

T. C.
İstanbul Üniversitesi
Sosyal Bilimler Enstitüsü
Felsefe Anabilim Dalı

Yüksek Lisans Tezi

Yapay Zekâ ve
Monoton-olmayan Mantık

Vedat Kamer
2501060036

Tez Danışmanı
Prof. Dr. Şafak Ural

Düzeltilmiş Tez

İstanbul 2009

T.C.
İSTANBUL ÜNİVERSİTESİ
SOSYAL BİLİMLER ENSTİTÜSÜ

TEZ ONAYI

Enstitümüz **FELSEFE ANABİLİM** Dalında **2501060036** numaralı **VEDAT KAMER'İN** hazırladığı “**YAPAY ZEKA VE MONOTON-OLMAYAN MANTIK**” konulu **YÜKSEK LİSANS/ DOKTORA-TEZİ** ile ilgili **TEZ SAVUNMA SINAVI**, Lisansüstü Öğretim Yönetmeliği'nin 15.Maddesi uyarınca **16.12.2009 ÇARŞAMBA** günü Saat:16.30'de yapılmış, sorulan sorulara alınan cevaplar sonunda adayın tezinin*Kabulü*.....'ne* **OYBİRLİĞİ /OYÇOKLUĞUYLA** karar verilmiştir.

JÜRİ ÜYESİ	KANAATI(*)	İMZA
PROF. DR. ŞAFAK URAL	<i>Kabul</i>	<i>Ş. Ural</i>
PROF. DR. AHMET FAHRİ ÖZOK	<i>Kabul</i>	<i>A. Özok</i>
DOÇ. DR. SEVİNÇ GÜLSEÇEN	<i>Kabul</i>	<i>S. Gülseçen</i>
YRD. DOÇ. DR. YÜCEL YÜKSEL	<i>Kabul</i>	<i>Y. Yüksel</i>
YRD. DOÇ. DR. NAZLI İNÖNÜ	<i>Kabul</i>	<i>N. İnönü</i>

DÜZELTMELER

1. Tüm metin yazım hataları için gözden geçirildi, mevcut hatalar düzeltildi.
2. İçindekiler bölümünde 1.5.1, 1.6, 2.3.6, 2.5.3, 4.2.2, 4.2.4 numaralı başlıklar sırasıyla “Bedene Sahip Olmanın Önemi: Nouvella Yapay Zekâ”, “1993'ten Günümüze Yapay Zekâ”, “Örüntü Tanıma”, “Olasılığa Bağlı Yöntemler”, “Eklentiler”, “Problemlere Çözümler ve Yarı-normal Öndeğerler” olarak düzeltildi.
3. 1.1, 1.2.4, 1.3.2, 1.3.4, 1.3.5, 1.4.1, 1.4.2, 1.6.5, 2.2.1, 2.5.1, 2.5.2, 2.5.6, 2.6, 3, 4, 4.1.1, 4.1.3, 4.1.4 numaralı bölümlerdeki anlatım bozuklukları düzeltildi.
4. 1.1.5, 1.1.6, 1.2.2, 1.3.1, 1.3.3, 1.4.2, 1.5.1, 1.6.1, 1.6.5, 2.3.1, 2.3.2, 2.3.6, 2.4.1, 2.4.3, 2.5.2, 2.5.6, 2.6, 4, 4.1.3, 4.2.4 numaralı bölümlerin dipnotları düzeltildi.
5. 2.2 numaralı bölümde El Cezerî ve Leonardo da Vinci ile ilgili bilgi eklendi.
6. 4.2.2 numaralı bölümün başlığının değiştirilmesiyle birlikte 4.2.1, 4.2.2, 4.2.3, 4.2.4 numaralı bölümlerde geçen “uzantı” ifadesi “eklenti” olarak değiştirildi.
7. Sonuç bölümündeki anlatım bozuklukları düzeltildi.
8. Kaynakça'daki çevrimiçi kaynaklar yeniden düzenlenip, değişiklikler yukarıda listelenen dipnotlara yansıtıldı.

ÖZ

Yapay zekâ 1956 yılında kurulmuş bir disiplindir. Bu tez çalışmasının amacı, bu yeni disiplinindeki mantık kullanımını ve mantığın işlevini açıklayarak, mantığın bir uygulama alanı olarak yapay zekânın olanaklarını serimlemektir. Bu çalışmada öncelikle yapay zekânın tarihçesi verilmiş, sonra da yapay zekânın temel problemleri ve yaklaşımları ele alınmıştır. Bu bilgiler ışığında yapay zekâ ile mantık ilişkisi değerlendirilerek, yapay zekânın sağduyu akılyürütmesini formelleştirmek için geliştirdiği monoton-olmayan mantık ele alınmıştır.

ABSTRACT

Artificial intelligence is a discipline established in 1956. The aim of this thesis is to expose the possibilities of artificial intelligence as an application of logic by explaining the use and the function of logic in this new discipline. First, the history of artificial intelligence is provided and then the main problems and approaches of artificial intelligence are addressed. In the light of this information, the relation between artificial intelligence and logic is evaluated, and non-monotonic logic, which artificial intelligence developed to formalize the common sense reasoning, is discussed.

ÖNSÖZ

Yapay zekâ bir disiplin olarak 1956 yılında kurulmuş olmasına rağmen disiplinlerarası özelliğinden ötürü büyük bir ivmeyle ilerlemiştir. Bu çalışmanın hazırlanmasında temel amaç yapay zekâdaki mantık kullanımını inceleyebilmektir. Yapay zekânın mantığa yeni uygulama alanları kazandırmak bakımından önemli bir birikime sahip olduğu düşünülerek, mantığın bu alandaki uygulama olanakları belirlenmeye çalışılmıştır.

Yapay zekâ ve monoton-olmayan mantığı ele aldığımız bu çalışmamızda 1. Bölüm'de yapay zekânın tarihçesini ele alarak yapay zekâyâ diğer disiplinlerin katkılarını ve bu katkılar sonucunda yapay zekânın gelişimini incelemeye çalıştık. 2. Bölüm'de yapay zekânın felsefi ve mantıksal arkaplanı ele alındıktan sonra yapay zekânın temel problemleri ele alınmış, bu problemlerin çözümlerine dair yaklaşımlar ve araçlar belirtilmiştir. 3. Bölüm'de yapay zekâ ve mantık arasındaki ilişki üzerine durulmuş ve yapay zekâda mantığın işlevi ele alınmıştır. 4. Bölümde sağduyu akılyürütmesine bir çözüm getirme çabası olarak monoton-olmayan mantık temel problemleriyle birlikte ele alınmıştır.

Çalışmamın her aşamasında yol göstericiliği, uyarı ve eleştirileri ile tezimi yöneten hocam Prof. Dr. Şafak Ural'a şükranlarımı sunar, çalışmalarım süresince beni destekleyen ve teşvik eden eşim Hilal Kamer'e teşekkür ederim.

İÇİNDEKİLER

Düzeltilmeler	iii
Öz	iv
Önsöz	v
Şekiller	ix
Giriş	1
1 Yapay Zekâ Tarihçesi	3
1.1 Yapay Zekânın Doğuşu (1943-1956)	3
1.1.1 Sibernetik ve İlk Yapay Sinir Ağları	3
1.1.2 Temel Yapay Zekâ Yaklaşımları	4
1.1.3 İlk Yapay Sinir Ağı Bilgisayarı	5
1.1.4 Turing Testi	6
1.1.5 Dartmouth Konferansı	7
1.1.6 Sembolik Akılyürütme ve Logic Theorist	10
1.2 Yapay Zekânın Altın Yılları (1956-1974)	11
1.2.1 Logic Theorist'in Takipçileri	12
1.2.2 LISP ve Advice Taker	13
1.2.3 Mikro-dünyalar ve Doğal Dil İşleme	15
1.2.4 Fiziksel Sembol Sistemleri Hipotezi	16
1.3 İlk Yapay Zekâ Kışı (1974-1980)	17
1.3.1 Yapay Zekânın Karşılaştığı Zorluklar	18
1.3.2 Filozofların Eleştirileri	21
1.3.3 Perseptronlar ve Bağlantıcılık	22
1.3.4 Tertipliler: PROLOG ve Uzman Sistemler	23
1.3.5 Pasaklılar: Çerçeveler ve Senaryolar	24
1.4 Yapay Zekâ Patlaması (1980-1987)	25
1.4.1 Uzman Sistemlerin Yükselişi	25
1.4.2 Fonların Geri Dönüşü	27

1.4.3	Bağlantıcılığın Geri Dönüşü	27
1.5	İkinci Yapay Zekâ Kışı (1987-1993)	28
1.5.1	Bedene Sahip Olmanın Önemi: Nouvella Yapay Zekâ	29
1.6	1993'ten Günümüze Yapay Zekâ	30
1.6.1	Dönüm Noktaları ve Moore Kanunu	31
1.6.2	Zeki Ajanlar	32
1.6.3	Tertiplilerin Zaferi	32
1.6.4	Sahne Arkasındaki Yapay Zekâ	33
1.6.5	HAL 9000 nerede?	34
2	Yapay Zekânın Özellikleri	36
2.1	Yapay Zekâ Nedir?	36
2.1.1	İnsanlar Gibi Davranan Sistemler	37
2.1.2	İnsanlar Gibi Düşünen Sistemler	37
2.1.3	Rasyonel Düşünen Sistemler	38
2.1.4	Rasyonel Davranan Sistemler	39
2.2	Yapay Zekânın Düşünsel Tarihi Geçmişi	39
2.2.1	Yapay Zekânın Felsefi ve Mantıksal Arkapları	41
2.3	Yapay Zekânın Problemleri	44
2.3.1	Arama Yaparak Problem Çözme	45
2.3.2	Bilgi Gösterimi	46
2.3.3	Planlama	48
2.3.4	Makine Öğrenmesi	49
2.3.5	Doğal Dil İşleme	51
2.3.6	Örüntü Tanıma	52
2.3.7	Hareket ve Manipülasyon	52
2.4	Yaklaşımlar	53
2.4.1	Sibernetik ve Beyin Simülasyonu	53
2.4.2	Sembolik Yapay Zekâ Yaklaşımı	54
2.4.3	Alt-sembolik Yapay Zekâ Yaklaşımı	56
2.4.4	İstatistiksel Yapay Zekâ Yaklaşımı	57

2.4.5	Yaklaşımların Entegrasyonu: Zeki Ajan Paradigması	57
2.5	Araçlar	58
2.5.1	Arama ve Optimizasyon	58
2.5.2	Mantık	59
2.5.3	Olasılığa Bağlı Yöntemler	60
2.5.4	Sınıfladıcılar ve İstatistiksel Öğrenme Yöntemleri	61
2.5.5	Yapay Sinir Ağları	62
2.5.6	Programlama Dilleri	62
2.6	Uygulamalar	64
3	Mantık ve Yapay Zekâ	68
4	Monoton-olmayan Mantık	73
4.1	Monoton-olmayan Mantığın Temel Problemleri	75
4.1.1	Çerçeve Problemi	75
4.1.2	Nitelik Problemi	79
4.1.3	Kapalı-dünya Varsayımı Akılyürütmesi	81
4.1.4	Kaldırılabilir Kalıtım Akılyürütmesi	82
4.2	Öndeğer Mantığı	84
4.2.1	Sembolleştirme	84
4.2.2	Eklentiler	86
4.2.3	Öndeğer Sonucu	91
4.2.4	Problemlere Çözümler ve Yarı-normal Öndeğerler	95
	Sonuç	100
	Kaynakça	102

ŞEKİLLER

Şekil 1: Kaldırılabilir Kalıtım Akılyürütmesi	83
Şekil 2: Nixon Elması	92

GİRİŞ

Bu çalışmada kuruluş tarihi bakımından oldukça yeni olan yapay zekâ disiplini ele alınmakta ve bu disiplindeki mantığın uygulama olanakları belirlenmeye çalışılmaktadır. Yapay zekâ disiplinlerarası bir çalışma alanıdır. Bu alana da en büyük katkı bilgisayar bilimlerinden gelmektedir. Bilgisayar bilimlerinin teorisinin temelini ise 20. yüzyıldan itibaren hızla gelişmekte olan mantık ve matematik çalışmaları oluşturmaktadır. Bu yüzden de mantık yapay zekânın hem teorisinde hem de uygulamalarında önemli araçtır. Bununla birlikte yapay zekânın en büyük hayali olan yapay insan fikrinin toplumsal ve kültürel pek çok yansıması vardır. Yapay zekânın sosyal ve etik kısmı çalışmamızın konusu dışında bırakılmıştır.

Yapay zekâ üzerine Türkçe yapılan çalışmalar maalesef oldukça azdır. Son dönemde konuyla ilgili araştırmalarda önemli bir yer tutan monoton-olmayan mantık üzerine herhangi bir Türkçe çalışmaya rastlanmamıştır. Bu yüzden çalışmamızda ağırlıklı olarak İngilizce kaynaklardan yararlanılmıştır.

Tezimiz dört bölümden oluşmuştur. Birinci bölümde yapay zekânın tarihçesi verilmektedir. Yapay zekânın disiplinlerarası bir çalışma alanı olması sebebiyle birbirinden çok farklı disiplinlerin katkısı mevcuttur. Olabildiğince kapsamlı tutmaya çalışacağımız tarihçe mantık temelli bir bakış açısıyla ele alınmıştır.

İkinci bölümde yapay zekânın özellikleri ele alınmıştır. Yapay zekânın disiplinlerarası özelliğinden ötürü tarihçesi içerisinde yapay zekânın özelliklerini sistematik bir şekilde ele alabilmek oldukça zordur. Bu bölümde yapay zekânın özellikleriyle birlikte problemleri ele alındıktan sonra, bu problemlerin çözümlerine dair yaklaşımlar ve araçlar ele alınmıştır.

Üçüncü bölümde mantık ve yapay zekâ arasındaki ilişki ve bu ilişkide yapay zekânın işlevi üzerine durulmuştur. Mantık temelli yapay zekâ yaklaşımının felsefe kökenli mantık çalışmalarıyla olan ilgisi incelenmeye çalışılmıştır.

İnsan seviyesinde bir akılyürütmenin yapay olarak gerçekleştirilebilmesi için sağduyu akılyürütmesinin formelleştirilmesi gerekir. Yapay zekâ çalışmalarının çözüm getirmek için en çok uğraştığı konulardan biri de budur. Düşüncenin

formelleřtirilmesi mantıđın yüzyıllardır uğrařtıđı bir alan olduđu için, dođal olarak sađduyu akılyürütmesinin formelleřtirilmesinde en önemli araç mantıktır. Dördüncü bölümde önermeler mantıđının monotonluk özelliđi ele alınarak, monoton-olmayan bir tarzda çalışan sađduyu akılyürütmesini formelleřtirmek için monoton-olmayan mantıđın bir türü olan öndeđer mantıđı gösterilmiřtir. Yapay zekânın çerçeve ve nitelik problemi gibi temel problemleri incelenerek, öndeđer mantıđında çözüm önerileri verilmeye çalışılmıřtır.

1. Yapay Zekânın Tarihçesi

1.1 Yapay Zekânın Doğuşu (1943-1956)

1940'li ve 1950'li yıllar boyunca değişik alanlardan (matematik, psikoloji, mühendislik, ekonomi ve siyasi bilimler) bir avuç bilim adamı yapay beyin üretilmesi olasılığını araştırmaya başladılar. Yapay zekâ disiplinin başlangıcı olarak kabul edilen bu çalışmalar sibernetik bünyesinde toplanmıştır.

1.1.1 Sibernetik ve İlk Yapay Sinir Ağları

Yapay zekâ olarak nitelendirilebilecek ilk çalışmalar, 1943 yılında **Warren Sturgis McCulloch** (1898-1969) ve **Walter Pitts** (1923-1969) tarafından yapılmıştır. “**A Logical Calculus of the Ideas Immanent in Nervous Activity**”¹ isimli ortak çalışmalarında **McCulloch** ve **Pitts**, yapay sinir ağlarını örnek olarak kullanan bir hesaplama modeli önerdiler. Modelin temelini 20. yüzyılın üç güçlü kuramının bir araya getirilmesi oluştuyordu: 1) **Bertrand Russell** (1872-1970) ve **Alfred North Whitehead**'in (1861-1947) önermeler mantığı, 2) **Charles Sherrington**'un (1857-1952) nöron teorisi, 3) **Alan Turing**'in (1912-1954) hesaplanabilirlik kuramı.²

McCulloch ve **Pitts** idealleştirilmiş basit sinir hücreleri kombinasyonlarının

1 Warren Sturgis McCulloch, Walter Pitts: **A Hierarchy of Values Determined by the Topology of Nervous Nets**, Bulletin of Mathematical Biophysics, Vol. 5, 1943, s. 115-133.

2 Margaret A. Boden: **Mind As Machine**, New York, Oxford University Press, 2006, s. 190; Stuart J. Russell, Peter Norvig: **Artificial Intelligence: A Modern Approach** (2nd ed.), New Jersey, Prentice Hall, 2003, s. 16; Margaret A. Boden: “Artificial Intelligence”, **Routledge Encyclopedia of Philosophy**, Ed: Edward Craig, London, Routledge, 1998.

'mantık geçitleri' olarak nitelendirilebileceğini göstermişlerdir. Örneğin, iki girdili bir **McCulloch-Pitts sinir hücresi** üç çeşit geçit olarak çalışabilmektedir: 1) **ve-geçidi**: iki girdi birlikte ateşleme yaparsa, 2) **veya-geçidi**: sadece bir girdi ateşleme yaparsa, 3) **değil-geçidi**: belirli bir girdi ateşleme yapmıyorsa. Böylece sinir ağları önermeler mantığının fonksiyonları ile ifade edilebilir hale gelmiştir. **Turing makinesi**³ tarafından hesaplanabilinen herhangi bir fonksiyon yapay sinir ağlarıyla da hesaplanabilmektedir.⁴

McCulloch ve **Pitts**, uygun şekilde tanımlanmış yapay sinir ağlarının öğrenme becerisi kazanabileceğini de ileri sürmüştür. **Donald Olding Hebb**'in (1904-1985) önerdiği kural (Hebbian learning) ile sinir hücreleri arasındaki bağlantıların şiddetini değiştirerek öğrenebilen yapay sinir ağları üretebilmek mümkün hale gelmiştir.⁵

1.1.2 Temel Yapay Zekâ Yaklaşımları

McCulloch ve **Pitts**'in bu çalışmasından yapay zekânın iki temel yaklaşımı filizlenmiştir: **sembolik yapay zekâ**⁶ ve **sibernetik yapay zekâ**⁷. Bu iki yaklaşım günümüzde sırasıyla **sayısal zihin teorisi** (computational theory of mind) ve **bağlantıcılık** (connectionism) temelinde ele alınmaktadır.⁸

Sayısal zihin teorisinde zihin bir bilgisayar olarak anlaşılmaktadır veya kabaca söylersek beynin yazılımı olarak tarif edilmektedir. Bunun temelinde insan

3 Karmaşık matematiksel hesapların belirli bir düzenek tarafından yapılmasını sağlayan hesap makinesi. Karmaşık hesapların belirli bir düzenek tarafından yapılıp yapılamayacağı 20. yy'ın başlarında büyük bir tartışma konusu olmuştur. **Bkz:** Alan M. Turing, "On computable numbers, with an application to the Entscheidungsproblem", **Proceedings of the London Mathematical Society**, Vol. 42, 1937, s. 230-265.

4 Margaret A. Boden: "Artificial Intelligence", **Routledge Encyclopedia of Philosophy**, Ed: Edward Craig, London, Routledge, 1998.

5 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 16.

6 Klasik yapay zekâ, geleneksel yapay zekâ ve GOFAI (Good Old-Fashioned Artificial Intelligence) olarak da bilinir.

7 Bağlantıcı yapay zekâ (connectionist AI) de denir.

8 Margaret A. Boden: "Artificial Intelligence", **Routledge Encyclopedia of Philosophy**, Ed: Edward Craig, London, Routledge, 1998.

zekâsının sembol manipölasyonuna indirgenebileceği fikri vardır. Bu yaklaşıma göre sembol manipölasyonunu gerçekleştirebilmek için beynin biyolojik yapısını taklit etmeye gerek yoktur. **Sembolik yapay zekâ**, açık sembol yapıları üretmek, karşılaştırmak ve değiştirmek için formel kurallardan oluşmuş programları kullanır. Programdaki formel yönergeler seri (tek tek sırayla) işlenir.⁹ Seri işlemede problemi çözecek algoritmadaki talimatlar sırasıyla ve teker teker çalıştırılır.

Bağlantıcılıkta ise zihnin işleyişinin beynin yapısı taklit edilerek ortaya konulabileceği savunulmuştur. Bunun için de yapay sinir ağları kullanılmaktadır. Bağlantıcılıkta hesaplamalar yapay sinir hücrelerini birbirine bağlayan ve yapay sinir ağını teşkil eden **düğüm**ler aracılığıyla yapılmaktadır. Bu hesaplamalar yapay sinir hücrelerinin bağlantı düğümlerinin ateşleme eşikleri, uyarıcı ve engelleyici görevdeki bağlantı yükleri değiştirilerek yapılır. Bağlantıcılıkta yapay sinir hücreleri arasında paralel (eş zamanlı) işleme yapılır. Paralel işlemede hesaplamadaki talimatlar (algoritma) eş zamanlı çalıştırılır.¹⁰

McCulloch ve **Pitts**'in çalışmasındaki mantık ve Turing hesaplanabilirliği kullanımı sembolik yapay zekâ yaklaşımını etkilemiştir. Bunun yanında **McCulloch** ve **Pitts**'in düğümler aracılığıyla komşu sinir hücresine mesaj iletebilen yapay sinir ağ kavramı **sibernetik yapay zekâ yaklaşımının** ilk örneğini oluşturur.

1.1.3 İlk Yapay Sinir Ağı Bilgisayarı

McCulloch ve **Pitts**'in öğrencisi olan **Marvin Minsky** (1927) ve **Dean Edmonds** ilk yapay sinir ağı bilgisayarını 1950'de geliştirmiştir. **SNARC** (Stochastic Neural Analog Reinforcement Calculator) olarak isimlendirilen bu bilgisayarda 40

⁹ Ned Block, Georges Rey: "Mind, computational theories of", **Routledge Encyclopedia of Philosophy**, Ed: Edward Craig, London, Routledge, 1998; Margaret A. Boden: "Artificial Intelligence", **Routledge Encyclopedia of Philosophy**, Ed: Edward Craig, London, Routledge, 1998.

¹⁰ Brain P. McLaughlin: "Connectionism", **Routledge Encyclopedia of Philosophy**, Ed: Edward Craig, London, Routledge, 1998; Margaret A. Boden: "Artificial Intelligence", **Routledge Encyclopedia of Philosophy**, Ed: Edward Craig, London, Routledge, 1998.

sinir hücresinden oluşan bir sinir ağını yönetmek için 3000 vakum tüpü ile B-24 bombardıman uçağının otomatik pilot mekanizması kullanılmıştır.

Bu mekanizma ile labirentte yolunu öğrenen bir farenin beyni taklit edilmiştir. Her bir sinir hücresi labirentteki bir konum ile haberleşiyordu ve ateşlendiği zaman “fare”nin kendini labirentte bu noktada olduğunu bildiğini gösteriyordu. Aktifleştirilmiş sinir hücresine bağlı diğer sinir hücreleri farenin hangi yolu seçebileceğini gösteriyordu. Bu sinir hücrelerinden hangisinin ateşleneceği aktifleştirilmiş sinir hücresine olan bağlantının kuvvetine dayanıyordu. “Fare” doğru hareketler dizisini gerçekleştirip labirentte yolunu bulduğunda bu hareketlerle ilgili bağlantılar kuvvetleniyordu. Böylece “fare” kademeli olarak labirentte yolunu öğreniyordu.¹¹ **Minsky** daha sonra sinir ağı araştırmalarının sınırlarını ortaya koyan önemli araştırmalar yapmıştır.¹²

1.1.4 Turing Testi

Yapay sinir ağlarındaki bu gelişmeleri **Alan Turing**'in 1950'de yayımladığı “**Computing Machinery and Intelligence**” isimli makalesi takip etmiştir.¹³ Yapay zekâ teorisinin temellerini oluşturan bu makalede “Makineler düşünebilir mi?” sorusu üzerinde durulmuştur. **Turing**'e göre “düşünmek” ve “makine” sözcükleri herkesi tatmin edecek açıklıkla tanımlanamaz. Bundan ötürü bu soruyu değiştirip, farklı şekilde sorulması gerektiğini savunur. Bir makinenin düşündüğünü tespit etmek yerine, **Turing** bir makinenin “Taklit Oyunu”nu kazanıp kazanamayacağını sormamızı önermektedir¹⁴:

Turing Testi olarak adlandırılan bu sınamaya göre bir bilgisayar gönüllü bir insanla birlikte sorgulayıcının görüş alanının dışında bir yere saklanır. Sorgulayıcı

11 Daniel Crevier: **AI: The Tumultuous Search for Artificial Intelligence**, New York, Basic Books, 1993, s. 35.

12 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 17.

13 Luciano Floridi: **Philosophy and computing: an introduction**, New York, Routledge, 1999, s. 26.

14 Alan Turing: “Computing Machinery and Intelligence”, *Mind*, Vol. LIX, Number 236, s. 433.

yalnız soru sormak suretiyle, hangisinin insan hangisinin bilgisayar olduğunu saptamaya çalışır. Sorgulayıcının soruları ve daha önemlisi aldığı yanıtlar, tamamen ses gizlenerek, yani ya bir klavye sisteminde yazılarak veya bir ekranda gösterilerek verilir. Sorgulayıcıya, bu soru-cevap oturumunda elde edilen bilgiler dışında, her iki taraf hakkında hiçbir bilgi verilmez. Dizi halinde tekrarlanan testler sonucunda sorgulayıcı, tutarlı bir şekilde, insanı saptayamadığı takdirde makine **Turing Testi**'ni geçmiş sayılmaktadır.¹⁵ Böylece “Makineler düşünebilir mi?” sorusu “Makineler insanın yaptıklarını yapabilir mi?” halini alır.¹⁶

Bu testi geçebilecek bir bilgisayarın aynı zamanda yapay zekânın alt alanları olan şu yeteneklere sahip olması gerekmektedir:¹⁷

- **Doğal dil işleme (natural language processing):** soruların sorulduğu dili anlayabilmek için.
- **Bilgi gösterimi (knowledge representation):** bildiklerini ve duyduklarını depolamak için.
- **Otomatikleşmiş akılyürütme (automated reasoning):** depolanan malumatları kullanarak sorulara cevap vermek ve yeni sonuçlar çıkarmak için.
- **Makine öğrenmesi (machine learning):** yeni sonuçlar uyarlamak ve modelleri tespit edip anlam çıkarmak için.

1.1.5 Dartmouth Konferansı

Mantık temelli yapay zekânın savunucu olan **John McCarthy** (1927), 1955 yazını IBM'de çalışarak geçirdikten sonra ilgisini dijital bilgisayarlar üzerine

15 Alan Turing: **a.y.**, s. 434; Nazlı İnönü: “Yapay Zekâ Felsefesi”, **Bilgisayar ve Beyin**, Ed. Ömer R. Akyüz, Mehmet Budak, Reşit Canbeyli, Ferruh Gençer, Işık Tabar Gençer, İstanbul, Nar Yayınları, 1997, s. 139.

16 Steven Harnard: “The Annotation Game: On Turing (1950) on Computing, Machinery, and Intelligence”, 2006, (Çevrimiçi), <http://cogprints.org/3322/1/turing.html>, 31 Temmuz 2009.

17 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 2-3.

kaydırmıştır. Turing makinesi ve otomat teorisi **McCarthy**'nin sadece kuramsal olarak zekâ üzerine çalışma imkânı sağlarken, bilgisayarlar yapay zekâlar üretebilmek için daha somut araçlar sağlamaktaydı.¹⁸ 1955-56 akademik yılı sırasında **McCarthy** ve Harvard Üniversitesi'nden arkadaşı **Mervin Minsky** zeki makineler fikri üzerine çalışan araştırmacıları bir araya getirmeye karar verdiler. Bu toplantıyı organize etmek üzere **Nathaniel Rochester** (1919-2001) ve **Claude Shannon** (1916-2001) da davet edildiler.¹⁹

IBM'de çalışan **Nathaniel Rochester**, genel amaca yönelik ve seri üretilen ilk bilgisayar olan IBM 701'in tasarımcısıdır. IBM'ın MIT'e hediye ettiği bilgisayar sayesinde **McCarthy** ile tanışmıştır. **Rochester** IBM'deki üç meslektaşı ile birlikte IBM 704 bilgisayarı üzerinde sinir ağlarını taklit etmeye çalışıyordu. **Rochester**'in ekipteki görevi büyük ölçekli bir sinir ağını tanımlayan sayısal eşitliği çözebilecek programı yazmaktır. Bilgisayarın bu amaca yönelik kullanımı **Minsky** ve **McCarthy**'nin ilgilendiği sembolik yapay zekâ yaklaşımından tamamen farklıdır. **Rochester**, **Minsky** ve **McCarthy**'nin mantık temelli yaklaşımının bilgisayarların problem çözme kabiliyetini arttıracaklarını ummaktadır.

1953 yazında **Minsky** ve **McCarthy** ile birlikte Bell Laboratuvarları'nda çalışmış olan **Claude Shannon** ise otomat teorisi üzerine çalışmalar yapmaktadır.²⁰

McCarthy 1955 tarihli konferans için davet yazısına aşağıdaki paragraf ile başlamıştır:

Bu çalışma insanların tekelinde olan bazı sorunları çözmek için dili kullanan, soyutlama yapabilen, kavramları işleyebilen ve kendini geliştirebilen makinelerin bunları nasıl yapacağını bulmaya yönelik olacaktır. Eğer dikkatlice seçilmiş bir grup bilimadamı, bir yaz boyunca bu uğurda çalışırsa, bu alanların bir veya birkaçında kayda değer bir gelişme olacağını düşünüyoruz. Çalışmada, öğrenmenin tüm yönlerini veya zekânın diğer özelliklerini ilkece açık, kesin ve net bir biçimde tanımlayabilmek suretiyle, benzerinin de makineler tarafından yapılabileceği varsayımıyla yapılmaktadır.²¹

Rochester ve **Shannon**'un yardımlarıyla, **McCarthy** ve **Minsky** düşünen

18 Daniel Crevier: **a.g.e.**, s. 38.

19 Daniel Crevier: **a.g.e.**, s. 39.

20 Daniel Crevier: **a.g.e.**, s. 40.

21 John McCarthy, Marvin Minsky, Nathan Rochester, Claude Shannon: "A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence", 1955, (Çevrimiçi), <http://www-formal.stanford.edu/jmc/history/dartmouth/dartmouth.html>, 31 Temmuz 2009.

makinelere üzerine bir yaz atölyesi düzenlemek için **Rockefeller Vakfı**'ndan masrafları karşılamak için \$7.500 bağış sağladılar.²² İki aylık konferans 1956 yılında “**The Dartmouth Summer Research Project on Artificial Intelligence**” isimiyle Dartmouth Koleji'nde gerçekleştirildi.

Konferansa, konferansı organize eden **McCarthy, Minsky, Rochester** ve **Shannon** ile birlikte toplam on kişi katılmıştır. Sonraki yirmi yılda konferansın bir araya getirdiği on katılımcı ve öğrencileri yapay zekânın önderliğini yapmıştır.²³

Beşinci katılımcı olan **Ray Solomonoff** (1926) konferansta “**An Inductive Inference Machine**”²⁴ isimli makalesini sunmuştur. Bu raporda makine öğrenmesi olasılıksal olarak incelenmektedir. Bu makale makine öğrenmesine ilişkin ilk yazı olma özelliğine sahiptir. **Ray Solomonoff** konferans süresince makinelere zekâlarını gösterebilmek için en karmaşık zihinsel problemlerin verilmesi eğilimine karşı çıkmıştır. **Solomonoff**'a göre makinelere daha basit problemler vermek, zihinsel süreçlerin analizini daha kolaylaştıracaktır.

Diğer katılımcı **Oliver Selfridge** (1926-2008) o dönemde MIT'teki Lincoln Laboratuvarları'nda **örüntü tanıma** (pattern recognition) üzerine araştırma yapıyordu. O sırada harfleri tanıyarak onları Mors koduna dönüştüren bilgisayarların programlaması üzerine çalışıyordu.²⁵

Katılımcılardan biri de **Doğal Tümdengelim** isimli teorem ispatlama tekniği üzerine çalışan **Trenchard More** idi.

Dördüncü katılımcı **Arthur Samuel** (1901-1990) ise bilgisayarlara dama oynamasını öğretmek suretiyle bir öğrenmenin nasıl sağlanacağı üzerine araştırmalar yapıyordu. **Samuel**'in çalışmaları daha sonra önemli bir etkiye sahip olacaktır.²⁶

Son iki katılımcı olan **Alan Newell** (1927-1992) ve **Herbert Simon**'ın da (1916-2001) yapay zekânın bu başlangıç döneminde doğrudan etkileri olmuştur. Konferansa tanıttıkları **Logic Theorist** isimli ilk yapay zekâ programını geliştirmişlerdir.

22 Daniel Crevier: **a.g.e.**, s. 40.

23 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 17.

24 Ray Solomonoff: “An Inductive Inference Machine”, 1956, (Çevrimiçi), <http://world.std.com/~rjs/indinf56.pdf>, 3 Ağustos 2009.

25 Daniel Crevier: **a.g.e.**, s. 40.

26 Daniel Crevier: **a.g.e.**, s. 41.

Otomat teorisi (automate theory), sinir ağırları ve zekâ çalışmaları konferans ile birlikte **McCarthy**'nin önerdiği **yapay zekâ** adı altında toplanmıştır.

