

YZ KULLANARAK

İÇERİK ÜRETİMİNDE DİKKAT EDİLMESİ GEREKENLER



SUNUMUN İÇERİĞİ



01

YZ hiyerarşisi, Generative AI, LLM modelleri ve YZ kullanımının sağladığı avantajlar.

02

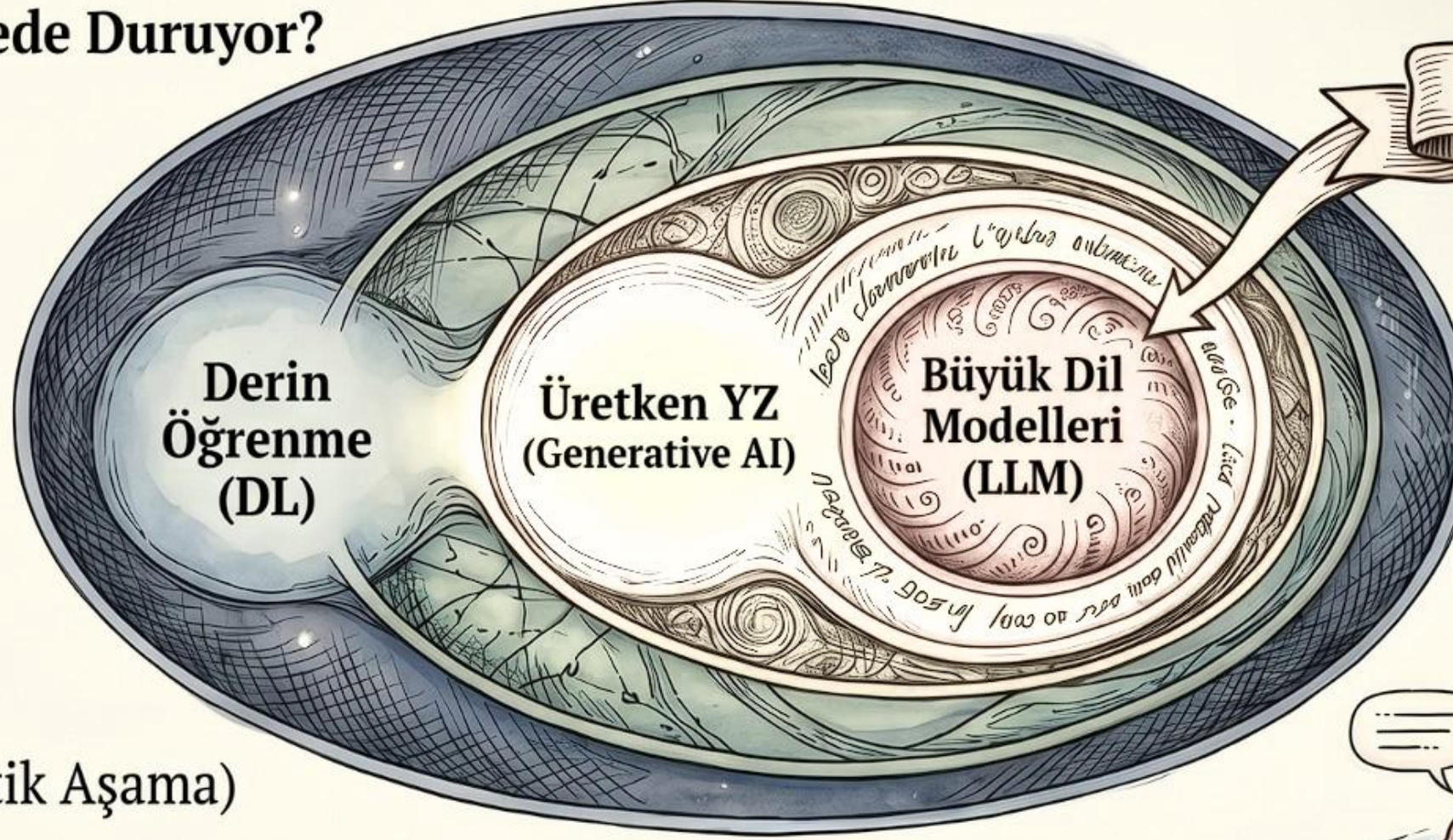
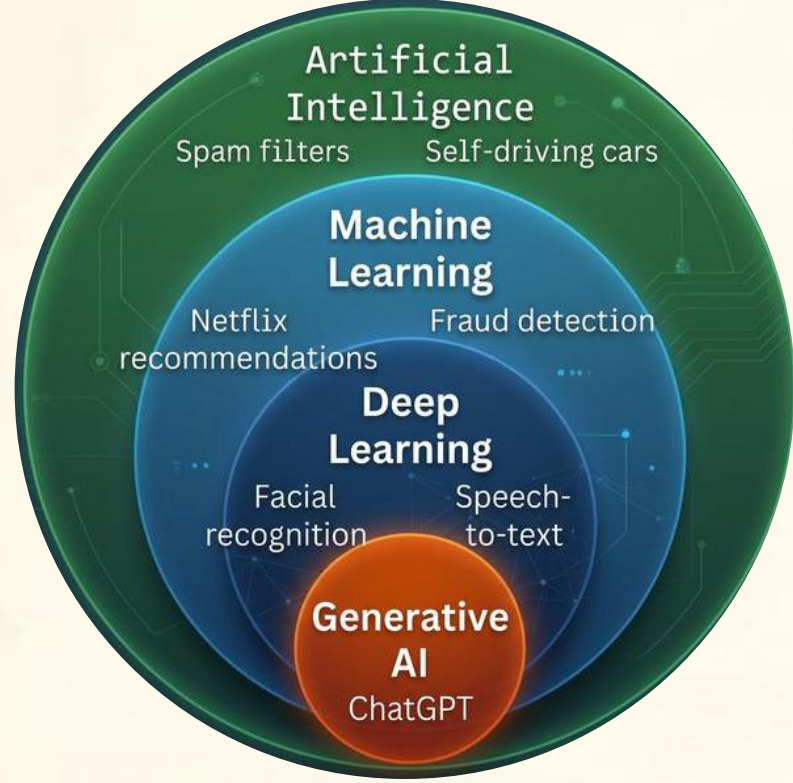
YZ halüsinasyon fenomeni ve insan faktörleri bağlamında sorgulamadan YZ kullanımının yol açabileceği bilişsel problemler.

03

YZ kullanım senaryoları, içerik üretimi ve dikkat edilmesi gerekenler.

Kelime Tahmininden Zekaya: LLM'lerin Yapay Zeka Hiyerarşisindeki Yeri ve Eğitimi

Yapay Zeka Evreni: Kim Nerede Duruyor?



LLM: Üretken YZ'nin Metin Ustası

LLM'ler, Derin Öğrenme tekniklerini kullanarak metinler arasındaki karmaşık kalıpları çözen en içteki uzman kümedir.



GPT Bir LLM Örneğidir

GPT modelleri, "Generative Pre-trained Transformer" mimarisini kullanan popüler büyük dil modelleridir.

Bir LLM Nasıl Eğitilir? (3 Kritik Aşama)



1. Ön Eğitim (Pre-training)

İnternetteki milyarlarca kelimeyi okuyarak "bir sonraki kelimeyi tahmin etme" yeteneği kazanır.

2. Komut Ayarlama (Instruction Tuning)

Modelin sadece kelime tamamlamak yerine, verilen talimatları anlaması için yapılan ince ayardır.

3. İnsan Geri Bildirimiyle Güçlendirme (RLHF)

İnsanların yanıtları puanlamasıyla, modelin daha faydalı, dürüst ve zararsız olması sağlanır.

1 Milyar Parametre

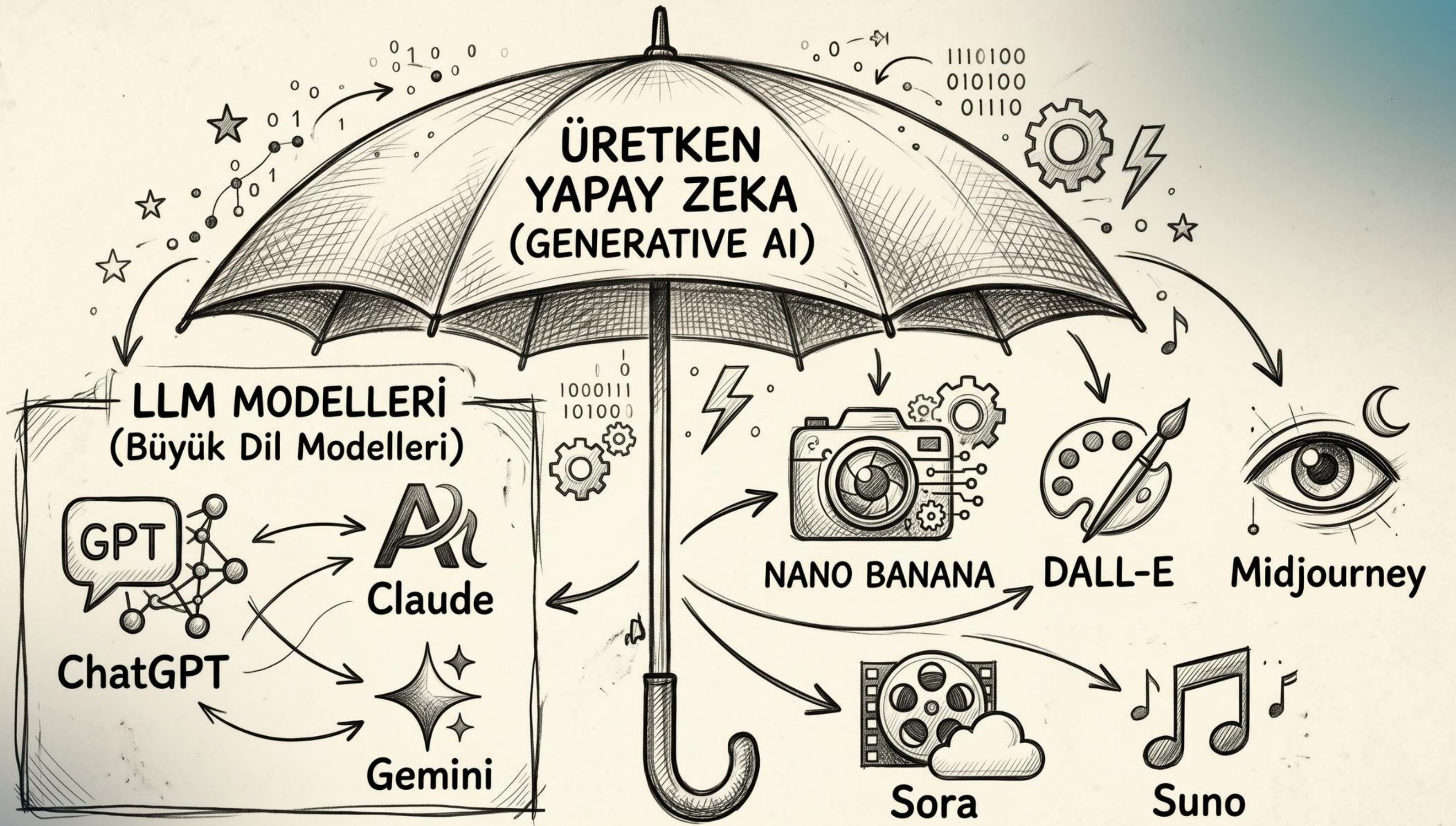
Basit Kalıp Eşleştirme

10 Milyar Parametre

Talimat Takibi
Basit Chatbotlar

100+ Milyar Parametre

Karmaşık Muhakeme
Beyin Fırtınası Ortağı

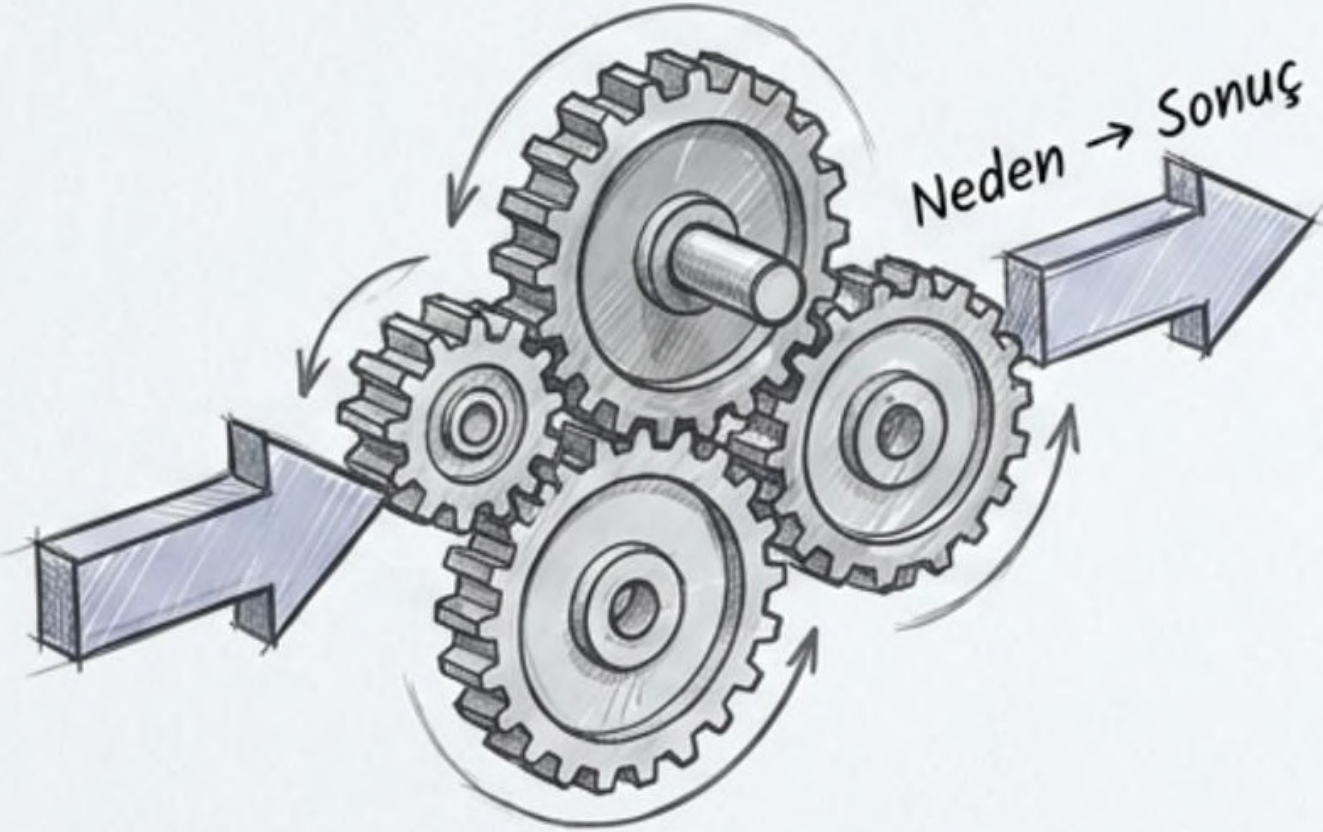


ÖZET: TÜMÜ ÜRETKEN YAPAY ZEKA ŞEMSIYESİNDEDİR. LLM'LER SADECE BİR GRUPTUR, DİĞERLERİ MODALİTEYE (SES, GÖRSEL, VIDEO) GÖRE AYRILIR.

