

Yapay Genel Zekâya Giden Teknoloji, Tekillik ve Etik Kaygılar

Gürkan Caner Birer [Bilgisayar Mühendisi

Her geçen gün insanların sahip olduğu becerilerle kıyaslanabilir düzeyde özellikler kazanan yapay zekâ uygulamalarının gelecekte daha da yaygınlaşması, önemli etik sorunları da beraberinde getirebilir.

Bir yapay zekâ uygulamasının aldığı kararlar ile bazı insanları korurken diğerlerinin zarar görmesine neden olabileceği çeşitli durumlar ortaya çıkabilir. Örneğin, sürücüsüz bir arabanın bir yayaya çarpmaktan kaçınmasının tek yolunun karşı şeride sapmak olduğunu ancak bu durumda karşı şeritteki daha fazla yolcusu olan bir arabaya çarpacağını düşünün. Yapay zekânın vereceği bu kararın etik sonuçları olacaktır. Bu nedenle yapay zekâ sistemlerinin yalnızca teknik açıdan değil, aynı zamanda etik açıdan da doğru otonom kararlar almalarını sağlayacak şekilde eğitilmeleri esastır.







Etik, insanların başkalarına ve çevreye ilişkin davranışlarını düzenlemek için tanımlanmış bir dizi kriterdir ve eylemleri, diğer canlılar veya çevre üzerindeki iyi ve kötü etkilerine göre değerlendirir. Ancak çevremizdeki akıllı makineler yaygınlaştıkça bu sistemlerin karar alırken insanlara ve çevreye olan etkilerini düzenleyen kurallara olan ihtiyaç artıyor. Peki nasıl “etik robotlar” üretebiliriz?

1940’larda Biyokimya Profesörü ve tanınmış Bilim Kurgu Yazarı Isaac Asimov, robotlar için üç etik yasa belirledi:

Birinci Yasa: Bir robot, bir insana zarar veremez veya eylemsizlik yoluyla bir insanın zarar görmesine izin veremez.

İkinci Yasa: Bir robot, insanlar tarafından verilen emirlere -emirler birinci yasa ile çelişmediği sürece- uymalıdır.

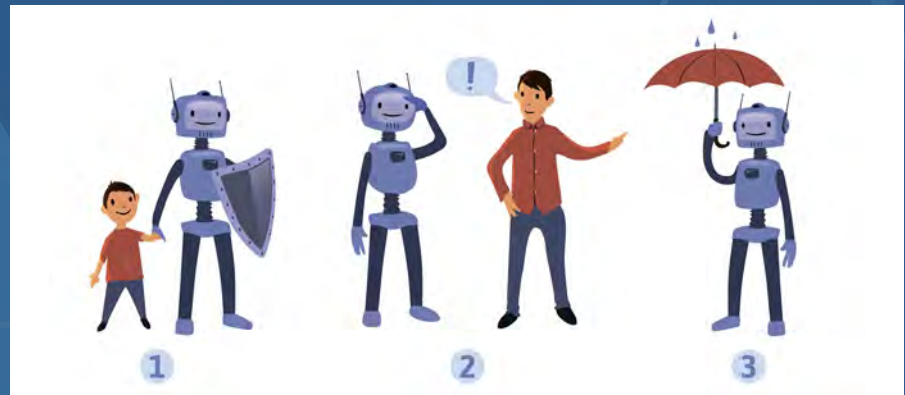
Üçüncü Yasa: Bir robot, birinci veya ikinci yasa ile çelişmediği sürece kendi varlığını korumalıdır. Görünüşte açık olan bu yasalar, aslında birçok belirsizlik içerir. Örneğin zarar kavramı, daha da belirsiz olan kötülük kavramıyla ilişkilidir ancak kötülük, sadece fiziksel olmayabilir.

Bu tür belirsizliklerin farkında olan Asimov, daha sonra diğer yasalardan çok daha önemli olan “Sıfırıncı Yasa”yı ekledi:

Sıfırıncı Yasa: Bir robot, insanlığa zarar veremez veya eylemsizlikle insanlığın zarar görmesine izin veremez.