1.1.6 Sembolik Akılyürütme ve Logic Theorist

Newell ve **Simon** 1955-56 yılında geliştirdikleri **Logic Theorist** isimli programı **Dartmouth Konferansı**'nda şu sözlerle tanıtmışlardır: “Sayısal-olmayan şekilde düşünebilen bir bilgisayar programı icat ettik ve dolayısıyla da maddeden oluşan bir sistemin zihnin özelliklerine nasıl sahip olacağını açıklayarak saygıdeğer zihin-beden problemini çözdük.”²⁷ İnsanın problem çözme kabiliyetini taklit eden **Logic Theorist**'in amacı mantık teoremlerini sembolik bir şekilde ispatlamaktır. Bunun için de insanın problem çözme kabiliyeti taklit edilmiştir. **Logic Theorist**, **Bertrand Russell** (1872-1970) ve **Alfred North Whitehead** (1861-1947) tarafından ortaklaşa kaleme alınan **Principia Mathematica** isimli eserin ikinci bölümündeki 52 teoremden 38'ini ispatlayabilmiştir.²⁸

Sembolik yapay zekâ yaklaşımı çerçevesinde geliştirilen ve ilk yapay zekâ programı olan **Logic Theorist**, bu yaklaşımın yirmi yıl kadar süren hükümdarlığının başlangıcını teşkil etmektedir.

Logic Theorist programında **Newell** ve **Simon** yapay zekâ araştırmalarının merkezine oturan üç yeni kavram sunmuştur:

- 1) **Arama olarak akılyürütme:** **Newell** ve **Simon** teorem ispatlamanın aramaya (search) indirgenebileceğini fark etmişlerdir. **Logic Theorist** bir arama ağacını kontrol etmektedir. Arama ağacı, düğüm noktalarında sıralı bir gruba ait veri değerlerinin depolanabildiği ağaç şeklinde bir veri yapısıdır. Arama ağacının kökünü başlangıç hipotezi oluşturur. Mantığın kuralları aracılığıyla hipotezin çok az değiştirilmiş

²⁷ Herbert A. Simon: **Models of My Life**, New York, Basic Books, 1991, s. 190. Filozof John Searle (1932) bilgisayarların insanlar gibi düşünebileceği yaklaşımını “Güçlü Yapay Zekâ” olarak isimlendirmiştir. **Bkz:** John Searle: Akıllar, Beyinler ve Bilim, İstanbul, Say Yayınları, 1996, s. 39.

²⁸ Daniel Crevier: **a.g.e.**, s. 46.

sürümleri açığa çıkmaktadır. Her manipölasyon kök düğümünü dallandırıp budaklandırmaktadır. Manipölasyon uygulamaları sonucunda ağacın en uç dalları programın ispatlamaya çalıştığı önermeleri ifade eder. Dallar boyunca giden yol da bu önermenin ispatını oluşturur. Dallarla temsil edilen ve her biri mantığın kurallarıyla türetilmiş olan bu ifadeler dizisi hipotezden önermeye kadar ispatı gösterir.

- 2) **Yol gösterici yöntemler**²⁹: **Newell** ve **Simon** arama ağacının üssel olarak büyüdüğünü ve gittikçe büyüyen bu ağacın bazı dallarının budanması gerektiğini fark ederler. Hangi dalın sonuca çıkacağını anlık olarak belirlemeye çalışarak sonuca varılmaya çalışır. Böylece sonuca gitmeyeceği düşünülen dallar budanmış olur. Aramaları hızlandırmak amacıyla yol gösterici yöntemler yapay zekânın sürekli gündemde tuttuğu bir alan haline gelir. Yol gösterici yöntemler özellikle seri bir şekilde çözülmek istenen problemler için en iyi cevabı vermeye çalışmaktadır.
- 3) **Liste işleme: Logic Theorist**'i bilgisayarda uygulayabilmek için **IPL** (Information Processing Language) isimli programlama dili geliştirilir. **IPL** sembolik listeleri işlemek için tasarlanmıştır ve yapay zekâ araştırmacıları tarafından hâlâ kullanılmaktadır.

1.2 Yapay Zekânın Altın Yılları (1956-1974)

Yapay zekânın ilk yılları basit ve ilkel bilgisayar ve programlama araçlarının örneklerinin verildiği yıllar olmuştur. Birkaç yıl öncesine kadar bilgisayarlar bir aritmetik makinesi olarak değerlendirilirken; **Turing**'in bilgisayarların yapamayacağı şeyler listesi, zamanla bilgisayarlar için imkânlı hale gelmeye başlamıştır. Yapay

²⁹ İng. Heuristics.

zekânın bu hızlı yükselişi yapay zekânın altın yılları³⁰ olarak adlandırılmıştır.³¹

1.2.1 Logic Theorist'in Takipçileri

Newell ve **Simon**'un başarılı **Logic Theorist** programını **General Problem Solver** (kısaca GPS) programı takip eder. Bu program doğrudan insanların akılyürütme tekniklerini taklit etmek için tasarlanmıştır. **GPS** ile sembolik problemler çözülebiliyordu: Örneğin teorem ispatlama, geometrik problemler ve satranç gibi. **Newell** ve **Simon**'un mantık makineleri üzerine teorik çalışmalarına dayanan **GPS**, problem bilgisini problemin nasıl çözülebileceği stratejisinden ayırabilen ilk bilgisayar programı olma özelliğini taşıyordu.

GPS, **Hanoi Kulesi**³² problemine benzer yeterince formalize edilmiş basit problemleri çözebiliyordu. Fakat gerçek dünyanın problemlerini çözemiyordu. **GPS**'in problem çözebilmesi için tüm nesnelerin ve bunların işleyişlerinin tanımlanması gerekmektedir. Kullanıcı nesneleri ve nesnelere uygulanabilecek işlemleri tanımlar ve **GPS** problemleri çözebilmek için neden-sonuç (means-end)

30 Yapay zekâ çalışmaları bu dönemde ilgili fonlar tarafından büyük meblağlarla desteklenmiştir. Haziran 1963'te ARPA (Advanced Research Projects Agency) tarafından, Minsky ve McCarthy önderliğindeki yapay zekâ araştırmalarını desteklemek için MIT'ye 2,2 milyon dolarlık fon sağlandı. ARPA 1970 yılının sonuna kadar MIT yapay zekâ projelerine ek olarak 3 milyon dolarlık fon daha ayırmıştır. ARPA'nın başkanı J. C. R. Licklider (1915-1990) 'projeleri değil kişileri desteklediklerini' vurgulayarak araştırmacıların ilgilendikleri alanda ilerlemelerine izin verdiklerini söylemiştir. ARPA benzer desteği CMU'da çalışmaları sürdüren Newell ve Simon'a ve Stanford Üniversitesi'nde McCarthy tarafından kurulan yapay zekâ laboratuvarına da yapmıştır. Edinburgh Üniversitesi'nde Donald Michie tarafından 1965 yılında kurulan yapay zekâ laboratuvarı da İngiliz Hükümeti'nce büyük destek görmüştür. Anılan dört üniversite yapay zekâ araştırmalarının önderliğini sürdürürken ilgili devlet kuruluşları tarafından sürekli desteklenmişlerdir. **Bkz:** Daniel Crevier: **a.g.e.**, s. 51, 64-65, 94.

31 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 18.

32 Hanoi Kulesi matematiksel bir oyun ya da bulmacadır. Üç çubuk ve bu çubuklara geçebilen farklı boyutlarda birkaç diskten oluşur. Bulmacaya diskler en küçüğü en üstte olacak şekilde boy sırasına göre, yani bir koni oluşturacak şekilde, bir çubukta dizilmiş olarak başlanır. Bulmacada amaç aşağıdaki kurallara uyarak tüm diskleri başka bir çubuğa aktarmaktır: 1) bir seferde sadece bir disk oynatılabilir; 2) her hareket çubukların birinin en üstündeki diski başka bir çubuktaki disklerin en üstüne koymaktan ibarettir; 3) hiçbir disk kendinden daha küçük bir diskin üzerine konamaz. **Bkz:** "Tower of Hanoi", (Çevrimiçi), <http://mazeworks.com/hanoi/index.htm>, 5 Ağustos 2009.

analizi³³ yaparak **yol gösterici yöntemler** oluşturur. Kullanabileceği işlemlere odaklanarak hangi girdilerin kabul edilebilir olduğunu ve bu girdilerin hangi çıktıları oluşturduğunu bulur. Sonra alt hedefler yaratarak asıl hedefe adım adım yaklaşmaya çalışır.³⁴

1952 yılında **Arthur Samuel** dama oynayan ve oynadıkça kendini geliştirebilen bir dizi program yazmıştır. Böylece bilgisayarların sadece kendilerine söyleneni yapabileceği fikri terk edilmiştir.

IBM'de **Nathaniel Rochester** tarafından 1958 yılında geliştirilen **Geometry Theorem Prover** matematik öğrencilerinin ispatlamakta zorlandığı geometri teoremlerini ispatlayabiliyordu.³⁵

Minsky'nin öğrencileri olan **James Slagle**'ın 1963 yılında geliştirdiği **SAINT** programı kapalı-form hesap entegrasyon problemlerini, **Tom Evans**'ın 1968 yılında geliştirdiği **ANALOGY** programı da IQ testlerinde kullanılan geometrik analogi problemlerini çözebiliyordu.³⁶

1.2.2 LISP ve Advice Taker

John McCarthy 1958'de MIT'e geçerek **AI Lab Memo No: 1** isimli yapay zekâ laboratuvarını kurmuştur. Aynı yıl yapay zekâ programlamasında en çok kullanılan **LISP** dilini yazmıştır. **LISP** hâlâ kullanılan en yaşlı ikinci³⁷ yüksek-seviyeli dildir. **McCarthy** yine 1958 yılında yayımladığı **“Programs with Common Sense”** isimli makalede ilk hipotetik bilgisayar programı olan **Advice Taker**'ı tanıtmıştır. **McCarthy**'nin makalesi, bilgi gösterimi için mantığı kullanmayı ve bir yapay zekâ yöntemi olarak sağduyu akılyürütmesini³⁸ öneren ilk çalışmadır.

33 Yapay zekâda problem çözme makinelerinin aramayı kontrol edebilmesini sağlayan teknik.

34 Daniel Crevier: **a.g.e.**, s. 52-53.

35 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 18.

36 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 19.

37 İlki de PASCAL isimli programlama dilidir.

38 **Bkz:** Yücel Yüksel: “Yapay Zekâ ve Puslu Mantık”, **Felsefe Arkivi**, İstanbul, İstanbul Üniversitesi Edebiyat Fakültesi Yayınları, 2008, Sayı: 32, s. 40

McCarthy'nin programı da **Logic Theorist** ve **Geometry Theorem Prover** programları gibi bilgiyi problem çözümlerinin araştırılması için kullanmak üzere tasarlanmıştır. Bu iki verilen bilgi ile ancak mikro-dünyalarda³⁹ çözüm getirilebilirken, **Advice Taker** yeni problemlerin yeni parametrelerine ayak uydurabilecek şekilde tanımlanmıştır. Örneğin, **McCarthy** basit aksiyomların nasıl havaalanına giden en kısa güzergâhı belirlemesini sağladığını göstermiştir. Program ayrıca operasyon gidişatında yeni aksiyomlar kabul edebilecek şekilde tasarlanmıştır; bu da onun yeniden programlanmaya gerek duymadan yeni alanlarda beceri kazanabilmesini sağlamaktadır. Dolayısıyla **Advice Taker** bilgi gösterimi ve akılyürütmenin temel ilkelerini içermektedir. Yani dünyanın ve bireyin davranışlarının dünyayı nasıl etkilediğinin formel, açık bir gösterimini elde etmek ve bu gösterimleri dedüktif işlemlerle manipüle edebilmek yararlıdır. 1958 tarihli bu makaledeki hususların günümüzde de hâlâ güncel olması dikkate değerdir.⁴⁰

1958 yılında **Marvin Minsky** de MIT'e geçmiştir. **McCarthy** ile **Minsky**'nin işbirliği sona ermemiş fakat **McCarthy**'nin bilgi gösterimi ve akılyürütmede formel mantığa olan vurgusu yerine **Minsky** mantık-karşıtı bir görüşe kaymıştır.⁴¹

1963 yılında **McCarthy** Standford Üniversitesi'nde bir yapay zekâ laboratuvarı kurmuştur. Standford Üniversitesi'ndeki çalışmalarda mantıksal akılyürütme için genel amaçlı metotlar üzerinde durulmuştur. Mantığın yapay zekâdaki uygulamalarına **Cordell Green**'in soru-cevaplama ile planlama sistemleri ve kendi hareketleri ile ilgili akılyürütebilen ilk robot olan **SHAKY** de dahildir. Sonraki robot projeleri mantıksal akılyürütme ile fiziksel aktivitenin tam entegrasyonunu sağlayan ilk örnekleri oluşturmaktadır.⁴²

39 Sınırlı alanlar

40 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 19.

41 Daniel Crevier: **a.g.e.**, s. 64.

42 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 19.

1.2.3 Mikro-dünyalar ve Doğal Dil İşleme

1960 yıllarının sonunda **Marvin Minsky** ve **Seymour Papert** (1928) MIT Yapay Zekâ Laboratuvarı'ndaki çalışmaların mikro-dünyalar olarak bilinen yapay, basit durumlara odaklanması gerektiğini düşünmüşlerdir. Bu model fizik gibi bilimlerde etkili olmuştur. Bu alandaki yapay zekâ çalışmalarının çoğu “bloklar dünyası”na odaklanmıştır. Bloklar dünyası bir yüzey üzerindeki değişik şekil ve boyuttaki renkli bloklardan oluşmaktadır.⁴³

Bloklar dünyası paradigması **makine görmesi** (machine vision) alanında yenilikçi çalışmaları beraberinde getirmiştir. **Gerald Sussman**'ın önderliğinde sürdürülen makine görmesi çalışmaları esnasında **Minsky** ve **Seymour Papert** “bloklar dünyası”ndaki blokları taşıyabilen bir robot kolu yapmışlardır.⁴⁴

Minsky'nin öğrencilerinden **Terry Winograd** (1946) **SHRDLU** isimli bir program geliştirmiştir. **SHRDLU** robot koluyla “bloklar dünyası”ndaki nesneleri algılayıp onları manipüle edebiliyordu. Sıradan İngilizce cümleler ile verilen talimat ve sorular doğrultusunda planladığı operasyonu uygulayabiliyordu⁴⁵:

İnsan: Kırmızı, büyük piramidi kaldır.
Bilgisayar: Tamam.
İnsan: Piramidi tut.
Bilgisayar: Hangi piramitten söz ettiğinizi anlayamadım.
İnsan: Şu an elinde tuttuğundan daha uzun bir blok bul ve onu kutuya koy.
Bilgisayar: Şu an elimde tuttuğumdan daha uzun bir bloktan bahsettiğinizi kabul ediyorum.
Bilgisayar: Tamam.
İnsan: Kutu ne ihtiva ediyor?
Bilgisayar: Mavi piramid ve mavi blok.
İnsan: Piramid ne tarafından taşıyor?
Bilgisayar: Kutu.
İnsan: Kaç tane blok kutunun içerisinde değil?
Bilgisayar: 4 tanesi..⁴⁶

Yapay zekâ araştırmalarının önemli bir amacı da doğal dil üzerinden iletişim kurabilmektir. **SHRDLU** her ne kadar mikro-dünyalar için bir örnek teşkil etse de,

43 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 19.

44 Daniel Crevier: **a.g.e.**, s. 87-88.

45 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 20.

46 SHRDLU, Terry Winograd'ın Kişisel Websitesi, (Çevrimiçi), <http://hci.stanford.edu/~winograd/shrdlu/>, (12 Ağustos 2009).

İngilizce iletişim kurabilme yeteneği bakımından **doğal dil işlemeye** de bir örnektir.

Doğal dil işlemede semantik ağlar yöntemi ilk defa **Ross Quillian** tarafından yazılan programda kullanılmıştır. Bu programın daha başarılı bir sürümü **Roger Schank** (1946) tarafından “**Conceptual Dependency**” ismiyle kodlanmıştır. Semantik ağlar yönteminde kavramlar birer düğüm olarak gösterilir ve aralarındaki ilgi oklar ile gösterilir.

Minsky'nin tez danışmanlığını yaptığı öğrencisi **Daniel Bobrow**'un 1964 yılında geliştirdiği **STUDENT** programı aşağıdakine benzer sözel cebir problemlerini çözebiliyordu:⁴⁷

Eğer Tom'un kazandığı müşterilerin sayısı, Tom'un yürüttüğü reklamların sayısının %20'sinin karesinin iki katı ise ve yürüttüğü reklamların sayısı 45 ise, Tom'un kazandığı müşterilerin sayısı kaçtır?

En çok ilgi görmüş İngilizce konuşan program olan **ELIZA**, **Joseph Weizenbaum** (1923-2008) tarafından 1964-1966 yıllarında geliştirilmiş bir sohbet robotu (chatbot) idi. **ELIZA** konuşanın insan olduğunu düşündürecek kadar konuşmaları gerçekçi bir şekilde sürdürebiliyordu. Fakat gerçekte **ELIZA**'nın konuştuğu konu hakkında bir “fikri” yoktu. **ELIZA** her soruyu ya genel cevaplar veriyor ya da ona söyleneni basit gramer kurallarıyla işleyip yeniden ifade ederek tekrarlıyordu.⁴⁸

1.2.4 Fiziksel Sembol Sistemleri Hipotezi

Logic Theorist ve **GPS** sonrasında 1976 yılında **Newell** ve **Simon** **fiziksel sembol sistemleri hipotezini** geliştirdiler. Bu hipoteze göre: “bir fiziksel sistem genel zekâ eylemi için gerekli ve yeterli araçlara sahiptir.”⁴⁹ Zihinlerimizin dünyaya

⁴⁷ Daniel Crevier: **a.g.e.**, s. 65, 75, 76.

⁴⁸ Daniel Crevier: **a.g.e.**, s. 133-134.

⁴⁹ Allen Newell, Herbert Simon, “Computer Science as Empirical Inquiry: Symbols and Search”, 1976, (Çevrimiçi), http://www.rci.rutgers.edu/~cfs/472_html/AI_SEARCH/PSS/PSSH4.html, 1

doğrudan erişimi yoktur. Biz sadece onun, bir grup sembol yapısından ibaret olan içsel bir temsili üzerinde işlem yapabiliriz. Bu yapılar herhangi bir fiziksel örüntünün şeklini alabilir. Dijital bir bilgisayarın içindeki elektronik anahtar dizilerinden ya da biyolojik bir beyindeki ateşlenen sinir hücresi ağlarından oluşabilir. Akıllı bir sistem (beyin ya da bilgisayar) bu yapılar üzerinde işlemler yaparak onları başka yapılara dönüştürebilir: Onları böler ve yeniden şekillendirir, bazılarını yok eder ve yenilerini oluşturur. Hipoteze göre, akıl sembolleri işleme becerisinden başka bir şey değildir. Akıl, onu destekleyen, onun işleminden geçen ve farklı fiziksel biçimlere girebilen donanımdan farklı bir alanda varlığını sürdürür.⁵⁰

1.3 İlk Yapay Zekâ Kışı (1974-1980)

1965 yılında **Herbert Simon** makinelerin yirmi yıl içinde bir insanın yapabileceği bütün işleri yapabileceğini söylemiştir.⁵¹ Yapay zekânın olanaklarının keşfedilmeye başlandığı bu ilk yılların rüzgârıyla pek çok kişi buna benzer kehanetlerde bulunmuştur. Bu kehanetlerin temelinde ilk yapay zekâ sistemlerinin basit örnekler üzerine gösterdiği başarılar vardı. Fakat iş karmaşık problemlere gelince yapay zekâ araştırmacıları tahmin etmedikleri zorluklarla karşılaşmıştır. Zorlukları aşamadıkları için de projeleri için parasal destek alamamışlardır. Bu dönemde bağlantıcılık (veya yapay sinir ağları) alanındaki çalışmalar **Minsky**'nin **perseptronlara** yönelik yıkıcı eleştirisi neticesinde neredeyse on yıl süreyle duraklamıştır. Bunlara rağmen **mantık programlama** (logic programming) ve **sağduyu akılyürütmesi** (commonsense reasoning) alanları başta olmak üzere pek çok alanda yeni fikirler ortaya atılmış ve yeni teoriler geliştirilmiştir. Yapay zekânın mikro-dünya problemlerini çözmede gösterdiği ilk yıldaki başarılarını, marko-dünya problemlerinde gösteremeyince başlayan duraklama dönemine yapay zekâ kışı⁵²

Ağustos 2009.

50 Daniel Crevier: **a.g.e.**, s. 48.

51 Daniel Crevier: **a.g.e.**, s. 109.

52 Bu dönemde İngiliz Hükümeti, DARPA (önceki ismiyle ARPA) ve NRC gibi yapay zekâyı

denilmiştir.

1.3.1 Yapay Zekânın Karşılaştığı Zorluklar

1970'li yılların başında yapay zekâ programlarının yetenekleri kısıtlıydı. En etkileyici programlar bile problemlerin ancak çok küçük bir bölümünü çözebiliyordu. Bu programlar bir tür “oyuncak” gibiydi. Yapay zekâ araştırmacıları 1970'li yıllar boyunca aşamadıkları pek çok problemle karşılaştılar. Kısıtlamaların bir kısmı sonraki yıllarda aşıldıysa bile, bir kısmı yapay zekâ alanının ilerlemesine hâlâ engel olmaktadır.⁵³ Bu kısıtlamaları kısaca belirtmeye çalışalım:

1. **Sınırlı bilgisayar gücü:** Herhangi bir şeyin üstesinden gelmek için yeterli hafıza kapasitesi ve işlem hızı mevcut değildi. Örneğin, **Ross Quillian**'ın doğal dil işlemesi üzerine başarılı çalışmaları sadece yirmi kelime ile çalışabiliyordu.⁵⁴ Çünkü hafıza kapasitesi daha fazlasını almaya müsait değildi. **Hans Moravec** (1948) 1976 yılında bilgisayarların insan zekâsından milyonlarca defa daha zayıf olduğunu tartışmış ve şu analogiyi yapmıştır: yapay zekâ, otomobillerin beygir gücüne ihtiyaç duyması gibi, bilgisayar gücü gerektirir.

İlk yapay zekâ programları çözüme gidecek adımları teker teker

destekleyen kuruluşlar alandaki başarısızlıklar yüzünden neredeyse bütün fonları durdurmuştur. Süreç 1966'da ALPAC'ın tercüme makineleri üzerine raporu ile başlar. Neredeyse 20 milyon dolar harcanan bu çalışmalardan NRC tüm desteğini çekmiştir. 1973 yılında Sir Jame Lighthill (1924-1998) tarafından hazırlanan “Yapay Zekâ: Genel Bir Araştırma” isimli rapor doğrultusunda İngiliz Hükümeti üç üniversite haricindeki bütün üniversitelerde yapay zekâ araştırmalarını durdurmuştur. Rapor özellikle kombinasyonel patlama problemi üzerine durmuştur. Aynı yıllarda DARPA, CMU'da sürdürülen konuşma algılama programı için ayırdığı fonu büyük hayal kırıklığı sonunda durdurmuştur. Hans Moravec krizi meslektaşlarının gerçekçi olmayan tahminlerine bağlamıştır: “Birçok araştırmacı gittikçe artan abartma ağına yakalanmış durumda.” DARPA 1969 yılından itibaren projeler yerine kişileri destekleyen politikasını sona erdirip, projeleri destekler bir politika benimsemiştir. 1960'lardaki yaratıcı ve serbest araştırmalar yerine DARPA desteğini otonom tanklar ve savaş yönetim sistemlerine kaydırmıştır. **Bkz:** Daniel Crevier: **a.g.e.**, s. 110, 115-117.

⁵³ Daniel Crevier: **a.g.e.**, s. 146.

⁵⁴ Daniel Crevier: **a.g.e.**, s. 146.

deneyerek ilerliyordu. Bu strateji çok az nesne ve eylem barındıran mikro-dünyalarda işe yaramıştır. Fakat karmaşıklık artıkça çözüm üretilemez olmuştur. Karmaşıklık teorimi geliştirilmeden önce, bunun bir ölçeklendirme (scaling) sorunu olduğu, daha hızlı donanım ve hafıza ile aşılabileceği düşünülmüştür. Teorem çözme araştırmalarında, birkaç olgu daha eklendiğinde araştırmacıların teoremleri ispatlayamaz hale geldikleri görülmüştür. Kuramsal olarak bir programın çözüm arayabilmesi, pratikte arayabilecek bir mekanizmaya sahip olduğu anlamına gelmemektedir.⁵⁵

2. **Tıkanma ve kombinasyonel patlama:** 1972 yılında **Richard Karp** (1935) pek çok problemin ancak üssel olarak ifade edilebilecek zamanlarda çözülebileceğini ortaya koymuştur. Yani olasılıklar arttıkça, denenmesi gereken olasılıklar üssel olarak artmaktadır. Dolayısıyla da karmaşık problemlerin çözümü için inanılmaz derecede fazla işlem yapılması gerekmektedir. Bu kadar çok sayıda işlemin tamamlanma süresi ise insan ömrünü geçebilecek kadar uzun olabilmektedir. Bu da mikro-dünyalarda başarılı olmuş yapay zekâ çözümlerinin muhtemelen hiçbir zaman kullanışlı sistemlere ölçeklendirilemeyeceği anlamına gelmektedir.⁵⁶
3. **Sağduyu bilgisi ve akılyürütmesi:** Doğal dil işleme ve görme gibi pek çok yapay zekâ alanı gerçek dünya hakkında çok fazla malumat gerektiriyordu. Programların neye baktıkları veya ne hakkında konuştukları hakkında fikir sahibi olmaları gerekiyordu. Bu da programın neredeyse bir çocuk kadar malumata sahip olması anlamına gelmektedir. 1970'li yıllarda hiç kimsenin bu büyüklükte bir veritabanı oluşturması mümkün değildi. Bundan öte, bir programın bu kadar malumatı nasıl öğrenebileceğini kimse bilmiyordu.⁵⁷

İlk yapay zekâ programlarının başarısı basit sentaktik işlemler vasıtasıyla olmuştur. **ABD Ulusal Araştırma Konseyi**'nin (NRC) önyak olduğu ilk makine tercüme çalışmaları **Sputnik**'in 1957'de uzaya fırlatılması öncesinde, ele geçirilen Rus bilimsel makalelerinin İngilizce'ye tercümesi

55 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 22.

56 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 9, 21-22.

57 Daniel Crevier: **a.g.e.**, s. 113-114.

amacıyla başlatılmıştır. Rusça ve İngilizce gramerlerinin basit sentaktik dönüştürülmesi ile elektronik sözlük kullanılarak sözcüklerin değiştirilmesinin cümlelerin tam anlamlarının çıkarılması için yeterli olacağı sanılmaktaydı. Sonuçlardan anlaşılmıştır ki konu hakkında genel bilgi sahibi olmadan cümlelerin muğlaklığını çözüp, içeriği cümle ile bağlamak mümkün değildi. 1966 yılında danışma komitesinin verdiği “genel bilimsel metinleri tercüme edebilecek bir makine henüz yoktur” raporu üzerine tüm akademik tercüme projeleri iptal edilmiştir.⁵⁸

4. **Moravec paradoksu:** **Hans Moravec**'a göre teoremleri ispatlamak ve geometri problemlerini çözmek bilgisayarlar için göreceli olarak kolaydı. Fakat yüz tanımak veya bir odayı hiçbir nesneye çarpmadan geçmek gayet güçtü. Bu 1970'li yıllarda görme ve robotbilimi alanında, neden bu kadar az ilerleme kaydedildiğini açıklamaktadır.⁵⁹
5. **Çerçeve⁶⁰ ve nitelik problemi⁶¹:** Yapay zekâ araştırmacıları, başta **John McCarthy** olmak üzere, planlama veya öndeğer akılyürütmesi (default

58 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 21.

59 Pamela McCorduck: **a.g.e.**, s. 456.

60 Yapay zekâda çerçeve problemi başlangıçta, hangi koşulların hareketten etkilenmediğini açıkça belirlemeden mantıkta dinamik bir etki alanı ifade etme problemi olarak formüle edilmiştir. John McCarthy ve Patrick J. Hayes bu problemi 1969 tarihli *Some Philosophical Problems from the Standpoint of Artificial Intelligence* başlıklı makalelerinde tanımlamıştır. Daha sonra, terim felsefe alanında daha geniş bir anlam kazanmış, hareketlere yanıt olarak güncellenmesi gereken inançları sınırlama problemi olarak formüle edilmiştir. “Çerçeve problemi” adı çizgi film yapımcılarının kullandığı çerçeveleme adı verilen bir teknikten türemiştir. Bu teknikte çizgi karakterin hareket eden kısımları sahnenin arkaplanının resmedildiği ve değişmeyen “çerçeve” üzerine üst üste konur. Mantık bağlamında hareketler tipik olarak neyi değiştirdikleriyle belirlenir, başka her şeyin (çerçeve) değişmeden kaldığını varsayılır. **Bkz:** John McCarthy: “Some Philosophical Problems from the Standpoint of Artificial Intelligence”, 1969, s. 30-31, (Çevrimiçi), <http://www-formal.stanford.edu/jmc/mcchay69.pdf>, 16 Ağustos 2009.

61 Felsefe ve yapay zekâda (özellikle bilgi temelli sistemlerde), nitelik problemi gerçek dünyadaki bir hareketin istenen etkiyi yaratabilmesi için gereken tüm önkoşulların listelenmesinin imkânsızlığıyla ilgilidir. Bu şekilde ortaya konabilir: İstediğim sonucu elde etmemi engelleyen şeylerle nasıl başa çıkacağım? Bu çerçeve probleminin dallanma kısmıyla yakından ilişkilidir ve buna zıttır. John McCarthy bir kayığın olağan işlevini yerine getirmesini engelleyebilecek bütün şartların tek tek sıralanmasının mümkün olmayışını şu örnekle veriyor: “Bir nehri geçmek için bir kayığın başarılı bir biçimde kullanılması, eğer kayık kürekle çalışıyorsa küreklerin ve kürek ıskarmozunun bulunması, sağlam olması ve birbirleriyle uyumlu olması gerekir. Birçok başka nitelik da eklenebilir, öyle ki bir kayığı kullanmak için gereken kuralların uygulanması neredeyse imkansız hale gelirken hala henüz dile getirilmemiş başka koşullar düşünülebilir.” **Bkz:** John McCarthy: “Circumscription – A Form of Nonmonotonic Reasoning”, 1986, s. 1-2, (Çevrimiçi), <http://www-formal.stanford.edu/jmc/circumscription.pdf>, 16 Ağustos 2009.

reasoning) içeren sıradan çıkarımları mantığın yapısında değişiklik yapmadan ifade edemediklerini keşfettiler. Problemleri çözebilmek için araştırmacılar monoton-olmayan mantığı geliştirmiştir.⁶²

1.3.2 Filozofların Eleştirileri

Birkaç filozof yapay zekâ araştırmacılarının iddialarına karşı güçlü eleştiriler getirmiştir. Bunlardan ilki **John Lucas** (1929) idi. **Lucas'a** göre **Kurt Gödel'in** (1906-1978) **eksiklik teoremi** bir formel sistemin (bilgisayar gibi) belirli durumların doğruluğunu asla göremeyeceğini ispatlamıştır. Bu yüzden eksiklik teorimi ile sınırlandırılmış formel sistemlerin mental olarak insanlara karşılaştırılması mümkün değildir.⁶³

Hubert Dreyfus (1929) 1960'lardan beridir yapay zekânın temel savlarından olan sembol işlemenin akılyürütmenin çok az bir kısmını teşkil ettiğini iddia etmiştir. Dreyfus'a göre insanın en önemli özelliği içerilmiş ve içgüdüsel bir şekilde neyi nasıl yapılacağını bilmesi, yani sağduyu akılyürütmesine sahip olmasıdır. Dreyfus bu özellik için “**know how**” kavramını kullanır.⁶⁴

John Searle (1932) 1980 yılında ortaya koyduğu **Çin Odası Argümanı**'nda⁶⁵, **Turing Testi**'ni eleştirilmiştir. Argümanda bir bilgisayar programının sembolleri anlamadığı gösterilmiştir. Eğer makine için sembollerin bir anlamı yoksa, **Searle'e** göre bir makine “düşünüyor” olarak tanımlanamaz.⁶⁶

62 Daniel Crevier: **a.g.e.**, s. 117-118.

63 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 949-950.

64 Daniel Crevier: **a.g.e.**, s. 128.

65 Bu argümanında Searle, biri Çince'yi bilen ve ismi C olan, diğeri de Çince'yi bilmeyen ve ismi D olan iki kişi varsayar. C bir oda içine kapatılarak, C'ye gelen ve C'den giden mesajlar kâğıtlar üzerine işaret dizileri ile yazılmıştır. Eğer odaya gönderilen işaretler, Çince yazılmış sorularsa ve gelen işaretler de Çince yazılmış cevaplarsa oda sakininin Çince bildiği manasına gelecektir. D'nin aynı odaya kapatılıp, eline de gönderilen işaret dizilerine karşılık gelecek, giden işaret dizilerini içeren bir tablo verildiğini farz edelim. Eğer D bu etkinlikte uzman olsaydı belki C'ye yakın bir performans sergileyecekti. C'nin ya da D'nin Çince bilip bilmediği anlaşılmayacaktır. Çin Odası Argümanında aynı “girişler” aynı “çıkışlar”ı vermiştir, fakat bunları oluşturan süreç ve işlemler aynı değildir.

66 Nazlı İnönü: **a.g.e.**, s. 144-145.

Bu eleştiriler yapay zekâ araştırmacılarınca ciddiye alınmamıştır. **Minsky**, **Dreyfus** ve **Searle** için “onlar yanlış anlamışlardır ve ihmal edilmelidirler”⁶⁷ demiştir. **ELIZA**'nın geliştiricisi **Joseph Weizenbaum**, **Dreyfus**'un eleştirilerini amatör ve çocukça olarak değerlendirmiştir.⁶⁸ Onlar için 'tıkanma' ve 'sağduyu bilgisi' problemi daha acil ve önemli görülmüştür.

