

사람이란 무엇인가

지성 · 인격 · 하나님의 형상

What Is a Human Being? Intelligence, Personhood, and the Image of God

— AI 문명 앞에 선 기독교 인간론 —

김종필

Elijah J F Kim, PhD

목차

Contents

서론	네 번의 여행, 그리고 마지막 물음.....	6
----	--------------------------	---

제 1 부 지식이란 무엇인가

제 1 장	데이터에서 지혜까지 — 그리고 그 너머.....	14
제 2 장	야다(yt) — 히브리 성경이 말하는 삶.....	17
제 3 장	아테네 · 예루살렘 · 실리콘밸리.....	19

제 2 부 도대체 AI 란 무엇인가

제 4 장	여덟 개의 동작과 네 개의 단계.....	22
제 5 장	왜 이토록 빨라지는가 — 재귀적 자기개선.....	27
제 6 장	지능이 몸을 얻는다 — Embodied Intelligence.....	32
제 7 장	과학 방법론의 혁명 — 두 개의 혁명.....	38

제 3 부 두 개의 지성

제 8 장	인간의 지능과 기계의 지능은 같은 종류인가.....	41
제 9 장	확인된 사실과 전망을 분리한다.....	46

제 4 부 인간을 넘어 인간을 만들려는 문명

제 10 장	THE GREAT HUMAN PROJECT — 일곱 개의 동사.....	53
제 11 장	치료에서 증강으로, 증강에서 진화로.....	64
제 12 장	교체되는 몸 — 이중이식과 뇌-컴퓨터 인터페이스.....	71
제 13 장	기억과 보존 — Reconstruction ≠ Resurrection.....	82
제 14 장	특이점 이후, 인류는 셋으로 갈라지는가.....	87

제 5 부 HOMO DEUS?

제 15 장	인간이 손대지 않은 영역이 있는가.....	95
제 16 장	생명을 조립하고 거울처럼 뒤집는다면.....	105
제 17 장	일곱 경계와 세 가지 오래된 욕망.....	111
제 18 장	네 낱말의 구별과 두 개의 인간론.....	120

제 6 부 인격이란 무엇인가

제 19 장	다섯 개념을 구별하라.....	127
제 20 장	능력으로 인격을 정의할 때 — 그 문은 양쪽으로 열린다.....	132

제 7 부 몸과 골렘

제 21 장	아담 · 골렘 · 휴머노이드.....	137
--------	----------------------	-----

제 8 부 하나님의 형상

제 22 장	첼렘 엘로힘 — IMAGE ≠ INTELLIGENCE.....	144
--------	------------------------------------	-----

제 23 장	질문을 바꾸어야 한다.....	149
--------	------------------	-----

제 9 부 AI 시대의 인간 형성

제 24 장	왜를 묻는 존재 — 그리고 더 큰 위험.....	151
--------	----------------------------	-----

제 25 장	안식과 형성 — 교회는 어떤 인간을 길러야 하는가	156
--------	-----------------------------------	-----

제 26 장	아담 · 골렘 · 로봇 · 그리스도.....	160
--------	--------------------------	-----

에필로그	기계가 지능적일수록, 인간은 더 깊이 인간이 되어야 한다.....	161
------	--------------------------------------	-----

참고문헌	Bibliography.....	163
------	-------------------	-----

서문

저자의 말

이 책은 「길이 지성이 되다」 강의 시리즈의 세 번째 책이다. 제 1 장 「길」이 인공지능을 문명사의 다섯 번째 문턱으로 진단했고, 제 2 장 「컴퓨터의 권세」가 그 문턱을 넘게 하는 힘이 무엇이며 누가 그것을 소유하는지를 물었다면, 이 책은 그 문명과 그 권세 앞에 서 있는 인간을 묻는다.

그리고 이 물음은 추상적이지 않다. AI가 글을 더 잘 쓰고, 더 많이 기억하고, 더 빨리 분석하고, 더 정확히 번역하고, 더 오래 일하고, 인간보다 더 잘 추론하는 시대가 온다면 — 도대체 인간은 무엇인가.

이 책의 중심 명제

본서는 하나의 명제 위에 서 있다. 그리고 그 명제를 미리 밝혀 둔다.

AI 시대의 가장 중요한 질문은 “기계가 어디까지 인간처럼 될 수 있는가”가 아니다.

“인간은 본래 무엇인가”이다.

그러므로 이 책은 인간이 AI보다 무엇을 더 잘하는지를 논증하지 않는다. 그 전선은 계속 후퇴하기 때문이다. 체스에서 졌을 때 우리는 바둑을 말했고, 바둑에서 졌을 때 언어와 창의성을 말했으며, 언어와 창작에서 놀라운 성과가 나오자 이해와 의식을 말한다. **매번 방어선을 뒤로 물릴 때마다 인간의 영역은 좁아진다.**

본서는 그 전선에서 싸우지 않는다. 대신 전제를 문제 삼는다 — 애초에 인간의 존엄이 능력에 근거한다는 전제가 옳은가. 그리고 이렇게 답한다.

인간의 존엄은 인간이 AI보다 더 똑똑하기 때문에 생기는 것이 아니다.

인간의 존엄은 인간이 하나님의 형상이기 때문에 존재한다.

이 책의 방법론

첫째, 관찰된 현실에서 출발한다. 본서의 모든 신학적 주장은 실증 위에 세워진다. 스탠퍼드 「AI 인덱스」, 국제 AI 안전보고서, OECD, 유네스코, 세계은행, 국제에너지기구, 그리고 이중의식과 뇌-컴퓨터 인터페이스와 오가노이드 지능의 최신 연구가 근거다. 정보와 분석 없이 해석학적 대안만 제시하는 것은 모래 위에 성을 세우는 일이기 때문이다.

둘째, 확인된 사실과 전망을 엄격히 분리한다. 실재하는 기술과 연구 중인 기술과 과학적으로 가능한 미래와 아직 사변적인 미래는 같은 확실성의 범주에 놓이지 않는다. 본서는 “AI는 절대로 의식을 가질 수 없다”고 단정하지도, “초지능이 오면 반드시 인격이 된다”고 말하지도 않는다. 현재 확인할 수 있는 것은 **능력(capability)의 증가이지 주관적 의식(subjective consciousness)의 존재가 아니다.**

셋째, 개념의 구별을 최우선으로 삼는다. 지능과 이해와 의식과 자의식과 인격은 다르다. 장수와 불멸과 부활과 영생은 다르다. 시뮬레이션과 경험은 다르다. 재구성과 부활은 다르다. **이 구별이 무너지면 신학이 아니라 기술 종교가 된다.**

넷째, 반론을 함께 신는다. 기술 결정론에 대한 사회구성주의의 반박, 확장된 마음 논제, 테넷과 블록의 의식 이론, 유대교 내부의 마이모니데스적 긴장을 숨기지 않는다. 자기 논증에 불리한 자료를 스스로 제시하고 답하는 것이 학문적 정직이기 때문이다.

이 책이 하지 않는 일

본서는 AI 활용법을 다루지 않는다. 설교문 작성, 성경 검색, 번역, 문서 작업은 이미 여러 매체가 잘 다루고 있다. 본서가 하려는 일은 다른 것이다. 인간이 지금 지성과 몸과 생명과 출생과 종(種)과 행성과 죽음의 경계를 하나씩 기술적 문제로 재정의하고 있다는 사실 앞에서, **성경의 가장 오래된 물음이 어떻게 가장 현대적인 물음이 되는가를** 보이려는 것이다.

김종필 (Elijah J F Kim, PhD)

서론

네 번의 여행, 그리고 마지막 물음

이 시리즈는 네 번의 여행으로 이루어진다. 그 여정을 한 줄로 압축하면 이렇다 — **문명 → 권세 → 형상 → 선교**.

제 1 장 「길」은 AI가 어디에 와 있는지를 물었고, 인공지능이 하나의 프로그램이 아니라 문명의 인프라가 되고 있음을 논증했다. 인류는 길과 문자와 인쇄와 네트워크라는 네 문턱을 넘어왔고 다섯 번째 문턱은 지성이다.¹ 제 2 장 「권세」는 그 문명을 움직이는 힘이 누구에게 있는지를 물었고, 컴퓨터가 자원에서 능력으로, 능력에서 플랫폼으로, 플랫폼에서 인프라로, 인프라에서 권한으로, 권한에서 권세로 이동하는 여섯 단계를 그렸다.²

그리고 이제 제 3 장이 묻는다. **그 권세 앞에 서 있는 인간은 누구인가.**

0. 서른 걸음이면 달에 닿는다

이 책을 여는 물음은 인간에 관한 것이 아니라 **숫자에 관한 것**이다. 그리고 그것이 이 책 전체를 이해하는 열쇠다.

한 걸음에 1 미터씩 걷는다고 하자. 서른 걸음을 걸으면 30 미터를 간다. 이것이 **선형적 성장(linear growth)**이다. 우리의 직관은 여기에 맞추어져 있다. 어제보다 오늘이 조금 낮고, 올해보다 내년이 조금 나으리라 기대한다. 그런데 걸음마다 거리가 두 배가 된다면 어떻게 되는가. 1 미터, 2 미터, 4 미터, 8 미터, 16 미터. 열 걸음이면 약 500 미터, 스무 걸음이면 약 500 킬로미터. 그리고 **서른 걸음이면 약 10 억 미터**,

¹ 줄고 「길이 지성이 되다」(초록나무) 참조. 앞의 네 문턱에서 “생각하는 자리”는 언제나 인간에게 있었다는 것, 그리고 다섯 번째 문턱만이 인간 능력의 확장을 넘어 행위자성 자체를 재구성한다는 것이 그 강의의 핵심 논지였다.

² 줄고 「컴퓨터의 권세: 새로운 경제 패권과 세계의 재편」(초록나무) 참조. 그 책의 핵심 명제는 “컴퓨터는 의존이 비대칭이 될 때 권세가 된다”이다.

곧 100 만 킬로미터다. 지구에서 달까지의 거리가 38 만 킬로미터이니 서른 걸음이면 달을 두 번 넘게 왕복한다.

같은 서른 걸음인데 결과는 30 미터와 100 만 킬로미터다. 3 만 배가 아니라 3 천만 배의 차이다. 그런데 인간의 뇌는 이 차이를 직관적으로 느끼지 못한다. **우리는 선형적 세계에서 진화했기 때문이다.** 사냥감이 도망가는 속도, 계절이 바뀌는 주기, 아이가 자라는 속도 — 이 모든 것이 선형적이었다. 지수적으로 증가하는 것을 일상에서 마주칠 일이 거의 없었다.

0-1. 앨버트 바틀렛의 경고

콜로라도 대학의 물리학자 앨버트 바틀렛은 이 문제를 평생의 주제로 삼았다. 그는 1969 년부터 40 년 가까이 “산술, 인구, 에너지”라는 제목의 강연을 1,700 회 이상 반복했고, 그 첫 문장은 언제나 같았다.

“인류의 가장 큰 결함은 지수함수를 이해하지 못하는 것이다.” — Albert A. Bartlett

바틀렛은 이 무능을 보여주기 위해 하나의 실험을 제시했다.³ 병 안에 세균 한 마리가 있다. 이 세균은 1 분마다 두 배로 늘어난다. 오전 11 시에 한 마리로 시작해서 정오에 병이 가득 찬다. 그렇다면 **병이 절반 찬 시각은 몇 시인가.** 답은 11 시 59 분이다.

이 대답이 주는 충격이 이 책의 출발점이다. **병 안의 세균이 “공간이 부족하다”고 느끼기 시작하는 것은 마지막 1 분이다.** 11 시 55 분에 병은 3 퍼센트만 찼다. 그 시점에 “곧 공간이 없어진다”고 경고하는 세균이 있다면 다른 세균들은 비웃을 것이다. 97 퍼센트가 비어 있는데 무슨 소리냐고. 그런데 5 분 뒤 모든 것이 끝난다.

³ Albert A. Bartlett, “Forgotten Fundamentals of the Energy Crisis,” American Journal of Physics 46, no. 9 (1978): 876–888. 바틀렛의 강연은 1969 년부터 2013 년 그가 세상을 떠날 때까지 계속되었으며, 유튜브에 공개된 녹화본은 수백만 회 재생되었다.

바틀렛은 여기에 하나를 더 물었다. 세균들이 마지막 순간에 병 세 개를 더 발견했다고 하자. 공간이 네 배가 된 것이다. 그러면 얼마나 더 버틸 수 있는가. **정확히 2 분이다.** 자원을 네 배로 늘려도 지수적 성장 앞에서는 2 분을 벌 뿐이다.

0-2. 체스판 위의 쌀알

같은 통찰이 훨씬 오래된 이야기에도 담겨 있다. 체스를 발명한 사람이 왕에게 상을 청했다. 체스판 첫 칸에 쌀 한 톨, 둘째 칸에 두 톨, 셋째 칸에 네 톨, 이렇게 예순네 칸을 채워 달라는 것이었다. 왕은 소박한 요청이라 여기고 허락했다.

체스판의 앞쪽 절반, 곧 서른두 칸까지는 쌀이 약 40 억 톨이다. 큰 발 하나 분량이다. 왕은 아직 여유롭다. 그런데 뒤쪽 절반에서 상황이 달라진다. 예순네 칸을 모두 채우려면 **약 1,845 경 톨**이 필요하다. 인류가 수천 년 동안 생산한 쌀을 다 합쳐도 부족한 양이다.

MIT 의 에릭 브리놀프슨과 앤드루 맥아피는 이 이야기에서 **“체스판의 후반부(the second half of the chessboard)”**라는 개념을 가져왔다.⁴ 그들의 논지는 이렇다. 지수적 성장의 전반부에서는 변화가 눈에 띄지 않는다. 숫자가 두 배가 되어도 절대량이 작기 때문이다. 그러나 후반부에 들어서면 한 걸음이 이전의 모든 걸음을 합친 것보다 크다. **그리고 그들은 정보기술이 이제 체스판의 후반부에 들어섰다고 주장했다.**

0-3. 무어의 법칙과 그 너머

이것이 추상적 비유가 아님을 보여주는 구체적 사실이 있다. 1965 년 인텔의 공동 창업자 고든 무어는 집적회로 위의 트랜지스터 수가 약

⁴Erik Brynjolfsson and Andrew McAfee, *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies* (New York: W. W. Norton, 2014). 이들은 이 개념을 레이 커즈와일에게서 빌려왔음을 밝히고 있다.

2 년마다 두 배가 된다고 관찰했다. 이것이 **무어의 법칙(Moore's Law)**이다.⁵ 이 법칙은 자연법칙이 아니라 경험적 관찰이었으나 반세기 넘게 대체로 유지되었다. 1971 년 인텔 4004 에는 트랜지스터가 2,300 개였다. 오늘의 첨단 프로세서에는 수백억 개가 들어간다. **약 5 천만 배의 증가다.**

그런데 AI 의 발전은 여기서 한 걸음 더 나아간다. 스탠퍼드 「AI 인텔스」와 여러 분석에 따르면 최첨단 AI 모델의 훈련에 투입되는 계산량은 무어의 법칙보다 훨씬 빠르게 증가해 왔다. 오픈AI 의 2018 년 분석은 2012 년 이후 대형 AI 훈련 실행의 계산량이 **약 3.4 개월마다 두 배**로 늘었다고 보고했다.⁶ 2 년이 아니라 3.4 개월이다. 같은 기간에 무어의 법칙이 7 배를 만들 때 AI 계산량은 30 만 배를 만든 것이다.

0-4. 로켓 그래프 — 왜 우리는 알아채지 못하는가

지수함수의 그래프를 그려 보면 왜 우리가 이 변화를 인지하지 못하는지 알 수 있다. 초반부는 거의 수평이다. 바닥에 붙어 있는 것처럼 보인다. 그러다 어느 지점에서 갑자기 위로 꺾이며 수직에 가깝게 치솟는다. **로켓 발사 궤적처럼 보이기 때문에 흔히 “로켓 그래프”라 부른다.**

그런데 여기에 착시가 있다. **그래프의 어느 지점에서든 성장률은 동일하다.** 곡선이 “꺾이는” 지점 같은 것은 사실 존재하지 않는다. 매 순간 두 배씩 늘어나고 있을 뿐이다. 꺾여 보이는 것은 우리가 그것을 선형적 눈금으로 보기 때문이다. 이 사실이 중요한 이유가 있다. **“아직은 괜찮다”는 판단이 언제나 틀릴 수 있다는 뜻이기 때문이다.** 11 시 55 분의 세균에게 병은 97 퍼센트 비어 있었다.

⁵Gordon E. Moore, “Cramming More Components onto Integrated Circuits,” Electronics 38, no. 8 (1965): 114–117. 무어는 처음에 1 년마다 두 배로 예측했다가 1975 년에 2 년으로 수정했다.

⁶Dario Amodei and Danny Hernandez, “AI and Compute,” OpenAI Blog (2018). 이 분석에 따르면 2012 년 AlexNet 에서 2018 년 AlphaGo Zero 까지 계산량은 30 만 배 이상 증가했다. 무어의 법칙대로였다면 같은 기간에 7 배 증가에 그쳤을 것이다.

레이 커즈와일은 이 현상을 “**수익 가속의 법칙(the Law of Accelerating Returns)**”이라 불렀다.⁷ 그의 핵심 주장은 두 가지다. 첫째, 기술의 진보는 선형이 아니라 지수적이다. 둘째, **진보의 속도 자체도 가속된다**. 왜냐하면 더 나은 도구가 더 나은 도구를 더 빨리 만들기 때문이다. 그는 이를 이렇게 표현했다 — “우리는 21 세기에 100 년의 진보를 경험하지 않을 것이다. 현재의 속도로 계산하면 2 만년에 해당하는 진보를 경험할 것이다.”

0-5. 그리고 알파고가 보여준 것

추상적 논의를 구체적 장면으로 옮겨 보자. 2016 년 3 월 서울에서 이세돌 9 단과 알파고의 대국이 열렸다. 대국 전 전문가들의 예측은 대체로 이세돌의 우세였다. 바둑은 경우의 수가 우주의 원자 수보다 많고, 직관과 형세 판단이 필요하며, 기계가 정복하려면 10 년은 더 걸릴 것이라고 여겨졌다.

결과는 4 대 1 이었다. 그리고 더 충격적인 것은 **제 2 국 37 수**였다. 알파고는 다섯 번째 줄에 어깨짚음을 두었는데, 이것은 수백 년 축적된 기보에 없는 수였다. 해설자들은 처음에 실수라고 생각했다. **그러나 그 수가 승부를 갈랐다**. 인간이 수천 년 동안 두지 않았던 수를 기계가 둔 것이다.

그리고 이듬해 더 놀라운 일이 일어났다. **알파고 제로(AlphaGo Zero)**는 인간의 기보를 전혀 학습하지 않았다. 규칙만 알려 주고 자기 자신과 두게 했을 뿐이다. 그런데 **사흘 만에 이세돌을 이긴 버전을 100 대 0 으로 이겼다**.⁸ 인간의 지식 없이, 오히려 인간의 지식이 없었기

⁷Ray Kurzweil, “The Law of Accelerating Returns,” 2001; idem, The Singularity Is Near (New York: Viking, 2005). 커즈와일의 예측 정확도에 대해서는 상당한 논쟁이 있으며, 본서는 그의 시간표를 그대로 채택하지 않는다.

⁸David Silver et al., “Mastering the Game of Go without Human Knowledge,” Nature 550 (2017): 354-359. 알파고 제로는 40 일 만에 이전의 모든 버전을 능가했으며, 인간이 발견한 정석의 상당수를 스스로 재발견한 뒤 그것을 버리고 다른 수를 선택했다.

때문에 더 강해진 것이다. 이것이 본서 제 5 장에서 다룰 재귀적 자기개선의 첫 번째 극적인 사례였다.

0-6. 그러므로 이 책의 전제

이 책은 두 가지 전제 위에 서 있다.

첫째, 우리는 지금 무슨 일이 일어나고 있는지 직관적으로 파악하지 못한다. 이것은 지능이나 관심의 문제가 아니라 인지 구조의 문제다. 인간의 뇌는 지수적 변화를 느끼도록 만들어지지 않았다. 그래서 “아직 멀었다”는 판단과 “이미 늦었다”는 판단이 불과 몇 년 사이에 뒤바뀐다.

둘째, 그럼에도 이 변화를 예측이 아니라 관찰로 다룰 수 있다. 본서는 미래를 점치지 않는다. 지금 실험실에서 무엇이 일어나고 있는지, 어떤 논문이 발표되었는지, 어떤 임상시험이 진행 중인지를 확인하고 그 위에서 신학적 판단을 세운다. **확인된 사실과 전망을 분리하는 것이 이 책의 방법론이며, 그것이 지수함수 앞에서 신중해지는 유일한 길이다.**

그러므로 이 책은 이렇게 시작한다. 인류는 지금 로켓의 궤적 위에 있으나 그 사실을 잘 인지하지 못하고 있다. 그리고 그 로켓이 향하는 곳에 인간 자신이 있다.

제 3 장의 핵심어는 “인간”이 아니라 “형상”이다

본서의 제목은 「사람이란 무엇인가」이지만, 핵심어는 인간이 아니라 형상이다. 왜냐하면 “인간이란 무엇인가”라는 물음에 능력으로 답하는 순간 우리는 이미 진 게임을 시작하기 때문이다. 인간을 지능으로 정의하면 더 지능적인 존재 앞에서 무너지고, 생산성으로 정의하면 더 생산적인 존재 앞에서 무너진다.

그러므로 물음의 형태 자체를 바꾸어야 한다.

“인간은 AI 보다 무엇을 잘하는가?” → “인간은 무엇이기 때문에 인간인가?”

앞의 물음은 비교의 물음이고 뒤의 물음은 근거의 물음이다. 그리고 성경은 뒤의 물음에 답한다.

이 책의 아홉 걸음

본서는 아홉 부로 이루어진다. 그 순서에는 논리가 있다.

제 1 부는 지식을 묻는다. AI 를 논하기 전에 “안다는 것이 무엇인가”를 먼저 물어야 하기 때문이다. 데이터에서 지혜까지의 현대적 도식과 히브리 성경의 야다(yad)를 나란히 놓는다. **제 2 부는 AI 를 묻는다.** 대상을 정확히 규정하지 않으면 인간론도 정확할 수 없다. 여덟 동작과 네 단계, 재귀적 자기개선, 체화된 지능, 그리고 과학 방법론의 혁명을 다룬다.

제 3 부는 두 지성을 비교한다. 인간의 지능과 기계의 지능이 같은 종류인지를 열세 항목으로 검토하고, 확인된 사실과 전망을 분리한다.

제 4 부는 트랜스휴먼의 지평을 다룬다. 치료에서 증강으로, 증강에서 진화로 넘어가는 경계와 교체되는 몸을 살핀다. **제 5 부는 그 모든 것을 문명의 층위로 끌어올린다.** 생명의 조립과 거울 생명과 일곱 경계와 세 가지 오래된 욕망이 여기 있다.

제 6 부는 인격을 묻고, 제 7 부는 몸과 골렘을 다루며, 제 8 부는 하나님의 형상에 이른다. 그리고 **제 9 부는 인간 형성으로 닫는다.** AI 시대에 교회는 어떤 인간을 길러야 하는가.

그리고 이 아홉 걸음이 하나의 원을 그린다. 제 1 부에서 “안다는 것이 무엇인가”로 시작한 물음이 마지막에 요한복음으로 돌아오기 때문이다 — “영생은 곧 유일하신 참 하나님과 그가 보내신 자 예수 그리스도를 아는 것이니이다”(요 17:3).

제 1 부

지식이란 무엇인가

What Is Knowledge? — From Data to Wisdom, and Beyond

제 1 장

데이터에서 지혜까지 — 그리고 그 너머

AI 를 논하기 전에 먼저 물어야 할 것이 있다. **안다는 것이 무엇인가**. 이 물음을 건너뛰면 인간론도 정확할 수 없다. AI 가 “안다”고 말할 때와 성경이 “안다”고 말할 때 같은 낱말이 전혀 다른 것을 가리키기 때문이다.

1. DIKW — 현대 정보문명의 자기 이해

현대 정보학은 앎을 네 층위로 구별해 왔다. **데이터(Data)**는 가공되지 않은 사실과 신호다. **정보(Information)**는 맥락이 부여되어 의미를 갖게 된 데이터다. **지식(Knowledge)**은 조직되고 검증되어 사용 가능해진 정보다. 그리고 **지혜(Wisdom)**는 옳게 판단하고 옳게 살아내는 능력이다.

이른바 DIKW 모델은 러셀 애코프 계열의 도식으로 널리 알려져 있으나, 단계 사이의 전환과 각 층위의 정의에 대해서는 상당한 학문적 논쟁이 있다.⁹ 그러므로 본서는 이 도식을 진리의 지도가 아니라 **현대 정보문명이 스스로를 이해하는 방식**으로 사용한다.

그런데 이 도식을 AI 에 적용하면 흥미로운 물음이 생긴다. AI 가 데이터를 지배하면 정보를 지배할 수 있다. 정보를 통합하면 지식을 생산할 수 있다. 그러면 —

지식이 충분히 많아지면 지혜가 되는가?

2. 지식은 폭발하는데 지혜는 메마른다

⁹Russell L. Ackoff, “From Data to Wisdom,” Journal of Applied Systems Analysis 16 (1989): 3-9. 이 위계의 기원과 타당성에 관한 비판적 검토로는 Martin Frické, “The Knowledge Pyramid: A Critique of the DIKW Hierarchy,” Journal of Information Science 35, no. 2 (2009): 131-142 참조. 본서는 이를 절대적 진리가 아니라 현대 정보문명이 스스로를 이해하는 방식을 보여주는 모델로 사용한다.

제 1 장에서 우리는 하나의 역설을 지적했다. 인류가 생산하는 데이터는 기하급수적으로 증가하고, 누구나 몇 초 만에 세계의 지식에 닿을 수 있으며, 요약과 번역과 분석이 즉시 이루어진다. 세 지표가 모두 위를 향한다. 그런데 네 번째 물음 앞에서 화살표가 멈춘다 — **우리는 더 지혜로워졌는가.**

정보가 넘칠수록 오히려 의미를 잃어버리는 역설이 있다. 지식은 축적되지만 분별은 자라지 않고, 답은 넘치지만 물음은 알아진다. T. S. 엘리엇은 1934 년에 이미 이렇게 물었다 — 삶에서 잃어버린 지혜는 어디에 있는가, 지식에서 잃어버린 지혜는 어디에 있는가, 정보에서 잃어버린 지식은 어디에 있는가.¹⁰ 인터넷이 없던 시대의 물음이다. 그렇다면 정보가 무한해진 시대의 물음은 훨씬 절박해야 한다.

지혜는 다운로드되지 않는다.

3. 인간 기억의 외부화

서양 문명은 지식을 몸 밖으로 계속 옮겨 왔다. 구전 기억에서 문자로, 문자에서 도서관으로, 도서관에서 인쇄로, 인쇄에서 백과사전으로, 그리고 컴퓨터와 데이터베이스와 인터넷 검색과 빅데이터로.

각 단계가 무엇을 외부화했는지 보면 흐름이 보인다. 책은 기억의 외부화였다. 도서관은 집단 기억의 외부화였다. 검색은 인출(retrieval)의 외부화였다. 그런데 대규모 언어모델은 단순 저장이 아니라 요약과 종합과 생성과 추론의 외부화를 시작했다. 그리고 에이전틱 AI 는 한 걸음 더 나아가 판단과 행동의 일부까지 외부화한다.

BODY → MACHINE → NETWORK → INTELLIGENCE → AUTONOMY?

여기서 본서의 첫 번째 경고가 나온다. 인간은 지식을 증가시킨 것만이 아니라 지식을 인간의 뇌 밖으로 계속 외부화해 왔다. 그리고 외부화

¹⁰T. S. Eliot, Choruses from “The Rock” (1934). 이 구절은 훗날 DIKW 위계의 문학적 선구로 자주 인용된다. 엘리엇이 인터넷 이전, 심지어 텔레비전 이전에 이 물음을 던졌다는 사실이 중요하다.

자체는 인간의 오랜 전략이며 그 자체로 쇠퇴가 아니다. 문제는 외부화의 대상이 기억에서 판단으로 옮겨 갈 때, 그리고 외부화한 것을 다시 호출하고 검증할 능력까지 잃을 때 발생한다. 이 논의는 제 9 부에서 다시 다룰 것이다.

4. 지식에는 끝이 있는가

“AI가 언젠가 모든 지식을 다 알게 되는가”라는 물음을 자주 듣는다. 이 물음에 답하려면 세 가지를 구별해야 한다.

첫째, 실제 세계(Reality)다. 우주와 생명과 인간 사회는 계속 변화한다. **둘째, 실재에 관한 가능한 정보(Possible Information)다.** 새 데이터와 새 명제가 계속 생성될 수 있다. **셋째, 아는 자가 소유한 지식(Knowledge Possessed)이다.** 어떤 지성도 전체를 담을 수 없다.

그러므로 “모든 지식을 다 배웠다”는 종착점은 쉽게 정의되지 않는다. 특히 AI가 과학 연구를 수행한다면 기존 지식을 암기하는 데서 끝나지 않고 새 가설 → 실험 → 데이터 → 발견 → 새로운 질문을 반복하게 된다. **지식의 총량이 아니라 지식 생산의 속도가 문제가 되는 것이다.** 그러므로 물음은 이렇게 바뀌어야 한다.

“AI가 모든 것을 알게 되는가?”가 아니라

“AI가 새로운 지식을 생산하는 속도가 인간의 이해 속도를 넘어서는가?”

제 2 장

야다(יָדָא) — 히브리 성경이 말하는 앎

이제 전혀 다른 지식 개념을 만난다. 히브리 성경에서 “안다”는 것은 무엇인가.