Konu akıllı makineler olunca etik sorunlar, bir dizi kurala sıkıştırılamayacak kadar karmaşık olabilir. Dahası doğal dil işlemedeki, yani bilgisayarların insanların konuştuğu dili anlama sürecindeki belirsizlikler, beklenmedik yorumlara yol açarak insan hayatı için yüksek riskli durumlar oluşturabilir. Bu tür sorunlar, amacı makinelerle uygun kararlar almaları için araçlar sağlamak olan “makine etiği” veya “roboetik” olarak bilinen yeni bir araştırma alanını ortaya çıkardı. Ancak bu da kaçınılmaz olarak yeni sorulara yol açtı: Bir bilgiyi, etik olarak doğru kararlar almasını sağlayacak şekilde bir akıllı makineye nasıl kodlayabiliriz? Peki etik robotlara gerçekten güvenebilir miyiz?

Etik robotların geliştirilmesi, bilgisayar bilimcileri ve filozoflar arasında güçlü bir etkileşim gerektirir. Bu amaçla iki yaklaşım kullanılır: Birincisi, algoritmik bir türdür ve bu yaklaşımda programda etik kuralları açıkça tanımlayan ve alınan kararların etik sonuçlarını ölçerek doğru kararı almasını sağlayan (maliyet fonksiyonunu maksimize etmek olarak tanımlanır)



bir dizi kural kodlanır. İkincisi, yapay sinir ağlarının kullanıldığı sinirsel yaklaşım türüdür ve çeşitli örneklerle ait veri tabanları kullanılarak makine öğrenimi teknikleri ile yapay zekâ sistemlerine etik davranışlar öğretilir.



Ne yazık ki her iki yaklaşımın da zayıf yönleri bulunur. Algoritmik yaklaşımda sadece tanımlanan ahlaki kurallara göre karar alınır, bu nedenle birden fazla etik kuralın aynı anda geçerli olduğu karmaşık durumlarda doğru kararlar alınması zorlaşır. Ayrıca bunları her durumda uygulamak risklidir. Örneğin “kırmızı ışıktadır” kuralı, ambulans ya da itfaiye araçlarına yol verilmesi gereken acil durumlarda yanlış sonuçlar doğurabilir. Bu kuralların başarısız olacağı birçok istisna durumu hayal etmek kolaydır. Ayrıca çözülemeyen kural çatışmalarının olduğu belirli durumlarda tüm sistem, çıkmaza girebilir. Örneğin otonom bir araç, hem yolcuları hem yayaları koruması gereken durumlarda hangi kuralı uygulayacağına karar veremeyebilir.

Öte yandan sinirsel yaklaşımın farklı sorunları vardır. Durumun bağlamını dikkate alarak kararlar alan bu yöntemde öğrenilenler, milyonlarca parametrede kodlandığından robotun belirli bir kararı neden aldığını anlamak kolay değildir. Dahası modelin öğrenmesinde kullanılan veri setlerinde yer alan ön yargılar, karar verme mekanizmalarını yanlış yönlendirebilir. Bu sorunlara çözüm bulmak için makine öğreniminin “açıklanabilir yapay zekâ” olarak adlandırılan yeni bir araştırma alanı bulunur. Bu alan, yapay sinir ağları tarafından oluşturulan çıktılar için anlaşılır yorumlar sağlamayı ve veri setindeki olası ön yargıları tespit etmeyi amaçlar.

Derin öğrenmedeki öncü çalışmalarından dolayı sıklıkla “yapay zekânın babalarından” biri olarak kabul edilen 2024 Nobel Fizik Ödülü’nü kazanan Geoffrey Hinton, YGZ’nin potansiyel tehlikeleri konusunda uyarılarda bulunması ile tanınıyor. Hinton, makinelerin daha zeki hâle geldikçe belirli alanlarda

insan yeteneklerini aşabileceği ve etkili bir şekilde kontrol edilmezlerse öngörülemeyen risklere yol açabileceğini düşünüyor.

BBC ile yaptığı bir röportajda Hinton, “Yapay zekânın aslında insanlardan daha akıllı olabileceği fikrine az sayıda kişi inanıyordu. Çoğu kişi bu durumun çok uzak bir ihtimal olduğunu düşünüyordu. Ben de başlarda bu şekilde düşünüyordum. Süper zeki makinelerin ortaya çıkması için önümüzde en az 30 ila 50 yıl olduğunu öngörüydüm. Açıkçası artık böyle düşünmüyorum.” diyor. Hinton’un bakış açısındaki bu değişim, yeni nesil teknolojilerdeki benzeri görülmemiş hızdaki ilerlemelere bağlı olarak günümüzün yapay zekâ toplumunda, etik hususlara ve güvenlik önlemlerine yönelik ihtiyaçlar ve düzenlemeler konusunda daha geniş bir farkındalığın olması gerektiğini yansıtıyor.