Kenneth Colby (1920-2001) terapi yapan **DOCTOR** isimli sohbet robotunu geliştirmiştir. **ELIZA**'nın geliştiricisi **Weizenbaum**, **DOCTOR** programı için ciddi etik endişeler taşımıştır. **Weizenbaum**, **Colby**'nin alanı hakkında hiçbir bilgisi olmayan **DOCTOR** programını ciddi bir terapi aracı olarak görmesinden tedirgin olmuştur. 1976 yılında **Weizenbaum** “**Computer Power and Human Reason: From Judgment To Calculation**” isimli kitabında yapay zekânın yanlış kullanımının insan hayatını riske atabileceğini tartışmıştır.⁶⁹

1.3.3 Perseptronlar ve Bağlantıcılık

Donald O. Hebb'in öğrenme metotları **Bernie Widrow** (1929) tarafından geliştirilmiştir. **Widrow**'un geliştirdiği yapay sinir ağları tek katmanlı **ADALINE** (ADaptive LINear Element) hücrelerinden oluşuyordu. **Frank Rosenblatt** (1928-1971) **Widrow**'un çalışmalarından ilham alarak **perceptron**⁷⁰ yakınsama teoremini geliştirmiştir. Bu teoremin önerdiği öğrenme algoritması aracılığıyla perseptronların bağlantı kuvvetleri girdi verilerine ayarlanabiliyordu.⁷¹

Birçok yapay zekâ araştırmacısı perseptronların sonunda öğrenme, karar verme ve tercüme konusunda çözüm getirme gücüne sahip olduğunu düşünüyordu.

⁶⁷ Daniel Crevier: **a.g.e.**, s. 143.

⁶⁸ Daniel Crevier: **a.g.e.**, s. 122.

⁶⁹ Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 961.

⁷⁰ Bir tür yapay sinir hücresi ağıdır. En basit geriye yayılma algoritması olarak görülebilir. **Bkz:** Sevinç Gülseçen: “Yapay Sinir Ağları, İşletme Alanında Uygulanması ve Bir Örnek Çalışma”, İstanbul Üniversitesi Sosyal Bilimler Enstitüsü, Yayınlanmamış Doktora Tezi, İstanbul, 1993, s.

80

⁷¹ Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 21.

1960'larda paradigma üzerine sürdürülen çalışmalar **Minsky** ve **Papert**'in 1969 yılında yayımladıkları “**Perceptrons**” isimli kitapla sona ermiştir. Perseptronların temsil edebildiği herşeyi öğrenebildiği söylenmişse de perseptronlar çok az şeyi temsil edebiliyordu. Özellikle, iki girdili perseptronlar, iki girdinin farklı olduğunu tespit edebilmek için eğitilemiyordu.⁷² Bu eleştiri sonrasında yapay sinir ağlarındaki araştırmalar on yıl süreyle duraklamıştır.

1.3.4 Tertipliler⁷³: PROLOG ve Uzman Sistemler

Mantığın bir araç olarak kullanımı 1958 yılında **John McCarthy**'nin geliştirdiği **Advice Taker** isimli programla yapay zekânın gündemine girmiştir. 1963 yılında **John Alan Robinson** (1930) bilgisayarlarda tümdengelimini uygulamaya sokmak için basit bir yöntem olan çözülme ve birleşme algoritmasını keşfetmiştir. Ancak, 1960'ların sonlarında **McCarthy** ve öğrencileri tarafından denenen doğrudan uygulamalar özellikle karmaşıktır. Programlar basit teoremleri bile ispat etmek için astronomik sayıda basamağa ihtiyaç duymuşlardır.⁷⁴

Fransa'nın Marseille Aux Üniversitesi'nde **Alain Colmerauer**'in (1941) ekibi tarafından geliştirilen **PROLOG** (PROgramming LOGic), programlama dilinde izlemesi kolay bir hesaplama izin veren mantık altkümeleri kullanılır. Böylece belirli **yol gösterici yöntemler** ile basamakların sayısı azaltılarak tıkanmanın önüne

⁷² Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 22.

⁷³ Yapay zekâ jargonunda araştırmalarda kullanılan temel iki yaklaşımın taraflarını belirtmek için kullanılır. Bunun temelinde insan akılyürütmesi ve yapay zekâ arasındaki ilişki yer almaktadır. Sembolik yapay zekâ yaklaşımı çerçevesinde araştırmalarını sürdüren tertipliler (neats), formel metotlar kullanarak insan zihninin rasyonelliğini taklit etmeye çalışırlar. Pasaklılar (scruffies) ise empirik bilgiye götüren bütün yolları kullanmaya eğilimindedirler. Bir tertipli (neat) için pasaklının (scruffy) yöntemleri gelişigüze, pasaklı sadece rastlantı sonucu başarılı olabilir ve o, zekânın gerçekte nasıl çalıştığı ile ilgili uğraşmaz. Bir pasaklı (scruffy) için de, bir tertipli (neat) formelleştirme ve “sağ duyuya” takılıp kalmıştır. Ünlü tertipliler şunlardır: John McCarthy, Alan Newell, Herbert Simon, Edward Feigenbaum, Robert Kowalski, Judea Pearl, David McAllester, Daphne Koller. Ünlü pasaklılar ise şunlardır: Rodney Brooks, Marvin Minsky, Roger Schank, Doug Lenat, Steve Grand. **Bkz:** Pamela McCorduck: **a.g.e.**, s. 486-488, Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 25.

⁷⁴ Daniel Crevier: **a.g.e.**, s. 190.

geçilebilir.⁷⁵ **PROLOG**, **Edward Feigenbaum**'ın (1936) uzman sistemlerine ve **Alan Newell** ile **Herbert Simon**'un o sıralar devam eden **SOAR** isimli projesine dayanak sağlayacaktır.

Mantık yaklaşımını eleştirenler, **Dreyfus**'un da belirttiği gibi insanların problemleri çözerken nadiren mantığı kullandığını belirtmişlerdir. **Peter Wason**, **Eleanor Rosch**, **Amos Tversky**, **Daniel Kahneman** ve diğer psikologların yaptıkları deneyler bunun kanıtı olmuştur. **McCarthy** buna insanların ne yaptığının önemsiz olduğunu söyleyerek karşılık vermiştir. Ona göre insanlar gibi düşünen makineler değil problemleri çözebilen makinelere ihtiyaç vardır.⁷⁶

1.3.5 Pasaklılar: Çerçeveler ve Senaryolar

McCarthy'nin yaklaşımını eleştirenler arasında MIT'deki meslektaşları da vardı. **Marvin Minsky**, **Seymour Papert** ve **Roger Schank** insan gibi düşünen bir makine gerektiren “hikâye anlama” ve “nesne tanıma” gibi problemleri çözmeye çalışıyorlardı. Bu makine “Sandalye” ya da “lokanta” gibi sıradan kavramları kullanmak için insanların normalde yaptıkları tüm mantık dışı varsayımları yapmak durumundaydı. Ne yazık ki, bunun gibi kesin olmayan kavramların mantıkta ifade edilmesi zordur. **Gerald Sussman**'ın şöyle bir gözlemi vardır: “özünde kesin olmayan kavramları anlatmak için kesin bir dil kullanmak onları daha kesin yapmaz”.⁷⁷ **Schank** “mantık-karşıtı” yaklaşımlarını **McCarthy**, **Kowalski**, **Feigenbaum**, **Newell** ve **Simon**'un kullandığı “**tertipli**” (neat) paradigmalara karşı “**pasaklı**” (scruffy) olarak nitelendirmiştir.⁷⁸

1975 yılında taslak halindeki bir makalesinde **Minsky**, “**pasaklı**” araştırmacı arkadaşlarının birçoğunun aynı tür aleti kullandığını belirtmiştir: Birşey hakkındaki

⁷⁵ Daniel Crevier: **a.g.e.**, s. 193.

⁷⁶ John McCarthy: “What Was Expected, What We Did, and AI Today”, (Çevrimiçi), http://www.engagingexperience.com/2006/07/ai50_ai_past_pr.html, 13 Ağustos 2009.

⁷⁷ Daniel Crevier: **a.g.e.**, s. 175.

⁷⁸ Daniel Crevier: **a.g.e.**, s. 168.

tüm sađduyulu varsayımlarımızı içeren bir yapı. Örneğın, kuş kavramını kullanırsak hemen akla gelen birtakım olgular vardır: uçtuğunu, solucan yediğini vs. varsayabiliriz. Bu olguların bütün kuşlar için doğıru olmadığını biliriz, fakat bu varsayım kümeleri söylediğimiz ve düşündüğümüz herşeyin bağlamının bir parçasıdır. **Minsky** bu yapılara “**çerçeveseler**” demiştir. **Schank**, İngilizce kısa hikâyeler hakkındaki soruları yanıtlamak için çerçeveselerin “senaryolar” olarak adlandırdığı bir versiyonunu kullanmıştır.⁷⁹ Yıllar sonra nesne yönelimli programlama, çerçeveseler üzerine yapılan yapay zekâ çalışmalarındaki “kalıtım” fikrini benimseyecektir.

1.4 Yapay Zekâ Patlaması (1980-1987)

1980’li yıllarda yapay zekâ programlarının bir biçimi olan “uzman sistemler” şirketler tarafından benimsenmiş ve bilgi yapay zekâ araştırmalarının odağına yerleşmiştir. Aynı yıllarda, Japon hükümeti **The Fifth Generation** projesini başlatmıştır. Yine 1980’li yılların başında **John Hopfield** (1933) ve **David Rumelhart**’ın (1942) çalışmalarıyla bağlantıcılıkta başarılı çalışmalar yapılmıştır. Yapay zekâ ilk yıllarındaki temposunu tekrar yakalamıştır.

1.4.1 Uzman Sistemlerin Yükselişi

Uzman kişilerin bilgisinden türetilmiş mantık kuralları kullanılarak, belirli uzmanlık alanındaki soruları cevaplayan veya problemleri çözen programlara uzman sistemler denir. Her sorunu çözecek genel amaçlı bir program yerine, belirli bir uzmanlık alanındaki bilgiyle donatılmış programlar kullanma fikri yapay zekâ

⁷⁹ Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 24.

alanında yeniden bir canlanmaya yol açmıştır.

Uzman sistemler daraltılmış özel bir bilgi alanıyla sınırlandırılmıştır. Böylece sağduyu bilgisi problemi aşılmıştır. Uzman sistemlerin üretilmesi ve değiştirilmesi basit tasarımlarından dolayı oldukça kolaylaşmıştır. Sonuç itibariyle, kullanışlı ve pratik programlar ortaya çıkmıştır. Süregelen yapay zekâ çalışmalarında daha önce bu noktaya ulaşılamamıştır.⁸⁰

Edward Feigenbaum, Bruce Buchanan, Jashua Lederberg (1925-2008) ve **Carl Djerassi** (1923) tarafından 1969 yılında Stanford Üniversitesi'nde geliştirilen **DENDRAL** programı, kütle spekturumu ve kimya bilgisini kullanarak, bilinmeyen organik molekülleri bulup organik kimyacılar yardımcı olmak amacıyla tasarlanmıştır. **DENDRAL** organik kimyacıların karar verme süreçlerini ve problem-çözme davranışlarını otomatikleştirdiği için ilk uzman sistem olarak kabul edilir.⁸¹

Feigenbaum önderliğinde Stanford Üniversitesi'nde uzman sistemlerin hangi alanlara uygulanabileceğini araştırmak için **Heuristic Programing Project** isimli çalışma başlatılmıştır. Sonraki çalışmaların ana konusunu tıbbi teşhis oluşturmuştu. **Feigenbaum, Buchanan** ve **Edward Shortliffe** (1947) kan enfeksiyonlarını teşhis edebilmek için **MYCIN** isimli programı geliştirdiler. 450 kuralla **MYCIN** bazı genç hekimler kadar iyi uygulamalar gerçekleştiriyordu. **MYCIN** programının **DENDRAL** programından iki temel farklılığı vardı: 1) Kurallarının sonuç olarak çıkartılabileceği genel teorik bir model bulunmuyordu. Bu kurallar uzmanlarla yapılan röportajlara dayanıyordu. Röportajdaki bilgiler ise uzmanların kitaplardan öğrendiklerinden, diğer uzmanlardan edindikleri bilgilerden ve vaka tecrübelerinden oluşuyordu. 2) Kurallar tıp bilgisindeki belirsizlikleri içermektedir.⁸²

Şirketler tarafından da uzman sistemler geliştirilmeye başlanmıştır. **R1** ilk başarılı ticari uzman sistemdir. **Digital Equipment Corporation** bünyesinde 1982 yılında faaliyete geçmiştir. Program yeni bilgisayar sistemleri için siparişlerin ayarlanmasına yardım ediyordu ve 1986 itibariyle her yıl 40 milyon dolar tasarruf sağlamıştı. 1988 itibariyle **DEC**'in yapay zekâ grubu 40 tane uzman sistem

80 Vasif V. Nabyev: **Yapay Zekâ**, Ankara, Seçkin Yayıncılık, 2005, s. 445.

81 Daniel Crevier: **a.g.e.**, s. 148-150.

82 Daniel Crevier: **a.g.e.**, s. 151.

uygulaması gerçekleştirmişti. **Du Pont** o zamanlar 100 tane kullanırken, 500 tanesi de geliştirilmekteydi ve her yıl 10 milyon dolar tasarruf sağlanıyordu.⁸³

1.4.2 Fonların Geri Dönüşü

1981 yılında Japonlar “**Fifth Generation Project**”i duyurmuşlardır. On yıllık projede PROLOG⁸⁴ çalıştıran zeki bilgisayarların üretilmesi planlamıştır.⁸⁵ ABD cevaben **Microelectronics and Computer Technology Corporation**’ı (kısaca MCC) araştırma konsorsiyumu olarak kurmuş ve ulusal rekabeti garanti altına almıştır.⁸⁶ İki girişimde de yapay zekâ ile birlikte çip tasarımı ve insan-arayüzü araştırmaları da birlikte yürütülmüştür. Ancak, **MCC** ve **Fifth Generation**’ın yapay zekâ bileşenleri hedeflerini ulaşamamıştır. İngiltere’de Japonların **Fifth Generation** projesine karşılık vermek için başlatılan **Alvey Programı** ile yapay zekâ araştırmaları tekrar finanse edilmeye başlanmıştır.

Bütün bunların sonunda yapay zekâ endüstrisi 1980’deki birkaç milyon dolarlık cirodan 1988’e kadar birkaç milyar dolarlık ciroya çıkmıştır. Bundan kısa bir süre sonra birçok şirketin yapay zekânın imkânlarını aşan vaatlerini yerine getirememesi yüzünden “ikinci yapay zekâ kışı” olarak adlandırılan dönem başlamıştır.⁸⁷

1.4.3 Bağlantıcılığın Geri Dönüşü

1970 sonlarında bilgisayar bilimi yapay sinir ağları alanından neredeyse

83 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 24.

84 Özellikle yapay zekâ ve sayısal dilbilimde kullanılan mantık programlam dili.

85 Daniel Crevier: **a.g.e.**, s. 211-212.

86 Daniel Crevier: **a.g.e.**, s. 240.

87 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 24.

tamamen çekilmişti. **John Hopfield** (1933) 1982'de sinir ağların özelliklerini ve depolamasını analiz etmek için istatistiksel mekanik tekniklerini kullanmıştır.⁸⁸ Bu da öğrenme ve malumat işlemede tamamen yeni yöntemler sağlamıştır. Aynı yıllarda **David Rumelhart** (1942) ve **Geoff Hinton**'un (1947) başını çektiği psikologlar sinir ağlarını eğitebilmek için “**geri yayılım**” (backpropagation) isimli yöntemi geliştirdiler. 1980 ortalarında, **Bryson** ve **Ho**'nun 1969 yılında geliştirdiği perseptron yöntemine büyük bir dönüş yapılarak, bu yöntem bilgisayar bilimi ve psikoloji alanındaki pek çok öğrenme problemine uygulanmıştır.⁸⁹

David Rumelhart ve **James McClelland**'ın (1948) editörlüğünde 1986'da yayımlanan “**Parallel Distributed Processing**”⁹⁰ isimli çalışma optik karakter tanıma ve ses tanıma alanının temelini oluşturmuştur.⁹¹

1.5 İkinci Yapay Zekâ Kışı (1987-1993)

İş dünyasının yapay zekâya olan ilgisi 1980 sonlarında ekonomik daralma ile birlikte iyice gerilemiştir. Özellikle hükümet kuruluşları ile yatırımcıların yapay zekâ algısındaki çöküntüsüne ve ağır eleştirilere rağmen alandaki ilerleme sürmüştür. **Rodney Brooks** (1954) ve **Hans Moravec** robotbiliminin ilgili alanları üzerine çalışıyorlardı ve yapay zekâ için tamamen yeni bir yaklaşım önermekteydiler.

“**Yapay zekâ kışı**” deyiimi 1974'te kesilen fonlar için kullanılmıştır. Bu dönemi yaşayan araştırmacılar uzman sistemlerin yaşadığı kontrol edilemez sarmalı endişe ile takip etmişlerdir. 1980 sonlarında ve 1990'lı yılların başlarında yapay zekâ benzer bir finansal gerileme yaşamıştır.

1987 yılında **Apple** ve **IBM**'in ürettiği kişisel bilgisayarlar, hız ve güç

88 Daniel Crevier: **a.g.e.**, s. 291.

89 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 25.

90 Günümüzde hakim olan bağlantıcı yaklaşım olarak bilinir. Sinirsel işlem biçiminin paralel doğasına ve sinirsel temsillerin dağılmış doğasına vurgu yapan bir sinirsel ağ yaklaşımıdır. Araştırmacılara içinde işlem yapabileceği genel matematiksel bir çerçeve sağlamıştır.

91 Daniel Crevier: **a.g.e.**, s. 215.

bakımından, **Symbolics**'in ve diğer şirketlerin ürettiği çok pahalı **LISP makinelerini** kat kat geçmişlerdir. Artık bu pahalı makineleri almanın bir anlamı kalmamıştır. Uzmanlaşmış yapay zekâ donanım pazarındaki bu daralma yüzünden yapay zekâ endüstrisi bir anda yarım milyar dolar kaybetmiştir.⁹²

Neticede **XCON** gibi ilk uzman sistemlerin bakımını sağlamanın çok pahalı olduğu anlaşılmıştır. Güncellenmesi oldukça güç bu sistemler öğrenemiyor, sıradışı girdiler üzerine de anlamsız hatalar üretiyor ve yıllar önce tanımlanmış **nitelik problemi** (qualification problem) gibi problemlere yenik düşüyordu. Uzman sistemlerin sadece belirli bağlamlarda kullanışlı olduğu anlaşılmıştır.⁹³

1980'li yılların sonunda **Strategic Computing Initiative** ve **DARPA**'nın değişen yönetimi yapay zekâ fonlarını durdurmuştur. **DARPA** o zamandan itibaren hemen sonuç verecek projeleri desteklemiştir.⁹⁴

1991 yılı itibariyle, 1981'de başlatılan Japonların **Fifth Generation Project**'indeki çoğu hedefe ulaşamamıştır.⁹⁵ Doğrusu istenirse, “gündelik bir sohbeti sürdürebilmek” 2009 yılı itibariyle hâlâ gerçekleştirilememiştir.

1.5.1 Bedene Sahip Olmanın Önemi: Nouvella Yapay Zekâ

1980'li yılların sonunda birçok araştırmacı robotbilim temelli yeni bir yapay zekâ yaklaşımını savunmuşlardır. **Nouvella Yapay Zekâ** olarak adlandırılan bu yaklaşıma göre zekânın gösterilebilmesi için makinenin bir bedene ihtiyacı vardır. Makinenin algılaması, hareket etmesi, hayatta kalması ve dünya ile ilişki kurması gerekmektedir. Bu yaklaşıma göre sensorimotor yetenekler sağduyu akılyürütmesi için gereklidir ve soyut akılyürütme insanın fazla önemli olmayan yeteneklerinden biridir. Yaklaşımın savunucuları zekânın basamak basamak, basitten zora doğru inşa

92 Daniel Crevier: **a.g.e.**, s. 209-210.

93 Daniel Crevier: **a.g.e.**, s. 204.

94 Pamela McCorduck: **Machines Who Think** (2nd ed.), AK Peters, Massachusetts, 2004, s. 430-431.

95 Daniel Crevier: **a.g.e.**, s. 212.

edilmesi gerektiğini savunurlar.⁹⁶

Bu yaklaşım 1960'lı yıllardan beri pek gündemde olmayan sibernetik ve kontrol teorisi üzerine inşa edilmiştir. Bu alanın diğer öncülerinden biri de MIT'e 1970'li yılların sonunda gelmiş **David Marr**'dır (1945-1980). Nöroloji eğitimi almış **Marr, görme** (vision) grubunun başına geçmiştir. **Marr** bütün sembolik yaklaşımları (McCarthy'nin mantığı ve Minsky'nin çerçeve problemi) reddederek, yapay zekânın önce görmenin aşağıdan yukarıya bütün fiziksel mekanizmalarını anlaması gerektiğini savunmuştur. Ancak böyle bir anlamadan sonra sembolik işleme devreye girebilir. Lösemi hastası olan **Marr**'ın 1980'deki ani ölümüyle çalışmaları sona ermiştir.

1990 yılında robotbilim araştırmacısı **Rodney Brooks**'un yayımladığı **“Elephants Don't Play Chess”** isimli makalede fiziksel sembol sistemi hipotezi eleştirilir. **Brooks** sembollerin her zaman gerekli olmadığını iddia etmektedir.⁹⁷ 1980 ve 1990'lı yıllarda birçok bilim adamı zihnin sembol işleme modelini reddederek, bedenin akılyürütme için gerekli olduğunu iddia etmişlerdir.

1.6 1993'ten Günümüze Yapay Zekâ

Bugün yarım yüzyıldan daha yaşlı olan yapay zekâ alanı sonunda ilk hedeflerinden bazılarına ulaşmıştır. Yapay zekâ, teknoloji endüstrisinde başarılı bir biçimde kullanılmaya başlanmıştır. Başarının bir kısmı artan bilgisayar gücüne bağlı olarak, bir kısmı da belli izole problemlere odaklanarak ve onları bilimsel sorumluluğun en yüksek standardıyla takip ederek kazanılmıştır. Yine de iş dünyası yapay zekâyâ olan temkinli tutumunu korumaktadır. Alan dahilinde ise yapay zekânın 1960'larda dünyanın hayaline giren insan düzeyinde zekâ düşünüyü gerçekleştirilememesinin nedenleri konusunda pek uzlaşma yoktu. Tüm bu etkenler

⁹⁶ Pamela McCorduck: **a.g.e.**, s. 456.

⁹⁷ Rodney Brooks: “Elephants Don't Play Chess”, **Robotics and Autonomous Systems**, Sayı 6, 1990, s. 3-4 (Çevrimiçi: <http://people.csail.mit.edu/brooks/papers/elephants.pdf>).

hep birlikte, yapay zekânın özel problemlere ve yaklaşımlara odaklanan alt alanlara bölünmesine neden olmuştur. Kimi zaman bunlar “yapay zekâ” isminin lekeli geçmişini örten yeni isimler altında ortaya çıkmıştır. Yapay zekâ şimdiye kadar olduğundan hem daha dikkatli hem de daha başarılı olmuştur.⁹⁸

1.6.1 Dönüm Noktaları ve Moore Kanunu

11 Mayıs 1997'de IBM'in geliştirdiği **Deep Blue** isimli bilgisayar altı oyunluk seri sonunda dünya satranç şampiyonu **Gary Kasparov**'u yenmiştir⁹⁹. 2005 yılında Stanford Üniversitesi'nin geliştirdiği insansız otomobil 131 mil yol katederek **DARPA**'nın düzenlediği **DARPA Grand Challenge** müsabakasını kazanmıştır.¹⁰⁰ 2009 yılında duyurulan **Blue Brain Project**¹⁰¹ bir fare beyninin bölümlerini başarıyla taklit etmiştir.

Bu dönemdeki başarılar bir çeşit devrimsel paradigmaya değil, daha çok mühendislik yeteneğine ve bugünkü bilgisayarların muazzam gücüne bağlıdır. Aslında **Deep Blue**'nun bilgisayarı **Christopher Strachey**'nin (1916–1975) 1951'de satranç oynamayı öğrettiği **Ferranti Mark 1** isimli bilgisayardan 10 milyon kat daha hızlıdır.¹⁰² Bu dramatik artış bilgisayarların hız ve bellek kapasitesinin her iki yılda iki katına çıkacağını öngören **Moore Kanunu**'yla ölçülür. “Kaba bilgisayar gücü” problemi yavaş yavaş çözülmüştür.

98 Pamela McCorduck: **a.g.e.**, s. 424.

99 Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 27.

100“Stanford Racing”, (Çevrimiçi), <http://cs.stanford.edu/group/roadrunner/old/index.html>, 14 Ağustos 2009.

101“Artificial brain '10 years away””, **BBC News**, (Çevrimiçi), <http://news.bbc.co.uk/2/hi/technology/8164060.stm>, 14 Ağustos 2009.

102Ferranti Mark 1'in döngü süresi 1,2 milisaniyeydi, bu yaklaşık 833 flop'a denk geliyordu. Deep Blue 11,38 gigaflopta çalışıyordu. Yaklaşık olarak bunlar arasında 10⁷ kat fark vardır.

1.6.2 Zeki Ajanlar

“**Zeki ajanlar**” denen yeni bir paradigma 1990’lı yıllar boyunca yaygın olarak kabul görmüştür. **Judea Pearl** (1936), **Alan Newell** ve diğerlerinin **karar teorisi** (decision theory) ve **ekonomiden** aldıkları kavramları yapay zekâ çalışmalarına uygulamalarıyla zeki ajanlar modern formuna ulaşmıştır.¹⁰³

Zeki ajan, bir ortamda gözlem yaparak eylemde bulunan (yani bir ajan) ve eylemlerini ulaşılmak istediği hedeflerine göre yönlendiren (yani akıllı) otonom bir varlıktır. Zeki ajanlar **bilgi gösterimi** ve/veya **makine öğrenmesi** sistemleri kullanır. Mevcut bilgileri ve/veya öğrenme algoritmaları sayesinde edindikleri yeni bilgileri kullanarak hedeflerine ilerler. Zeki ajan yaklaşımında canlılar da zeki ajan olarak değerlendirilir. En kompleks yapıya sahip olan zeki ajan ise insandır. Diğer yandan termostat gibi basit refleks makinelerine de ajan denir. Dolayısıyla ajanlar çok basit yapılar olabileceği gibi, bir hedef doğrultusunda birlikte çalışan bir insan topluluğu gibi çok kompleks yapılar da ihtiva edebilir.¹⁰⁴

Zeki ajanlar, yazılım ajanlarıyla (kullanıcıların yerine çeşitli görevleri yapan otonom yazılım programları) yakından ilgilidir. Bilgisayar biliminde zeki ajan terimi, ajanın yapısının ne kadar kompleks olduğuna bakmadan, yazılım ajanı için de kullanılabilir. Örneğin, işlem asistanı ya da veri madenciliği için kullanılan otonom programlar da “zeki ajanlar” olarak adlandırılırlar.¹⁰⁵

1.6.3 Tertiplilerin Zaferi

Yapay zekâ araştırmacıları geçmişte yaptıklarından daha sık biçimde sofistike matematik araçları geliştirmeye ve kullanmaya başlamıştır. Yapay zekânın çözmesi gereken problemlerin birçoğunun matematik, ekonomi veya yöneylem araştırması

103Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 27, 54-55.

104Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 32.

105Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 27, 32, 55.

gibi alanlardaki araştırmacılarca zaten çalışıldığı yaygın olarak anlaşılmıştır. Paylaşılan matematik dil, hem daha köklü ve başarılı alanlarla daha üst düzey bir işbirliği yapılmasını hem de ölçülebilen ve kanıtlanabilen sonuçlara ulaşılmasını sağlamıştır. Yapay zekânın bir bilim dalı haline gelmesi bir “devrim” olarak görülmüş ve “tertiplilerin zaferi” olarak nitelendirilmiştir. Bu süreç içerisinde yapay zekânın alt branşlarında uzmanlaşmalar görülmektedir. Örneğin robotbilimi ve görme, yapay zekâ çatısı altından giderek uzaklaşmıştır.¹⁰⁶

Judea Pearl'ün 1988'deki “**Probabilistic Reasoning in Intelligent Systems**” isimli kitabı yapay zekâyâ olasılık ve karar teorisini getirmiştir. Kullanımda olan birçok yeni araç arasında **Bayes ağları, gizli Markov modelleri, bilgi teorisi, stokastik modelleme** ve **klasik optimizasyon** da vardır. “Bilgisayar zekâsı” paradigmaları için sinir ağları ve evrimsel algoritmalar gibi kesin matematiksel tanımlar da geliştirilmiştir.¹⁰⁷

1.6.4 Sahne Arkasındaki Yapay Zekâ

Yapay zekâ araştırmacıları tarafından geliştirilen algoritmalar daha büyük sistemlerin bir parçası olarak görünmeye başlanmıştır. Yapay zekâ birçok zor problemi çözmüş ve bu çözümler teknoloji endüstrisinde kullanılmıştır. Bu alanların başında veri madenciliği¹⁰⁸, endüstriyel robotlar, lojistik, konuşma tanıma, bankacılık programları, medikal tanı ve **Google**'ın arama motoru gelmektedir.¹⁰⁹

Yapay zekânın bu başarılarında neredeyse hiç adı geçmez. Yapay zekânın getirdiği yeniliklerin birçoğu bilgisayar bilimlerinin alet çantasına katılan bir başka nesne durumuna düşürülmüştür.¹¹⁰ **Nick Bostrom** (1973) 2006'daki CNN demecinde şunları söyler: “Birçok ileri yapay zekâ ürünü genel uygulamalara geçirilmiş,

106Pamela McCorduck: **a.g.e.**, s. 486-487; Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 25-26.

107Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 26.

108Veri madenciliği, veriden modeller çıkarma işlemidir.

109Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 27.

110Pamela McCorduck: **a.g.e.**, s. 423.

çoğunda da yapay zekâ adıyla anılmamıştır. Çünkü bir şey yeterince kullanışlı ve yaygın hale gelirse, artık yapay zekâ olarak adlandırılmaz.”¹¹¹

Bugün birçok yapay zekâ araştırmacısı yaptıkları işleri bilerek başka isimlerle anıyorlar: enformatik, bilgi-temelli sistemler, bilişsel sistemler ya da bilgisayar zekâsı gibi. Bunun bir nedeni, alanlarının yapay zekâdan temelde ayrı olarak görmeleri olsa da, yeni isimlerin finansman temin etme konusundaki faydası da ihmal edilmemelidir. En azından ticari dünyada “yapay zekâ kışı”, yapay zekâ araştırmalarını gölgelemeyi sürdürmektedir. **New York Times**’ın 2005’te bildirdiği gibi: “Bilgisayar uzmanları ve yazılım mühendisleri çılgın hayalperestler olarak görülmekten korktuklarından yapay zekâ teriminden kaçınıyorlar.”¹¹²

1.6.5 HAL 9000 nerede?

1968 yılında **Arthur C. Clarke**¹¹³ (1917-2008) tarafından yazılan ve **Stanley Kubrick**¹¹⁴ (1928-1999) tarafından filme alınan “**2001: A Space Odyssey**” isimli yapıtta, 2001 yılına gelindiğinde insanların kapasitesine denk veya onu aşan zekâyâ sahip makineler olacağı hayal edilmiştir. **HAL-9000** kitaptaki zeki bilgisayarın adıydı. Birçok önde gelen yapay zekâ araştırmacısı da 2001 yılına gelindiğinde **HAL-9000** gibi bir makinenin olacağına inanmıştır.¹¹⁵

Marvin Minsky şöyle sorar: “Peki 2001’de neden HAL’ı göremedik?” **Minsky** yanıtın sağduyu akılyürütmesi gibi temel problemlerin göz ardı edilmesi, araştırmacıların çoğunun sinir ağları ya da genetik algoritmalar gibi ticari

111“AI set to exceed human brain power”, (Çevrimiçi),

<http://www.cnn.com/2006/TECH/science/07/24/ai.bostrom>, 14 Ağustos 2009.

112“Behind Artificial Intelligence, a Squadron of Bright Real People”, The New York Times,

(Çevrimiçi), <http://www.nytimes.com/2005/10/14/technology/14artificial.html>, 14 Ağustos 2009.

113İngiliz mucit ve bilimkurgu yazarı. Yazdığı bilimkurgu romanı 2001: Bir Uzay Destanı ve yönetmen Stanley Kubrick ile birlikte çalıştığı aynı isimli film ile meşhurdur. Clarke, Robert A. Heinlein ve Isaac Asimov’la birlikte, “bilim-kurgunun üç büyük yazarı”ndan biri olarak kabul edilmektedir.

114ABD’li film yönetmeni.

115Daniel Crevier: **a.g.e.**, s. 108, 317.

uygulamalara yönelmesi olduğuna inanır.¹¹⁶ Öte yandan **John McCarthy** insan seviyesinde bir yapay zekâyâ ne zaman sahip olacağımız sorusunu yanlış bir soru olarak değerlendirir. McCarthy'ye göre insan-seviyesinde bir yapay zekâyâ ancak aşağıdaki üç temel problem çözüldüğünde ulaşılabilir:

1. **çerçeve problemi** (frame problem): bir eylem oluştuğunda nelerin gerçekleşmediğinin belirtilmesi ,
2. **nitelik problemi** (qualification problem): bir eylemin başarıya ulaşabilmesi için tüm niteliklerin belirtilmesi,
3. **çatallaşma problemi** (ramification problem): bir eylemin tüm etkilerinin belirtilmesi.

Bu problemler önemli uygulamalar için çözümler getirilmiş olsa da bunlar insan zekâsı seviyesinde değildir.¹¹⁷ Bu da göstermektedir ki, insanlığın yüzyıllardır düşlediği zeki makineler, bazı problemler aşılmadan gerçekleştirilebilir değildir. Fakat yapay zekâ araştırmaları sayesinde insanların çalışma ortamlarındaki angaryalar azalmış, daha güvenli üretim imkânları sağlanmış, bilgiye ulaşılabilirlik artmıştır. **McCarthy**'nin de dediği gibi, yapay zekâyı insan zekâsını taklit etmek olarak tanımlamak yanlıştır. Yapay zekâ belki de daha çok problem çözme kabiliyetinin iyileştirilmesidir.¹¹⁸

Yapay zekâ disiplini 1956 yılında kurulmuş oldukça yeni bir disiplindir. Bu alanda yapılan çalışmalar teknolojinin eliyle hayatımızdaki pek çok otomatikleştirilmiş uygulamanın arkaplanını oluşturmuştur. Yapay zekânın bunu hangi özellikleriyle sağladığı problematik bir şekilde ikinci bölümde ele alınacaktır.

116Marvin Minsky, "It's 2001. Where Is HAL?", **Dr. Dobb's Technetcast**, (Çevrimiçi), <http://www.ddj.com/hpc-high-performance-computing/197700454>, 14 Ağustos 2009.

