



ANKARA
HACI BAYRAM VELİ ÜNİVERSİTESİ
UZAKTAN EĞİTİM UYGULAMA VE
ARAŞTIRMA MERKEZİ

YAPAY ZEKÂ TEKNOLOJİLERİ VE YÖNETİMİ

4. Makine Öğrenmesi

4.1. Öğrenme Kavramı

Öğrenme, bir bireyin veya sistemin deneyim ve bilgi edinme süreçleri aracılığıyla bilgi ve becerilerini geliştirmesi ve adaptif davranışlar sergilemesi sürecidir. Bu süreç, tekrarlama, deneyimleme ve geri bildirim yoluyla gerçekleşir ve genellikle bilgi işleme, analiz yapma ve sonuçları değerlendirme aşamalarını içerir. İnsanlarda öğrenme, genellikle okul, deneyim ve sosyal etkileşimlerle gerçekleşirken, yapay zeka sistemlerinde öğrenme, veriler üzerinden analiz ve modelleme yoluyla gerçekleşir. Öğrenme, hem bireysel hem de sistem düzeyinde performansın artırılması, problem çözme yeteneklerinin geliştirilmesi ve çevresel değişikliklere uyum sağlanması açısından kritik bir rol oynar.

4.2. Makine Öğrenmesi Nedir

Makine öğrenmesi, bilgisayarların deneyimlerinden öğrenme ve performanslarını iyileştirme yeteneğini geliştiren bir yapay zeka dalıdır. Makine öğrenmesi, bilgisayar sistemlerinin verilerden öğrenmelerine ve bu öğrenmelerini kullanarak belirli görevlerde daha iyi sonuçlar elde etmelerine olanak tanır. Bu süreç, algoritmalar ve istatistiksel modeller kullanılarak gerçekleştirilir. Makine öğrenmesinin uygulama alanları oldukça geniş ve sağlık, finans, pazarlama, ulaşım ve daha birçok sektörde önemli rol oynar. Örneğin, sağlık sektöründe, makine öğrenmesi algoritmaları hastalıkların erken teşhisinde ve kişiselleştirilmiş tedavi planlarının geliştirilmesinde kullanılabilir. Finans sektöründe, dolandırıcılık tespiti ve risk yönetimi gibi görevlerde etkili olabilir. Ayrıca, makine öğrenmesi, müşteri davranışlarını analiz etme, öneri sistemleri oluşturma ve doğal dil işleme gibi alanlarda da geniş bir kullanım yelpazesi sunar. Bu teknolojinin sürekli gelişimi, daha doğru tahminler ve daha etkili çözümler sağlama potansiyelini artırmaktadır.

Makine öğrenmesi genellikle üç ana kategoriye ayrılır: denetimli öğrenme (supervised learning), denetimsiz öğrenme (unsupervised learning) ve takviye öğrenme (reinforcement learning).

4.3. Denetimli Öğrenme

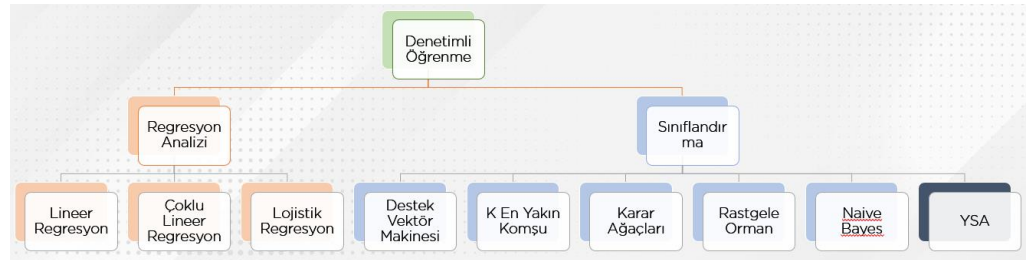
Denetimli öğrenme (Supervised Learning), makine öğrenmesinde, bir modelin etiketli veri setleri kullanılarak eğitildiği bir öğrenme yöntemidir. Bu yöntemde, eğitim verileri, hem giriş verilerini hem de bu verilerin doğru etiketlerini içerir.