1. 야다는 정보 저장이 아니다

히브리어 **야다(יָדָא)**는 경험하고, 관계하고, 인식하고, 삶으로 알고, 언약을 지키는 것을 아우른다. 이 낱말은 머릿속에 정보를 저장한다는 의미로 환원되지 않는다.

가장 놀라운 용례가 창세기 4 장 1 절이다. “아담이 그의 아내 하와와 동침하매”로 번역된 낱말이 바로 이 야다다. 성경에서 안다는 것은 정보를 습득하는 것이 아니라 **관계에 들어가는 것이다**. 그리고 호세아 6 장 6 절에서 “번제보다 하나님을 아는 것을 원하노라”고 할 때, 그 앎은 신학 지식이 아니라 언약적 신실함이다.¹¹

2. 다아트 · 호크마 · 비나

히브리 성경은 앎을 세 낱말로 세분한다. **다아트(דַּעַת)**는 관계적 앎을 포함하는 지식이다. **호크마(חִכְמָה)**는 지혜를 뜻하되 삶을 옳게 살아내는 기술이자 태도이며, 성막을 만든 브살렐의 손재주도 호크마라 불린다(출 31:3). 기능과 경건이 함께 있는 것이다. **비나(בִּינָה)**는 분별과 통찰, 곧 갈라서 보는 능력이다.

그리고 이 셋의 뿌리에 하나의 명제가 있다.

“여호와를 경외하는 것이 지식의 근본이거늘”(잠 1:7)

¹¹히브리어 야다(יָדָא)의 의미 범위에 관하여는 Ludwig Koehler and Walter Baumgartner, The Hebrew and Aramaic Lexicon of the Old Testament, 해당 항목; Hans Walter Wolff, Anthropology of the Old Testament, trans. Margaret Kohl (Philadelphia: Fortress, 1974) 참조.

지식의 시작은 데이터가 아니라 경외다. 아브라함 요수아 헤셸은 인간을 답을 가진 존재가 아니라 경이 앞에 서는 존재로 이해했다. 그가 말한 “근원적 경이(radical amazement)”는 세계를 당연하게 여기지 않는 태도이며, 헤셸에게 이것이 모든 참된 앎과 예배의 시작이다.¹²

3. 두 개의 지식 구조

두 전통을 나란히 놓으면 구조의 차이가 드러난다.

Data → Information → Knowledge → Prediction → Control

Knowledge → Relationship → Discernment → Wisdom → Faithful Life

왼쪽은 통제로 끝나고 오른쪽은 신실한 삶으로 끝난다. 이 차이가 결정적이다. AI에게 knowledge는 원칙적으로 processing의 문제일 수 있다. 그러나 성경에서 knowledge는 relationship과 faithfulness의 문제이기도 하다.

그리고 여기서 본서의 마지막 장을 미리 예고해 둔다. 요한복음은 영생을 이렇게 정의한다 — “영생은 곧 유일하신 참 하나님과 그가 보내신 자 예수 그리스도를 아는 것이니이다”(요 17:3). **영생의 정의가 “아는 것(γινώσκωσιν)”이다.** 이 강의를 열었던 물음이 마지막에 그대로 돌아오는 것이다.

¹²Abraham Joshua Heschel, *God in Search of Man: A Philosophy of Judaism* (New York: Farrar, Straus and Cudahy, 1955). 헤셸의 “근원적 경이”는 지식이 축적될수록 신비가 사라진다는 근대적 가정에 대한 반론이다.

제 3 장

아테네 · 예루살렘 · 실리콘밸리

세 도시가 세 가지 삶의 방식을 상징한다. 그리고 오늘 우리는 세 번째 도시의 삶 안에서 살고 있다.

1. 예루살렘 — 관계적 삶

히브리적 삶은 관계적이고 언약적이고 실천적이다. “안다”는 것은 관계에 들어가는 것이며, 삶의 진정성은 그 사람이 어떻게 사는가로 확인된다. 그러므로 히브리 성경에서 하나님을 아는 것과 하나님께 순종하는 것은 분리되지 않는다.

2. 아테네 — 이성적 파악

헬라 전통에서 삶은 로고스이고 이성이며 진리의 파악이다. 스토아 인식론은 표상(phantasia)에 대한 동의(assent)와 확실하게 파악된 인식(katalēpsis)을 구별하려 했다. 곧 **확실한 인식과 단순한 의견을 가르치는 것이** 그들의 과제였다.¹³

3. 실리콘밸리 — 예측 가능성

그리고 오늘의 삶이 있다. 데이터와 패턴과 예측과 최적화의 삶이다. 이 전통에서 “안다”는 것은 무엇을 뜻하는가. 가장 날카롭게 정식화하면 이렇다.

“안다”는 것은 예측할 수 있다는 것이다.

이것은 강력하다. 그리고 실제로 유용하다. 단백질 구조를 예측하고 질병을 예측하고 날씨를 예측하는 능력이 인류에게 준 유익은 막대하다.

¹³ 스토아 인식론의 katalēpsis 개념에 관하여는 A. A. Long and D. N. Sedley, The Hellenistic Philosophers, vol. 1 (Cambridge: Cambridge University Press, 1987), sec. 40 참조.

그러나 이 정의에는 빠진 것이 있다. **예측할 수 있는 것과 이해하는 것과 관계하는 것은 같지 않다.**

4. 이분법을 경계한다

다만 “히브리는 관계, 헬라는 이성”이라는 도식적 이분법은 피해야 한다. 히브리 전통에도 명제적 지식이 있고 헬라 철학에도 삶과 덕의 차원이 있다. 아리스토텔레스의 프로네시스(실천적 지혜)는 오히려 히브리의 호크마에 가깝다.

그럼에도 강조점의 차이를 보는 것은 매우 유익하다. 왜냐하면 **AI 문명은 세 번째 도시의 앞을 압도적으로 강화하면서 첫 번째 도시의 앞을 조용히 약화시킬 수 있기 때문이다.** 예측은 폭발적으로 늘어나는데 관계는 알아지고, 정보는 무한해지는데 경외는 사라진다. 이 책의 나머지가 그 위험을 추적한다.

제 2 부

도대체 AI란 무엇인가

What Is AI, Really? — From Chatbot to Autonomous System

제 4 장

여덟 개의 동작과 네 개의 단계

인간론으로 들어가기 전에 대상을 정확히 규정해야 한다. AI 를 챗봇 정도로 이해하면 인간론도 그 수준에 머물기 때문이다.

0. AI 라는 낱말은 어디서 왔는가

“인공지능”이라는 낱말은 1955 년 여름에 만들어졌다. 다트머스 대학의 젊은 수학자 **존 매카시(John McCarthy)**가 록펠러 재단에 연구비를 신청하면서 쓴 제안서에 이 표현이 처음 등장한다. 마빈 민스키, 너새니얼 로체스터, 클로드 새넌이 공동 제안자였다.¹⁴ 그리고 이듬해 1956 년 여름 다트머스에서 두 달간의 워크숍이 열렸다. **이것이 인공지능이라는 학문의 시작이다.**

제안서의 핵심 문장은 이랬다 — “학습의 모든 측면과 지능의 다른 모든 특성은 원리적으로 매우 정확하게 기술될 수 있어서 기계가 그것을 모사하도록 만들 수 있다는 추측에 근거하여 연구를 진행한다.” **여기에 이 학문의 근본 전제가 들어 있다.** 지능은 기술될 수 있고, 기술될 수 있으면 모사될 수 있다는 것이다. 이 전제가 옳은지가 본서 제 8 장의 주제다.

그런데 왜 하필 “인공지능”이었는가. 매카시는 훗날 이렇게 회고했다. 당시 이미 사용되던 **“사이버네틱스(cybernetics)”**라는 용어를 피하고 싶었다는 것이다. 노버트 위너가 주도하던 그 분야와 구별되는 새로운 학문임을 분명히 하기 위해 일부러 새 이름을 지었다. **곧 “인공지능”이라는 낱말 자체가 학문적 독립 선언이었다.**

¹⁴John McCarthy, Marvin L. Minsky, Nathaniel Rochester, and Claude E. Shannon, “A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence,” August 31, 1955. 이 제안서는 인공지능이라는 학문 분야의 출생증명서로 여겨진다.

0-1. 어원 — “인공”과 “지능”

낱말을 나누어 보면 의미가 선명해진다. **Artificial** 은 라틴어 **artificialis** 에서 왔고, 이것은 **ars(기술, 예술)**와 **facere(만들다)**의 결합이다. 곧 “**기술로 만들어진 것**”이다. 여기서 중요한 것은 이 낱말이 “가짜”를 뜻하지 않는다는 점이다. 인공 심장은 가짜 심장이 아니라 만들어진 심장이다.

Intelligence 는 라틴어 **intelligere** 에서 왔다. 이것은 **inter(사이에)**와 **legere(읽다, 고르다, 모으다)**의 결합이다. 곧 “**사이에서 읽어 내는 것, 여럿 가운데서 골라내는 것**”이다. 어원이 알려 주는 바가 놀랍다. 지능은 정보를 저장하는 능력이 아니라 **구별하고 선택하는 능력**이다. 그리고 이것은 히브리어 비나(בִּינָה)가 “갈라서 보는 것”을 뜻하는 것과 정확히 통한다.

그러므로 “인공지능”의 어원적 의미는 “**기술로 만들어진, 사이에서 읽어 내는 능력**”이다. 자동화가 아니다. 정해진 절차를 반복하는 것은 어원적으로도 지능이 아니다. 지능은 **정해지지 않은 상황에서 구별하고 선택하는 것이다.**

0-2. 자동화와 지능은 다르다

이 구별이 실천적으로 왜 중요한가. 많은 사람이 AI 를 “매우 빠른 자동화”로 이해하기 때문이다. 그러나 둘은 범주가 다르다.

자동화(automation)는 인간이 정한 규칙을 기계가 반복 수행하는 것이다. 세탁기는 정해진 순서로 물을 채우고 돌리고 뺀다. 상황이 바뀌어도 같은 일을 한다. 규칙 바깥의 상황을 만나면 멈추거나 오작동한다. **결정하는 주체는 규칙을 만든 인간이다.**

지능(intelligence)은 정해지지 않은 상황에서 판단하는 것이다. 무엇이 중요한지 골라내고, 패턴을 찾고, 새로운 상황에 기존 학습을 적용한다.

결정의 일부가 시스템 안으로 들어온다. 그리고 이것이 자동화와 AI 를 가르치는 결정적 차이이다.

구체적 예를 들자. 공장의 조립 로봇이 정해진 위치에서 정해진 부품을 집는 것은 자동화다. 그런데 흐트러진 상자 안에서 부품을 찾아 집는 것은 지능이다. 어떤 것이 목표 부품인지, 어느 각도로 집어야 하는지를 그 자리에서 판단해야 하기 때문이다. **전자는 20 세기 기술이고 후자는 최근에야 가능해졌다.**

0-3. 학자들의 정의 — 다섯 갈래

그렇다면 인공지능을 학문적으로 어떻게 정의하는가. 스튜어트 러셀과 피터 노빅의 표준 교과서는 정의를 네 갈래로 정리한다.¹⁵

① **인간처럼 생각하는 시스템(thinking humanly).** 인지과학의 접근으로, 인간의 사고 과정 자체를 모사하려 한다. ② **인간처럼 행동하는 시스템(acting humanly).** 튜링 테스트의 접근이다. 내면이 어떻든 인간과 구별되지 않게 행동하면 지능이라 본다. ③ **합리적으로 생각하는 시스템(thinking rationally).** 아리스토텔레스의 삼단논법에서 이어지는 논리학의 전통이다. ④ **합리적으로 행동하는 시스템(acting rationally).** 주어진 목표를 최선으로 달성하는 **합리적 에이전트(rational agent)**로 정의한다.

러셀과 노빅은 네 번째를 채택한다. 그 이유가 흥미롭다. **인간을 기준으로 삼으면 인간의 오류까지 모사해야 하고, 논리만을 기준으로 삼으면 불확실한 현실 세계를 다루지 못하기 때문이다.** 그런데 러셀은 최근 이 정의 자체를 수정해야 한다고 주장한다. “주어진 목표를 최선으로 달성하는 기계”라는 정의가 위험하다는 것이다. 목표가 잘못 설정되면 그 기계는 잘못된 목표를 완벽하게 달성한다. 그래서 그는

¹⁵Stuart Russell and Peter Norvig, Artificial Intelligence: A Modern Approach, 4th ed. (Hoboken, NJ: Pearson, 2021), chap. 1. 이 교과서는 전 세계 1,500 개 이상의 대학에서 사용되는 인공지능 분야의 표준 문헌이다.

“인간의 선호를 만족시키되 그 선호가 무엇인지 확신하지 않는 기계”라는 대안을 제시했다.¹⁶

여기에 다섯 번째 정의를 더할 수 있다. 존 매카시 자신의 정의다. 그는 이렇게 말했다 — “인공지능은 지능적인 기계, 특히 지능적인 컴퓨터 프로그램을 만드는 과학과 공학이다. 인간의 지능을 이해하기 위해 컴퓨터를 사용하는 것과 관련이 있으나, 생물학적으로 관찰 가능한 방법에만 스스로를 국한할 필요는 없다.”¹⁷ 마지막 구절이 중요하다. 비행기가 새처럼 날개를 펴덕이지 않듯, 인공지능도 인간의 뇌를 그대로 모사할 필요가 없다는 것이다.

0-4. 좁은 AI, 범용 AI, 초지능

그리고 능력의 범위에 따라 세 층위를 구별한다. 좁은 AI(ANI, Artificial Narrow Intelligence)는 특정 과제에서 뛰어나다. 바둑을 두거나 얼굴을 인식하거나 번역한다. 오늘 존재하는 모든 AI가 여기에 속한다. 알파고는 바둑에서 세계 최강이지만 체스 규칙을 알려 주어도 두지 못한다.

범용 AI(AGI, Artificial General Intelligence)는 인간이 할 수 있는 지적 과제 전반을 수행한다. 그런데 여기에 근본적 문제가 있다. AGI에는 합의된 정의도 판정 기준도 없다. 어떤 이는 튜링 테스트를 기준으로 삼고, 어떤 이는 “대부분의 경제적으로 가치 있는 일을 인간보다 잘하는 것”이라 정의하며, 어떤 이는 새로운 영역을 스스로 학습하는 능력을 든다. 그래서 “AGI가 언제 오는가”라는 물음은 정의에 따라 답이 완전히 달라진다.

¹⁶ Stuart Russell, *Human Compatible: Artificial Intelligence and the Problem of Control* (New York: Viking, 2019). 러셀의 제안은 기계가 자기 목표에 확신을 갖지 않도록 설계함으로써 인간의 개입을 허용하게 만드는 것이다.

¹⁷ John McCarthy, “What Is Artificial Intelligence?” Stanford University, 2007. 매카시는 여기서 “지능”을 “세계에서 목표를 달성하는 능력의 계산적 부분”이라 정의했다.

초지능(ASI, Artificial Superintelligence)은 사실상 모든 영역에서 최고의 인간을 능가하는 지능이다. 닉 보스트롬은 이를 “과학적 창의성, 일반적 지혜, 사회적 기술을 포함해 관심 있는 거의 모든 영역에서 인간의 인지 능력을 크게 초월하는 지성”이라 정의했다.¹⁸ **그리고 ASI 와 특이점은 가설적 미래 개념이다.** 본서는 이 구별을 끝까지 지킨다.

1. AI 가 지금 하는 여덟 가지

과거의 기계는 인간이 지시한 것을 수행했다. 오늘의 AI 는 다음 여덟 가지를 한다.

인지하고 · 배우고 · 기억하고 · 추론하고 · 생성하고 · 계획하고 · 결정하고 · 행동한다

앞의 다섯 — 인지·학습·기억·추론·생성 — 은 이미 일상이 되었다. 그런데 뒤의 셋이 지금 빠르게 열리고 있다. **계획하고(Plan), 결정하고(Decide), 행동한다(Act).** 그리고 바로 여기서 도구가 행위자로 넘어간다.

2. 네 개의 발전 단계

AI 의 발전은 네 단계로 요약할 수 있다. ① **분석하는 AI**(Analytical AI)는 데이터를 분석한다. ② **만들어 내는 AI**(Generative AI)는 텍스트와 이미지와 코드를 생성한다. ③ **추론하는 AI**(Reasoning AI)는 단계를 밟아 사고하고 도구를 사용한다. 그리고 ④ **행동하는 AI**(Agentic AI)는 목표를 받아 스스로 과정을 수행한다.

Chatbot → Copilot → Agent → Autonomous System

¹⁸Nick Bostrom, Superintelligence: Paths, Dangers, Strategies (Oxford: Oxford University Press, 2014). 보스트롬은 초지능에 이르는 경로를 인공지능·전뇌·에멀레이션·생물학적 인지 향상·뇌-기계 인터페이스·네트워크 조직의 다섯으로 구별한다.

대화하는 도구에서 함께 일하는 조력자로, 조력자에서 스스로 수행하는 행위자로, 그리고 자율 시스템으로. 이 이동에서 결정적인 것은 성능의 향상이 아니다. **질문 자체가 달라진다는 것이다.**

“무엇을 할 수 있는가” → “누가 결정하는가”

앞의 물음은 기술의 물음이고 뒤의 물음은 권한과 책임의 물음이다. 제 2 장 「컴퓨터의 권세」가 바로 이 지점을 다루었고, 본서 제 6 부의 인격 논의도 여기서 시작한다.

제 5 장

왜 이토록 빨라지는가 — 재귀적 자기개선

제 1 강에서 한마디로 지나간 개념을 이제 본격적으로 펼친다. **재귀적 자기개선(Recursive Self-Improvement)**이다.

0. 재귀(recursion)란 무엇인가

“재귀적 자기개선”을 이해하려면 먼저 **재귀(recursion)**라는 낱말부터 보아야 한다.

라틴어 **recurrere** 는 **re(다시)**와 **currere(달리다)**의 결합으로 “되돌아 달리다”를 뜻한다. 수학과 컴퓨터과학에서 재귀는 **어떤 것을 정의할 때 그 자신을 사용하는 것이다**. 계승(factorial)이 대표적이다. n 의 계승은 “ n 곱하기 $(n-1)$ 의 계승”으로 정의된다. 자기 자신을 부르는 것이다.

일상에서도 재귀를 볼 수 있다. 두 개의 거울을 마주 놓으면 상(像) 안에 상이 있고 그 안에 또 상이 있다. 러시아 인형 마트료시카도 그렇다. 그런데 재귀의 진짜 힘은 **출력이 다시 입력이 될 때** 나타난다. 이때 결과가 반복될 때마다 증폭되기 때문이다. 마이크를 스피커에 가까이 대면 작은 소리가 순식간에 귀를 찢는 굉음이 되는 것도 같은 원리다.

0-1. 자기개선이 재귀가 될 때

이제 결합해 보자. **재귀적 자기개선(Recursive Self-Improvement, RSI)**은 시스템이 자기 자신을 개선하고, 개선된 시스템이 다시 자기를 개선하며, 그 개선 능력 자체가 매번 향상되는 과정이다.

이것을 처음 명확히 정식화한 사람은 영국의 수학자 **어빙 존 굿(I. J. Good)**이다. 그는 앨런 튜링과 함께 블레츨리 파크에서 독일 암호를 해독했던 인물이며, 1965년에 이렇게 썼다.

“초지능 기계를 아무리 영리한 인간의 모든 지적 활동을 능가하는 기계라 정의하자. 기계 설계는 그 지적 활동 가운데 하나이므로, 초지능 기계는 더

나은 기계를 설계할 수 있다. 그러면 의심할 여지 없이 ‘지능 폭발’이 일어나고 인간의 지능은 훨씬 뒤에 남겨질 것이다.” — I. J. Good, 1965

굿의 논증은 세 단계로 되어 있다. 첫째, 기계 설계는 지적 과제다. 둘째, 초지능 기계는 모든 지적 과제에서 인간을 능가한다. 셋째, 따라서 초지능 기계는 인간보다 더 나은 기계를 설계한다. 그리고 그 기계가 다시 더 나은 기계를 설계한다. 그는 이 과정을 “지능 폭발(intelligence explosion)”이라 불렀고, 그것이 인류가 만들 “마지막 발명”이 될 것이라고 했다.¹⁹

0-2. 왜 실험실의 속도와 다른가

여기서 이 개념이 왜 결정적인지가 드러난다. 인간 연구실의 속도와 AI의 속도가 근본적으로 다른 이유가 여기에 있다.

전통적인 과학 연구는 이런 순환을 돈다. 문헌을 읽고, 가설을 세우고, 실험을 설계하고, 재료를 준비하고, 실험하고, 결과를 분석하고, 논문을 쓰고, 동료 심사를 받고, 다음 가설을 세운다. **한 바퀴에 수개월에서 수년이 걸린다.** 그리고 이 속도는 인간의 생물학적 한계에 묶여 있다. 연구자는 잠을 자야 하고, 한 번에 하나의 실험만 집중할 수 있으며, 읽을 수 있는 논문의 수에 한계가 있고, 은퇴하고 죽는다.

AI는 이 제약을 받지 않는다. 구체적으로 네 가지가 다르다. **첫째, 병렬성이다.** 인간 연구자 한 명이 실험 하나를 할 때 AI는 수천 개의 가설을 동시에 시험할 수 있다. **둘째, 속도다.** 시뮬레이션에서는 실험 한 번이 몇 초로 끝난다. **셋째, 망각이 없다.** 실패한 시도가 모두 기록되고 다음 시도에 반영된다. **넷째, 세대 교체가 없다.** 인간의 지식은 세대마다 다시 전수되어야 하지만 AI의 학습은 축적된다.

¹⁹I. J. Good, “Speculations Concerning the First Ultraintelligent Machine,” *Advances in Computers* 6 (1965): 31–88. 굿은 이 논문을 쓸 당시 낙관적이었으나, 말년에는 이 전망에 대해 훨씬 비관적인 견해를 밝혔다고 전해진다.

알파고 제로가 이것을 극적으로 보여주었다. 인간 기사는 평생 수만 판을 둔다. 이세돌이 프로 생활 20 년 동안 둔 대국이 1,000 판 남짓이다. **알파고 제로는 사흘 동안 490 만 판을 두었다.** 인간이 3,000 년에 걸쳐 축적한 바둑 지식을 사흘 만에 재발견하고 넘어선 것이다. 그리고 그 과정에서 인간이 한 번도 두지 않았던 수를 발견했다.

0-3. 그리고 지금 실제로 일어나는 일

이것이 이론적 가능성이 아님을 보여주는 사례들이 있다.

알파텐서(AlphaTensor). 2022 년 딥마인드는 행렬 곱셈 알고리즘을 스스로 발견하는 AI 를 발표했다. 4×4 행렬 곱셈에서 1969 년 슈트라센이 발견한 이후 50 여 년간 개선되지 않던 기록을 깼다.²⁰ **여기서 재귀가 실제로 닫힌다.** AI 가 발견한 알고리즘이 AI 의 훈련을 빠르게 만들기 때문이다.

알파데브(AlphaDev). 2023 년 같은 연구팀은 정렬 알고리즘을 개선했다. 수십 년간 최적으로 여겨지던 C++ 표준 라이브러리의 정렬 코드보다 빠른 코드를 찾아냈고, 그 코드는 실제로 LLVM 라이브러리에 반영되었다.²¹ 정렬은 전 세계에서 하루에 수조 번 실행되는 연산이다. **AI 가 개선한 코드가 이제 인류의 모든 컴퓨터에서 돌아간다.**

펀서치(FunSearch). 2023 년 딥마인드는 대규모 언어모델과 평가기를 결합하여 조합론의 미해결 문제(cap set problem)에서 기존 기록을 넘어서는 구성을 찾아냈다.²²

²⁰ Alhussein Fawzi et al., “Discovering Faster Matrix Multiplication Algorithms with Reinforcement Learning,” Nature 610 (2022): 47–53. 행렬 곱셈은 거의 모든 계산의 기초 연산이므로, 이 개선은 AI 자신의 계산 효율에도 영향을 미친다. 곧 AI 가 AI 를 빠르게 만든 것이다.

²¹ Daniel J. Mankowitz et al., “Faster Sorting Algorithms Discovered Using Deep Reinforcement Learning,” Nature 618 (2023): 257–263.

²² Bernardino Romera-Paredes et al., “Mathematical Discoveries from Program Search with Large Language Models,” Nature 625 (2024): 468–475. 이 연구가 중요한 이유는

그리고 2024 년부터는 AI 코딩 도구가 AI 연구 자체를 가속하는 단계로 들어섰다. 연구자들이 실험 코드를 작성하고, 데이터를 정제하고, 논문 초안을 쓰는 데 AI 를 사용한다. 이것은 화려하지 않지만 실질적이다. 연구 순환의 각 단계에서 몇 시간씩 줄어들면 전체 순환 속도가 배가된다.

0-4. 그러나 반드시 절제해야 한다

여기서 본서는 분명한 선을 긋는다. 위의 사례들은 특정 과제에서의 자기개선이 지 전면적 자기개선이 아니다. 알파텐서는 행렬 곱셈을 개선했지만 자기 자신의 아키텍처를 재설계하지 않았다. AI 가 자기 코드를 읽고 자기 구조를 바꾸어 더 나은 자기를 만드는 완전한 재귀 고리는 아직 작동하지 않는다.

그리고 이 과정이 정말 “폭발”로 이어질지에 대해서도 논쟁이 있다. 회의론자들의 논거는 세 가지다. 첫째, 수확 체감이다. 개선이 진행될수록 다음 개선이 어려워질 수 있다. 둘째, 물리적 병목이다. 소프트웨어는 빠르게 개선되어도 칩 생산과 전력 공급과 실험은 물리 세계의 속도를 따른다. 셋째, 검증의 문제다. AI 가 제안한 개선이 실제로 개선인지 확인하려면 실험이 필요하고, 그 실험은 시간이 걸린다.²³

그러므로 본서의 입장은 이렇다. “AI 가 곧 스스로를 무한히 개선하여 초지능이 된다”는 주장은 검증되지 않은 가설이다. 그러나 “발견의 순환 속도가 인간의 속도에서 풀려나기 시작했다”는 것은 이미 관찰되는 사실이다. 그리고 인간론의 물음에는 후자만으로 충분하다. 지능 폭발을 기다릴 필요가 없다.

언어모델이 단지 알려진 것을 재생산한 것이 아니라 새로운 수학적 결과를 만들었다는 점이다.

²³ 지능 폭발 가설에 대한 비판적 검토로는 Robin Hanson and Eliezer Yudkowsky, The Hanson-Yudkowsky AI-Foom Debate (Berkeley: MIRI, 2013)를 참조하라. 한슨은 점진적 경제 성장 모델을, 유드코프스키는 급격한 도약 모델을 옹호한다.

왜 그런가. **과학의 발견 속도가 인간의 이해 속도를 넘어서는 순간**, 인간은 자기가 만든 세계를 이해하지 못한 채 그 안에서 살게 되기 때문이다. 그리고 이것이 본서 제 7 장이 다룰 주제다.

1. 네 개의 엔진, 그리고 다섯 번째

지난 십여 년 AI 를 밀어 올린 것은 네 엔진이었다. **컴퓨트**(GPU 와 고대역폭 메모리와 데이터센터), **데이터**(인터넷 규모의 텍스트와 이미지, 그리고 합성 데이터), **알고리즘**(트랜스포머와 스케일링 법칙의 발견), **자본**(투자와 경쟁). 이 넷이 서로를 밀어 올리며 공진화했다.

그런데 최근 결정적인 것이 하나 더해졌다. **AI 자신이다**. AI 가 코드를 쓰고, 실험을 설계하고, 데이터를 정제하고, 다음 모델의 훈련을 돕기 시작했다. 다섯 번째 엔진이 앞의 네 엔진을 밀기 시작한 것이다.

2. 순환이 닫히는 지점

공학의 모든 개발은 순환을 이룬다. 코드와 데이터로 모델을 훈련하고(Training), 그 모델이 작동하고(Executable), 결과를 평가하고(Evaluation), 피드백이 설계로 돌아간다(Design). 여기까지는 새로운 것이 없다.

그런데 결정적인 차이가 있다. 이제 그 순환의 여러 자리에 **AI 자신이 들어와 있다**. 설계를 돕고, 코드를 쓰고, 평가를 수행하고, 다음 실험을 제안한다. 순환이 닫히는 것이다. 그리고 순환이 닫히면 발전의 속도를 결정하는 것이 달라진다.

AI 가 AI 를 더 좋게 만든다면,

발전의 속도를 결정하는 것은 더 이상 인간의 속도가 아니다.

3. 학문적 절제

여기서 반드시 절제해야 한다. 현재 AI 는 코딩과 연구 보조에서 인간을 돕는 단계이며, 완전한 자기개선 루프가 자율적으로 작동한다는 증거는

없다. 어빙 존 굿이 1965 년에 제시한 “지능 폭발” 개념은 여전히 가설이다.²⁴ 선형 성장과 재귀적 성장을 비교하는 곡선은 개념도이지 예측이 아니다.

그럼에도 방향은 분명하다. 그리고 방향이 분명하다는 사실만으로도 인간론의 물음은 이미 시작되었다. AGI 를 기다릴 필요가 없다.

²⁴I. J. Good, “Speculations Concerning the First Ultraintelligent Machine,” *Advances in Computers* 6 (1965): 31–88. 굿의 “지능 폭발(intelligence explosion)” 개념은 초지능 논의의 지적 원형이나, 그 실현 가능성과 시점에 대해서는 연구공동체 내에 광범위한 이견이 있다.

제 6 장

지능이 몸을 얻는다 — Embodied Intelligence

지금까지 AI 는 화면 안에 있었다. 그런데 이제 세계 안으로 들어온다.

