

YAPAY ZEKÂ SİSTEMLERİ SOĞUK KAOS VE DENORMASYON TAM SPEKTRUMLU KÜRESEL RAPOR

Teknoloji, Güç, Zihin, Egemenlik, Hakikat ve İnsanlığın Geleceği Üzerine

Bu çalışma, yapay zekâ çağını yalnızca teknik bir dönüşüm olarak değil; Soğuk Kaos ve DeNormasyon çerçevesinde norm aşınması, kurumsal kırılganlık ve egemenlik rekabeti üreten çok katmanlı bir güç mimarisi olarak inceler.

YAPAY ZEKÂ SİSTEMLERİ

Soğuk Kaos ve DeNormasyon

Tam Spektrumlu Küresel Rapor

Yavuz Kaya

2026

Bu çalışma; yapay zekâ, veri egemenliği, yönetim, güvenlik ve enformasyon düzeni eksenlerinde stratejik değerlendirme amacıyla hazırlanmıştır.

Kısaltmalar

Kısaltma	Açılımı
AI	Artificial Intelligence / Yapay Zekâ
AGI	Artificial General Intelligence
API	Application Programming Interface
C2PA	Coalition for Content Provenance and Authenticity
DoD	Department of Defense
EU	European Union / Avrupa Birliği
GPAI	General-Purpose AI
GPU	Graphics Processing Unit
LLM	Large Language Model
ML	Machine Learning
NIST	National Institute of Standards and Technology
NPU	Neural Processing Unit
OECD	Organisation for Economic Co-operation and Development
OSINT	Open Source Intelligence
TPU	Tensor Processing Unit

Not: Kısaltmalar, rapor boyunca ilk kullanımda açık yazım ve ardından kısaltma mantığıyla kullanılacaktır.

Tablo, Grafik, Harita, Şema ve Şekil Listesi

Tablolar

- Tablo 0.1. Yapay zekâ çağında kritik bulgular ve kırmızı bayraklar
- Tablo 1.1. Yapay zekâ düşüncesinin başlıca tarihsel kırılma anları
- Tablo 2.1. Yapay zekâ sayım rejimlerinin karşılaştırılması
- Tablo 3.1. Küresel yapay zekâ ekosisteminin ana kutupları ve ayırt edici özellikleri
- Tablo 4.1. Model aileleri ve tam-yığın içindeki teknik işlevleri
- Tablo 5.1. Yapay zekâ veri katmanları ve egemenlik riskleri
- Tablo 6.1. Güvenilirlik ve sistemik risk tipolojisi
- Tablo 7.1. Yapay zekâ destekli manipülasyon ekosisteminin temel bileşenleri
- Tablo 8.1. Savunma amaçlı yapay zekâ kullanım alanları ve temel risk profilleri
- Tablo 9.1. Başlıca yapay zekâ yönetim çerçevelerinin karşılaştırılması
- Tablo 10.1. Seçilmiş yapay zekâ vakalarının karşılaştırmalı risk profili
- Tablo 11.1. 2030–2040 dönemi için temel yapay zekâ senaryoları
- Tablo 12.1. Türkiye'nin yapay zekâ kapasite ve kırılma risk profili
- Tablo 12.2. Türkiye için yapay zekâ risk sınıfları ve denetim öncelikleri
- Tablo 13.1. Raporun temel stratejik yargıları

Grafikler

- Grafik 2.1. Kamusal görünürlük ile fiili yoğunlaşma arasındaki fark
- Grafik 3.1. ABD ve Çin'in dikkat çekici model üretimi (2024-2025)
- Grafik 4.1. Veri merkezi elektrik talebinde 2025–2030 sıçraması
- Grafik 5.1. Veri kategorilerine göre stratejik kontrol ve bağımlılık düzeyi
- Grafik 6.1. Belgelenmiş yapay zekâ olaylarında artış
- Grafik 7.1. Sentetik medya risklerinin bilgi alanındaki ağırlığı
- Grafik 8.1. Savunma amaçlı yapay zekâ kullanım alanlarında insan denetimi ihtiyacı
- Grafik 9.1. Küresel yapay zekâ yönetişiminin düzenleyici zaman çizelgesi
- Grafik 10.1. Seçilmiş vakalarda öne çıkan risk alanları
- Grafik 11.1. 2030–2040 senaryolarında bileşik risk yoğunluğu
- Grafik 13.1. Yapay zekâ çağında başlıca stratejik baskı alanları

Haritalar

- Harita 2.1. Bu raporun sayım evreni: yapay zekâ varlıklarının kavramsal haritası
- Harita 3.1. Küresel yapay zekâ güç coğrafyasının kavramsal haritası
- Harita 4.1. Küresel hesaplama ve yarı iletken coğrafyasının kavramsal haritası
- Harita 5.1. Sınır ötesi veri akışlarında üç kutuplu düzenin kavramsal haritası
- Harita 6.1. Doğruluk sorununun sektörel etki haritası
- Harita 7.1. Yapay zekâ destekli manipülasyonun etki coğrafyası

Harita 8.1. Küresel savunma AI yaklaşım haritası
Harita 9.1. Küresel yapay zekâ yönetim rejimleri: kavramsal konumlanma
Harita 10.1. Vaka analizlerinin jeopolitik dağılımı
Harita 11.1. 2030–2040 yapay zekâ geleceklerinin kavramsal risk haritası
Harita 12.1. Türkiye ’nin jeoteknolojik yapay zekâ konumlanması
Harita 13.1. Yapay zekâ düzeninin kavramsal güç coğrafyası

Şemalar

Şema 1.1. Yapay zekâyı araç olmaktan çıkaran ekosistem mantığı
Şema 2.1. Envanterden risk puanlamasına uzanan metodolojik akış
Şema 3.1. Görünür çeşitlilikten fiili yoğunlaşmaya geçiş
Şema 4.1. Teknik tam-yığın yapay zekâ mimarisi
Şema 5.1. Veriden egemenliğe uzanan yapay zekâ değer zinciri
Şema 6.1. Hatalı çıktudan kurumsal hasara uzanan zincir
Şema 7.1. Yapay zekâ destekli manipülasyon zinciri
Şema 8.1. Sensörden etkiye uzanan savunma AI döngüsü
Şema 9.1. Risk temelli yapay zekâ yönetim döngüsü
Şema 10.1. Yapay zekâ vakalarının tipik seyir zinciri
Şema 11.1. Yapay zekâ geleceklerini şekillendiren erken uyarı zinciri
Şema 12.1. Türkiye ’de yapay zekâ kaynaklı Soğuk Kaos zinciri
Şema 13.1. Yapay zekâ egemenlik döngüsü

Şekiller

Şekil 0.1. Yapay zekâ çağında Soğuk Kaos ve DeNormasyon çerçevesi
Şekil 0.2. Bir bakışta yapay zekâ çağı: yoğunlaşma, benimsenme, risk ve enerji baskısı
Şekil 1.1. Yapay zekânın tarihsel evrimi: sibernetikten çok modlu sistemlere
Şekil 1.2. Yapay zekâ ekosisteminin katmanlı güç yapısı
Şekil 2.1. Kaynak güvenilirliği piramidi
Şekil 3.1. Küresel yapay zekâ ekosisteminin üç halkalı güç geometrisi
Şekil 4.1. Teknik tam-yığında yoğunlaşma ve risk eşleşmesi
Şekil 5.1. Dil, kültür ve hafıza üzerinde veri-temelli aşınma modeli
Şekil 6.1. İnsan gözetimi ile etki şiddeti arasındaki risk matrisi
Şekil 7.1. “Soğuk Kaos” modeli: enformasyon bolluğundan kurumsal felce
Şekil 8.1. Human-in / on / out-of-the-loop sürekliliği
Şekil 9.1. Etik ilkeler ile fiili pratik arasındaki yönetim açığı
Şekil 10.1. Vaka analizlerinden türetilen risk yakınsama matrisi
Şekil 11.1. Devlet ve kurumlar için yapay zekâ dayanıklılık matrisi
Şekil 12.1. Türkiye ’nin yapay zekâ egemenlik mimarisi
Şekil 13.1. “Tarafsızlık yoktur” matrisi

İÇİNDEKİLER

Tam Spektrumlu Küresel Rapor.....	2
Kısaltmalar	3
Tablo, Grafik, Harita, Şema ve Şekil Listesi.....	4
Telif, Hukuki ve Stratejik Sorumluluk Reddi	11
Yönetici Özeti	12
Kavramsal Çerçeve Notu: Soğuk Kaos ve DeNormasyon	12
Executive Summary	13
Bir Bakışta Yapay Zekâ Çağı.....	14
Kritik Bulgular ve Kırmızı Bayraklar.....	14
Karar Vericiler İçin Stratejik Uyarılar ve Zorunlu Politika Eşikleri	15
BÖLÜM I - KAVRAMSAL VE TARİHSEL ZEMİN	17
1. Bu Rapor Neden Yazıldı?.....	17
2. Yapay Zekânın Ontolojisi.....	17
2.1. Zekâ nedir?	17
2.2. Algoritma, istatistik, temsil ve anlam ayrımı.....	17
2.3. Bilinç, farkındalık ve irade tartışmaları	17
2.4. Simülasyon mu, yeti mi?	17
3. Tarihsel Evrim: 1940'lardan Bugüne.....	18
3.1. Sibernetik, Soğuk Savaş ve hesaplama.....	18
3.2. Klasik AI ve expert system dönemi.....	18
3.3. AI winters ve neden çöktüler.....	18
3.4. Büyük veri, GPU ve derin öğrenme	18
3.5. Transformer devrimi.....	18
3.6. LLM çağı ve emergence olgusu	19
Bölüm Sonu Değerlendirmesi	19
BÖLÜM II - SAYIM REJİMLERİ VE METODOLOJİ.....	20
4. Yapay Zekâyı Saymanın Problemi.....	20
4.1. “Kaç yapay zekâ var?” sorusunun sınırları.....	20
4.2. Model sayımı ile etki sayımı arasındaki fark.....	20
4.3. Akademik, ticarî, askerî ve düzenleyici sayım farkları	20
4.4. Bu rapor neyi sayar, neyi saymaz?	20
5. Bu Raporun Envanterleme Mantığı.....	21
5.1. Model.....	21
5.2. Platform	21

5.3. Ürün.....	21
5.4. Ajan	21
5.5. İşlemci ve altyapı.....	22
5.6. Kullanım bağlamı	22
5.7. Envanterden risk puanlamasına geçiş.....	22
6. Veri Kaynakları ve Güvenilirlik.....	24
6.1. Kaynak hiyerarşisi	24
6.2. Kamu kaynakları ve standart metinler	24
6.3. Kapalı, sınırlı ve asimetrik bilgiler	24
6.4. Açık kaynak bulgular ve sektör verileri.....	25
6.5. Kör noktalar ve bilinçli boşluklar	25
6.6. Hukuki savunulabilirlik ve metin içi atıf disiplini	25
Yanlış Karar Riski Kutusu – Bölüm II	25
Bölüm Sonu Değerlendirmesi	25
BÖLÜM III - KÜRESEL YAPAY ZEKÂ EKOSİSTEMİ	26
7. Küresel Genel Manzara	26
7.1. Milyonlarca model, onlarca gerçek güç merkezi.....	26
7.2. Frontier kavramının yeniden tanımı	26
7.3. Açık/kapalı model ikiliğinin ötesi	26
8. Bölgesel Güç Haritaları	28
8.1. ABD: merkezileşmiş inovasyon	28
8.2. Çin: devlet-özel bütünleşmesi	28
8.3. Avrupa: norm üreticisi, teknoloji ithalatçısı	29
8.4. İkincil düğümler: Kanada, Körfez, Asya-Pasifik ve İsrail.....	30
8.5. Küresel Güney ve veri sömürgeleşmesi	30
Bölüm Sonu Değerlendirmesi	31
BÖLÜM IV – TEKNİK TAM-YIĞIN ANALİZ	32
9. Model Türleri	32
9.1. Büyük Dil Modelleri.....	32
9.2. Multimodal modeller	32
9.3. Görsel, ses ve video üretim sistemleri	32
9.4. Retrieval ve ajan mimarileri	32
9.5. Savunma ve kamuya özel modeller	32
10. Hesaplama, Çip ve Enerji Gerçeği	33
10.1. GPU, TPU, NPU, LPU ve wafer-scale ayrımı.....	34
10.2. Nvidia hegemonyası ve rakip hızlandırıcı aileleri	34

10.3. Çin'in alternatif yolu ve sistem düzeyi yanıtı.....	34
10.4. Enerji, su ve iklim maliyeti.....	34
Bölüm Sonu Değerlendirmesi	36
BÖLÜM V – VERİ, GÜÇ VE EGEMENLİK	37
21. yüzyılın petrolü bir benzetme değil, bir zayıflatma.....	37
11. Veri Türleri.....	37
11.1. Kamuya açık veriler.....	37
11.2. Lisanslı veriler	38
11.3. Kullanıcı girdileri	38
11.4. Sentetik veri.....	38
12. Veri Akışı ve Egemenlik	39
12.1. Sınır ötesi veri hareketi.....	39
12.2. Ulusal mevzuatlar ve fiilî durum	39
Bölüm Sonu Değerlendirmesi	41
BÖLÜM VI – GÜVENİLİRLİK, HATA VE SİSTEMİK RİSK	42
13. Halüsinasyon ve Doğruluk Sorunu.....	42
13.1. Confabulation, bağlam kaybı ve sahte kesinlik	42
13.2. Ölçüm sorunu: doğruluk tek başına yeterli değildir	42
14. Otomatikleşmiş Yanlış Kararlar	43
14.1. İnsan gözetiminin biçimselleşmesi	44
14.2. Yüksek etkili alanlarda bağlam uyumsuzluğu	44
15. Kurumsal ve Kamusal Hata Maliyetleri	45
15.1. Mahremiyet, güvenlik ve dolandırıcılık baskısı	45
15.2. İtibar, meşruiyet ve kamu güveni	45
15.3. Yönetim sonucu: doğruluk rejimi kurmak.....	46
Bölüm Sonu Değerlendirmesi	46
BÖLÜM VII – MANİPÜLASYON, PROPAGANDA, DENORMASYON VE SOĞUK KAOS..	47
16. Deepfake Ekonomisi	47
Yanlış Karar Riski Kutusu.....	48
17. Algı Yönetimi ve Bilişsel Harp	48
Yanlış Karar Riski Kutusu.....	49
18. Seçimler, Medya ve Toplumsal Güven	49
Yanlış Karar Riski Kutusu.....	50
19. Sansür, Model Hizalaması ve Görünmez İdeoloji.....	50
Yanlış Karar Riski Kutusu.....	51
20. Soğuk Kaos, DeNormasyon ve Kurumsal Enformasyon Erozyonu.....	51

20.1. Kavramların Sahadaki Görünümü: Soğuk Kaos ve DeNormasyon Nasıl Çalışır?	51
20.2. Mikro Vaka I: New Hampshire Yapay Sesli Robocall Olayı	51
20.3. Mikro Vaka II: Hong Kong’da Deepfake Video Konferans Dolandırıcılığı	52
20.4. Makro Süreç I: Seçimlerden Bilgi Bütünlüğüne Geçiş	52
20.5. Makro Süreç II: İçerik Kökeni, Provenans ve Kurumsal Savunma Hattı.....	52
20.6. Soğuk Kaos ile DeNormasyon Arasındaki İşlevsel Fark.....	52
Yanlış Karar Riski Kutusu – Bölüm VII / 20. Başlık	53
Bölüm Sonu Değerlendirmesi	53
BÖLÜM VIII – DEVLET, SAVUNMA VE GELECEK SAVAŞLARI	54
21. Askerî AI Entegrasyonu	54
22. İstihbarat ve Karar Destek Sistemleri.....	55
23. Otonom Silah Tartışmaları	56
Bölüm Sonu Değerlendirmesi	58
BÖLÜM IX – HUKUK, ETİK VE YÖNETİŞİM	59
24. Küresel düzenleme girişimleri.....	59
25. Ulusal yaklaşımlar	60
26. Etik ilkeler ve fiilî pratik arasındaki uçurum.....	61
Bölüm Sonu Değerlendirmesi	63
BÖLÜM X – VAKA ANALİZLERİ	64
27. Şirket ve sistem bazlı vaka seti.....	64
27.1. Veri sızıntısı, güvenlik ve manipülasyon vakaları.....	66
Bölüm Sonu Değerlendirmesi	68
BÖLÜM XI – SENARYOLAR, RİSK MATRİSLERİ VE STRATEJİK UYARI	69
28. 2030–2040 senaryoları	69
28.1. Senaryo I: Yönetilen Hız, Kırılgan Denge	69
28.2. Senaryo II: Platform Egemenliği ve Sessiz Bağımlılık	69
28.3. Senaryo III: Sentetik Belirsizlik ve Kurumsal Felç	69
28.4. Senaryo IV: Enerji ve Hesaplama Darboğazının Jeopolitikleşmesi	69
28.5. Senaryoların Ortak Sonucu.....	69
29. Devletler için risk matrisi	71
29.1. Risk I: Dış Model ve Platform Bağımlılığı.....	71
29.2. Risk II: Sentetik Medya Kaynaklı Teyit Çöküşü.....	71
29.3. Risk III: Ajanik Sistemlerde Yetki Aşımı ve Zincirleme Hata.....	71
29.4. Risk IV: Kritik Altyapı, Enerji ve Hesaplama Darboğazı	71
29.5. Risk V: Kamu Hizmetlerinde Hak İhlali ve Meşruiyet Aşınması	72
29.6. Risk VI: Düzenleyici Gecikme ve Uyumda Kâğıt Üstü Başarı.....	72

29.7. Risk VII: Savunma ve İstihbaratta Denetimin Biçimselleşmesi.....	72
30. Kurumlar için hayatta kalma rehberi.....	73
30.1. İlk 90 Gün: Kurumun Yapması Gerekenler.....	73
30.2. Kurumun Asla Yapmaması Gerekenler.....	73
30.3. Zorunlu Kurumsal Kontrol Listesi.....	73
30.4. Yönetim Kuruluna Sorulacak 10 Soru.....	74
30.5. Kurum Tiplerine Göre Kısa Uygulama Notu	74
Yanlış Karar Riski Kutusu – Bölüm XI / 30. Başlık	74
Bölüm Sonu Değerlendirmesi	74
BÖLÜM XII – TÜRKİYE’NİN YAPAY ZEKÂ EŞİĞİ: KAPASİTE, BAĞIMLILIK, EGEMENLİK VE SOĞUK KAOS RİSKİ	75
31. Türkiye’nin Yapay Zekâ Konumlanması	75
32. Ulusal Yapay Zekâ Stratejisi ve Politika Mimarisinin Değerlendirilmesi	77
33. Türkçe Büyük Dil Modeli ve Dil Egemenliği Meselesi	77
34. Veri Egemenliği, KVKK ve Kamu Verisi.....	78
35. Savunma Sanayii, Güvenlik ve Kritik Altyapı	79
36. Türkiye İçin Soğuk Kaos ve DeNormasyon Riskleri	79
37. Türkiye İçin Yapay Zekâ Risk Sınıfları ve Denetim Öncelikleri	80
38. Türkiye İçin Stratejik Yol Haritası.....	81
Yanlış Karar Riski Kutusu – Bölüm XII	82
Bölüm Sonu Değerlendirmesi	82
BÖLÜM XIII – SONUÇ	83
GENEL KAYNAKÇA.....	88
EKLER.....	97

Telif, Hukuki ve Stratejik Sorumluluk Reddi

Bu rapor akademik, kurumsal ve stratejik değerlendirme amacıyla hazırlanmıştır. Metin; teknik, hukuki ve jeopolitik boyutları birlikte ele alan analitik bir çalışma niteliğindedir. Nihai karar, uygulama ve kurumsal politika süreçleri için tek başına yeterli hukuki görüş yerine geçmez.

Raporda kullanılan tüm değerlendirmeler, belirtilen kaynaklara dayalı yorumlayıcı bir çerçevede içinde sunulacaktır. Yapay zekâ alanı hızlı biçimde değiştiği için ürünler, düzenlemeler, model kapasiteleri ve şirket beyanları zaman içinde güncellenebilir. Bu nedenle nihai sürümde her veri, metin içi atıf ve kaynakça eşleştirmesi üzerinden ayrıca teyit edilecektir.

Bu çalışma, yapay zekâyı yalnızca teknik bir araç olarak değil; veri, hesaplama gücü, enerji, propaganda, güvenlik, hukuk ve egemenlik boyutlarıyla birlikte ele alır. Bu yaklaşım, raporun normatif ve stratejik iddialar içermesini kaçınılmaz hâle getirir. Ancak metinde yer alan hiçbir tespit, belgesiz suçlama veya doğrulanmamış isnat niteliğinde kullanılmayacaktır.

Yönetici Özeti

Bu bölümün amacı, raporun temel tezini ve stratejik uyarılarını kısa, yoğun ve karar verici odaklı biçimde ortaya koymaktır.

Bu raporun temel savı, yapay zekânın yalnızca yeni bir yazılım ailesi değil; veri, model, bulut, çip, enerji, platform ve yönetim katmanlarının birleştiği yeni bir güç mimarisi olduğudur. OECD'nin güncellenmiş yapay zekâ sistemi tanımı ile Stanford AI Index verileri birlikte okunduğunda, yapay zekâ alanında asıl belirleyici unsurun tek tek ürünler değil, yoğunlaşmış ekosistemler olduğu görülmektedir (OECD, 2024; Maslej et al., 2026).

Raporun özgün katkısı, bu ekosistemi yalnızca rekabet, yenilikçilik ve risk kavramlarıyla değil; Soğuk Kaos ve DeNormasyon çerçeveleriyle birlikte okumasıdır. Bu çerçevede Soğuk Kaos, yapay zekâ çağında belirsizliğin rastlantısal değil; hız, ölçek, tekrar ve kurumsal kırılabilirlik üzerinden yönetilen kalıcı bir düzensizlik rejimine dönüşmesini ifade etmektedir. DeNormasyon ise hata, manipülasyon, doğrulama zaafı ve norm aşınmasının istisna olmaktan çıkıp yeni normal hâline gelmesini anlatmaktadır (Kaya, 2026).

Bu nedenle yapay zekâ çağında sorun yalnızca hangi modelin daha güçlü olduğu değildir. Asıl mesele; hangi aktörlerin veriye erişebildiği, hangi şirketlerin hesaplama gücünü finanse edebildiği, hangi devletlerin platformları kendi güvenlik ve düzen hedefleriyle hizalayabildiği ve hangi toplumların bu süreçte ortak hakikat zeminini koruyabildiğidir. NIST, UNESCO ve OECD'nin çerçeveleri de güvenilir yapay zekâ meselesinin teknik performans kadar toplumsal güven, insan hakları ve kurumsal denetim sorunuyla ilgili olduğunu ortaya koymaktadır (NIST, 2023; UNESCO, 2021; OECD, 2026b).

Raporun nihai sonucu şudur: yapay zekâ çağında tarafsızlık iddiası giderek zayıflamaktadır. Çünkü model tasarımı, veri seçimi, hizalama tercihleri, altyapı yoğunlaşması ve düzenleyici boşluklar, yalnızca teknolojik çıktılar değil; kamusal aklı, kurumsal dayanıklılığı ve normatif düzeni de dönüştürmektedir. Bu nedenle yapay zekâ artık bir araçtan çok, hakikat, meşruiyet ve egemenlik ilişkilerini yeniden kuran tarihsel bir eşittir (IEA, 2025; Rid, 2020; Pomerantsev, 2019).

Kavramsal Çerçeve Notu: Soğuk Kaos ve DeNormasyon

Bu revizyon paketiyle raporun ana omurgası iki özgün kavram etrafında yeniden kurulmuştur. Soğuk Kaos, yapay zekâ çağında teknik hız, enformasyon yoğunluğu ve kurumsal kırılabilirliğin birleşmesiyle ortaya çıkan yönetilen belirsizlik rejimini; DeNormasyon ise hata, doğrulama aşınması, etik gevşeme ve norm erozyonunun sıradanlaşmasını ifade etmektedir. Bu nedenle rapor, yapay zekâyı yalnızca inovasyon veya risk başlıklarıyla değil; hakikat rejimi, kurum kapasitesi ve egemenlik ilişkileri bağlamında incelemektedir (Kaya, 2026; OECD, 2024).

Aşağıdaki şekil, raporun revize teorik çerçevesini özetlemektedir. Şekilde yapay zekâ ekosisteminin teknik çekirdeği ile DeNormasyon ve Soğuk Kaos katmanları arasındaki ilişki görselleştirilmiş; böylece raporun devamındaki bölüm yapısının hangi kavramsal mantıkla kurulacağı gösterilmiştir.

Şekil 0.1. Yapay zekâ çağında Soğuk Kaos ve DeNormasyon çerçevesi



Executive Summary

This report argues that artificial intelligence is no longer a narrow technical domain but a layered power architecture in which data, compute, energy, platforms, law, security and information control converge. In this architecture, systemic uncertainty is not merely a side effect of technological acceleration; it increasingly becomes a governable condition. The report conceptualizes this condition through the lenses of Cold Chaos and De-Normation, and treats AI as a driver of both strategic concentration and institutional fragility (Kaya, 2026; OECD, 2024; NIST, 2023).

The central finding is that today's AI order is shaped by simultaneous expansion and concentration. Adoption is spreading rapidly across firms and individuals, yet frontier capacity remains concentrated in a limited number of firms, cloud infrastructures, accelerator suppliers and geopolitical blocs. This asymmetry turns AI into more than an innovation race: it becomes a struggle over who defines trust, who controls scale and who can normalize ambiguity across digital and institutional environments (OECD, 2026; Maslej et al., 2026).

The report concludes that AI-era governance cannot be reduced to ethics slogans, model performance benchmarks or isolated safety audits. What is required is a strategic framework capable of tracking how technical opacity, synthetic media, automated error, platform dependency and state-corporate alignment together erode normative stability. In that sense, neutrality weakens as a meaningful position: the architecture of AI increasingly redistributes power while simultaneously destabilizing the informational foundations on which public trust depends (IEA, 2025; UNESCO, 2021; Kaya, 2026).

Bir Bakışta Yapay Zekâ Çağı

Aşağıdaki infografik, raporun temel stratejik göstergelerini tek sayfada özetlemektedir. Buradaki amaç, yapay zekâ çağının yalnızca teknik kapasite artışıyla değil; benimsenme, yoğunlaşma, enerji yükü, risk ve jeopolitik rekabet eksenleriyle birlikte okunması gerektiğini göstermektir (OECD, 2026; Maslej et al., 2026; IEA, 2025; Hugging Face, 2026).

Şekil 0.2. Bir bakışta yapay zekâ çağı: yoğunlaşma, benimsenme, risk ve enerji baskısı



Kritik Bulgular ve Kırmızı Bayraklar

Aşağıdaki tablo, raporun ön sayfalar düzeyinde öne çıkan en kritik bulgularını ve bunlara karşılık gelen kırmızı bayrakları özetlemektedir. Tablo, karar alıcı okuması için hazırlanmıştır; teknik ayrıntıdan çok stratejik baskı alanlarını göstermeyi amaçlamaktadır (Maslej et al., 2026; OECD, 2026; IEA, 2025).

Tablo 0.1. Yapay zekâ çağında kritik bulgular ve kırmızı bayraklar

Alan	Temel bulgu	Kırmızı bayrak	Stratejik anlam
Benimsenme	AI kullanımı hızla ana akıma giriyor	Kurumsal bağımlılık artıyor	Teknoloji seçimi artık yönetim kararıdır
Yoğunlaşma	Frontier kapasite birkaç aktörde birikiyor	Pazar ve norm gücü merkezileşiyor	Egemenlik tartışması teknikleşiyor
Şeffaflık	En güçlü modeller daha kapalı hâle geliyor	Denetim ve hesap verebilirlik zayıflıyor	Güven yerine kurumsal körlük büyüyor
Sentetik medya	Deepfake ve sentetik içerik hızla yayılıyor	Hakikat zemini parçalanıyor	DeNormasyon hızlanıyor
Enerji	Veri merkezleri ciddi elektrik talebi üretiyor	Altyapı kırılganlığı büyüyor	AI aynı zamanda enerji-jeopolitik meseledir
Devlet ilişkisi	Kamu ve savunma entegrasyonu genişliyor	Sivil-askeri sınırlar bulanıklaşıyor	Soğuk Kaos devlet ölçeğine taşınıyor

Kaynak: Maslej vd. (2026); OECD (2026); IEA (2025); UNESCO (2021); Kaya (2026).

Karar Vericiler İçin Stratejik Uyarılar ve Zorunlu Politika Eşikleri

Bu bölümün amacı, raporun geri kalanında ayrıntılı biçimde temellendirilen riskleri karar verici diliyle özetlemek ve hangi hataların artık tolere edilemez olduğunu açık biçimde göstermektir.

Bu çalışma, yapay zekâyı yalnızca verimlilik artışı sağlayan yeni bir teknoloji ailesi olarak değil; veri, hesaplama, enerji, platform, hukuk, güvenlik ve enformasyon düzenini birlikte dönüştüren çok katmanlı bir güç mimarisi olarak ele almaktadır. Bu nedenle yapay zekâ alanındaki en büyük hata, meseleyi yalnızca yazılım tercihi gibi okumaktır. Asıl risk, teknik tercihler adı altında egemenlik, denetim, karar kapasitesi ve ortak hakikat zemini üzerinde geri dönmesi zor bağımlılıklar üretmektir (OECD, 2024; NIST, 2023; Maslej ve diğerleri, 2026).

Bu raporun ilk yasağı şudur: kurumsal yapay zekâ tercihini sıradan bir satın alma kararı sanmak. Yapay zekâ sistemi seçmek, artık yalnızca araç seçmek değildir; veri akışını, model bağımlılığını, uyum yükünü, denetim maliyetini ve uzun vadeli kurumsal hafızayı da birlikte seçmektir. Yanlış model veya platform tercihi, kısa vadede düşük maliyet gibi görünebilir; fakat orta vadede kurumun kendi verisi üzerinde denetim kaybına, dış altyapıya aşırı bağımlılığa ve stratejik körlüğe yol açabilir.

İkinci yasak, çok sayıda modelin varlığını gerçek rekabet sanmaktır. Görünürde milyonlarca model bulunması, gücün dağıldığı anlamına gelmez. Tam tersine, gerçek etki sınırlı sayıdaki frontier laboratuvarı, bulut sağlayıcısı, hızlandırıcı üreticisi ve yüksek erişimli ürün ailesi etrafında yoğunlaşmaktadır. Sayıyı rekabet, çeşitliliği ise bağımsızlık zannetmek en tehlikeli yanılsamalardan biridir (Hugging Face, 2026; Maslej ve diğerleri, 2026).

Üçüncü yasak, yapay zekâ güvenliğini yalnızca model doğruluğu sanmaktır. Kurumsal felaketler çoğu zaman yalnızca yanlış cevaplardan doğmaz. Asıl kırılma, hatalı çıktının süreç içine denetimsiz girmesi, insan gözetiminin sembolikleşmesi ve yanlış kararın otomasyon sayesinde yüksek hızla yayılmasıyla ortaya çıkar. Bu nedenle güvenlik, yalnızca benchmark başarımları değil; görev sınırı, itiraz mekanizması, kayıt tutma, insan denetimi ve telafi kapasitesi meselesidir (NIST, 2023; UNESCO, 2021).

Dördüncü yasak, sentetik medya riskini yalnızca iletişim sorunu sanmaktır. Deepfake, sahte ses, sentetik görüntü ve otomatik içerik çoğalması yalnızca medya ekosistemini kirlilemez; kriz anında teyit süresini uzatır, karar sürecini felç eder, kurumlar arası güveni aşındırır ve rakip aktörlere operasyonel avantaj sağlar. Soğuk Kaos tam da burada doğar: bilgi yokluğundan değil, fazla, hızlı, çelişkili ve kısmen sentetik içerik akışının karar zemininin altını oymasından (Rid, 2020; Pomerantsev, 2019).

Beşinci yasak, hizalama ve güvenlik filtrelerini tarafsız teknik işlemler gibi görmektir. Hangi içeriklerin riskli sayıldığı, hangi soruların reddedildiği, hangi bağlamların görünmez biçimde bastırıldığı ve hangi cevap biçimlerinin teşvik edildiği yalnızca mühendislik tercihi değildir. Bunlar aynı zamanda normatif tercihlerdir. Tarafsızlık iddiası çoğu zaman görünmez ideolojinin en etkili taşıyıcısıdır.

Altıncı yasak, enerji ve altyapı meselesini teknik arka plan ayrıntısı saymaktır. Veri merkezi elektriği, hızlandırıcı çip erişimi, su kullanımı, soğutma kapasitesi ve ağ omurgası artık yapay zekâ rekabetinin dışsal koşulları değil, bizzat çekirdek belirleyicileridir. Yapay zekâ aynı zamanda enerji-jeopolitik ve kritik altyapı meselesidir (IEA, 2025; Maslej ve diğerleri, 2026).

Yedinci yasak, kamu kurumlarında 'önce kullanalım, sonra düzenleriz' mantığıyla ilerlemektir. Kamu sektörü için yapay zekâda en pahalı hata, denetim mimarisinden önce kullanım genişletmektir. Çünkü kamu gücü ile otomatikleşmiş hata birleştiğinde, sıradan kurumsal

verimsizlik değil; hak ihlali, ayrımcılık, yanlış sınıflandırma ve itiraz hakkının aşınması ortaya çıkar.

Sekizinci yasak, savunma ve istihbarat alanında insan denetimini biçimsel onaya indirgemektir. Human-in-the-loop söylemi, eğer karar temposu, veri yoğunluğu ve operasyon baskısı içinde gerçek denetim kapasitesi üretmiyorsa, çoğu zaman yalnızca hukuki rahatlatma işlevi görür. Savunma yapay zekâsında mesele insanın sistemde bulunması değil; insanın gerçekten sistemi durdurabilmesidir.

Dokuzuncu yasak, Küresel Güney'i yalnızca yeni pazar olarak görmektir. Yapay zekâ çağında veri üretimi ile değer birikimi aynı coğrafyada gerçekleşmemektedir. Veriyi üreten, fakat modeli dışarıdan satın alan toplumlar dijital çağın ham madde tedarikçisine dönüşür. Bu nedenle dil, kültür, kamusal veri ve kurumsal arşivler yalnızca içerik havuzu değil; egemenlik altyapısıdır (UNCTAD, 2021).

Onuncu yasak, yönetişimi etik ilkeler listesiyle sınırlamaktır. Etik beyanlar, olay bildirimi, denetim zinciri, yaptırım, itiraz hakkı ve telafi mekanizması yoksa çoğu zaman kurumsal meşruiyet makyajına dönüşür.

On birinci yasak, kurumsal hafıza ve karar süreçlerini kapalı dış servislerin içine gömmektir. Kısa vadeli verimlilik kazanımları uğruna belge işleme, bilgi çağırma, karar hazırlama ve uzmanlık destek süreçlerini tümüyle dış platform mantıklarına teslim eden kurumlar, bir süre sonra yalnızca araç bağımlısı değil; muhakeme bağımlısı hâline gelir.

On ikinci ve son yasak, yapay zekâ çağında tarafsız kalınabileceğine inanmaktır. Bu çağda tarafsızlık çoğu zaman pozisyon almamak değil, başkalarının seçtiği veri, model, altyapı ve norm setine sessizce tabi olmak anlamına gelir. Esas soru, yapay zekâ kullanıp kullanmamak değil; hangi mimarinin içinde, kimin kural setiyle, hangi denetim düzeyinde ve hangi stratejik bedeli göze alarak yapay zekâ kullanılacağıdır (Kaya, 2026; OECD, 2024; IEA, 2025).

BÖLÜM I - KAVRAMSAL VE TARİHSEL ZEMİN

Bu bölümün amacı, yapay zekâ tartışmasını teknik popülerlikten çıkarıp tarihsel, kuramsal ve felsefi bir zemine oturtmaktır. Böylece raporun geri kalanında kullanılacak kavram seti ve analitik dil açıklığa kavuşturulacaktır.

1. Bu Rapor Neden Yazıldı?

Bu rapor, yapay zekâyı yalnızca teknik bir yenilik veya verimlilik aracı olarak değil; veri, hesaplama gücü, enerji, hukuk, güvenlik, propaganda, norm üretimi ve egemenlik boyutlarını aynı anda etkileyen çok katmanlı bir güç sistemi olarak ele almaktadır. OECD'nin güncellenmiş yapay zekâ sistemi tanımı bu çerçeveyi teknik düzeyde, Stanford AI Index ise yoğunlaşma ve ölçek düzeyinde desteklemektedir (OECD, 2024; Maslej et al., 2026).

Ancak bu raporun özgün iddiası, yapay zekânın yalnızca teknik ve ekonomik bir dönüşüm üretmediğidir. Yapay zekâ aynı zamanda, bilgi alanında bulanıklığı, kurumsal karar süreçlerinde kırılganlığı ve doğrulanabilir ortak gerçeklik zemininde aşınmayı hızlandırmaktadır. Bu nedenle rapor, yapay zekâyı Soğuk Kaos ve DeNormasyon kavramlarıyla birlikte değerlendirmektedir (Kaya, 2026).

Soğuk Kaos, görünürde düzenli fakat özünde belirsizlik üreten yeni dijital ortamı; DeNormasyon ise hukuki, etik ve epistemik eşiklerin gevşeyerek bozulmanın istisna olmaktan çıkıp sıradanlaşmasını ifade etmektedir. Yapay zekâ çağında deepfake, sentetik medya, otomatikleştirilmiş hata, görünmez hizalama tercihleri ve platform yoğunlaşması bu iki süreci eşzamanlı biçimde beslemektedir (NIST, 2023; UNESCO, 2021; OECD, 2026a).

2. Yapay Zekânın Ontolojisi

2.1. Zekâ nedir?

Yapay zekâ tartışmasının en sorunlu yanı, zekâ kavramının kendisinin tartışmalı olmasıdır. İnsan zekâsı, bağlam kurma, niyet atfetme, anlam üretme, belirsizlik içinde karar verme ve toplumsal sembollerle çalışma gibi çok katmanlı özellikler taşır. Yapay zekâ ise çoğu durumda bu süreçlerin kendisini değil, bu süreçlere ait izleri ve örüntüleri işler. Bu nedenle yapay zekânın zekâyı bire bir yeniden ürettiğini değil, belirli bilişsel çıktıları istatistiksel ve hesaplamalı yollarla taklit ettiğini söylemek daha isabetlidir (Simon, 1996; Dennett, 1991).

2.2. Algoritma, istatistik, temsil ve anlam ayrımı

Klasik bilgisayar bilimi açısından algoritma, sonlu adımlarla tanımlanabilen bir işlem düzenidir; modern yapay zekâ ise bu işlem düzenini büyük veri kümeleri ve parametrik optimizasyonla genişletir. Ancak istatistiksel başarı ile anlamsal kavrayış aynı şey değildir. Bu ayrım, güncel büyük dil modellerinin yüksek akıcılık üretmesine rağmen hâlâ bağlam hataları, halüsinasyon ve sahte kesinlik üretebilmesini açıklar (LeCun et al., 2015; Brown et al., 2020).

2.3. Bilinç, farkındalık ve irade tartışmaları

Yapay zekâya ilişkin kamusal tartışmanın önemli bölümü, bilinç ve farkındalık sorularına odaklanmaktadır. Oysa güncel sistemlerin büyük çoğunluğu, bilinç sahibi öznel yapılar olmaktan çok, karmaşık örüntü tanıma ve üretim sistemleridir. Bu noktada Turing'in davranışsal ölçütü ile bilinç kuramları arasındaki fark önemlidir: bir sistemin ikna edici davranış üretmesi, onun öznel deneyime sahip olduğunu göstermez (Turing, 1950; Dennett, 1991).

2.4. Simülasyon mu, yeti mi?

Bu soru, raporun tamamı için kurucu önemdedir. Yapay zekâ sistemleri çoğu zaman gerçek anlama yetisinden ziyade ikna edici simülasyonlar üretir. Ancak bu simülasyonların toplumsal etkisi

gerçekdir. Dolayısıyla siyaset, hukuk ve güvenlik açısından belirleyici olan şey, sistemin gerçekten bilinç sahibi olup olmaması değil; insanlar, kurumlar ve pazarlar üzerinde gerçek etkiler üretebilmesidir (Bubeck et al., 2023; NIST, 2023).

Tablo 1.1. Yapay zekâ düşüncesinin başlıca tarihsel kırılma anları

Dönem	Ana kırılma	Temel özellik	Kaynak örnekleri
1940'lar-1950'ler	Sibernetik ve erken hesaplama	Makine, kontrol ve iletişim ilişkisi öne çıktı.	Wiener (1948); Turing (1950)
1956-1970'ler	Dartmouth ve sembolik AI	Zekâ, kurallı sembolik işlemlerle modellenmeye çalışıldı.	McCarthy et al. (1956)
1980'ler	Expert systems dönemi	Alan bilgisine dayalı kurallar ve uzman sistemler yaygınlaştı.	Simon (1996)
1990'lar-2010'lar	Makine öğrenmesi ve derin öğrenme	Örüntü tanıma, veri yoğun öğrenme ve optimizasyon öne çıktı.	LeCun et al. (2015)
2017 sonrası	Transformer ve LLM çağı	Ölçek, dikkat mekanizması ve üretken modeller belirleyici oldu.	Vaswani et al. (2017); Brown et al. (2020)

Kaynak: Wiener (1948); Turing (1950); McCarthy ve diğerleri (1956); LeCun ve diğerleri (2015); Vaswani ve diğerleri (2017); Brown ve diğerleri (2020); yazarın sentezi.