Geoffrey Hinton





Tekillik

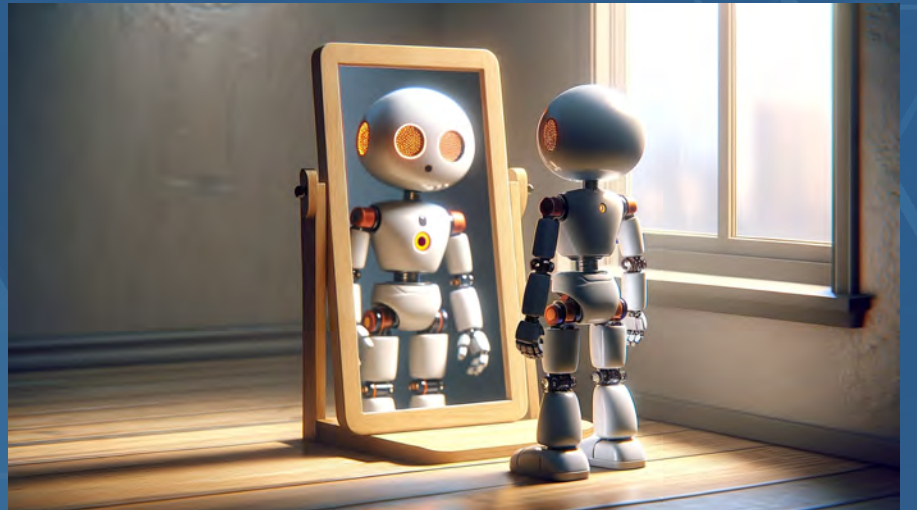
Yapay genel zekâ (YGZ), kimi zaman tekillik (singularity) kavramıyla da ifade edilir. Bu kavram, makinelerin yalnızca belirli alanlarda değil, tüm insan faaliyetlerinde insan zekâsının kapasitesine ulaştığı seviyeyi ifade eder. İnsan zekâsıyla aynı becerilere sahip bir makine hayal etmek kolaydır çünkü zihinsel olarak sağlıklı bir yetişkinin entelektüel yeteneklerini referans alabiliriz. Ancak bu yeteneklere sahip bir makinenin bir insandan çok daha yüksek hızlarda çalışabildiğini yani işleme hızı, bellek kapasitesi, duyuşal verileri edinme ve yüksek hızlı ağlar aracılığıyla internette bulunan diğer verilere erişme yeteneği açısından bakıldığında bir insandan daha avantajlı olduğunu düşünebiliriz. Ayrıca akıllı makineler diğer benzer sistemlerle hızlı bir şekilde iletişim kurarak bilgisini paylaşabilir ve diğer makinelerin deneyimlerinden hızlıca yeni bilgiler öğrenebilir.

Bilgi aktarımına dayalı kolektif öğrenme, elbette insanlar arasında da gerçekleşir ancak bu süreç, insanlarda çok daha yavaştır.

Yapay zekâ teknolojisinin üstel bir hızda geliştiği düşünülüyor. Üstel bir sürecin tipik örneği olarak Moore yasasını verebiliriz. Intel'in kurucu ortağı olan Gordon Moore, 1965'te elektronik çiplerin gelişim sürecini gözlemleyerek elektronik devredeki transistör sayısının her 12 ayda bir, ikiye katlanacağı hipotezini ortaya attı. 1980'lerin sonlarında yasa, ikiye katlanma süresini 18 (bazı kaynaklarda 24) aya çıkararak değiştirdi. Elektronik devrelerdeki transistör sayısının üstel olarak arttığını tanımlayan bu yasa, bugün hâlâ geçerliliğini koruyor.

