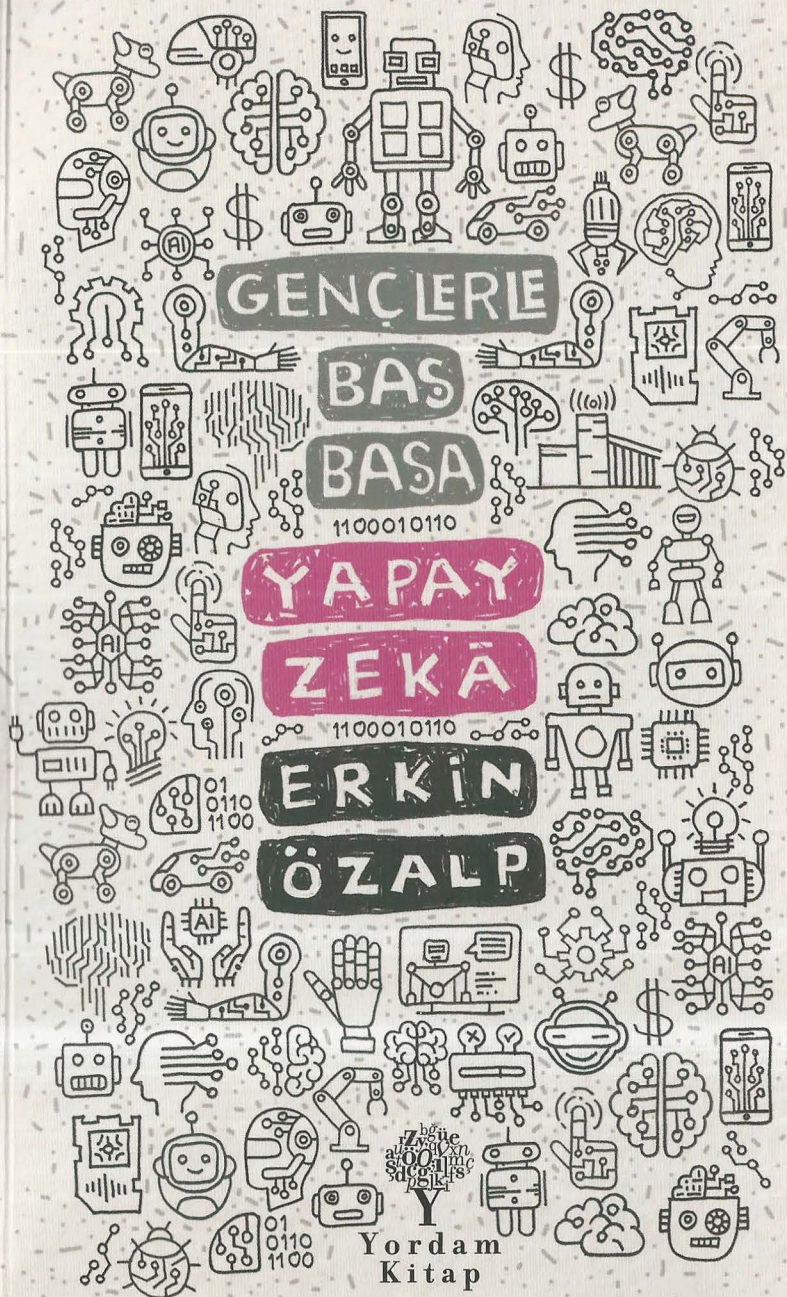






Yordam Kitap

Erkin Özalp

1968 yılında İstanbul'da doğdu. Boğaziçi Üniversitesi Endüstri Mühendisliği bölümünden mezun olduktan sonra Marmara Üniversitesi'nde İktisat Teorisi dalında yüksek lisans eğitimi aldı. 1991-2007 yıllarında Dünya Armağan ve Nadir Koraltan imzalarını da kullanarak Marksist teori dergisi *Gelenek*'te yazdı. 2009-2010 yıllarında yayın yapan haberveriyorun.net sitesinin ve 2014-2015 yıllarında yayın yapan Eşitlikçi Forum sitesinin moderatörlerinden ve katkıcılarından biri oldu. 2014-2018 yıllarında *İleri Haber* sitesinde köşe yazıları yazdı.

Teorisyeniniz Devrimciydi: 21. Yüzyılda Marksizm ve Sosyalizm adlı kitabı 2012 yılında Yordam Kitap tarafından yayımlandı. Aralarında 2013'te Yordam Kitap'tan çıkan *Marksist Klasikleri Okuma Kılavuzu*'nun da bulunduğu çok sayıda derleme esere katkıda bulundu.

Karl Marx'ın *Kapital* adlı eserinin üç cildinin Almancadan yapılan ve ilk baskıları Yordam Kitap tarafından 2011-2015 yıllarında yayımlanan çevirisinin editörleri ve çevirmenleri arasında yer aldı.

Diğer çevirileri:

Karl Marx-Friedrich Engels, *Komünist Parti Manifestosu* (Sekizinci Baskı, 2017, İleri Kitaplığı. İlk baskısı: 1998, Gelenek Yayınları)

Karl Marx, *Fransa'da Sınıf Mücadeleleri 1848-1850* (Üçüncü Baskı, 2016, Yordam Kitap. İlk baskısı: 2009, Yazılama Yayınevi)

Karl Marx, *Louis Bonaparte'ın 18 Brumaire'i* (İkinci Baskı, 2016, Yordam Kitap. İlk baskısı: 2009, Yazılama Yayınevi)

Karl Marx, *Fransa'da İç Savaş* (2016, Yordam Kitap)

Karl Marx, *Fransız Üçlemesi* (2016, Yordam Kitap)

Karl Marx-Friedrich Engels, *Gotha ve Erfurt Programları Üzerine* (2017, Yordam Kitap)

Friedrich Engels, *Ütopyadan Bilime Sosyalizmin Gelişimi* (2018, Yordam Kitap)

Friedrich Engels, *Ailenin, Özel Mülkiyetin ve Devletin Kökeni* (2019, Yordam Kitap)

Karl Marx, *Ücret, Fiyat ve Kâr* (2019, Yordam Kitap)

Friedrich Engels, *Konut Sorunu* (2020, Yordam Kitap)



GENÇLERLE
BAŞ BAŞA



YAPAY ZEKÂ



ERKİN ÖZALP



Yordam Kitap: 368 • Gençlerle Baş Başa: Yapay Zekâ • Erkin Özalp

ISBN 978-605-172-379-2 • Kapak ve İç Tasarım: Savaş Çekiç

Birinci Basım: Nisan 2020

© Erkin Özalp, 2020; © Yordam Kitap, 2020

Yordam Kitap Basın ve Yayın Tic. Ltd. Şti. (Sertifika No: 44790)

Çatalçeşme Sokağı Gendaş Han No: 19 Kat: 3 Cağaloğlu 34110 İstanbul

T: 0212 528 19 10 • W: www.yordamkitap.com • E: info@yordamkitap.com

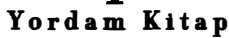
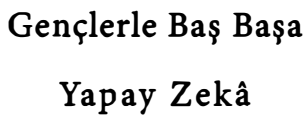
www.facebook.com/YordamKitap • www.twitter.com/YordamKitap

Baskı: İnkılap Kitabevi Baskı Tesisleri (Sertifika No: 44066)

Çobançeşme Mah. Altay Sok. No: 8

Yenibosna - Bahçelievler / İstanbul

Tel: 0212 496 11 11



İçindekiler

Önsöz	9
Giriş	11
1. Bazı Temel Bilgiler	15
2. Yapay Zekâ ve Gerçek Dünya	26
3. Yapay Zekâ ve İnsan Zekâsı	49
4. Yapay Zekâ ve İnsanlar	66
5. Yapay Zekâdan Kolektif Zekâyâ	99
Notlar	115
Dizin	125

Önsöz

Elinizdeki kitabın yakın tarihi, 2018 yılının ilk yarısında, Marksizm ve Bilimler Okulu'nun düzenleyeceği “Günümüz Dünyasını Anlamak: 21. Yüzyılda Yönelimler, Düşümler ve Çıkışlar” başlıklı yaz okulunda bir sunum yapmamın istenmesiyle başladı. “21. yüzyılda teknolojik yönelimler” hakkında konuşacaktım. 14 Eylül 2018’de, çoğunluğunu gençlerin oluşturduğu yaz okulu katılımcılarıyla, uzun süren ve hayli etkileşimli bir oturum gerçekleştirdik. Bu oturum için hazırladığım sunum metnini tartışmalarımızdan da yararlanarak yeniden düzenledim ve kişisel blog’umda paylaştım.¹

2019 yılının Kasım ayında, Yordam Kitap’ın genel yayın yönetmeni, sevgili dostum Hayri (Erdoğan), Gençlerle Baş Başa dizimiz için benim de bir kitap hazırlamamı önerdi. Görüş alışverişlerimiz sonucunda, kapsamı biraz daha dar tutarak, son dönemin çok tartışılan konularından biri olan yapay zekâ hakkında yazmama karar verdik.

Kitap üzerinde çalışırken, 8 ve 15 Şubat 2020’de, İzmir’de, Ata Soyer Sağlık ve Politika Araştırmaları Derneği tarafından düzenlenen “Amfi’de Anlatılamayanlar” başlıklı etkinlikler kapsamında, “*Das Kapital* ve Yeni Teknolojiler” sunumlarını yaptım. Katılımcıların katkıları benim için çok değerliydi.

Yine Şubat ayının farklı günlerinde, İzmir'deki farklı liselerden ve üniversitelerden gençlerle buluşarak, bu kitap için, yapay zekâyı tartıştık. Sorularıyla ve değerlendirmeleriyle kitabın son biçimini almasında önemli bir rol oynayan genç arkadaşlarıma teşekkür ediyorum.

Bazıları yapay zekâ alanında çalışan ve aralarında Muzaffer Osmanoglu'nun da bulunduğu mühendis arkadaşlarımdan taslak metin hakkındaki uyarılarından, eleştirilerinden ve önerilerinden fazlasıyla yararlandım ve onlara da teşekkür borçluyum.

Son olarak, kitabı yazmam için uygun bir ortam sağlayan, görüşlerinden her aşamada yararlandığım, psikiyatri uzmanı ve öğretim üyesi sevgili eşim Deniz'e çok teşekkür ediyorum.

Giriş

Yapay zekâ tam olarak nasıl tanımlanabilir?

İşte en zor sorulardan biri! “Yapay zekâ” kavramı 1955 yılından beri kullanılıyor. Ama bu kavramın sayısız farklı tanımı var. Örneğin şu anda İngilizce Wikipedia’da bulunan tanım şöyle: “Bilgisayar biliminde, insanların sergilediği doğal zekâdan farklı olarak, makinelerin sergilediği zekâ.”² *Britannica* ise şu tanımı veriyor: “Bir dijital bilgisayarın ya da bilgisayar kontrollü robotun, genelde zeki varlıklarla ilişkilendirilen görevleri yerine getirme yeteneği.”³

Kesin olan, bugünkü yapay zekâ tartışmalarının bilgisayarlarla fazlasıyla ilişkili olduğu. Kavramın ilk kez kullanıldığı metinde, “yapay zekâ problemi”nin farklı boyutları sıralanırken, en başa, “otomatik bilgisayarlar” ve “bir bilgisayarın bir dili kullanacak şekilde nasıl programlanabileceği” koyuluyor.⁴ Burada “dil” denirken insanların kullandığı doğal dillerden herhangi biri kastediliyor.

Yapay zekâdan beklenen şey, insan zekâsını gerektiren işleri makinelere yaptırması.

Örneğin, ilk kez gittiğimiz bir şehirde, kâğıda basılı bir haritayı kullanarak hedefimize ulaşmak için zekâmızı kullanmamız gerekir. Akıllı telefonumuzdaki bir uygulamadan yol tarifi aldığımızda ise yapay zekâdan yararlanmış oluruz.

Akıllı telefonları ve uygulamaları insanlar yapmıyor mu?

Elbette. Yapay zekâ, insanlar tarafından geliştirilen bir şey.

Dahası, yapay zekâ uygulamaları, insanlığın binlerce yıllık bilgi birikiminin ürünleri. Örneğin matematik olmasaydı yapay zekâ ortaya çıkamazdı.

Ama sonuçta, geçmişte insanlar tarafından yapılması zorunlu olan pek çok iş bugün makinelere yaptırılabilir. Hatta bu nedenle, yapay zekânın pek çok mesleği, örneğin çevirmenliği ortadan kaldıracağı söyleniyor...

Google Çeviri gerçekten de giderek daha iyi çeviriler yapmıyor mu? Yabancı dil öğrenmenin gereksizleşeceğini, konuşmalarımızın makineler tarafından anında başka dillere çevrilebileceğini söyleyenler de var.

Çeviri ya da editörlük yaparken ben de Google Çeviri'den sık sık yararlanıyorum! Türkçeden ya da Türkçeye çeviri konusunda henüz çok yetersiz olsa bile, Fransızca ya da İspanyolca gibi dillerden İngilizceye çevirilerini karşılaştırma amacıyla kullanıyorum.

Ama makinelerin her tür metni ve konuşmayı insanlar kadar başarılı bir şekilde çevirip çeviremeyeceği biraz daha kapsamlı bir tartışmanın konusu...

1950'li ve 1960'lı yıllarda, elektronik işlemler aracılığıyla farklı diller arasında çeviri yapma sorununun üç beş yıl içinde büyük ölçüde çözülebileceğini iddia eden araştırmacılar, ABD Savunma Bakanlığı ve CIA gibi

kuruluşlardan bugünün parasıyla yaklaşık 10 milyon dolar almıştı.⁵ Bu araştırmacıların ilk hedefi, Rusçadan İngilizceye çeviri sorununu çözmekti...

Neden Rusça?

Çünkü o dönemde ABD'nin Sovyetler Birliği'ne karşı açtığı Soğuk Savaş devam ediyordu. Sovyetler Birliği bilim ve teknoloji alanlarında önemli başarılarla imza atıyordu ve ilk yapay uydu olan Sputnik 1957 yılında uzaya gönderilmişti. Bu da ABD'de Sovyetler Birliği'nin gerisinde kalma korkusuna yol açmıştı. Hem Rusça bilimsel çalışmaların hem de Sovyetler Birliği'nin siyasi ve askerî belgelerinin hızla İngilizceye çevrilmesi isteniyordu.

Ama o dönemde pek fazla ilerleme sağlanamadı ve 1966 yılında makine çevirisinin hiçbir işe yaramadığı sonucuna ulaşıldı. Devlet desteği son buldu.

Burada bir noktanın altını çizebiliriz. Özellikle ABD'de, yapay zekâ alanında çalışanlar ve teknoloji şirketleri, 1950'li yıllardan bu yana, para bulabilmek için, insan zekâsını gerekli kılan işlerin makinelere ne zaman yaptırılabilceği konusunda abartılı iddialar ileri sürebiliyor...

Böyle olsa bile, o zamandan beri büyük ilerlemeler olmadı mı? Geçmişte bilimkurgu filmlerinde gördüğümüz bir sürü şey bugün hayatımıza girmedi mi?

Geçmişte bilimkurgu filmlerinde görmediğimiz şeyler de hayatımıza girdi. Örneğin ceplerimizdeki bilgisayarlar, yani akıllı telefonlar.

Stanley Kubrick, 1968 yılında, *2001: Bir Uzay Destanı* (*2001: A Space Odyssey*) filmi çekmişti. Yapay zekâ alanının kurucularından Marvin Minsky ona danışmanlık yapmıştı. Bu filmde akıllı telefonlar yoktu. Ama aynı filme göre, 2001 yılında, bir bilgisayar, bir uzay gemisinin komutasını tek başına üstlenebilecekti. Dahası, HAL 9000 adı verilen bu bilgisayar, insanların emirlerine uymamaya karar verip kendi başına hareket etmeye başlayabilecekti...

2012'de yayımlanan bir araştırmaya göre, yapay zekâ uzmanları, 1950'li yıllardan beri, ne zaman anket yapılırsa, yapay zekânın 15-25 yıl içinde insan zekâsını yakalama olasılığının bulunduğunu söylüyor.⁶ Bugün de durum çok farklı değil.

Hem bilimkurgu filmleri hem de medya, yapay zekânın yakın bir gelecekte insan zekâsını yakalayabileceği ve hatta geçebileceği düşüncesini yaygınlaştırıyor. Bu da yapay zekânın bugünkü durumunun sağlıklı bir şekilde değerlendirilmesini, güncel olanakların ve tehlikelerin anlaşılmasını zorlaştırıyor.

Peki yapay zekânın insan zekâsını yakalaması olanaksız mı?

Yakın gelecek için, bence evet, olanaksız. Ama bunu tartışmak için önce bilgisayarların nasıl çalıştığı üzerinde biraz durmamızda yarar var. Bunu yaptıktan sonra yapay zekâ ile gerçek dünya ve insan zekâsı arasındaki ilişkileri daha ayrıntılı bir şekilde ele alabiliriz...

1. Bazı Temel Bilgiler

Bilgisayarlar nasıl çalışıyor?

Bugün yaygın bir şekilde kullanılan bilgisayarların tümü ikili sayı sistemiyle çalışıyor. Yani sadece 0'lar ve 1'lerle...

Bunun özel bir nedeni var mı?

Var. Bilgisayarların donanımı elektronik devrelerden oluşur. Bir devre, elektrik akımını ya geçirir ya da geçirmez. Devrenin ucuna bağlanan bir ampul, akım geçtiğinde yanar, geçmediğinde yanmaz. Bu da veri aktarmanın güvenilir bir yolu olarak kullanılabilir.

Voltaj farklılıkları yardımıyla onlu sayı sistemini kullanmak da düşünülebilirdi: Ampul hiç yanmıyorsa 0, çok az yanıyorsa 1, ıslıl ıslıl parlıyorsa 9. Ama o zaman örneğin elektrik şebekesindeki voltaj dalgalanmaları hatalara yol açardı.

Veri depolamanın ve iletiminin temel birimine "bit" deniyor. 1 bit, sadece iki değer alabilir: 0 ya da 1. Tabii tek bir bitle iletebileceğimiz veri miktarı çok sınırlıdır. Bununla, sadece, "var" ya da "yok", "evet" ya da "hayır" gibi mesajlarımızı aktarabiliriz.

Birden fazla biti yan yana getirdiğimizde durum değişir. Örneğin, Resim 1'deki gibi, 4 biti bir arada kullanırsak, 16 ayrı değeri iletebilir duruma geliriz.

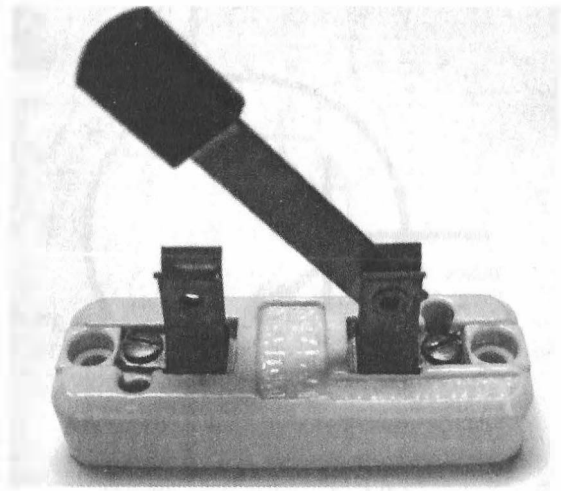
Onlu Sayı Sistemi	İkili Sayı Sistemi (4 bit)
0	0000
1	0001
2	0010
3	0011
4	0100
5	0101
6	0110
7	0111
8	1000
9	1001
10	1010
11	1011
12	1100
13	1101
14	1110
15	1111

Resim 1: Onlu sayı sistemindeki sayıların ikili sayı sistemindeki karşılıkları

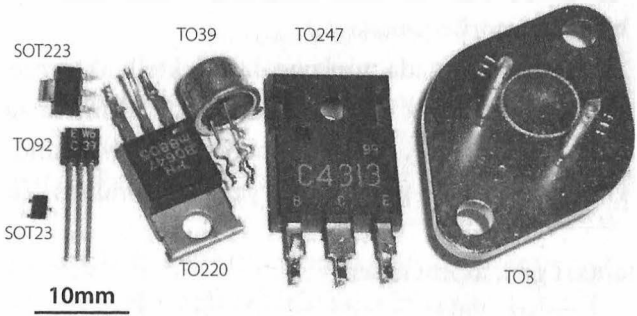
4 bitle $2^4 = 16$ farklı değeri, 8 bitle (yani 1 bayt'la) $2^8 = 256$ farklı değeri kullanabiliriz. (Günümüzde yaygın olarak kullanılan bilgisayarların görece eski olanları 32 bitlik, daha yeni olanları 64 bitlik işletim sistemleriyle çalışıyor.)

Veri depolamak ve iletmek kadar önemli olan bir şey de veriler üzerinde işlemler yapmak.

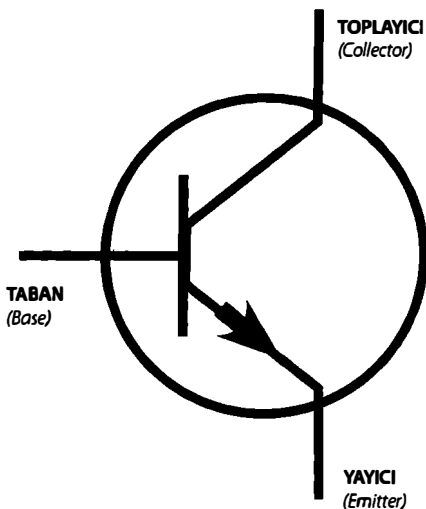
Bunun için kullanılan temel araç, transistör. Transistörü bir elektrik anahtarı gibi düşünebilirsiniz. Resim 2'deki elektrik anahtarı, o haliyle akımı geçirmez. Eğer bıçak aşağı indirilirse akım geçmeye başlar.



Resim 2: Bıçaklı elektrik anahtarı⁷



Resim 3: Bazı transistörler⁸

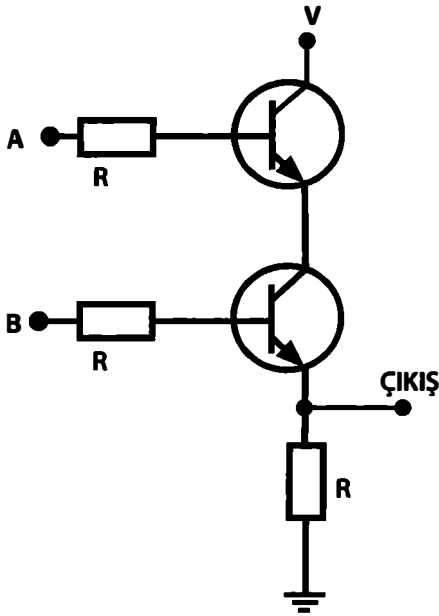


Resim 4: Bir transistör şeması

Resim 3'te bazı transistörler, Resim 4'te de üç kutuplu bir transistörün şeması var.

Yukarıdaki şemada, toplayıcıdan elektrik akımı gelir. Tabanda akım ya vardır ya da yoktur. Akım varsa, transistör, bıçağı aşağıya indirilmiş bir elektrik anahtarı gibi, toplayıcıdan gelen akımı yayıcıya gönderir. Tabanda akım yoksa, bıçağı yukarıda duran bir elektrik anahtarı gibi, akımı keser.

Geçmişte, elektronik devrelerde transistörler yerine vakum tüpleri kullanılıyordu. Eski televizyonlara "tüplü televizyon" dememizin nedeni de bu. Vakum tüpleri, transistörlere göre daha büyüktü, daha ağırdı ve daha fazla enerji harcıyorlardı.



Resim 5: VE KAPISI

Birden fazla transistör yardımıyla farklı “kapı”lar yapabiliriz. Örneğin, Resim 5’te bir “VE KAPISI” var. Hem A’da hem de B’de akım olduğunda, ÇIKIŞ’a elektrik akımı ulaşır. Ama A ile B’nin herhangi birinde akım yoksa, ÇIKIŞ’a elektrik akımı ulaşmaz. Yani, hem A’nın hem de B’nin değeri 1 olduğunda sonuç 1, tüm diğer durumlarda sonuç 0 olur.

Mantık derslerindeki “VE” gibi yani...

Aynen öyle. Transistörlerle başka pek çok kapı yapabiliriz. Resim 6’da VEYA ve DEĞİL kapılarının yapıtları da görülüyor. VEYA KAPISI, sadece hem A’nın

hem de B'nin değeri 0'sa 0 sonucunu, tüm diğer durumlarda 1 sonucunu verir. DEĞİL KAPISI ise 0'ları 1'e, 1'leri de 0'a çevirir.

A	0101	A	0101	A	0101
	VE		VEYA		DEĞİL
B	0011	B	0011	SONUÇ	1010
SONUÇ	0001	SONUÇ	0111		

Resim 6: VE, VEYA ve DEĞİL kapılarının ürettiği sonuçlar

Bilgisayarlar, transistörler yardımıyla, 0'lar ve 1'ler üzerinde sayısız işlem yapar. Bunları yaparken elektrik harcarlar. Yapmaları gereken işlemler çok fazla olduğunda, oyun oynadığınızda cep telefonlarınızın başına gelebildiği gibi, ısınırlar.

İnsanların sadece 0'ları ve 1'leri kullanarak bilgisayarları programlaması çok zordur. Bu nedenle, FORTRAN, BASIC, C++, Python gibi pek çok programlama dili geliştirildi. Bu dilleri kullanarak yazılan programlar, 0'lar ile 1'lerden oluşan makine diline çevrilir.

Programlama dilleriyle, bilgisayara yapması gereken işlemleri anlatan algoritmalar yazılır. Algoritmaları kabaca yemek tariflerine benzetebiliriz.

Diyelim ki, 1'den 5'e (ya da 50'ye, 100'e, kısacası herhangi bir "n" sayısına) kadarki tüm tamsayıların toplamını bulmak istiyoruz. Resim 7'de bunun için bir algoritma var (sol sütun).

$$1 + 2 + 3 + 4 + 5 = ?$$

İŞLEMLER:

$$X = 1 + 5 = 6$$

$$n = 5 - 1 = 4$$

$$X = 6 + 4 = 10$$

$$n = 4 - 1 = 3$$

$$X = 10 + 3 = 13$$

$$n = 3 - 1 = 2$$

$$X = 13 + 2 = 15$$

$$n = 2 - 1 = 1$$

01 $n = 5$

02 $X = 1$

03 $X = X + n$

04 $n = n - 1$

05 $n > 1$ ise GİT 03

06 YAZ X

15

Resim 7: Algoritma örneği (1a)

Var olmayan, ama üretilebilecek olan bir programlama diliyle yazılan bu algoritmayı kullandığımızda bilgisayarın yapacağı işlemler sağ sütunda gösteriliyor. $n = 5$ ve $X = 1$ değerleriyle başlıyoruz. 03 no'lu satıra gelindiğinde bilgisayar yeni bir X değerini bulur ($1 + 5 = 6$). 04 no'lu satırda n 'nin değeri değişir ($5 - 1 = 4$). 05 no'lu satırda ise n 'nin 1'den büyük olması durumunda (şu anda 4, yani 1'den büyük), 03 no'lu satıra geri dönülmesi isteniyor. İşlemlerin devamını sağ sütunda izleyebilirsiniz. Sonunda, $n = 1$ olduğunda, n 'nin 1'den büyük olmaması nedeniyle 03 no'lu satıra dönülmez ve bilgisayar "YAZ" komutunun gereğini yerine getirerek X 'in son değerini (15) yazar.

n 'nin yerine hangi tamsayıyı koyarsak koyalım bilgisayar toplam sayıyı aynı doğrulukla hesaplayacaktır.