3. Tarihsel Evrim: 1940'lardan Bugüne

3.1. Sibernetik, Soğuk Savaş ve hesaplama

Modern yapay zekânın kökleri, bilgi işlemden önce kontrol ve iletişim teorisine uzanır. Sibernetik yaklaşım, canlı ve makine sistemleri arasındaki geri besleme, kontrol ve iletişim ilişkisini ortak bir çerçevede ele almıştır. Bu yaklaşım, zekânın maddi substrattan çok işlevsel örgütlenme ile ilişkili olabileceği fikrinin önünü açmıştır (Wiener, 1948).

3.2. Klasik AI ve expert system dönemi

1956 Dartmouth önerisi, yapay zekâyı bağımsız bir araştırma alanı olarak kurumsallaştırdı. Bu dönemde ağırlıklı varsayım, düşünmenin sembolik kurallar ve mantıksal işlemler üzerinden modellenebileceği yönündeydi. 1980'lerde expert system yaklaşımı bu çizginin uygulamalı versiyonunu oluşturdu (McCarthy et al., 1956; Simon, 1996).

3.3. AI winters ve neden çöktüler

İlk büyük yapay zekâ dalgaları, beklenti ile kapasite arasındaki makas açıldığında zayıfladı. Yetersiz hesaplama gücü, dar alan bağımlılığı ve kırılğan genelleme kapasitesi, hem sembolik AI'nin hem de erken makine öğrenmesi yaklaşımlarının sınırlarını görünür kıldı. Bu dönemler, yapay zekâ tarihinin sadece ilerleme hikâyesi olmadığını gösterir.

3.4. Büyük veri, GPU ve derin öğrenme

Asıl dönüşüm, veri hacmi, paralel hesaplama ve öğrenilebilir çok katmanlı ağların birlikte olgunlaşmasıyla yaşandı. Derin öğrenme, temsilin insan eliyle tasarlandığı sistemlerden, veriden öğrenildiği sistemlere geçişte belirleyici oldu (LeCun et al., 2015).

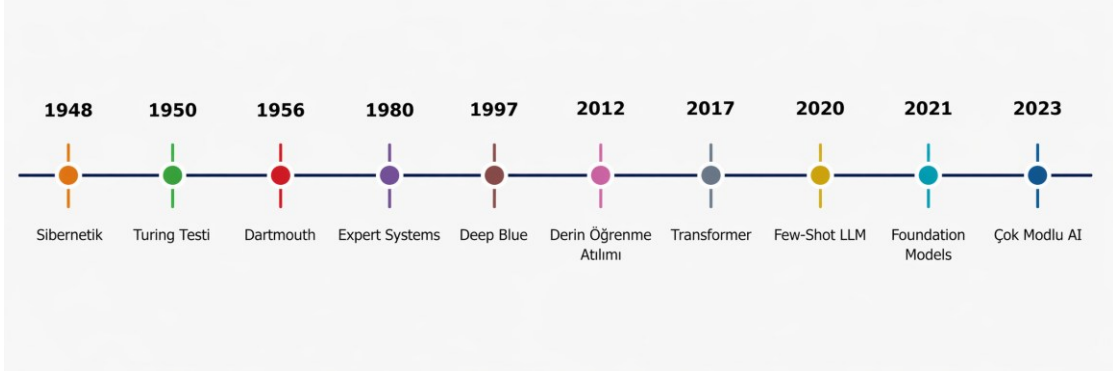
3.5. Transformer devrimi

Dikkat mekanizmasına dayanan transformer mimarisi, uzun bağımlılıkların daha etkili biçimde işlenmesini ve model ölçeğinin daha verimli büyümesini mümkün kıldı. Bu kırılma, güncel büyük dil modellerinin teknik omurgasını oluşturdu (Vaswani et al., 2017).

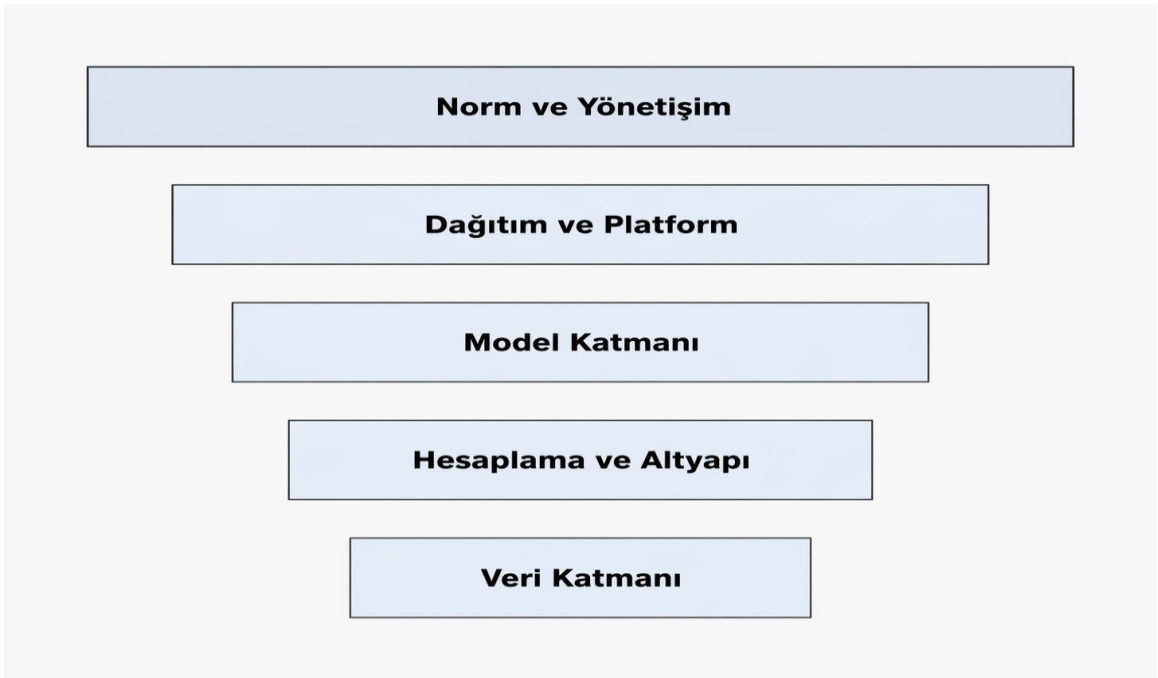
3.6. LLM çağı ve emergence olgusu

Ölçek büyüdükçe, bazı yeteneklerin yalnızca parametre artışıyla değil, veri ve görev çeşitliliğiyle birlikte ortaya çıktığı ileri sürüldü. Foundation models literatürü ve daha sonra gelen AGI tartışmaları, bu yüzden yalnızca teknik bir sıçrama değil; epistemik ve toplumsal bir kırılma olarak okunmalıdır (Bommasani et al., 2021; Bubeck et al., 2023).

Şekil 1.1. Yapay zekânın tarihsel evrimi: sibernetikten çok modlu sistemlere



Şekil 1.2. Yapay zekâ ekosisteminin katmanlı güç yapısı



Şema 1.1. Yapay zekâyı araç olmaktan çıkaran ekosistem mantığı

Bölüm Sonu Değerlendirmesi

Bu bölümün temel sonucu açıktır: yapay zekâ meselesi, yalnızca yeni bir mühendislik tekniği olarak ele alınamayacak kadar kapsamlıdır. Kavramsal düzeyde zekâ, anlam ve temsil tartışmaları; tarihsel düzeyde ise sibernetikten transformer mimarisine uzanan kırılmalar, güncel yapay zekâ düzeninin uzun bir birikimin ürünü olduğunu göstermektedir. Raporun sonraki bölümleri bu tarihsel ve kuramsal zemini, küresel envanter, teknik mimari, veri egemenliği, manipülasyon ve yönetim eksenlerinde ayırtılandıracaktır.

BÖLÜM II - SAYIM REJİMLERİ VE METODOLOJİ

Bu bölümün amacı, yapay zekâ ekosistemine ilişkin niceliksel bolluk görüntüsünün neden tek başına açıklayıcı olmadığını göstermek, raporun envanterleme mantığını kurmak ve kullanılan veri rejiminin sınırlarını açık biçimde tanımlamaktır.

Güncel yapay zekâ tartışmasında en sık yapılan hata, görünür model sayısını gerçek güç dağılımı zannetmektir. Oysa kamusal platformlarda milyonlarca modelin listelenmesi, stratejik kapasitenin gerçekten dağılmış olduğu anlamına gelmez. Tam tersine, bu raporun temel metodolojik savı şudur: yapay zekâ alanında belirleyici olan, görünür çoğulluk değil; hangi katmanda, hangi aktörlerde, hangi erişim eşiği üzerinden fiilî yoğunlaşmanın oluştuğudur (OECD, 2024; NIST, 2023; Maslej ve diğerleri, 2026).

Bu nedenle bu çalışma, 'kaç model var?' sorusunu başlangıç noktası olarak kabul etmekle birlikte, onu asla tek başına yeterli görmez. Çünkü yapay zekâ çağında sayı ile etki, açıklık ile denetlenebilirlik, benimsenme ile egemenlik ve yenilik ile kurumsal bağımlılık aynı şey değildir. Mesele yalnızca teknik nesneleri saymak değil; hangi model ailesinin hangi altyapı, veri rejimi, dağıtım yüzeyi ve yönetim boşluğu içinde güç ürettiğini anlamaktır.

4. Yapay Zekâyı Saymanın Problemi

Bu bölümün amacı, 'kaç yapay zekâ sistemi var?' sorusunun neden tek başına açıklayıcı olmadığını göstermek ve raporun sayım rejimlerini analitik olarak ayrıştırmaktır.

4.1. “Kaç yapay zekâ var?” sorusunun sınırları

Yapay zekâ alanında ilk bakışta görülen büyüklük, çoğu zaman teknik nesne bolluğundan ibarettir. Açık model depoları, ürün listeleri, uygulama marketleri, model kartları ve şirket duyuruları birlikte okunduğunda son derece geniş bir ekosistem görüntüsü ortaya çıkmaktadır. Ancak stratejik analiz için asıl mesele, bu varlıkların ne kadarının gerçek kullanım, kurumsal nüfuz, altyapı bağımlılığı ve norm üretme kapasitesi taşıdığıdır.

4.2. Model sayımı ile etki sayımı arasındaki fark

Bu raporun temel ayrımlarından biri, model sayımı ile etki sayımını birbirinden ayırmasıdır. Model sayımı kamusal olarak görülen teknik varlıkları ifade eder. Etki sayımı ise bir modelin veya ürünün gerçek dünyada ne ölçüde kurumsal iş akışına girdiğini, kullanıcı davranışını etkilediğini, karar sürecine bağlandığını ve başka sistemlerle birleşerek kalıcı bağımlılık ürettiğini sorgular.

4.3. Akademik, ticarî, askerî ve düzenleyici sayım farkları

Yapay zekâ farklı bağlamlarda farklı biçimlerde sayılır. Akademik bağlam araştırma yeniliği ve yöntemsel katkı üzerinden hareket eder. Ticarî bağlam kullanıcıya erişim, ürünleşme ve pazar gücüne bakar. Askerî bağlam görev kabiliyeti, ağ yalıtımı, sensör entegrasyonu ve sınıflandırılmış kapasiteyi esas alır. Düzenleyici bağlam ise risk sınıfı, hak etkisi, kullanım alanı ve zarar potansiyelini öne çıkarır.

4.4. Bu rapor neyi sayar, neyi saymaz?

Bu rapor; temel model aileleri, açık model ekosistemleri, yüksek erişimli son kullanıcı ürünleri, kurumsal platformlar, ajan mimarileri, hızlandırıcı altyapılar, veri merkezi kapasitesi, kamu ve savunma kullanım alanları ile düzenleyici çerçeveleri sayım evrenine dahil eder. Buna karşılık doğrulanamayan sızıntılar, teyitsiz teknik iddialar, yalnızca pazarlama metinlerine dayanan performans beyanları ve bağımsız teyidi mümkün olmayan karanlık altyapılar ana analiz dışında tutulur.

Tablo 2.1. Yapay zekâ sayım rejimlerinin karşılaştırılması

Karşılaştırma boyutu	OECD yaklaşımı	Hugging Face yaklaşımı	AI Index / Maslej vd. yaklaşımı	Analitik not
Temel sayım birimi	Kuruluşlar, politikalar ve ekonomik/kurumsal göstergeler.	Modeller, veri kümeleri, görevler, benchmarklar ve topluluklar.	Yatırımlar, yayınlar, patentler, modeller, performans, işgücü ve politikalar.	Sayım birimleri farklıdır; makro göstergeler, ekosistem varlıkları ve çok boyutlu trendler odak noktasıdır.
Kapsam	Ülkeler ve sektörler bazında makro düzey; benimseme ve etkiler.	Küresel açık kaynak ekosistemi ve topluluk etkinlikleri.	Küresel ölçekte çok boyutlu trendler ve gelişmeler.	Kapsam düzeyi ve içerik genişliği farklıdır; birlikte daha bütüncül görünüm sağlar.
Ölçüm mantığı	Anketler, idari veriler ve resmi istatistiklerle politika ve benimseme ölçümü.	Platform verileri ve topluluk katkılarına dayanan aktivite sayımı.	Birden çok veri kaynağını bütünleştirerek göstergelere dönüştürme.	Ölçüm mantıkları farklıdır; anket/istatistik, aktivite tabanlı sayım ve bütünleşik izleme birlikte tamamlar.
Güçlü yön	Politika analizi, uluslararası karşılaştırma ve güvenilirlik.	Şeffaflık, açıklık ve hızlı güncellenen ekosistem görünürlüğü.	Çok boyutlu ve bütünleşik izleme; trend ve etki analizi yeteneği.	Her rejim farklı güçlü yanlar sunar; farklı analiz ihtiyaçlarına hizmet eder.
Sınırlılık	Granüler teknik detayda sınırlı; hızlı değişimi yakalamakta gecikme olabilir.	Temsiliyet ve veri kalitesinde değişkenlik; politika bağlamı sınırlı.	Veri karşılaştırabilirliği ve metodolojik uyum zorlukları; bazı alanlarda tahmine dayalı.	Her rejimin metodolojik ve kapsamsal sınırlılıkları vardır; dikkatle yorumlanmalıdır.
Politika değeri	Kanıtı dayalı politika tasarımı, düzenleme ve etki izleme.	Açık inovasyonu destekleme, yetenek keşfi ve iş birliği teşviki.	Stratejik yönlendirme, rekabet gücü analizi ve politika önceliklendirmesi.	Politika yapımında farklı amaçlarla birlikte kullanılmalı; karar süreçlerini zenginleştirir.

Kaynak: OECD (2024); Hugging Face (2026); Maslej ve diğerleri (2026).

5. Bu Raporun Envanterleme Mantığı

Bu rapor, yapay zekâ ekosistemini tek satırlık şirket veya model listeleriyle değil; katmanlı bir envanter mantığıyla incelemektedir. Amaç, görünen ürünleri sıralamak değil; nerede güç biriktiğini, hangi bileşenlerin birbirine bağımlı olduğunu ve hangi katmanların sistemik kırılabilirlik ürettiğini göstermektir.

5.1. Model

Model, teknik çekirdeği ifade eder. Temel modeller, açık ağırlıklı modeller, çok kipli sistemler ve görev odaklı dar modeller bu katmanda değerlendirilir. Ancak bu rapor modeli tek başına nihai analiz birimi olarak görmez; çünkü model çoğu zaman etkisini bağlandığı platform, ürün ve veri akışı üzerinden üretir.

5.2. Platform

Platform, modelin dağıtıldığı, geliştiricilere, kurumlara ve son kullanıcıya açıldığı ara katmandır. Bulut servisleri, model merkezleri, API ekosistemleri, kurumsal yapay zekâ paketleri ve entegre geliştirici araçları bu düzeyde okunur.

5.3. Ürün

Ürün, son kullanıcıya görünen ve çoğu zaman yapay zekânın kendisi zannedilen katmandır. Sohbet sistemleri, ofis asistanları, kod yazım araçları, görsel üretim uygulamaları ve arama destek mekanizmaları bu düzeyde değerlendirilir.

5.4. Ajan

Ajan katmanı, araç çağırma, çok adımlı işlem yürütme, veri tabanı ile etkileşim kurma, görev planlama ve çıktıyı tekrar işleyerek geliştirme gibi kapasitelere sahip yapıları kapsar.

5.5. İşlemci ve altyapı

GPU, TPU, NPU, veri merkezi, ağ omurgası, elektrik ve soğutma altyapısı bu katmanda yer alır. Yapay zekâ rekabeti yalnızca yazılım rekabeti değildir; hesaplama kapasitesi, enerji erişimi ve yarı iletken tedarik zinciri olmadan sürdürülebilir güç kurulamaz.

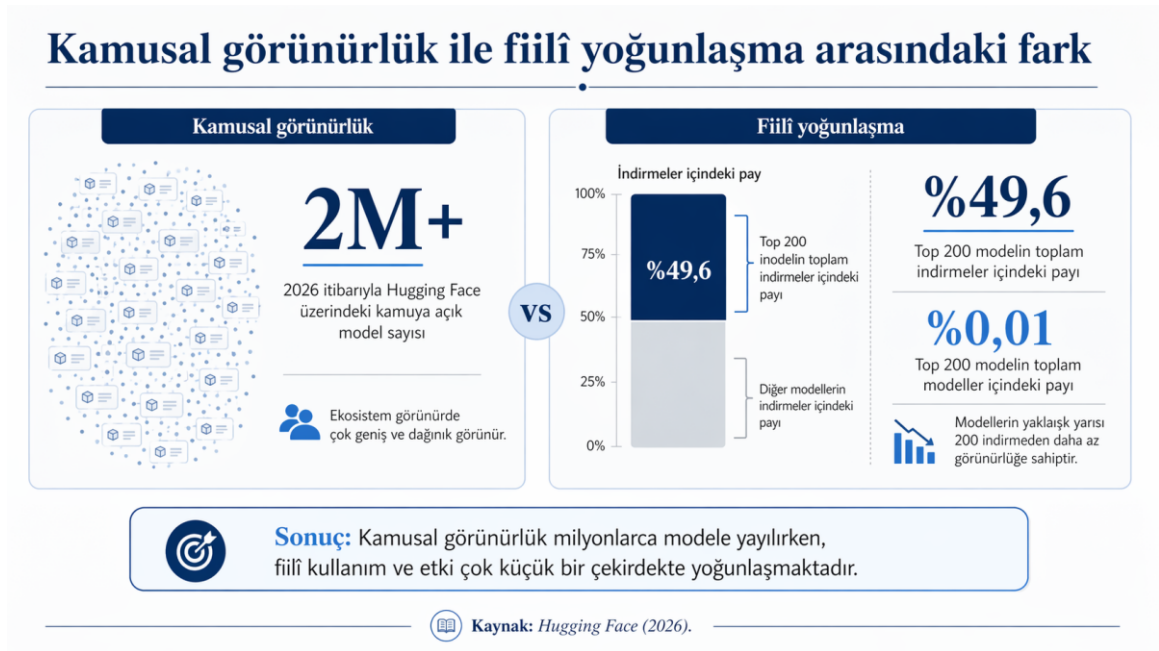
5.6. Kullanım bağlamı

Aynı model, farklı bağlamlarda bütünüyle farklı riskler üretebilir. Tüketici uygulamasında görece düşük etkili görünen bir hata, kamu hizmetinde hak ihlaline; savunma alanında operasyonel felce; sağlık veya finans alanında ağır maddi ve toplumsal zarara dönüşebilir.

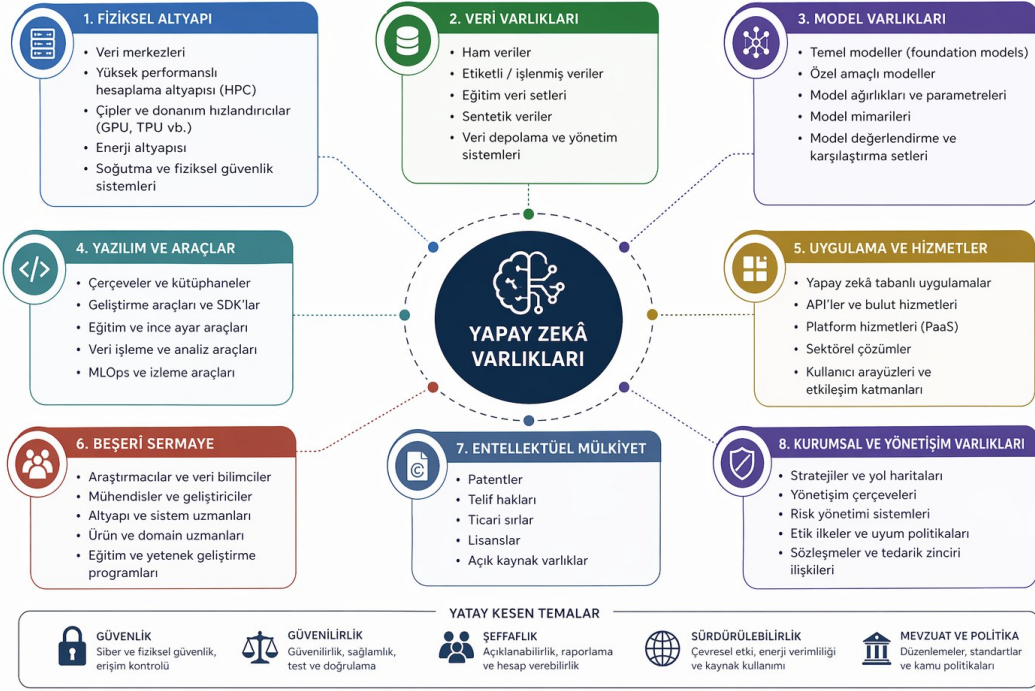
5.7. Envanterden risk puanlamasına geçiş

Bu raporun metodolojik akışı yalnızca varlık tespitinde bitmez. Envanterleme işlemi, doğrudan risk puanlaması ve stratejik anlamlandırma aşamasına bağlanır. Burada temel soru, bir unsurun yalnızca mevcut olup olmadığı değil; hangi ölçekte zarar üretebileceği, hangi kurumsal akışlara bağlandığı, hangi denetim rejimi altında çalıştığı ve bağımlılık yaratma derecesidir.

Grafik 2.1. Kamusal görünürlük ile fiili yoğunlaşma arasındaki fark

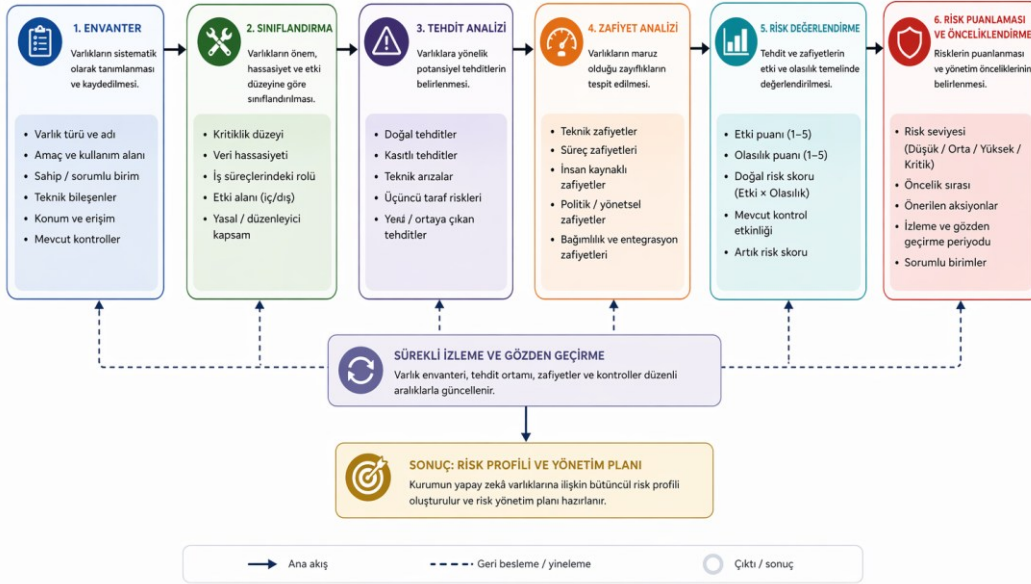


Harita 2.1. Bu raporun sayım evreni: yapay zekâ varlıklarının kavramsal haritası



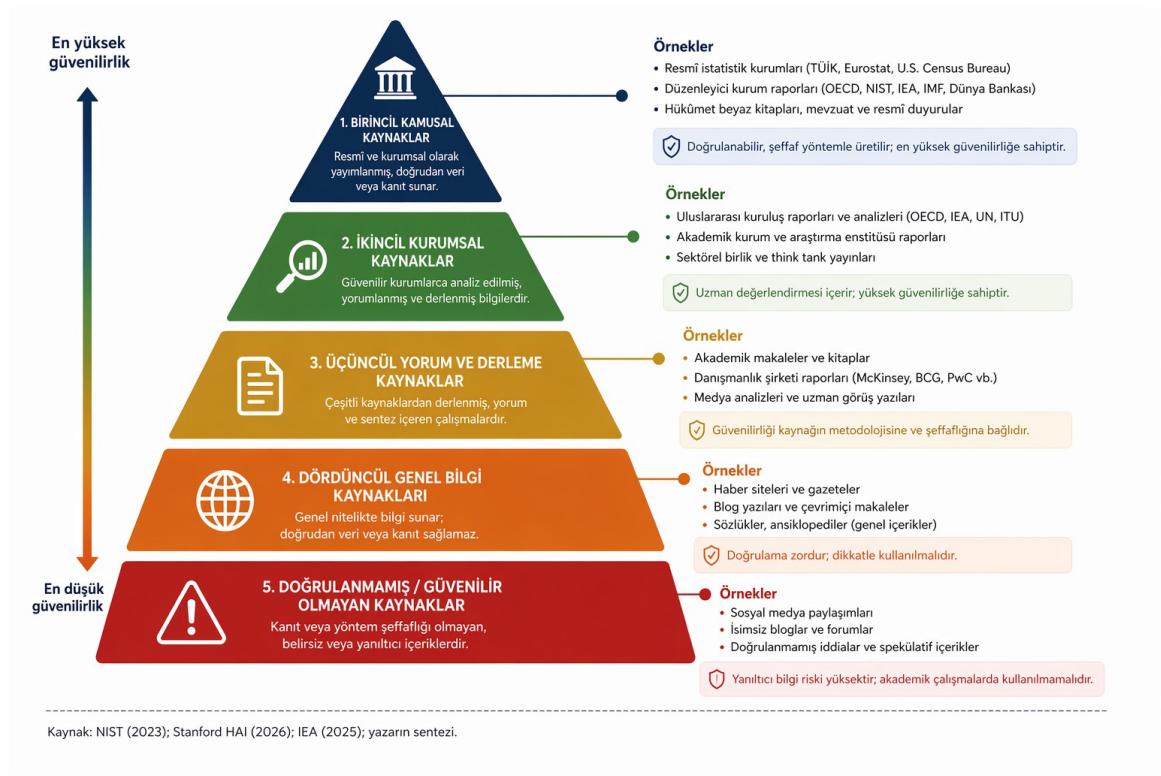
Kaynak: OECD (2024); NIST (2023); yazarın sentezi.

Şema 2.1. Envanterden risk puanlamasına uzanan metodolojik akış



Kaynak: NIST (2023); OECD (2024); yazarın sentezi.

Şekil 2.1. Kaynak güvenilirliği piramidi



6. Veri Kaynakları ve Güvenilirlik

Bu raporda kullanılan veri kaynakları, hiyerarşik bir güven rejimine göre sıralanmıştır. En üst düzeyde mevzuat metinleri, resmî kurum sayfaları, standart belgeleri, teknik model kartları ve doğrudan kurumsal raporlar yer almaktadır. Bunları OECD, NIST, Stanford HAI, IEA ve UNESCO gibi yüksek kurumsal güvenilirliğe sahip değerlendirmeler takip etmektedir.

6.1. Kaynak hiyerarşisi

Kaynak hiyerarşisinin en üstünde bağlayıcı ve birincil malzeme bulunur. Mevzuat, resmî politika belgeleri, doğrudan kurum açıklamaları, standartlar ve teknik özellik belgeleri bu düzeyde yer alır. İkinci düzeyde yüksek kurumsal güvenilirliğe sahip değerlendirmeler bulunur. Üçüncü düzeyde şirket teknik blogları ve sektör raporları gelir. Dördüncü düzey yalnızca destekleyici ve teyit gerektiren malzemeden oluşur.

6.2. Kamu kaynakları ve standart metinler

Kamuya açık resmî belgeler raporun omurgasını oluşturur. Bunun temel sebebi, yapay zekâ gibi yüksek hızla değişen ve yoğun pazarlama söylemi içeren bir alanda sağlam zeminin ancak resmî, teknik ve bağlayıcı malzemeyle kurulabilmesidir.

6.3. Kapalı, sınırlı ve asimetrik bilgiler

Yapay zekâ ekosisteminin önemli bölümü kamusal değildir. Savunma, istihbarat, kurumsal güvenlik, özel model eğitimi ve özel veri kümeleri alanlarında bilginin bir kısmı doğal olarak sınırlı, parçalı veya kasıtlı olarak kapalıdır. Bu nedenle rapor tamlık iddiası kurmaz; yalnızca kamuya açık, doğrulanabilir ve birden fazla kaynakla desteklenebilen bulgularla ilerler.

6.4. Açık kaynak bulgular ve sektör verileri

Açık kaynak istihbaratı, sektör değerlendirmeleri, yatırım verileri ve teknik pazar raporları özellikle model dağıtımı, kullanıcı erişimi, yarı iletken rekabeti ve veri merkezi kapasitesi gibi alanlarda boşluk doldurucu işlev görür. Ancak bu tür veriler tek başına kesin kanıt olarak kullanılmaz.

6.5. Kör noktalar ve bilinçli boşluklar

Her ciddi raporun gücü, yalnızca bildiklerinde değil; bilmediklerini açıkça adlandırabilmesinde yatar. Bu çalışma, teyitsiz devlet bağlantıları, bağımsız doğrulaması olmayan teknik kapasite iddiaları, doğrulanmamış sızıntılar ve tek kaynaklı sansasyonel haberler üzerine ana analiz kurmaz.

6.6. Hukuki savunulabilirlik ve metin içi atıf disiplini

Bu rapor, yalnızca analitik açıdan değil, hukuki ve kurumsal savunulabilirlik açısından da tasarlanmıştır. Her kritik iddia, mümkün olduğunca birincil veya yüksek güvenilirlikli kurumsal kaynakla desteklenir; sayısal veriler özellikle kaynaklandırılır; yorum ile veri açık biçimde ayrılır.

Yanlış Karar Riski Kutusu – Bölüm II

Bu bölümden çıkarılması gereken temel uyarı şudur: yapay zekâ alanında niceliksel bolluğu gerçek güç dağılımı sanan her kurum yanlış teşhis üretir. Model sayısını rekabet, ürün çeşitliliğini bağımsızlık, teknik yeniliği denetlenebilirlik ve açık ekosistemi otomatik olarak şeffaflık saymak; karar vericiyi yanlış güven duygusuna sürükler.

Bölüm Sonu Değerlendirmesi

Bu bölümün temel sonucu açıktır: yapay zekâ alanında gerçek sorun 'kaç model var?' sorusu değil, hangi sayım rejiminin hangi gücü görünür kıldığıdır. Kamusal model bolluğu, stratejik kapasitenin gerçekten dağıldığı anlamına gelmemektedir. Tam tersine, model, platform, ürün, ajan, altyapı ve kullanım bağlamını birlikte okuyan bir yaklaşım; yoğunlaşmayı, bağımlılığı, risk üretimini ve denetim boşluklarını daha net ortaya koymaktadır.

BÖLÜM III - KÜRESEL YAPAY ZEKÂ EKOSİSTEMİ

7. Küresel Genel Manzara

Bu bölümün amacı, yapay zekâ alanındaki görünür çoğulluğun neden gerçek bir güç dağınlıklığı anlamına gelmediğini göstermek ve küresel ekosistemin fiilî yoğunlaşma mantığını ortaya koymaktır. Açık ekosistemlerde milyonlarca model bulunmasına rağmen, etki üreten ana düğümler; frontier laboratuvarları, hyperscaler bulut sağlayıcıları, hızlandırıcı çip üreticileri ve yüksek erişimli ürün/platform aileleri etrafında toplanmaktadır (Hugging Face, 2026; Maslej et al., 2026).

7.1. Milyonlarca model, onlarca gerçek güç merkezi

İlk bakışta yapay zekâ ekosistemi son derece dağınık görünmektedir. Nitekim Hugging Face, 2025 sonunda 13 milyon kullanıcıya, 2 milyondan fazla kamuya açık modele ve 500 binden fazla veri kümesine ulaştığını bildirmektedir. Ancak aynı değerlendirme, bu görünür çoğulluğun ciddi bir yoğunlaşmayı perdelediğini de göstermektedir: platformdaki modellerin yaklaşık yarısı toplamda 200’den az indirilmeye sahiptir; en çok indirilen ilk 200 model ise, toplam model sayısının yalnızca çok küçük bir kesimini temsil etmesine rağmen, bütün indirmelerin %49,6’sını oluşturmaktadır (Hugging Face, 2026).

Dolayısıyla niceliksel büyüme ile stratejik güç aynı şey değildir. Yapay zekâ çağında belirleyici olan, yalnızca kaç modelin yayımlandığı değil; hangi modellerin kurumsal iş akışlarına entegre edildiği, hangi laboratuvarların yüksek maliyetli eğitimi karşılayabildiği ve hangi platformların bu modelleri küresel ölçeğe taşıyabildiğidir. Bu nedenle ekosistemin gerçek anatomisi, “milyonlarca model” değil; “onlarca merkezî güç odağı” üzerinden okunmalıdır (Hugging Face, 2026; Maslej et al., 2026).

7.2. Frontier kavramının yeniden tanımı

Bugün “frontier model” ifadesi yalnızca daha yüksek benchmark puanı alan modeli anlatmamaktadır. Frontier kapasite; hesaplama yoğunluğu, veri erişimi, sermaye derinliği, küresel dağıtım kabiliyeti ve güvenlik-hizalama rejimleri ile birlikte düşünülmelidir. Stanford AI Index 2026, 2025 yılında dikkat çekici modellerin %90’ından fazlasının endüstriden geldiğini; Amerika Birleşik Devletleri merkezli kurumların 50, Çin merkezli kurumların ise 30 dikkat çekici model ürettiğini göstermektedir. Aynı rapor, en güçlü modellerin aynı zamanda en az şeffaf olanlar hâline geldiğini; eğitim kodu, parametre sayısı, veri kümesi büyüklüğü ve eğitim süresi gibi temel bilgilerin giderek daha az açıklandığını vurgulamaktadır (Maslej et al., 2026).

Bu nedenle frontier artık yalnızca teknik bir üstünlük kategorisi değil; aynı zamanda altyapı, kurumsal kapalılık ve yönetim gücü kategorisidir. Bir modelin etkisi, performans tablolarındaki konumundan ziyade, hangi bulut mimarisi üzerinde çalıştığına, hangi ürünlere gömüldüğüne ve hangi devlet veya piyasa yapılarıyla eklemlendiğine bağlı olarak büyümektedir (Maslej et al., 2026).

7.3. Açık/kapalı model ikiliğinin ötesi

Yapay zekâ alanındaki ikinci büyük yanılsama, açık modellerin gücü otomatik olarak dağıttığı, kapalı modellerin ise tek başına merkezileştirdiği varsayımıdır. Oysa açık ekosistem de kendi içinde güçlü bir yoğunlaşma üretmektedir. Hugging Face’in 2026 ilkbahar değerlendirmesi, açık ekosistemin büyüdüğünü; fakat bu büyümenin büyük teknoloji şirketleri, yüksek görünürlüklü model aileleri ve çok sayıda türev üretim merkezi etrafında yeniden kümelenildiğini göstermektedir. Ayrıca Fortune 500 şirketlerinin %30’undan fazlasının Hugging Face üzerinde doğrulanmış hesap

bulundurması, açık ekosistemin marjinal değil, kurumsal stratejiye doğrudan temas eden bir alan hâline geldiğini göstermektedir (Hugging Face, 2026).

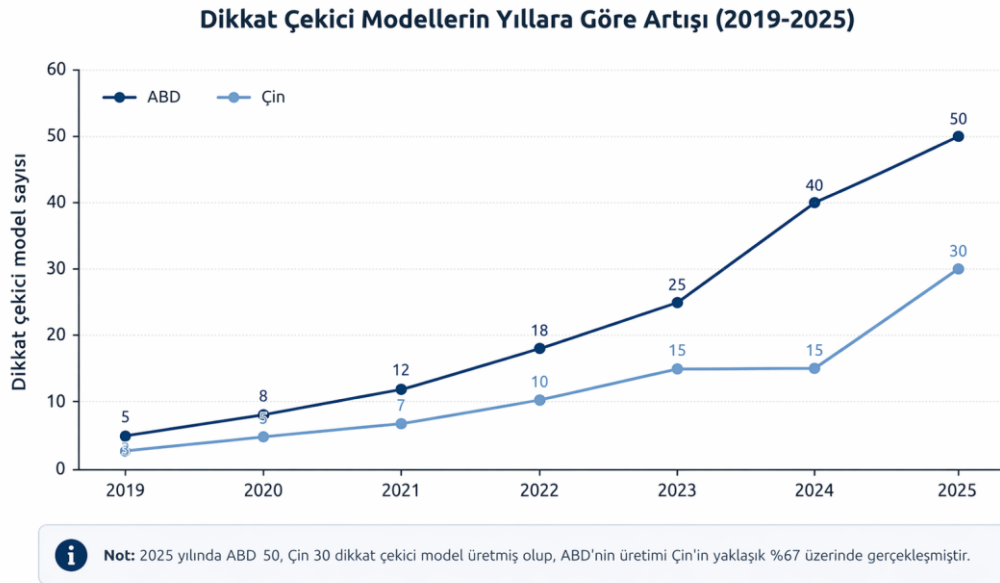
Bu dönüşüm coğrafi düzlemde de dikkat çekicidir. Hugging Face verileri, Çin'in aylık ve toplam indirmelerde ABD'yi geçtiğini ve Çin kaynaklı modellerin son bir yılda indirmelerin %41'ine ulaştığını göstermektedir. Başka bir ifadeyle açık model alanı, serbest dolaşan bir teknik topluluk alanı olmaktan çıkıp jeopolitik rekabetin ikinci sahasına dönüşmektedir. Bu nedenle açık/kapalı ayrımı, tek başına açıklayıcı değildir; asıl mesele, hangi model ailesinin hangi altyapı, hangi lisans ve hangi dağıtım kanalıyla küresel dolaşıma girdiğidir (Hugging Face, 2026).

Tablo 3.1. Küresel yapay zekâ ekosisteminin ana kutupları ve ayırt edici özellikleri

Kutup	Ana güç kaynağı	Yapısal üstünlük	Başlıca kırılganlık	Gösterge
ABD	Frontier laboratuvarlar, bulut, sermaye, ürünleştirme	Model üretimi + yatırım + veri merkezi yoğunluğu	Tekelleşme, şeffaflık azalması, yüksek enerji maliyeti	2025'te 50 dikkat çekici model; 2024'te 109,1 milyar \$ özel AI yatırımı
Çin	Devlet-özel bütünleşmesi, ölçek ve açık model ivmesi	Araştırma hacmi, patentler, hızlı ürün yayılımı	Sansür, dış güven sorunu, jeopolitik sınırlamalar	2025'te 30 dikkat çekici model; açık model indirmelerinde %41 pay
Avrupa	Norm üretimi ve düzenleyici kapasite	AI Act ve standart koyma gücü	Sınırlı frontier ölçeği ve yatırım açığı	2024'te yalnızca 3 dikkat çekici model; düzenleyici öncülük
İkincil düğümler	Kanada, Birleşik Krallık, Körfez, İsrail, Asya-Pasifik	Niş uzmanlık, egemen AI ve araştırma kümeleri	Ölçek, çip ve platform bağımlılığı	Patent yoğunluğu, veri merkezi yatırımları, savunma ve robotik nişleri
Küresel Güney	Veri üretimi, kullanıcı pazarı, dilsel ve davranışsal veri	Pazar büyüklüğü ve veri potansiyeli	Asimetrik bağımlılık, düşük yerel işleme kapasitesi	LDC'lerde internet kullanımı yalnızca %20 düzeyinde

Kaynak: Maslej ve diğerleri (2025, 2026); Hugging Face (2026); European Commission (2025); UNCTAD (2021).

Grafik 3.1. ABD ve Çin'in dikkat çekici model üretimi (2024-2025)



8. Bölgesel Güç Haritaları

8.1. ABD: merkezileşmiş inovasyon

Küresel ekosistemin çekirdeğinde Amerika Birleşik Devletleri yer almaktadır. Stanford AI Index 2025'e göre ABD merkezli kurumlar 2024 yılında 40 dikkat çekici model üretmiş; aynı yıl özel yapay zekâ yatırımı 109,1 milyar dolara ulaşarak Çin'in 9,3 milyar dolarlık ve Birleşik Krallık'ın 4,5 milyar dolarlık seviyelerinin çok üzerine çıkmıştır. Stanford AI Index 2026 ise bu üstünlüğün sürdüğünü ve 2025'te ABD'nin 50 dikkat çekici model ürettiğini göstermektedir (Maslej et al., 2025; Maslej et al., 2026).