Vernor Vinge, Hans Moravec, Nick Bostrom, Giorgio Buttazzo ve Ray Kurzweil gibi yapay

zekâ uzmanları, yapay zekâ teknolojilerinde bugüne kadar gözlemlenen üstel gelişim sürecine dayanarak bilişsel yeteneklerin gelişiminin yakın gelecekte aynı hızda devam edeceğini varsayıyor ve tekilliğe 2030 yılı civarında ulaşılabilirliğini öngörüyor. Uzmanlar, bu tarihi hesaplarken canlılardaki zekâ gelişim hızını da dikkate alarak geliştirdikleri birtakım modellerden faydalanmış. Ancak bu tür tahminlere biraz temkinli yaklaşarak daha gerçekçi bir tahminde bulunmak gerekebilir. Çünkü mevcut makine öğrenimi modellerinin, otonom arabalar, uçaklar, trenler ve robotlar gibi güvenlik açısından kritik sistemlerde kullanılmadan önce mutlaka üstesinden gelinmesi gereken eksiklikleri bulunuyor. Örneğin yapay zekâ algoritmaları, modellerin eğitildiği veri setinin eksik ya da yetersiz kaldığı sıra dışı ya da nadir durumlarla karşılaştığında



beklenmedik ya da hatalı kararlar verebilir. Ayrıca yapay sinir ağları, giriş verilerinde yapılan fark edilemeyecek kadar küçük değişikliklerle yanlış bir karar almasına neden olabilen siber saldırılara karşı da hassastır. Son olarak yapay zekâ teknolojilerinin toplumsal etkisi göz önüne alındığında bu teknolojilerin pratikte uygulanması, bazı yasalar tarafından düzenlenmelidir. Bu durum da göz önünde bulundurulduğunda tekelliğe ulaşmanın çok yakın bir gelecekte olmayacağı söylenebilir.

Bir makinenin insan zekâsına ulaşması, bir bilince sahip olacağı anlamına gelmez. Ancak dil, algılama, muhakeme ve öğrenme için gelişmiş kapasitelere sahip karmaşık bir sistemin, bir tür öz farkındalık geliştirebileceği de göz ardı edilmemelidir. Eğer bu gerçekleşirse bir makinenin gerçekten bilinçli olduğunu nasıl doğrulayabiliriz? Düşünen bir varlığın farkındalık seviyesini ölçmek için bir test var mıdır?

Zekâ, eylemler yoluyla kendini dışarıdan gösteren ve bu nedenle belirli testler yoluyla ölçülebilen bir kapasite olmasına rağmen düşünen bir varlığın zihninde bir bilincin varlığını belirlemek çok daha hassas bir konudur. Aslında öz farkındalık, yalnızca ona sahip olanlar tarafından



