

Deepfake Teknolojisi ve Siyasi Güven Krizi: Seçim Süreçlerinde Dezenformasyonun Sosyolojik Etkileri

Ali Kiriktaş¹

Özet

21. yüzyılın dijital iletişim ekosistemi, yapay zekâ tabanlı sentetik medya üretim teknolojilerinin, özellikle deepfake (derin kurgu) fenomeninin yükselişiyle epistemolojik bir kırılma noktasına ulaşmıştır. Bu çalışma, 2014 yılında Üretken Çekişmeli Ağlar (GANs) ile başlayan ve günümüzde Difüzyon Modelleri ile hiper-gerçekçi bir seviyeye ulaşan deepfake teknolojisinin, teknik altyapısından toplumsal sonuçlarına uzanan çok katmanlı etkilerini incelemektedir. Araştırmanın odak noktası, seçim süreçlerinde kullanılan görsel ve işitsel dezenformasyonun, demokratik meşruiyet, siyasal güven ve kolektif bellek üzerindeki aşındırıcı etkileridir. Jean Baudrillard'ın *simülakr teorisi* ve Jürgen Habermas'ın *kamusal alan* kavramları çerçevesinde temellendirilen sosyolojik analiz, gerçeğin ve kurgunun ayırt edilemez hale geldiği *hakikat sonrası* (post-truth) çağda, seçmen algısının nasıl manipüle edildiğini ortaya koymaktadır. Çalışma, Slovakya (2023), Hindistan (2024), ABD (2024) ve Türkiye (2023) seçimlerinden elde edilen ampirik veriler ve vaka analizleri ışığında, *yalancının kârı* (liar's dividend) kavramını tartışmakta; siyasetçilerin bu teknolojiyi sadece saldırı aracı olarak değil, aynı zamanda gerçek skandalları örtbas etmek için bir savunma kalkını olarak nasıl kullandıklarını irdelemektedir. Ayrıca, medya okuryazarlığı eğitimlerinin paradoksal bir şekilde *sinizm* ve kurumsal güvensizliği artırabileceğine dair bulgular (geri tepme etkisi) ele alınmaktadır. Sonuç olarak, Avrupa Birliği Yapay Zekâ Yasası ve Türk Ceza Kanunu kapsamındaki mevcut yasal düzenlemelerin yetersizlikleri tartışılarak, teknolojik tespit mekanizmaları (C2PA) ile sosyolojik direnç stratejilerinin entegre edildiği hibrit bir çözüm modeli önerilmektedir.

1 Doktora Öğrencisi, Kapadokya Üniversitesi Lisansüstü Eğitim- Öğretim ve Araştırma Enstitüsü Siyaset Bilimi ve Uluslararası İlişkiler, ali_kiriktas_44@hotmail.com, ORCID ID 0000-0003-4209-0094

1. Giriş

İnsanlık tarihi boyunca bilgi ve belge, gerçeğin kanıtı olarak kabul edilmiştir. Fotoğrafın icadından bu yana görsel materyaller, görmek inanmaktır anlayışıyla, yaşanan olayların tartışmasız şahitleri olarak konumlanmıştır. Ancak dijital çağın getirdiği en yıkıcı teknolojik gelişmelerden biri olan ve literatürde *deepfake* olarak adlandırılan sentetik medya teknolojisi, bu kadim epistemolojik kabulü kökünden sarsmaktadır. *Deep learning* (derin öğrenme) ve *fake* (sahte) kelimelerinin birleşiminden türetilen bu kavram, yapay zekâ algoritmaları kullanılarak, gerçekte var olmayan olayların, kişilerin veya konuşmaların, insan algısını yanıltacak düzeyde gerçekçi bir şekilde üretilmesini ifade etmektedir (Akool, t.y.).

Deepfake teknolojisinin tarihsel kökenleri, akademik olarak 2014 yılında Ian Goodfellow ve ekibinin *Üreten Çekişmeli Ağlar* (Generative Adversarial Networks- GANs) mimarisini geliştirmesiyle başlasa da, toplumsal bir fenomen haline gelmesi 2017 yılında *deepfakes* takma adlı bir Reddit kullanıcısının, ünlülerin yüzlerini pornografik videolara monte etmesiyle gerçekleşmiştir (Acim vd., 2025; Westerlund, 2019). Başlangıçta bireysel mahremiyeti tehdit eden ve rıza dışı pornografi (non-consensual pornography) üretiminde kullanılan bu teknoloji, kısa süre içinde evrilerek ulusal güvenlik, finansal piyasalar ve demokratik süreçler için varoluşsal bir tehdide dönüşmüştür. Akademik araştırmalar, 2023 yılı itibarıyla internette dolaşan deepfake videolarının sayısının 100.000'i aştığını ve bu sayının her altı ayda bir katlanarak arttığını göstermektedir (Akool, t.y.). Bu hızlı bir şekilde artmış, teknolojinin sadece teknik bir yenilik olmaktan çıkıp, sosyolojik ve politik sonuçları olan küresel bir krize dönüştüğünü kanıtlamaktadır.

Siyaset sosyolojisi perspektifinden bakıldığında, deepfake teknolojisi demokratik süreçlerin temel taşı olan *bilgilendirilmiş seçmen* (informed voter) idealini doğrudan tehdit etmektedir. Demokrasi, özünde vatandaşların doğru bilgiye erişebildiği, adayları ve politikaları rasyonel bir şekilde değerlendirebildiği ve hür iradesiyle karar verebildiği varsayımına dayanır. Ancak seçim süreçlerinde dolaşıma sokulan, liderlerin hiç söylemedikleri sözleri söylüyormuş gibi gösteren videolar veya ses kayıtları, seçmenin gerçeklik algısını bozmakta ve bilgi kaynaklarına olan temel güveni zedelemektedir. Bu durum, siyasal iletişim literatüründe *epistemik güven krizi* veya *epistemik anarşi* olarak tanımlanan durumu derinleştirmekte; toplumları kutuplaşmaya, komplo teorilerine ve radikalizme daha açık hale getirmektedir (İslam vd., 2024: 460-478; Hancock ve Bailenson, 2021).

Bu çalışmanın temel amacı, deepfake teknolojisinin teknik altyapısından başlayarak, yarattığı sosyolojik ve politik sarsıntıları kapsamlı bir şekilde

analiz etmektir. Çalışma, teknolojinin nasıl çalıştığını, hangi teorik çerçevelerle (Baudrillard, Habermas) anlamlandırılabilirliğini, dünya genelindeki seçimlerde (Slovakya, Hindistan, ABD, Türkiye) nasıl kullanıldığını ve hukuk sistemlerinin bu yeni tehdide nasıl yanıt verdiğini detaylı bir şekilde inceleyecektir. Çalışmanın kapsamı, sadece teknolojinin yıkıcı etkilerini betimlemekle sınırlı kalmayıp, aynı zamanda *yalancının kârı* gibi siyasi stratejilerin ve medya okuryazarlığı eğitimlerinin paradoksal sonuçlarının (geri tepme etkisi) derinlemesine analizini de içermektedir. Bu bağlamda, deepfake sadece bir siber güvenlik sorunu değil, modern demokrasilerin karşı karşıya olduğu en büyük *hakikat sınavı* olarak ele alınmaktadır.