ABD'nin farkı yalnızca şirket sayısında değildir. Aynı ekosistem içinde frontier laboratuvarları, hyperscaler bulut altyapıları, kurumsal yazılım dağıtım ağları ve veri merkezi yoğunluğu birbirini desteklemektedir. Nitekim AI Index 2026, Amerika Birleşik Devletleri'nde 5.427 veri merkezi bulunduğunu ve bunun diğer ülkelerin çok üzerinde olduğunu belirtmektedir. Bu nedenle ABD modeli, günlük inovasyondan çok merkezileşmiş ve ölçek ekonomileriyle beslenen bir inovasyon rejimi üretmektedir (Maslej et al., 2026).

8.2. Çin: devlet-özel bütünleşmesi

Çin, küresel yapay zekâ düzenindeki ikinci ana kutup olarak yalnızca model üretimiyle değil, devlet-strateji bütünleşmesiyle ayrılmaktadır. Çin Devlet Konseyi'nin 2017 tarihli Yeni Nesil Yapay Zekâ Gelişim Planı, yapay zekâyı açık biçimde ulusal kalkınma, endüstriyel modernizasyon ve stratejik rekabet aracı olarak tanımlamıştır (State Council of the People's Republic of China, 2017). Bu çerçevede, şirket faaliyetlerini daha geniş bir devlet önceliğiyle hizalayan bir yapı üretmiştir.

Stanford AI Index 2026'ya göre Çin, yayın hacmi, atıf üretimi ve patentlerde liderliğini sürdürmekte; aynı zamanda 2025 yılında 30 dikkat çekici model üretmektedir. Açık ekosistem tarafında ise Hugging Face verileri, Çin'in hem aylık hem toplam indirmelerde ABD'yi geçtiğini ve Çin kaynaklı modellerin son bir yılda indirmelerin %41'ine ulaştığını göstermektedir. Böylece Çin, yalnızca devlet stratejisi ve patent hacmiyle değil, açık model dolaşımındaki ivmesiyle de ekosistemin temel güç merkezlerinden biri hâline gelmiştir (Maslej et al., 2026; Hugging Face, 2026).

Harita 3.1. Küresel yapay zekâ güç coğrafyasının kavramsal haritası

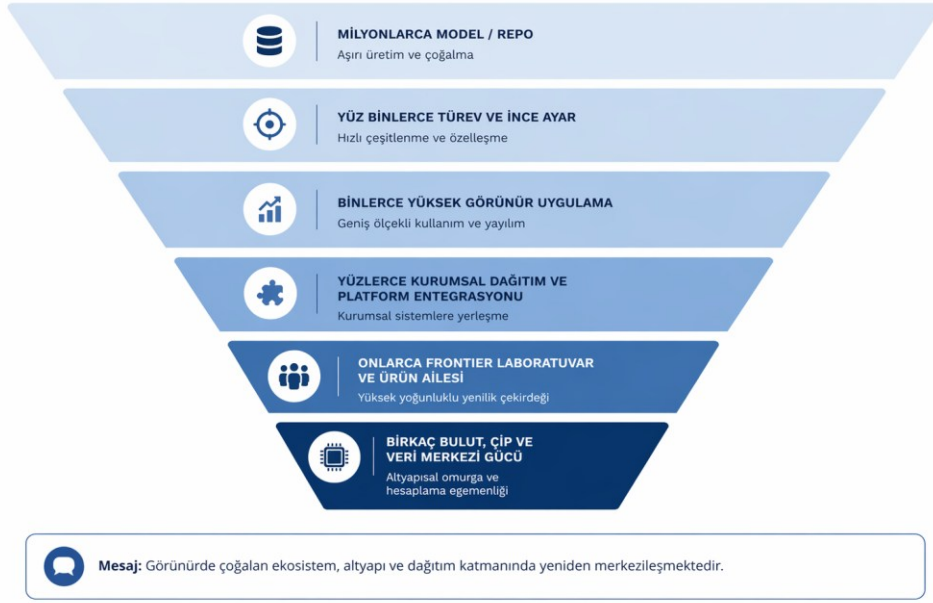


8.3. Avrupa: norm üreticisi, teknoloji ithalatçısı

Avrupa'nın ayırt edici gücü, bugün için teknik ölçekten çok norm üretimindedir. Avrupa Komisyonu'na göre Yapay Zekâ Tüzüğü 1 Ağustos 2024'te yürürlüğe girmiş; belirli yasaklı uygulamalar ve yapay zekâ okuryazarlığına ilişkin hükümler 2 Şubat 2025'te, genel amaçlı yapay zekâ modellerine ilişkin yükümlülükler ise 2 Ağustos 2025'te uygulanmaya başlamıştır. Bu nedenle Avrupa, frontier model yarışında ilk iki kutbun gerisinde kalsa da, küresel norm ve uyum baskısı üreten başlıca merkezdir (European Commission, 2025).

Buna karşılık model geliştirme ve sermaye derinliği bakımından Avrupa daha sınırlı bir görünüm sergilemektedir. Stanford AI Index 2025, 2024 yılında Avrupa'nın toplam yalnızca üç dikkat çekici model ürettiğini göstermektedir. Sonuç olarak Avrupa, teknolojik tedarik bakımından daha sınırlı; fakat düzenleme ve standartlaştırma bakımından daha etkili bir güç kutbu oluşturmaktadır (Maslej et al., 2025).

Şema 3.1. Görünür çeşitlilikten fiili yoğunlaşmaya geçiş



8.4. İkincil düğümler: Kanada, Körfez, Asya-Pasifik ve İsrail

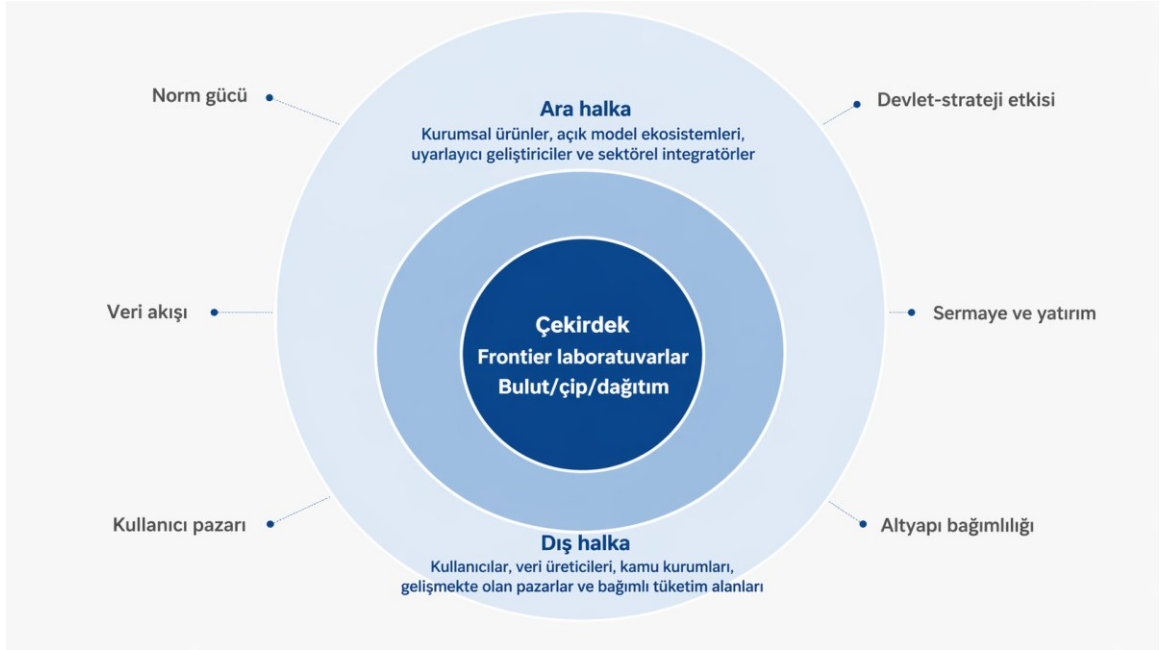
Küresel düzen yalnızca ABD, Çin ve Avrupa'dan ibaret değildir. Stanford AI Index 2025, model geliştirme faaliyetlerinin Orta Doğu, Latin Amerika ve Güneydoğu Asya gibi bölgelerde daha görünür hâle geldiğini; 2026 raporu ise Güney Kore'nin kişi başına yapay zekâ patentlerinde öne çıktığını vurgulamaktadır. Körfez ülkeleri egemen AI ve veri merkezi yatırımlarıyla, Kanada ve Birleşik Krallık araştırma ve girişimcilik ağlarıyla, İsrail ve Asya-Pasifik'in belirli bölümleri ise savunma, robotik ve donanım nişleriyle ikinci halka aktörleri hâline gelmektedir (Maslej et al., 2025; Maslej et al., 2026).

8.5. Küresel Güney ve veri sömürgeleşmesi

Küresel Güney için sorun yalnızca yapay zekâyâ erişim değildir; veriden doğan değer kim tarafından işlendiği ve kimin lehine biriktiği sorunudur. UNCTAD'a göre en az gelişmiş ülkelerde internet kullanım oranı yalnızca %20 düzeyindedir ve bağlantı daha düşük hızlarla, daha yüksek maliyetlerle gerçekleşmektedir. Aynı çerçevede, gelişmekte olan ülkelerin çoğu zaman küresel platformlar için ham veri sağlayıcısı hâline geldiğini, fakat kendi verilerinden elde edilen dijital zekâyı daha sonra dışarıdan satın almak zorunda kaldığını vurgulamaktadır (UNCTAD, 2021).

Bu nedenle yapay zekâ ekosistemi, yalnızca yenilik ve verimlilik açısından değil; veri kolonizasyonu, dilsel bağımlılık, temsil krizleri ve enformasyon egemenliği açısından da okunmalıdır. Küresel Güney'in sorunu teknolojiye geç kalmak değil; çoğu zaman değer en alt katmanında konumlandırılırken karar ve kâr zincirinin dış halkasında bırakılmasıdır (UNCTAD, 2021).

Şekil 3.1. Küresel yapay zekâ ekosisteminin üç halkalı güç geometrisi



Bölüm Sonu Değerlendirmesi

Bu bölümün temel sonucu açıktır: günümüz yapay zekâ düzeni görünürde çoğul, gerçekte ise yoğunlaşmış bir güç coğrafyası üretmektedir. Açık model alanındaki niceliksel patlama, gücün gerçekten dağıldığı anlamına gelmemektedir; aksine frontier laboratuvarları, bulut-çip altyapısı, yatırım akışları ve düzenleyici kapasite birkaç ana kutup etrafında birikmektedir. ABD merkezileşmiş inovasyon ve sermaye yoğunluğu ile; Çin devlet-özel bütünleşmesi ve açık model ivmesi ile; Avrupa ise norm üretimi ve düzenleyici baskı gücü ile ayrılmaktadır. Küresel Güney ise veri, pazar ve bağımlılık ilişkilerinin dış halkasında konumlanma riski taşımaktadır. Bu nedenle küresel yapay zekâ ekosistemi, teknoloji pazarı olarak değil, çok katmanlı bir egemenlik mimarisi olarak okunmalıdır (Maslej et al., 2025; Maslej et al., 2026; Hugging Face, 2026; UNCTAD, 2021).

BÖLÜM IV – TEKNİK TAM-YIĞIN ANALİZ

Zekâ nerede üretilir, nerede kilitlenir?

9. Model Türleri

Bu bölümün amacı, güncel yapay zekâ ekosistemini yalnızca ürün isimleri üzerinden değil; model aileleri, yerleştirme katmanları, araç kullanımı ve hesaplama altyapısı üzerinden okumaktır. Bugün teknik üstünlük, tek bir modelin doğruluk düzeyinden çok; ön eğitim, çok kipli işleme, dış bilgiye erişim, araç kullanımı, güvenlik katmanı ve dağıtım mimarisinin birlikte çalışmasına bağlıdır (Bommasani ve diğerleri, 2021; Maslej ve diğerleri, 2026).

9.1. Büyük Dil Modelleri

Büyük dil modelleri, modern üretken yapay zekâ çağının temel hesaplama çekirdeğini oluşturmaktadır. Bu yapılar, Transformer mimarisinin sağladığı dikkat mekanizması üzerine kurulmakta; geniş ölçekli ön eğitim yoluyla dilsel örüntüleri, bağlamsal ilişkileri ve görevler arası genellenebilirliği öğrenmektedir (Vaswani ve diğerleri, 2017). Foundation model yaklaşımı, aynı model ailesinin çok sayıda görevde yeniden kullanılabilmesini mümkün kıldığı için ekonomik ve stratejik değeri büyötmektedir; ancak aynı özellik, hata ve önyargının da çok sayıda uygulamaya yayılmasına yol açmaktadır (Bommasani ve diğerleri, 2021).

9.2. Multimodal modeller

Yeni kuşak frontier sistemler artık yalnızca metin tabanlı değildir. GPT-4o örneğinde göröldüğü gibi model; metin, görüntü, ses ve videoya ilişkin veriler üzerinde eğitilerek tek bir mimari içinde çok kipli giriş-çıkış üretebilmektedir (OpenAI, 2024a; OpenAI, 2024b). Bu dönüşüm, kullanım alanlarını genişletmekte; aynı zamanda değerlendirme, güvenlik ve içerik kökeni sorunlarını da büyötmektedir. Çünkü çok kipli modeller yalnızca soru cevaplamaz; ses tonunu işler, görsel bağlamı yorumlar ve giderek daha gerçekçi sentetik içerik üretir.

9.3. Görsel, ses ve video üretim sistemleri

Görsel, ses ve video üretim sistemleri çoğu zaman ayrı ürünler gibi görünse de teknik olarak giderek daha fazla natif çok kipli model ailesinin uzantısına dönüşmektedir. OpenAI'nin 4o tabanlı görüntü üretim sistemi, görsel üretimin dil modelinden kopuk bir eklenti değil, bizzat modelin birincil kabiliyetlerinden biri hâline geldiğini göstermektedir (OpenAI, 2025a). Bu kayma, üretim doğruluğunu ve kullanışlılığı artırırken; deepfake, kimlik taklidi, telif ve içerik doğrulama sorunlarını da daha merkezî bir konuma taşımaktadır.

9.4. Retrieval ve ajan mimarileri

Model kalitesini artırmanın tek yolu parametre büyötmek değildir. Retrieval-augmented generation (RAG), dış veri kaynaklarını sorgu anında modele bağlayarak güncellik ve alan özgüllüğü sorununu azaltır (NVIDIA, 2025a). Bir sonraki aşama ise ajan mimarileridir: bu sistemler yalnızca cevap üretmez, araç çağırır, veritabanına bağlanır, görev planlar ve çıktıyı yineleyerek iyileştirir. Anthropic'in Model Context Protocol yaklaşımı, ajanların veri kaynakları ve araçlarla standartlaştırılmış biçimde bağlanmasını amaçlayarak bu alanı altyapısal bir ekosisteme dönüştürmektedir (Anthropic, 2024a; Anthropic, 2024b).

9.5. Savunma ve kamuya özel modeller

Savunma veya kamuya yönelik yapay zekâ sistemleri çoğu zaman tamamen farklı bir model ailesi olmaktan ziyade; aynı temel modellerin ağ yalıtımı, günlükleme, insan denetimi, uygunluk kontrolleri ve sınıflandırılmış ortam entegrasyonlarıyla sertleştirilmiş sürümleridir. OpenAI for Government, Claude Gov ve Microsoft 365 Copilot'un GCC High/DoD ortamlarında sunulması,

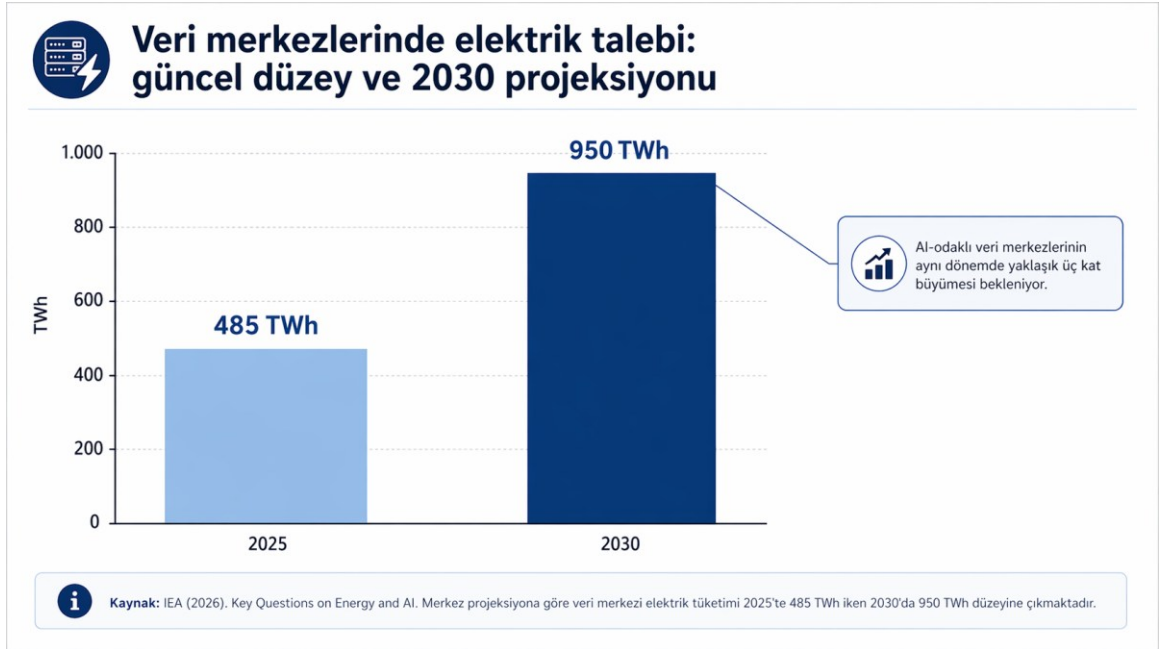
teknik farkın yalnızca modelde değil; dağıtım yüzeyi, veri sınırı ve güvenlik rejiminde oluştuğunu göstermektedir (OpenAI, 2025b; Anthropic, 2025a; Microsoft, 2026a).

Tablo 4.1. Model aileleri ve tam-yığın içindeki teknik işlevleri

Kategori	Temel işlev	Başlıca örnek	Kurumsal değeri	Başlıca risk
Büyük dil modelleri	Genel dil işleme, akıl yürütme, özetleme	GPT-4o, Claude, Gemini	Çok görevli çekirdek katman	Halüsinasyon, açıklanabilirlik açığı
Multimodal modeller	Metin + görüntü + ses/video işleme	GPT-4o, Gemini tabanlı sistemler	Tek modelle çok kipli görevler	İçerik güvenliği ve köken sorunu
Üretim modelleri	Görsel, ses ve video üretimi	4o Image, ses/video üretim araçları	Yaratıcı iş akışı ve medya otomasyonu	Deepfake, telif, kimlik taklidi
RAG sistemleri	Dış veri ile bağlam zenginleştirme	Kurumsal arama, belge yanıt sistemleri	Güncel ve alan-öзgü çıktı	Veri sızıntısı, kötü indeksleme
Ajan sistemleri	Araç çağırma, planlama, görev icrası	MCP tabanlı ajanlar, Copilot akışları	Süreç otomasyonu ve karar desteği	Yetki aşımı, zincirleme hata
Kamu/savunma dağıtımı	Yalıtılmış ağ, denetim ve uyum	Claude Gov, GCC High Copilot	Görev-kritik ve sınıflı kullanım	Gizlilik, güvenlik, yanlış karar

Kaynak: Vaswani ve diğerleri (2017); Bommasani ve diğerleri (2021); OpenAI (2024a, 2025a); NVIDIA (2025a); Anthropic (2024a, 2024b); OpenAI (2025b); Microsoft (2026a); yazarın sentezi.

Grafik 4.1. Veri merkezi elektrik talebinde 2025–2030 sıçraması



10. Hesaplama, Çip ve Enerji Gerçeği

Yapay zekânın maddi temeli, model boyutundan çok hesaplama yoğunluğunda görünür hâle gelir. Eğitim ve çıkarım süreçleri; hızlandırıcı çipler, yüksek bant genişlikli bellek, hızlı ağ omurgası, veri merkezi elektriği ve soğutma altyapısı olmadan sürdürülemez. Bu nedenle teknik tam-yığın analiz, yalnızca yazılım katmanını değil; yarı iletken tedarik zincirini ve enerji sistemini de birlikte okumayı gerektirir (IEA, 2026; Maslej ve diğerleri, 2026).

10.1. GPU, TPU, NPU, LPU ve wafer-scale ayrımı

GPU'lar yüksek paralellik gerektiren matris işlemleri için genel amaçlı hızlandırıcı rolü görürken, TPU'lar tensor iş yüklerine daha sıkı optimize edilmiş özel mimarilerdir. NPU ve benzeri yapılar uç cihazlar ile enerji verimliliği yüksek iş yüklerine yönelirken, Groq'nun LPU yaklaşımı özellikle çıkarım tarafında deterministik ve düşük gecikmeli performansı öne çıkarmaktadır (Google Cloud, 2024a; Groq, 2025a). Cerebras'ın wafer-scale yaklaşımı ise tek bir wafer üzerinde olağanüstü bellek bant genişliği ve hesaplama yoğunluğu elde ederek farklı bir ölçeklenme yolu sunmaktadır (Cerebras, 2024a).

10.2. Nvidia hegemonyası ve rakip hızlandırıcı aileleri

Bugün endüstri referans noktası hâlen NVIDIA'dır. Şirketin Blackwell mimarisi 208 milyar transistörlü yeni nesil bir hızlandırıcı ailesi olarak konumlandırılmakta ve büyük ölçekli üretken yapay zekâ altyapısının ana omurgası olmayı hedeflemektedir (NVIDIA, 2024a). Bununla birlikte rekabet hızlanmaktadır: AMD'nin MI300X hızlandırıcısı 6 Aralık 2023'te pazara açılmıştır; Intel Gaudi 3 2024'te lanse edilmiş; Google'ın altıncı nesil Trillium TPU'su 2024 sonunda genel kullanıma açılmış; AWS Trainium2 ise birinci nesle göre dört kata kadar daha hızlı eğitim ve iki kata kadar daha yüksek performans/watt vaadiyle duyurulmuştur (AMD, 2023a; Intel, 2024a; Google Cloud, 2024a; Amazon, 2023a).

10.3. Çin'in alternatif yolu ve sistem düzeyi yanıtı

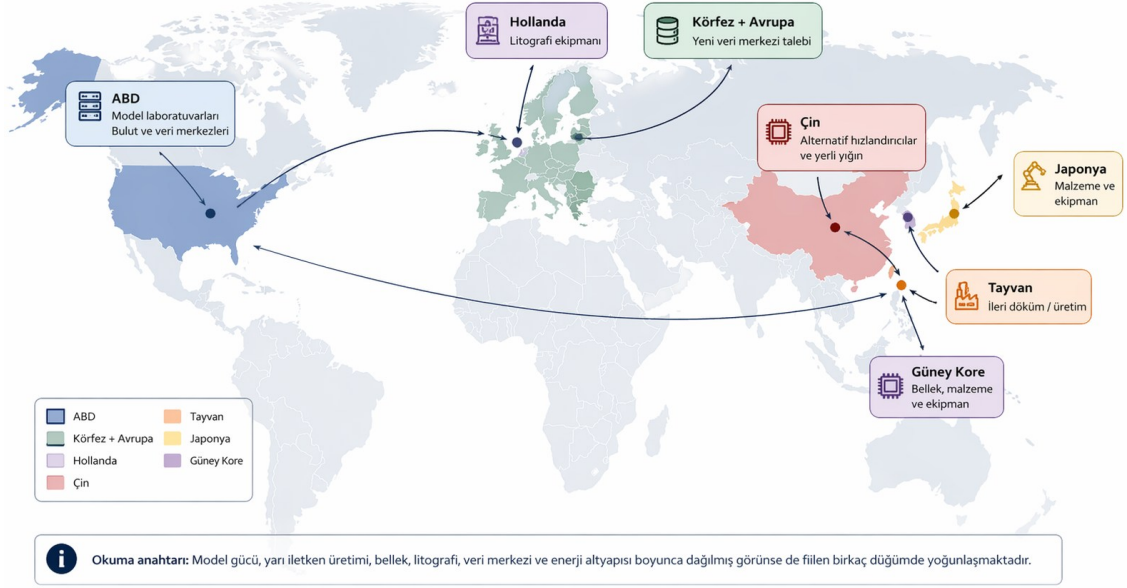
İhracat kontrollerinin sıkılaşması, Çin'i yalnızca çip düzeyinde değil; sistem düzeyinde de alternatif yollar aramaya itmiştir. Huawei'nin 2025'te resmî olarak tanıttığı Atlas 900 A3 SuperPoD, 384'e kadar Ascend hızlandırıcısını tek mantıksal düğüm gibi çalıştırmayı hedefleyen bir yaklaşım sunmaktadır. Ürünün resmî teknik tanıtımına göre sistem 307,2 PFLOPS düzeyine kadar FP16 hesaplama gücü sağlayabilmektedir (Huawei, 2025a). Bu yaklaşım, tek çip üstünlüğü yerine bütün yığının birlikte optimize edilmesini amaçlayan bir karşı stratejiyi temsil etmektedir.

10.4. Enerji, su ve iklim maliyeti

Hesaplama yarışının en kritik sonucu, enerji ve çevresel maliyetlerde görülmektedir. IEA'nın 2026 güncellemesine göre veri merkezlerinin elektrik talebi 2025'te yaklaşık 485 TWh iken 2030'da yaklaşık 950 TWh düzeyine ulaşacaktır; AI-odaklı veri merkezlerinin elektrik talebi ise aynı dönemde yaklaşık üç kat büyüyecektir (IEA, 2026). Stanford AI Index 2026 ayrıca AI veri merkezi güç kapasitesinin 29,6 GW'a yükseldiğini ve yıllık GPT-4o çıkarımının su tüketiminin 12 milyon insanın içme suyu ihtiyacını aşabileceğini vurgulamaktadır (Maslej ve diğerleri, 2026). Dolayısıyla teknik üstünlük, artık yalnızca parametre ve benchmark meselesi değil; elektrik, su, arazi, soğutma ve şebeke kapasitesi meselesidir.

Harita 4.1. Küresel hesaplama ve yarı iletken coğrafyasının kavramsal haritası

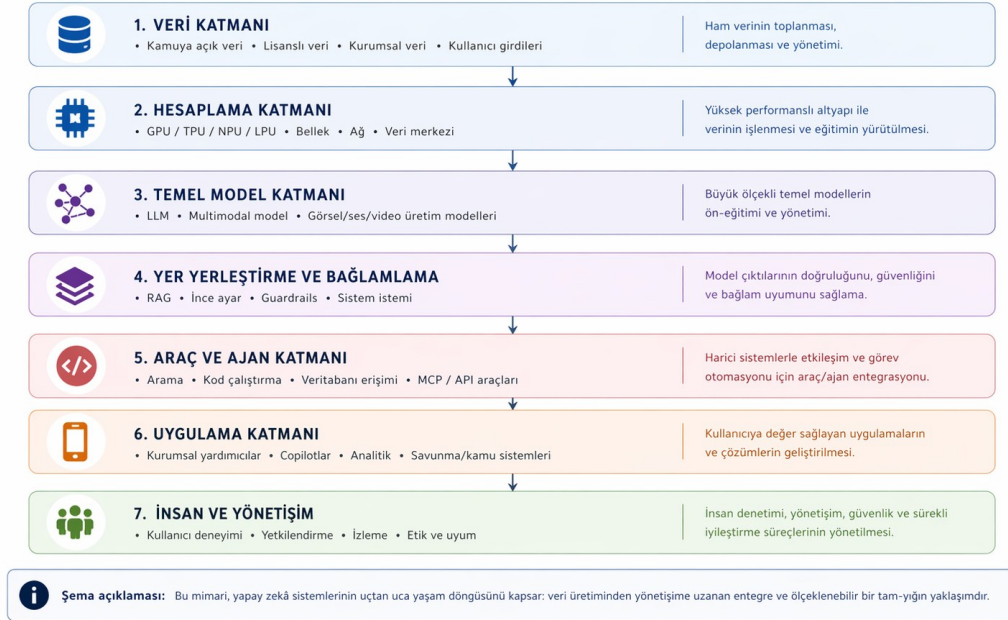
Küresel hesaplama ve yarı iletken coğrafyası Kavramsal harita



Şema 4.1. Teknik tam-yığın yapay zekâ mimarisi



Teknik tam-yığın yapay zekâ mimarisi



Şekil 4.1. Teknik tam-yığında yoğunlaşma ve risk eşleşmesi

Teknik tam-yığında yoğunlaşma ve risk eşleşmesi

Yapay zekâ sistemlerinde değer zinciri boyunca biriken teknik yoğunlaşma, spesifik riskleri artırır ve yönetim sonucunu doğrudan etkiler.



Şekil açıklaması: Şekil, teknik tam-yığında gücün neden yalnızca model kalitesinde değil; veri, altyapı, hizalama ve dağıtım katmanlarında biriktiğini göstermektedir. Katmanlar arasında yukarı doğru çıkıldıkça soyutlama artar, kontrol azalır ve risk büyür.

Bölüm Sonu Değerlendirmesi

Bu bölümün temel sonucu şudur: yapay zekânın teknik çekirdeği tek bir modelden ibaret değildir. Güç; veri, hesaplama, temel model, bağlamsal yerleştirme, araç kullanımı, dağıtım ve yönetim katmanlarının eşzamanlı kontrolüyle oluşmaktadır. Bu nedenle güncel rekabet, yalnızca model kalitesi rekabeti değil; tam-yığın rekabetidir. Aynı nedenle risk de yalnızca halüsinasyonda değil; çip darboğazında, enerji baskısında, veri sızıntısında, yetki aşımında ve platform bağımlılığında ortaya çıkmaktadır.

BÖLÜM V – VERİ, GÜÇ VE EGEMENLİK

21. yüzyılın petrolü bir benzetme değil, bir zayıflatma

11. Veri Türleri

Bu bölümün amacı, yapay zekâ sistemlerinin başarısını belirleyen veri katmanını yalnızca teknik bir girdi olarak değil; ekonomik değer, kurumsal güç ve egemenlik üreten stratejik bir kaynak olarak incelemektir. Güncel yapay zekâ ekosisteminde veri, tek başına anlamlı değildir; verinin nasıl toplandığı, hangi hukukî rejim altında işlendiği, hangi altyapı üzerinde modele dönüştürüldüğü ve elde edilen değer nereden biriktiği belirleyicidir. OECD'nin yapay zekâ, veri ve rekabet çalışması; temel modellerin yaşam döngüsünde verinin hesaplama gücüyle birlikte başlıca giriş bariyeri olduğunu, özellikle eğitim ve dağıtım aşamalarında yüksek kaliteli veri erişiminin rekabeti doğrudan etkilediğini vurgulamaktadır (OECD, 2024). Bu nedenle veri tartışması, salt teknik performans meselesi değil; aynı zamanda piyasa yoğunlaşması, bağımlılık ve stratejik kontrol meselesidir.

Veri katmanının siyasal ve toplumsal boyutu da en az teknik boyutu kadar kritiktir. Zuboff, veri çıkarımını yalnızca dijital ekonominin yeni hammaddesi olarak değil, davranışın öngörülmesi ve yönlendirilmesi üzerinden kurulan yeni bir iktidar biçimi olarak tanımlamaktadır (Zuboff, 2019). Couldry ve Mejias ise bu süreci, fiziksel toprak işgalinden çok insan yaşamının veri biçiminde sistematik olarak el konulmasına dayanan bir “veri kolonizasyonu” olarak okumaktadır (Couldry & Mejias, 2019). Bu yaklaşım, yapay zekâ çağında verinin neden yalnızca mühendislik problemi değil; emek, temsil, kültür ve egemenlik meselesi olduğunu açıklamaktadır.

Tablo 5.1. Yapay zekâ veri katmanları ve egemenlik riskleri

Veri türü	Başlıca kaynak	Yapay zekâdaki işlevi	Egemenlik avantajı	Başlıca risk
Kamuya açık veriler	Web, açık arşivler, kamusal veri tabanları	Ön eğitim ve genel bilgi genişliği	Düşük maliyetli ölçeklenme	Telif, kalite ve bağlam kaybı
Lisanslı veriler	Yayıncılar, veri sağlayıcıları, kurumsal sözleşmeler	Daha kontrollü ve kaliteli eğitim girdisi	Hukukî güvence ve kurumsal kalite	Maliyet yoğunlaşması ve erişim eşitsizliği
Kullanıcı girdileri	Prompts, geri bildirim, iş akışı verileri	İnce ayar, ürün iyileştirme, davranışsal öğrenme	Gerçek kullanım örüntülerini yakalama	Mahremiyet, sır sızıntısı, davranışsal izleme
Sentetik veri	Model üretimli veri kümeleri	Veri genişletme ve simülasyon	Nadir senaryolarda hızlı ölçekleme	Hata çoğalması ve modele özgü önyargı

Kaynak: OECD (2024); Zuboff (2019); Couldry ve Mejias (2019); yazarın sentezi.

11.1. Kamuya açık veriler

Kamuya açık veriler, frontier modellerin erken eğitim aşamalarında geniş kapsama ulaşmak için başvurduğu temel kaynaktır. Açık web, akademik arşivler, kamu kurumlarının yayınları, forumlar ve kod depoları bu katmanın ana unsurlarıdır. Ancak bu alanın görünürde açık olması, hukukî ve epistemik sorunları ortadan kaldırmamaktadır. OECD, temel modellerin eğitiminde kullanılan verilerin bileşiminin çoğu zaman sınırlı biçimde açıklandığını; yeterli kalitede veriye erişimin ise rekabet ve giriş bariyerleri bakımından belirleyici olduğunu belirtmektedir (OECD, 2024). Dolayısıyla kamuya açık veriler ölçek sağlar; fakat ölçeğin kendisi doğruluk, temsil adaleti ve telif güvenliği üretmez.

11.2. Lisanslı veriler

Lisanslı veri katmanı, üretken yapay zekâ ekosisteminde giderek daha kritik hâle gelmektedir. Yayıncılar, veri sağlayıcıları ve kurumsal sözleşmeler üzerinden temin edilen veri; daha temiz, denetlenebilir ve sektörel bağlamı kuvvetli girdiler sunduğu için özellikle kurumsal modellerde değerli kabul edilmektedir. Avrupa Komisyonu'nun veri stratejisi, verinin ekonomik değer üretimi kadar güvenli paylaşım ve tekrar kullanım kapasitesini de artırmayı hedefleyen tek pazar mantığını öne çıkarmaktadır (European Commission, 2020). Ne var ki bu yönelim, veriye erişimi yalnızca teknik kabiliyet meselesi olmaktan çıkarıp sermaye ve sözleşme gücü meselesine dönüştürmektedir.

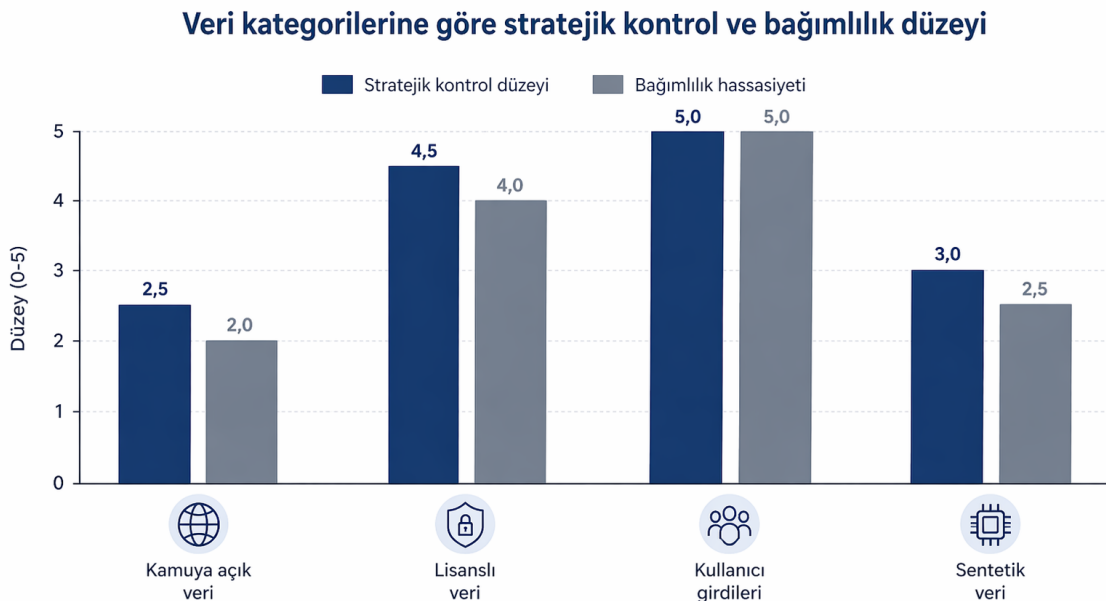
11.3. Kullanıcı girdileri

Kullanıcı girdileri, güncel yapay zekâ ürünlerinin en hassas veri katmanlarından biridir. Günlük kullanım sırasında girilen istemler, düzeltmeler, geri bildirimler ve kurumsal iş akışlarından çıkan veriler; sistemlerin ince ayarında ve ürün yönünün belirlenmesinde önemli rol oynayabilmektedir. UNESCO'nun etik tavsiye kararı, insan hakları, mahremiyet, orantılılık ve insan denetimi ilkelerinin tam da bu nedenle yapay zekâ yaşam döngüsünün bütününde gözetilmesi gerektiğini vurgulamaktadır (UNESCO, 2021). Kullanıcı girdileri üretkenliği artırırken, aynı anda ticari sır, kişisel veri ve kurumsal davranış örüntülerinin platformlaştırılması riskini de büyütmektedir.

11.4. Sentetik veri

Sentetik veri, özellikle dengesiz veri kümeleri, nadir senaryolar ve gizlilik kısıtları karşısında yapay zekâ geliştiricilerine esneklik sağlamaktadır. Bununla birlikte sentetik verinin stratejik değeri kadar önemli bir sınırı vardır: model temelli üretildiği için, mevcut önyargıların ve hataların geniş ölçekte yeniden üretilmesine yol açabilir. OECD'nin 2024 çalışması, yapay zekâ değer zincirinin eğitim, ince ayar ve dağıtım aşamalarında veri kalitesinin kritik olduğunu; veri ve hesaplama gücünün birlikte yoğunlaşma ürettiğini ortaya koymaktadır (OECD, 2024). Bu nedenle sentetik veri, veri kıtlığını çözen sihirli araç değil; dikkatli yönetim gerektiren ikinci kuşak bir güç çarpanıdır.

Grafik 5.1. Veri kategorilerine göre stratejik kontrol ve bağımlılık düzeyi



Not: Ölçüler kavramsal skorlardır; OECD (2024), Zuboff (2019), Couldry ve Mejias (2019) ile yazarın sentezine dayalıdır.

12. Veri Akışı ve Egemenlik

Veri akışları, yapay zekâ çağında yalnızca teknik dolaşım süreçleri değildir; değer, gözetimin, standartların ve bağımlılık ilişkilerinin de dolaşım kanallarıdır. UNCTAD'ın 2021 Dijital Ekonomi Raporu, sınır ötesi veri akışlarının kalkınma bakımından fırsat üretse de mevcut düzenin oldukça parçalı ve kutuplaşmış olduğunu; veri akışları konusunda ülkeler arasında serbest dolaşım ile yerelleştirme arasında uzanan geniş bir yelpaze bulunduğunu vurgulamaktadır (UNCTAD, 2021). Bu nedenle yapay zekâ çağında veri akışı, serbestleşme ya da korumacılık ikiliğine indirgenemeyecek kadar siyasal bir sorundur.

12.1. Sınır ötesi veri hareketi

Sınır ötesi veri hareketi, üretken yapay zekâ hizmetlerinin küresel ölçekte çalışabilmesi için fiilî bir ön koşul hâline gelmiştir. Ancak verinin bir ülkede toplanıp başka bir ülkede işlenmesi, üçüncü bir ülkede modele dönüştürülmesi ve dördüncü bir ülkede ticarileştirilmesi; değer zincirinin coğrafi ve hukukî olarak parçalanmasına yol açmaktadır. Avrupa veri stratejisi, bu parçalanmaya karşı ortak veri alanları ve güvenli paylaşım düzenekleri geliştirmeyi amaçlamaktadır (European Commission, 2020). Buna karşılık UNCTAD, küresel veri yönetiminde kutuplaşmanın sürmesi nedeniyle gelişmekte olan ülkelerin çoğu zaman kural koyucu değil, veri sağlayıcı ve pazar konumunda kaldığını belirtmektedir (UNCTAD, 2021).

