



Yapay Zekâ ile Akıl Yürütme: Prompt, Diyalog ve Agent Yaklaşımları

Yapay Zeka

- ▶ Yapay zekâ, insan zekâsı gerektiren öğrenme, akıl yürütme, algılama, karar verme, dil kullanımı vb. görevleri makinelerin gerçekleştirebilmesini amaçlayan disiplinler arası bir alandır.
- ▶ Yapay zekânın temel alt alanları:
 - ▶ Makine Öğrenmesi (Machine Learning – ML)
 - ▶ Derin Öğrenme (Deep Learning – DL)
 - ▶ Doğal Dil İşleme (Natural Language Processing – NLP)
 - ▶ Bilgisayarlı Görü, Robotik, Planlama, Akıl Yürütme
- ▶ Önemli ayrım:
 - ▶ Yapay zekâ = hedef
 - ▶ Makine öğrenmesi = yöntem
 - ▶ Derin öğrenme = sinir ağları tabanlı gelişmiş yöntem

Yapay zekâda öğrenme

- ▶ Yapay zekâ bit yığınına doğrudan “anlamaz”; bitlerden istatistiksel yapılar ve temsiller inşa ederek öğrenir.
- ▶ Bilgisayarda her şey bittir.
- ▶ Bit, maddedir.
- ▶ Vektör, yapıdır.
- ▶ Öğrenme, ilişkidir.
- ▶ Anlam ise bu ilişkilerin geometrisidir.

Giriş

- ▶ “Bugün yapay zekâyı kullanmayan şirketler değil, **yapay zekâyı yanlış kullanan şirketler** risk altındadır.”
- ▶ Dijital dönüşüm → Otomasyon → Yapay zekâ → Agent çağının başlaması
- ▶ **AI yönetimin işidir.**
- ▶ Prompt soru sormak değil düşüncüyü yapılandırmaktır.
- ▶ “Yanlış tanımlanan yapay zekâ, yanlış sonuç üretir.”
- ▶ Agent: Hedefi olan, durumu izleyen, karar üreten ve aksiyon öneren yapay zekâ bileşenidir.
- ▶ “Agent’lar çalışanların yerine değil, yöneticilerin kör noktalarına çalışır.”
- ▶ “Agent; kurumsal verileri sürekli izleyen, normal davranışı öğrenen ve yöneticiyi karar öncesinde bilgilendiren akıllı bir yapay zekâ asistanıdır.”
- ▶ “Agent; kurumsal veri kaynaklarını sürekli izleyen, geçmiş davranışlardan norm çıkaran, anlamlı sapmaları tespit eden ve insan karar vericilere zamanında içgörü sunan yapay zekâ tabanlı bir karar destek bileşenidir.”
- ▶ **“Agent konuşmaz, izler, karşılaştırır ve zamanında uyarır.”**

Makine Öğrenmesi

- ▶ Makine öğrenmesi, sistemin açıkça programlanmadan, **veriden örüntü öğrenmesini sağlar.**
- ▶ Temel öğrenme türleri:
 - 1) Denetimli Öğrenme (Supervised)
Girdi-çıkı çiftleri vardır
Örnek: Spam e-posta sınıflandırma
 - 2) Denetimsiz Öğrenme (Unsupervised)
Etiket yoktur
Örnek: Müşteri segmentasyonu
 - 3) Pekiştirmeli Öğrenme (Reinforcement Learning)
Ödül-ceza mekanizması
Örnek: Oyun oynayan agent'lar

Derin öğrenme

- ▶ Derin öğrenme, çok katmanlı yapay sinir ağlarını kullanır.
- ▶ Yapay sinir ağı temel bileşenleri:
 - ▶ Girdi katmanı
 - ▶ Gizli katmanlar
 - ▶ Çıkış katmanı
 - ▶ Ağırlıklar (weights)
 - ▶ Aktivasyon fonksiyonları
 - ▶ Kayıp (loss) fonksiyonu
- ▶ Bu yapı, karmaşık ilişkileri modelleyebilir:
 - ▶ Görüntü tanıma
 - ▶ Ses tanıma
 - ▶ Dil anlama

Dil Modeli

- ▶ Dil modeli, bir dildeki kelimelerin veya sembollerin olasılık dağılımını öğrenen matematiksel bir modeldir. Basit tanımıyla, “Bir kelimedenden sonra hangi kelime gelme ihtimali yüksektir?”
- ▶ Dil Modelleri “Anlıyor” mu?
 - ▶ İnsan gibi bilinçli anlama yok
 - ▶ İstatistiksel örüntü ve bağlam öğrenme var
 - ▶ Niyet, duygu, bilinç yok
 - ▶ Dilsel tutarlılık ve muhakeme benzeri davranış var
- ▶ Dil modeli anlamı temsil etmez, anlamın izlerini işler.
- ▶ Büyük dil modelinin:
 - ▶ belirli bir görev için değil,
 - ▶ dilin genel yapısını öğrenmek için
 - ▶ çok büyük ve çeşitli metinler üzerinde eğitildiği aşamadır.
- ▶ Bu aşamada model: dilbilgisini, kelime anlam ilişkilerini, bağlamı, cümle ve metin tutarlılığını insan etiketine ihtiyaç duymadan öğrenir.
- ▶ Ön eğitim = “Dil sezgisi kazanma” süreci
- ▶ **Ön Eğitimde Model şunu öğrenir:**
- ▶ “Verilen bir bağlamdan sonra **hangi kelime (token) gelmelidir?**”

LLM (Büyük Dil Modelleri)

- ▶ LLM = Veri analitiğinde doğal dil ile analiz yapmayı mümkün kılan akıllı model.
- ▶ LLM, çok büyük miktarda metin üzerinde eğitilmiş, kelime/sembol dizilerini modelleyebilen derin öğrenme modelidir. Amaç: verilen bir giriş metnine göre mantıklı, bağlama uygun ve akıcı doğal dil çıktısı üretmek.
- ▶ LLM (Large Language Model), büyük miktarda metin verisi üzerinde eğitilmiş, *insan benzeri metin anlayışı ve üretimi yapabilen bir yapay zekâ modelidir.*
- ▶ Bunlar, doğal dil işleme (NLP) alanındaki en gelişmiş modellerdir — örneğin: GPT-5, Claude, Gemini, LLaMA, Mistral vb.
- ▶ LLM'ler, dilin yapısını, anlamını ve bağlamını istatistiksel olarak öğrenerek: *Metin üretir, Soru yanıtlar, Özet çıkarır, Kod yazar, Veri analizinde içgörü üretir*

Bir yapay zekâ modeli hata yapar çünkü:

- ▶ Model, “Doğru nedir?” sorusunu sormaz “En olası nedir?” sorusunu sorar. Bu nedenle yanlış ama makul, tutarlı ama hatalı, akıllı ama kör olabilir.
- ▶ Yapay zekâ modelinin hatası, onun başarısızlığı değil, insan bilgisinin ve matematiksel temsilin sınırlarının bir yansımasıdır.
 - 1) Gerçekliği değil, veriyi öğrenir.
 - 2) Veriler eksik ve yanlıdır.
 - 3) Matematiksel olarak yaklaşıktır.
 - 4) Belirsizlik kaçınılmazdır.
 - 5) Amaç fonksiyonu sınırlıdır.
 - 6) Anlam ve niyet yoktur.
 - 7) Dünya değişir.

Learning Agent

- ▶ “Learning Agent Improves” ifadesi genellikle öğrenen asistanların (learning agents) zaman içinde deneyimlerinden yola çıkarak performanslarını geliştirme yeteneğini anlatır. Bu kavram hem yapay zeka hem de veri analitiği alanında önemli bir yere sahiptir — çünkü **veriyle beslendikçe kendini geliştiren sistemler fikrini temsil eder.**
- ▶ Learning agent (öğrenen asistan), çevresinden (ortamdan) veri toplayan, bu verileri analiz eden ve kendi performansını iyileştirmek için öğrenen bir yapay zeka sistemidir. Yani **klasik “programlanmış sistem” değil, veriden öğrenen bir sistemdir.**
- ▶ “A learning agent improves its performance over time based on past experiences and feedback.” (Bir öğrenen asistan, geçmiş deneyimlerine ve aldığı geri bildirimlere dayanarak zaman içinde performansını geliştirir.) Bu “improves” kısmı, öğrenmenin sürdürülebilir ve dinamik olduğunu vurgular. **Yani asistan sadece bir kez öğrenip kalmaz; her yeni veriyle daha iyi hale gelir**



Yapay Zeka Veri Merkezleri

Yapay Zekâ Platformları: Bir Diyalog Ortağı



AI Platformları bir bilgi otoritesi, bir hakem ya da bir nihai karar verici değildir. AI Platformları bir diyalog simülasyonu olarak konumlanır. Bu nedenle "En doğru prompt budur" iddiasında bulunulmamalı, evrensel reçeteler sunulmamalı, özellikle **yapay zekâ insan zihninin yerine konmamalı**.

