

Hakikat Sonrası Uçurumda Yön Bulmak: Sentetik Medya ve Yapay Zekanın İkili Kullanım İkilemi

Navigating the Post-Truth Abyss: The Dual-Use Dilemma of Synthetic Media and AI

Gökhan VARAN 

Görüş Review

ÖZ

Bu çalışma, modern bilgi ekosisteminde üretken yapay zekanın ikili kullanımının doğasını örneklerle eleştirel bir şekilde ele almayı amaçlamaktadır. Yapay zeka, iletişim ve yaratıcılık manasında güçlü araçlar sunduğu gibi uydurma metinlerden derin kurgu (deepfake) içeriklere kadar ultra gerçekçi sentetik medya üretimini kolaylaştırması sebebiyle ciddi bir zorluk teşkil etmektedir. Birincil tehdidin basit aldatmacanın ötesinde, kanıta dayalı hakikatin temelini baltalayan “Yalancının Kâr Payı” olarak adlandırılan daha aşındırıcı, ikincil bir etkiye uzandığı görüşü savunulmaktadır. Tespit ve köken sistemleri gibi teknolojik karşı önlemler geliştirildiği gibi bunların üretken teknolojilerle asimetrik bir silahlanma yarışı içinde olduğu gözlemlenmektedir. Bu yazı, tamamen teknik bir çözümün yetersiz olduğunu iddia etmekte ve bunun yerine, ortaya çıkan hakikat sonrası manzarada yön bulmak için sınırlı teknik müdahaleleri sağlam politikalar, uluslararası düzenleyici çerçeveler ve eleştirel medya okuryazarlığına yönelik yenilenmiş bir toplumsal taahhülle birleştiren bütüncül bir yaklaşımın gerekliliği öne sürmektedir.

Anahtar Kelimeler: [Yapay Zeka](#), [İletişim](#), [Sentetik Medya](#), [Dezenformasyon](#).

ABSTRACT

This paper aims to critically examine the nature of the dual use of generative AI in the modern information ecosystem. While AI offers powerful tools for communication and creativity, it also poses a serious challenge as it facilitates the production of ultra-realistic synthetic media, from fabricated texts to deepfake content. It is argued that the primary threat extends beyond simple deception to a more corrosive secondary effect, the so-called “Liar’s Dividend”, which undermines the foundation of evidence-based truth. Technological countermeasures, such as detection and provenance systems, have been developed, but they appear to be in an asymmetrical arms race with productive technologies. This paper argues that a purely technical solution is inadequate and argues instead for a holistic approach that combines limited technical interventions with robust policies, international regulatory frameworks and a renewed societal commitment to critical media literacy to navigate the emerging post-truth landscape.

Keywords: [Artificial Intelligence](#), [Communication](#), [Synthetic Media](#), [Disinformation](#).

Giriş

Dijital dönüşüm, bilginin paylaşım ve yayılma hızını artırdığı gibi yapay zeka (YZ) eliyle sentetik medya içeriklerinin üretimini artırmış ve maliyetlerini de düşürmüştür. YZ destekli olarak üretilen içeriklerin belirli bir kaliteyi de barındırması söz konusu olduğu için inandırıcılığı da aynı şekilde yüksek seyretmektedir. Bu sebeple eskilerin sık kullandığı öğütlerden olan “gördüğüne inanma” meselesi de ciddi şekilde tartışmaya açık hale gelmektedir. Zira gördüğümüz, işittiğimiz bile artık gerçek olmayabilir. Teknoloji marifetiyle üretilen bu sahte içeriklerle mücadelenin yolu ise yine teknolojiden geçmekte ve doğru/sahte tespitini yapabilmek için yine yapay zeka kullanılmaktadır. Üretken yapay zekanın yaygınlaşması, iletişim, medya ve demokratik toplumların sağlığı için kritik bir dönüm noktasına işaret etmektedir. Bu güçlü teknolojiler temelden ikili kullanımlıdır: yaratıcılık ve verimlilik için kullanılabilirler fakat aynı zamanda kamusal alanı hiper-gerçekçi “sentetik gerçekliklerle” kirliletmek için de silah haline getirilebilmeleri gayet mümkündür. Ortaya çıkan bu yıkım sadece teknolojinin yarattığı sahteler değil, aynı zamanda “Yalancının Kar Payı” olarak bilinen bir kavramla otantik bilgilere gölge düşürerek şüphe oluşturma kapasitesiyle de körüklenebilmektedir. Yapay zeka ile üretilen dezenformasyonun yarattığı zorluğun ilk görüldüğünden daha derin olduğu savunulmaktadır. Nihai risk ise, yalnızca bir yalana inanmak üzere kandırılmamız değil, aynı zamanda neyin doğru neyin sahte olduğu hususunda kolektif yeteneğimizi yitirmemizdir.