2. Teknolojik Şecere: GAN'lardan Difüzyon Modellerine Evrim

Deepfake teknolojisinin yarattığı toplumsal etkiyi tam olarak kavrayabilmek için, bu manipülasyonun arkasındaki teknik mimarinin ve tarihsel evrimin anlaşılması elzemdir. Sentetik medyanın evrimi, basit fotoğraf düzenlemelerinden, insan gözünün ve kulağının ayırt edemeyeceği karmaşıklıkta video ve ses sentezlerine doğru hızlı ve ürkütücü bir seyir izlemiştir. Bu evrim süreci, akademik literatürde genellikle üç ana dönemde incelenmektedir: Erken dönem derin öğrenme teknikleri, GAN devrimi ve modern difüzyon modelleri çağı.

Deepfake teknolojisinin ilk nüveleri, derin öğrenme algoritmalarının, özellikle de *oto kodlayıcıların* (autoencoders) görüntü işleme alanında kullanılmasıyla atılmıştır. Oto kodlayıcılar, veriyi sıkıştırmak (encode) ve ardından yeniden oluşturmak (decode) için eğitilen sinir ağlarıdır. Yüz değiştirme (face-swapping) uygulamalarında, iki farklı kişinin yüzü aynı kodlayıcı (encoder) ile eğitilir, ancak farklı kod çözücüler (decoder) kullanılır. Bu sayede, A kişinin yüz ifadeleri ve mimikleri, B kişinin yüzüne yapıştırılabilir. Bu yöntem, 2017 civarında Reddit ve benzeri platformlarda ortaya çıkan ilk amatör deepfake videolarının temelini oluşturmuştur (Akool, t.y.). Ancak bu erken dönem teknolojisi, özellikle göz kırpması, dudak senkronizasyonu ve ışıklandırma gibi detaylarda ciddi hatalar (artifacts) barındırıyordu. Videolardaki titremeler veya yüz hatlarındaki bulanıklıklar, sahteliğin çıplak gözle anlaşılmasını kolaylaştırıyordu.

Deepfake teknolojisinin miladı ve asıl sıçrama noktası, 2014 yılında Ian Goodfellow ve ekibinin *Üretken Çekişmeli Ağlar* (Generative Adversarial Networks- GANs) mimarisini literatüre kazandırmasıyla başlamıştır (Akool, t.y.). GAN mimarisi, yapay zekâ tarihinde bir devrim niteliğindedir çünkü makinenin *hayal etme* yeteneğine en yaklaştığı noktayı temsil eder. Sistem, birbirleriyle sürekli rekabet eden iki ana sinir ağından oluşur:

1. *Üretici (Generator)*: Rastgele gürültüden veriler (görüntü, ses vb.) üreterek, bunları gerçek veriye benzetmeye çalışan ağıdır. Üretici ağın amacı, sahte veriyi o kadar gerçekçi üretmektir ki, karşı tarafı kandırabilsin. Bir nevi *kalpazan* rolündedir.
2. *Ayırt Edici (Discriminator)*: Kendisine sunulan verinin gerçek veri setinden mi (örneğin gerçek insan yüzleri veritabanı) yoksa Üretici tarafından mı oluşturulduğunu tespit etmeye çalışan ağıdır. Bir nevi *polis* veya *sanat uzmanı* rolündedir (Stroebel vd., 2023: 85-109).

Bu iki ağ arasındaki çekişme (adversarial game), sıfır toplamalı bir oyun teorisi mantığıyla işler. Üretici her seferinde Ayırt Edici'yi kandırmak için daha iyi sahteler üretir; Ayırt Edici ise her seferinde sahteleri yakalamak için kriterlerini daha da hassaslaştırır. Binlerce, hatta milyonlarca tekrar sonucunda, Üretici ağ, Ayırt Edici'nin bile gerçek sandığı *hiper-gerekçi* sentetik veriler üretmeyi öğrenir. Bu süreç, sadece görüntü kalitesini artırmakla kalmamış, aynı zamanda insan yüzünün mikro-mimiklerini, cilt dokusunu ve ışık yansımalarını da kusursuz bir şekilde taklit etmeyi mümkün kılmıştır (Stroebel vd., 2023: 89-95).

Son yıllarda, özellikle 2020 sonrası dönemde, sentetik medya üretiminde GAN'ların tahtını sallayan ve kaliteyi bir üst seviyeye taşıyan yeni bir mimari ortaya çıkmıştır: *Difüzyon Modelleri* (Diffusion Models). Midjourney, DALL-E 3 ve Stable Diffusion gibi popüler metinden-görüntüye (text-to-image) araçlarının temelini oluşturan bu modeller, termodinamikteki difüzyon prensibinden ilham almıştır.

Difüzyon modelleri, bir görüntüye kademeli olarak rastgele gürültü (Gaussian noise) ekleyerek görüntüyü tamamen bozma (forward process) ve ardından bu süreci tersine çevirerek (reverse process) gürültüden net bir görüntü elde etme prensibine dayanır (Babaei vd., 2025). Yapılan teknik karşılaştırmalar, Difüzyon Modellerinin GAN'lara kıyasla daha çeşitli, daha yüksek çözünürlüklü ve anlamsal bütünlüğü daha güçlü görüntüler üretebildiğini göstermektedir. Görüntü gerçekçiliğini ölçen standart bir metrik olan *Frechet Inception Distance* (FID) skorlarında, Difüzyon Modelleri (örneğin FID: 31.3), GAN modellerine (örneğin FID: 40.2) göre daha düşük ve dolayısıyla daha başarılı sonuçlar vermektedir (Luo, 2024).

Bu teknolojik geçişin siyasi dezenformasyon açısından kritik bir önemi vardır: Difüzyon modelleri, metin komutlarıyla çalıştığı için, teknik bilgi gerektirmeyen herhangi bir kullanıcının, *Donald Trump tutuklanırken* veya *Papa rapçi kıyafetiyle* gibi fotorealistik görüntüleri saniyeler içinde üretmesine

olanak tanır. Bu durum, dezenformasyonun üretim maliyetini sıfıra indirirken, hızını ve ölçeğini maksimize etmektedir.

Teknolojinin teknik kapasitesi kadar, erişilebilirliği de sosyolojik etkiyi belirleyen temel faktördür. 2018-2020 yılları arası, deepfake teknolojisinin *demokratikleşme* sürecine işaret eder. Başlangıçta yüksek işlemci gücü ve kodlama bilgisi gerektiren bu süreçler, DeepFaceLab, FakeApp ve daha sonra mobil uygulamalar (Zao, Reface) sayesinde tabana yayılmıştır (Acim vd., 2025).

Bu erişilebilirlik, *tehdit yüzeyini* (threat landscape) genişletmiş; eskiden sadece devlet destekli aktörlerin veya yüksek bütçeli prodüksiyon şirketlerinin tekelinde olan manipülasyon kapasitesini, bireysel aktivistlere, trol ordularına ve yerel siyasi gruplara kadar indirgemıştır. Bugün bir akıllı telefon kullanıcısı, birkaç dakika içinde bir siyasetçinin yüzünü değiştirerek veya sesini klonlayarak inandırıcı bir içerik üretebilmektedir (Bohacek, 2024). Bu durum, dezenformasyonla mücadeleyi merkezi bir kaynaktan ziyade, dağınık ve amorf bir yapıya karşı verilen bir savaşa dönüştürmüştür.