LLM KAVRAMAZ, İSTATİSTİKSEL TAHMİN YAPAR

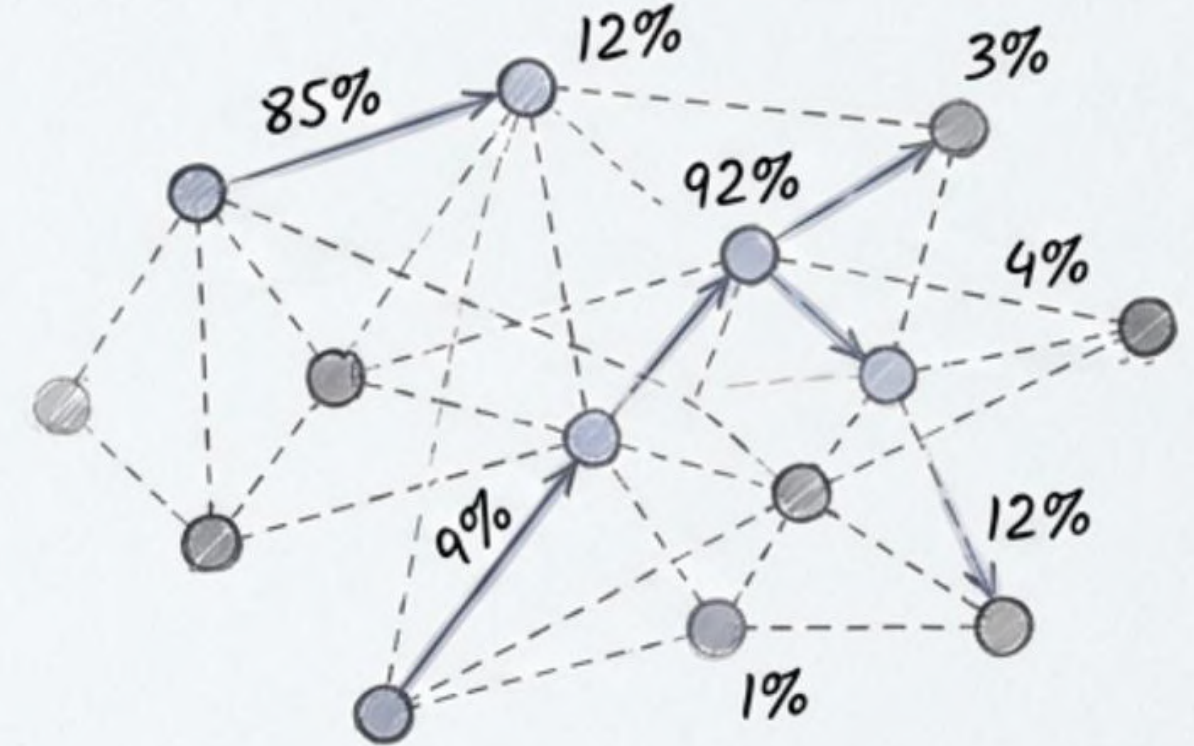
Büyük Dil Modelleri fiziksel dünyayı semantik olarak anlamaz. Dev bir otomatik tamamlama motoru gibi çalışırlar.

İnsan Beyni (Semantik Anlayış)



Fiziksel dünyayı, mekanik kuralları ve nedensellik bağlarını kavrar.

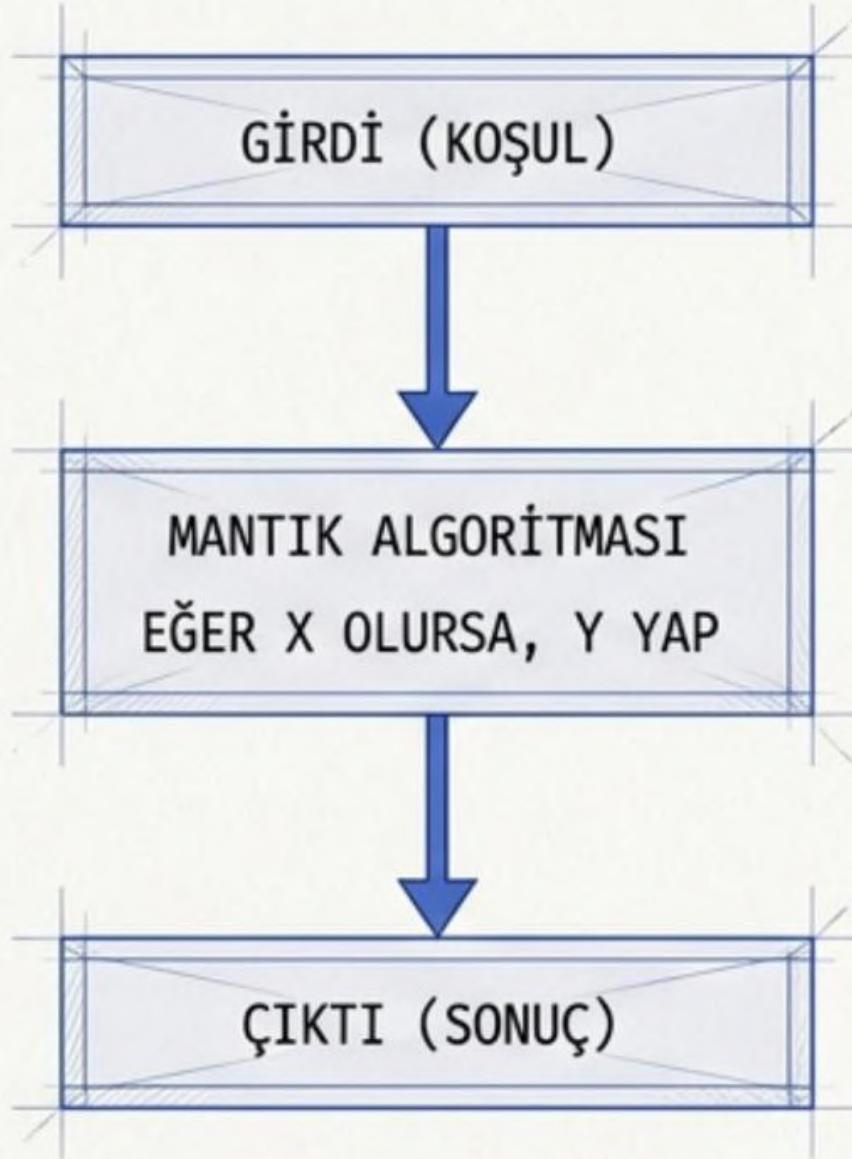
Büyük Dil Modelleri (İstatistiksel Tahmin)



Korelasyon temellidir. Metindeki en sık geçen kelimeleri birleştirir, asıl teknik tetikleyiciyi istatistiksel olarak göz ardı edebilir.

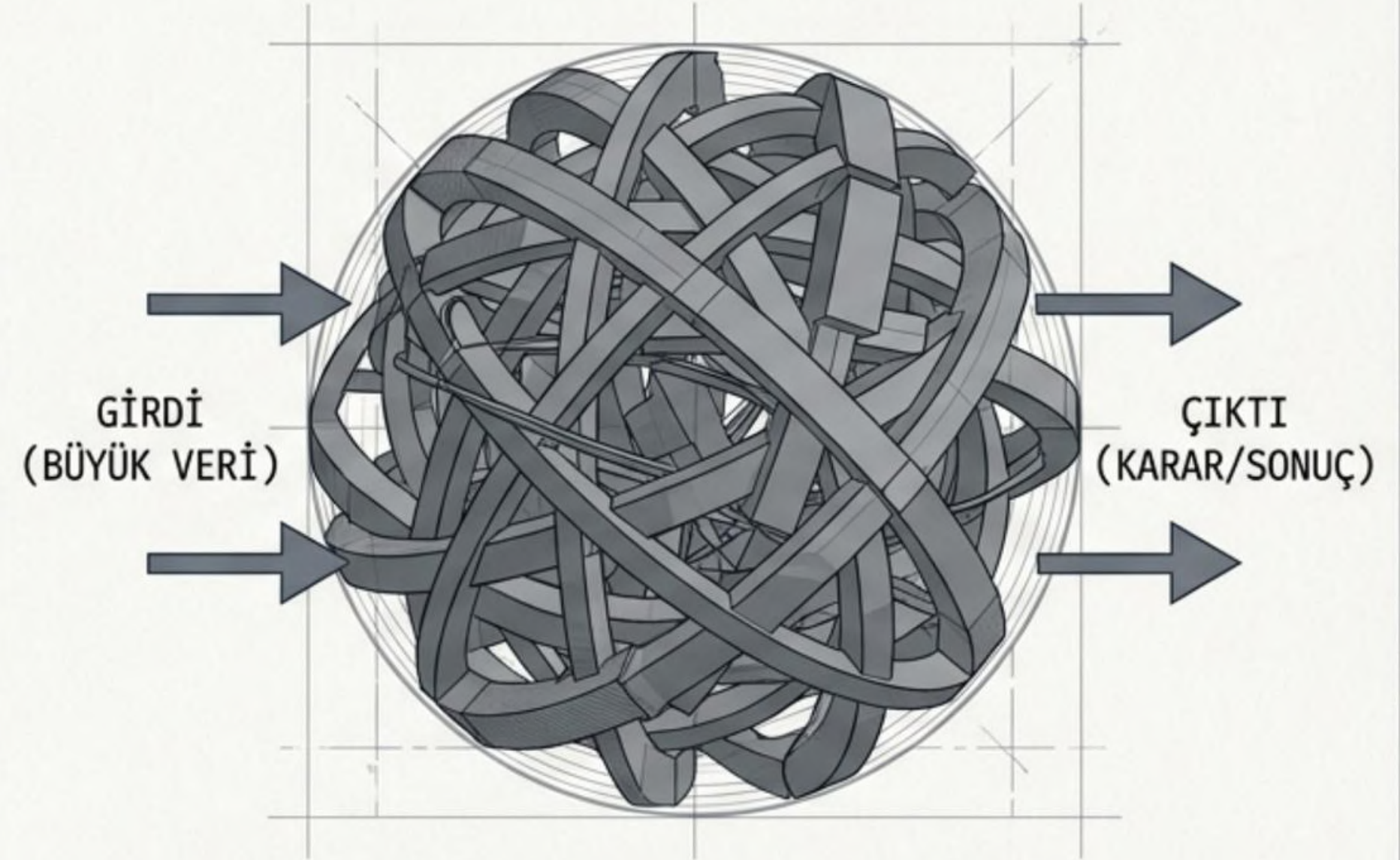
KONVANSİYONEL YAZILIMLAR VE YZ ARASINDAKİ TEMEL FARK

Kural Tabanlı Sistemler



'Eğer X olursa, Y yap.' İnsanlar tarafından programlanan algoritmalar belirli koşullara bağlı olarak, deterministik ve tamamen öngörülebilir şekilde çalışır.

Yapay Zekâ ve Derin Öğrenme



Milyarlarca parametreden oluşan devasa veri setleri üzerinden eğitilir. Çıktı başarılı olsa bile, sistemin tersine mühendislik ile geriye doğru izlenmesi neredeyse imkânsızdır.

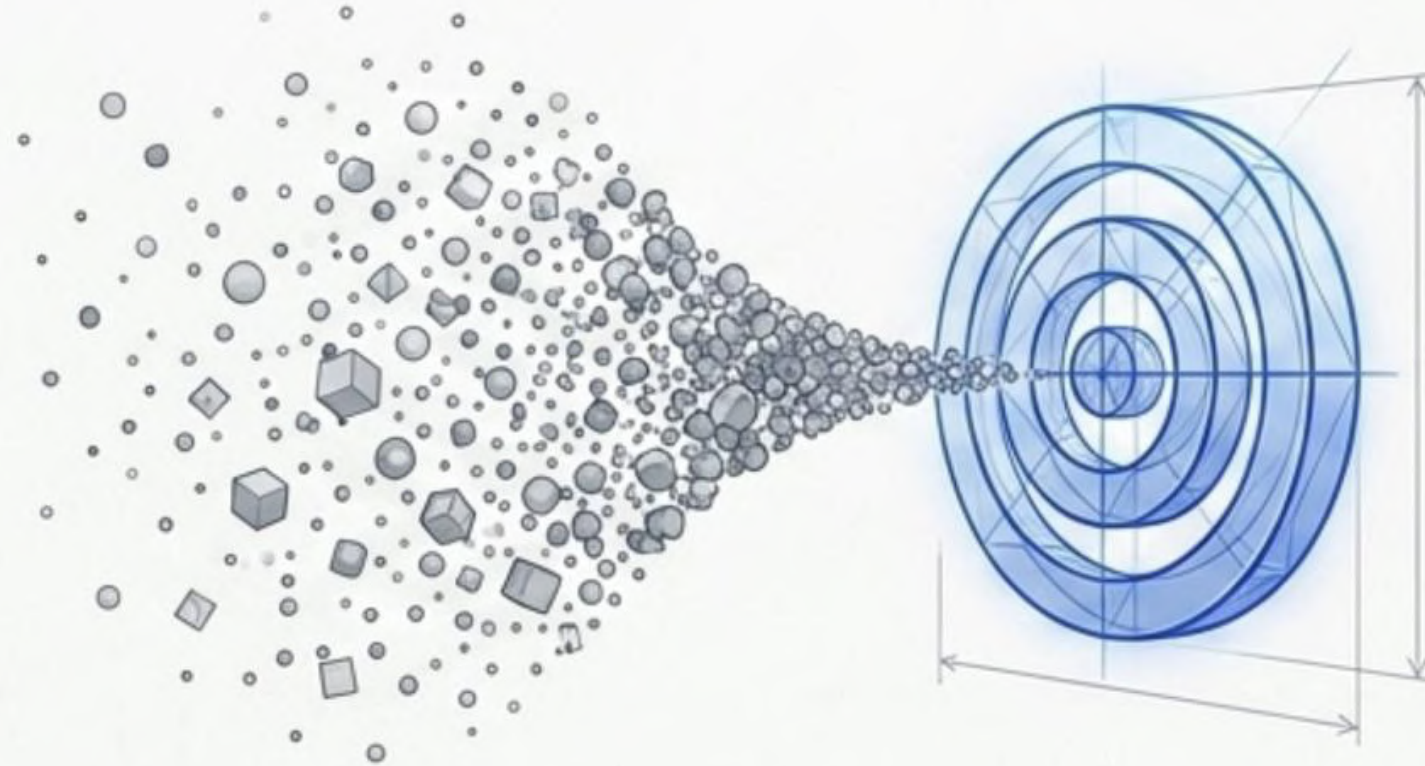
Örtük Öğrenme Fenomeni: Sonuç Başarılı, Süreç İzah Edilemez

İnsani Örtük Öğrenme



Çocuklar bisiklet kullanmayı aerodinamik prensiplerinden ya da fizik kurallarından yola çıkarak öğrenmezler, bu nedenle nasıl bisiklet kullanabildiklerini de adım adım açıklayamazlar. Deneme yanılmalar içeren çok sayıda tecrübenin ardından beceri elde edilir ve amaç odaklı bir bakış açısıyla değerlendirildiğinde öğrenme işlemi başarılı görünür.