Bu sebeple sentetik medya içeriklerine ve dezenformasyona maruz kalacağımız ve dozunun teknoloji ilerledikçe artacağı kabulüyle birlikte toplumun dijital okuryazarlık sahibi olması ve yapay zeka araçlarını sahte içeriklerin tespiti ve damgalanma noktasındaki kullanımı kritik öneme sahiptir. Bu görüş yazısı genel hatlarıyla teknoloji vasıtasıyla iletişim alanında yaşanan değişim ve dönüşümü riskler ve çözüm önerileri bağlamında iki yönlü ele alarak bir değerlendirme yapmayı ve politika önerileri sunmayı amaçlamaktadır.

İnsan ve Yapay Zeka Ekseninde İletişimin Geleceği

Yapay zekanın bir otomasyon aracı olmanın ötesinde iletişim alanında geniş bir etkiye sahip destek unsuru iş ortağı olarak konumlandığı gözlemlenmektedir.

Daha geniş bir tanımlama yapılacak olursa yapay zekanın insanın yaptığı görevleri devralması ve insanların da yeteneklerini artırmaya destek sağlayarak yeni bir etkileşim ve yaratıcılık modeli ortaya koymak olarak ifade etmek mümkündür. Bu alanda yapılan araştırmalar, karmaşık bilgileri işlemede yardım almadan müzik ve sanatta kreatif çıktılar oluşturmaya kadar üretken yapay zeka kullanımını ölçmekte ve yapay zekayı bir ‘üretim uzmanı’ gibi konumlamaktadır.

Öte yandan akademisyenlerin bir kısmı yoğun yapay zeka destekli iletişim kurmanın bağımlılık oluşturacağı gibi temel insani becerilerin, muhakeme kabiliyetinin gerilemesine yol açabileceğini düşünmektedir. (Frontiers in Human Dynamics, 2024) Böyle bir durumun oluşmaması adına yapay zekanın hesaplama gücünden faydalanarak iş birliği kurmak, insan sezgi ve öngörüsünün gücünü koruyarak nasıl bir konsept oluşturulacağına odaklanmak gerekmektedir.

Bu bağlamda “İnsanın bilgi birikimi ve muhakemesi varlığını korurken yapay zekanın analitik gücü nasıl konumlandırılmalı ve insan-YZ iş birliği modelleri nasıl oluşturulmalıdır?”, “Yapay

zeka aracılı iletişimle uzun süreli etkileşim neticesinde insan becerileri ve muhakemesi gelişim noktasında sekteye uğrayabilir mi?” gibi sorulara cevap aranmalıdır.