Aşağıdaki tablo, iki temel mimari arasındaki farkları ve deepfake üretimindeki rollerini özetlemektedir:

Tablo 1: GAN ve Difüzyon Modellerinin Karşılaştırmalı Analizi

Özellik	Üretken Çekişmeli Ağlar (GANs)	Difüzyon Modelleri (Diffusion Models)
Çalışma Prensibi	Üretici ve Ayırt Edici ağların rekabeti (Min-Max oyunu).	Gürültü ekleme ve tersine çevirme (Denoising).
Güçlü Yanları	Hızlı örnekleme, video üretiminde (yüz değiştirme) yüksek başarı.	Yüksek görüntü kalitesi, çeşitlilik, metinden üretim kolaylığı.
Zayıf Yanları	Eğitim dengesizliği (Mode Collapse), ince detaylarda hata.	Yavaş üretim süreci, yüksek hesaplama maliyeti.
Dezenformasyon Kullanımı	Video manipülasyonu, yüz değiştirme, dudak senkronizasyonu.	Sahte belge, fotoğraf ve <i>olay yeri</i> görüntüsü üretimi.
Gerçekçilik Skoru (FID)	Orta-Yüksek (Örn: 40.2)	Çok Yüksek (Örn: 31.3) (Luo, 2024).

Kaynak: Tablo yazar tarafından oluşturulmuştur.

3. Teorik Çerçeve: Simülasyon, Kamusal Alan ve Hakikat Sonrası

Deepfake teknolojisinin yarattığı kriz, sadece teknik bir güvenlik sorunu veya yasal bir boşluk olarak ele alınamaz (Gaur, 2022). Bu kriz, özünde ontolojik (varlık bilimsel) ve epistemolojik (bilgi kuramsal) bir kırılmadır.

Bu kırılmayı anlamlandırmak ve sosyolojik derinliğini kavramak için Jean Baudrillard'ın simülasyon teorisi, Jürgen Habermas'ın kamusal alan kavramı ve güncel *Post-truth* literatürüne başvurmak gerekir.

Fransız sosyolog Jean Baudrillard'ın *Simülakrlar ve Simülasyon* (1981) adlı eseri, bugünkü yapay zekâ destekli medya ortamını anlamak için bir çerçeve sunar. Baudrillard'a göre modern toplum, gerçeğin yerini *gerçeğin kopyalarının* (simülakr) aldığı, işaretlerin ve sembollerin gerçeğin önüne geçtiği bir *hiper-gerçeklik* (hyperreality) evrenine dönüşmüştür (Baudrillard, 1994).

Baudrillard, imgenin tarihsel dönüşümünü dört evrede inceler:

1. İmge, derin bir gerçekliğin yansımasıdır (Sakramental düzen).
2. İmge, derin bir gerçekliği maskeler ve tahrif eder (Kötü niyetli düzen).
3. İmge, derin bir gerçekliğin yokluğunu maskeler (Büyü düzeni).
4. İmge, gerçeklikle hiçbir ilişkisi olmayan, kendi saf simülakrı olan görüntüdür (Debrix, 2024).

Deepfake teknolojisi, Baudrillard'ın tanımladığı *bu dördüncü evre simülakra* tam olarak denk gelmektedir. Geleneksel sahtecilik (örneğin Photoshop ile bir fotoğraftan bir kişiyi silmek) bir orijinalin tahrif edilmesine dayanırken (ikinci evre); YZ ile *üretilen olmayan bir insanın yüzü veya hiç söylenmemiş bir sözün videosu*, referanssız bir gerçeklik üretir. Bu görüntülerin bir aslı yoktur. Baudrillard'ın ifadesiyle bu durum “kusursuz cinayettir”; gerçeğin, kendi kopyası tarafından, kopyanın aslından daha gerçekçi görünmesi suretiyle öldürülmesidir (Debrix, 2024).

Siyaset arenasında bu durum, liderlerin veya olayların aslı ile sureti arasındaki farkın silinmesi anlamına gelir. Baudrillard'ın öngördüğü gibi, kitleler artık gerçeği aramaktan vazgeçip, en ikna edici, en eğlenceli veya kendi önyargılarını en çok okşayan simülasyona inanmayı tercih etmektedir. “Yapay zeka ile üretilmiş bir politik skandal”, gerçekleşmiş bir olaydan daha gerçek etkiler (öfke, oy değişimi, protesto) doğurabilir (Nosta, 2023).

Jürgen Habermas'ın “İletişimsel Eylem Kuramı” ve “Kamusal Alan” kavramları, demokrasinin sağlıklı işleyişi için rasyonel tartışmanın ve müzakerenin önemini vurgular. Habermas'a göre, ideal bir kamusal alan, katılımcıların üzerinde uzlaştığı asgari müşterekler ve olgusal gerçekler zemininde yükselir. Vatandaşlar, argümanlarını kanıtlara dayandırarak bir fikir birliğine (konsensus) ulaşmaya çalışırlar (Łuczewski vd., 2013: 337-347).

Deepfake teknolojisi, Habermasçı idealdeki *iletişimsel rasyonaliteyi* sistematik olarak zehirlemektedir. Kamusal tartışmanın temel malzemesi olan kanıt

(bir video, bir ses kaydı, bir fotoğraf), güvenilirliğini yitirdiğinde, rasyonel tartışma zemini ortadan kalkar. Tartışma, *politika önerileri* veya *toplumsal sorunlar* üzerinden değil, kanıtın *gerçek olup olmadığı* üzerinden yürütülür hale gelir. Habermas'ın modernite eleştirisinde kullandığı “sistem dünyasının yaşam dünyasını sömürgeleştirmesi” tezi, bugün algoritmik manipülasyonların insan algısını ve iletişim süreçlerini sömürgeleştirmesi şeklinde tezahür etmektedir (Łuczewski vd., 2013: 336-348).

Bir seçim sürecinde, sahte bir kasetin kamusal alana sürülmesi, Habermas'ın “iletişimsel eylem” dediği karşılıklı anlama sürecini, stratejik ve manipülatif bir eyleme dönüştürür. Hakikat (truth) iddiası, yerini doğruluk veya samimiyet iddialarının da sorgulandığı kaotik bir şüphe ortamına bırakır (Łuczewski vd., 2013). Post-truth kavramı, nesnel hakikatlerin kamuoyunu şekillendirmede duygulara ve kişisel inançlara göre daha az etkili olduğu durumu ifade eder. Deepfake, post-truth çağının en güçlü katalizörü ve hızlandırıcısıdır (Kochetkova, 2024: 145-148).

Sosyolojik açıdan deepfake, bir *epistemik anarşi* durumu yaratır. Buradaki temel sorun, sadece insanların bir yalana inanması (misinformation) değildir; asıl tehlike, insanların artık *gerçeğe ulaşılabilir olduğuna* dair inançlarını yitirmeleridir (Patarlapati, 2023: 884-887). Hannah Arendt'in totaliterizm analizinde belirttiği gibi, ideal totaliter özne, doğru ile yanlış ayırt edemeyen değil, bu ayrımın var olduğuna dair inancını kaybetmiş olandır (The Origins of Totalitarianism).

Deepfakes, toplumları tam da bu noktaya sürüklemektedir. *Her şey sahte olabilir* düşüncesi, siyasal apatiyi (ilgisizlik) ve kurumlara (medya, yargı, devlet) olan güvensizliği derinleştirir. Vatandaşlar, gerçeği ayırt etme çabasını maliyetli ve beyhude bularak, bilişsel bir tembelliğe sığınır ve sadece kendi kabilelerinin (parti, ideoloji) sunduğu anlatıyı kabul ederler. Bu durum, toplumsal kutuplaşmayı (polarization) geri döndürülemez bir noktaya taşır (İslam vd., 2024).

4. Sosyolojik ve Politik Etkiler: Güven, Hafıza ve Yalancının Kârı

Deepfake teknolojisinin toplumsal doku üzerindeki etkileri, bireysel psikolojiden kolektif hafızaya, seçim stratejilerinden kurumsal meşruiyete kadar geniş bir yelpazeye yayılmaktadır. Bu bölümde, teknolojinin yarattığı özgün sosyopolitik fenomenler ele alınacaktır.