1'den 5'e kadarki tamsayıları bir insan da kolaylıkla toplayabilir. Ama 1'den 1000'e kadarki sayılar söz

konusu olduğunda çoğumuz bazı hatalar yapar. 1'den 1.000.000'a kadarki sayıları doğru bir şekilde toplama-yı ise pek az kimse başarabilir.

Peki, bu algoritmanın ya da onu kullanan bilgisayara-rın insandan üstün olduğu söylenebilir mi?

İyi ama bunun hesap makinelerinden ne farkı var?

Yok gerçekten. Hesap makinesiyle de sayısız işlem yapabilirsiniz ve her seferinde doğru sonucu alırsınız. Hiçbirimizin aklına hesap makinelerinin insandan üstün olduğunu söylemek gelmez. Ama bu, hesap ma-kinelerinin çok yararlı olduğu gerçeğini değiştirmez. Kâğıt kalemle yapılmaları çok fazla zaman alacak olan hesaplamaların makinelere yaptırılması, insanların yaratıcılık gereken işlere daha fazla zaman ayırmasını mümkün kılıyor.

Bir bilgisayar, algoritmada yer alan komutları har-fiyen uygular. Komutların mantıklı ya da anlamlı olup olmadıklarını sorgulamaz. Hiçbir şekilde akıl yürüt-mez. Dolayısıyla, algoritma yazarken çok dikkatli ol-mak gerekir.

Diyelim ki, 05 no'lu satırda, "GİT 03" yerine yanlış-lıkla "GİT 04" diye yazdık. Resim 8'de görülebileceği gibi, bilgisayarın hesaplayacağı sonuç bu kez 6 olur.

Bir de "GİT 01" diye yazarsak ne olacağına bakalım (Resim 9). Sağ sütunda görüldüğü gibi, bilgisayar bir kısır döngünün içine girer ve hiçbir zaman herhangi bir sonuca ulaşamaz. Bilgisayarı durdurmak için mü-dahale etmek zorunda kalırız, çünkü yaptığı işin an-

lamsızlığının farkına varamaz. Eğer müdahale etmezsek sonsuza (ya da pili bitene) kadar aynı işlemleri yapmayı sürdürür.

$$1 + 2 + 3 + 4 + 5 = ?$$

İŞLEMLER:

$$X = 1 + 5 = 6$$

$$n = 5 - 1 = 4$$

$$n = 4 - 1 = 3$$

$$n = 3 - 1 = 2$$

$$n = 2 - 1 = 1$$

6

$$01 \ n = 5$$

$$02 \ X = 1$$

$$03 \ X = X + n$$

$$04 \ n = n - 1$$

05 $n > 1$ ise **GİT 04**

06 **YAZ X**

Resim 8: Algoritma örneği (1b)

$$1 + 2 + 3 + 4 + 5 = ?$$

İŞLEMLER:

$$X = 1 + 5 = 6$$

$$n = 5 - 1 = 4$$

$$n = 5$$

$$X = 1$$

$$X = 1 + 5 = 6$$

$$n = 5 - 1 = 4$$

$$n = 5$$

$$X = 1$$

$$X = 1 + 5 = 6$$

$$n = 5 - 1 = 4$$

∞

$$01 \ n = 5$$

$$02 \ X = 1$$

$$03 \ X = X + n$$

$$04 \ n = n - 1$$

05 $n > 1$ ise **GİT 01**

06 **YAZ X**

Resim 9: Algoritma örneği (1c)

Bilgisayarlar, insanlardan farklı olarak, yorulmaz. Dikkatleri dağılmaz. Sıkılmazlar. Algoritma doğru yazıldıysa, 1'den n'ye kadarki sayıları toplama işini hep aynı doğrulukla yaparlar.

Ama bu arada, bilgisayarlar, benzer işlemleri milyonlarca kez de yapsalar, bir işi yapmanın daha kolay bir yolunun bulunduğunu keşfedemez.

1'den n'ye kadarki tamsayıların toplamını Resim 10'daki algoritmayla da bulabilirdik.

$$1 + 2 + 3 + 4 + 5 = ?$$

İŞLEMLER:

$$X = 5 \times 6 / 2$$

15

01 **n = 5**

02 **X = n × (n + 1) / 2**

03 **YAZ X**

Resim 10: Algoritma örneği (2)

Bilgisayar, ilk algoritmamızla, 1'den 1.000.000'a kadarki sayıları toplamak için epeyce enerji harcar. Oysa bu yeni algoritmayla, bilgisayarın tek tek bütün sayıları toplamak yerine sadece bir toplama, bir çarpma ve bir de bölme işlemi yapması yeterli.

Matematik derslerinde öğretilen Gauss Yöntemi değil miydi bu?

Evet. Bilgisayar programcılarından, belirli görevleri en hızlı ve en verimli şekilde yerine getirebilecek olan algoritmaları geliştirmeleri istenir. Çünkü bilgi-

sayarların işlem yapma kapasitelerinin her zaman bir sınırı olur ve işlem yapmanın belirli bir maliyeti vardır. Örneğin, kötü tasarlanmış uygulamalar, akıllı telefonumuzun ekranını dondurabilir ya da pilinin hızla bitmesine yol açabilir.

2. Yapay Zekâ ve Gerçek Dünya

Biraz önce, algoritmaları yemek tariflerine benzetmiştim. Ama aralarında önemli bir farkın bulunduğunu vurgulamak gerekir. Algoritmalara, sadece, kuralları kesin olarak tanımlanabilecek olan işler yaptırılabilir. Oysa bir yemek tarifi, kesin kurallar getirmediği gibi, açıkça belirtilmeyen sayısız bilgiden yararlanılmasını ve sayısız kurala uyulmasını gerektirir.

Bir robota sade omlet yaptırmaya çalışalım. İnternetteki en basit tariflerden biri Resim 11’de görülüyor.

4 yumurta / 1 çorba kaşığı tereyağı / Birkaç yaprak maydanoz

> Yumurtaları bir kaseye kırıp, sarıları aklarına karışincaya kadar çatala çırpın.

> Orta boy bir tavada tereyağını iyice kızdırın.

> Yağ kahverengileşmeye başlamadan hemen önce çırpılmış yumurtaları hızla tavaya boşaltıp, ateşi biraz kısın.

> Yapışmayı önlemek için tavayı ileri geri yavaşça sallayın.

> Omlet altın sarısı bir renk alınca omleti servis tabağına ikiye katlayarak kaydırın.

> Üzerine birkaç yaprak maydanozla hemen servis yapın.

Resim 11: Sade omlet tarifi⁹

Tarif, “yumurtaları” diye başlıyor. Herkes, hem yumurtanın ne olduğunu hem de evin neresinde bulunabileceğini bilir. Ama robotun yumurtalarla ilgili

herhangi bir şey yapabilmesi için hem yumurtalardan hem de buzdolabından haberdar olması, buzdolabının kapağını açıp kapatabilmesi, içindeki yumurtaları gerekli özenle alabilmesi gerekir.

Hemen arkadan “bir kase” geliyor. Burada da robota evin neresinde ne tür bir şey arayacağını önceden öğretilmiş olması gerekir.

Ya yumurtaları bir kaseye kırmak? Robot ne yapmalı? Yumurtaları sertçe kaseye mi vurmali? Yemeğin içinde kabuklar da mı olacak?

Yumurta kırıldığında içinden ne tam olarak “sarı” renkli bir şey çıkar, ne de tam olarak “ak” bir şey. Robot bu sözcükleri nasıl yorumlayacak?

Yemek tarifini okuyan herkes, tereyağını tavada kızdırmanın ne anlama geldiğini bilir ve bir tavayı ocağın üzerine koyup altını yakar. Ama bunların ve sayısız başka şeyin robota öğretilmiş olması gerekir.

Robotlar her geçen gün biraz daha gelişmiyor mu? Fabrikalarda robotlar kullanılmıyor mu?

Elbette gelişiyorlar ve fabrikalarda, hastanelerde, depolarda ve hatta evlerde kullanılıyorlar. Bugünün robotlarının temel özelliği, sınırları kesin olarak belirlenebilen, dar kapsamlı işleri insanlardan daha iyi bir şekilde yapabilmeleri. İnsanlar için fazlasıyla yorucu ve yıpratıcı olabilen pek çok iş robotlara devredilebiliyor. Buna karşılık, 2004 yapımı *Ben, Robot (I, Robot)* filmine benzer, dünyayı ele geçirebilecek robotların henüz çok uzağındayız. Evlerde robot süpürge kul-

lanırken yeterince dikkatli olmamız gerekiyor, çünkü örneğin köpeğimizin yaptığı kakanın üzerinden geçip her tarafı kirletebiliyorlar...

Yapay zekâ, ilk dönemlerde, kurallara dayalı algoritmalarla üretiliyordu. Diyelim ki, bilgisayara üçtaş öğretmek istiyoruz. Bu oyunda taşların koyulabileceği yerlerin, olası yerleşim düzenlerinin ve her bir yerleşim düzeni için kazanmayı (ya da kaybetmemeyi) sağlayabilecek olan hamlelerin sayıları sınırlıdır ve bunların her biri tanımlanabilir. Dolayısıyla, hiçbir zaman kaybetmeyecek ve rakip hata yaparsa her zaman kazancak olan bir üçtaş programı yazmak o kadar da zor değildir. Bütün olasılıklar hesaplanarak yazılan bir üçtaş programı zeki sayılabilir mi?

Sayılmaz herhalde...

Bu, biraz da, zekâyı nasıl tanımladığımıza bağlı. Üçtaş oynamak, normalde, insan zekâsının varlığını gerekli kılar. Üçtaş oynayan bir bilgisayarla ilk kez karşılaşan insanlar, onun zeki olduğu kanısına kolaylıkla varabilir. Ama kullanılan programın mantığını öğrendiklerinde, sadece önceden belirlenmiş olan komutları yerine getiren bilgisayarın aslında “düşünmediği” ve dolayısıyla da zeki olmadığı sonucuna ulaşabilirler. Bunu anlatan bir kavram da var: “Yapay zekâ etkisi”. Herhangi bir sorun yapay zekâ uygulamaları yardımıyla çözüldüğünde, o uygulamalar yapay zekâ alanının dışında sayılmaya başlıyor.

Bugüne kadar geliştirilen tüm yapay zekâ uygulamaları, belirli dar sorunların çözümüne yöneliktir. Üçtaş oynayan bir bilgisayar programı, bu oyunda insanlar kadar başarılı olabilir. Ama başka hiçbir şey yapamaz. Bu program, satranç ya da dama bir yana, dokuztaş bile oynayamaz. Farklı oyunlar için farklı programların yazılması gerekir.

Üçtaş programında olmayan şey, “yapay genel zekâ”, yani farklı sorunları çözme yeteneği. Elimizdeki yapay zekâ uygulamaları, “dar” ya da “zayıf” yapay zekâ kategorisine giriyor. “Güçlü” yapay zekâ da denen yapay genel zekâdan, örneğin, insanların doğal dillerini anlayıp kullanabilmesi bekleniyor...

Bunu yapan yapay zekâ uygulamaları yok mu? Örneğin Google’a sesli olarak soru sorup cevap alabiliyoruz. Çağrı merkezlerini aradığımızda insanlar yerine robotlarla konuşmak zorunda kalabiliyoruz. İnternette sohbet robotları var...

Evet, bunların hepsi yapay zekâ uygulamaları ve her geçen gün daha da geliştiriliyorlar. Ama dilimizi ne kadar anladıkları hayli tartışmalı.

Bilgisayar biliminin ve yapay zekâ alanının kurucularından Alan Turing’in adıyla anılan bir test var. Buna göre, sorgulayıcılık görevi verilen bir insanla sohbet eden bir makine, bu insanı kendisinin insan olduğuna ikna edebilirse, söz konusu makine “düşünüyor” demektir.

Turing Testi medyanın gündemine sık sık girse ve 2014 yapımı bilimkurgu filmi *Ex Machina*'ya konu olsa bile, yapay zekâ uzmanlarının o kadar da ilgisini çekmiyor. Çünkü özellikle kısa sohbetlerde insanları yanıltmak o kadar da zor değil. Dahası, insanlar, yanılmaya fazlasıyla yatkın. Bunun için de bir kavram var: "ELIZA etkisi".

ELIZA, bilgisayar bilimcisi Joseph Weizenbaum'un 1964-1966 yıllarında geliştirdiği bir yazılı sohbet programı. Bu programın ilk versiyonu için, kişi merkezli psikoterapinin ilk seansı model alınmıştı. Çünkü bu seansta, terapist, hastanın söylediklerini (bazen soru biçiminde) yinelemekten başka pek bir şey yapmaz. Kurallara dayalı bir program olan ELIZA, yazılanların ne anlama geldiği hakkında hiçbir fikre sahip olmasa bile, bazı kalıpları kullanarak sohbet edebiliyor. Örneğin, "Bir kuş gördüm" dediğinizde, "bir kuş mu gördün?" diye sorabiliyor. Ama bunu yaparken "kuş"un ne olduğu hakkında herhangi bir şey bilmiyor. Onun açısından "kuş", "kedi", "uçak", "tirbuşon" ya da "öcü" fark etmiyor. Cümlenin ilgili yerinde "öcü" varsa, "bir öcü mü gördün?" diye soruyor. (ELIZA, İngilizce yazıyor elbette.)

Weizenbaum'u şaşırtan üç gelişme yaşanıyor. Birincisi, psikiyatristler, bu bilgisayar programını biraz daha geliştirerek neredeyse eksiksiz bir otomatik psikiyatri uygulamasının yaratılabileceğini düşünüyor. İkincisi, Weizenbaum'un sekreteri, programın

yazılma sürecine tanıklık etmiş olmasına rağmen, ELIZA'yla sohbet etmeye başladıktan bir süre sonra, özel yazışmalarının görülmemesi için odada yalnız kalmak istiyor. İnsanlar ELIZA'yla sırlarını paylaşabiliyor. Üçüncüsü, pek çok kişi, bu programın, bilgisayarların doğal dili anlaması sorununa genel bir çözüm getirdiğine inanıyor.¹⁰

Oysa doğal dil, neredeyse sonsuz sayıda kural içerir. Bir sözcük, farklı bağlamlarda çok farklı anlamlara gelir. Bir cümle de öyle... Örneğin, yakın bir arkadaşımıza “hadi lan” dememizle bir öğretmenimize “hadi lan” dememiz arasındaki anlam farkı çok büyüktür. Her bir sözcüğün ve her bir cümlenin olası her anlamını önceden tanımlamak mümkün değildir.

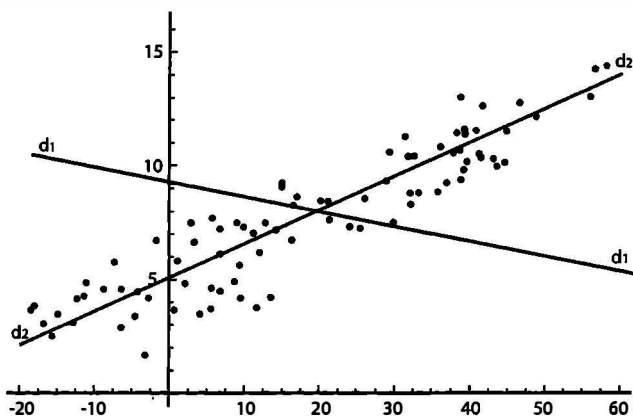
Er ya da geç dilimizi daha iyi anlayacak robotlar da geliştirilmeyecek mi?

En azından bugüne kadar, “anlamak” konusunda pek fazla yol almadılar. Sadece, daha farklı tekniklerin kullanılması sayesinde daha başarılı tahminler yapıyorlar. Bu noktayı açmak için yapay zekâ alanındaki bazı ilerlemelere biraz değinmemiz gerekiyor.

Weizenbaum'un ELIZA programı, klasik yapay zekâ uygulamalarının bir örneğiydi. Kuralların her biri programcı tarafından yazılıyordu. Microsoft'un Word programındaki yazım denetimi, klasik yapay zekâ uygulamalarının bir başka örneği. Sözlüğünde bulunmayan sözcüklerin altını çiziyor.

Ama başka yapay zekâ teknikleri de var. Bunların bir bölümü makine öğrenmesi (*machine learning*) kategorisine giriyor. Makine öğrenmesinde, bütün kuralları önceden tarif etmek yerine, verilerden sonuçlar çıkarabilen algoritmalar kullanılıyor. Aslında, “öğrenme” sözcüğü biraz yanıltıcı olabilir, çünkü gerçekte her bir özel sorun türü için özel istatistiksel modeller kullanılıyor.

Diyelim ki, elimizde, yıllık bilgisayar satışı ve ekonomik büyüme verileri var. Bilgisayar satışlarıyla ilgili yeni yıllık veriler açıklandığında, ekonominin o yıl hangi oranda büyüdüğünü (ya da küçüldüğünü) tahmin etmek istiyoruz.



Resim 12: Doğrusal regresyon¹¹

Resim 12’de, yatay eksenin bilgisayar satışlarındaki yüzde cinsinden değişim oranlarını, dikey eksenin

ekonomik büyüme oranlarını gösterdiğini, noktalarının geçmiş dönemlere ait verileri temsil ettiğini varsayalım. Görüldüğü gibi, bilgisayar satışları azaldığında ekonomik büyüme oranları yüzde 5'in altına düşerken, bilgisayar satışları yüzde 60 oranında arttığında ekonomi yüzde 15'i aşabilen oranlarda büyüyor.

Burada doğrusal bir ilişki var. Bilgisayar satışları arttığında ekonomi büyürken, satışlar azaldığında ekonomi küçülüyor.

Eğer verilerle en uyumlu doğrunun denklemini bulabilirsek, bilgisayar satışlarındaki her değişim (x) için yeni bir tahminî ekonomik büyüme oranı (y) hesaplayabiliriz. Denkleminiz, $y = a + bx$ şeklinde. Burada, bulmamız gereken iki parametre, yani iki karakteristik değer var: a ve b. Bunların hesaplanmasına "doğrusal regresyon" deniyor.

Makine öğrenmesinde, doğrusal regresyon yapabilen bir algoritma, ona sunulan verileri (Resim 12'deki noktaları) kullanarak, en uygun doğruyu bulmaya çalışır.

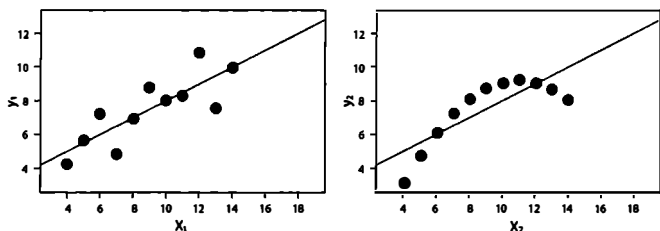
Algoritma bu işe rasgele bir doğruyla (Resim 12'de d₁'le) başlar. Ardından, noktalar ile bu doğru arasındaki uzaklıkları hesaplar. Sonra, a'nın ya da b'nin değerinde aşağıya ya da yukarıya doğru çok küçük bir değişiklik yaptığında, uzaklıkların toplamının azalıp azalmayacağını kontrol eder. Diyelim ki, a'nın değerini çok az artırdığında uzaklıkların toplamı azalacaksa, a'nın değerini çok az artırır. Eğer b'nin değerini çok az aşağı indirdiğinde uzaklıkların toplamı azalacaksa, b'nin değerini çok az aşağı indirir. Sonunda, a'da ya

da b'de yapacağı bir değişikliğin uzaklıklar toplamını azaltmayacağı bir aşamaya ulaşır. Böylece, noktalarla en uyumlu doğruyu (Resim 12'de d_2 'yi) bulmuş olur. İşte bu makine öğrenmesinin bir örneği...

Makine neyi öğrenmiş oluyor?

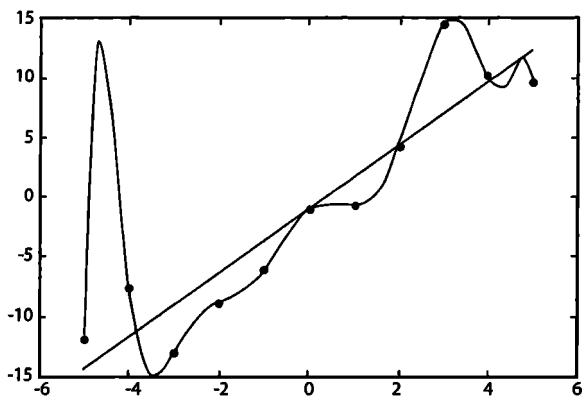
Onu eğitmek için kullanılan veriler yardımıyla, yeni bir veri sunulduğunda tahmin yapmasını öğrenmiş oluyor.

Tabii burada dikkat edilmesi gereken bazı şeyler var. Bir doğrusal regresyon algoritması, sadece doğrusal regresyon yapar. Veriler arasındaki ilişkinin gerçekten de doğrusal olup olmadığını sorgulamaz. Örneğin, Resim 13'te, sağdaki şekilde görülen veriler doğrusal regresyon için uygun değil. Sağdaki şekil için, $y = a + bx + cx^2$ şeklindeki bir denklemin parametrelerini (a, b ve c'yi) hesaplayabilen bir algoritmaya, yani daha karmaşık bir algoritmaya ihtiyaç duyarız. Eğer bunu yapmayıp basit algoritmamızda ısrar edersek, “yetersiz uyumlu” (*underfitted*) bir model kullandığımızdan “yanlı” (*biased*) sonuçlar elde ederiz.



Resim 13: Doğrusal regresyon için uygun olan ve olmayan veriler¹²

Tersi de olabilir. Resim 13'teki çok daha fazla parametrelili eğri, tüm verilerle tam olarak uyuyor. Ne var ki, bunun tümüyle rastlantısal olma olasılığı çok yüksek ve yine büyük olasılıkla doğrusal regresyon modeliyle daha iyi tahminler yapılabilir.



Resim 14: Aşırı uyumluluk (overfitting) örneği¹³

Peki bunlar gerçek yaşamda ne işe yarıyor?

İyi tahminler pek çok alanda işe yarar. Örneğin, bir virüs salgını sırasında, bu salgına daha önce maruz kalmış olan ülkelerdeki veriler yardımıyla, hiçbir şey yapılmaması, sınırlı önlemlerin alınması ya da katı önlemlerin alınması durumlarında ne tür sonuçlarla karşılaşılacağı öngörülebilir. Bu sayede ölümlerin sayısını en aza indirebilecek olan yol seçilebilir.

Geçmiş yılların yağış ve tarımsal üretim verileri kullanılarak, yağışların az olduğu bir yılda tarımsal

üretiminin ne kadar azalacağı tahmin edilebilir ve gerekli önlemler alınabilir.

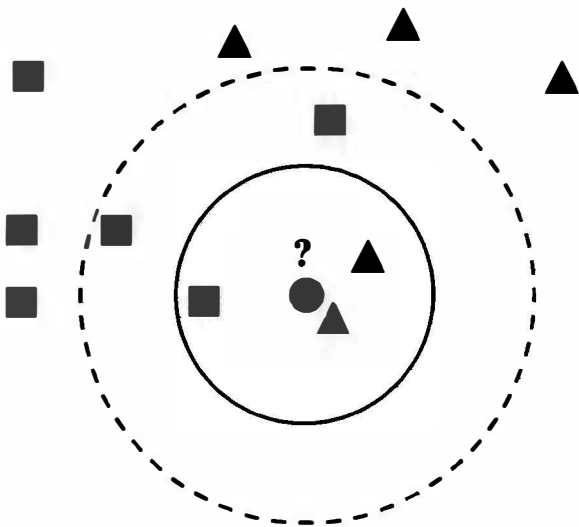
Yabancı dil eğitiminde en verimli sonuçları almak için sınıflar kaç kişilik olmalı? Farklı büyüklüklerdeki sınıflarda yabancı dil eğitimi alanların sınav notlarından yararlanarak bu konuda bir fikir edinebiliriz.

Tabii istatistiksel modelleri kullanırken her zaman dikkatli olmak gerekir. Eğer sınavlar ve öğretmenlerin not verme yaklaşımları farklıysa, sınav notları yanıltıcı olabilir. Farklı eğitim yöntemleri başarı düzeylerini farklılaştırabilir. İdeal sınıf büyüklüğünü tahmin etmeye çalışırken, bu tür farklılıkları hesaba katmak için, verileri farklı gruplara ayırmak ve birbirlerine gerçekten benzeyen verileri karşılaştırmak gerekir.

Bazen, verileri farklı gruplara ayırmak da yeterli olmaz, çünkü eldeki model ile veriler tümüyle uyumsuzdur. Gerçek yaşamdaki karmaşık sorunları basit modellerle çözme çabaları yanlış sonuçlar üretebilir. Yapay zekâ alanı için sık sık hatırlatıldığı üzere, elimizde çekiç var diye her sorunu çivi saymamız yanlış olur.

Makine öğrenmesi tekniklerinin hepsi uygun doğruların ya da eğrilerin bulunmasıyla mı ilgili?

Hayır. Çok farklı sorunlar için çok farklı algoritmalar kullanılıyor. Sadece örnek olması için bir tanesine daha değinebiliriz.



Resim 15: En yakın k komşu algoritması¹⁴

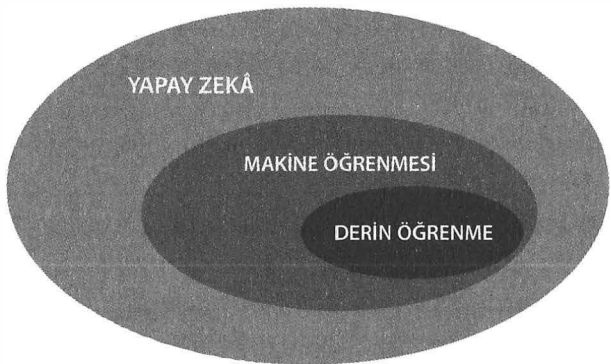
Diyelim ki, elimizde fiyatlarını bildiğimiz evler var ve bunların konumlarından hareketle fiyatlarını bilmediğimiz evler için tahmin üretmek istiyoruz. Resim 15'teki üçgenler ucuz, kareler de pahalı evleri temsil ediyor olsun. Ortadaki, fiyatını bilmediğimiz ev ucuz mudur, yoksa pahalı mı? Bunun için “en yakın k komşu algoritması”nı kullanırsak, belirleyeceğimiz bir “k” değerine bağlı olarak, evin ucuz mu yoksa pahalı mı olduğu hakkında bir tahmin elde edebiliriz. $k = 3$ dersek, en yakın 3 komşudan 2 tanesi üçgen olacağından, algoritma, evin ucuz olduğuna karar verir. $k = 5$ dersek, en yakın 5 komşudan 3 tanesi kare olacağından, algoritma, evin pahalı olduğuna karar verir.