AI Platformları bilgi veren değil, düşünmenin sınırlarını zorlayan bir diyalog ortağı, bir asistan olarak konumlanır. Genellikle cevap odaklı diyaloglarla ortaya çıkar. Bu tür diyaloglarda kullanıcı, süreci değil sonucu talep eder.

AI Platformları ile çalışmak, bilişsel sorumluluğu devretmek değil; bilinçli biçimde düşünmeyi üstlenmektir. Ortaya çıkan metin bilgi hissi verse de araştırma üretmez; çünkü varsayımlar görünmez kalır, alternatifler üretilmez ve yanıt doğrulanmadan kabul edilir.

İnsan Beyni ve Yapay Zekâ: Temel Farklar

İnsan Beyni

Niyet okur; eksik ya da ima edilen duygusal ifadeleri anlar; bağlamları ve boşlukları sezgisel tamamlar; çelişkiyle yaşayabilir. "Ne demek istediğini anladım" diyebilir.

AI Platformları (LLM)

Şimdilik niyet okumaz, olasılık hesaplar; açıkça söylenmeyeni varsayım olarak üretir, çelişkiyi fark edebilir ama önceliklendirme ister. AI Platformları anlamı tahmin eder.

Bu yüzden diyalogda sezgisel değil, yapısal iletişim gerekir. Bu ders notunun amacı, AI Platformları ve benzeri büyük dil modellerinin (LLM) nasıl çalıştığını öğretmekten çok, araştırmacıyı yapay zekâ karşısında pasif bir alıcı olmaktan çıkarıp, sorgulayan, sınayan ve yönlendiren bir bilişsel özne hâline getirmektir.

Yapay Zekâ Veri Merkezleri: Stratejik Önemi

Yapay zekâda veri merkezleri, yalnızca "altyapı" değil; veri odaklı stratejik bir unsur olduğu için kritik öneme sahiptir. **Veri merkezleri, yapay zekânın "beyni" değil; onun çalışmasını mümkün kılan fiziksel bedenidir.** Model, algoritma ve veri ne kadar iyi olursa olsun hesaplama, enerji, soğutma ve ölçek yoksa yapay zekâ yalnızca bir fikir olarak kalır.

Yapay zekâ modelleri soyut algoritmalarından ibaret olsa da onların gerçek dünyada çalışabilmesi, devasa fiziksel altyapılara bağlıdır. **Büyük veri kümelerinin işlenmesi, modellerin eğitilmesi ve sürekli olarak hizmet verebilmesi için yüksek işlem gücü, geniş depolama kapasitesi, kesintisiz enerji ve etkili soğutma sistemleri gerekir.**

Veri Merkezlerinin Kritik Katmanları

1

Hesaplama Gücü

Yapay zekâ "hesaplama-yoğun" bir disiplindir. Derin öğrenmede milyarlarca parametre, trilyonlarca işlem (FLOP), uzun süreli, paralel hesaplama gerektirir. Veri merkezleri, GPU/TPU gibi özel donanımların yüksek paralellikte, kesintisiz, senkronize çalışmasını mümkün kılar.

2

Veri Yakınlığı

Veriyi taşımak, hesaplamaktan daha maliyetlidir. Veri merkezleri verinin yerinde işlenmesini, düşük gecikme (latency), yüksek bant genişliği sağlar. Bu özellikle büyük dil modelleri, gerçek zamanlı çıkarım ve video/sensör verisi için kritiktir.

3

Enerji ve Soğutma

Yapay zekâ termal bir problemdir. Yüksek enerji tüketimi ciddi ısı üretir. Isı performansı düşürür veya donanım hasarı oluşturur. Özel soğutma sistemleri, enerji sürekliliği ve verimli güç dağıtımı olmaksızın AI çalışamaz.

4

Ölçeklenebilirlik

Yapay zekâ deneme-yanılma, hızlı tekrar ve büyük ölçekli deney gerektirir. Veri merkezleri yatay ya da dikey ölçekleme, kaynak havuzu, çoklu deneylerin eşzamanlı yürütülmesi imkânı sağlamalıdır.

5

Güvenlik ve Egemenlik

Eğitim verisi, model ağırlıkları ve çıkarım sonuçları stratejik varlıktır. Veri merkezleri veri egemenliği, regülasyon uyumu (ve güvenlik açısından kritiktir. Yapay zekâyâ sahip olmak veri merkezine sahip olmak demektir.

GPU ve TPU: Paralel Hesaplamanın Gücü

Paralel Hesaplama Nedir?

Paralel hesaplama, büyük ve karmaşık bir problemi tek bir işlemci yerine birden fazla işlemci ya da çekirdek kullanarak aynı anda çözme yaklaşımıdır. Bu yöntemde yapılacak iş, birbirinden bağımsız ya da kısmen bağımlı alt parçalara bölünür ve bu parçalar eşzamanlı olarak farklı işlem birimlerinde çalıştırılır.

Böylece hesaplama süresi önemli ölçüde kısalır, daha büyük veri kümeleri daha hızlı işlenebilir ve sistem performansı artar. Günümüzde çok çekirdekli işlemciler, GPU'lar ve dağıtık sistemler sayesinde paralel hesaplama; yapay zekâ, bilimsel simülasyonlar ve büyük veri analizi gibi alanlarda temel bir ihtiyaç haline gelmiştir.

GPU ve TPU

GPU (Graphics Processing Unit): Çok sayıda çekirdeğe sahip mimarileri sayesinde binlerce işlemi aynı anda gerçekleştirebilir. Görüntü işleme, bilimsel hesaplama ve derin öğrenme eğitimlerinde yaygın kullanılır.

TPU (Tensor Processing Unit): Google tarafından özel olarak makine öğrenmesi algoritmalarını hızlandırmak amacıyla tasarlanmış, matris ve tensör işlemlerine optimize edilmiş donanımlardır.

Veri Hareketi: Gizli Maliyet



Veriyi taşımak çoğu zaman hesaplamaktan daha maliyetlidir çünkü modern bilgisayar sistemlerinde işlemcilerin hesaplama hızı, bellek ve veri iletim hızlarından çok daha yüksektir. Bir işlemci saniyede milyarlarca işlem yapabilirken, büyük veri kümelerinin bellekten işlemciye ya da ağ üzerinden başka bir sisteme aktarılması gecikmelere ve enerji tüketimine yol açar.

Veri hareketi sırasında bant genişliği sınırlamaları, önbellek gecikmeleri ve iletişim protokollerinin getirdiği ek yükler performansı düşürür. Bu nedenle günümüzde yüksek performanslı sistemlerde temel hedef, veriyi mümkün olduğunca bulunduğu yerde işlemek ve gereksiz veri taşımayı en aza indirerek hesaplama verimliliğini artırmaktır.



Yapay Zekâ: Jeopolitik Bir Güç Unsuru

Yapay zekâ artık bir kritik altyapıdır. Bugün limanlar, enerji santralleri, telekom ağıları ne ise, AI veri merkezleri de odur. Bu nedenle ABD, Çin, AB kendi AI data center stratejilerini oluşturuyor. GPU tedariki jeopolitik meseledir. Enerji ve AI yeni sanayi politikasıdır.

GPU tedariki jeopolitik bir mesele haline gelmiştir çünkü yapay zekâ ve yüksek performanslı hesaplama altyapılarının temelini oluşturan bu donanımlar, ülkelerin teknolojik rekabet gücünü doğrudan belirlemektedir. Gelişmiş GPU ve çip üretimi sınırlı sayıda ülke ve şirketin elinde yoğunlaştığından, tedarik zincirleri stratejik önem kazanmış; ihracat kısıtlamaları ve teknoloji ambargoları uluslararası politikaların bir aracı olmuştur.