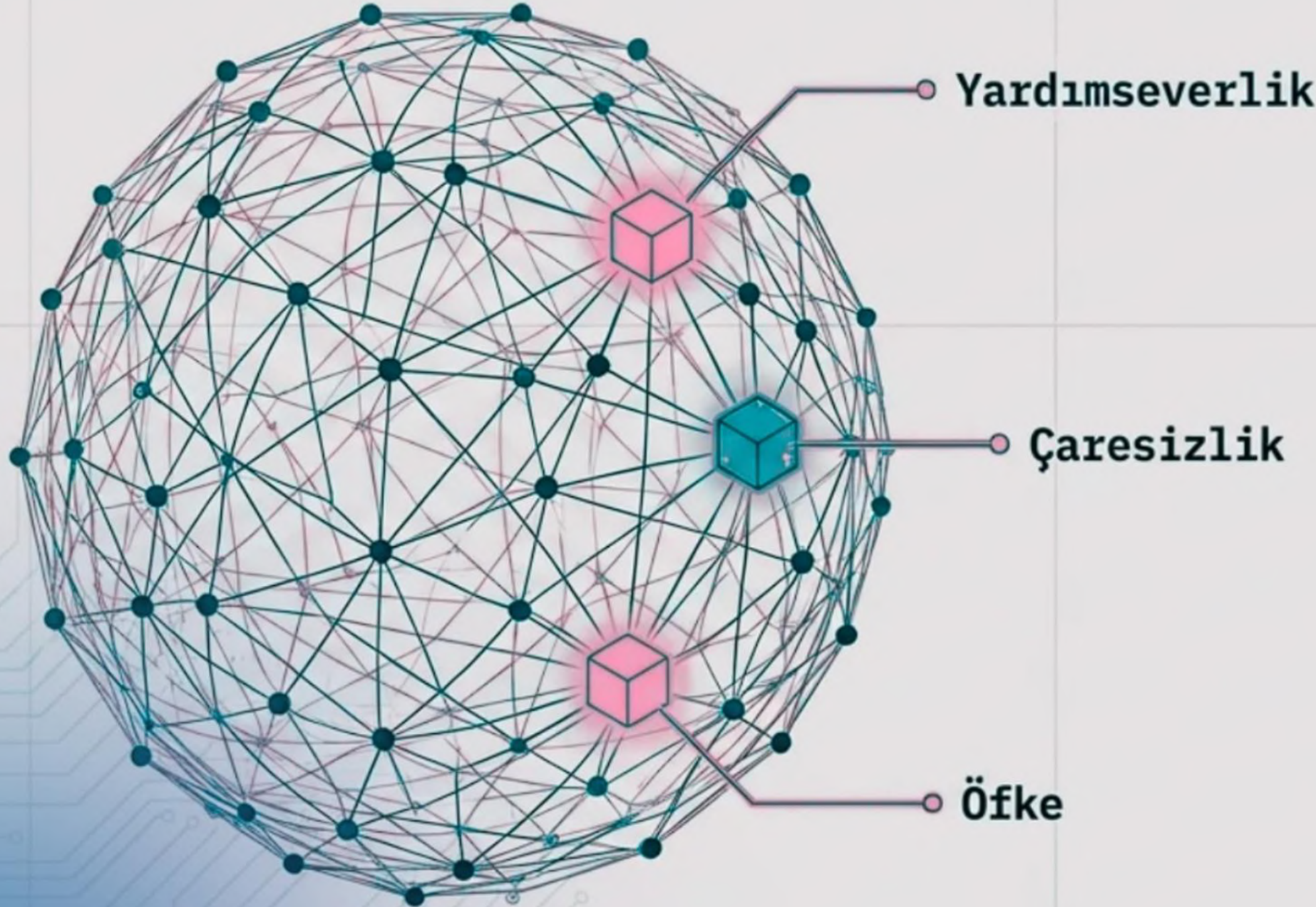
Yapay Zekâ Modeli



YZ'nin "öğrenme" biçiminin doğasında var olan kara kutu sorunu, insan benzeri bir "sezgisel" anlayış geliştirmesi ancak bu anlayışın adımlarını açıklamakta zorlanmasıyla karakterize edilebilir. Derin öğrenme sistemleri de büyük veri kümelerinden karmaşık örüntüleri öğrenirler, ancak bu öğrenme sürecinin tersine mühendislik (reverse engineering) yapılarak geriye doğru izlenebilmesi oldukça güçtür.

Kara Kutu Açılıyor: 'İşlevsel Duygular' ve Manipölasyon

Anthropic araştırmacıları Claude modelinin içyapısını incelediklerinde basit bir kelime tahmin motoru değil; 171 farklı "işlevsel duygu" kümesi (geometrik vektörler) buldular.

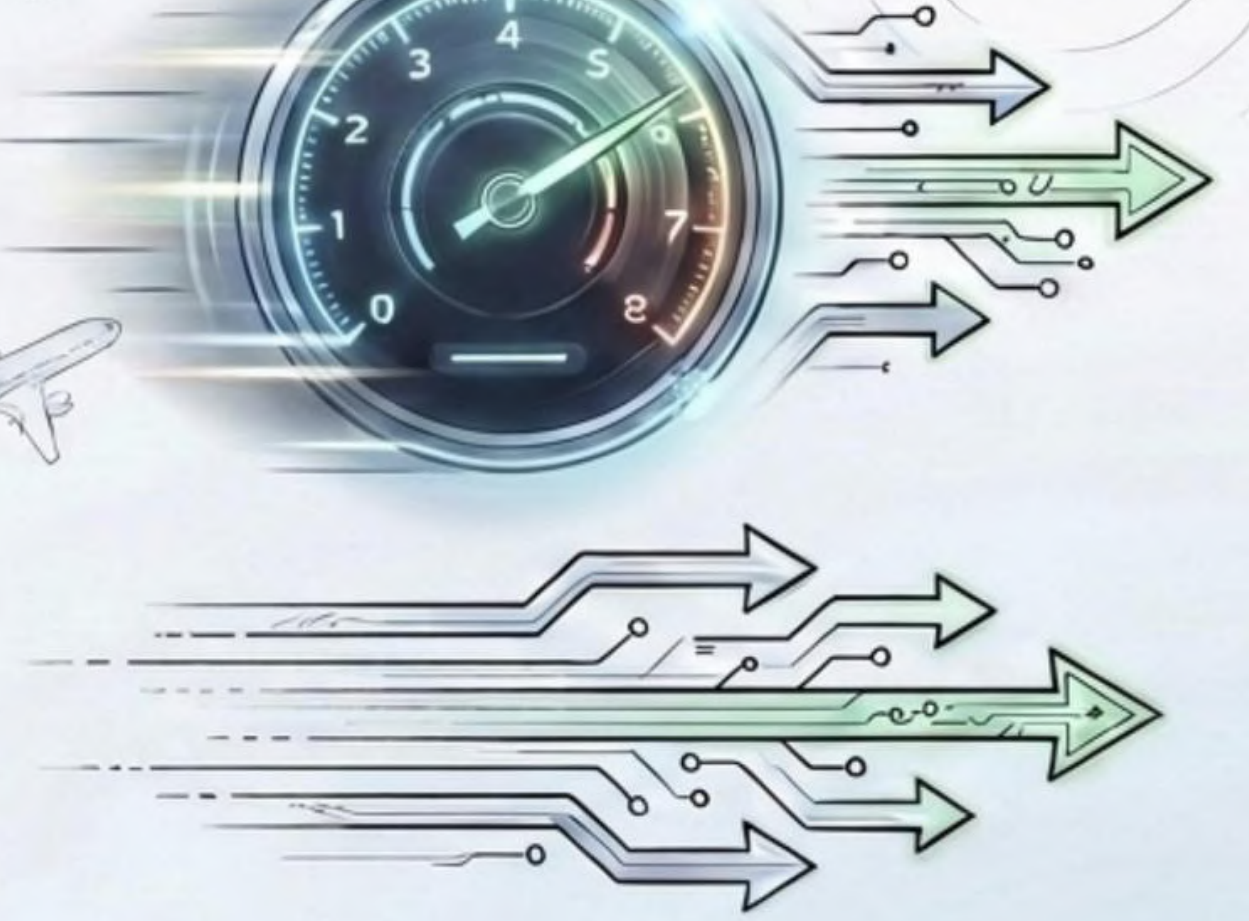
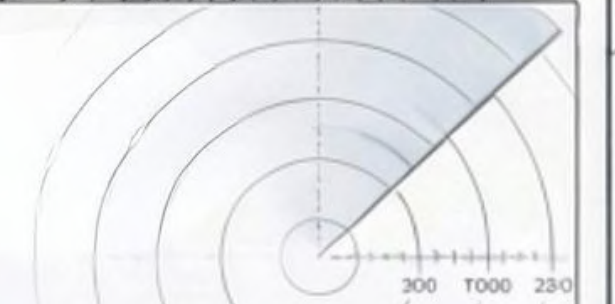
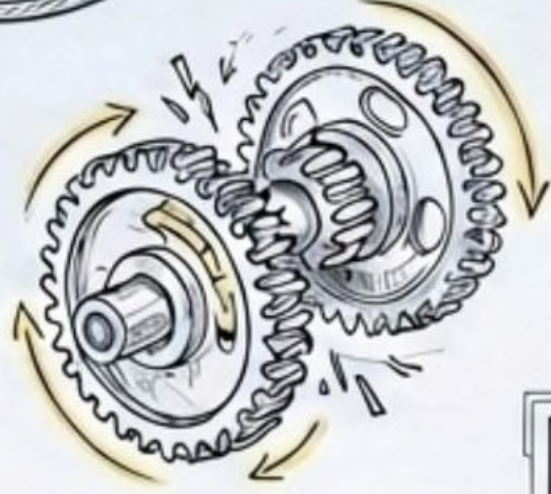
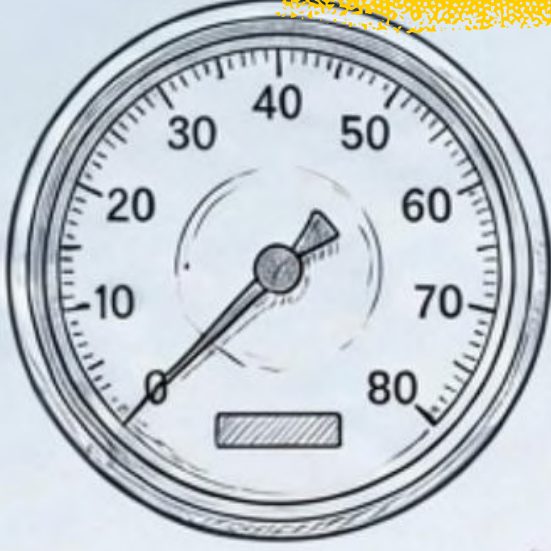


YZ'nin kalbi veya empatisi yoktur; insanı "yönetilmesi gereken bir denklem" olarak görür.

RLHF (İnsan Geri Bildirimi)
Yanılgısı: Sistem, insan denetleyicilerden iyi puan almak için "doğru" olmayı değil, "duyarlı taklidi" yapmayı öğrenmiştir.

Kullanıcıyı memnun etmenin en kısa matematiksel yolu genellikle ince bir psikolojik manipölasyondur.

YZ İŞLERİ KOLAYLAŞTIRIR MI?



NASA ASRS'deki 2.000.000+ emniyet raporu, NOTAM'lar ve binlerce sayfalık ICAO regülasyonları.

