

AI Report

We predict this text is

Human Generated

AI Probability

0%

This number is the probability that the document is AI generated, not a percentage of AI text in the document.

Plagiarism



The plagiarism scan was not run for this document. Go to gptzero.me to check for plagiarism.

Eğitsel Yapay Zeka (AI-ED) Sistemlerinde _Algoritmik Şeffaflık_ ve Öğrenci Özerkliği_ Karar Mekanizmalarının Etik Denetimi_ (2).docx - 5/12/2026

Eğitsel Yapay Zeka (AI-ED) Sistemlerinde "Algoritmik Şeffaflık" ve Öğrenci Özerkliği: Karar Mekanizmalarının Etik Denetimi

Beyza Naz Kavuş

Bilgisayar ve Öğretim Teknolojileri Eğitimi, Marmara Üniversitesi

Özet

Eğitimsel yapay zeka sistemleri, kişiselleştirilmiş öğrenme deneyimleri sunma konusunda büyük bir potansiyele sahip olsa da, bu sistemlerin karar alma süreçleri kullanıcılardan gizlenmektedir, bu da "kara kutu" sorununa ve "algoritmik paternalizm" riskine yol açmaktadır.

İncelendiğinde, buradaki sistemin öğrenci adına kararlar aldığı ancak bu kararların gerekçelerini sunmadığı, bunun sonucunda da öğrenci otonomisinin zayıfladığı görülmektedir.

Bu çalışma, ilgili etik problem için çözüm olarak geliştirilen "Etik Bildirim ve Onay Modeli"ni sunmaktadır.

Model XAI (Açıklanabilir Yapay Zeka: Explainable Artificial Intelligence) prensiplerine dayanan 2 aşamalı bir yapı önermektedir: 1. aşamada yapay zekanın yönlendirmeleri somut veriler ile açıklanmakta ve öğrenciye "şeffaflık etiketleri" aracılığı ile sunulmaktadır.

İkinci aşamada ise, geliştirilen "özerklik kontrol mekanizması" ile öğrenciye sistemin kararlarına itiraz etme ya da alternatif öğrenme yollarını seçme hakkı verilmektedir.

Bu model, literatürdekine benzer güven ve şeffaflık dengesine sahiptir; sistem odaklı karar verme sürecinin yapay zekanın bir "otorite" olarak konumlanması öğrenciden çıkarması ve bunun yerine kararlarını açıklayan, öğrenciyle karşılıklı etkileşimde olan şeffaf bir "asistan" rolüne getirmesi için tasarlanmıştır.

Bu dönüşüm, öğrencinin karşısında pasif bir öğrenici yerine kararlarının bilinçli takipçisi ve uygulayıcısı olan aktif öğrenici konumuyla güçlenmesini sağlamayı hedeflemektedir.

Anahtar Kelimeler: Eğitsel Yapay Zeka, Etik Bildirim, Algoritmik Şeffaflık, Öğrenci Otonomisi, Açıklanabilir Yapay Zeka (XAI).

Giriş

Günümüzde yapay zeka yardımıyla eğitim programları, her öğrenci için bireysel ihtiyaçlarına göre uyarlanmış özel bir ders programı oluşturabilir.

Bu sistemler, bir öğrencinin öğrenme hızı ve performans düzeyi gibi bilgileri analiz ederek, hangi konuların hangi zorluk seviyesinde öğrenilmesi gerektiğini belirler.

Ancak süreç içerisinde yapay zeka, öğrenci hakkında önemli kararlar alırken kararlarının mantığını kullanıcıyla paylaşmaz.

Öğrenci, neden belirli bir seviyede tutulduğunu, önüne neden belirli içeriklerin konduğunu bile bilmeden sistemi takip etmek zorunda kalmaktadır.

Bu durum, yapay zekanın destekleyici rolünü aşarak, eğitim sürecinde öğrencinin yolculuğunu tek başına yönlendiren kapalı bir sisteme dönüşmesine neden olur.

Kararların açıklanmaması, öğrencinin kendi öğrenme süreci üzerindeki onay hakkını elinden alır.

