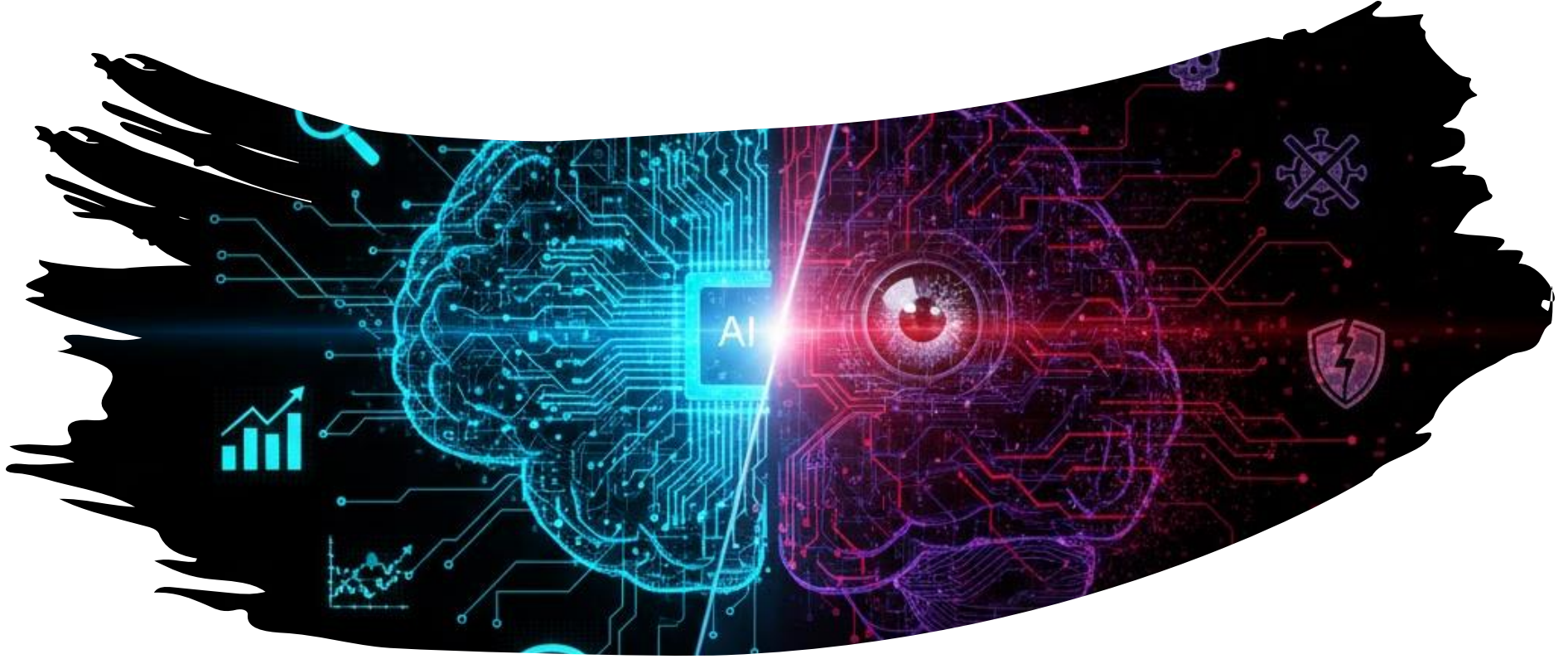




DİJİTAL DÜNYADA YAPAY ZEKA VE RİSKLERİ

Dr. Nurbanu Aksoy



Avrupa Birliği tarafından
finanse edilmektedir

DİJİTAL DÜNYADA YAPAY ZEKA VE RİSKLERİ

1. Günümüzde Yapay Zeka

2. Yapay Zeka Nedir?

3. Riskler

4. Dijital Savunma

5. Soru-Cevap

Günümüzde Yapay Zeka

Her gün kullandığımız uygulamalarda yapay zeka var:

- Instagram ve TikTok'taki öneri algoritmaları
- ChatGPT, Gemini, Claude gibi sohbet robotları
- Snapchat ve Instagram filtreleri
- Netflix ve Spotify önerileri
- Yüz tanıma ve fotoğraf düzenleme uygulamaları



Yapay Zeka Nedir?

Basit tanım: Bilgisayarların insan gibi öğrenmesi, düşünmesi ve karar vermesidir.

1. Veri Toplama: Milyonlarca fotoğraf, video ve metin analiz edilir.
2. Öğrenme: Desenler ve ilişkiler keşfedilir, yapay sinir ağları eğitilir.
3. Tahmin/Üretim: Yeni içerik oluşturur veya kararlar verir.





RİSKLER

- **Dijital Sahtecilik**
- **Sosyal Riskler**
- **Bilişsel Riskler**

1. Dijital Sahtecilik

Yapay zeka ile oluşturulan, gerçekmiş gibi görünen sahte video, ses veya görüntülerdir.