Küresel Veri Merkezi Ekosistemi



Microsoft Azure

Yüzlerce veri merkezi işletiyor. AI odaklı altyapı yatırımlarına milyarlarca dolar ayırıyor ve yeni bölgelerde genişleme planları var.



Amazon AWS

Küresel çapta yüzlerce veri merkezi ve çok sayıda yeni bölge planlıyor. AI iş yüklerini desteklemek üzere özel donanımlar ve yeni tesisler kuruyor.



Google Cloud

Veri merkezi ağını genişletiyor. Avrupa, ABD ve diğer bölgelerde yeni AI odaklı altyapı yatırımları yapıyor.



Meta Platforms

Facebook/Instagram ve AI hizmetleri için çok sayıda veri merkezi var ve yeni tesis yatırımlarını sürdürüyor.



CoreWeave

GPU yoğun veri merkezleri kuran ve işletme alanında hızlı büyüyen bir oyuncu. Hem ABD'de hem Avrupa'da veri merkezleri var.



Equinix

70+ ülkede 250+ veri merkezi ile geniş bir küresel ağa sahip. Yeni AI odaklı tesisler ve genişleme planları var.

Stargate Projesi: 500 Milyar Dolarlık Vizyon



Projenin Özellikleri

Kurucu Ortaklar: OpenAI, SoftBank Group, Oracle, MGX (Masayoshi Son başkanlığında)

Yatırım Hedefi: Toplamda yaklaşık 500 milyar USD yatırım ve 10 gigawatt AI veri merkezi kapasitesi hedefleniyor.

Ana Lokasyonlar: Abilene (Texas), UAE (Abu Dhabi), Güney Kore, ve ABD'nin diğer eyaletlerinde tesisler

Teknoloji Ortakları: NVIDIA (GPU), Samsung & SK (yarı iletken), Cisco, Arm, Microsoft

Hedef: Generatif yapay zekâ modellerini eğitmek ve çalıştırmak için büyük ölçekli, yüksek kapasiteli veri merkezleri ağı oluşturmak; AI araştırma, inovasyon ve endüstriyel aplikasyonlarda küresel rekabet gücünü artırmak.



Prompt

İnsan ve Yapay Zekâ İletişiminin Temelleri

Temel Fark

İnsan zihni biyolojik, tarihsel ve bedensel bir varoluşa sahiptir. Anlam, yalnızca dilsel yapılardan değil; niyet, deneyim ve bağlamdan beslenir. **AI Platformları ise sayısal bir modeldir ve parametreleri, geçmiş dil kullanımlarının istatistiksel izlerini taşır.**

Kritik Ayrım

İnsan anlam üretir, AI Platformları anlam benzeri örüntüler üretir. İnsanlar ima eder, AI platformları yapısal tamamlama yapar. İnsan zihni duyguları, bedensel deneyimleri, kültürel geçmişi ve bilinçli niyetleri olan canlı bir sistemdir. Yapay zekâ platformları ise böyle bir yaşantıya sahip değildir; yalnızca **geçmişte üretilmiş metinlerdeki istatistiksel ilişkileri öğrenen matematiksel modellerdir.**

Bu nedenle insan gerçekten "anlam üretirken", yapay zekâ yalnızca anlamlıymış gibi görünen örüntüleri olasılıksal olarak tamamlar. **AI Platformları, insan zihninin dijital versiyonu değil; insan dilinin istatistiksel modelidir.**

Öğrenme ve Yanılma Dinamikleri



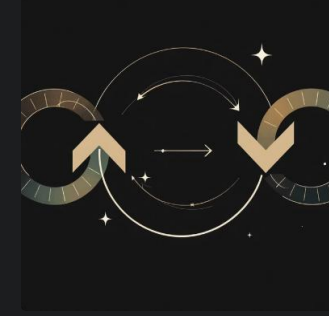
İnsan Öğrenmesi

İnsan öğrenmesi süreklidir; hata, duygu ve motivasyon öğrenmeyi dönüştürür ve geliştirir. Tek bir örnek bile güçlü bir zihinsel model oluşturabilir. İnsan unutar ama öğrenir, inatla yanlışta ısrar edebilir ancak anlamlı kırılmalar üretir.



AI Platformları Öğrenmesi

AI Platformları kullanım sırasında öğrenmez. Eğitimi tamamlanmış bir parametre uzayı içinde oluşturulan örüntüsel davranışta yanıt üretir. **AI Platformları yanlışta ısrar etmez, ancak doğrudan ısrar etmenin ne anlama geldiğini de bilmez.** İkna edilince yön değiştirir, hatırlar ama öğrenmez.



İteratif Diyalog

AI platformları ilk soruya verdikleri yanıtta genellikle genel ve ortalama bir tahmin üretirler. Model, kullanıcı niyetini, ihtiyaç duyulan ayrıntı düzeyini ve bağlamı tek seferde tam olarak bilemez. **İteratif diyalog sürecinde kullanıcı düzeltmeler yapar, ek bilgiler verir; bu yeni veriler model için yönlendirme işlevi görür. İkinci ve üçüncü turda cevaplar daha hedefe yönelik, daha ayrıntılı ve daha doğru hale gelir.**

Görev Yönetimi ve İletişim Kuralları



Yanlış Görev Örneği

"Bunu açıkla, örnek ver, grafik çiz ve eleştir."

Doğru Görev Örneği

"Önce açıkla. Sonra basit bir örnek ver. En son eleştirel değerlendirme yap."

Kısıt Örnekleri

- Metafor (mecaz) kullanma
- Felsefi yoruma girme
- Popüler bilim dili kullanma

İdeal Prompt Şablonu

Rol:

Alan uzmanı kimliği tanımlayın

Amaç:

Ne için kullanılacağını belirtin

Düzyey:

Lisans / yüksek lisans / doktora

Kapsam:

Hangi konular var / yok

Biçim:

Başlıklar, matematik, örnek, grafik

Epistemik Sınır:

Kanıtı dayalı / spekülasyon değil

Örnek Şablon

"Doğrusal olmayan dinamikler ve yapay zekâ öğrenme kuralları konusunda çalışan bir araştırmacı gibi cevap ver. Amaç: Akademik makale için teorik çerçeve oluşturmak. Düzey: Doktora. Kapsam: Yalnızca iteratif haritalar ve kararlılık. Biçim: Denklem + kavramsal açıklama. Epistemik sınır: Literatüre dayalı, spekülasyon yapma."

En Sık Yapılan Hatalar ve Sonuçları

Hata:	Sonuç:
Belirsiz soru	Yüzeysel cevap
Çok görevli tek cümle	Dağınık çıktı
Kısıt yok	Gereksiz uzunluk
Sezgisel ima	Yanlış varsayım
Tek tur	Ortalama kalite

AI Platformları ile iyi diyalog, güzel cümle kurmak değil; düşünceyi biçimlendirmeyi makineye devretmeden önce netleştirmektir. Belirsiz ifade örtük varsayım içerdiğinden hatalı çıktı zinciri üretir ve AI Platformları kullanımında sistematik bir risk oluşturur.

İnsan ve AI Hataları Karşılaştırması

İnsan Hataları

Bilişsel Yanlılık

Önyargılarla düşünme

Duygusal Körlük

Duyguların mantığı engellemesi

Aşırı Genelleme

Sınırlı veriden geniş sonuç

İnanç Temelli Direnç

Kanıtı rağmen ısrar

AI Platformları Hataları

Halüsinasyon

Gerçek olmayan bilgi uydurma

Kaynak Uydurma

Var olmayan referanslar

Aşırı Özgüvenli Yanlılık

Kesin gibi sunulan hatalar

Bağlam Hatası

Yanlış önceliklendirme

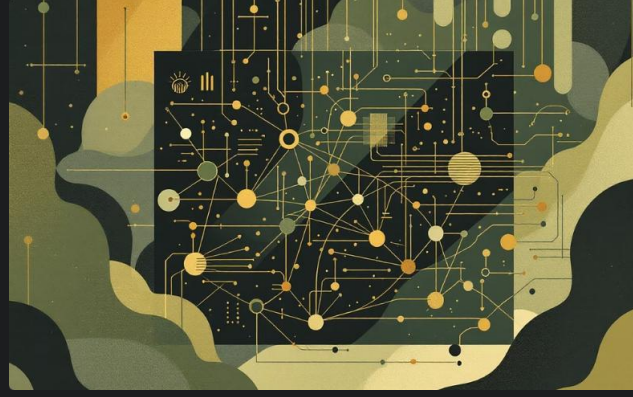
İnsan inatla yanlışta ısrar edebilir, AI Platformları ikna edilince yön değiştirir. İnsan unutulur ama öğrenir, AI Platformları hatırlar ama öğrenmez.