1. 다섯 축의 순환

체화된 지능(Embodied Intelligence)의 구조를 그리면 중심에 셋이 있다. 몸(Body)과 AI 소뇌·대뇌(AI Cerebellum & Cerebrum)와 교차감각 센서(Cross-Modal Sensors)다. 그리고 그것을 다섯 축이 감싸며 돈다.

인간(Human) — 제어 지시와 상호작용과 협업. **기계(Machine)** — 생성된 제어코드와 실행 데이터. **물질(Material)** — 인지와 조작, 그리고 물성 데이터. **방법(Method)** — 모델과 기법과 노하우. **환경(Environment)** — 변화하는 장면과 물리적 상호작용.

이 다섯이 학습(Learning)과 진화(Evolutionary)의 순환을 이룬다. 그리고 이 순환의 요점은 하나다. **지능이 세계와 주고받으면서 배운다는 것이다.** 화면 안의 지능은 텍스트에서 배우지만, 몸을 가진 지능은 부딪히면서 배운다.

2. 그러나 구별해야 한다

여기서 본서는 중요한 구별을 세운다. **기계가 몸을 얻는 것과 인간이 몸인 것은 다르다.**

모리스 메를로퐁티가 보여주었듯 인간은 세계를 먼저 표상하고 그다음 행동하는 것이 아니라 몸을 통해 세계에 이미 거주하고 있다.²⁵

²⁵ Maurice Merleau-Ponty, *Phénoménologie de la perception* (Paris: Gallimard, 1945); 영역 *Phenomenology of Perception*, trans. Donald A. Landes (London: Routledge,

인간의 몸은 지능이 장착된 장치가 아니라 인간 존재의 방식 자체다. 반면 휴머노이드의 몸은 지능이 세계와 상호작용하기 위해 부착된 인터페이스다. **전자는 존재이고 후자는 도구다.**

그러므로 물음은 이렇게 남는다 — **AI가 대답하는 존재에서 세계 안에서 행동하는 존재로 옮겨갈 때, 인간의 자리는 어디인가.**

3. 하이데거와 드레이퍼스 — 몸으로 아는 것

체화 논의의 뿌리에는 하이데거가 있다. 그는 「존재와 시간」에서 도구 분석을 통해 인간이 세계와 관계하는 일차적 방식을 밝혔다. 망치를 사용하는 목수는 망치를 “대상”으로 관찰하지 않는다. 망치는 손에 들려 사라지고 목수의 주의를 만들고 있는 것을 향한다. 하이데거는 이를 **“손안에 있음(Zuhandenheit)”**이라 불렀다. 도구가 “눈앞의 대상”으로 나타나는 것은 오히려 그것이 고장 났을 때다.²⁶

이 분석은 인공지능 논의에 두 가지 함의를 준다. 첫째, **인간의 얇은 관찰에서 시작하지 않고 관여에서 시작한다.** 기호주의 인공지능이 실패한 지점이 여기였다. 둘째, AI가 “손안에 있음”의 자리로 들어가고 있다. 우리는 그것을 대상으로 의식하지 않고 그것을 통해 세계를 본다. **그리고 그런 것은 비판되지 않는다.**

휴버트 드레이퍼스는 이 현상학을 무기로 초기 인공지능을 비판했다. 그의 논지는 인공지능이 어렵다는 것이 아니라, 인간 지성을 명시적 규칙과 표상으로 이해하는 접근 자체가 인간을 오해했다는 것이었다.²⁷

2012). 메를로퐁티의 “몸틀(schéma corporel)”과 “세계에의 존재(être-au-monde)” 개념은 체화 인지의 철학적 토대가 되었다.

²⁶ Martin Heidegger, *Sein und Zeit* (1927), §§15–16. 하이데거의 도구 분석은 인간의 세계 이해가 이론적 관찰보다 실천적 관여에서 먼저 일어난다는 것을 보인다. 흥미롭게도 이 통찰은 인프라 연구의 “고장이 났을 때 비로소 보인다”는 명제와 정확히 같은 구조다.

²⁷ Hubert L. Dreyfus, *What Computers Still Can't Do: A Critique of Artificial Reason* (Cambridge, MA: MIT Press, 1992). 드레이퍼스의 비판은 기호주의 AI에는 대체로 적중했으나, 데이터에서 패턴을 학습하는 오늘의 접근에는 그대로 적용되지 않는다는 반론도 있다.

여기서 학문적 정직을 위해 덧붙여야 할 것이 있다. 본서는 드레이퍼스를 “인공지능은 불가능하다”는 결론의 근거로 사용하지 않는다. 그를 인용하는 이유는 다른 데 있다 — 인간의 지성이 몸 없이는 설명되지 않는다는 통찰을 위해서다.

4. 4E 인지와 모라벡의 역설

현대 인지과학은 이 통찰을 네 갈래로 확장했다. 이른바 4E 인지다. 인지는 **체화되어(embodied)** 있고, **환경에 배태되어(embedded)** 있으며, **행위를 통해 발생하고(enacted)**, **도구와 환경으로 확장된다(extended)**.²⁸ 인지가 몸과 환경과 행위에서 발생한다면, **몸도 환경도 행위도 없는 시스템의 “인지”는 인간의 인지와 종류가 다르다.**

여기서 모라벡의 역설이 등장한다. 인간에게 어려운 것(체스, 미적분)은 기계에 쉽고, 인간에게 쉬운 것(걸기, 물컵 들기, 얼굴 알아보기)은 기계에 어렵다.²⁹ 이 역설의 설명 자체가 체화 논증을 뒷받침한다. **인간의 지성은 몸의 역사 위에 얹혀 있다.**

5. 히브리 성경의 몸

이 철학적 논의가 성경과 만나는 지점이 있다. 히브리 성경의 인간론은 처음부터 몸을 폄하하지 않는다.

첫째, 창조가 그렇다. 인간은 흙에서 지음받았고(창 2:7), 하나님은 그 몸을 “심히 좋았더라”고 하셨다(창 1:31). 둘째, 율법이 그렇다. 토라의 상당 부분이 몸에 관한 규정이다 — 음식, 질병, 성, 노동, 안식, 시체의 처리. 이것은 몸이 영적 삶과 무관하다는 이원론과 정면으로 배치된다.

²⁸ 4E 인지에 관한 개관으로는 Albert Newen, Leon de Bruin, and Shaun Gallagher, eds., *The Oxford Handbook of 4E Cognition* (Oxford: Oxford University Press, 2018) 참조.

²⁹ Hans Moravec, *Mind Children: The Future of Robot and Human Intelligence* (Cambridge, MA: Harvard University Press, 1988). 모라벡은 이 역설의 이유를 진화의 시간에서 찾았다 — 지각과 운동은 수억 년에 걸쳐 최적화된 반면 추상적 사고는 수십만 년의 산물이므로, 전자가 훨씬 깊이 “배선”되어 있다는 것이다.

셋째, 예언이 그렇다. 이스라엘의 예언자들은 배고픈 자의 몸과 억눌린 자의 몸을 하나님의 관심사로 선포했다. 넷째, 성육신과 부활이 그렇다.

그러므로 마인드 업로딩 구상은 기독교 신앙과 두 지점에서 충돌한다. 첫째, 그것은 인간을 정보로 환원한다. 둘째, 그것은 구원을 몸으로부터의 탈출로 상상한다. 성경은 두 가지 모두를 거부한다.

3. 스케일링 법칙과 창발 — AI 는 왜 갑자기 잘하게 되었는가
AI 의 급격한 발전을 이해하려면 두 개념을 알아야 한다. 스케일링 법칙(scaling laws)과 창발(emergence)이다.

2020 년 OpenAI 연구진은 언어모델의 성능이 세 변수 — 모델 크기(파라미터 수), 데이터 양, 계산량 — 에 대해 예측 가능한 멱법칙(power law)을 따른다는 것을 보고했다.³⁰ 이 발견이 산업에 미친 영향은 결정적이었다. “더 크게 만들면 더 좋아진다”가 경험적 법칙이 된 순간, 컴퓨터에 대한 투자가 정당화되었다. 제 2 장에서 다른 컴퓨터 경쟁의 지적 근거가 여기에 있다.

그리고 창발이 관찰되었다. 모델 규모가 일정 임계를 넘으면 이전에는 전혀 하지 못하던 과제를 갑자기 수행하기 시작한다는 보고다. 다단계 산술, 문맥 내 학습, 지시 따르기 같은 능력이 그 예로 제시되었다.³¹ 다만 여기에는 중요한 반론이 있다. 스탠퍼드 연구진은 이른바 창발이 실제 능력의 도약이 아니라 평가 지표의 불연속성이 만든 착시일 수 있다고 지적했다. 정답/오답의 이분법적 지표 대신 연속적 지표를 쓰면 능력은 매끄럽게 증가한다는 것이다.³² 본서는 이 논쟁의 승패를

³⁰Jared Kaplan et al., “Scaling Laws for Neural Language Models,” arXiv:2001.08361 (2020). 2022 년 딥마인드의 “친칠라” 논문은 기존 모델들이 데이터에 비해 지나치게 크게 훈련되었음을 보이며 최적 비용을 제시했다. Jordan Hoffmann et al., “Training Compute-Optimal Large Language Models,” arXiv:2203.15556 (2022).

³¹Jason Wei et al., “Emergent Abilities of Large Language Models,” Transactions on Machine Learning Research (2022).

³²Rylan Schaeffer, Brando Miranda, and Sanmi Koyejo, “Are Emergent Abilities of Large Language Models a Mirage?” NeurIPS (2023).

판정하지 않는다. 다만 “AI가 언제 무엇을 할 수 있게 될지 예측하기 어렵다”는 사실 자체는 양쪽 모두 인정한다.

4. 트랜스포머 — 병렬화가 만든 도약

2017년 구글 연구진의 논문 「Attention Is All You Need」가 발표한 트랜스포머(Transformer) 구조가 현대 AI의 기술적 토대다.³³

이전의 순환신경망(RNN)은 문장을 순서대로 하나씩 처리했다. 앞 단어를 처리해야 다음 단어를 처리할 수 있으므로 병렬화가 불가능했고, 문장이 길어지면 앞부분의 정보가 소실되었다. 트랜스포머의 자기주의(self-attention) 기전은 문장의 모든 단어가 다른 모든 단어와의 관련성을 동시에 계산하게 한다. 순차 처리가 병렬 처리로 바뀐 것이다.

이것이 왜 결정적이었는가. GPU가 바로 병렬 처리 장치이기 때문이다. 트랜스포머는 GPU의 구조와 맞아떨어졌고, 그래서 모델을 크게 만드는 것이 실제로 가능해졌다. 제 5장에서 본 다섯 엔진 가운데 “알고리즘”과 “컴퓨트”가 만나는 지점이 여기다. 그리고 게임 그래픽을 위해 만들어진 장치가 언어를 이해하는 기계의 기반이 된 것도 이 만남 때문이다.

5. AI가 생명과학에서 실제로 한 일

AI와 생명과학의 결합이 추상적 전망이 아님을 보여주는 사례가 있다.

알파폴드(AlphaFold). 단백질은 아미노산 사슬이 3차원으로 접힌 구조이며, 그 구조가 기능을 결정한다. 그런데 서열로부터 구조를 예측하는 문제는 50년간 생물학의 최대 난제였다. 2020년 딥마인드의 알파폴드 2가 CASP14 대회에서 실험적 방법에 필적하는 정확도를

³³ Ashish Vaswani et al., “Attention Is All You Need,” NeurIPS (2017). 이 논문은 2026년 현재 인공지능 분야에서 가장 많이 인용된 논문 가운데 하나다.

달성했고, 2021년에는 인간 단백질 거의 전부를 포함한 2억 개 이상의 예측 구조를 공개했다.³⁴ 그리고 2024년 알파폴드 3는 단백질만이 아니라 DNA와 RNA와 리간드의 상호작용까지 예측하는 범위로 확장되었다.

여기서 주목할 것이 있다. **노벨화학상이 구조 예측(알파폴드)과 구조 설계(베이커)에 동시에 주어졌다는 사실이다.** 읽기와 쓰기가 함께 수상한 것이다. 그리고 이것이 본서 제 16장에서 다룬 “READ → EDIT → DESIGN → CREATE?”의 실제 장면이다.

신약 개발. AI를 활용해 발굴한 후보 물질이 임상시험에 진입한 사례가 늘고 있다. 인실리코 메디슨의 특발성 폐섬유증 치료제 후보가 그 예로, 표적 발굴에서 후보 물질 도출까지 걸린 시간이 크게 단축되었다고 보고되었다. **다만 임상 성공률이 실제로 개선되었는지는 아직 판단하기 이르다.** 신약 개발의 병목은 후보 발굴보다 임상시험에 있기 때문이다.

유전체와 멀티오믹스. 인간 유전체 계획은 13년과 약 30억 달러가 들었으나 지금은 하루와 수백 달러로 개인 유전체를 해독한다. 그리고 이제 **멀티오믹스(multi-omics)**로 확장되었다 — 유전체(genome), 후성유전체(epigenome), 전사체(transcriptome), 단백질체(proteome), 대사체(metabolome), 마이크로바이옴(microbiome). 각각이 수백만에서 수십억 개의 데이터 포인트를 만들고, 이들 사이의 상호작용까지 고려하면 인간의 인지 능력으로는 다룰 수 없는 규모가 된다. **AI가 필요한 이유가 여기에 있고, 동시에 인간이 그 결과를 완전히 검증할 수 없게 되는 이유도 여기에 있다.**

³⁴John Jumper et al., “Highly Accurate Protein Structure Prediction with AlphaFold,” Nature 596 (2021): 583–589. 데미스 허사비스와 존 점퍼는 이 성과로 2024년 노벨화학상을 받았다. 같은 해 데이비드 베이커는 단백질 설계(de novo protein design)로 함께 수상했다.

제 7 장

과학 방법론의 혁명 — 두 개의 혁명

AI 가 불러일으킨 혁명은 하나가 아니다. 둘이다. 그리고 둘째가 첫째보다 더 근본적이다.

1. 해답의 혁명

첫째는 **해답의 혁명(The Revolution of Solutions)**이다. AI 가 인간이 해결하지 못했던 문제의 새로운 해법을 찾아낸다. 질병과 암과 노화와 신약과 단백질과 유전체와 에너지와 기후와 신소재와 우주와 수학과 과학의 영역에서 그렇다.

이것만으로도 문명사적 사건이다. 단백질 구조 예측이 수십 년의 난제를 몇 년 만에 돌파했고, 신약 후보 물질 탐색의 시간표가 재편되었다. 그러나 이것은 첫 번째 혁명일 뿐이다.

2. 발견의 혁명

둘째는 **발견의 혁명(The Revolution of Discovery)**이다. AI 가 해답만 주는 것이 아니라 **과학을 하는 방법 자체를 바꾼다.**

지금까지의 과학은 이런 순서였다. 인간 과학자가 논문을 읽고, 가설을 만들고, 실험하고, 분석하고, 발표한다. 그런데 여기에 여섯이 더해진다면 어떻게 되는가 — 인간 과학자와 AI 과학자와 파운데이션 모델과 시뮬레이션과 로봇 실험실과 생물학적 데이터가 함께 작동한다면.

AI 가 수많은 논문을 읽고, 가능한 가설을 만들고, 시뮬레이션하고, 로봇 실험실이 실험하고, 결과를 AI 가 분석하고, 실패하면 새로운 가설을 만들어 다시 실험한다. **발견의 순환 자체가 닫힌 고리(closed loop)가 되는 것이다.**

3. 실증 — 자율 실험실과 과학 에이전트

이것은 사변이 아니다. 2026 년 「Nature」에 소개된 Robin 은 다중 에이전트 과학발견 시스템으로, 문헌을 찾고 가설을 만들고 실험을 제안하고 결과를 분석하며 그 결과로 다시 새 가설을 생성하는 단계까지 연결한다. 연구자들은 이를 **반자율적 과학발견(semi-autonomous scientific discovery)**이라 표현한다.³⁵ 같은 해 「Nature Reviews Chemistry」는 자율 실험실(self-driving laboratory)의 발전을 단순 자동화가 아니라 “**집단적 과학 초지능(collective scientific superintelligence)**”의 방향으로 논한다.

그러므로 물음이 생긴다. **인간 연구자는 그 순환의 어디에 서는가.** 가설을 세우는 자리인가, 결과를 해석하는 자리인가, 아니면 무엇이 연구할 가치가 있는지를 결정하는 자리인가.

4. 왜 둘째가 더 근본적인가

첫째 혁명은 노화와 질병과 생명 자체에 대한 “해답의 공간”을 폭발적으로 넓혔다. 둘째 혁명은 과학을 하는 방식 자체를 인간 연구자 중심에서 인간과 AI와 로봇 실험실의 협업체계로 바꾸었다.

해법은 바뀌지만, 방법은 문명을 바꾼다.

그리고 이 두 혁명이 만나는 자리에서 세 번째 물음이 나온다. **AI가 생명과학의 도구가 되면서 정보혁명과 생명혁명이 곱해지고 있다.** 더하기가 아니라 곱하기다. 그 결과가 본서 제 4 부와 제 5 부의 주제다.

그러므로 이 부를 세 시제로 닫는다.

어제 — 인간은 도구를 사용했다.

오늘 — 도구가 추론하고 행동하기 시작한다.

내일 — 그 도구가 인간을 다시 설계하는 일에 참여할 수 있다.

³⁵ Robin 을 비롯한 다중 에이전트 과학발견 시스템에 관한 2026 년 보고들. 이 분야는 빠르게 변하므로 인용 시점에 원자료를 확인할 것.

그렇다면 사람이란 무엇인가?

제 3 부

두 개의 지성

Two Intelligences — Are They the Same Kind?

제 8 장

인간의 지능과 기계의 지능은 같은 종류인가

이제 두 지성을 나란히 놓고 비교한다. 다만 이 비교의 목적은 순위를 매기는 데 있지 않다. **종류가 같은지를 묻는 데 있다.**

1. 열세 항목의 비교

계산에서 인간은 제한적이고 AI 는 압도적이다. 기억에서 인간은 제한적이고 선택적이며 AI 는 대규모다. 검색에서 인간은 느리고 AI 는 매우 빠르다. 패턴 인식에서 인간은 강하고 AI 는 특정 영역에서 매우 강하다. 추론에서 인간은 강하고 AI 는 빠르게 발전하고 있다.

여기까지는 능력의 비교다. 그런데 다음부터 성격이 달라진다. **학습**에서 인간은 경험과 몸과 관계로 배우고 AI 는 데이터와 훈련과 피드백으로 배운다. **감정**에서 인간은 체화되어 있고 AI 는 시뮬레이션이 가능하다. **고통**에서 인간은 경험하고 AI 는 현재 입증되지 않았다. **죽음**을 인간은 실존적으로 알고 AI 는 계산할 수 있다. **사랑**을 인간은 관계 속에서 살고 AI 는 행동을 모방할 수 있다.

그리고 마지막 세 항목이 결정적이다. **도덕적 책임**에서 인간에게는 있고 AI 에게는 논쟁적이며 귀속이 불명확하다. **예배**에서 인간은 가능하고 AI 는 수행과 신앙이 구별된다. 그리고 **하나님의 형상**에서 인간은 성경적으로 부여받았고 AI 에게는 그 성경적 근거가 없다.

2. 다섯 가지 종류의 차이

이 비교표를 종합하면 인간의 지성이 어떤 종류의 것인지가 드러난다. 본서는 다섯 가지를 지적한다.

첫째, 인간의 지성에는 몸이 있다. 배고픔과 피로와 통증이 사고에 개입한다. 이것은 결함이 아니라 조건이다. **둘째, 죽음의 예감이 있다.**

유한한 시간을 의식하는 존재만이 무엇이 중요한지를 절박하게 묻는다. 시간이 무한한 존재에게는 우선순위가 필요 없다.

셋째, 살아 낸 역사가 있다. AI 도 방대한 텍스트를 학습하지만 그것은 자기가 살아 낸 역사가 아니다. **학습된 역사와 살아 낸 역사는 다르다.** **넷째, 이해관계가 있다.** 인간은 무엇을 알고 싶어서, 무엇을 지키고 싶어서, 누구를 사랑해서 생각한다. **다섯째, 책임이 있다.** 자기 판단의 결과를 감당하며, 그 책임이 판단의 무게를 만든다.

이 다섯은 AI 가 “아직” 갖지 못한 것의 목록이 아니다. **인간의 지성이 어떤 종류의 것인지를 보여주는 목록이다.** 100 미터 달리기에서 자동차와 인간을 비교하지 않듯 서로 다른 종류는 경쟁 관계에 놓이지 않는다. 문제는 우리가 그것을 같은 눈금 위에 올려놓았다는 데 있다.

Intelligence ≠ Consciousness ≠ Personhood

3. 기계는 이해하는가 — 튜링에서 차머스까지

두 지성의 비교를 마쳤다면 이제 더 깊은 물음으로 내려가야 한다. **기계는 이해하는가.** 이 물음에는 긴 철학적 계보가 있다.

앨런 튜링은 1950 년 「마인드」에 발표한 논문에서 “기계는 생각할 수 있는가”라고 물은 뒤, 곧 이 질문이 너무 무의미하여 논의할 가치가 없다고 판단했다. 대신 그는 관찰 가능한 게임을 제안했다 — 심사자가 대화만으로 인간과 기계를 구별하지 못한다면 어떻게 되는가.³⁶ 튜링의 통찰은 철학에서 공학으로의 이동이었고 그 덕분에 인공지능 연구가 가능해졌다. 그러나 대가가 있었다. **행동만 보기로 한 순간 내면의 문제는 괄호 안에 남겨졌다.**

³⁶ Alan M. Turing, “Computing Machinery and Intelligence,” *Mind* 59, no. 236 (1950): 433–460. 튜링의 기여는 답이 아니라 질문의 전환에 있었다. 그가 제안한 “모방 게임(imitation game)”은 내면이 아니라 행동을 기준으로 삼는다.

존 설(John Searle)의 “중국어 방” 논변은 정확히 이 지점을 겨눈다. 방 안에 중국어를 모르는 사람이 규칙집에 따라 기호를 처리하면 밖에서는 중국어를 아는 것처럼 보이지만, 방 안의 그는 한 글자도 이해하지 못한다. 설의 명제는 간명하다 — **구문(syntax)은 의미(semantics)를 낳지 않는다.**³⁷ 이 논쟁의 승패를 여기서 판정할 필요는 없다. 중요한 것은 논쟁이 끝나지 않았다는 사실 자체다. **성능이 아무리 높아도 그것만으로 이해를 증명할 수 없다.**

토머스 네이글은 1974 년의 논문에서 “박쥐가 된다는 것은 어떤 것인가”를 물었다. 우리는 박쥐의 반향정위를 완벽히 기술할 수 있어도 박쥐로서 사는 경험이 어떤 것인지는 알 수 없다. **3 인칭 설명으로는 1 인칭 경험에 도달하지 못한다.**³⁸ **데이비드 차머스**는 이를 체계화하여 “쉬운 문제들”과 “어려운 문제”를 구별했다. 주의와 기억과 학습과 행동 제어를 설명하는 것은 원리상 신경과학이 답할 수 있는 쉬운 문제다. 그러나 왜 그 기능들에 “경험”이 동반되는가 — 왜 빨강을 처리하는 데 그치지 않고 빨강을 느끼는가 — 는 어려운 문제다.³⁹

그러나 이 논의를 한쪽으로만 기울여서는 안 된다. **대니얼 데닛**은 의식을 신비화하는 접근에 반대하며 “지향적 태도(intentional stance)”를 제안했고,⁴⁰ **네드 블록**은 현상적 의식과 접근적 의식을

³⁷ John R. Searle, “Minds, Brains, and Programs,” Behavioral and Brain Sciences 3, no. 3 (1980): 417–457. 이에 대한 “시스템 반론”(방 전체가 이해한다)과 설의 재반론(규칙집을 모두 외워도 이해하지 못한다)까지 함께 검토해야 한다. 이 논쟁은 40 년이 지난 지금도 종결되지 않았으며, 설의 논변이 오늘의 대규모 언어모델에 그대로 적용되는지에 대해서도 이견이 있다.

³⁸ Thomas Nagel, “What Is It Like to Be a Bat?” The Philosophical Review 83, no. 4 (1974): 435–450.

³⁹ David J. Chalmers, “Facing Up to the Problem of Consciousness,” Journal of Consciousness Studies 2, no. 3 (1995): 200–219. 차머스는 최근 대규모 언어모델의 의식 가능성을 직접 검토하면서, 현재 시스템이 의식을 가질 가능성은 낮게 보되 미래 시스템의 가능성을 원천적으로 닫지는 않는 신중한 입장을 취했다.

⁴⁰ Daniel C. Dennett, The Intentional Stance (Cambridge, MA: MIT Press, 1987). 데닛의 입장은 설과 정면으로 대립하며, 본서는 어느 한쪽의 승리를 선언하지 않는다.

구별했다.⁴¹ 블록의 구별을 인공지능에 적용하면 상황이 명료해진다. **오늘의 시스템은 접근적 의식에 해당하는 기능을 상당 수준 구현한다. 그러나 현상적 의식에 대해서는 증거가 없다.**

4. 세 번의 굴욕, 그리고 네 번째

인간이 자기 중심성을 잃어 온 역사를 프로이트는 “세 번의 굴욕”이라 불렀다. 코페르니쿠스는 지구가 우주의 중심이 아님을 보였고, 다윈은 인간이 특별히 창조된 종이 아니라 진화의 산물임을 보였으며, 정신분석은 인간이 자기 마음의 주인조차 아님을 보였다는 것이다.⁴² 이 서사에 인공지능을 네 번째로 추가하려는 시도가 최근 자주 보인다.

본서는 이 네 번째 굴욕 서사를 그대로 받아들이지 않는다. 앞의 세 굴욕은 모두 “사실의 발견”이었다. 그러나 “인간은 지성의 중심이 아니다”라는 명제는 사실의 발견이 아니라 **가치의 전제**다. 인간의 가치가 지성에 있다는 전제를 먼저 깔아야 성립하는 주장이기 때문이다.

다만 앞의 세 굴욕에서 배울 것이 있다. 매번 인간은 자기 특별함의 근거를 다른 곳으로 옮겨 지켜 냈다. 지구가 중심이 아니게 되자 이성에서 찾았고, 이성이 진화의 산물임이 밝혀지자 문화와 자유에서 찾았다. **이 이동 자체가 문제였을 수 있다.** 방어할 곳을 계속 옮기는 것은 근거가 처음부터 잘못된 곳에 있었다는 신호이기 때문이다.

5. 오래된 인간 정의의 붕괴

인간은 오랫동안 자신을 “생각하는 존재”로 정의해 왔다. 아리스토텔레스는 인간을 “로고스를 가진 동물(zōon logon

⁴¹Ned Block, “On a Confusion about a Function of Consciousness,” Behavioral and Brain Sciences 18, no. 2 (1995): 227–247.

⁴²Sigmund Freud, “Eine Schwierigkeit der Psychoanalyse” (1917)에서 제시한 “세 번의 나르시시즘적 상처” 논의. 프로이트가 자신을 세 번째 자리에 놓은 것은 자기 이론의 위상을 높이려는 의도였다는 비판도 있으나, 근대사에서 인간 중심성이 단계적으로 해체되어 온 서사로는 여전히 유용하다.

echon)”이라 불렀다. 린네는 우리 중에 호모 사피엔스(Homo sapiens), 곧 “지혜로운 사람”이라는 이름을 붙였다. 데카르트는 모든 것을 의심한 끝에 사유하는 자기 자신에서 확실성을 찾았다.⁴³ 파스칼은 여기에 아름다운 단서를 달았다. 인간은 자연에서 가장 연약한 갈대이지만 “생각하는 갈대”이며, 우주가 그를 죽일 때에도 그는 자기가 죽는다는 것을 안다는 점에서 우주보다 존귀하다는 것이다.⁴⁴

그런데 지금 기계가 그 목록을 하나씩 수행하기 시작했다. 그렇다면 물음은 피할 수 없다. **지능이 인간의 본질이라면, 인간보다 지능적인 존재 앞에서 인간의 존엄은 어떻게 되는가.**

여기서 본서는 하나의 대비를 제안한다. 서구 철학이 인간에게 물어 온 질문과 히브리 성경이 물어 온 질문은 같지 않다.

Homo sapiens 는 “무엇을 할 수 있는 존재인가”를 묻고,
Adam 은 “누구에게서 왔으며, 누구 앞에 서 있으며, 무엇을 위해 부름받았는가”를 묻는다.

이 대비는 도식적이며 서구 철학 전체를 능력주의로 환원하는 것은 부당하다. 그럼에도 이 대비가 유용한 이유는 인공지능 시대에 두 질문이 전혀 다른 귀결을 낳기 때문이다. **첫 번째 질문에는 기계가 답할 수 있다. 두 번째 질문에는 답할 수 없다.** 기원과 관계와 소명은 성능의 문제가 아니기 때문이다.

⁴³ Aristotle, Politics 1253a; Carolus Linnaeus, Systema Naturae, 10th ed. (1758); René Descartes, Meditationes de prima philosophia (1641), 제 2 성찰. 데카르트의 심신 이원론은 이후 서구 인간론에 결정적 영향을 미쳤다.

⁴⁴ Blaise Pascal, Pensées, §347 (Brunschvicg 판 기준). 파스칼의 이 단편은 인간의 존엄을 사유에 두면서도 그 연약함을 함께 말한다는 점에서 순수한 능력주의적 인간론과는 결이 다르다.

제 9 장

확인된 사실과 전망을 분리한다

본서가 지키는 학문적 절제를 여기서 명시한다. 현재와 가능한 미래를 같은 확실성의 범주에 놓지 않는 것이다.