117John McCarthy: "What Was Expected, What We Did, and AI Today", (Çevrimiçi), http://www.engagingexperience.com/2006/07/ai50_ai_past_pr.html, 13 Ağustos 2009.

118John McCarthy: "What Was Expected, What We Did, and AI Today", (Çevrimiçi), http://www.engagingexperience.com/2006/07/ai50_ai_past_pr.html, 13 Ağustos 2009.

2. Yapay Zekânın Özellikleri

2.1 Yapay Zekâ Nedir?

Yapay zekâ, insanlarda zekâ ile ilgili zihinsel fonksiyonları bilgisayar modelleri yardımıyla inceleyip bunları formel hale getirdikten sonra yapay sistemlere uygulamayı amaçlayan bir araştırma alanıdır.¹¹⁹

Yapay zekâ oluşum sürecinde bilgisayar biliminin, özellikle de bu alanın mühendislik kısmının, bir yan alanı olarak yorumlanmıştır. Fakat yapay zekâ araştırmaları aynı zamanda felsefe, psikoloji, dilbilim ve nöroloji ile yakın işbirliği içindedir; çünkü bu disiplinler yapay zekânın bilişsel modelini oluşturmaktadır. Bu model yapay zekânın disiplinlerarası karakterini belirginleştirir. Yapay zekânın bu disiplinlerarası karakteri, onun, bilgisayar biliminin diğer dallarının aksine, iyi formüle edilmemiş ve genelgeçer çözümleri olmayan problemlerle de uğraşmasına imkân sağlar. Sınırları akışkandır ve problem sahası sürekli gelişmektedir.¹²⁰

Yapay zekâ araştırmacıları yapay zekâ sistemlerini **dört** yaklaşım çerçevesinde ele alır: 1) insanlar gibi davranan sistemler, 2) insanlar gibi düşünen sistemler, 3) rasyonel düşünen sistemler, 4) rasyonel davranan sistemler. 2. ve 3. yaklaşımlar düşünme süreçleri ve akılyürütmeyi önplana taşırken, 1. ve 4. yaklaşımlar davranışa vurgu yaparlar. 1. ve 2. sistemler insan performansına uygunluğuna göre başarıyı ölçmektedir, 3. ve 4. yaklaşımlar ideal zekâ kavramına göre başarıyı ölçmektedir. Yapay zekâda bu dört yaklaşım da takip edilmiştir, fakat genel karşıtlık insan temelli yaklaşım ile rasyonalite temelli yaklaşım arasında sürmektedir. İnsan temelli yaklaşım hipotezler ve deneysel doğrulamalar gerektiren

¹¹⁹Şakir Kocabaş: **Yapay Zekâya Giriş**, s. 3, (Çevrimiçi),
http://www.sakirkocabas.com/files/yzgir_1n.rtf, 27 Ağustos 2009.

¹²⁰Günther Görz, Bernhard Nebel: **Yapay Zekâ**, İstanbul, İnkilâp Kitabevi, 2005, s. 12.

empirik bir bilim anlayışında sürdürülürken, rasyonalite temelli yaklaşım ise matematik ve mühendisliğin bir bileşimini oluşturmaktadır.¹²¹

2.1.1 İnsanlar Gibi Davranan Sistemler (Turing Testi Yaklaşımı)

Yapay zekâ araştırmacılarının baştan beri ulaşmak istediği ideal, insan gibi davranan sistemler üretmektir. **Alan Turing** zeki davranışı, bir sorgulayıcıyı kandıracak kadar bütün bilişsel görevlerde insan düzeyinde başarı göstermek olarak tanımlamıştır. Bunu ölçmek için **Turing Testi** olarak bilinen bir test önermiştir. Yapay zekâ araştırmacıları **Turing Testi**'ni geçecek bilgisayarlar/sistemler üretmek için çok az gayret göstermişlerdir. Zeki davranışı taklit etmek yerine, zekânın prensiplerinin araştırılması araştırmacılar için daha önemli olmuştur. Örneğin “yapay uçuş” çalışmaları **Wright Kardeşlerin** kuşları taklit etmeyi bırakıp, aerodinamiği öğrenmeleri ile başarılı olmuştur. Uçak mühendisleri, alanlarının amacını “diğer yaraları kandırsın diye aynen yarasalar gibi uçabilen makineler” yapmak olarak tanımlazlar. İnsan gibi davranan sistemlerde önemli olan zeki davranışın üretilmesidir. Bunun hangi modelle üretildiği ikinci planda yer almaktadır.¹²²

2.1.2 İnsanlar Gibi Düşünen Sistemler (Bilişsel Modelleme Yaklaşımı)

İnsan gibi düşünen bir program üretmek için, insanların nasıl düşündüğünü belirlememiz gerekir. İnsan zihninin çalışma süreçlerini iki yöntem ile izleyebiliriz: 1) içebakış, 2) psikolojik deneyler. Açık bir şekilde zihnin teorisini ortaya koyduktan sonra, teorinin bir bilgisayar programıyla ifade edilebilmesi mümkün hale gelir. Örneğin **General Problem Solver** programını geliştiren **Allen Newell** ve **Herbert**

¹²¹Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 1-2.

¹²²Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 2-3.

Simon problemlerin doğru bir şekilde çözülmesiyle ilgilenmiyorlardı. Onlar daha çok programın akılyürütme sırasında takip ettiği yolu, aynı problemi çözen bir insanın takip ettiği yolla karşılaştırmanın derdindeydiler. Bilişsel bilimin disiplinlerarası araştırmaları yapay zekâdaki bilgisayar modelleriyle psikolojinin deneysel tekniklerini bir araya getirerek, insan zihninin çalışmasına ilişkin kesin ve test edilebilir teoriler oluşturmaya çalışır.¹²³

2.1.3 Rasyonel Düşünen Sistemler (Düşünmenin Yasaları Yaklaşımı)

Aristoteles'ten (M.Ö. 384-322) beri “doğru düşünme”, yani akılyürütme süreçleri mantık çatısı altında sistemleştirilmeye çalışılır. Düşünmenin yasaları ortaya koyularak zihnin işleyişi belirlenmek istenir.

19. yüzyılda mantıkçılar, akılyürütme işlemlerini ifade edilebilecek mantık sistemleri üzerinde durmuşlardır. 1965 yılı itibariyle mantık notasyonunda tanımlanmış herhangi bir problem bilgisayar programlarıyla çözülebilir hale gelmiştir. Yapay zekâda çok önemli yer tutan mantıkçı gelenek zeki sistemler üretmek için bu çeşit programlar üretmeyi amaçlamaktadır. Bu yaklaşımı kullanarak gerçek sorunları çözmeye çalışınca iki önemli engel karşımıza çıkmaktadır. Gündelik yaşamla içiçe olan, çoğu kez de belirsizlik, kaypaklık ve çokanlamlılık içeren konuşma dilini mantığın işleyebileceği bir şekilde formelleştirmek hiç de kolay değildir. Bir başka güçlük de problemlerin çözüm yolları arttıkça bilgisayar kapasitesi ihtiyacının üssel olarak artmasıdır.¹²⁴

¹²³Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 3-4.

¹²⁴Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 4.

2.1.4 Rasyonel Davranan Sistemler (Rasyonel Ajan Yaklaşımı)

“Rasyonel ajan”, çevresini algılayan ve en iyi sonucu getirecek davranışı üretmeye çalışan sistemlerin genel adıdır. Bu yaklaşımda yapay zekâ, rasyonel ajanların incelenmesi ve oluşturulması olarak tanımlanmaktadır. Böylece bir önceki yaklaşımdan daha genel bir model sunarak, girdileri doğru algılayacak ve doğru çıktıları üretecek sistemlerin geliştirilmesi hedeflenir.

“Düşünmenin yasaları yaklaşımı”ndan farklı olarak, rasyonel ajanların çıkarımdan kaynaklanmayan bazı rasyonel davranışlar da üretebilmesi beklenir. Örneğin, sıcak bir şeye değen insanın elini çekmesi bir refleks harekettir ve uzun düşünme süreçlerine girmeden yapılır. **Turing Testi** için gerekli tüm yetenekler rasyonel hareketler üretebilmek içindir. Rasyonel ajanlar görme, işitme gibi algılama süreçleriyle elde edilen verileri düşünmenin yasaları ile işleyerek zeki davranışı üretmeye çalışmaktadır.

Yapay zekânın rasyonel ajan tasarımı olarak ele alınmasının iki avantajı vardır. Birincisi, “düşüncenin yasaları” yaklaşımından daha genel olmasıdır; çünkü doğru çıkarım rasyonaliteyi sağlamak için birkaç olası mekanizmadan sadece biridir. İkincisi, bilimsel gelişme bakımından insan davranışı veya insan düşüncesine dayanan yaklaşımlardan daha uygundur; çünkü rasyonalitenin standardı açıkça tanımlanmıştır ve tamamen geneldir. Öte yandan insan davranışı özel bir ortama iyi uyum sağlamıştır; kısmen karmaşık, büyük ölçüde bilinmeyen ve halen mükemmele ulaşmaya çok uzaktır.¹²⁵

2.2 Yapay Zekânın Düşünsel Tarihi Geçmişi

İlk bölümde ele aldığımız yapay zekâ tarihçesinde yapay zekânın modern dönemdeki gelişim çizgisini ortaya koymaya çalıştık. Dijital bilgisayarlara kadar

¹²⁵Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 4-5.

uzanan bu sürecin kökeninde iki temel fikir yer almaktadır.

Bunlardan ilki yapay insan yaratma mitidir. Yapay insan fikri, Grek mitolojisindeki metalin tanrısı **Hephaistos**'un bronzdan yaptığı **Talos** isimli mekanik adama kadar uzanır. “**Artes mechanicae**” diye adlandırılan mekanik sanatlar Antik Çağ'dan beri insana benzeyen, insana yardımcı olan otomatlar üretmek amacıyla. **Aristoteles** otomatlarla ilgili olarak **Politika** adlı yapıtında emirlere itaat eden ve gelecekle ilgili planlar yapan, inşaat ustasını çıraqlardan, efendiyi de kölelerinden vazgeçirebilecek aletlerden bahseder.

Antik dönemin doğa bilimleri, tıp, felsefe ve teknik konusundaki mirası İslam medeniyeti tarafından da alınılmıştır. Saat yapımı başta olmak üzere, pek çok otomat geliştirilmiştir. 1153-1233 arasında Diyarbakır'da yaşamış **El Cezerî** robotik biliminin babası kabul edilen ve sibernetik üzerine çalışmalar yapan ilk müslüman bilim adamı ve mühendistir. İslam medeniyetinin otomat üretmek konusundaki başarıları Arap şiirine de konu olmuş, pek çok mit ve öyküde android ve otomatlar yer almıştır.

Geç Antik dönemdeki yüksek teknik standartlara ancak Rönesans döneminde tekrar ulaşılmıştır. **Leonardo da Vinci**'nin (1452-1519), hareket süreçlerinin daha iyi anlaşılmasını sağlayan anatomik beden çizimleri ve **William Harvey** (1578-1657) tarafından kan dolaşımının keşfedilişiyle insan bedenine yönelik araştırmalar, insan organizmasının tamamen farklı modellerine yol açmıştır. Bu süreç **Descartes**'in (1596-1650) düalizmin anlayışıyla birlikte hızlanmıştır.¹²⁶

Zaman içinde ikinci bir insan türü yaratma hayaliyle zaman ölçen çark mekanizmalarından, endüstri makinelerine kadar sayısız otomat üretilmiştir.¹²⁷ **Mary Shelley**'in (1797-1851) 1818'de kaleme aldığı “**Frankenstein ya da Modern Prometheus**” romanından beri de zeki makineler, robotlar, androidler ve otomatlar bilim kurgunun en sevilen temalarını oluşturur.¹²⁸ Fakat ilk bölümde ortaya koyduğumuz üzere bilim-kurgu kaynaklı izlenimler yapay zekâ alanının bulunduğu konumu göstermek bakımından yanıltıcı olabilmektedir. Yapay zekânın hali

¹²⁶Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 939

¹²⁷Günther Görz, Bernhard Nebel: **a.g.e.**, s. 17-19

¹²⁸Blay Whitby: **a.g.e.**, s. 21

hazırdaki durumu bilim-kurgunun hayallerinin dışında kalmaktadır.

İkinci fikir olan insan düşüncesinin formülleştirilmesinin kökleri, Çin, Hint ve Grek filozofların mantık çalışmalarına kadar uzanır. 17. yüzyılda **Hobbes**, **Descartes** ve **Leibniz**'in önderliğini sürdürdüğü düşüncenin matematik ve geometriye indirgenebileceği anlayışını **Leviathan** isimli kitabında **Thomas Hobbes** (1588-1679) şöyle özetler: “Akılyürütme sadece bir hesaplama”¹²⁹.

Matematik ve mantık geleneğinden gelen birçok düşünür ise tarih boyunca insanın tinsel edimlerini matematiksel fonksiyonlarla taklit etmek istemişlerdir. Bu düşünürler tarih boyunca otomat ve saat ustalarından farklı olarak, yaşayan töz yerine madde koymakla ilgilenmemişlerdir.¹³⁰ Bunun yerine düşünmeyi yapısal olarak betimlemeye ve böylelikle otomatikleştirilebilir hale getirmeye çalışmışlardır.

Tarihi geçmişi içinde insan davranışlarını taklit eden ve insan gibi düşünen makine yapımı bir hayal olarak kalırken, günümüzde bu konuda şaşırtıcı gelişmeler birbirini izlemektedir.

2.2.1 Yapay Zekânın Felsefi ve Mantıksal Arkapları

Yapay zekâ denilince makinelerin yapımının yanı sıra, formel bir dil inşa ederek insan gibi düşünen araçların yapımı da anlaşılmaktadır. Bu araçların yapılmasında felsefi ve mantıksal çalışmalar da söz konusu olur.

Yapay zekânın geçmişi içinde aşağıda kısaca işaret edeceğimiz dönüm noktalarından ilki, kuşkusuz **Aristoteles**'tir. **Aristoteles** (M.Ö. 384-322) zihnin rasyonel kısmına hükmeden kesin yasaları formüle eden ilk kişidir. Geçerli akılyürütmenin ifadesi ve denetlenmesi için bir kıyas sistemi geliştirmiştir.¹³¹ Öyle ki bu sistem kişinin belli öncüllerden mekanik olarak sonuçlar çıkarmasına olanak verir. Çok daha sonra **Ramon Llull** (1232-1315) akılyürütmenin mekanik bir insan yapısı

¹²⁹David Wootton: **Modern political thought: readings from Machiavelli to Nietzsche**, Indiana, Hackett Publishing, 1996, s. 137.

¹³⁰Günther Görz, Bernhard Nebel: **a.g.e.**, s. 23.

¹³¹**Bkz:** Şafak Ural: **Temel Mantık**, İstanbul, Çantay Kitabevi, 1995, s. 81.

tarafından da yürütülebileceği fikrini ortaya atmıştır. **Thomas Hobbes** (1588-1679) akılyürütmenin sayısal hesap işlemine benzediğini ileri sürmüştür.¹³²

M.Ö. 2400'de Çin'de abaküs ile başlayan hesap makinesi çalışmaları, 1500 civarında **Leonardo da Vinci**'nin (1452-1519) tasarladığı hesap makinesi ile mekanikleştirilmiştir. Yapılan son rekonstrüksiyonlarda bu tasarımın çalıştığı görülmüştür. Bilinen ilk hesap makinesi 1623 yılı civarında Alman bilim adamı **Wilhelm Schickard** (1592-1635) tarafından yapılmıştır, ancak **Blaise Pascal**'ın (1623-1662) 1642'de yaptığı **Pascaline** daha ünlüdür. **Gottfried Wilhelm Leibniz** (1646-1716) sayılardan ziyade kavramlar üzerinde işlem yapmak üzere mekanik bir aygıt yapmıştır, ancak bu aygıt oldukça sınırlı bir faaliyet alanına sahiptir.¹³³

Zihnin formel ve rasyonel yönünü ortaya koymaya yönelik kurallar fikrinden sonraki adım, zihnin fiziksel bir sistem olarak tasarlanmasıdır. **Descartes** (1596-1650) zihin-madde ayrımını ve ortaya çıkan problemleri açıkça tartışan ilk isimdir. Zihnin sadece fiziksel olduğu düşüncesi özgür irade problemine sebep olmuştur. Eğer zihin tamamen fiziksel kuralların hükmü altındaysa, bir taşın dünyanın merkezine düşmeye “karar vermesi”nden daha fazla özgür iradeye sahip olamayacaktır. Akılyürütmenin gücüne oldukça inanan biri olmasına karşın **Descartes** beden-ruh ayrımını yaparak, insan zihninin (ya da ruhunun) doğadışı ve fiziksel kurallara tabi olmayan bir parçası olduğunu savunur. **Descartes** bu **düalist** yapının sadece insanlarda olduğunu, hayvanların ruhu olmadığı için de onların birer makine olarak sayılabileceğini söyler. **Materyalizme** göre beynin fizik kurallarına göre yaptığı işlemler zihni oluşturur. **Descartes**'in doğadışı unsurlarını reddeden bu anlayış, özgür iradeyi ise mümkün seçenekler arasından birini seçebilme özgürlüğü olarak tanımlar.¹³⁴

“Bilgi işleyen fiziksel bir zihin” yaklaşımının bir sonraki adımı, bilginin kaynağını belirlemektir. **Francis Bacon**'ın (1561-1626) **Novum Organum** adlı eseriyle başlayan **emprizim** akımı, **John Locke**'un (1632-1704) şu sözüyle tanımlanmıştır: “Önce duyularda olmayan hiçbir şey anlıkta değildir.” **David Hume**

¹³²Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 6.

¹³³Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 6.

¹³⁴Pamela McCorduck: **a.g.e.**, s. 38-39.

(1711-1776) 1739 yılında yayımladığı **A Treatise of Human Nature** adlı eserinde **tümevarımı** şu şekilde tanımlar: Ögeleri arasındaki tekrarlanan ilişkiler ile karşılaşınca elde edinilen genel kurallar.¹³⁵

Mantığın sembolleştirilmesi düşüncesi antik dönemdeki Çinli, Hintli ve Grek filozoflara kadar uzanır. Ancak matematiksel gelişimi gerçekte önerme ya da Boolean mantığının detaylarını oluşturan **George Boole**'un (1815-1864) çalışmalarıyla başlamıştır. 1879'da **Gottlob Frege** (1848-1925) Boole'un mantığını nesneleri ve ilişkileri de içerecek şekilde genişletmiş, bugün en temel bilgi gösterim sistemi olarak kullanılan önermeler mantığını oluşturmuştur. **Alfred Tarski** (1902-1983) mantıktaki nesnelerle gerçek dünyadaki nesneler arasındaki ilişkiyi gösteren bir semantik teori kurmuştur. Bir sonraki adım ise mantık ve hesaplama ile neler yapılabileceğinin sınırlarını belirlemek olmuştur.¹³⁶

Boole ve **Frege** mantıksal tümdengelim için algoritmalar oluşturmaya çalışmıştır. 19. yüzyılın sonlarında çalışmaları genel matematiksel akılyürütmeyi mantıksal tümdengelim olarak formelleştirmeye çok yaklaşmıştır. 1900'de **David Hilbert** (1862-1943) yüzyıl boyunca matematikçileri uğraştıracak doğru şekilde tahmin ettiği 23 problemlik bir liste oluşturmuştur. Son problemde (Entscheidungsproblem ya da karar verme problemi) doğal sayıları içeren herhangi bir mantık önermesinin doğruluğuna karar verecek bir algoritma olup olmadığını sorar. **Hilbert** temelde, etkili kanıt süreçlerinin sınırları olup olmadığını soruyordu. 1930'da **Kurt Gödel** (1906-1978), 1931'de mantıkta gerçek sınırların olduğunu ispatlamıştır. Onun **eksiklik teoremi** doğal sayıların özelliklerini ifade edebilen herhangi bir dilde, doğrulukları herhangi bir algoritmayla ispatlanamayan; yani karar verilemez doğru yargılar olduğunu göstermiştir.¹³⁷

Eksiklik teoremi hesaplanamayan fonksiyonlar olduğunun gösterilmesi olarak da yorumlanabilir. Bu **Alan Turing**'i (1912-1954) tam olarak hangi fonksiyonların hesaplanabilir olduğunu tanımlamaya itmiştir. Bu nosyon aslında biraz problematiktir, çünkü bir hesaplama ya da verimli süreç nosyonu formel bir biçimde

135Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 6.

136Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 7-8.

137Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 8.

tanımlanamaz. Buna rağmen, Turing makinesinin hesaplanabilir herhangi bir fonksiyonu hesaplayabildiğini savunan **Church-Turing** tezinin yeterli bir tanım verdiği genel olarak kabul görmüştür. **Turing** hiçbir Turing makinesinin hesaplayamayacağı fonksiyonlar olduğunu da göstermiştir. Örneğin, hiçbir makine bir programın bir girdiye karşılık bir yanıt mı vereceğini yoksa sonsuza kadar mı çalışacağını söyleyemez.¹³⁸

Frege'nin mantıkta açtığı yeni çağır yeni felsefi akımların üzerinde de etkili olmuştur. **Rudolf Carnap**'ın (1891-1970) da dahil olduğu ünlü Viyana Çevresi'nin **mantıksal pozitivizm** öğretisini geliştirmiştir.¹³⁹ Bu öğretiye göre mantık teorileri, bütün bilgilerin empirik verilere karşılık gelen gözlem cümleleriyle açıklanabileceğini iddia eder. **Carnap** ve **Carl Hempel**'ın (1905-1997) **doğrulama teorisi**, bilginin gözlem aracılığıyla nasıl elde edildiğini modellemeye çalışmıştır. **Carnap**'ın **The Logical Structure of the World** (1928) adlı eseri temel deneyimlerden bilgi oluşturmaya ait açık bir hesaplama süreci tanımlamıştır; zihnin süreçlerini birer hesaplama olarak gören ilk zihin teorisidir.¹⁴⁰

Yapay zekânın son hedefi, eylemde bulunabilen makinelerin inşasıdır. Bu eylem, yukarıda yapılan açıklamalar da dikkate alındığında, ilkin algı (perception) adımı; daha sonra bu algının anlamlandırılması (cognition) ve son olarak da reaksiyon gösterebilen, yani karar veren, eylemde (action) bulunan makinelerin tasarlanması demektir. İnsan ve bilinç sahibi tüm canlıların ortak özelliği, bu üç adımda işaret edilen aşamalara sahip olmalarıdır.

2.3 Yapay Zekânın Problemleri

Zekâyı taklit etme problemi belirli alt-problemlere ayrılmıştır. Bu problem alanları araştırmacıların yapay zekâdan ne anladıklarıyla da yakından ilgilidir.

138Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 8.

139Şafak Ural: **Pozitivist Felsefe**, İstanbul, Say Yayınları, 2006, s. 56-57.

140Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 6.

Aşağıda belirtilen problemler, yapay zekâ araştırmacılarının yapay zekâyı nasıl tanımladıklarını da göstermektedir.

2.3.1 Arama Yaparak Problem Çözme

Yüksek derecede karmaşık arama ağaçlarında hızlı çözüm yolları bulma amacına hizmet eden arama yöntemleri yapay zekânın hemen hemen bütün yan alanlarında büyük öneme sahiptir. İlk yapay zekâ araştırmacıları insan zekâsının aşamalarını adım adım takip eden algoritmalar geliştirmişlerdir. **Arama ağacı, yol gösterici yöntemler ve liste işleme** gibi algoritmalar basit yapbozların çözümü, mantıksal ve geometrik önermelerin kanıtlanması ve satrançla dama gibi oyunlarda kullanılmıştır. Oyunlarda kullanılan yöntemler arama yöntemlerinden ayrılır. Çünkü ortada bir çözüm yolu yoktur. Fakat tıpkı arama yönteminde olduğu gibi burada da karmaşık bir arama ağacı keşfedilmelidir.¹⁴¹ 80'li yılların sonunda ve 90'lı yıllarda yapay zekâ araştırmacıları olasılık ve ekonomideki kavramları kullanarak belirsiz veya eksik malumat ile başa çıkabilecek başarılı yöntemler ortaya koymuşlardır.

Çözüm yolları birden çok olan problemler için pek çok algoritma devasa hesaplama kaynaklarına ihtiyaç duyar. Daha etkili problem çözme algoritmaları ister istemez yapay zekânın öncelikli araştırmalarından birisi değil midir?¹⁴²

İnsanlar birçok problemi çözerken tümdengelim adım adım uygulamak yerine, hızlı sezgisel yargılar kullanır. Yapay zekâ bu tarzda “alt-sembolik”¹⁴³

141Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 185.

142Günther Görz, Bernhard Nebel: **a.g.e.**, s. 100.

143İnsanları zeki yapan şeyin ne olduğu konusu çok tartışılmıştır. Zihnin gücünü sembolik olarak akıl yürütebilme özelliğinden aldığı iddia edilmiştir. Bu görüş, bilginin bir semboller kümesine doğrudan denk bir biçimde depolandığı ve akılyürütmenin bu sembollerin manipülasyonu olduğu düşüncesine dayanır. Alternatif bir görüş, akılyürütmenin bu kadar basit bir şeye bağlı olmadığını. Bu görüşe göre, bilgi dağılmış bir biçimde depolanır ve bir tür ağ ile en iyi şekilde temsil edilmektedir. Bu durumda akılyürütme genellikle ağın düğümlerindeki ağırlıkların ayarlanması olarak temsil edilir. Bu tür akılyürütme modelleri bazen “sembolik olmayan” olarak tanımlanır. Newell’in *Unified Theories of Cognition* eserinde açıkladığı görüş, bu ikisinin bileşimidir. Yani zekâ hem sembolik hem de sembolik olmayan süreçleri içerir. Sembolik akılyürütme bilinçli düşüncüyü gerektiren görevlerle ilgilidir; sembolik olmayan akılyürütme ise bilinçsiz bir düzeydeki

problem çözümlerinin taklitlerinde az da olsa ilerleme kaydetmiştir. İnsanlardaki sensorimotor yetenekleri ileri akılyürütmeye olan etkileri bu yaklaşım içerisinde değerlendirilir. Yapay sinir ağları araştırmaları insan ve hayvan beyinleri içerisindeki ileri akılyürütmeyi sağlayan yapıları taklit etmeyi hedeflemiştir.¹⁴⁴

2.3.2 Bilgi Gösterimi

Bilgi gösterimi, nasıl formel olarak düşünüleceğini, yani hakkında konuşulabilecek bir “konu evreni”nin nesneleri hakkında nasıl çıkarım yapılacağını (formelleştirilmiş akılyürütme) ve konu evreninin fonksiyonları ile birlikte sembolik bir sistemle nasıl ifade edilebileceğini araştırır. Hem akılyürütme fonksiyonlarının konu evrenindeki sembollere uygulanma biçimine ait formel bir semantik sağlayan, hem de yorumlama teorisiyle beraber mantık cümlelerine anlam verecek niceleyiciler, kipsel operatörler öneren bir tür mantık kullanılır. **Bilgi gösterimi** ve **bilgi mühendisliği** (knowledge engineering) yapay zekâ araştırmalarının merkezinde yer almaktadır. Buradaki amaç insan bilgisinin bir örneğini makinede, yapay olarak oluşturmaktır.¹⁴⁵

Dünya hakkında gereksinim duyulan bilginin çoğunlukla elde hazır olmaması ya da modellenememesi yüzünden makineler çözmeye çalıştığı pek çok problemde başarısız olmuştur. Bunun nedeni, bağlamsal değişkenliğe ve rasyonel ajanların yüzleşebileceği durumların çokluğudur. Yapay zekânın modellenmesi gerekenler arasında şunlar vardır: nesneler, nesnelerin özellikleri ve nesneler arasındaki ilişkiler; durumlar, olaylar, şartlar ve zaman; nedenler ve etkileri; bilgi hakkındaki bilgi¹⁴⁶ vb. “Nelerin varolduğu” felsefeden alınan ontoloji¹⁴⁷ kavramıyla ifade edilir. “En genel

süreçlerle ilişkilidir. Bu görüşe göre sembolik olmayan yaklaşım, alt-sembolik süreçlerin işlediği durumlar için güçlü bir adaydır.

144Günther Görz, Bernhard Nebel: **a.g.e.**, s. 101.

145Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 195.

146İng. knowledge about knowledge. Diğer insanların ne bildiği hakkında bildiklerimiz.

147Bilgisayar bilimlerinde ve bilişimde ontoloji, bir alana ait kavramlar kümesinin ve bu kavramlar arasındaki ilişkilerin formel bir gösterimidir. O alanı tanımlamak ve özellikleri üzerine

olanlarına” da üst-ontoloji¹⁴⁸ denir.¹⁴⁹

Bilgi gösteriminin üç temel problemi şunlardır:

1. **Öndeğer akılyürütmesi ve nitelik problemi:** Birçok insanın bildikleri genellemelerden ibarettir. Örneğin, bir kuştan bahis açıldığında insanların aklına yumruk büyüklüğünde, öten ve uçan bir hayvan gelir. Fakat bunlardan hiçbiri bütün kuşlar için doğru değildir. John McCarthy bu problemi 1969'da **nitelik problemi** olarak adlandırmıştır. Yapay zekâ araştırmacılarının göstermek istedikleri herhangi bir sağduyu kuralı muazzam sayıda istisna ihtiva ediyordu. Yapay zekâ araştırmacıları bu probleme geliştirdikleri monoton-olmayan mantık ve monoton-olmayan mantığın bir türü olan öndeğer mantığı (default logic) ile çözüm üretmeye çalışmışlardır.¹⁵⁰
2. **Sağduyu bilgisinin genişliği:** Ortalama bir insanın bildiği tek tek olguların sayısı olağanüstü büyüktür. Sağduyu bilgisi için bir bilgi tabanı geliştirmeyi amaçlayan projeler büyük ölçüde ontolojik mühendisliğe ihtiyaç duyar. Amaç, bir bilgisayarın internette okuduğunu anlamasıdır. Fakat bunun için yeterli sayıda kavramın ontolojik bağının kurulmasına ihtiyaç vardır. Böylece bilgisayarın kendine ait bir ontolojisi olabilir. Bu henüz gerçekleştirilememiştir.¹⁵¹
3. **Sağduyu bilgisinin alt-sembolik biçimi:** Bildiklerimizin birçoğu “olgular” veya “durumlar” olarak temsil edilmez. Örneğin bir satranç ustası bazı hamlelerden açık vereceği için uzak duracaktır; ya da bir

akılyürütme için kullanılabilir. Teoride ontoloji, “ortak bir kavramlaştırmanın formel, açık bir belirtisidir”. Bir ontoloji bir alanı modelleyecek ortak bir sözcük dağarcığı sağlar. Yani, var olan nesnelerin ve/veya kavramların türleri, özellikleri ve ilişkilerini belirler. Ontolojiler yapay zekâ, semantik ağ, yazılım mühendisliği, biyomedikal bilişim, kütüphane bilimi ve bilgi mimarisi alanlarında dünya veya bir bölümü hakkında bir bilgi gösterim biçimi olarak kullanılır. **Bkz:** Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 261, 264.

148Bilişimde üst-ontoloji (en yüksek seviye ontoloji ya da temel ontoloji) tüm alanlarda aynı olan çok genel kavramları anlatan bir ontolojidir. Bir üst ontolojinin en önemli işlevi, erişilebilen çok sayıda ontoloji arasında geniş bir semantik işlerliği desteklemektir. Metaforun öne sürdüğü gibi o, belirli bir problem alanına dahil olmayan genel varlıkları tanımlamaya çalışan, varlıkların ve onlarla ilişkili kuralların (hem teorilerin hem düzenlemelerin) bir hiyerarşisidir. **Bkz:** Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 321, 326.

149Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 320-321.

150Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 360.

151Daniel Crevier: **a.g.e.**, s. 113-114.

sanat eleştirmeni bir heykele bakar bakmaz onun sahte olduğunu anlayabilir. Bu sezgi ve eğilimler, beynimizde bilincin ve sembolik gösterimlerin ulaşamayacağı bir yerde bulunmaktadır. Bu tarz bilgiler sembolik düşünüş için gerekli koşulları sağlar ve destek verirler. **Geleneksel yapay zekâ** veya **hesaplanabilir zekâ** (computational intelligence) çalışmalarının bu tip konulara çözüm getirmesi umulmaktadır.¹⁵²

2.3.3 Planlama

Planlama (planning), formel bir yolla şekillendirilmiş bir problemten, onu çözecek bir plan türetmekle uğraşır. Zeki ajanlardan yalnızca tepki vermelerini değil, aynı zamanda önleyici tedbirler almalarını da bekleriz. Ajanların, gerçekleştirdikleri eylemlerin ne etkide bulunacağını öngörmeleri ve amaçlanan hedeflere ulaşmalarını sağlayacak eylem zincirlerini kurabilmeleri gerekmektedir. Diğer bir deyişle eylemlerini planlayabilmeli ve imkânlı seçenekler arasında en fazla fayda sağlayacak olanı seçebilmelidir.¹⁵³

Planlamada çözülecek problemlerin mantık diliyle betimlenmesi gerekmektedir. Burada hem eylemler hem de hedefi yerine getirecek koşullar mantık formülleriyle betimlenmelidir. Planlamada bilgi gösterim teknikleri ve arama yoluyla problem çözüme teknikleri kullanılır.¹⁵⁴

Planlama problemlerinde, ajan kendini dünyada hareket eden tek şey olarak kabul eder ve eylemlerinin sonuçlarının ne olacağı önceden tanımlanmıştır. Ajanların periyodik olarak dünyanın kendi tahminlerine uyduğunu kontrol etmesi gerekir. Gerekirse de planlarını değiştirebilmelidir. Böylece ajan belirsizlik altında

152Paul Brna: “Symbolic and Subsymbolic Processing”, (Çevrimiçi), http://www.comp.lancs.ac.uk/computing/research/aai-aid/people/paulb/old243prolog/subsection3_1_2.html, 28 Ağustos 2009.

153Günther Görz, Bernhard Nebel: **a.g.e.**, s. 115-116.