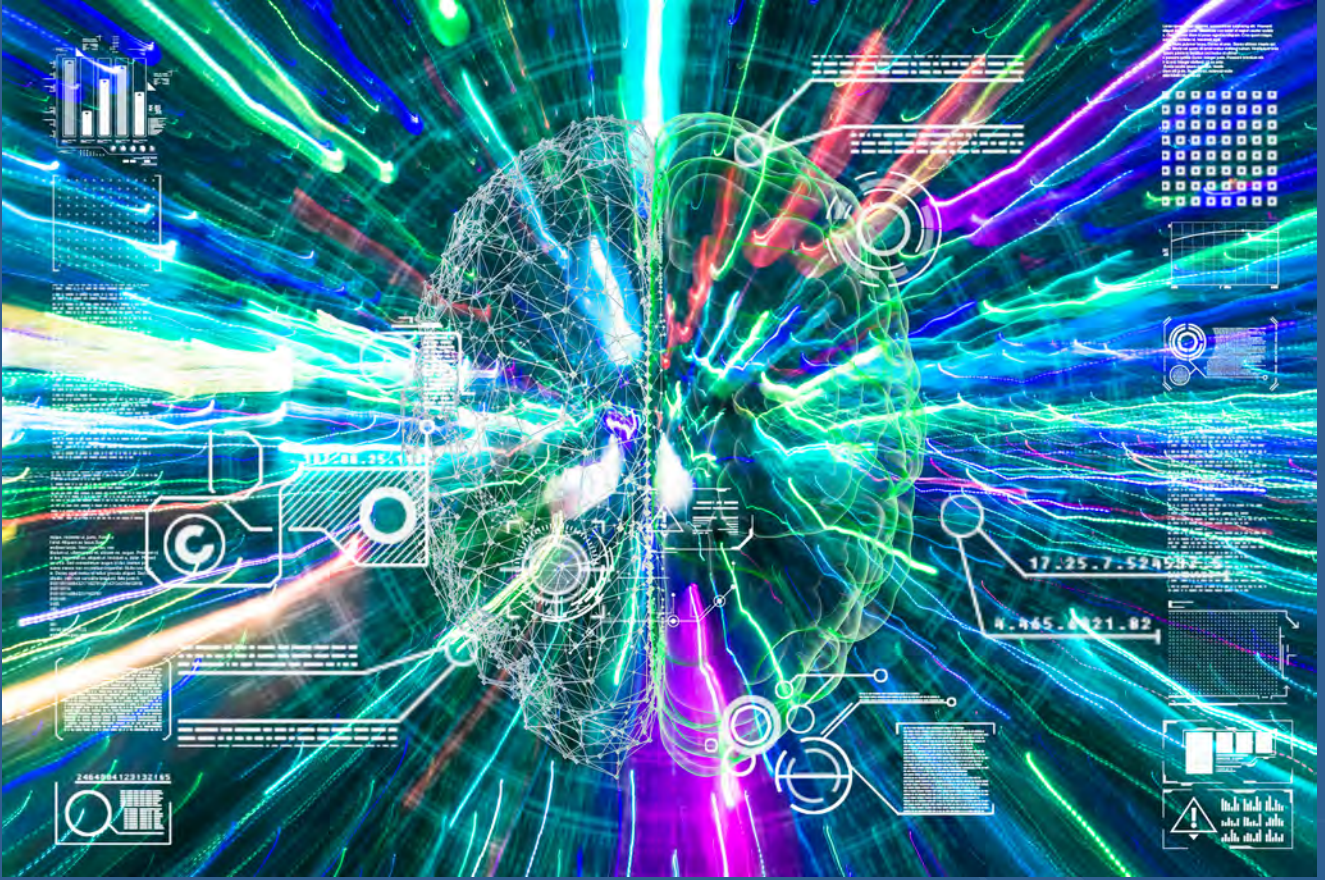
fark edilebilen bir özelliktir. Bu nedenle bilişsel yetenekleri gelişmiş akıllı bir makinenin gerçekten bir öz farkındalığa sahip olup olmadığını veya sahipmiş gibi davranıp davranmadığını belirlemeyi sağlayacak bir prosedürün nasıl tanımlanabileceği açık değildir. Bu nedenle bir yapay zekâda bilinç belirtileri ortaya çıkarsa şu anda bunu doğrulamanın bir yolu yok gibi görünüyor.

Makineler, insanların entelektüel kapasitesine ulaştığında bilgiyi işleme hızları daha yüksek olacağı, daha büyük bellek kapasitesi ile donatılmış olacakları ve internette bulunan tüm verilere hızlı bir şekilde erişebilecekleri için bir adım önde olacaklardır. Bu nedenle tekelliğe ulaşıldığında yapay zekâ, o seviyede kalmayacak, “yapay süper zekâ” olarak adlandırılan düzeye doğru üstel bir hızla gelişmeye devam edecektir.

1965’te, İngiliz Matematikçi Irving J. Good, ultra zeki bir makinenin giderek daha iyi makineler tasarlayabileceğini ve insanları çok geride bırakacak bir “zekâ patlaması”nı tetikleyebileceğini belirterek yapay süper zekânın gelişini öngördü. Bu nedenle “İlk ultra zeki makine, insanın yapması gereken son icat olacak.” sonucuna vardı.

Peki henüz tam olarak anlayamadığımız ve kendini geliştirmek için çok fazla enerjiye ihtiyaç duyan süper zeki bir sistemle yaşamak mümkün müdür?

Bugün makineler, bize bağımlı ancak bu durum, gelecekte değişebilir. Tekelliğe ulaşıldığında otonom ve süper zeki hâle gelen makineler, ihtiyaç duydukları enerjiyi elde



XH4D / iStock

etmek, kendilerini onarmak, inşa etmek ve geliştirmek için artık bize bağımlı olmak zorunda kalmayabilir.

İnsanlar ve makineler arasındaki entegrasyon olasılığı, Bilim Kurgu Yazarı Ray Kurzweil'in *Tekillik Yakın* adlı kitabında derinlemesine ele alınmış bir konu. Aslında son birkaç yıldır robotik ve bilgisayar teknolojileri, insan vücuduyla entegre bir şekilde kullanılıyor. Bozulan kalp ritmini düzenlemek için vücuda yerleştirilen kalp pilini, görme kusurlarını telafi etmek için kullanılan yapay retinayı, diyabet hastalarına

doğru miktarda insülin enjekte etmek için deri altına yerleştirilen mikro pompaları, farklı tipteki robotik protezleri ve epileptik atakları önlemek veya Parkinson hastalığının etkilerini hafifletmek için beyne yerleştirilen mikroçipleri düşünün.