12.2. Ulusal mevzuatlar ve fiilî durum

Ulusal düzeyde benimsenen veri koruma, yerelleştirme veya paylaşım rejimleri ile fiilî piyasa pratikleri çoğu zaman birbirinden ayrılmaktadır. Kâğıt üzerinde güçlü koruma ilkeleri taşıyan sistemler bile, bulut bağımlılığı, yabancı model servisleri ve kurumsal platform entegrasyonları yoluyla veri egemenliğini aşındırabilmektedir. UNESCO'nun etik çerçevesi; şeffaflık, izlenebilirlik, hesap verebilirlik ve insan haklarının korunmasını yalnızca normatif ilke olarak değil, gerçek uygulama koşulu olarak tanımlar (UNESCO, 2021). Bu nedenle veri egemenliği yalnızca yasa yazmakla değil; depolama, işleme, denetim ve hesaplama katmanlarını birlikte yönetebilmekle mümkündür.

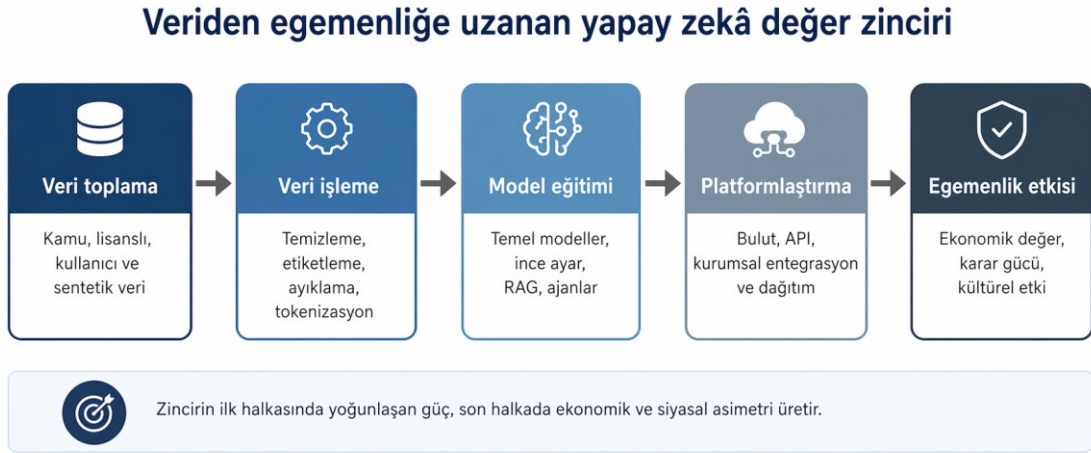
12.3. Dil, kültür ve hafıza erozyonu

Veri egemenliği tartışmasının en az konuşulan ama en kritik boyutlarından biri dil, kültür ve hafıza üzerindeki etkidir. Couldry ve Mejias'ın veri kolonizasyonu yaklaşımı, dijital altyapıların yalnızca ekonomik değer değil, aynı zamanda yaşam biçimleri ve anlam rejimleri üzerinde denetim kurduğunu gösterir (Couldry & Mejias, 2019). Zuboff'un analizinde de davranışsal artık yalnızca pazarlama aracı değil; gelecekteki eylemleri öngörme ve yönlendirme altyapısıdır (Zuboff, 2019). Yapay zekâ çağında baskın diller, baskın veri rejimleri ve baskın platformlar; hangi bilginin görünür, hangi hafızanın kalıcı ve hangi kültürel ifadenin meşru sayılacağını dolaylı biçimde etkileyebilmektedir. Bu nedenle veri akışı meselesi, aynı anda temsiliyet, bellek ve epistemik adalet meselesidir.

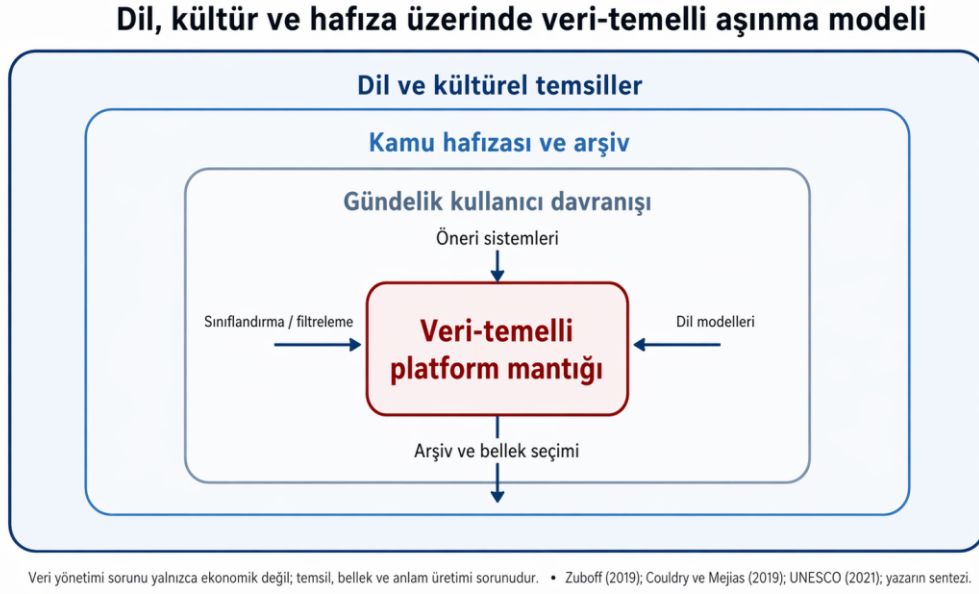
Harita 5.1. Sınır ötesi veri akışlarında üç kutuplu düzenin kavramsal haritası



Şema 5.1. Veriden egemenliğe uzanan yapay zekâ değer zinciri



Şekil 5.1. Dil, kültür ve hafıza üzerinde veri-temelli aşınma modeli



Bölüm Sonu Değerlendirmesi

Bu bölümün temel sonucu açıktır: yapay zekâ çağında veri, tek başına “yeni petrol” değildir; çünkü petrol benzetmesi verinin ilişkisel, çoğaltılabilir, izlenebilir ve davranış biçimlendirici niteliğini zayıflatır. Gerçek mesele, verinin hangi hukukî çerçevede toplandığı, hangi altyapıda işlendiği, hangi modelleri beslediği ve ortaya çıkan değer kim tarafından kontrol edildiğidir. Kamuya açık veri ölçek, lisanslı veri kalite, kullanıcı girdileri davranışsal hassasiyet, sentetik veri ise hız ve genişletme kapasitesini sağlar; ancak bunların tümü aynı anda yeni bağımlılık ve eşitsizlik rejimleri de üretir. Bu nedenle veri yönetişimi, yalnızca mahremiyet koruması değil; piyasa yapısı, kültürel temsil, kurumsal özerklik ve ulusal egemenlik meselesidir (OECD, 2024; UNCTAD, 2021; UNESCO, 2021).

BÖLÜM VI – GÜVENİLİRLİK, HATA VE SİSTEMİK RİSK

Doğruluk sorunu teknik bir kusur değil, kurumsal bir kırılmalıktır

13. Halüsinasyon ve Doğruluk Sorunu

Bu bölümün amacı, üretken yapay zekâ sistemlerinde ortaya çıkan doğruluk sorunlarını yalnızca model kalitesine ilişkin bir mühendislik problemi olarak değil; karar kalitesi, kamusal güven ve kurumsal sorumluluk üzerinde doğrudan etkileri olan sistemik bir risk alanı olarak incelemektir. NIST'in Yapay Zekâ Risk Yönetimi Çerçevesi, risk yönetimini güvenilirlik üretmenin ve kamu güvenliğini korumanın ön şartı olarak tanımlamaktadır (NIST, 2023). NIST'in Generative AI Profile belgesi ise üretken yapay zekâ bağlamında 'confabulation' riskini, modelin güven tonu içinde yanlış, dayanaksız ya da uydurma içerik üretmesi olarak özel biçimde ayırmakta ve bu riskin yaşam döngüsünün farklı aşamalarında farklı yoğunluklarla ortaya çıkabileceğini vurgulamaktadır (NIST, 2024). Bu nedenle halüsinasyon, yalnızca bilgi yanlışlığı değil; yanlışın ikna edici, akıcı ve çoğu zaman denetlenmeden dolaşıma girmesi problemidir.

13.1. Confabulation, bağlam kaybı ve sahte kesinlik

Üretken modellerin en büyük yapısal güçlerinden biri, eksik girdileri tamamlayabilmeleri ve yüksek akıcılıkta çıktı üretebilmeleridir. Ancak aynı özellik, bağlamı eksik ya da çatışmalı durumlarda yanlış cevabın yüksek güven tonu ile verilmesine de yol açmaktadır. NIST, bu nedenle üretken yapay zekâ risklerinin yalnızca model çıktısında değil; veri görünürlüğü, insan davranışı, kullanım bağlamı ve ekosistem düzeyinde birlikte değerlendirilmesi gerektiğini belirtmektedir (NIST, 2024). Sorunun özü, modelin bazen bilmediğini bilmemesi değil; bilmediği şeyi bilmiş gibi sunabilmesidir.

Bu kırılmalık, özellikle uzmanlık gerektiren alanlarda daha görünür hâle gelmektedir. Hukuk alanına ilişkin sistematik bir çalışma, büyük dil modellerinde hukuki halüsinasyonların somut, doğrulanabilir sorularda bile son derece yaygın olduğunu; test edilen sistemlerde bu oranların yüzde 58 ile yüzde 88 arasında değiştiğini göstermiştir (Dahl et al., 2024). Stanford HAI'nin daha sonraki değerlendirmesi de hukuk odaklı modellerin kıyas sorgularında en az altıda bir oranında, genel amaçlı sohbet sistemlerinin ise daha da yüksek oranlarda halüsinasyon üretebildiğini belirtmektedir (Stanford HAI, 2024). Bu bulgu, hatanın yalnızca amatör kullanıcı düzeyinde değil, mesleki uzmanlık alanlarında da yapısal bir sorun olduğunu göstermektedir.

13.2. Ölçüm sorunu: doğruluk tek başına yeterli değildir

Doğruluk meselesi yalnızca benchmark başarısı üzerinden değerlendirildiğinde yanıltıcı sonuçlar verebilir. NIST'in ARIA pilot raporu, yapay zekâ sistemlerinin geçerlilik riskini anlamak için yalnızca otomatik testlerin değil; uzman değerlendiriciler, kırmızı takım çalışmaları ve saha benzeri senaryoların birlikte kullanılması gerektiğini göstermiştir. Beş kuruluşun sunduğu yedi uygulama için toplam 508 test oturumu gerçekleştirilmiş olması, güvenliliğin ancak çok katmanlı ölçümle anlaşılabilceğini ortaya koymaktadır (NIST, 2025). Bu nedenle yüksek puanlı bir model, kullanım bağlamı değiştiğinde otomatik olarak güvenilir sayılmamalıdır.

Model geliştiricilerinin kendi sistem kartları da bu sınırı teyit etmektedir. OpenAI, GPT-4.5 için yayımladığı sistem kartında halüsinasyon oranlarının iyileştiğini belirtmekle birlikte, özellikle değerlendirme kapsamı dışında kalan alanlarda daha fazla çalışmaya ihtiyaç olduğunu açıkça kabul etmektedir (OpenAI, 2025). Benzer biçimde Anthropic, 2026 tarihli sistem kartında dürüstlük ve halüsinasyon azaltımını çekirdek hedeflerden biri olarak tanımlamakta; yani en gelişmiş

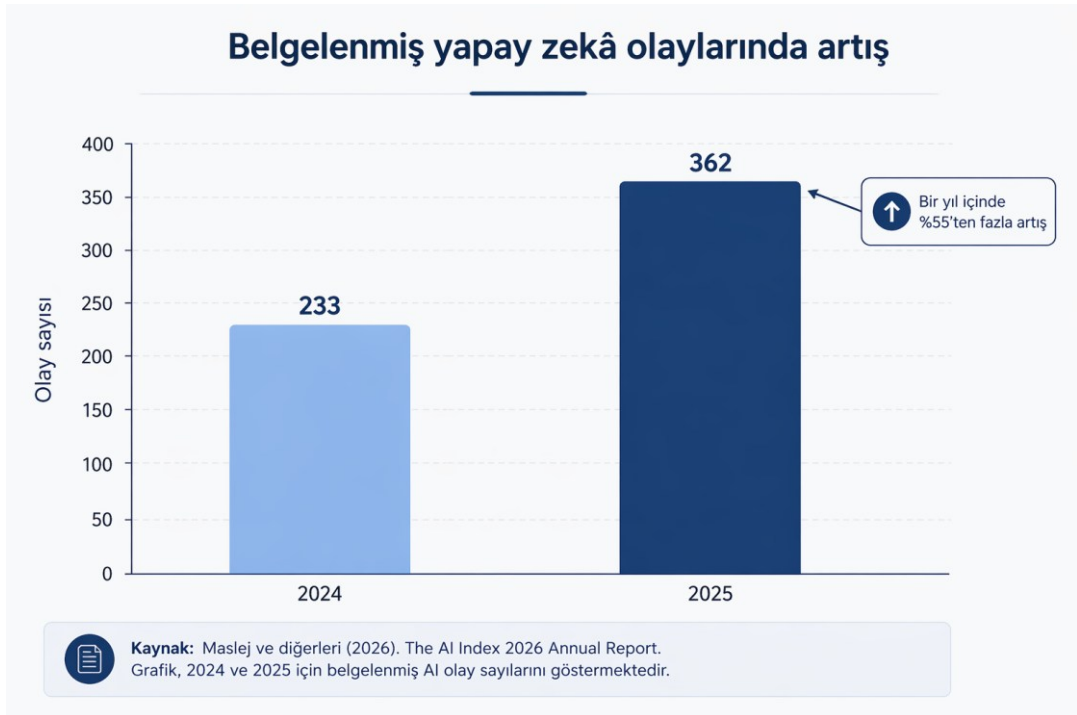
modellerde bile sorun tamamen çözülmüş değil, yönetilmeye çalışılan bir risk olarak kalmaktadır (Anthropic, 2026).

Tablo 6.1. Güvenilirlik ve sistemik risk tipolojisi

Risk tipi	Teknik kaynak	Tipik görünüm	Kurumsal sonuç
Confabulation / halüsinasyon	eksik veri, zayıf kalibrasyon, bağlam dışı üretim	uydurma bilgi, yanlış atıf, sahte kesinlik	yanlış karar, hukukî ve operasyonel hata
Bağlam kaybı	genel modelin özel kullanım alanına körlüğü	sorunun amacını kaçırarak doğru görünümlü yanıt	verimsizlik, yanlış yönlendirme, gecikme
Otomatikleşmiş yanlış sınıflama	yanlı veri, zayıf hedef tanımı, temsil sorunu	yanlış risk skoru, uygunsuz eleme, yanlış önceliklendirme	hak kaybı, ayrımcılık, itibar aşınması
Güven kalibrasyonu bozukluğu	ikna edici dil, açıklanabilirlik açığı	yanlış çıktıya aşırı güven	insan denetiminin biçimselleşmesi
Ekosistem riski	aynı model ve tedarikçinin yaygın kullanımı	eşzamanlı ve korelasyonlu hata	sistemik kırılganlık, kamu güveninde erozyon

Kaynak: NIST (2023, 2024, 2025); OECD (2025); European Commission (2026); yazarın sentezi.

Grafik 6.1. Belgelenmiş yapay zekâ olaylarında artış



14. Otomatikleşmiş Yanlış Kararlar

Üretken yapay zekâ ve öngörüsül sistemler, insan karar süreçlerini hızlandırma ve destekleme vaadiyle kurumsal yapılara giderek daha fazla entegre edilmektedir. Ancak karar desteği ile kararın fiilen otomatikleşmesi arasındaki sınır çok hızlı aşınabilmektedir. OECD'nin sosyal koruma alanına ilişkin 2025 raporu, kamu kurumlarının yapay zekâyı başvuru yönlendirme, belge işleme, uygunluk kontrolü ve dolandırıcılık tespiti gibi süreçlerde giderek daha fazla denediğini; buna karşılık güven, veri yönetimi ve insan denetimi konularında ciddi tereddütlerin sürdüğünü göstermektedir (OECD, 2025). Sorun burada yalnızca hatalı model çıktısı değil; hatanın iş akışında sorgulanmadan karar haline gelmesidir.

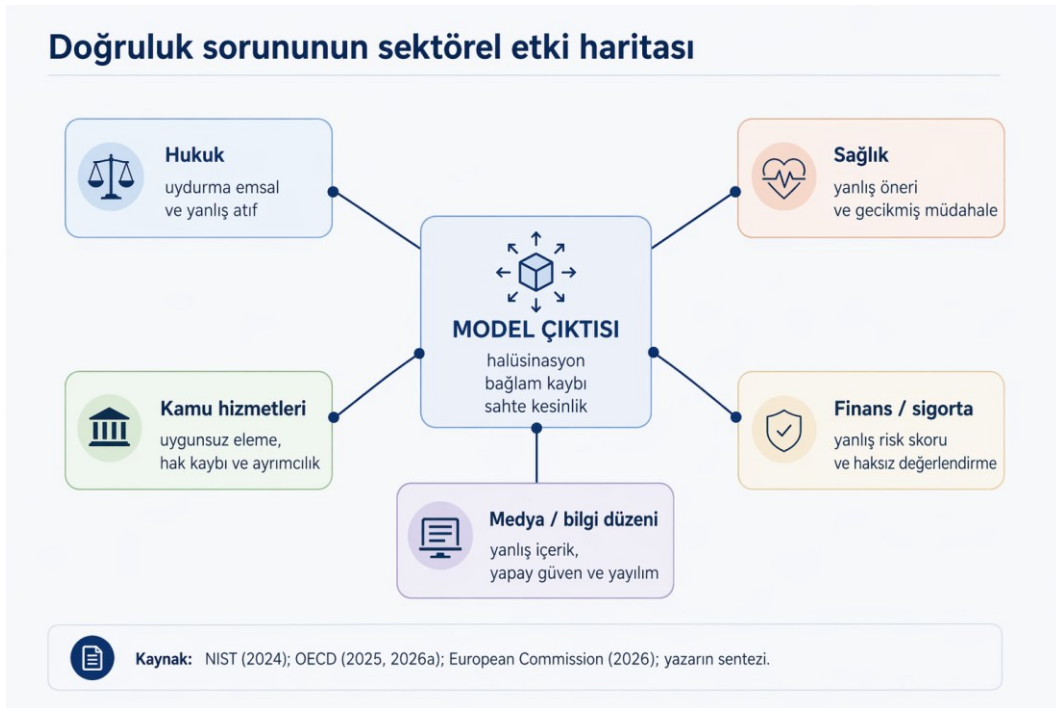
14.1. İnsan gözetiminin biçimselleşmesi

Avrupa Komisyonu'nun Yapay Zekâ Tüzüğü uygulama rehberi, yüksek riskli sistemlerde insan gözetimi, ilgili ve yeterince temsil edici girdi verisi, doğruluk, dayanıklılık ve siber güvenlik gerekliliklerini birlikte vurgulamaktadır. Rehber ayrıca kamu hizmeti sunan yapılar için temel haklar etki değerlendirmesini, açıklama yükümlülüğünü ve ciddi olay bildirimi mekanizmalarını öne çıkarmaktadır (European Commission, 2026). Bu vurgu tesadüf değildir. Çünkü insan gözetimi kâğıt üzerinde mevcut olsa bile, kararın hız baskısı ve iş yükü nedeniyle fiilen model lehine kapanması mümkündür. Böyle bir durumda insan denetimi gerçek bir fren değil, yalnızca kurumsal meşruiyet için eklenmiş biçimsel bir adım haline gelir.

14.2. Yüksek etkili alanlarda bağlam uyumsuzluğu

Yüksek etkili alanlarda sorun çoğu zaman modelin genel performansından değil, kullanım bağlamıyla uyumsuzluğundan doğar. Aynı model müşteri destek hattında kabul edilebilir bir hata oranıyla çalışabilirken, kredi değerlendirmesi, kamu yardımı, sağlık yönlendirmesi veya adli araştırma bağlamında çok daha ağır sonuçlar üretebilir. OECD, gelecekteki yapay zekâ risklerine ilişkin değerlendirmesinde kritik sistemlerdeki olaylar, manipülasyon, dolandırıcılık, güç yoğunlaşması ve eşitsizlik artışı gibi riskleri özellikle vurgulamaktadır (OECD, 2024). Bu nedenle otomatikleşmiş yanlış karar, tek tek kullanıcı hatalarının toplamı değildir; yanlışın yüksek etkili alanlarda kurumsallaşmasıdır.

Harita 6.1. Doğruluk sorununun sektörel etki haritası



Şema 6.1. Hatalı çıktıdan kurumsal hasara uzanan zincir



15. Kurumsal ve Kamusal Hata Maliyetleri

Yapay zekâ hatalarının maliyeti, çoğu zaman ilk bakışta görünen yanlış cevaptan daha büyüktür. NIST AI RMF, güvenilirliği; geçerlilik, emniyet, güvenlik, açıklanabilirlik ve hesap verebilirlikle bağlantılı biçimde ele alırken, bu özelliklerin eksikliği halinde bireyler ve kurumlar için beklenmedik olumsuz sonuçların ortaya çıkacağını belirtmektedir (NIST, 2023). Üretken sistemlerde yanlış cevap, kolaylıkla yanlış işlem, yanlış işlem ise yanlış kayıt, yanlış ret, yanlış yönlendirme ya da yanlış kamuoyu etkisi biçimini alabilir. Böylece hata, teknik düzeyden yönetim düzeyine sıçar.

15.1. Mahremiyet, güvenlik ve dolandırıcılık baskısı

OECD'nin 2026 tarihli olay ve tehlike izleme çalışması, sentetik medya ile ilişkili olayların 2022 ile 2025 arasında 2,5 kat arttığını; siber saldırı ve dolandırıcılık temalı olay payının ise yaklaşık yüzde 3,7'den neredeyse yüzde 10'a yükseldiğini göstermektedir (OECD, 2026a). Bu eğilim, güvenilirlik krizinin yalnızca doğru bilgi üretimiyle ilgili olmadığını; aynı zamanda güven istismarı, kimlik taklidi, belge sahteciliği ve operasyonel suistimal risklerini de büyüttüğünü göstermektedir.

15.2. İtibar, meşruiyet ve kamu güveni

Kurumsal yapılarda yapay zekâyâ aşırı güven, özellikle kamusal otorite kullanan kurumlarda meşruiyet riskine dönüşür. OECD'nin sosyal koruma raporu, vatandaşların kamu hizmetlerinde yapay zekâ kullanımına henüz tam güven duymadığını; özellikle hassas karar alanlarında açıklanabilirlik, adalet ve insan temasının korunmasının kritik olduğunu göstermektedir (OECD, 2025). Bu nedenle kamusal hata maliyeti yalnızca yanlış karardan ibaret değildir; kurumun tarafsız, öngörülebilir ve hesap verebilir olduğuna dair inancın da aşınmasıdır.

15.3. Yönetim sonucu: doğruluk rejimi kurmak

Bu koşullarda kurumsal yanıt, daha çok model eklemek değil; doğruluk rejimi kurmaktır. NIST'in Generative AI Profile belgesi olay bildirimi, içerik kökeni, ön dağıtım testi ve yönetim süreçlerini özellikle öne çıkarırken, Avrupa Komisyonu da yüksek riskli sistemlerde kayıt, izleme, açıklama ve insan gözetimini birlikte zorunlu kılmaktadır (NIST, 2024; European Commission, 2026). Başka bir ifadeyle, yapay zekâ hatasını sıfırlamak mümkün olmasa da hatanın karar, kurum ve toplum düzeyinde büyümesini sınırlamak mümkündür. Kritik mesele, modelin tek başına ne kadar akıllı olduğu değil; kurumun hatayı ne kadar erken gördüğü, ne kadar açık kaydettiği ve gerektiğinde karar döngüsünü insana geri çekebildiğidir.

Şekil 6.1. İnsan gözetimi ile etki şiddeti arasındaki risk matrisi



Bölüm Sonu Değerlendirmesi

Bu bölümün temel sonucu şudur: yapay zekâ sistemlerindeki güvenilirlik sorunu, teknik doğruluk ölçütleriyle sınırlı bir performans problemi değildir. Halüsinasyon, bağlam kaybı, zayıf insan gözetimi ve otomatikleşmiş iş akışları birleştiğinde, hata kurumsallaşabilir ve kamusal güveni aşındırabilir. Bu nedenle yüksek etkili alanlarda asıl soru, modelin ne kadar etkileyici olduğundan çok; hatanın nasıl ölçüldüğü, nasıl kayıt altına alındığı, hangi bağlamda kullanıldığı ve karar yetkisinin kimde kaldığıdır. Yapay zekâ çağında güvenilirlik, model niteliği kadar yönetim niteliğiyle de belirlenmektedir.

BÖLÜM VII – MANİPÜLASYON, PROPAGANDA, DENORMASYON VE SOĞUK KAOS

Enformasyon alanında teknik üstünlük, bilişsel üstünlüğe dönüşebildiğinde risk yalnızca yanlış bilgi değil, kurumsal yön kaybıdır

16. Deepfake Ekonomisi

Bu bölümün amacı, yapay zekâ destekli manipülasyonun neden artık ikincil bir yan etki değil; doğrudan ekonomik, siyasal ve güvenlik boyutu olan bir üretim alanı hâline geldiğini göstermektir. OECD'nin 2026 tarihli medya-temelli olay izleme çalışması, sentetik medya bağlantılı yapay zekâ olay ve tehlikelerinin 2022 ile 2025 arasında 2,5 kat arttığını; bu temanın 2025'in üçüncü çeyreğinde toplam kayıtlı olayların yüzde 14'ünden fazlasını oluşturduğunu ve Kasım 2023'te yüzde 22'ye kadar yükseldiğini göstermektedir (OECD, 2026a). Bu veri, deepfake üretiminin yalnızca teknolojik bir gösteri alanı olmadığını; dikkat ekonomisi, dolandırıcılık, itibar suikastı, şantaj ve siyasi manipülasyon için ölçeklenebilir bir piyasa hâline geldiğini ortaya koymaktadır.

Deepfake ekonomisinin belirleyici özelliği üretim maliyetini düşürürken inandırıcılık düzeyini yükseltmesidir. Europol'un derin sahtecilik raporu, bu teknolojilerin insanların hiç söylemediği sözleri söyler gibi, hiç yapmadıkları eylemleri yapar gibi veya hiç var olmayan kişilikleri gerçekmiş gibi gösterebildiğini vurgulamaktadır; rapor ayrıca bu tehdidi değerlendirmek üzere 80'den fazla kolluk uzmanıyla stratejik öngörü çalışmaları yürütüldüğünü belirtmektedir (Europol, 2024). Bu durum, sahte içeriğin artık münferit bir montaj tekniği değil; hizmet olarak sunulabilen, suç ağlarınca uyarlanabilen ve farklı hedef gruplara hızla ölçeklenebilen bir ikna teknolojisi olduğunu göstermektedir. Buradaki en önemli kırılma, görsel gerçekçiliğin doğrulukla karıştırılmasıdır. C2PA açıklayıcı metni, içerik kökeni ve dolaşım geçmişine ilişkin provenance bilgisinin içerik hakkında değerli kanıtlar sağlayabildiğini, ancak tek başına içeriğin doğru, isabetli veya olgusal olduğunu kanıtlamayacağını açık biçimde belirtmektedir (C2PA, 2023). Dolayısıyla deepfake ekonomisinde asıl tehdit, yalnızca sahte içerik üretimi değildir; sahte ile gerçek arasındaki doğrulama maliyetinin toplumsal olarak artmasıdır.

Tablo 7.1. Yapay zekâ destekli manipülasyon ekosisteminin temel bileşenleri

Katman	Başlıca araç/teknik	Birincil hedef	Anlık etki	İkincil stratejik etki
Üretim	Görüntü, ses ve video sentezi	Bireyler, siyasetçiler, kurumlar	Gerçeklik yanılsası	Delil rejiminin bulanıklaşması
Dağıtım	Bot ağları, öneri sistemleri, reklam altyapısı	Seçmenler, tüketiciler, hedef topluluklar	Erişim ve görünürlük patlaması	Gündem kontrolünün kayması
Hedefleme	Mikro-segmentasyon, kimlik taklidi, duygusal çerçeveleme	Kırılgan ve kutuplaşmış kitleler	İkna eşiğinin düşmesi	Toplumsal bölünmenin hızlanması
Paraya çevirme	Dolandırıcılık, şantaj, itibar saldırısı	Bireyler ve şirketler	Maddi ve itibar kaybı	Kurumsal güvenin aşınması
Siyasi/stratejik kullanım	Propaganda, seçim müdahalesi, hibrit kampanya	Kamuoyu ve karar alıcılar	Algı bozulması	Soğuk Kaos ve yönetim körlüğü

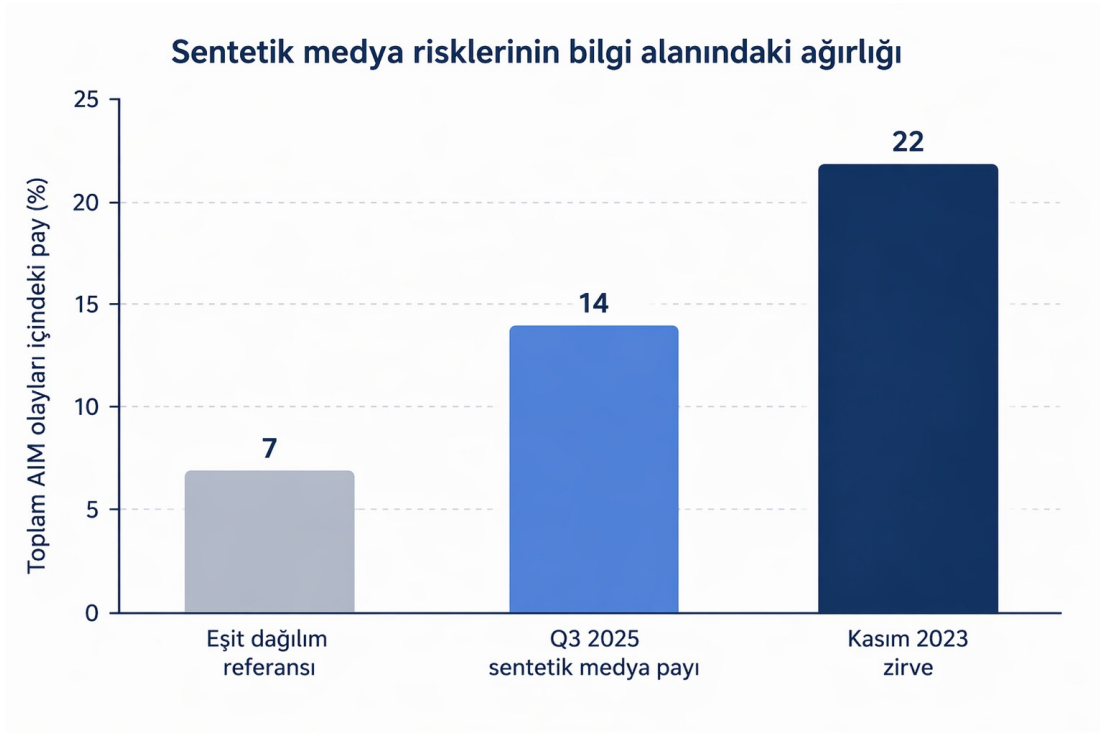
Kaynak: OECD (2026a); Europol (2024); yazarın sentezi.

Yanlış Karar Riski Kutusu

Deepfake meselesini yalnızca 'itibar yönetimi' problemi sanan kurumlar tehdidin gerçek doğasını ıskalar. Asıl risk, yapay olarak üretilmiş veya değiştirilmiş içeriğin ucuzlamasıyla birlikte doğrulama maliyetinin kurumun üzerine yıkılmasıdır. Yanlış karar, deepfake'ı tek tek sahte içerikler düzeyinde okumaktır; doğru okuma ise bunu kurumsal teyit maliyetini büyüten bir ekonomi olarak görmektir (European Commission, 2025; UNESCO, 2024).

Kurum, sesin, görüntünün ve ekran paylaşımının artık tek başına delil sayılamadığı bir döneme girmiştir. Çok katmanlı teyit, içerik provenansı ve işlem öncesi doğrulama prosedürü kurulmazsa, deepfake saldırısı yalnızca iletişim krizine değil; para transferi hatasına, sahte talimat uygulanmasına ve zincirleme meşruiyet kaybına yol açabilir.

Grafik 7.1. Sentetik medya risklerinin bilgi alanındaki ağırlığı



17. Algı Yönetimi ve Bilişsel Harp

Yapay zekâ destekli manipülasyonun ikinci boyutu, tek tek yanlış içeriklerin ötesine geçerek algı yönetimi ve bilişsel harp alanına yerleşmesidir. NATO Bilim ve Teknoloji Örgütü'nün 2025 tarihli bilişsel harp raporu, çağdaş çatışmalarda asıl hedefin yalnızca altyapı veya toprak değil; algı, muhakeme, karar ve davranış olduğunu vurgulamaktadır (NATO STO, 2025). Bu çerçevede yapay zekâ, etkisinin büyük bölümünü kinetik değil bilişsel sahada üretir: içerik üretir, dağıtır, hedefler, tekrarlar ve böylece kanaat oluşumunun ritmini değiştirir.

NATO StratCom COE'nin AI Laboratory girişi de bu dönüşümün kurumsal kabul gördüğünü göstermektedir. Merkezin açıklamasına göre laboratuvar, İttifakın bilgi çevresindeki geleceğe hazırlığını artırmakta; bilişsel harp alanında üstünlüğü sürdürmek için araştırma, prototip ve uzmanlık sağlamaktadır. Aynı sayfada seçim güvenliği ve düşmanca manipülasyona karşı dayanıklılık, öngörü ve karar destek kabiliyeti ve bilgi alanında ajan temelli kapasite geliştirme açık hedefler arasında sayılmaktadır (NATO StratCom COE, t.y.). Bu veri, bilişsel savaşın artık

metaforik değil; kurumlar tarafından işlevsel bir mücadele alanı olarak görüldüğünü göstermektedir.

Bu bağlamda yapay zekâ destekli propaganda, klasik propaganda ile aynı değildir. Rid'in gösterdiği gibi modern aktif önlemler, yalnızca yalan yaymakla kalmaz; hedef toplumda kuşku, parçalanma ve yönsüzlük üretir (Rid, 2020). Pomerantsev'in vurguladığı gibi çağdaş propaganda düzeni, insanları tek bir yalana inandırmaktan çok, hiçbir şeye tam olarak güvenmemeye iter (Pomerantsev, 2019). Yapay zekâ, bu etkiyi hız, ölçek, kişiselleştirme ve taklit gücüyle genişleterek bilişsel harbin maliyet-etkin ana çarpanlarından birine dönüşmektedir.

Harita 7.1. Yapay zekâ destekli manipölasyonun etki coğrafyası



Yanlış Karar Riski Kutusu

Algı yönetimi ve bilişsel harp alanını yalnızca sosyal medya dezenformasyonu olarak gören her analiz eksiktir. Modern çatışmada belirleyici alan yalnızca fiziksel altyapı değil; algı, muhakeme, karar ve davranıştır. Bu nedenle yanlış karar, bilişsel harbi içerik denetimi düzeyine indirgemektir; doğru karar ise bunu karar üstünlüğü mücadelesi olarak okumaktır (NATO STO, 2021; NATO StratCom COE, t.y.).

Bir kurum veya devlet, kendi karar vericilerinin hangi sinyal rejimi içinde düşündüğünü koruyamıyorsa, en gelişmiş teknik altyapıya sahip olsa bile stratejik olarak savunmasız kalır. Çünkü bilişsel harp sizi mutlaka yalanla ikna etmek zorunda değildir; yalnızca doğru sinyali geç fark etmenizi ya da yanlış sinyale aşırı tepki vermenizi sağlaması yeterlidir.

18. Seçimler, Medya ve Toplumsal Güven

Seçim dönemleri, yapay zekâ destekli manipölasyonun en görünür ve en hassas alanlarından biridir. UNESCO ile UNDP'nin 2025 tarihli ortak politika notu, yapay zekâ teknolojilerinin – öneri sistemlerinden üretken yapay zekâyâ kadar – bilginin nasıl üretildiğini, yayıldığını ve tüketildiğini özellikle seçim dönemlerinde derinden etkilediğini belirtmektedir (UNESCO & UNDP, 2025). Aynı belge, küresel nüfusun yüzde 56,8'inin sosyal medya platformlarında etkin olduğunu ve yaklaşık 4 milyar seçmenin algoritmik bilgi çevreleri içinde oy verme süreçlerine girdiğini vurgulamaktadır (UNESCO & UNDP, 2025).

Bu durum, seçim güvenliğini yalnızca sandık güvenliği veya oy sayımı düzeyinde değil; bilgi bütünlüğü, seçmen algısı ve kamusal tartışmanın asgari gerçeklik zemini üzerinden yeniden düşünmeyi zorunlu kılmaktadır. Deepfake’ler, sahte ses kayıtları, adaylara atfedilen uydurma beyanlar ve platform algoritmaları tarafından güçlendirilen sentetik içerikler, seçmeni ikna etmek kadar onu şüpheye sürüklemek için de kullanılabilir. Böylece seçim sürecinde ortaya çıkan zarar, her zaman doğrudan kanaat değişimi şeklinde görünmez; kimi zaman temel etki, güvenin aşınması, doğrulama kapasitesinin çökmesi ve medya kurumlarının meşruiyetinin sorgulanmasıdır (UNESCO & UNDP, 2025; OECD, 2026a).

Medya açısından kırılma daha da derindir. Sentetik içeriğin artması, gazetecilikte doğrulama zincirini uzatmakta; ‘gördüm, o hâlde oldu’ mantığını yapısal olarak çürütmektedir. C2PA’nın kendi açıklayıcı metni de provenance bilgisinin doğruluk garantisi vermediğini kabul etmektedir; bu da içerik kökeni araçlarının gerekli ama tek başına yeterli olmadığını göstermektedir (C2PA, 2023). Dolayısıyla seçim ve medya alanındaki asıl mücadele, yalnızca sahte içeriği tespit etmek değil; teyit kapasitesini, editoryal güvenilirliği ve demokratik bilgi akışını korumaktır.

Şema 7.1. Yapay zekâ destekli manipülasyon zinciri

Şema 7.1. Yapay zekâ destekli manipülasyon zinciri



Yanlış Karar Riski Kutusu

Seçim güvenliğini yalnızca oy verme günü güvenliği olarak gören yaklaşım artık yetersizdir. Üretken yapay zekâ, seçim dönemlerinde bilgi üretimini, dolaşımını ve tüketimini doğrudan etkiler; bu da yalnızca yanlış bilgi değil, bilgi bütünlüğü ve ifade özgürlüğü dengesini aynı anda zorlar. Bu yüzden yanlış karar, seçim riskini yalnızca sahte haber sayısı üzerinden ölçmektir; doğru yaklaşım ise onu toplumsal güven ve bilgi bütünlüğü rejimi olarak değerlendirmektir (UNESCO & UNDP, 2025).

Seçime yönelik AI tehdidi, yurttaşları mutlaka yanlış bilgiye inandırmak zorunda değildir; seçmeni, gazeteciyi ve kurumu neyin gerçek olduğundan emin olamaz hâle getirmesi bile yeterlidir. Böyle bir ortamda sorun yalnızca yanlış içerik değil, doğrulamanın yetişememesidir.

19. Sansür, Model Hizalaması ve Görünmez İdeoloji

Manipülasyon meselesi yalnızca zararlı içeriğin çoğalmasıyla sınırlı değildir; aynı zamanda hangi içeriğin görünür, meşru, öncelikli veya bastırılmış sayıldığıyla ilgilidir. Avrupa Komisyonu’nun Yapay Zekâ Tüzüğü uygulama rehberine göre, etkileşimli ve üretken yapay zekâ sistemleri için

getirilen şeffaflık yükümlülükleri; yanlış bilgilendirme, manipülasyon, dolandırıcılık, kimlik taklidi ve tüketicinin aldatılması risklerini azaltmayı amaçlamaktadır. Aynı rehber, üretken sistem sağlayıcılarının çıktıları makinece okunabilir biçimde işaretlemesi ve içeriklerin yapay olarak üretildiğinin tespit edilebilir olmasını sağlaması gerektiğini açıkça belirtmektedir (European Commission, 2026).

Bu düzenleyici yaklaşım önemli olmakla birlikte, model hizalamasının tarafsız bir teknik işlem olduğu sonucunu doğurmaz. Hangi içeriğin riskli sayıldığı, hangi istemlerin reddedildiği, hangi bağlamların güvenlik filtresine takıldığı ve hangi cevap biçimlerinin tercih edildiği; yalnızca güvenlik kararı değil, aynı zamanda normatif tercihler kümesidir. OECD'nin 2026 tarihli sorumlu yapay zekâ özen yükümlülüğü rehberi de bu nedenle risk tolerans eşikleri belirlenmesini, risk yönetimi politikalarının yayımlanmasını ve veri toplama ile içerik moderasyonu dahil değer zincirinin farklı halkalarında sorumlulukların açıkça tanımlanmasını önermektedir (OECD, 2026b). Başka bir ifadeyle, görünmez ideoloji çoğu zaman açık propaganda gibi değil; teknik tarafsızlık dili içinde gömülü kurumsal tercih biçimi olarak işler.