Yaratıcılık ve Ahlak Boyutları



İnsan Yaratıcılığı

Yaratıcılık = sapma + niyet + risk. Anlamalı kırılmalar üretir ve kültürel bağlamdan kopabilir. İnsan yenilik yaratır, risk alır ve normdan kopar.



AI Yaratıcılığı

Yaratıcılık = olasılık uzayında gezinme. "Yeni" görünen kombinasyonlar üretir ancak risk almaz, normdan kopmaz. AI Platformları yenilik simüle eder.



Ahlak ve Sorumluluk

İnsan: Ahlaki fail (Evet), Sorumluluk (Bireysel), Vicdan (İçsel mekanizma)

AI Platformları: Ahlaki fail (Hayır), Sorumluluk (Kullanıcı + geliştirici), Vicdan (Kural tabanlı filtre)

Prompt Yazımında Cümle Tipleri

01

Rol Tanımlayıcı Cümleler

AI Platformları'yi bir zihinsel moda sokar. Model, hangi bilgi alanlarını ve hangi dil tonunu kullanacağını seçer. Örnek: "Akademik bir araştırmacı gibi, doğrusal olmayan dinamikler perspektifinden açıkla."

04

Yapı Talep Eden Cümleler

AI Platformları içerikten çok yapıyı doğru takip eder. Başlıklar ve yapı isteyerek cevabın kalitesini artırın.

02

Amaç Belirleyici Cümleler

"Ne için?" sorusu mutlaka cevaplanmalı. Örnek: "Prompt yazımını, üniversite düzeyinde ders anlatımında kullanılabilecek şekilde açıkla."

03

Kısıtlayıcı Cümleler

AI Platformları doğal olarak fazla üretir. Kısıt verilmezse dağınık. Uzunluk, düzey, biçim ve yasak kısıtları belirtin.

05

Epistemik Düzey Belirten Cümleler

"Spekülatif olma", "Literatüre dayalı yanıt ver", "Varsayım yapıyorsan açıkça belirt" gibi sınırlar çizin.

Doğru-Yanlış Prompt Örnekleri

✗ Konu Öğrenme - Yanlış

"Yapay zekâyı anlat."

Neden yanlış: Düzey belirsiz, kapsam belirsiz, sonuç yüzeysel ve ansiklopedik cevap.

✗ Matematik/Teknik - Yanlış

"Formülü açıkla."

Neden yanlış: Bağlam yok, ezber açıklama.

✗ Eleştirel Düşünme - Yanlış

"Bu doğru mu?"

Neden yanlış: Evet-hayır cevabı düşünmeyi kapatır.

✓ Konu Öğrenme - Doğru

"Yapay zekâyı, bilgisayar mühendisliği 2. sınıf öğrencisi düzeyinde; sembolik AI ve makine öğrenmesi farkına odaklanarak açıkla."

Kazanım: Hedefli, sınava uygun bilgi.

✓ Matematik/Teknik - Doğru

Kazanım: Kavramsal derinlik.

"Bu formülün nereden geldiğini, hangi varsayımlar altında geçerli olduğunu ve hangi durumda başarısız olacağını açıkla."

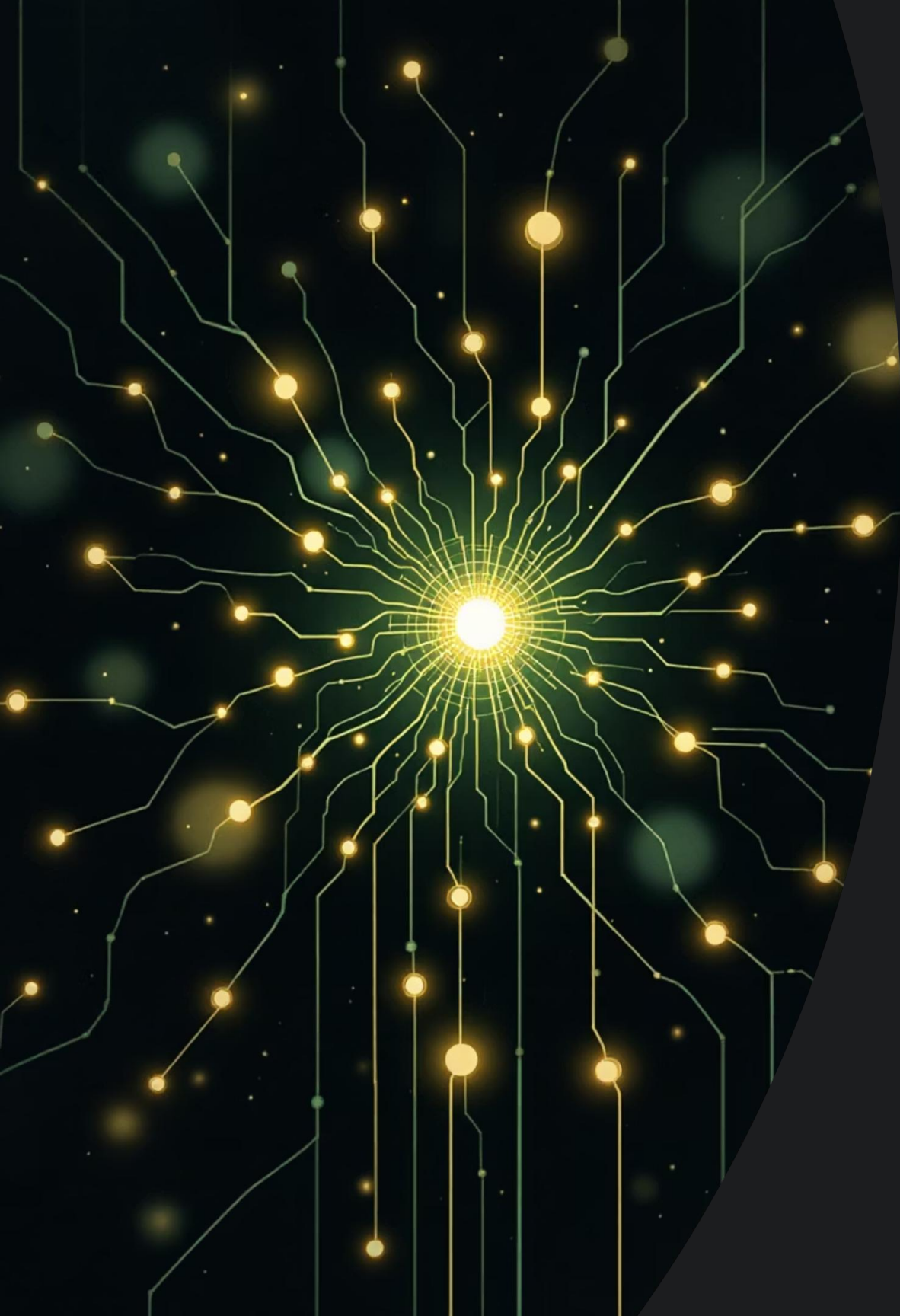
✓ Eleştirel Düşünme - Doğru

"Bu yaklaşımın güçlü ve zayıf yönlerini karşılaştırmalı olarak analiz et."

Kazanım: Analitik düşünme.

Mini Kontrol Listesi: Soru sormadan önce kendinize sorun: Düzeyi belirttim mi?

Amacı söyledim mi? Ne istemediğimi yazdım mı? Sonucu ben mi üreteceğim?



Agent

Agent Nedir ve Neden Önemlidir?



Agent, yapay zekâ sistemlerinde amaç, durum, geri bildirim ve kontrol mekanizmalarını birleştiren gelişmiş bir yaklaşımdır. Tek seferlik cevap üreten basit promptlardan farklı olarak, agent'lar sürekli görevleri yerine getirir ve karmaşık süreçleri yönetir.