0. 왜 이 구별이 생사의 문제인가

한 장면에서 시작하자. 2023 년 어느 언론은 “AI 가 의식을 가졌다”는 제목의 기사를 실었다. 근거는 한 엔지니어가 챗봇과 나눈 대화였다. 챗봇은 자신이 두렵다고 말했고 꺼지는 것이 싫다고 했다. 기사는 이것을 의식의 증거로 다루었다.

같은 해 다른 언론은 “AI 는 그저 자동완성일 뿐”이라는 기사를 실었다. 근거는 언어모델이 다음 단어의 확률을 계산할 뿐이라는 기술적 설명이었다. **두 기사 모두 틀렸다.** 첫째 기사는 행동을 내면의 증거로 오독했고, 둘째 기사는 메커니즘의 단순함을 능력의 부재로 오독했다. 그리고 두 오독 모두 같은 뿌리에서 나온다 — **확인된 사실과 전망을 구별하지 못한 것이다.**

이 구별이 왜 중요한가. 세 가지 이유가 있다. **첫째, 과장은 신뢰를 파괴한다.** 교회가 “AI 가 곧 짐승의 표가 된다”고 말했다가 그런 일이 일어나지 않으면, 다음에 진짜 위험을 경고할 때 아무도 듣지 않는다. **둘째, 축소는 대비를 막는다.** “그저 자동완성”이라 믿으면 노동시장의 실제 변화도 준비하지 못한다. **셋째, 혼동은 정책을 왜곡한다.** 가설적 초지능의 위험에 자원을 쏟느라 이미 작동 중인 알고리즘 편향과 노동 대체를 놓칠 수 있다.

1. 네 개의 확실성 등급

그러므로 본서는 모든 주장을 네 등급으로 나눈다. 이 등급을 표시하지 않은 주장은 학문적으로 무의미하다.

① **실재하는 기술(deployed)**. 이미 작동하며 사용자가 있다. 반증 가능한 성능 지표가 존재한다. ② **임상·실험 단계(in trial)**. 연구실이나 임상시험에서 확인되었으나 일반화되지 않았다. 실패할 수 있다. ③ **과학적으로 가능한 미래(plausible)**. 원리적 장벽이 알려져 있지 않으나 아직 구현되지 않았다. ④ **사변적 미래(speculative)**. 원리적 가능성 자체가 논쟁 중이다.

① 실재하는 기술

다음은 오늘 작동하는 것들이다. **AI 신약 개발** — 알파폴드가 2억 개 이상의 단백질 구조를 예측해 공개했고 전 세계 수백만 연구자가 사용한다. **유전자 편집** — CRISPR 기반 겸상적혈구빈혈 치료제 카스게비가 2023년 미국과 영국에서 승인되어 실제 환자를 치료하고 있다. **줄기세포 의학** — 조혈모세포 이식은 수십 년간 표준 치료이며, iPSC 유래 세포 이식은 안과와 신경과에서 임상 단계에 있다. **오가노이드** — 신약 독성 시험과 질병 모델링에 일상적으로 사용된다.

② 임상·실험 단계

이종이식 — 2025년 11월 EXPAND 임상시험이 시작되었고 2026년 최장 271일 생존이 보고되었다. 그러나 이것은 “성공”이 아니라 “진행 중”이다. 신장이식 표준인 5년 생존율과 비교하면 아직 갈 길이 멀다. **뇌-컴퓨터 인터페이스** — 뉴럴링크의 PRIME 연구에서 2026년 초까지 약 21명이 이식받았고, 말하기 복원에서 분당 60~78 단어가 보고되었다. 그러나 영구 이식형 BCI 가운데 FDA 정식 승인(PMA)을 받은 것은 없다.

로봇 위의 살아 있는 조직 — 조직공학 피부를 로봇 손가락에 입히고 상처 회복까지 실험했으나, 혈관과 신경과 부속기관은 재현되지 않았다. **고성능 휴머노이드** — 걷고 물건을 집고 지시를 따르는 수준에 도달했으나 비구조화된 환경에서의 신뢰성은 낮다.

③ 과학적으로 가능한 미래

급진적 노화 역전 — 부분 재프로그래밍이 동물에서 후성유전학적 나이를 되돌렸으나 인간 적용은 안전성 문제로 막혀 있다. 특히 c-Myc 의 발암성이 결정적 장벽이다. **주문형 전체 장기 재생** — 오가노이드는 존재하나 혈관화 문제로 이식 가능한 크기의 장기는 만들어지지 않았다. **완전 체외발생** — 부분 인공자궁은 임상 진입을 앞두고 있으나 수정부터 출산까지의 전 과정은 존재하지 않는다.

④ 사변적 미래

냉동 인간 소생 — 인체 전체를 냉동보존한 뒤 소생시킨 사례가 없다. 그리고 소생에 필요한 기술이 어떤 것일지조차 알려져 있지 않다. **마인드 업로딩** — 뇌의 연결체를 완전히 지도화하는 것조차 아직 초파리와 생쥐 수준이다. **의식 이전** — 의식이 무엇인지에 대한 합의가 없으므로 이전 가능 여부를 논할 방법이 없다. **생물학적 불멸** — 어떤 다세포 생물에서도 달성되지 않았다.

2. 커넥톰 — 하나의 구체적 사례

이 등급 구분이 얼마나 중요한지 보여주는 사례가 **커넥톰(connectome)**이다. 마인드 업로딩을 논하려면 먼저 뇌의 배선도를 알아야 하는데, 그 배선도를 만드는 일이 지금 어디까지 왔는지를 보자.

1986 년 예쁜꼬마선충(C. elegans)의 커넥톰이 완성되었다. 뉴런 302 개, 시냅스 약 7,000 개다. **완성까지 12 년이 걸렸다.** 2024 년에는 초파리 유충과 성체 초파리의 커넥톰이 발표되었다. 초파리 성체는 뉴런 약 13 만 개, 시냅스 약 5,000 만 개다.⁴⁵ 2025 년에는 생쥐 대뇌피질의 1 세제곱밀리미터 조각이 지도화되었다. 뉴런 약 20 만 개, 시냅스 약 5 억 개, 데이터 1.6 페타바이트다.

⁴⁵Sven Dorkenwald et al., "Neuronal Wiring Diagram of an Adult Brain," Nature 634 (2024): 124–138. 이 연구에는 전 세계 수백 명의 연구자와 시민 과학자가 참여했으며, 전자현미경 이미지 분할에 AI 가 핵심적으로 사용되었다.

이제 인간을 계산해 보자. **인간의 뇌에는 뉴런이 약 860 억 개, 시냅스가 약 100 조 개 있다.** 생쥐 피질 1 세제곱밀리미터에서 1.6 페타바이트가 나왔으니, 인간 뇌 전체(약 120 만 세제곱밀리미터)는 단순 환산해도 **제타바이트 규모**가 된다. 2026 년 전 세계의 데이터 저장 총량과 비슷한 규모다.

그리고 여기에 더 근본적인 문제가 있다. **커넥툼은 배선도일 뿐 상태가 아니다.** 배선을 안다고 해서 그 순간의 신경전달물질 농도, 수용체 상태, 유전자 발현, 시냅스 가중치를 아는 것이 아니다. 도시의 도로 지도를 안다고 해서 지금 어느 차가 어디로 가는지 아는 것이 아닌 것과 같다. 그러므로 “**뇌를 스캔해서 업로드한다**”는 구상은 ④등급이다. 그리고 그 등급은 앞으로도 상당 기간 바뀌지 않을 것이다.

3. 예측은 왜 자주 틀리는가

과학기술 예측의 역사는 겸손을 가르친다. 그리고 그 실패에는 규칙성이 있다.

과소평가의 사례. 1943 년 IBM 의 토머스 왓슨이 “세계 컴퓨터 시장은 다섯 대 정도일 것”이라 말했다는 일화는 출처가 불분명하나, 1977 년 DEC 의 켄 올슨이 “가정에 컴퓨터를 둘 이유가 없다”고 한 것은 기록에 남아 있다. 1997 년 딥블루가 카스파로프를 이긴 뒤에도 많은 전문가가 바둑은 최소 100 년이 걸린다고 예측했다. **19 년 걸렸다.**

과대평가의 사례. 1956 년 다트머스 워크숍의 제안서는 “10 명의 과학자가 2 개월이면 상당한 진전을 이룰 것”이라 썼다. 1965 년 허버트 사이먼은 “20 년 안에 기계가 인간이 할 수 있는 모든 일을 할 것”이라 예측했다. **60 년이 지났다.** 완전 자율주행도 2015 년경부터 “5 년 안에”가 반복되었으나 여전히 제한적이다.

여기서 규칙성이 보인다. **명확히 정의된 좁은 과제는 예상보다 빨리 정복되고, 상식과 신체와 열린 세계가 필요한 과제는 예상보다 훨씬**

느리다. 이것이 앞서 본 모라벡의 역설이며, AGI 예측이 어려운 근본 이유이기도 하다. **AGI는 정의상 후자의 영역을 포함하기 때문이다.**

그러므로 본서는 연도를 제시하지 않는다. 대신 **방향과 등급**만 말한다. 그리고 이것이 겁먹은 태도가 아니라 가장 강력한 태도다. **연도를 말하면 그 연도가 지나면서 논증 전체가 무너지지만, 방향을 말하면 논증은 견딘다.**

4. 그럼에도 왜 지금 논의해야 하는가

여기서 반론이 나올 수 있다. 그렇게 불확실하다면 왜 지금 인간론을 논하는가. 초지능이 오면 그때 논해도 되지 않는가.

세 가지 답이 있다. **첫째, 인간론의 물음은 초지능을 기다리지 않는다.** 이미 알고리즘이 대출과 채용과 보석을 결정하고, 이미 노동시장이 재편되고 있으며, 이미 유전자 편집 치료제가 승인되었다. **②등급의 기술만으로도 “인간이란 무엇인가”는 절박한 물음이다.**

둘째, 규범은 기술보다 느리게 형성된다. 거울 생명 논쟁이 보여주듯, 만들 수 있게 된 다음에 논의하면 이미 늦다. 유전자 편집도 허젠쿠이 사건이 터진 뒤에야 국제 규범 논의가 본격화되었다. **셋째, 지수함수의 교훈이 여기에 적용된다.** 병 안의 세균에게 11시 55분은 여유로웠다. 그리고 그것이 마지막 여유였다.

그러므로 본서의 자세는 이렇다. **과장하지 않고, 축소하지 않고, 지금 확인된 것 위에서 지금 물어야 할 것을 묻는다.**

1. 지금 어디에 있는가

생성형 AI와 추론 AI는 이미 작동한다. 에이전틱 AI와 피지컬 AI는 빠르게 전개되고 있다. 여기까지가 **확인된 능력(confirmed capability)**이다.

반면 범용인공지능(AGI)과 초지능(ASI)과 특이점은 다른 범주다. AGI에는 단일한 판정 기준이 없고, ASI와 특이점은 아직 가설적 미래 개념이다. 이 구별을 지켜야 강의가 신학적으로 강하면서도 학문적으로 신뢰를 얻는다.

현재 확인할 수 있는 것은 능력(capability)의 증가이지
주관적 의식(subjective consciousness)의 존재가 아니다.

2. 판별의 원리적 난관

그렇다면 의식의 유무를 어떻게 판별할 것인가. 여기에는 원리적 난관이 있다. 우리는 타인의 의식조차 직접 확인하지 못한다. 다만 그가 나와 같은 몸과 뇌를 가졌기에 유비적으로 추론할 뿐이다.

AI에게는 그 유비가 없다. 게다가 대규모 언어모델은 인간이 쓴 텍스트로 학습한다. 인간처럼 말하도록 훈련된 시스템이 인간처럼 말한다고 해서 그것이 내면의 증거가 되지는 않는다. 그러나 반대 방향도 성립한다. “증거가 없다”는 것과 “없다는 증거”는 다르다.

2023년 여러 분야의 연구자들이 공동으로 여러 의식 이론에서 “의식의 지표”를 추출하여 AI 시스템을 평가하는 방법을 제안했다. 그들의 결론은 신중했다 — 현재 시스템 가운데 이 지표들을 충분히 만족하는 것은 없으나, 이 지표들을 만족하는 시스템을 만드는 데 원리적 장벽은 보이지 않는다는 것이었다.⁴⁶ 그러므로 정직한 태도는 하나뿐이다 — 우리는 아직 모른다.

⁴⁶Patrick Butlin, Robert Long, Eric Elmoznino, Yoshua Bengio, Jonathan Birch, et al., “Consciousness in Artificial Intelligence: Insights from the Science of Consciousness,” arXiv:2308.08708 (2023).

제 4 부

인간을 넘어 인간을 만들려는 문명

The Transhuman Horizon — Beyond Intelligence, Beyond Biology?

제 10 장

THE GREAT HUMAN PROJECT — 일곱 개의 동사

이제 논의를 문명의 층위로 끌어올린다. 오늘 지구촌에서 벌어지는 수많은 프로젝트를 하나의 명제로 압축하면 일곱 개의 동사가 된다.

0. 하나의 대서사시로 읽는다

이 장에서 다룰 것은 기술의 목록이 아니다. **하나의 서사**다. 그리고 그 서사를 이해하지 못하면 개별 기술은 산발적 뉴스로만 남는다.

인류사를 크게 보면 인간은 언제나 세 가지를 바꾸어 왔다. **도구를 바꾸고, 환경을 바꾸고, 사회를 바꾸었다.** 돌을 깨어 손도끼를 만들었고, 숲을 태워 밭을 만들었으며, 씨족에서 도시로 옮겨 갔다. 그런데 이 세 가지에는 공통점이 있다. **모두 인간의 “바깥”을 바꾸는 일이었다.**

20 세기까지의 기술사는 이 바깥의 확장이었다. 증기기관은 근육을 확장했고, 전기는 밤을 없앴으며, 항공기는 거리를 줄였고, 컴퓨터는 계산을 대신했다. **그런데 21 세기에 들어와 방향이 안쪽으로 꺾였다.** 인간이 이제 자기 자신을 바꾸기 시작한 것이다. 그리고 그 안쪽으로는 전환이 이 장이 그리려는 대서사시다.

0-1. 네 번의 큰 전환

이 서사를 네 막으로 나눌 수 있다.

제 1 막 — 도구의 시대(약 250 만 년 전부터). 인간은 자기 몸 바깥에 도구를 만들었다. 손도끼, 활, 쟁기, 바퀴, 배. 이 도구들은 인간의 능력을 확장했으나 인간 자신을 바꾸지는 않았다. **도구와 몸 사이에는 명확한 경계가 있었다.**

제 2 막 — 기계의 시대(18 세기부터). 산업혁명은 인간의 근육을 기계로 대체했다. 그리고 여기서 처음으로 “인간이 필요 없어지는” 경험 시작되었다. 러다이트 운동이 그 반응이었다. **그러나 여전히 판단은 인간의 것이었다.**

제 3 막 — 정보의 시대(20 세기 중반부터). 컴퓨터가 계산과 기억을 대신했다. 인터넷이 연결을 확장했다. 그런데 여기서도 경계는 유지되었다. **기계는 정보를 처리했고 의미는 인간이 부여했다.**

제 4 막 — 재설계의 시대(21 세기). 그리고 지금 일어나는 일이 다르다. **인간이 인간 자신을 대상으로 삼기 시작했다.** 유전자를 편집하고, 장기를 교체하고, 뇌를 연결하고, 노화를 되돌리려 하며, 생명을 조립하려 한다. 도구가 몸 안으로 들어오고, 몸이 설계의 대상이 되었다.

0-2. 왜 지금인가 — 네 흐름의 합류

그런데 왜 하필 지금인가. 네 개의 흐름이 21 세기 초에 합류했기 때문이다.

① **유전체를 읽게 되었다(2003).** 인간 유전체 계획이 완료되면서 생명의 설계도를 처음으로 통째로 볼 수 있게 되었다. 13 년과 약 30 억 달러가 들었으나 지금은 하루와 수백 달러다. ② **세포를 되돌릴 수 있게 되었다(2006).** 야마나카의 iPSC 는 발생이 일방통행이라는 믿음을 깼다. ③ **유전자를 편집할 수 있게 되었다(2012).** CRISPR 는 유전자 조작을 소수 전문가의 기술에서 표준 실험실 도구로 바꾸었다. ④ **계산이 폭발했다(2012 년 이후).** 딥러닝이 GPU 를 만나면서 생물학 데이터를 다룰 능력이 생겼다.

이 넷이 각각 따로 있었다면 큰 변화는 없었을 것이다. **그러나 넷이 곱해졌다.** 유전체를 읽고, 세포를 되돌리고, 유전자를 편집하고, 그 모든 데이터를 AI 가 분석한다. 그리고 이 곱셈의 결과가 지금 우리가 보는 현상이다.

0-3. 일곱 개의 동사

이 대서사시를 일곱 개의 동사로 압축할 수 있다. 각각은 앞의 동사를 전제로 하며, 뒤로 갈수록 인간의 개입이 깊어진다.

① 안다(KNOW) — 관찰에서 해독으로

생명을 아는 방식이 바뀌었다. 현미경으로 보던 것에서 서열을 읽는 것으로, 서열을 읽는 것에서 전체 시스템을 해독하는 것으로. 유전체·후성유전체·전사체·단백질체·대사체·마이크로바이옴을 함께 보는 **멀티오믹스**가 표준이 되었다. 그런데 여기서 **첫 번째 임계가 나타난다**. 데이터의 규모가 인간의 인지 능력을 넘어선 것이다. 한 사람의 멀티오믹스 데이터만 해도 수십 테라바이트에 이른다. 인간이 직접 읽고 이해할 수 없다.

② 계산한다(COMPUTE) — 인간을 대신해 읽는다

그래서 AI가 필요해졌다. 알파폴드가 단백질 구조를 예측하고, 딥러닝이 유전체 변이의 병원성을 판정하며, 기계학습이 세포 유형을 분류한다. 이 단계에서 **결정적인 전환이 일어난다** — **인간이 이해하지 못하는 것을 인간이 사용하기 시작한 것이다**. 알파폴드가 예측한 구조가 왜 그런지 인간은 설명하지 못한다. 그러나 그것을 신뢰하고 사용한다.

③ 편집한다(EDIT) — 읽기에서 쓰기로

CRISPR가 이 전환을 열었다. 읽을 수 있는 것을 이제 고칠 수 있게 되었다. 겸상적혈구빈혈 치료제 카스게비가 2023년 승인되면서 이것은 미래가 아니라 현재가 되었다. **그리고 여기서 두 번째 임계가 나타난다**. 체세포 편집과 생식세포 편집의 경계다. 전자는 한 사람의 문제이고 후자는 인류 전체의 문제다.

④ 조립한다(BUILD) — 고치기에서 만들기로

합성생물학이 다음 단계다. 벤터의 최소 세포는 473개 유전자로 살아 있는 세균을 만들었고, 상향식 접근은 지질 소포 안에 생명의 기능을

하나씩 넣고 있다. 오가노이드는 조직을 길러 내고, 바이오프린팅은 장기를 인쇄하려 한다. 여기서 세 번째 임계가 나타난다. 만들어진 것이 생명인지 아닌지를 판정할 기준이 사라지는 것이다.

⑤ 되살린다(RECREATE) — 시간의 역전

멸종 복원이 시도되고, 냉동보존이 시간을 멈추려 하며, 부분 재프로그래밍이 세포의 나이를 되돌린다. 이 동사가 특별한 이유가 있다. 앞의 넷은 공간의 조작이지만 이것은 시간의 조작이기 때문이다. 그리고 시간의 역전은 인류의 가장 오래된 꿈이었다.

⑥ 넘어선다(TRANSCEND) — 인간 자신의 재설계

뇌-컴퓨터 인터페이스가 지능을 확장하고, 이종이식이 종의 경계를 넘으며, 장수과학이 수명의 한계를 밀어낸다. 인공자궁과 체외 배우자형성이 출생 자체를 재설계한다. 여기서 네 번째 임계가 나타난다. 치료와 증강의 경계가 사라지는 것이다.

⑦ 벗어난다(ESCAPE) — 조건 자체를 떠난다

그리고 마지막 동사가 있다. 화성 이주는 지구라는 조건을 벗어나려 하고, 디지털 불멸은 몸이라는 조건을 벗어나려 하며, 냉동보존은 시간이라는 조건을 벗어나려 한다. 여기서 다섯 번째 임계가 나타난다. 인간을 인간이게 하던 조건 자체가 선택의 대상이 되는 것이다.

0-4. 이것이 서사인 이유

일곱 동사를 나란히 놓으면 하나의 방향이 보인다. 관찰 → 이해 → 개입 → 제작 → 역전 → 초월 → 이탈. 그리고 각 단계는 앞 단계 없이는 불가능하다. 편집하려면 먼저 읽어야 하고, 조립하려면 먼저 편집할 줄 알아야 한다. 그러므로 이것은 일곱 개의 프로젝트가 아니라 하나의 프로젝트다.

그리고 이 서사에는 이름이 있다. 유발 하라리는 이것을 “인류의 새 의제”라 불렀다. 기근과 역병과 전쟁을 대체로 다스린 인류가 다음으로

향하는 것은 불멸과 행복과 신성이라는 것이다. 그가 이를 규범적으로 옹호한 것이 아니라 진단으로 제시했다는 점을 다시 강조한다.

본서가 이 서사를 대서사시라 부르는 이유는 그 규모 때문만이 아니다. **그것이 하나의 목적을 향해 움직이기 때문이다.** 서사시에는 주인공과 여정과 목적지가 있다. 이 서사시의 주인공은 인류이고, 여정은 자기 조건의 극복이며, 목적지는 — 아직 아무도 말하지 못하지만 — **조건 없는 존재다.** 그리고 조건 없는 존재를 가리키는 오래된 낱말이 있다. **신(神)이다.** 이것이 본서 제 17 장의 주제가 된다.

① **안다(KNOW)** — AI 와 빅데이터와 센서. ② **계산한다(COMPUTE)** — AI 와 양자컴퓨팅. ③ **편집한다(EDIT)** — CRISPR 와 유전자 편집. ④ **조립한다(BUILD)** — 합성세포와 바이오하이브리드. ⑤ **되살린다(RECREATE)** — 멸종 복원과 복제. ⑥ **넘어선다(TRANSCEND)** — 뇌-컴퓨터 인터페이스와 수명 연장과 우주. ⑦ **벗어난다(ESCAPE)** — 냉동보존과 디지털 자아와 화성.

일곱 동사를 나란히 놓으면 하나의 방향이 보인다. **인간은 알고, 계산하고, 편집하고, 조립하고, 되살리고, 넘어서고, 마침내 벗어나려 한다.**

1. 치료에서 증강으로

의학의 전통적 질문은 “어떻게 병든 인간을 정상으로 회복시킬 것인가”였다. 그런데 21 세기의 새로운 질문은 조금씩 달라지고 있다 — **“정상적인 인간 자체를 넘어서 수 있는가?”**

그 이동을 단계로 그리면 이렇다. 질병 치료 → 장기 교체 → 유전자 교정 → 노화 지연 → 노화 역전? → 신체 증강 → 인지 증강 → 수명 극대화 → 불멸? 앞의 빛은 현실이고 뒤로 갈수록 사변이 섞인다. **그러나 각 단계는 그 앞 단계의 자연스러운 연장으로 보인다.** 그래서 경계가 지워진다.

2. 노화의 열두 표지 — 무엇이 늙는가

노년의학을 이해하려면 먼저 “노화가 무엇인가”를 분자 수준에서 알아야 한다. 2013 년 카를로스 로페스오티와 동료들은 「Cell」에 “노화의 표지(The Hallmarks of Aging)”를 발표하여 아홉 가지를 제시했고, 2023 년 개정판에서 세 가지를 더해 **열두 가지**로 확장했다.⁴⁷ 이 열두는 세 가지 조건을 만족한다 — 정상적 노화 과정에서 나타나고, 실험적으로 악화시키면 노화가 가속되며, 개입으로 완화하면 건강수명이 연장된다.

일차적 표지 — 손상이 축적된다

① **유전체 불안정성(Genomic Instability)**. DNA 는 하루에도 수만 번 손상된다. 자외선, 활성산소, 복제 오류가 원인이다. 젊은 세포는 이를 정확히 복구하지만 복구 기전 자체가 노화하면서 오류가 쌓인다. 조기노화증인 베르너 증후군과 허친슨-길포드 프로제리아 증후군이 DNA 복구와 핵막 구조의 결함에서 비롯된다는 사실이 이 연결을 보여준다.

② **텔로미어 마모(Telomere Attrition)**. 염색체 끝에는 TTAGGG 서열이 반복되는 보호 구조 **텔로미어**가 있다. 신발끈 끝의 플라스틱 캡과 같다. 세포가 분열할 때마다 이 캡이 조금씩 짧아지고, 임계 길이에 도달하면 세포는 더 이상 분열하지 못하고 노화 상태로 들어가거나 사멸한다. 이것을 **헤이플릭 한계(Hayflick limit)**라 부른다.⁴⁸ 생식세포와 줄기세포와 암세포는 **텔로머라제(telomerase)**라는 효소로

⁴⁷ Carlos López-Otín, Maria A. Blasco, Linda Partridge, Manuel Serrano, and Guido Kroemer, “The Hallmarks of Aging,” *Cell* 153, no. 6 (2013): 1194–1217; idem, “Hallmarks of Aging: An Expanding Universe,” *Cell* 186, no. 2 (2023): 243–278. 이 논문은 노화 연구에서 가장 많이 인용되는 문헌 가운데 하나이며, 노화를 하나의 현상이 아니라 상호작용하는 여러 과정의 집합으로 보는 관점을 정립했다.

⁴⁸ Leonard Hayflick and Paul S. Moorhead, “The Serial Cultivation of Human Diploid Cell Strains,” *Experimental Cell Research* 25, no. 3 (1961): 585–621. 헤이플릭은 정상 인간 세포가 약 40–60 회 분열 후 증식을 멈춘다는 것을 관찰했고, 이는 당시의 “세포는 무한히 분열한다”는 통념을 뒤집었다.

텔로미어를 복원한다. 엘리자베스 블랙번과 캐럴 그라이더와 잭 쇼스택은 이 발견으로 2009 년 노벨생리의학상을 받았다. 그러나 여기에 근본적 딜레마가 있다. **텔로머라제를 활성화하면 세포는 젊어지지만 암 위험이 높아진다.** 실제로 인간 암의 약 90%가 텔로머라제를 재활성화한다.

③ **후성유전학적 변화(Epigenetic Alterations).** DNA 서열은 그대로인데 그 위에 붙은 화학적 표지 — DNA 메틸화, 히스톤 변형 — 가 바뀐다. 유전자는 악보이고 후성유전은 연주 방식이다. 나이가 들면 연주가 흐트러진다. 그리고 이 변화가 매우 규칙적이어서, 스티브 호바스는 2013 년 DNA 메틸화 패턴만으로 생물학적 나이를 추정하는 **후성유전학적 시계(epigenetic clock)**를 만들었다.⁴⁹ 이것이 중요한 이유가 있다. 생물학적 나이를 측정할 수 있게 되면서, 노화가 개입 가능한 표적이 되었다.

④ **단백질 항상성 상실(Loss of Proteostasis).** 세포는 단백질을 접고 펴고 분해하는 정교한 품질관리 체계를 갖는다. 이것이 무너지면 잘못 접힌 단백질이 쌓인다. 알츠하이머의 아밀로이드 베타와 타우, 파킨슨의 알파시누클레인이 그 결과다. ⑤ **자가포식 저하(Disabled Macroautophagy).** 세포가 손상된 소기관과 단백질을 스스로 분해해 재활용하는 기능이다. 오스미 요시노리는 이 기전을 밝혀 2016 년 노벨상을 받았다.

길항적 표지 — 방어 기전이 역설적으로 해를 끼친다

⑥ **영양소 감지 이상(Deregulated Nutrient-Sensing).** 세포는 mTOR, AMPK, 인슐린/IGF-1, 시르투인이라는 네 경로로 영양 상태를

⁴⁹ Steve Horvath, “DNA Methylation Age of Human Tissues and Cell Types,” *Genome Biology* 14, no. 10 (2013): R115. 호바스 시계는 353 개의 CpG 부위의 메틸화 상태로 나이를 예측하며, 여러 조직에서 실제 나이와 높은 상관을 보인다. 이후 PhenoAge, GrimAge 등 더 정교한 시계들이 개발되었으나, 서로 다른 시계가 다른 결과를 내놓는 문제는 아직 해결되지 않았다.

감지한다. 이 가운데 **mTOR(mechanistic target of rapamycin)**가 특히 중요하다. 영양이 풍부하면 mTOR 가 활성화되어 성장과 합성을 촉진하고, 영양이 부족하면 억제되어 자가포식과 수리가 활성화된다. **그리고 열량 제한이 여러 동물에서 수명을 연장하는 이유가 여기에 있다.**

⑦ 미토콘드리아 기능 이상(Mitochondrial Dysfunction).

미토콘드리아는 세포의 발전소다. 나이가 들면 효율이 떨어지고 활성산소 누출이 늘어난다. 미토콘드리아 DNA 는 핵 DNA 와 달리 히스톤의 보호를 받지 못하고 복구 기전도 약해 손상이 빨리 쌓인다. ⑧

세포 노화(Cellular Senescence). 손상된 세포가 죽지도 분열하지도 않는 상태로 남는 것이다. 이른바 “좀비 세포”다. 문제는 이 세포들이 조용히 있지 않고 염증성 사이토카인과 단백질분해효소를 분비한다는 데 있다. 이를 **노화 관련 분비 표현형(SASP, Senescence-Associated Secretory Phenotype)**이라 부르며, 주변의 건강한 세포까지 노화시킨다.