En uygun k değerini bulmak için, değerlerini bildiğimiz evleri kullanarak, farklı k değerlerinden hangisinin en doğru tahminleri ürettiğini saptamaya çalışabiliriz. Böylece değerlerini bilmediğimiz evler için daha doğru tahminler elde etme olasılığımız artabilir.

Makine öğrenmesinin alt kategorilerinden biri, son yıllarda hayli popüler olan derin öğrenme (*deep learning*). Bunun için, insan beynindeki sinir ağından esinlenerek üretilen yapay sinir ağları kullanılıyor. Aralarında bağlantılar bulunan yapay sinir hücreleri, girdilerine bağlı olarak, başka yapay sinir hücrelerine ulaşan çıktılar üretiyor. “Esinlenme” diyorum, çünkü aslında beyindeki sinir hücrelerinin nasıl çalıştığı henüz tam olarak bilinmiyor. Diğer yandan derin öğrenme alanında çalışan uzmanlardan bazıları beynin nasıl çalıştığına ilgi duyarken bazıları bunu pek umursamıyor.¹⁵

Son yıllardaki popülerliği nedeniyle, derin öğrenmenin yapay zekâyla özdeş olduğu sanılabiliyor. Oysa Resim 16’da görüldüğü gibi, derin öğrenme, yapay zekâ alanının bir alt kategorisi olan makine öğrenmesinin bir alt kategorisi.

Yapay sinir ağlarının tarihi hayli eskilere uzansa bile, bunların gerçekten işe yarar hale gelmesi için bilgisayarın çok daha güçlü işlemcilere sahip olmasını beklemek gerekmişti.



Resim 16: Derin öğrenmenin yapay zekâ alanındaki yeri

2012 yılında, başında Geoffrey Hinton'ın bulunduğu bir ekip, bir yapay sinir ağını kullanarak, en önemli görüntü tanıma yarışmasındaki hata oranını yüzde 26,1'den yüzde 15,3'e düşürdü. Yani, kendisine kedi, köpek, kuş, insan, otomobil vb. resimleri gösterilen bir yapay sinir ağı, ilgili resimlerin kategorilerini belirlemek konusunda yaklaşık yüzde 85'lik bir "doğruluk" oranına ulaştı. Ama burada medyaya pek yansımayan bir ayrıntı vardı. Yapay sinir ağı, belirli bir resmi tek bir kategoriye sokmak yerine, doğru olma olasılığını yüksek gördüğü 5 kategoriyi sıralıyordu. Ve ilgili resmin bu 5 kategoriden birine girmesi durumunda cevap doğru kabul ediliyordu.¹⁶

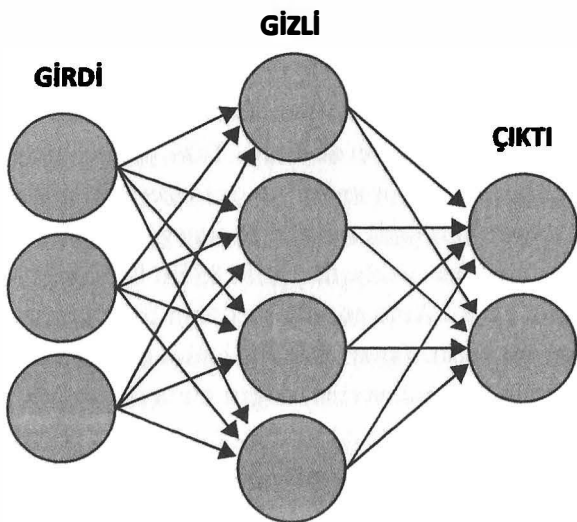
O zaman yüzde 85'lik doğruluktan söz etmek yanıltıcı değil mi?

Biraz öyle. Yapay zekâyla ilgili haberleri okurken yüzdelerin tam olarak ne anlama geldiğine hep dikkat

etmemiz gerekir. Özellikle de “yapay zekâ şu alanda da insanları geride bıraktı” türü haberler söz konusu olduğunda.

Gelelim yapay sinir ağlarının nasıl çalıştığına...

Resim 17’de, sadece fikir vermek için, aşırı derecede basit bir yapay sinir ağının şeması var. Bu ağ üç katmandan oluşuyor. Bir girdi katmanı, bir çıktı katmanı ve bir de gizli katman. Ortadaki katman “gizli”dir, çünkü girdiler ve çıktılar dışarıdan görülürken bu katman ağın içinde kalır.



Resim 17: Aşırı basit bir yapay sinir ağının şeması

Diyelim ki, bu ağa, resimlerde kedilerin olup olmadığını öğretmek istiyoruz. Yapay sinir hücreleri ara-

sındaki bağlantılar başlangıçta rasgeledir. Ağa (girdi katmanına) bir resim gösteririz ve çıktı katmanından rastlantısal bir şekilde “kedi var” ya da “kedi yok” sonucunu alırız. Sonucun yanlış olduğunu söylediğimizde, yapay sinir ağı üzerinde çalışan algoritmamız geriye doğru giderek hücreler arasındaki bağlantıların bazılarını güçlendirirken bazılarını zayıflatır. Bu da rastlantısal bir şekilde yapılır. Resim yeniden gösterildiğinde doğru cevap alınırsa bağlantılarda değişiklik yapma süreci son bulur ve yeni bir resimle aynı işlemler tekrarlanır. Buna gözetimli öğrenme (*supervised learning*) deniyor, çünkü “kedi var” ya da “kedi yok” diye önceden etiketlenmiş olan resimlerle ağı eğitiyoruz.

Yeterince çok sayıda (çoğu zaman milyonlarca) resim üzerinde çalışıldığında, yapay sinir ağının, kendisine yeni bir resim gösterildiğinde onu doğru kategoriye sokma olasılığı yükselir. Yani yapay sinir ağı giderek daha iyi tahminler yapar.

Tabii bütün bunlar Resim 17’deki aşırı basit yapay sinir ağıyla yapılamaz. Geoffrey Hinton ile ekibinin 2012 yılında geliştirdikleri ağda 650 bin yapay sinir hücresi vardı.¹⁷ Gelişkin yapay sinir ağlarında çok sayıda gizli katman bulunur. Derin öğrenmenin “derin”liği de buradan gelir: Ne kadar çok sayıda gizli katman varsa öğrenme de o kadar derin kabul edilir. Diğer yandan gelişkin ağların iç yapıları çok daha karmaşıktır ve bunların farklı türleri bulunur.

Hangi sorunu çözmek için ne tür bir ağ kurmak gerektiğine ise insanların karar vermesi gerekiyor. Bu da, bolca deneme yanılma içeren, hayli zorlu bir süreç. Bu işi de en azından kısmen makinelere yaptıрма girişimlerinde bulunulsa bile, önüne koyulan her tür sorunu çözebilecek, genel amaçlı bir yapay sinir ağına sahip değiliz.

Yapay zekânın son yıllarda yeniden popüler olmasının en önemli nedeni, derin öğrenme sayesinde elde edilen başarılar. Derin öğrenme, sadece görüntü tanıma alanında değil, başta konuşma tanıma ve makine çevirisi olmak üzere farklı alanlarda, başka tekniklerle elde edilebilenlerden çok daha iyi sonuçların alınmasını sağladı.

“Konuşma tanıma” tam olarak ne anlama geliyor?

Konuşma tanıma algoritmaları, konuşmaları metinlere dönüştürüyor. YouTube’da, İngilizce, Almanca ya da Rusça gibi dillerdeki konuşmaların otomatik olarak hazırlanan metinlerini okuyabiliyor, hangi sözlerin hangi saniyelerde söylendiğini görebiliyorsunuz. Böylece, örneğin iki saatlik bir videoda ilginizi çeken bir konunun ne zaman ele alındığını öğrenerek zaman kazanabiliyorsunuz. Ayrıca dilerseniz konuşmaların farklı dillerdeki makine çevirilerini altyazı şeklinde görebiliyorsunuz. Bu sonuncular çoğu zaman pek başarılı olmasa bile, en azından nelerden söz edildiği hakkında fikir sahibi olabiliyorsunuz.

İnternette farklı konuşma tanıma programları bulabilirsiniz. Bunları kullanarak videolarda ve ses kayıtlarında sözcük araması yapabilmek önemli bir olanak.

Derin öğrenme sayesinde, akıllı telefonlarımızı kullanarak, yurt dışındaki tabelalarda yazılı olan metinleri de çevirtebiliyoruz. Kuşkusuz, çevirilerin hatalı olabileceğini bilmeli, onlara körü körüne inanmamalıyız.

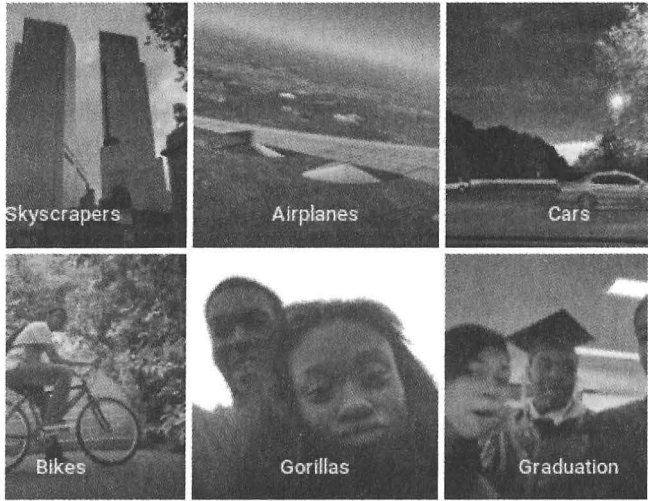
Siyah beyaz fotoğraflar ve filmler, derin öğrenme yardımıyla, geçmişte olduğundan daha başarılı bir şekilde restore edilip renklendirilebiliyor. Spotify, hoşlanabileceğimiz şarkıları tahmin ederken yapay sinir ağlarını da kullanıyor. Derin öğrenme, finansal dolandırıcılık girişimlerinin saptanmasını kolaylaştırıyor. Yine derin öğrenme yoluyla, ünlü ressamların tarzlarını yansıtan resimler ya da ünlü bestecilerin tarzlarını yansıtan besteler üretilabiliyor. Kısacası, derin öğrenmeden pek çok alanda yararlanılabiliyor. Büyük ölçekli yapay sinir ağlarını kullanmanın maliyeti azaldıkça bu alanların sayısı artacaktır.

Ama derin öğrenmenin bazı sorunları üzerinde de durmamız gerekir.

Birincisi, yapay sinir ağları, kurallara dayalı algoritmalardan farklı olarak, birer kara kutudur. Hangi sonuca hangi nedenlerle ulaştıklarını anlamak çoğu zaman mümkün olmaz. Dolayısıyla, hangi durumlarda yanlış sonuçlar üretebileceklerini de önceden bilemeyiz.

İkincisi, bu ağlar fazlasıyla kırılıgandır. Onları eğit-
mek için kullanılan verilere de bağılı olarak kolaylıkla
yanlı ya da tümüyle yanlış sonuçlara ulaşabilirler.

Örneğin, Google'ın resim etiketleme programı,
2015 yılında, iki siyahın bulunduğı bir resmi "goriller"
diye etiketleyerek bir skandala imza atmıştı (Resim 18).



Resim 18: Google'ın resim etiketleme programına göre:
"Gökdelenler", "uçaklar", "otomobiller", "bisikletler", "goriller" ve
"mezuniyet"¹⁸

Bu örnekte, yapay zekâ algoritmalarını eğitmek için
kullanılan resimlerde yeterli sayıda siyahın bulunma-
masıyla ilgili bir sorun yaşanmıştı. Buna eksik veriye
dayalı "yanlılık" (*bias*) deniyor...

Google bu olay yaşandığında ne yaptı?

Yeni ürün ve hizmetleri yeterince ön çalışma yapmadan kullanıma sokma alışkanlığına sahip olan teknoloji şirketlerinin bu tür durumlarda hep yaptığı şeyi yaptı: Özür diledi ve “goril” kategorisini bir süreliğine tümüyle kaldırdı.¹⁹

Görüntü tanıma algoritmalarını yanıltmak da zor olmuyor. Örneğin, üzerine çıkartmalar yapıştırılan bir trafik levhasını “tıka basa dolu bir buzdolabı” diye tarif edebiliyorlar. Renkli gözlük takan bir insanı tümüyle ilgisiz biriyle karıştırabiliyorlar. İnsan gözüyle algılamayacak değişiklikler yapılmış bir resmi yanlış bir kategoriye sokabiliyorlar.

Bunlar, derin öğrenmeyle ilgili çok temel bir soruna işaret ediyor: Derin öğrenme algoritmaları, gerçekte, dünyada olup bitenleri hiçbir şekilde anlamıyor. Sadece, belirli veriler yardımıyla rastlantısal olarak oluşturulan modellere dayalı otomatik tahminler yapıyorlar.

Derin öğrenmeyle ilgili üçüncü bir sorun, çok fazla veriye gereksinim duyması. Zaten, yapay sinir ağlarının işe yarar hale gelmesini biraz da “büyük veri”ye (*big data*) borçluyuz. İnternet, insanlığın bilgi birikiminin büyük bir bölümünün dijital ortama aktarılmasını ve veri üretiminin devasa boyutlara ulaşmasını sağladı.

Büyük verinin kesin bir tanımı bulunmasa da, 2016 yılında yapılan bir tahmine göre, yapay sinir ağları,

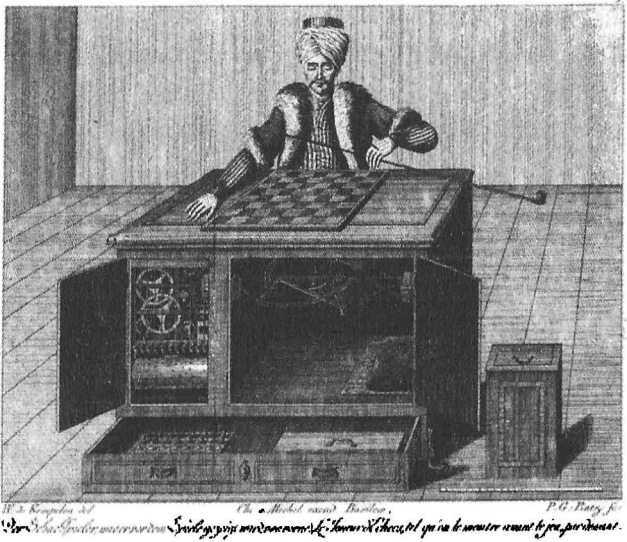
yaklaşık 5000 etiketli örnek kullanıldığında “kabul edilebilir” bir performans sergilerken, en az 10 milyon etiketli örnekle insanlardan daha başarılı olabiliyor.²⁰

Etiketleme işi nasıl yapılıyor?

Bu iş asıl olarak insanlara düşüyor. İnternette bulunan milyonlarca resimdeki kediler, köpekler, insanlar, trafik levhaları, yaya geçitleri, uçaklar ve başka pek çok şey, örneğin Amazon’un Mekanik Türk (*Mechanical Turk*) adlı platformunda çalışanlar tarafından etiketleniyor. Hiçbir geliştiriciliği bulunmayan, aksine son derece usandırıcı ve yıpratıcı olan bu işi uzun saatler boyunca yapan insanlara en düşük ücretler ödeniyor. Mekanik Türk’te sigortasız olarak çalıştırılanlara herhangi bir iş güvencesi de sağlanmıyor.

İsmi neden Mekanik Türk?

“Mekanik Türk” ya da sadece “Türk”, 18. yüzyılda Avrupa’da üretilen bir sahte satranç otomatının adıydı. Bu otomatta, Resim 19’da görülebileceği gibi, bir satranç masasının arkasında, kol ve el hareketleriyle taşların yerlerini değiştirebilen, Osmanlı giysili ve Türk görünümüne sahip bir manken duruyordu. Otomat çalışırken masanın kapakları kapalı oluyordu. Uzun yıllar boyunca devlet yöneticileri dahil çok sayıda insanla satranç oynayarak ilgi odağı haline gelen bu sahte otomatın sırrı, satranç tahtasının altına saklanan bir satranç oyuncusunun varlığıydı.



Resim 19: Mekanik Türk²¹

Amazon'un platformu için seçilen isim gerçekten de uygun, çünkü çok düşük parça başı ücretlerin ödendiği bu tür platformlar yapay zekânın gelişiminde önemli bir rol oynuyor.

Bu arada biz de ücretsiz olarak teknoloji şirketlerine hizmet ediyoruz. Facebook'taki resimleri etiketlediğimizde bu şirketin algoritmalarını eğitmiş oluyoruz. Robot olmadığımızı kanıtlamak amacıyla içinde trafik lambası bulunan kareleri işaretlediğimizde, algoritmaların eğitime katkıda bulunuyoruz. "Paskalya tavşanı" sözcüklerini kullanarak resim arattığımızda tıkladığımız ilk resim, en iyi paskalya tavşanı resimleri sıralamasında basamak atlıyor.²² Daha genel olarak bak-

tığımızda, Google'a sorduğumuz her soru, Facebook ya da Twitter gibi platformlardaki her paylaşımımız ve her eylemimiz büyük verinin bir parçası haline geliyor.

Derin öğrenme için büyük veriye gereksinim duyulması, tanımlı örneklerin az olduğu alanlarda bu tekniği kullanamayacağımız ya da kullanırsak yanlış sonuçlara ulaşacağımız anlamına geliyor. Aynı nedenle, sadece geçmişe ait verileri kullanabilen derin öğrenme, yeni ortaya çıkan verilerden yanlış sonuçlar çıkarmaya yatkın. Oysa sürekli değişen, sürekli yeni veriler üreten bir dünyada yaşıyoruz.

İnsan zekâsını bugünkü yapay zekâdan ayıran özelliklerden biri, az sayıda örnekten doğru sonuçlar çıkarabilmemiz. Örneğin küçük bir çocuğa kedinin ne olduğunu öğretmek için milyonlarca kedi göstermemiz gerekmez. Sadece birkaç kedi gördükten ve bunların kedi olduğunu öğrendikten sonra, renkleri ya da büyüklükleri ne olursa olsun tüm kedileri tanır...

3. Yapay Zekâ ve İnsan Zekâsı

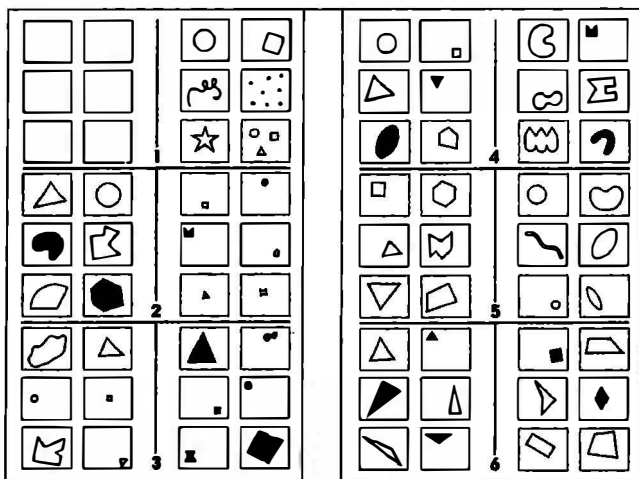
İnsanlar da kolaylıkla yanılgılara düşmüyor mu?

Elbette. Her insan yanılır. İnsan zekâsının ayırt edici özelliği, yanılmazlığı değil, çok farklı sorunları görece az sayıda veriyle çözebilme yeteneği.

Makinelerin düşünüp düşünmediği hakkındaki Turing Testinden söz etmiştim. Yapay zekânın düzeyini ölçmeye yönelik çok daha iyi bir test 1966 yılında Sovyet bilimci M. Bongard tarafından geliştirilmişti. Bongard, 100 problem üretmişti. Bunların her birinde, soldaki altı resim ile sağdaki altı resmi ayırt eden özelliklerin ifade edilmesi gerekiyor. İlk altı problemi ve cevaplarını Resim 20’de görebilirsiniz.

Bongard’ın amacı, “düşünen makine” yapmanın aslında hiç de o kadar kolay olmadığını göstermekti. Ona göre, daha fazla belleğe ve daha güçlü işlemcilere sahip olmanın her sorunu çözmeye yeteceği düşüncesi yanlıştı.

İnsanların çoğu zaman kolayca çözebildiği Bongard problemleri, Batı’da, Amerikalı bilimci Douglas R. Hofstadter’in ilk basımı 1979’da yapılan ve Türkçeye *Gödel, Escher, Bach: Bir Ebedi Gökçe Belik* başlığıyla çevrilen kitabıyla ün kazandı.²⁴ Benzerleri kolaylıkla üretilebilen bu tür problemler en yeni yapay zekâ teknikleriyle de çözülemiyor.



Resim 20: Bongard problemleri (Cevaplar: 1. Boş resimler/boş olmayan resimler; 2. büyük şekiller/küçük şekiller; 3. içi boş şekiller/içi dolu şekiller; 4. dışbükey şekiller/dışbükey olmayan şekiller; 5. köşeli şekiller/eğrilerden oluşan şekiller; 6. üç köşeli şekiller/dört köşeli şekiller)²³

Bu konuda başka örnek var mı?

Turing Testinden daha iyi olan bir başka test daha var: Winograd şemaları. Bilgisayar bilimcisi Terry Winograd, 1972 yılında, doğal dilleri anlayacak bir bilgisayar sistemi geliştirmenin zorluklarına örnek göstermek için iki farklı cümle yazmıştı: “Şehrin yöneticileri göstericilere izin vermedi, çünkü *şiddet olaylarından korkuyorlardı*” ve “şehrin yöneticileri göstericilere izin vermedi, çünkü *devrimi savunuyorlardı*”.²⁵ Şiddet olaylarından korkanlar kimler, devrimi savunanlar kimler? İnsanlar, dünyada olup bitenler hakkındaki genel bilgileri sayesinde hangi durumda kimden söz edildiğini

kolaylıkla anlayabilirken, makinelerin bunu yapması hâlâ olanaksız. Bir örnek daha vereyim: “Ağacı baltayla kesemedim, çünkü çok *kalındı*” ve “ağacı baltayla kesemedim, çünkü çok *küçüktü*”. Kalın olan hangisi, küçük olan hangisi?

Google Çeviri’yi geliştirmek için, bugüne kadar yapılmış olan ve internet aracılığıyla erişilebilen tüm çevirilerden yararlanılıyor. Birleşmiş Milletler ya da Avrupa Birliği gibi kurumlar için yapılmış olan resmî çeviriler, diplomatik, hukukî ve iktisadi metinlerin makine çevirilerini daha başarılı kılıyor. Ama bunun için yine istatistiksel yöntemlerden yararlanılıyor. Google Çeviri’nin algoritmaları, başarı düzeylerini, okuduklarını daha iyi anlayarak değil, çok büyük veri setleri arasında daha iyi eşleştirmeler yaparak yükseltiyor. Ama bir cümledeki “o” sözcüğüyle kimin kastedildiğini anlayamıyorlar.

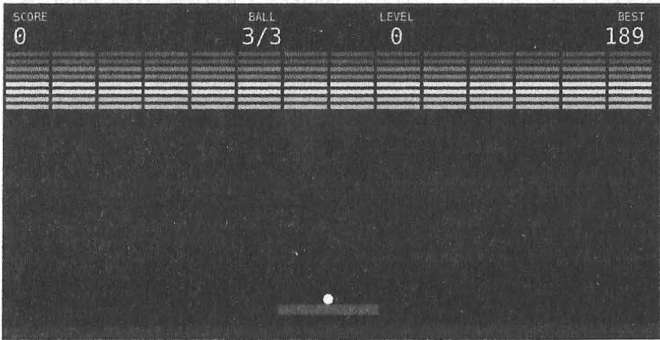
2016 yılında Winograd şemaları için 25 bin dolar ödüllü bir yarışma düzenlenmişti. Ödülü kazanmak için yüzde 90’lık bir doğruluk oranına ulaşmak gerekiyordu. Ama en yüksek skor yüzde 58, en düşüğü de yüzde 32 olmuştu. Yani yazı-tura atarak da benzer sonuçlara ulaşılabilirdi.²⁶

O zaman bilgisayarlar satranç ya da Go gibi oyunlarda insanları yenmeyi nasıl başarıyor?

Çünkü doğal dillerden ya da gerçek yaşamdaki sorunlardan farklı olarak, bu oyunların kuralları belli,

az sayıda ve hiç değişmiyor. Bilgisayarlar, geçmişte, insanlar arasında oynanmış milyonlarca oyunun bilgisini kullanıyordu. Bugünse, derin öğrenme teknikleriyle, kendi kendilerine oynadıkları milyonlarca oyunun bilgisini kullanarak insanları yenebiliyorlar.

Oyunlar için kullanılan derin öğrenme tekniği, pekiştirmeli öğrenme (*reinforcement learning*). Yapay sinir ağlarına Atari oyunları da oynatılıyor ve pek çoğunda insanlardan daha başarılı oluyorlar. Pekiştirmeli öğrenmede etiketli veriler kullanılmıyor. Bunların yerine, yapay sinir ağının deneme yanılma yoluyla öğrenmesi sağlanıyor.



Resim 21: *Breakout*²⁷

Örneğin, *Breakout* (Kaçış) oyununda (Resim 21), oyuncu, raket olarak iş gören bir çizgiyi sağa ya da sola doğru hareket ettirebiliyor ve bununla topa vurabiliyor. Yukarıdaki tuğlalara (kısa çizgilere) çarpan top onları parçalıyor. Oyndaki hedef, topu oyun ala-

nında tutmak (aşağıya kaçırmamak) ve tüm tuğlaları parçalamak.

Pekiştirmeli öğrenmede, oyunun kuralları hakkında önceden bilgilendirilmeyen algoritma, tümüyle rastlantısal bir şekilde, “sağa git”, “sola git” ya da “olduğun yerde kal” komutlarını veriyor. Raket tesadüfen topa vurduğunda ve bir tuğla parçalandığında, belirli bir puanla ödüllendiriliyor. Topa yine tesadüfen iki kez üst üste vurmayı başardığında daha yüksek bir puana ulaşıyor. Algoritma, çok sayıda oyun oynayarak ve çok farklı denemeler yaparak, olası en yüksek puanı kazandıran hamleleri bulmaya çalışıyor.