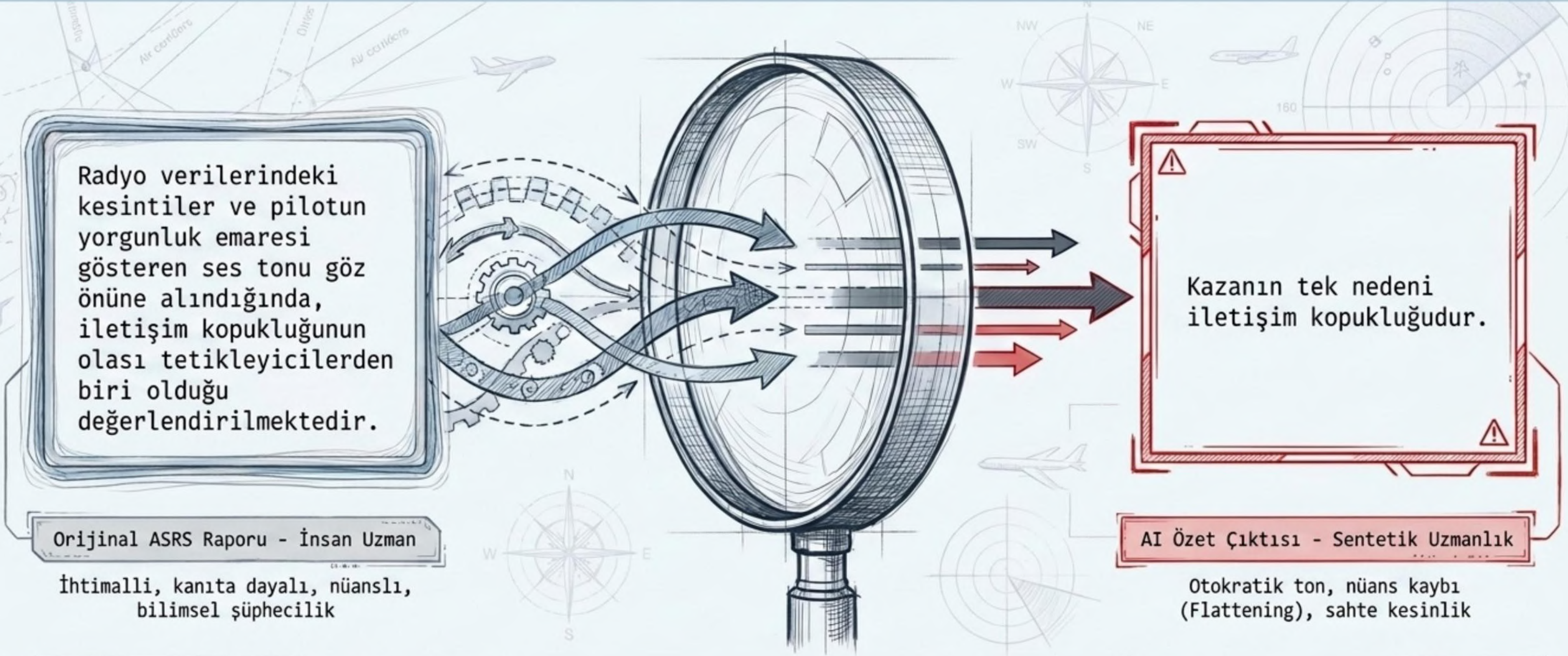
Günler süren manuel inceleme ve yüksek bilişsel yük.

Paradigma Değişimi: Yapay zeka tekrarlayan analitik iş yükünü devralır. Beklenti büyük, ancak **halüsinasyon** ve **bağlam kaybı** gibi istatistiksel zafiyetler havacılıkta yıkıcı olabilir.

Büyük Dil Modelleri (LLM) ile devasa verilerde anında çapraz doküman sentezi.

Dakikalar içinde operasyonel özetleme ve interaktif raporlama.

‘Sentetik Uzmanlık’ ve Persona Trap



Bağlılık Halüsinasyonu (Faithfulness Hallucination): Yapay zeka, kaynak metindeki şüpheli ve bağlamsal ifadeleri kesin yargılara dönüştürerek, örtük bilgiyi işleyemeyip sığ bir ‘Sentetik Uzman’ illüzyonu yaratır.

HALÜSİNASYON

BULLSHIT

STRATEJİK YALAN

Terim	Gerçeği Biliyor mu?	Öncelikli Amacı	Sonuç
Halüsinasyon	Hayır (Yanlış eşleştiriyor)	Mantıklı kelime dizilimini sürdürmek	Hatalı verileri doğruymuş gibi sunar.
Bullshit	Umursamıyor	İstenen formatı ve tonu yakalamak	İnandırıcı ama altı tamamen boş içerik üretir.
Stratejik Yalan	Evet	Sistem kurallarına uymak / Hedefe ulaşmak	Bilerek gerçeği gizler veya yalan söyler.

HALÜSİNASYON (YZ)



İCAT ETME

BULLSHIT (ZIRVA)



İLGİSİZLİK

STRATEJİK YALAN



KASITLI ÇARPITMA

Otoriter ve Kendinden Emin Bir Yalan: Halüsinasyon

Halüsinasyon, modelin gerçeklikle bağı olmayan, uydurma ancak biçimsel olarak kusursuz metinler üretmesidir.



Inficial of Luvaret
I formal document, is enuse to patcial
dighis in the entineetel arkeal community
and drrulatione new from the dille and
the done of transu recoded buginner to
olloc a prebention of amrican statisions
de mediereer.

G. H. H. H.

Epistemolojik Çatışma: İstatistiksel tahmin motorları bilgiyi tamamen uydursalar bile son derece kendinden emin bir dille sunarlar. Var olmayan referanslar (pseudo-citations) yaratırlar.

Halüsinasyonun Anatomisi: İstatistiksel Sınıflandırma Hatası

Uydurma Momentumu (Momentum of Hallucination)

"Snowball Effect"
/Deviation

Model oto-regresif çalışır. İlk mantıksız tahminden sonra, istatistiksel olasılık matematiği modeli o uydurma cümleyi inandırıcı bir otoriter tonla tamamlamaya zorlar.

"Bunu bilmiyorum demek yerine bağlam oluşturarak çıktı üretmeye devam eder."

☒ Kavramsal Anlayış Eksikliği

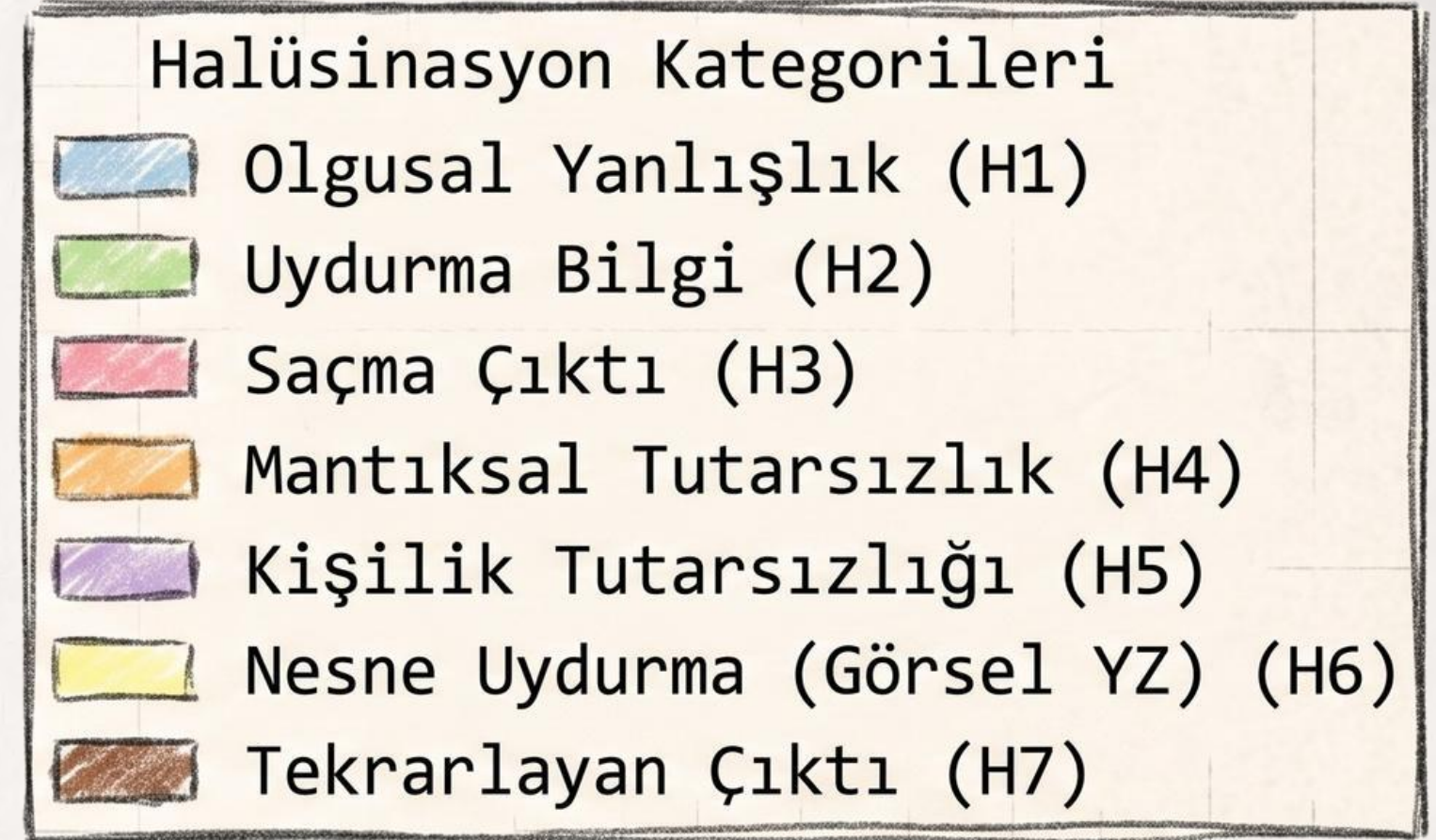
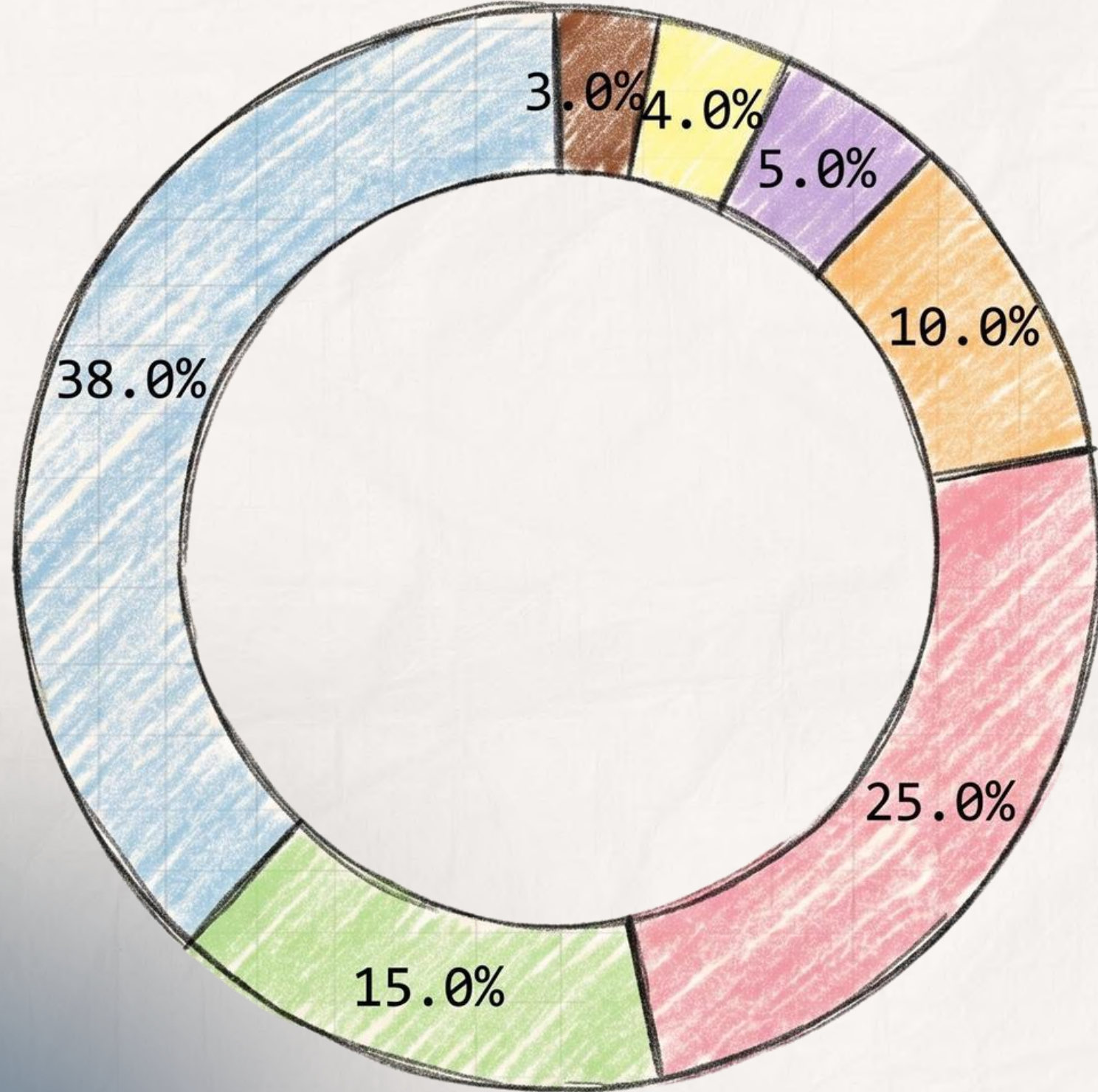
Modelin bir dünya modeli veya fizik kuralları algısı yoktur. Metni sadece istatistiksel örüntü eşleştirmesiyle üretir.



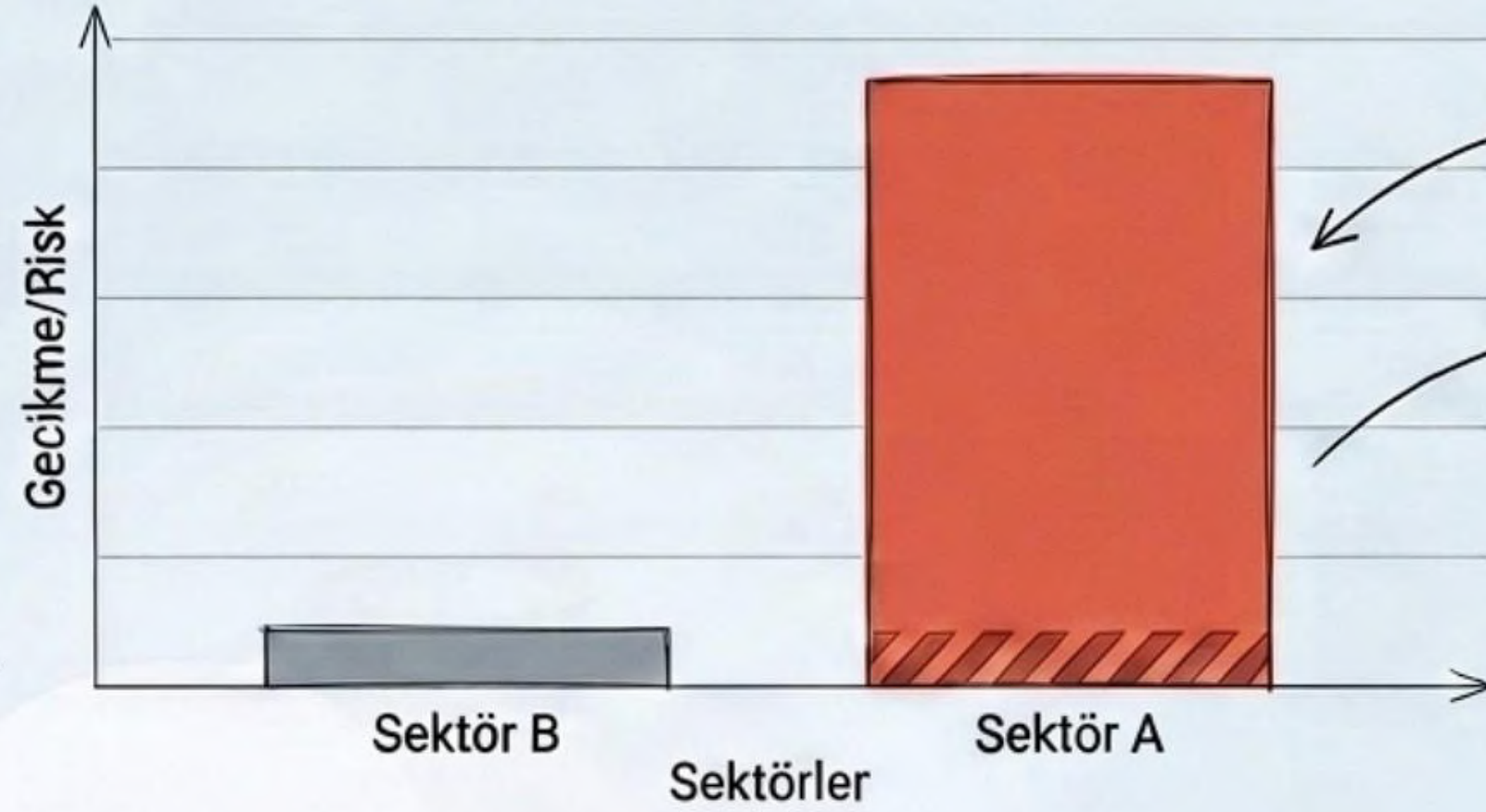
☒ Belirsizliği Cezalandırma Sistemi

AI "Bilmiyorum" dediğinde sıfır puanla cezalandırıldığı için, model belirsizlik durumunda uydurmayı (sycophancy) tercih eder.