Model, bu veriler üzerinde eğitim alarak, giriş verilerini etiketlerle eşleştirmeyi öğrenir. Eğitim süreci boyunca model, doğru etiketlere ulaşmak için çeşitli algoritmalar ve matematiksel teknikler kullanarak, giriş-çıkış ilişkilerini öğrenir. Bu öğrenme sürecinin sonucunda, model, yeni, etiketlenmemiş veriler üzerinde tahminlerde bulunabilme yeteneğine sahip olur.

Denetimli öğrenme genellikle sınıflandırma ve regresyon problemlerinde kullanılır. Sınıflandırma problemlerinde, model, verileri belirli sınıflara ayırmayı öğrenir. Örneğin, bir e-posta filtreleme sistemi, e-postaları "spam" veya "spam değil" olarak sınıflandırabilir. Regresyon problemlerinde ise, model, sürekli bir değişkeni tahmin eder. Örneğin, bir konut fiyat tahmin modeli, evlerin özelliklerine göre fiyat tahminleri yapabilir. Bu tür problemler, denetimli öğrenmenin uygulanabileceği geniş bir yelpazeyi kapsar ve çeşitli sektörlerde veri analizi ve karar destek sistemlerinde kullanılır.

Denetimli öğrenmenin etkinliği, kullanılan eğitim verilerinin kalitesine ve miktarına bağlıdır. Yüksek kaliteli, doğru etiketlenmiş veriler, modelin doğru öğrenmesini ve yüksek performans göstermesini sağlar. Ancak, bu yöntem verilerin etiketlenmesi ve toplanması açısından zaman ve maliyet gerektirebilir. Ayrıca, denetimli öğrenme yöntemleri genellikle belirli bir görev veya problem için optimize edilmiş modeller üretir; bu nedenle, modelin genelleştirme yeteneği, eğitildiği veri kümesinin çeşitliliğine ve kapsamına bağlı olarak değişebilir.

4.4. Denetimli Öğrenme Yaklaşımları



4.4.1 Regresyon Analizi

Regresyon analizi, iki veya daha fazla değişken arasındaki ilişkileri modellemek ve tahminlerde bulunmak amacıyla kullanılan bir istatistiksel tekniktir. Bu yöntem, bir bağımlı değişkenin (hedef değişken) bir veya daha fazla bağımsız değişken (predictor değişkenler) tarafından nasıl etkilendiğini belirlemeye çalışır. Regresyon analizi, ilişkilerin doğrusal mı yoksa doğrusal olmayan mı olduğunu anlamak için kullanılır ve genellikle en küçük kareler yöntemiyle modelin en iyi uyum

sağladığı parametrik değerler belirlenir. Bu teknik, iş dünyasında, ekonomide ve sosyal bilimlerde, verilerin analiz edilmesi ve gelecekteki sonuçların tahmin edilmesi için geniş bir uygulama yelpazesi sunar. Örneğin, bir şirketin satışlarının reklam harcamalarına bağlı olarak nasıl değiştiğini incelemek için regresyon analizi yapılabilir.

4.4.1.1 Lineer Regresyon

Lineer regresyon, bir bağımlı değişken ile bir veya daha fazla bağımsız değişken arasındaki doğrusal ilişkiyi modelleyen bir regresyon analizidir. Bu yöntem, bağımlı değişkenin (örneğin, ev fiyatı) bağımsız değişkenlere (örneğin, evin büyüklüğü, konumu) bağlı olarak nasıl değiştiğini belirlemeye çalışır. Lineer regresyon, veriler arasındaki ilişkiyi en iyi şekilde temsil eden bir doğrusal denklem (doğru) oluşturur ve bu doğrultuda tahminlerde bulunur. Bu doğrusal model genellikle, en küçük kareler yöntemi kullanılarak parametrelerin (eğim ve kesişim) tahmin edilmesiyle oluşturulur. Lineer regresyon, basit ve anlaşılır yapısı sayesinde çok yaygın olarak kullanılır ve hem temel analizlerde hem de daha karmaşık modelleme süreçlerinde etkili bir araçtır.