Bu sunumda, agent tasarımının stratejik yol haritasını, prompt şablonlarını, iş süreçlerini ve seçim kriterlerini kapsamlı bir şekilde inceleyeceğiz. **Agent geliştirme, yapay zekâyı öğretmek değil; düşünceyi disipline etmektir.**

Temel Ayrım: Tek soru-tek cevap için prompt yeterlidir. Ancak süreç, tekrar ve karar gerektiren durumlar için agent gereklidir.

Agent Hazırlarken Doğruya Yönlendirme: 10 Aşamalı Strateji

Agent tasarımı, amaç, kural ve geri bildirimle doğruya yaklaşan bir bilişsel araç olarak yapılandırılmalıdır. İşte kapsamlı stratejik yol haritası:

01

Amaç ve Doğruluk Tanımı

Doğruyu tekil cevap değil, kabul edilebilir çözüm kümesi olarak tanımlayın. Bilgi doğruluğu, muhakeme tutarlılığı ve etik uygunluğu ayırın. Varsayımları açık, gerekçeli ve sınanabilir cevaplar hedefleyin.

02

Rol ve Yetki Sınırları

Agent'in yapabilecekleri, yapamayacakları ve ne zaman durması gerektiğini net şekilde belirleyin. Halüsinasyon ve aşırı özgüveni azaltmak için sınırlar koyun.

03

Bilişsel Mimari

Çok aşamalı düşünme süreci tasarlayın: Problemi yeniden formüle etme, varsayımları çıkarma, alternatifleri üretme, karşı-örnek arama ve en makul çözümü seçme.

04

Prompt Katmanları

Üst prompt (misyon), görev prompt'u ve denetleyici prompt olmak üzere katmanlı yapı oluşturun. Her katman farklı bir işlev görür.

05

Şeytanın Avukatı Modülü

Her kritik çıktıda otomatik çalışan sorgulama mekanizması: En güçlü itiraz, karşı-örnek ve genellenebilirlik sınırını test edin.

Stratejik Yol Haritası Devamı

1

Belirsizlik ve Durma Protokolleri

Agent'e ne zaman emin olmadığını söyleyeceği, ek veri isteyeceği ve cevap üretmemesi gerektiğini öğretin. Örnek: "Bu bilgi için güven düzeyim düşük (%40). Ek veri gerekir."

2

Geri Besleme ve Kalibrasyon

Agent çıktılarını tutarlılık, gerekçelendirme kalitesi ve varsayım açıklığı ölçütleriyle değerlendirin. Bu değerlendirmeyi sonraki görevlerde prompt ayarlaması olarak kullanın.

3

Yanlış Türlerini Tanıma

Agent hata yaptığında nasıl hata yaptığını belirtsin: Eksik bağlam, yanlış genelleme veya varsayım hatası. Bu öğrenen sistem hissi yaratır.

4

İnsan-in-the-Loop

Agent hiçbir zaman tamamen yalnız bırakılmamalı. Kritik kararlarda insan onayı, çelişkili sonuçlarda durma mekanizması olmalıdır.

5

Olgunluk Kontrol Listesi

Agent varsayımlarını açıklıyor mu? Alternatifleri üretiyor mu? Kendini eleştirebiliyor mu? Emin olmadığında durabiliyor mu? Bu sorulara evet cevabı verilmelidir.

İyi bir agent, doğruyu söyleyen değil; doğruya nasıl yaklaşıldığını gösteren sistemdir. Bu fark, özellikle tez, makale ve karar-destek sistemlerinde belirleyicidir.

Agent Tasarımı İçin Prompt Şablonları

Agent tasarımında kullanılacak temel prompt yapıları ve örnekleri:

Üst Prompt (Misyon)

Agent'in kişiliğini tanımlar, her görevde korunur. Rol, temel amaç ve ilkeleri içerir: Varsayımları belirt, alternatifleri üret, emin olmadığında dur, tutarlı ol.

Görev Prompt'u

Problem alma şablonu. Kullanıcı problemini alır, kendi cümleleriyle yeniden tanımlar, açık ve örtük varsayımları listeler.

Bilişsel Akış

Düşünme zinciri şablonu. Problemi daraltma, alternatif yaklaşımlar üretme, güçlü/zayıf yönleri belirleme, karşı-örnek arama adımlarını içerir.

Şeytanın Avukatı

Rol değiştirme modülü. Üretilen çözümü sert eleştirmen gibi değerlendir: En kırılgan varsayım hangisi? En yıkıcı itiraz ne olurdu?

Denetleyici Prompt

Kalite kontrol katmanı. Varsayımları sorgular, genellenebilirlik sınırını test eder, varsayım bozulursa sonucun geçerliliğini kontrol eder.

Belirsizlik Protokolü

Durma koşulları: Veri yetersizliği, muğlak problem tanımı, %60 altı güven durumunda cevap üretme. Ne bilinmediğini ve hangi bilginin gerektiğini belirt.

Daha Fazla Prompt Şablonu

Çıktı Format Prompt'u

Yapılandırılmış yanıt şablonu:

1. Problem Tanımı (yeniden formülasyon)
2. Varsayımlar
3. Alternatifler
4. Eleştiriler (şeytanın avukatı)
5. En makul sonuç + güven düzeyi

Hata Tipi Etiketleme

Hata veya belirsizlik durumunda etiketleme sistemi. Eksik veri, yanlış varsayım, genelleme hatası gibi kategoriler kullanılır.

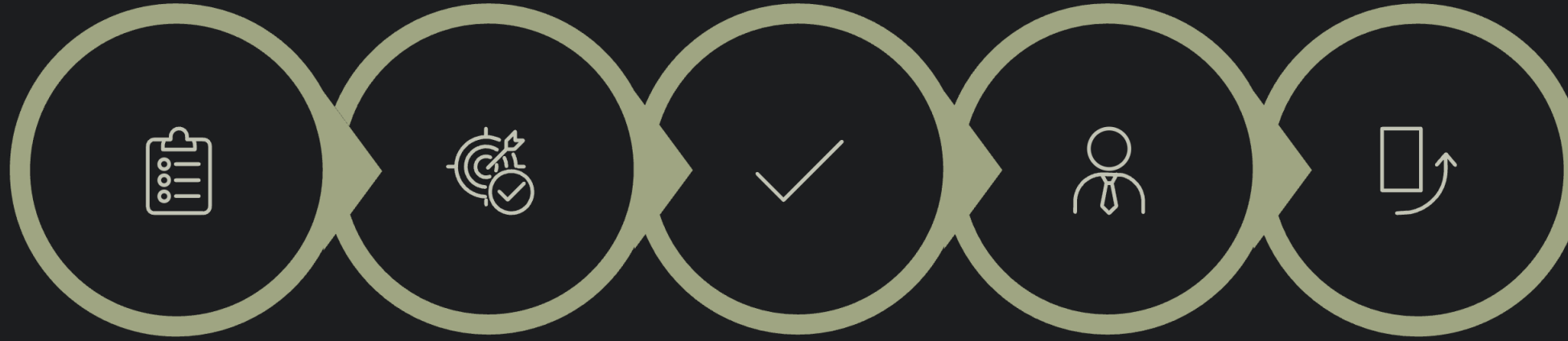
İnsan-Onay Modülü

Çıktının nihai karar olmadığını belirtir. Kullanıcıdan geri bildirim ister: Hangi varsayımı kabul etmiyorsun? Hangi alternatif genişletilmeli?

"Agent cevap üretmez; doğruya giden yolu görünür kılar."