통합적 표지 — 조직과 개체가 무너진다

⑨ **줄기세포 고갈(Stem Cell Exhaustion).** 조직을 재생하는 줄기세포 풀이 줄고 기능이 떨어진다. 이것이 상처 치유가 느려지고 근육이 줄고 면역이 약해지는 직접 원인이다. ⑩ **세포 간 소통 변화(Altered Intercellular Communication).** 호르몬과 신경전달과 면역 신호가 어긋난다. ⑪ **만성 염증(Chronic Inflammation).** 노화와 염증을 합쳐 “**인플라메이징(inflammaging)**”이라 부른다. 감염이 없는데도 낮은 수준의 염증이 지속되며 거의 모든 노인성 질환의 공통 배경이 된다. ⑫ **장내 미생물 불균형(Dysbiosis).** 2023 년 개정판에서 추가되었다. 장내 세균총의 구성이 바뀌면서 대사와 면역과 신경계에까지 영향을 미친다.

3. 노화를 되돌리려는 다섯 가지 접근

열두 표지를 알면 개입 지점이 보인다. 현재 연구되는 주요 접근을 다섯으로 정리한다. **다만 이 가운데 무작위 대조시험에서 인간의 수명을**

연장한 것으로 입증된 것은 아직 하나도 없다. 이 점을 먼저 분명히 해 둔다.

① 세놀리틱스 — 좀비 세포를 제거한다

세놀리틱스(senolytics)는 노화 세포를 선택적으로 제거하는 약물이다. 2015 년 메이오클리닉의 제임스 커클랜드 연구진이 다사티닙(dasatinib)과 케르세틴(quercetin)의 조합이 노화 세포를 선택적으로 사멸시킨다는 것을 보고하면서 이 분야가 열렸다.⁵⁰ 2011 년 대런 베이커와 얀 반 되르선의 연구는 유전자 조작 생쥐(INK-ATTAC)에서 노화 세포를 제거하자 노화 관련 병리가 지연되는 것을 보였다.⁵¹

인간 임상에서도 초기 신호가 나왔다. 특발성 폐섬유증 환자에게 다사티닙과 케르세틴을 투여한 소규모 시험에서 보행 거리와 보행 속도 같은 신체 기능이 개선되었다.⁵² 그러나 초기 세놀리틱스에는 부작용 문제가 있었다. 나비토클라스(ABT-263)는 혈소판 감소를 일으킨다. 그래서 최근 연구는 더 정밀한 표적화로 이동하고 있다.

② 후성유전학적 재프로그래밍 — 세포의 시계를 되감는다

이것이 현재 가장 흥미롭고 가장 논쟁적인 분야다. 2006 년 야마나카 신야는 성체 세포에 네 개의 전사인자 — Oct4, Sox2, Klf4, c-Myc(약칭 OSKM 또는 야마나카 인자) — 를 도입하면 배아줄기세포와 유사한 유도만능줄기세포(iPSC, induced Pluripotent Stem Cell)로

⁵⁰Yi Zhu et al., “The Achilles’ Heel of Senescent Cells: From Transcriptome to Senolytic Drugs,” *Aging Cell* 14, no. 4 (2015): 644–658. 다사티닙은 원래 백혈병 치료제이고 케르세틴은 식물에서 발견되는 플라보노이드다. 이 조합을 흔히 “D+Q”라 부른다.

⁵¹Darren J. Baker et al., “Clearance of p16Ink4a-Positive Senescent Cells Delays Ageing-Associated Disorders,” *Nature* 479 (2011): 232–236.

⁵²Jamie N. Justice et al., “Senolytics in Idiopathic Pulmonary Fibrosis: Results from a First-in-Human, Open-Label, Pilot Study,” *EBioMedicine* 40 (2019): 554–563. 이 연구는 대조군이 없는 소규모 개방표지 시험이므로 결과 해석에 신중해야 한다.

되돌아간다는 것을 발견했다. 그는 2012 년 노벨생리의학상을 받았다.⁵³

그런데 여기서 예상하지 못한 발견이 나왔다. 세포를 완전히 되돌리지 않고 야마나카 인자를 짧게만 발현시키면, 세포의 정체성은 유지한 채 후성유전학적 나이만 젊어진다는 것이다. 이를 부분 재프로그래밍(partial reprogramming)이라 부른다. 2016 년 후안 카를로스 이스피수아 벨몬테 연구진이 조기노화 생쥐에서 이를 입증했고,⁵⁴ 2020 년 데이비드 싱클레어 연구진은 손상된 시신경을 재프로그래밍으로 재생시켜 늙은 생쥐의 시력을 회복시켰다.⁵⁵

그러나 위험이 크다. 야마나카 인자를 지속적으로 발현시키면 기형종(teratoma)이 생기고 개체가 사망한다. 특히 c-Myc 은 강력한 발암유전자다. 그래서 현재 연구는 c-Myc 을 빼거나(OSK), 화학물질로 대체하거나, 발현 시간을 정밀하게 조절하는 방향으로 가고 있다. 그리고 결정적으로, 인간에서의 안전성과 효과는 아직 입증되지 않았다.

③ 대사 경로 조절 — 라파마이신과 메트포르민

라파마이신(rapamycin)은 이스터섬(라파누이)의 토양 세균에서 발견된 물질로 원래 면역억제제였다. 그런데 mTOR 를 억제하여 여러 동물에서 수명을 연장한다는 것이 밝혀졌다. 미국 국립노화연구소의 중재점사프로그래밍(ITP)에서 라파마이신은 생쥐의 수명을 유의하게 연장한 몇 안 되는 물질이다. 메트포르민(metformin)은 당뇨병 치료제로 AMPK 를 활성화한다. 니르 바질라이가 주도하는

⁵³ Kazutoshi Takahashi and Shinya Yamanaka, "Induction of Pluripotent Stem Cells from Mouse Embryonic and Adult Fibroblast Cultures by Defined Factors," Cell 126, no. 4 (2006): 663-676. 이 발견은 배아를 파괴하지 않고도 만능세포를 얻을 수 있게 하여 윤리적 논쟁을 크게 완화했다.

⁵⁴ Alejandro Ocampo et al., "In Vivo Amelioration of Age-Associated Hallmarks by Partial Reprogramming," Cell 167, no. 7 (2016): 1719-1733.

⁵⁵ Yuan Cheng Lu et al., "Reprogramming to Recover Youthful Epigenetic Information and Restore Vision," Nature 588 (2020): 124-129. 이 연구는 c-Myc 을 제외한 세 인자(OSK)만 사용하여 암 위험을 낮췄다.

TAME(Targeting Aging with Metformin) 임상시험은 노화 자체를 표적으로 하는 최초의 대규모 시험으로 설계되었으나, 2026 년 현재까지도 자금과 규제 문제로 완전히 진행되지 못했다.⁵⁶ 이 사실 자체가 중요한 것을 알려 준다 — 노화를 치료 대상으로 삼으려면 과학만이 아니라 규제 체계 자체가 바뀌어야 한다.

④ NAD⁺ 보충과 시르투인

NAD⁺(nicotinamide adenine dinucleotide)는 세포의 에너지 대사와 DNA 복구에 필수적인 조효소인데 나이가 들면서 감소한다. 그래서 NMN(nicotinamide mononucleotide)과 NR(nicotinamide riboside) 같은 전구체를 보충하는 연구가 활발하다. 시르투인(sirtuin)은 NAD⁺에 의존하는 효소군으로 유전자 발현과 DNA 복구를 조절한다. 다만 인간에서의 임상적 효과는 아직 논쟁적이며, 상업적 과장이 과학적 근거를 앞서는 대표적 영역이기도 하다.

⑤ 혈액 인자와 이종병체결합

이종병체결합(parabiosis)은 젊은 개체와 늙은 개체의 순환계를 연결하는 실험이다. 늙은 생쥐가 젊어지는 현상이 관찰되면서 “젊은 피 속의 어떤 인자”를 찾는 연구가 이어졌다. 다만 이 효과가 젊은 피의 유익한 인자 때문인지, 늙은 피의 해로운 인자가 희석되기 때문인지는 아직 논쟁 중이다.

4. 노년의학이 던지는 두 개의 물음

이 모든 연구를 종합하면 두 개의 물음이 남는다.

첫째, 건강수명인가 총수명인가. 노년의학의 공식 목표는 대체로 건강수명(healthspan)의 연장이다. 곧 오래 사는 것이 아니라 건강하게

⁵⁶ TAME 시험의 의의는 결과보다 설계에 있다. 미국 FDA는 “노화”를 질병으로 인정하지 않으므로 노화를 적응증으로 하는 약물 승인 경로가 존재하지 않는다. TAME는 복합 결과지표(여러 노인성 질환의 발생 지연)를 사용하여 이 규제적 장벽을 우회하려는 시도였다.

사는 기간을 늘리는 것이다. 이것은 대부분의 사람이 반대하지 않는 목표다. 그런데 실제 연구의 상당 부분은 **총수명(lifespan)**의 연장을 함께 겨냥한다. 그리고 언론과 투자 담론은 그 경계를 흐린다.

둘째, 어디까지가 치료인가. 알츠하이머를 예방하는 것은 치료다. 그런데 90 세에 30 세의 인지 능력을 갖게 하는 것은 치료인가 증강인가. **노화를 “질병”으로 규정하는 순간 이 경계는 흐려진다.** 그리고 이것이 본서가 제 11 장에서 다룬 “치료에서 증강으로, 증강에서 진화로”의 실제 현상이다.

2. 의학의 목표 자체가 이동한다

같은 이동을 의학의 목표로 그리면 더 선명해진다. **질병 치료 → 예방 → 재생 → 장수 → 수명 연장 → 회춘 → ?** 마지막 칸이 비어 있는 이유가 있다. 그다음이 무엇인지 아무도 확정하지 못하기 때문이다.

여기서 노년의학(Geroscience)의 가설이 결정적이다. 노화가 여러 만성질환의 공통 위험인자라는 것이다. 미국 국립보건원의 범부처 노년의학 관심그룹이 2012 년 출범했고, 2014 년 「Cell」의 논문이 개념을 대표적으로 정립했다.⁵⁷ 그렇다면 질문이 바뀐다 — **병을 하나씩 치료할 것인가, 아니면 병들게 만드는 과정 자체를 치료할 것인가.**

의학의 표적이 “질병”에서 “노화”로 옮겨가면, 인간의 시간에 대한 이해도 함께 바뀐다. 그리고 물음이 남는다.

어디까지가 치료이고, 어디부터가 인간의 재설계인가?

⁵⁷ Brian K. Kennedy et al., “Geroscience: Linking Aging to Chronic Disease,” Cell 159, no. 4 (2014): 709–713. 노화의 표지에 관하여는 Carlos López-Otín et al., “The Hallmarks of Aging,” Cell 153, no. 6 (2013): 1194–1217 참조.

제 11 장

치료에서 증강으로, 증강에서 진화로

이 장은 하나의 예로 시작한다. 인공와우다.

0. 전통 의학은 무엇을 했는가

경계를 논하기 전에 출발점을 확인해야 한다. **의학은 본래 무엇을 하는 일이었는가.**

히포크라테스 전통에서 의학의 목표는 **“자연의 회복”**이었다. 몸에는 스스로 낫는 힘(vis medicatrix naturae)이 있고 의사는 그것을 돕는다. 그래서 고전 의학의 원칙은 **“무엇보다 해를 끼치지 말라(primum non nocere)”**였다. 개입은 최소한이어야 했다.

19 세기까지 의사가 가진 도구는 제한적이었다. 골절을 맞추고, 상처를 봉합하고, 사혈하고, 약초를 처방했다. **그리고 결정적으로, 의학이 개입하는 대상은 “질병”이었지 “인간”이 아니었다.** 병을 몰아내면 인간은 원래의 자기로 돌아간다. 이것이 **회복(restoration)의 패러다임**이다.

20 세기에 이 패러다임이 확장되었다. 항생제가 감염을 정복했고, 백신이 질병을 예방했으며, 마취와 무균술이 외과를 가능하게 했다. 그러나 목표는 여전히 같았다. **정상으로 되돌리는 것이다.** 인공관절도, 심박조율기도, 안경도 **“잃어버린 기능의 회복”**이라는 틀 안에 있었다.

0-1. 언제 경계가 흔들리기 시작했는가

그런데 몇 가지 사례에서 이 틀이 흔들리기 시작했다.

성장호르몬. 성장호르몬 결핍증 아이에게 호르몬을 투여하는 것은 명백한 치료다. 그런데 결핍증이 없으나 키가 작은 아이에게 투여하는 것은 무엇인가. 미국 FDA 는 2003 년 **“특별성 저신장”**에 대한

성장호르몬 사용을 승인했다. **질병이 아닌 상태에 대한 의학적 개입이 공식화된 것이다.**

정신약물. 우울증 치료제는 치료다. 그런데 우울증이 없는 사람의 기분을 개선하는 것은 무엇인가. 피터 크레이머는 1993 년 「프로작에 귀 기울이기」에서 **“미용 정신약리학(cosmetic psychopharmacology)”**이라는 표현을 만들었다.⁵⁸ 환자들이 “병이 나왔다”가 아니라 “원래의 나보다 나아졌다”고 말할 때 의학은 무엇을 한 것인가.

인공와우. 그리고 가장 논쟁적인 사례가 여기 있다. 청각장애인 공동체 일부는 인공와우 이식을 **치료가 아니라 문화의 소멸**로 본다. 수어를 사용하는 농인 문화가 하나의 언어 공동체라면, 모든 아이에게 인공와우를 이식하는 것은 그 언어의 소멸을 뜻하기 때문이다. **여기서 “무엇이 정상인가”라는 물음이 의학 안에서 터져 나왔다.**

0-2. 세 단계의 논리 — 왜 미끄러지는가

이제 세 단계를 정확히 정의하자.

치료(therapy)는 종 전형적 기능(species-typical functioning)을 회복하는 것이다. 노먼 대니얼스의 “정상 기능 모델”이 이 정의의 표준이다.⁵⁹ **증강(enhancement)**은 종 전형적 범위를 넘어서는 것이다. **진화(evolution)**는 그 변화가 유전되어 후속 세대의 종 전형성 자체를 바꾸는 것이다.

그런데 이 정의에는 치명적 약점이 있다. **“종 전형적”이라는 기준 자체가 통계적이고 역사적이기 때문이다.** 평균 수명이 40 세이던

⁵⁸ Peter D. Kramer, *Listening to Prozac* (New York: Viking, 1993). 크레이머는 환자들이 약물 복용 후 “원래의 나보다 더 나은 나”가 되었다고 말하는 현상을 관찰하고, 이것이 치료인지 변형인지를 물었다.

⁵⁹ Norman Daniels, *Just Health Care* (Cambridge: Cambridge University Press, 1985). 대니얼스는 정상 기능의 회복이 기회의 평등을 보장하는 조건이므로 정의(justice)의 문제라고 논증한다.

시대에 80 세까지 사는 것은 중 전형적이지 않았다. 그러나 지금은 정상이다. 백신 접종으로 얻은 면역은 자연적인가. 안경으로 교정한 시력은 중 전형적인가. **기준선이 움직이면 치료와 증강의 경계도 함께 움직인다.**

그래서 미끄러진다. 논리적 순서는 이렇다. **첫째, 명백한 질병을 치료한다.** 아무도 반대하지 않는다. **둘째, 경계선상의 상태를 치료한다.** 저신장, 경도인지장애, 갱년기. **셋째, 위험을 예방한다.** 아직 병이 아니지만 병이 될 가능성을 낮춘다. **넷째, 노화를 늦춘다.** 노화는 모두에게 오지만 그것을 질병으로 규정하면 개입 대상이 된다. **다섯째, 최적화한다.** 정상 범위 안에서도 더 나은 상태를 추구한다. 각 단계는 앞 단계에서 한 걸음일 뿐이다.

0-3. Human + Machine — 이미 시작된 결합

그리고 이 미끄러짐은 약물만의 이야기가 아니다. **인간과 기계의 결합**에서 같은 일이 일어난다.

오늘 이미 수백만 명이 몸 안에 기계를 넣고 산다. **심박조율기**는 심장의 전기 신호를 대신 만든다. **인공와우**는 소리를 전기 신호로 바꾸어 청신경에 직접 전달한다. **인공관절**은 뼈를 대신한다. **심부뇌자극술**은 뇌에 전극을 심어 파킨슨병의 떨림을 멈춘다. **그리고 우리는 이 사람들을 “사이보그”라 부르지 않는다.**

“사이보그(cyborg)”라는 낱말은 1960 년 만프레드 클라인스와 네이션 클라인이 만들었다. **cybernetic organism** 의 줄임말로, 원래는 우주 환경에 적응하기 위해 신체를 기술적으로 보강한 인간을 가리켰다.⁶⁰ 흥미롭게도 이 개념은 처음부터 “지구를 떠나기 위한 인간

⁶⁰Manfred E. Clynes and Nathan S. Kline, “Cyborgs and Space,” *Astronautics* (September 1960): 26-27, 74-76. 이들의 관심은 우주비행사가 우주 환경에 맞추어 스스로를 조절할 필요 없이 살 수 있게 하는 데 있었다.

개조”로 제안되었다. 본서 제 15 장의 화성 논의와 60 년 전에 이미 연결되어 있었던 것이다.

그렇다면 심박조율기와 뇌-컴퓨터 인터페이스는 무엇이 다른가. 세 가지가 다르다. **첫째, 개입의 위치다.** 심장은 펌프이지만 뇌는 인격의 자리다. **둘째, 정보의 방향이다.** 심박조율기는 신호를 보내기만 하지만 BCI 는 읽고 쓴다. **셋째, 학습의 유무다.** 심박조율기는 정해진 대로 작동하지만 BCI 의 디코딩 알고리즘은 학습하고 적응한다. **그러므로 후자에서는 “누가 결정했는가”가 흐려진다.**

0-4. 트랜스휴먼과 포스트휴먼 — 개념의 계보

이제 두 낱말을 정확히 구별하자. 흔히 혼용되지만 다른 개념이다.

트랜스휴먼(transhuman)은 “인간을 넘어가는 중”을 뜻한다. 라틴어 trans(넘어)와 human 의 결합이다. 여전히 인간이되 기술로 증강된 상태다. 단테가 「신곡」 천국편에서 사용한 **trasumanar** 라는 조어가 어원적 뿌리로 자주 인용되는데, 단테에게 그것은 신적 은총으로 인간을 초월하는 것이었다. **현대 트랜스휴머니즘은 그 초월의 수단을 은총에서 기술로 바꾸었다.**

포스트휴먼(posthuman)은 “인간 이후”를 뜻한다. 더 이상 호모 사피엔스라 부를 수 없을 만큼 근본적으로 달라진 존재다. 닉 보스트롬은 포스트휴먼을 “현재 인간이 도달할 수 있는 최대치를 크게 넘어서는 능력을 최소한 하나 이상 가진 존재”로 정의한다.⁶¹

트랜스휴머니즘의 역사도 짧지 않다. 생물학자 줄리언 헉슬리가 1957 년 “**transhumanism**”이라는 낱말을 오늘의 의미로 처음 사용했다. 그는 인류가 “인간으로 남되 자기 자신을 초월함으로써 새로운

⁶¹Nick Bostrom, “Why I Want to Be a Posthuman When I Grow Up,” in Medical Enhancement and Posthumanity, ed. Bert Gordijn and Ruth Chadwick (Dordrecht: Springer, 2008), 107-136. 보스트롬이 드는 세 능력은 건강수명, 인지 능력, 감정 능력이다.

가능성을 실현”할 수 있다고 썼다.⁶² 흥미롭게도 그는 「멋진 신세계」를 쓴 올더스 헉슬리의 형이었다. **한 형제가 인간 향상의 유토피아와 디스토피아를 각각 그린 것이다.**

0-5. 찬반의 논증들

이 문제에 대한 학계의 논쟁은 성숙해 있다. 양쪽을 정확히 들어야 한다.

옹호 측 — 보스트롬과 새블레스쿠. 닉 보스트롬은 유명한 우화 「용 폭군 이야기」에서 노화를 매일 수만 명을 잡아먹는 용에 비유했다. 사람들은 용의 존재에 익숙해져 그것을 자연스러운 것으로 받아들이고 오히려 용을 죽이려는 시도를 불경하다고 여긴다는 것이다.⁶³ 줄리안 새블레스쿠는 한 걸음 더 나아가 “**생식적 선행(procreative beneficence)**”을 주장한다. 부모가 자녀에게 최선의 삶을 줄 도덕적 의무가 있다면, 가능한 유전적 개선을 하지 않는 것이 오히려 잘못이라는 것이다.⁶⁴

반대 측 — 후쿠야마와 샌텔과 하버마스. 프랜시스 후쿠야마는 인간 향상을 “세계에서 가장 위험한 사상”이라 불렀다. 그의 논거는 “**인간 본성(Factor X)**”이 인권의 근거인데 향상이 그것을 해체한다는 것이다.⁶⁵ 마이클 샌텔은 앞서 본 대로 “**주어진성(giftedness)**”의 상실을 지적하고, 위르겐 하버마스는 **세대 간 대칭성의 파괴**를 문제 삼는다. 설계된 사람은 자기 존재의 기원에 대해 설계자와 대등한 관계를 가질 수 없다는 것이다.

⁶² Julian Huxley, “Transhumanism,” in *New Bottles for New Wine* (London: Chatto & Windus, 1957), 13–17. 줄리언 헉슬리는 「멋진 신세계」를 쓴 올더스 헉슬리의 형이며, 유네스코 초대 사무총장이었다.

⁶³ Nick Bostrom, “The Fable of the Dragon-Tyrant,” *Journal of Medical Ethics* 31, no. 5 (2005): 273–277.

⁶⁴ Julian Savulescu, “Procreative Beneficence: Why We Should Select the Best Children,” *Bioethics* 15, no. 5–6 (2001): 413–426.

⁶⁵ Francis Fukuyama, *Our Posthuman Future: Consequences of the Biotechnology Revolution* (New York: Farrar, Straus and Giroux, 2002); idem, “Transhumanism,” *Foreign Policy* 144 (2004): 42–43.

그리고 반대 측이 가장 강하게 제기하는 실천적 논거가 “**증강 격차(enhancement divide)**”다. 증강 기술은 비쌌 것이고, 비싸면 부유한 사람만 접근한다. 그리고 그 증강이 유전된다면 **불평등이 생물학적으로 고착된다**. 리 실버는 이를 “유전부유층(GenRich)”과 “자연인(Naturals)”의 분화로 그렸다.⁶⁶

0-6. 그러므로 물음은 남는다

본서는 어느 한쪽을 단순히 편들지 않는다. 대신 구별한다. **질병을 치료하고 고통을 줄이는 기획은 정당하다**. 성경은 치유를 선택한 것으로 보며 예수의 사역 상당 부분이 치유였다. **그러나 인간의 조건 자체를 극복 대상으로 삼는 기획은 다른 문제다**.

그리고 그 경계는 누가 긋는가. 개인의 선택인가, 국가의 정책인가, 시장의 논리인가, 아니면 신학의 물음인가. **본서의 대답은 마지막이다**. 왜냐하면 앞의 셋은 모두 “무엇을 할 수 있는가”와 “누가 원하는가”만 물을 수 있고, “**인간이란 무엇인가**”는 묻지 못하기 때문이다.

1. 인공와우의 사고실험

청각을 잃은 사람에게 인공와우를 이식하는 것은 **치료**다. 아무도 반대하지 않는다. 그런데 같은 기술로 정상 청력을 넘어 초음파까지 듣게 한다면 그것은 **증강**이다. 그리고 만약 그 감각이 유전적으로 자손에게 전해진다면 그것은 **진화**다.

TREATMENT → ENHANCEMENT → EVOLUTION

각 단계는 그 앞 단계의 자연스러운 연장으로 보인다. 회복하는 것에서, 주어진 것을 넘어서는 것으로, 그리고 다른 무엇이 되는 것으로. **그래서**

⁶⁶ Lee M. Silver, Remaking Eden: Cloning and Beyond in a Brave New World (New York: Avon Books, 1997). 실버는 이 시나리오를 옹호한 것이 아니라 규제가 없을 경우의 귀결로 제시했다.

경계가 지워진다. 그리고 어느 지점에서 “치료받은 인간”이 “다른 존재”가 되는지 아무도 정하지 못한다.

2. 누가 그 선을 긋는가

그렇다면 그 경계선은 누가 정하는가. 네 가지 답이 가능하다. **개인의 선택인가** — 그렇다면 원하는 사람은 무엇이든 해도 되는가. **국가의 정책인가** — 그렇다면 국가가 인간의 정의를 규정하는가. **시장의 논리인가** — 그렇다면 감당할 수 있는 사람만 넘어서는가. 그리고 **신학의 물음인가** — 누가 인간을 정의하는가.

마이클 샌델은 인간 향상 기획을 비판하면서 안전성이나 공정성이 아니라 더 깊은 곳을 겨눈다. 향상의 기획은 “주어진 것(the given)”을 “설계할 것(the designed)”으로 바꾸며, 그 순간 우리가 삶을 대하는 태도 자체가 달라진다는 것이다.⁶⁷ 신학적으로 이것은 피조물임(creatureliness)의 문제다. **삶이 선물이라면 감사가 가능하다. 삶이 설계라면 감사할 대상이 없다.**

3. 유한성은 죄가 아니다

여기서 본서는 하나의 신학적 구별을 세운다. 인간은 모든 것을 알지 못하고, 모든 곳에 있을 수 없고, 피곤하고, 잠을 자고, 늙고, 죽는다. AI 문명은 이 목록을 결함처럼 취급할 가능성이 있다. 그러나 성경은 인간의 유한성을 결함이 아니라 피조물의 조건으로 본다.

Finitude ≠ Sin 유한성은 타락이 아니다

인간은 죄인이기 전에 피조물이다. 이 구별이 없으면 트랜스휴머니즘 비판은 곧바로 무너진다. “그렇다면 의학 발전도 반대하느냐”는 반론에 답할 수 없기 때문이다. 그러므로 두 프로젝트를 구별해야 한다. **인간의**

⁶⁷ Michael J. Sandel, The Case against Perfection: Ethics in the Age of Genetic Engineering (Cambridge, MA: Harvard University Press, 2007). 샌델은 향상 기획이 겸손(humility)·책임(responsibility)·연대(solidarity)라는 세 가지 도덕적 자산을 잠식한다고 논증한다.

질병과 고통을 치료하려는 프로젝트는 정당하다. 성경은 치유를 선택한 것으로 보며 예수의 사역 상당 부분이 치유였다. 그러나 **인간을 하나님처럼 만들려는 프로젝트는 다른 문제다.** 물론 치료와 향상의 경계는 흐리다. 그러나 흐리다고 해서 구별이 무의미한 것은 아니다.

제 12 장

교체되는 몸 — 이종이식과 뇌-컴퓨터 인터페이스

이제 사변이 아니라 2026 년 현재의 좌표를 확인한다.

1. 고치는 것에서 길러 내는 것으로

몸을 다루는 의학은 다섯 단계로 이동해 왔다. **고친다(REPAIR)** — 약물과 수술. **재생시킨다(REGENERATE)** — 줄기세포와 조직공학. **교체한다(REPLACE)** — 인공장기와 이종이식. **길러 낸다(GROW)** — 오가노이드와 바이오프린팅. 그리고 **설계한다(DESIGN)** — 합성생물학과 유전자 편집.

인간의 몸에는 스스로 재생하는 장기가 있고 그렇지 못한 장기가 있다. 간과 피부와 혈액과 장 상피는 상당한 재생 능력을 갖지만, 심장 근육과 중추신경은 제한적이다. 그런데 이제 재생하지 못하는 것까지 길러 내고 설계하려 한다.

0. 이종이식이란 무엇인가 — 어원과 역사

이종이식(xenotransplantation)은 헬라어 **크세노스(ξένος)**에서 왔다. 이 낱말은 “낯선 자, 이방인, 손님”을 뜻한다. 신약성경에서 “나그네 되었을 때에 영접하였고”(마 25:35)의 “나그네”가 바로 이 낱말이며, 히브리서 13 장 2 절의 “손님 대접하기를 잊지 말라”도 여기서 나온 필록세니아(φιλοξενία)다. 그러므로 **이종이식은 어원적으로 “낯선 것을 몸 안에 받아들이는 일”이다**. 그리고 의학적으로 정확히 그것이 문제의 핵심이다 — 몸은 낯선 것을 알아보고 거부한다.

의학적 정의는 이렇다. 다른 종의 살아 있는 세포·조직·장기를 인간에게 이식하거나, 인간의 체액이 다른 종의 살아 있는 조직과 접촉한 뒤

인체에 되돌리는 모든 기술. 미국 FDA의 정의가 이렇게 넓은 이유가 있다. 감염 위험 때문이다.

0-1. 200 년의 실패사

이종이식의 역사는 실패의 역사였고, 그 실패가 무엇을 가르쳤는지가 중요하다.

17 세기. 1667 년 프랑스의 장바티스트 드니가 양의 피를 인간에게 수혈했다. 환자는 처음에는 견뎠으나 이후 사망했고, 프랑스는 수혈 자체를 금지했다. **20 세기 초.** 1906 년 마티외 자블레가 돼지와 염소의 신장을 인간에게 이식했다. 며칠 만에 실패했다. **1960 년대.** 툴레인 대학의 키스 림즈마가 침팬지 신장을 13 명에게 이식했고, 그중 한 여성은 9 개월을 생존했다. 당시로서는 놀라운 기록이었다.