Bu yazı, eğitimin tüm aşamalarında yapay zekanın kullanılacağı bir ortamda eğitime yön veren yapay zeka kararlarının nasıl şeffaf hale getirilebileceğini ve öğrencinin bu sürece nasıl taraf olabileceğini incelemektedir.

Temel hedef, teknolojinin öğrenciyi yönettiği değil, öğrencinin teknolojiye yön verdiği bir durumun gerekli olduğu argümanını ortaya koymaktır.

Eğitsel Yapay Zeka Algoritmalarının Karar Alma Süreçleri

Eğitsel zeka sistemlerinin temelinde, her öğrenciyi temsil eden ve tüm dijital etkileşimlerini işleyen bir "Öğrenci Modeli" bulunur.

Sistem, yalnızca doğru/yanlış cevaplar hakkında değil, aynı zamanda "tıklama akışı" (arayüzdeki tıklama davranışı), görevde geçirilen süre ve öğrenme materyalleriyle etkileşim sıklığı hakkında da veri toplar.

Bu veriler daha sonra, (UNESCO, 2021)'de açıklanan gibi, bir öğrencinin bir kavram hakkında ne kadar bilgi sahibi olduğuna dair olasılıksal bir tahmin üreten Bilgi İzleme algoritmaları tarafından işlenir.

Teknik olarak, bu dizinin tamamı, arka planda öğrencinin bilgi durumunu sürekli olarak güncelleyen matematiksel bir model tarafından desteklenmektedir.

Böylece sistem, bir öğrencinin bir sonraki içeriğe hazır olup olmadığını tahmin etmek için dinamik bir veri profili oluşturur.

Toplanan veriler, sistemin karar verme birimine aktarılır.

Bu kurallar, öğrenci modelinden elde edilen verileri pedagojik kurallarla veya bazı durumlarda makine öğrenimi modelleriyle karşılaştırır.

Öğrencinin hata yapma eğilimi belirli bir eşiği aşarsa, algoritma aynı zamanda "başarısızlık olasılığı" kararı verir ve gelişmiş içeriğe erişimi otomatik olarak kısıtlar.

Bu aşamada sistem, öğrenciye sunulacak tüm içeriği bir olasılık hesaplamasına dayanarak seçer ve öğrencinin "kişiselleştirilmiş" bir öğrenme yolunu izlemesini sağlar.

Teknik olarak bu bir "öneri sistemi" olarak çalışır, ancak bir öğretim bağlamında bu, kişinin akademik ilerlemesini doğrudan engelleyen veya yönlendiren bir müdahaleye dönüştürülür ki bu, genel bir öneri sistemi için oldukça sıra dışıdır.

Arayüz, kararın hangi parametrelere dayandığını ve bu parametrelere hangi ağırlıkların verileceğini açıklayana kadar, algoritma öğrenci üzerinde mutlak ve denetlenemeyen bir egemenliğe sahiptir.

Algoritmik Paternalizm

Algoritmik paternalizm, söz konusu sistemin kullanıcının tercihlerini hiçe sayarak ne onun için "en doğru" karar olduğunu hesaplaması, ve bu kararı bir otorite olarak kullanıcıya dayatmasıdır (Yeung, 2017).

Eğitim yazılımları bağlamında bu durum, yapay zekanın öğrenciyi bir veri profili olarak değerlendirmesi ve onun "iyiliği için" öğrenme yolunu tek taraflı olarak dikte etmesi şeklinde tezahür eder.

Öğrencinin geçmişte düşük performans sergilemiş olması dikkate alınarak, sistem "bu öğrenci ileri seviye konuları öğrenemez" tahminini yapabilir ve sıkıntılı materyali onunla paylaşmama kararı alabilir .

Buradaki etik mesele, yapay zekanın pedagojik bir destek aracı olma özelliğini kaybederek öğrenciye karşı denetlenemeyen vasis ya da koruyucu bir rol üstlenmesi olarak görülür (UNESCO, 2021).

Bu koruyucu tavır, öğrencinin başarısız olmasını önlemeye çalışırken kendi sınırlarını deneyimlemesinin de önüne geçmektedir.

Dolayısıyla paternalizm, algoritmanın öğrencinin akademik kaderini gizli bir şekilde belirlemesi ve kararını şeffaf olmayan bir yöntemle uygulamasıdır. .