Siyaset bilimi literatürüne Robert Chesney ve Danielle Citron (2019) tarafından kazandırılan *Yalancının Kârı* (The Liar's Dividend) kavramı, deepfake teknolojisinin en paradoksal ve tehlikeli sonuçlarından biridir (Schiff vd., 2024: 71-88). Bu kavram, deepfake'lerin varlığının yarattığı genel şüphe

ortamının, dürüst olmayan siyasetçiler tarafından gerçek skandalları örtbas etmek için stratejik bir silah olarak kullanılmasını ifade eder.

Mekanizma şu şekilde işler:

1. Toplum, videoların ve seslerin yapay zekâ ile kusursuzca taklit edilebileceğini öğrenir (Farkındalık artışı).
2. Bir siyasetçinin yolsuzluk, ırkçılık veya uygunsuz davranışını belgeleyen *gerçek* ve *otantik* bir kayıt ortaya çıkar.
3. Siyasetçi, bu kaydın içeriğini reddetmek yerine, kaydın kendisinin “yapay zekâ ürünü bir komplo”, “montaj” veya “deepfake” olduğunu iddia eder (Schuh, 2024).
4. Medya ve kamuoyu, teknolojinin varlığını bildiği için bu iddiaya bir “şüphe payı” (benefit of the doubt) bırakır (UNESCO, t.y.).
5. Teknik analizler gerçeği ortaya çıkarana kadar (ki bu genellikle aylar sürer ve seçim geçmiş olur), siyasetçi kendi tabanını konsolide eder ve hesap vermekten kaçınır (Schiff ve Schiff, 2024).

Yale Üniversitesi’nden Schiff ve arkadaşlarının (2024) 15.000’den fazla Amerikan yetişkini üzerinde yaptığı deneysel araştırma, bu teoriyi ampirik olarak doğrulamıştır. Araştırma sonuçlarına göre, bir skandalla karşılaşan siyasetçilerin “bu bir deepfake/yalan haber” savunması yapması, sessiz kalmaya veya özürlü dilemeye göre seçmen desteğini korumada anlamlı derecede daha etkilidir. Hatta bu strateji, *ateşli destekçiler* (core supporters) arasında bir kenetlenme (rallying effect) yaratmaktadır (Schiff vd., 2024). Bu durum, gerçekliğin kendisinin kanıtlara dayalı bir olgu olmaktan çıkıp, partizan bir inanç meselesine indirgenmesine neden olmaktadır.

Deepfake teknolojisi, sadece bugünü değil, geçmişini de manipüle etme ve yeniden yazma potansiyeline sahiptir. Toplumsal bellek ve tarih bilinci, büyük ölçüde görsel ve işitsel arşivler üzerine kuruludur. Yapay zeka ile üretilen sahte tarihi belgeler, Holokost gibi tarihsel travmaların inkar edilmesinde, ulusal kahramanların itibarsızlaştırılmasında veya tam tersine, geçmişteki diktatörlerin aklanmasında kullanılabilir (Makhortykh, 2024).

Murphy ve Flynn (2021) tarafından yapılan bir araştırma, insanların sahte video içeriklerine maruz kaldıklarında, olayları yanlış hatırlama eğilimlerinin arttığını göstermektedir. *Mandela Etkisine* benzer bir mekanizmayla, hiç gerçekleşmemiş bir olayın deepfake videosunu izleyen deneklerin, bu olayı hatırladıklarını beyan etme oranları, sadece metin okuyanlara göre anlamlı derecede yüksektir (Tas, 2023).

Literatürde sentetik bellek (synthetic memory) olarak adlandırılan bu olgu, ulusların tarih anlatısını, kültürel mirasını ve ortak geçmiş algısını tehdit etmektedir. Örneğin, bir kurucu liderin geçmişte söylemediği sözleri söylemiş gibi gösteren tarihi dokulu siyah-beyaz bir deepfake video, o liderin mirasını ve dolayısıyla üzerine kurulu olduğu devletin meşruiyetini sorgulatabilir (Orquidea, 2025). Bu durum, Habermas'ın *geçmişin kamusal kullanımı* kavramını, *geçmişin algoritmik manipülasyonuna* dönüştürür.

Deepfake içerikler, sosyal medya platformlarının algoritmik yapısı içinde, *doğrulama yanlışlığı* mekanizmasını tetiklemek üzere tasarlanmaktadır. Bireyler, halihazırda sevmedikleri bir siyasetçiyi aşağılayan veya kendi görüşlerini destekleyen sahte bir videoya inanmaya psikolojik olarak daha meyillidirler (Etienne, 2021: 555-560). Yankı odalarında yayılan bu içerikler, toplumun farklı kesimlerinin farklı gerçekliklerde yaşamasına neden olur. Bir grup için kesin kanıt niteliğindeki video, diğer grup için *dış güçlerin deepfake operasyonudur*. Bu epistemik kopuş, toplumsal diyalogu inkânsız hale getirir ve siyasi kutuplaşmayı (polarization) radikalize eder. Öteki, artık sadece farklı düşünen biri değil, apaçık gerçeği (bizim gerçeğimizi) inkâr eden bir hain veya cahil olarak kodlanır (Islam vd., 2024).

Aşağıdaki tablo, *Yalancının Kârı* mekanizmasının nasıl işlediğini ve toplumsal etkilerini özetlemektedir:

Tablo 2: Yalancının Kârı (Liar's Dividend) Mekanizması

Aşama	Eylem	Sosyolojik Etki
1. Skandal Patlak Verir	Gerçek bir video/ses kaydı yayınlanır.	Kamuoyunda şok ve öfke oluşur.
2. İnkâr Stratejisi	Siyasetçi Bu AI üretimi, deepfake der.	Belirsizlik tohumları ekilir.
3. Kabileci Refleks	Taraftarlar lidere inanır, karşıtlar videoya.	Epistemik kutuplaşma derinleşir.
4. Medya Felci	Basın teknik doğrulama yapamaz, tartışma uzar.	Habermasçı rasyonel tartışma çöker.
5. Sonuç	Gerçek kanıt etkisizleşir, lider güçlenir.	Hakikat sonrası siyaset kurumsallaşır.

Kaynak: Tablo yazar tarafından oluşturulmuştur.

5. Küresel Vaka Analizleri: Seçim Sandığında Yapay Zekâ

2023 ve 2024 yılları, dünya nüfusunun büyük bir kısmının sandığa gittiği ve deepfake teknolojisinin ilk kez kitlesel ölçekte, organize bir şekilde seçimlere

müdahale ettiği bir *laboratuvar dönemi* olmuştur. Aşağıdaki vaka analizleri, teknolojinin farklı siyasi kültürlerde nasıl farklı tezahür ettiğini göstermektedir.

Eylül 2023 Slovakya parlamento seçimleri, deepfake'in bir seçim sonucunu doğrudan etkileme potansiyelini gösteren en net ve erken örneklerden biridir. Seçimlere sadece 48 saat kala, ülkedeki seçim yasaklarının (moratorium) başladığı ve medyanın siyasi içerik yayınlamadığı bir anda, Facebook ve WhatsApp üzerinden bir ses kaydı viral olmuştur. Bu kayıтта, anketlerde önde giden Batı yanlısı İlerici Slovakya partisinin lideri Michal Šimečka'nın, ülkenin önde gelen gazetecilerinden Monika Tódová ile konuştuğu iddia ediliyordu. Ses kaydında ikili, "seçim hilesi yaparak Roman azınlığın oylarını satın alma" ve "seçimden sonra bira fiyatlarını iki katına çıkarma" gibi halkı provoke edecek konuları konuşuyordu (Schneier ve Sanders, 2024).