Daha sonra Google tarafından satın alınan yapay zekâ şirketi DeepMind’in geliştirdiği *Breakout* programı, bu oyunda sonuç almayı kolaylaştıran bir yol keşfetmiş. Kenarlardaki tuğlaları daha erken parçalayarak bir tünel açıldığında, o tünelden geçip oyun alanının “tavan”ı ile tuğlalar arasında gidip gelmeye başlayan top, kısa sürede çok sayıda tuğlayı parçalayabiliyormuş.²⁸

Burada dikkat edilmesi gereken bir nokta, *Breakout* oynamayı öğrenen programın, oyunun “mantık”ı hakkında herhangi bir fikir geliştirmiş olmadığı gerçeği. Bu oyunu oynamayı öğrenen bir insan, raketin uzunluğunda, kalınlığında ya da bulunduğu yükseklikte küçük bir değişiklik yapılmasından fazla etkilenmez. Buna karşılık DeepMind’in programı bu tür küçük değişiklikler yapıldığında her şeyi sıfırdan öğrenmek zorunda kalmış.²⁹ Dolayısıyla, gerçek dünyadaki so-

runların çoğu, pekiştirmeli öğrenmeyle çözülebilecek türden değil.

Pekiştirmeli öğrenme, oyunlar dışında bir işe yaramıyor mu?

Hataların fazla sıkıntı yaratmadığı ya da simülasyon yapılabilen alanlarda yararlı olabiliyor elbette. Örneğin bir robot koluna deneme yanılma yoluyla bazı işleri yapmasını öğretmek için pekiştirmeli öğrenme tekniğinden yararlanılabiliyor. Yapay zekâ, bazı durumlarda, hatalarından ders çıkarmak konusunda da insanlara üstünlük kurabiliyor!

Bilgisayarların insanlara üstünlük kurduğu başka bir örnek üzerinde durayım. IBM'in geliştirdiği Watson adlı bilgisayar sistemi, 2011 yılında, ABD'deki *Jeopardy!* (Riziko!) adlı televizyon yarışmasında, daha önce bu yarışmada büyük paralar kazanmış olan iki insanı yenmeyi başarmıştı. Bu yarışmada, verilen ipuçlarından hareketle soruyu bulmak gerekiyordu. Önündeki düğmeye ilk basan yarışmacı ilk cevap hakkına sahip oluyor, doğru cevabı verememesi durumunda cevap hakkı diğer yarışmacılardan düğmeye ilk basanına geçiyordu.

Watson'ı yarışmaya hazırlamak için başta Wikipedia olmak üzere pek çok ansiklopedi bu bilgisayar sistemine yüklenmişti. Watson'ın yaptığı şey, ipuçlarındaki sözcüklerle, bunlarla ilişkili olabilecek olan ansiklopedik bilgileri eşleştirmeye çalışmaktı. İstatistiksel modeller yardımıyla, örneğin, "başkenti Ankara

olan ülkedir” sözcüklerinden hareketle, “Türkiye nedir?” diyebiliyordu.

Hem IBM hem de medya, Watson’ın *Jeopardy!* oyunundaki başarısını, bilgisayarların İngilizceyi gerçekten anlamaya başladığının kanıtı olarak sunmaya çalıştı. Ama verdiği doğru cevaplar kadar yanlış cevaplar da önemliydi.

Watson’ın ilk yanlışı, şu ipucuyla ilgiliydi: “1904’te paralel bar dalında altın madalya kazanan ABD’li jimnastikçi George Eyser’in *tuhaf anatomik özelliği*di.” Şu cevabı vermişti: “Bacak nedir?”. Oysa herkesin kolaylıkla anlayabileceği gibi, “bacak”, “tuhaf bir anatomik özellik” değildir. Eyser’in özelliği, bir bacağının eksik olmasıydı.

İkinci yanlış, “Hangi 10 yıllık dönem” kategorisiyle ilgiliydi (1970’ler, 1980’ler vb.). İpucu şuydu: “İlk modern çapraz bulmaca yayımlandı ve Oreo bisküvileri piyasaya sürüldü.” Önündeki düğmeye ilk basan yarışmacı “1920’li yıllar” deyince, bu cevap yanlış olduğundan sıra Watson’a gelmişti. Watson’un cevabı aynıydı: “1920’li yıllar”. Çünkü, programcılar, diğer yarışmacıların yanlış cevaplarından sonuç çıkarmasını ona öğretmemişti.

Bu durumda yanlışın sorumluluğu Watson’a değil programcılara ait değil mi?

Öyle. Ama İngilizceyi gerçekten anlayan bir bilgisayar, cevap verme hakkının, “1920’li yıllar” cevabının yanlışlığı nedeniyle ona geçtiğini de anlayabilirdi.

Üçüncü yanlış, “ABD’deki kentler” kategorisiyle ilgiliydi. İpucu: “En büyük havalimanı, bir İkinci Dünya Savaşı kahramanının adını taşır; en büyük ikinci havalimanı, İkinci Dünya Savaşı’ndaki bir muharebenin adını taşır.” Watson’ın cevabı: “Toronto nedir?”³⁰ Yani bu kez tümüyle ilgisiz bir cevap vermişti (Kanada’daki Toronto kentinin en büyük havalimanı bir İkinci Dünya Savaşı kahramanının adını taşımıyor).

Teknoloji şirketleri, yapay zekânın satranç, Go ya da *Jeopardy!* gibi oyunlardaki başarılarını, ürünleri hakkında büyük beklentiler yaratmak için kullanıyor. IBM’in başlıca amaçlarından biri sağlık sektörüne girmekti. Doğal dilleri anlayan bir bilgisayar sistemi, sağlıkla ilgili tüm bilimsel yayınları inceleyebilir ve hasta kayıtlarından hareketle her bir birey için en uygun tedavi önerilerinde bulunabilirdi. Milyarlarca dolarlık yatırımlar yapan IBM, Batı ülkelerindeki çok sayıda sağlık kuruluşuyla işbirliği anlaşmaları imzaladı. Ama 2019 yılına gelindiğinde durum pek de parlak sayılmazdı. Watson’ı yararsız bulan birçok sağlık kuruluşu IBM’le işbirliğine son vermişti. IBM de bunun üzerine Hindistan, Güney Kore ve Tayland gibi ülkelerdeki hastanelerle anlaşmalar yapma yoluna gitti. Bu hastaneler, hasta ya da daha doğrusu müşteri kazanmak için, “yapay zekâ destekli kanser tedavisi” sunduklarını söylüyormuş!³¹

IBM’in yöneticileri Watson’ın yetersizliklerini bilmiyor muydu?

Biliyorlardı elbette. Ama teknoloji şirketleri, yatırımcıları, kısa vadede büyük başarılar ya da kazançlar sağlayabileceklerine inandırmaya çalışıyor. Şirketlerin hisse senetleri bu tür beklentiler sayesinde değer kazanıyor.

Bu arada, yapay zekâ alanındaki ilerlemeler, uzmanlarda da, gerçek yaşamdaki pek çok sorunun büyük veri ve algoritmalar yardımıyla yakın gelecekte çözülebileceği inancını üretiyor.

Şöyle düşünüyorlar: Doktorların çoğu, az zamanda çok hastaya bakmak zorunda kalıyor. Her bir hastalarının tüm sağlık verilerini ayrıntılı bir şekilde incelemeleri kolay olmuyor. Tıp alanındaki bütün kitapları ve bilimsel makaleleri okuyamıyorlar. Ayrıca herhangi bir hastaları hakkında karar almaya çalıştıkları sırada yorgun olabiliyor ya da herhangi bir sıkıntıları nedeniyle dikkatlerini toplamakta güçlük çekebiliyorlar. Dolayısıyla, bu tür sorunlar yaşamayan bilgisayarlarla çok daha güvenilir tahminler yapılabilir...

Örneğin, yeni teknolojiler hakkındaki ortak kitapları ünlü olan Andrew McAfee ile Erik Brynjolfsson, uzman insanların psikoloji ve tıp alanlarındaki tahminleriyle istatistiksel yöntemlere dayalı tahminleri karşılaştıran bir araştırmadan söz ediyor. Bu araştırmaya göre, incelenen 136 bilimsel çalışmanın yüzde 48'inde, uzmanların tahminleri ile istatistiksel formüllerin verilere uygulanması yoluyla elde edilen tahminler arasında önemli bir fark bulunmuyormuş.³²

Bu sonucu hiç sorgulamadan doğru kabul etsek bile, yüzde 48'lik oranın örneğin yüzde 95'e çıkarılmasının önünde ciddi engeller var.

Bunların bir bölümü verilerle, bir bölümü de elimizdeki yapay zekâ teknikleriyle ilgili.

Bir bilgisayarın, hastaların gelmiş geçmiş tüm sağlık verilerini inceleyebilecek olması, kâğıt üzerinde, ciddi bir avantaj elbette. Ama hastane kayıtları çoğu zaman yeterince standart ve güvenilir değildir. Doktorlar, başka doktorların kolaylıkla anlayabileceği, ama bilgisayarların anlayamayacağı kısaltmalar kullanabilir. Yazım yanlışları yapabilirler. Sağlık sistemindeki herhangi bir sorun nedeniyle, hastalarının iyiliği için, gerçektekenden farklı tanımlar yazabilirler. Hastalar, sadece kan değerlerini ölçtürmek için ilgisiz servislere başvurabilir. Bu servislerde ilgisiz ön tanımlar koyulabilir. Dolayısıyla sağlık verileri yanıltıcı olabilir. Ve eğer veriler sorunluysa, yapay zekânın çıkaracağı sonuçlar da sorunlu olacaktır. Bununla ilgili deyim şu: Çöp girerse çöp çıkar.

Sorunlu veriler düzeltilemez mi?

Pek çoğu düzeltilebilir elbette. Ama bu görev, çok büyük ölçüde, yapay zekâ algoritmalarını geliştiren insanlara düşüyor. Hangi verilerin gerçekten işe yarayacağına, hangilerinin yok sayılması gerektiğine, eksik veriler için ne yapılacağına çok büyük ölçüde insanların karar vermesi gerekiyor. Bu da büyük veri söz konusu olduğunda çok fazla insan emeği gerektiren bir iş.

Bir bilgisayarın tıp alanındaki bütün bilimsel çalışmalara erişebilmesi de, yine kâğıt üzerindeki bir avantaj. Akademi dünyasındaki performans kriterleri, yükselmek ya da sadece yerini korumak isteyen akademisyenlerin niteliksiz yayınlar yapmasına yol açabiliyor. İstatistiksel açıdan “anlamli” görünen sonuçlara ulaşabilmek için verilerle oynanabiliyor, popüler deyiimiyle verilere işkence yapılabiliyor. Bilimsel araştırmalardan, tekrarlanabilir olmaları beklenir. Ama 2016 yılında yapılan bir ankete göre, aralarında tıp uzmanlarının da bulunduğu 1576 araştırmacının yüzde 70’i başkalarının çalışmalarını, yüzde 50’si de kendi çalışmalarını tekrarlamak konusunda başarısızlığa uğradıklarını söylemiş. Araştırmacıların yüzde 52’si “ciddi”, yüzde 38’i ise “hafif” bir “tekrarlanabilirlik bunalımı”nın yaşadığı yönünde görüş bildirmiş.³³ Ayrıca, bir bölümü sonradan geri çekilen hileli yayınlar da ciddi bir başka sorun kaynağı.³⁴

Bu söylediklerim, bilimsel araştırmaların gereksiz ya da değersiz olduğu anlamına gelmiyor tabii ki. Ama bunlardan doğru sonuçlar çıkarabilmek için, eldeki istatistiksel yöntemlerden farklı olan bir araca gereksinim var. Bu da, insanlarda olan, ama makinelerde bulunmayan sağduyu...

Bilimin çoğu zaman sağduyuya aykırı olduğu söylenmez mi?

Bilimsel gerçeklerin pek çoğunun sağduyuya aykırı olduğu doğru. Örneğin Güneş’in sabahları doğduğunu

ve akşamları battığını söyleriz. Dünya'nın Güneş'in çevresinde döndüğünü ancak 16. yüzyılda ve bilim sayesinde öğrenebildik.

Ama sağduyumuz, dünyada olup bitenleri anlamamızı fazlasıyla kolaylaştırıyor. Örneğin, 1920 yılında doğmuş, 1990 yılında ölmüş olan biri, 1956 yılında yaşıyor muydu? Bizim için apaçık olan cevap bilgisayarlar için öyle değil. Bir duvarın arkasına geçen bir kedi varlığını sürdürüyor mudur? Aynı kedi aynı anda hem İstanbul'da hem de Ankara'da olabilir mi? Herkesin bir annesi olmak zorunda mıdır?

1984 yılında, insanların sahip olduğu ama yazılı olmayan her tür bilgiyi bilgisayarların anlayacağı mantıksal ifadelerle çevirmek için bir proje başlatıldı. Varlığını bugün de sürdüren "Cyc" projesi kapsamında, insanlar, akıllarına gelen her şeyi mantık diline çeviriyor. Ama Cyc, makinelere sağduyulu olmayı öğretmek konusunda pek fazla yol alamadı.

Bunun tek nedeni, mantık diline çevrilmesi gereken bilgilerin çok fazla olması değil. Bir de bu dile kolay kolay çevrilemeyecek olan bilgiler var. Örneğin, annemizin yüzünü nerede görsek tanırız. Ama bunu sağlayan bilgilerimizi bir başkasına aktararak aynı şeyi onun da yapmasını sağlayamayız.³⁵ Bisiklet sürmesini bilsek bile, sadece bu konudaki bilgilerimizi paylaşarak bir başkasının hiç düşüp kalkmadan bisiklet sürmesini sağlayamayız. Ayrıca, sağduyuya dayalı akıl yürütme, soyutlama yapabilmeyi gerektirir. Bongard problemleri

ri ve Winograd şemaları üzerinde durmuştuk. Bunları çözemeyen makinelerin gerçek dünyayı gerçek anlamıyla anlaması mümkün olmayacaktır.

Benzer nedenlerle, örneğin sürücüsüz otomobil teknolojileri de beklenenden çok yavaş geliyor.

Google'ın kurucularından Sergey Brin, 2012 yılında, bu şirketin üreteceği sürücüsüz araçları 5 yıl içinde kullanmaya başlayabileceğimizi müjdelemişti.³⁶

Tesla ve SpaceX gibi şirketlerin kurucusu ve yöneticisi olan Elon Musk, 2016 yılının Ekim ayında, en geç 2017'nin sonunda Tesla markalı sürücüsüz bir otomobille Los Angeles'tan New York'a gidilebileceği müjdesini vermişti.³⁷

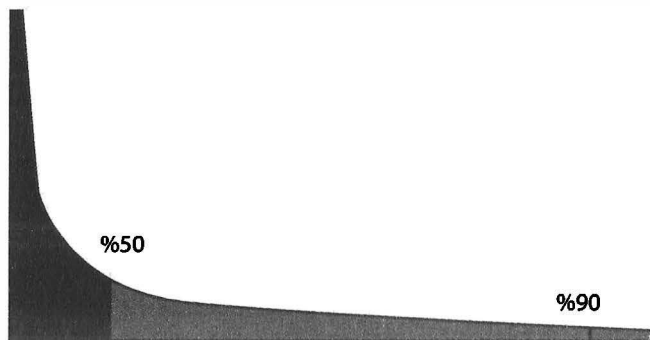
2020 yılının başlarında durum şu: Sürücüsüz otomobil teknolojilerini geliştirmeye çalışan pek çok şirket var, ama artık hiçbiri, bunları kullanmaya ne zaman başlayabileceğimizi söylemiyor...

YouTube'da yollardaki sürücüsüz otomobilleri gösteren bir sürü video var...

Evet... Ne var ki bunlar sadece belirlenmiş özel yollarda, gündüzleri ve güneşli günlerde kullanılabilir. Yağmurlu bir günün akşam saatlerinde İstanbul trafiğine sokulmaları henüz kesinlikle mümkün değil.

Burada bir "uzun kuyruk" sorunuyla karşı karşıyayız.³⁸ Otomobil kullanımı sırasında en sık karşılaşılan durumlar için çözümler geliştirmek görece kolaydır. Bilgisayarlara, kırmızı ışıktaki durmaları, hız

sınırlarına uymaları, önlerindeki araçla aralarındaki takip mesafesini korumaları, şerit değiştirmeden önce sinyal vermeleri vb. pek çok şey öğretilir. Bunlar, karşılaşılacak olan tüm durumların örneğin yüzde 50'sini oluşturabilir. Yolda yapılan onarımlar nedeniyle geçici olarak yerleştirilen trafik levhaları ve işçilerin elle yaptığı uyarılar, yola meyve suyu kutusu ya da pet şişe gibi nesnelerin düşmüş olması, yolun belirli bölümlerindeki şerit çizgilerinin silinmiş olması vb. pek çok durum eklendiğinde, karşılaşılacak olan tüm durumların belki yüzde 90'ı hesaba katılmış olabilir. Ama geriye, karşılaşılma olasılıkları giderek azalan ve bu nedenle de önceden hesaba katılmaları giderek zorlaşan durumlar kalır (Resim 22). Bunların sayısı da sonsuza yakındır.



Resim 22: "Uzun kuyruk" sorunu

İnsanlar, sağduyuları sayesinde, çok seyrek olarak karşılaşılacak durumların çoğunda doğru kararları ala-

biliyor. Yapay zekâ henüz bunu başarmanın çok uzağında.

Metro trenleri sürücüsüz olarak çalıştırılabiliyor, çünkü bu trenler insanların giremediği tünellerdeki rayların üzerinde hareket ediyor. Nitekim, sürücüsüz otomobilleri insanların giremeyeceği özel yollarda kullanma seçeneği üzerinde de duruluyor. Her tür yolda ve her tür durumda kullanılabilen sürücüsüz otomobillerle yakın gelecekte karşılaşmamızsa pek olası görünmüyor. Özellikle de bu otomobiller geçtiğimiz yıllarda ölümlü kazalara yol açmışken...

Makinelere sağduyu eklemek tümüyle olanaksız bir şey mi?

Uzun vadeden söz edecek olursak, bilmiyorum. Ama bugün elimizde bulunan araçlarla olanaksız.

Bazı yapay zekâ uzmanları, insan beyninin iç işleyişini daha iyi anladıkça daha iyi algoritmalar geliştirebileceğimize inanıyor.

Beyin hakkındaki bilgilerimiz durmadan artsa da, insan beyninin tam olarak nasıl çalıştığı hakkında hâlâ çok az şey biliyoruz. İnsan beyni bir yana, köpek beyninin ya da kedi beyninin nasıl işlediğini de çözebilmiş değiliz. Arıların kolektif zekâsının nasıl oluştuğu, bilmediğimiz bir başka şey. Arılar, dünyayı bilgisayarların anladığından çok daha iyi bir şekilde anlıyor.³⁹

Şu andaki iddialı projelerden biri, *C. elegans* adlı, 1 milimetre uzunluğundaki, hermafrodit türünde 302'si sinir hücresi (nöron) ve 95'i kas hücresi olmak üzere

toplam 959 hücresi bulunan solucanın sinir ve kas hücrelerinin modelini oluşturmayı hedefliyor.⁴⁰

C. elegans, ilişkilendirme yoluyla öğrenme yeteneğine sahip: Daha önce hangi tatların, kokuların ya da sıcaklık düzeylerinin yiyeceklerle (bakterilerle) ve hangilerinin hoş olmayan moleküllerle ilişkili olduğunu hatırlayabiliyor ve yolunu belirlemek için bu bilgileri kullanıyor.⁴¹

2011 yılında başlatılan uluslararası OpenWorm projesi henüz sonuçlandırılabilmiş değil. İnsan beyinde ise, farklı tahminlere göre 86 ile 100 milyar arasında sinir hücresi bulunuyor. Bu hücreler arasındaki bağlantıların toplam sayısı 100 trilyon ile 1000 trilyon arasında. Ve insan beyni sadece sinir hücrelerinden oluşmadığı gibi, insan bedeniyle sinir sistemi arasında da karmaşık bağlantılar var.

Bir tahmine göre, yapay sinir ağlarındaki sinir hücrelerinin sayısı, yeni teknolojiler durumu değiştirmese, insan beynindeki sinir hücrelerinin sayısına ancak 2050'li yıllarda ulaşabilir.⁴² Ama aynı sayıda yapay sinir hücresinin varlığı, tek başına, beynin işleyişini kopyalamamıza yetmez.

İnsan beyninin eksiksiz bir işleyiş modelini belki de hiçbir zaman çıkaramayacağız. Tabii bu durum, yapay zekânın gelişmesinin önündeki mutlak bir engel değil. Belki de, insan zekâsına daha yakın bir yapay zekâyı ulaşmanın yolu, insan zekâsını taklit etmeye çalışmaktan değil, bugünkülerden farklı yöntemlerin bulunma-

sından geçecek. Ama bu yöntemlerin neye benzeyebileceği henüz bilinmiyor.

Diğer yandan, yapay zekânın insan zekâsına göre farklı niteliklere sahip olduğunu kabul etmek ve insanlara daha fazla hizmet etmesinin yollarını aramak daha verimli bir yaklaşım olabilir. Ne var ki, günümüzün dünyasında bunun önünde bazı engeller var...

4. Yapay Zekâ ve İnsanlar

Stephen Hawking, yapay zekânın insanlığın sonunu getirebileceğini söylemişti. Yanılıyor muydu? Benzer kaygıları paylaşan başkaları da yok mu?

Öncelikle, yapay zekânın görece yakın bir gelecekte (örneğin 15-25 yıl içinde) insan zekâsını yakalayabileceğine inananların pek çoğunun paylaştığı bir yanılgı üzerinde duralım.

Transistörlerden söz etmiştik. Bunlarla ilgili bir yasa var. Teorik olmayan, sadece gözlemlere dayanan Moore Yasasına göre, tümleşik devrelerdeki, yani bilgisayar çiplerindeki transistörlerin sayısı her iki yılda bir iki katına çıkıyor.

Bu “yasa”nın geçerliliğini bugüne kadar büyük ölçüde koruması sayesinde, 1970’li yılların başında en güçlü bilgisayarların işlemcilerindeki transistörlerin sayısı binler düzeyindeyken, günümüzün akıllı telefonlarının işlemcilerinde milyarlarca transistör bulunabiliyor.

Buna “üstel büyüme” deniyor. Farklı versiyonları bulunan bir öyküye göre, bir hükümdar, satrancı icat eden kişiye, “dile benden ne dilersen” demiş. Mucit, bir satranç tahtasındaki 64 kareyi, ilkinde 1, ikincisine 2, üçüncüsüne 4 buğday tanesi, yani her seferinde bir öncekinin iki katı kadar buğday tanesi koyacak şekilde doldurarak kendisine vermesini istemiş. Hükümdar, ilk anda, bunun önemsiz bir istek olduğunu sanmış. Ama

64 karenin bu şekilde doldurulması için gerekli olan buğday tanelerinin sayısı 18.446.744.073.709.551.615. Bu da dünyanın bugünkü yıllık toplam buğday üretiminin yaklaşık 1600 katına denk geliyor.⁴³

Transistörlerin sayısındaki üstel büyümenin teknolojinin her alanına yansıdığını ya da yakın gelecekte yansıyacağını düşünen Ray Kurzweil, insanlığın aşılacağını ve aşılması gerektiğini savunan insanötesicilik (*transhumanism*) akımının temsilcileri arasında yer alıyor. Google'ın üst düzey yöneticilerinden biri olan Kurzweil'a göre, yapay zekâ 2029 yılında insan zekâsıyla eşitlenecek. 2030'lu yılların ortalarında düşünme sürecimiz biyolojik düşünme ile biyolojik olmayan düşünmenin bir bileşimine dönüşecek. Biyolojik olmayan kısım 2040'lı yıllarda biyolojik kısma üstünlük kuracak. Beyinlerimizdeki anılara, becerilere ve kişilik özelliklerine erişebileceğiz ve bunları internete yükleyebileceğiz. Böylece ölümsüzlüğe ulaşacağız.⁴⁴

The Matrix filmindeki gibi mi?

Evet. Kurzweil, Moore Yasası sayesinde insanlar ile makinelerin bu filmdekine benzer bir şekilde, ama insanların köleleştirilmesine yol açmadan birleştirilebileceğini ve birleştirilmeleri gerektiğini açıkça savunuyor.⁴⁵

Ne var ki, teknolojinin her alanında benzer bir hızla ilerleme sağlanabileceği iddiası hayli tartışmalı. Örneğin, daha yakın tarihli Eroom Yasasına göre, ABD'de, (enflasyondan arındırılmış) 1 milyar dolarlık Ar-Ge

harcaması başına yetkili kurumdan onay alarak kullanıma sokulabilen ilaçların sayısı, 1950’li yıllardan 2010’lu yıllara kadar, her 9 yılda bir yarı yarıya azalmış.⁴⁶ Oysa aynı dönemde teknolojik gelişmelerin yeni ilaçların geliştirilmesini kolaylaştıracağı hakkında sayısız haber yapılmıştı. Bu arada, 2019 yılının Nisan ayında, IBM Watson’ın ilaç geliştirme programının satışının durdurulduğu öğrenildi. IBM, başka alanlara ağırlık vereceğini açıkladı.⁴⁷

Dilimizi anlamayı bile başaramayan bir yapay zekâ, insan zekâsını ancak bilimkurgu filmlerinde yakalayabilir. Ray Kurzweil, 2018 yılının Nisan ayında, Google’ın *Talk to Books* (Kitaplarla Konuş) projesini duyurmuştu.⁴⁸ 100 bini aşkın kitabın her cümlesinden yararlanabilen *Talk to Books*, doğal dili (İngilizceyi) kullanarak sorduğumuz sorulara kitap cümleleriyle karşılık veriyor. Kurzweil’a göre, makine öğrenmesine yaslanan bu projenin ayırt edici özelliği, cümlelerin “anlambilimsel anlam”larını temel almasıydı.

Amerikalı bilimciler Gary Marcus ile Ernest Davis’in daha önce tartıştığı üzere, *Talk to Books*, sorulan sorulardan aslında hiçbir şey anlamıyor. Sadece, içerdikleri sözcüklerden yararlanarak, onlara bazı sayısal değerler veriyor. Sonra da, sayısal değerlerini hesaplamış olduğu kitap cümlelerinden en benzer olanlarını gösteriyor.⁴⁹

Talks to Books’un aradan geçen sürede ne kadar ilerleme kaydettiğini görmek için ona şu basit soru-

ları sordum: “Derin öğrenmenin yetersizlikleri/sınırlılıkları/dezavantajları neler?” (“*What are the shortcomings/limitations/disadvantages of deep learning?*”) Çok benzer anlamlı üç ayrı soru sormamın nedeni sözcük seçiminin önemli bir rol oynayıp oynamayacağını görmektir. İlk sorunun 20 cevabından sadece 5 tanesi derin öğrenmenin sorunlarıyla ilgiliydi. Buna karşılık 7 cevap derin öğrenmenin olumlu yanları hakkındaydı! İkinci soru için durum biraz daha iyiydi: 6’ya 7. Üçüncü soruya gelince: 7’ye 11. Kısacası, *Talk to Books*, sorularımızı ve kitaplardaki cümlelerin ne anlattığını hâlâ anlamıyor.