Öğrenci Otonomisi

Öğrenci otonomisi öğrenenin öz öğrenme süreci üzerinde kontrolü olması, kendi öğrenme hedeflerini belirleyebilmesi ve hata yapma özgürlük alanına sahip olmasıdır.

Yapay zeka sistemleri, "en verimli yol" önerileri ile öğrenciyi gizlice belirli tercihlere doğru yönlendirerek bu bağımsızlığı sarsabilir (Yeung, 2017).Yapay zeka sistemleri, kişiselleştirme vaadiyle öğrenciyi sürekli belirli bir yöne doğru "dürtmekte" (nudging), ve bu durumda bireyin bağımsız karar alma mekanizmaları zayıflamaktadır (Selwyn, 2019) Gerçek bir özerklik seçeneklerden birini seçmek değildir, sunduğu opsiyonun arkasındaki mekanizmayı anlamak ve bu yönlendirmeyi kabul etmeme hakkına sahip olmaktır (UNESCO, 2021).

Öğrenenlerin kendi öğrenme stratejilerini geliştirebilecekleri alan daralmakta ve öğrenenlerin bu algoritmanın öngörülerine bağımlı hale geldiği görülmektedir.

Eğitimde özerklik teknolojinin birey üzerine dayatılması değil, teknoloji ile kendi gelişim hedefleri doğrultusunda bir aracı olarak ilişki kurabilmesidir.

Bu bağlamda otonomi, teknolojinin öğrenciye ne yapacağını söylemesi değil, öğrencinin teknoloji üzerinde kendi hedefleri doğrultusunda kontrolü sağlayacak bilinçte ve güçte olmasıdır.

Algoritmaların Görünmez İşleyişi (Kara Kutu)

Yapay zeka sistemleri, karmaşık veri setlerini işleyerek bir sonuç üretir ancak bu sonuca hangi kriterlerle ulaştığını kullanıcıya açıklayamaz.

Sistemin girdi ve çıktıları görünürken, aradaki karar verme sürecinin gizli kalması "kara kutu" (black box) problemi olarak tanımlanır (Adadi ve Berrada, 2018).

Eğitimde bu kapalılık, öğrencinin neden belirli bir konuya yönlendirildiğini anlayamamasına ve dolayısıyla sistemi sorgulayamamasına neden olur.

Şeffaflık ve Güven Dengesi

Adil ve anlamlı kararların verilmesi, güvene dayalı bir öğrenme ortamının oluşması için gereklidir. Ancak opak bir algoritma yaptığı hataları ya da sahip olduğu önyargıları saklayabilir. Eğer öğrenci sistemin mantığını kavrayamıyor ise, ona tam olarak güvenmesi de beklenemez. Bu noktada şeffaflık teknik bir özellik değil öğrencinin sistemle sağlıklı iletişim kurabilmesi için temel bir gerekliliktir. Açık ve şeffaf eğitim teknolojileri, kullanıcı dostu bir dilde algoritmanın "neden" bu kararı verdiğini kullanıcıya açıklamasıyla sağlanabilir.

Mevcut Dijital Öğrenme Ortamlarında Eğitsel Yapay Zeka Kullanımı

Günümüzde popüler olan birçok öğrenme platformu, öğrenciyi yönlendirmek için gelişmiş algoritmalar kullanmaktadır.

Ancak bu sistemler incelendiğinde, şeffaflık ve öğrenci otonomisi açısından ciddi eksiklikler barındırdıkları görülmektedir.

Duolingo ve Birdbrain Algoritması

Duolingo, bir dizimleme kuramı modeli olan "Birdbrain" adlı yapay zeka modeli üzerinden çalışır, ve dünyanın en yaygın dil öğrenme platformlarının biridir.

Bu model, her bir öğrenci için çevrimiçi olarak bir soruyu doğru cevaplama olasılığını hesaplayarak ona en uygun soruyu (genellikle %70 %80 arasında başarı olasılığı olan) sunmaktadır.