154Günther Görz, Bernhard Nebel: **a.g.e.**, s. 116.

akılyürütme yapmak zorunda kalır. Bu da koşullu planlama ve olasılıksal planlamayı beraberinde getirmiştir.¹⁵⁵

2.3.4 Makine Öğrenmesi

Makine öğrenmesi yöntemleri, “deneyim” yoluyla öğrenen, yani olgular ve kurallar hakkında yeni bilgiler kazanabilen ya da öncelikleri belirleyip, karar veren sistemlerinin temelidir.¹⁵⁶ “Makine öğrenmesi” yapay zekâ araştırmalarının başlangıcından beridir odak noktasında yer almıştır. Marvin Minsky'nin iddiasına göre: yapay zekâ, ancak yapay zekâ sistemleri öğrenme becerisine sahip olabildiği zaman olanaklı olacaktır. Pek çok yapay zekâ araştırmacısı bu düşünceye katılmaktadır.¹⁵⁷

Rasyonel ajanlar sensörleri aracılığıyla çevresini gözlemler. Bu gözlemleri değerlendirerek, verili problemleri çözmek ve görevleri tamamlamak amacıyla kendi bilgi dağarcığına kaydeder. Makine öğrenmesinde ajanın mimarisi dört bileşen gerektirir: 1) ajanın performansını iyileştirmeden sorumlu **öğrenme bileşeni**, 2) sensörlerden gelen verilere göre dışsal eylemlerin seçimini belirleyen **performans bileşeni**, 3) ajanın eylemlerinde ne derecede başarılı olduğunu ölçen **kritik bileşeni**¹⁵⁸, 4) **öğrenme bileşeninden** öğrenme hedeflerini içeren verileri alan ve performans bileşenine, yeni “deneyimlere” yol açacak davranış ve eylemleri öneren **problem jeneratörü**. Problem jeneratörü olmadan da performans bileşeni her zaman, verili bilgi durumundaki en iyi şekilde değerlendirilmesi yapılmış eylemleri seçecektir. Fakat hiçbir zaman uzun vadede daha büyük başarılarla yol açması olası eylemleri seçmeyi düşünmeyecektir.¹⁵⁹

155Günther Görz, Bernhard Nebel: **a.g.e.**, s. 118.

156Günther Görz, Bernhard Nebel: **a.g.e.**, s. 50.

157Günther Görz, Bernhard Nebel: **a.g.e.**, s. 124.

158Bu değerlendirmeyi yapabilmek için, başarı ölçütlerinin önsel olarak verilmiş olması gerekmektedir.

159Günther Görz, Bernhard Nebel: **a.g.e.**, s. 124-125.

Makine öğrenmesi, sınıflandırma ve sayısal regresyon içermektedir. Sınıflandırma, birkaç kategoriden bir miktar şey gördükten sonra, bir şeyin hangi kategoriye ait olduğunu belirlemek için kullanılır. Regresyon bir dizi sayısal girdi/çıktı örneği alır ve girdilerden uygun çıktıları üretecek fonksiyonu keşfetmeye çalışır. Makine öğrenmesi algoritmalarının ve performanslarının matematiksel analizi **sayısal öğrenme teorisi** olarak bilinen teorik bilgisayar biliminin bir branşıdır. **Çıkarım yöntemleri**¹⁶⁰, **puslu mantık** ve **Bayes istatistiği** gibi araçlar teknolojinin ihtiyacı olan öğrenme algoritmalarını geliştirmek için sıklıkla kullanılır.

Makinesel öğrenmenin üç farklı türü vardır: 1) Eğer temelde bir ögenin girdi ve çıktıları gözlemlenebiliyorsa, buna **gözetim altında öğrenme** denir. Bu durumda sisteme, bir adet girdi ve bir adet çıktıyı içeren rastgele seçilmiş bir örnek verilir. Sistem bu örnek sayesinde “öğrenir” ve öğrenmiş olduğu bu yeni bilgiyi sonradan karşısına çıkacak yeni ve başka test verilerini sınıflandırmada temel alır. 2) Eğer koşul-eylem sürecini öğrenirken ajana eylemin değerlendirmesi verildiği halde o durumdaki olası en iyi eylemin ne olacağı tanıtılmazsa, bu ajan için bir ödül gerektiren öğrenme yöntemidir. Buna **ödüllü öğrenme** denir. 3) Doğru çıktıların ne olduğuna dair hiçbir ipucu verili olmadan öğrenme söz konusu olabilir. Bu, **gözetimsiz öğrenme** olarak adlandırılır. Gözetimsiz öğrenme en zayıf öğrenme çeşididir. Bu türlerin ortak özelliği, hepsinde matematiksel bir fonksiyonun temsiline öğrenilmesidir.¹⁶¹

2009 yılında, İsviçre'deki bir yapay zekâ laboratuvarında yapılan deney sırasında, yararlı bir kaynağı aramak ve zehirli olanından uzak durmak için birbirleriyle işbirliği yapmaya programlanmış robotlar zamanla yararlı kaynağı istifleme amacıyla birbirlerine yalan söylemeyi öğrenmişlerdir.¹⁶²

160Mantıksal açıdan doğru olan bir önermeden başka bir doğru önerme daha kazanmak için uygulanan yöntemler.

161Günther Görz, Bernhard Nebel: **a.g.e.**, s. 126-127.

162“Evolving Robots Learn To Lie To Each Other”, (Çevrimiçi), <http://www.popsci.com/scitech/article/2009-08/evolving-robots-learn-lie-hide-resources-each-other>, 29 Ağustos 2009.

2.3.5 Doğal Dil İşleme

Doğal dil¹⁶³ **işlemeye** yönelik yöntemler, dilin mekanizması (dilin yapısı, işlenişi ve kullanımını) hakkında bilgiler kazanma ve bunları insan-makine etkileşimi içinde kullanma üzerine yoğunlaşır.¹⁶⁴

Rasyonel ajanların insanların gündelik dille verdiği komutları algılamasını ve insanın alışık olduğu dilsel biçimde karşılık vermesini bekleriz. Bu yüzden doğal dil işleme yöntemleri yapay zekânın önemli bir problem alanını teşkil etmektedir. Doğal dil işleme çalışmaları, rasyonel ajanlarla insanlar arasında dilsel, olabildiğince günlük dile yakın, bir etkileşimi gerçekleştirecek yöntemler geliştirmeye çalışmaktadır. Bu çalışmaların dayanağını “dilsel ifadelerin formelleştirilerek işleme tabi tutulması” çabası oluşturmaktadır. Dilsel ifadelerin formelleştirilebilmesi için dilin yapısının bilinmesi, konuşmanın gerçekleştiği durum ve şartların anlama verdiği etkinin belirlenebilir olması gerekmektedir.¹⁶⁵

Dil işleme sistemlerinin alt bileşenleri şunlardır: 1) **Dili tanımak**: konuşmacının sesini algılayarak söylediğini anlayabilme yeteneği; 2) **Sözlük ve biçimbilim**: sözcüklerin veritabanı ve bunların çekimlerinin bilgisi; 3) **Dilbilgisel yapı çözümleme**: sözcüklerin cümledeki kullanım türlerinin belirlenerek dilbilgisel olarak analiz edilmesi; 4) **Anlambilimsel yorumlama**: sözcüklerin cümledeki anlamına uygun şekilde anlaşılması, yani çokanlamlılıkların ayıklanması süreci; 5) **Edimdilbilim ve söylem**: bir dinleyicinin, bir cümlesini duyduğu zaman onun sözcüksel anlamının ötesine geçerek, onu dil yoluyla ifade eden konuşucunun amacı ve konusal bağlantıları kavrayabilmesi; 6) **Dil üretimi**: verili bir hedef için doğal dilsel bir söylemin oluşturulmasıdır.¹⁶⁶

Dil işleme sistemlerinin gündelik hayatımızdaki uygulamalarına şu örnekler verilebilir: 1) girdi olarak sunulan dilsel ifadeleri veritabanlarında sorgulayıp, dilsel olarak çıktı üretmek; 2) doğal diller arasında çeviri yapmak; 3) göreve yönelik doğal

163İnsanların bildirişimde kullandıkları dillere, doğal dil denir. Türkçe, İngilizce ve diğer diller doğal dile örnektir.

164Günther Görz, Bernhard Nebel: **a.g.e.**, s. 50.

165Günther Görz, Bernhard Nebel: **a.g.e.**, s. 53-54.

166Günther Görz, Bernhard Nebel: **a.g.e.**, s. 60, 62, 63, 65, 69, 72.

dil üzerinden diyalog yürütebilmek, örneğin danışma hizmeti ya da sohbet etmek gibi; 4) verilen bir metni açıklamak, özetlemek ya da anlatmak; 5) doğal dilin dilbilgisine göre işlem yapmak, örneğin kelime işlem programlarının yazım hatalarını ve dilbilgisi hatalarını düzeltmesi gibi.¹⁶⁷ Doğal dil işleme uygulamaları daha etkileşimli kullanıcı arayüzleri geliştirilmesine olanak sağlar. Böylece insanlara doğal dilde verilen hizmetlerin, insanlara gerek kalmadan sunulabilmesi olanaklı hale getirilmeye çalışılır. Günümüzde ise pek çok araştırmacı okuduğunu anlayarak kendiliğinden bilgi toplayabilen doğal dil işleme sistemleri tasarlamak için uğraşmaktadır.

2.3.6 Örüntü Tanıma

Örüntü tanıma, sayısal makinelerin resimleri, sesleri ve diğer ortam içeriklerini algılama ve yorumlama yeteneği olarak tanımlanır. Örüntü tanıma sensörlerden (kamera, mikrofon, sonar gibi) aldığı girdileri kullanarak dünyanın görünümü hakkında çıkarımda bulunur. **Bilgisayar görmesi** görsel bir girdiyi analiz edebilir. Bazı alt-problemleri şunlardır: **konuşma tanıma**, **yüz tanıma** ve **nesne tanıma**.¹⁶⁸

2.3.7 Hareket ve Manipülasyon

Robotbilim alanı yapay zekâ ile yakın ilişkidir. Robotlar **nesne manipülasyonu**, **yer belirleme** (nerede olduğunu bilme), **eşleme** (etrafında ne olduğunu öğrenme) ve **hareket planlama** (gideceği yeri belirleyebilmek) gibi görevleri gerçekleştirebilmek için örüntü tanıma ve bu algılamalar sonucunda

¹⁶⁷Günther Görz, Bernhard Nebel: **a.g.e.**, s. 73.

¹⁶⁸Vasif V. Nابیev: **a.g.e.**, s. 525.

hareketin üretilmesinde yapay zekâ yöntemlerini kullanır.¹⁶⁹

2.4 Yaklaşımlar

Yapay zekâ araştırmalarına yön veren yerleşmiş bir birleştirici teori ya da paradigma yoktur. Araştırmacılar birçok konu hakkında farklı görüşlere sahiptir. Bu farklı görüşlerin temelinde şu sorular yatmaktadır: Yapay zekâ psikoloji ya da nörolojiyi inceleyerek doğal zekâyı mı taklit etmelidir? Yoksa kuş biyolojisinin havacılık mühendisliğiyle ne kadar ilgisi varsa insan biyolojisiyle yapay zekâ araştırmaları da o kadar mı ilgilidir? Akıllı davranış mantık ya da optimizasyon ile açıklanabilir mi? Yoksa başka alanlardan çok sayıda problemin çözülmesini mi gerektirir? Zekâ sembol (kelimeler ve fikirler) işleme midir? Yoksa zihin semboller dışında (alt-sembolik) başka şeyler de ihtiva eder mi?

Çalışmamızın ilk bölümünü oluşturan yapay zekâ tarihçesinin “*Temel Yapay Zekâ Yaklaşımları*” alt-bölümünde, “sibernetik yapay zekâ” ve “sayısal yapay zekâ” yaklaşımlarını incelemiştik. Burada bu iki yaklaşımla birlikte, yapay zekânın diğer yaklaşımları da aşağıda ele alınmaktadır.

2.4.1 Sibernetik ve Beyin Simülasyonu

Sibernetik, insan beynini taklit ederek yapay zekâ oluşturmaya çalışmaktır. Beynin hangi birimlerinin, ne ölçüde taklit edilmesi gerektiğine dair bir uzlaşma yoktur. 1940’larda ve 1950’lerde bir grup araştırmacı nöroloji, bilgi teorisi ve sibernetik arasındaki bağlantıyı keşfetmiştir. İlkel beyinleri taklit etmek için elektronik ağlar kullanan makineler geliştirmişlerdir. **W. Grey Walter**’in (1910-1977)

¹⁶⁹**Bkz:** Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 907-915

Turtle isimli robotları ve **Johns Hopkins Beast** isimli robot bunlara örnek gösterilebilir. 1960'a gelindiğinde bu yaklaşım sembolik yapay zekânın başarılı uygulamalarından sonra büyük oranda terk edilmiştir. 1980'lerde ortaya konan yeni yapay sinir ağları çalışmalarıyla sibernetik yeniden canlanmıştır.¹⁷⁰

2.4.2 Sembolik Yapay Zekâ Yaklaşımı

1950'lerin ortalarında dijital bilgisayarların geliştirilmesi ile birlikte yapay zekâ araştırmaları insan zekâsının sembol manipölasyonuna indirgenmesi olanağını keşfetmeye başlamıştır. Araştırmaların merkezini üç kurum oluşturuyordu: Carnegie Mellon Üniversitesi, Stanford Üniversitesi ve Massachusetts Institute of Technology. Her biri kendine özgü bir araştırma geliştirmiştir. Klasik yapay zekânın temelini oluşturan bu yaklaşım dört ana kolda ilerlemiştir:

1. **Bilişsel Simölasyon: Herbert Simon** ve **Alan Newell** insanın problem çözme yeteneklerini incelemiş ve onları formeletirmeyi denemişlerdir. Çalışmaları yapay zekâyla birlikte bilişsel bilim, yöneylem araştırması ve yönetim bilimi alanlarının temelini oluşturmaktadır. Carnegie Mellon Üniversitesi'nde merkezileşen bu gelenek 1980'lerin ortasında birleşmiş biliş teorilerini ortaya çıkarmak olan **SOAR** mimarisinin geliştirilmesiyle doruğa çıkacaktır.¹⁷¹
2. **Mantık Temelli Yaklaşım: Newell** and **Simon**'dan farklı olarak **John McCarthy**, makinelerin insan düşüncesini taklit etmelerinin gerekmediğini, bunun yerine insanların aynı algoritmaları kullanıp kullanmadığından bağımsız olarak soyut akılyürütme ve problem çözmenin özünü bulmaya çalışmaları gerektiğini düşünmüştür. Stanford Üniversitesi'nde başkanlığını yürüttüğü laboratuvarında çeşitli

¹⁷⁰Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 940.

¹⁷¹Pamela McCorduck: **a.g.e.**, s. 139, 245, 322.

problemleri çözmek için formel mantığı kullanmaya odaklanmıştır. Bu problemlerin başında **bilgi temsili, planlama ve öğrenme** gelmektedir. **Mantık**, Edinburg Üniversitesi ve Avrupa'nın diğer bölgelerindeki çalışmaların da odak noktasıdır. **LISP** ve **PROLOG** programlama dilleri bu yaklaşımın en önemli araçlarındandır.¹⁷²

3. **Mantık Karşıtı Yaklaşım:** MIT'teki araştırmacılar (Marvin Minsky ve Seymour Papert gibi) zor problemleri **makine görmesi** ve **doğal dil işleme** ile çözmenin geçici (ad-hoc) çözümler gerektirdiğini bulmuşlardır. Akıllı davranışın her yönünü kapsayacak basit ve genel bir ilkenin (mantık gibi) olmadığını savunmuşlardır. **Sağduyu bilgi tabanları**¹⁷³ bu yaklaşımın örneklerindedir.¹⁷⁴
4. **Bilgi Temelli Yaklaşımı:** 1970 civarında yüksek hafızalı bilgisayarların geliştirilmesiyle birlikte, belirli alanlar konusunda uzmanlaşmış sistemler geliştirilmeye başlanmıştır. Her sorunu çözecek genel amaçlı bir program yerine belirli bir uzmanlık alanındaki bilgiyle donatılmış programlar kullanma fikri, yapay zekâ alanında yeniden bir canlanmaya yol açmıştır. Bu “bilgi devrimi” uzman sistemlerin hızla gelişmesine ve yayılmasına yol açmıştır. Uzman sistemler, yapay zekâ programlarının gerçekten başarılı ilk şeklidir.¹⁷⁵

172Pamela McCorduck: **a.g.e.**, s. 100-101, 251.

173Yapay zekâ araştırmalarında sağduyu bilgisi sıradan bir insanın bilmesi beklenen gerçekler ve bilgiler yığındır. Sağduyu bilgisi problemi bilgi gösterimi alanında yürütülen sağduyu bilgisi için bir veritabanı oluşturma projesidir. Bu veritabanı çoğu insanın sahip olduğu tüm genel bilgileri saklar. Bu bilgiler doğal dili kullanan ve sıradan dünya hakkında çıkarımlar yapan yapay zekâ programlarının ulaşabileceği bir biçimde gösterilir. Böyle bir veritabanı en genelinin üst-ontolojiler olarak adlandırıldığı bir tür ontolojidir. Bkz: Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 363-364.

174Daniel Crevier: **a.g.e.**, s. 168.

175Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 22-23.

2.4.3 Alt-sembolik Yapay Zekâ Yaklaşımı

1960'lar sırasında sembolik yaklaşımlar küçük uygulama programlarında yüksek seviye düşünmenin taklit edilmesinde büyük başarı yakalamıştır. Siberetik veya sinir ağlarına dayanan yaklaşımlar terk edilmiş ya da arka plana itilmiştir. Ancak 1980'lere gelindiğinde, sembolik yapay zekâdaki ilerleme durmuş ve sembolik sistemlerin hiçbir zaman insan bilişinin tüm süreçlerini taklit edemeyeceği görüşü kabul edilir olmuştur. Özellikle de **robotlar**, **makine öğrenmesi** ve **örüntü tanıma** gibi süreçleri buna örnek gösterilebilir. Bir grup araştırmacı belli yapay zekâ problemleri için “alt-sembolik” yaklaşımlara kaymaya başlamıştır.¹⁷⁶ Bunlar iki ana kolda toplanmıştır:

1. **Nouvella Yapay Zekâ:** Rodney Brooks gibi robotikle uğraşan araştırmacılar sembolik yapay zekâyı reddetmişler, robotların hareket edebilmesini ve hayatta kalabilmesini sağlayacak temel mühendislik problemlerine odaklanmışlardır. Çalışmaları 1950'lerin ilk siberetik araştırmacılarının sembolik olmayan görüş açılarını yeniden canlandırmış ve yapay zekâda kontrol teorisi kullanımını yeniden gündeme getirmiştir.¹⁷⁷
2. **Hesaplama Zekâsı:** Sinir ağları ve bağlantıcılığa olan ilgi David Rumelhart ve diğerleri tarafından 1980'lerin ortalarında canlandırılmıştır. **Puslu mantık sistemler** ve **evrimsel hesaplama** gibi diğer alt-sembolik yaklaşımlar yeni ortaya çıkan hesaplama zekâsı disiplini içinde yer almaktadır.¹⁷⁸

176Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 25.

177Pamela McCorduck: **a.g.e.**, s. 456.

178Daniel Crevier: **a.g.e.**, s. 214-215.

2.4.4 İstatistiksel Yapay Zekâ Yaklaşımı

1990’larda yapay zekâ araştırmacıları belirli alt problemleri çözmek için gelişmiş matematiksel araçlar geliştirmişlerdir. Matematiksel araçların sonuçları hem ölçülebilir hem de doğrulanabilir özelliktedir. Bu yüzden de bilimsel karakteri diğer alanlara göre daha ileridedir. Bu araçlar yapay zekânın son zamanlardaki başarılarının birçoğunun arka planını oluşturmaktadır. Paylaşılan matematiksel dil, matematik, ekonomi ve yöneylem araştırması gibi diğer alanlarda yüksek seviyede işbirliğine izin vermektedir.¹⁷⁹

2.4.5 Yaklaşımların Entegrasyonu: Zeki Ajan Paradigması

Zeki ajanlar çevresini algılayarak başarı ihtimali maksimum olacak şekilde hareket eden sistemlerdir. Araştırmacılara izole problemler üzerinde çalışma ve tek bir yaklaşımda hemfikir olmadan, hem doğrulanabilir hem de kullanışlı çözümler bulmaları için gerekli koşulları sağlar. Belirli bir problemi çözen bir ajan, işe yarayan herhangi bir yaklaşımı kullanabilir. Bazı ajanlar sembolik ve mantıksaldır; bazıları alt-sembolik sinir ağlarıdır ve diğerleri de daha farklı yaklaşımlar kullanabilir. Bu yaklaşım araştırmacılara soyut ajan kavramlarını kullanan diğer alanlarla (karar teorisi ve ekonomi gibi) iletişim kurmak için ortak bir dil de sağlar. Zeki ajan paradigması 1990’larda büyük oranda kabul görmüştür.¹⁸⁰

Araştırmacılar bir birbirleriyle etkileşim içinde olan ve zeki ajanlardan oluşan “çok ajanlı akıllı sistemler” tasarlamışlardır. Hem sembolik hem de alt-sembolik öğeleri olan bu sistemler ile iki yapay zekâ yaklaşımı birleştirilmiştir.¹⁸¹

179Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 25-26.

180Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 27.

181Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 932.

2.5 Araçlar

Elli yılı aşan araştırma sürecinde yapay zekâ çalışmaları, bilgisayar bilimlerindeki en zor problemleri çözecek çok sayıda araç geliştirebilmiştir. Araçlardan en çok kullanılanları aşağıda ele alınmaktadır.

2.5.1 Arama ve Optimizasyon

Yapay zekâ alanındaki birçok problem, problemlerin sayısız olası çözüm yolları arasından arama yapılarak çözülmeye çalışılır.¹⁸² Böylece akılyürütme arama yapmaya indirgenmiştir. Aramanın her adımı bir çıkarım kuralının uygulanmasıdır. Planlama algoritmaları hedef ve alt-hedef ağaçlarında arama yaparak, asıl hedefe giden bir yol bulmaya çalışır.¹⁸³ Bu sürece, araç-amaç analizi adı verilir.¹⁸⁴ Uzuvarları hareket ettirmek ve nesneleri tutmak için yazılan robot algoritmaları lokal aramalardan faydalanır. Yine birçok öğrenme algoritması optimizasyona dayanan arama algoritmalarını kullanır.

Basit, kapsamlı aramalar gerçek dünyadaki çoğu problem için nadiren yeterli olur: Arama uzayı (aranacak yerlerin sayısı) çabucak muazzam rakamlara çıkar. Ortaya çıkan sonuç çok yavaş ya da hiçbir zaman bitmeyen bir aramadır. Birçok problem için çözüm **yol gösterici yöntemler** kullanmaktır. Bunlar hedefe ulaşacak gibi görünmeyen seçenekleri eler. Bu işlem “**arama ağacını budamak**” olarak adlandırılır. **Yol gösterici yöntemler** programa çözüme giden yol için “en iyi tahmini” verir.¹⁸⁵

Çok farklı bir arama türü, 1990’larda öne çıkmıştır. Bu yöntem matematiksel optimizasyon teorisine dayanır. Burada problem için çözüm aramaya bir çeşit

¹⁸²Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 59.

¹⁸³Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 280.

¹⁸⁴Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 382.

¹⁸⁵Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 94-95.

tahminle başlanmaktadır. Ardından da yapılan tahmini, aşama aşama sürdürmek ve artık bir iyileştirme yapılamayana kadar iyileştirmek mümkündür.¹⁸⁶

2.5.2 Mantık

Mantık, yapay zekâ araştırmalarına **John McCarthy** tarafından 1958’de yazdığı **Advice Taker** isimli kurumsal programıyla dahil edilmiştir. Mantık, özellikle **bilgi temsili** ve **problem çözme** için kullanılır. Ancak diğer problemlere de uygulanabilir. Örneğin, **satplan algoritması mantığı**¹⁸⁷ planlama için kullanır ve tümevarımsal mantık programlamada bir öğrenme yöntemidir.¹⁸⁸

Yapay zekâ araştırmalarında kullanılan birkaç farklı mantık türü vardır. **Önermeler mantığı** ve **eklemler mantığı** doğru ya da yanlış olabilen yargıların mantığıdır.¹⁸⁹ **Niceleme mantığı** da ifadelerde niceleyici ve yüklemelerin kullanımına izin verir. Niceleyiciler, nesneler, nesnelerin özellikleri ve nesnelerin birbirleriyle ilişkileri hakkında yargılar ifade eder.¹⁹⁰ Önermeler mantığının bir çeşidi olan **puslu mantık** bir yargının doğruluk değerinin, sadece Doğru (1) ya da Yanlış (0) olarak değil, 0 ile 1 arasında bir değer olarak gösterilmesine izin verir. Puslu sistemler kesin olmayan akılyürütmelerde kullanılabilir. Bunlar modern endüstriyel ve tüketim ürünleri kontrol sistemlerinde geniş kullanım alanı bulmuştur.¹⁹¹ **Monoton-olmayan mantığın en çok kullanılan iki tipi olan öndeğer mantığı ve sınırlama mantığı**¹⁹²

186Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 110.

187Satplan bir otomatik planlama yöntemidir. Bir planlama problemini Boolean sağlanabilirlik problemine dönüştürür. Sonra problem DPLL algoritması ya da WalkSAT gibi sağlanabilirliği ispatlayan yöntemler kullanılarak çözülür.

188Daniel Crevier: **a.g.e.**, s. 113-114.

189Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 204.

190Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 240.

191Şafak Ural: “Puslu (Fuzzy) Mantık”, **Mantık, Matematik ve Felsefe**, I. Ulusal Sempozyumu 26-28 Eylül 2003 Assos-Çanakkale, Ed. Şafak Ural, Mehmet Özer, Ayten Koç, Arzu Şen, Gürsel Hacıbekiroğlu, İstanbul, İstanbul Kültür Üniversitesi Yayınları, İstanbul, 2004, s. 45-46.

192Sınırlama mantığı, şeylerin aksi belirtilmedikçe beklendiği gibi olduğu şeklindeki sağduyu varsayımını formelleştirmek için John McCarthy tarafından ortaya atılmış monoton-olmayan bir mantıktır. McCarthy daha sonra sınırlama mantığını çerçeve problemini çözmek için kullanmıştır. Orijinal önermeler mantığı formülasyonunda sınırlama mantığı, bazı yüklemelerin eklentilerinin en

(circumscription), mantığın öndeğer akılyürütmesi, çerçeve ve niteleme problemine yardımcı olmak için tasarlanan bir biçimdir.¹⁹³ Mantığın birkaç uzantısı belirli bilgi alanlarında kullanılmak üzere tasarlanmıştır: **tanım mantığı** kategorileri ve aralarındaki ilişkileri gösterebilmeyi; **durum analizi**, **olay analizi** ve **akıcı analiz** olayları ve zamanı temsil edebilmeyi; **nedensel analiz**; **inanç analizi** ve **kipler mantığı** bilgi hakkındaki bilgiyi¹⁹⁴ gösterebilmeyi amaçlamaktadır.

2.5.3 Olasılığa Bağlı Yöntemler

Yapay zekâda birçok problem (akılyürütme, planlama, öğrenme, algı ve robotbilimde) rasyonel ajanın eksik ya da kesin olmayan bilgiyle işlem yapmasını gerektirir. 1980'lerin sonlarında 1990'ların başlarında başlayarak **Judea Pearl** (1936) ve diğerleri, bu problemleri çözecek birtakım güçlü araçlar tasarlamak için olasılık teorisi ve ekonomiden alınan yöntemlerin kullanımını savunmuştur.¹⁹⁵

Bayes ağları, akılyürütme (Bayes çıkarım algoritmasını kullanarak), planlama (karar ağlarını kullanarak) ve örüntü tanıma (dinamik Bayes ağlarını kullanarak) gibi değişik tipteki problemler için kullanılabilecek genel bir araçtır.¹⁹⁶ **Olasılık ifade eden algoritmalar**, süzgeçleme, öngörme, düzleme ve veri akıntıları için açıklamalar bulmada da kullanılabilir. Örüntü tanıma sistemlerinin zamanla ortaya çıkan süreçleri analiz etmesine yardımcı olur.¹⁹⁷

Ekonomi biliminden alınan kilit bir kavram **yararlılıktır**: Bir şeyin bir zeki ajan için ne kadar değerli olduğunun ölçüsüdür. **Karar teorisi**, **karar analizi**¹⁹⁸ ve

aza indirir. Bir yüklemın eklentisi o yüklemın doğru olarak uygulanabildiği çok öğeli değerler kümesidir. Bu minimizasyon, 'bir şeyin doğru olduğu bilinmiyorsa o şey yanlıştır' kabulü bakımından kapalı-dünya varsayımına benzer.

193Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 354, 358.

194Ing. Knowledge about knowledge. Diğer insanların ne bildiği hakkında bildiklerimiz.

195Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 487-488.

196Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 492.

197Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 537.

198Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 584.

bilgi değeri teorisi¹⁹⁹ kullanılarak bir ajanın nasıl karar verdiğini ve plan yaptığını analiz eden kesin matematiksel araçlar geliştirilmiştir. Bu araçlar arasında **Markov karar süreçleri**²⁰⁰, **dinamik karar ağları**, **oyun teorisi**²⁰¹ sayılabilir.

2.5.4 Sınıflandırıcılar ve İstatistiksel Öğrenme Yöntemleri

En basit yapay zekâ uygulamaları iki türe ayrılabilir: **sınıflandırıcılar** (“parlaksa elmadır”) ve **kontrolcüler** (“parlaksa topla”). **Kontrolcüler** hareketler çıkarsamadan önce koşulları da sınıflandırır. Dolayısıyla **sınıflandırma** birçok yapay zekâ sisteminin merkezi bir parçasını oluşturur. **Sınıflandırıcılar** en yakın eşleşmeyi belirlemek için **örüntü eşleştirmesini** (pattern matching) kullanırlar. “Örüntü eşleştirme” örneklerle göre ayarlanabilir. Böylece onların yapay zekâda kullanılması oldukça cazip hale gelir. Bu örnekler **gözlemler** ve **desenler** olarak bilinir. Makine öğrenmesinde her örüntü belirli bir öntanımlı sınıfa dahildir. Bir sınıf, verilmesi gereken bir karar olarak görülebilir. Sınıf etiketleriyle birleştirilmiş bütün gözlemler bir veri kümesi olarak bilinir. Yeni bir gözlem yapıldığında bu gözlem geçmiş deneyime dayanarak sınıflandırılır.²⁰²

Bir sınıflandırıcı çeşitli yollarla eğitilebilir. Birçok istatistiksel ve makine öğrenme yaklaşımı vardır. En yaygın sınıflandırıcılar **sinirsel ağ**²⁰³, **destek vektörü makinesi** gibi **kernel yöntemleri**²⁰⁴, **k-en yakın komşu algoritması**²⁰⁵, **Gauss karışımı modeli**²⁰⁶, **sade Bayes sınıflandırıcı**²⁰⁷ ve **karar ağacıdır**²⁰⁸. Bu sınıflandırıcıların performansı çok çeşitli görevlerde karşılaştırılmıştır. Sınıflandırıcı

199Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 600.

200Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 613.

201Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 631.

202Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 712, 755.

203Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 736.

204Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 749.

205Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 733.

206Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 725.

207Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 718.

208Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 631.

performansı önemli ölçüde sınıflandırılacak verinin özelliklerine bağlıdır. Verilen bütün problemleri en iyi şekilde işleyen tek bir sınıflandırıcı yoktur.

2.5.5 Yapay Sinir Ağları

Yapay sinir ağları insan beynindeki dev sinir hücreleri ağına benzeyen birbirine bağlı düğümler grubudur. Yapay sinirsel ağların çalışması, yapay zekâ araştırma alanı oluşturulmadan önceki on yılda **Walter Pitts** ve **Warren McCullough**'un çalışmalarıyla başlamıştır. İlk zamanlardaki diğer önemli araştırmacılar, algılayıcıyı icat eden **Frank Rosenblatt** ve geri yayılım algoritmasını geliştiren **Paul Werbos** olmuştur.²⁰⁹

Yapay sinir ağları iki ana kategoride toplanır: **ileri besleme sinirsel ağları** (sinyalin sadece tek yönde geçtiği) ve **ynelenen sinirsel ağlar** (geri beslemeye izin veren). En popüler **ileri besleme ağları** arasında **algılayıcılar**, **çok katmanlı algılayıcılar** ve **radyal bazlı ağlar** vardır. **Yinelenen sinirsel ağlar** arasında ise en ünlüsü **Hopfield ağıdır**. Bu bir tür atraktör ağıdır ve ilk olarak 1982'de **John Hopfield** tarafından tanımlanmıştır. Sinirsel ağlar **Hebbian öğrenmesi** ve **rekabetçi öğrenme** gibi teknikler kullanılarak akıllı kontrol (robotbilim için) ya da öğrenme problemine uygulanabilir.²¹⁰

2.5.6 Programlama Dilleri

Programlama dilleri yapay zekâda öğrenme, akılyürütme ve anlama gibi zekâ süreçlerini taklit etmek için kullanılan bir araçtır. Bilgisayar dillerinin tasarlandığı ilk yıllarda, bilgisayarlar sayılarla hesaplama yapmak için kullanılmaktaydı. Kısa bir

²⁰⁹Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 631.