Analiz:

- *Zamanlama Stratejisi:* Kaydın, adayların ve medyanın yasal olarak açıklama yapamayacağı sessizlik döneminde yayınlanması, yalanlamayı ve teknik doğrulamayı imkânsız kılmıştır. Bu, hukuk sisteminin teknolojinin hızına yenildiği andır.
- *Format:* Ses (audio) deepfake kullanılması bilinçlidir; zira sesin manipülasyonu videodan daha kolaydır, daha az veri gerektirir ve WhatsApp gibi kapalı devre (dark social) uygulamalarında hızla yayılır (Bohacek, 2024).
- *Sonuç:* Seçimi az bir farkla Rusya yanlısı popülist lider Robert Fico kazanmıştır. Analistler, bu son dakika manipülasyonunun kararsız seçmen üzerinde belirleyici olduğunu değerlendirmektedir (Schneier ve Sanders, 2024).

Dünyanın en büyük demokrasisi olan Hindistan'da 2024 seçimleri, YZ'nin hem bir kampanya aracı hem de bir dezenformasyon silahı olarak kullanıldığı devasa bir sahneye dönüşmüştür (Sharma, 2024; Mirza, 2024).

- *Ölülerin Diriltilmesi (Resurrection):* Tamil Nadu eyaletinde, Dravida Munnetra Kazhagam (DMK) partisi, yıllar önce ölen ikonik lider M. Karunanidhi'yi YZ kullanarak diriltmiştir. Karunanidhi'nin deepfake avatari, videolarda güneş gözlükleri ve karakteristik duruşuyla belirerek, kendi partisinin mevcut adayına oy istemiştir (Rebelo, 2024). Bu durum, etik tartışmaları (ölünün rızası) beraberinde getirmiştir.
- *Dil Bariyerinin Aşılması:* Başbakan Narendra Modi'nin konuşmaları, yapay zeka araçları (Kharvi, 2024) (Bhashini vb.) kullanılarak gerçek zamanlı olarak farklı Hint lehçelerine (Tamilce, Teluguca vb.)

çevrilmiş ve seçmenin kendi dilinde hitap eden videolar üretilmiştir. Bu, teknolojinin kapsayıcı ve faydalı kullanımına bir örnektir.

- *Dezenformasyon ve Bollywood:* Seçim sürecinde, ünlü Bollywood yıldızları Aamir Khan ve Ranveer Singh'in, iktidardaki BJP partisini eleştirdiği ve Kongre partisine oy istediği sahte videolar geniş kitlelere yayılmıştır. Bu videolar, aslında oyuncuların başka konularda konuştukları eski videoların dudak senkronizasyonu (lip-sync) ve ses klonlama ile değiştirilmesiyle üretilmiştir (Lakshané, 2024). Hindistan örneği, basit kurgu (cheapfake) ile deepfake arasındaki çizginin bulanıklaştığını ve kitlelerin teknolojiye olan hayranlığı ile korkusunun iç içe geçtiğini göstermiştir.
- *ABD:* Ocak 2024 New Hampshire ön seçimlerinde, Başkan Joe Biden'ın sesini taklit eden bir otomatik arama (robocall) binlerce seçmene ulaşmıştır. Sahte Biden sesi, "Oylarınızı Kasım ayına saklayın, Salı günü oy kullanmak sadece Cumhuriyetçilerin işine yarar" diyerek Demokrat seçmenleri sandığa gitmemeleri konusunda manipüle etmeye çalışmıştır (Voter Suppression) (Schneier ve Sanders, 2024). Bu olay, kişisel telefon aramalarının bile güvenilirmez olduğu bir dönemi başlatmış ve Federal İletişim Komisyonu'nun (FCC) YZ seslerini robocall'larda yasaklamasına yol açmıştır (Loux ve Loux, 2025; Schiff vd., 2025).

Türkiye: Türkiye siyaseti, uzun süredir montaj kaset tartışmalarına aşinadır ancak deepfake teknolojisi bu tehdidi daha sofistike bir boyuta taşımıştır.

- *2023 Cumhurbaşkanlığı Seçimleri:* Seçim sürecinde, iktidar partisinin mitinglerinde muhalefet lideri Kemal Kılıçdaroğlu'nun terör örgütü PKK liderleriyle aynı videoda gösterildiği görüntüler izletilmiştir. Bu görüntülerin montaj olduğunun kabul edilmesine rağmen (Ama montaj, ama şu, ama bu ifadesiyle), miting meydanlarında siyasi bir argüman olarak kullanılması, post-truth siyasetin Türkiye'deki en net tezahürü olmuştur (Polishchuk, 2024).
- *Muharrem İnce Vakanı:* Daha dramatik olanı, Cumhurbaşkanı adaylarından Muharrem İnce'nin, seçimlere günler kala kendisine ait olduğu iddia edilen deepfake pornografik görüntüler ve sahte dekontlar nedeniyle adaylıktan çekilmesidir. İnce, bu görüntülerin "İsrail menşeli sitelerden alınan görüntülerin üzerine kendi yüzünün yapılandırılmasıyla" oluşturulduğunu belirtmiştir (Gehring vd., 2024). Bu olay, TCK kapsamındaki "özel hayatın gizliliği" ve "şantaj" suçlarının dijital boyutunu ve teknolojinin bir siyasi tasfiye aracı olarak yıkıcı gücünü göstermiştir (Hukuk ve Bilişim, t.y.).

6. Hukuki Çerçeve ve Düzenleme Arayışları

Deepfake teknolojisinin hızı, geleneksel hukuk sistemlerinin adaptasyon kapasitesini zorlamaktadır. Küresel ve ulusal düzeyde çeşitli düzenleme girişimleri mevcut olsa da “yasa her zaman teknolojinin gerisinden gelir” prensibi geçerliliğini korumaktadır.

Dünyanın ilk kapsamlı YZ yasası olan bu düzenleme, deepfake’leri risk temelli bir yaklaşımla ele almaktadır.

- *Şeffaflık Yükümlülüğü (Transparency Obligations)*: Yasa (Madde 52), YZ ile üretilen veya manipüle edilen içeriklerin (deepfakes), kullanıcıların yapay bir içerikle etkileşimde olduğunu anlayabileceği şekilde açıkça etiketlenmesini zorunlu kılar (Future of Life Institute, 2024; European Parliament, 2023).
- *Risk Kategorizasyonu*: Deepfake sistemleri genellikle sınırlı risk kategorisinde değerlendirilirken, seçimleri manipüle etme potansiyeli taşıyan kullanımlar veya biyometrik kategorizasyon gibi alanlar yüksek risk veya kabul edilemez risk statüsüne girebilmektedir.
- *Yaptırımlar*: Yasa, ihlaller durumunda şirketlere küresel cirolarının %7’sine veya 35 milyon Euro’ya varan ağır para cezaları öngörmektedir (Fragale ve Grilli, 2024). Ancak yasanın uygulanabilirliği, özellikle AB dışındaki (örneğin Rusya veya Çin merkezli) anonim aktörlere karşı sınırlıdır.