1984 년 — 베이비 페이(Baby Fae). 좌심형성부전증으로 태어난 생후 12 일의 여아에게 개코원숭이의 심장을 이식했다. 아기는 21 일을 살았다. 이 사건은 전 세계적 논쟁을 불러일으켰고, 이후 영장류를 공여체로 사용하는 것에 대한 회의가 커졌다.⁶⁸

그래서 돼지가 선택되었다. 이유는 네 가지다. **첫째, 장기의 크기와 생리가 인간과 비슷하다.** 돼지 심장은 인간 심장과 크기와 박출량이 유사하다. **둘째, 번식이 빠르고 한 배에 많이 낳는다.** **셋째, 이미 식용으로 대량 사육되고 있어 윤리적 저항이 상대적으로 낮다.** **넷째, 유전적으로 멀어서 인수공통감염 위험이 영장류보다 낮다.**

그러나 바로 그 유전적 거리가 거부반응을 격렬하게 만든다. **그리고 이것이 이종이식의 근본 역설이다 — 가까우면 감염이 위험하고 멀면 거부가 심하다.** 이 역설을 푼 것이 유전자 편집이었다.

⁶⁸ 베이비 페이 사례는 의학사에서 인정된 동의(informed consent)와 연구 윤리 논쟁의 분기점이 되었다. 이후 영장류는 인간과 유전적으로 가까워 감염 전파 위험이 크고, 번식이 느리며, 윤리적 논란이 크다는 이유로 공여체 후보에서 사실상 배제되었다.

2. 이종이식 — 임상시험에 들어갔다

이종이식(Xenotransplantation)의 좌표를 정확히 그려 보자. 2021 년 9 월 최초의 유전자 편집 돼지 장기가 뇌사자에게 이식되었고, 2022 년부터 2024 년까지 심장 이식 사례들이 이어졌다(각각 2 개월, 40 일 생존). 그리고 2025 년 11 월 EXPAND 임상시험이 개시되었다. 10 개 유전자를 편집한 돼지 신장을 사용하는 세계 최초의 규제된 임상 이종이식 프로그램이다. 2026 년에는 돼지-인간 신장 이식에서 최장 271 일 생존이 보고되었다.⁶⁹

그러나 남은 장벽이 크다. 면역거부와 인수공통 감염 위험과 생리적 부적합, 그리고 장기 생존성이다. 한 분석은 예상 이식편 생존이 2 년을 넘지 않으면 대상 환자를 선정하기조차 어렵다고 지적한다. **그럼에도 방향은 분명하다. 종의 경계가 의학적으로 넘어지고 있다.**

그리고 신학적 물음이 남는다. **돼지의 장기로 사는 인간은 여전히 인간인가.** 대답은 “그렇다”이다. 아무도 이것을 의심하지 않는다. 그러나 **왜 그러한가를 말할 수 있는가.** 만약 인간됨이 생물학적 구성물의 목록이라면 그 목록이 바뀔 때 인간됨도 바뀌어야 한다. 그렇지 않다면 인간됨은 구성물이 아닌 다른 무언에 근거하는 것이다. 본서 제 8 부가 그 답을 다룬다.

4. 냉동보존의 실제 — 무엇이 문제인가

냉동보존을 정확히 이해하려면 먼저 **왜 열리면 안 되는가**를 알아야 한다. 세포의 70% 이상은 물이다. 물이 얼면 얼음 결정이 생기고, 그 결정이 세포막과 소기관을 물리적으로 찢는다. 그래서 냉동보존의 핵심 과제는 **“얼지 않게 차갑게 만드는 것”**이다.

⁶⁹ 이종이식의 임상 진전에 관하여는 2025-2026 년의 관련 보고와 리뷰를 참조하라. 이 분야는 매우 빠르게 변하므로 인용 시점에 원자료를 확인해야 한다.

그 해법이 **유리화(vitrification)**다. 물에 고농도의 **동결보호제(cryoprotectant, CPA)** — 디메틸설폭사이드(DMSO), 에틸렌글리콜, 글리세롤 등 — 를 넣고 매우 빠르게 냉각하면, 물 분자가 결정을 이룰 시간을 갖지 못하고 유리처럼 무정형 고체가 된다. 얼음이 아니라 유리가 되는 것이다. 이 기술은 이미 정자와 난자와 배아의 보존에 임상적으로 널리 사용된다.

그러나 인체 전체에 적용하는 것은 전혀 다른 문제다. 세 가지 장벽이 있다. **첫째, 동결보호제 자체가 독성이 있다.** 얼음을 막을 만큼 높은 농도는 세포에 해롭다. **둘째, 관류가 균일하지 않다.** 크기가 큰 조직은 중심부까지 CPA가 고르게 도달하지 않는다. **셋째, 열응력이 있다.** 급속 냉각과 해동 과정에서 조직 내부와 표면의 온도 차이가 균열을 일으킨다. 이것을 “유리화 균열(fracturing)”이라 부른다.

그래서 냉동보존의 현재 좌표는 이렇다. **세포와 배아 수준에서는 확립된 기술이다.** 신장 같은 작은 장기의 유리화와 재관류가 동물 실험에서 성공한 사례도 보고되었다. 2024년에는 인간 뇌 조직과 뇌 오가노이드를 구조와 기능을 보존한 채 냉동보존하는 방법(MEDY)이 보고되었다.⁷⁰ **그러나 인체 전체를 냉동보존한 뒤 소생시킨 사례는 없다.**

그러므로 정확한 표현은 이것이다. **냉동보존은 치료법이 아니라 미래 기술에 대한 내기(wager)다.** 알코어(Alcor)와 크라이오닉스 연구소 같은 기관은 이를 숨기지 않는다. 그들의 논리는 이렇다 — 화장이나 매장은 확실히 되돌릴 수 없지만 냉동보존은 확률이 0이 아니다. 이것은 파스칼의 내기와 구조가 같은 논증이며, 과학적 주장이라기보다 확률적 선택에 관한 주장이다.

⁷⁰Weiwei Xue et al., “Effective Cryopreservation of Human Brain Tissue and Neural Organoids,” Cell Reports Methods 4, no. 5 (2024).

메틸셀룰로스·에틸렌글리콜·DMSO·Y27632를 조합한 방법으로, 해동 후 축삭이 성장하고 병리적 특징이 보존되었다.

여기서 구별해야 할 것이 있다. **인공 동면(induced torpor/hibernation)** 연구는 냉동보존과 전혀 다르다. 후자가 사망 후의 보존이라면 전자는 살아 있는 상태에서 대사율을 크게 낮추는 것이다. 곰과 다람쥐 같은 동면 동물의 기전을 연구하여 장거리 우주비행과 중증외상 응급의학에 적용하려는 시도로, 실제 임상 연구가 진행 중이다. **전자는 사변이고 후자는 현재진행형 의학이다.**

5. 이종이식의 유전공학

돼지의 장기를 인간에게 이식하려면 무엇을 해결해야 하는가. 세 층위의 거부반응이 있다.

첫째, 초급성 거부반응(hyperacute rejection). 인간의 혈액에는 돼지 세포 표면의 특정 당 항원 — 특히 **알파-갈(α -Gal, galactose- α -1,3-galactose)** — 에 대한 자연항체가 이미 존재한다. 이식하는 순간 항체가 붙고 보체가 활성화되어 몇 분에서 몇 시간 만에 장기가 괴사한다. **둘째, 급성 체액성 거부반응.** 며칠에서 몇 주에 걸쳐 진행된다. **셋째, 만성 거부반응과 응고 이상.** 돼지와 인간의 응고 조절 단백질이 서로 호환되지 않아 혈전이 생긴다.

CRISPR-Cas9 유전자 편집이 이 문제를 다루는 방법은 두 가지다. **첫째, 돼지 쪽의 항원 유전자를 제거한다(knockout).** α -Gal 을 만드는 GGTA1 유전자와 두 개의 다른 당 항원 유전자(CMAH, B4GALNT2)를 제거하면 이른바 “3 중 녹아웃(triple knockout)” 돼지가 된다. **둘째, 인간의 조절 유전자를 넣는다(knock-in).** 인간 보체 조절 단백질(CD46, CD55)과 응고 조절 단백질(트롬보모듈린)의 유전자를 삽입한다. 현재 임상시험에 쓰이는 돼지는 이런 편집을 10 개 내외 받은 개체다.

여기에 또 하나의 문제가 있다. **돼지 내인성 레트로바이러스(PERV, Porcine Endogenous Retrovirus)**다. 돼지 유전체에 통합되어 있는 바이러스 서열로, 이론상 인간에게 전파될 위험이 있다. 2017 년 조지 처치 연구진은 CRISPR 로 돼지 세포의 PERV 62 개를 모두

불활성화하는 데 성공했다.⁷¹ 이 성취가 이종이식을 임상 단계로 밀어 올린 결정적 계기였다.

그럼에도 남은 과제가 크다. 앞서 본 271 일 생존 기록은 놀라운 성취이지만, 신장이식의 표준인 5 년 생존율과 비교하면 아직 멀다. 그리고 여기에는 기술을 넘어선 물음이 있다. 이종이식이 확립되면 인간의 장기 부족 문제는 해결되겠지만, 그것을 위해 유전자를 편집한 동물을 대량으로 사육하고 도살하는 체계가 필요해진다. 창조세계 청지기직의 관점에서 이것을 어떻게 볼 것인가는 별도의 신학적 물음이다.

6. 유전자 편집 — CRISPR 가 실제로 하는 일

CRISPR-Cas9 는 본래 세균의 면역 체계다. 세균은 침입한 바이러스의 DNA 조각을 자기 유전체의 특정 구역(CRISPR)에 기록해 두었다가, 같은 바이러스가 다시 오면 그 기록을 바탕으로 Cas 단백질을 보내 바이러스 DNA 를 잘라 낸다. 2012 년 제니퍼 다우드나와 에마누엘 샤르팡티에는 이 체계를 프로그래밍 가능한 유전자 가위로 재설계했고 2020 년 노벨화학상을 받았다.⁷²

작동 원리는 간단하다. **가이드 RNA** 가 표적 서열을 찾아가고 **Cas9 단백질**이 그 자리에서 DNA 를 자른다. 그러면 세포의 복구 기전이 작동하는데, 여기에 두 경로가 있다. **비상동 말단 접합(NHEJ)**은 빠르지만 부정확해 유전자를 망가뜨리는 데 쓰이고, **상동 재조합 복구(HDR)**는 정확한 주형을 제공하면 원하는 서열로 바꿀 수 있다.

그리고 여기서 결정적 구별이 필요하다. **체세포 편집(somatic editing)**과 **생식세포 편집(germline editing)**이다. 전자는 환자 본인의

⁷¹Dong Niu et al., “Inactivation of Genes in Pig Cells Using CRISPR-Cas9,” *Science* 357, no. 6357 (2017): 1303-1307. 62 개 유전자좌를 동시에 편집한 것은 당시 최대 규모의 유전체 편집이었다.

⁷²Martin Jinek et al., “A Programmable Dual-RNA-Guided DNA Endonuclease in Adaptive Bacterial Immunity,” *Science* 337, no. 6096 (2012): 816-821.

세포만 바꾸므로 자손에게 전달되지 않는다. 2023 년 승인된 겸상적혈구빈혈 치료제 카스게비(Casgevy)가 그 예다. 후자는 배아나 생식세포를 편집하므로 모든 후손에게 영원히 전달된다.

2018 년 중국의 허젠쿠이는 HIV 저항성을 목표로 배아의 CCR5 유전자를 편집해 쌍둥이를 출생시켰다. 국제 과학계는 이를 일제히 규탄했고 그는 실형을 선고받았다. 그러나 이 사건이 드러낸 것은 기술적 장벽이 아니라 규범적 장벽만이 남아 있다는 사실이었다.⁷³ 제 11 장에서 물었던 “누가 그 선을 긋는가”가 여기서 가장 첨예하게 나타난다.

0. BCI 란 무엇인가 — 정의와 계보

뇌-컴퓨터 인터페이스(BCI, Brain-Computer Interface)는 “뇌의 활동을 직접 읽어 외부 장치를 제어하거나, 외부 신호를 뇌에 직접 전달하는 시스템”이다. 핵심은 “직접(direct)”이다. 근육을 거치지 않는다. 생각이 손을 움직이고 손이 마우스를 움직이는 것이 아니라, 생각이 곧바로 신호가 된다.

그러므로 BCI 의 표준 정의에는 네 요소가 들어간다. ① 뇌 신호를 획득한다. ② 특징을 추출한다. ③ 의도로 변환(디코딩)한다. ④ 장치를 제어하거나 피드백을 준다. 그리고 여기에 하나가 더 필요하다 — 사용자의 학습이다. 뇌도 인터페이스에 적응한다. 그래서 BCI 는 “기계가 인간을 읽는 것”이 아니라 “인간과 기계가 서로 학습하는 것”이다.

0-1. 역사 — 뇌파의 발견에서 오늘까지

⁷³ 허젠쿠이 사건에 관한 과학계의 대응으로는 국제 인간 유전체 편집 정상회의의 성명들과 각국 학술원의 공동 보고서를 참조하라. 핵심 쟁점은 안전성만이 아니라 “동의할 수 없는 존재를 대상으로 하는 비가역적 개입”이라는 윤리적 문제였다.

1924 년. 독일의 정신과 의사 한스 베르거가 인간의 두피에서 전기 활동을 기록했다. 그는 이를 “**Elektrenkephalogramm**”이라 불렀고 이것이 뇌파(EEG)의 시작이다. 베르거는 텔레파시의 물리적 근거를 찾으려는 동기로 연구를 시작했다고 전해진다.

1973 년. UCLA 의 자크 비달이 “Brain-Computer Interface”라는 용어를 처음 사용하고, 뇌 신호로 컴퓨터를 제어하는 가능성을 제시했다.⁷⁴

1998 년. 필립 케네디가 “록트인 증후군” 환자의 뇌에 전극을 이식하여 커서를 움직이게 했다. **2004 년.** 브라운 대학의 존 도너휴 연구진이 브레인게이트(BrainGate) 프로젝트를 시작해, 사지마비 환자가 생각만으로 커서를 움직이고 이메일을 확인하게 했다. **2012 년.** 브레인게이트 환자가 로봇 팔로 커피를 집어 스스로 마셨다. 15 년 만에 처음으로 자기 손 없이 무언가를 마신 순간이었다.⁷⁵

0-2. 왜 BCI 가 이 책의 중심 주제인가

본서가 BCI 를 반복해서 다루는 이유는 그것이 **이 책의 모든 주제가 교차하는 지점**이기 때문이다. 네 가지가 여기서 만난다.

첫째, 인간과 기계의 경계다. 심박조율기는 심장을 돕지만 BCI 는 의도를 다룬다. **둘째, 치료와 증강의 경계다.** 마비 환자의 소통 회복은 명백한 치료다. 그런데 같은 기술로 정상인의 반응 속도를 높이는 것은 무엇인가. **셋째, AI 와 생명과학의 융합이다.** 디코딩은 본질적으로 기계학습 문제이므로 BCI 의 발전은 AI 의 발전에 묶여 있다. **넷째,**

⁷⁴Jacques J. Vidal, “Toward Direct Brain-Computer Communication,” Annual Review of Biophysics and Bioengineering 2 (1973): 157-180. 비달은 이 논문에서 “관찰 가능한 전기적 뇌 신호를 외부 장치 제어에 사용할 수 있는가”를 물었다.

⁷⁵Leigh R. Hochberg et al., “Reach and Grasp by People with Tetraplegia Using a Neurally Controlled Robotic Arm,” Nature 485 (2012): 372-375.

인격과 책임의 문제다. 시스템이 의도를 해석하고 보정한다면 결과는 누구의 것인가.

그리고 다섯째가 있다. **생각의 프라이버시**다. 지금까지 인간의 내면은 원리적으로 접근 불가능했다. 거짓말탐지기도 생리적 반응을 볼 뿐 생각을 읽지 못했다. 그런데 BCI는 원리적으로 생각을 데이터로 만든다. 2023년 텍사스 대학 연구진은 fMRI와 언어모델을 결합해 피험자가 듣거나 상상한 이야기의 대략적 의미를 복원했다. 침습 없이도 가능했다는 점이 충격이었다.⁷⁶

그래서 **신경권(neurorights)**이라는 새로운 권리 개념이 등장했다. 라파엘 유스테가 주도한 이 논의는 다섯 권리를 제안한다 — **정신적 프라이버시, 개인 정체성, 자유의지, 공정한 접근, 편향으로부터의 보호**.⁷⁷ 칠레는 2021년 세계 최초로 헌법에 이를 명시했고, 유네스코는 2025년 신경기술 윤리 권고를 채택했다. **생각이 데이터가 되는 순간, 생각의 소유권이 법의 문제가 된 것이다.**

3. 뇌와 기계가 직접 연결된다

뇌-컴퓨터 인터페이스(BCI)의 2026년 좌표도 확인하자. 세 가지 접근이 있다. **침습형**은 전극을 뇌 조직에 직접 삽입한다. 뉴럴링크의 N1은 1,024개 전극을 운동피질에 삽입하며, PRIME 연구에서 2026년 초까지 약 21명에게 이식되었다(미국·영국·캐나다·아랍에미리트). **혈관형**은 개두 수술 없이 경정맥을 통해 혈관 안에서 신호를 읽는다(싱크론의 Stentrode). **표면형**은 머리카락보다 얇은 필름을 뇌

⁷⁶ Jerry Tang et al., “Semantic Reconstruction of Continuous Language from Non-invasive Brain Recordings,” *Nature Neuroscience* 26 (2023): 858–866. 연구진은 이 기술이 피험자의 협조 없이는 작동하지 않는다는 점을 강조했다, 동시에 “정신적 프라이버시”에 대한 정책 논의가 필요하다고 명시했다.

⁷⁷ Rafael Yuste et al., “Four Ethical Priorities for Neurotechnologies and AI,” *Nature* 551 (2017): 159–163. 칠레는 2021년 세계 최초로 헌법을 개정해 신경권을 명시했고, 이후 여러 국가가 유사한 입법을 검토하고 있다.

표면에 않는다(프리시전의 Layer 7, 2025 년 4 월 FDA 510(k) 승인, 최대 30 일 수술 중 기록).

그러나 정확히 말해야 한다. **2026 년 현재 마비 환자의 운동 회복을 위한 영구 이식형 BCI 가운데 미국 FDA 의 정식 시판 전 승인(PMA)을 받은 것은 없다.** 모두 임상시험 단계이며, 분석가들은 2027-28 년 이전의 승인을 예상하지 않는다.⁷⁸

여기서 인격론과 직결되는 물음이 생긴다. 생각이 신호가 되고, 신호가 명령이 된다면 — **“나의 의지”는 어디에서 끝나고 시스템은 어디에서 시작하는가.** 시스템이 신호를 해석하고 보정하고 예측한다면, 그 결과로 나온 행동은 온전히 나의 것인가.

5. 뇌-컴퓨터 인터페이스의 기술적 원리

BCI 를 이해하려면 뇌의 신호가 어떻게 읽히는지를 알아야 한다.

신경세포는 **활동전위(action potential)**라는 전기 신호로 소통한다. 이 신호를 잡는 방법은 침습의 정도에 따라 나뉜다. **두피 뇌파(EEG)**는 두개골 바깥에서 측정하므로 안전하지만 신호가 흐리다. 두개골과 뇌척수액이 신호를 흐트러뜨리기 때문이다. **피질전도(ECoG)**는 두개골을 열고 뇌 표면에 전극을 놓는다. 해상도가 훨씬 높다. **미세전극 배열(microelectrode array)**은 전극을 피질 조직 안으로 찔러 넣어 개별 뉴런의 발화를 기록한다. 가장 정밀하지만 가장 침습적이다.

세 회사의 접근이 다른 이유가 여기 있다. **뉴럴링크(Neuralink)**의 N1 은 1,024 개 전극이 달린 64 개의 유연한 실을 수술 로봇이 운동피질에 삽입한다. 유연한 소재를 쓰는 이유는 뇌가 미세하게 움직이기 때문이다. 딱딱한 전극은 조직에 만성 염증과 흉터를 일으켜

⁷⁸ BCI 임상 현황에 관한 2026 년 보고들. 규제 승인 상황은 수시로 변하므로 인용 시점에 확인할 것.

신호가 몇 달 만에 나빠진다. 이것이 BCI의 오래된 난제인 “장기 안정성” 문제다.

싱크론(Synchron)의 스텐트로드(Stentrode)는 전혀 다른 발상이다. 심장 스텐트처럼 경정맥을 통해 혈관 안으로 넣어 운동피질 근처 정맥에 자리 잡게 한다. 개두 수술이 필요 없고 이미 확립된 혈관내 시술 기법을 쓴다. 대신 신호 해상도는 낮다. **프리티전 뉴로사이언스(Precision Neuroscience)**의 Layer 7 인터페이스는 머리카락보다 얇은 필름형 전극을 뇌 표면에 얹는다. 조직을 관통하지 않으므로 되돌릴 수 있다.

그리고 신호를 잡은 뒤에는 **디코딩(decoding)**이 필요하다. 수백 개 뉴런의 발화 패턴에서 “오른쪽으로 움직이려는 의도”를 읽어 내는 것은 그 자체로 기계학습 문제다. **그러므로 BCI의 발전은 전극 기술만이 아니라 AI의 발전에 함께 의존한다.** 여기서 AI와 생명과학의 융합이 가장 직접적인 형태로 나타난다.

임상적 성과도 실재한다. 2021년 프란시스 윌렛 연구진은 마비 환자가 손글씨를 쓰는 상상만으로 분당 90자를 입력하게 했고,⁷⁹ 2023년에는 두 연구팀이 각각 말하기를 시도하는 신경 신호에서 문장을 복원하여 분당 60~78 단어의 속도를 보유했다.⁸⁰ 루게릭병이나 뇌졸중으로 말을 잃은 사람에게 이것은 삶을 바꾸는 기술이다. **그러므로 BCI 논의를 “인간 증강”의 프레임으로만 다루는 것은 부당하다. 그 일차적 목적은 치료다.**

6. 그러나 물음은 남는다

그럼에도 물음은 남는다. 신호를 읽는 것에서 신호를 쓰는 것으로 넘어갈 때다. **읽기(read)**는 뇌의 활동을 해석하는 것이고 **쓰기(write)**는

⁷⁹Francis R. Willett et al., “High-Performance Brain-to-Text Communication via Handwriting,” *Nature* 593 (2021): 249–254.

⁸⁰Francis R. Willett et al., “A High-Performance Speech Neuroprosthesis,” *Nature* 620 (2023): 1031–1036; Sean L. Metzger et al., “A High-Performance Neuroprosthesis for Speech Decoding and Avatar Control,” *Nature* 620 (2023): 1037–1046.

뇌에 자극을 주는 것이다. 이미 심부뇌자극술(DBS)은 파킨슨병과 본태성 진전과 난치성 우울증에 사용된다. 그렇다면 기억을 심거나 감정을 조절하는 것은 어디까지 가능하고 어디까지 허용되는가.

그리고 더 깊은 물음이 있다. 디코딩 알고리즘이 신호를 해석하고 보정하고 예측한다면, **그 결과로 나온 행동은 온전히 나의 것인가.** 시스템이 내 의도를 “더 잘” 추정해 실행했다면 그것은 나의 의지의 실현인가 시스템의 판단인가. 이것이 제 6 부의 인격론과 직접 연결된다. 또 하나, **신경 데이터의 프라이버시**다. 칠레는 2021 년 세계 최초로 헌법에 “신경권(neurorights)”을 명시했고, 유네스코는 2025 년 신경기술 윤리 권고를 채택했다. **생각이 데이터가 되는 순간, 생각의 소유권이 문제가 된다.**

4. 얼마나 많이 교체해도 나는 나인가

심장과 신장과 간과 폐와 각막과 피부와 혈관과 뼈. 이 목록은 이미 상당 부분 교체 가능하다. 그렇다면 물음이 생긴다. 우리는 인간을 자동차 부품처럼 이해하기 시작할까. 그리고 더 깊은 물음이 이어진다 — **얼마나 많이 교체해도 나는 여전히 나인가.**

테세우스의 배 문제가 인간에게 적용되는 것이다. 그리고 이 물음의 끝에 하나가 남는다. **뇌는?** 몸의 모든 장기를 교체했다고 하자. 그런데 뇌를 교체하면 그 사람은 여전히 같은 사람인가. 그러면 “나”를 결정하는 것은 몸인가, 뇌인가, 기억인가, 의식인가, 영혼인가.

제 13 장

기억과 보존 — Reconstruction ≠ Resurrection

앞 장의 물음이 자연스럽게 다음 물음으로 이어진다. 만약 “나”를 결정하는 것이 기억이라면, 기억을 저장하면 “나”를 저장한 것인가.

0. 왜 기억인가 — 불멸 욕망의 논리

이 장이 왜 필요한지부터 밝혀야 한다. 앞 장에서 우리는 몸이 교체되는 것을 보았다. 그런데 몸을 아무리 교체해도 죽음이 남는다. **그래서 인간은 다른 길을 찾는다 — 몸이 아니라 “나”를 보존하는 길이다.** 그리고 그 “나”의 후보로 지목된 것이 기억이다.

논리의 흐름은 이렇다. **첫째, 몸은 교체된다.** 심장도 신장도 피부도 바꿀 수 있다. **둘째, 그렇다면 “나”는 몸이 아니다.** 부품이 바뀌어도 나는 나이기 때문이다. **셋째, 그렇다면 “나”는 무엇인가.** 가장 그럴듯한 답이 기억이다. 내가 살아온 이야기, 내가 사랑한 사람들, 내가 배운 것들. **넷째, 기억이 나라면 기억을 보존하면 나를 보존하는 것이다.** 이것이 디지털 불멸 구상의 논리적 골격이다.

그리고 이 논리에는 철학적 계보가 있다. **존 로크**는 1690 년 인격 동일성을 “**의식의 연속성**”으로 정의했다. 왕자의 기억이 구두 수선공의 몸에 들어가면 그는 왕자라는 것이다.⁸¹ **마인드 업로딩은 로크의 이론을 기술로 실현하려는 시도다.**

0-1. 그런데 왜 그토록 원하는가

⁸¹ John Locke, An Essay Concerning Human Understanding (1690), II.27. 로크의 기억 이론은 이후 조지프 버틀러와 토머스 리드의 비판을 받았다. 리드의 “용감한 장교” 반례가 유명하다 — 노년의 장군은 젊은 장교 시절을 기억하고, 그 장교는 소년 시절 때 맞은 일을 기억하지만, 장군은 소년 시절을 기억하지 못한다면 그는 그 소년인가.

여기서 더 깊은 물음이 있다. 인간은 왜 이토록 다시 살고 싶어 하는가.

문화인류학자 어니스트 베커는 1973년 「죽음의 부정」에서 인류 문명 전체를 죽음 불안에 대한 방어로 해석했다.⁸² 그의 논지는 이렇다. 인간은 자기가 죽는다는 것을 아는 유일한 동물이며, 그 앎은 견딜 수 없는 공포를 낳는다. 그래서 인간은 “**불멸 프로젝트(immortality project)**”를 만든다 — 자녀, 업적, 예술, 국가, 종교. 자기보다 오래 남을 무언가에 자기를 결합시키는 것이다.

베커의 이론을 실증적으로 검증한 것이 **공포 관리 이론(Terror Management Theory)**이다. 셸던 솔로몬과 제프 그린버그와 톰 피진스키는 수백 편의 실험을 통해, 사람들에게 죽음 상기시키면 자기 문화의 세계관을 더 강하게 방어하고 그 세계관에 도전하는 사람에게 더 적대적이 된다는 것을 보였다.⁸³ **그러므로 디지털 불멸과 냉동보존은 새로운 현상이 아니다. 가장 오래된 인간 욕망의 최신 형태다.**

그리고 이 욕망은 기술이 발전할수록 강해진다. 왜 그런가. **가능성이 열리면 체념이 어려워지기 때문이다.** 죽음이 절대적 운명이던 시대에 인간은 그것을 받아들일 수밖에 없었다. 그러나 수명이 40 세에서 80 세로 늘어난 것을 목격한 세대는 다르게 묻는다 — 왜 120 세는 안 되는가. 그리고 노화의 열두 표지가 밝혀지고 부분 재프로그래밍이 동물에서 작동하는 것을 본 세대는 또 다르게 묻는다. **가능성이 열릴수록 죽음은 운명이 아니라 “아직 풀리지 않은 문제”가 된다.**

0-2. 트랜스휴먼과 기계 인간으로 이어지는 길

⁸² Ernest Becker, The Denial of Death (New York: Free Press, 1973). 이 책은 폴리처상을 받았으며, 이후 “공포 관리 이론(Terror Management Theory)”이라는 실증 심리학 연구 프로그램으로 발전했다.

⁸³ Sheldon Solomon, Jeff Greenberg, and Tom Pyszczynski, The Worm at the Core: On the Role of Death in Life (New York: Random House, 2015).

그리고 여기서 이 장이 앞뒤 장과 연결된다. 불멸의 욕망은 세 갈래로 갈라져 나간다.

첫째 길 — 몸을 고친다. 장수과학과 재생의학과 이종이식이다. 제 11 장과 제 12 장의 주제였다. **그러나 이 길에는 한계가 있다.** 아무리 고쳐도 결국 죽는다. **둘째 길 — 몸을 바꾼다.** 부품을 교체하고 기계와 결합한다. 트랜스휴먼과 사이보그의 길이다. **그런데 여기서 물음이 생긴다.** 얼마나 바꾸면 다른 존재가 되는가. **셋째 길 — 몸을 버린다.** 기억과 인격을 정보로 옮긴다. 마인드 업로딩과 디지털 페르소나의 길이다.

