] Çelik S, Bozkurt ÖÇ, Ekşili N. "Çalışan Performansı Ölçeğindeki İfadelerin Karar Ağacı Algoritması ile Belirlenmesi". Mehmet Akif Ersoy Üniversitesi İktisadi ve İdari Bilimler Fakültesi Dergisi, 9(1), 561-584, 2022.
- [12] Ülgen T, El Sayed A, Poyrazoğlu G. "Prediction Algorithm & Learner Selection for European Day-Ahead Electricity Prices". 2020 2nd Global Power, Energy and Communication Conference (GPECOM), Antalya, Türkiye, 285-286, IEEE, 8-10 Ekim 2020.
- [13] Breiman L. "Random Forests". Machine Learning, 45, 5-32, 2001.
- [14] Liao Z, Ju Y, Zou Q. "Prediction of G Protein-Coupled Receptors with SVM-Prot Features and Random Forest". Scientifica, 2016(1), 8309253, 2016.
- [15] Coşar M, Deniz E. "Makine Öğrenimi Algoritmaları Kullanarak Kalp Hastalıklarının Tespit Edilmesi". Avrupa Bilim ve Teknoloji Dergisi, (28), 1112-1116, 2021.
- [16] Chen T, Guestrin C. "XGBoost: A Scalable Tree Boosting System". Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, 785-794, San Francisco, USA, 13-17 Ağustos 2016.
- [17] XGBoost Documentation. "Introduction to Boosted Trees". <https://xgboost.readthedocs.io/en/stable/tutorials/model.html> (14.12.2023).
- [18] Kuş İ, Keser SB, Yolaçan E. "Saldırı Tespit Sistemlerinde Topluluk Öğrenme Yöntemlerinin Kıyaslanması". Avrupa Bilim ve Teknoloji Dergisi, (31), 725-734, 2021.
- [19] Aydın, Z. E., Erdem, B. İ., Cıcek, Z. I. E. "Prediction bike-sharing demand with gradient boosting methods". Pamukkale Üniversitesi Mühendislik Bilimleri Dergisi, 29(8), 824-832, 2023.
- [20] Bilgin M. "Gerçek Veri Setlerinde Klasik Makine Öğrenmesi Yöntemlerinin Performans Analizi". Breast, 2(9), 683, 2017.
- [21] Erdoğmuş P, Çolak B, Durdağ Z. "K-Means Algoritması ile Otomatik Kümeleme". El-Cezeri, 3(2), 2016.
- [22] Liyeu CM, Di Nardo E, Ferraris S, Meo R. "Hyperparameter Optimization of Machine Learning Models for Predicting Actual Evapotranspiration". Machine Learning with Applications, 20, 100661, 2025.
- [23] Souza, Afd, Verri FAN, Campos PhDs, Silva PHd, Balestrassi PP. "Predicting tool life and sound pressure levels in dry turning using machine learning models". International Journal of Advanced Manufacturing Technology, 135, 3777-3793, 2024.
- [24] Hossin M, Sulaiman MN. "A Review on Evaluation Metrics for Data Classification Evaluations". International Journal of Data Mining & Knowledge Management Process, 5(2), 1, 2015.
- [25] Akın P, Terzi Y. "Dengesiz Veri Setli Sağkalım Verilerinde Cox Regresyon ve Rastgele Orman Yöntemlerin Karşılaştırılması". Veri Bilimi Dergisi, 3(1), 21-25, 2020.
- [26] Aslanyürek M, Mesut A. "Kümeleme Performansını Ölçmek İçin Yeni Bir Yöntem ve Metin Kümeleme İçin Değerlendirmesi". Avrupa Bilim ve Teknoloji Dergisi, (27), 53-65, 2021.
- [27] Aslanyürek M, Mesut A. "Kümeleme Performansını Ölçmek İçin Yeni Bir Yöntem ve Metin Kümeleme İçin Değerlendirmesi". Avrupa Bilim ve Teknoloji Dergisi, (27), 53-65, 2021.
- [28] Caliński T, Harabasz J. "A Dendrite Method for Cluster Analysis". Communications in Statistics - Theory and Methods, 3(1), 1-27, 1974.
- [29] Davies DL, Bouldin DW. "A Cluster Separation Measure". IEEE Transactions on Pattern Analysis and Machine Intelligence, (2), 224-227, 1979.