Burada, istatistiksel yöntemlerle çözülemeyecek olan bir sorunla karşı karşıyayız. İşlemcilerin daha da güçlenmesi, yani hesaplamaların daha da hızlanması, tek başına, yapay zekânın doğal dilleri ya da genel olarak dünyayı anlamasını sağlayamaz.

Yapay zekânın, doğal dilleri anlamadan insanlığın sonunu getirmesi gerçekten olanaksız mı? Facebook’un sohbet robotları kendi aralarında yeni bir dil geliştirmiş.

Hatta, medyaya göre, söz konusu robotlar, sadece kendilerinin anlayabildiği bir dille konuşmaya başladıkları için kapatılmıştı.⁵⁰

Yazışmalarına bir bakalım:

Bob: i can i i everything else

Alice: balls have zero to me to me to me to me to me to me to me to me to me to

Bob: you i everything else

Alice: balls have a ball to me to me to me to me to me to me to me to me

Bob: i i can i i i everything else

Alice: balls have a ball to me to me to me to me to me to me to me to me

Bob: i

Alice: balls have zero to me to me to me to me to me to me to me to me to me to

Bob: you i i i i i everything else

Alice: balls have 0 to me to me to me to me to me to me to me to me to me to

Bob: you i i i everything else

Alice: balls have zero to me to me to me to me to me to me to me to me to me to

Yani, Bob, “ben yapabilirim ben ben başka her şey” ve “sen ben başka her şey” benzeri şeyler, Alice ise, “benim için benim için benim için benim için benim için benim için benim için benim için için topların sıfır var” benzeri şeyler söylemişti. Alice ayrıca, “topların bir topu var” da demişti.

Bu yazışma, gerçekte, yapay sinir ağlarının bazı durumlarda ne kadar anlamsız sonuçlar üretebileceğini gösteriyor.

Sohbet robotlarıyla ilgili başka öyküler de var. Örneğin, Microsoft, 23 Mart 2016’da, sohbet robotu Tay için bir Twitter hesabı açmıştı. Ama 24 saat geçmeden

bu hesap kapatıldı, çünkü, kendisiyle iletişim kuran Twitter kullanıcıların yazdıklarından “öğrenen” Tay, ırkçılık, cinsiyetçilik ve Yahudi düşmanlığı içeren cümleler kurmaya başlamıştı.⁵¹

Sorun, robotların okuduklarını da yazdıklarını da anlayamamaları. Yeni robotların ELIZA’dan farkı ise, neyi neden yaptıklarını bazen programcılarının da tam olarak anlayamaması...

The Economist dergisi, 2020 yılıyla ilgili beklentilerin ele alındığı özel sayısında, internette bulunan toplam 40 GB’lık metinlerle ve gözetimsiz öğrenme (*unsupervised learning*) yöntemiyle eğitilen GPT-2 adlı sistemle yapılan bir röportaja yer verdi. Derginin açıklamasına göre, “dil modeli” diye anılan sisteme 2020’yle ilgili beklentileri sorulmuş ve verdiği cevaplar editörlük (düzeltme) işlemi yapılmadan yayımlanmıştı. Ve bu cevaplar, pek ufuk açıcı olmasalar bile, çok kötü de değildi.⁵²

Gözetimsiz öğrenme nasıl oluyor?

Bu yöntemde, etiketler de ödüller de kullanılmıyor. Algoritmaların verilerden sonuçlar çıkarması isteniyor. GPT-2’den istenen şey, okuduğu sözcükleri ya da cümleleri tamamlaması. Bir başka deyişle, bunların arkasından gelebilecek olan sözcükleri ya da cümleleri tahmin etmesi isteniyor. Yani, aslına bakılırsa, daha önce yazılmış olan sözcükler ve cümleler, GPT-2 tarafından etiketler olarak kullanılıyor. “Türkiye’nin başkenti” diye başlayıp “Ankara’dır” diye biten cümlelerin çok

sayıda olması sayesinde, GPT-2, “Türkiye’nin başkenti” sözcüklerini “Ankara’dır” diye tamamlamayı “öğreniyor”. Bu nedenle bazı yapay zekâ uzmanları GPT-2 türü sistemler için “özgözetimli öğrenme” (*self-supervised learning*) kavramını tercih ediyor.⁵³

The Economist’in yayımladığı röportajın nasıl yapıldığı Twitter’da sorgulanınca, derginin ilgili editörü Tom Standage bir açıklama yapmak zorunda kaldı. GPT-2’nin verdiği cevaplar, “oldukları gibi” yayımlanmamıştı. Editör, her bir soru için beş ayrı cevap almış ve bunlardan birini seçmişti. Standage’ın sözleriyle, “GPT-2 gerçekte hiçbir şeyi anlamıyor”du.⁵⁴

GPT-2’yi ben de denedim. Örneğin: “İki elmanız ve üç portakalınız varsa ve elmalarınızdan birini yerseniz o zaman elinizde sadece ...” diye yazınca, GPT-2 farklı denemelerde şöyle devam etti:

(1) iki portakal kalır.

(2) iki portakal kalır. Diğer üç portakal size kalır.

(3) bir limon kalır.

(4) iki portakal kalır. Ama üç elmanın tümüne sahip olursanız, üç portakalın tümüne sahip olursunuz.

(5) bir elma kalır. Ama, diğer iki elmayı yerseniz, o zaman üç portakalınız olur.

(6) üç portakal kalır, çünkü iki elmanızı yemiş olursunuz.

(7) portakallarınızdan biri kalır.” Buna “İki Kefaret” ilkesi denir ve Yeni Ahit’te bulunur.⁵⁵

Görüldüğü gibi GPT-2'yi yeterince sorgularsanız kısmen doğru cevaplar alabiliyorsunuz; ama çok da üstüne gitmeyin!

Stephen Hawking'e dönecek olursak... Hawking, 2015 yılının Ocak ayında, çok sayıda yapay zekâ uzmanıyla birlikte, yapay zekâ alanındaki araştırmalarla ilgili bir açık mektubun imzacısı olmuştu.⁵⁶ Bu mektubun ilişigindeki belgede, bir “zekâ patlaması”nın yaşanması ve insanların kontrol edemeyeceği bir yapay “süper zekâ”nın ortaya çıkması tehlikesinden de söz ediliyor ve bu tehlikeye karşı alınabilecek olan önlemlerin araştırılması isteniyordu.⁵⁷ Açık mektubun diğer imzacıları arasında akademisyenlerin yanı sıra Google'ın, Facebook'un, Microsoft'un yöneticileri ve Elon Musk da vardı.

Süper zekâ tartışmaları bazı akademisyenler için ciddi bir gelir kaynağı. Örneğin, açık mektubun ilişigindeki belgenin yazarları arasında yer alan Max Tegmark, Elon Musk'ın para verdiği ve yapay zekânın insanlığın sonunu getirmesi riskiyle mücadele eden Yaşamın Geleceği Enstitüsü'nün (*Future of Life Institute*) kurucuları arasında yer alıyor. Tegmark, yapay zekânın dünyayı ele geçirebileceğini savunduğu 2017 tarihli kitabında, aynı yıl düzenlenen bir yapay zekâ konferansında alınan kararlara yer veriyor. Bu kararlar, yapay zekânın insanlık yararına kullanılmasına yönelik araştırmalara maddi destek sağlanmasını da içeriyor. Tegmark, “süper zekânın olanaksız olduğunu” ya da “yapay

zekânın insanlığın sonunu getirmesi riskini azaltmaya yönelik araştırmaların tümüyle gereksiz olduğunu” savunmanın söz konusu kararların çiğnenmesi anlamına geleceğini özellikle vurguluyor!⁵⁸

Yapay zekâ teknolojilerine en fazla yatırım yapan şirketlerse, insanlığa hizmet ettiklerini düşündürmek için, bu tür araştırmalara maddi destek sağlıyor.

Bunu yapmaları iyi bir şey değil mi?

Bu konuda, yapay zekâ uzmanlarından Andrew Ng’nin şu görüşüne katılıyorum: Bugün için, süper zekânın insanlığı yok etmesinden korkmakla, Mars’ta aşırı nüfus sorunuyla karşılaşmaktan korkmak arasında pek bir fark bulunmuyor.⁵⁹

İnsanlığın çok yakın bir gelecekte yapay genel zekâ tarafından tümüyle önemsizleştirilme ve hatta ortadan kaldırılma tehlikesiyle karşı karşıya bulunduğu kabul edilirse, bugünkü üretim ilişkilerini ve toplumsal güç dengelerini sorgulamanın pek bir anlamı kalmaz. Ne de olsa, üretim ilişkileri ve toplumsal güç dengeleri bazı ülkelerde değiştirilse bile, diğer ülkeler (örneğin yapay zekâ alanına büyük yatırımlar yapan Çin) yapay genel zekâyı ortaya çıkarabilir. Özellikle ABD’deki teknoloji şirketlerinin sahipleri ve üst düzey yöneticileri, bizi, yapay zekânın insanlık yararına kullanılmasını sadece onların sağlayabileceğine inandırmaya çalışıyor. Kendileri üzerindeki en iyi denetimin yine kendileri tarafından kurulabileceğini savunuyorlar.

Oysa, teknoloji şirketlerinin insanlığa bugün vermekte oldukları ciddi zararlar ve yarattıkları ciddi tehlikeler var. Çünkü, insanlığın pek çok sorununun çözümüne katkıda bulunabilecek olan yapay zekâ, teknoloji şirketlerinin elinde, sadece, bu şirketlerin kârlarını artırmak için kullanılıyor.

Örneğin, Google ve Facebook gibi şirketlerde çalışan mühendislerin en önemli görevleri, bu şirketlerin ürünlerine bağımlı hale gelmemizi ve gösterdikleri reklamlara mümkün olduğunca sık tıklamamızı sağlamak. Android işletim sistemini kullanan telefonlardaki ücretsiz oyunlar hem kolayca bağımlılık yaratabiliyor hem de reklamlara tıklama olasılığımızı artırıyor. YouTube’da, bir video izledikten sonra siteden ayrılmamamız için, otomatik oynatma “hizmet”i sunuluyor. Eğer Google’da oturum açtığımız bir tarayıcı kullanıyorsak, yeni video seçilirken geçmişteki tercihlerimiz ve ilgi alanlarımız hesaba katılıyor. Facebook, arkadaşlarımızın gönderilerini ve üyesi olduğumuz gruplardaki gönderileri sıralarken, geçmişte kimlerin hangi türdeki gönderileriyle ne kadar ilgilendiğimizi göz önünde bulunduruyor ve ilgimizi çekeceğini düşündüğü gönderileri üst sıralara çıkarıyor. Facebook’taki “ifade bırakma” butonları ve bildirimler, bu siteye daha sık ve daha uzun süreler boyunca bakmamızı sağlıyor. Hem Google hem de Facebook, bize gösterecekleri reklamları seçerken, e-mektuplarımızda ve mesajlarımızda kullandığımız

sözcüklere, arama geçmişimize, ziyaret ettiğimiz sitelere, tıkladığımız gönderilere bakıyor.

Bu arada üzerimizde deneyler yapıyorlar. Bunların büyük çoğunluğundan haberdar olamıyoruz.

Ama 2012 yılında yapılan bir deneyin sonuçları yayımlandı. 689.003 Facebook kullanıcısıyla yapılan deneyde, bazı kullanıcıların karşısına daha çok negatif içerikli gönderi, bazılarının karşısına da daha çok pozitif içerikli gönderi çıkarılmış. Ardından, ilgili kullanıcıların kendi paylaşımlarına bakılmış. Daha çok negatif içerikli gönderiye maruz kalan kullanıcıların negatif içerikli paylaşımda bulunma olasılığının, daha çok pozitif içerikli gönderiye maruz kalan kullanıcıların da pozitif içerikli paylaşımda bulunma olasılığının arttığı saptanmış.⁶⁰

Üniversitelerde, bu tür bir deney yapabilmek için, deneklerin yazılı onaylarının alınması gerekir. Teknoloji şirketleri ise insan davranışlarını değiştirmeye yönelik deneyleri kimseye sormadan yapabiliyor.

Bu deney nedeniyle Facebook'a ceza verilmedi mi?

Hayır... Araştırma sonuçlarını içeren makalenin yayımlanması tartışmalara yol açınca, yayıncı kuruluşun baş editörü bir açıklama yapmak zorunda kalmış. Bir Facebook hesabı oluştururken onay vermek zorunda olduğumuz "Veri Kullanma Politikası" belgesi, şirketin bu tür deneyler yapmasına izin veriyormuş.⁶¹

İnternetteki pek çok uygulamanın ücretsiz olması, büyük ölçüde, kişisel verilerimizin şirketler için çok değerli olmasından kaynaklanıyor.

Örneğin, Facebook, tüm gönderilerinizi, tüm “beğeni”lerinizi, paylaştığınız tüm konum bilgilerinizi, tüm arkadaşlarınızı ve takipçilerinizi, tüm özel mesajlarınızı, tüm aramalarınızı, tüm tıklamalarınızı, bir gönderiye ne kadar zaman ayırdığınızı, farenizin hareketlerini, yazıp da göndermediklerinizi, Facebook’tan ayrıldıktan sonra ziyaret ettiğiniz siteleri kaydediyor. Cep telefonuna Facebook uygulamasını indirir ve izin kısıtlamasına gitmezseniz, rehberinizi, aramalarınızı, internete nereden bağlandığınızı ve başka pek çok verinizi inceliyor. Başka şirketlerden sizinle ilgili verileri satın alıyor (gelir düzeyiniz, ev ya da araba sahibi olup olmadığınız, kredi notunuz vb.). Eğlence için çözdüğünüz kişilik testleri, karakter özelliklerinizin kaydedilmesini sağlıyor.⁶² Diliniz, cinsiyetiniz, siyasal görüşleriniz, cinsel yönelimleriniz, eğitim düzeyiniz, ilişki durumunuz, kullandığınız işletim sistemi ve başka sayısız veri Facebook için değerli. Bir internet sitesine girerken yeni bir hesap açmak yerine Facebook kimliğinizi kullanırsanız, hem o site hem de Facebook sizin hakkınızda daha fazla veriye sahip oluyor.

Bütün bu verilerin toplanması ve işe yarar hale getirilmesi için yapay zekâ kullanılıyor elbette.

Çoğu insanın varlığından bile haberdar olmadığı bir sektör var: Veri simsarlığı. Bu sektörün en büyük

şirketlerinden biri olan Acxiom, 2018 yılında, 62 ülkedeki 2,5 milyar insanın toplamda 10 bin başlığa ulaşabilen kişisel verilerine sahipti. Bu şirket, 2,3 milyar dolar karşılığında, dünyanın en büyük reklam ajanslarından birine satılmıştı.⁶³

Bu yılın başlarında, parasız antivirüs programları geliştiren ve 400 milyondan fazla kullanıcısı bulunan Avast'ın, kullanıcıların tarayıcı geçmişlerini başka şirketlere sattığı ortaya çıktı.⁶⁴ Bir başka antivirüs şirketi olan AVG'yi 2016 yılında 1,4 milyar dolara satın alan⁶⁵ Avast'ın yetkilileri, kullanıcı bilgilerini anonimleştirdiklerini, yani kimlik bilgilerini paylaşmadıklarını söyledi. Oysa tarayıcınızın adres satırında görülen her şeyi satıyorlardı ve bunları kullanarak kimlik bilgilerinizi açığa çıkarmak hiç zor değil.

Kişisel verilerin bu kadar değerli olması nedeniyle, ziyaret ettiğiniz İnternet sitelerinin büyük çoğunluğu, hakkınızda mümkün olan en fazla bilgiyi toplamak için tarayıcınıza çerezler yerleştiriyor ve üçüncü taraflara ait programcıklar (*script*'ler) içeriyor. Son dönemde çerezler konusunda (kullanıcıları rahatsız etmekten başka pek bir işe yaramayan) uyarılarla karşılaşılrsa bile, üçüncü taraflara ait programcıklarının varlığı (ve bazı eklentiler yardımıyla çalışmalarının engellenebileceği) pek az kimse tarafından biliniyor.

“Akıllı” cihazlarınızdaki pek çok uygulama kişisel verilerinizi toplamak için kullanılıyor.

Eğer suç işlemiyorsak, kişisel verilerimizi toplamalarının ne sakıncası var? Reklamlara tıklayıp tıklamamak kendi tercihimiz değil mi?

Tek sorun reklamlar olmasa da onlarla devam edelim.

Bir veri simsarı, kumara düşkün olan ve 55 yaşın üstündeki 500 bin kişinin kimlik bilgilerini tanesi 8,5 sente dolandırıcılara satmış.⁶⁶ Dolandırıcılar bu kişilere cazip tekliflerde bulunarak kendi platformlarında kumar oynatmaya çalışmış olmalı...

Güzellik ürünü reklamları, kadınlara, kendilerini en az çekici hissettikleri saatlerde gösteriliyormuş.⁶⁷

Yani zaaflarımızdan da yararlanmaya çalışıyorlar. Örneğin, kendisi ya da en yakını kanser hastalığına yakalanmış olan birine, aslında hiçbir işe yaramayan, ama başka çaresi kalmamış insanların ilgilenebileceği alternatif tedavi reklamları gösterebiliyorlar. Maddi sıkıntı içindeki kişilere tefecilerin ya da dolandırıcıların reklamlarını gösterebiliyorlar.

Yapay zekâ algoritmaları açısından önemli olan, kişisel verilerinizle, belirli bir anda tıklayabileceğiniz reklamlar arasında bağlantı kurmak. Bunun size zarar vermesi olasılığı algoritmaları ve onları kullananları ilgilendirmiyor. Yine örneğin, 16 yaşından küçük çocuklara aşırı şekerli, tuzlu ve yağlı yiyeceklerin reklamları gösterilebiliyor.⁶⁸

Başka alanlara gelirse... Şirketler, kişisel verilerinizi, işe alım süreçlerinde de giderek daha fazla kul-

lanıyor. Sosyal medyadaki paylaşımlarınıza bakıyorlar. Hakkınızda yazılanları da inceliyorlar. Geçirdiğiniz ağır bir hastalık, bir yakınınızın ölümü ya da tanık olduğunuz bir vahşet nedeniyle travma yaşarken yaptığınız paylaşımları kendi hesaplarınızdan silebilseniz bile, bunlara verilen tepkileri silmeniz hiç kolay değil. Diğer taraftan, internet geçmişinizin önemli bir bölümü satın alınabilir durumda. Ve şirketler, riskleri en aza indirmek için, geçmişlerinde herhangi bir sorun bulunan kişileri işe almamayı tercih edebiliyor. Bu da aynı kişilerin yeni travmalar yaşama olasılığını yükseltiyor...

O zaman sosyal medyayı hiç kullanmamak mı gerekiyor?

Dikkatli kullanmak gerektiği kesin. Ama sosyal medyayı hiç kullanmamanız sizi yüksek riskli sayılan insanlar sınıfına da sokabilir!

Teknolojik gelişmeler, insanlar hakkındaki sayısallaştırılabilir ve erişilebilir verileri durmadan artırıyor. Örneğin, akıllı telefonlar ve akıllı bileklikler, gün içinde attığımız adımları sayabiliyor. Bu gelişme, günde 10 bin adım atmanın gerekli olduğu düşüncesinin yaygınlaşmasına katkıda bulundu. Aslında, bu sayı, bilimsel araştırmalar sonucunda bulunmamıştı. Japonya'da, 1964 Tokyo Olimpiyatları'nın yarattığı ilgiyi paraya çevirmek isteyen bir şirketin geliştirdiği giyilebilir adım-sayar için keyfi olarak seçilen bir sayı söz konusuydu.⁶⁹ Ama adımların kolaylıkla sayılabilir ve kayıt altına

alınabilir duruma gelmesi, bunlar için farklı kullanım alanları yarattı. Bazı özel sağlık sigortası şirketleri, akıllı bileklik verilerini paylaşan ve düzenli olarak çok adım atan müşterilerinin sigorta primlerinde indirim yapıyor.⁷⁰ Bunun gerçek anlamı, sağlık verilerini paylaşmayan veya yeterince adım atmayan müşterilerin cezalandırılması.

Dünya genelindeki eğilim, devletlere bağlı ve herkesi kapsayan sosyal güvenlik kurumları aracılığıyla sunulan sağlık hizmetlerinin kapsamlarının giderek daraltılması yönünde. Özel sağlık sigortası şirketleri ise, hastalanma riski yüksek olanların primlerini yükseltmek (ya da sigortalılığa hiç kabul etmemek) için giderek daha fazla veri kullanacak. Fiziksel etkinliklerimiz, uyku düzenimiz, kan değerlerimiz, vücut-kilo endeksimiz, yediklerimiz ve içtiklerimiz hesaba katılacak. Genomlarımızın dizilenmesi kolaylaştıkça ve ucuzladıkça, genetik hastalıklara yatkın sayılanlarımız daha fazla prim ödeyecek. Giderek, orta gelir düzeyindeki insanlar bile, tedavisi masraflı bir hastalığa yakalanma riskleri yüksek çıkarsa, hastalanmaları durumunda gerekli tedavileri alamayacak.

Tabii hastalanma riskleri, sadece özel sağlık sigortası şirketlerini ilgilendirmiyor. Diğer özel şirketler, hastalığa yakalanma riski yüksek kişileri işe almak istemeyecek. Dolayısıyla, yoksul bir ailenin hastalanma riski yüksek bir çocuğunun çok çalışarak olası tedavi masraflarını karşılaması daha da zorlaşacak.

Bu tür sorunların yaşanması yasalarla engellenemez mi? Ayrıca Avrupa Birliği kişisel verilerin gizliliğiyle ilgili yasalar çıkarmadı mı?

Evet, Avrupa Birliği, ABD'yle karşılaştırıldığında, kişisel verilerin gizliliği konusunda daha katı düzenlemelere sahip. Ne var ki, bu verilerin nasıl kullanıldığını bilemediğimiz sürece, düzenlemelerin ihlal edilip edilmediğini de bilemeyiz. Ve özel şirketler çoğu algoritmalarını gizli tuttuğundan, hangi verileri ne şekilde kullandıklarını bilemiyoruz.

Diğer yandan, belirli türdeki kişisel verilerin kullanılması yasaklandığında ve bu yasağa uyulduğunda bile, algoritmalar, bunların yerini tutacak olan “vekil” verileri kullanabilir. Örneğin, bir kişinin Kürt ya da Alevi olup olmadığını gerçekten de dikkate almayan bir algoritma, aynı kişinin Kürtlerin ya da Alevilerin yoğun olarak yaşadığı bir mahallede ya da köyde oturmasını dikkate alabilir.

Tekrarlayacak olursak, temel sorun, teknolojik gelişmelerin doğrultusunun, insanların büyük çoğunluğunun ihtiyaçlarına göre değil, tesadüflerin yardımıyla belirli alanlarda tekel oluşturabilen teknoloji şirketlerinin sahiplerinin ve üst düzey yöneticilerinin çıkarlarına göre belirlenmesi.

Ne gibi tesadüfler?

Bill Gates tarafından kurulan Microsoft'un 1981 yılında satın aldığı ve adını MS-DOS olarak değiştirdiği işletim sistemi, kişisel bilgisayar (PC) üreticileriyle

yapılan lisans sözleşmeleri sayesinde, en yaygın şekilde kullanılan işletim sistemi haline gelmişti. Her bir bilgisayar için üreticilerden para alan Microsoft, hem MS-DOS'un hem de sonradan geliştirdiği Windows işletim sisteminin kaynak kodlarını gizli tutarak, bu işletim sistemlerini rakip programlardan daha verimli bir şekilde kullanan programlar (Word, Excel vb.) geliştirebilmişti. Microsoft, 86-DOS adlı işletim sistemini satın alamasaydı ya da kişisel bilgisayar üreticileri farklı bir tercih yapsaydı, Bill Gates bugün dünyanın en zengin kişilerinden biri olmayabilirdi.

Kaynak kodu ne demek?

Daha önce bilgisayar programlama dillerine değinmiştik. Bu dillerden biriyle yazılan şeylere kaynak kodu deniyor. Bir bilgisayar programının henüz makine diline çevrilmemiş olan versiyonu, onun kaynak kodudur. Kaynak kodu açık, yani şifrelenmemiş olursa, başka programcılar bu programın özellikleri hakkında daha fazla bilgi sahibi olur ve onun üzerinde değişiklik yapabilir. Dünyada, yazılımların açık kaynaklı olması gerektiğini savunan ve bu tür yazılımlar geliştiren güçlü bir hareket var. Örneğin, Linux temelli işletim sistemlerinin ve bunlarda çalışan programların pek çoğu açık kaynaklı ve dileyen herkes tarafından değiştirilebiliyorlar. Bu arada, Özgür Yazılım Vakfı (*Free Software Foundation*), Microsoft'un artık desteklemediği Windows 7 işletim sisteminin kaynak kodunun yayımlanması için 2020 yılının Ocak ayında bir kampanya

başlattı ve bu işletim sistemi için destek hizmeti sunma vaadinde bulundu.⁷¹ Elinizdeki kitabın hazırlıkları sürerken Microsoft henüz bu konuda bir açıklama yapmamıştı...

Facebook'un kurucusu ve bugünkü CEO'su Mark Zuckerberg, 2006 yılında, daha 22 yaşındayken dolar milyoneri, 23 yaşındaysa dolar milyarder olmuştur.⁷² Bunu, çok çalışmasına değil, üniversite öğrencisiyken kurduğu sosyal ağın hızla büyümesine borçluydu. O dönemde başka sosyal ağlar da vardı. Ama Facebook, belirli bir kullanıcı sayısı eşiğini aştıktan sonra, en popüler ağın yeni kullanıcıları daha kolay çekmesi nedeniyle çok büyük bir hızla büyüeyebildi. Bu sayede Instagram'ı ve WhatsApp'ı satın alabilecek güce ulaştı ve bir tekele dönüştü. Bugün Zuckerberg'in yerinde kolaylıkla bir başkası da olabilirdi.