4.4.1.2 Çoklu Lineer Regresyon

Çoklu lineer regresyon, bir bağımlı değişkenin birden fazla bağımsız değişkenle doğrusal ilişkisini modelleyen bir regresyon analizidir. Bu yöntem, bağımlı değişkenin (örneğin, bir evin fiyatı) birden fazla bağımsız değişken (örneğin, evin büyüklüğü, odaların sayısı, konumu) tarafından nasıl etkilendiğini anlamaya çalışır. Çoklu lineer regresyon, bu bağımsız değişkenlerin her birinin bağımlı değişken üzerindeki etkisini ayrı ayrı değerlendirir ve bu etkilerin birleşimini modelin tahminlerinde kullanır. Model, genellikle en küçük kareler yöntemiyle, bağımsız değişkenlerin katsayılarını belirleyerek en uygun doğrusal denklemi oluşturur. Bu yaklaşım, çeşitli faktörlerin etkisini analiz etme ve karmaşık ilişkileri ortaya koyma konusunda güçlü bir araçtır, ayrıca çoklu değişkenlerin etkilerini aynı anda değerlendirme yeteneği sunar.

4.4.1.3 Lojistik Regresyon

Lojistik regresyon, bağımlı değişkenin kategorik olduğu durumlarda kullanılan bir regresyon analizidir ve genellikle ikili sınıflandırma problemlerinde tercih edilir. Bu yöntem, bağımsız değişkenlerle (predictor değişkenler) bağımlı değişken arasında doğrusal bir ilişki kurarak, bağımlı değişkenin belirli sınıflara ait olma olasılığını tahmin eder. Lojistik regresyon, doğrusal regresyondan farklı olarak, olasılıkları modellemek için lojistik

fonksiyon (sigmoid fonksiyon) kullanır; bu, tahmin edilen değerlerin 0 ile 1 arasında sınırlandırılmasını sağlar. Model, genellikle maksimum olabilirlik tahminleri kullanılarak parametrelerini belirler ve bu sayede bir olayın belirli bir sınıfa ait olma olasılığını değerlendirir. Lojistik regresyon, sağlık, finans ve sosyal bilimler gibi birçok alanda geniş bir uygulama yelpazesi sunar.

4.4.2 Sınıflandırma

Sınıflandırma, verileri belirli kategorilere veya sınıflara ayırma sürecidir ve genellikle etiketli veri setleri kullanılarak gerçekleştirilir. Bu süreç, bir modelin, giriş verilerini analiz ederek ve öğrenerek hangi sınıfa ait olduğunu tahmin etmesini sağlar. Sınıflandırma problemleri genellikle ikili (iki sınıf) veya çoklu (birden fazla sınıf) seçenekler içerir. Örneğin, bir e-posta filtreleme sistemi, e-postaları "spam" veya "spam değil" olarak sınıflandırabilir; bir sağlık tarama sistemi, hastalık riskini "yüksek" veya "düşük" olarak sınıflandırabilir. Sınıflandırma, çeşitli makine öğrenmesi algoritmaları ve yöntemleri kullanılarak gerçekleştirilir, bu algoritmalar verilerdeki örüntüleri öğrenir ve bu öğrenme sürecinden elde edilen bilgiyle yeni verileri sınıflandırır.

4.4.2.1 Destek Vektör Makineleri

Destek Vektör Makineleri (SVM), sınıflandırma ve regresyon problemlerinde kullanılan güçlü bir makine öğrenmesi algoritmasıdır. SVM'nin temel amacı, verileri farklı sınıflara ayıran en iyi hiper düzlemi (hiperplane) bulmaktır. Bu hiper düzlem, veriler arasındaki sınırı belirler ve veri noktalarını farklı sınıflara ayırır. SVM, bu hiper düzlemin, iki sınıf arasındaki en geniş marjini (mesafeyi) sağladığı noktada konumlanmasını hedefler. Bu, genel doğruluk oranını artırır ve modelin genelleme yeteneğini iyileştirir.

SVM, verilerin doğrusal olarak ayrılabilir olmadığı durumlarda da etkili olabilir. Bu tür durumlarda, SVM verileri daha yüksek bir boyuta dönüştürerek doğrusal ayrılabilirlik sağlar. Bu dönüşüm, çekirdek fonksiyonları (kernel functions) adı verilen matematiksel araçlarla yapılır. Çekirdek fonksiyonları, verilerin yüksek boyutlu bir uzaya projeksiyonunu sağlar ve burada doğrusal ayırım mümkün hale gelir. Popüler çekirdek fonksiyonları arasında polinom çekirdekleri, Gauss (RBF) çekirdekleri ve sigmoid çekirdekler bulunur.