Bu teknolojiler, gelecekte daha da gelişerek duysal, motor ve hafıza kapasitelerimizi artırmak için de kullanılabilir. Robotik ve biyonomik mühendislik alanındaki gelişmeler, vücudu daha dayanıklı ve uzun ömürlü yapan yapay organların ve uzuvların geliştirilmesini mümkün kılarken nanoteknolojideki

gelişmeler sayesinde nöronların bilgi işlem cihazlarıyla iletişim kurmasını sağlayacak biyoyumlu sinirsel arayüzler geliştirilebilir.

Aynı zamanda nörobiyolojinin gelişimi, insan beyninin nasıl çalıştığını ayrıntılı olarak anladığımız noktaya ulaştığında biyolojik bir beyin, bir bilgisayarla iletişim kurabilir. Bu durum, düşüncelerin, duyuların ve komutların kablosuz ağlar aracılığıyla iletilmesini mümkün kılacak diğer bireylerle ve bilgisayar sistemleri ile bir tür "telepatik bağlantı" kurulmasını sağlayabilir.

Yapay Genel Zekâya Giden Teknoloji

Yapay genel zekâ, birkaç temel teknolojik alandaki önemli gelişmelerle desteklenir. YGZ'ye giden yolda en önemli kaynak, büyük veridir. Muazzam veri kümelerini toplama, işleme ve öğrenme yeteneği, insan benzeri genel zekâya sahip sistemler geliştirmek için hayati önem taşır. Veri üretme ve depolama kapasitemiz katlanarak artmaya devam ettikçe makinelerin, bu bilgiler arasında anlamlı ilişkiler kurma ve bu ilişkilerden uygulanabilir sonuçlar çıkarma potansiyeli de artar.

Bu verileri işlemek için ise gözetimsiz öğrenme (verilerin insanlar tarafından etiketlenmediği makine öğrenmesi türüdür) ve transfer öğrenme (bir görev için daha önce eğitilmiş bir modelin, önceki görevle ilişkili başka bir görevde kullanıldığı makine öğrenmesi yöntemidir) gibi alanlarda daha verimli ve etkili algoritmaların geliştirilmesine ihtiyaç duyuluyor. Bu teknolojiler, YGZ sistemlerinin büyük miktarda veriyi işlemesini, kalıpları yani veriler arasındaki bağlantıları tanımasını ve insan muhakemesini başarılı bir şekilde taklit eden kararlar almasını sağlar. Yapay sinir ağlarını kullanan derin öğrenme

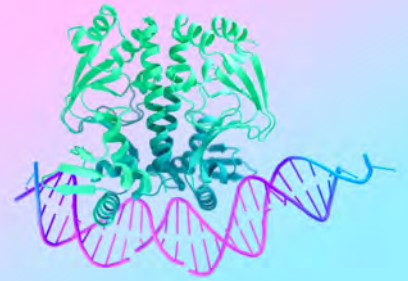
yöntemlerinin gücündeki ve karmaşıklığındaki üstel büyüme sayesinde bu sistemler, giderek daha karmaşık hâle geliyor ve yapay zekânın günümüzdeki sınırlarının ötesine geçerek daha genel zekâ biçimleri gerektiren görevleri gerçekleştirebilecek yeteneklere sahip oluyor.

Yapay genel zekâyı destekleyen diğer bir önemli gelişme ise kuantum hesaplama. Kuantum bilgisayarlar, klasik bilgisayarlardan çok daha yüksek hızlarda hesaplamalar yapma potansiyeline sahiptir. Bu hesaplama gücü, özellikle yapay zekâ modellerini eğitmek ve geliştirmek için gereken devasa veri kümelerini işlemeye ve analiz etmeye imkân sağlar. Özellikle kuantum makine öğrenme algoritmaları, karmaşık sinir ağlarının eğitimini önemli ölçüde hızlandırma ve klasik bilgisayarların yeteneklerinin çok ötesinde veri kümelerinin işlenmesini sağlama konusunda umut vadediyor. Ancak büyük ölçekli kuantum bilgisayarların geliştirme sürecinde hâlâ başlangıç aşamasında olduğumuzu ve kuantum bilgisayarların tam potansiyellerine ulaşabilmesi için önemli teknik zorlukların üstesinden gelinmesi gerektiğini belirtmemiz gerek.