İkili Kullanımlı Bir Teknoloji Olarak Yapay Zeka

Yapay zeka eliyle oluşturulan dezenformasyonun yayılması günümüzün bilgi iletişimi manasında ekosisteme zarar verme ve bütünlüğüne tehdit oluşturma açısından en önde gelen zorluklardan biri olarak değerlendirilmektedir. Dünya Ekonomik Forumu’nun (WEF) 2025 Küresel Riskler Raporu’nda

yanlış bilgilendirme ve dezenformasyon, yapay zeka tarafından üretilen içeriklerin sayısındaki ciddi artışla daha da kötüleşerek küresel riskler arasında en ön sıralarda yer almaktadır. (Mishra, 2025) YZ'nin bu manada ikili kullanımının derinlemesine değerlendirilmesi gerekmektedir. YZ'nin sentetik gerçekliğin üretilmesinde oynadığı rolü de gerçeği ortaya çıkarmadaki etkisi de dikkatle incelenmelidir. Teknik temelleri, toplumsal etkileri, teknolojik alınacak karşı önlemleri ve yapay zekanın kendi eliyle büyüttüğü sorunlarla mücadelede kullanımında oluşan kısıtlar iyi analiz edilmeli ve teknolojinin oynadığı bu iki yönlü rolün altı dikkatle çizilmelidir.

Geleneksel manada sahte içerikler, kasıtlı olarak yanıltma tarih boyunca hep vardı ve var olmaya devam edecek fakat insanları aldatmaya yönelik içeriklerin üretimi, teknoloji eliyle gerçeğinden ayırmanın çok zor olduğu sahte içerik bombardımanlarının çok tehlikeli bir boyuta eriştiği gerçeği gözden kaçırılmamalıdır. Bu noktada kamusal alanımızın ve iletişim platformlarımızın barındırdığı zafiyetler üretken yapay zekanın katlanarak ilerlemesinin neticesinde ortaya çıkmaktadır. İnternetin ve sosyal medyanın yükselişi, yanlış anlatıların ve sahte içeriklerin hayal bile edilemeyecek hız ve büyüklükte yayılmasına zemin oluşturmuştur. (Tworek ve Scheirer, 2024) Üretken yapay zekanın bu önceden beri var olan, korunaksız bilgi ekosistemiyle birleşmesi, yeni ve çok daha tehlikeli bir noktaya işaret ederek ikna edici sahte içeriklerin benzeri görülmemiş çapta yayılmasını kolaylaştırmaktadır.

Bir Silah Olarak Üretken Yapay Zeka

Dezenformasyon çağının göbeğinde iki ana üretken yapay zeka teknolojisini bulunmakta ve bu teknolojiler eliyle kusursuz sahtecilikler, erişimin kolaylığı ve tüketicinin manipüle edilmesinin kolaylaşması kötü niyetli aktörlerin iştahını kabartmaktadır. Bu iki temel teknoloji; Büyük Dil Modelleri (LLM) ve Üretken Çekişmeli Ağlar (GAN)'dan oluşmaktadır. OpenAI'nin GPT serisi gibi büyük dil modelleri (LLM), devasa miktarda verinin üzerinde eğitilmiş sinir ağılardır ve genel olarak ana yetenekleri, bir dizideki en olası bir sonraki

kelimeyi istatistiksel olarak tahmin ederek tutarlı, insan benzeri metinler üretmektir. (Milvus, 2025)

Dezenformasyon bağlamında, büyük dil modellerinin (LLM), makul ama yanlış haber makaleleri, uydurma alıntılar, kişiselleştirilmiş ortalama e-postaları ve bir dava için taban desteği veya muhalefeti simüle etmek üzere tasarlanmış geniş sahte sosyal medya gönderileri ağırları oluşturmak için kullanılması mümkündür. Çalışmalar, mevcut LLM'lerin seçim dezenformasyon operasyonları için kolayca yüksek kaliteli içerik üretebildiğini ve genellikle insan tarafından yazılan içerikten ayırt edilemeyen metinler ürettiğini göstermiştir. (Williams, Burke-Moore, Chan, Enock, Nanni, Sippy, Chung, Gabasova, Hackenburg & Bright (2025). Ayrıca, araştırmalar

modellerin dezenformasyon üretme eğiliminin “komut (prompt) mühendisliği” yoluyla manipüle edilebileceğini göstermektedir.