²¹⁰Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 758.

süre sonra da iki sayı sistemiyle yalnız sayıların değil, aynı zamanda görelî nesnelerin de temsil edilebileceği anlaşıldı. Bunun üzerine sembol oluşturmak, ilişkilendirmek ve manipüle etmek için kuralları sembol olarak kullanılabilir. Böylece sembolik hesaplama kavramı herhangi türden malumatları işleyebilen algoritmaları tanımlamak anlamında kullanılmaya başlanmıştır. Bu yüzden de insan zekâsını taklit etmek amacıyla kullanılabilir. Çok geçmeden sembollerle de programlamanın, sayı işleme için tasarlanmış programlama dillerine göre daha üst seviyede bir soyutlamaya ihtiyaç duyduğu anlaşılmıştır.²¹¹

Belirli bir yapay zekâ probleminin programlaması başlangıcı sırasında problem çok kesin bir biçimde belirlenememektedir. Sadece etkileşimli ve giderek artan iyileştirme ile daha kesin bir belirleme mümkün olabilmektedir. Bunun bir nedeni de tipik yapay zekâ problemlerinin çok fazla alana özgü olmasıdır. Dolayısıyla yol gösterici stratejilerin deneysel olarak oluştur-ve-test et yaklaşımıyla geliştirilmesi gerekir. Bu yüzden yapay zekâ programlaması, genelde programlamanın ayrıntılı bir formel tanımıyla başlayan standart yazılım mühendisliği yaklaşımlarından önemli oranda ayrılır. Yapay zekâ programlamada uygulama çabası aslında problemi tanımlama sürecinin bir parçasıdır.²¹²

Bildirimsel (declarative) programlama stili, yüksek seviye veri yapıları (örneğin, listeler ya da ağaçlar) ve işlemler (örneğin, desen eşleştirme) kullanması sayesinde daha elverişli olur, böylelikle sembolik hesaplama Fortran, Pascal ve C gibi standart buyurucu dillerden²¹³ çok daha soyut bir düzeyde desteklenmiş olur. Yapay zekâ programlarının derlenmesi buyurucu dillerde olduğu kadar etkin yapılamaz. Buna rağmen, belli bir yapay zekâ problemi (kısmen de olsa) anlaşıldığında, buyurucu bir dil kullanarak yeniden uygulanmasına temel oluşturmak için ayrıntılı tanımlamalar şeklinde yeniden formüle edilmesi mümkündür.²¹⁴

Sembolik hesaplama ve yapay zekânın gereksinimlerinden buyurucu stile alternatif olarak iki yeni temel programlama paradigması ortaya çıkmıştır:

211Günter Neumann: "Programming Languages in Artificial Intelligence", s. 3, (Çevrimiçi), <http://www.dfki.de/~neumann/publications/new-ps/ai-pgr.pdf>, 29 Ağustos 2009.

212Günter Neumann: **a.g.e.**, s. 3.

213Yerine koyma ile değişkenlerin durumunu değiştirmek için tasarlanmış programlama dili.

214Günter Neumann: **a.g.e.**, s. 4.

fonksiyonel ve mantık programlama stilleri. Her ikisi de matematiksel formelleştirmelere, yani tekrarlı fonksiyon teorisi ve formel mantığa dayanır. Hâlâ en yaygın olarak kullanılan ilk pratik yapay zekâ programlama dili **John McCarthy** tarafından 1950'lerin sonlarında geliştirilen LISP'tir. LISP matematiksel fonksiyon teorisi ve lambda soyutlamasına dayanır. Birkaç önemli ve etkili yapay zekâ uygulamaları LISP dilinde yazılmıştır.²¹⁵

İlk mantık programlama dili PROLOG'dur. PROLOG'daki problemler olgular, aksiyomlar ve yeni olgular çıkarmaya yarayan mantık kuralları olarak ifade edilir. PROLOG matematiksel olarak yüklem değişkenler hesabı ve 1960'ların sonlarında otomatik teori ispatı alanında elde edilen teorik sonuçlar üzerine kurulmuştur. PROLOG, programların ilişkiler açısından ifade edildiği ve bu ilişkiler üzerinde sorgular çalıştırılarak iş gören bildirimsel bir dildir. PROLOG özellikle sembolik akılyürütme, veri ve dil çözümleme uygulamalarında kullanışlıdır. Günümüzde yapay zekâ alanında PROLOG hâlâ kullanılmaktadır.

STRIPS dili otomatik planlama problemlerini ifade etmek için kullanılmıştır. Her hareket için önkoşullar (örneğin, hareket yapılmadan önce ne gerekiyor) ve artkoşullar (örneğin, hareket yapıldıktan sonra ne oluyor) belirlenir.

Planner, prosedürel ve mantık dillerinin bir karışımıdır. Söylenmek istenenin örüntü-yönelimli çıkarımlarla yorumlandığı mantık cümleleri için süreç olarak yorumlarını verir.²¹⁶

2.6 Uygulamalar

Yapay zekâ tıbbi tanı, borsa, robot kontrolü, hukuk, bilimsel keşif gibi çok farklı alanlarda kullanılmıştır. Birçok üstün yapay zekâ ürünü genel uygulamalara yayılmıştır. Binlerce yapay zekâ uygulaması her türlü endüstrinin altyapısının derinlerine nüfuz etmiştir. 1990'ların sonu itibariyle yapay zekâ teknolojisi büyük

²¹⁵Günter Neumann: **a.g.e.**, s. 4.

²¹⁶Daniel Crevier: **a.g.e.**, s. 190-196.

sistemlerin ögeleri olarak kullanılmaya başlamıştır. Fakat bu başarılarda yapay zekâ olarak pek anılmaz. Aşağıda değişik endüstri alanlarındaki yapay zekâ uygulamaları kısaca belirtilmiştir:²¹⁷

Bilgisayar Bilimi: Yapay zekâ araştırmacıları bilgisayar biliminde karşılaşılan en zor problemleri çözecek birçok araç geliştirmiştir. Buluşlarının çoğu bilgisayar bilimi tarafından benimsenmiştir ve artık yapay zekânın bir parçası olarak görülmezler. Yapay zekâ laboratuvarlarında geliştirilmiş bazı teknolojiler şunlardır: zaman paylaşımı, interaktif yorumlayıcılar, grafik tabanlı kullanıcı arabirimleri ve bilgisayar faresi, bağlantılı liste veri türü, otomatik depo yönetimi, sembolik programlama, işlevsel programlama, dinamik programlama, nesne-yönelimli programlama.

Finans: Bankalar yapay zekâ sistemlerini işlemleri düzenlemek, hisse senetlerine yatırım yapmak ve varlık (mal-mülk) yönetimi için kullanırlar. Ağustos 2001’de bir ticaret yarışı simülasyonunda robotlar insanları yenmiştir. Mali kuruluşlar insanlar için yatırım aracı olarak kullanmak üzere norm dışı ücretleri ya da alacakları belirlemek için uzun süredir yapay sinir ağı sistemlerini kullanmaktadır.

Tıp: Bir tıbbi klinik yapay zekâ sistemlerini yatak programlarını düzenlemek, personel vardiyası oluşturmak ve tıbbi bilgi sağlamak için kullanabilir. Yapay sinir ağları tıbbi tanı alanında bilgisayarlı ayırıcı tanı işlevi görürler.

Ağır Sanayi: Robotlar birçok sanayide yaygınlaşmıştır. Robotlara genellikle insanlar için tehlikeli olan görevler verilir. Robotlar, konsantrasyon kaybının hatalara ya da kazalara neden olduğu tekrarlı işlerde ve insanların aşağılayıcı bulduğu işlerde verimlidir. General Motors boya, kaynak ve montaj gibi görevlerde yaklaşık 16.000 robot kullanır. Japonya dünyada robot kullanımı ve üretiminde liderdir. 1995’te dünya çapında 700.000 robot kullanımdadır, bunların 500.000’i Japonya’dadır.

Ulaşım: Otomobillerdeki otomatik vites kutularında kullanılmak üzere puslu mantık denetim birimleri geliştirilmiştir.

Telekomünikasyon: Birçok telekomünikasyon şirketi işgücü yönetiminde buluşsal aramadan faydalanır. Örneğin British Telecom 20.000 mühendisin çalışma

²¹⁷Stuart J. Russell, Peter Norvig: **a.g.e.**, s. 27-28.

çizelgelerini düzenleyen bir programda buluşsal arama yöntemini kullanmıştır.

Müzik: Müziğin evrimi her zaman teknolojiden etkilenmiştir. Bilim adamları yapay zekâ kullanarak bir bilgisayarın yetenekli bir müzisyenin hareketlerini taklit etmesine çalışmaktadır. Beste, performans, müzik teorisi ve ses işleme müzik ve yapay zekâ alanındaki araştırmaların odak noktalarından bazılarıdır.

Havacılık: Hava Harekat Birimi (AOD) uzman sistemlerden faydalanır. AOD yapay zekâyı savaş ve eğitim simülatörlerinde, görev yönetimine yardımda, taktiksel kararlar için destek sistemlerinde ve simülatör verilerinin sembolik özetlerini çıkarmada kullanır. Simülatörlerde yapay zekâ kullanımının AOD için çok faydalı olduğu görülmektedir. Uçak simülatörleri yapay uçuşlardan elde edilen verileri işlemek için yapay zekâ kullanmaktadır. Uçuş simülasyonundan başka hava hareketi simülasyonları da vardır. Bilgisayarlar bu tür durumlarda en başarılı senaryoları bulabilmektedirler. Ayrıca bilgisayarlar dost kuvvetlerle düşman kuvvetlerin konumlarına, hızlarına ve güçlerine dayanarak stratejiler oluşturabilmektedir. Pilotlar savaş esnasında bilgisayarlardan yardım alabilmektedir. Yapay zekâ programları verileri düzene sokup pilota mümkün olan en iyi manevraları gösterebilmektedir.

Yapay zekâ yöntemlerinin kullanıldığı tipik problemler²¹⁸:

- Desen tanıma
 - Optik karakter tanıma
 - El yazısı tanıma
 - Konuşma tanıma
 - Yüz tanıma
- Yapay yaratıcılık
- Bilgisayar görüntüsü, sanal gerçeklik ve resim işleme
- Tanı (yapay zekâ)
- Oyun teorisi ve stratejik planlama
- Oyun yapay zekâsı ve bilgisayar oyunu robotu
- Doğal dil işleme, çeviri ve konuşan robot
- Lineer-olmayan kontrol ve robot teknolojisi

Yapay zekâ yöntemlerinin uygulandığı diğer alanlar²¹⁹:

218Bkz: "Applications", **AITopics**, (Çevrimiçi),
<http://www.aaai.org/AITopics/pmwiki/pmwiki.php/AITopics/Applications>, 14 Eylül 2009.

219Bkz: "Applications", **AITopics**, (Çevrimiçi),

- Yapay yaşam
- Otomatik akılyürütme
- Otomasyon
- Biyolojiden esinlenen hesaplama
- Kavram madenciliği
- Veri madenciliği
- Bilgi gösterimi
- Semantik ağ
- İstenmeyen e-posta filtreleme
- Robotbilimi
- Melez akıllı sistem
- Akıllı ajan
- Akıllı kontrol
- Dava takibi

Yapay zekânın olanaklarının geliştirilebilmesi için insanın akılyürütme süreçlerinin olabildiğince açılanması gerekmektedir. Her ne kadar yapay insan yaratma fikrinin rüzgârı ile bu alana gelinse de, sembolik yapay zekâ yaklaşımının yapay zekâyı bir problem çözme alanı görerek, akılyürütme süreçlerini modellemeyi ana hedefi olarak seçmesinin ne kadar isabetli olduğu uygulamalardaki başarılarından anlaşılmaktadır. Akılyürütme süreçlerinin modellenmesinde ve programlama dillerinde mantık en önemli araçlardan biridir.

3. Mantık ve Yapay Zekâ

Yapay zekâ çalışmalarında mantığın bir araç olarak kullanımını ilk öneren **John McCarthy** olmuştur.²²⁰ Yapay zekânın isim babası olan **McCarthy**, çalışmaları boyunca mantık temelli yapay zekâ yaklaşımını savunarak, bu alanda çalışmalar yapılmasına önayak olmuştur. Mantık temelli yapay zekâ yaklaşımının temelinde karşılaşılan problemlerin formelleştirilmesi fikri yer almaktadır. Eğer bu problemler mantık ifadelerine dönüştürülebilirse, mantık kurallarına göre kolayca çözülebileceği varsayılmıştır.²²¹

Yapay zekâ, mantığı araç olarak kullanma fikrini felsefeden ödünç almıştır. **Aristoteles**'ten (M.Ö. 384-322) beridir “doğru düşünme” yani akılyürütme süreçleri mantık çatısı altında sistemleştirilmeye çalışılır. **Aristoteles**'in ortaya koyduğu mantık anlayışı 19. yüzyılın başına kadar ana çizgisini koruyarak sürmüştür. Mantığın matematiksel gelişimi **George Boole**'un (1815-1864) çalışmalarıyla başlamıştır. **Gottlob Frege** (1848-1925) Boole'un mantığını nesneleri ve ilişkileri de içerecek şekilde genişletmiş, bugün en temel bilgi gösterim sistemi olarak kullanılan önermeler mantığını oluşturmuştur. **Frege**'nin çalışmalarıyla temellendirilen modern mantık çalışmalarıyla birlikte mantık matematiksel bir temele oturtulmuş ve sembolleştirme olanakları genişletilmiştir. Böylece sembollerden ve bu sembollerin kullanım kurallarından oluşan bir formelleştirme dili inşa edilmiştir. **Alan Turing** ise soyut mantık ve fiziksel mekanizmalar arasındaki bağlantının öncüsü olmuştur.²²²

Diğer yandan matematiği mantık ile temellendirme üzerine de önemli

220**Programs with Common Sense** isimli makalesinde ilk hipotetik bilgisayar programı olan **Advice Taker**'da bilgi gösterimi için mantık kullanımını önererek sağduyu akılyürütmesini bir yapay zekâ yöntemi olarak kullanmıştır.

221Richmond Thomason: “Logic and Artificial Intelligence”, **Stanford Encyclopedia of Philosophy**, (Çevrimiçi), <http://plato.stanford.edu/entries/logic-ai/>, 18 Eylül 2009.

222G. Aldo Antonelli: “Logic”, **The Blackwell Guide to the Philosophy of Computing and Information**, Ed. Luciano Floridi, Cornwall, Blackwell Publishing, 2004, s. 263.

çalışmalar yapılmıştır. Mantık alanına böylece felsefe kökenli araştırmacılar ile birlikte matematik kökenli araştırmacılar da katılmıştır. Araştırmacıların aldıkları profesyonel eğitim mantık yaklaşımlarını, dolayısıyla mantık kullanımı ve amaçlarını da etkilemiştir. Matematik kökenli mantıkçılar hem mantığın matematik yönüyle incelenmesiyle araştırılması hem de formel mantığın matematiğin diğer alanlarına uygulanmasıyla uğraşmışlardır. Yeni ve farklı mantık sorunları üzerine düşünmeyi gerektiren ve matematiksel olmayan bilimleri formalize etmeye yönelik uğraşlar, matematik kökenli mantıkçıların gündeminde olmamıştır. Felsefe kökenli mantıkçıların ana hedefi ise mantık yöntemlerinin matematiksel olmayan akılyürütme alanlarına uygulanmasıdır. Buradaki gaye dünyanın mantıksal olarak ifade edilmesinin olanaklarını genişletmektir.²²³

Matematik kökenli mantık ile felsefe kökenli mantık arasındaki bu görüş farklılığı güncel bilimsel çalışmaların buluşma noktası olan dergilerden de kolayca takip edilebilir. Örneğin **Journal of Symbolic Logic** (Sembolik Mantık Dergisi) isimli dergi 1936'da yayımlanan ilk sayısında matematikçilerle felsefecileri bir araya getirmeyi amaçladığını duyurmuştur. Böylece kendi disiplinleri içerisinde marjinal olarak değerlendirilen çalışmalar ayrı bir çatı altında toplanmış olacaktır. 1960'lı yıllara kadar dergiye iki disiplinden de aynı oranda katkı sağlanmıştır. 1960'lara gelindiğinde ise derginin sayılarında artık felsefe kökenli çalışmalar tek tük yer bulabilmiştir. Bunun temelinde ise yukarıda değindiğimiz iki grubun hedeflerindeki önemli derecedeki farklılık yatmaktadır.²²⁴

Matematikçiler giderek daha teknik ve karmaşık bir yöntemler ve teoriler yığını geliştirmeyi amaçlıyordu. Oysa birçok felsefeci bu amacın aydınlatıcı felsefi konulardan giderek uzaklaştığını düşünmekteydi. Bu ayrılıklar 1972'de **Journal of Philosophical Logic** (Felsefi Mantık Dergisi) adıyla yeni bir derginin kurulmasına neden olmuştur.²²⁵

Yapay zekânın temel alanlarından biri olan bilgi gösteriminin alt-konuları

223Richmond Thomason: **a.y.**

224Richmond Thomason: **a.y.**

225Richmond Thomason: **a.y.**

incelendiğinde, bu alandaki çalışmaların matematik kökenli mantık çalışmalarından ziyade felsefe kökenli mantık çalışmalarının konularıyla örtüştüğü gözlemlenebilir. Bunun temelinde ise yapay zekânın makro-dünya problemlerini çözmek konusundaki hedefi, yani matematiksel olmayan akılyürütme alanlarında da çözüm getirebilme isteği yer almaktadır.²²⁶

Günlük yaşamda kullanılan akılyürütme modellenerek geniş bir alandaki yapay zekâ problemlerine çözüm getirilmek istenmiştir. Bu hedefin arkaplanını, felsefe kökenli mantığın, mantık temelli yapay zekâyâ yaptığı etki vardır. Yapay zekâdaki mantıkçı geleneği temsil eden araştırmacıların felsefe kökenli mantık literatürünü okuduğu ve ondan etkilendiği aşikârdır. Örneğin, 1969 yılında **McCarthy** ve **Hayes**'in ortaklaşa kaleme aldığı **Some Philosophical Problems from the Standpoint of Artificial Intelligence** isimli makalede 58 atıf vardır ve bu atıfların 35'i felsefe kökenli mantık literatürüne yapılmıştır.²²⁷

Felsefi Mantık Dergisi'ndeki teorik çalışmalar ile örtüşen bu çalışmalar, uygulamada ise aynı paralelliği sağlayamamaktadır. Bunun temelinde felsefenin teorik yapısıyla birlikte, mantık teorilerinin küçük ölçekli ve yapay örneklerle uygulanması yer almaktadır.

Yapay zekânın ilk yıllarından sonra felsefe kökenli mantık ile mantık temelli yapay zekâ yaklaşımı birbirinden uzaklaşmıştır. Yapay zekânın mantık temeli üzerinde incelenmesi, felsefe kökenli mantıkta pek bilinmeyen yeni mantık teorilerinin (en ünlü örnek 'monoton-olmayan mantık'tır) geliştirilmesine sebep olmuştur. Bununla birlikte yapay zekâ araştırmacıları daha performanslı uygulamalar geliştirebilmek için algoritmaların teorik analizine büyük önem vermişlerdir. Bunun sebebini ise eşi görülmemiş büyüklükteki mantıksal aksiyom öbeklerini kullanan iddialı uygulamalarla gittikçe gelişen bilgisayar bilimi oluşturur. Bu uygulamaların büyüklükleri yeni problemler ve yeni metodolojilerin ortaya çıkmasına neden olmuştur.²²⁸

Uygulamalar hakkında düşünmek, araştırmanın nasıl yürütüleceği ve

226Richmond Thomason: **a.y.**

227Richmond Thomason: **a.y.**

228Richmond Thomason: **a.y.**

sunulacağı üzerinde çok etkili olabilir. Felsefe kökenli mantık geleneği otomatikleştirilmiş akılyürütme uygulamalarından eskidir ve bugüne kadar bu tür uygulamalara karşı ilgisiz kalmıştır. Genelde felsefe literatürü akılyürütmenin uygulanabilirliği, verimliliğiyle ya da akılyürütme sürecinin herhangi bir özelliği ile ilgilenmez. Felsefe teorilerinin gerçekçi, büyük çapta akılyürütme problemleriyle resmedildiği ya da test edildiği durumlar bulmak zordur. Fakat temel teorik konuların (kipler, şartlı ve geçici mantık, inanç revizyonu ve bağlam mantığı) mantık temelli yapay zekâ çalışmalarına çok benzer olması ve nihai hedefin (matematiksel olmayan akılyürütmenin formelleştirilmesi) aynı olması nedeniyle yapay zekâda kullanılan mantık, felsefe kökenli mantık geleneğinin sürekli bir uzantısı olarak görünmektedir.²²⁹

Mantık içinde birbirinden ayrı matematiksel ve felsefi alt-disiplinlerin oluşması sağlıklı bir gelişme olarak değerlendirilmemiştir. Günümüzde mantık alanı, matematiğin çoğu alanları kadar önemli, zorlu problemler ve sonuçlardan oluşmaktadır. Son tahlilde, mantık temel olarak akılyürütmeyle uğraşır. Yaptığımız akılyürütmelerin görece az bir kısmı matematikselidir. Matematikçi olmayanların yaptığı matematiksel akılyürütmelerin ise neredeyse tümü sadece basit hesaplardan ibarettir. Hem kesinliğe hem de geniş bir etki alanına sahip olması için, mantığın matematiksel ve felsefi taraflarını bir disiplin altında birleştirilmesi gerekir. Son yıllarda ne matematikle ne de felsefeyle uğraşanlar bu birleşmeyi geliştirecek çok fazla şey yapmamışlardır. Fakat bilgisayar biliminin ihtiyaçları kuvvetli bir birleştirici etki oluşturur. Bilgisayar bilimindeki mantık araştırmaları uygulayıcılarını mutlaka matematiksel olması gerekmeyen akılyürütme alanlarıyla da temasa sokar ve yenilikçi mantık teorileri üretme ihtiyacı yaratır.²³⁰

Felsefe kökenli mantığa en yakın alan yapay zekâdaki mantık çalışmalarıdır. Özellikle de bilgisayar biliminin akılyürütme konusundaki yenilikçi çalışmaları mantıkta yeni yöntemlerin geliştirilmesine önayak olmuştur. Bir benzetme ile söylersek: Newton fiziği ile mekanik mühendisliği arasında nasıl bir bağlantı varsa,

229Richmond Thomason: **a.y.**

230Richmond Thomason: **a.y.**

mantık ile yapay zekâ arasında da benzer bir bağlantı vardır. Bu anlamda, yapay zekâ mantığının başlıca uygulama alanıdır.²³¹

Zekânın daha iyi bir şekilde modellenenebilmesi için gündelik hayattaki zeki davranışlarımızı oluşturan sağduyu akılyürütmesinin formelleştirilmesi gerekir. Önermeler mantığı bu formelleştirmenin üstesinden gelemeyince 1970'li yılların sonunda monoton-olmayan mantık geliştirilmiştir.²³² Monoton-olmayan mantık sağduyu akılyürütmesi dediğimiz gündelik hayattaki zeki davranışı ve karar vermeyi modellemeyi amaçlamaktadır.²³³

231Jack Minker: "Introduction to Logic-based Artificial Intelligence", Ed. Jack Minker, **Logic-based Artificial Intelligence**, Boston, Kluwer Academic Publishers, 2000, s. 9.

232Jack Minker: **a.g.e.**, s. 11.

233Alexander Bochman: "Nonmonotonic Reasoning", **Handbook of the History of Logic**, Vol. 8, Ed. Dov M. Gabbay, John Woods, Amsterdam, Elsevier, 2007, s. 336.

4. Monoton-olmayan Mantık

Önermeler mantığında bir çıkarım bağıntısına, yeni önerme eklenmesi ile önceden ulaşılan sonuç değişmiyorsa **monotondur**. Monotonluğu şöyle de ifade edebiliriz: Eğer φ önermesi, Γ önerme kümesinden çıkarılabiliyorsa, Γ alt-kümesini içeren²³⁴ Δ kümesinden de çıkarılabilir. Dolayısıyla $\Gamma \models \varphi$ ve $\Gamma \subseteq \Delta$ ise $\Delta \models \varphi$ dir. Yani bir önerme kümesinden çıkarılan herhangi bir sonuç, önerme kümesine ek bir önerme eklense da sonuç değişmeyecektir. Dolayısıyla öncül kümesi büyüdükçe sonuçlar kümesi de monoton bir biçimde büyür. Bu durum içerme bağıntısının doğası gereğidir. $\Gamma \models \varphi$ için, Γ kümesindeki tüm önermelerin doğru olduğu durumlar için φ doğrudur.²³⁵

Monotonluk özelliği sadece önermeler mantığında değil, kanıt teorisi ve model teorisinde de vardır. Kanıt teorisi açısından monotonluk, Γ önerme kümesinden φ önermesinin her türlü çıkarımının aynı zamanda genişletilmiş önerme kümesi olan Δ 'dan de yapılabildiği gerçeğinden çıkar. Monotonluğun doğrulanması model teorisi açısından daha dolaysızdır: Δ 'nın her modeli Γ 'nın da bir modeli olduğundan, φ önermesi Γ 'nin her modelinde bulunuyorsa, aynı zamanda Δ 'nın her modelinde bulunuyor olmalıdır.²³⁶

Bunun tersine, gündelik akılyürütme çoğunlukla monoton değildir, çünkü risk içerir: Çıkarım yapmak için yetersiz öncüllerden sonuçlara varırız. Ne zaman risk almaya degeceğini, hatta risk almak gerektiğini (örneğin, tıbbi tanı koyarken) biliriz.

234K önerme kümesinin p gibi bir önermeyi **içerme**'si, (p 'yi de yorumlayan) K 'nın her doğrulayıcı yorumunda p 'nin doğru olmasıdır.

235G. Aldo Antonelli: "Non-monotonic Logic", **Stanford Encyclopedia of Philosophy**, (Çevrimiçi), <http://plato.stanford.edu/entries/logic-nonmonotonic/>, 21 Eylül 2009; Karl Schlechta: "Nonmonotonic Logics: A Preferential Approach", **Handbook of the History of Logic**, Ed. Dov M. Gabbay, John Woods, Amsterdam, Elsevier, 2007, s. 451.

236John F. Horty: "Nonmonotonic Logic", **The Blackwell Guide to Philosophical Logic**, Ed. Lou Goble, London, Blackwell Publishers, 2001, s. 336.

Ama bu tür bir çıkarımının geçersiz olabileceğini de biliriz. Çünkü yeni bilgiler eski sonuçları değiştirebilir. Bu durumu karşılamak için monoton-olmayan mantık geliştirilmiştir.²³⁷

Monoton-olmayan mantık, basit olarak sonuç ilişkisi monotonluk özelliğini karşılamayan mantıktır. Daha fazla öncülün eklenmesi, önceden yapılmış bir çıkarımın geçersiz kalmasına yol açabilir. Dolayısıyla sonuç kümesi öncül kümesiyle birlikte monoton olarak artmak zorunda değildir. Monoton-olmayan mantık genellikle yapay zekâ kökenli mantık ailesi için kullanılır ve zeki davranışlarımızı yönlendiren **öndeğer akılyürütmesi** (default reasoning) düzeneklerini formelleştirmeyi amaçlar.²³⁸

Öndeğer akılyürütmesi, malumatın varlığı kadar yokluğuna da dayanan bir akılyürütmedir. Malumat yokluğunu şu şekilde ifade edebiliriz: *P* önermesi sonucunda *Q* önermesi çıkar, aksine bir malumat olmadıkça. Bu tür bir akılyürütme mantıksal açıdan monoton-olmayan bir sonuç ilişkisi gerektirir. Örneğin, “Kuşlar uçar” önermesi sonucunda, aksine bir bilgi olmadıkça, *x* bir kuş ise *x*’in uçtuğu sonucuna varılır. Bize **Tux**’un²³⁹ bir kuş olduğu söylendiğinde, yukarıdaki akılyürütme sonucunda Tux’un uçtuğunu düşünürüz. Çünkü önerme kümesinde aksine bir malumat yoktur. Fakat önerme kümesine Tux’un bir penguen olduğu malumatı eklendiğinde, “Tux uçar” önermesi geri çekilmek zorunda kalınacaktır. Çünkü bu sonuca varan olağan çıkarım, aksine malumatın olmamasına dayanmaktadır. Fakat yeni önerme kümesinde artık böyle bir malumat vardır. Bunu niceleme mantığı içerisinde ifade edebilmek için tüm istisnaların tanımlanması gerekmektedir: “Penguenler, deve kuşları, tavuklar... hariç bütün kuşlar uçar”. Monoton-olmayan mantık, tüm istisnai durumları listelemeyi gerektirmeyen bir çatı geliştirmeyi amaçlar.²⁴⁰

237Yücel Yüksel: **a.g.e.**, s. 41.

238Patrick Doherty, Witold Lukaszewicz, Andrzej Skowron, Andrzej Szalas: **Knowledge Representation Techniques**, Berlin, Springer, 2006, s. 77.

239Linux işletim sisteminin resmî maskotu olan penguen.

240John F. Horty: **a.g.e.** s. 337.

4.1 Monoton-olmayan Mantığın Temel Problemleri

Monoton-olmayan mantığın gelişmesine yol açan temel problemler aşağıda ele alınmaktadır. Bu problemlere dair çözüm önerileri monoton-olmayan mantığın en işlevsel türü olan **öndeğer mantığı** anlatılırken verilecektir.

4.1.1 Çerçeve Problemi

Yapay zekâda ele alınan en önemli akılyürütme görevlerinden biri planlamadır. Planlama en basit şekliyle belirli bir başlangıç konumundan belli bir hedefe ulaşmak için gereken bir dizi hareketi bulma problemidir. Mantıksal bir çerçeve dahilinde planlama problemi genellikle değişkenler hesabı (calculus) bakış açısından çalışılır. Yani Φ önermesinin s durumuna ait olduğu gerçeğini temsil eden $Sağlar[\Phi, s]$ formunda ifadeler içeren, çeşitli hareketlerin etkilerinin tanımına izin veren önermeler mantığı kullanılarak yapılan bir formelleştirme şeklidir.²⁴¹

Bu formelleştirmenin kullanımına örnek verecek olursak, bir masa üzerinde dört tane blok düşünelim: A , B , C ve D . A , C ve D blokları masanın yüzeyinde yer almaktadır, B bloğu ise A bloğunun üzerindedir. Diğerlerinin üzerinde herhangi bir şey yoktur. Bu durum $s1$ olarak tanımlanırsa, bu duruma bağlı bazı olgular şu formüllerle ifade edilebilir:

$$Sağlar[\dot{U}zerinde(B, A), s1]$$

$$Sağlar[\dot{U}zeri_boş(B), s1]$$

$$Sağlar[\dot{U}zeri_boş(C), s1]$$

$$Sağlar[\dot{U}zeri_boş(D), s1]$$

(1)

²⁴¹John F. Horty: **a.g.e.** s. 338.

Yani B bloğu A bloğunun üzerindedir önermesi ile B , C , ve D bloklarının üzeri boştur önermeleri $s1$ durumuna aittir. $\ddot{U}zerinde(B, A)$ ve $\ddot{U}zeri_boş(B)$ ifadeleri önermeleri ve olguları temsil eden karmaşık terimler olarak dilbilgisel bir yapıya sahiptir.²⁴²

Bu blokların bir robot kolu tarafından hareket ettirilebildiğini düşünelim. Bu robot kolu sadece iki hareket yapabiliyor olsun: 1) bir bloğu diğerinin üzerine koymak, 2) bir bloğu diğerinin üzerinden alıp masanın üzerine koymak. $\ddot{U}zerine_koy(X, Y)$ ve $\ddot{U}zerinden_al(X, Y)$, X 'i Y 'nin üzerine koyma ve X 'i Y 'nin üzerinden alma hareketlerini temsil etsin. Bu hareketlerin etkileri şu aksiyomlar ile bulunabilir:

$$(Sağlar[\ddot{U}zeri_boş(X), s] \wedge Sağlar[\ddot{U}zeri_boş(Y), s] \wedge X \neq Y) \supset \\ Sağlar[\ddot{U}zerinde(X, Y), Res(\langle \ddot{U}zerine_koy(X, Y) \rangle, s)]$$

$$(Sağlar[\ddot{U}zerinde(X, Y), s] \wedge Sağlar[\ddot{U}zeri_boş(X), s]) \supset \\ Sağlar[\ddot{U}zeri_boş(Y), Res(\langle \ddot{U}zerinden_al(X, Y) \rangle, s)]$$

(2)

Buradaki tüm değişkenlerin niceleyicisinin tümel olduğu varsayılır. α bir dizi hareketi ifade eder. $Res(\alpha, s)$ ise s durumu ile başlanıp α 'daki hareketler sırayla uygulandığında ortaya çıkan son durumu ifade eder. Öyleyse bu iki aksiyomdan ilkinin şunu ifade etmektedir: s durumunda X ve Y bloklarının üzeri boş iken; X , Y 'nin üzerine konduğunda s durumundan ortaya çıkan son durum X 'in Y 'nin üzerinde olduğu durumdur. İkinci aksiyomda ise X , Y 'nin üzerindeyse ve s durumunda X 'in üzeri boş ise; X , Y 'nin üzerinden alındığında s 'den ortaya çıkan durumda Y 'nin üzeri boştur.²⁴³

Bu iki aksiyom, gerçekleştiren bir eylem sonucunda, eylemin yol açtığı etkileri tanımlar. Daha uzun dizilerin etkileri tümevarımsal olarak şu koşulla tanımlanabilir:

²⁴²John F. Horty: **a.g.e.** s. 338.

²⁴³John F. Horty: **a.g.e.** s. 339.

$$Res(\langle A_1, \dots, A_n \rangle, s) = Res(\langle A_n \rangle, Res(\langle A_1, \dots, A_n \rangle, s)) \quad (3)$$

n birden büyük olmak şartıyla; s durumunda n tane hareket içeren bir dizinin uygulanmasının sonucu, sonuncusu dışındaki tüm hareketlerin uygulanmasıyla ortaya çıkan durumda, bu hareketlerin sonuncusunun tek başına uygulanmasının sonucuna eşittir.²⁴⁴

Varsayalım ki Γ , s başlangıç durumunun tanımı ile mümkün hareketlerin etkilerini belirleyen aksiyomları içermektedir. Örneğin **Res fonksiyonunun** (Res function) tümevarımsal tanımını da içeren bir önerme kümesi olsun ve Φ ise hedef olarak belirlenen önermeyi temsil etsin. Buradaki planlama problemi, bir α dizisi bulunması problemidir. Bu α dizisinin s başlangıç durumundaki uygulamasında Φ hedef önermesinin içinde bulunduğu durum, Γ tarafından sağlanan malumat ile ispatlanır. Daha formel bir ifadeyle:²⁴⁵

$$\Gamma \vdash \text{Sağlar}[\Phi, Res(\alpha, s)]$$

Örneğin, (1)'deki $s1$ başlangıç durumu olsun ve Γ , (1), (2) ve (3)'teki şu ifadeleri kapsasın: 1) başlangıç durumunu tarif edilen dört önerme, 2) *Üzerine_koy* ve *Üzerinden_al* hareketlerini tanımlayan aksiyomlar, 3) Res fonksiyonunun tümevarımsal özelliğini tanımlayan önermeler. Şimdi A bloğunun C bloğunun üzerinde olduğu bir durum hedeflensin, yani *Üzerinde*(A, C). Bu hedefe şu plan doğrultusunda ulaşılabilir: önce B 'yi A 'nın üzerinden al, sonra A 'yı C 'nin üzerine koy: $\langle \text{Üzerinden_al}(B, A), \text{Üzerine_koy}(A, C) \rangle$. Planın formel ifadesi aşağıdaki gibidir:²⁴⁶

$$\Gamma \vdash \text{Sağlar}[\text{Üzerinde}(A, C), Res(\langle \text{Üzerinden_al}(B, A), \text{Üzerine_koy}(A, C) \rangle,$$

²⁴⁴John F. Horty: **a.g.e.** s. 339.