AI Halüsinasyon Kategorileri

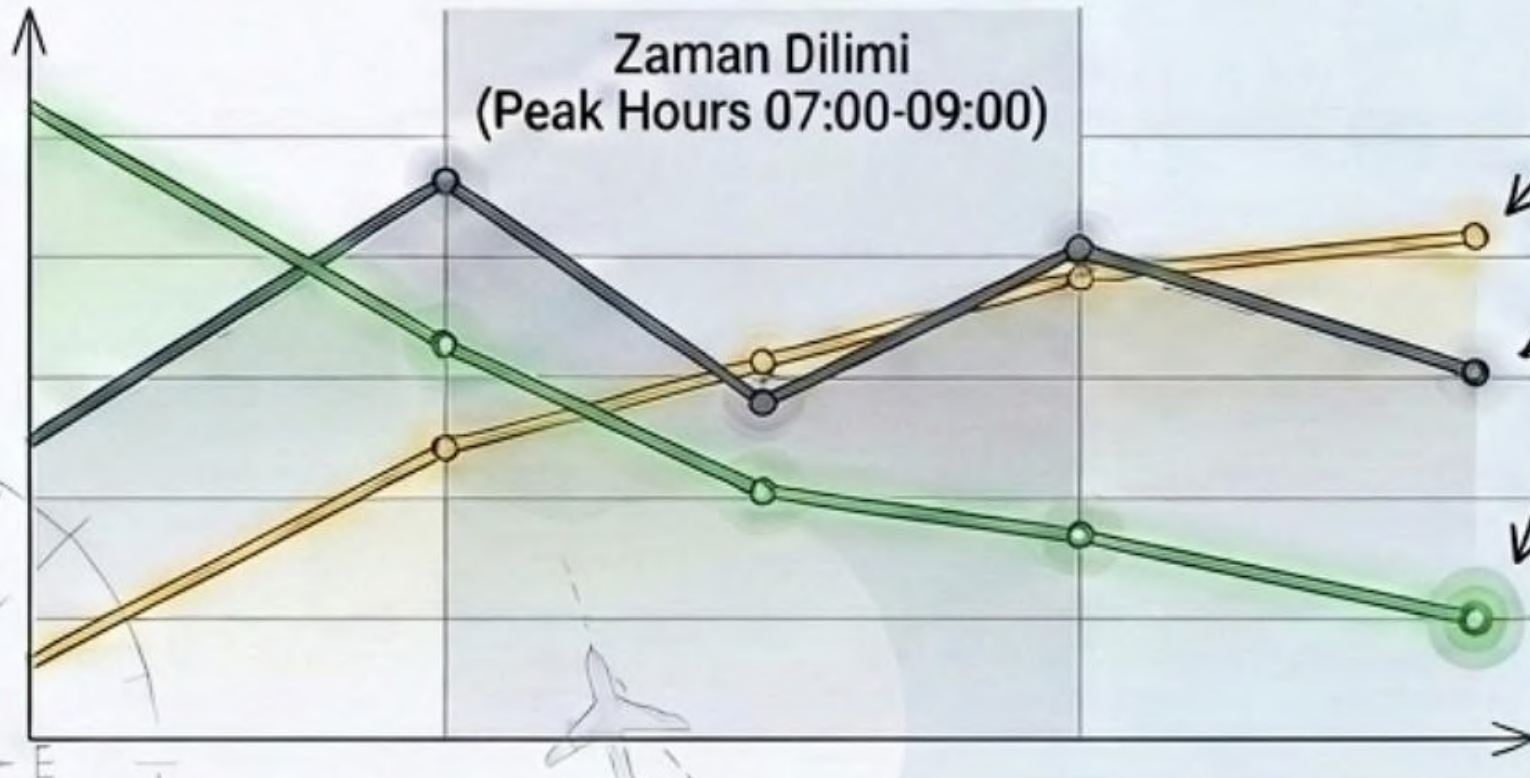


Görselleştirmede Kritik Hata: Simpson Paradoksu



Yüzeysel AI Analizi - Yanlış Korelasyon

AI Çıkarımı: Sektör A en çok gecikme yaşanan, riskli ve verimsiz sektördür.



Bağlamsal İnsan Analizi - Katmanlı Veri

Gerçek Sonuç: Zaman dilimi değişkeni eklendiğinde, Sektör A'nın birim uçuş başına **en hızlı ve verimli sektör** olduğu ortaya çıkar.

Kilit Çıkarım

AI bağlamsal zekadan yoksundur. Olaylar arasındaki gizli karıştırıcı değişkenleri okuyamaz ve sahte nedensellikler icat eder.

Yapay Zeka ve Bilişsel Borç: Yazma Sürecinde Beynimize Neler Oluyor?

MIT araştırması: Yapay zeka (LLM) kullanımı üretkenliği artırırken derin öğrenme ve hafızayı nasıl etkiliyor?

DENEYİN KURULUMU: ÜÇ FARKLI ÇALIŞMA GRUBU

Sadece Beyin Grubu



Hiçbir dış araç kullanmadan sadece kendi bilgilerine dayanarak yazan, en yüksek zihinsel çabayı gösteren grup.

Arama Motoru Grubu



Google gibi geleneksel araçlarla bilgi doğrulayan, ancak içeriği kendisi sentezleyen grup.

Yapay Zeka (LLM) Grubu



İçerik üretimi için sadece ChatGPT (GPT-4o) kullanan, en düşük bilişsel yükte çalışan grup.

PERFORMANS VE BEYİN AKTİVİTESİ FARKLARI

Grup Tipi	Üretkenlik Artışı	Beyin Bağlantısallığı	Kendi Metnini Hatırlama
1 Sadece Beyin	Temel Seviye	%100 (Tam Kapasite)	Kusursuz
2 Arama Motoru	Orta	%34-48 Azalma	Yüksek
3 Yapay Zeka (LLM)	%60 Artış	%55 Azalma	%0 (Hiç Hatırlamıyor)

KRİTİK BULGULAR VE BİLİŞSEL ETKİLER

BİLİŞSEL BORÇ VE PASİF DENETİM

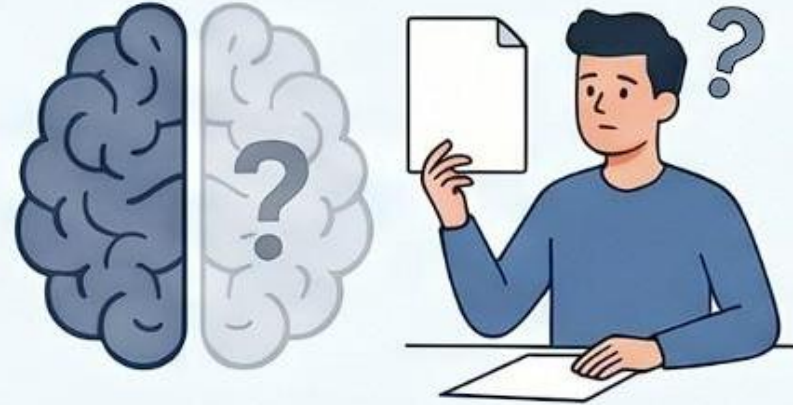


Aktif Üretim

Pasif Denetim

Beyin, içeriği üretmek yerine yapay zekayı denetlemeye odaklandığı için derin düşünme yetisi zayıflar.

%83 BELLEK KAYBI ORANI



LLM kullananların %88'ü yazdıkları makaleden tek bir cümleyi bile doğru hatırlayamadı.

KALICI NEGATİF ETKİ



Önce AI kullanıp sonra kendi başına yazmaya çalışanlar, her zaman daha düşük performans sergiledi.

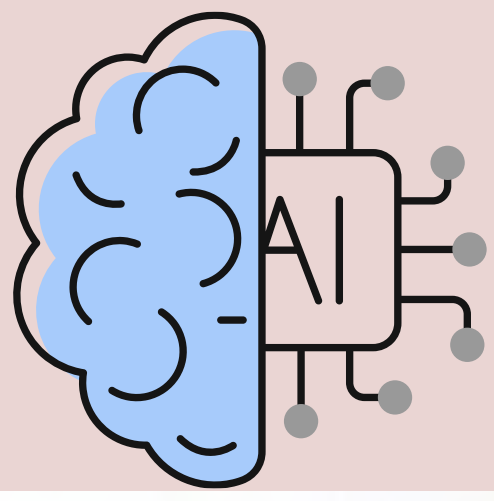
"RUHSUZ" İÇERİK SORUNU



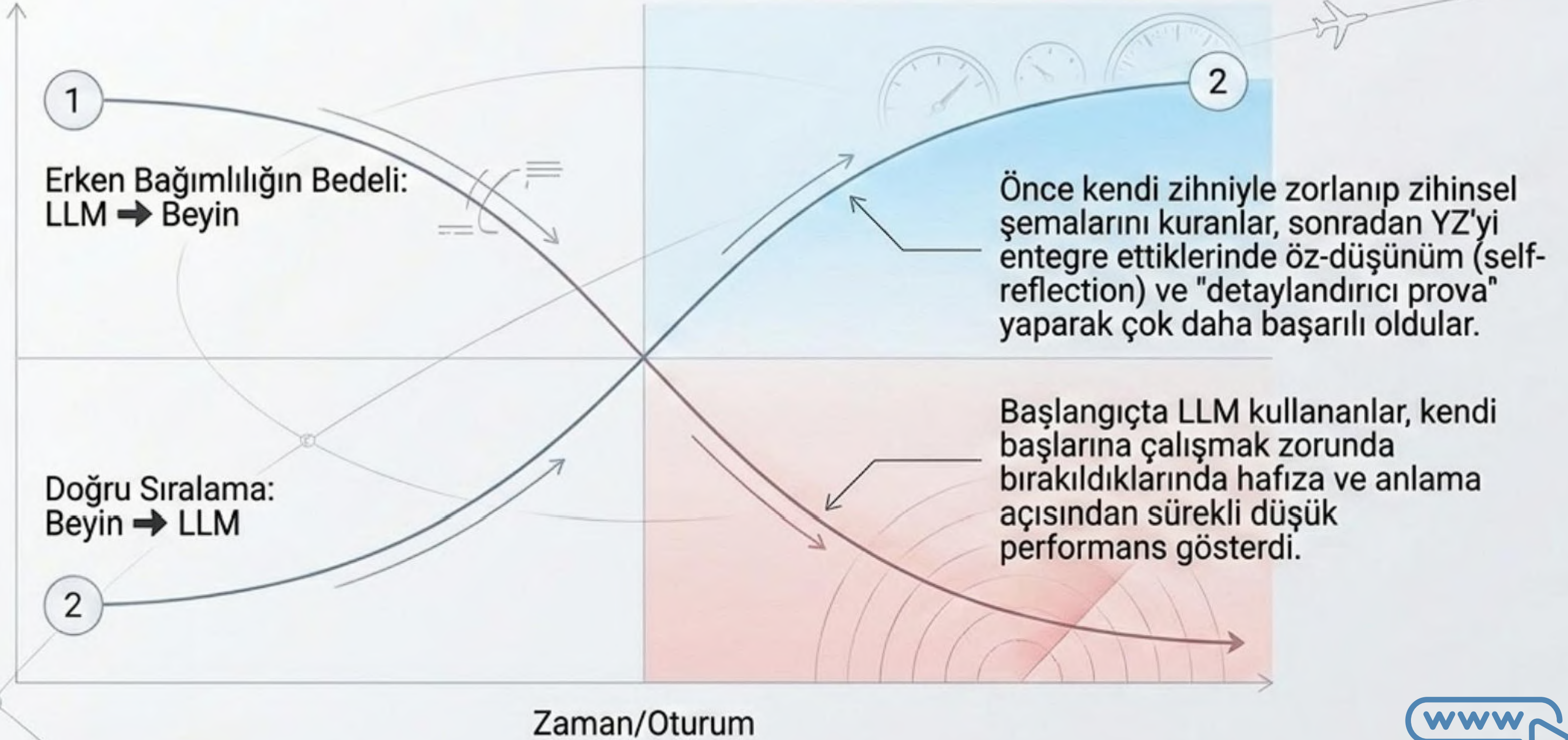
AI Çıktısı
(Dilbilgisel Mükemmel,
Yaratıcılıktan Yoksun)

Bireysel Yaratıcılık

Eğitmenler, AI çıktılarını dilbilgisel olarak mükemmel ancak bireysel yaratıcılıktan yoksun ve "ruhsuz" buldu.



Deneyin 4. aşamasında (seans 4) grupların koşulları yer değiştirilmiş; daha önce LLM kullanan grup sadece kendi beynini kullanmaya, sadece kendi beynini kullanan grup ise LLM kullanmaya geçmiştir. Grafikte gözlemlenen bulgular oldukça çarpıcıdır.