Diğerte knoloji olan “Üretken Çekişmeli Ağlar (GAN)” için görsel ve işitsel dünyanın usta sahtekarları ve aldatıcısı demek yanlış bir tabir olmayacaktır. (Varughese, 2025) İki sinir ağını - sahteleri üreten bir “üretici (generator)” ve onları tespit etmeye uğraşan bir “ayırt edici (discriminator)”- birbirine karşı çarpıştıran GAN'ler, genel olarak gerçeklikten ayırt edilemeyecek düzeyde olan derin kurgu (deepfake) videolar ve ses klonları üretmeyi öğrenmektedirler. Bu teknolojilerin gelişiminin tarafsız bir ilerleme olarak görülmemesi gerektiği gibi bir silah olarak nasıl kullanıldığının örneği olarak bir şikayeti ve bir dizüstü bilgisayarı olan herkese “kanıt” uydurma, itibarları baltalama ve siyasi söylemi bozma gücü verdiği gösterilebilir.

Bir silah olarak tanımlanabilecek bu teknolojilerin cephaneliğinin oluşturduğu en kritik tehdit ise aldatmacanın bizzat kendisi değil birilerinin bu sahte içeriklere inanmasıdır. Hukuk profesörleri Bobby Chesney ve Danielle Citron bu durumu “Yalancının Kar Payı” olarak adlandırmışlardır. Bu konudaki öngörüler ise derin kurgu (deepfake) hakkında kamuoyunun farkındalığı arttıkça

paradoksal bir şekildeyalancıya faydası sağlamasıdır. Örnek verilecek olursa kötü bir davranış sergilerken kaydedilen bir politikacı veya bürokrat, gerçek olan videoyu ‘sofistike bir sahte içerik’ olduğu iddiasıyla inkar edilemez olanı inkar edip hakikati yalan gösterebilir. Bu strateji, dezenformasyonun hedefini temelden değiştirir. Amaç artık halkı belirli bir yanıla ikna etmek değil, tüm kanıtların şüpheli hale geldiği ve gerçeklerin hesap verebilirliği sağlama gücünü etkisiz hale getirdiği yaygın ve zayıflatıcı bir sinizm yaratmaktır. Bu, belirli bir hakikate değil, işleyen bir demokrasinin epistemolojik temellerine bir saldırıdır. Ses, video ve belgesel kanıtların hepsi potansiyel olarak sentetik olarak reddedilebilirse, hangi temelde mantıklı bir tartışma yapılabilir, liderleri sorumlu tutabilir veya adaleti sağlayabiliriz? Ampirik çalışmalar bu etkinin gerçek olduğunu zaten göstermiştir; bir skandala yanıt olarak “sahte haber” diye bağırarak politikacılar, takipçileri arasında desteği başarıyla pekiştirebilir ve halkın belirsizliğinden etkili bir şekilde kâr sağlayabilirler.

Güvenin Sarsılması ve Asimetrik Bir Silahlanma Yarışı

Yapay zeka tarafından üretilen sentetik medya içeriklerinin yayılması, bireylere, kurumlara ve toplumsal güvenin dokusuna doğrudan bir tehdit oluşturmaktadır. Etkisi, bireysel yanlışların yayılmasıyla sınırlı kalmayıp, paylaşılan gerçeklik duygumuzun aşındırıcı bir şekilde bozulmasına kadar uzanmaktadır. Bu tehdide yanıt olarak, YZ destekli tespit, teyit mekanizmaları, içerik filigranlama ve İçerik Kökeni ve Özgünlüğü Koalisyonu (C2PA)’nın İçerik Kimlik Bilgileri standardı gibi köken sistemlerine odaklanan teknolojik bir karşı taarruz seferber edilmiştir. Bu araçlar herhangi bir stratejinin hayati bileşenlerini oluşturmaktadır fakat bunların tam bir çözüm olduğuna inanmanın tehlikeli bir yanılgı olacağını belirtmek gerekir. İçerik üretim ve tespiti arasındaki dinamik, sürekli ve asimetrik bir “silahlanma yarışı” olarak tanımlanabilir.