Bu, bilişim etiğinde bir kara kutu örneğidir.. Öğrenci, kendisine gösterilen soruların neden o kadar zorlu olduğunu ya da sistemin kendi performansı hakkında hangi görüşü paylaştığını bilmez kararlar tamamen kendi kendine kapalı devre bir sistemde oluşturur, öğrenci ise pasif bir halde sistemin kendisine gösterdiği yolu takip eder.

ALEKS ve Algoritmik Paternalizm

Diğer önemli örnek olan ALEKS sistemi,Bilgi Uzayı Teorisi"ni uygulayarak öğrenmeye hazır olduğu şeyi tanımlayan katı bir yönlendirme mekanizmasına sahiptir.

Algoritma bu sistemde, belli bir konuyahenüz hazır olmadığına karar verirse,öğrencinin o konuya erişimini tamamen engeller.

Bu, daha önce bu bölümde betimlenen algoritmik paternalizmin somut bir örneğidir.

Sistem, öğrencilerin hata yapma ya da kendi sınırlarını zorlama hakkını ellerinden alır, çünkü kendisi öğrencinin adına en iyisini biliyormuş gibi davranır.

Duolingo benzeri kapalı yönlendirmeler ve ALEKS gibi kısıtlayıcı otoriteler öğrenciyi teknolojinin pasif kullanıcısı olma hâline koymaktadır.

Önerilen "Etik Bil

FAQs

What is GPTZero?

GPTZero is the leading AI detector for checking whether a document was written by a large language model such as ChatGPT. GPTZero detects AI on sentence, paragraph, and document level. Our model was trained on a large, diverse corpus of human-written and AI-generated text with support for English, Spanish, French, German, and other languages. To date, GPTZero has served over 17 million users around the world, and works with over 100 organizations in education, hiring, publishing, legal, and more.

When should I use GPTZero?

Our users have seen the use of AI-generated text proliferate into education, certification, hiring and recruitment, social writing platforms, disinformation, and beyond. We've created GPTZero as a tool to highlight the possible use of AI in writing text. In particular, we focus on classifying AI use in prose. Overall, our classifier is intended to be used to flag situations in which a conversation can be started (for example, between educators and students) to drive further inquiry and spread awareness of the risks of using AI in written work.

Does GPTZero only detect ChatGPT outputs?

No, GPTZero works robustly across a range of AI language models, including but not limited to ChatGPT, GPT-5, GPT-4, GPT-3, Gemini, Claude, and AI services based on those models.

What are the limitations of the classifier?

The nature of AI-generated content is changing constantly. As such, these results should not be used to punish students. We recommend educators to use our behind-the-scenes [Writing Reports](#) as part of a holistic assessment of student work. There always exist edge cases with both instances where AI is classified as human, and human is classified as AI. Instead, we recommend educators take approaches that give students the opportunity to demonstrate their understanding in a controlled environment and craft assignments that cannot be solved with AI. Our classifier is not trained to identify AI-generated text after it has been heavily modified after generation (although we estimate this is a minority of the uses for AI-generation at the moment). Currently, our classifier can sometimes flag other machine-generated or highly procedural text as AI-generated, and as such, should be used on more descriptive portions of text.

I'm an educator who has found AI-generated text by my students. What do I do?

Firstly, at GPTZero, we don't believe that any AI detector is perfect. There always exist edge cases with both instances where AI is classified as human, and human is classified as AI. Nonetheless, we recommend that educators can do the following when they get a positive detection: Ask students to demonstrate their understanding in a controlled environment, whether that is through an in-person assessment, or through an editor that can track their edit history (for instance, using our [Writing Reports](#) through Google Docs). Check out our list of [several recommendations](#) on types of assignments that are difficult to solve with AI.

Ask the student if they can produce artifacts of their writing process, whether it is drafts, revision histories, or brainstorming notes. For example, if the editor they used to write the text has an edit history (such as Google Docs), and it was typed out with several edits over a reasonable period of time, it is likely the student work is authentic. You can use GPTZero's Writing Reports to replay the student's writing process, and view signals that indicate the authenticity of the work.

See if there is a history of AI-generated text in the student's work. We recommend looking for a long-term pattern of AI use, as opposed to a single instance, in order to determine whether the student is using AI.