Türkiye’de henüz deepfake teknolojisini spesifik olarak düzenleyen müstakil bir yasa bulunmamaktadır. Mevcut durumda Türk Ceza Kanunu’nun (TCK) ilgili maddeleri kıyas yoluyla uygulanmaktadır (Hukuk ve Bilişim, t.y.):

- *TCK Madde 134 (Özel Hayatın Gizliliğini İhlal)*: Rıza dışı üretilen sahte pornografik içerikler veya özel hayatı ifşa eden deepfake’ler bu kapsamda değerlendirilir. Yargıtay kararları, “görüntülerin hukuka aykırı olarak ifşa edilmesini” suç saymaktadır (Doğan, 2025).
- *TCK Madde 125 (Hakaret)*: Kişinin şeref ve saygınlığını zedeleyen, onu küçük düşüren sahte içerikler bu madde kapsamında cezalandırılabilir.
- *TCK Madde 217/A (Halkı Yıltıcı Bilgiyi Alenen Yayma)*: 2022 yılında yürürlüğe giren ve kamuoyunda *Dezenformasyon Yasası* olarak bilinen bu madde, deepfake yoluyla yayılan sahte haberler için en uygun ceza normu olarak görünmektedir. Ancak *gerçeğe aykırı bilgi* tanımının YZ bağlamında nasıl yorumlanacağı ve *kamu düzenini bozma* suçunun nasıl ispatlanacağı, yargı içtihatlarına bağlıdır.

- *Yeni Kanun Teklifi (2/2234)*: 2024 yılında TBMM'ye sunulan Yapay Zeka Kanun Teklifi, Türkiye'nin bu alandaki ilk somut yasal girişimidir. Teklif, deepfake içeriklerin etiketlenmesi zorunluluğunu, izinsiz kullanımın cezalandırılmasını ve sosyal ağ sağlayıcılarına (Twitter, Meta vb.) içerik kaldırma konusunda yükümlülükler getirilmesini öngörmektedir. Teklifin, AB Yapay Zekâ Yasası ile uyumlu olması hedeflenmektedir (Çakmakcı, 2024).

Hukuki düzenlemelerin önündeki en büyük engel anonimlik ve sınıraşan suçlar meselesidir. Bir deepfake videosunun sunucusu yurtdışında olduğunda, fail VPN veya Tor ağları arkasına gizlendiğinde veya video Dark Web üzerinden yayıldığında, ulusal yasaların caydırıcılığı azalmaktadır (Trombevski, 2025). Ayrıca, *ifade özgürlüğü* ile *dezenformasyonla mücadele* arasındaki ince çizgi, özellikle siyasi hiciv (satire) veya parodi amaçlı deepfake'lerin yasaklanması durumunda tartışmalı hale gelmektedir (Kapoor ve Narayanan, 2024).

7. Çözüm Arayışları ve Medya Okuryazarlığı Paradoksu

Teknolojik tehditlere karşı hem teknolojik (tespit algoritmaları) hem de eğitsel (medya okuryazarlığı) savunma mekanizmaları geliştirilmektedir. Ancak bu çözümlerin her biri, kendi içinde yeni riskler ve paradokslar barındırmaktadır.

Teknoloji şirketleri (Adobe, Microsoft, Intel), *İçerik Kökeni ve Doğruluğu Koalisyonu* (C2PA) (National Security Agency vd, 2025) gibi inisiyatiflerle, dijital içeriğin kaynağını ve değiştirilme geçmişini belgeleyen standartlar geliştirmektedir. Bu standartlar, içerik üretildiği anda dosyaya kriptografik bir imza (dijital filigran) eklemeyi hedefler (Netizen, 2025).

Ancak bu yöntemin ciddi zaafı vardır:

- *Açık Kaynak Sorunu*: Stable Diffusion gibi açık kaynaklı modellerde, kullanıcılar filigran ekleme kodunu yazılımdan çıkarabilir.
- *Veri Kaybı*: Sosyal medya platformları, yüklenen görüntüleri sıkıştırdığında veya formatını değiştirdiğinde, bu meta veriler (metadata) genellikle silinmektedir.
- *Kedi-Fare Oyunu*: Tespit algoritmaları ile üretim algoritmaları arasında bitmeyen bir yarış vardır. Genellikle üretim algoritmaları, tespit algoritmalarını kandırarak şekilde eğitildiği için, tespit sistemleri her zaman bir adım geriden gelmektedir (EY & FICCI, 2024; Laurier vd., 2024: 2-7).

Geleneksel yaklaşım, toplumu deepfake konusunda eğitmeyi ve farkındalık yaratmayı savunur. *İnternette gördüğünüz her şeye inanmayın* uyarısı, medya

okuryazarlığının temelidir (El Mokadem, 2023: 120-133). Ancak son yapılan akademik çalışmalar, bu eğitimin paradoksal ve tehlikeli bir sonucu olduğunu göstermektedir: Sinizm ve Geri Tepme Etkisi (Deng & Ahmed, 2025).

İnsanlara sürekli olarak *teknolojinin her şeyi taklit edebileceği, hiçbir videonun güvenilir olmadığı* öğretildiğinde, vatandaşların *gerçek* kanıtlara (gerçek savaş suçları, gerçek yolsuzluk videoları, gerçek insan hakları ihlalleri) inanma oranları da düşmektedir. Bu durum, yukarıda bahsedilen *Yalancının Kârını* besler. Seçmen, bilişsel yükten kurtulmak için *gözlerime güvenemem, o halde sadece liderime güvenirim* diyerek, eleştirel düşüncüyü terk eder. Bu nedenle, yeni nesil medya okuryazarlığı eğitimleri, sadece *şüphe etmeyi değil, nasıl doğrulama yapılacağını* (proaktif), güvenilir kaynak hiyerarşisini ve teyit platformlarını (Fact-checking organizations) kullanmayı öğretmelidir. Aksi takdirde eğitim, toplumu aydınlatmak yerine, onu nihilist bir güvensizlik durumuna itebilir (Luitse, 2019).

8. Sonuç

Deepfake teknolojisi, modern siyasal iletişimde ve toplumsal yapıda geri döndürülemez bir dönüşümü tetiklemiştir. Bu teknoloji, sadece seçim dönemlerinde kullanılan bir *dezenformasyon aracı* değil, demokratik sistemlerin dayandığı *ortak gerçeklik zeminini* ve *rasyonel müzakere imkanını* ortadan kaldıran bir *güven çözücüdür*.

Bu kapsamlı analizden elde edilen temel çıkarımlar şunlardır:

Asimetrik Savaş: Deepfake üretiminin maliyeti ve teknik bariyeri sıfıra yaklaşırken (difüzyon modelleri sayesinde), savunmanın maliyeti (tespit, doğrulama, hukuki takip) katlanarak artmaktadır. Bu asimetri, kötü niyetli aktörleri avantajlı kılmaktadır.

Sosyolojik Parçalanma: Toplamlar, sadece siyasi görüşleri farklı olan gruplara değil, gerçeğin ne olduğu konusunda anlaşılamayan, farklı epistemolojik evrenlerde yaşayan kabilelere bölünmektedir. Bu durum, Habermasçı kamusal alanın sonunu getirmekte ve Baudrillard'ın simülasyon evrenini kalıcı hale getirmektedir.

Politik Fırsatçılık: Siyasetçiler, deepfake tehdidini, kendi gerçek hatalarını ve suçlarını örtmek için bir kalkan (Yalancının Kârı) olarak kullanmakta ustalaşmaktadır. Bu, demokratik hesap verilebilirliği (accountability) yok etmektedir.

Hukuki Yetersizlik ve Hibrit Çözüm: Ulusal yasalar ve cezai yaptırımlar, teknolojinin hızına ve sınır tanımazlığına yetişememektedir. Çözüm, sadece yasaklamada değil; teknolojik standartların (C2PA) (National Security Agency

vd, 2025), platformların algoritmik sorumluluğunun ve *doğrulama mimarisinin* internetin altyapısına entegre edilmesindedir.