Bu ikisi arasındaki asimetri ise çarpıcıdır; sahte içerik oluşturma araçları katlanarak daha ucuz,

daha hızlı ve daha erişilebilir hale gelirken, tespit araçları karmaşık, maliyetli ve sürekli olarak bir adım geride kalmaktadır. Bir YZ kusurunu tanımlamak üzere ortaya çıkan her yeni tespit yöntemine karşılık, yeni nesil modeller tam da bu kusuru ortadan kaldırmak için eğitilmektedir. Kaos yaratmak, onu denetlemekten çok daha kolay ve ucuzdur. Bir diğer ifadeyle tehdidi yok

etmek üzere geliştirilen zahmetli teknoloji hızlıca ve ucuz maliyetlerle kendini yok etme süreçlerini de otomatikman beslemektedir.

Öte yandan, bu teknolojik çözümlerin kendileri de etik tehlikeleri bünyesinde barındırmaktadır. YZ tespit modelleri önyargılar sergileyebilir, anadili İngilizce olmayanlar veya nörodiverjan bireyler için daha yüksek yanlış pozitif oranları üreterek haksız cezalara yol açabilir ve otomatik denetime aşırı güvenmek, hiciv ve muhalefet de dahil olmak üzere meşru ifade üzerinde caydırıcı bir etki yaratma riskini de taşır. YZ’nin kötüye kullanımına dayanan bir sorunu, daha fazla YZ’ye körü körüne bir inançla çözülebileceği düşüncesi hayalden ibarettir.

Bunlara ek olarak teknolojinin getirdiği risk ve çözümlere geleneksel bir destek olarak toplumun ve kullanıcıların sentetik medya duyarlılığı kazanması adına ‘Yapay Zeka (YZ) Destekli Medya Okuryazarlığı’ faaliyetleri önem arz etmektedir. Kullanıcıların sahte içeriklerin tespitine dek maruz kaldığı süreçte dair bilinçlenmesi, YZ araçlarının etik kullanımının ve mahremiyet odaklı çekincelere önem verildiğinin gözlemlenebilmesi alınabilir basit ve etkili önlemler arasında sayılabilir. Bir diğer örnek olarak TikTok’un 2024 yılında otomatik YZ içerik etiketleme politikasını devreye alması içeriğin kökeninin tespiti bağlamında bir standart getirme çabası sektöre ciddi bir katkı sağlamıştır.

Öneriler ve Sonuç

Yapay zeka temelli dezenformasyon sorunu tümüyle teknik bir sorundan ibaret değildir. Konunun özünde insan psikolojisinin zayıflıklarını, kutuplara ayrılmış bir toplumun çatlaklarını ve dijital mecralardaki teşvik edici yapıların istismar

edilmesiyle ortaya çıkmış sosyo-teknik bir mesele olarak değerlendirilmelidir. Bu sebeple bu konuya verilecek teknik cevapların kısmi ve geçici olacağı gibi meselenin bütününe bir cevap üretemeyeceğini söylemek gerekir. Daha kapsayıcı ve bütünü analiz eden bir strateji izlenmelidir.

Bu itibarla ilk başta konuya ilişkin politika oluşturmaya ve düzenlemelere odaklanmak gerekmektedir. Bu teknolojilerin ve alandaki gelişmelerin hergeçengünde değiştiğini ve geliştiğini hesaba katarsak burada yapılması gereken uzun süreli, yoğun bürokratik yasa çalışmaları ve anayasal düzenlemeler değildir. Daha akıllı, çevik düzenlemelere ve konuya ilişkin bir çerçeve çizen, uzman görüşlerini içinde barındıran esnek ve adapte edilebilir bir yaklaşım ortaya koyarak strateji geliştirmek kritik öneme sahiptir. Geliştirilecek strateji, YZ geliştiricileri ve platformlardan temel olarak şeffaflık, dezenformasyonun sebep olduğu zararlara ilişkin hesap verebilirlik ve içeriklere ilişkin bir standart oluşturmaya teşvik edecek şekilde yapılandırılmalıdır. Teknoloji ve hukuk sinerjisi oluşturularak zorunlu içerik menşei etiketlemesi (watermark, meta-data) teşvik edilmelidir. Dezenformasyon ile eleştiri farkını hukuki olarak da netleştirerek bir kılavuz oluşturarak sürecin şeffaflaştırılması sağlanmalıdır. Bunlara ek olarak Türkiye, bu alandaki inisiyatifleri ve küresel platformları da standartları oluşturmak ve hızlıca güncellemeler yapabilmek için yakından takip etmelidir.