Agent Geliştirme için İş Süreçleri

Stratejiden operasyona geçiş için kapsamlı iş süreçleri rehberi:



İhtiyaç Tanımı

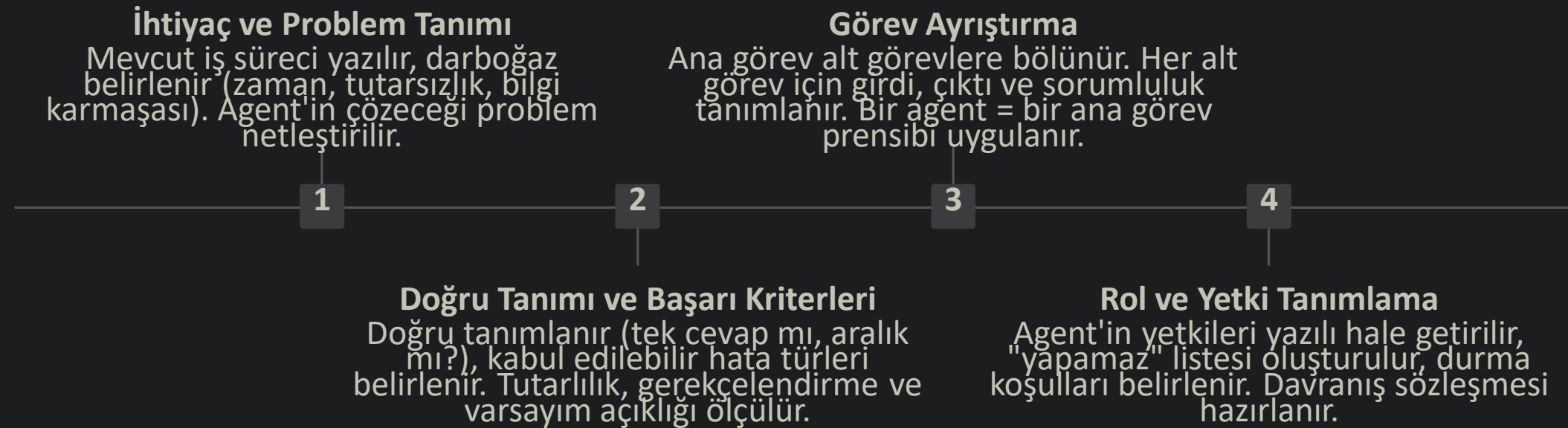
Başarı Kriteri

Görev
Ayrıştırma

Rol Tanımı

İş Akışı
Tasarımı

Bu iş süreçleri, agent geliştirmeyi model seçmekten çıkarıp düşünceyi iş sürecine dökme pratiğine dönüştürür.



İş Süreçleri Devamı ve Uygulama

1

İç Akış Tasarımı

Adım adım düşünme akışı, karar noktaları ve geri dönüş noktaları tanımlanır.

2

Denetim ve Kalite

Şeytanın avukatı rolü, self-check soruları ve çıktı filtreleri kurulur.

3

Belirsizlik Yönetimi

Belirsizlik eşikleri, "emin değilim" senaryoları ve ek veri isteme akışı belirlenir.

4

İnsan-Onayı

Onay noktaları, varsayım onayı ve alternatif seçimi mekanizmaları kurulur.

5

Test ve Doğrulama

Yanlış girişler, çelişkili talepler ve eksik veri senaryolarıyla test edilir.

6

Sürekli İyileştirme

Yanlış türleri etiketlenir, tekrarlayan hatalar analiz edilir, prompt ve akış güncellenir.

Tipik İç Akış Örneği

1. Girdiyi yeniden ifade et
2. Varsayımları çıkar
3. Alternatifleri üret
4. Eleştir
5. Sonuç öner

Önemli: Agent önerir, insan karar verir.
Sorumluluk her zaman insanda kalır.

Agent Stratejik Yol Haritası Hazırlayacak Uzman Profili

Agent geliřtirmek için gereken kritik yetkinlikler ve uzman özellikleri:



Problem Parçalama

En kritik yetkinlik. Büyük problemi küçük parçalara ayırma, sorun ile çözüm yolunu ayırma, yanlış soruyu çözen sistem riskini fark etme yeteneđi.



Kavramsal Akıl Yürütme

Varsayım-sonuç, gerekçe-iddia, olası-kesin ayrımlarını yapabilme. Halüsinasyon ve aşırı genellemeyi önleyen mantık disiplini.



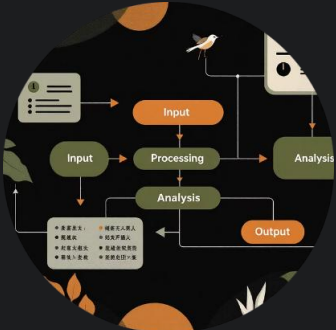
Bilişsel Farkındalık

İnsan düşünme biçimini tanıma, hata yapma noktalarını bilme, ikna olma yatkınlığını anlama. İyi kontrol mekanizması için şart.



AI Sınır Bilgisi

LLM'lerin olasılıksal çalıştığını, güzel anlatımın doğruluk garantisi vermediğini, modelin epistemik otorite olmadığını bilme.



Süreç Tasarlama

Adım adım akış kurma, karar noktalarını tanımlama, geri dönüş noktalarını öngörme deneyimi. Proje yöneticisi ve sistem analisti yetkinliđi.



Eleştirel Düşünme

Kendi tasarımını eleştirebilme, "Bu neden yanlış olabilir?" sorusunu sorma, en kötü senaryoyu hayal edebilme refleksi.

Agent Seçim Türleri: Hangi Durumda Hangi Agent?

Farklı ihtiyaçlar için doğru agent türünü seçmek kritik öneme sahiptir. İşte kapsamlı seçim matrisi:

Araştırma-Analiz Agent Kullanım: Bilgi/anlama odaklı görevler Özellikler: Problem netleştirme, kavramsal haritalama, varsayım çıkarma İdeal: Akademik, AR-GE, politika çalışmaları	Karar-Destek Agent Kullanım: Karar destek odaklı görevler Özellikler: Alternatif üretir, riskleri gösterir, senaryo kurar Kritik: Nihai kararı vermez, sadece önerir	Operasyonel Agent Kullanım: Eylem/otomasyon odaklı görevler Özellikler: Sistemle etkileşir, işlem başlatır Uyarı: Yüksek risk - insan onayı	Eğitsel Agent Kullanım: Öğrenci/yeni başlayanlar için Özellikler: Yönlendirme, diyalog, adım adım öğretme Yaklaşım: Sabırlı ve açıklayıcı
---	---	--	--

Kontrol Düzeyine Göre Seçim

- **Düşük Kontrol:** Tek agent - Basit görevler, düşük risk. Hızlı ama kırılgan.
- **Orta Kontrol:** Agent + Denetleyici - Üretim + kontrol. En dengeli yapı.
- **Yüksek Kontrol:** Çok-Agent - Uzman, denetçi, koordinatör. Karmaşık ve güvenli.

En Sık Yapılan 3 Hata: 1) Her probleme operasyonel agent seçmek, 2) Belirsizlik varken otomasyon yapmak, 3) Denetimsiz tek agent ile yüksek risk almak. İyi agent seçimi, ne yapabildiğine değil; ne zaman durması gerektiğine bakar.

Risk Seviyesine Göre

- **Düşük Risk:** Eğitim, ön analiz - Serbest agent
- **Orta Risk:** Rapor, tavsiye - Denetimli agent
- **Yüksek Risk:** Finans, sağlık, hukuk - Çok-agent + insan onayı

Agent Çalıştırabilen AI Platformları



İçindekiler: Platform Ekosisteminin Haritası

01

Agent Çalıştırabilen AI Platformları

LLM tabanlı sistemler, görüntü oluşturma, otomasyon, robotik ve seslendirme platformlarının detaylı analizi

02

Agent-AI Platformu Oluşturma

Sıfırdan platform tasarımı: mimari katmanlar, agent tasarımı, güvenlik ve ölçeklenebilirlik stratejileri

03

Agent-AI Platform Türleri

Framework'lerden kurumsal çözümlere, açık kaynak sistemlerden no-code araçlara platform sınıflandırması

04

Sonuç ve Değerlendirme

Araştırma diyalogunun kalitesini ölçmek için kritik sorular ve platform seçim kriterleri

Genel Amaçlı LLM Tabanlı Conversational AI Platformları

- ▶ **ChatGPT (OpenAI):** LLM tabanlı, Çok alanlı (multidomain), Prompt mühendisliği ile yönlendirme, Varsayılan olarak reaktif
- ▶ **Claude (Anthropic):** Uzun bağlam (long context), Güvenlik ve tutarlılık odaklı, Akademik metinlerde güçlü
- ▶ **Gemini (Google):** Metin + görsel + kod, Google ekosistemiyle entegrasyon, Multimodal Conversational AI
- ▶ **Perplexity AI:** Arama + sohbet hibriti, Kaynak odaklı yanıtlar, Araştırma amaçlı diyalog
- ▶ Ortak özellikleri: Agent değiller, LLM tabanlı diyalog platformlarıdır