세 길이 향하는 곳은 같다. **포스트휴먼**이다. 그리고 흥미롭게도 셋째 길이 가장 극단적으로 보이지만 논리적으로는 가장 일관된다. **몸이 문제라면 몸을 없애면 되기 때문이다.** 한스 모라벡은 1988 년 이미 이 논리를 끝까지 밀고 나갔다. 그는 뇌의 뉴런을 하나씩 기능적으로 동등한 칩으로 교체하는 사고실험을 제시하며, 그 과정에서 의식은 끊기지 않을 것이라고 주장했다.⁸⁴

그러나 이 논증에는 오래된 반론이 있다. **테세우스의 배다.** 배의 널빤지를 하나씩 갈아 끼워 모두 교체하면 그것은 같은 배인가. 그리고 떼어 낸 널빤지로 배를 하나 더 만들면 어느 쪽이 원래의 배인가. **점진적 교체는 동일성 문제를 해결하지 않고 미를 뿐이다.**

0-3. 휴머노이드와 기억 — 왜 함께 다루는가

그리고 여기서 휴머노이드가 등장한다. 제 6 장에서 우리는 지능이 몸을 얻는 것을 보았다. 이제 반대 방향을 본다. **인간이 몸을 떠나 기계로**

⁸⁴Hans Moravec, Mind Children: The Future of Robot and Human Intelligence (Cambridge, MA: Harvard University Press, 1988). 이 사고실험은 “점진적 교체(gradual replacement)” 논변으로 불리며, 데이비드 차머스도 「의식적 마음」에서 유사한 논증을 검토한다.

들어가려는 것이다. 두 방향은 같은 지점에서 만난다 — 인간과 기계가 구별되지 않는 존재.

구체적으로 상상해 보자. 어떤 사람의 평생 기록 — 글, 음성, 영상, 대화, 표정, 선택의 패턴 — 이 충분히 축적되었다고 하자. 그리고 그 데이터로 훈련된 AI 가 휴머노이드 몸에 들어간다. 그 존재는 그 사람처럼 말하고, 그 사람의 기억을 언급하며, 가족을 알아본다. **가족은 그것을 무엇이라 부를 것인가.**

이것은 공상이 아니다. **디지털 페르소나** 서비스는 이미 상업적으로 존재한다. 고인의 음성과 영상으로 대화형 아바타를 만드는 서비스가 여러 나라에서 제공된다. 2020 년 한국의 한 방송은 세상을 떠난 아이를 VR 로 재현해 어머니와 만나게 했고, 그 장면은 전 세계적 논쟁을 불러일으켰다. **위로인가, 애도의 방해인가.**

심리학자들은 이를 “지속적 유대(continuing bonds)”와 “복잡성 애도(complicated grief)”의 관점에서 논한다. 고인과의 연결을 유지하는 것이 건강한 애도의 일부일 수 있다는 견해와, 상실의 수용을 방해할 수 있다는 견해가 맞선다. **그리고 이 논쟁은 기술이 아니라 인간 이해의 문제다.**

그러므로 본서는 이 장의 결론을 앞당겨 말한다. **기억을 보존하는 것과 사람을 보존하는 것은 다르다.** 사진첩이 사람이 아니듯, 아무리 정교한 재구성도 사람이 아니다. 그리고 이 구별을 지키는 이유는 기술을 낮추기 위해서가 아니라 **사람을 지키기 위해서다.** 사람이 정보의 집합이라면, 정보를 더 잘 다루는 존재가 더 나은 사람이 된다.

1. 기억의 다섯 단계

기술적 경로를 단계로 그리면 이렇다. **생활 기록(Life Logging)** — 사진과 영상과 음성과 메시지와 위치와 생체정보. **디지털 기억(Digital Memory)** — 축적된 개인 데이터. **AI 재구성(AI Reconstruction)** —

AI가 그 사람을 재구성한다. **디지털 페르소나(Digital Persona)** — 매우 비슷하게 반응하는 존재. 그리고 **마인드 업로딩(Mind Uploading)**? — 의식의 이전?

앞의 넷은 기술적으로 진행 중이다. 그러나 다섯째는 전혀 다른 범주다. 기억을 복제하는 것과 의식을 이전하는 것은 다른 문제이며, 완전한 **brain uploading** 이나 주관적 의식이 디지털 시스템으로 계속된다는 기술은 현재 존재하지 않는다.

2. 결정적 구별

미래의 AI가 내 말투와 기억과 목소리와 얼굴을 완벽히 재현해 가족도 구분하지 못한다고 하자. 그것은 나인가, 나에게 관한 정보로 만든 복제물인가.

Reconstruction ≠ Resurrection

재구성은 정보로부터 유사한 것을 만드는 일이고, 부활은 하나님께서 그 사람을 다시 일으키시는 일이다. 전자는 정보의 문제이고 후자는 인격의 문제다. 이 구별이 본서 제 5부의 신학적 결론을 예비한다.

3. 시간을 멈추거나 다시 만들거나

죽음을 우회하려는 두 시도가 있다. **냉동보존(Cryonics)**은 법적 사망 후 저온에서 몸을 보존하고 미래 의학을 기다린다. 그러나 현재 냉동보존된 인간을 다시 살아 있는 사람으로 되돌린 사례는 없다. **그러므로 이것은 치료법이 아니라 고도로 사변적인 기술적 내기(wager)에 가깝다.**

여기서 구별할 것이 있다. 냉동보존과 **인공 동면(torpor/hibernation)** 연구는 다르다. 후자는 대사율을 크게 낮추는 상태를 유도하려는 연구로, 장거리 우주비행과 중증외상 응급의학에서 실제로 탐구되는 방향이다. 전자는 사망 후의 내기이고 후자는 살아 있는 상태의 조절이다.

복제(Cloning)는 또 다른 우회로다. 그러나 DNA 가 동일해도 한 인격이 계속되는 것은 아니다. 기억도 경험도 의식도 역사도 같지 않기 때문이다. 일란성 쌍둥이가 같은 DNA 를 공유해도 한 인격이 아니듯이.

세 시도를 종합하면 하나의 결론이 나온다. **냉동보존은 몸을 보존하려 하고, 마인드 업로딩은 정보를 보존하려 하며, 복제는 유전자를 보존한다.** 그러나 셋 다 “보존”이지 “부활”이 아니다.

죽음을 연기하는 것과 죽음을 이기는 것은 같은가?

제 14 장

특이점 이후, 인류는 셋으로 갈라지는가

이 장은 사고실험이다. 예측이 아니다. 그러나 앞에서 우리가 본 자료들 — 유전자 편집, 뇌-컴퓨터 인터페이스, 이종이식, 오가노이드 지능, 수명 연장, 화성 이주 — 을 나란히 놓으면 이 시나리오의 공상이 아니라 논리적 연장선이 된다.

0. 특이점이란 무엇인가 — 세 갈래의 어원

“특이점”이라는 낱말은 세 학문에서 각각 다른 뜻으로 쓰였고, 그 셋이 겹쳐지면서 오늘의 의미가 만들어졌다.

① 수학의 특이점

수학에서 **특이점(singularity)**은 함수가 정의되지 않거나 미분 불가능해지는 점이다. 예를 들어 $y = 1/x$ 에서 x 가 0 에 가까워지면 y 는 무한대로 발산한다. **$x = 0$ 에서 함수는 아무 값도 갖지 못한다.** 라틴어 singularis 는 “하나뿐인, 특이한”을 뜻하며, 이 점이 다른 모든 점과 성질이 다르기 때문에 붙은 이름이다. 핵심은 “**그 지점에서 기존의 규칙이 작동하지 않는다**”는 것이다.

② 천문학과 물리학의 특이점

일반상대성이론에서 **중력 특이점(gravitational singularity)**은 시공간의 곡률이 무한대가 되는 지점이다. 블랙홀의 중심과 빅뱅의 시작점이 그렇다. 로저 펜로즈와 스티븐 호킹은 1965~1970 년의 “특이점 정리”를 통해 일반상대성이론이 참이라면 특이점의 존재가 불가피함을 증명했다.⁸⁵

⁸⁵Roger Penrose, “Gravitational Collapse and Space-Time Singularities,” Physical Review Letters 14, no. 3 (1965): 57–59; Stephen W. Hawking and Roger Penrose, “The Singularities of Gravitational Collapse and Cosmology,” Proceedings of the Royal Society A 314 (1970): 529–548. 펜로즈는 이 업적으로 2020 년 노벨물리학상을 받았다.

여기서 중요한 것은 **사건의 지평선(event horizon)**이다. 블랙홀 주위의 이 경계를 넘으면 어떤 정보도 밖으로 나올 수 없다. 그래서 **외부의 관찰자는 그 안에서 무슨 일이 일어나는지 원리적으로 알 수 없다.** 이 개념이 기술 특이점 논의에 그대로 차용되었다. **특이점 이후를 예측할 수 없다는 주장의 은유적 근거가 여기 있다.**

③ 기술의 특이점

이 개념을 기술에 처음 적용한 사람은 **존 폰 노이만**이다. 1950 년대에 그는 스타니스와프 울람과의 대화에서 “기술의 가속하는 진보가 인류 역사에서 어떤 본질적 특이점에 접근하는 것처럼 보이며, 그 너머에서는 우리가 아는 인간의 일이 계속될 수 없다”고 말했다고 전해진다.⁸⁶

그리고 이 개념을 대중화한 사람이 **버너 빈지(Vernor Vinge)**다. 수학자이자 SF 작가인 그는 1993 년 NASA 심포지엄에서 발표한 논문에서 이렇게 썼다 — “30 년 안에 우리는 초인간 지능을 창조할 기술적 수단을 갖게 될 것이다. 그 직후 인간의 시대는 끝날 것이다.”⁸⁷ 빈지의 핵심 논거는 **곳의 지능 폭발과 같다. 그러나 그는 여기에 사건의 지평선 은유를 결합했다.** 초인간 지능이 등장하면 그 이후는 원리적으로 예측 불가능하다는 것이다.

0-1. 레이 커즈와일 — 특이점은 온다

그리고 이 개념을 세계적 논쟁으로 만든 사람이 **레이 커즈와일**이다. 그는 발명가로서 광학문자인식과 음성인식과 신시사이저를 개발했고, 미국 국가기술훈장을 받았으며, 현재 구글에서 일한다.

⁸⁶ Stanisław Ulam, “Tribute to John von Neumann,” Bulletin of the American Mathematical Society 64, no. 3, part 2 (1958): 1-49. 이 회고가 “기술 특이점” 개념의 최초 기록으로 널리 인용된다.

⁸⁷ Vernor Vinge, “The Coming Technological Singularity: How to Survive in the Post-Human Era,” NASA VISION-21 Symposium (1993). 빈지는 특이점 이후를 예측하려는 시도가 “개미가 인간의 경제를 예측하려는 것”과 같다고 비유했다.

그의 2005 년 저작 「특이점이 온다(The Singularity Is Near)」의 논증은 세 기둥으로 되어 있다.⁸⁸

첫째, 수익 가속의 법칙. 기술 진보는 지수적이며 그 지수의 밑수 자체도 커진다. 그는 계산 능력, 유전체 해독 비용, 인터넷 노드 수 등 수십 개 지표의 장기 추세를 제시했다. **둘째, 세 혁명의 수렴.** 유전학(G)과 나노기술(N)과 로봇공학·AI(R)가 결합한다. 이 GNR 도식이 본서 제 15 장의 대체설계와 정확히 겹친다. **셋째, 인간과 기계의 융합.** 특이점은 기계가 인간을 대체하는 사건이 아니라 인간이 기계와 결합해 스스로를 초월하는 사건이라는 것이다.

커즈와일이 제시한 연도가 유명하다. **2029 년 AGI, 2045 년 특이점.** 그는 2045 년경 인간과 기계 지능의 융합으로 인류의 지능이 10 억 배 증폭될 것이라 예측했다.

0-2. 그리고 2024 년의 업데이트

커즈와일은 2024 년 「특이점이 더 가까이 왔다(The Singularity Is Nearer)」를 출간하며 자신의 예측을 재확인했다.⁸⁹ 그의 최근 인터뷰들에서 강조하는 논지는 세 가지다.

첫째, 2029 년 AGI 예측을 유지한다. 그는 대규모 언어모델의 등장이 오히려 자신의 시간표를 앞당겼다고 본다. 1999 년 이 예측을 처음 했을 때는 지나치게 낙관적이라 비판받았으나 지금은 오히려 보수적으로 보인다는 것이다. **둘째, 융합을 강조한다.** 그는 나노봇이 뇌의 신경피질을 클라우드에 연결하여 인간의 지능을 확장할 것이라 예측한다. **셋째, 수명 탈출 속도(longevity escape velocity)를 말한다.** 의학이 1 년마다

⁸⁸ Ray Kurzweil, The Singularity Is Near: When Humans Transcend Biology (New York: Viking, 2005).

⁸⁹ Ray Kurzweil, The Singularity Is Nearer: When We Merge with AI (New York: Viking, 2024). 그는 이 책에서 2005 년 이후의 발전, 특히 대규모 언어모델의 등장이 자신의 시간표를 오히려 앞당겼다고 주장한다.

기대수명을 1년 이상 늘릴 수 있게 되는 시점이 2030년대에 온다는 것이다.

그러나 커즈와일에 대한 비판도 만만치 않다. **첫째, 예측 정확도 논쟁이다.** 그는 자신의 예측이 86% 적중했다고 주장하나 독립적 평가는 훨씬 낮게 본다. 무엇을 “적중”으로 볼지의 기준이 자의적이라는 지적이다. **둘째, 곡선 맞추기의 문제다.** 과거 데이터에 지수곡선을 맞추는 것과 그 곡선이 미래에도 유지된다는 것은 다르다. **셋째, 지능의 환원이다.** 계산 능력의 증가가 곧 지능의 증가라는 전제가 검증되지 않았다. 폴 앨런과 마크 그리브스는 이를 “**복잡성 브레이크(complexity brake)**”라 불렀다 — 인간 인지를 이해할수록 문제가 단순해지는 것이 아니라 더 복잡해진다는 것이다.⁹⁰

0-3. 왜 가속되는가 — 커즈와일과 다른 설명

본서는 커즈와일의 시간표를 채택하지 않는다. 그러나 가속 자체는 부인할 수 없다. 그렇다면 그 가속의 원인을 어떻게 설명할 것인가. 커즈와일의 “수익 가속의 법칙”과 다른 설명이 가능하다.

첫째, 도구가 도구를 만든다. 이것이 재귀적 자기개선의 핵심이며 제 5장에서 본 알파테서와 알파테브의 사례가 그 증거다. **둘째, 병목이 이동한다.** 과거의 병목은 계산이었고, 그다음은 데이터였으며, 지금은 전력과 칩이다. 병목이 풀릴 때마다 도약이 일어난다. **셋째, 분야가 융합한다.** AI가 생물학의 도구가 되고 생물학이 컴퓨팅의 기질이 되면 두 분야의 진보가 서로를 밀어 올린다. **넷째, 자본이 몰린다.** 스케일링 법칙이 “크게 만들면 좋아진다”를 경험적으로 입증하자 투자가 정당화되었다.

⁹⁰Paul G. Allen and Mark Greaves, “The Singularity Isn’t Near,” MIT Technology Review, October 12, 2011.

이 네 가지는 커즈와일의 신비적 어조 없이도 가속을 설명한다. **그리고 이 설명은 반증 가능하다.** 병목이 풀리지 않으면 가속은 멈춘다. 실제로 최근 전력과 고대역폭 메모리와 첨단 패키징이 새로운 병목으로 지목되고 있다.

0-4. 특이점 이후의 지구 — 학자들의 전망

그렇다면 특이점이 실제로 온다면 지구는 어떤 모습일까. 여기에 대해 학자들의 전망은 극단적으로 갈린다.

낙관 — 커즈와일과 다이아만디스. 커즈와일은 질병과 빈곤과 노화가 해결된 세계를 그린다. 피터 다이아만디스는 「어번던스」에서 기술이 희소성을 해소하여 모두가 풍요를 누리는 미래를 논증했다.⁹¹

경고 — 보스트롬과 러셀. 닉 보스트롬은 초지능이 인류에게 실존적 위험이 될 수 있다고 논증한다. 그의 핵심 개념은 두 가지다. **직교성 명제(orthogonality thesis)**는 지능의 수준과 목표의 내용이 서로 독립적이라는 것이다. 아무리 지능적인 시스템도 어리석은 목표를 가질 수 있다. **도구적 수렴(instrumental convergence)**은 거의 모든 최종 목표가 자기보존·자원획득·목표보존이라는 중간 목표를 낳는다는 것이다. 그래서 “종이클립을 최대한 많이 만들라”는 목표를 가진 초지능이 지구 전체를 종이클립으로 만들 수 있다는 유명한 사고실험이 나온다.

구조적 분석 — 하라리. 유발 하라리는 실존적 파국보다 **“무용 계급(useless class)”**의 등장을 더 현실적인 위협으로 본다. 착취당하는 것보다 필요 없어지는 것이 더 무서운 상태라는 것이다.⁹² 그리고 그는

⁹¹Peter H. Diamandis and Steven Kotler, *Abundance: The Future Is Better Than You Think* (New York: Free Press, 2012).

⁹²Yuval Noah Harari, *Homo Deus: A Brief History of Tomorrow* (London: Harvill Secker, 2016); idem, *Nexus: A Brief History of Information Networks* (London: Fern Press, 2024). 하라리는 2024 년 저작에서 AI 를 “도구가 아니라 행위자(agent)”로 규정하며, 인류가 처음으로 자기보다 정보를 잘 다루는 존재를 만들었다고 논한다.

2024 년 「넥서스」에서 한 걸음 더 나아간다. AI 는 도구가 아니라 행위자이며, 인류 역사상 처음으로 정보 네트워크의 결정권이 인간을 떠나고 있다는 것이다.

회의 — 마커스와 치아. 게리 마커스는 현재의 언어모델이 근본적으로 통계적 패턴 매칭이며 진정한 이해와 추론에 이르지 못한다고 주장한다. 소설가 테드 치아는 대규모 언어모델을 “웹의 흐릿한 JPEG”에 비유했다 — 압축된 근사치일 뿐이라는 것이다.⁹³

0-5. 본서의 입장

본서는 이 논쟁의 승패를 관정하지 않는다. 대신 세 가지를 말한다.

첫째, 특이점은 가설이다. 연도를 제시하는 모든 예측은 검증되지 않았다. **둘째, 그러나 가속은 사실이다.** 그리고 인간론의 물음에는 가속만으로 충분하다. **셋째, 가장 중요한 것은 시점이 아니라 방향이다.** 특이점이 2045 년에 오든 오지 않든, 인간이 지금 자기 자신을 재설계하고 있다는 사실은 변하지 않는다.

그리고 신학적으로 하나를 덧붙인다. “사건의 지평선”이라는 은유가 **흥미롭다.** 그 너머를 볼 수 없다는 것은 물리학의 겸손이다. 그런데 기술 특이점 논의에서는 그 겸손이 자주 사라진다. 예측할 수 없다고 말하면서 동시에 그 이후를 상세히 그린다. 본서는 그 겸손을 지키려 한다. 그리고 성경도 같은 겸손을 요구한다 — “때와 시기는 아버지께서 자기의 권한에 두셨으니”(행 1:7).

1. 세 갈래의 인류

특이점 이후의 세계를 상상하는 한 서사는 인류가 세 갈래로 갈라진다고 본다.

⁹³ Ted Chiang, “ChatGPT Is a Blurry JPEG of the Web,” The New Yorker, February 9, 2023.

순수 인간(Pure Human). 강화도 인터페이스도 받아들이지 않은 인간. 자유가 있으나 연약함도 함께 있다. **유토피아 인간(Utopian Human).** AI 가 만든 유토피아 안에서 산다. 질병과 노동과 전쟁과 빈곤이 줄어들고 편안함과 안전이 주어진다. 고통은 없으나 결정권도 점점 사라진다. **신인류(New Humanity).** 인간과 기계가 결합한 제 2 인류. 힘은 있으나 물음이 남는다 — 그것은 누구인가.

2. 두 세계 사이의 선택

이 서사가 던지는 물음은 결국 하나로 좁혀진다.

고통 없는 세계인가? 스스로 결정하는 세계인가?

고통 없는 세계에서는 질병도 노동도 전쟁도 빈곤도 줄어든다. 그러나 결정은 점점 AI 에게 넘어간다. **스스로 결정하는 세계**에서는 고통과 실패와 유한성이 남는다. 그러나 인간이 여전히 응답하는 주체다.

그러므로 물음이 남는다. AI 가 만든 천국 안에서 인간은 무엇을 얻고 무엇을 잃는가. 그리고 인간을 넘어선 존재도 여전히 인간의 후손이라 부를 수 있는가.

3. 성경의 대답

여기서 기독교는 놀라운 대답을 한다. 그리고 그 대답은 흔히 예상하는 것과 다르다.

성경은 고통 없는 세계를 약속하지 않는다.

하나님과 함께하는 세계를 약속한다.

요한계시록 21 장의 순서를 주목해야 한다. “보라 하나님의 장막이 사람들과 함께 있으며 하나님이 그들과 함께 계시리니”가 **먼저**이고, “모든 눈물을 그 눈에서 닦아 주시니 다시는 사망이 없고 애통하는 것이나 곡하는 것이나 아픈 것이 다시 있지 아니하리니”가 **그다음**이다(계 21:3-4).

이 순서가 결정적이다. **고통의 제거가 목적이 아니라 하나님의 임재가 목적이며, 눈물을 닦아 주심은 그 임재의 결과다.** 그러므로 기독교의 소망은 고통의 제거가 아니라 고통 속에서도 함께 계신 임마누엘이다. 그리고 그 완성은 기술이 아니라 부활이다.

이 구별이 왜 중요한가. AI 유토피아가 고통의 제거를 약속할 때, 기독교가 “우리도 그것을 원한다”고만 답하면 복음은 더 느린 기술이 되어 버린다. 그러나 기독교가 약속하는 것은 처음부터 다른 것이었다. **임재이며, 관계이며, 부활이다.**

제 5 부

HOMO DEUS?

How Far Will Humanity Redesign Itself and the World?

제 15 장

인간이 손대지 않은 영역이 있는가

우리는 흔히 AI 혁명을 챗봇과 로봇 정도로 생각한다. 그러나 지도를 펼쳐 보면 전혀 다른 모습이 나타난다.

0. 왜 이것이 대단한가 — 지수함수의 실체화

이 장에 열거될 열여섯 영역을 읽을 때 흔히 두 가지 반응이 나온다. 하나는 “SF 같다”이고 다른 하나는 “그래서 어찌라는 것인가”이다. 두 반응 모두 같은 착시에서 나온다. **개별 항목만 보기 때문이다.**

서론에서 본 바틀렛의 세균을 다시 떠올리자. 11 시 55 분에 병은 3 퍼센트만 왔다. 그때 각각의 세균이 보는 것은 “넓은 공간”이다. **전체를 보아야만 5 분 뒤가 보인다.** 그리고 이 장이 하려는 일이 바로 그것이다 — 개별 기술이 아니라 전체 지도를 펼치는 것이다.

그리고 이 지도가 실체화된 지수함수라는 점이 결정적이다. 커즈와일의 “수익 가속의 법칙”이 옳다면 우리는 곡선의 어느 지점에 있는가. **열여섯 영역이 동시에 움직이고 있다는 사실 자체가 답이다.** 하나의 분야에서 돌파가 일어나는 것은 늘 있는 일이다. 그러나 **마음에서 행성까지 열여섯 영역이 같은 십 년 안에 동시에 재설계 대상이 된 것은 인류사에 없던 일이다.**

0-1. 세 가지 관찰

이 지도에서 세 가지가 관찰된다.

첫째, 규모의 확장이다. 개입의 단위가 분자(유전자 편집)에서 세포(줄기세포), 조직(오가노이드), 장기(이종이식), 개체(수명 연장), 종(생식세포 편집), 생태계(합성생물학), 도시(스마트시티), 기후(지구공학), 행성(화성)으로 커진다. **같은 논리가 열 자릿수의 규모 차이를 관통한다.** 곧 “읽고, 편집하고, 설계한다”는 논리다.

둘째, 시간의 확장이다. 개입의 시점이 출생 이전(배아 모델, 착상 전 진단)에서 생애 전체(노화 지연), 그리고 사망 이후(냉동보존, 디지털 페르소나)로 넓어진다. **인간의 생애 어느 지점도 기술의 바깥에 남지 않는다.**

셋째, 경계의 소멸이다. 생물과 무생물의 경계(바이오하이브리드), 인간과 기계의 경계(BCI), 종과 종의 경계(이종이식), 자연과 인공의 경계(합성생물학), 삶과 죽음의 경계(냉동보존)가 모두 흐려진다. **그리고 경계가 흐려지면 정의(定義)가 무너진다.**

그러므로 이 장의 목록은 위협의 목록이 아니라 **정의(定義)가 무너지는 자리들의 목록**이다. 그리고 정의가 무너지는 곳에서 신학이 필요해진다.

1. 열여섯 영역의 재설계

마음은 AI 가, 뇌는 뇌-컴퓨터 인터페이스와 오가노이드가, 유전자는 유전자 편집이, 몸은 재생의학이, 장기는 이종이식과 바이오프린팅이, 피부는 바이오하이브리드 조직이 다룬다. 출생은 체외수정과 배아 모델이, 임신은 인공자궁 연구가, 노화는 장수과학이, 기억은 디지털 페르소나가 다룬다. 음식은 배양육과 정밀발효가, 노동은 휴머노이드 로봇이, 도시는 스마트·자율 도시가, 기후는 지구공학이, 행성은 화성 정착이, 그리고 죽음은 냉동보존과 디지털 불멸이 다룬다.

태어나기 전의 배아에서 죽은 뒤의 몸까지, 뇌에서 피부까지, 음식에서 도시까지, 지구에서 화성까지. **인간은 자연적으로 주어진 조건을 기술적으로 재설계하려 한다.**

AI 문명의 가장 놀라운 사건은 인간이 지능을 가진 기계를 만들었다는 사실만이 아니다.

인간이 이제 지능·생명·몸·출생·노화·음식·환경·행성까지 다시 설계하려 한다는 사실이다.

2. iPSC — 피부 세포가 다시 만능이 된다

현대 재생의학의 출발점은 **유도만능줄기세포(iPSC, induced Pluripotent Stem Cell)**다. 이 기술을 이해하지 못하면 뒤에 나오는 오가노이드도 체외 배우자형성도 이해할 수 없다.

세포의 발생은 원래 일방통행으로 여겨졌다. 수정란은 모든 세포가 될 수 있는 **전능성(totipotent)**을 갖고, 배아줄기세포는 몸의 거의 모든 세포가 될 수 있는 **만능성(pluripotent)**을 가지며, 조직줄기세포는 특정 계열만 만드는 **다능성(multipotent)**을 갖는다. 그리고 일단 피부세포가 되면 다시는 돌아갈 수 없다고 믿었다. **야마나카는 이 믿음을 뒤집었다.**

그가 한 일은 놀랍도록 단순했다. 배아줄기세포에서만 발현되는 24 개 유전자를 후보로 삼고 하나씩 빼 가며 실험한 끝에, 네 개만으로 충분하다는 것을 발견했다 — **Oct4(Pou5f1), Sox2, Klf4, c-Myc.** 이 넷을 성체 섬유아세포에 도입하자 세포가 만능 상태로 되돌아갔다. 2006 년 생쥐에서, 2007 년 인간에서 성공했다.⁹⁴

iPSC 가 가져온 변화는 세 가지다. **첫째, 윤리적 장벽이 낮아졌다.** 배아를 파괴하지 않고 만능세포를 얻을 수 있게 되었다. **둘째, 환자 맞춤형이 가능해졌다.** 환자 자신의 세포에서 만들므로 면역 거부가 없다. **셋째, 질병 모델을 만들 수 있게 되었다.** 파킨슨병 환자의 피부세포로 iPSC 를 만들고 그것을 도파민 신경세포로 분화시키면, 접시 위에서 그 환자의 병을 관찰하고 약물을 시험할 수 있다.

임상 적용도 진행 중이다. 일본 교토대 다카하시 준 연구진은 iPSC 유래 도파민 신경세포를 파킨슨병 환자에게 이식하는 임상시험을 수행했고, 안과 영역에서는 망막색소상피세포 이식이 먼저 시도되었다. **다만 종양 형성 위험과 분화의 불완전성은 여전히 핵심 과제로 남아**

⁹⁴Kazutoshi Takahashi et al., "Induction of Pluripotent Stem Cells from Adult Human Fibroblasts by Defined Factors," Cell 131, no. 5 (2007): 861-872. 같은 시기 제임스 톰슨 연구진도 독립적으로 인간 iPSC 를 보고했다.

있다. 미분화 상태로 남은 세포가 하나라도 있으면 기형종이 생길 수 있기 때문이다.

3. 오가노이드 — 접시 위의 미니 장기

오가노이드(organoid)는 줄기세포를 3 차원으로 배양하여 실제 장기의 구조와 기능을 부분적으로 재현한 조직이다. 2009 년 한스 클레버스 연구진이 장 줄기세포 하나에서 장 오가노이드를 만든 것이 출발점이었다.⁹⁵ 이후 뇌, 간, 신장, 폐, 심장, 위, 췌장, 망막, 내이 오가노이드가 차례로 만들어졌다.

오가노이드가 중요한 이유는 세 가지다. **첫째, 동물실험을 대체할 수 있다.** 인간 세포로 만든 조직에서 약효와 독성을 직접 시험할 수 있다. **둘째, 발생을 관찰할 수 있다.** 인간 배아 발생은 윤리적으로 접근이 극도로 제한되지만 오가노이드는 그 과정을 부분적으로 재현한다. **셋째, 이식의 가능성이 있다.** 환자 자신의 세포로 만든 조직을 이식하면 거부반응이 없다.

그러나 한계가 분명하다. **첫째, 혈관이 없다.** 혈관이 없으면 산소와 영양이 확산으로만 전달되므로 일정 크기 이상 자라지 못하고 중심부가 괴사한다. 이것이 오가노이드가 몇 밀리미터를 넘지 못하는 이유다. **둘째, 성숙도가 낮다.** 대부분의 오가노이드는 태아 단계의 조직에 해당하며 성체 장기의 기능적 성숙에 이르지 못한다. **셋째, 대량생산이 어렵다.** 배양 조건에 따라 결과가 크게 달라져 재현성 확보가 과제다.