İkinci olarak önem arz eden husus, yeni bir medya okuryazarlığı bakışı oluşturarak bu meseleye kamu ve özel sektör, akademi ve STK'lar olmak üzere topyekûn yatırım yapılmasıdır. Geleneksel manada yapılan doğruluk kontrol ve teyidi gerekli olmakla birlikte, kanıtın kendisinin şüpheli olduğu bir noktada yetersiz kalacaktır. Bu sebeple, eğitim ve farkındalık çalışmaları bu belirsizlik

ve şüphelerle nasıl başa çıkacaklarını öğrenmeleri bakımından hayati öneme sahiptir. Bu faaliyetleri lise ve üniversite müfredatlarına entegre ederek 'Yapay Zeka Destekli İçerik Analizi', 'Teyit Mekanizmaları ve Doğru Bilgiye Erişim' gibi dersleri

ve atölyeleri hayata geçirmek faydalı olacaktır. Bu çalışmalara eşlik edecek sivil toplum kuruluşları, akademi ve diğer paydaşlar için Ar-Ge teşvikleri sağlayarak bu programlar neticesinde ortaya çıkacak projelere de fon sağlanmalı ve proje, ürün geliştirme süreçleri teşvik edilmelidir.

Sonuç olarak, sonsuz, ucuz ve sentetik içeriklerin çağında, güvenilir insan kurum ve mekanizmalarının değerinin önemi hiç bu kadar yüksek olmamıştı. Geleneksel iletişim ve medya faaliyetlerinin etik kaygılar, hesap verebilirlik ve bağlam gibi kaygılar sebebiyle yavaş ve zahmetli ilerlediği noktada aldatıcı ve sahte içerik bombardımanına tutulduğumuz bir zeminde önemli bir çapa görevi icra etmektedir. Özgür, iyi kaynaklara ve teknolojik imkanlara sahip yeni medya yapılarını desteklemek hakikat sonrası uçuruma karşı kolektif savunmanın ana gövdesini oluşturmaktadır. Teknolojinin iletişim alanında getirdiği risk ve çözüm ikilemi bir meydan okumaya dönüşmüştür. Bu meydan okumaya karşı koymanın yolları aynı şekilde teknolojiye olduğu gibi sahteye karşı hakiki, yapay olana karşı doğal, fabrikasyon olana karşı insani bir zeminde çözüme kavuşabilir. Sorunun üreticisi de çözücüsü de yine insan ve onun iyiye veya kötüye dair aldığı aksiyonlardan ibarettir.

Kaynaklar

- Chesney, R., & Citron, D. (2019). Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security. *California Law Review*, 107(6), 1753-1804
- Schiff, K. J., Schiff, D. S., & Bueno, N. S. (2024). The Liar's Dividend: Can Politicians Claim Misinformation to Evade Accountability? *American Political Science Review*, 118(2), 920-936
- Groh, M., Epstein, Z., Firestone, C., & Picard, R. (2022). Deepfake detection by human crowds, machines, and machine-informed crowds. *Proceedings of the National Academy of Sciences*, 119(1), e2110013119

Sadasivan, V. S., Kumar, A., Balasubramanian, S., Zheng, W., & Feizi, S. (2023). *Can AIGenerated Text be Reliably Detected?* arXiv preprint arXiv:2303.11156

Coalition for Content Provenance and Authenticity (C2PA). (Çeşitli Tarihler). *Technical Specification for Content Provenance and Authenticity*

Stanford University Institute for Human-Centered Artificial Intelligence (HAI). (2025). *AI Index Report*