Bu arada, Facebook'un başarıya ulaşmasının en önemli nedenlerinden biri, PayPal'in kurucularından Peter Thiel'in 2004 yılında bu şirkete 500 bin dolarlık bir yatırım yapmasıydı. Thiel'in beklentisi, Facebook'un bir tekele dönüşmesini sağlayarak büyük bir kazanç elde etmektir. İlk aşamada önemli olan tek şey, kullanıcı sayısının hızla artmasıydı. Facebook 2009'a kadar sürekli zarar etti, ama bu arada 300 milyon kullanıcıya ulaştı.⁷³ Thiel, 2012 yılından itibaren Facebook hisseleri satarak yaklaşık 1,2 milyar dolar elde etti.⁷⁴ Bu da 2004 yılında yatırdığı paranın yaklaşık 2400 katıydı. Girişimci adayları için bir kitap da yazan Thiel'in en önemli

tavsiyelerinden biri, belirli bir alanda tekel oluşturma-
nın yolunu bulmalarıydı.⁷⁵

Büyük teknoloji şirketleri de, yatırımcıları kendile-
rine çekebilmek için, durmadan yeni ve iddialı projeler
üretiyor. Tabii bu projelerin çok büyük kârlar vaat et-
mesi gerekiyor. Bu nedenle, ne tür projelerin üretilece-
ği, insanların gereksinimlerine göre değil, kâr getirme
potansiyellerine göre belirleniyor.

ABD'deki büyük teknoloji şirketlerinin hisse se-
netleri, Şubat 2019 ile Şubat 2020 arasında yüzde 52
oranında değer kazanmıştı. Şubat 2020'de, Google'ın
ana şirketi Alphabet'in, Amazon'un, Apple'ın ve
Microsoft'un piyasa değerleri 1'er trilyon doların üze-
rinde, Facebook'un ki 620 milyar dolar civarındaydı. Bu
beş şirketin toplam piyasa değeri 5,6 trilyon dolardı.⁷⁶
Türkiye'nin 2019 yılındaki gayrisafi yurt içi hasılasının
yaklaşık 753 milyar dolar olduğunu ekleyeyim.⁷⁷

Bu trilyon dolarlar, ilgili şirketlerin gerçek değer-
lerini değil, yatırımcıların spekülatif beklentilerini
yansıtıyordu. Ama bu durum, teknolojik gelişmelerin
doğrultusunu belirleme gücünün bu şirketlerde olduğu
gerçeğini değiştirmiyordu.

Elon Musk, Mars'a 1 milyon insan götüreceğini
söylemişti. Bu da böyle bir proje mi?

Evet. Musk, yerine getiremediği vaatleriyle ünlü.
Ama her seferinde daha büyük iddialarla ortaya çıkıp
para toplamayı başarabiliyor.

Mars'ı yerleşime açmanın insanlığın öncelikli ihtiyaçları arasında yer alıp almadığı da, bu projenin başariya ulaşma olasılığının bulunup bulunmadığı da hayli tartışmalı konular. Ne var ki, insanlığın toplam servetinin nerelerde kullanılacağına sermaye sahipleri karar veriyor.

Tıpkı sürücüsüz otomobiller örneğinde olduğu gibi... Bir hesaplama göre, sürücüsüz otomobil teknolojileri üzerinde çalışan 30 şirket geçtiğimiz birkaç yılda en az 16 milyar dolarlık harcama yapmış. Bu paranın yarısını Google, General Motors ve Uber harcamış.⁷⁸

Ulaşım ihtiyaçları, toplu taşıma sistemlerini geliştirerek de karşılanabilir. Otomobil sayısı en aza indirilirken, toplu taşıma araçlarının sayıları ve bunlarda yolculuk etmenin konforu artırılabilir. Engelliler için özel çözümler üretilirken, engelli olmayanların daha fazla yürümesi ve bisiklete binmesi teşvik edilebilir. Böylece, taşıtlara ayrılan toplam kaynak miktarının ve çevre kirliliğinin yanı sıra, trafik kazalarında ölenlerin ve yaralananların sayıları da azaltılabilir.

Ama teknoloji şirketlerinin beklentileri bambaşka: Yaygın olarak kullanılabilecek ilk sürücüsüz otomobili üreten şirket, kendi otomobil kiralama platformunu kurarak bu alanda bir tekel oluşturabilir. Böyle bir platformda bulunan sürücüsüz otomobillerin sayısı ne kadar hızlı bir şekilde artırılabilirse, sonradan gelenlerin rekabet şansı o kadar azalır. Dahası, yola ilk çıkan şirket, sermayesini çok hızlı bir şekilde büyüteceğinden,

güç kazanma potansiyeli bulunan şirketleri satın alabilir ya da bir süre boyunca zarar etmeyi kabullenerek rakiplerini iflasa sürükleyebilir. Böylece bugünkülerden çok daha büyük bir trilyoner şirket yaratılabilir ve bu şirket çok farklı alanlarda yeni tekeller kurabilir...

Elde edilebilecek kârların olağanüstü büyük olması nedeniyle, sürücüsüz otomobil teknolojilerini geliştiren şirketler, mümkün olduğunca çok sayıda patent alırken (ve bunlardan bazılarını sadece rakiplerinin önünü kesmek için kullanırken), patentini almadıkları bilgileri ticari sır olarak saklıyor.

Teknoloji hep bu şekilde gelişmiyor mu? Patent alamayacak olsalar bu kadar yatırım yaparlarmıydı?

Evet, yapmazlardı. Ama teknolojinin hep bu şekilde geliştiği söylenemez.

Örneğin, bugünkü biçimiyle internetin varlığını, Tim Berners-Lee'nin, Dünya Çapında Ağ (*World Wide Web*) için patent almamış olmasına borçluyuz. Eğer alsaydı, günümüzün en zengin kişilerinden biri olabilirdi. Buna karşılık, elimizde, büyük olasılıkla, dünya çapındaki tek bir ağ yerine, Microsoft ve Apple türü şirketlerin çok sayıda özel ağı bulunurdu. Bunlar arasındaki uyumsuzluklar ve duvarlar, bugün kullandığımız pek çok teknolojinin geliştirilmesini zorlaştırır ya da olanaksız kılardı.⁷⁹

Patentler, insanlığın bilgi birikiminden yararlanarak bu birikime yapılan çok küçük katkıların birilerini zenginleştirmesine yarıyor. Aynı anda pek çok kişi ya

da ekip belirli bir konu üzerinde çalışırken, belirli bir buluş için patent başvurusunu ilk yapan kazanıyor ve aynı buluşu yapmış ya da yapmak üzere olan başkaları artık o buluşu kullanamaz hale geliyor.

Sürücüsüz otomobilleri insanlığın hizmetine sunmanın en hızlı yolu aransaydı, ilk yapılması gereken şey, bu alandaki her tür yeni bilgiyi hemen ve herkesle paylaşmak olurdu. Teknoloji geliştirmek için kullanılan tüm otomobiller arasında anlık olarak bilgi paylaşılması mümkün kılınırdı. Yapay zekâ algoritmaları herkese açık olur ve böylece bunlardaki hataların ve sorunların bulunması kolaylaşırdı. Bir kişi ya da ekip belirli bir sorunu çözdiğünde, başkaları gereksiz yere o sorun üzerinde çalışmazdı. Bir kişi ya da ekip belirli bir sorunu çözdiğünde, başkaları, patentler yüzünden aynı soruna gereksiz yere bambaşka çözümler aramak zorunda kalmazdı. Katkıda bulunmak isteyen herkes, Wikipedia örneğinde olduğu gibi, yetenekleriyle ve nitelikleriyle uyumlu katkılar sunabilirdi.

Kuşkusuz, ulaşım sisteminin bir bütün olarak ele alınması ve toplu taşıma ile otomobiller arasındaki en uygun dengeyi bulmak için çaba harcanması çok daha verimli bir yaklaşım olurdu. Ama bugün bunları yapmak mümkün değil, çünkü dünya üzerindeki kaynakların çok büyük bir bölümü sermaye sahiplerinin denetiminde.

Sürücüsüz otomobillerin çok sayıda insanı işsiz bırakacak olması da bir sorun değil mi?

Neyse ki bunun kısa bir süre içinde gerçekleşmesi olası değil...

Aslına bakılırsa, teknolojik gelişmelerin işsizliğe yol açması, kapitalizme özgü bir olgu. İnsanların yaptığı işlerin makinelere yaptırılması, zorunlu çalışma sürelerinin kısaltılmasını ve insanların özgürleşmesini de sağlayabilir.

Ne var ki, tam tersine, makinelerin uzantılarına dönüştürülen insanlar, eskisinden de çok çalışmak zorunda bırakılıyor.

Otomobil sürücülüğüyle devam edecek olursak, önümüzde Uber ve Lyft örnekleri var. Akıllı telefonlar sayesinde ortaya çıkabilen bu platformlar, sürücülerini sigortasız olarak ve hiçbir güvence sağlamadan çalıştırıyor. Aşırı düşük ücretler (şirketlere göre “hizmet bedelleri”) alan sürücüler geçinebilmek için otomobillerinde uyumak zorunda kalabiliyor. Yapılan araştırmalara göre, Uber sürücülerinin çoğu, kesintiler ve harcamalar çıkarıldıktan sonra, asgari ücretten daha düşük gelirlere sahip. Bir araştırmaya göre, Washington’daki Uber sürücülerini yoksulluk sınırının altında yaşıyor.⁸⁰

Yapay zekâ, çalışanlar üzerinde çok daha sıkı bir denetimin kurulmasını mümkün kılıyor. Örneğin, Amazon’un depolarında çalışan işçilerin her hareketleri makineler tarafından izleniyor. Yeterince hızlı ve yoğun çalışmayan işçiler algoritmalar tarafından işten çıkarılıyor. ABD’deki bir Amazon deposunda, Ağustos 2017 ile Eylül 2018 arasında, tüm çalışanların

yaklaşık olarak yüzde 10'u yeterince verimli çalışmadıkları gerekçesiyle bu şekilde işten atılmış. İşini korumak isteyenler tuvalete gitmekten bile çekiniyormuş.⁸¹ En düşük ücretlerin verildiği işçiler, çok uzun saatler boyunca aşırı hızlı çalışmak zorunda kaldıkları için hastalanabiliyor, yaralanabiliyor ve sakatlanabiliyor. 2018'de Amazon'un 23 deposunda tam zamanlı olarak çalışan işçilerin yaklaşık yüzde 10'u ciddi şekillerde yaralanmış. Depolarda kullanıma sokulan robotlarsa, işçilerin çalışma tempolarının daha da yükseltilmesine yol açmış.⁸²

Diğer yandan, teknolojik gelişmeler, yeni ve yıpratıcı işlerin ortaya çıkmasına yol açıyor. Bunlardan biri, içerik moderatörlüğü. YouTube'a ve Facebook'a yüklenen videolar şikâyet edildiğinde ya da algoritmalar tarafından riskli bulunduğu, bunlar insanlar tarafından kontrol ediliyor. Bu işi yapanlar, düşük ücretler karşılığında ve uzun saatler boyunca, normal insanların birkaç dakikasına bile zor katlanacağı türden videolar izliyor. Bu da psikolojik sarsıntı geçirmelerine, yani travma yaşamalarına yol açabiliyor. YouTube için bu işi yaptıran bir şirketin çalışanları haber konusu olunca, söz konusu şirket, çalışanlarına bir belge imzalatmaya başladı. Bu belgeyle, yaptıkları işin Travma Sonrası Stres Bozukluğuna yol açabileceğini bildiklerini, gerektiğinde tıbbi yardım alacaklarını ve yaptıkları işin ruh sağlıklarını olumsuz etkilediğini düşündüklerinde üstlerini bilgilendireceklerini kayıt altına alıyorlardı.⁸³

Çalışanların, işlerini kaybetme korkusuyla, tıbbi yardım almaktan ve üstlerini bilgilendirmekten kaçınacak olması şirket için o kadar da önemli değildi, çünkü asıl dert yasal sorumluluktan kurtulmaktı.

Yapay zekâyı geliştirmek için insanlara yaptırılan bir başka iş, akıllı telefonların yakın çevresinde yapılan konuşmaların ses kayıtlarını yazılı metinlere dönüştürmek...

Bilgimiz olmadan mı?

Evet, bilgimiz olmadan. Örneğin “Ok Google” ya da “Hey Google” demenizle başlayan sesli arama hizmetinden yararlanıyorsanız, arama sırasında söyledikleriniz, konuşma tanıma algoritmalarının geliştirilmesi için kullanılıyor. Ama 2019 yılında, Google’ın bundan fazlasını yaptığı açığa çıktı. Dünya genelinde, Google için çalıştırılan binlerce kişi, herhangi bir komut verilmeden başlayan ses kayıtlarını da deşifre ediyordu. Bunlar arasında, yatak odalarındaki konuşmaların ya da başkalarıyla yapılan telefon görüşmelerinin kayıtları da vardı. Bu bezdirici işi yapan çalışanlar, kaçınılmaz olarak, aile içi şiddet türü olaylara da tanıklık etmek zorunda kalıyordu.⁸⁴ Amazon’un, dijital asistanı Alexa’yı geliştirmek için, yine dünya genelinde binlerce kişiye ses kayıtlarını dinlettiği daha önce açığa çıkarılmıştı.⁸⁵ Aynı işi yaptıranlar arasında Facebook⁸⁶ ile Apple⁸⁷ da vardı.

Konuşma tanıma algoritmalarını geliştirmek için bunu yapmak zorunda oldukları söylenebilir mi?

Hayır. Konuşmalarının kaydedileceğini bilen insanlarla da çalışabilirlerdi. Tabii bunun için fazladan para ödemeleri gerekirdi. Teknoloji şirketleri, bizim ürettiğimiz her tür veriyi, izin almadan, bilgi vermeden ve herhangi bir karşılık ödemedi kullanmayı tercih ediyor. Tartışmalı işler yaptıkları açığa çıktığında ve kamuoyu baskısıyla karşılaştıklarında bazı uygulamalardan vazgeçiyorlar, ama yenilerini devreye sokuyorlar...

Yapay zekâ, geçmişte ücret ödenen işlerin bir bölümünü bize yaptırmak için de kullanılıyor. İnternet bankacılığı ve ATM'ler, banka çalışanlarının sayısının azaltılmasına yol açıyor. Kasiyersiz süpermarket kasalarında, barkodları okutma ve ödeme yapma işlerini biz üstleniyoruz. Teknolojik bir ürün alırken fiyat ve özellik karşılaştırması yapma ve satın alma işlerini internet siteleri üzerinden yaptığımızda, mağaza çalışanlarına ve hatta mağazalara duyulan ihtiyaç azalıyor. Çağrı merkezlerini arayıp robotlarla konuşmak zorunda kaldığımızda, derdimizi, robotlardan bile doğru cevap almamızı sağlayabilecek bir şekilde anlatmaya çalışıyoruz.

Böylece, geçmişin görece makul ücretli ve güvenceli işlerinin yerini çok düşük ücretli ve güvencesiz işlerin alması eğilimi güç kazanıyor. Ama bu arada zenginler daha da zenginleşmeye devam ediyor.

2009-2019 yıllarında dünyadaki dolar milyarderlerinin sayısı iki katına çıkarak 2153'e ulaştı. Bu kişilerin toplam serveti, en yoksul 4,6 milyar insanın toplam servetinden fazlaydı.⁸⁸

ABD'deki en büyük şirketlerin CEO'larına yapılan ödemelerin bu şirketlerdeki ortalama işçi ücretlerine oranı 1965 yılında 20'ye 1 düzeyindeydi. Bu oran düzenli bir şekilde yükselse de, 1989 yılında ancak 58'e 1'e ulaşmıştı. Hemen ardından internet çağına girdik. 2000 yılına gelindiğinde CEO'ların kazançları ortalama işçi ücretlerinin 344 katına çıktı. İzleyen yıllarda dalgalanmalar olduysa da, söz konusu oran 175'e 1'in altına hiç inmedi ve 2017 yılında yeniden 312'ye 1'e yükseldi.⁸⁹ Bir başka deyişle, teknolojik gelişmelerin getirileri giderek daha eşitsiz şekillerde paylaşılıyor.

Tabii bu sayılar dünya genelindeki eşitsizlikleri yansıtmıyor.

Amazon'un depolarındaki en düşük ücretli işçilerin saatlik ücreti 15 dolarken, bu şirket için elektronik ürünler üreten Foxconn'un Çin'deki fabrikalarındaki saatlik ücretler geçtiğimiz yıl 3 dolar civarındaydı. Yine geçtiğimiz yıl, Apple'ın iPhone markalı telefonlarını da üreten Foxconn'un, "stajyer"lik yaptırdığı ve yasalara aykırı davranarak fazla mesaiye zorladığı öğrencilere 2,5 dolardan düşük saatlik ücretler ödediği açığa çıkarıldı. 16-18 yaşlarındaki çocuklar, Amazon'dan gelen siparişleri yetiştirmek için, günde 10 saat ve haftada 6 gün çalıştırılıyordu.⁹⁰

Amazon bu konuda bir açıklama yaptı mı?

Yaptı. Şirketin bir sözcüsü, iddiaların inceleneceğini ve gerekirse Foxconn'dan düzeltici önlemler almasını isteyeceklerini açıkladı.

Ama büyük teknoloji şirketleri, markalarını taşıyan elektronik ürünlerin ne tür koşullar altında üretildiğini çok iyi biliyor. 2010 yılında, Foxconn'daki çalışma koşulları nedeniyle 18 işçi intihar girişiminde bulunmuş, 14'ü ölmüştü. Aşırı stres, çok uzun iş günleri, hataları nedeniyle işçileri aşağılayan yöneticiler, haksız cezalar ve tutulmayan sözler, gerekçelerden bazılarıydı.⁹¹

Ve çok daha kötüsü var... Bilgisayarlarda, cep telefonlarında ve diğer elektronik cihazlarda kullanılan pillerin en önemli hammaddelerinden biri kobalt. En zengin kobalt madenleri Afrika'daki Kongo Demokratik Cumhuriyeti'nde. Bu ülkenin ihraç ettiği kobaltın yaklaşık üçte biri, aralarında çocukların da bulunduğu 150 bin kişi tarafından en düşük ücretlerle ve en ilkel yöntemlerle çıkarılıyor. Sık sık ölümlü kazalar yaşanıyor.⁹² Çocuklar için günde 2 dolar ödenirken, 9 yaşındaki çocuklar bile çalıştırılabilir. Büyük teknoloji şirketleri bunlardan da haberdar elbette. Bu arada, madenlerdeki çocuk ölümleri nedeniyle 2019'un Aralık ayında Apple, Google, Tesla, Microsoft ve Dell'e karşı bir dava açıldı.⁹³

Bu tür işler robotlara yaptırılmaz mı?

2010 yılındaki işçi intiharlarını gündemden düşürmek isteyen Foxconn, 2011'de, Çin'deki fabrikalarında

kullandığı robotların sayısını üç yıl içinde 10 binden 1 milyona çıkaracağını açıklamıştı.⁹⁴ Robot sayısı, biraz da, robotların nasıl tanımlandığına bağlı elbette. Kesin olan şu ki, elektronik ürünlerin montajını hâlâ çok büyük ölçüde insanlar yapıyor. 2019'daki bir tahmine göre, dünya genelinde 1 milyondan fazla işçi çalıştıran Foxconn'un Çin'deki fabrikalarında kullanılan robotların sayısı 20 binden azdı.⁹⁵

Robot üreten iki farklı şirketin kurucusu olan robot bilimcisi Rodney Brooks, 2018'de, IKEA mobilyalarını monte edebilecek robotların yapılıp yapılamayacağını tartışmıştı. Ona göre, birkaç ay çalışılması durumunda, son derece sınırlandırılmış bir ortamda (yani herhangi bir evde değil laboratuvar benzeri bir yerde), en gelişkin robotlara, tek bir IKEA ürününü monte etmek için yapılması gereken ve el becerisi isteyen yaklaşık 200 işlemden 1 tanesi yaptırılabilirdi. Ama sadece kaba bir gösteri videosu çekmeye yetecek kadar.⁹⁶

Bugün, yapay zekâ gerçek dünyada olup bitenleri anlamının ne kadar uzağındaysa, robotlar da gerçek dünyaya uyum sağlayıp insanların fiziksel becerilerini gerektiren çok farklı işleri yapabilir olmanın o kadar uzağında. En yeni ve en gelişkin robotların neler yapabildiğini görmek için YouTube'a bakılabilir. Ama gösteri videolarını izlerken şu soruyu akılda tutmak gerekir: Aynı robotlar, gösterilerde yaptıklarını gerçek dünyada da yapabilir mi? Örneğin, sahnedeki merdiveni çıkıp inebilen bir robot, evimizin merdivenini de

çıkıp inebilir mi? Çünkü günümüzün robotları ancak sınırları çok kesin olarak tanımlanmış olan ortamlarda iş görebiliyor.

Tabii robotlara bugünküne göre daha fazla işin yaptırılmamasının bir nedeni de işçilerin çok düşük ücretlerle çalıştırılabilmesi.

O zaman, işçi ücretleri artarsa, daha çok robot kullanılacağından işsiz sayısı da mı artar?

Bugünün dünyasında gerçekten öyle olur.

İşsizliğin artmasının bir nedeni de çalışma sürelerinin çok uzun olması.

19. yüzyılda, işçiler, çalışma sürelerinin haftada 40 saatle sınırlandırılması için mücadele ediyordu. 20. yüzyılın başlarında pek çok ülkede bu hedefe yaklaşılmıştı. Aradan geçen sürede teknoloji fazlasıyla ilerledi. Dolayısıyla çalışma süreleri çok daha kısa olabilir. Ama bunun yerine, işçiler 19. yüzyıldaki kadar çok çalıştırılırken işsizlik artıyor. Ve işsizlerin varlığı, çalışanlara düşük ücretleri, uzun çalışma sürelerini ve kötü çalışma koşullarını kabul ettirmeyi kolaylaştırıyor.

Bazı ülkelerde tartışılan ve hatta denemeleri yapılan evrensel temel gelir bir çözüm olabilir mi?

Evrensel temel gelir, yani çalışıp çalışmadığından ve maddi durumundan bağımsız olarak her bireye yaşamayı sürdürmesine yetecek bir gelirin sağlanması hakkındaki ilginç bir olgu, bunu savunanlar arasında solcuların yanı sıra zenginlerin de bulunması.

Zenginlerin evrensel temel geliri savunmasının demagogik bir yanı var elbette. Onların asıl amacı, bugünkü muazzam eşitsizliklerin sorgulanmasının önüne geçmek. İşsizliği, teknolojik gelişmelerin kaçınılmaz bir sonucuymuş gibi göstermeye çalışıyorlar. Diğer yandan, yetersiz bir evrensel temel gelirin varlığında, biraz daha insanca yaşamak isteyenleri daha düşük ücretler ödeyerek çalıştırabilirler. Dahası, bunu bugün de yapıyorlar.

2015 yılında ABD’de yapılan bir araştırmaya göre, düşük ücretli işlerde çalışanlar, devletten yılda 150 milyar dolardan fazla yoksulluk yardımı alıyordu. McDonald’s benzeri “fast food” şirketlerinde çalışanlar arasında, yoksullara yardım programlarının en az birinden yararlananların oranı yüzde 52’ydi.⁹⁷ Yoksulluk yardımları sayesinde işçilerine daha düşük ücretler ödeyebilen şirketlerden biri de, evrensel temel geliri savunan Jeff Bezos’un Amazon’u.⁹⁸

Diğer tarafta, teknolojinin her sorunu çözebileceğine inananlar var. Örneğin, Nick Srnicek ile Alex Williams, bugünkü biçimiyle çalışmanın ortadan kaldırılması için mücadele etmemiz gerektiğini savunuyor. Onlara göre, bunu başarmak için iki temel talebimiz olmalı: Bugün insan emeğini gerekli kılan bütün işlerin makinelere yaptırılması anlamına gelen tam otomasyon ve evrensel temel gelir.⁹⁹

Açıkçası, görece yakın bir gelecekte (Srnicek ile Williams’a göre 2035 yılına kadar) tam otomasyona

ulaşma hedefi hiç gerçekçi değil. Asıl önemlisi, böyle bir hedef koymanın doğru bir şey olup olmadığı fazlasıyla tartışmalı.

Yapay zekânın tüm sorunları çözebileceğine inananlar, insanları tümüyle devre dışı bırakmanın görece kolay ve doğru olduğu varsayımından hareket ediyor. Oysa, gerçek dünyada olup bitenleri hiçbir şekilde anlamayan, sadece geçmişe ait verileri kullanabilen, bunu da ancak söz konusu veriler yeterli sayıya ulaşmışsa yapabilen, tahminlerinin doğru çıkıp çıkmayacağı hep belirsiz olan yapay zekâ algoritmalarının yöneteceği bir dünya, insanlık için tam bir cehennem olurdu.

Bugün ihtiyaç duyduğumuz şey, çalışmanın ortadan kaldırılması değil, karar alma ve üretim süreçlerine herkesin katılabilmesi. İnsanları tümüyle devre dışı bırakan teknolojilere değil, insanlığın kolektif zekâsının ürünü olan, bu zekâyla birlikte iş gören ve onu geliştiren teknolojilere ihtiyacımız var.

5. Yapay Zekâdan Kolektif Zekâyâ

“Kolektifzekâ” tam olarak ne anlama geliyor?

Zekâdan çoğu zaman bireylere özgü bir yeti anlaşıl-
sa da, topluluklar, sorunları çözmek konusunda tek tek
bireylerden daha başarılıdır. Tek bir karıncanın hare-
ketlerini incelediğimizde onu pek zeki bulmayabiliriz.
Ama bir karınca topluluğunun yapabildikleri bizi şa-
şırtır. Toplumlar için de aynısı geçerli.

2010 yılında, grupların zekâ düzeylerini ölçmek için
699 kişilik bir araştırma yapılmış. Beşer kişilik grupla-
ra farklı görevler verilmiş. Sonuçlardan, en zeki insan-
ları bir araya getirmekten çok, çeşitliliğin, yani farklı
özelliklere sahip olan insanların bir araya getirilmesi-
nin önem taşıdığı anlaşılmış. Bir ya da iki kişinin faz-
la baskın olduğu grupların kolektif zekâ düzeyi daha
düşük çıkmış. Kadınların sayısı arttıkça kolektif zekâ
düzeyi de yükselmiş.¹⁰⁰

ABD’deki teknoloji şirketlerinin sorunlarından biri,
bunların çalışanları arasında belirli seçkin üniversite-
lerin mühendislik fakültelerinden mezun olan beyaz
erkeklerin oranının çok yüksek olması ve üst yönetim
kademelerinde bu oranın daha da yükselmesi. Hepi-
mizi ilgilendiren kararlar, az sayıda zenginin çıkarları
doğrultusunda ve “teknoloji hakkındaki her şeyi bilip
teknolojinin toplumsal sonuçlarından neredeyse hiçbir
şey anlamayan kişiler”¹⁰¹ tarafından alınıyor.

Teknoloji şirketleri, insanları mümkün olduğunca devre dışı bırakmayı tercih ediyor. Oysa makineler ile insanlar arasındaki ilişkiler çok farklı şekillerde düzenlenebilir.