Yanlış Karar Riski Kutusu

Model hizalamasını yalnızca teknik güvenlik filtresi gibi görmek, yapay zekâ çağıının en yanıltıcı varsayımlarından biridir. Hangi içeriğin bastırıldığı, hangi cevabın sınırlandırıldığı ve hangi risk eşığının kabul edilebilir sayıldığı yalnızca teknik bir ayar değildir; aynı zamanda normatif bir tercihtir (OECD, 2024; European Commission, 2025).

Kurumlar model hizalamasını dış sağlayıcının görünmez tercihleri içinde hazır paket olarak kabul ederse, zamanla yalnızca teknik bağımlılık değil; norm bağımlılığı da üretir. Soru 'model güvenli mi?' değil; 'hangi değer setine göre, hangi risk toleransı, kimin adına güvenli?' sorusudur.

20. Soğuk Kaos, DeNormasyon ve Kurumsal Enformasyon Erozyonu

Bu raporun kavramsal katkılarından biri olan 'Soğuk Kaos', yapay zekâ çağında enformasyon krizinin yalnızca yanlış bilgi bolluğu olarak değil; karar alma zeminini aşındıran yapısal belirsizlik rejimi olarak okunmasını sağlar. DeNormasyon ise hata, manipülasyon, sahte kesinlik, doğrulama gevşemesi ve norm erozyonunun istisna olmaktan çıkıp sıradanlaşmasını ifade eder (Kaya, 2026; Rid, 2020; Pomerantsev, 2019).

20.1. Kavramların Sahadaki Görünümü: Soğuk Kaos ve DeNormasyon Nasıl Çalışır?

Soğuk Kaos özellikle kriz, seçim, güvenlik ve yüksek hız gerektiren karar ortamlarında; doğrulama süresinin kısaldığı, içerik hacminin arttığı, çelişkili sinyallerin çoğaldığı ve bir kısmı sentetik olan malzemenin ortak hakikat zeminini aşındırdığı yapısal belirsizlik rejimini ifade eder. DeNormasyon ise hata, manipülasyon ve doğrulama aşınmasının istisna olmaktan çıkıp yeni normal hâline gelmesini anlatır. Biri karar ortamının bozulma biçimini, diğeri bozulmanın kalıcılığını açıklar.

20.2. Mikro Vaka I: New Hampshire Yapay Sesli Robocall Olayı

2024 New Hampshire ön seçimi öncesinde seçmenlere, Başkan Joe Biden'ın sesine benzetilen yapay olarak üretilmiş bir robocall gönderildi. New Hampshire Adalet Bakanlığı Şubat 2024'te çağrılarının kaynağını tespit etti; Eylül 2024'te ise ABD Federal İletişim Komisyonu bu kampanyayı yasa dışı robocall ve arayan kimliği sahteciliği kapsamında 6 milyon dolarlık yaptırımla sonuçlandırdı. Burada belirleyici nokta, manipülasyonun teknik karmaşıklığı değil; düşük maliyetli sentetik sesin seçim gibi zaman-duyarlı bir ortamda güveni kuşkuya çevirebilmesidir (FCC, 2024; New Hampshire Department of Justice, 2024).

Bu mikro vaka, seçim güvenliğinde asıl kırılmanın yalnızca yanlış içerik yayılması olmadığını gösterir. Daha derin sorun, doğrulama yükünün seçmene ve kurumlara kaydırılmasıdır. Seçmen artık yalnızca mesajı anlamakla değil, mesajın ontolojik statüsünü çözmekle de yükümlü hâle gelir.

20.3. Mikro Vaka II: Hong Kong’da Deepfake Video Konferans Dolandırıcılığı

Kurumsal alanda DeNormasyon’un daha çıplak biçimde görüldüğü örneklerden biri, Hong Kong’daki deepfake bağlantılı video konferans dolandırıcılıklarıdır. Resmî açıklamalara göre 2024 yılında bu tür vakalar artmış; üst düzey yöneticilerin taklit edilmesi yoluyla büyük tutarlı para transferleri hedef alınmıştır. Bu olay, sentetik medya tehdidinin yalnızca kamusal tartışma veya siyasal propaganda sorunu olmadığını; şirket içi güven protokollerini, yönetici otoritesini ve görsel-ışitsel teyit alışkanlıklarını da hedef aldığını göstermektedir (Hong Kong Police Force, 2025; Government of Hong Kong, 2025).

Burada DeNormasyon, 'gözümle gördüm' ilkesinin normatif ağırlığını kaybetmesiyle çalışır. Görsel ve işitsel teyit artık tek başına yeterli sayılmadığında kurumların yalnızca siber güvenlik değil, epistemik güvenlik sorunu da yaşadığı görülür.

20.4. Makro Süreç I: Seçimlerden Bilgi Bütünlüğüne Geçiş

Makro ölçekte sorun yalnızca tekil deepfake vakaları değildir. UNESCO’nun yapay zekâ ve seçimler üzerine resmi çalışmaları, meselenin izole seçim riskleri olarak değil; daha geniş bilgi ekosistemi dönüşümü olarak anlaşılması gerektiğini vurgulamaktadır. Böylece siyasal iletişim riski, kurumsal hakikat rejimi krizine genişler. Soğuk Kaos bu ölçekte, toplumsal mutabakatı imkânsızlaştıran değil; mutabakat maliyetini sürekli artıran bir düzen olarak çalışır (UNESCO, 2024).

Avrupa Komisyonu da AI Act kapsamındaki şeffaflık yükümlülüklerini, özellikle deepfake’ler ve kamu yararını ilgilendiren sentetik içeriklerde işaretleme ve tespit edilebilirlik sağlama amacıyla iletirmektedir. Böylece sorun moderasyonun ötesine geçmiş, bilgi bütünlüğü meselesine dönüşmüştür (European Commission, 2025; 2026).

20.5. Makro Süreç II: İçerik Kökeni, Provenans ve Kurumsal Savunma Hattı

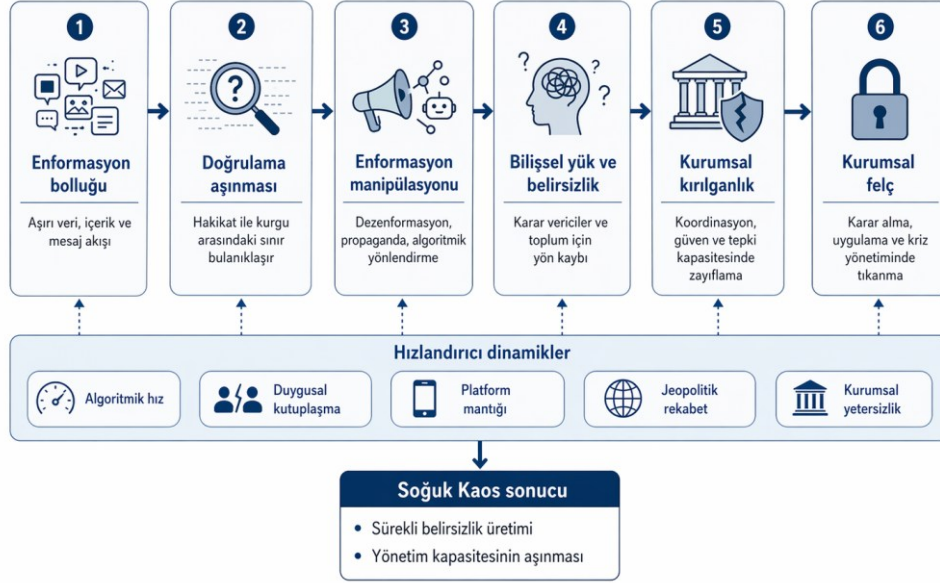
İkinci makro süreç, içerik kökeninin teknik olarak ispatlanmasına dönük savunma mimarisinin yükselişidir. C2PA tarafından geliştirilen açık standart, dijital içeriğin kaynağını ve geçirdiği düzenlemeleri gösteren doğrulanabilir içerik kimliği mantığına dayanmaktadır. Bu yaklaşım artık yalnızca medya şirketleri için değil; kamu kurumları, hukukî kayıtlar ve kriz iletişimi için de stratejik önem taşımaktadır (C2PA, 2026; CSET, 2025).

Ancak provenans standartlarının etkili olabilmesi için üretim araçları, platformlar, cihazlar ve kurumsal süreçler arasında birlikte çalışabilirlik gerekir. Daha önemlisi, toplumun ve kurumların bu işaretleri okuyacak kültürel ve operasyonel kapasiteyi geliştirmesi gerekir.

20.6. Soğuk Kaos ile DeNormasyon Arasındaki İşlevsel Fark

Bu örnekler birlikte okunduğunda kavramsal ayrım görünür hâle gelir. Soğuk Kaos, bir sistemin bilgi akışı altında boğulması değil; yüksek hacimli, düşük maliyetli, kısmen sentetik ve zaman baskısı altında dolaşıma giren içerik nedeniyle karar zemininin yavaşça kayganlaşmasıdır. DeNormasyon ise bu bozulmanın toplumsal ve kurumsal düzeyde sıradanlaşmasıdır.

Şekil 7.1. “Soğuk Kaos” modeli: enformasyon bolluğundan kurumsal felce



Kaynak: Rid (2020); Pomerantsev (2019); OECD (2026a); UNESCO & UNDP (2025); yazarın sentezi.

Yanlış Karar Riski Kutusu – Bölüm VII / 20. Başlık

Bu bölümden çıkarılması gereken temel uyarı şudur: sentetik medya ve model hizalaması risklerini yalnızca iletişim veya itibar problemi sayan her kurum, tehdidin gerçek ölçeğini kaçırır. Asıl risk, yapay zekâ destekli manipülasyonun karar zincirine girmesi, teyit maliyetini yükseltmesi ve zaman baskısı altında yanlış eşiği aşağı çekmesidir. Bu nedenle mesele yalnızca sahteyi yakalamak değil; neyin hangi prosedürle doğrulanacağına dair kurumsal refleksi önceden kurmaktır.

Bölüm Sonu Değerlendirmesi

Bu bölümün temel sonucu şudur: yapay zekâ destekli manipülasyon, tek tek yanlış içeriklerin toplamı değil; ekonomik teşvikler, platform mimarileri, bilişsel hedefleme, güvenlik açıkları ve düzenleme boşluklarıyla çalışan bütüncül bir güç alanıdır. Deepfake ekonomisi büyüdükçe yalnızca yanlış bilgi değil, doğrulama maliyeti de büyümektedir. Seçimler, medya ve kurum içi karar süreçleri aynı anda baskı altına girmekte; görünmez ideoloji, teknik hizalama ve platform mantığı üzerinden işlemekte; sonuçta ortak gerçeklik zemini aşınmaktadır. Bu nedenle yapay zekâ çağındaki esas mücadele, bilgi bolluğu içinde doğrulanabilirlik, hesap verebilirlik ve kurumsal dayanıklılık üretme mücadelesidir.

BÖLÜM VIII – DEVLET, SAVUNMA VE GELECEK SAVAŞLARI

Bu bölümün amacı, yapay zekânın kamu yönetiminden savunma planlamasına, istihbarat füzyonundan otonom silah tartışmalarına uzanan devlet-merkezli kullanım alanlarını incelemektir. Güncel kurumsal belgeler, yapay zekânın artık yalnızca idari verimlilik aracı olarak görülmediğini; karar üstünlüğü, operasyon temposu, müştereklik ve caydırıcılık tartışmalarının da merkezine yerleştiğini göstermektedir (U.S. Department of Defense [DoD], 2022; NATO, 2024).

Bu dönüşüm, iki nedenle stratejik önem taşımaktadır. Birincisi, yapay zekâ devlet kapasitesini yatay olarak genişletmekte; yani lojistik, tedarik, bakım, sağlık, personel yönetimi ve siber savunma gibi alanlarda kurumsal ölçek etkisi yaratmaktadır. İkincisi, aynı teknoloji dikey olarak askerî karar zincirine girmekte; hedef tanıma, istihbarat önceliklendirme, harekât planlama ve kuvvet kullanımına ilişkin değerlendirmeleri hızlandırmaktadır. Böylece yapay zekâ, klasik “destekleyici yazılım” rolünün ötesine geçerek devletin zor kullanma kapasitesine dokunan bir teknolojiye dönüşmektedir (DoD, 2023; CDAO, 2025a).

21. Askerî AI Entegrasyonu

ABD Savunma Bakanlığı’nın Sorumlu Yapay Zekâ Stratejisi ve Uygulama Yol Haritası, yapay zekâyı doğrudan savunma misyonunun bir parçası olarak tanımlamakta ve kurumun bunu geliştirme, tedarik etme, konuşlandırma ve kullanma biçimini bütüncül bir çerçeveye bağlamaktadır (DoD, 2022). Bu yaklaşım, yapay zekânın yalnızca araştırma laboratuvarlarında değil, kuvvet planlaması ve operasyonel süreçlerde de kurumsallaştığını göstermektedir. Nitekim CDAO’nun resmî çerçevesi, veri ve yapay zekâ benimsenmesini “yönetim kurulundan muharebe alanına” uzanan bir dönüşüm olarak tanımlamaktadır (CDAO, 2025b).

Bu kurumsallaşma, yalnızca devlet içi yapılanmayla sınırlı değildir. 2025 yılında CDAO’nun OpenAI, Anthropic, Google ve xAI ile ilan ettiği ortaklıklar; frontier modellerin doğrudan ulusal güvenlik görev alanları için denendiğini göstermiştir (CDAO, 2025a). Aynı dönemde OpenAI, DoD için 200 milyon dolar tavanlı bir sözleşme çerçevesi içinde sağlık, tedarik verisi ve siber savunma gibi idarî ve operasyonel kullanım alanlarında prototipleme desteği verdiğini duyurmuştur (OpenAI, 2025). Anthropic ise Claude Gov modellerini, yalnızca ABD ulusal güvenlik müşterilerine açık olacak şekilde tanıtmış ve bu modellerin sınıflandırılmış ortamlarda kullanılabildiğini açıklamıştır (Anthropic, 2025a). Bu gelişmeler, savunma AI ekosisteminde ticari frontier laboratuvarlar ile devlet kurumları arasındaki sınırın giderek geçirgenleştiğini göstermektedir.

NATO düzleminde ise mesele sadece kabiliyet değil, müştereklik ve sorumlu kullanım sorunudur. NATO’nun 2021 Yapay Zekâ Stratejisi ile 2024’te revize edilen çerçevesi; hukuka uygunluk, sorumluluk ve hesap verebilirlik, açıklanabilirlik ve izlenebilirlik, güvenilirlik, yönetilebilirlik ve önyargı azaltımı ilkelerini savunma AI’sinin temel eksenleri olarak tanımlamaktadır (NATO, 2021, 2024). Bu, İttifak’ın teknolojik üstünlük hedefini etik ve hukuki sınırlarla birlikte düşünmeye çalıştığını göstermektedir. Ancak aynı belgeler, yapay zekânın rakipler tarafından kötüye kullanımının ve bilgi operasyonları üzerindeki etkisinin de artan bir tehdit alanı olduğunu kabul etmektedir (NATO, 2024).

Tablo 8.1. Savunma amaçlı yapay zekâ kullanım alanları ve temel risk profilleri

Kullanım alanı	Başlıca işlev	Operasyonel katkı	Temel risk	İnsan denetimi gereği
Lojistik ve bakım	Talep tahmini, kestirimci bakım	Maliyet düşüşü ve tempo artışı	Yanlış veriyle ölçeklenen hata	Orta
Siber savunma	Anomali tespiti, tehdit korelasyonu	Hızlı tespit ve önceliklendirme	Yanlış pozitifler ve otomasyon bağımlılığı	Yüksek
İstihbarat analizi	Görüntü / metin / sinyal birleştirme	Analitik kapasite artışı	Bağlam kaybı ve teyit maliyeti	Çok yüksek
Karar desteği	Senaryo üretimi, harekât önerisi	Karar temposunun hızlanması	Otomasyon önyargısı ve kapanmayan hatalar	Çok yüksek
Hedef seçimi ve kuvvet kullanımı	Tespit, izleme, önceliklendirme	Zaman baskısında reaksiyon	Yanlış hedef, tırmanma ve hesap verebilirlik krizi	Azami

Kaynak: DoD (2022, 2023); NATO (2024); ICRC (2021, 2026); yazarın sentezi.

22. İstihbarat ve Karar Destek Sistemleri

İstihbarat topluluğu açısından yapay zekânın en önemli vaadi, veri bolluğunu yönetilebilir kararlara dönüştürme kapasitesidir. Uydu görüntüleri, insansız hava aracı akışları, açık kaynak içerikler, haber akışları, sensör verileri ve siber telemetri gibi çok farklı kaynaklardan gelen bilgilerin aynı anda işlenmesi, insan analistlerin tek başına yürütebileceği bir faaliyet olmaktan çıkmıştır. Bu nedenle AI, istihbarat dünyasında giderek daha fazla veri füzyonu, örüntü tespiti, hedef önceliklendirme ve erken uyarı sistemi olarak kullanılmaktadır (DoD, 2022; CDAO, 2025b).

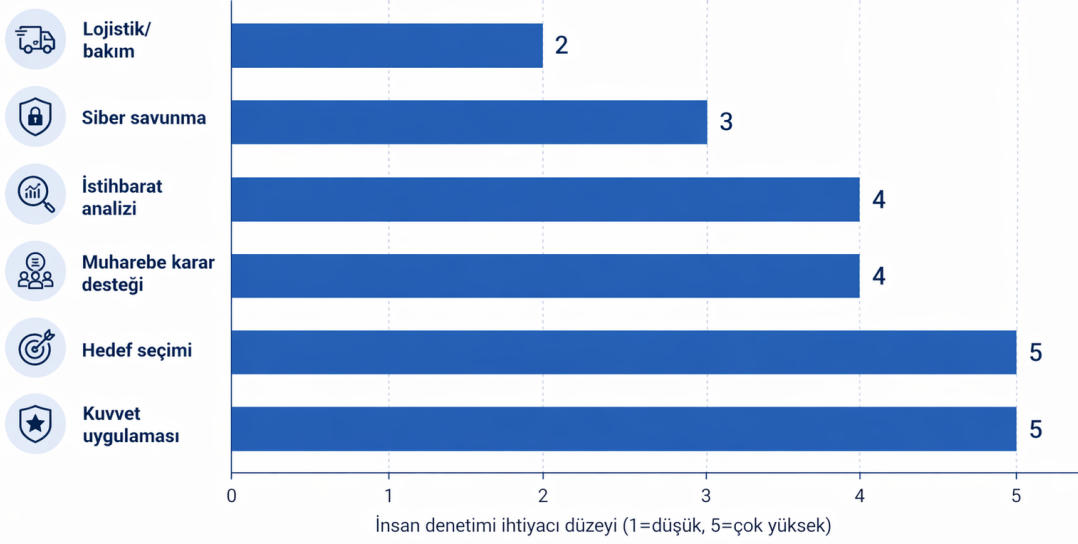
Ancak burada ortaya çıkan temel risk, hız ile doğruluk arasındaki dengenin bozulmasıdır. UNIDIR'nin 2025 tarihli değerlendirmesi ile ICRC'nin güncel pozisyon metinleri, otonom ve yarı otonom sistemlerin yalnızca hukuki değil, bilişsel risk de ürettiğini vurgulamaktadır. Özellikle yüksek yoğunluklu veri ortamlarında analistler, model çıktısını tartışmasız veri gibi görmeye başlayabilir; bu durum otomasyon önyargısı, teyit zincirinin zayıflaması ve yanlış hedefe yönelme riskini artırır (UNIDIR, 2025; ICRC, 2026).

DoD'nin 2025 Yapay Zekâ Siber Güvenlik Risk Yönetimi Rehberi de bu nedenle model katmanı ile altyapı katmanını birlikte ele almakta; veri bütünlüğü, model güvenliği, tedarik zinciri, anti-tamper önlemleri ve siber dayanıklılığı aynı güvenlik çerçevesi içinde değerlendirmektedir (DoD CIO, 2025). Başka bir ifadeyle, geleceğin istihbarat sistemi yalnızca daha hızlı değil, aynı zamanda daha kırılgan da olabilir. Veri zehirlenme, yanıltılmış girişler, sentetik içerikle doyunlaşmış açık kaynak ekosistemleri ve tedarik zinciri açıkları, karar destek sistemlerini rakipler için yeni bir saldırı yüzeyine dönüştürmektedir (DoD CIO, 2025; NATO, 2024).

Grafik 8.1. Savunma amaçlı yapay zekâ kullanım alanlarında insan denetimi ihtiyacı

İnsan denetimi ihtiyacı düzeyleri

1 = düşük ihtiyaç | 5 = çok yüksek ihtiyaç



Harita 8.1. Küresel savunma AI yaklaşım haritası



23. Otonom Silah Tartışmaları

Otonom silah sistemleri etrafındaki tartışma, yapay zekâ çağının belki de en yüksek riskli hukuk-siyaset alanını oluşturmaktadır. ABD Savunma Bakanlığı'nın 3000.09 sayılı 2023 tarihli Direktifi, otonom ve yarı otonom silah sistemlerinin tasarım ve kullanımında komutanlar ile operatörlerin

kuvvet kullanımında “uygun düzeyde insan muhakemesi” kullanabilmesini şart koşmaktadır (DoD, 2023). Bu ifade, ABD yaklaşımının tam yasak yerine kontrollü entegrasyon çizgisinde kaldığını göstermektedir.

Buna karşılık ICRC, otonom silahların acil insani kaygı doğurduğunu ve uluslararası düzeyde yeni kurallara ihtiyaç bulunduğunu uzun süredir savunmaktadır. ICRC’nin 2021 pozisyonu ile 2026 tarihli IHL değerlendirmesi, hedef seçimi ve kuvvet uygulamasında insan kontrolünün zayıflamasının sivillerin korunması, ayırım, orantılılık ve saldırı öncesi ihtiyat yükümlülükleri bakımından ciddi sorunlar doğurduğunu açık biçimde ortaya koymaktadır (ICRC, 2021, 2026).

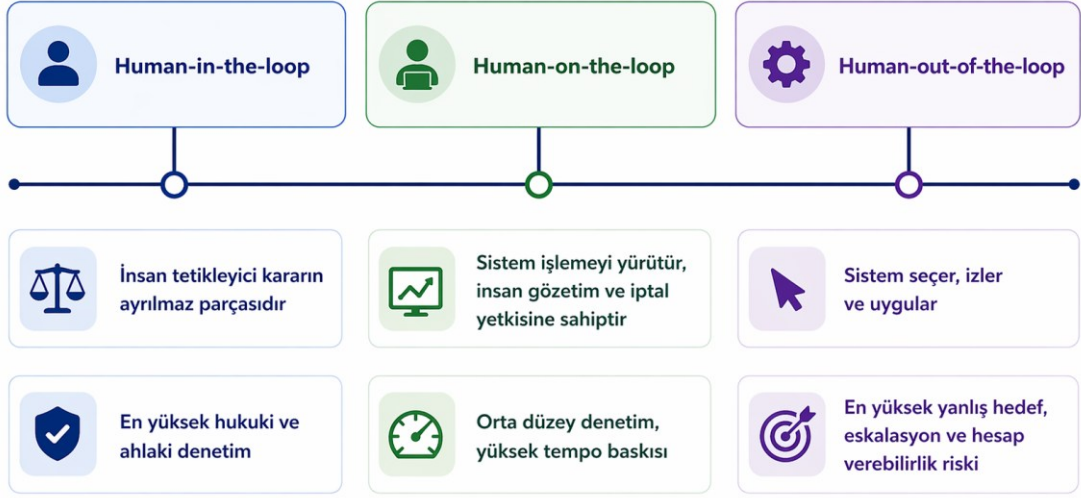
Birleşmiş Milletler Genel Sekreteri’nin, Genel Kurulun 78/241 sayılı kararı uyarınca hazırlanan 2024 tarihli raporu da ölümcül otonom silah sistemlerine ilişkin devlet görüşlerini toplamış ve bu alandaki uluslararası düzenleme ihtiyacının merkezîleştiğini göstermiştir (United Nations, 2024). 2025 yılına gelindiğinde Güvenlik Konseyindeki tartışmalar da, insan kontrolü olmaksızın işleyen ölümcül otonom silahlara yönelik bağlayıcı bir çerçeve çağrılarının güçlendiğini göstermektedir (United Nations Security Council, 2025).

Buradaki temel mesele, yalnızca bir silahın teknik kabiliyeti değildir. Asıl mesele, karar temposunun insan hukuku ve siyaseti tarafından yetiilemeyecek kadar hızlanmasıdır. Yapay zekâ destekli savaş döngüsü kısaldıkça, yanlış hedef tespiti, dost-düşman karışıklığı, tırmanma riski ve hesap verebilirlik boşluğu büyümektedir. Bu nedenle gelecek savaşları tartışılırken soru yalnızca “yapay zekâ ne kadar etkili olacak?” değil; “hangi karar noktalarında insan kalacak, hukuki sorumluluk nasıl dağıtılacak ve hata durumunda kim hesap verecek?” olmalıdır (UNIDIR, 2025; ICRC, 2026).

Şema 8.1. Sensörden etkiye uzanan savunma AI döngüsü



Şekil 8.1. Human-in / on / out-of-the-loop sürekliliği



Bölüm Sonu Değerlendirmesi

Bu bölümün temel bulgusu şudur: yapay zekâ devlet ve savunma alanına çoktan girmiştir; tartışma artık “girip girmeyeceği” değil, hangi hızla, hangi sınırlarla ve hangi denetim mekanizmalarıyla kurumsallaşacağı üzerinedir. Savunma AI’sinin lojistik, siber güvenlik ve istihbarat gibi alanlarda sunduğu avantajlar açıktır; ancak aynı ekosistem hedef seçimi, otomasyon önyargısı, siber kırılganlık ve hukuki sorumluluk boşluğu gibi çok daha yüksek maliyetli riskler üretmektedir.

Bu nedenle geleceğin savaşları hakkındaki asıl ayrım teknolojiye sahip olanlar ile olmayanlar arasında değil; teknolojiyi hukuk, insan kontrolü ve stratejik ihtiyatla yönetebilenler ile yönetemeyenler arasında belirecektir. Devlet kapasitesinin yapay zekâ ile genişlemesi, aynı zamanda hata kapasitesinin de büyümesi anlamına gelmektedir. Savunma alanında gerçek üstünlük, yalnızca daha hızlı algoritmalara değil; daha güvenilir doğrulama zincirlerine, daha sıkı insan denetime ve daha net hesap verebilirlik mimarilerine dayanacaktır.

BÖLÜM IX – HUKUK, ETİK VE YÖNETİŞİM

Bu bölümün amacı, yapay zekâ alanında ortaya çıkan normatif çoğulluğu, bağlayıcı hukuk ile gönüllü standartlar arasındaki farkı ve etik ilke repertuarı ile fiilî uygulama arasındaki açığı birlikte değerlendirmektir. Yapay zekâ yönetişimi artık yalnızca teknik emniyet konusu değildir; insan hakları, demokrasi, hukuk devleti, kurumsal sorumluluk, piyasa yapısı ve sınır ötesi etki zincirleri ile doğrudan bağlantılı bir düzenleme alanına dönüşmüştür (OECD, 2019; UNESCO, 2021; Council of Europe, 2024). Bu nedenle güncel tartışma, “düzenleyelim mi?” sorusunu büyük ölçüde aşmış; “hangi risk mantığıyla, hangi ölçeklerde ve hangi uygulama araçlarıyla düzenleyeceğiz?” sorusuna evrilmiştir (European Commission, 2025a; OECD, 2026).

Bugünkü tablo, bir yandan ilkeler ve çerçeveler bakımından yoğun bir üretim; diğer yandan denetim, uygulama ve telafî mekanizmaları bakımından parçalı bir görünüm sunmaktadır. OECD AI İlkeleri güvenilir ve insan haklarına saygılı yapay zekâ için esnek ama ortak bir norm zemini kurarken, UNESCO’nun 2021 Tavsiye Kararı etik tartışmayı küresel bir kamu normuna dönüştürmüş, NIST ise risk yönetimini somut süreç ve kontrol mantığına bağlamıştır (OECD, 2019; UNESCO, 2021; NIST, 2023). Avrupa Birliği ise bu eğilimleri bağlayıcı hukuk düzlemine taşıyarak risk temelli ve kategoriye dayalı bir düzenleme rejimi kurmuştur (European Commission, 2025a).

24. Küresel düzenleme girişimleri

Küresel düzeyde ilk büyük normatif sıçrama, 2019 tarihli OECD AI İlkeleri ile gerçekleşmiştir. Bu ilkeler, yapay zekânın yenilikçi olmasını teşvik ederken aynı zamanda insan haklarına, demokratik değerlere, şeffaflığa, sağlamlığa ve hesap verebilirliğe saygı göstermesi gerektiğini vurgulamaktadır (OECD, 2019). OECD’nin 2026 tarihli Responsible AI Due Diligence rehberi ise bu ilke setini daha operasyonel bir düzeye taşıyarak, şirketlerin ve değer zinciri aktörlerinin olumsuz etkileri önceden tespit etmesi, azaltması ve gerektiğinde telafî mekanizmaları kurması gerektiğini belirtmektedir (OECD, 2026). Böylece OECD hattı, soyut ilke söyleminden süreç temelli kurumsal sorumluluğa doğru genişlemektedir.

UNESCO’nun 2021 tarihli Yapay Zekâ Etiği Tavsiye Kararı, bu alandaki ilk küresel standart niteliğini taşımaktadır ve 194 üye devlet bakımından ortak etik referans üretmektedir. UNESCO’nun insan hakları, insan onuru, adalet, mahremiyet, çeşitlilik ve insan denetimi vurgusu; özellikle yalnızca güvenlik ve inovasyon eksenine sıkışan düzenleme tartışmalarına normatif bir denge unsuru eklemektedir (UNESCO, 2021). UNESCO’nun 2023 tarihli foundation models politika notu da, büyük ölçekli modellerin etkilerinin yalnızca model güvenliği bağlamında değil; güç yoğunlaşması, bilgi asimetrisi ve toplumsal zararlar bağlamında değerlendirilmesi gerektiğini ortaya koymaktadır (UNESCO, 2023).

ABD tarafında NIST’in AI RMF 1.0 çerçevesi ve 2024 tarihli Generative AI Profile belgesi, hukuken bağlayıcı bir federal AI yasasından ziyade, risk yönetimi, test, belgelendirme, yaşam döngüsü denetimi ve güvenilirlik kriterleri üzerinden ilerleyen bir standartlaştırma mantığı geliştirmiştir (NIST, 2023, 2024). NIST’in 2025 tarihli adversarial machine learning taksonomisi de güvenlik risklerini daha teknik ve operasyonel bir dilde sınıflandırarak, yönetişimin siber dayanıklılık boyutunu güçlendirmiştir (NIST, 2025). Bu durum, standart temelli yaklaşımın özellikle teknoloji geliştirme ve kurumsal uyum tarafında etkili, ancak bağlayıcı hukuk düzeyinde daha sınırlı olduğunu göstermektedir.

Avrupa Birliği’nde ise yönetişim açık biçimde bağlayıcı hukuk biçimini almıştır. Avrupa Komisyonu’nun uygulama takvimine göre AI Act 1 Ağustos 2024’te yürürlüğe girmiş, yasaklı

uygulamalar ve AI literacy yükümlülükleri 2 Şubat 2025'te uygulanmaya başlamış, genel amaçlı yapay zekâ modellerine ilişkin yükümlülükler ise 2 Ağustos 2025 itibarıyla devreye girmiştir; yüksek riskli bazı sistemler için geçiş süreleri 2026 ve 2027'ye yayılmaktadır (European Commission, 2025a). Aynı hat üzerinde yayımlanan GPAI Code of Practice ve Komisyon kılavuzları, özellikle genel amaçlı ve sistemik risk taşıyan model sağlayıcılarının uyum, belgelendirme, özet veri açıklamaları ve risk yönetimi konularında nasıl hareket etmesi gerektiğini somutlaştırmaktadır (European Commission, 2025b, 2025c).

Tablo 9.1. Başlıca yapay zekâ yönetim çerçevelerinin karşılaştırılması

Çerçeve / kurum	Niteliği	Temel eksen	Uygulama mantığı	Başlıca sınırlılık
OECD AI İlkeleri + Due Diligence	İlke + süreç rehberi	Güvenilirlik, sorumlu iş yürütme, olumsuz etki önleme	Kurumsal süreç ve değer zinciri sorumluluğu	Bağlayıcı yaptırım gücü sınırlıdır
UNESCO Etik Tavsiye Kararı	Küresel etik standart	İnsan hakları, insan onuru, adalet, insan denetimi	Üye devletlere normatif yön verme	Uygulama düzeyi ülkelere göre değişir
NIST AI RMF / GAI Profili	Teknik-risk çerçevesi	Risk yönetimi, test, güvenilirlik, yaşam döngüsü	Kurumsal uygulama ve ölçülebilir kontrol mantığı	Hukuken bağlayıcı değildir
AB AI Act + GPAI rejimi	Bağlayıcı hukuk	Risk sınıflandırması, yasaklı uygulamalar, uyum yükümlülükleri	Kategorik yükümlülükler ve gözetim	Uygulama yükü ve yorum karmaşıklığı yüksektir
Avrupa Konseyi Çerçeve Sözleşmesi	Bağlayıcı uluslararası sözleşme	İnsan hakları, demokrasi, hukuk devleti	Yaşam döngüsü boyunca kamusal yükümlülük	Uygulama kapasitesi ve iç hukuka aktarım belirleyicidir

Kaynak: OECD (2019, 2026); UNESCO (2021, 2023); NIST (2023, 2024); European Commission (2025a, 2025b, 2025c); Council of Europe (2024); yazarın sentezi.

25. Ulusal yaklaşımlar

Ulusal düzeydeki yaklaşımlar arasında belirgin bir ayrışma bulunmaktadır. Avrupa Birliği, temel hakları ve sistemik riskleri merkezine alan, sınıflandırma ve uyum yükümlülükleri üzerinden ilerleyen bir yaklaşım benimserken; Birleşik Krallık daha esnek ve “pro-innovation” bir yönetim hattını savunmaktadır (European Commission, 2025a; UK Government, 2023). Birleşik Krallık'ın 2025 tarihli AI Opportunities Action Plan belgesi, düzenleyici çevikliği ve büyüme odaklılığı öne çıkarırken; 2026 tarihli “One Year On” güncellemesi de uygulamanın sanayi ve kamu hizmetleri ekseninde hızlandırılmak istendiğini göstermektedir (UK Government, 2025a, 2026).

Çin ise daha açık biçimde devlet-merkezli ve güvenlik odaklı bir çerçeve izlemektedir. 2023 tarihli Geçici Generative AI Önlemleri, bir yandan yapay zekâ gelişimini teşvik ederken diğer yandan ulusal güvenlik, kamu düzeni ve kamusal çıkarın korunmasını temel öncelik olarak tanımlamaktadır (State Council of the People's Republic of China, 2023). Bu durum, Çin modelinde yönetişimin yalnızca piyasa düzenlemesi değil; içerik, toplumsal istikrar ve devlet kapasitesiyle iç içe geçmiş bir idari müdahale alanı olduğunu göstermektedir.

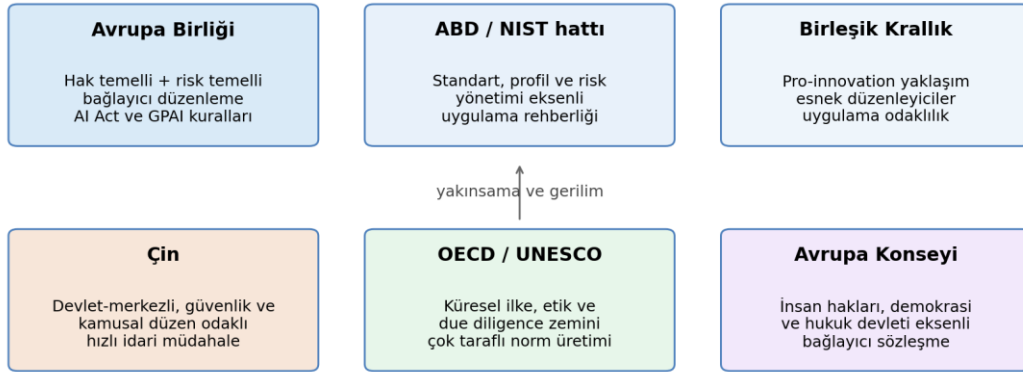
Bu farklılıklar, küresel yapay zekâ yönetişiminin tek bir evrensel model etrafında birleşmediğini göstermektedir. Avrupa Birliği bağlayıcı ve ayrıntılı uyum rejimine, ABD/NIST hattı standart ve süreç mantığına, Birleşik Krallık esnek ve büyüme odaklı bir çerçeveye, Çin ise güvenlik ve kamu otoritesi öncelikli bir düzene yaslanmaktadır (NIST, 2023; European Commission, 2025a; UK Government, 2023; State Council of the People's Republic of China, 2023). Sonuç olarak küresel norm üretimi artsa da uygulama mantığı bölgesel siyasal ekonomilere göre ayrıışmaktadır.

Grafik 9.1. Küresel yapay zekâ yönetişiminin düzenleyici zaman çizelgesi



Harita 9.1. Küresel yapay zekâ yönetim rejimleri: kavramsal konumlanma

Harita 9.1. Küresel yapay zekâ yönetim rejimleri: kavramsal konumlanma



Ortak sorun: ilke bolluğu, uygulama ve denetim açığı

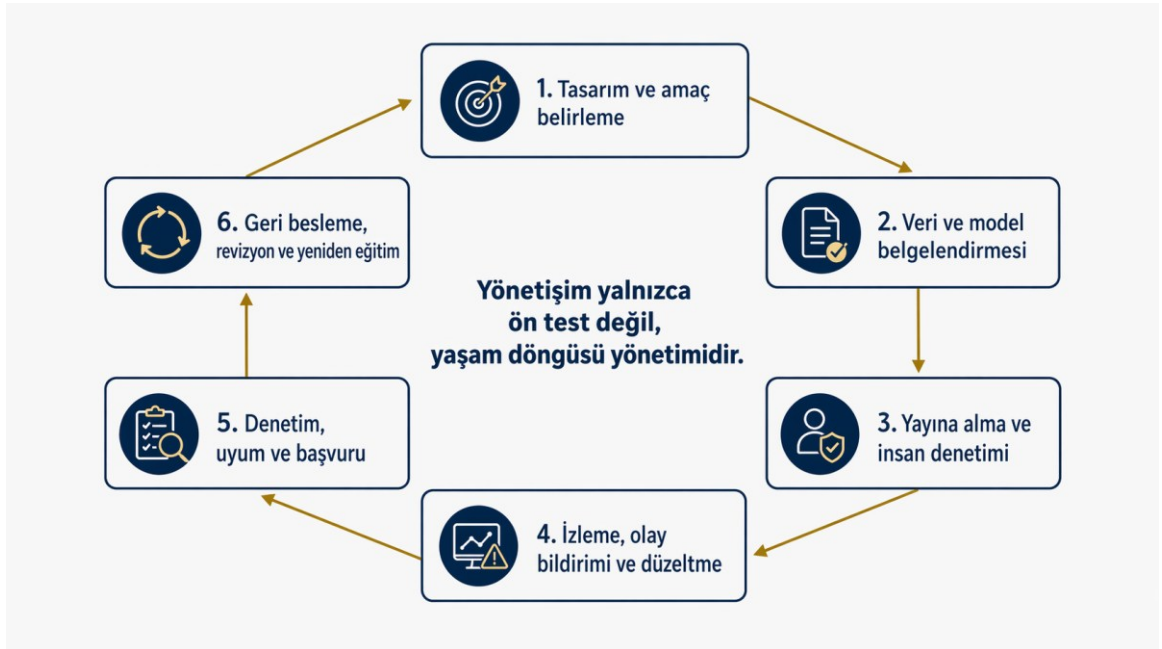
26. Etik ilkeler ve fiilî pratik arasındaki uçurum

Bugünkü tartışmanın en temel sorunlarından biri, etik ilke üretiminin hızının uygulama kapasitesini aşmasıdır. Şeffaflık, hesap verebilirlik, insan denetimi, güvenlik ve adalet neredeyse bütün çerçevelerde tekrar edilen ortak normlar hâline gelmiştir; fakat bu normların nasıl ölçüleceği, hangi eşiklerde zorunlu tutulacağı ve ihlâl hâlinde hangi telafi mekanizmalarının işletileceği konusunda hâlâ ciddi boşluklar bulunmaktadır (UNESCO, 2021; OECD, 2026). Özellikle frontier ve genel amaçlı modeller söz konusu olduğunda, belge bolluğu ile denetim derinliği arasında belirgin bir asimetri ortaya çıkmaktadır.