²⁴⁵John F. Horty: **a.g.e.** s. 339.

²⁴⁶John F. Horty: **a.g.e.** s. 339.

s1)]

Fakat formelleştirilmesine rağmen bu sonuç sağlanamaz. Çünkü Γ , (1)'deki $\ddot{U}zerinde(B, A)$ ve $\ddot{U}zeri_boş(B)$ yargılarını içerdiğinden, $\ddot{U}zerinde_al$ aksiyomundan hareketle şu sonuca varılabilir:²⁴⁷

$$Sağlar[\ddot{U}zeri_boş(A), Res(\langle \ddot{U}zerinden_al(B, A), s1 \rangle)] \quad (4)$$

$s1$ durumunda iken B , A 'nın üzerinden alındığında A bloğunun üzeri boş olduğu ifade edilir. Γ , $Sağlar[\ddot{U}zeri_boş(C), s1]$ ifadesini içerdiğinden, başlangıç durumunda C bloğunun üzerinin zaten boş olduğu bilinir. Şimdi A 'nın üzeri de boş olduğundan, sadece A bloğunu C bloğu üzerine koyarak hedef duruma ulaşılabilceğini düşünmek mantıklıdır, yani $\ddot{U}zerine_koy$ aksiyomu şu ifadeyi çıkaracak şekilde kullanılabilir.²⁴⁸

$$Sağlar[\ddot{U}zerinde(A, C), Res(\langle \ddot{U}zerine_koy(A, C) \rangle, Res(\langle \ddot{U}zerinden_al(B, A) \rangle, s1))]]$$

Buradan Res fonksiyonunun tanımıyla istenen sonuç ortaya çıkar. $\ddot{U}zerine_koy$ aksiyomunun bu uygulaması, sadece C 'nin başlangıçta üzerinin boş olduğunun bilinmesini değil, $\ddot{U}zerinden_al(B, A)$ hareketiyle ortaya çıkan durumda da C 'nin üzerinin boş kaldığının bilinmesini de gerektirir. Yani bir ara basamak olarak şu ifadenin sağlanması gerekir:²⁴⁹

$$Sağlar[\ddot{U}zeri_boş(C), Res(\langle \ddot{U}zerinden_al(B, A) \rangle, s1)]$$

Bu ara basamak sıradan akılyürütme açısından son derece doğaldır: C

²⁴⁷John F. Horty: **a.g.e.** s. 339.

²⁴⁸John F. Horty: **a.g.e.** s. 339.

²⁴⁹John F. Horty: **a.g.e.** s. 339.

başlangıçta üzeri boş olduğundan; B , A 'nın üzerinden alındığında da üzerinin boş kalacağı düşünülür. Fakat aslında I 'de bu çıkarımı yapmaya olanak verecek bir bilgi yoktur. Gerçekten de mantıksal açıdan bu basamak çıkarılabilir değildir. Çünkü en azından B 'nin A 'nın üzerinden alınmasının C 'nin üzerinin boş olmasıyla çatışması her zaman olasıdır. Örneğin: B ve D blokları birbirine bir tel ile bağlı olabilir, öyle ki B 'nin A 'nın üzerinden alınması D 'nin C 'nin üzerine çıkmasına neden olur. Bu olasılık I 'deki malumatlarla tutarlıdır. Karşımızdaki problem ilk defa **McCarthy** ve **Hayes** tarafından fark edilen **çerçeve problemidir**. Bir hareket uygulandığında bazı olgular değişir, bazıları değişmez. Çerçeve problemi, oluşan hareket sonucu değişen tüm olguların tespit edilememesi problemi olarak açıklanabilir.²⁵⁰

4.1.2 Nitelik Problemi

$\dot{U}zerine_koy(X, Y)$ hareketinden ortaya çıkacak her durumda X 'in Y 'nin üzerinde olacağını değil, sadece başlangıçta X ve Y üzerileri boş ve ayrı bloklar olduğu sürece X 'in Y 'nin üzerinde olacağını bildirmektedir. Bu nitelikler planın uygulanması için gereklidir. Çünkü robot kolu üzeri boş olmayan bloklara erişemez ve bir bloğun kendi üzerine konması imkânsızdır.²⁵¹

Ancak eğer bu nitelikler sağlanırsa, $\dot{U}zerine_koy$ aksiyomu doğru olacak mıdır? Hayır. Ya X bloğu robot kolunun tutamayacağı kadar kaygansa ya da X , Y bloğunu kırarak kadar ağırsa? Y üzerine başka bir blok konduğunda patlayan bir bombaysa? Bu garip olasılıkların ortaya koyduğu problem **nitelik problemi** olarak bilinir. Gerçek dünyadaki bir hareketin istenen etkiyi yaratabilmesi için gereken tüm önkoşulların listelenmesi imkânsızdır. Sorun şudur: hareketleri yöneten aksiyomların doğru ve uygun nitelikli formüllerine nasıl ulaşacağız?²⁵²

Olası tüm nitelikleri açık ya da örtük biçimde aksiyomların öncülleri olarak

²⁵⁰John F. Horty: **a.g.e.** s. 340.

²⁵¹John F. Horty: **a.g.e.** s. 340.

²⁵²John F. Horty: **a.g.e.** s. 340.

koymak bir çözüm olabilir. Örneğimizde, *Üzerine_koy* hareketiyle çatışacak bir durumu temsil eden *Garip* adında bir sabit önerme tanımlanabilir. Hareketi yöneten aksiyom da böyle bir durumun oluşmaması önkoşuluyla değiştirilebilir:

$$(Sağlar[Üzeri_boş(X), s] \wedge Sağlar[Üzeri_boş(Y), s] \wedge X \neq Y \wedge \neg Garip) \supset Sağlar[Üzerinde(X, Y), Res(\langle Üzerine_koy(X, Y) \rangle, s)]$$

(5)

Bir önceki paragrafta düşünülen çatışma yaratacak durumlar *Garip* olarak sınıflandırılabilir:

$$\begin{aligned} Kaygan(X) &\supset Garip \\ Ağır(X) &\supset Garip \\ Bomba(Y) &\supset Garip \end{aligned}$$

(6)

Ancak bu çözümde iki problem çıkar. Birincisi üzerine koyma hareketiyle çatışabilecek durumların listesi açık uçludur. Düşünülebilecek herhangi bir liste asla tam olmayacaktır. İkinci problem daha zordur ve görece yeterli bir nitelikler listesi olsa bile ortaya çıkar. Bir hareket aksiyomunun öncülü olarak önkoşullar tanımlanmadığında, hareketin başarılı olduğu sonucuna varmadan önce önkoşulların sağlandığının doğrulanması gerekir. *Üzerine_koy* aksiyomu örnek dahilinde mantıklı görünür. Fakat *X*'in *Y*'nin üzerine konmasının sonucunun başarılı olduğunun bilinmesi için *X* ve *Y*'nin üzerilerinin boş olduğunun doğrulanması gerekir. Hareketle çatışacak hiçbir *garip* durumun olmadığına doğrulanması pratik olmamaktan öte, imkânsıza yakındır.²⁵³

²⁵³John F. Horty: **a.g.e.** s. 341.

4.1.3 Kapalı-dünya Varsayımı Akılyürütmesi²⁵⁴

Türk Hava Yolları'nın İstanbul'dan Isparta'ya direkt uçuşu olup olmadığını müşteri temsilcisine sorduğumuzda, müşteri temsilcisi uçuş bilgilerinin yer aldığı veritabanında İstanbul-Isparta uçuşu için sorgulama yapar. Bu veritabanı aşağıdakine benzer mantık cümlelerinin bir kümesi olarak düşünülebilir:

Uçuş(TK191, Ankara, İstanbul)

Uçuş(TK192, İstanbul, İzmir)

Uçuş(TK193, İzmir, Ankara)

Uçuş(TK194, İstanbul, Antalya)

...

(7)

Müşteri temsilcisi bu cümlelere bakarak bir sonuca varır ve bize İstanbul'dan Isparta'ya direkt uçuşunun olmadığını söyler. Fakat müşteri temsilcisi bu sonuca nasıl varmıştır? Uçuş veritabanı hangi şehirlere uçuş olmadığını içermemektedir. Müşteri temsilcisi veritabanında *Uçuş(TKXXX, İstanbul, Isparta)* cümlesinin olmamasına göre bu çıkarımı yapmaktadır. Oysa tipik mantık cümlelerinde sonuç bu tarzdaki olumsuz malumatlardan değil, olumlu malumatlardan türetilebilmektedir. Bu uzlaşım gereğince olumlu önerme veritabanında açıkça listelenmediğinde yanlış olarak sonuçlandırılır. Müşteri temsilcisi de veritabanında İstanbul-Isparta uçuşunu bulamadığında aynı şekilde sonuç çıkartmıştır. Yani doğru olarak bilinmeyen yanlış olduğuna inanılmıştır.²⁵⁵

Kapalı-dünya varsayımı sadece havayolları veritabanlarına değil, olumsuz malumat tarafından ezilmiş bütün olumlu malumat durumları için geçerlidir. Örneğin, Nazlı Hanım vereceği partinin davet listesine baktığında kendisinin listede yer almadığı, dolayısıyla da davet edilmediği sonucuna ulaşabilir. Ya da Yücel Bey masaüstü takvimine baktığında Perşembe saat 15:00'da doktor randevusu kaydı

²⁵⁴İng. Closed world reasoning.

²⁵⁵John F. Horty: **a.g.e.** s. 341.

bulunmadığı için, bu saatte doktor randevusu olmadığı sonucunu çıkarabilir. Kapalı-dünya varsayımı temelli akılyürütme, malumat eksiliğine dayanmak bakımından **öndeğer akılyürütmesi** (default reasoning) modellerini örneklenmektedir.²⁵⁶

4.1.4 Kaldırılabilir Kalıtım Akılyürütmesi²⁵⁷

Başlangıç örneğine geri dönersek: kuşlar uçar, Tux bir kuştur, bu nedenle Tux uçar. Bu tarzdaki akılyürütmelere yapay zekâda **kalıtım akılyürütmesi** (inheritance reasoning) denir ve taksonomik malumatın etkili bir şekilde gösterimi ve erişimi ihtiyacına cevap olarak geliştirilmiştir. Her bir bireyselin özelliklerinin listesini tutmak yerine, sınıflar ve özellikler taksonomik bir hiyerarşide sıralanır. Bireysellerin ait oldukları sınıfların özelliklerini kalıtsal olarak alırlar. Tux'un uçtuğu özelliğini açıkça belirtmek gerekmez. Çünkü bu özellik kuşlar sınıfından alınan kalıtsal bir özelliktir.²⁵⁸

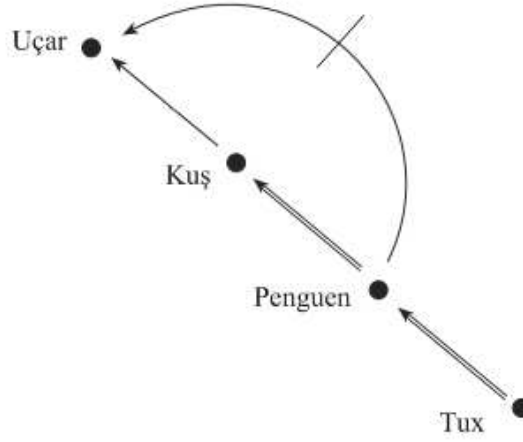
Buna benzer bir **kaldırılabilir kalıtım ağı** Şekil 1'de gösterilmiştir. Değişmez bağlantılar çift çizgili ok ile ($\Rightarrow$), kaldırılabilir bağlantılar ise tek çizgili ok ile gösterilmiştir ($\rightarrow$). Şekil 1'de gösterilen ağdan şu bilgiyi sağlanır: Tux bir penguendir; penguenler kuştur; kural gereği kuşlar uçma eğilimindedir fakat penguenler uçmaz.²⁵⁹

²⁵⁶John F. Horty: **a.g.e.** s. 342.

²⁵⁷İng. Defeasible Inheritance Reasoning.

²⁵⁸John F. Horty: **a.g.e.** s. 342.

²⁵⁹John F. Horty: **a.g.e.** s. 342.



(Şekil 1: Kaldırılabilir Kalıtım Akılyürütmesi)

$Tux \Rightarrow Penguin$ gibi bir bağlantı $Penguin(Tux)$ şeklinde atomik bir yargı olarak, $Penguin \Rightarrow Kuş$ gibi bir bağlantı ise $\forall[Penguin(x) \supset Kuş(x)]$ şeklinde tümel bir yargı olarak gösterilebilir. Ancak sıradan mantıkta kaldırılabilir bağlantıları, yani kuşlar uçar ve penguenler uçmaz anlamına gelen bağlantıları gösterecek bir biçim yoktur.²⁶⁰

Yukarıdaki problemler aslında tek tek sağduyu akılyürütmesinin özellikleridir. Gündelik faaliyetlerimizde kullandığımız bu akılyürütmede pek çok aşama örtüktür veya önkabul gerektirir. Sağduyu akılyürütmesinin bu özelliklerinin modellenmesi için monoton-olmayan mantık adı altında değişik mantık çeşitleri geliştirilmiştir.

²⁶⁰John F. Horty: **a.g.e.** s. 342.

4.2 Öndeğer Mantığı

Monoton-olmayan akılyürütme için kullanılan formelleştirmelerin en bilineni ve en çok kullanılanı **Raymond Reiter** (1939-2002) tarafından geliştirilen **öndeğer mantığıdır** (default logic). Bu formelleştirme sıradan mantığın **öndeğer kuralları** olarak bilinen yeni çıkarım kurallarıyla desteklenmesi, ardından standart sonuç nosyonunun bu yeni kurallara uygun şekilde değiştirilmesinden doğar.²⁶¹

4.2.1 Sembolleştirme

Öndeğer mantığı iki tür cümle üzerine kurulmuştur: D ve W . W dünya (world) hakkındaki kesin olguları barındıran kümeyi, D ise bu dünyaya uygulanacak **öndeğer kuralları** (default rules), ya da kısaca **öndeğer** (default), kümesini temsil eder. Öndeğer teorisi tarafından onaylanmış inançlara **eklenti** (extension) denir.²⁶² Öndeğer kurallarının her biri aşağıdaki biçimde gösterilir:

$$\frac{\text{Önkoşul} : \text{Doğrulama}_1, \dots, \text{Doğrulama}_n}{\text{Sonuç}}$$

Sıradan (tek öncüllü) bir çıkarım kuralı basit olarak (A/B) gibi öncül/sonuç çifti şeklinde gösterilebilir. Bu kural akılyürütenin A sağlandığında B 'ye götürür. Bunun tersine öndeğer bir kural bir üçlüdür, kısaca $(A : C/B)$ şeklinde gösterilir. Öndeğer kuralında A sağlandığında ve buna ek olarak C akılyürütenin sonuç kümesiyle tutarlı olduğunda B 'ye götürür. A **önkoşul**, B **sonuç**, C ise **doğrulama** olarak kabul edilir. Yani öndeğer kuralına göre, *Önkoşul*'un doğru olduğuna inanırsak ve her bir *Doğrulama_i* bizim mevcut inançlarımıza uyuyorsa *Sonuç*'un doğru

²⁶¹Grigoris Antoniou, Kewen Wang: "Default Logic", **Handbook of the History of Logic**, Ed. Dov M. Gabbay, John Woods, Amsterdam, Elsevier, 2007, s. 517.

²⁶²Judea Pearl: **Probabilistic Reasoning in Intelligent Systems**, San Fransisco, Morgan Kaufmann Publishers, 1998, s. 468-469.

olduđuna inanırız.²⁶³

Öndeđer kuralındaki (*D*) bütün formüller ve *W*'nin barındırdığı kesin olgular niceleme mantığı formülleri olarak düşünölmüştür. Fakat diđer formel mantık türleri de kullanılabilir. *D* ve *W* bir araya gelerek **öndeđer teorisini** oluşturur: $\Delta = (W, D)$. *W* aynı zamanda **arkaplan teorisi** (background theory) olarak da adlandırılır.

264

“Kuşlar genellikle uçar” önermesi bir öndeđer kuralı olarak aşğıdaki gibi gösterilir:

$$D = \left\{ \frac{Kuş(x) : Uçar(x)}{Uçar(x)} \right\}$$

Bu ifade şu şekilde okunabilir: Eğer *x* bir kuşsa ve tutarlı bir şekilde uçuyorsa, *x* uçar sonucunu çıkarabiliriz. Arkaplan teorisi kuşlar hakkında aşğıdaki olguları içermektedir:

$$W = \{ Kuş(Serçe), Kuş(Penguen), \neg Uçar(Penguen), Uçar(Kartal) \}$$

Öndeđer kuralına göre, 'serçe uçar' önermesi $Kuş(Serçe)$ önkoşuluna göre doğrudur ve $Uçar(Serçe)$ mevcut malumat ile bağdaşmaz değildir. Bunun aksine $Kuş(Penguen)$, $Uçar(Penguen)$ sonucuna izin vermez. $Kuş(Penguen)$ önermesi doğru olsa da, $Uçar(Penguen)$ doğrulaması verilen bilgiyle bağdaşmaz. Bu öndeđer kuralının arkaplan teorisi açısından $Kuş(Kartal)$ sonucuna varılmaz. Çünkü öndeđer kuralı sadece $Kuş(x)$ 'ten $Uçar(x)$ 'in türetimine izin vermektedir, tersi geçerli değildir.

Şimdi Tux örneğini göstermekle öndeđer mantığının kullanışını gösterelim: Tux'un konu alınan tek nesne olduđu örneğimizde, tek gerekli öndeđer ($Kuş(Tux) : Uçar(Tux)/Uçar(Tux)$)'dir. Yani eđer Tux bir kuşsa, bilinenlerle tutarlı olduđu süreçte Tux'un uçtuđu sonucuna varılabilir. Öyleyse bu ilk örnekteki malumat şu öndeđer

263Judea Pearl: **a.g.e.**, s. 469.

264G. Aldo Antonelli: **a.g.e.**, s. 271.

teoriyle ifade edilebilir: $\Delta_1 = (W_1, D_1)$. Burada $W_1 = \{Kuş(Tux)\}$ ve $D_1 = \{(Kuş(Tux) : Uçar(Tux)/Uçar(Tux))\}$ ifadesidir. Bu örnekte, $Kuş(Tux)$ bilindiği ve $Uçar(Tux)$ kişinin bilgisiyle tutarlı olduğundan öndeğer kural $Kuş(Tux)$ sonucunun çıkarılmasını doğruluyor. Δ_1 'e dayanan uygun sonuç kümesi dolayısıyla $Th(\{Kuş(Tux), Uçar(Tux)\})$ dir. Yani birlikte başlanması söylenen şeyin mantıksal kapatılması ve uygulanabilir öndeğerlerin sonuçlarıdır. Eğer ek olarak Tux'un uçmadığı bildirilirse, $\Delta_2 = (W_2, D_2)$ öndeğer teorisine gidilir. Burada $D_2 = D_1$ ve $W_2 = W_1 \cup \{\neg Uçar(Tux)\}$ şeklindedir. Burada öndeğer kural $(Kuş(Tux) : Uçar(Tux)/Uçar(Tux))$ artık uygulanamaz. Çünkü doğrulaması artık kişinin bilgisiyle tutarsızdır ve dolayısıyla Δ_2 'ye dayanan uygun sonuç kümesi basitçe $Th(W_2)$ dir.²⁶⁵

4.2.2 Eklentiler

Reiter **eklenti**yi W 'yi içeren yeni bir mantık teorisi olarak tanımlamaktadır. T (öndeğer teorisi), W 'deki (başlangıç bilgisi) olguları makul sonuçlar ile genişletir. T tarafından temsil edilen olgulara ait inançlara eklenti denir. Bu eklentiler T 'den çıkarılabilecek sonuçlardır. Eklentiler için şu kurallar geçerlidir: 1) Öndeğer kurallarının hiçbirisi eklentide olmayan sonuçları elde etmek için uygulanamaz; 2) Eklenti ilk koşula minimal şekilde bağlıdır. İlk koşul tüm tutarlı öndeğer sonuçlarını düzenlemeyi garanti etmektedir. İkinci koşul da mantık teorisine sebepsiz yere bir şey eklenmemesini sağlamaktadır.²⁶⁶

Aşağıdaki örnek bir eklentinin nosyonu örneklemektedir. $T = (W, D)$ teorisi aşağıda gibi verilmiş olsun:

$$W = \{ E, E \supset F \}$$

²⁶⁵John F. Horty: **a.g.e.** s. 343-344.

²⁶⁶Judea Pearl: **a.g.e.**, s. 469.

$$D = \left\{ \frac{E:C}{\neg B}, \frac{F:B}{\neg C} \right\}$$

E ve $E \supset F$, W dünyası hakkında bildiğimiz iki şeydir. İlk öndeğere C , W ile tutarlı olduğunda başvurulabilir. Böylece $\neg B$ sonucuna varırız. $\neg B$ ikinci öndeğerinin uygulanmasını engellemektedir. Dolayısıyla ek çıkarımlar yapılamaz. Bu $(W \cup \{\neg B\})^*$ eklentisini sağlamaktadır. $*$ sembolü eklentinin parantezdeki cümlelerin içermelerini (logical implications) ihtiva ettiğini gösterir. Örneğin: F , E $\neg B$ ve tüm totolojileri. İkinci eklenti olan $(W \cup \{\neg C\})^*$ benzer bir şekilde ikinci öndeğer tarafından çağrılabilir.²⁶⁷

Bu örnek göstermektedir ki, bu eklentiler D operatörünün (minimal) sabit-noktalarıdır. Şöyle ki D içindeki öndeğer kurallarının ek uygulamalarının bir eklenti içindeki cümlelere etkisi yoktur. Bu aynı zamanda çoklu eklentilerin imkânlı olduğunu, bunların herbirinin farklı, hatta bazen çelişkili, inanç kümelerini temsil ettiğini örneklemektedir.²⁶⁸

Sıradan mantıkta W formül setiyle ilişkili sonuç kümesi sadece $Th(W)$ dir, W 'nin mantıksal kapanması (logical closure). Öndeğer bir teoriyle $\Delta = (W, D)$ ilişkili sonuç kümesi W 'nin kapanması ile birlikte şu şekilde ifade edilebilir:

$$D = Th(W) \cup [C : (A : B/C) \in D, A \in Th(W), \neg B \notin Th(W)]$$

Teori önkoşulları W tarafından zorunlu kılınan, doğrulamaları ise W ile tutarlı olan öndeğer kuralların sonuçlarıdır. Bir an düşünecek olursak bu önermenin yetersiz olduğunu görebiliriz. Öncelikle bu şekilde tanımlanan E kümesi mantıksal sonuç altında kapanmış bile değildir: öndeğer bir kuraldan çıkan sonucun E kümesine eklenmesi sonuç kümesine dahil olması gereken yeni mantıksal imalar doğurur. Daha da kötüsü öndeğer bir kuraldan çıkan sonucun eklenmesi bir diğerinin ortaya

²⁶⁷Judea Pearl: **a.g.e.**, s. 470.

²⁶⁸Judea Pearl: **a.g.e.**, s. 470.

çıkmasına yol açabilir. Örnek verecek olursak, şu öndeğer teorisini düşünelim: $\Delta_3 = (W_3, D_3)$, burada $W_3 = \{A\}$ ve $D_3 = \{(A : B/C), (C : D/E)\}$ dir. Yukarıdaki tanım ilk öndeğer kuralın C sonucunu doğru biçimde E sonuç kümesine ekliyor. C 'nin varlığı ikinci kuralın ortaya çıkmasına yol açacaktır. Bu da E 'nin de sonuç kümesine eklenmesine neden olacaktır, ancak bu ifade içerilmemektedir.²⁶⁹

Bu örneğin anlattığı, öndeğer bir kural için uygun sonuç kümesinin tanımının yinelemeli olması gerektiğidir. Belki de öndeğer teori $\Delta = (W, D)$ 'nin sonuç kümesini şu şekilde almak gerekir:

$$E = \bigcup_{i=0}^{\infty} E_i$$

ile

$$\begin{aligned} E_0 &= W \\ E_{i+1} &= Th(E_i) \cup [C : (A : B/C) \in D, A \in Th(E_i), \neg B \notin Th(W)] \end{aligned}$$

Bu önerme bir önceki problemin yanıtıdır, Δ_3 öndeğer teorisinin sonuç kümesi olarak istenen şekilde $Th(\{A, C, E\})$ yi verir. Ancak şimdi yeni bir problem vardır: $\Delta_4 = (W_4, D_4)$ teorisi, burada $W_4 = \{A, B \supset \neg C\}$ ve $D_4 = \{(A : C/B)\}$ dir. Yinelemeyi izlersek $(A : C/B)$ kuralının ilk aşamada uygulanabilir olduğu görülebilir, zira önkoşulu $Th(W_4)$ 'e aittir ve doğrulaması bu kümeyle tutarlıdır. Dolayısıyla E_1 'de B vardır. Sadece birazcık daha akılyürütme ile $\neg C$ 'nin E_2 'ye, dolayısıyla E 'ye ait olması gerektiği görülebilir, zira bu formül E_1 'deki bilginin mantıksal bir sonucudur. $(A : C/B)$ kuralı başlangıçta uygulanabilir görünür. Zira uygulanmasından önce, $\neg C$ sonucuna varmak için bir neden yoktur. Fakat kural uygulanırsa verdiği bilgi $\neg C$

²⁶⁹John F. Horty: **a.g.e.** s. 344.

sonucunu çıkarmamıza izin verir. Böylelikle kural kendi uygulanabilirliğini sarsar.²⁷⁰

Tabii ki, bunun gibi öndeğer bir kuralın sarsıldığı bir akılyürütme zinciri amaçsızca uzun olabilir, dolayısıyla öndeğer bir kuralın belli bir bağlamda uygulanabilirliğinden, onu uygulanabilir görünen tüm diğer kurallarıyla beraber bizzat uygulamadan ve sonucun mantıksal kapanması sağlanmadan emin olunamaz. Bu nedenle, öndeğer bir teoriyle ilişkili sonuç kümesi sıradan yinelemeli yoldan, yani çıkarımın uygulanabilir kurallarının sonuçlarını orijinal veriye başarıyla ekleyerek ve sonra bu sürecin limitini alarak tanımlanamaz.²⁷¹

Reiter öndeğer teorilerin uygun sonuç kümelerinin belirlenmesinde sabit-nokta yaklaşımını benimsemek zorunda kalmıştır ve yukarıda da bahsettiğimiz gibi bu kümeler eklentiler olarak tanımlamıştır. Aslında eklenti kavramına ait iki tanım verir. Burada ele alınan ilk tanım, resmi tanım olmasa da, hem daha sezgisel hem de pratikte daha kullanışlıdır. Bu özel tanımın arkasındaki düşünce şudur: Bir öndeğer teorisi verildiğinde önce teori için bir aday eklenti varsayılır, sonra –bu adayı kullanarak– bir sonuç kümesine ilişkin tahminler dizisi tanımlanır. Bu tahmini dizinin limiti orijinal aday ise, bu aday öndeğer teorisi için bir eklenti olarak onaylanır.²⁷²

Tanım 1: E kümesi $\Delta = (W, D)$ öndeğer teorinin bir eklentisidir, ancak ve ancak $E_0, E_1, E_2, \dots$ gibi öyle bir küme dizisi varsa ki

$$E = \bigcup_{i=0}^{\infty} E_i$$

$$E_0 = W$$

$$E_{i+1} = Th(E_i) \cup [C : (A : B/C) \in D, A \in Th(E_i), \neg B \notin Th(W)] \quad 273$$

Burada E kümesi adaydır ve E 'nin tahmini dizi olan $E_0, E_1, E_2 \dots$ 'nin

²⁷⁰John F. Horty: **a.g.e.** s. 345.

²⁷¹John F. Horty: **a.g.e.** s. 345.

²⁷²John F. Horty: **a.g.e.** s. 345.

²⁷³John F. Horty: **a.g.e.** s. 345.

birleşimiyle çakıştığı ortaya çıkarsa Δ 'nın gerçek bir eklentisi olarak onaylanır. E_{i+1} tahmini dizisinin tanımında kullanılan E sembollerinin orijinal aday açısından tanımlandığına dikkat edin.

Eklentilerin sabit-nokta doğası **Reiter**'in resmi tanımında (Tanım 2) daha belirgindir. Bu tanım formül kümelerini formül kümeleriyle eşleştirmek için belli bir öndeğer teorisinden aldığı bilgiyi kullanan Γ operatörüne dayanır.²⁷⁴

Tanım 2: $\Delta = (W, D)$ öndeğer bir teori ise ve S bir formül kümesi ise, $\Gamma_{\Delta}(S)$ üç koşulu sağlayan minimum kümedir²⁷⁵:

1. $W \subseteq \Gamma_{\Delta}(S)$
2. $Th(\Gamma_{\Delta}(S)) = \Gamma_{\Delta}(S)$
3. Herbir $(A : B/C) \in D$ için, eğer $A \in \Gamma_{\Delta}(S)$ ve $\neg B \notin S$ ise $C \in \Gamma_{\Delta}(S)$

Bu tanımdaki ilk iki koşul $\Gamma_{\Delta}(S)$ 'nin orijinal teorisinin verdiği bilgiyi içerdiğini ve mantıksal sonuç altında kapandığını söyler. Üçüncü koşul S 'de uygulanabilen öndeğer kurallarının sonuçlarını içerdiğini ve minimalite kısıtlamasının garanti olmayan sonuçların içeri sızmasını engellediğini söyler. $\Delta = (W, D)$ öndeğer bir teori ise Γ_{Δ} operatörü herhangi bir S formül kümesini, hem sıradan mantıksal sonuç altında hem de S 'de uygulanabilen D 'deki öndeğer kuralları altında kapalı olan W minimal süperkümesiyle eşleştirir. Bir öndeğer teorisi eklentilerini bu operatörün sabit noktaları olarak tanımlar.²⁷⁶

Kural: E kümesi Δ öndeğer teorisinin bir eklentisidir ancak ve ancak $\Gamma_{\Delta}(E) = E$ ise.²⁷⁷

²⁷⁴John F. Horty: **a.g.e.** s. 346.

²⁷⁵John F. Horty: **a.g.e.** s. 346.

²⁷⁶John F. Horty: **a.g.e.** s. 346. Patrick Doherty, Witold Lukaszewicz, Andrzej Skowron, Andrezej Szalas: **a.g.e.**, s. 81.

²⁷⁷John F. Horty: **a.g.e.** s. 346.

Δ_1 ve Δ_2 öndeğer teorileri eklentileri olarak istendiği gibi sırasıyla $Th(\{Kuş(Tux), Uçar(Tux)\})$ ve $Th(\{Kuş(Tux), \neg Uçar(Tux)\})$ kümelerini alır. Burada tanımlanan eklenti nosyonunun sıradan mantıktaki dengi olan sonuç kümesi nosyonunun tutucu bir genellemesi olduğu açıktır: 1) D 'nin boş olduğu (W, D) öndeğer teorisinin eklentisi $Th(W)$ dir. 2) Öndeğer kurallarının tutarsızlık göstermeyecekleri de gösterilebilir: (W, D) öndeğer teorisinin herhangi bir eklentisi, teorisinin sıradan ögesi olan W tutarlı olduğu sürece tutarlı olacaktır.²⁷⁸

4.2.3 Öndeğer Sonucu

Sıradan mantıktaki durumun tersine her öndeğer teorisi tek bir eklentiye, tek bir uygun sonuç kümesine götürmez. Bazı öndeğer teorilerin eklentileri yoktur: Δ_4 buna bir örnektir. Bu teorisinin eklentisinin olmadığı en kolay biçimde nosyonun birinci tanımıyla (Tanım 1) çalışmakta ve Δ_4 'ün bir eklentisi olduğunu, örneğin E , varsaymaktır. Açıkça ya $\neg C \in E$ ya da $\neg C \notin E$ elde edilecektir. Önce $\neg C \in E$ olduğunu varsayalım. $\neg C \notin E_0$ olduğuna göre ve $\neg C \in E$ varsayımı altında tahmini dizi tanımından yola çıkarak $\neg C \notin E_1$, $\neg C \notin E_2$, vd. olduğunu görmek kolaydır. Ancak E basitçe E_0, E_1, E_2 vd. birleşimi olduğundan varsayımın tersine $\neg C \notin E$ sonucu çıkar. Şimdi de $\neg C \notin E$ olduğunu varsayalım. Bu durumda $\neg C \in E_2$ olduğunu ve E_2, E 'nin alt kümesi olduğuna göre $\neg C \in E$ olduğu görülebilir, ki bu da varsayıma ters düşer.²⁷⁹

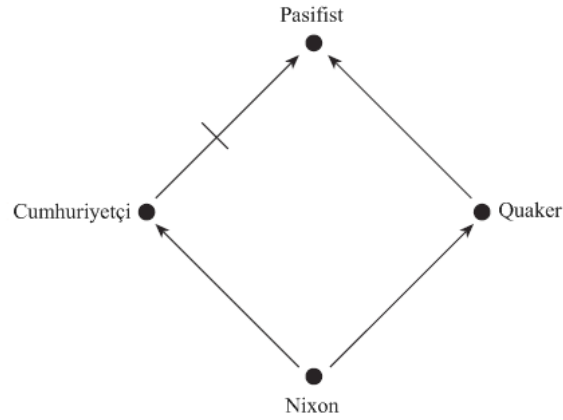
Eklentisi olmayan öndeğer teoriler genellikle anlamsız görülür ve muhtemelen kuraldışı diye dışlanır. Ancak aynı zamanda birden çok eklentiye izin veren tamamen anlamlı öndeğer teorileri de vardır. Bunlara Nixon Elması²⁸⁰ örneği gösterilebilir.

²⁷⁸John F. Horty: **a.g.e.** s. 346.