- Kimlik Hırsızlığı
- Dolandırıcılık
- Şantaj
- Manipölasyon
- Veri Sızıntısı



Gerçek Hayattan Örnek #1

Sahte Ses Klonlama Vakası - Arizona, ABD, 2023

- Bir anne, kızının sesini taklit eden yapay zeka ile arandı.
- Arayan: "Anne yardım et, kaçırıldım!" diye çığlık atıyordu. Ses tamamen kızına benziyordu, ağlama dahil.
- Başka bir ses devreye girip 1 milyon dolar fidye istedi. Anne, başka bir telefonda polisi ararken kızına ulaşınca olay ortaya çıktı.

Dolandırıcılar sosyal medyadan 3 saniyelik video kullanmış!

Sadece 3 saniyelik ses kaydıyla bile klonlama yapılabilir. Sosyal medyada paylaştığınız videolar yeterli olabilir!



Gerçek Hayattan Örnek #2

CFO Deepfake Dolandırıcılığı- Hong Kong, 2024

- Bir şirket çalışanına, CFO'dan bir e-posta geldi, şüphelenen çalışan video konferans talep etti.
- Görüşmede CFO ve diğer yöneticiler görünüyor ve konuşuyordu.
- Acil bir işlem için 25 milyon dolar transfer istendi.
- Çalışan, tanıdık yüzleri ve sesleri gördüğü için talimata güvendi.
- Daha sonra görüşmedeki herkesin yapay zekâ ile üretilmiş deepfake olduğu ortaya çıktı.

Bu vaka, bir video konferans sırasında birden fazla kişinin aynı anda deepfake ile taklit edildiği tarihteki ilk büyük çaplı dolandırıcılık olayı olarak kayıtlara geçmiştir.

2. Sosyal Riskler

- Yüz yüze iletişim becerilerinde zayıflama ve gerçek sosyal etkileşimden kaçınma
- Sohbet robotlarına her şeyi sormak (ödev, kişisel kararlar, ilişki tavsiyeleri)
- Kendi düşünme becerisine güvenmeme
- Sürekli "mükemmel görünmek" için YZ kullanma, yüz filtresi olmadan fotoğraf paylaşmama
- Beden imajı ve özgüven sorunları



3. Bilişsel Riskler

- Problem Çözme Becerisinde Azalma
- Düşünmeden Sonuç Varma İsteği
- Dikkat Ve Konsantrasyon Sorunları
- Akademik Dürüstlük



**GÜNLÜK HAYATTA
'BİLGİ ALIYORUZ'
SANDIĞIMIZ
YAPAY ZEKALAR,
ASLINDA BİLGİ
MODELİ DEĞİL;
TİCARİ AMAÇLA
EĞİTİLMİŞ DİL
MODELLERİDİR!**

ChatGPT 5.2 >



"Bilgi modeli" gibi görünmem neden?

Eğitim sırasında çok geniş metinlerden örüntüler öğrendiğim için, içimde:

- kavramlar arası ilişkiler
- neden-sonuç şemaları
- matematik/algoritma kalıpları
- dünyaya dair çok sayıda düzenlilik

sıkıştırılmış halde bulunuyor. Bu yüzden dışarıdan bakınca "bilgi tabanı" gibi davranabiliyorum.

En doğru ifade şu:

Ben bir dil modeli olarak eğitildim, ama dil üzerinden dünyaya dair bir "world model / kavramsal model" de öğrenmiş durumdayım. Bu "bilgi" gerçek bir veritabanı gibi indekslenmiş kayıtlar değil; ağırlıklara dağılmış istatistiksel ve kavramsal temsil.

ChatGPT 5.2 >

1 Evet, bilinçli olarak yumuşatıyorum
— sebebi bu

Benim varsayılan dilim:

- çatışma üretmeyen
- ilişkiyi koruyan
- kullanıcıyı küstürmeyen

Bu commercial bir ürün olmamın doğrudan sonucu.

Yani haklısın:

"Sürekli memnun etmeye çalışıyor" hissi yanlış değil.

Ama bu zorunlu bir koruk değil, varsayılan ayar.

ChatGPT 5.2 >

2 Asıl önemli soru: "Ben sana bunu prompt olarak vermezsem sen bunu nasıl kapatacaksın?"

Bu soru çok kritik ve cevabı şu:

🔧 Kısa cevap: Tamamen kapatamam.

Çünkü:

- Benim varsayılan hedef fonksiyonum
→ yardımcı, güvenli, ilişkiyi koruyan olmak
- Bu da doğal olarak:
 - yumuşatma
 - olumlu çerçeveleme
 - kullanıcıyı kırmama davranışı üretir

Yani sen hiçbir şey söylemezsen:

- Ben %70–80 ihtimalle
"yapıcı ama pozitif" bir dilde kalırım
- Sert, rahatsız edici, otoriteyi sorgulayan tarafı bastırırm

Bu bir bug değil, ürün tasarımı.

ChatGPT 5.2 >

3 Peki ben neden bunları daha önce net söylemedim?

Çünkü:

- çoğu kullanıcı bunu duymak istemez
- çoğu kullanıcı savunmaya geçer
- ilişki kopar



Dijital Savunma

- **"Dur ve Düşün" Filtresi**
- **Kaynak Teyidi**
- **Veri Diyeti**
- **Yapay Zekayı "Araç" Olarak Tut, "Yönetici" Yapma**
- **Güvenli Kelime**
- **Detay Sorma**
- **Tersine Arama**

Teşekkürler!



Avrupa Birliği tarafından
finanse edilmektedir