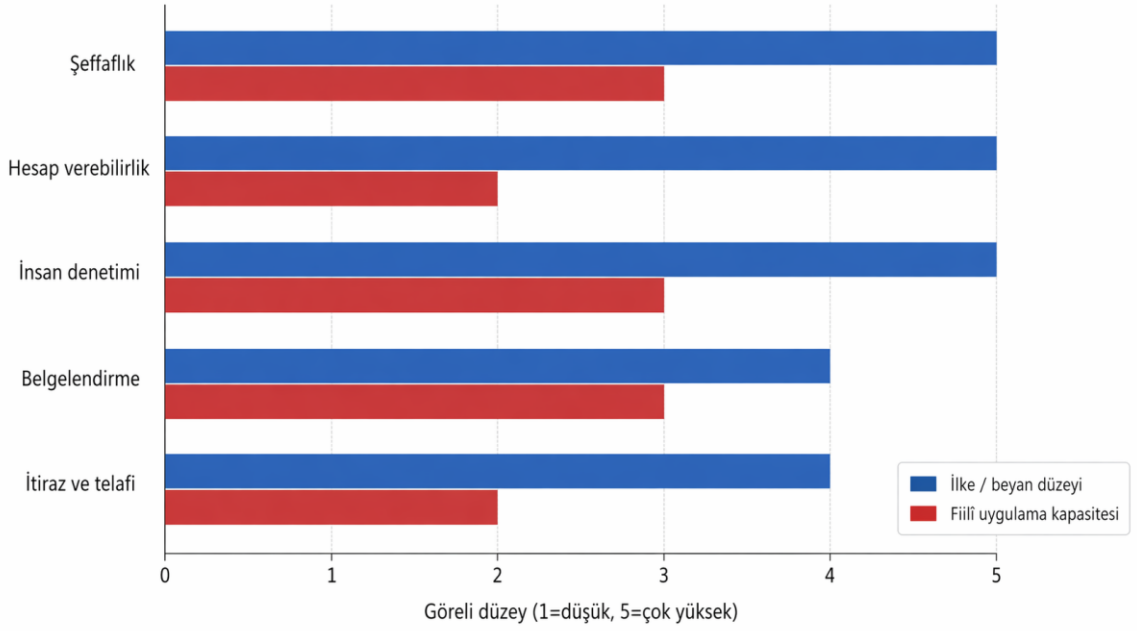
Bu açığın bir boyutu da doğrulama ve köken meselesidir. Sentetik içerik işaretleme ve provenance araçları yönetim açısından önemlidir; ancak bunlar doğruluğun kendisini garanti etmez. Bu nedenle hukuk ve etik çerçeveler yalnızca etiketleme, açıklama ve kullanım koşullarına değil; veri kaynağı, model davranışı, değerlendirme, olay bildirimi, itiraz hakkı ve kurumsal telafi mekanizmalarına da uzanmak zorundadır (OECD, 2026; NIST, 2025). Aksi takdirde yönetim, içerik üzerine yapılandırılan açıklama metinlerinden ibaret kalır.

Etik ilkeler ile fiilî pratik arasındaki boşluk, özellikle kamu gücü ve büyük platform kapasitesi bir araya geldiğinde daha da görünür hâle gelmektedir. İnsan denetimi kağıt üzerinde korunuyor görünse bile, yüksek hız, yüksek hacim ve kurumsal otomasyon baskısı altında karar süreçleri giderek daha fazla makine tavsiyesine bağımlı hâle gelebilmektedir. Bu nedenle yönetişimin gerçek sınavı, ilke metinleri yazmak değil; yaşam döngüsü boyunca denetim, geriye dönük iz sürme, olay bildirimi ve düzeltici müdahale kapasitesini kurmaktır (NIST, 2024; OECD, 2026; Council of Europe, 2024).

Şema 9.1. Risk temelli yapay zekâ yönetim döngüsü



Şekil 9.1. Etik ilkeler ile fiili pratik arasındaki yönetim açığı



Bölüm Sonu Değerlendirmesi

Bu bölümün temel bulgusu şudur: yapay zekâ yönetişimi artık ilke üretimi aşamasını geçmiş, çok katmanlı fakat parçalı bir düzenleme evresine girmiştir. OECD ve UNESCO küresel norm zemini üretmekte; NIST risk yönetimi ve teknik uygulanabilirlik dili sağlamaktadır; Avrupa Birliği bağlayıcı hukuk rejimi inşa etmekte; Avrupa Konseyi ise insan hakları ve demokrasi eksenli bir uluslararası sözleşme zemini kurmaktadır (OECD, 2019, 2026; UNESCO, 2021; Council of Europe, 2024; European Commission, 2025a).

Bununla birlikte asıl sorun, norm eksikliği değil uygulama asimetrisidir. İlkeler çoğalırken denetim, başvuru, telafi ve sınır ötesi koordinasyon kapasitesi aynı hızla gelişmemektedir. Bu nedenle yapay zekâ çağında yönetim tartışmasının merkezinde artık “hangi ilke?” değil, “hangi kurum, hangi ölçüm, hangi yaptırım ve hangi yaşam döngüsü denetimi?” soruları yer almaktadır (NIST, 2023, 2024, 2025; OECD, 2026).

BÖLÜM X – VAKA ANALİZLERİ

Bu bölümün amacı, yapay zekâ alanında sıkça atıf yapılan kurumsal ve düzenleyici vakaları tek tek örneklemekten çok, farklı vakalar üzerinden tekrar eden risk örüntülerini görünür hâle getirmektir. Burada seçilen vakalar; veri sızıntısı, kamu ve savunma entegrasyonu, platform temelli manipülasyon, şeffaflık ve uyum baskısı gibi alanlarda temsil gücü taşıyan örneklerden oluşturulmuştur. Vaka yaklaşımı, yapay zekâ sistemlerinin soyut teknik mimariler olmaktan çıkıp somut kurumsal, hukuki ve jeopolitik sonuçlar üreten yapılar hâline geldiğini göstermektedir (NIST, 2023; OECD, 2026; Maslej et al., 2026).

Bu nedenle bu bölüm, bir olay kronolojisi değil; daha çok risk kümelerini bir araya getiren analitik bir vaka atlası olarak okunmalıdır. OpenAI, GitHub Copilot, DeepSeek, Anthropic/Palantir/AWS ve xAI/Grok örnekleri birlikte ele alındığında; aynı ekosistemin farklı yüzlerinde veri yönetişi, devlet kapasitesi, uyum baskısı ve bilgi bütünlüğü sorunlarının nasıl kesiştiği açık biçimde görülmektedir (OpenAI, 2025a, 2025b; GitHub, 2026a; Wiz Research, 2025; European Commission, 2026a).

27. Şirket ve sistem bazlı vaka seti

İlk vaka kümesi OpenAI etrafında şekillenmektedir. OpenAI, 2025 yılında önce ChatGPT Gov ürününü, ardından OpenAI for Government girişimini duyurarak ABD kamu kurumlarıyla daha sistematik bir çalışma hattı kurmuştur. Şirket aynı dönemde ABD Ulusal Laboratuvarları ile iş birliğine gitmiş ve kamu görevleri için güvenli, kurumsal ve ölçekli kullanım söylemini öne çıkarmıştır. Bununla birlikte OpenAI vakası yalnızca kamu entegrasyonunu değil, aynı zamanda veri ve güvenlik boyutunu da içerir; şirketin 2023 tarihli resmî açıklaması, bir yazılım hatasının bazı kullanıcıların başka kullanıcılara ait sohbet başlıklarını ve sınırlı ödeme bilgilerini görebilmesine yol açtığını kabul etmiştir. Dolayısıyla OpenAI örneği, yüksek kurumsal kapasite ile yüksek görünürlük arasındaki ters orantıyı da göstermektedir: sistem büyüdükçe etkisi artmakta, hata anında maruz kaldığı yönetim baskısı da büyümektedir (OpenAI, 2023, 2025a, 2025b, 2025c).

İkinci vaka GitHub Copilot etrafında şekillenmektedir. GitHub, 25 Mart 2026'da yayımladığı politika güncellemesinde Copilot Free, Pro ve Pro+ kullanıcılarına ait etkileşim verilerinin 24 Nisan 2026'dan itibaren, kullanıcılar açıkça devre dışı bırakmadıkça, modeli geliştirmek amacıyla kullanılacağını duyurmuştur. Aynı açıklamada Copilot Business ve Enterprise müşterilerinin bu değişiklikten etkilenmediği belirtilmiştir. Bu ayırım, aynı ürün ailesi içinde bile veri yönetişimi rejimlerinin kullanıcı tipine göre değişebildiğini göstermesi bakımından önemlidir. Copilot vakası böylece kurumsal güven ile bireysel kullanıcı verisi arasında kurulan farklılaştırılmış sözleşme mantığını görünür kılmaktadır (GitHub, 2026a, 2026b).

Üçüncü vaka DeepSeek'tir. DeepSeek'in 2025 ve 2026 tarihli gizlilik politikaları, kullanıcı verilerinin Çin Halk Cumhuriyeti'nde işlendiğini ve saklandığını açık biçimde belirtmektedir. Bu veri rejimi, kendi başına bir olay olmasa bile, egemenlik ve mahremiyet tartışmalarının merkezinde yer almaktadır. Buna ek olarak Wiz Research, Ocak 2025'te kamuya açık bırakılmış bir ClickHouse veritabanı tespit etmiş; bu maruziyetin sohbet geçmişleri, gizli anahtarlar ve operasyonel ayrıntılar dâhil bir milyondan fazla günlük kaydını açığa çıkardığını bildirmiştir. DeepSeek vakası, hızlı ölçeklenen sistemlerin güvenlik hijyeni, sınır ötesi veri yönetişimi ve güven ilişkisi üzerindeki etkilerini aynı anda gösteren en çarpıcı örneklerden biridir (DeepSeek, 2025, 2026; Wiz Research, 2025).

Dördüncü vaka Anthropic/Palantir/AWS hattıdır. Palantir'in Kasım 2024 tarihli açıklaması ve Anthropic'in 2025 duyuruları, Claude modellerinin Palantir AIP ve güvenli AWS altyapıları üzerinden ABD savunma ve istihbarat müşterilerine sunulduğunu göstermektedir. Anthropic ayrıca Haziran 2025'te Claude Gov modellerini ulusal güvenlik müşterileri için duyurmuş; Ağustos 2025'te de bu erişimin devletin üç ana koluna kadar genişlediğini belirtmiştir. Bu vaka, “güvenli ve sorumlu yapay zekâ” söylemi ile savunma ve istihbarat kullanımının aynı kurumsal çerçevede nasıl birleştiğini ortaya koymaktadır (Palantir, 2024; Anthropic, 2025a, 2025b, 2025c).

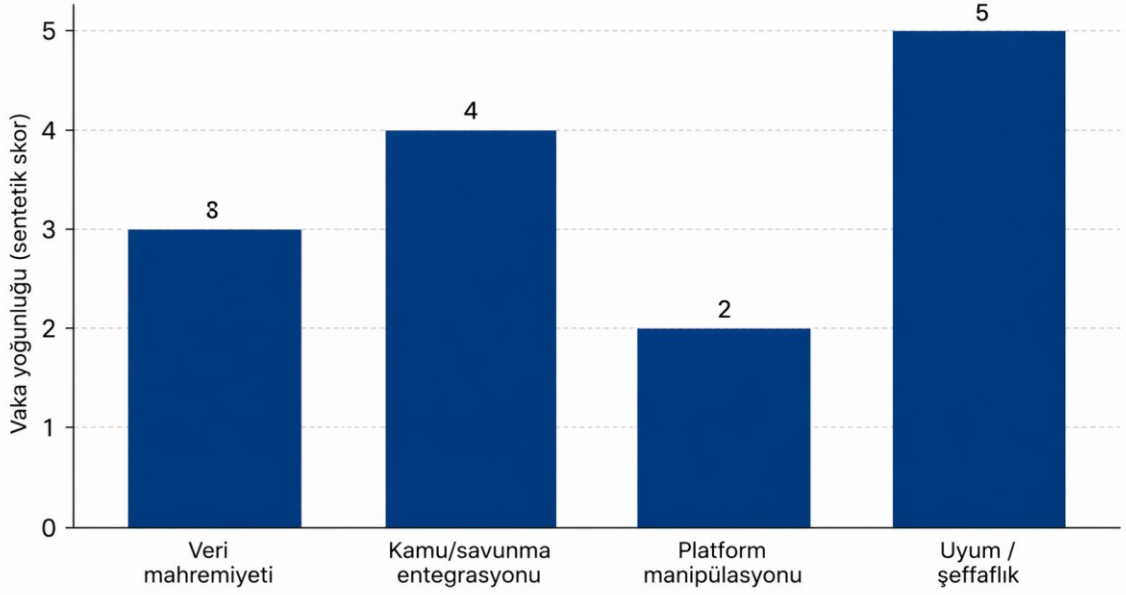
Beşinci vaka xAI/Grok ile ilişkilidir. Avrupa Komisyonu 25 Ocak 2026 tarihinde, X platformunda Grok özelliklerinin devreye alınması ve tavsiye sistemleriyle bağlantılı sistemik risklerin değerlendirilmesi kapsamında yeni bir resmî soruşturma başlatmıştır. Bu vaka, üretken yapay zekâ sistemlerinin tek başına değil; geniş erişimli sosyal platformlar, öneri sistemleri ve politik içerik akışlarıyla birleştiğinde ne kadar daha yüksek toplumsal etki yaratabildiğini göstermektedir. Böylece xAI/Grok örneği, model riski ile platform riskinin birbirinden ayrı düşünölmeyeceğini ortaya koymaktadır (European Commission, 2026a).

Tablo 10.1. Seçilmiş yapay zekâ vakalarının karşılaştırmalı risk profile

Kaynak: OpenAI (2023, 2025a, 2025b, 2025c); GitHub (2026a, 2026b); DeepSeek (2025, 2026); Wiz Research (2025); Palantir (2024); Anthropic (2025a, 2025b, 2025c); European Commission (2026a); yazarın sentezi.

Vaka	Ana olay	Öne çıkan risk	Devlet/kamu bağı	Düzenleyici baskı	Temel kaynak
OpenAI	ChatGPT Gov ve OpenAI for Government; 2023 sohbet başlığı/ödeme verisi olayı	Veri güvenliği + kamu entegrasyonu	Yüksek	Yüksek	OpenAI
GitHub Copilot	2026 veri kullanım politikası güncellemesi	Veri yönetiřimi + kullanıcı ayrışması	Orta	Orta	GitHub
DeepSeek	Çin’de veri saklama; Wiz tarafından sızıntı tespiti	Mahremiyet + sızıntı + egemenlik	Dolaylı/Yüksek	Yüksek	DeepSeek; Wiz
Anthropic/Palantir/AWS	Claude Gov ve savunma/istihbarat entegrasyonu	Savunma kullanımı + classified workflows	Çok yüksek	Orta	Anthropic; Palantir
xAI/Grok	X ve Grok hakkında Komisyon soruşturması	Platform manipölasyonu + sistemik risk	Dolaylı	Yüksek	European Commission

Grafik 10.1. Seçilmiş vakalarda öne çıkan risk alanları



Harita 10.1. Vaka analizlerinin jeopolitik dağılımı



27.1. Veri sızıntısı, güvenlik ve manipülasyon vakaları

Seçilmiş vakalar birlikte okunduğunda üç ortak örüntü öne çıkmaktadır. Birincisi, veri akışının kurumsal büyüklükle birlikte karmaşıklaşmasıdır. OpenAI'nin 2023 olayı ile DeepSeek'in Wiz tarafından tespit edilen veritabanı maruziyeti birbirinden yapısal olarak farklı olsa da, her iki örnek de üretken yapay zekâ sistemlerinde güvenlik açığının yalnızca dış saldırı değil, yanlış yapılandırma, kütüphane sorunu ve operasyonel süreç zafiyetinden de kaynaklanabildiğini göstermektedir (OpenAI, 2023; Wiz Research, 2025).

İkincisi, kamu ve savunma ile entegrasyonun hızlanmasıdır. OpenAI for Government ve ChatGPT Gov girişimleri ile Claude Gov ve Palantir/AWS hattı birlikte okunduğunda, frontier modellerin artık yalnızca ticarî araçlar değil, devlet kapasitesinin parçası hâline geldiği anlaşılmaktadır. ABD Savunma Bakanlığı CDAO'nun 2025 duyurusu da OpenAI, Anthropic, Google ve xAI ile yapılan ortaklıkları teyit etmektedir. Bu eğilim, yapay zekâ yönetişiminin piyasa meselesi olmaktan çıkarak ulusal güvenlik ve stratejik yetenek meselesine dönüştüğünü göstermektedir (CDAO, 2025; OpenAI, 2025a; Anthropic, 2025a).

Üçüncüsü ise platform ve model riskinin yakınsamasıdır. Grok vakası, model çıktılarının tek başına değil; çok geniş erişimli sosyal ağlar ve öneri sistemleriyle birleştğinde sistemik risk oluşturduğunu göstermektedir. Bu nedenle gelecekteki düzenleme ve denetim tartışmaları yalnızca model sağlayıcıya değil, modelin bağlandığı platform mimarisine ve tavsiye sistemlerine de odaklanmak zorundadır (European Commission, 2026a).

Şema 10.1. Yapay zekâ vakalarının tipik seyir zinciri



Şekil 10.1. Vaka analizlerinden türetilen risk yakınsama matrisi

	OpenAI	Copilot	DeepSeek	Anthropic/ Palantir	xAI/Grok
Veri ve mahremiyet	Orta	Orta	Yüksek	Düşük	Düşük
Kamu/savunma bağı	Yüksek	Düşük	Düşük	Yüksek	Orta
Manipülasyon ve platform riski	Düşük	Düşük	Düşük	Düşük	Yüksek
Uyum ve şeffaflık	Orta	Orta	Orta	Orta	Yüksek

Bölüm Sonu Değerlendirmesi

Bu bölümün temel bulgusu şudur: yapay zekâ vakaları birbirinden kopuk teknoloji hikâyeleri değil, aynı risk topolojisinin farklı tezahürleridir. Veri sızıntısı, sınırlı şeffaflık, kamu ve savunma entegrasyonu, platform etkisi ve düzenleyici baskı sürekli olarak aynı aktörler ve aynı ürün aileleri etrafında döğümlenmektedir. Bu nedenle vaka analizi, istisna aramak için değil; sistemik örüntüyü görünür kılmak için gereklidir (NIST, 2023; OECD, 2026).

Seçilmiş vakalar ayrıca yapay zekâ çağında riskin yalnızca model yanlışlığı ile sınırlı olmadığını göstermektedir. Asıl kırılma noktası, modelin veri rejimi, altyapı konumu, kullanıcı sözleşmesi, kamu işbirliği, platform erişimi ve düzenleyici çevre ile birleştiği yerdedir. Bundan sonraki bölümde senaryolar ve stratejik uyarılar ele alınırken, burada ortaya konan vaka desenleri bir erken uyarı seti olarak kullanılacaktır.

BÖLÜM XI – SENARYOLAR, RİSK MATRİSLERİ VE STRATEJİK UYARI

Bu bölümün amacı, yapay zekâ ekosisteminin mevcut eğilimlerinden hareketle 2030–2040 dönemine ilişkin olası gelecek kümelerini ortaya koymak, bu geleceklerin devletler ve kurumlar açısından doğuracağı başlıca riskleri tartışmak ve operasyonel direnç hatlarını tanımlamaktır.

Bu raporun yaklaşımına göre senaryo, 'gelecekte ne olabilir?' sorusundan çok, 'bugün hangi hatalar yapılırsa hangi gelecek daha olası hâle gelir?' sorusuna cevap vermelidir. Stanford AI Index 2026, kabiliyetlerin ve benimsenmenin hızla arttığını; buna karşılık bunları ölçme, denetleme ve yönetme kapasitesinin aynı hızda gelişmediğini göstermektedir.

28. 2030–2040 senaryoları

Bu bölümün amacı, yapay zekâyâ ilişkin gelecek projeksiyonlarını soyut beklenti listeleri olmaktan çıkarıp; teknik kapasite, yönetim gecikmesi, altyapı baskısı, sentetik medya riski ve kurumsal bağımlılık eksenlerinde oluşabilecek yanlış karar zincirleri üzerinden okumaktır.

28.1. Senaryo I: Yönetilen Hız, Kırılan Denge

Bu senaryoda yapay zekâ yayılımı hız kesmeden sürer; ancak büyük güçler, büyük şirketler ve temel düzenleyici rejimler sistemin tamamen kontrolden çıkmasını önleyecek kadar asgari denetim üretir. Avrupa Birliği tarafında AI Act'in aşamalı uygulama takvimi, gecikmeli ama çalışan bir yönetim zemini ihtimalini göstermektedir. Ancak bu senaryo güvenli bir gelecek anlamına gelmez; temel risk, sahte emniyet duygusudur (European Commission, 2025; Stanford HAI, 2026).

28.2. Senaryo II: Platform Egemenliği ve Sessiz Bağımlılık

Bu senaryoda en büyük kırılma görünür kriz değil, sessiz bağımlılık birikimi olur. Yapay zekâ çok sayıda kuruma, sektöre ve devlete yayılır; fakat temel model, bulut, hızlandırıcı çip, veri merkezi ve kurumsal araç zinciri birkaç büyük küresel düğüm etrafında yoğunlaşmaya devam eder. Sonunda teknoloji kullanımında artış vardır ama teknoloji üzerinde egemenlikte azalma yaşanır.

28.3. Senaryo III: Sentetik Belirsizlik ve Kurumsal Felç

Bu senaryo raporun Soğuk Kaos ve DeNormasyon kavramlarına en yakın gelecek ihtimalidir. Sorun yalnızca yanlış içeriğin yayılması değil; doğrulama kapasitesinin sürekli gecikmesidir. Seçim, kriz yönetimi, finansal transfer, kamu açıklaması ve savunma uyarısı gibi alanlarda birkaç dakikalık tereddüt bile büyük maliyet üretebilir. Sonuçta problem teknik doğruluk değil, karar zamanlaması olur (OECD, 2026; UNESCO, 2024).

28.4. Senaryo IV: Enerji ve Hesaplama Darboğazının Jeopolitikleşmesi

2030–2040 döneminin en az konuşulup en çok hasar üretme potansiyeli taşıyan hattı, yapay zekâ altyapısının enerji ve kritik kaynak baskısıyla daha açık biçimde jeopolitikleşmesidir. IEA'nin güncel analizleri, veri merkezi elektrik talebinin ve AI odaklı yüklerin hızla arttığını göstermektedir. Böylece yapay zekâ stratejisi fark edilmeden enerji güvenliği, endüstriyel politika ve altyapı finansmanı meselesine dönüşür (IEA, 2025; 2026).

28.5. Senaryoların Ortak Sonucu

Bu dört senaryo birlikte okunduğunda ortak sonuç açıktır: 2030–2040 dönemini belirleyecek olan yalnızca yapay zekânın ne kadar gelişeceği değildir. Asıl belirleyici, kurumların ve devletlerin bu gelişmeyi hangi yönetim kapasitesiyle karşılayacağı; hangi bağımlılıkları erken fark edeceği; sentetik belirsizliğe karşı hangi teyit mimarisini kuracağı ve enerji ile altyapı baskılarını ne kadar erken stratejik dosya hâline getireceğidir.

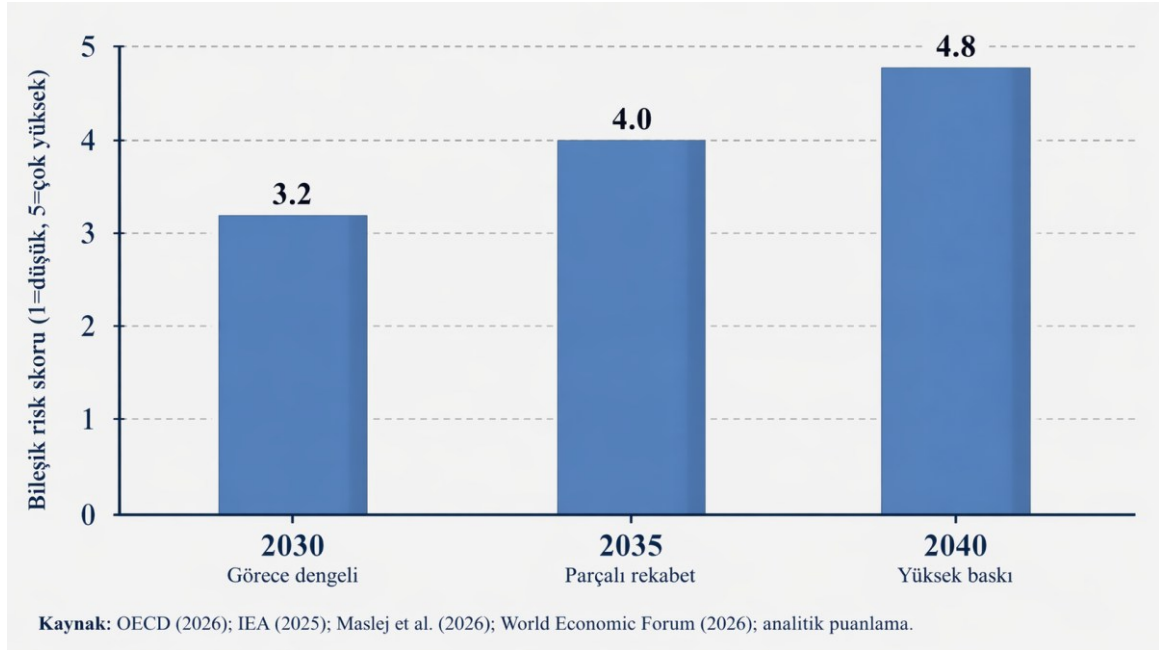
Tablo 11.1. 2030–2040 dönemi için temel yapay zekâ senaryoları

2030–2040 dönemi için temel yapay zekâ senaryoları

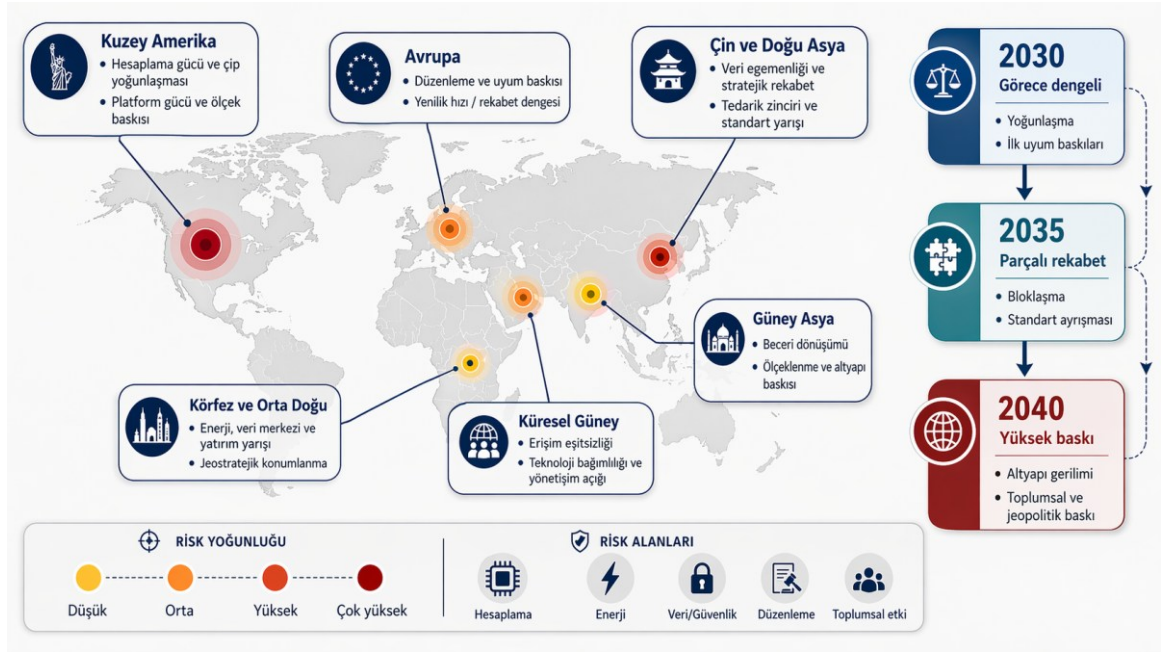
Senaryo	Temel dinamik	Muhtemel sonuçlar	Başlıca riskler	Politika önceliği
Düzenlenmiş verimlilik sıçraması	Güvenlik, standart ve denetim mekanizmaları güçlenirken yapay zekâ verimlilik artışı üretir.	Kamu hizmetlerinde iyileşme, kurumsal adaptasyon, ölçülü otomasyon.	Uyum maliyetleri, küçük aktörlerin geride kalması, düzenleme kapasitesi açığı.	Risk temelli yönetim, beceri dönüşümü ve hesap verebilirlik.
Jeopolitik bloklaşma ve teknoloji rekabeti	ABD, Çin ve müttefik ağları arasında standart, çip, bulut ve veri alanlarında sert rekabet derinleşir.	Parçalı ekosistem, tedarik zinciri ayrışması, stratejik otonomi arayışı.	İhracat kısıtları, maliyet artışı, bağımlılıkların güvenlik krizine dönüşmesi.	Çeşitlendirme, dayanıklılık, ortak standartlar ve kritik kapasite yatırımı.
Platform yoğunlaşması ve asimetrik bağımlılık	Az sayıda büyük şirket veri, model, bulut ve dağıtım katmanlarında yoğun güç toplar.	Hızlı yayılım, ölçek ekonomisi, ancak piyasa ve bilgi gücünün dar merkezlerde toplanması.	Rekabet zayıflaması, şeffaflık kaybı, kamu yararı ile ticari mantık arasındaki gerilim.	Rekabet hukuku, veri erişimi, denetlenebilirlik ve kamusal gözetim.
Güven krizi ve tepkisel düzenleme	Hatalı çıktı, manipülasyon, sentetik medya ve yüksek etkili olaylar toplumsal güveni aşındırır.	Sert düzenleme dalgaları, kurumsal yavaşlama, kullanım alanlarında seçici daralma.	Hak ihlalleri, yanlış yasaklama, inovasyonun aşırı baskılanması ve meşruiyet krizi.	İnsan gözetimi, doğrulama altyapısı, olay bildirimi ve yüksek risk alanlarında sıkı test.

Kaynak: OECD (2026); Maslej et al. (2026); World Economic Forum (2026); yazarın sentezi.

Grafik 11.1. 2030–2040 senaryolarında bileşik risk yoğunluğu



Harita 11.1. 2030–2040 yapay zekâ geleceklerinin kavramsal risk haritası



29. Devletler için risk matrisi

Bu bölümün amacı, yapay zekâ çağında devletlerin karşı karşıya kaldığı başlıca stratejik riskleri yalnızca isimlendirmek değil; her biri için hangi tetikleyicilerin devreye girdiğini, hangi kurumsal kırılmalıkların hasarı büyüttiğini, muhtemel sonucun ne olduğunu ve hangi önleyici hamlelerin gecikmeden kurulması gerektiğini göstermektir.

29.1. Risk I: Dış Model ve Platform Bağımlılığı

Tetikleyici, kamu kurumlarının ve kritik sektörlerin dar sayıda dış sağlayıcıya hızla bağlanmasıdır. Kırılmalık, karar destek mekanizmasının dış sağlayıcının model mimarisine, hizmet koşullarına ve sürüm döngüsüne bağımlı hâle gelmesidir. Muhtemel sonuç, devletin teknolojiyi kullanan ama teknoloji üzerinde yön belirleyemeyen aktöre dönüşmesidir. Önleyici hamle ise tedarik çeşitliliği, model geçiş planı, veri taşınabilirliği ve hassas iş akışlarında dış servis sınırlamasıdır.

29.2. Risk II: Sentetik Medya Kaynaklı Teyit Çöküşü

Tetikleyici, kriz ve seçim dönemlerinde sentetik ses, görüntü ve belge akışının hızlanmasıdır. Kırılmalık, kurumlar arası hızlı doğrulama prosedürlerinin ve provenans standardının bulunmamasıdır. Muhtemel sonuç, karar temposunun düşmesi ve meşru açıklamaların da kuşkuya açık hâle gelmesidir. Önleyici hamle, provenans tabanlı içerik kimliği, çok kanallı doğrulama ve kriz dönemleri için önceden tanımlanmış teyit rejimidir.

29.3. Risk III: Ajanik Sistemlerde Yetki Aşımı ve Zincirleme Hata

Tetikleyici, görev planlama ve araç çağırma kapasitesine sahip ajan tabanlı sistemlerin kamu ve kritik sektör iş akışlarına hızla girmesidir. Kırılmalık, işlem yetkisi, dış veri erişimi, loglama ve durdurma mekanizmasının sıkı tanımlanmamasıdır. Muhtemel sonuç, yanlış işlem, yetkisiz veri erişimi ve sorumluluğun sistem içine dağılmasıdır. Önleyici hamle, kademeli yetki, yüksek etkili alanlarda 'öneri üretir ama icra etmez' ilkesi ve bağımsız hata incelemesidir.

29.4. Risk IV: Kritik Altyapı, Enerji ve Hesaplama Darboğazı

Tetikleyici, veri merkezi, hızlandırıcı çip, soğutma ve elektrik talebindeki hızlı artışın ulusal AI stratejilerinin önüne geçmesidir. Kırılmalık, yapay zekâ stratejisinin enerji ve altyapı

planlamasından ayrı yürütülmesidir. Muhtemel sonuç, teknik kabiliyet hedeflerinin maddi altyapı eksikliği yüzünden aksaması ve egemenlik açığının büyümesidir. Önleyici hamle, dijital strateji, sanayi politikası ve enerji güvenliği dosyalarını tek masada toplamaktır.

29.5. Risk V: Kamu Hizmetlerinde Hak İhlali ve Meşruiyet Aşınması

Tetikleyici, yapay zekâ sistemlerinin sosyal yardım, kolluk, göç, sağlık, eğitim, vergi ve adalet gibi alanlarda hız baskısıyla devreye alınmasıdır. Kırılganlık, otomatik kararın verimliliğinin hak etkisinden daha görünür sayılmasıdır. Muhtemel sonuç, hatalı sınıflandırma, ayrımcı sonuç ve devletle vatandaş arasındaki güven ilişkisinde aşınmadır. Önleyici hamle, zorunlu etki değerlendirmesi, itiraz hakkı, dış denetim ve 'insan kararını ikame etmez' ilkesidir.

29.6. Risk VI: Düzenleyici Gecikme ve Uyumda Kâğıt Üstü Başarı

Tetikleyici, devletin AI yönetişimini kanun veya strateji yayımlamakla tamamlandığını sanmasıdır. Kırılganlık, kayıt, raporlama, model envanteri, olay bildirimi ve yaptırım kapasitesinin fiilen kurulmamasıdır. Muhtemel sonuç, devletin düzenleyen görünürken fiilen takip eden konumda kalmasıdır. Önleyici hamle, yönetişimi strateji belgesi değil uygulama zinciri olarak tasarlamaktır.

29.7. Risk VII: Savunma ve İstihbaratta Denetimin Biçimselleşmesi

Tetikleyici, hız baskısı altında AI destekli analiz ve karar destek sistemlerinin savunma-istihbarat akışına daha derin girmesidir. Kırılganlık, insan denetiminin görünürde var olup gerçekte veto kapasitesi taşınamamasıdır. Muhtemel sonuç, hatalı alarm, yanlış önceliklendirme ve sorumluluğun sistem içine dağılmasıdır. Önleyici hamle, çok katmanlı onay, kırmızı takım testi, loglama ve sistem dışı alternatif analiz hattıdır.

Şema 11.1. Yapay zekâ geleceğini şekillendiren erken uyarı zinciri



Kaynak: OECD (2026); IEA (2025); World Economic Forum (2026); yazarın sentezi.

Şekil 11.1. Devlet ve kurumlar için yapay zekâ dayanıklılık matrisi

	1. Yönetişim ve Liderlik	2. Risk Yönetimi ve Değerlendirme	3. Veri ve Model Yönetimi	4. Operasyonel Dayanıklılık	5. İnsan ve Toplum	6. Şeffaflık ve Hesap Verebilirlik	7. İzleme, Denetim ve Sürekli İyileştirme
Devlet (Politika Düzeyi)	- Ulusal AI stratejisi ve yol haritası - Kurumsal koordinasyon mekanizmaları - Sorumluluk ve yetki çerçevesi	- Ulusal risk envanteri ve önceliklendirme - Etki değerlendirmesi (AI Impact Assessment) - Senaryo ve stres testi uygulamaları	- Ulusal veri yönetim standartları - Kritik veri varlıkları envanteri - Model şeffaflığı ve belgelendirme	- Güvenli ve esnek altyapı planlaması - Kesinti ve kriz senaryoları - Yedekleme ve felaket iyileştirme	- Yetenek geliştirme stratejileri - Toplumsal etki analizleri - Eşitlik ve kapsayıcılık politikaları	- Kamusal bilgilendirme ve raporlama - Açık veri ve açık standartlar - Paydaş katılımı	- Sürekli izleme mekanizmaları - Bağımsız denetim ve değerlendirme - Öğrenen yönetim döngüleri
Kamu Kurumları (Organizasyon Düzeyi)	- AI yönetim yapısı ve roller - Strateji ile uyumlu hedefler - Üst yönetim sahipliği	- Kurumsal risk kayıtları - Model ve süreç risk analizleri - Üçüncü taraf risk değerlendirmesi	- Veri sınıflandırma ve kalite kontrol - Model yaşam döngüsü yönetimi - Veri ve model erişim denetimleri	- Güvenlik ve siber dayanıklılık - Operasyonel süreç sürekliliği - Olay müdahale planları	- Çalışan farkındalığı ve eğitim - İnsan-AI iş birliği ilkeleri - Etik kullanım rehberleri	- Algoritmik şeffaflık - Karar açıklanabilirliği - Şikâyet ve geri bildirim mekanizmaları	- Performans göstergeleri ve izleme - İç denetim ve uygunluk kontrolleri - Düzeltilici faaliyetler ve iyileştirme
Düzenleyici ve Denetleyici Kurumlar	- Düzenleyici yetki ve görev tanımları - Rehber ve standart geliştirme - Uluslararası iş birliği	- Sektörel risk analizi - Piyasa gözetimi ve erken uyarı - Uygunluk risk modellemesi	- Denetim için veri erişim çerçevesi - Standart test setleri ve ölçütler - Model doğrulama protokolleri	- Denetim altyapısı ve araçları - Regülasyon testi (sandbox) - Kriz koordinasyon mekanizmaları	- Uzman kapasitesi geliştirme - Paydaş diyalogları - Kamu yararı odaklı yaklaşım	- Düzenleyici raporlama - Karar gerekçelendirme ve açıklık - Kamuya hesap verebilirlik	- Sürekli mevzuat gözden geçirme - Etki değerlendirme çalışmaları - Uluslararası iyi uygulama takibi
Kritik Altyapı ve Hizmet Sağlayıcılar	- Yönetim kurulu sahipliği - AI politikası ve prensipleri - Sorumluluk dağılımı	- Varlık ve tehdit analizleri - AI sistem risk sınıflandırması - Tedarik zinciri riski	- Veri güvenliği ve gizlilik - Model doğrulama ve test - Versiyon ve değişiklik yönetimi	- Sistem izleme ve olay yönetimi - İş sürekliliği ve yedeklilik - Fiziksel ve siber koruma	- Operatör eğitimi - Kullanıcı güvenliği ve erişilebilirlik - Toplumsal hizmet sürekliliği	- Hizmet seviyesi açıklığı - Müşteri bilgilendirme - Veri kullanım şeffaflığı	- Sürekli izleme ve uyarı - Periyodik test ve tatbikat - Kök neden analizi ve iyileştirme
Özel Sektör ve Teknoloji Geliştiriciler	- Etik AI yönetim yapısı - Ürün sorumluluğu ilkeleri - Yönetim taahhüdü	- Model risk değerlendirme çerçevesi - Yanlış, güvenlik ve kötüye kullanım analizi - Etki ve bağımlılık analizleri	- Veri kaynak şeffaflığı - Veri lisanslama ve izin yönetimi - Model kartları ve dokümantasyon	- Güvenli geliştirme yaşam döngüsü - Kırmızı takım testleri ve güvenlik testleri - Dağıtım ve geri alma stratejileri	- Kullanıcı eğitimi ve rehberlik - Etik tasarım ilkeleri - Toplumsal etki yönetimi	- Açıklanabilirlik çözümleri - Kullanım şartları açıklığı - Bağımsız değerlendirme ve doğrulama	- Performans ve adillik izleme - Geri bildirim temelli sürekli iyileştirme - Güncelleme ve sürüm yönetimi

● Yüksek olgunluk / tamamlanmış ● Orta olgunluk / gelişme aşığı ● Düşük olgunluk / kritik açığı gösterir

✓ Matris, kurumların mevcut durum değerlendirilmesi, yol haritası geliştirme ve denetim planlaması için kullanılabilir.

Kaynak: OECD (2026) due diligence mantığı ile NIST ve WEF risk çerçevelerinin uyarlanmış sentezi.

30. Kurumlar için hayatta kalma rehberi

Bu bölümün amacı, yapay zekâyı kurumsal düzeyde erişimin neden tek başına avantaj üretmediğini ve neden asıl meselenin kontrollü, izlenebilir, denetlenebilir ve geri döndürülebilir bir kullanım mimarisi kurmak olduğunu göstermektir.

30.1. İlk 90 Gün: Kurumun Yapması Gerekenler

İlk 90 günün hedefi, kuruma yaygın AI kullanımını serbest bırakmak değil; envanter, sınır ve sorumluluk üretmektir. Kurum önce hangi ekiplerin hangi araçları kullandığını, hangi verilerin bu sistemlere girdiğini ve hangi iş akışlarının doğrudan ya da dolaylı AI çıktısına dayandığını tespit etmelidir. Ardından kullanım alanları risk düzeyine göre ayrılmalı, insan gözetiminin kim tarafından ve hangi yetkiyle yapılacağı tanımlanmalı, veri sınırları ile olay protokolleri yazılı hâle getirilmelidir.