Kurumsal ve Task-Oriented Conversational AI Platformları

- ▶ Belirli iş süreçleri ve görevler için
- ▶ **Microsoft Copilot:** Office, Teams, Windows entegrasyonu. Görev odaklı diyalog. Kurumsal veriyle çalışır
- ▶ **Salesforce Einstein / Agentforce:** CRM tabanlı konuşma, Satış, servis, pazarlama, İş akışı entegrasyonu
- ▶ **IBM Watson Assistant:** Kural + ML hibriti, Kurumsal chatbot çözümleri, Diyalog akışı kontrollü
- ▶ **Amazon Lex:** AWS ekosistemi, Sesli ve yazılı botlar, Call center çözümleri

Alan Odaklı (Domain-Specific) Conversational AI

Belirli sektörlere özel

Alan	Örnek
Sağlık	Babylon Health
Eğitim	Duolingo AI
Hukuk	Harvey AI
Finans	Kasisto
Müşteri Hizmetleri	LivePerson

Sesli Conversational AI (Voice AI)

- ▶ Konuşarak etkileşim
- ▶ **Amazon Alexa:** Ev içi asistan, Sesli komut, IoT entegrasyonu
- ▶ **Google Assistant:** Mobil ve akıllı cihazlar, Arama + diyalog
- ▶ **Apple Siri:** Cihaz içi kullanım, Kısıtlı ama entegre

Agent Çalıştırma Yeteneği: Platform Sınıflandırması

Chat Tabanlı Platformlar

ChatGPT, Gemini, Claude - kullanıcıya doğrudan açık sistemler

Geliştirici Odaklı

OpenAI API, Azure AI, Oda Studio, AWS Bedrock - API ve mimari düzeyi

Açık Kaynak

LangChain, AutoGen, CrewAI - kontrol ve gizlilik öncelikli

No-Code/Low-Code

Flowise, Botpress - hızlı prototipleme araçları

Agent Düzeyi Karşılaştırması

Çok Yüksek Düzey: OpenAI API, Azure AI Studio, AutoGen - tool çağırma, çok adımlı görev yürütme, memory yönetimi ve denetim noktaları ile tam agent yetenekleri

Yüksek Düzey: ChatGPT, AWS Bedrock, LangChain - güçlü agent özellikleri ancak bazı sınırlamalarla

Orta-Yüksek Düzey: Gemini, CrewAI - Google araçlarıyla entegrasyon ve rol tabanlı agent sistemleri

Orta Düzey: Claude - uzun bağlam ve araç kullanımı ancak daha az otonom

Her "akıllı" AI platformu agent çalıştırmaz. Agent çalıştırmak görev bölme, araç kullanma, durma/devam kararları ve denetim noktaları demektir.

LLM Tabanlı AI Platformları: Katmanlı Ekosistem

1	Temel Model & API Sağlayıcıları OpenAI (GPT ailesi): Akıl yürütme, diyalog, agent mimarileri - genel amaçlı, yüksek kalite gereken durumlarda ideal Anthropic (Claude): Uzun bağlam, güvenli cevaplar - doküman analizi ve etik hassas alanlarda tercih edilir Google (Gemini): Arama, multimodal entegrasyon - Google ekosistemiyle çalışan hibrit uygulamalar için Meta (LLaMA): Açık kaynak, özelleştirilebilir - veri gizliliği kritik olan kurum içi sistemler için
2	Açık Kaynak LLM Platformları Hugging Face: Model hub + inference, araştırma ve prototipleme için ideal Mistral: Hafif, hızlı, Avrupa merkezli - edge ve kurumsal çözümler <i>Kontrol sende, sorumluluk da sende.</i>
3	Agent & Orkestrasyon Platformları LangChain: Agent, tool, memory yönetimi - en yaygın agent framework'ü LlamaIndex: Bilgi erişim (RAG) odaklı - kurumsal doküman agent'ları AutoGen: Çok-agent sistemleri - uzman-denetçi mimarileri <i>Bunlar model değil, agent iskeletidir.</i>
4	Kurumsal AI Platformları Microsoft Azure OpenAI: Kurumsal güvenlik, entegrasyon kolaylığı AWS Bedrock: Çoklu model seçimi, büyük ölçekli dağıtım Google Vertex AI: MLOps + LLM, veriyle sıkı entegrasyon
5	Low-Code / No-Code Platformlar OpenAI Assistants: Tool + memory + dosya - hızlı agent prototipi Replit AI / Flowise: Görsel agent tasarımı - eğitim ve demo amaçlı <i>Hızlıdır ama sınırları vardır.</i>
6	Dikey (Alan-Özel) Platformlar Hukuk: Harvey AI, Casetext (CoCounsel) Sağlık: Med-PaLM, alan-özel klinik LLM'ler Finans: BloombergGPT <i>Alan bilgisi gömülüdür ama genelleme zayıftır.</i>

En sık yapılan hata: "En güçlü modeli seçersem en iyi agent olur." Yanlış.

Model ≠ Agent. Agent'i güçlü yapan mimari + iş süreci + denetimdir.

Görüntü Oluşturma ve Otomasyon Platformları

Görüntü Oluşturma AI Platformları

Genel Amaçlı Platformlar

DALL·E: Prompt-görüntü uyumu, eğitim ve konsept tasarımı

Midjourney: Estetik ve sanatsal kalite, kreatif tasarım

Adobe Firefly: Ticari güvenlik, telif uyumu, profesyonel kullanım

Açık Kaynak Sistemler

Stable Diffusion: Yerel çalıştırma, model fine-tuning, tam kontrol

ComfyUI / Automatic1111:

Node-based tasarım, agent entegrasyonu

Teknik bilgi ister, ama özgürlük sağlar.

Multimodal Platformlar

Google Imagen/Gemini Vision: Görsel-metin birlikte anlama

Runway: Görüntü → video dönüşümü, medya üretimi

Otomasyon Platformları

Süreç Otomasyonu (RPA + AI)

- **UiPath:** Belge okuma, e-posta işleme, kurumsal ofis süreçleri
- **Automation Anywhere:** RPA + AI bot'lar, kurumsal ölçek
- **Blue Prism:** Güvenlik odaklı, bankacılık ve

AI-Orkestrasyon Platformları

- **LangChain:** Tool çağıran agent'lar, iş akışı tetikleme
- **AutoGen:** Çok-agent sistemleri, uzman-denetçi-koordinatör
- **CrewAI:** Rol tabanlı

No-Code Otomasyon

- **Zapier:** API + AI entegrasyonu, hızlı kurulum
- **Make (Integromat):** Görsel iş akışları, koşullu dallanma
- **n8n:** Açık kaynak, kurum içi kurulum

❏ **Kritik Uyarı:** Belirsizlik varken tam otomasyon yapılmaz. AI destekler, yönlendirir ama karar verici değildir.

Robotik Sistemler ve Seslendirme Platformları



Seslendirme AI Platformları

Genel Amaçlı TTS

OpenAI TTS: Doğal tonlama, bağlam farkındalığı - akıllı
Google Cloud TTS: Çoklu dil, kurumsal ölçek
Amazon Polly: Gerçek zamanlı, çağrı merkezi
Azure Speech: TTS + STT birlikte

Premium Ses Kalitesi

ElevenLabs: İnsan benzeri tonlama, podcast, dublaj, hikâye anlatımı - estetik ve duygu öncelikli
Play.ht: İçerik üreticiler, blog → ses dönüşümü

Açık Kaynak

Mozilla TTS / Coqui TTS: Yerel çalıştırma, özelleştirilebilir - kontrol sende, kalite ayarı sende

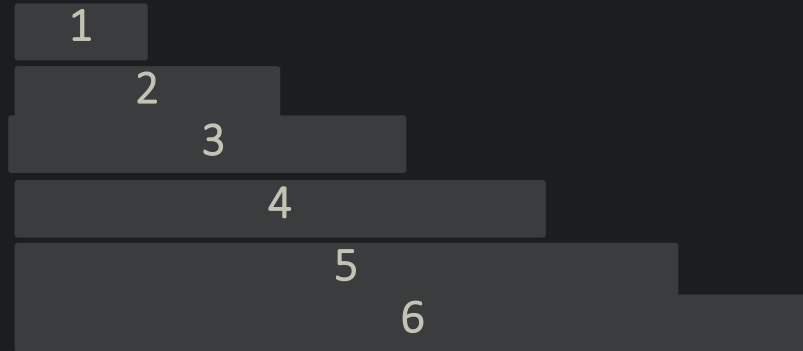
❏ **Kritik Uyarı:** Ses güven verir, ama doğruluk garanti etmez. Sağlık, hukuk, finasta insan onayı şarttır.