Gelecekte, seçim süreçlerinin güvenliği, sadece sandıkların fiziksel güvenliğine veya oyların doğru sayılmasına değil, seçmenin zihinsel ve algısal güvenliğine de bağlı olacaktır. Demokrasilerin hayatta kalması, *görmenin inanmak olmadığı* bu yeni çağda, gerçeği yeniden inşa edebilecek, bağımsız, şeffaf ve güvenilir kurumsal mekanizmaların (teyitçiler, bağımsız yargı, bilimsel kurullar) kurulmasına bağlıdır. Aksi takdirde, siyaset, en iyi simülasyonu üretenin kazandığı, halkın ise sadece seyirci kaldığı bir gösteriden ibaret kalacaktır.

Tablo 3: Deepfake Teknolojisinin Seçimlere Etkisi- Vaka Karşılaştırması

Ülke	Seçim/Olay Yılı	Kullanılan Medya Tipi	Hedef/İçerik	Politik Sonuç/Etki
Slovakya	2023 Genel Seçim	Ses (Audio)	Muhalefet liderinin seçim hilesi yaptığı ve bira fiyatlarını artıracığı iddiası.	Seçim yasaklarında yayıldı, yalanlanamadı, popülist parti kazandı.
Hindistan	2024 Genel Seçim	Video ve Ses (Avatar)	Ölmüş liderlerin (M. Karunanidhi) diriltilmesi, muhalefetin sahte itirafları.	Kitlesel erişim, “cheapfake” ile karışık yoğun bilgi kirliliği ve etik tartışmalar.
ABD	2024 Ön Seçim	Ses (Robocall)	Biden’ın sesinden “oy vermeyin” çağrısı.	Seçmen baskılama girişimi, FCC yasakları getirdi.
Türkiye	2023 Genel Seçim	Video (Montaj) Görüntü	Muhalefetin terörle iltisaklı gösterilmesi, M. İnce’ye kaset kumpası.	Siyasal kutuplaşmayı derinleştirdi, aday çekilmesine yol açtı, TCK tartışmalarını başlattı.

Kaynak: Tablo Yazar tarafından oluşturulmuştur.

Kaynakça

- Acim, B., Boukhelif, M., Ouhnni, H., Kharmoum, N., & Ziti, S. (2025). A decade of deepfake research in the generative AI era, 2014–2024: A bibliometric analysis. *Publications*, 13(4), Article 50. <https://doi.org/10.3390/publications13040050>
- Akool. (n.d.). *History of deepfake technology*. <https://akool.com/knowledge-base-article/history-of-deepfake-technology>
- Babaei, R., Cheng, S., Duan, R., & Zhao, S. (2025). Generative artificial intelligence and the evolving challenge of deepfake detection: A systematic analysis. *Journal of Sensor and Actuator Networks*, 14(1), Article 17. <https://doi.org/10.3390/jsan14010017>
- Baudrillard, J. (1994). *Simulacra and simulation* (S. F. Glaser, Çev.). University of Michigan Press.
- Bohacek, M. (2024). Slovakia as the precursor to deepfake-enabled election interference: Lessons learned and pathways forward. *Workshop Proceedings of the 18th International AAAI Conference on Web and Social Media*. <https://doi.org/10.36190/2024.67>
- Çakmakçı, R. (2024, 24 Temmuz). *Yapay zeka kanun teklifi üzerine ChatGPT-5.1 ile yapılan hukuk röportajı*. Legal Blog. <https://legal.com.tr/blog/anayasa-hukuku/yapay-zeka-kanun-teklifi-uzerine-chatgpt-5-1-ile-yapilan-hukuk-roportaji/>
- Debrix, F. (2024, 31 Aralık). *AI's perfect crime*. Baudrillard Now. <https://baudrillard-scijournal.com/ais-perfect-crime/>
- Deng, R., & Ahmed, S. (2025). Breaking the illusion: Impact of news literacy on deepfake identification and sharing. *Online Information Review*, 49(6), 1211–1230. <https://doi.org/10.1108/OIR-10-2024-0613>
- Doğan, B. (2025, 20 Şubat). *TCK madde 134 özel hayatın gizliliğini ihlal suçu*. <https://barandogan.av.tr/blog/mevzuat/tck-madde-134-ozel-hayatin-gizliliğini-ihlal-sucu.html>
- El Mokadem, S. S. (2023). The effect of media literacy on misinformation and deep fake video detection. *Arab Media & Society*, 35, 115–138. <https://doi.org/10.70090/SM23EMLM>
- Etienne, H. (2021). The future of online trust (and why Deepfake is advancing it). *AI Ethics*, 1(4), 553–562. <https://doi.org/10.1007/s43681-021-00072-1>
- European Parliament. (2023, 1 Haziran). *EU AI Act: First regulation on artificial intelligence*. <https://www.europarl.europa.eu/topics/en/article/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence>
- EY, & FICCI. (2024, Eylül). *Identifying AI generated content in the digital age: The role of watermarking*. <https://www.ey.com/content/dam/ey-unified-si>

- te/ey-com/en-in/insights/ai/documents/ey-identifying-ai-generated-content-in-the-digital-age-the-role-of-watermarking.pdf
- Fragale, M., & Grilli, V. (2024, 11 Kasım). *Deepfake, deep trouble: The European AI Act and the fight against AI-generated misinformation*. Columbia Journal of European Law. <https://cjel.law.columbia.edu/preliminary-reference/2024/deepfake-deep-trouble-the-european-ai-act-and-the-fight-against-ai-generated-misinformation/>
- Future of Life Institute. (2024, 30 Mayıs). *High-level summary of the AI Act*. <https://artificialintelligenceact.eu/high-level-summary>
- Gaur, L. (Ed.). (2022). *DeepFakes: Creation, detection, and impact* (1. baskı). CRC Press. <https://doi.org/10.1201/9781003231493>
- Gehringer, F. A., Nehring, C., & Łabuz, M. (2024, 10 Mayıs). *The influence of deep fakes on elections*. Konrad-Adenauer-Stiftung. <https://www.kas.de/en/monitor/detail/-/content/the-influence-of-deep-fakes-on-elections>
- Hancock, J. T., & Bailenson, J. N. (2021). The social impact of deepfakes. *Cyberpsychology, Behavior, and Social Networking*, 24(3), 149–152.
- Hukuk ve Bilişim. (t.y.). *Deepfake ve cezai boyutu*. <https://hukukvebilisim.org/deepfake-ve-cezai-boyutu/>
- Islam, M. B. E., Haseeb, M., Batool, H., Ahtasham, N., & Muhammad, Z. (2024). AI threats to politics, elections, and democracy: A blockchain-based deepfake authenticity verification framework. *Blockchains*, 2(4), 458–481. <https://doi.org/10.3390/blockchains2040020>
- Kapoor, S., & Narayanan, A. (2024, 13 Aralık). *We looked at 78 election deepfakes: Political misinformation is not an AI problem*. Knight First Amendment Institute at Columbia University. <https://knightcolumbia.org/blog/we-looked-at-78-election-deepfakes-political-misinformation-is-not-an-ai-problem>
- Kharvi, P. L. (2024). Understanding the impact of AI-generated deepfakes on public opinion, political discourse, and personal security in social media. *IEEE Security & Privacy*, 22(4), 115–122. <https://doi.org/10.1109/MSEC.2024.3405963>
- Kochetkova, N. P. (2024). Vliyaniye dipfeykov i postpravdy na obshchestvennoye soznaniye: fenomenologicheskii analiz [The influence of deepfakes and post-truth on public consciousness: Phenomenological analysis]. *Obshchestvo: Filosofiya, Istoriya, Kul'tura*, (7), 144–150. <https://doi.org/10.24158/fik.2024.7.19>
- Lakshané, R. (2024, 9 Nisan). *Indian elections 2024: Social media, misinformation, and regulatory challenges*. Heinrich Böll Stiftung | Regional Office New Delhi. <https://in.boell.org/en/elections-2024-media>
- Laurier, L., Giulietta, A., Octavia, A., & Cleti, M. (2024). *The cat and mouse game: The ongoing arms race between diffusion models and detection methods*. arXiv. <https://arxiv.org/abs/2410.18866>