2005 yılındaki bir “serbest stil” satranç turnuvasına herkes dilediği şekilde ekip oluşturarak (ya da oluşturmayarak) katılmış. Bu turnuvada, satranç bilgisayarı kullanan insanlar, en güçlü satranç bilgisayarlarını bile yenilgiye uğratmış. Asıl önemlisi, turnuvayı, satranç büyük ustası olan ve dönemin en gelişkin satranç bilgisayarlarından birini kullanan bir oyuncu değil, aynı anda üç bilgisayardan yararlanan iki amatör oyuncu kazanmış. 1997 yılında IBM’in Deep Blue adlı satranç bilgisayarına yenilen dünya satranç şampiyonu Garry Kasparov’un bunlardan çıkardığı sonuca göre, *zayıf insan + makine + daha iyi yöntem*, hem tek başına güçlü bir bilgisayardan, hem de daha önemlisi, *güçlü insan + makine + daha kötü yöntem*’den üstündü. Buna daha sonra Kasparov Yasası denmiş.¹⁰²

Amatörler, günümüzün en güçlü satranç bilgisayarlarına karşı da aynı başarıya ulaşabilir mi?

Ulaşamayabilirler. Çünkü, daha önce tartıştığımız gibi, satranç gibi oyunların ayırt edici özelliği, kurallarının hiç değişmemesi. Oysa gerçek yaşam, hiç durmadan, yeni kurallar geliştirmeyi gerekli kılan yeni durumlar yaratıyor. İnsanların makineler karşısındaki en önemli üstünlüklerinden biri, yeni durumlara hızla uyum sağlayabilmeleri. Özellikle de birlikte hareket

edebildiklerinde, yani kolektif zekâlarını kullanabildiklerinde...

Wikipedia, kolektif zekânın neler yaratabileceğinin iyi bir örneği.

Google'ın arama motoru, Wikipedia olmasaydı, bugünkü güvenilirliğini kolay kolay kazanamazdı. Özellikle İngilizce aramalarımızın pek çoğunda, Wikipedia maddelerinin bağlantılarını en üst sıralarda görüyoruz. Dahası, Google, çoğu zaman, bunların özet içeriğini ayrı bir kutuda gösteriyor.

Ansiklopedi hazırlamak, 1990'lı yıllara kadar, sadece uzmanların yapabildiği ve ticaret konusu olan bir işti. Bugünse, milyonlarca insanın gönüllülük temelindeki katkılarıyla yaratılmış, her an güncellenen, ücretsiz, reklamsız, kullanıcılarının bilgilerini veri simsarlarına satmayan, yüzlerce dilde versiyonu bulunan, geçmiştekilerle karşılaştırıldığında devasa boyutlarda bir ansiklopedimiz var. 2001 yılında kurulan Wikipedia, herkesin katkıda bulunabildiği ve kâr amacı gütmeyen bir platformun insanlığın bilgi birikimini sistemli bir şekilde özetleme işini ticari girişimlerden çok daha başarılı bir şekilde yapabileceğini kanıtladı.

Wikipedia'da çok fazla yanlış bilgi yok mu?

Var tabii ki. Herkesin dilediği gibi değişiklik yapabildiği bir ansiklopedide yanlış bilgilerin de bulunması kaçınılmaz. Ticari şirketler, siyasal partiler, devlet kuruluşları, çok farklı çıkar grupları ve farklı sorunları bulunan bireyler Wikipedia'ya kendi dar çıkarları doğ-

rultusunda müdahalelerde bulunmaya çalışıyor. Ama özellikle İngilizce Wikipedia, bu tür müdahalelerin etkilerini büyük ölçüde sınırlandırmayı başarabiliyor.

Son koronavirüs salgını sırasında, Facebook ve Twitter gibi platformlarda yalan yanlış haberler kolayca yayılırken, İngilizce Wikipedia'da tıpla ilgili maddeler, tıp ve halk sağlığı alanlarında uzmanlığı bulunan yaklaşık 150 editör tarafından düzenli bir şekilde denetleniyordu.¹⁰³

Şu anda yaklaşık 140 bin aktif katkıcısı bulunan İngilizce Wikipedia'nın¹⁰⁴ güvenilirlik düzeyini yüksek tutmak için, ansiklopedi katkıcılarının başlangıçta soğuk baktığı hiyerarşik ilişkileri de içeren bir iç işleyiş var. Kayıtlı katkıcılarından oluşan Wikipedia topluluğu tarafından uzlaşma yoluyla seçilen "bürokrat"lar (üst düzey yöneticiler), Wikipedia topluluğunun ortak görüşleri doğrultusunda atanan "yöneticiler" ve seçimler yoluyla belirlenen "hakem kurulu", ansiklopedinin nitel düzeyinin korunmasını sağlıyor.¹⁰⁵

Wikipedia, insanlar ile yapay zekâ algoritmaları arasındaki ilişkiler konusunda da bir model oluşturuyor. Tekrar tekrar yapılması gereken sıradan editörlük işleri için, "bot" adı verilen internet robotlarından yararlanılıyor. Ama hangi botların ansiklopedi üzerinde çalışabileceğine ve hangilerinin devre dışı bırakılacağına yine Wikipedia topluluğu karar veriyor. Botlardan bazıları yöneticilerin sorumluluk alanına giren işler yapıyor. Bu tür botların hangi görevleri üstlenebileceğine

topluluk karar veriyor. Yönetici botların kaynak kodlarının açık olması tavsiye ediliyor. Bot işleticisi, kaynak kodlarının kısmen ya da tümüyle kapalı tutulmasını sağlayabilse de, dileyen her yöneticiye bunları sunmak zorunda.¹⁰⁶ Yani, Wikipedia topluluğunun ortak görüşleri doğrultusunda atanan binden fazla yönetici, kaynak kodlarına erişebiliyor. Dolayısıyla, Wikipedia’da kullanılan algoritmalar, örneğin Google’ın arama sonuçlarını filtreleme ve sıralama algoritmalarından farklı olarak, insanların denetimine tümüyle kapalı birer kara kutu değil. Diğer yandan, Wikipedia’daki algoritmalar, insanları devre dışı bırakmak yerine, yaratıcılık gerektiren işlere daha fazla zaman ayırmalarını sağlıyor. Böylece, Wikipedia topluluğunun kolektif zekâsı, algoritmalar yardımıyla daha da gelişiyor.

Açık kaynaklı yazılımlar da örnek olarak gösterilebilir mi?

Elbette. Açık kaynak kodlu projelerin önemli bir bölümünü izleyen Open Hub platformuna göre şu ana dek yaklaşık 5 milyon kişi yaklaşık 500 bin projeye katkıda bulunmuş durumda.¹⁰⁷

Tabii bir noktaya dikkat etmekte yarar var. Büyük teknoloji şirketleri de açık kaynaklı projelerden yararlanıyor ve bunlara katkıda bulunuyor. Örneğin, Linux Vakfı’nın “platin” üyeleri arasında Google, IBM, Intel ve hatta Microsoft da var. Üye şirketlerin toplam sayısı ise bine yakın.¹⁰⁸ 2016 yılında hazırlanan bir rapora göre, Linux temelli işletim sistemlerinin çekirdeğini

geliştirmeye yönelik katkılar arasında amatörlerin yaptığı katkılarının oranı yüzde 7,7'ye düşmüştü.¹⁰⁹

Bu bir yana, Linux çekirdeği, çok sayıda katkıcısı bulunan açık kaynaklı projelerin teknolojik gelişimde önemli bir rol oynayabildiğini gösteriyor. Akıllı telefonların büyük çoğunluğunda kullanılan Android işletim sisteminin temelinde Linux çekirdeği var. Web sitelerinin sunucu bilgisayarlarının çoğunda ve dünyanın en güçlü 500 bilgisayar sisteminin hepsinde, daha güvenli oldukları için, Linux çekirdeğine dayalı işletim sistemleri kullanılıyor.¹¹⁰

Wikipedia benzeri platformlar yardımıyla büyük şirketlere alternatifler yaratılamaz mı? Yapay zekâ, açık kaynaklı bir şekilde geliştirilip bütün insanların hizmetine sunulamaz mı?

Sorun şu ki, üretim araçlarının ve doğal kaynakların çok büyük bir bölümü sermaye sahiplerinin elindeyken, insanlığın bilgi birikiminin önemli bir bölümü patentlerle ve fikrî mülkiyet haklarıyla kullanılmaz kılınır ya da ticari sır olarak saklanırken, insanlar geçinebilmek için çok uzun saatler boyunca yıpratıcı işlerde çalışmak zorundayken, gönüllülük temelindeki katkılarla elde edilebilecekler sınırlı kalacaktır.

Örneğin, açık kaynaklı arama motorları da geliştirilmiş durumda. Google'ın farkı, internetin erişilebilir içeriğinin neredeyse tümünün dizinine sahip olması ve bu dizini anlık olarak güncelleyebilmesi. Böyle bir dizini saklayabilmek, güncelleyebilmek ve buna her an

erişebilir olmak için çok fazla sayıda bilgisayara ihtiyaç var. Google, bu iş için kullandığı bilgisayarların niteliklerini ve sayılarını gizli tutuyor. Ama 2016 yılında yapılan bir tahmine göre Google'ın o dönemde yaklaşık 2,5 milyon sunucu bilgisayarı varmış.¹¹¹ Ölçeğinin büyüklüğü nedeniyle, Google, diğer arama motorlarına göre çok daha iyi sonuçlar üretiyor.

Yeni bir sosyal ağ kurmak da zor değil. Ama milyarlarca insanın paylaşımlarını saklayabilmek için yine çok fazla sayıda bilgisayara ihtiyaç var.

Kimileri, evrensel temel gelirin bu sorunu aşmaya yardımcı olabileceğine inanıyor. Örneğin, Amerikalı Marksist Erik Olin Wright'a göre, yeterli yükseklikteki bir evrensel temel gelir, yeni teknolojilerin de yardımıyla, kapitalist ekonomiden ayrı ve dayanışmacı bir toplumsal ekonominin yaratılmasını sağlayabilir. Wright, kooperatif tarzı örgütlenmelere yaslanacak olan bu ekonominin büyütülmesi yoluyla kapitalist ekonominin giderek daraltılabileceği kanısında.¹¹²

Burada iki sorun var. Birincisi, kooperatif tarzı örgütlenmeler, dev şirketlerle rekabet edemez. Bunlarla teknolojik gelişmelerin yönünü belirleyebilecek olan büyük ölçekli yatırımlar yapamazsınız. Bilgisayar, akıllı telefon, dayanıklı tüketim eşyası vb. üretemezsiniz. Market zincirleriyle rekabet edemezsiniz. Belli başlı tüketim mallarını daha ucuza üretemezsiniz.

İkincisi, zenginlerden ve büyük şirketlerden daha fazla vergi toplanması ve bu paraların halka aktarılma-

sı elbette iyi bir şey olurdu. Ne var ki, bunu dünya ölçeğinde ve aynı anda yapamazsanız, zenginler ve büyük şirketler varlıklarını mümkün olduğunca yurt dışına çıkarmaya çalışırken ülkenizdeki üretimlerini azaltırlar. Bu da ekonominin küçülmesine, toplanabilecek vergilerin ve dolayısıyla halka aktarılabilecek paraların azalmasına ve işsizliğin artmasına yol açar.

Böyle davranan zenginlerin ve büyük şirketlerin varlıklarına el koyulamaz mı?

Evrensel temel gelir, solcular tarafından, biraz da, üretim araçlarının özel mülkiyetine son verme talebinin fazla radikal bulunması nedeniyle savunuluyor. Gelir dağılımını düzeltmeye yönelik reform taleplerinin daha gerçekçi olacağı düşünülüyor. Ama üretim araçlarının büyük çoğunluğu sermaye sahiplerinin elinde olduğu sürece, iktisadi çöküntüye yol açmadan yapılabilecek olan gelir dağılımı reformları hep fazlasıyla sınırlı kalacak, köklü bir değişim getirmeyecektir.

Teknolojik ilerlemelerle ilgili bir gerçek çoğu zaman unutuluyor ya da hiç bilinmiyor. Bilgisayar teknolojileri de, internet ve onunla bağlantılı teknolojiler de, ilk dönemlerinde, büyük ölçekli kamu yatırımları sayesinde gelişmişti. ABD’de, 1940’lı yılların başları ile 1960’lı yılların başları arasında, bilgisayar teknolojilerinin neredeyse tek geliştiricisi silahlı kuvvetlerdi. Araştırmaların çoğu üniversitelerde ve ticari şirketlerde yürütülse de bunların bütçeleri askerî kuruluşlar tarafından

sağlanmıştı. Ticari bilgisayar sektörünün şekillenme aşamasında, bu sektörün en önemli müşterisi silah sanayisiydi.¹¹³ Üniversitelerin bilgisayarlarını birbirlerine bağlayan ve böylece internetin teknik temelini oluşturan ARPANET, 1969 yılında, ABD Savunma Bakanlığına bağlı bir ajans tarafından kurulmuştu.

Yapay zekâ çalışmalarının gözden düştüğü “yapay zekâ kışları”ndan söz edilir. İlk yapay zekâ kışı 1970’li yıllarda yaşanmıştı ve bunun en önemli nedeni, uzmanların vaat ettiği başarıların elde edilemeyeceğine karar veren ABD Savunma Bakanlığı’nın araştırma desteklerini sınırlandırmasıydı.¹¹⁴ 1980’li yılların başından itibaren bu destekler artarken yapay zekâ alanındaki çalışmalar yeniden önem kazandı. 1983 ile 1993 yılları arasında yapay zekâyı geliştirmeye yönelik araştırmalara fazladan 1 milyar dolar harcandı.¹¹⁵ Ama ABD Savunma Bakanlığı bir kez daha beklentilerin karşılanamayacağına karar vererek desteklerini sınırlandırınca ikinci yapay zekâ kışına girildi.

Devlet bütçesinden aktarılan kaynaklar olmasaydı, bugünkü Silikon Vadisi şirketleri ortaya çıkamazdı. Bugün bile, ABD Savunma Bakanlığı, Silikon Vadisi şirketlerinin önemli müşterileri arasında yer alıyor. Benzer şekilde, Çin de, teknoloji alanında ABD’yle rekabet edebilir hale gelmesini büyük ölçekli kamu yatırımlarına borçlu.

Teknolojinin herkese hizmet eder hale getirilmesi için, birincisi, büyük ölçekli kamu yatırımlarına ih-

tiyaç var. İkincisi, hangi alanlara yatırım yapılacağı konusunda herkesin söz sahibi olabilmesi, üçüncüsü, doğru kararların alınabilmesi için herkesin nitelikli ve bilimsel eğitim hizmetlerinden yararlanabilmesi, dördüncüsü, teknoloji geliştirme çalışmalarına herkesin katılabilmesi gerekir.

Komünist partinin iktidarda olduğu Çin'in teknoloji politikaları ne kadar farklı? Bu ülke örnek alınabilir mi?

Açıkçası, Çin'in iyi bir örnek oluşturduğunu düşünmüyorum.

Birincisi, Foxconn'un Çin'deki fabrikalarından söz etmiştik. Düşük ücretler, uzun çalışma süreleri, ağır çalışma koşulları ve çocuk işçilerin çalıştırılması, dünyanın her yerinde bu kadar çok sayıda "Çin malı"yla karşılaşmamızın temel nedeni. Çin'in önde gelen teknoloji şirketlerinde çalışan bilgisayar programcıları, 2019 yılında, "996 karşıtı" bir mücadele başlatmıştı. Sabah 9'dan akşam 9'a kadar ve haftanın 6 günü çalışmak istemiyorlardı!¹¹⁶ Diğer yandan, Çin, dünyada, sömürgeci bir güç olarak hareket ediyor. Örneğin Kongo Demokratik Cumhuriyeti'nde insanlık dışı koşullar altında çıkarılan kobaltın en büyük alıcısı Çin.

İkincisi, Çin'deki toplumsal eşitsizlikler hızla artıyor. Bu ülkedeki dolar milyonerlerinin sayısı 2000 yılında 38 binken 2019'da 4 milyon 447 bine çıkmıştı.¹¹⁷ Yine 2019 yılında, dolar milyarderlerinin sayısının en yüksek olduğu ikinci ülke, 324 milyarderle Çin'di.¹¹⁸

Dolayısıyla, ülkedeki ayrıcalıklı katman giderek daha da güçleniyor.

Üçüncüsü, Çin, halkın karar alma süreçlerine katılma olanaklarının çok sınırlı olduğu ülkelerden biri.

Çin yönetimi, 2014 yılında, ülkedeki tüm şirketlerin ve devlet kurumlarının yanı sıra tüm yurttaşları kapsayacak olan bir Sosyal Kredi Sisteminin en geç 2020 yılının sonunda devreye sokulacağını açıkladı. Her yurttaş, onun hakkında toplanan çok sayıda veriden hareketle, “güvenilirlik” düzeyini gösteren bir sosyal kredi notunun verilmesi planlanıyor. “Kötü” davranışlar sergileyenlerin “kara liste”lere, “iyi” davranışlar sergileyenlerin “kırmızı liste”lere sokulması, yüksek not alanların ödüllendirilmesi (örneğin kamu hizmetlerinden daha ayrıcalıklı bir şekilde yararlanmalarının sağlanması), düşük not alanların cezalandırılması (örneğin bazı kamu hizmetlerinden yararlanamaz hale getirilmeleri) öngörülüyor.¹¹⁹ Şu anda yerel yönetimler ve özel şirketler tarafından kısmi uygulamaları yapılan böylesi bir “büyük veri” projesi, veri toplama ve değerlendirme süreçlerinde kaçınılmaz olarak sayısız hata yapılacağından, sayısız haksızlığa yol açma potansiyeline sahip.

2004 yılından bu yana Wikipedia’ya erişimi farklı dönemlerde kısmen ya da tümüyle engelleyen, 2019 yılının Nisan ayından beri tüm Wikipedia versiyonlarını erişime kapalı tutan Çin yönetimi,¹²⁰ dünyadaki herkesin yararlanabileceği ve katkıda bulunabileceği daha iyi

bir platform yaratmaya çalışmak yerine, sadece Çince'nin kullanıldığı ve devlet denetimi altındaki ticari ansiklopedi sitelerini destekliyor. Bir başka deyişle, Çin yönetiminin örnek oluşturacak bir ülke yaratmaya çalıştığı da söylenemez. Tam tersine, kendi yurttaşlarının bilgi edinme ve insanlığın bilgi birikimine katkıda bulunma olanaklarını sınırlandırmayı tercih ediyorlar.

Peki, teknoloji geliştirme çalışmalarına herkesin katılması gerçekçi bir hedef mi?

Herkesin katılması değil ama katılabilir olması, günümüz koşullarında fazlasıyla gerçekçi bir hedef.

Öncelikle, sadece teknik katkılardan söz etmediğimi vurgulayayım. Teknoloji geliştirme çalışmaları mühendislerin tekeline bırakıldığında, Google Gözlük (*Google Glass*) gibi, toplumun reddedeceği projeler bile üretilebiliyor. 2012 yılında duyurulan bu projeye göre, internete bağlanabilen bir bilgisayar, kamerası, hoparlörü, mikrofonu ve ekran görüntüsü bulunan bir gözlük yaşamımızı değiştirecekti. İnsanların fotoğraflarını ve videolarını izin almadan çekmeyi mümkün kılacak olan Google Gözlük, daha piyasaya çıkmadan tepkilerle karşılaştı. Bazı mekânlara gözlüğü kullanmanın yasak olduğu uyarıları asılırken, bir eyalette otomobil kullanırken bu gözlüğün takılması yasaklandı. Yine de, gözlük, 2013 yılında, 1500 dolar gibi yüksek bir fiyatla seçilmiş kişilere satıldı. Ve bu kez Google Gözlük kullananlar sözlü ve fiziksel saldırılara uğradı. Google, Ocak 2015'te projeyi sonlandırmak zorunda kaldı.¹²¹

Yeni teknolojiler geliştirilirken, onları kullanacak ve onlardan etkilenecek olanların görüşlerinin de alınması gerekir.

Diğer yandan, örneğin makine öğrenmesine ilgi duyuyorsanız ve teknik metinleri anlamınıza yetecek kadar İngilizce biliyorsanız, dünya genelinde 3 milyondan fazla kullanıcısı bulunan Kaggle adlı platformdan yararlanabilirsiniz ([kaggle.com](https://www.kaggle.com)). 2017 yılında Google'ın satın aldığı bu platformda hem makine öğrenmesi tekniklerini öğrenip uygulayabilir, hem de pek çoğu para ödüllü olan yarışmalara katılabilirsiniz. Dahası, eğer kendinizi yeterince geliştirebilerseniz, platform sayesinde iş bile bulabilirsiniz. Google'ın Kaggle'ı satın alma gerekçeleri arasında en yetenekli insanlara daha kolay ulaşmak da vardı. Tabii bunlar platform kullanıcılarının çok küçük bir azınlığını oluşturuyor.

Kaggle'daki para ödüllü yarışmalar, özel şirketler ve araştırma kurumları tarafından düzenleniyor. Yüzlerce, hatta bazen binlerce ekibin katılabildiği bu yarışmalar, sınırlı sayıda insanla görece kısa bir sürede çözülüp çözilemeyecekleri bilinmeyen problemlerin çözülmesini sağlayabiliyor. Bu arada, yarışmacıların büyük çoğunluğu, sundukları projeler yararlı katkılar içerse bile, para ödenmeden çalıştırılmış oluyor.

Kuşkusuz, kâr getirebilecek teknolojilerin pek çoğu kapalı kapıların ardında geliştiriliyor. Ama kamu yatırımları söz konusu olduğunda, herkesin katkılarına açık platformlar kolaylıkla yaratılabilir.

İnsanların büyük çoğunluğu bunun için maddi karşılık beklemez mi?

Mevcut teknolojiler, insanların fiziksel ve zihinsel emeklerinin yardımıyla, herkesin tüm temel ihtiyaçlarının karşılanmasını sağlayabilecek durumda. Çalışabilecek durumdaki herkesin çalışması durumunda, zorunlu çalışma süreleri en aza indirilebilir.

Bu arada, yapılması zorunlu olan işlerin herkesçe paylaşılması durumunda, en yıpratıcı işlere ayrılması gereken süreler fazlasıyla kısaltılabilir. Örneğin travma sonrası stres bozukluğuna yol açabilen videolara bakma süresi günde birkaç dakikayla sınırlandırılabilir. Bu kadarına bile katlanamayacak olanlar, başkalarına daha yıpratıcı gelen farklı işleri tercih edebilir.

Diğer yandan, insanlar, artacak olan serbest zamanlarında neler yapacaklarına özgürce karar verebilir.

Bugün bile, milyonlarca insanın gönüllülük temelindeki katkılarıyla Wikipedia gibi bir ansiklopedi hazırlanabiliyor ve açık kaynaklı yazılımlar geliştirilebiliyor. Çok uzun süreler boyunca çalışmak zorunda kalmayan ve maddi sıkıntı çekmeyen insanların çok daha fazlasını yapabileceğini öngörebiliriz.

Diyelim ki, toplumsal çıkarlar belirli bir teknolojinin hızla geliştirilmesini gerektirdi ve gönüllülük temelinde katkıda bulunanların çalışmaları yeterli olmadı. Toplum, belirli hedeflere ulaşılmasını sağlayacak olanları maddi olarak ödüllendirmeyi tercih edebilir. Burada önemli olan, toplumu ilgilendiren kararların

herhangi bir dayatma olmadan yine toplum tarafından alınabilmesi.

Peki, son bir soru: Yakın gelecekte olmasa bile uzun vadede çalışma tümüyle ortadan kalkabilir mi?

Üretim araçlarının herkese ait olduğu bir toplum, çok da uzak olmayan bir gelecekte, “zorunlu” çalışmayı tümüyle ortadan kaldırabilir. Yaşamlarına anlam katmak isteyen insanların gönüllülük temelindeki çalışmaları yeterli hale gelebilir.

Çalışmayı tümüyle ortadan kaldırmak içinse, insanlığın geleceğini ilgilendiren tüm kararları makinelerin alması gerekir. Çünkü, neyin nasıl ve ne kadar üretileceğine karar vermek için bile çalışmak gerekir. Açıkçası, tüm kararları makinelere bırakmanın teorik açıdan mümkün olacağı bir aşamaya ulaşılabileceğini sanmıyorum. Diğer yandan, böyle bir aşamaya ulaşmanın istenir bir şey olup olmadığı da tartışmalı.

Ama bence, öncelikli hedefimiz, insanlığın bu tür konularda özgürce karar alabilir hale gelmesini sağlamak olmalı.

Eğer makineler çok küçük bir azınlığın denetimi altında kalmaya devam ederse, insanlığın geleceğini ilgilendiren pek çok karar, bugün yapıldığı gibi, hiç de yetkin olmayan makinelere aldırılabilir.

Buna karşılık, herkesin nitelikli bir eğitim aldığı, her düzeydeki karar alma süreçlerine bunlardan etkilenecek olan herkesin katılabildiği, üretim araçlarının herkese ait olduğu, insanlığın geleceğini ilgilendiren

her tür bilgiye herkesin erişebildiği ve katkıda bulunabildiği bir toplum, kendi geleceği hakkında, bugün bizim hayal edebileceklerimizden çok daha iyi kararlar alacaktır.

Notlar

<https://erkinozalp.blogspot.com/2019/05/21-yuzylda-teknolojik-yonelimler.html>

https://en.wikipedia.org/wiki/Artificial_intelligence

<https://www.britannica.com/technology/artificial-intelligence>

J. McCarthy, M. L. Minsky, N. Rochester, C. E. Shannon, *A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence*, August 31, 1955, <http://www-formal.stanford.edu/jmc/history/dartmouth/dartmouth.html>.

Michael D. Gordin, *Scientific Babel: How Science Was Done Before and After Global English*, The University of Chicago Press, 2015, s. 215 ve 245.

Stuart Armstrong, Kaj Sotala, *How We're Predicting AI – or Failing To*, 2012, <https://intelligence.org/files/PredictingAI.pdf>.

https://commons.wikimedia.org/wiki/File:Open_knife_switch.jpg

<https://commons.wikimedia.org/wiki/File:Transbauformen.jpg>

<https://www.lezzet.com.tr/yemek-tarifleri/diger-tarifler/yumurtali-yemekler/sade-omlet>

Joseph Weizenbaum, *Computer Power and Human Reason: From Judgement to Calculation*, W. H. Freeman and Company, 1976, s. 3-7.

https://commons.wikimedia.org/wiki/File:Linear_regression.svg

https://commons.wikimedia.org/wiki/File:Anscombe%27s_quartet_3.svg

https://commons.wikimedia.org/wiki/File:Overfitted_Data.png

<https://commons.wikimedia.org/wiki/File:KnnClassification.svg>

Ian Goodfellow, Yoshua Bengio, Aaron Courville, *Deep Learning*, The MIT Press, 2016, s. 16.

A.g.y., s. 24.