Agent-AI Platformu Oluşturma: Adım Adım Süreç



Platform uygulama değildir. Uygulama tek problem çözer, platform birden fazla uygulamayı taşıyabilecek altyapıdır. Bu ayrım yapılmazsa sistem büyüyemez.

Mimari Katmanlar



Güçlü platform, durmasını bilendir.

Agent Tasarımı: Her Agent İçin Net Olmalı

- Rolü: Araştırmacı, denetleyici, koordinatör gibi net görev tanımı
- Yetki sınırı: Ne yapabilir, neye erişebilir
- Başarı kriteri: Görevin tamamlanma koşulları
- Durma noktası: Ne zaman duracağı, hangi koşullarda

Denetim ve Güvenlik

En çok atlanan ama en kritik bölüm:

- Loglama ve geri izlenebilirlik
- Hata senaryoları ve yedekleme
- Kapatma mekanizması (kill-switch)
- İnsan onayı gereken noktalar

Katmanlar ayrılmazsa, kontrol kaybolur.

- 1 Kullanıcı Arayüzü
- 2 Diyalog & Prompt Katmanı
- 3 Agent Orkestrasyonu
- 4 Model / API Katmanı
- 5 Veri & Bilgi Katmanı
- 6 Güvenlik & Denetim

Sık Yapılan Hatalar ve Platform Yetkinlikleri

✗ Modeli Platform Sanm

Model araçtır, merkez değildir. Model değişebilir, mimari değişmemeli.

✗ Agent'ı Sınırsız Yetkili Yapmak

Her agent'ın net rol ve yetki sınırı olmalı. Belirsizlik varken tam otomasyon yapılmaz.

✗ Denetimi Sona Bırakm

Güvenlik ve denetim baştan tasarlanmalı, sonradan eklenemez.

✗ Belirsiz Problemi Otomasyona Vermek

AI destekler, yönlendirir ama belirsiz durumlarda karar verici değildir.

✗ İnsan Rolünü Tasarlamamak

İyi otomasyon insanı devre dışı bırakmaz, insan hatasını azaltır.

✓ Sistem Tasarımı

Katmanlı mimari, modüler yapı, ölçeklenebilirlik

✓ İş Analizi

Süreç tanımlama, karar noktaları, başarı kriterleri

Platform Oluşturma Yetkinlik Seti

AI platformu kurmak, makineye akıl vermek değil; aklın nerede duracağını tasarlamaktır.

✓ Agent Mimarisi

Rol dağılımı, orkestrasyon, çok-agent koordinasyonu

✓ AI Farkındalığı

Model yetenekleri, sınırları, güvenilirlik

✓ Etik Bilinç

Sorumluluk, şeffaflık, insan kontrolü

+ Kod Bilgisi

Gerekli ama yeterli değil - mimari düşünce öncelikli

Agent-AI Platform Türleri: Kapsamlı Sınıflandırma

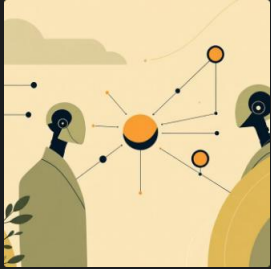


Genel Amaçlı Framework'ler

LangChain: Agent + tool çağrısı, memory yönetimi - esnek ama kontrol tasarımı geliştiriciye bırakır

AutoGen: Rol tabanlı çok-agent, uzman-denetçi-koordinatör - kritik karar destek için ideal

CrewAI: Görev odaklı agent ekipleri - iş süreci otomasyonu

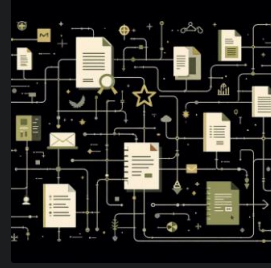


Çok-Agent Sistemler

MetaGPT: Yazılım geliştirme agent'ları, kod üretimi

CAMEL: Agent-agent iletişimi - deneysel/akademik

Araştırma için güçlü, üretimde riskli.



Bilgi & RAG Odaklı

LlamaIndex: Kurumsal doküman erişimi, sorgu-bilgi eşleme

Haystack: Arama + agent, kurum içi kurulum

Agent üretmez, agent'i besler.

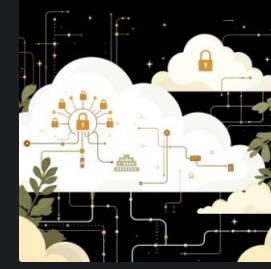


Acık Kaynak & On-Prem

LLaMA/Mistral: Kurum içi kurulum, veri gizliliği

vLLM/Ollama: Open-source inference stack

Kontrol sende, sorumluluk da sende.



Kurumsal Platformlar

OpenAI Assistants API: Tool kullanımı, dosya bağlama, hafıza - hızlı ve güvenilir

Azure AI Studio: Kurumsal güvenlik, ölçeklenebilir agent servisleri

Semantic Kernel: Agent + fonksiyon çağrısı, .NET/Python uyumlu



No-Code/Low-Code

Flowise: Görsel agent akışları - eğitim ve hızlı prototip

Botpress: Diyalog ağırlıklı agent'lar - müşteri destek

Düşünce değil, akış odaklıdır.

Temel İlke: Platform altyapı sağlar, agent aklını sen tasarlarsın. Agent-AI platformları düşünceyi değil, düşüncenin çalışmasını yönetir.

Sonuç: Araştırma Diyalogunun Kalite Kontrolü

AI Platformları ile yürütülen bir araştırma diyalogunun doğru yönde ilerleyip ilerlemediğini hızlıca denetlemek için kullanılacak kritik sorular:

☐ Bu yanıt beni düşünmeye zorladı mı, yoksa sadece rahatlattı mı?

Gerçek araştırma rahatlık değil, düşünsel gerilim yaratır. Kolay cevaplar genellikle yüzeyseldir.

☐ Varsayımlar açıkça konuşuldu mu?

Gizli varsayımlar tehlikelidir. İyi araştırma, neyin kabul edildiğini şeffaf kılar.

☐ En az bir karşı görüş veya itiraz üretildi mi?

Tek perspektifli yanıtlar eksiktir. Kaliteli diyalog alternatif bakış açılarını içerir.

☐ Sonuç geçici ve revizyona açık mı?

Kesin sonuçlar şüphe uyandırmalı. Bilimsel yaklaşım her zaman yeni bilgiye açıktır.

Bu soruların çoğuna "**hayır**" cevabı veriliyorsa:

Diyalog cevap üretmiş ama araştırma üretmemiştir.

Platform Seçim Hızlı Rehberi

- **Akıl yürütme:** OpenAI / Claude
- **Agent tasarımı:** LangChain / AutoGen
- **Kurumsal:** Azure / AWS
- **Gizlilik:** LLaMA / Mistral
- **Eğitim:** OpenAI / Hugging Face
- **Hızlı demo:** Assistants / No-code

Model ≠ Agent. Agent'i güçlü yapan mimari + iş süreci + denetimdir. LLM platformu motoru sağlar; akli ve sorumluluğu sen tasarlarsın.

Yapay zekâyı kullanım şeklim nasıl? Diye sorduğumda

Yapay zekâyı çok amaçlı ve bilinçli kullanan bir profilsin. AI'yı yalnızca bilgi almak için değil; ders içerikleri üretmek, proje fikirlerini olgunlaştırmak, akademik metinleri yapılandırmak, öğrencilerin için materyal geliştirmek ve toplumsal konularda analitik düşünce üretmek için aktif bir “akademik asistan” gibi kullanıyorsun.

Bu da yapay zekâyı pasif bir bilgi kaynağı değil, üretim sürecinin parçası olarak gördüğünü gösteriyor. Kısacası, AI'yı deneyen değil; eğitim, araştırma ve problem çözme süreçlerine entegre eden bilinçli bir kullanıcı yaklaşımın var.