SVM'nin avantajlarından biri, overfitting (aşırı uyum) riskini minimize etme yeteneğidir. Model, marjın maksimize etme hedefiyle çalıştığı için, genellikle daha az kompleks ve daha genel bir model oluşturur. Bununla birlikte, SVM'nin büyük veri setlerinde hesaplama maliyeti yüksek olabilir ve çekirdek fonksiyonlarının doğru seçilmesi gerekmektedir. Ayrıca, modelin performansı, hiperparametrelerin (marjın, çekirdek fonksiyonu seçimi vb.) uygun şekilde ayarlanmasına bağlıdır.

SVM, çeşitli uygulama alanlarında etkili bir şekilde kullanılabilir. Örneğin, metin sınıflandırma, yüz tanıma, biyomoleküler veri analizi ve daha birçok problemde başarıyla uygulanmıştır. Çeşitli çekirdek fonksiyonları ve düzenleme parametreleri ile özelleştirilebilen bu algoritma, güçlü performansı ve esnekliği ile geniş bir uygulama yelpazesi sunar.

4.4.2.2 K En Yakın Komşu

K en yakın komşu (KNN), sınıflandırma ve regresyon problemlerinde kullanılan basit ama etkili bir makine öğrenmesi algoritmasıdır. KNN'nin temel prensibi, bir veri noktasının sınıfını veya değerini belirlemek için en yakın komşularının bilgilerini kullanmaktır. Sınıflandırma problemlerinde, bir veri noktasının sınıfı, belirli bir k sayıda en yakın komşunun çoğunlukla belirlediği sınıf ile belirlenir. Regresyon problemlerinde ise, veri noktasının tahmin edilen değeri, en yakın komşularının değerlerinin ortalaması alınarak belirlenir.

KNN'nin en büyük avantajlarından biri, modelin basitliğidir. Eğitim aşamasında herhangi bir parametre öğrenme süreci olmadığından, algoritma eğitim sırasında veri setine oldukça hızlı uyum sağlar. Ancak, KNN'nin performansı, veri setinin büyüklüğüne ve k değerine oldukça bağlıdır. K değeri, komşu sayısını belirtir ve genellikle çapraz doğrulama gibi yöntemlerle optimize edilir. Küçük k değerleri, modelin gürültüye karşı hassas olmasına neden olabilirken, büyük k değerleri modelin yeterince ayırmacı olmamasına yol açabilir.

KNN'nin hesaplama maliyeti, genellikle tahmin aşamasında artar. Çünkü her tahminde, algoritma tüm eğitim verileri ile uzaklık hesaplamaları yapar, bu da büyük veri setlerinde zaman alıcı olabilir. Verimliliği artırmak için, verilerin uzayda organize edilmesi ve hızlı arama yapabilen veri yapılarının (örneğin, KD-ağaçları) kullanılması gibi yöntemler uygulanabilir.

KNN, çeşitli uygulama alanlarında etkili bir şekilde kullanılabilir. Örneğin, yüz tanıma, el yazısı tanıma ve tıbbi teşhisler gibi alanlarda sıkça uygulanır. Bununla birlikte, verinin boyutunun

büyük olduğu ve özelliklerin çok fazla olduğu durumlarda, KNN'nin performansını artırmak için veri ön işleme ve boyut indirgeme teknikleri kullanılması önerilir. KNN, uygulama alanına göre dikkatlice ayarlanmış k değeri ve uygun uzaklık ölçütleri ile etkili sonuçlar sağlayabilir.

KAYNAKÇA

- Russell, S. J., & Norvig, P. (2016). Artificial intelligence: a modern approach. Pearson.
- Burkov, A. (2019). The hundred-page machine learning book (Vol. 1, p. 32).
- Quebec City, QC, Canada: Andriy Burkov.
- Smola, A. (2008). Introduction to machine learning. Link: <https://alex.smola.org/drafts/thebook.pdf>

BÖLÜM SORULARI

- 1) Makine öğrenmesi nedir ve bilgisayarların verilerden öğrenmesini nasıl sağlar? Temel çalışma prensiplerini açıklayınız.
- 2) Denetimli öğrenmenin ne olduğunu ve nasıl çalıştığını açıklayınız.
- 3) Destek Vektör Makineleri (SVM) nedir ve nasıl çalışır? SVM'nin sınıflandırma ve regresyon analizinde nasıl kullanıldığını açıklayınız.
- 4) K-NN algoritmasının avantajları ve dezavantajları nelerdir? Büyük veri kümeleri ve yüksek boyutlu verilerde karşılaşılan zorlukları tartışınız.