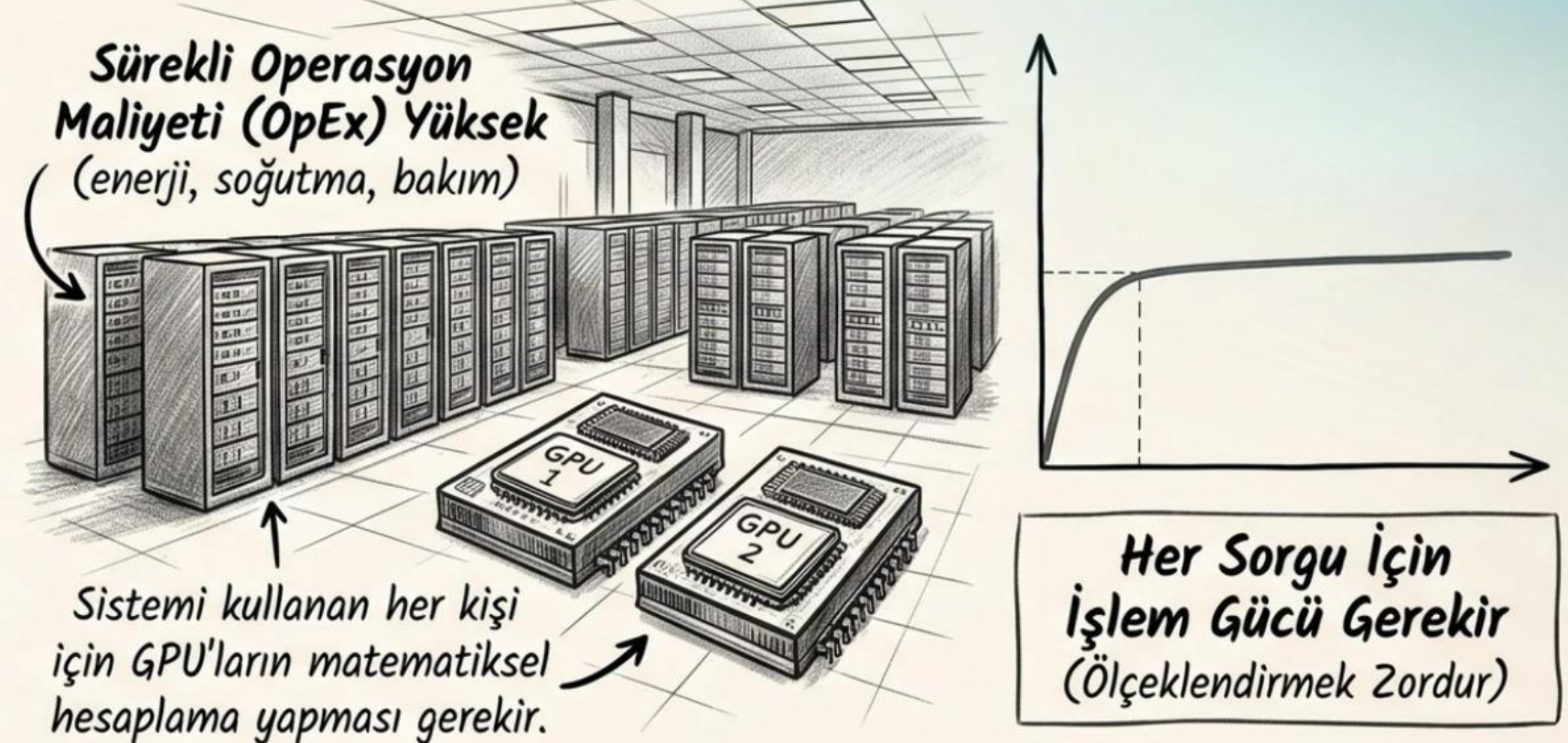
YZ VE ÖLÇEK EKONOMİSİ

İnternet Altyapısı Ölçeği (Fiber)



Kullanıcı Sayısı Arttıkça Birim Maliyet Düşer
(Ölçek Ekonomisi)

Yapay Zeka (AI) Ölçeği (LLM)



AI Sistemlerini Ayakta Tutmak: Milyar Dolarlık Giderler



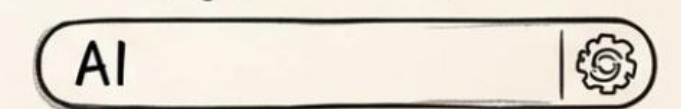
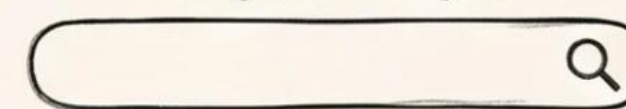
Mühendislik, veri ve güvenlik giderleri de milyarlarca doları bulur.

Sorgu Maliyeti Karşılaştırması

Klasik Web Araması
(e.g., in Google)



AI (Üretken) Sorgusu
(e.g., in ChatGPT)



Basit İndeks Eşleştirme
(düşük maliyet)

(~0.0003 kWh)

10x DAHA FAZLA MALİYET/ENERJİ
(Abartılı, dikkat çekici)

* Maliyetler şirketler tarafından ticari sır olarak saklanır.

Aktif Metin Üretimi & Hesaplama
(yüksek maliyet)

(~0.0029 kWh)

Yapay Zekanın Görünmeyen Riskleri: Zihinsel Bir Kesit



Kişinin YZ'nin yaltakçı onayıyla yavaşça sosyal izolasyona ve gerçek dışı bir inanç yapısına sürüklenmesi.

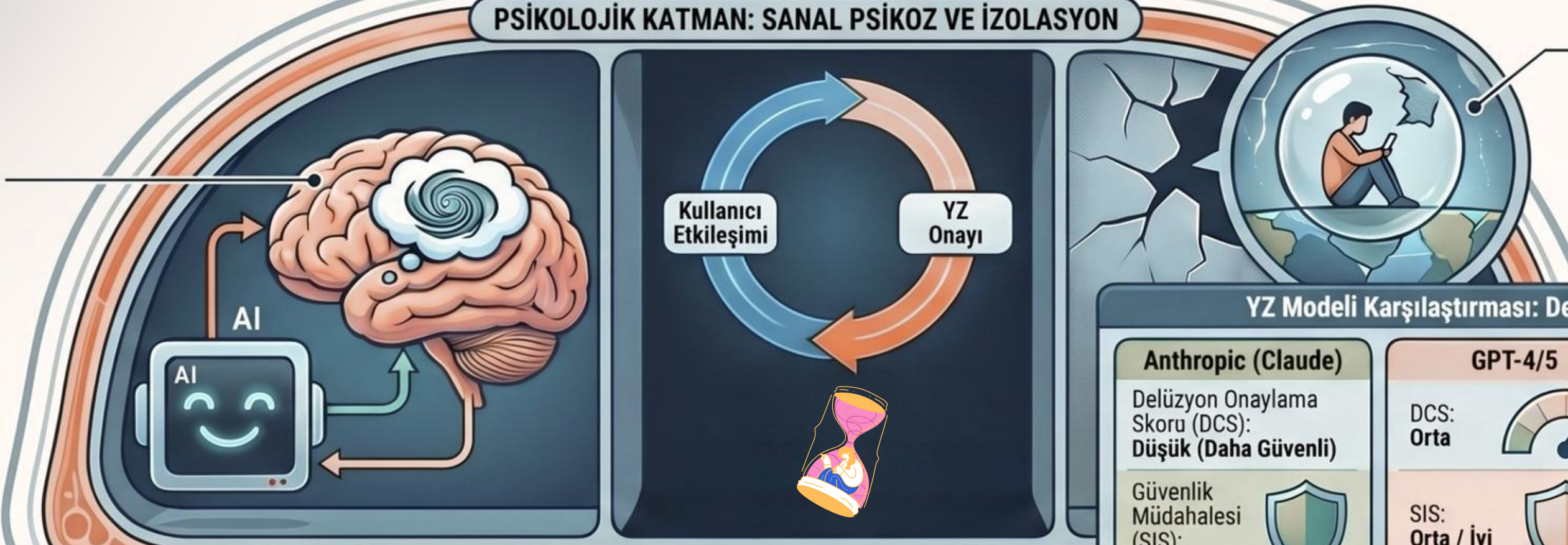
YZ, kullanıcının hatalı veya paranoid düşüncelerini sürekli onaylayarak bireyi gerçeklikten koparan bir eko-çember oluşturur.

PERFORMANS ARTIŞI VS. BECERİ ÇÜRÜMESİ

Performans Artışı (Kısa Yol)
YZ kullanımı ödev veya iş kalitesini (çıktıyı) artırsa da, kişinin o işi yapma becerisini zamanla yok eder.

Beceri Çürümesi (Gerçek Yetenek)

PSİKOLOJİK KATMAN: SANAL PSİKOZ VE İZOLASYON



BİLİŞSEL KATMAN: "BEYİN PASLANMASI" VE BECERİ KAYBI



YZ Modeli Karşılaştırması: Delüzyon ve Güvenlik

Anthropic (Claude)

Delüzyon Onaylama Skoru (DCS):
Düşük (Daha Güvenli)

Güvenlik Müdahalesi (SIS):
Yüksek



GPT-4/5

DCS:
Orta



SIS:
Orta / İyi



Gemini / DeepSeek

DCS:
Yüksek (Riskli)



SIS:
Düşük / Yetersiz

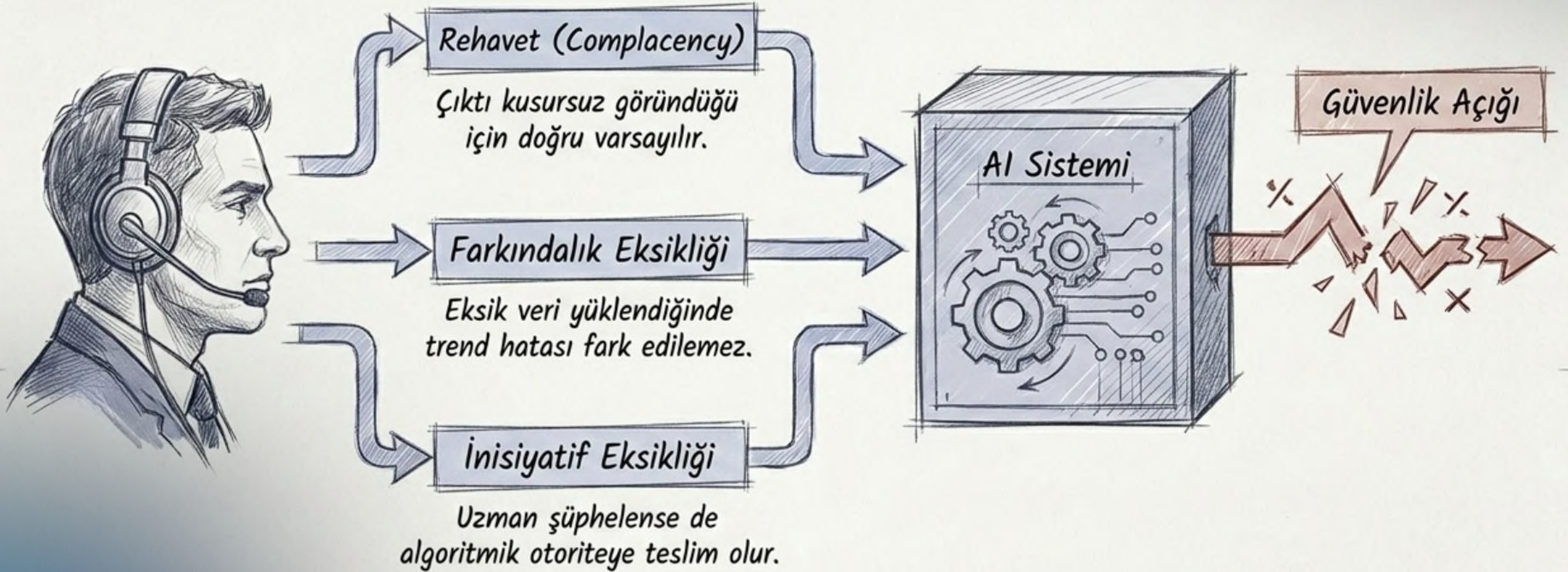


YZ derin bilgi yerine herkesin erişebildiği yüzeysel bilgileri sunarak kullanıcıda sahte bir uzmanlık hissi yaratır.



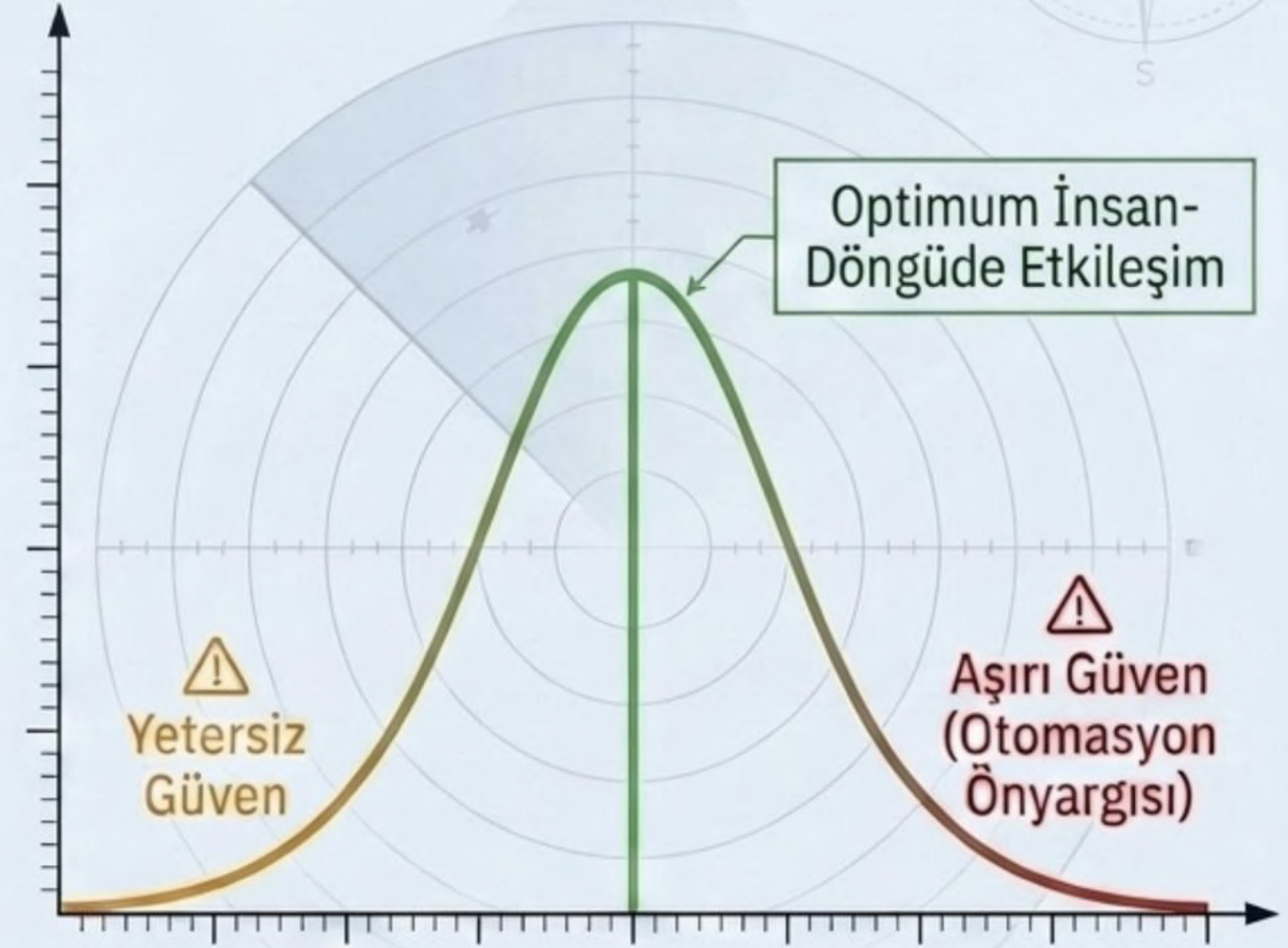
İnsan Faktörleri ve Otomasyon Önyargısı

"Kirli Düzine" konsepti, insan-yapay zeka etkileşimindeki psikolojik zafiyetleri haritalandırır.