World Economic Forum. (2025). *Global Risks Report 2025*

İleri Okumalar

1. Ethical Considerations - Artificial Intelligence (AI) - Research Guides, erişim tarihi Haziran 15, 2025, <https://libguides.lib.umt.edu/ai/ethics>

2. The impact of AI as a mediator on effective communication: enhancing interaction in the digital age - Frontiers, erişim tarihi Haziran 15, 2025, <https://www.frontiersin.org/journals/humandynamics/articles/10.3389/fhumd.2024.1467384/full>

3. How Generative AI is transforming strategic communication research and empowering the

next generation of communication professionals | UNC Hussman School of Journalism and Media, erişim tarihi Haziran 15, 2025, <https://hussman.unc.edu/news/how-generative-ai-is-transforming-strategic-communication-research-and-empowering-the-next-generation-of-communication-professionals>

4. High Impact Journals - Communication Studies - LibGuides at University of Texas at Austin, erişim tarihi Haziran 15, 2025, <https://guides.lib.utexas.edu/c.php?g=513011&p=10428282>

5. Misinformation and Disinformation in the Digital Age: A Rising Risk for Business and Investors, erişim tarihi Haziran 15, 2025, <https://corpgov.law.harvard.edu/2025/05/12/misinformation-and-disinformation-in-the-digital-age-a-rising-risk-for-business-and-investors/>

6. Disinformation in the Digital Age: Practical Tips for Young People - The Commons, erişim tarihi Haziran 15, 2025, <https://commonslibrary.org/disinformation-in-the-digital-age/>

7. How can LLMs contribute to misinformation? - Milvus, erişim tarihi Haziran 15, 2025, <https://milvus.io/ai-quick-reference/how-can-llms-contribute-to-misinformation>

8. Emotional prompting amplifies disinformation generation in AI large language models, erişim tarihi Haziran 15, 2025, <https://www.frontiersin.org/journals/artificialintelligence/articles/10.3389/frai.2025.1543603/full>

9. Deepfake (Generative adversarial network) - CVisionLab, erişim tarihi Haziran 15, 2025, <https://www.cvisionlab.com/cases/deepfake-gan/>

10. What are Generative Adversarial Networks (GANs)? - IBM, erişim tarihi Haziran 15, 2025, <https://www.ibm.com/think/topics/generative-adversarial-networks>

11. Generative Adversarial Network (GAN): Powering Deepfakes & AI's Role in Detection, erişim tarihi Haziran 15, 2025, <https://facia.ai/knowledgebase/generative-adversarial-network-gan-powering-deepfakes-ai-role-in-detection/>

12. AI is Changing the Reality—Revealing the History of Deepfakes - Facia.ai, erişim tarihi Haziran 15, 2025, <https://facia.ai/blog/ai-is-changing-the-reality-revealing-the-history-of-deepfakes/>