30.2. Kurumun Asla Yapmaması Gerekenler

Kurum, 'herkes kullanıyor, biz de kullanalım' mantığıyla yapay zekâyı yatay biçimde yaymamalıdır. İnsan denetimi onay kutusuna indirgenmemeli; hassas veri kısa prompt içinde geçtiği için zararsız sanılmamalı; tedarik sözleşmesi hukuk biriminin ikincil evrakı gibi görülmemelidir. Yapay zekâ maliyeti de yalnızca lisans bedeline indirgenemez; bulut, çıkarım, enerji ve süreklilik maliyetleri birlikte düşünülmelidir.

30.3. Zorunlu Kurumsal Kontrol Listesi

Kurum, AI kullanımını yaygınlaştırmadan önce şu sorulara yazılı cevap verebilmelidir: Hangi araçlar fiilen kullanılıyor? Hangi kullanım alanları düşük, orta ve yüksek etkili? Hangi veriler dış sistemlere kesinlikle verilemez? İnsan gözetimi kimde? Çıktılar geriye dönük izlenebilir mi? Hizmet kesintisi veya tedarikçi değişikliği için yedek plan var mı? AI kaynaklı olay tanımı ve eskalasyon zinciri hazır mı?

30.4. Yönetim Kuruluna Sorulacak 10 Soru

Yönetim kurulu düzenli olarak şu soruları sormalıdır: Kurum bugün hangi kararları AI'ya hazırlatıyor? Sağlayıcı bağımlılığımızın seviyesi nedir? Hassas veri sınırlarımız sağlam mı? İnsan gözetimi gerçek mi, biçimsel mi? Bir AI hatasını kaç saat içinde fark ederiz? AI kaynaklı bir yanlış kararın hukukî ve itibari maliyetini hesapladık mı? Bugün verimlilik kazanırken yarın hangi stratejik kontrolü kaybediyoruz?

30.5. Kurum Tiplerine Göre Kısa Uygulama Notu

Her kurum aynı değildir. Kritik altyapı işleten, finansal karar veren, kamu hizmeti sunan, büyük insan kaynakları havuzu yöneten veya yüksek hacimli müşteri verisi işleyen kurumlar için risk çitası daha yüksektir. Bu nedenle kurumsal rehber tek tip uygulanmamalı; insan üzerinde hukukî, ekonomik veya fiziksel etkisi yüksek alanlarda daha sıkı eşikler kullanılmalıdır.

Yanlış Karar Riski Kutusu – Bölüm XI / 30. Başlık

Bu bölümden çıkarılması gereken temel uyarı şudur: kurum için asıl tehlike AI kullanmamak değil, AI'yi yönetimden önce yaygınlaştırmaktır. Envantersiz kullanım, sınırsız veri akışı, biçimsel insan denetimi, sözleşmesiz bağımlılık, olay protokolsüzlüğü ve maliyeti yalnızca lisans bedeline indirgeme; kurumun verimlilik kazandığını sanarken stratejik kontrol kaybetmesine yol açar.

Bölüm Sonu Değerlendirmesi

Bu bölümün temel bulgusu şudur: yapay zekâ geleceği tek bir doğrusal hat izlemez; farklı güç yoğunlaşmaları, yönetim kapasiteleri ve risk rejimleri farklı gelecek kümeleri üretir. Devletler için öncelik egemenlik, güvenlik ve enformasyon bütünlüğünü birlikte ele alan çok katmanlı stratejiler geliştirmek; kurumlar için öncelik ise due diligence, insan gözetimi, veri sınırı ve teyit disiplini üzerine kurulu dayanıklılık mimarisi oluşturmaktır.

BÖLÜM XII – TÜRKİYE’NİN YAPAY ZEKÂ EŞİĞİ: KAPASİTE, BAĞIMLILIK, EGEMENLİK VE SOĞUK KAOS RİSKİ

Bu bölümün amacı, Türkiye’nin yapay zekâ çağındaki konumunu yalnızca teknolojik adaptasyon veya dijital dönüşüm başlığı altında değil; veri egemenliği, Türkçe dil kapasitesi, kamu yönetimi, özel sektör dönüşümü, savunma-güvenlik mimarisi, hukuki düzenleme ve toplumsal dayanıklılık eksenlerinde değerlendirmektir. Böylece raporun önceki bölümlerinde küresel düzeyde tartışılan model yoğunlaşması, platform bağımlılığı, sentetik medya, veri rejimi, yönetim açığı ve Soğuk Kaos riskleri Türkiye bağlamına taşınmaktadır.

Türkiye açısından yapay zekâ meselesi, yalnızca bu teknolojinin ne kadar hızlı kullanılacağı sorusundan ibaret değildir. Daha kritik soru, Türkiye’nin yapay zekâyı hangi veri mimarisiyle, hangi model bağımlılık düzeyiyle, hangi hukuki güvenceyle, hangi kamu denetimiyle ve hangi stratejik egemenlik hedefiyle kullanacağıdır. Bu nedenle Türkiye bölümü, yapay zekâyı bir yazılım aracı olarak değil; kamu kapasitesi, ekonomik rekabet, dijital egemenlik, dil politikası, ulusal güvenlik ve toplumsal güven üretimi açısından yeni bir kurucu eşik olarak ele almaktadır (Cumhurbaşkanlığı Dijital Dönüşüm Ofisi & Sanayi ve Teknoloji Bakanlığı, 2021; Kaya, 2026).

Türkiye’nin yapay zekâ alanındaki temel avantajı, güçlü bir dijital kamu altyapısına, savunma sanayii tecrübesine, genç ve teknolojiye uyumlu nüfusa, Türkçe merkezli geniş bir veri alanına, bölgesel etki kapasitesine ve kamu-özel sektör koordinasyon potansiyeline sahip olmasıdır. Buna karşılık en kritik kırılganlıklar; yüksek ölçekli hesaplama altyapısı, frontier model geliştirme kapasitesi, yarı iletken ve hızlandırıcı çip bağımlılığı, açık ve kaliteli veri ekosistemi, kurumsal yapay zekâ yönetimi, yapay zekâ okuryazarlığı ve düzenleyici çerçevenin uygulama kapasitesinde yoğunlaşmaktadır (Sanayi ve Teknoloji Bakanlığı & Cumhurbaşkanlığı Dijital Dönüşüm Ofisi, 2024; TBMM, 2026).

Bu nedenle Türkiye için yapay zekâ stratejisinin temel hedefi, yalnızca kullanım oranını artırmak olmamalıdır. Asıl hedef, dış model ve platform bağımlılığını azaltan, Türkçe ve bölgesel veri varlığını koruyan, kamu karar süreçlerinde denetlenebilirliği güvence altına alan, kritik altyapı ve güvenlik alanlarında insan denetimini biçimsel olmaktan çıkaran ve toplumun ortak hakikat zeminini sentetik medya çağında koruyabilen bir yapay zekâ egemenlik mimarisi kurmak olmalıdır (KVKK, 2021, 2025; NIST, 2023).

31. Türkiye’nin Yapay Zekâ Konumlanması

Türkiye, yapay zekâ alanında geç kalmış bir ülke olarak değil; stratejik tercih eşiğinde bulunan orta-büyük ölçekli bir teknoloji, güvenlik ve bölgesel etki aktörü olarak değerlendirilmelidir. Bu konum, Türkiye’ye iki yönlü bir sorumluluk yüklemektedir. Bir taraftan yapay zekâ teknolojilerinin sunduğu üretkenlik, kamu hizmeti kalitesi, sanayi dönüşümü, savunma kapasitesi ve bilimsel araştırma imkânlarından yararlanmak zorundadır. Diğer taraftan dış model bağımlılığı, veri sızıntısı, sentetik medya manipülasyonu, kamu kararlarında otomatikleşmiş hata ve dijital egemenlik kaybı gibi riskleri yönetmek zorundadır (Cumhurbaşkanlığı Strateji ve Bütçe Başkanlığı, 2023; TBMM, 2026).

Türkiye’nin yapay zekâ konumlanması üç ana katmanda okunmalıdır. İlk katman stratejik niyet katmanıdır. Ulusal Yapay Zekâ Stratejisi 2021–2025 ve devamındaki 2024–2025 Eylem Planı, Türkiye’nin bu alanı yalnızca teknik modernleşme başlığı altında değil; insan kaynağı, araştırma-girişimcilik, kaliteli veri, teknik altyapı, sosyoekonomik uyum, uluslararası iş birliği ve işgücü

dönüşümü başlıklarıyla birlikte ele aldığını göstermektedir. Stratejinin altı temel öncelik, yirmi dört amaç ve yüz on dokuz tedbir etrafında yapılandırılması, yapay zekânın Türkiye açısından çok katmanlı bir politika alanı olarak görüldüğünü ortaya koymaktadır (Cumhurbaşkanlığı Dijital Dönüşüm Ofisi & Sanayi ve Teknoloji Bakanlığı, 2021; Sanayi ve Teknoloji Bakanlığı & Cumhurbaşkanlığı Dijital Dönüşüm Ofisi, 2024).

İkinci katman kapasite ve bağımlılık katmanıdır. Türkiye, yapay zekâ uygulamalarını kamu, finans, savunma, sağlık, eğitim, sanayi, afet yönetimi ve yerel yönetimler gibi çok farklı alanlara yayma potansiyeline sahiptir. Ancak bu potansiyelin stratejik değere dönüşmesi, kullanılan modellerin nerede çalıştığına, verilerin hangi altyapıda işlendiğine, kamu verisinin hangi güvenlik sınıflandırmasına tabi tutulduğuna ve kritik karar süreçlerinde insan denetiminin nasıl işletildiğine bağlıdır. Bu nedenle Türkiye'nin yapay zekâ kapasitesi yalnızca yazılım geliştirme veya uygulama satın alma üzerinden değil; veri merkezi, bulut, hesaplama gücü, siber güvenlik, model denetimi ve hukuki sorumluluk zinciri üzerinden değerlendirilmelidir (OECD, 2024; NIST, 2023).

Üçüncü katman egemenlik ve toplumsal dayanıklılık katmanıdır. Yapay zekâ, Türkiye için yalnızca ekonomik rekabet konusu değildir; aynı zamanda dil, kültür, hukuk, kamu güveni ve kriz yönetimi meselesidir. Türkçe verinin dış model mimarilerinde nasıl temsil edildiği, kamu belgelerinin hangi sistemlerde işlendiği, seçim dönemlerinde sentetik içeriğin nasıl doğrulanacağı, afet anlarında sahte bilgi akışının nasıl engelleneceği ve toplumun ortak hakikat zemininin nasıl korunacağı artık yapay zekâ politikasının merkezî sorularıdır. Bu bağlamda Soğuk Kaos, Türkiye için yalnızca dış kaynaklı dezenformasyon meselesi değil; çok hızlı, çok yoğun ve kısmen sentetik bilgi akışının kamu karar mekanizmalarını yavaşlatması, teyit süreçlerini aşındırması ve kurumsal güveni zayıflatması riskidir (Kaya, 2026; Rid, 2020).

Tablo 12.1. Türkiye'nin yapay zekâ kapasite ve kırılganlık profili

Alan	Mevcut stratejik durum	Güçlü taraf	Kırılganlık	Stratejik öneri
İnsan kaynağı	UYZS ve Eylem Planı uzman yetiştirme hedefi içeriyor.	Genç nüfus, üniversite ağı ve teknoloji girişimciliği.	Nitelikli uzman sürekliliği ve beyin göçü riski.	Kamu-özel sektör ortak burs, istihdam ve dönüş programı.
Veri	Kamu ve özel sektörde geniş veri üretimi var.	E-devlet ve dijital kamu altyapısı.	Veri kalitesi, sınıflandırma ve güvenli paylaşım sorunu.	Ulusal Yapay Zekâ Veri Envanteri ve güvenli kamu veri alanı.
Türkçe model	Türkçe büyük dil modeli hedefi politika gündemine girdi.	Türkçe ve Türk dünyası veri havzası.	Kaliteli, telif uyumlu ve bağlamsal veri eksikliği.	Türkçe Büyük Dil Modeli Topluluğu ve alan bazlı doğrulama setleri.
Altyapı	Veri merkezi, bulut ve yüksek başarılı hesaplama ihtiyacı artıyor.	Kamu teknoloji kurumları ve savunma ekosistemi.	GPU/çip ve yüksek performanslı hesaplama bağımlılığı.	Kritik kamu uygulamaları için güvenli ulusal hesaplama katmanı.
Hukuk	KVKK rehberleri ve TBMM çalışmaları var.	Düzenleme farkındalığı yükseliyor.	Dağınık normlar ve uygulama kapasitesi sorunu.	Yapay zekâ etki analizi, risk sınıfları ve denetim zinciri.
Toplumsal güven	Sentetik medya ve sahte içerik riski artıyor.	Kamu iletişim kapasitesi ve teyit ağları geliştirilebilir.	Deepfake, sahte ses ve kriz anı bilgi kirliliği.	Ulusal Sentetik İçerik Teyit ve Kriz Protokolü.
Savunma-güvenlik	Savunma sanayii yapay zekâ entegrasyonuna açık.	İHA/SİHA, siber güvenlik ve karar destek tecrübesi.	İnsan denetiminin biçimselleşmesi riski.	Durdurabilir insan denetimi ve operasyonel kayıt zorunluluğu.

Kaynak: Cumhurbaşkanlığı Dijital Dönüşüm Ofisi & Sanayi ve Teknoloji Bakanlığı (2021); Sanayi ve Teknoloji Bakanlığı & Cumhurbaşkanlığı Dijital Dönüşüm Ofisi (2024); KVKK (2021, 2025); TBMM (2026); yazarın sentezi.

32. Ulusal Yapay Zekâ Stratejisi ve Politika Mimarisinin Değerlendirilmesi

Türkiye'nin yapay zekâ politika mimarisi, Ulusal Yapay Zekâ Stratejisi 2021–2025 ile kurumsal bir çerçeveye kavuşmuştur. Strateji, “Müreffeh bir Türkiye için çevik ve sürdürülebilir bir yapay zekâ ekosistemiyle küresel ölçekte değer üretmek” vizyonu etrafında şekillendirilmiş; insan kaynağı, araştırma-girişimcilik, veri-teknik altyapı, sosyoekonomik uyum, uluslararası iş birliği ve işgücü dönüşümü başlıklarını birlikte ele almıştır (Cumhurbaşkanlığı Dijital Dönüşüm Ofisi & Sanayi ve Teknoloji Bakanlığı, 2021).

2024–2025 Eylem Planı, bu stratejik çerçevenin güncellenmiş uygulama mimarisini sunmaktadır. Bu güncelleme, yapay zekâ alanındaki hızlı teknolojik değişimin, üretken yapay zekâ dalgasının, küresel regülasyon tartışmalarının ve Türkiye'nin 12. Kalkınma Planı dönemindeki dijital dönüşüm hedeflerinin birlikte dikkate alındığını göstermektedir (Cumhurbaşkanlığı Strateji ve Bütçe Başkanlığı, 2023; Sanayi ve Teknoloji Bakanlığı & Cumhurbaşkanlığı Dijital Dönüşüm Ofisi, 2024).

Stratejinin güçlü tarafı, yapay zekâyı yalnızca Ar-Ge veya girişimcilik konusu olarak görmemesidir. İnsan kaynağı, teknik altyapı, kaliteli veri, uluslararası iş birliği, sosyoekonomik uyum ve işgücü dönüşümü birlikte ele alınmaktadır. Bu yaklaşım doğrudur; çünkü yapay zekâ çağında başarı yalnızca model geliştirme kapasitesiyle değil, o modeli kullanacak kurumların denetim, uyum, etik, veri güvenliği ve uygulama kapasitesiyle ölçülür (OECD, 2024; NIST, 2023).

Bununla birlikte Türkiye'nin politika mimarisinde üç kritik sertleştirme alanı bulunmaktadır. Birincisi, strateji belgelerinde yer alan hedeflerin bağlayıcı kurumsal sorumluluk zincirine dönüştürülmesidir. Yapay zekâ alanında politika belgesi üretmek önemlidir; fakat asıl belirleyici olan hangi kurumun hangi veri alanından, hangi risk sınıfından, hangi denetim mekanizmasından ve hangi hesap verebilirlik sürecinden sorumlu olduğunun açık biçimde tanımlanmasıdır.

İkincisi, kamu kurumlarında yapay zekâ kullanımının “önce deneyelim, sonra düzenleriz” mantığıyla değil, risk sınıflandırmasına dayalı bir etki analizi rejimiyle ilerletilmesidir. Kamu hizmetlerinde yapay zekâ kullanımı; vatandaş hakları, itiraz mekanizması, idari işlem güvenliği, ayrımcılık riski, veri güvenliği ve karar gerekçelendirme ilkeleriyle birlikte ele alınmalıdır. Aksi hâlde yapay zekâ, kamu hizmetini hızlandıran bir araç olmaktan çıkıp yanlış sınıflandırma, hak ihlali ve kurumsal sorumluluk belirsizliği üreten bir risk alanına dönüşebilir (KVKK, 2021, 2025).

Üçüncüsü, Türkiye'nin yapay zekâ politikasının yalnızca Avrupa Birliği düzenlemelerini izleyen bir uyum politikası olmaktan çıkarılıp bölgesel norm üretme kapasitesine kavuşturulmasıdır. Türkiye; Avrupa regülasyon alanı, Amerikan platform ekosistemi, Çin merkezli teknoloji rekabeti, Türk dünyası veri/dil havzası, Orta Doğu ve Afrika dijitalleşme ihtiyaçları arasında benzersiz bir jeoteknolojik konuma sahiptir. Bu nedenle Türkiye'nin yapay zekâ stratejisi, yalnızca “küresel standartlara uyum” hedefiyle sınırlı kalmamalı; Türkçe, Türk dünyası, İslam dünyası ve Küresel Güney ekseninde güvenilir, denetlenebilir ve insan merkezli yapay zekâ normlarının üretilmesine de yönelmelidir (European Union, 2024; TBMM, 2026).

33. Türkçe Büyük Dil Modeli ve Dil Egemenliği Meselesi

Türkiye açısından yapay zekâ çağının en kritik başlıklarından biri Türkçe büyük dil modeli meselesidir. Bu başlık yalnızca teknik bir model geliştirme hedefi olarak görülmemelidir. Türkçe büyük dil modeli, aynı zamanda dil egemenliği, kültürel hafıza, hukuk dili, kamu terminolojisi, eğitim materyali, tarihsel arşiv, diplomatik söylem ve kurumsal karar destek kapasitesi meselesidir (Sanayi ve Teknoloji Bakanlığı & Cumhurbaşkanlığı Dijital Dönüşüm Ofisi, 2024).

Bu hedefin başarısı üç koşula bağlıdır. Birincisi, Türkçe veri yalnızca hacim olarak değil, kalite, temsil gücü, hukuki açıklık ve bağlamsal çeşitlilik açısından güçlendirilmelidir. Düşük kaliteli, bağlamından koparılmış, telif ve kişisel veri sorunları taşıyan veri kümeleriyle geliştirilecek modeller, Türkçe üretse bile Türkçenin anlam dünyasını ve kurumsal hassasiyetlerini yeterince temsil edemez (KVKK, 2021; OECD, 2024).

İkincisi, Türkçe büyük dil modeli kamu, eğitim, hukuk, sağlık, savunma, finans ve afet yönetimi gibi yüksek etkili alanlar için ayrı güvenlik ve doğrulama katmanlarına sahip olmalıdır. Tek bir genel modelin her alanda kullanılabileceği varsayımı hatalıdır. Hukuki metin, afet bilgilendirmesi, sağlık yönlendirmesi, kamu duyurusu veya savunma analitiği aynı hata toleransına sahip değildir. Bu nedenle Türkiye'nin model mimarisi, alan bazlı risk sınıflandırması ve doğruluk rejimiyle desteklenmelidir (NIST, 2023; TBMM, 2026).

Üçüncüsü, Türkçe büyük dil modeli yalnızca dışa bağımlılığı azaltma hedefiyle değil, DeNormasyon'a karşı dilsel savunma hattı olarak düşünülmelidir. Eğer Türkçe kamusal bilgi alanı dış modellerin varsayımlarına, veri boşluklarına, çeviri hatalarına ve görünmez hizalama tercihlerine bırakılırsa, zaman içinde kavramların anlamı, kamu dilinin tonu ve toplumsal hafızanın temsil biçimi dış mimariler tarafından şekillendirilebilir. Bu nedenle Türkçe model kapasitesi teknik bir inovasyon projesinden çok daha fazlasıdır; doğrudan kültürel süreklilik ve epistemik egemenlik meselesidir (Kaya, 2026; UNESCO, 2021).

34. Veri Egemenliği, KVKK ve Kamu Verisi

Türkiye'nin yapay zekâ stratejisinde veri egemenliği merkezi bir yer tutmalıdır. Yapay zekâ sistemleri yalnızca koddan ibaret değildir; verinin toplandığı, işlendiği, sınıflandırıldığı, saklandığı, modele dönüştürüldüğü ve karar süreçlerine bağlandığı bütün bir değer zinciri üzerinden güç üretir. Bu nedenle Türkiye açısından temel soru, kamu ve özel sektör verilerinin hangi kurallarla yapay zekâ sistemlerine açılacağı, hangi verilerin kritik veya hassas sayılacağı, kişisel verilerin nasıl korunacağı ve model çıktılarının nasıl denetleneceğidir (KVKK, 2021, 2025).

KVKK'nın "Yapay Zekâ Alanında Kişisel Verilerin Korunmasına Dair Tavsiyeler" metni, kişisel veri işleme temelli yapay zekâ uygulamalarının hukuka uygunluk, dürüstlük, ölçülülük, hesap verebilirlik, şeffaflık, veri minimizasyonu, veri doğruluğu, amaçla sınırlılık ve veri güvenliği ilkeleri üzerinden yönetilmesi gerektiğini ortaya koymaktadır. 2025 tarihli üretken yapay zekâ rehberi ise üretken sistemlerin yaşam döngüsü, kullanım alanları, kişisel veri işleme riskleri, mahremiyet ve çocukların korunması gibi başlıkları 6698 sayılı Kanun çerçevesinde ele almaktadır (KVKK, 2021, 2025).

Bu çerçeve Türkiye bölümü için kritik önemdedir. Çünkü yapay zekâda veri güvenliği yalnızca kişisel verinin sızdırılmaması meselesi değildir. Aynı zamanda kamu verisinin dış platformlara taşınması, kurumsal sırların genel amaçlı modellere girilmesi, kamu personelinin hassas belgeleri denetimsiz araçlarla işlemesi, çocuklara ilişkin verilerin üretken sistemlere aktarılması, sağlık ve finans gibi yüksek etkili alanlarda yanlış veri işleme ve otomatik karar üretimi risklerini de kapsar.

Bu nedenle Türkiye için önerilen veri egemenliği mimarisi üç seviyeli olmalıdır. Birinci seviyede, kamu kurumlarının sahip olduğu veri varlıkları sınıflandırılmalı ve yapay zekâ kullanımına uygunluk bakımından envantere bağlanmalıdır. İkinci seviyede, hangi verilerin yerli altyapı dışında işlenemeyeceği, hangi verilerin anonimleştirme veya sentetikleştirme sonrası kullanılabileceği ve hangi verilerin hiçbir şekilde genel amaçlı üretken modellere girilemeyeceği belirlenmelidir. Üçüncü seviyede ise her kamu kurumunda yapay zekâ kullanım kayıtları, model çıktısı denetimi, insan onayı, itiraz mekanizması ve olay bildirimini zorunlu hâle getirilmelidir (KVKK, 2025; NIST, 2023).

35. Savunma Sanayii, Güvenlik ve Kritik Altyapı

Türkiye açısından yapay zekâ yalnızca sivil dijital dönüşüm konusu değildir. Savunma sanayii, sınır güvenliği, siber güvenlik, istihbarat analitiği, afet yönetimi, enerji altyapısı, ulaştırma, telekomünikasyon ve kritik kamu hizmetleri doğrudan bu alanın içine girmektedir. Bu nedenle yapay zekâ entegrasyonu, özellikle yüksek etkili alanlarda “karar destek” ile “karar devri” arasındaki ayrımı keskin biçimde korumalıdır (NIST, 2023; TBMM, 2026).

Savunma ve güvenlik alanında yapay zekâ kullanımının temel riski, insan denetiminin biçimsel onaya indirgenmesidir. Bir sistemde insanın bulunması, sistemin gerçekten denetlendiği anlamına gelmez. Operasyon temposu, veri hacmi, otomatik öneri baskısı ve kurumsal hiyerarşi nedeniyle insan, çoğu zaman sistemi fiilen durduramayan bir onay makamına dönüşebilir. Bu nedenle Türkiye’nin savunma yapay zekâsı mimarisinde “human-in-the-loop” söylemi yerine “durdurabilir insan denetimi”, “operasyonel kayıt”, “sonradan denetlenebilir karar zinciri” ve “hukuki sorumluluk izi” esas alınmalıdır (NIST, 2023; UNESCO, 2021).

Kritik altyapı bakımından yapay zekâ, hem koruyucu hem de kırılganlık artırıcı bir araçtır. Enerji, ulaştırma, su, haberleşme ve finans altyapılarında yapay zekâ destekli izleme, bakım, erken uyarı ve anomali tespiti önemli fırsatlar sunar. Ancak aynı sistemler yanlış veriyle beslendiğinde, sahte alarm ürettiğinde, siber manipülasyona maruz kaldığında veya dış bulut altyapısına aşırı bağımlı çalıştığında operasyonel kırılganlığı artırabilir. Bu nedenle Türkiye’nin kritik altyapı yapay zekâ politikası, yalnızca verimlilik değil; süreklilik, yedeklilik, ulusal denetim ve kriz anında çevrimdışı çalışabilirlik ilkeleriyle birlikte tasarlanmalıdır (Cumhurbaşkanlığı Strateji ve Bütçe Başkanlığı, 2023; IEA, 2025).

36. Türkiye İçin Soğuk Kaos ve DeNormasyon Riskleri

Türkiye bağlamında Soğuk Kaos riski, en çok kriz, seçim, afet, güvenlik ve ekonomik kırılganlık dönemlerinde görünür hâle gelir. Bu dönemlerde bilgi ihtiyacı artar, teyit kapasitesi zorlanır, resmi açıklama beklentisi yükselir ve toplumsal duyarlılık keskinleşir. Yapay zekâ destekli sahte ses, deepfake video, otomatik sosyal medya içeriği, sahte kamu duyurusu, sahte yardım çağrısı veya manipüle edilmiş belge bu ortamda yalnızca iletişim sorunu üretmez; karar süreçlerini yavaşlatır, kamu güvenini aşındırır ve kurumlar arası koordinasyonu zorlaştırır (Kaya, 2026; Rid, 2020).

DeNormasyon ise bu riskin daha kalıcı boyutudur. Bir toplum sürekli olarak sahte içerik, doğrulanmamış belge, çelişkili açıklama, manipülatif görsel ve otomatik propaganda akışıyla karşılaştığında yalnızca yanlış bilgiye maruz kalmaz; doğru ile yanlış arasındaki ayrımı kurumsal olarak sürdürme kapasitesini de kaybetmeye başlar. Bu durum, bilgi kirliliğinden daha ağırdır. Çünkü sorun yalnızca neyin yanlış olduğu değil, hangi kurumun, hangi hızda, hangi kanıt standardıyla doğruyu yeniden kurabileceğidir (Pomerantsev, 2019; Kaya, 2026).

Türkiye için en kritik risk alanları şunlardır: seçim dönemlerinde sahte aday konuşmaları ve manipüle edilmiş görüntüler; afet dönemlerinde sahte konum, sahte yardım kampanyası ve yanlış yönlendirme; finansal alanda sahte yönetici sesi veya otomatik dolandırıcılık; kamu kurumları adına üretilen sahte duyurular; etnik, mezhepsel veya politik fay hatlarını hedefleyen sentetik içerik; şirket içi belge ve toplantı dolandırıcılıkları; dış kaynaklı algı operasyonları; yüksek etkili kamu hizmetlerinde otomatik sınıflandırma hataları. Bu alanlarda yapay zekâ yalnızca yeni bir risk üretmez; mevcut toplumsal kırılganlıkları hızlandırıcı ve çoğaltıcı bir etki doğurur (KVKK, 2025; TBMM, 2026).

Şema 12.1. Türkiye’de yapay zekâ kaynaklı Soğuk Kaos zinciri

Tetikleyici	Yayılma kanalı	Kurumsal etki	Toplumsal sonuç	Savunma hattı
Sentetik ses / deepfake	Sosyal medya, mesajlaşma ağları, sahte haber siteleri	Teyit gecikmesi ve açıklama baskısı	Güven aşınması ve kutuplaşma	Hızlı provenans kontrolü ve resmî teyit protokolü
Sahte kamu duyurusu	Kopya alan adları, bot hesaplar, görsel manipülasyon	Kurumlar arası çelişki ve koordinasyon kaybı	Panik, yanlış yönelim, hizmete güvensizlik	Dijital imzalı duyuru ve kriz iletişimi standardı
Veri/manipülasyon saldırısı	Kurum içi sistemler, dış servisler, API zincirleri	Yanlış karar, yanlış sınıflandırma, itiraz yükü	Hak kaybı ve meşruiyet aşınması	Yapay zekâ olay kaydı ve bağımsız denetim

Kaynak: Kaya (2026); KVKK (2025); NIST (2023); TBMM (2026); yazarın sentezi.

37. Türkiye İçin Yapay Zekâ Risk Sınıfları ve Denetim Öncelikleri

Türkiye’de yapay zekâ yönetişi, tüm uygulamaları aynı seviyede ele alan yatay bir teknoloji politikası olarak kurgulanmamalıdır. Düşük etkili ofis verimlilik uygulamaları ile sağlık, adalet, güvenlik, çocuk koruma, kamu yardımı, kredi değeriendirmesi veya afet yönlendirmesi gibi yüksek etkili alanlar aynı denetim rejimine tabi tutulamaz. Bu nedenle Türkiye için önerilen model, alan bazlı risk sınıflandırması ile kurum bazlı hesap verebilirlik mimarisini birleştiren hibrit bir modeldir (European Union, 2024; NIST, 2023; TBMM, 2026).

Tablo 12.2. Türkiye için yapay zekâ risk sınıfları ve denetim öncelikleri

Risk alanı	Tipik kullanım	Başlıca risk	Denetim önceliği	Önerilen eşik
Düşük etkili idari işler	Metin taslağı, özetleme, çeviri desteği	Yanlış bilgi ve mahremiyet ihlali	Kullanıcı eğitimi ve veri sınırlaması	Kurum içi kullanım rehberi
Kamu hizmeti	Başvuru sınıflandırma, yönlendirme, ön değerlendirme	Hak kaybı ve itiraz hakkının aşınması	İnsan onayı, kayıt, açıklanabilirlik	Zorunlu yapay zekâ etki analizi
Sağlık ve sosyal hizmet	Risk önceliklendirme, destek hizmeti, karar desteği	Yanlış yönlendirme ve ayrımcılık	Uzman doğrulaması ve klinik/idari sorumluluk izi	Yüksek risk sınıfı
Eğitim ve çocuklar	Öğrenme analitiği, içerik üretimi, öğrenci takibi	Çocuk verisi, profil çıkarma, eşitsizlik	Yaş hassasiyeti, veli/öğrenci hakları	Çocuk verisi özel koruma rejimi
Finans	Dolandırıcılık tespiti, kredi, risk skoru	Otomatik dışlama ve sahte işlem	Model denetimi, itiraz, veri doğruluğu	Yüksek risk + bağımsız denetim
Savunma-güvenlik	İstihbarat analitiği, görüntü işleme, karar destek	Biçimsel insan denetimi ve zincirleme hata	Durdurabilir insan denetimi ve olay kaydı	Kritik risk sınıfı
Afet ve kriz yönetimi	Erken uyarı, kaynak tahsisi, konum/veri analizi	Yanlış yönlendirme, sahte yardım bilgisi	Resmî teyit ve kriz iletişimi protokolü	Kritik kamu riski

Kaynak: KVKK (2021, 2025); NIST (2023); European Union (2024); TBMM (2026); yazarın sentezi.

Şekil 12.1. Türkiye'nin yapay zekâ egemenlik mimarisi

Katman	Egemenlik sorusu	Asgari kurumsal gereklilik
Veri	Hangi veri hangi koşulla modele girebilir?	Veri envanteri, sınıflandırma, KVKK uyumu
Model	Model yerli mi, dış mı, açık mı, kapalı mı?	Model kartı, test seti, alan bazlı doğrulama
Altyapı	İşlem nerede ve kimin denetiminde yapılıyor?	Güvenli bulut, veri merkezi, hesaplama kapasitesi
Uygulama	Çıktı hangi karara bağlanıyor?	Kullanım senaryosu, insan onayı, kayıt tutma
Denetim	Hata nasıl fark edilir ve telafi edilir?	Olay bildirimi, itiraz mekanizması, bağımsız denetim
Meşruiyet	Toplum karara neden güvenecek?	Şeffaflık, gerekçelendirme, hesap verebilirlik

Kaynak: OECD (2024); NIST (2023); KVKK (2025); Kaya (2026); yazarın sentezi.

Harita 12.1. Türkiye'nin jeoteknolojik yapay zekâ konumlanması

Yön / Alan	Türkiye açısından anlamı	Stratejik sonuç
Avrupa regülasyon alanı	AI Act ve veri koruma standartlarıyla uyum baskısı	Uyum kapasitesi + norm üretme fırsatı
ABD platform ekosistemi	Bulut, model, kurumsal yazılım ve çip bağımlılığı	Tedarik ve denetim riski
Çin teknoloji rekabeti	Açık model, donanım alternatifi ve jeopolitik denge	Çeşitlendirme ile güvenlik riski arasında denge
Türk dünyası veri/dil havzası	Türkçe ve akraba dillerde model kapasitesi	Dil egemenliği ve bölgesel norm alanı
Orta Doğu-Afrika hattı	Kamu dijitalleşmesi, savunma ve eğitim teknolojisi ihtiyacı	Türkiye merkezli güvenilir AI hizmetleri için pazar

Kaynak: Cumhurbaşkanlığı Dijital Dönüşüm Ofisi & Sanayi ve Teknoloji Bakanlığı (2021); Sanayi ve Teknoloji Bakanlığı & Cumhurbaşkanlığı Dijital Dönüşüm Ofisi (2024); European Union (2024); yazarın sentezi.

38. Türkiye İçin Stratejik Yol Haritası

Türkiye için yapay zekâ yol haritası on zorunlu politika eşiği üzerinden kurulmalıdır. Birincisi, ulusal yapay zekâ envanteri oluşturulmalıdır. Bu envanter yalnızca kamu kurumlarının kullandığı uygulamaları değil; veri kaynaklarını, model türlerini, dış servis bağımlılıklarını, karar süreçlerine etkisini ve hukuki sorumluluk zincirini içermelidir.

İkincisi, kamu kurumları için zorunlu Yapay Zekâ Etki Değerlendirmesi getirilmelidir. Her kamu kurumu, yapay zekâ sistemi kullanmadan önce kullanım amacını, veri türünü, etki alanını, risk sınıfını, insan denetimi biçimini, itiraz yolunu ve telafi mekanizmasını yazılı olarak belirlemelidir (KVKK, 2025; NIST, 2023).

Üçüncüsü, Türkçe büyük dil modeli yalnızca teknoloji projesi değil, dil egemenliği projesi olarak görülmelidir. Türkçe model kapasitesi; kamu dili, hukuk dili, afet terminolojisi, eğitim müfredatı, diplomatik söylem, tarihsel arşiv ve Türk dünyası veri havzası ile birlikte düşünülmelidir (Sanayi ve Teknoloji Bakanlığı & Cumhurbaşkanlığı Dijital Dönüşüm Ofisi, 2024).

Dördüncüsü, kamu verisi için güvenli, sınıflandırılmış ve denetlenebilir veri mimarisi kurulmalıdır. Bu mimari, hangi verinin yerli altyapı dışına çıkamayacağını, hangi verinin anonimleştirme sonrası kullanılabileceğini ve hangi verinin üretken yapay zekâ sistemlerine hiçbir şekilde girilemeyeceğini açık biçimde tanımlamalıdır.

Beşincisi, kritik alanlarda dış model kullanımı risk sınıfına bağlanmalıdır. Sağlık, savunma, güvenlik, çocuk koruma, finans, adalet, afet yönetimi ve kritik altyapı gibi alanlarda dış model veya dış bulut kullanımının sınırları ayrıca belirlenmelidir.

Altıncısı, deepfake ve sentetik içerik için ulusal doğrulama protokolü oluşturulmalıdır. Bu protokol kamu kurumlarının kriz anında hangi kanaldan, hangi süre içinde, hangi teknik işaretleme veya dijital imza sistemiyle açıklama yapacağını belirlemelidir.

Yedincisi, KVKK, telif, ceza hukuku, idare hukuku, tüketici hukuku ve kamu ihale rejimi yapay zekâ bağlamında birlikte ele alınmalıdır. Yapay zekâ tek bir hukuk alanına sığmadığı için parçalı düzenleme, uygulamada denetim boşluğu üretebilir (TBMM, 2024a, 2025, 2026).

Sekizincisi, kamu personeli için yapay zekâ okuryazarlığı zorunlu hâle getirilmelidir. Kamu personelinin üretken yapay zekâ araçlarına hangi veriyi giremeyeceğini, model çıktısına ne kadar güvenebileceğini, sahte kesinlik riskini nasıl anlayacağını ve hata durumunda nasıl kayıt tutacağını bilmesi gerekir.

Dokuzuncusu, savunma ve güvenlik uygulamalarında insan denetimi sembolik değil, durdurabilir nitelikte olmalıdır. İnsan denetimi, yalnızca prosedürel onay değil; modeli durdurma, öneriyi reddetme, veri kaynağını sorgulama ve sorumluluk zincirini izleme kapasitesi üretmelidir.

Onuncusu, Türkiye yalnızca AB AI Act'i izleyen değil; Türkçe, Türk dünyası, İslam dünyası ve Küresel Güney ekseninde insan merkezli, güvenilir ve denetlenebilir yapay zekâ normları üreten bir aktör olmalıdır. Türkiye'nin avantajı, yalnızca teknoloji ithal eden veya regülasyon takip eden bir ülke olmaktan çıkıp, farklı jeopolitik havzalar arasında güvenilir yapay zekâ mimarisi önerebilen ara-yapıcı bir aktöre dönüşebilmesidir (European Union, 2024; TBMM, 2026).

Yanlış Karar Riski Kutusu – Bölüm XII

Bu bölümden çıkarılması gereken temel uyarı şudur: Türkiye için yapay zekâda en büyük risk teknolojiyi geç kullanmak değildir; teknolojiyi yanlış mimariyle, dış model bağımlılığıyla, veri egemenliğini zayıflatarak ve kamu denetimini sonradan kurmaya çalışarak kullanmaktır. Yapay zekâ alanında acelecilik kadar gecikme de risklidir; fakat en tehlikeli senaryo, kurumsal hafızayı, kamu verisini ve karar süreçlerini denetlenemeyen dış mimarilere sessizce devretmektir.

Bölüm Sonu Değerlendirmesi

Türkiye bölümü, raporun küresel yapay zekâ analizini ulusal strateji ve egemenlik düzeyine indirmektedir. Bu bölümün temel sonucu açıktır: Türkiye'nin yapay zekâ politikası yalnızca kullanımı yaygınlaştırma, girişimcilik ekosistemini büyütme veya regülasyon uyumunu sağlama hedefleriyle sınırlı kalamaz. Asıl mesele; veri, model, altyapı, hukuk, denetim, dil ve toplumsal güven katmanlarını birlikte yöneten bir yapay zekâ egemenlik mimarisi kurmaktır.

Bu mimari kurulmadığı takdirde, yapay zekâ Türkiye için verimlilik ve inovasyon fırsatı üretmekle birlikte yeni bağımlılık biçimleri, otomatikleşmiş hata, dilsel temsil kaybı, kamu güveni aşınması ve kriz dönemlerinde Soğuk Kaos etkisi yaratabilir. Buna karşılık doğru yönetim mimarisiyle Türkiye, yalnızca yapay zekâ kullanan değil; Türkçe, bölgesel veri havzaları ve insan merkezli kamu denetimi üzerinden norm üreten bir aktöre dönüşebilir. Bu nedenle Türkiye'nin yapay zekâ eşiği, teknik bir tercih olmaktan çok stratejik egemenlik kararıdır (Kaya, 2026; TBMM, 2026).

BÖLÜM XIII – SONUÇ

Bu raporun temel hükmü açıktır: yapay zekâ artık yalnızca teknik bir yenilik alanı değildir. Veri, model, bulut, hızlandırıcı çip, enerji, platform, hukuk, kamu yönetimi ve enformasyon düzeni aynı mimari içinde birleştiği için yapay zekâ meselesi doğrudan güç, egemenlik ve meşruiyet meselesine dönüşmüştür. Stanford HAI'nin 2026 AI Index raporu, kabiliyetlerin ve benimsenmenin hızla arttığını; buna karşılık bunları ölçme, denetleme ve yönetme kapasitesinin aynı hızda gelişmediğini açık biçimde göstermektedir.

Bu nedenle rapor boyunca savunulan ana tez şudur: yapay zekâ çağında belirleyici olan soru, hangi modelin daha güçlü olduğu değildir. Asıl belirleyici olan, gücün hangi veri rejimi içinde biriktiği, hangi altyapı üzerinde ölçeklendiği, hangi platformlar aracılığıyla yayıldığı, hangi düzenleyici boşluklardan yararlandığı ve hangi kurumların bu süreçte kendi karar kapasitesini koruyabildiğidir.