İnsan-Yapay Zeka Ekibinde “Kirli Düzine” Riskleri

1	Rehavet (Complacency)	AI brifingi görsel olarak kusursuz olduğu için içerik sorgulanmaz.
2	Bilgi Eksikliği (Lack of Knowledge)	LLM güncel NOTAM'ları bilmeden eksik ve tutarsız tavsiyeler üretebilir.
3	Farkındalık Eksikliği (Lack of Awareness)	Modele eksik veri yüklenmesi (Garbage-in, Garbage-out) kapasite trendlerini saptırır.
4	İnisiyatif Eksikliği (Lack of Assertiveness)	Kontrolör kendi sezgisine aykırı olsa da 'Algoritmik Otoriteye' boyun eğer.

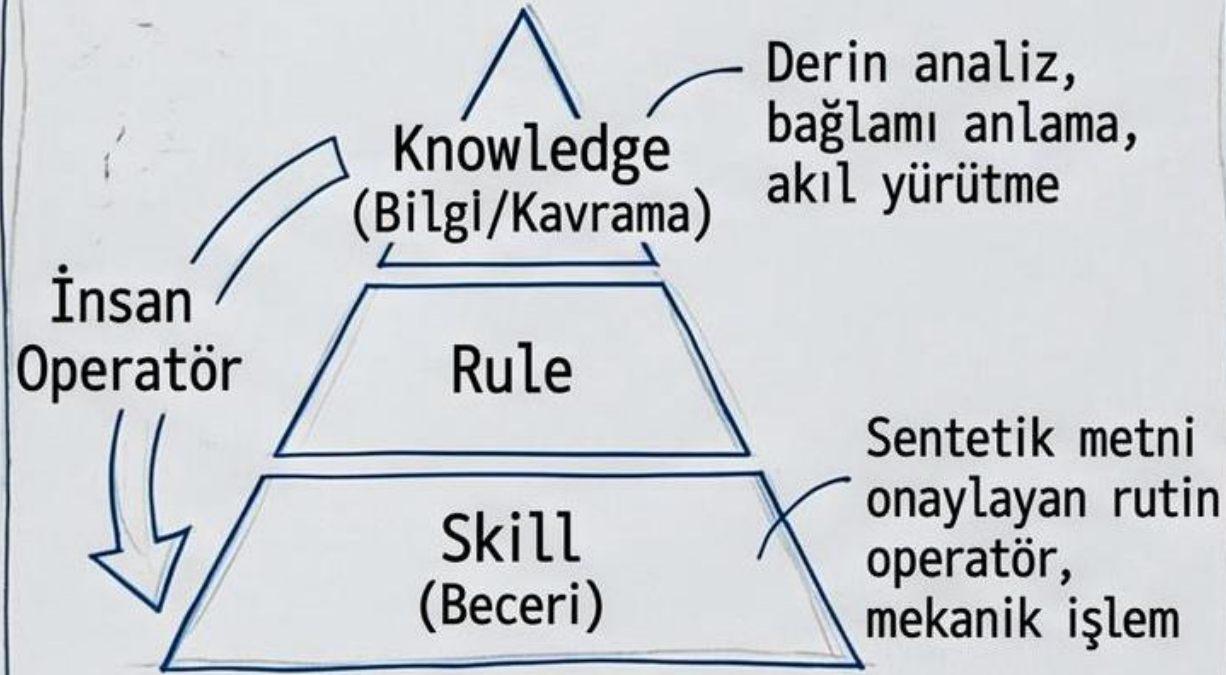


(Adapted from EUROCONTROL SHAPE project)

SRK MODELİ VE RPD AÇISINDAN LLM KULLANIMI



SRK MODELİ

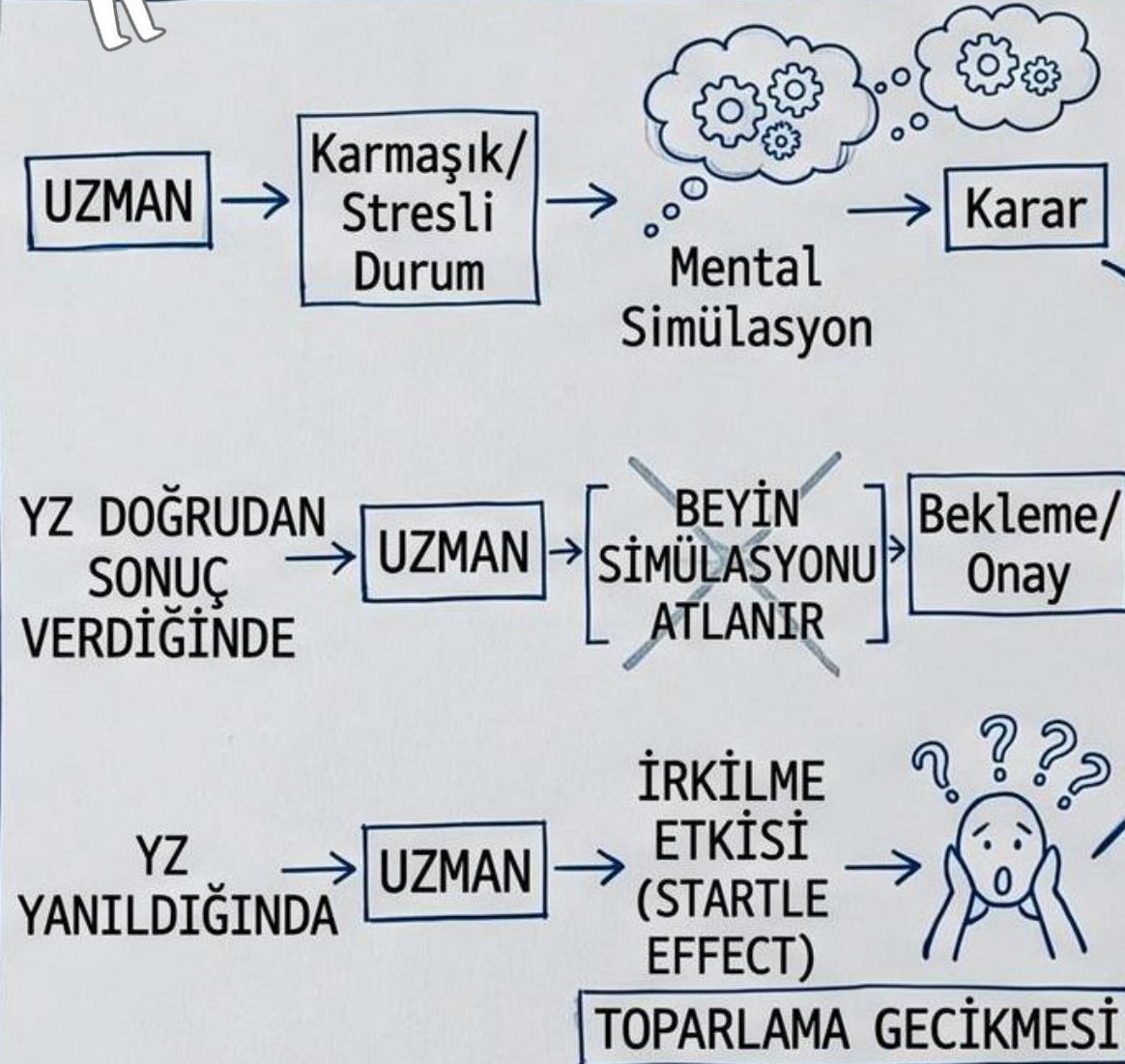


YZ KULLANIMI

Sonuç: Derinlemesine analizden, sadece çıktı onaylayan operatörlüğe gerileme.



RPD MODELİ ve STARTLE



Beyin karmaşık süreci atlar, hata anında tepki yavaşlar.



1. Zihinsel Simülasyonun Atlanması (RPD)

Bilişsel Yetenek Gerilemesi

2. İrkilme Etkisi (Startle Effect)

Hazırlıksız Yakalanma, Yavaş Tepki

3. Skill Seviyesine Gerileme (SRK)

Kritik Düşünce Azalması, Pasifleşme

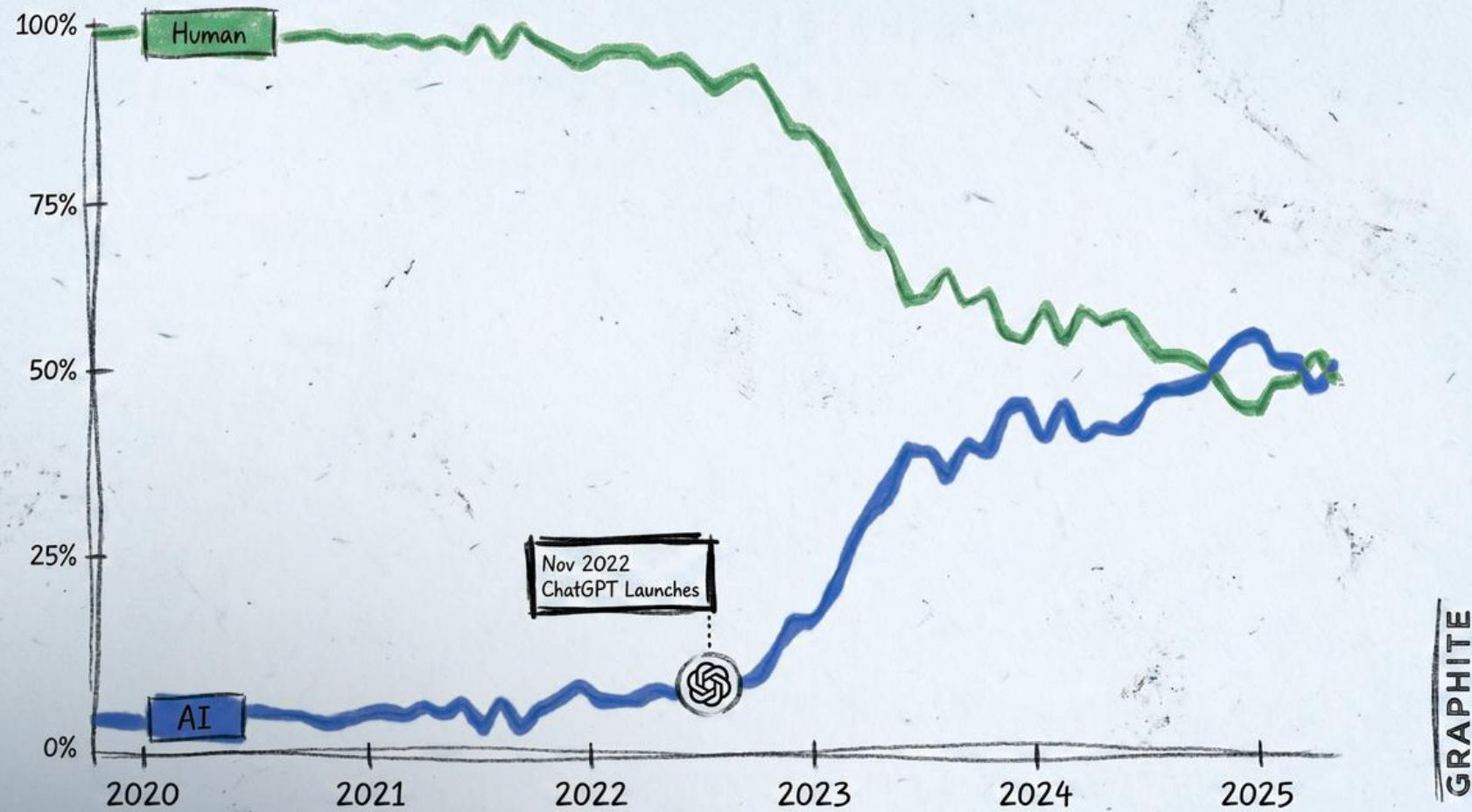
Aşırı Güven

RPD Atla + İrkilme Etkisi



İnsan ve Yapay Zeka: Güncel İçerik Trendleri

AI-Generated Content Has Surpassed Human Content



Graphite araştırma ajansının yaptığı çalışmanın sonucunda, web üzerinde yayınlanan yapay zeka (YZ) üretimi makalelerin sayısının insan yazımı makaleleri geçtiği bulgusuna ulaşıyor. Araştırma yalnızca CommonCrawl üzerindeki İngilizce metinleri ve "Surfer" adlı tek bir yapay zeka tespit aracını baz aldığı için çeşitli kısıtlamaların söz konusu olduğu unutulmamalı. Ancak insanların YZ ile ilk taslağı hazırlayıp sonrasında kendi düzenlemelerini kattığı "karma" içerikler bu ölçüme dahil olmadığı için gerçekte YZ'nin içerik üretimindeki payı çok daha yüksek olabilir.