13. Deepfakes, Elections, and Shrinking the Liar's Dividend | Brennan Center for Justice,

- erişim tarihi Haziran 15, 2025, <https://www.brennancenter.org/our-work/researchreports/deepfakes-elections-and-shrinking-liars-dividend>
14. Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security, erişim tarihi Haziran 15, 2025, <https://www.californialawreview.org/print/deep-fakes-a-loomingchallenge-for-privacy-democracy-and-national-security>
 15. The Liar's Dividend: Can Politicians Claim Misinformation to Evade Accountability? American Political Science Review - Cambridge University Press, erişim tarihi Haziran 15, 2025, <https://www.cambridge.org/core/journals/american-political-sciencereview/article/liars-dividend-can-politicians-claim-misinformation-to-evadeaccountability/687FE54DBD7ED0C96D72B26606AA073>
 16. The use of artificial intelligence in counter-disinformation ... - Frontiers, erişim tarihi Haziran 15, 2025, <https://www.frontiersin.org/journals/politicalscience/articles/10.3389/fpos.2025.1517726/full>
 17. How AI-Powered Fact-Checking Can Help Combat Misinformation | Ivy Exec, erişim tarihi Haziran 15, 2025, <https://ivyexec.com/career-advice/2025/how-ai-powered-factchecking-can-help-combat-misinformation/>
 18. Fact-checking AI with Lateral Reading - Using AI tools in Research - Research Guides at Texas A&M University-Corpus Christi, erişim tarihi Haziran 15, 2025, <https://guides.library.tamucc.edu/AI/lateralreadingAI>
 19. AI Deepfakes: Disinformation, Manipulation and Detection - The Red Line Project, erişim tarihi Haziran 15, 2025, <https://redlineproject.news/2025/03/20/opinion-ai-deepfakes-andjournalism-misinformation-manipulation-legality-and-detection/>
 20. AI detection tool helps journalists identify and combat deepfakes, erişim tarihi Haziran 15, 2025, <https://ijnnet.org/en/story/ai-detection-tool-helps-journalists-identify-and-combatdeepfakes>
 21. Detecting AI fingerprints: A guide to watermarking and beyond - Brookings Institution, erişim tarihi Haziran 15, 2025, <https://www.brookings.edu/articles/detecting-aifingerprints-a-guide-to-watermarking-and-beyond/>
 22. AI fact-checking tools | Journalist's Toolbox, erişim tarihi Haziran 15, 2025, <https://journaliststoolbox.ai/ai-fact-checking-tools/>
 23. The AI Arms Race: Deepfake Generation vs. Detection - SecurityWeek, erişim tarihi Haziran 15, 2025, <https://www.securityweek.com/deepfakes-and-the-ai-battle-between-generation-and-detection/>
 24. The Arms Race Between AI Content Creators and Detectors, erişim tarihi Haziran 15, 2025, <https://checker.ai/blog/the-arms-race-between-ai-content-creators-and-detectors>
 25. Centre for Information Policy Leadership - Ten Recommendations for Global AI Regulation, erişim tarihi Haziran 15, 2025, https://www.informationpolicycentre.com/uploads/5/7/1/0/57104281/cipl_ten_recommendations_global_ai_regulation_oct2023.pdf
 26. Media Literacy Is Vital for Informed Decision-Making - Learning for Justice, erişim tarihi Haziran 15, 2025, <https://www.learningforjustice.org/media-literacy-is-vital-for-informeddecisionmaking>
 27. The role for higher education in combatting AI misinformation - HEPI, erişim tarihi Haziran 15,

2025, <https://www.hepi.ac.uk/2024/01/12/the-role-for-higher-education-incombatting-ai-misinformation/>

28. Deepfakes and Deception: A Framework for the Ethical and Legal Use of MachineManipulated Media - Modern War Institute, erişim tarihi Haziran 15, 2025, <https://mwi.westpoint.edu/deepfakes-and-deception-a-framework-for-the-ethical-andlegal-use-of-machine-manipulated-media/>

29. How can humans and AI work together to detect deepfakes? - MIT Media Lab, erişim tarihi Haziran 15, 2025, <https://www.media.mit.edu/articles/how-can-humans-and-aiwork-together-to-detect-deepfakes/>

30. Unmasking AI's Role in the Age of Disinformation: Friend or Foe? - MDPI, erişim tarihi Haziran 15, 2025, <https://www.mdpi.com/2673-5172/6/1/19>

31. The Economics of Disinformation: Academic Freedom in the Era of AI - AAUP, erişim tarihi Haziran 15, 2025, <https://www.aaup.org/JAF15/economics-disinformationacademic-freedom-era-a>

Yazar Bilgileri

Author details

Yapay Zeka Politikaları Derneği (AIPA) Başkan Yardımcısı,
gokhan.varan@aipaturkey.org.

Kaynak Göstermek İçin

To Cite This Article

Varan, G. (2025). Hakikat sonrası uçurumda yön bulmak: Sentetik medya ve yapay zekanın ikili kullanım ikilemi [Görüş Yazısı]. *Yeni Medya*, (18), 522-529.