- Loux, M., & Loux, B. (2025, 11 Eylül). *What is deepfake technology? Understanding its broad impact*. American Military University. <https://www.amu.apus.edu/area-of-study/information-technology/resources/what-is-deepfake-technology/>
- Łuczewski, M., Małlanka, T., & Bednarz-Łuczewska, P. (2013). Bringing Habermas to memory studies. *Polish Sociological Review*, 183(3), 335–349. <https://polish-sociological-review.eu/pdf-125612-53634?filename=Bringing%20Habermas%20to.pdf>
- Luitse, D. (2019). *Does digital literacy have a role in mitigating the spread of misinformation? A critical analysis of deepfake videos and their comment sections* [Pre-master tezi, Utrecht University]. https://studenttheses.uu.nl/bitstream/handle/20.500.12932/823/BAEW_DieuwertjeLuitse_6558372_FinalVersion.pdf?sequence=1
- Luo, A. (2024, 18 Nisan). *GANs vs. diffusion models: Putting AI to the test*. Aurora Solar. <https://aurorasolar.com/blog/putting-ai-to-the-test-generative-adversarial-networks-vs-diffusion-models/>
- Makhortykh, M. (2024). *AI and the Holocaust: Rewriting history? The impact of artificial intelligence on understanding the Holocaust*. UNESCO; World Jewish Congress. <https://doi.org/10.54675/ZHJC6844>
- Mirza, R. (2024, 16 Şubat). *How AI deepfakes threaten the 2024 elections*. The Journalist’s Resource. <https://journalistsresource.org/home/how-ai-deepfakes-threaten-the-2024-elections/>
- National Security Agency, Federal Bureau of Investigation, Cybersecurity and Infrastructure Security Agency, Australian Signals Directorate, Canadian Centre for Cyber Security, New Zealand National Cyber Security Centre, & National Cyber Security Centre. (2025, Ocak). *Content credentials: Strengthening multimedia integrity in the generative AI era*. <https://media.defense.gov/2025/Jan/29/2003634788/-1/-1/0/CSI-CONTENT-CREDENTIALS.PDF>
- Netizen. (2025, 27 Mayıs). *How C2PA, watermarking, and nightshade are shaping the battle against deepfakes*. <https://www.netizen.net/news/post/6377/how-c2pa-watermarking-and-nightshade-are-shaping-the-battle-against-deepfakes>
- Nosta, J. (2023, 7 Haziran). *The ‘scientific simulacra’: When AI and hyperreality collide*. Medium. <https://johnnosta.medium.com/the-scientific-simulacra-when-ai-and-hyperreality-collide-b676eb265045>
- Orquidea. (2025, September 2). *Synthetic memory and the future of collective recall*. <https://orquidea.ai/synthetic-memory-and-the-future-of-collective-recall/>
- Patarlapati, N. (2023). Unmasking reality: Exploring the sociological impacts of deepfake technology. *International Journal for Research in Applied Science & Engineering Technology*, 11(11), 882–889. <https://doi.org/10.22214/ijraset.2023.56639>

- Polishchuk, A. (2024, 26 Ocak). *AI poses risks to both authoritarian and democratic politics*. Wilson Center. <https://www.wilsoncenter.org/blog-post/ai-poses-risks-both-authoritarian-and-democratic-politics>
- Rebelo, K. (2024, Ekim). *India's generative AI election pilot shows artificial intelligence in campaigns is here to stay* (Generative Artificial Intelligence and Elections, Ed. I. Trauthig & S. Woolley). Center for Media Engagement, The University of Texas at Austin. <https://mediaengagement.org/wp-content/uploads/2024/10/Indias-Generative-AI-Election-Pilot-Shows-Artificial-Intelligence-In-Campaigns-Is-Here-To-Stay.pdf>
- Schiff, K. J., & Schiff, D. (2024, 25 Eylül). *Research on the "liar's dividend" gains attention*. Purdue University College of Liberal Arts. <https://www.cla.purdue.edu/news/college/2024/liars-dividend-research.html>
- Schiff, K. J., Schiff, D. S., & Bueno, N. S. (2025). The liar's dividend: Can politicians claim misinformation to evade accountability? *American Political Science Review*, 119(1), 71–90. <https://doi.org/10.1017/S0003055423001454>
- Schneier, B., & Sanders, N. E. (2024, 4 Aralık). *The apocalypse that wasn't: AI was everywhere in 2024's elections, but deepfakes and misinformation were only part of the picture*. Ash Center for Democratic Governance and Innovation. <https://ash.harvard.edu/articles/the-apocalypse-that-wasnt-ai-was-everywhere-in-2024s-elections-but-deepfakes-and-misinformation-were-only-part-of-the-picture/>
- Schneier, B., & Sanders, N. E. (2024, 4 Aralık). *The apocalypse that wasn't: AI was everywhere in 2024's elections, but deepfakes and misinformation were only part of the picture*. Ash Center for Democratic Governance and Innovation. <https://ash.harvard.edu/articles/the-apocalypse-that-wasnt-ai-was-everywhere-in-2024s-elections-but-deepfakes-and-misinformation-were-only-part-of-the-picture/>
- Schuh, J. (2024). AI as artificial memory: A global reconfiguration of our collective memory practices? *Memory Studies Review*, 1, 231–255. <https://doi.org/10.1163/29502422-01010012>
- Sharma, Y. (2024, 20 Şubat). *Deepfake democracy: Behind the AI trickery shaping India's 2024 election*. Al Jazeera. <https://www.aljazeera.com/news/2024/2/20/deepfake-democracy-behind-the-ai-trickery-shaping-indias-2024-elections>
- Stroebel, L., Llewellyn, M., Hartley, T., Ip, T. S., & Ahmed, M. (2023). A systematic literature review on the effectiveness of deepfake detection techniques. *Journal of Cyber Security Technology*, 7(2), 83–113. <https://doi.org/10.1080/23742917.2023.2192888>
- Tas, R. (2023). *Changing memories via deepfakes: Are deepfakes a threat to our own memories? Examining the influence of deepfake videos on memory retrieval* [Yüksek lisans tezi, Ghent University]. Ghent University Library. https://libstore.ugent.be/fulltxt/RUG01/003/146/162/RUG01-003146162_2023_0001_AC.pdf

- Trombevski, D. (2025, 27 Ocak). *Deepfakes and the future of AI legislation: Ethical and legal challenges*. GDPRLocal. <https://gdprlocal.com/deepfakes-and-the-future-of-ai-legislation-overcoming-the-ethical-and-legal-challenges/>
- UNESCO. (t.y.). *Case study: Can we believe what we hear?*. International Media and Information Literacy e-Platform. <https://www.unesco.org/mil4teachers/en/toolkit-media/indicator-5/activities/ai-disinformation/case-study-2>
- Westerlund, M. (2019). The emergence of deepfake technology: A review. *Technology Innovation Management Review*, 9(11), 39–52. <https://doi.org/10.22215/timreview/1282>