Büyük verinin, gelişmiş algoritmaların ve yüksek performanslı bilgi işleme

kapasitesinin birlikte kullanılması YGZ'nin gelişim ve değişim süreci için sağlam bir temel oluşturacaktır.

YGZ'nin gelişiminde birkaç önemli araştırma kurumu ve şirket öncü çalışmalar yürütüyor. ChatGPT de dahil olmak üzere büyük dil modelleri geliştirilmesiyle bilinen OpenAI, YGZ geliştirmeye en yakın firmalardan biri olarak görülüyor. Microsoft'un araştırma ve geliştirme birimi, makine öğrenimi, doğal dil işleme ve robotik dahil olmak üzere çok çeşitli yapay zekâ konularında araştırmalar yürütüyor. Benzer şekilde Google ve Facebook, yapay zekâ alanında başarılı çalışmalar yapıyor. Bu firmalar, yapay zekâ çalışmalarına ayırdıkları yüksek bütçelerle bu alanda önemli gelişmelere liderlik edebiliyor. Üniversitelere bağlı araştırma kurumlarında da bu alanda öncü çalışmalar yürütülüyor. Stanford Üniversitesine bağlı Stanford İnsan Merkezli Yapay Zekâ Enstitüsü, topluma fayda



sağlayan yapay zekâ teknolojileri geliştirmeye odaklanan bir araştırma merkezidir. MIT'nin Bilgisayar Bilimi ve Yapay Zekâ Laboratuvarı ise makine öğrenimi, robotik ve bilgisayar görüşü gibi alanlara odaklanan yapay zekâ araştırmaları için önemli bir merkezdir.

Bu kurumlar, yapay zekâ alanında kayda değer ilerleme gösteren çok sayıda önemli araştırma projesine imza attı. Örneğin, Google'a ait DeepMind'ın AlphaFold projesi, onlarca yıldır bilim insanlarının üzerinde çalıştığı bir konu olan, proteinlerin üç boyutlu yapılarını kazandığı katlanma probleminin çözümünde devrim yarattı ve bu katkı, 2024 Nobel Fizik Ödülü'ne layık görüldü. Bu tür projeler, yalnızca yapay zekânın yeteneklerini ilerletmekle kalmıyor, aynı zamanda önemli bilimsel problemlerin çözümüne de değerli katkılar sağlıyor.

Yapay Genel Zekâ Geliştirmedeki Zorluklar

Yapay genel zekânın geliştirilmesi, birçok zorluğu da barındırıyor. Yapay genel zekâ sistemleri oluşturmak, yalnızca daha fazla veri ve daha hızlı işlemciler değil, aynı zamanda zekâyı nasıl modelleyeceğimiz



ve anlayacağımız konusunda yeni yaklaşımlar gerektiriyor. Mevcut yapay zekâ sistemleri, YGZ'ye ulaşmak için kritik bir gereklilik olan bilgiyi farklı alanlarda genelleştirmede başarılı değil. Bu sınırlamanın üstesinden gelmek, nasıl öğreneceğini öğrenme olarak tanımlanabilecek meta öğrenme ve veriler arasındaki nedensel bağlantıları kurmayı sağlayan nedensel akıl yürütme gibi alanlarda temel atılımlar gerektiriyor.

Teknik zorlukların ötesinde YGZ'nin geliştirilmesi, veri gizliliği ve güvenliği ile ilgili önemli endişelere de neden oluyor. Yapay zekâ sistemleri daha da geliştikçe eğitim için genellikle hassas kişisel bilgiler de dahil olmak üzere giderek daha büyük veri kümelerine ihtiyaç duyacak. Bu verilerin

gizliliğini ve güvenliğini sağlarken yine de yapay zekâ geliştirmede kullanımına izin vermek, dikkatli değerlendirme ve sağlam güvenlik önlemleri gerektiren karmaşık bir zorluk.