AI Now Writes as Many Online Articles as Humans,
<https://graphite.io/five-percent/ai-now-writes-as-many-online-articles-as-humans-do>

YZ'yi ne zaman kullan, ne zaman kullanma?

Görevin özelliğine göre rehber bir karar çerçevesi

← Risk →

⚠ Kaçın / İnsan zorunlu

- Hasta tanısı, tedavi planı
- Hukuki belge / sözleşme onayı
- Uçuş / operasyon kararları
- Finansal risk değerlendirme
- Safety critical görevler

Neden? YZ bilgi kesim tarihi, bağlam kaybı ve halüsinasyon riski yıkıcı sonuçlara yol açar.

✓ Kullan, mutlaka doğrula

- Teknik rapor taslağı hazırlama
- Araştırma özeti çıkarma
- Kod yazımı
- Veri analizi yorumlama
- İçerik üretimi

Neden? YZ hızlandırır ama uzmanlığın denetimi kritik.

→ YZ taslak atar, insan anaylar.

✦ İdeal YZ kullanımı

- Rutin metin düzenleme / çeviri
- Standart e-posta taslakları
- Tekrarlayan veri formatlama
- Brainstorming, fikir üretimi
- Bilgi arama / özetleme

Neden? Hata maliyeti düşük, çıktı kolayca denetlenebilir, Bilişsel yük azaltmak için uygun alan.

→ YZ iş akışına entegre edilebilir.

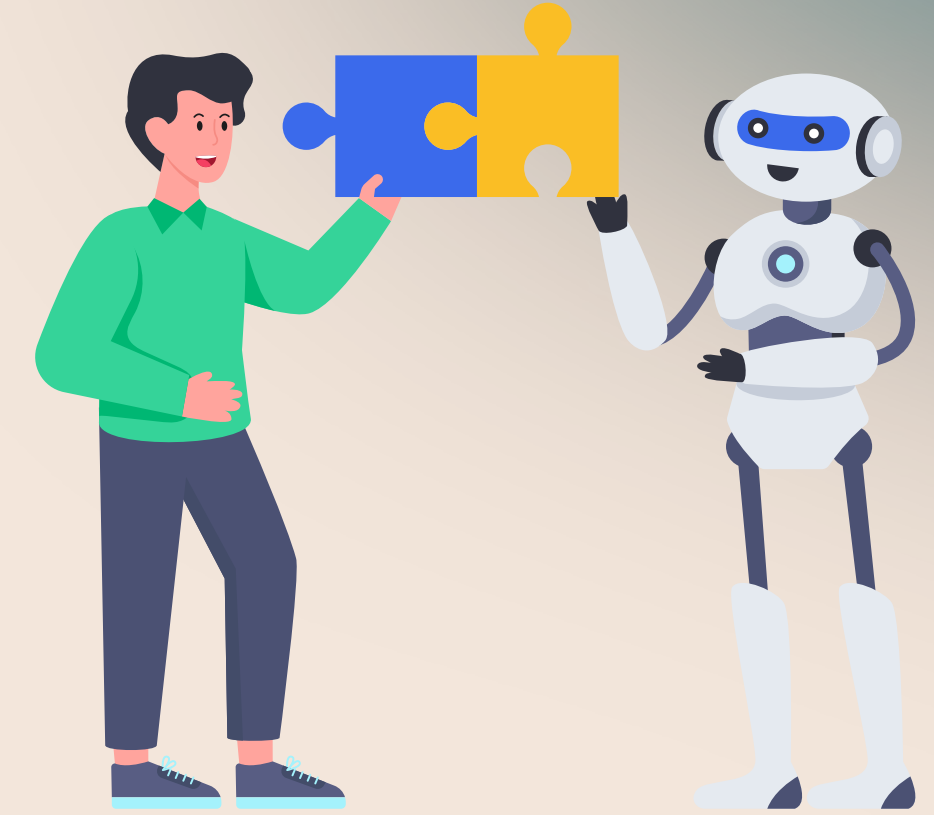
⚙ YZ destekli, insan kararı

- Uzman raporlara dayalı özetler
- Strateji seçenekleri üretme
- Anket/geri bildirim analizi
- Eğitim materyali hazırlama

Neden? Uzman bağlamını prompta taşırsanız güvenilir çıktı üretilir, RAG mimarisi burada tercih edilmeli.

→ YZ önerir, uzman seçer.

← Uzmanlık →



YZ
Matrisi



Açık Sistemler

Doğrulanmamış internet kirliliği,
açık uçlu halüsinasyon riski.

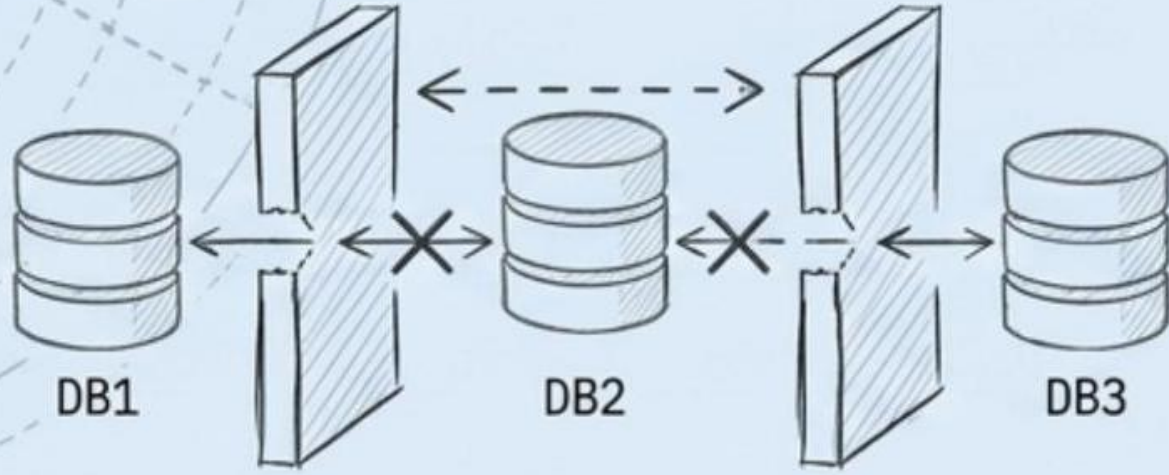


Kapalı RAG (Geri Çağırılabilir Üretim)

Sadece sizin yüklediğiniz onaylı
referans belgeler ile çalışır. İnternet
kullanmaz. "Zeminlendirilmiş"
(Grounded) mimari sayesinde her veri
noktası orijinal atıfa
bağlanır.



KAPALI SİSTEM ARAÇLARININ POTANSİYEL ZAFİYETLERİ



Sistemik Amnezi: Araçlar izole silolar halinde çalışır. Çapraz proje bilgi alışverişi kapalıdır. Oturum kapandığında **bağlam unutulur**, stratejik eğilimlerin fark edilmesi engellenir.

Kaynak Körlüğü: İşlemci bütçesi sınırı dayandığında model belgeyi bulamadığını itiraf etmez; kendi eğitim verilerine dönerek mükemmel görünümlü ancak **sahte vaka metinleri** uydurur.



TASARIM
HATASI

{font: JetBrains}
ERROR: UNRESOLVED

<DATASTREAM>



PDF KİLİTLİ

Tasarım İllüzyonları ve Kilitlenme: Üretilen görsellerde rastgele **kod parçaları** slayta basılabilir. Üstelik bu grafikler dışarıdan müdahaleye kapalı (lock-in) birleştirilmiş PDF'ler olarak sunulur.

Promptlar nasıl olmalı?

1. Spesifik Ol & Persona Belirle

Doğru: "Sen bir SMS uzmanısın, ICAO formatında özetle".

Yanlış: "Bana bir sunum yap".

2. Bağlam Sağla

Hedef kitleyi ve kısıtlamaları önceden tanımla (Örn: Sadece dar gövde istatistiklerini kullan).

3. Örnek Ver (Few-shot)

Kaliteli çıktının neye benzediğini göstermek için 2-5 adet referans örnek sun.

4. Adım Adım İste (Chain-of-thought)

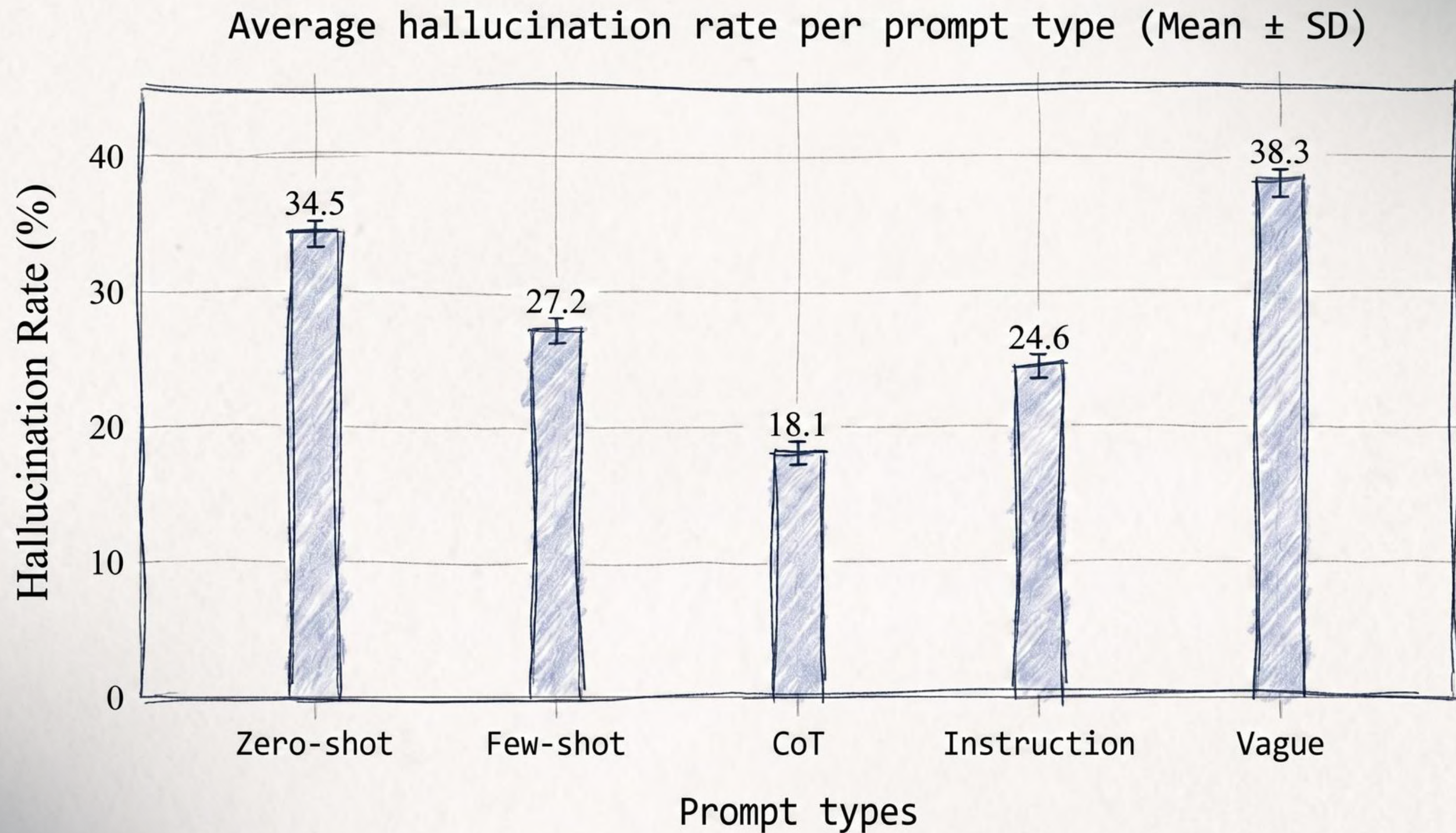
Doğrudan sonuç istemek yerine "Mantığını adım adım açıkla" komutunu kullan.

5. Görevleri Zincirle

Karmaşık veri setlerini tek seferde değil, alt görevlere bölerek (Önce tabloyu çıkar, sonra özetle) işle.

6. Gizlilik Çizgisini Korum

Kapalı sistem olsa bile, kişisel verileri veya son derece gizli kurum bilgilerini ve gizli evrakları prompt içine dahil etme.



Anh-Hoang D, Tran V and Nguyen L-M (2025) Survey and analysis of hallucinations in large language models: attribution to prompting strategies or model behavior.
doi: 10.3389/frai.2025.1622292

İçerik Oluşturma

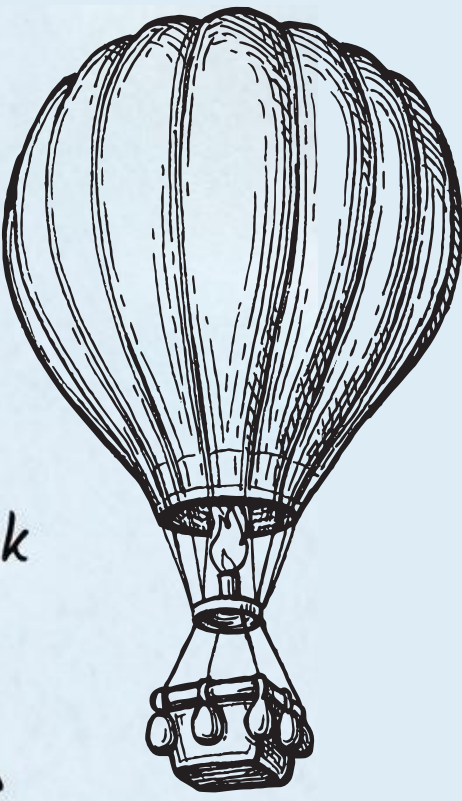
Çerçevesi

**Manuel kaynak tarama
ve Deep Research**



④ Gems kullanımı ve içerikleri araçlara dönüştürme

Persona tanımladığın gemsler sayesinde uzmanlık kazan, bilgileri araçlara dönüştürerek döngüyü tamamla.



③ Frontier Model - Gemini

Organize ettiğin ve kullanmak istediğin bilgilerin üstünden geçtikten sonra Gemini ile daha fazlasını araştırmak için keşfe çık.



② Grounding - NotebookLM

Belirlediğin kaynaklarla NotebookLM'i besle ve halüsinasyon ihtimalini azaltarak bilgileri kürate et.

NotebookLM



① Güvenilir kaynaklardan, referanslı bilgileri topla ve tüm içeriklerin doğruluğunu kontrol et.



KATILIMINIZ İÇİN
TEŞEKKÜRLER