²⁷⁹John F. Horty: **a.g.e.** s. 347.

²⁸⁰John F. Horty: **a.g.e.** s. 347.

- Nixon bir Quaker'dır²⁸¹.
- Nixon bir Cumhuriyetçidir
- Quaker'lar pasifist olmaya eğilimlidirler.
- Cumhuriyetçiler pasifist olmaya eğilimli değildir.



(Şekil 2: Nixon Elması)

Burada Quaker'lar için ifade edilen önermeler Nixon üzerinde kanıtlamak istenirse ortaya çıkan teori şudur: $\Delta_5 = (W_5, D_5)$, öyle ki

$$W_5 = \{Quaker(Nixon), Cumhuriyetçi(Nixon)\}$$

ve

$$D_5 = \{(Quaker(Nixon) : Pasifist(Nixon)/Pasifist(Nixon)), \\ (Cumhuriyetçi(Nixon) : \neg Pasifist(Nixon)/\neg Pasifist(Nixon))\}$$

Bu teori eklenti olarak hem $Th(W_5 \cup \{Pasifist(Nixon)\})$ hem de $Th(W_5 \cup$

²⁸¹“Dostlar Derneği” denen protestan mezhebi üyesi. George Fox’un hâkim Bennet’i “Tanrı’ya saygı göstermeye ve Kelâmı önünde titremeye” (ingilizce *to quake*) davet eden sözlerinden kinaye ile “Dostlar Derneği” üyelerine takılan lakaptır.

$\{\neg Pasifist(Nixon)\}$) sonucuna izin verir. Başlangıçta yeni bir sonuç çıkarmadan önce, D_5 'ten gelen iki öndeğer kural da uygulanabilir. Fakat bunlardan biri benimsendiğinde diğerinin uygulanabilirliği engellenmiş olur.²⁸²

Bir öndeğer teorisinin birden fazla eklentiye götürdüğü bu tür durumlarda, teoride verilen bilgiden hangi sonuçların çıkarılacağına karar vermek zordur. Bu zorluğu aşmak için üç seçenek önerilmiştir: 1) Teorinin birkaç eklentisinden birinin rastgele seçilmesi ve bunun barındırdığı sonuçların kabul edilmesidir. 2) Bir sonucu öndeğer teorisinin bir eklentisinde içerildiği sürece kabul etmeye istekli olmaktır. Bazen bu iki seçeneğin herşeye inanan akılyürütme stratejisini yansıttığı söylenir. 3) Bir sonucu sadece öndeğer teorisinin tüm eklentisinde içeriliyorsa kabul etmektir.²⁸³

Bu seçeneklerden birincisi çelişen bilgilerle akılyürütme için akılcı bir tutum gibi görünmektedir. Sıkça, kişi bu tür bilgi verildiğinde, çelişkilerin belli bir biçimde çözüldüğü içsel olarak tutarlı bir bakış açısını benimser, çelişkilerin başka bir biçimde çözüldüğü başka tutarlı bakış açıları olmasına karşın. Bu akılyürütme tutumu akılcı olsa da, bu tür bir tutumun nasıl formel bir sonuç ilişkisi olarak kodlanacağını anlamak güçtür. Eklenti seçimi gerçekten de rastgeleyse, farklı kişiler kolaylıkla farklı eklentiler seçebilir veya aynı kişi farklı zamanlarda farklı eklentiler seçebilir. Bu durumda teorinin sonuç kümesini temsil eden eklenti hangisidir?²⁸⁴

İkinci seçenek bir sonuç ilişkisi olarak kodlanabilse de bu garip bir kodlama olacaktır. Bu tutuma göre bir öndeğer teorisinin sonuç kümesinin standart mantıksal sonuçta kapanmasına gerek yoktur ve aslında tutarsız da olabilir. Örneğin Δ_5 'nin sonuç kümesi hem *Pasifist(Nixon)* hem de $\neg Pasifist(Nixon)$ sonucunu içerecektir, zira bu formüllerin ikisi de öndeğer teorisinin bir eklentisine aittir. Ancak *Pasifist(Nixon) $\wedge$ $\neg Pasifist(Nixon)$* 'yi içermez. Bu ikinci seçenek, bir öndeğer teorisine dayanarak inanılması gereken formüllerin değil, inanılabilir formüllerin bir tarifini verir.²⁸⁵

Sadece üçüncü seçenek doğal bir sonuç ilişkisi ortaya çıkarır. Aşağıda

282John F. Horty: **a.g.e.** s. 347. Patrick Doherty, Witold Lukaszewicz, Andrzej Skowron, Andrej Szalas: **a.g.e.**, s. 81.

283John F. Horty: **a.g.e.** s. 347.

284John F. Horty: **a.g.e.** s. 348.

285John F. Horty: **a.g.e.** s. 348.

gösterirsek:

Tanım 3: $\Delta = (W, D)$ bir öndeğer teorisi ve A da bir formül olsun. O zaman A Δ 'nın şüpheli bir sonucudur, $\Delta \vdash A$ olarak gösterilir. Δ 'nın her E eklentisi için $A \in E$ durumunda $\Delta \in A$ şeklinde yazılır.²⁸⁶

Şimdi formel bir sonuç ilişkisi tanımlandığına göre, açıkça belirtmek gerekir ki, bu iki açıdan monoton değildir: Hem bir öndeğer teorisinin W -ögesi'ne yeni bir olgu bilgisi eklenmesi, hem de D -ögesi'ne yeni öndeğer bilgisi eklenmesi kişiyi daha önce benimsediği sonuçlardan vazgeçmeye zorlar. İlk olasılık öndeğer teorileri Δ_1 ve Δ_2 'ye geri dönerek açıklamaktır. Burada, $\Delta_1 \vdash \text{Pasifist}(\text{Nixon})$ dır. Ancak Δ_2 , Δ_1 'in W -ögesi'ne yeni olgu bilgisi $\Delta_2 \vdash \text{Pasifist}(\text{Nixon})$ 'nin eklemesiyle elde edilmiş olmasına rağmen, $\Delta_2 \vdash \text{Pasifist}(\text{Nixon})$ değildir. İkinci durumu açıklamak için şu öndeğer teorisini düşünelim:

$$\begin{aligned}\Delta_6 &= (W_6, D_6), \text{ öyle ki } W_6 = W_5 \\ D_6 &= \{(\text{Quaker}(\text{Nixon}) : \text{Pasifist}(\text{Nixon})/\text{Pasifist}(\text{Nixon}))\}.\end{aligned}$$

Bu teori Cumhuriyetçilerin pasifist olmaya meyilli olmaması öndeğeri dışında Δ_5 , Nixon Elmasına benzemektedir. Kolaylıkla görülebilir ki Δ_6 'nın tek eklentisi $\text{Th}(W_6 \cup \{\text{Pasifist}(\text{Nixon})\})$ dır, öyle ki $\Delta_6 \vdash \text{Pasifist}(\text{Nixon})$. Buna rağmen Δ_5 teorisinin iki eklentisi vardır. Biri $\text{Pasifist}(\text{Nixon})$ içermez. Diğerinde ise Δ_5 'in yeni öndeğer bilgisi $(\text{Cumhuriyetçi}(\text{Nixon}) : \neg \text{Pasifist}(\text{Nixon})/\neg \text{Pasifist}(\text{Nixon}))$, Δ_6 'nın D -ögesi'ne eklenmesiyle ortaya çıkmış olmasına rağmen $\Delta_5 \vdash \text{Pasifist}(\text{Nixon})$ değildir.²⁸⁷

²⁸⁶John F. Horty: **a.g.e.** s. 348.

²⁸⁷John F. Horty: **a.g.e.** s. 348.

4.2.4 Problemlere Çözümler ve Yarı-normal Öndeğerler

Öndeğer mantığında çerçeve probleminin dolaysız bir çözümü varmış gibi görünmektedir. İlk durumun standart mantık ifadesi ve mevcut hareketleri, basitçe olguların değişmeme eğilimi olduğunu söyleyen öndeğer kuralları ile tamamlanırsa çözüm ortaya çıkacaktır. Şöyle ki, problem öndeğer teorisi $\Delta_7 = (W_7, D_7)$ olarak şu şekilde kodlanabilir: Önce, olgusal öge W_7 (1)'den (3)'e kadar formülleri içerir. Bunlar ilk durumu, *Üzerine_koy* ve *Üzerinden_al* hareketlerinin etkilerini ifade eden aksiyomları ve hareket dizilerinin tümevarımsal tanımlarını ifade eder. Sonra, D_7 öndeğer ögesi, öndeğer kuralı şemasının tüm hallerini içerir:

$$(Sağlar[\phi, s] : Sağlar[\phi, Res(\alpha, s)]/Sağlar[\phi, Res(\alpha, s)])$$

Yani, ne zaman ki α hareketinin uygulanmasından sonra hâlâ ϕ 'nin geçerli olduğu sonucu doğruysa, bir ϕ olgusu s durumunda geçerlidir. α uygulanmasından sonra ϕ 'nin hâlâ geçerli olduğu sonucuna bir öndeğer olarak (by default) varılabilir.²⁸⁸

Bu öndeğer teorisinin (4) no'lu formülünü içeren tek bir eklentisi olduğunu, ki bu daha önce çerçeve aksiyomlarının yardımı olmadan çıkarılamayan bir ara adımdır, doğrulamak kolaydır. B, A 'nın üzerinden alındığında C 'nin üzerinin hâlâ boş olduğu önermesi sadece (1)'den (3)'e kadarki olgusal bilgiden çıkarılamazsa da, aksine bilgi olmadıkça olguların değişmemeye meyilli olduğu sonucuna varılmasını söyleyen öndeğer kuralının yardımıyla çıkarılabilir.²⁸⁹

Nitelik problemine dönersek yine öndeğer mantığını kullanarak kısmi bir çözüm bulunabilir. Çözüm basitçe hareketlerin sonucunu değiştirecek garip olayların gerçekleşmemeye meyilli olduğunu söyleyen öndeğer kurallarının eklenmesidir. Verilen örnekte geçen bilgi $\Delta_8 = (W_8, D_8)$ teorisiyle formüle edilebilir. Burada W_8 arkaplan bilgisine ek olarak, değiştirilmiş *Üzerine_koy* aksiyomunu (5) ve (6)'da

²⁸⁸John F. Horty: **a.g.e.** s. 349.

²⁸⁹John F. Horty: **a.g.e.** s. 349.

geçen hareketin sonucunu değiştirecek çeşitli garip olayların özelliklerini içerir. D_8 tek öndeğer olan aşağıdaki ifadeyi içerir:

(Doğru : $\neg\text{Garip}/\neg\text{Garip}$)

Bu ifade, aksine bilgi olmadıkça bu tür garip olayların olmayacağını varsaymamızı söyler (T , evrensel olarak doğru önerme anlamındadır). Bu gösterim nitelik probleminin ortaya attığı iki konunun ilkinin çözmeye yetmeyecektir: *Üzerine_koy* hareketinin sonucunu değiştirebilecek koşulların listesi açık uçludur. Buna rağmen gösterim, bu konuların ikincisi için bir çözüm sunmaktadır. *Üzerine_koy* hareketinin sonucunu değiştirebilecek çeşitli garip olayların bir listesi verilmiş olsa, kişi artık *Üzerine_koy* hareketinin istenen etkileri sağladığı sonucuna varmak için bu koşullardan her birinin geçersiz olduğunu doğrulamaz. Öndeğer kuralı aksine bilgi olmadıkça bu koşulların geçersiz olduğunu varsaymaya izin vermektedir.²⁹⁰

Çerçeve ve nitelik problemleri gibi, kapalı-dünya varsayımı akılyürütmesinin zorlukları da öndeğer mantığına dayanarak çözülebilir görünmektedir. Bir başlangıç önerisi olarak, uçuş bilgisi $\Delta_9 = (W_9, D_9)$ öndeğer teorisiyle gösterilebilir. Öyle ki W_9 (7)'deki olgu bilgilerini ve D_9 şu öndeğer kuralı şemasının her halini içerir:

(Doğru : $\neg\text{Uçuş}(x, y, z)/\neg\text{Uçuş}(x, y, z)$)

Bu ifade, aksine bilgi olmadıkça şehirlerin arasında direk uçuş olmadığını varsaymamız gerektiğini söylemektedir. Bu teorinin tek bir eklentisi olacaktır, o da (mantıklı varsayımlarla, mevcut tüm uçuşların isimlendirildiği gibi) İstanbul ile Isparta arasında direk uçuş olmadığı sonucuna varmaya izin vermektedir.²⁹¹

Çerçeve problemi, nitelik problemi ve kapalı-dünya problemi ile Nixon Elmasını açıklayan örneklerin öndeğer mantığı ile gösterimlerinde ortak bir özelliği

²⁹⁰John F. Horty: **a.g.e.** s. 349.

²⁹¹John F. Horty: **a.g.e.** s. 350.

vardır. Bu durumların her biri tamamen, özel bir form olan $(A : B/B)$ formunun öndeğer kurallarına dayanmaktadır. Bu formda hem doğrulama hem de sonuç olarak aynı formüller geçmektedir. Bu tür öndeğer kuralları **normal öndeğerler**, sadece normal öndeğerler içeren teoriler de **normal öndeğer teorileri** olarak bilinir. Normal öndeğer teoriler genelde öndeğer teorileri tarafından paylaşılmayan bir takım çekici özelliklere sahiptir. En dikkat çeken şudur ki normal öndeğer teorilerin eklentilerinin olması kesindir. Bu çekici özellikleri nedeniyle ve görüldüğü gibi birçok önemli örneğin normal teoriler olarak kodlanabilmesinden dolayı, Reiter öndeğer mantığının tam ifade edici gücünün gerçekçi uygulamalarda gerekmeyebileceğini ve normal teorilerle sınırlandırılabilirliğini öne sürmüştür.²⁹²

Oysa ki çok geçmeden bu tahminin doğru olmadığı anlaşılmıştır. Kaldırılabilir kalıtım akılyürütmesi'ndeki Tux Üçgeni örneğinde görüldüğü üzere. Sadece normal öndeğerler göz önüne alındığında, Tux Üçgeni'ndeki malumat $\Delta_{10} = (W_{10}, D_{10})$ teorisiyle gösterilebilir. Öyle ki W_{10} , $Penguen(Tux)$ ve $\forall x(Penguen(x) \supset Kuş(x))$ cümlelerini içerir. Yani Tux bir penguenlerdir ve bütün penguenler kuştur. D_{10} da $(Kuş(Tux) : Uçar(Tux)/Uçar(Tux))$ ve $(Penguen(Tux) : \neg Uçar(Tux)/\neg Uçar(Tux))$ öndeğerlerini içerir. Yani Tux için şu genel gerçekleri ifade ederler: Kuşlar uçmaya meyillidir ama penguenler değil. Bu öndeğer teorisi, Nixon Elması'nın Δ_4 şeklinde gösterilmesi gibi, iki çelişen öndeğer kuralı içerir ve dolayısıyla iki eklentiye gider.²⁹³

$$Th(W_{10} \cup \{Uçar(Tux)\}) \text{ ve } Th(W_{10} \cup \{\neg Uçar(Tux)\})$$

Peki bu doğru mu? Nixon Elması örneğinde çoklu eklentiler akla yatkındır, çünkü Quaker'lar ile Cumhuriyetçiler hakkındaki öndeğerler görünüşte eşit ağırlıktadır. Ancak Tux Üçgeni örneğinde, penguenler hakkındaki öndeğerin kuşlar hakkındaki öndeğere tercih edilmesi gerekir. Çünkü penguenler özel bir tür kuştur ve her zaman en özel bilgiye dayanarak akılyürütmek en uygundur. Öndeğerler arasındaki bu tür tercihleri yakalamanın bir yolu gösterimi değiştirmektir. İlk olarak

292John F. Horty: **a.g.e.** s. 350. Patrick Doherty, Witold Lukaszewicz, Andrzej Skowron, Andrezej

Szalas: **a.g.e.**, s. 83-84.

293John F. Horty: **a.g.e.** s. 350.

David W. Etherington ve **Reiter** (1983) tarafından keşfedilmiştir. Öyle ki bir öndeğer kuralının uygulanmasına baskın gelebilecek nedenler o kuralın ifadesinde açıkça belirtilmiş olsun. Bu yaklaşıma göre, örneğin Tux Üçgeni'nde kuşlar hakkındaki öndeğer, normal öndeğer kuralı ($Kuş(Tux) : Uçar(Tux)/Uçar(Tux)$) ile değil, fakat **yarı-normal**²⁹⁴ şu kural ile gösterilebilir:

$$(Kuş(Tux) : [Uçar(Tux) \wedge \neg Penguen(Tux)]/Uçar(Tux))$$

Bu kurala göre, Tux'nin kuş olduğu biliniyorsa ve bu Tux'nin uçtuğu ve onun bir penguen olmadığı bilgisiyle tutarlı ise, o zaman onun uçtuğu varsayılmalıdır.²⁹⁵

Yarı-normal kurallara başvurmak Tux Üçgeni'nin ortaya koyduğu ilk problemi çözer: Bu yeni, yarı-normal öndeğer daha önceki Δ_{10} 'daki normal öndeğerlerle yer değiştirdiğinde ortaya çıkan teoremin sadece tek bir eklentisi vardır: $Th(W_{10} \cup \{\neg Uçar(Tux)\})$. Bu teoride Tux açık bir şekilde uçmaz. Bu teorie sadece ($Penguen(Tux) : \neg Uçar(Tux)/\neg Uçar(Tux)$) öndeğer kuralı uygulanabilir. Yeni ($Kuş(Tux) : [Uçar(Tux) \wedge \neg Penguen(Tux)]/Uçar(Tux)$) öndeğer kuralı $Penguen(Tux)$ bilindiği için devreye girmez.²⁹⁶

Bir önceki problemi çözerken, kaldırılabilir kalıtım ağlarından birbiriyle yarışan öndeğerler arasında tercihleri ifade ederken yarı-normal kurallar kullanma stratejisi yeni bir zorluğa yol açar: bilginin kalıtım ağlarından öndeğer kurallarına aktarılması holistikdir; belli bir ifadenin çevirisi bağlamına göre değişebilir. Örnek vermek gerekirse, Tux Üçgeni'ne başka bir tür kuşun, örneğin çok genç kuşların, uçmadığı bilgisini eklediğimizi düşünelim. Gösterime $\forall x(Genç(x) \supset Kuş(x))$ formülünü eklemek gerekir. Yani genç kuşlar kuştur, aynı zamanda ($Genç(Tux) : \neg Kuş(Tux)/\neg Kuş(Tux)$) öndeğeri de eklenmeli, yani Tux için genç kuşların uçmaya meyilli olmadığını ifade eder. Ancak ayrıca, artık kuşların uçmaya meyilli olduğu öndeğerine baskın gelecek olası başka bir nedenimiz olduğuna göre, bu öndeğerin

294A : $B \wedge \neg E/B$ şeklindeki ifadelere de yarı-normal öndeğer denir.

295John F. Horty: **a.g.e.** s. 351.

296John F. Horty: **a.g.e.** s. 351.

önceki gösterimi şu yeni kural ile değiştirilmelidir:²⁹⁷

$$(Kuş(Tux) : [Uçar(Tux) \wedge \neg Genç(Tux)]/Uçar(Tux))$$

Hesaplama yönünden bakıldığında, bu sonuç çekici değildir. Çünkü bir malumat öbeğini güncelleme işini son derece karmaşık bir hale getirir. Zira sadece yeni bilginin gösterimini değil, zaten gösterilmiş olan bilginin yeniden formülasyonunu gerektirir. Felsefi bir bakış açısından da sonuç benzer nedenle çekici değildir, zira holizm²⁹⁸ genelde istenmez: Kuşların uçmaya meyilli olduğu önermesinin anlamının bağlamdan bağlama değişkenlik göstereceği düşünülmez, dolayısıyla çevirisinin de değişmesi gerektiği fikri tuhaf karşılanacaktır.²⁹⁹

Öndeğer mantığı bilgi gösterimi ve akılyürütme için önemli yöntemlerden biridir. Çünkü eksik malumatlı akılyürütmeye izin vermektedir. Bu özelliğiyle özellikle teşhis problemi (diagnostic problems), malumat edinme (information retrieval), düzenleme (regulation) ve mantık programlamada kullanılmaktadır.

297John F. Horty: **a.g.e.** s. 351.

298Felsefede bütün kendisini oluşturan parçaların toplamından daha fazla olan bir şeydir diyen görüştür.

299John F. Horty: **a.g.e.** s. 351.

SONUÇ

Felsefe Bölümü Mantık Anabilim Dalı'nda yapılan bu çalışma mantığın uygulamalarına yöneliktir. Mantığın uygulamaya yönelik kısmı günümüzde mühendislerle birlikte mantıkçıların da dikkatini çekmektedir ve onlar için de bir ilgi alanı oluşturmaktadır. Yapay zekâ ile ilgili bir çalışma söz konusu olduğunda yalnız mantığı değil, sosyoloji, antropoloji ve psikoloji gibi sosyal bilimleri de ilgilendirmektedir. Çünkü yapay zekâ kolayca anlaşılabilceği gibi yalnızca biyolojik ve fizyolojik özellikleri değil, toplumsal ve kültürel yönü de taklit etme yolundadır. Bu özellik mantık disiplininin felsefenin geleneksel sorunlarını bazı uygulama alanlarına taşıması anlamına gelmektedir. Çünkü bilinç gibi, toplumsal ve kültürel değerler gibi, algı gibi felsefi sorunlar yapay zekânın da konusu içindedir ve çözüm için uğraştığı sorunlar arasındadır.

Yapay zekâ oldukça yeni bir çalışma alanıdır. Bu alandaki çalışmaların temelinde 20. yüzyılda mantığın geçirdiği büyük değişim yatmaktadır. Günümüzde en çok kullandığımız araçlardan biri olan bilgisayarın teorisi bu gelişmeye dayanmaktadır. Bununla birlikte yapay zekâ disiplininin ilk yıllarında yapay insan düşünün ne kadar uzak olduğu anlaşılmıştır. Bizce yapay zekâ, John McCarthy'nin önerdiği gibi, akılyürütme ve problem çözmenin özünü bulma çalışmaları olarak anlaşılmalıdır. Bu süreçte de en önemli problem sağduyu akılyürütmesinin formelleştirilmesidir. Bunun da en önemli aracı mantıktır.

Son yüz yıllık gelişimi süresince mantığa, felsefe kökenli araştırmacıların katkıları kadar matematik kökenli araştırmacıların da katkıları olmuştur. Görebildiğimiz kadarıyla yapay zekâda kullanılan mantık çeşitleri köken bakımından felsefe kökenli mantık çalışmalarına daha yakındır. Bunun temelinde ise felsefenin dünyayı açıklama konusundaki iddiası vardır. Bu durum yapay zekânın makro-dünyalar ölçeğinde çözüm getirebilme arzusuyla örtüşmektedir. Bunun gerçekleştirilebilmesi için de yapay zekâ insanın öngörü yeteneğinin temelini oluşturan sağduyu akılyürütmesini formelleştirmeye çalışmaktadır.

Sağduyu akılyürütmesi eksik bilgiye dayanan bir akılyürütme şekli olarak

düşünülebilir. Bu akılyürütme şeklinin modellenebilmesi için sağduyu akılyürütmesinin örtük aşamalarının da formelleştirilebilmesi gerekmektedir. Yapay zekâ bu formelleştirmeyi yapabilmek için monoton-olmayan mantığı geliştirmiştir.

Monoton-olmayan mantık, sağduyu akılyürütmesini formelleştirmede büyük bir problem olarak önümüzde duran çerçeve problemi ve nitelik problemine çözüm üretmektedir. John McCarty'nin de belirtmiş olduğu gibi bu problemler çözülmeden insan benzeri bir zekâ üretebilmek mümkün değildir. Monoton-olmayan mantık üzerine yapılan çalışmalar bu çabaya yönelik bir yakınsama sağlayacaktır. Bununla birlikte, felsefe kökenli mantık araştırmacılarının monoton-olmayan mantığa yönelmesinin felsefeye yeni imkânlar kazandıracağı umulmaktadır.

KAYNAKÇA

- Antonelli, G. Aldo: “Non-monotonic Logic”, **Stanford Encyclopedia of Philosophy**, (Çevrimiçi), <http://plato.stanford.edu/entries/logic-nonmonotonic/>, 21 Eylül 2009.
- Antoniou, Grigoris, Wang, Kewen: Grigoris Antoniou, Kewen Wang: “Default Logic”, **Handbook of the History of Logic**, Ed. Dov M. Gabbay, John Woods, Amsterdam, Elsevier, 2007.
- Block, Ned, Rey, Georges: “Mind, computational theories of”, **Routledge Encyclopedia of Philosophy**, Ed: Edward Craig, London, Routledge, 1998.
- Bochman, Alexander: “Nonmonotonic Reasoning”, **Handbook of the History of Logic**, Vol. 8, Ed. Dov M. Gabbay, John Woods, Amsterdam, Elsevier, 2007.
- Boden, Margaret A.: “Artificial Intelligence”, **Routledge Encyclopedia of Philosophy**, Ed: Edward Craig, London, Routledge, 1998.
- Boden, Margaret A.: **Mind As Machine**, New York, Oxford University Press, 2006.
- Brna, Paul: “Symbolic and Subsymbolic Processing”, (Çevrimiçi), http://www.comp.lancs.ac.uk/computing/research/aai-aid/people/paulb/old243prolog/subsection3_1_2.html, 28 Ağustos 2009.
- Brooks, Rodney: “Elephants Don't Play Chess”, **Robotics and Autonomous Systems**, Sayı 6, 1990. (Çevrimiçi: <http://people.csail.mit.edu/brooks/papers/elephants.pdf>).
- Crevier, Daniel: **AI: The Tumultuous Search for Artificial Intelligence**, New York, Basic Books, 1993.

- Doherty, Patrick,
Lukaszewicz, Witold,
Skowron, Andrzej,
Szalas, Andrezj:
Floridi, Luciano: **Knowledge Representation Techniques**, Berlin, Springer, 2006.
- Görz, Günther,
Nebel, Bernhard:
Grünberg, Teo,
Onart, Adnan,
Grünberg, David,
Turan, Halil:
Gülseçen, Sevinç: **Philosophy and computing: an introduction**, New York, Routledge, 1999.
- Görz, Günther,
Nebel, Bernhard:
Grünberg, Teo,
Onart, Adnan,
Grünberg, David,
Turan, Halil:
Gülseçen, Sevinç: **Yapay Zekâ**, İstanbul, İnkılâp Kitabevi, 2005.
- Görz, Günther,
Nebel, Bernhard:
Grünberg, Teo,
Onart, Adnan,
Grünberg, David,
Turan, Halil:
Gülseçen, Sevinç: **Mantık Terimleri Sözlüğü**, Ankara, METU Press, 2003.
- Harnard, Steven: “Yapay Sinir Ağları, İşletme Alanında Uygulanması ve Bir Örnek Çalışma”, İstanbul Üniversitesi Sosyal Bilimler Enstitüsü, Yayınlanmamış Doktora Tezi, İstanbul, 1993.
- Harnard, Steven: “The Annotation Game: On Turing (1950) on Computing, Machinery, and Intelligence”, 2006, (Çevrimiçi), <http://cogprints.org/3322/1/turing.html>, 31 Temmuz 2009.
- Horty, John F.: “Nonmonotonic Logic”, **The Blackwell Guide to Philosophical Logic**, Ed. Lou Goble, London, Blackwell Publishers, 2001.
- İnönü, Nazlı: “Yapay Zekâ Felsefesi”, **Bilgisayar ve Beyin**, Ed. Ömer R. Akyüz, Mehmet Budak, Reşit Canbeyli, Ferruh Gençer, Işık Tabar Gençer, İstanbul, Nar Yayınları, 1997.
- Kocabaş, Şakir: Yapay Zekâyâ Giriş, s. 3, (Çevrimiçi), http://www.sakirkocabas.com/files/yzgir_1n.rtf, 27 Ağustos 2009.
- McCarthy, John: “A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence”, 1955, (Çevrimiçi), <http://www-formal.stanford.edu/jmc/history/dartmouth/dartmouth.html>, 31 Temmuz 2009.

- McCarthy, John: “What Was Expected, What We Did, and AI Today”, (Çevrimiçi), http://www.engagingexperience.com/2006/07/ai50_ai_past_pr.html, 13 Ağustos 2009.
- McCarthy, John: “Artificial Intelligence, Logic and Formalizing Common Sense”, 1990, (Çevrimiçi), <http://www-formal.stanford.edu/jmc/ailogic.pdf>, 15 Eylül 2009.
- McCarthy, John: “Circumscription – A Form of Nonmonotonic Reasoning”, 1986, (Çevrimiçi), <http://www-formal.stanford.edu/jmc/circumscription.pdf>, 16 Ağustos 2009.
- McCarthy, John: “Some Philosophical Problems from the Standpoint of Artificial Intelligence”, 1969, (Çevrimiçi), <http://www-formal.stanford.edu/jmc/mcchay69.pdf>, 16 Ağustos 2009.
- McCorduck, Pamela: **Machines Who Think** (2nd ed.), AK Peters, Massachusetts, 2004.
- McCulloch, Warren Sturgis, Pitts, Walter: A Hierarchy of Values Determined by the Topology of Nervous Nets, **Bulletin of Mathematical Biophysics**, Vol. 5, 1943, s. 115-133.
- McLaughlin, Brain P.: Brain P. McLaughlin: “Connectionism”, **Routledge Encyclopedia of Philosophy**, Ed: Edward Craig, London, Routledge, 1998.
- Minker, Jack: “Introduction to Logic-based Artificial Intelligence”, Ed. Jack Minker, **Logic-based Artificial Intelligence**, Boston, Kluwer Academic Publishers, 2000.
- Minsky, Marvin: “It's 2001. Where Is HAL?”, Dr. Dobb's Technetcast, (Çevrimiçi), <http://www.ddj.com/hpc-high-performance-computing/197700454>, 14 Ağustos 2009.
- Nabiyev, Vasıf V.: **Yapay Zekâ**, Ankara, Seçkin Yayıncılık, 2005.

- Neumann, Günter: “Programming Languages in Artificial Intelligence”, (Çevrimiçi), <http://www.dfki.de/~neumann/publications/new-ps/ai-pgr.pdf>, 29 Ağustos 2009.
- Newell, Allen, Simon, Herbert: “Computer Science as Empirical Inquiry: Symbols and Search”, 1976, (Çevrimiçi), http://www.rci.rutgers.edu/~cfs/472_html/AI_SEARCH/PS/S/PSSH4.html, 1 Ağustos 2009.
- Pearl, Judea: **Probabilistic Reasoning in Intelligent Systems**, San Fransisco, Morgan Kaufmann Publishers, 1998.
- Russell, Stuart J., Norvig, Peter: **Artificial Intelligence: A Modern Approach** (2nd ed.), New Jersey, Prentice Hall, 2003.
- Schlechta, Karl: “Nonmonotonic Logics: A Preferential Approach”, **Handbook of the History of Logic**, Ed. Dov M. Gabbay, John Woods, Amsterdam, Elsevier, 2007.
- Simon, Herbert: **Models of My Life**, New York, Basic Books, 1991.
- Solomonoff, Ray: “An Inductive Inference Machine”, 1956, (Çevrimiçi), <http://world.std.com/~rjs/indinf56.pdf>, 3 Ağustos 2009.
- Thomason, Richmond: “Logic and Artificial Intelligence”, **Stanford Encyclopedia of Philosophy**, (Çevrimiçi), <http://plato.stanford.edu/entries/logic-ai/>, 18 Eylül 2009.
- Turing, Alan: “On computable numbers, with an application to the Entscheidungsproblem”, **Proceedings of the London Mathematical Society**, Vol. 42, 1937, s. 230-265.
- Turing, Alan: “Computing Machinery and Intelligence”, **Mind**, Vol. LIX, Number 236, s. 433-460
- Ural, Şafak: “Puslu (Fuzzy) Mantık”, **Mantık, Matematik ve Felsefe**, I. Ulusal Sempozyumu 26-28 Eylül 2003 Assos-Çanakkale, Ed. Şafak Ural, Mehmet Özer, Ayten Koç, Arzu Şen, Gürsel Hacıbekiroğlu, İstanbul, İstanbul Kültür Üniversitesi Yayınları, İstanbul, 2004.

- Ural, Şafak: **Pozitivist Felsefe**, İstanbul, Say Yayınları, 2006.
- Ural, Şafak: **Temel Mantık**, İstanbul, Çantay Kitabevi, 1995.
- Winograd, Terry: SHRDLU, (Çevrimiçi),
<http://hci.stanford.edu/~winograd/shrdlu/>, 12 Ağustos 2009.
- Wootton, David: **Modern political thought: readings from Machiavelli to Nietzsche**, Indiana, Hackett Publishing, 1996.
- Yüksel, Yücel: “Yapay Zekâ ve Puslu Mantık”, **Felsefe Arkivi**, İstanbul, İstanbul Üniversitesi Edebiyat Fakültesi Yayınları, 2008, Sayı: 32, s.29-48
-
- “AI set to exceed human brain power”, (Çevrimiçi),
<http://www.cnn.com/2006/TECH/science/07/24/ai.bostrom>, 14 Ağustos 2009.
-
- “Applications”, AITopics, (Çevrimiçi),
<http://www.aaai.org/AITopics/pmwiki/pmwiki.php/AITopics/Applications>, 14 Eylül 2009.
-
- “Artificial brain '10 years away'”, BBC News, (Çevrimiçi),
<http://news.bbc.co.uk/2/hi/technology/8164060.stm>, 14 Ağustos 2009.
-
- “Behind Artificial Intelligence, a Squadron of Bright Real People”, The New York Times, (Çevrimiçi),
<http://www.nytimes.com/2005/10/14/technology/14artificial.html>, 14 Ağustos 2009.
-
- “Evolving Robots Learn To Lie To Each Other”, (Çevrimiçi), <http://www.popsoci.com/scitech/article/2009-08/evolving-robots-learn-lie-hide-resources-each-other>, 29 Ağustos 2009.
-
- “Stanford Racing”, (Çevrimiçi),
<http://cs.stanford.edu/group/roadrunner//old/index.html>, 14 Ağustos 2009.

“Tower of Hanoi”, (Çevrimiçi),

<http://mazeworks.com/hanoi/index.htm>, 5 Ağustos 2009.