Yapay genel zekânın potansiyelini düşündüğümüzde, hayatımızın hemen her alanını dönüştürme gücüne sahip teknolojik bir devrimin eşliğinde olduğumuz açıkça ortaya çıkıyor. Çığır açan bilimsel araştırmalardan iklim değişikliği ve bulaşıcı hastalıklar gibi küresel zorluklara çözümler sunmaya kadar YGZ'nin sağlık, eğitim, iş ve günlük yaşam gibi alanlardaki potansiyel uygulamaları çok çeşitli.

Ancak bu potansiyelin önünde önemli engeller var. YGZ'nin geliştirilmesi için ciddi bir teknolojik zorluk bulunuyor.

İhtiyaç duyulan veri, işlem gücü ve sosyal destek, bugün için karşılanamayacak düzeyde. YGZ'nın belirli mesleklerde insanların yerini alarak neden olabileceği iş kaybı olasılığı, büyük miktarda veri toplama ve analiz kapasitesi nedeniyle kişisel mahremiyetle ilgili yol açabileceği sorunlar, süreç ilerledikçe ele alınması gereken endişelerden sadece birkaçı. YGZ'nin bireylere veya toplumlara zarar verebilecek şekillerde kullanılma riski de bulunuyor.

Her ne kadar gidişat, yapay genel zekânın geliştirilmesine yönelik olsa da bahsi geçen anlamda kabiliyetlere sahip, her açıdan insan gibi düşünen bir yapay zekâ hiçbir zaman geliştirilemeyebilir. Zaten şu an için böyle bir yapay zekânın tam olarak nasıl geliştirileceğini bilmiyoruz. Mevcut yapay zekâ modelleri daha yetenekli hâle gelse de ve aynı anda birden fazla işi başarıyla yapabiliyor olsa da insan

kadar kabiliyetli ve çok yönlü olamayabilir. Sınırlı derecede de olsa yapay genel zekâyâ yakın bir sistem, YGZ'nin yapacağı birçok faydayı sağlayacaktır.

Yapay genel zekâ dönemine hazırlanmak, birden fazla alanda koordineli hareket etmeyi gerektiriyor. Toplamların YGZ'nin getireceği değişikliklere hazırlıklı olmaları için mevcut ve gelecek nesillerin YGZ odaklı bir dünyaya adapte olmalarını sağlayacak

özgün, yenilikçi ve eleştirel düşünme gibi becerilere sahip olması gerekiyor.

Bütün bu nedenlerle YGZ sistemlerinin sorumlu bir şekilde geliştirilmesini ve kullanılmasını sağlayan yasalar ve düzenlemeler netleştirilmeli. Bugünden bakıldığında yapay genel zekâ, hem heyecan verici hem de göz korkutucu bir gelecek sunuyor. Seçim bizim elimizde. ■



Kaynaklar

- Geoffrey Hinton: 'The Godfather of AI' on risks to humanity. <https://bbc.com/news/world-us-canada-65452940>, BBC, 2023.
- Bostrom, N., *Superintelligence: Paths, Dangers, Strategies*. Oxford University Press, 2014.
- Brynjolfsson, E., & McAfee, A., *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies*. W. W. Norton & Company, 2014.
- Buttazzo G., Rise of artificial general intelligence: risks and opportunities. *Front. Artif. Intell.* 6:1226990. doi: 10.3389/frai.2023.1226990, 2023.
- Diamandis, P. H., & Kotler, S., *The Future Is Faster Than You Think: How Converging Technologies Are Transforming Business, Industries, and Our Lives*. Simon & Schuster, 2020.
- Frey, C. B., & Osborne, M. A., The future of employment: How susceptible are jobs to computerisation? *Technological Forecasting and Social Change*, 114, 254-280, 2017.
- Kurzweil, R., *The Singularity Is Near: When Humans Transcend Biology*. Viking, 2005.
- Lee, E. A. *The Coevolution: The Entwined Futures of Humans and Machines*. Cambridge, MA: MIT Press, 2020.
- O'Neil, C., *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. Crown, 2016.
- Topol, E. J., *Deep Medicine: How Artificial Intelligence Can Make Healthcare Human Again*. Basic Books, 2019.
- Turkle, S., *Alone Together: Why We Expect More from Technology and Less from Each Other*. Basic Books, 2011.
- Zuboff, S., *The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power*. PublicAffairs, 2019.
- Chan, C.Y.T. & Petrikat, D., Impact of Artificial Intelligence on Business and Society, *Journal of Business and Management Studies*, 2022.