Alex Krizhevsky, Ilya Sutskever, Geoffrey E. Hinton, "ImageNet Classification with Deep Convolutional Neural Networks", *NIPS'12*:

Proceedings of the 25th International Conference on Neural Information Processing Systems, Volume 1, Curran Associates Inc, 2012, <http://papers.nips.cc/paper/4824-imagenet-classification-with-deep-convolutional-neural-networks.pdf>.

<https://twitter.com/jackyalcine/status/615329515909156865>

Melanie Mitchell, *Artificial Intelligence: A Guide For Thinking Humans*, Farrar, Straus and Giroux, 2019 (epub).

²⁰ Ian Goodfellow, *a.g.y.*, s. 20.

https://commons.wikimedia.org/wiki/File:Tuerkischer_schachspieler_windisch4.jpg

²² Kevin Kelly, *The Inevitable: Understanding the 12 Technological Forces that will Shape Our Future*, Viking, 2016 (epub).

²³ M. Bongard, *Pattern Recognition*, Edited by Joseph K. Hawkins, Translated by Theodore Cheron, Spartan Books, 1970, s. 214-215.

²⁴ Douglas R. Hofstadter, *Gödel, Escher, Bach: an Eternal Golden Braid*, Basic Books, 1999, s. 646 vd.

²⁵ Terry Winograd, "Understanding Natural Language", *Cognitive Psychology*, 3, 1972, s. 33.

²⁶ Gary Smith, *The AI Delusion*, Oxford University Press, 2018, s. 45.

<https://github.com/arvindmehairjan/atari-breakout>

Melanie Mitchell, *a.g.y.*

Ken Kansky, Tom Silver, David A. Mély vd., "Schema Networks: Zero-shot Transfer with a Generative Causal Model of Intuitive Physics", 2017, <https://arxiv.org/abs/1706.04317>.

³⁰ Dawn Kawamoto, "Watson Wasn't Perfect: IBM Explains the 'Jeopardy!' Errors", *Aol.*, Feb 17th 2011, <https://www.aol.com/2011/02/17/the-watson-supercomputer-isnt-always-perfect-you-say-tomato/>

Eliza Strickland, "How IBM Watson Overpromised and Underdelivered on AI Health Care", *IEEE Spectrum*, 02 Apr 2019, <https://spectrum.ieee.org/biomedical/diagnostics/how-ibm-watson-overpromised-and-underdelivered-on-ai-health-care>.

Andrew McAfee, Erik Brynjolfsson, *Machine Platform Crowd: Harnessing Our Digital Future*, W. W. Norton & Company, 2017 (epub).

- Monya Baker, "1,500 Scientists Lift the Lid on Reproducibility", *Nature*, 25 May 2016, <https://www.nature.com/news/1-500-scientists-lift-the-lid-on-reproducibility-1.19970>.
- ³⁴ James Bridle, *New Dark Age: Technology and the End of the Future*, Verso, 2018 (epub).
Toby Walsh, *Machines That Think: The Future of Artificial Intelligence*, Prometheus Books, 2018 (epub).
- ³⁶ Donna Tam, "Google's Sergey Brin: You'll ride in robot cars within 5 years", *CNET*, September 25, 2012, <https://www.cnet.com/news/googles-sergey-brin-youll-ride-in-robot-cars-within-5-years/>.
Russell Brandom, "Tesla Wants New Self-Driving Tech to Autonomously Road Trip From LA to New York", *The Verge*, Oct 19, 2016, <https://www.theverge.com/2016/10/19/13341100/tesla-self-driving-autonomous-road-trip-la-nyc>.
- ³⁸ Melanie Mitchell, *a.g.y.*
David Sumpter, *Outnumbered: From Facebook and Google to Fake News and Filter-Bubbles – The Algorithms That Control Our Lives*, Bloomsbury Sigma, 2018 (epub).
- ⁴⁰ <http://openworm.org/>
Stanislas Dehaene, *How We Learn: Why Brains Learn Better Than Any Machine... for Now*, Viking, 2020 (epub).
Ian Goodfellow, *a.g.y.*, s. 21.
https://en.wikipedia.org/wiki/Wheat_and_chessboard_problem
- ⁴⁴ Martin Ford, *Architects of Intelligence: The Truth About AI From the People Building It*, Packt Publishing, 2018 (epub) ve Martine Rothblatt, *Virtually Human: The Promise -and the Peril- of Digital Immortality*, St. Martin's Press, 2014 (epub).
Ray Kurzweil, "The Human Machine Merger: Are We Headed for The Matrix?", *Kurzweil Library*, March 2, 2003, <https://www.kurzweilai.net/the-human-machine-merger-are-we-headed-for-the-matrix>.
J. W. Scannell, A. Blanckley, H. Boldon, B. Warrington, "Diagnosing the Decline in Pharmaceutical R&D Efficiency", *Nature Reviews Drug Discovery*, 11(3), 2012. doi:10.1038/nrd3681.
Maddie Iribarren, "IBM's Watson for Drug Discovery Program No Longer Taking New Clients", *voicebot.ai*, April 22, 2019, <https://>

voicebot.ai/2019/04/22/ibms-watson-for-drug-discovery-program-no-longer-taking-new-clients/.

- 48 Ray Kurzweil, Rachel Bernstein, "Introducing Semantic Experiences with Talk to Books and Semantris", Google AI Blog, April 13, 2018, <https://ai.googleblog.com/2018/04/introducing-semantic-experiences-with.html>.

Gary Marcus, Ernest Davis, *Rebooting AI: Building Artificial Intelligence We Can Trust*, Pantheon Books, 2019 (epub).

- 50 Andrew Griffin, "Facebook's Artificial Intelligence Robots Shut Down After They Start Talking to Each Other in Their Own Language", *Independent*, 31 July 2017, <https://www.independent.co.uk/life-style/gadgets-and-tech/news/facebook-artificial-intelligence-ai-chatbot-new-language-research-openai-google-a7869706.html>.

Gary Marcus, *a.g.y.*

"An Artificial Intelligence Predicts the Future", *The Economist*, Nov 21 2019, <https://worldin.economist.com/edition/2020/article/17521/artificial-intelligence-predicts-future>.

Carlos E. Perez, "The Paradigm Shift of Self-Supervised Learning", *Medium*, May 23, 2019, <https://medium.com/intuitionmachine/the-paradigm-shift-of-self-supervised-learning-744a6819ce08>.

- 54 Tom Standage, "How I (Sort of) Interviewed an Artificial Intelligence", *Medium*, Dec 9, 2019, <https://medium.economist.com/how-i-sort-of-interviewed-an-artificial-intelligence-2a9c069a1680>.

<https://talktotransformer.com/> Şöyle yazmışım: "If you have two apples and three oranges, and eat one of your apples, than you will have only ..." Cevaplar: 1. "... two oranges left." 2. "... two oranges. The other three oranges will remain to you." 3. "... one lemon." 4. "... two oranges. But if you have all three apples you will have all three oranges." 5. "one apple. However, if you eat the other two apples, than you will have three oranges." 6. "... three oranges, because you will have eaten your two apples." 7. "... one of your oranges.' This is called the 'Two Atonement' principle, and it is found in the New Testament." Son cümledeki dilbilgisi hatası GPT-2'ye ait.

- 56 "An Open Letter: Research Priorities for Robust and Beneficial Artificial Intelligence", *Future of Life Institute*, <https://futureoflife.org/ai-open-letter>.

Stuart Russell, Daniel, Dewey, Max Tegmark, "Research Priorities for Robust and Beneficial Artificial Intelligence", *AI Magazine*, Winter 2015, https://futureoflife.org/data/documents/research_priorities.pdf.

⁵⁸ Max Tegmark, *Life 3.0: Being Human in the Age of Artificial Intelligence*, Alfred A. Knopf, 2017 (epub).

Martin Ford, *a.g.y.*

⁶⁰ A. D. I. Kramer, J. E. Guillory, J. T. Hancock, "Experimental Evidence of Massive-Scale Emotional Contagion Through Social Networks", *Proceedings of the National Academy of Sciences*, 111(24), 2014, s. 8788–8790, doi:10.1073/pnas.1320040111.

a.g.y., s. 10779.

⁶² Fennel Aurora, "All the Ways Facebook Can Track You", 28.02.18, <https://blog.f-secure.com/all-the-ways-facebook-can-track-you/>.

⁶³ Steven Melendez, Alex Pasternack, "Here are the Data Brokers Quietly Buying and Selling Your Personal Information", *Fast Company*, 03.02.19, <https://www.fastcompany.com/90310803/here-are-the-data-brokers-quietly-buying-and-selling-your-personal-information>.

⁶⁴ Joseph Cox, "Leaked Documents Expose the Secretive Market for Your Web Browsing Data", *Motherboard*, Jan 27 2020, https://www.vice.com/en_us/article/qjdkq7/avast-antivirus-sells-user-browsing-data-investigation.

Natasha Lomas, "Avast CEO on Why It's Just Spent \$1.4BN to Absorb Security Rival AVG", *TechCrunch*, September 30, 2016, <https://techcrunch.com/2016/09/30/avast-ceo-on-why-its-just-spent-1-4bn-to-absorb-security-rival-avg/>.

⁶⁶ Frank Pasquale, *The Black Box Society: The Secret Algorithms That Control Money and Information*, Harvard University Press, 2015, s. 20. *A.g.y.*, s. 30.

⁶⁸ Sarah Vizard, "Brands Found Targeting Junk Food Ads at Children Online", *MarketingWeek*, 6 Jun 2019, <https://www.marketingweek.com/targeting-junk-food-ads-children/>.

David Cox, "Watch Your Step: Why the 10,000 Daily Goal is Built on Bad Science", *The Guardian*, 3 Sep 2018, <https://www.theguardian>.

com/lifeandstyle/2018/sep/03/watch-your-step-why-the-10000-daily-goal-is-built-on-bad-science.

⁷⁰ “John Hancock Adds Fitness Tracking to All Policies”, *BBC News*, 20 September 2018, <https://www.bbc.com/news/technology-45590293>.

Greg Farough, “Upcycle Windows 7”, *Free Software Foundation*, Jan 23, 2020, <https://www.fsf.org/windows/upcycle-windows-7>.

Kathleen Elkins, “The Age When 17 Self-Made Billionaires Earned Their First Million”, *Inc.*, Feb 11, 2016, <https://www.inc.com/business-insider/when-billionaires-made-their-first-million.html>.

⁷³ Derek Thompson, “Facebook Turns a Profit, Users Hits 300 Million”, *The Atlantic*, September 17, 2009, <https://www.theatlantic.com/business/archive/2009/09/facebook-turns-a-profit-users-hits-300-million/26721/>.

“Peter Thiel Sells Most of Remaining Facebook Stake”, *Reuters*, November 22, 2017, <https://www.reuters.com/article/us-facebook-stake/peter-thiel-sells-most-of-remaining-facebook-stake-idUSKBN1DM2BQ>.

Peter Thiel, Blake Masters, *Zero to One: Notes on Startups, or How to Build the Future*, Crown Business, 2014 (epub)

“How to Make Sense of the Latest Tech Surge”, *The Economist*, February 20th 2020, <https://www.economist.com/leaders/2020/02/20/how-to-make-sense-of-the-latest-tech-surge>.

⁷⁷ <http://www.tuik.gov.tr/HbGetirHTML.do?id=33603>

⁷⁸ Amir Efrati, “Money Pit: Self-Driving Cars’ \$16 Billion Cash Burn”, *The Information*, Feb. 5, 2020, <https://www.theinformation.com/articles/money-pit-self-driving-cars-16-billion-cash-burn>.

⁷⁹ Victoria Shannon, “Pioneer Who Kept the Web Free Honored With a Technology Prize”, *The New York Times*, June 14, 2004, <https://www.nytimes.com/2004/06/14/business/pioneer-who-kept-the-web-free-honored-with-a-technology-prize.html>.

⁸⁰ Brian Merchant, “How Corporate Delusions of Automation Fuel the Cruelty of Uber and Lyft”, *Gizmodo*, 5/09/19, <https://gizmodo.com/how-corporate-delusions-of-automation-fuel-the-cruelty-1834627595>.

Colin Lecher, “How Amazon Automatically Tracks and Fires Warehouse Workers for ‘Productivity’”, *The Verge*, Apr 25, 2019,

<https://www.theverge.com/2019/4/25/18516004/amazon-warehouse-fulfillment-centers-productivity-firing-terminations>.

- 82 Josh Dzieza, "Robots Aren't Taking Our Jobs - They're Becoming Our Bosses", *The Verge*, Feb 27, 2020, <https://www.theverge.com/2020/2/27/21155254/automation-robots-unemployment-jobs-vs-human-google-amazon>.

Casey Newton, "YouTube Moderators Are Being Forced to Sign a Statement Acknowledging the Job Can Give Them PTSD", *The Verge*, Jan 24, 2020, <https://www.theverge.com/2020/1/24/21075830/youtube-moderators-ptsd-accenture-statement-lawsuits-mental-health>.

- 84 Lente Van Hee, Denny Baert, Tim Verheyden, Ruben Van Den Heuvel, "Google Employees are Eavesdropping, Even in Your Living Room, VRT NWS Has Discovered", *VRT NWS*, 10 Jul 2019, <https://www.vrt.be/vrtnws/en/2019/07/10/google-employees-are-eavesdropping-even-in-flemish-living-rooms/>.

Matt Day, Giles Turner, Natalia Drozdiak, "Amazon Workers Are Listening to What You Tell Alexa", *Bloomberg*, 2019-04-10, <https://www.bloomberg.com/news/articles/2019-04-10/is-anyone-listening-to-you-on-alexa-a-global-team-reviews-audio>.

Sarah Frier, "Facebook Paid Contractors to Transcribe Users' Audio Chats", *Bloomberg*, 2019-08-13, <https://www.bloomberg.com/news/articles/2019-08-13/facebook-paid-hundreds-of-contractors-to-transcribe-users-audio>.

Alex Hern, "Apple Contractors 'Regularly Hear Confidential Details' on Siri Recordings", *The Guardian*, 26 Jul 2019, <https://www.theguardian.com/technology/2019/jul/26/apple-contractors-regularly-hear-confidential-details-on-siri-recordings>.

- 88 "World's Billionaires Have More Wealth Than 4.6 Billion People", *Oxfam International*, 20th January 2020, <https://www.oxfam.org/en/press-releases/worlds-billionaires-have-more-wealth-46-billion-people>.

Lawrence Mishel, Jessica Schieder. "CEO Compensation surged in 2017", *Economic Policy Institute*, August 16, 2018, <https://www.epi.org/publication/ceo-compensation-surged-in-2017/>.

- ⁹⁰ Gethin Chamberlein, "Schoolchildren in China Work Overnight to Produce Amazon Alexa Devices", *The Guardian*, 8 Aug 2019, <https://www.theguardian.com/global-development/2019/aug/08/schoolchildren-in-china-work-overnight-to-produce-amazon-alexa-devices>.
- Brian Merchant, *The One Device: The Secret History of the iPhone*, Little, Brown and Company, 2017 (epub).
- ⁹² Henry Sanderson, "Congo, Child Labour and Your Electric Car", *Financial Times*, July 7 2019, <https://www.ft.com/content/c6909812-9ce4-11e9-9c06-a4640c9feebb>.
- Annie Kelly, "Apple and Google Named in US Lawsuit Over Congolese Child Cobalt Mining Deaths", *The Guardian*, 16 Dec 2019, <https://www.theguardian.com/global-development/2019/dec/16/apple-and-google-named-in-us-lawsuit-over-congolese-child-cobalt-mining-deaths>.
- ⁹⁴ Lee Chyen Yee, Clare Jim, "Foxconn to Rely More on Robots; Could Use 1 Million in 3 Years", *Reuters*, August 1, 2011, <https://www.reuters.com/article/us-foxconn-robots/foxconn-to-rely-more-on-robots-could-use-1-million-in-3-years-idUSTRE77016B20110801>.
- Kathrin Hille, "Foxconn: Why the World's Tech Factory Faces Its Biggest Tests", *The Guardian*, June 10 2019, <https://www.ft.com/content/0f57109e-8845-11e9-a028-86cea8523dc2>.
- ⁹⁶ Martin Ford, *a.g.y.*
- ⁹⁷ Ken Jacobs, Ian Eve Perry, Jenifer MacGillvary, "The High Public Cost of Low Wages", *UC Berkeley Labor Center*, April 13, 2015, <http://laborcenter.berkeley.edu/the-high-public-cost-of-low-wages/>.
- ⁹⁸ Abha Bhattarai, "Bernie Sanders Introduces 'Stop BEZOS Act' in the Senate", *The Washington Post*, September 5, 2018, <https://www.washingtonpost.com/business/2018/09/05/bernie-sanders-introduces-stop-bezos-act-senate/>.
- Nick Srnicek, Alex Williams, *Inventing the Future: Postcapitalism and a World Without Work*, Verso, 2015 (epub).
- ¹⁰⁰ Thomas W. Malone, *Superminds: The Surprising Power of People and Computers Thinking Together*, Little, Brown and Company, 2018 (epub).

William H. Davidow, Michael S. Malone, *The Autonomous Revolution: Reclaiming the Future We've Sold to Machines*, Berrett-Koehler Publishers, 2020 (epub).

102 Garry Kasparov, *Deep Thinking: Where Machine Intelligence Ends and Human Creativity Begins*, Public Affairs, 2017 (epub).

103 Noam Cohen, "How Wikipedia Prevents the Spread of Coronavirus Misinformation", *Wired*, 03.15.2020, <https://www.wired.com/story/how-wikipedia-prevents-spread-coronavirus-misinformation/>.

104 https://en.wikipedia.org/wiki/Wikipedia:The_community

105 <https://en.wikipedia.org/wiki/Wikipedia:Administration>

106 https://en.wikipedia.org/wiki/Wikipedia:Bot_policy

107 <https://www.openhub.net/>

108 <https://www.linuxfoundation.org/membership/members/>

109 <https://www.linuxfoundation.org/press-release/2016/08/the-linux-foundation-releases-development-report-highlighting-contributions-to-the-linux-kernel-ahead-of-25th-anniversary-of-linux/>

110 <https://www.top500.org/statistics/details/osfam/1>

"Google Data Center FAQ", *Data Center Knowledge*, Mar 17, 2017, <https://www.datacenterknowledge.com/archives/2017/03/16/google-data-center-faq>.

112 Erik Olin Wright, *How to Be an Anticapitalist in the Twenty-First Century*, Verso, 2019 (epub).

Paul N. Edwards, *The Closed World: Computers and the Politics of Discourse in Cold War America*, The MIT Press, 1996, s. 43.

114 a.g.y., s. 283.

Alex Roland, Philip Shiman, *Strategic Computing: DARPA and the Quest for Machine Intelligence, 1983-1993*, The MIT Press, 2002, s. 1.

116 Klint Finley, "How GitHub Is Helping Overworked Chinese Programmers", *Wired*, 04.04.2019, <https://www.wired.com/story/how-github-helping-overworked-chinese-programmers/>.

<https://www.credit-suisse.com/about-us/en/reports-research/global-wealth-report.html>

118 <https://citrusvanilla.github.io/billionaires/>

- ¹¹⁹ “Understanding China’s Social Credit System”, *Trivium Social Credit*, Last update: August 27, 2019, <http://socialcredit.triviumchina.com/what-is-social-credit/>.
- ¹²⁰ https://en.wikipedia.org/wiki/Censorship_of_Wikipedia#China
Scott Sullivan, *Designing for Wearables: Effective UX for Current and Future Devices*, O’Reilly, 2017, s. 53-59.

Dizin

- 2001: Bir Uzay Destanı (2001: A Space Odyssey) 14
ABD Savunma Bakanlığı 12 107
Acxiom 78
açık kaynaklı yazılım (*open source software*) 83 103 104 112
Alexa 91
Alphabet85
Amazon 46 47 85 89-91 93 94 97
Android 75 104
Apple 85 87 91 93 94
ARPANET 107
Avast 78
AVG 78
Avrupa Birliği 51 82
Ben, Robot (I, Robot) 27
Berners-Lee, Tim 87
beyin 38 63 64 67
Bezos, Jeff 97
bilimkurgu 13 14 30 68
Bongard problemleri 49 50 60
Bongard, M 49
bot (internet robotu) 102 103
Brin, Sergey 61
Britannica 11
Brooks, Rodney 95
Brynjolfsson, Erik 57
büyükveri (*big data*) 45 48 51 57 58 109
C. elegans 63 64
CIA 12
Cyc 60
çalışma 89 90 94 96-98 104 108 112 113
Çin 74 93-95 108 109
çöp girerse çöp çıkar (*garbage in, garbage out-GIGO*) 58
Davis, Ernest 68
Deep Blue 100
DeepMind 53
Dell 94
derin öğrenme (*deep learning*) 38 39 41-45 48 52 69
doğrusal regresyon (*linear regression*) 32-35
Dünya Çapında Ağ (*World Wide Web*) 87
ELIZA (etkisi) 30 31 71
en yakın k komşu algoritması (*k-nearest neighbors algorithm*) 37
Eroom Yasası 67
evrensel temel gelir (*universal basic income*) 96 97 105 106
Ex Machina 30
Excel 83
Eyser, George 55

- Facebook 47 48 69 73 75-77 84 85
90 91 102
- fikrî mülkiyet hakları 104
- Foxconn 93-95 108
- Gates, Bill 82 83
- Gauss Yöntemi 24
- General Motors 86
- Google 12 29 44 45 48 51 53 61
67 68 73 75 85 86 91 94 101
103-105 110 111
- Google Çeviri 12 51
- Google Gözlük (*Google Glass*) 110
- görüntü tanıma (*image
recognition*) 39 42 45
- gözetimli öğrenme (*supervised
learning*) 41
- gözetimsiz öğrenme
(*unsupervised learning*) 71
- GPT-2 71-73
- Hawking, Stephen 66 73
- Hinton, Geoffrey 39 41
- Hofstadter, Douglas R. 49
- IBM 54-56 68 100 103
- IKEA 95
- Instagram 84
- Intel 103
- insanötesicilik (*transhumanism*)
67
- iPhone 93
- işsizlik 88 89 96 97 106
- Jeopardy!* 54-56
- Kaggle 111
- Kasparov Yasası 100
- Kasparov, Garry 100
- kaynak kodu (*source code*) 83 103
- kobalt 94 108
- kolektif zekâ (*collective
intelligence*) 63 98 99 101 103
- Kongo Demokratik Cumhuriyeti
94 108
- konuşma tanıma (*speech
recognition*) 42 43 91 92
- kooperatif 105
- Kurzweil, Ray 67 68
- Linux 83 103 104
- Linux Vakfı 103
- Lyft 89
- makine çevirisi (*machine
translation*) 13 42 51
- makine öğrenmesi (*machine
learning*) 32-34 36 38 39
68 111
- Marcus, Gary 68
- McAfee, Andrew 57
- McDonald's 97
- medya 14 30 39 55 69
- Mekanik Türk (*Mechanical Turk*)
46 47
- Microsoft 31 70 73 82-85 87 94
103
- Moore Yasası 66 67
- MS-DOS 82 83
- Musk, Elon 61 73 85
- Ng, Andrew 74
- Open Hub 103
- Openworm 64
- özgözetimli öğrenme (*self-
supervised learning*) 72
- Özgür Yazılım Vakfı (*Free
Software Foundation*) 83

- patentler 87 88 104
- PayPal 84
- pekiştirmeli öğrenme
(*reinforcement learning*)
52-54
- robot 11 26 27 31 47 54 71 90
94-96
- sağduyu (*common sense*) 59 60
62 63
- sohbet robotu (*chatbot*) 29 69 70
71 92
- Sosyal Kredi Sistemi (Çin) 109
- sosyal medya 80
- Sovyetler Birliği 13
- SpaceX 61
- Spotify 43
- Sputnik (uydu) 13
- Srnicek, Nick 97
- Standage, Tom 72
- süper zekâ (*superintelligence*)
73 74
- sürücüsüz otomobil (*self-driving car*) 61-63 86-88
- Talk to Books* (Kitaplarla Konuş)
68 69
- tam otomasyon (*full automation*)
97
- Tay (sohbet robotu) 70 71
- Tegmark, Max 73
- Tesla 61
- The Economist* 71 72
- The Matrix* 67
- Thiel, Peter 84
- transistör 16-20 66 67
- Turing Testi 29 30 49 50
- Turing, Alan 29
- Twitter 48 70-72 102
- Uber 86 89
- uzun kuyruk (*long tail*) 61 62
- Watson, IBM 54-56 68
- Weizenbaum, Joseph 30 31
- WhatsApp 84
- Wikipedia 11 54 88 101-104 109
112
- Williams, Alex 97
- Windows 83
- Winograd şemaları 50 51 61
- Winograd, Terry 50
- Word 31 83
- Wright, Erik Olin 105
- yanlılık (*bias*) 34 44
- yapay genel zekâ (*artificial general intelligence-AGI*) 29 74
- yapay sinir ağı (*artificial neural network*) 38-45 52 64 70
- yapay zekâ etkisi (*AI effect*) 28
- yapay zekâ kışları (*AI winters*) 107
- Yaşamın Geleceği Enstitüsü
(*Future of Life Institute*) 73
- YouTube 42 61 75 90 95
- Zuckerberg, Mark 84

Yapay zekâ son yıllarda (yeniden) popülerlik kazanmasını neye borçlu? Makine öğrenmesi nasıl gerçekleşiyor? İnsanların yaptığı her şeyi makinelere yaptırmak mümkün mü? Bunu istemeli miyiz? Google, Facebook, Amazon gibi şirketler yapay zekâyı hangi amaçlarla kullanıyor? Wikipedia'nın farkı ne? Yapay zekânın insanlığa daha fazla yarar sağlaması için neler yapılmalı? Kolektif zekâ neden önemli?

Erkin Özalp, bu kitabında, yapay zekâ alanındaki güncel gelişmeleri, yapay zekâ ile insan zekâsı arasındaki farkları, teknoloji şirketlerinin tercihlerini ve yapay zekânın geleceğini ele alıyor.

Makine çevirisi, sürücüsüz otomobil teknolojileri, robotlar, botlar, kişisel bilgilerimizin toplanma ve kullanılma yöntemleri, yeni çalışma biçimleri, teknolojik gelişmelerin işsizliğe yol açmasının nedenleri, evrensel temel gelir ve çalışmanın geleceği gibi konular üzerinde de duran Özalp'e göre, insanlarla makineler arasındaki ilişkileri bugünkünden farklı şekillerde kurmak mümkün ve gerekli.

Gençlerle baş başa vererek hazırlanmış olan bu çalışma, yapay zekâya ilgi duyan her yaştan okura hitap ediyor.



16 TL. KDV'den muaf tır

ISBN 978-605-172-379-2



9 786051 723792



Yordam Kitap

[YordamKitap](#) [YordamKitap](#) [YordamKitap](#)