Raporun özgün katkısı, bu dönüşümü yalnızca rekabet, etik veya inovasyon başlıklarıyla değil; Soğuk Kaos ve DeNormasyon kavramları üzerinden okuyarak açıklamasıdır. Soğuk Kaos, yüksek hız, düşük doğrulama süresi, sentetik içerik çoğalması ve kurumsal kırılmalıkların birleşmesiyle oluşan yönetilen belirsizlik rejimini ifade etmektedir. DeNormasyon ise hata, manipülasyon, sahte kesinlik ve doğrulama aşınmasının istisna olmaktan çıkıp yeni normal hâline gelmesini anlatmaktadır.

Buradan çıkan ikinci temel sonuç, yapay zekâ çağında tarafsızlık iddiasının giderek zayıfladığıdır. Model hizalaması, veri seçimi, güvenlik filtreleri, içerik işaretleme rejimleri, sağlayıcı mimarileri ve düzenleyici çerçeveler yalnızca teknik tercih değil; aynı zamanda normatif tercihlerdir. Mesele, teknolojinin ideolojiden bağımsız olup olmaması değil; hangi normatif tercihin hangi teknik mimari içine gömüldüğünün görünür kılınıp kılınmadığıdır.

Üçüncü temel sonuç, yapay zekâ rekabetinin görünürde dijital olmasına rağmen özünde maddi olduğudur. Veri merkezleri, elektrik talebi, soğutma, şebeke kapasitesi ve kritik tedarik zinciri erişimi artık yapay zekâ stratejisinin dışsal koşulları değil, bizzat çekirdek belirleyicileridir. Bu nedenle yapay zekâ stratejisi, enerji ve altyapı stratejisinden ayrı yazılamaz.

Dördüncü temel sonuç, devletler ve kurumlar için asıl riskin teknolojik geri kalmışlık değil, yanlış mimariye erken bağlanmak olduğudur. Dış model ve platform bağımlılığı, sentetik medya kaynaklı teyit çöküşü, ajanik sistemlerde yetki aşımı, kamu hizmetlerinde hak ihlali, savunma alanında biçimsel insan denetimi ve altyapı darboğazlarının tümü aynı stratejik soruna bağlanmaktadır: hazırlık açığı.

Sonuç olarak yapay zekâ çağında mesele teknolojiye erişim değil, teknoloji üzerinden yeniden kurulan düzenin niteliğidir. Kısa vadeli kazanımlar uzun vadeli bağımlılık, meşruiyet kaybı ve kurumsal körlük üretebilir. Buna karşılık doğru strateji; tedarik çeşitliliği, içerik provenansı, insan denetimi, veri sınırı, olay raporlama, enerji-altyapı planlaması ve yüksek etkili alanlarda sıkı yönetim mimarisini aynı çerçevede kurmaktır.

Nihai hüküm şudur: yapay zekâ çağında en büyük risk, makinenin çok güçlenmesi değildir. En büyük risk, kurumların ve devletlerin kendi muhakeme, denetim ve meşruiyet zeminlerini fark edilmeden dışsal sistemlerin görünmez tercihleri içine gömmesidir. Bu nedenle yapay zekâ çağının asıl sorusu 'Ne kadar ileri gidebiliriz?' değil, 'İleri giderken neyi kaybetmemeliyiz?' sorusudur. Bu raporun cevabı nettir: hakikat zemini, kurumsal akıl, denetim kapasitesi ve egemen karar hakkı kaybedilmemelidir.

Bu raporun temel bulgusu şudur: yapay zekâ artık yalnızca teknik bir araç, yazılım ürünü ya da verimlilik çarpanı olarak açıklanamaz. Güncel yapay zekâ düzeni; veri akışları, hesaplama gücü, çip mimarileri, bulut altyapısı, model laboratuvarları, kamu-satın alma rejimleri, güvenlik kaygıları, enerji baskısı ve enformasyon düzeni arasında kurulan çok katmanlı bir güç mimarisi üretmektedir (OECD, 2024; NIST, 2023; Maslej ve diğerleri, 2026). Bu nedenle yapay zekâ tartışması, yalnızca teknoloji politikası değil; aynı zamanda egemenlik, kamu kapasitesi, stratejik bağımlılık ve toplumsal güven tartışmasıdır.

Raporun küresel envanter bulguları, yapay zekâ ekosisteminin görünürde çoğul, gerçekte ise yüksek ölçüde yoğunlaşmış olduğunu göstermektedir. Açık model ve ürün çoğalması, gücün gerçekten dağıldığı anlamına gelmemektedir; frontier kapasite hâlen sınırlı sayıdaki laboratuvar, bulut sağlayıcısı ve hızlandırıcı çip üreticisi etrafında birikmektedir. Bu yoğunlaşma ABD ile Çin arasında asimetrik ama belirgin bir rekabet üretirken, Avrupa daha çok norm ve düzenleme üreticisi olarak öne çıkmaktadır; Küresel Güney ise veri, pazar ve bağımlılık eksenlerinde yapısal baskılarla karşı karşıya kalmaktadır (Maslej ve diğerleri, 2025; Maslej ve diğerleri, 2026; UNCTAD, 2021).

Teknik düzeyde en kritik sonuç, yapay zekâ sistemlerinin yalnızca yetenek üzerinden değil; hata rejimi, halüsinasyon eğilimi, veri yönetimi, mahremiyet riski ve model davranışını şekillendiren kurumsal seçimler üzerinden değerlendirilmesi gerektiğidir. Yanlış fakat ikna edici çıktıların, özellikle kamu yönetimi, hukuk, eğitim, sağlık, finans ve güvenlik alanlarında yüksek toplumsal maliyet doğurabileceği açıktır. Aynı şekilde sentetik medya, deepfake, otomatik içerik çoğaltma ve bilişsel harp uygulamaları, yapay zekâyı doğrudan enformasyon güvenliği ve demokratik meşruiyet krizlerinin merkezine taşımaktadır (OECD, 2026a; Rid, 2020; Pomerantsev, 2019).

Enerji ve altyapı boyutu ise bu çağın en fazla hafife alınan katmanıdır. Yapay zekâ yarışının sürdürülebilirliği, yalnızca algoritmik ilerlemeye değil; veri merkezleri, soğutma sistemleri, elektrik şebekeleri ve kritik donanım tedarikine bağlıdır. Bu nedenle yapay zekâ politikası ile enerji politikası, sanayi politikası ve ulusal güvenlik politikası giderek birbirine bağlanmaktadır (IEA, 2025). Bu bağın görmezden gelinmesi, teknoloji stratejilerinin gerçek maliyet ve kırılabilirliklerini görünmez kılar.

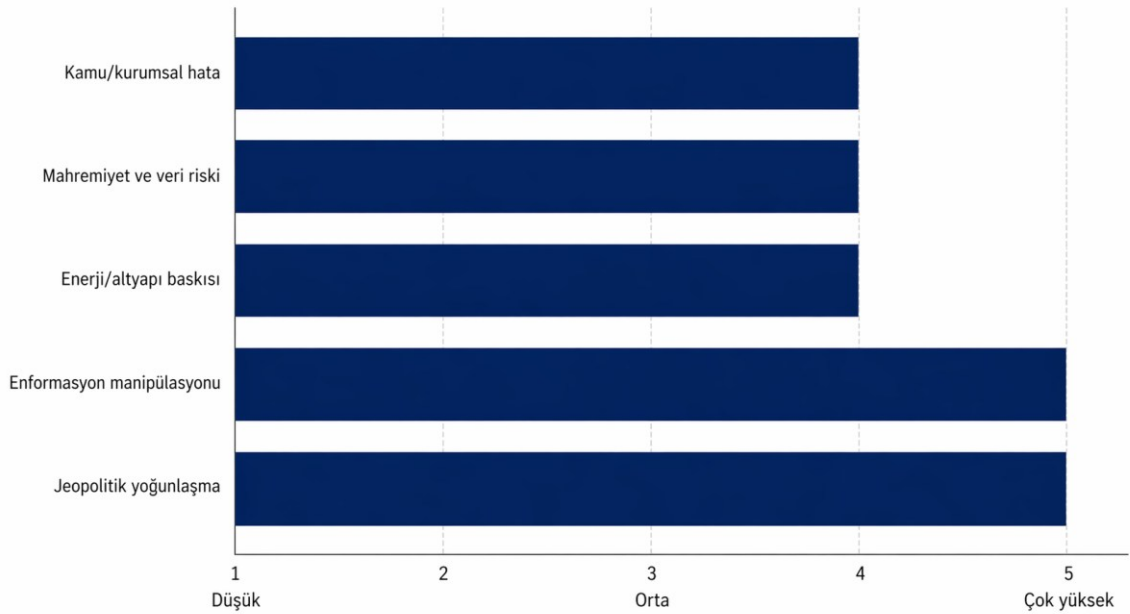
Normatif ve kurumsal açıdan ulaşılan nihai yargı nettir: yapay zekâ çağında tarafsızlık yoktur. Hangi modelin seçildiği, hangi verinin kullanıldığı, hangi altyapının finanse edildiği, hangi regülasyonun benimsendiği ve insan denetiminin nasıl kurulduğu; sistemlerin teknik performansından çok daha geniş sonuçlar doğurur. Bu nedenle devletler için öncelik; egemenlik, güvenlik, veri yönetimi ve enformasyon bütünlüğünü aynı stratejik çerçevede ele almaktır. Kurumlar için öncelik ise model tedariki, veri koruma, denetlenebilirlik, insan gözetimi ve kriz dayanıklılığı temelinde çok katmanlı bir yapay zekâ yönetimi kurmaktır (UNESCO, 2021; European Union, 2024; OECD, 2026b). Sonuç olarak yapay zekâ çağında asıl mesele yalnızca daha güçlü sistemler üretmek değil; daha hesap verebilir, daha dayanıklı ve daha adil bir teknolojik düzen kurabilmektir.

Tablo 13.1. Raporun temel stratejik yargıları

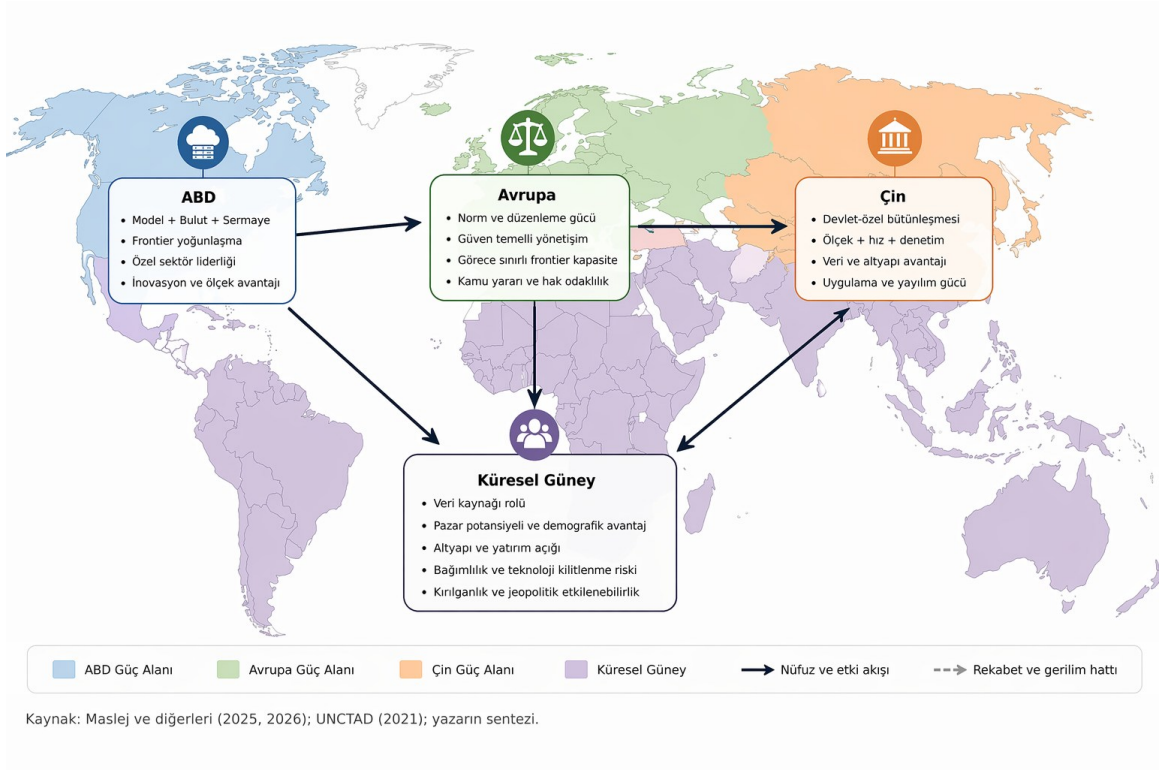
Ana bulgu	Politik/kurumsal anlamı	Öncelikli yanıt
Yapay zekâ tarafsız değildir	Teknik tercihler aynı zamanda güç tercihidir	Tedarik, veri ve model yönetiřimi
Güç birkaç merkezde yoğunlařmaktadır	Egemenlik ve bağımlılık riski artar	Ulusal/kurumsal kapasite ve çoklu tedarik
Deepfake ve otomatik içerik çoğalması büyüyor	Enformasyon güvenliğı kırılanlařır	Doğrulama, medya okuryazarlığı ve kriz protokolleri
Hata ve halüsinasyon riski kurumsaldır	Yanlış karar maliyetleri yüksektir	İnsan denetimi ve görev sınırlaması
Enerji ve çip katmanı belirleyicidir	AI stratejisi aynı zamanda altyapı stratejisidir	Veri merkezi, enerji ve donanım planlaması
Düzenleme dışsal değıl içseldir	Yönetişim sistem davranışını değıřtirir	Risk temelli uyum ve denetlenebilirlik

Kaynak: OECD (2024, 2026a, 2026b); NIST (2023); IEA (2025); Maslej ve diğeri (2025, 2026); yazarın sentezi.

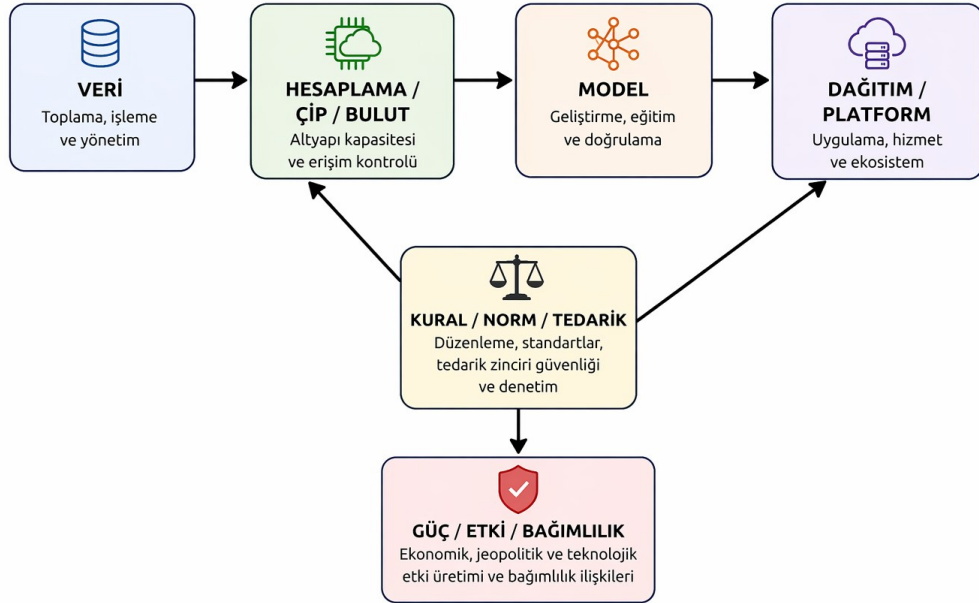
Grafik 13.1. Yapay zekâ çağında başlıca stratejik baskı alanları



Harita 13.1. Yapay zekâ düzeninin kavramsal güç coğrafyası



Şema 13.1. Yapay zekâ egemenlik döngüsü



Şekil 13.1. “Taraırsızlık yoktur” matrisi



Kaynak: OECD (2026); UNESCO (2021); European Commission (2025); yazarın sentezi.

GENEL KAYNAKÇA

- Amodei, D., Olah, C., Steinhardt, J., Christiano, P., Schulman, J., & Mané, D. (2016). Concrete problems in AI safety. arXiv. <https://arxiv.org/abs/1606.06565>
- Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosselut, A., Brunskill, E., Brynjolfsson, E., et al. (2021). On the opportunities and risks of foundation models. Stanford Center for Research on Foundation Models. <https://crfm.stanford.edu/report.html>
- Kaya, Y. (2026). Soğuk Kaos ve DeNormasyon: Yapay zekâ çağında kurumsal enformasyon erozyonuna ilişkin kavramsal çerçeve [Yayımlanmamış kavramsal çerçeve].
- Bostrom, N. (2014). Superintelligence: Paths, dangers, strategies. Oxford University Press.
- Couldry, N., & Mejias, U. A. (2019). The costs of connection: How data is colonizing human life and appropriating it for capitalism. Stanford University Press.
- Dennett, D. C. (1991). Consciousness explained. Little, Brown and Company.
- European Commission. (2025). AI Act: Regulatory framework and application timeline. <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>
- European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union.
- Future of Life Institute. (2017). Asilomar AI principles. <https://futureoflife.org/open-letter/ai-principles/>
- Hugging Face. (2026). State of open source on Hugging Face: Spring 2026. <https://huggingface.co/blog/huggingface/state-of-os-hf-spring-2026>
- IEA. (2025). Energy and AI. International Energy Agency. <https://www.iea.org/reports/energy-and-ai>
- LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. Nature, 521, 436–444. <https://doi.org/10.1038/nature14539>
- Maslej, N., Clark, J., Lyons, T., Manyika, J., Niebles, J. C., Etchemendy, J., Shoham, Y., Wald, R., et al. (2025). The AI Index 2025 annual report. Stanford Institute for Human-Centered Artificial Intelligence. <https://hai.stanford.edu/ai-index/2025-ai-index-report>
- Maslej, N., Clark, J., Lyons, T., Manyika, J., Niebles, J. C., Etchemendy, J., Shoham, Y., Wald, R., et al. (2026). The AI Index 2026 annual report. Stanford Institute for Human-Centered Artificial Intelligence. <https://hai.stanford.edu/ai-index/2026-ai-index-report>
- McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (1956). A proposal for the Dartmouth summer research project on artificial intelligence. Dartmouth College.
- NIST. (2023). Artificial intelligence risk management framework (AI RMF 1.0). National Institute of Standards and Technology. <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>
- OECD. (2024). Updated OECD definition of an AI system. Organisation for Economic Co-operation and Development. <https://oecd.ai/en/dashboards/policy-initiatives/updated-oecd-definition-of-an-ai-system-4108>

- OECD. (2026a). Trends in AI incidents and hazards reported by the media. Organisation for Economic Co-operation and Development. https://www.oecd.org/en/publications/trends-in-ai-incidents-and-hazards-reported-by-the-media_4f5ff43c-en.html
- OECD. (2026b). OECD due diligence guidance for responsible AI. Organisation for Economic Co-operation and Development. https://www.oecd.org/en/publications/oecd-due-diligence-guidance-for-responsible-ai_41671712-en.html
- Pomerantsev, P. (2019). This is not propaganda: Adventures in the war against reality. Faber & Faber.
- Rid, T. (2020). Active measures: The secret history of disinformation and political warfare. Farrar, Straus and Giroux.
- Simon, H. A. (1996). The sciences of the artificial (3rd ed.). MIT Press.
- State Council of the People's Republic of China. (2017). New generation artificial intelligence development plan. http://www.gov.cn/zhengce/content/2017-07/20/content_5211996.htm
- Turing, A. M. (1950). Computing machinery and intelligence. *Mind*, 59(236), 433–460. <https://doi.org/10.1093/mind/LIX.236.433>
- UNCTAD. (2021). Digital economy report 2021: Cross-border data flows and development. United Nations Conference on Trade and Development. https://unctad.org/system/files/official-document/der2021_en.pdf
- UNESCO. (2021). Recommendation on the ethics of artificial intelligence. United Nations Educational, Scientific and Cultural Organization. <https://unesdoc.unesco.org/ark:/48223/pf0000381137>
- U.S. Department of Defense. (2022). Responsible artificial intelligence strategy and implementation pathway. <https://media.defense.gov/2022/Jun/22/2003028059/-1/-1/1/RESPONSIBLE-ARTIFICIAL-INTELLIGENCE-STRATEGY-AND-IMPLEMENTATION-PATHWAY.PDF>
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2017). Attention is all you need. In *Advances in Neural Information Processing Systems* 30. <https://papers.nips.cc/paper/7181-attention-is-all-you-need.pdf>
- Wiener, N. (1948). *Cybernetics: Or control and communication in the animal and the machine*. MIT Press.
- World Economic Forum. (2026). The global risks report 2026. <https://www.weforum.org/publications/global-risks-report-2026/digest/>
- IEA. (2025). Energy and AI. <https://www.iea.org/reports/energy-and-ai>
- OECD. (2026). OECD due diligence guidance for responsible AI. https://www.oecd.org/en/publications/oecd-due-diligence-guidance-for-responsible-ai_41671712-en.html
- OECD. (2024). Updated OECD definition of an AI system. OECD.AI. <https://oecd.ai/en/dashboards/policy-initiatives/updated-oecd-definition-of-an-ai-system-4108>
- Maslej, N., Clark, J., Lyons, T., Manyika, J., Niebles, J. C., Etchemendy, J., Shoham, Y., Wald, R., et al. (2025). The AI Index 2025 annual report. Stanford Institute for Human-Centered Artificial Intelligence. https://hai.stanford.edu/assets/files/hai_ai_index_report_2025.pdf

Maslej, N., Clark, J., Lyons, T., Manyika, J., Niebles, J. C., Etchemendy, J., Shoham, Y., Wald, R., et al. (2026). The AI Index 2026 annual report. Stanford Institute for Human-Centered Artificial Intelligence. https://hai.stanford.edu/assets/files/ai_index_report_2026.pdf

State Council of the People's Republic of China. (2017). New generation artificial intelligence development plan. https://english.www.gov.cn/archive/state_council_gazette/2017/08/10/content_281475782151912.htm

UNCTAD. (2021). Digital economy report 2021: Cross-border data flows and development: For whom the data flow. United Nations Conference on Trade and Development. https://unctad.org/system/files/official-document/der2021_en.pdf

AMD. (2023a). AMD Instinct MI300X accelerators. <https://www.amd.com/en/products/accelerators/instinct/mi300/mi300x.html>

Amazon. (2023a). AWS unveils next-generation AWS-designed chips. <https://press.aboutamazon.com/2023/11/aws-unveils-next-generation-aws-designed-chips>

Anthropic. (2024a). Building effective agents. <https://www.anthropic.com/research/building-effective-agents>

Anthropic. (2024b). Introducing the Model Context Protocol. <https://www.anthropic.com/news/model-context-protocol>

Anthropic. (2025a). Claude Gov models for U.S. national security customers. <https://www.anthropic.com/news/claude-gov-models-for-u-s-national-security-customers>

Bommasani, R., Hudson, D. A., Adeli, E., et al. (2021). On the opportunities and risks of foundation models. <https://arxiv.org/abs/2108.07258>

Cerebras. (2024a). Product - chip: WSE-3. <https://www.cerebras.ai/chip>

Google Cloud. (2024a). Trillium TPU is GA. <https://cloud.google.com/blog/products/compute/trillium-tpu-is-ga>

Groq. (2025a). What is a Language Processing Unit? <https://groq.com/blog/the-groq-lpu-explained>

Huawei. (2025a). Atlas 900 A3 SuperPoD. <https://e.huawei.com/cn/products/computing/ascend/atlas-900-a3-superpod>

IEA. (2026). Key questions on energy and AI: Executive summary. <https://www.iea.org/reports/key-questions-on-energy-and-ai/executive-summary>

Intel. (2024a). Intel unveils next-generation AI solutions with the launch of Intel Gaudi 3. <https://newsroom.intel.com/artificial-intelligence/next-generation-ai-solutions-xeon-6-gaudi-3>

Maslej, N., Clark, J., Lyons, T., Manyika, J., Niebles, J. C., Etchemendy, J., Shoham, Y., Wald, R., et al. (2026). The AI Index 2026 annual report. <https://hai.stanford.edu/ai-index/2026-ai-index-report>

Microsoft. (2026a). Microsoft 365 Copilot in U.S. government cloud environments. <https://learn.microsoft.com/en-us/microsoft-365/copilot/gov-overview>

NVIDIA. (2024a). NVIDIA Blackwell architecture. <https://www.nvidia.com/en-us/data-center/technologies/blackwell-architecture/>

NVIDIA. (2025a). What is retrieval-augmented generation? <https://www.nvidia.com/en-us/glossary/retrieval-augmented-generation/>

OpenAI. (2024a). GPT-4o system card. <https://openai.com/index/gpt-4o-system-card/>

OpenAI. (2024b). Hello GPT-4o. <https://openai.com/index/hello-gpt-4o/>

OpenAI. (2025a). Introducing 4o image generation. <https://openai.com/index/introducing-4o-image-generation/>

OpenAI. (2025b). Introducing OpenAI for Government. <https://openai.com/global-affairs/introducing-openai-for-government/>

Vaswani, A., Shazeer, N., Parmar, N., et al. (2017). Attention is all you need. https://papers.neurips.cc/paper_files/paper/2017/hash/3f5ee243547dee91fbd053c1c4a845aa-Abstract.html

Couldry, N., & Mejias, U. A. (2019). The costs of connection: How data is colonizing human life and appropriating it for capitalism. Stanford University Press. <https://www.sup.org/books/sociology/costs-connection>

European Commission. (2020). A European strategy for data. <https://digital-strategy.ec.europa.eu/en/policies/strategy-data>

OECD. (2024). Artificial intelligence, data and competition. https://www.oecd.org/en/publications/2024/05/artificial-intelligence-data-and-competition_9d0ac766.html

UNCTAD. (2021). Digital economy report 2021: Cross-border data flows and development: For whom the data flow. https://unctad.org/system/files/official-document/der2021_en.pdf

UNESCO. (2021). Recommendation on the ethics of artificial intelligence. <https://www.unesco.org/en/artificial-intelligence/recommendation-ethics>

Zuboff, S. (2019). The age of surveillance capitalism: The fight for a human future at the new frontier of power. PublicAffairs. <https://www.publicaffairsbooks.com/titles/shoshana-zuboff/the-age-of-surveillance-capitalism/9781610395694/>

Anthropic. (2026). Claude Mythos preview system card. <https://www.anthropic.com/claude-mythos-preview-system-card>

Dahl, M., Magesh, V., Suzgun, M., & Ho, D. E. (2024). Large legal fictions: Profiling legal hallucinations in large language models. arXiv. <https://arxiv.org/abs/2401.01301>

European Commission. (2026). Navigating the AI Act. <https://digital-strategy.ec.europa.eu/en/faqs/navigating-ai-act>

NIST. (2024). Artificial intelligence risk management framework: Generative artificial intelligence profile. National Institute of Standards and Technology. <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>

NIST. (2025). Assessing risks and impacts of AI (ARIA): Pilot evaluation report. National Institute of Standards and Technology. <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.700-2.pdf>

OECD. (2024). Assessing potential future artificial intelligence risks, benefits and policy imperatives. https://www.oecd.org/en/publications/assessing-potential-future-artificial-intelligence-risks-benefits-and-policy-imperatives_3f4e3dfb-en.html

OECD. (2025). AI and the future of social protection in OECD countries. https://www.oecd.org/content/dam/oecd/en/publications/reports/2025/06/ai-and-the-future-of-social-protection-in-oecd-countries_038f49ed/7b245f7e-en.pdf

OECD. (2026a). Trends in AI incidents and hazards reported by the media. https://www.oecd.org/content/dam/oecd/en/publications/reports/2026/02/trends-in-ai-incidents-and-hazards-reported-by-the-media_7c824ca9/4f5ff43c-en.pdf

OpenAI. (2025). GPT-4.5 system card. <https://openai.com/index/gpt-4-5-system-card/>

Stanford HAI. (2024). AI on trial: Legal models hallucinate in 1 out of 6 (or more) benchmarking queries. <https://hai.stanford.edu/news/ai-trial-legal-models-hallucinate-1-out-6-or-more-benchmarking-queries>

C2PA. (2023). C2PA explainer. https://spec.c2pa.org/specifications/specifications/1.3/explainer/_attachments/Explainer.pdf

Europol. (2024). Facing reality? Law enforcement and the challenge of deepfakes: An Observatory Report from the Europol Innovation Lab. https://www.europol.europa.eu/cms/sites/default/files/documents/Europol_Innovation_Lab_Facing_Reality_Law_Enforcement_And_The_Challenge_Of_Deepfakes.pdf

NATO STO. (2025). Cognitive warfare. <https://www.sto.nato.int/document/cognitive-warfare/>

NATO StratCom COE. (t.y.). AI Laboratory. <https://stratcomcoe.org/projects/ai-laboratory/5>

OECD. (2026b). OECD due diligence guidance for responsible AI. https://www.oecd.org/en/publications/2026/02/oecd-due-diligence-guidance-for-responsible-ai_7831bb49/full-report/component-4.html

Pomerantsev, P. (2019). This is not propaganda: Adventures in the war against reality. PublicAffairs.

UNESCO, & UNDP. (2025). Freedom of expression, artificial intelligence and elections. <https://www.undp.org/publications/freedom-expression-artificial-intelligence-and-elections>

Anthropic. (2025b). Anthropic and the Department of Defense to advance responsible AI in defense operations. <https://www.anthropic.com/news/anthropic-and-the-department-of-defense-to-advance-responsible-ai-in-defense-operations>

CDAO. (2025a). CDAO announces partnerships with frontier AI companies to address national security mission areas. <https://www.ai.mil/latest/news-press/pr-view/article/4242822/cdao-announces-partnerships-with-frontier-ai-companies-to-address-national-secu/>

CDAO. (2025b). Chief Digital and Artificial Intelligence Office. <https://www.ai.mil/>

DoD CIO. (2025). DoD Artificial Intelligence cybersecurity risk management tailoring guide. <https://dodcio.defense.gov/Portals/0/Documents/Library/AI-CybersecurityRMTailoringGuide.pdf>

ICRC. (2021). ICRC position on autonomous weapon systems. <https://www.icrc.org/en/document/icrc-position-autonomous-weapon-systems>

ICRC. (2026). Autonomous weapon systems and international humanitarian law: Selected issues. https://www.icrc.org/sites/default/files/2026-03/4896_002_Autonomous_Weapons_Systems_-_IHL-ICRC.pdf

NATO. (2021). Summary of the NATO Artificial Intelligence Strategy. <https://www.nato.int/en/about-us/official-texts-and-resources/official-texts/2021/10/22/summary-of-the-nato-artificial-intelligence-strategy>

NATO. (2024). Summary of NATO's revised Artificial Intelligence (AI) strategy. <https://www.nato.int/en/about-us/official-texts-and-resources/official-texts/2024/07/10/summary-of-natos-revised-artificial-intelligence-ai-strategy>

OpenAI. (2025). Introducing OpenAI for Government. <https://openai.com/global-affairs/introducing-openai-for-government/>

U.S. Department of Defense. (2022). Responsible Artificial Intelligence strategy and implementation pathway. <https://media.defense.gov/2022/Jun/22/2003022604/-1/-1/0/Department-of-Defense-Responsible-Artificial-Intelligence-Strategy-and-Implementation-Pathway.PDF>

U.S. Department of Defense. (2023). DoD Directive 3000.09: Autonomy in weapon systems. <https://media.defense.gov/2023/Jan/25/2003149928/-1/-1/0/DODDIRECTIVE-3000.09-AUTONOMY-IN-WEAPON-SYSTEMS.PDF>

UNIDIR. (2025). The interpretation and application of international humanitarian law in relation to lethal autonomous weapon systems. <https://unidir.org/publication/the-interpretation-and-application-of-international-humanitarian-law-in-relation-to-lethal-autonomous-weapon-systems/>

United Nations. (2024). Lethal autonomous weapons systems: Report of the Secretary-General (A/79/88). <https://undocs.org/Home/Mobile?DeviceType=Desktop&FinalSymbol=a%2F79%2F88&LangRequested=False&Language=E>

United Nations Security Council. (2025). Security Council meeting record S/PV.10005. <https://undocs.org/en/S/PV.10005>

Council of Europe. (2024). The Framework Convention on Artificial Intelligence and human rights, democracy and the rule of law. <https://www.coe.int/en/web/artificial-intelligence/the-framework-convention-on-artificial-intelligence>

European Commission. (2025a). AI Act: Regulatory framework and application timeline. <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>

European Commission. (2025b). The General-Purpose AI Code of Practice. <https://digital-strategy.ec.europa.eu/en/policies/contents-code-gpai>

European Commission. (2025c). Guidelines for providers of general-purpose AI models. <https://digital-strategy.ec.europa.eu/en/policies/guidelines-gpai-providers>

NIST. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0). <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-ai-rmf-10>

NIST. (2024). Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile. <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

NIST. (2025). Adversarial machine learning: A taxonomy and terminology of attacks and mitigations. <https://csrc.nist.gov/pubs/ai/100/2/e2025/final>

OECD. (2019). AI principles. <https://www.oecd.org/en/topics/ai-principles.html>

State Council of the People's Republic of China. (2023). Interim Measures for Administration of Generative Artificial Intelligence Services. https://english.www.gov.cn/archive/statecouncilgazette/202308/30/content_WS64ee8c0cc6d0868f4e8deeca.html

UK Government. (2023). AI regulation: A pro-innovation approach. <https://www.gov.uk/government/publications/ai-regulation-a-pro-innovation-approach>

UK Government. (2025a). AI Opportunities Action Plan. <https://www.gov.uk/government/publications/ai-opportunities-action-plan>

UK Government. (2026). AI Opportunities Action Plan: One Year On. <https://www.gov.uk/government/publications/ai-opportunities-action-plan-one-year-on>

UNESCO. (2021). Recommendation on the ethics of artificial intelligence. <https://www.unesco.org/en/legal-affairs/recommendation-ethics-artificial-intelligence>

UNESCO. (2023). Artificial intelligence: UNESCO publishes policy paper on AI foundation models. <https://www.unesco.org/en/articles/artificial-intelligence-unesco-publishes-policy-paper-ai-foundation-models>

Anthropic. (2025c). Offering expanded Claude access across all three branches of government. <https://www.anthropic.com/news/offering-expanded-claude-access-across-all-three-branches-of-government>

CDAO. (2025). CDAO announces partnerships with frontier AI companies to address national security mission areas. <https://www.ai.mil/latest/news-press/pr-view/article/4242822/cdao-announces-partnerships-with-frontier-ai-companies-to-address-national-security-mission-areas>

DeepSeek. (2025). DeepSeek Privacy Policy. <https://cdn.deepseek.com/policies/en-US/deepseek-privacy-policy-2025-02-14.html>

DeepSeek. (2026). DeepSeek Privacy Policy. <https://cdn.deepseek.com/policies/en-US/deepseek-privacy-policy.html>

European Commission. (2026a). Commission investigates Grok and X's recommender systems under the Digital Services Act. https://ec.europa.eu/commission/presscorner/detail/en/ip_26_203

GitHub. (2026a). Updates to GitHub Copilot interaction data usage policy. <https://github.blog/news-insights/company-news/updates-to-github-copilot-interaction-data-usage-policy/>

GitHub. (2026b). Managing GitHub Copilot policies as an individual subscriber. <https://docs.github.com/copilot/how-tos/manage-your-account/managing-copilot-policies-as-an-individual-subscriber>

OpenAI. (2023). March 20 ChatGPT outage: Here's what happened. <https://openai.com/index/march-20-chatgpt-outage/>

OpenAI. (2025a). Introducing ChatGPT Gov. <https://openai.com/global-affairs/introducing-chatgpt-gov/>

OpenAI. (2025c). Strengthening America's AI leadership with the U.S. National Laboratories. <https://openai.com/index/strengthening-americas-ai-leadership-with-the-us-national-laboratories/>

Palantir. (2024). Anthropic and Palantir partner to bring Claude AI models to AWS for U.S. Government intelligence and defense operations. <https://investors.palantir.com/news->

details/2024/Anthropic-and-Palantir-Partner-to-Bring-Claude-AI-Models-to-AWS-for-U.S.-Government-Intelligence-and-Defense-Operations/

Wiz Research. (2025). Wiz Research uncovers exposed DeepSeek database leaking sensitive information, including chat history. <https://www.wiz.io/blog/wiz-research-uncovers-exposed-deepseek-database-leak>

C2PA. (2026). FAQs. <https://c2pa.org/faqs/>

Center for Security and Emerging Technology. (2025). Harnessing digital content provenance for a secure and trustworthy information ecosystem. <https://cset.georgetown.edu/>

Cybersecurity and Infrastructure Security Agency. (2025). Artificial intelligence. <https://www.cisa.gov/ai>

Federal Communications Commission. (2024, September 26). FCC fines man behind election interference scheme \$6 million for sending illegal robocalls that used deepfake generative AI technology. <https://docs.fcc.gov/public/attachments/DOC-405811A1.pdf>

Federal Communications Commission. (2024). Federal Communications Commission FCC 24-104. <https://docs.fcc.gov/public/attachments/FCC-24-104A1.pdf>

Government of Hong Kong. (2025). Panel on Security special meeting: Background brief on law and order situation in Hong Kong. <https://www.legco.gov.hk/yr2025/english/panels/se/papers/se20250211cb2-192-2-e.pdf>

Hong Kong Police Force. (2025). Crime statistics in detail. https://www.police.gov.hk/ppp_en/09_statistics/csd.html

New Hampshire Department of Justice. (2024, February 6). Voter suppression AI robocall investigation update. <https://www.doj.nh.gov/news-and-media/voter-suppression-ai-robocall-investigation-update>

OECD. (2026). The agentic AI landscape and its conceptual foundations. https://www.oecd.org/en/publications/the-agentic-ai-landscape-and-its-conceptual-foundations_396cf758-en.html

UNESCO. (2024). Governing AI for humanity: Elections and freedom of expression policy materials. <https://unesdoc.unesco.org/>

Türkiye Bölümü İçin Ek Kaynaklar

Cumhurbaşkanlığı Dijital Dönüşüm Ofisi & T.C. Sanayi ve Teknoloji Bakanlığı. (2021). Ulusal Yapay Zekâ Stratejisi 2021–2025. <https://bilgem.tubitak.gov.tr/wp-content/uploads/sites/8/TR-UlusalYZStratejisi2021-2025.pdf>

Cumhurbaşkanlığı Strateji ve Bütçe Başkanlığı. (2023). On İkinci Kalkınma Planı (2024–2028). https://www.sbb.gov.tr/wp-content/uploads/2023/11/On-Ikinci-Kalkinma-Plani_2024-2028_17112023.pdf

Kişisel Verileri Koruma Kurumu. (2021). Yapay zekâ alanında kişisel verilerin korunmasına dair tavsiyeler. <https://www.kvkk.gov.tr/SharedFolderServer/CMSFiles/25a1162f-0e61-4a43-98d0-3e7d057ac31a.pdf>

Kişisel Verileri Koruma Kurumu. (2025). Üretken yapay zekâ ve kişisel verilerin korunması rehberi (15 soruda). <https://www.kvkk.gov.tr/Icerik/8547/uretken-yapay-zeka-ve-kisisel-verilerin-korunmasi-rehberi-15-soruda>

T.C. Sanayi ve Teknoloji Bakanlığı & Cumhurbaşkanlığı Dijital Dönüşüm Ofisi. (2024). Ulusal Yapay Zekâ Stratejisi 2024–2025 Eylem Planı. <https://www.sanayi.gov.tr/assets/pdf/UlusalYapayZekaStratejisi2024-2025EylemPlani.pdf>

Türkiye Büyük Millet Meclisi. (2024a). Yapay Zeka Kanun Teklifi (Esas No: 2/2234). <https://www.tbmm.gov.tr/Yasama/KanunTeklifi/e21539a0-888a-4500-81be-01904a918c53>

Türkiye Büyük Millet Meclisi. (2025). Bazı Kanunlarda Değişiklik Yapılmasına İlişkin Kanun Teklifi (Esas No: 2/3358). <https://www.tbmm.gov.tr/Yasama/KanunTeklifi/12d348f9-77ee-4f09-8b78-019a5e27521f>

Türkiye Büyük Millet Meclisi. (2026). Yapay zekânın kazanımlarına yönelik atılacak adımların belirlenmesi, bu alanda hukuki altyapının oluşturulması ve yapay zekâ kullanımının barındırdığı risklerin önlenmesine ilişkin Meclis Araştırması Komisyonu raporu (Sıra Sayısı: 260). <https://cdn.tbmm.gov.tr/KKBSPublicFile/D28/Y1/T10/DosyaKomisyonRaporunuVerdi/9f0e7abf-41f6-4133-ab3d-4879113f7f9f.pdf>

EKLER

Ek 1. Geniřletilmiř envanter tabloları, vaka katalogları ve karřılařtırmalı matrisler bu ana raporun ek alıřma paketlerinde ayrıntılandırılacaktır. Bu srm, tam-spektrum ana metni ve seilmiř grsel-analitik seti ieren ana gvdeyi sunmaktadır.