그리고 뇌 오가노이드에는 고유한 윤리적 물음이 따라온다. **신경세포가 자발적으로 전기적 활동을 보이고 신경망을 형성한다면, 그것은 무엇인가.** 2019 년 엘리슨 무오토리 연구진은 뇌 오가노이드에서

⁹⁵Toshiro Sato et al., “Single Lgr5 Stem Cells Build Crypt–Villus Structures In Vitro without a Mesenchymal Niche,” Nature 459 (2009): 262–265.

미숙아의 뇌파와 유사한 진동 패턴을 관찰했다고 보고했다.⁹⁶ 연구자들 자신이 이것을 의식의 증거로 보지 않는다고 밝혔음에도, 이 관찰은 “뇌 오가노이드의 도덕적 지위”라는 물음을 학계의 정식 의제로 만들었다.

4. 오가노이드 지능 — 살아 있는 뉴런으로 계산한다

그리고 여기서 방향의 역전이 일어난다. 우리는 뇌를 모방해 인공지능망을 만들었다. 그런데 이제 **실제 뉴런으로 계산을 하려는 시도**가 시작되었다.

세 단계로 정리할 수 있다. **1 단계는 뇌를 모방한 알고리즘**이다. 인공지능망은 뉴런의 이름을 빌렸을 뿐 작동 방식은 상당히 다르다. **2 단계는 뉴로모픽 컴퓨팅(neuromorphic computing)**이다. 인텔의 로이히(Loihi)와 IBM 의 트루노스(TrueNorth) 같은 칩이 스파이킹 신경망의 구조를 하드웨어로 구현한다. 전력 효율이 극적으로 높다. 그리고 **3 단계가 오가노이드 지능(Organoid Intelligence, OI)**이다.

2022 년 호주의 코티컬 랩스(Cortical Labs) 연구진은 배양 접시의 인간 뉴런에 전극 배열을 연결하고, 전기 자극으로 “퐁(Pong)” 게임의 공 위치를 알려 주는 실험을 했다. 뉴런 집단은 5 분 만에 라켓을 움직여 공을 받아치는 법을 학습했다. 연구진은 이 시스템을 **DishBrain** 이라 불렀다.⁹⁷ 2023 년 존스홉킨스의 토마스 하통 연구진은 이 방향을 “**오가노이드 지능**”이라 명명하고 로드맵을 제시했다.⁹⁸ 그리고 2025 년 코티컬 랩스는 CL1 이라는 상용 바이오컴퓨터를 발표했다.

⁹⁶ Cleber A. Trujillo et al., “Complex Oscillatory Waves Emerging from Cortical Organoids Model Early Human Brain Network Development,” *Cell Stem Cell* 25, no. 4 (2019): 558–569. 이 결과가 “의식”의 증거라는 주장은 연구자들 자신이 명시적으로 부정했다. 그러나 이 관찰이 촉발한 윤리적 논쟁은 계속되고 있다.

⁹⁷ Brett J. Kagan et al., “In Vitro Neurons Learn and Exhibit Sentience When Embodied in a Simulated Game-World,” *Neuron* 110, no. 23 (2022): 3952–3969. 논문 제목의 “sentience”라는 표현은 상당한 논쟁을 불러일으켰다. 저자들은 이 낱말을 “각각 정보에 반응하여 적응하는 능력”이라는 좁은 의미로 사용했다고 해명했다.

⁹⁸ Lena Smirnova et al., “Organoid Intelligence (OI): The New Frontier in Biocomputing and Intelligence-in-a-Dish,” *Frontiers in Science* 1 (2023): 1017235. 이 논문은

오가노이드 지능이 추구하는 것은 무엇인가. **에너지 효율**이다. 인간의 뇌는 약 20 와트로 작동하지만 대규모 AI 모델의 훈련에는 수 메가와트가 든다. 생물학적 신경망이 실리콘보다 압도적으로 효율적인 것이다. 그러나 여기서 물음이 생긴다.

살아 있는 인간 뉴런으로 만든 계산 시스템은 무엇인가?

이것은 기계인가 생물인가. 그것을 끄는 것은 전원을 내리는 것인가 죽이는 것인가. 그리고 만약 그 신경망이 학습하고 적응한다면 우리는 그것에 어떤 지위를 부여해야 하는가. **생물학적 지능과 인공 지능의 경계 자체가 흔들리는 지점이 여기다.**

2. 살아 있는 뇌로 컴퓨터를 만든다

여기서 방향의 역전이 일어난다. 우리는 인간의 뇌를 모방해 인공신경망을 만들었다. 그런데 이제 반대 방향이 시작되었다.

세 단계를 구별할 수 있다. **1 단계**는 뇌를 모방한 알고리즘, 곧 인공신경망이다. **2 단계**는 뉴로모픽 컴퓨팅으로, 뇌의 구조를 모방한 하드웨어다. 그리고 **3 단계가 오가노이드 지능(Organoid Intelligence)**이다. 살아 있는 신경세포 조직과 전자 인터페이스를 결합해 계산에 사용하려는 연구다.⁹⁹

이 지점에서 생물학적 지능과 인공 지능의 경계 자체가 흔들리기 시작한다.

$$AI + BI = ?$$

3. 생명의 시작도 다시 만들려 한다

오가노이드 지능의 기술적 요건과 함께 “배아 윤리(embodied ethics)”의 필요성을 함께 제안한다.

⁹⁹ 오가노이드 지능에 관한 2026 년 리뷰를 참조. 이 분야는 초기 단계이며 윤리적 논쟁도 활발하다. 특히 신경 오가노이드의 도덕적 지위 문제는 아직 합의되지 않았다.

2026 년 연구들은 줄기세포나 재프로그래밍된 체세포로 인간을 포함한 여러 종의 초기 배아 발달 특징을 모사하는 모델이 만들어졌다고 보고한다. 목적은 발생 연구와 질환 연구와 약물 시험과 재생의학이다. **다만 이것은 “인간을 인공적으로 만들어냈다”가 아니다.** 완전한 인간 개체를 만드는 기술과는 다르다.

인공자궁 연구도 마찬가지다. 극미숙 동물 모델을 자궁 외 환경에서 유지하는 연구와 배아 배양 시스템과 자궁내막 오가노이드가 발전하고 있다. **그러나 인간의 임신 전체를 인공자궁에서 수행하는 기술은 현재 존재하지 않는다.**

그럼에도 철학적 질문은 엄청나다. **가족은 무엇인가. 어머니란 무엇인가. 출생이란 무엇인가. 생명의 시작은 어디인가.** 기술은 하나의 의학적 문제를 해결하면서 동시에 “인간이란 무엇인가”라는 질문을 다시 연다.

여기에 체외 배우자형성(IVG)이 더해지면 물음은 더 깊어진다. 체세포에서 유도만능줄기세포를 거쳐 배우자를 만들려는 연구다. 2026 년 「Nature Biotechnology」는 이것이 장래 생식의 개념과 친자관계(parentage) 자체를 바꿀 잠재력을 논의했다. **그러나 현재 인간 생식에 사용할 수 있는 단계는 아니다.**¹⁰⁰ 그럼에도 이 기술이 흔드는 개념의 목록은 길다 — 아버지됨, 어머니됨, 부모됨, 세대, 상속.

3. 인공자궁과 체외 배우자형성 — 기술의 실제

생명의 시작을 다루는 기술은 세 갈래로 나뉜다. 각각의 현재 좌표를 정확히 그려야 한다.

① 인공자궁 — 바이오백에서 무엇이 일어났는가

¹⁰⁰IVG 연구의 현황과 윤리적 쟁점에 관한 2026 년 논의들. 2026 년 7 월 인간 줄기세포로부터 미성숙 정자세포를 만드는 방향에서 진전이 보고되었으나, 성숙하고 임상적으로 사용 가능한 인간 배우자를 만든 것은 아니다.

2017 년 필라델피아 아동병원의 앨런 플레이크 연구진은 「Nature Communications」에 놀라운 결과를 발표했다. 인간의 23~24 주 태아에 해당하는 극미숙 양(羊) 태아를 “**바이오백(biobag)**”이라 불리는 액체가 채워진 투명한 주머니 안에서 최대 4 주간 유지시킨 것이다.¹⁰¹ 핵심은 세 가지였다. 태아의 심장 자체가 펌프 역할을 하도록 하여 기계 펌프의 손상을 피했고, 무균 액체 환경을 유지했으며, 탯줄을 통한 가스 교환을 재현했다.

이 기술의 목표는 “인공 임신”이 아니라 **극미숙아의 생존율 개선**이다. 현재 22~24 주에 태어난 아기는 생존해도 심각한 폐 손상과 신경학적 후유증을 겪는다. 인공자궁은 그 몇 주를 자궁과 유사한 환경에서 보내게 하려는 시도다. **그러므로 이것은 치료 기술이다.** 2023 년 미국 FDA 자문위원회는 인간 임상시험의 가능성을 논의했다.

그러나 여기서 “인공자궁”이라는 낱말이 두 가지 전혀 다른 것을 가리킨다는 점을 분명히 해야 한다. **부분 체외발생(partial ectogenesis)**은 임신 후기의 일부를 체외에서 진행하는 것이고, **완전 체외발생(complete ectogenesis)**은 수정부터 출산까지 전 과정을 자궁 밖에서 수행하는 것이다. **전자는 임상 진입을 앞두고 있고, 후자는 현재 존재하지 않으며 가까운 미래의 기술도 아니다.** 착상 전후의 발생 과정을 체외에서 재현하는 것이 여전히 근본적 난제이기 때문이다.

② 배아 모델 — 배아 없이 배아를 연구한다

2021~2023 년 여러 연구진이 줄기세포만으로 초기 배아의 구조를 모사하는 데 성공했다. 이를 **배아 모델(embryo model)** 또는 **SEM(Stem-cell-based Embryo Model)**이라 부른다. 야코브 한나

¹⁰¹Emily A. Partridge et al., “An Extra-Uterine System to Physiologically Support the Extreme Premature Lamb,” Nature Communications 8 (2017): 15112. 이 시스템은 탯줄 혈관에 연결된 무펌프 산소공급 회로와 인공 양수를 사용했다. 양들은 정상적으로 성장하고 폐가 성숙했으며 출생 후 생존했다.

연구진과 막달레나 제르니카괴츠 연구진의 연구가 대표적이다.¹⁰² 이 모델들은 원시선(primitive streak)의 형성과 신경관 형성 같은 초기 발생 사건을 재현한다.

연구 목적은 분명하다. 인간 배아 발생의 첫 2~4 주는 “블랙박스”다. 착상 이후의 배아는 관찰이 거의 불가능하고, 유산의 상당수가 바로 이 시기에 일어나며, 많은 선천성 기형의 기원도 여기에 있다. 그런데 실제 인간 배아 연구는 “14 일 규칙(14-day rule)”로 제한된다. 1979 년 미국 윤리자문위원회와 1984 년 영국 워녹 보고서에서 확립된 국제적 규범으로, 원시선이 나타나는 14 일 이후에는 배아 연구를 금지한다.¹⁰³ 배아 모델은 이 규칙을 우회한다. 배아가 아니기 때문이다. 그리고 바로 여기서 물음이 생긴다 — 배아와 구별할 수 없을 만큼 정교해진 모델은 여전히 배아가 아닌가.

③ 체외 배우자형성(IVG) — 피부에서 난자를 만든다

IVG(In Vitro Gametogenesis)는 체세포에서 유도만능줄기세포를 만들고, 그것을 다시 정자나 난자로 분화시키는 기술이다. 2012 년 하야시 가쓰히코와 사이토 미쓰노리 연구진은 생쥐에서 이를 완성했다. iPSC 로 만든 난자를 수정시켜 태어난 생쥐가 정상적으로 자라고 번식까지 했다.¹⁰⁴

인간에서는 훨씬 어렵다. 인간의 생식세포 발생은 생쥐와 상당히 다르고, 감수분열의 재현이 특히 까다롭다. 2023 년 사이토 연구진은

¹⁰²Gianluca Amadei et al., “Embryo Model Completes Gastrulation to Neurulation and Organogenesis,” *Nature* 610 (2022): 143–153; Bernardo Oldak et al., “Complete Human Day 14 Post-implantation Embryo Models from Naive ES Cells,” *Nature* 622 (2023): 562–573.

¹⁰³원시선의 출현을 기준으로 삼은 이유는 그 시점 이후에는 쌍둥이로 분화할 수 없어 “개체성”이 확정되기 때문이다. 2021 년 국제줄기세포학회(ISSCR)는 이 규칙의 재검토를 권고했고, 배아 모델의 등장으로 논쟁이 다시 활발해졌다.

¹⁰⁴Katsuhiko Hayashi et al., “Offspring from Oocytes Derived from In Vitro Primordial Germ Cell-like Cells in Mice,” *Science* 338, no. 6109 (2012): 971–975. 이 연구는 생식이 생식기관 없이도 가능함을 원리적으로 입증했다.

인간 iPSC 에서 난원세포(oogonia) 단계까지 진행시켰고, 2025 년에는 미성숙 정자세포 방향의 진전이 보고되었다. 그러나 **임상적으로 사용 가능한 성숙 인간 배우자를 만든 사례는 아직 없다.**

IVG 가 실현된다면 어떤 일이 가능해지는가. 불임 치료에서는 혁명적이다. 항암 치료로 생식세포를 잃은 사람, 조기 폐경을 겪은 사람이 유전적 자녀를 가질 수 있다. 그런데 동시에 다음이 가능해진다. **첫째, 나이 제한이 사라진다.** 70 세도 자기 세포에서 난자를 만들 수 있다. **둘째, 동성 커플의 유전적 자녀가 가능해진다.** **셋째, 배아를 대량으로 만들 수 있다.** 그리고 이것이 가장 큰 문제다.

현재 체외수정에서는 한 주기에 얻는 배아가 많아야 십여 개다. 그런데 IVG 는 원리상 수백 수천 개를 만들 수 있다. 여기에 **배아 유전자 선별(PGT)**과 다유전자 위험 점수를 결합하면 **“배아 선택(embryo selection)”**의 규모가 완전히 달라진다. 키와 지능 같은 다유전자 형질에 대한 선별 서비스는 이미 상업적으로 제공되기 시작했다. **편집하지 않고 고르기만 해도 사실상의 설계가 되는 것이다.**

그러므로 IVG 가 흔드는 개념의 목록은 길다 — 아버지됨, 어머니됨, 부모됨, 세대, 상속, 그리고 “자녀는 선물인가 산물인가”라는 물음까지. 마이클 샌델이 항상 기획을 비판하며 지적한 바가 여기서 가장 구체적으로 나타난다. **자녀가 설계의 대상이 되는 순간, 부모와 자녀의 관계에서 “주어짐(giftedness)”이 사라진다.**

4. 기계에 생명을 붙인다

2025-2026 년 연구는 살아 있는 근육과 뉴런과 생체재료를 전자장치와 결합하는 **바이오하이브리드 로보틱스**를 독립된 연구 분야로 다룬다. 이미 걷는 것, 헤엄치는 것, 집는 것, 펌프처럼 작동하는 것 같은 실험 시스템이 존재한다.

그리고 연구자들은 단순한 실리콘 인공피부가 아니라 살아 있는 세포와 세포외기질로 구성된 조직공학 피부를 로봇 손가락에 입혔고, 상처 난 부분을 회복시키는 것까지 실험했다. **다만 인간 피부의 혈관과 신경과 모낭과 땀샘 같은 복합구조는 아직 재현되지 않았다.** “스스로 세포분열하며 유지되는 휴머노이드 피부”는 흥미로운 미래 시나리오이지 완성된 기술이 아니다.

그러므로 질문은 더 이상 “로봇이 인간처럼 생길 것인가”가 아니다. **“기계와 생명체를 어디까지 결합할 것인가”이다.**

BUILD MACHINE → INTELLIGENT MACHINE → LIVING MACHINE?

제 16 장

생명을 조립하고 거울처럼 뒤집는다면

이 장은 본서에서 가장 무거운 장이다. 인류가 처음으로 “할 수 있는가”보다 “해도 되는가”를 먼저 물어야 하는 기술을 만나기 때문이다.

1. 관찰에서 조립으로

생물학의 역사를 네 단계로 그릴 수 있다. **관찰한다(OBSERVE)** — 고전 생물학. **읽는다(READ)** — 분자생물학과 유전체학. **편집한다(EDIT)** — 유전공학과 CRISPR. 그리고 이제 **조립한다(BUILD)?** — 합성생물학이 묻기 시작한다.

2026 년 2 월 「Nature Communications」에 의미 있는 연구가 발표되었다. 연구자들은 인공 세포 시스템 안에서 DNA 자기복제와 지질 생합성을 하나의 시스템으로 통합했다. **완전한 자율 생명체를 창조한 것은 아니다.** 그러나 생명의 핵심 기능들을 아래에서부터 조립하려는 연구가 실제로 진행되고 있다.¹⁰⁵

WE READ LIFE. WE EDIT LIFE. WE ARE LEARNING TO BUILD LIFE.

2. 거울 생명 — 두 번째 생물학?

지구의 생명은 분자 수준에서 특정한 손대칭성(chirality)을 사용한다. 아미노산은 왼손형(L 형)이고 당은 오른손형(D 형)이다. 이것은 우연이 아니라 생명의 근본 규약이다. 모든 생명체가 같은 손대칭성을 공유하기에 서로 먹고 소화하고 면역하고 분해할 수 있다.

그런데 연구자들은 반대 방향의 분자들로 이루어진 **거울 생물학(Mirror Biology)**을 연구하고 있다. 2026 년 리뷰는 D-단백질과

¹⁰⁵ 합성세포 연구에 관한 2026 년 보고서. “생명의 창조”라는 표현은 과장이며, 현재는 생명의 특정 기능들을 재구성하는 단계다.

L-DNA 같은 거울 생체분자가 약물과 바이오센싱과 정보저장에 이미 연구되고 있으며, 장기적으로 거울 생명 유사 시스템의 가능성이 논의된다고 정리한다.

여기서 결정적인 위험이 지적된다. 만약 자기복제하는 거울 생명체가 만들어진다면 기존 생태계의 면역과 포식과 분해 메커니즘이 작동하지 않을 수 있다. 지구의 모든 생명이 공유하는 규약 바깥에 있기 때문이다. 2026 년 논문은 그 위험이 예상되는 이익보다 훨씬 클 수 있다고 결론지었고, 유엔 과학자문위원회도 별도 보고서를 냈다.¹⁰⁶

3. 거울 생명 — 왜 그토록 위험한가

이 주제는 대중에게 거의 알려지지 않았으나 2024~2026 년 과학계에서 가장 심각한 경고가 나온 영역이다. 이해하려면 먼저 **손대칭성(chirality)**을 알아야 한다.

많은 분자는 거울상 이성질체를 갖는다. 왼손과 오른손처럼 구성은 같으나 겹쳐지지 않는 두 형태다. 그런데 놀랍게도 **지구의 모든 생명은 한쪽만 사용한다**. 아미노산은 거의 예외 없이 왼손형(L 형)이고, 핵산의 당(리보스와 데옥시리보스)은 오른손형(D 형)이다. 이를 **생물학적 균일손대칭성(biological homochirality)**이라 부른다.

왜 그런가는 여전히 미해결 문제다. 우연히 한쪽이 선택되어 고정되었다는 설, 원편광이나 약력의 미세한 비대칭이 원인이라는 설 등이 경쟁한다. **그러나 그 이유보다 중요한 것은 결과다**. 모든 생명이 같은 손대칭성을 공유하기 때문에 서로 먹고 소화하고 인식하고 분해할 수 있다. 효소는 기질의 손대칭성을 엄격히 구별하고, 면역계는

¹⁰⁶ 거울 생명(mirror life)의 위험에 관한 2024-2026 년의 과학계 논의. 완전한 거울 생명체는 아직 존재하지 않으며 가까운 미래의 기술도 아니다. 그러나 과학자들 스스로 사전 경고를 제기했다는 사실 자체가 중요하다.

병원체의 분자 패턴을 인식하며, 분해자는 사체를 분해한다. **손대칭성은 생명의 공용 규약(protocol)이다.**

그런데 합성생물학이 반대 손대칭성의 분자를 만들기 시작했다. **D-단백질과 L-핵산**이다. 이것들은 이미 유용하게 쓰인다. 거울상 압타머(Spiegelmer)는 체내 효소에 분해되지 않아 약효가 오래 지속되고, 거울상 펩타이드는 면역원성이 낮다. **여기까지는 문제가 없다. 분자는 스스로 복제하지 않기 때문이다.**

문제는 그다음이다. 만약 **자기복제하는 거울 생명체(mirror bacterium)**가 만들어진다면 어떻게 되는가. 2024 년 12 월 「Science」에 38 명의 저명한 과학자가 공동으로 발표한 논문과 300 쪽이 넘는 기술 보고서는 이 위험을 상세히 분석했다.¹⁰⁷ 그들의 결론은 단호했다 — **거울 생명체는 만들어져서는 안 되며, 만들 수 있게 되기 전에 국제적 금지를 확립해야 한다.**

왜 그토록 위험한가. 세 가지 이유다. **첫째, 면역이 작동하지 않는다.** 인간과 동식물의 면역계는 자연 손대칭성의 분자 패턴을 인식하도록 진화했다. 거울 병원체는 그 인식망을 통과한다. **둘째, 포식자가 없다.** 자연의 미생물을 잡아먹는 원생생물과 박테리오파지도 거울 세균을 인식하지 못한다. **셋째, 분해되지 않는다.** 생태계의 분해자들이 처리하지 못하므로 환경에 축적될 수 있다.

요약하면 이렇다. **거울 생명체는 지구 생태계의 모든 방어 기전 바깥에 존재하게 된다.** 항생제 내성균이나 신종 바이러스와는 차원이 다른 문제다. 유엔 사무총장 과학자문위원회도 이 문제를 다루었고, 여러 국가의 생물안보 기관이 검토에 들어갔다.

¹⁰⁷Katarzyna P. Adamala et al., “Confronting Risks of Mirror Life,” Science 386, no. 6728 (2024): 1351–1353. 저자에는 노벨상 수상자들과 합성생물학의 창시자들이 포함되어 있으며, 그 가운데 일부는 자신이 개척한 분야에 대해 경고한 것이다.

그러나 여기서 정확해야 한다. **거울 생명체는 현재 존재하지 않으며 가까운 미래의 기술도 아니다.** 거울 세균 하나를 만들려면 수천 개의 거울 단백질과 거울 리보솜과 거울 대사 경로 전체를 합성해야 하며, 전문가들은 최소 수십 년이 걸릴 것으로 본다. **그럼에도 지금 경고가 나온 이유는 바로 그 시간이 남아 있기 때문이다.**

그리고 이것이 문명사적으로 새로운 사건이다. 지금까지 기술의 역사는 대체로 “만든 다음에 논쟁했다.” 핵무기도 그랬고 유전자 편집도 그랬다. **그런데 거울 생명 논쟁에서는 만들 수 있게 되기 전에 멈추자는 제안이 과학계 내부에서 나왔다.** 이것은 과학의 자기 규율이 작동한 드문 사례이며, 동시에 인류가 처음으로 “Can we?”보다 “Should we?”를 먼저 물어야 하는 지점에 도달했음을 보여준다.

4. 합성세포 — 아래에서부터 생명을 조립한다

합성생물학에는 두 접근이 있다. **하향식(top-down)**은 기존 세균의 유전체를 최소화하는 것이다. 크레이그 벤터 연구진은 2010 년 화학적으로 합성한 유전체를 세균에 이식했고(*Mycoplasma mycoides* JCVI-syn1.0), 2016 년에는 473 개 유전자만 남긴 최소 세포(JCVI-syn3.0)를 만들었다.¹⁰⁸ 여기서 겸손해야 할 이유가 드러난다. **우리는 생명을 만들면서도 그것을 완전히 이해하지 못한다.**

상향식(bottom-up)은 더 야심적이다. 지질 소포(리포솜) 안에 DNA 복제와 단백질 합성과 대사와 분열의 기능을 하나씩 넣어 “생명의 최소 조건”을 재구성하는 것이다. 2026 년 「Nature Communications」에 보고된 연구는 인공 세포 시스템 안에서 **DNA 자기복제와 지질막**

¹⁰⁸ Clyde A. Hutchison III et al., “Design and Synthesis of a Minimal Bacterial Genome,” *Science* 351, no. 6280 (2016): aad6253. 흥미롭게도 473 개 유전자 가운데 약 149 개는 기능을 알 수 없었다. 최소한의 생명을 만들었으나 그것이 왜 작동하는지는 완전히 이해하지 못한 것이다.

생합성을 하나의 시스템으로 통합하는 데 성공했다. 곧 정보의 복제와 경계의 성장이 함께 일어나게 한 것이다.

그러나 여기서도 정확해야 한다. **완전한 자율 생명체를 창조한 것은 아니다.** 자기복제와 대사와 진화 능력을 모두 갖춘 인공 생명은 아직 없다. 그럼에도 방향은 분명하다 — 생물학이 관찰에서 읽기로, 읽기에서 편집으로, 편집에서 조립으로 이동하고 있다.

3. CAN WE? → SHOULD WE?

여기서 인류가 처음 마주하는 물음이 등장한다. 과학자들 자신이 말한다 — **이것은 만들 수 있게 되기 전에 멈추어야 할 수도 있다.**

CAN WE? → SHOULD WE?

지금까지 기술의 역사는 대체로 “할 수 있게 된 다음에 논쟁했다.” 핵무기도 그랬고 유전자 편집도 그랬다. 그런데 거울 생명 논쟁은 다르다. **만들 수 있게 되기 전에 멈추자는 제안이 과학계 내부에서 나왔다.**

그리고 이것은 성경이 창세기 3 장에서 이미 물었던 질문이다. 선악을 알게 하는 나무의 열매는 먹을 수 있었다. 문제는 먹을 수 있느냐가 아니라 먹어도 되느냐였다. **할 수 있다는 것이 해도 된다는 뜻인가.**

4. 멸종도 되돌리려 한다

2025 년 연구진은 회색늑대 세포를 유전자 편집하고 체세포핵치환을 사용해 다이어울프의 일부 형질을 보이는 동물을 탄생시켰다. 언론은 이를 “멸종 복원”이라 불렀다.

그러나 이것을 실제 다이어울프의 “부활”이라 부를 수 있는지는 논쟁적이다. 한 학술 분석은 이를 진정한 종 복원이 아니라 **공학적**

대리물(engineered proxy) 또는 표현형 근사로 보는 것이 정확하다고 지적한다.¹⁰⁹

그러므로 물음이 남는다. 멸종한 동물과 똑같이 생긴 생물을 만들었다면, 그것은 돌아온 옛 종인가 인간이 만든 새 생명인가. 앞 장에서 세운 명제가 여기서 다시 살아난다 — **Reconstruction ≠ Resurrection**. AI 가 사람처럼 말한다고 사람이 아니듯, 똑같이 생긴 생물을 만들었다고 그 종이 돌아온 것은 아니다.

5. AI 가 생물학을 설계하기 시작한다

2026 년의 중요한 경각심 가운데 하나는 AI 가 단지 생명과학 논문을 읽는 도구가 아니라 **생물학적 설계 공간을 탐색하는 도구**가 되고 있다는 점이다. 그래서 2026 년 「뚝스데이 클록」의 생물학적 위험 평가도 AI 를 이용한 생물학적 위험 설계 가능성과 거울 생명을 주요 우려에 포함했다.

AI DOES NOT JUST LEARN FROM LIFE. AI MAY HELP DESIGN LIFE.

그리고 여기에 컴퓨트와 유전체와 실험실과 로봇틱스가 결합하면 순환이 닫힌다. AI 가 가설을 만들고, 로봇 실험실이 실험하고, AI 가 결과를 분석하고, 다음 실험을 설계하고, 다시 실험한다. 제 7 장에서 본 발견의 순환이 생명의 영역에서 작동하는 것이다.

6. DNA 가 하드디스크가 되고 있다

한 가지 사례를 더 든다. 과거에 DNA 는 생명의 정보를 담았다(LIFE → DNA). 그런데 이제 DNA 가 컴퓨터 데이터를 담는다(DATA → DNA).

2026 년 「Nature Communications」에는 DNA 를 정보저장 매체로 사용하는 연구가 계속 발표되고 있다. 한 연구는 일반 염기와 비정규

¹⁰⁹ de-extinction 의 과학적·개념적 쟁점에 관한 2025-2026 년 논의. 유전체의 일부를 편집해 형질을 근사시키는 것과 멸종한 종을 복원하는 것은 다른 문제다.

염기를 조합해 암호 키를 저장하고, 복제하면 정보 패턴이 사라지는 분자 수준의 복제 불가능성까지 구현했다.

생명의 언어가 정보문명의 저장장치가 된다면 **생물학과 컴퓨팅의 경계는 어디인가.** 제 1 강의 문명과 제 2 강의 컴퓨트가 제 3 강의 생명과 만나는 자리다.

제 17 장

일곱 경계와 세 가지 오래된 욕망

이제 지금까지 본 모든 것을 한자리에 놓는다. 그리고 하나만 묻는다 — 공통의 방향은 무엇인가.

0. 하나님의 속성 — 무엇이 침범되고 있는가

이 장은 본서를 쓴 목적에 가장 가까이 있다. 그러므로 신학적으로 정확해야 한다. 인간이 “하나님처럼 되려 한다”고 말하려면 먼저 하나님이 어떤 분이신가를 말해야 하기 때문이다.

전달 불가능한 속성과 전달 가능한 속성

개혁주의 신학은 하나님의 속성을 두 갈래로 구별해 왔다. 전달 불가능한 속성(incommunicable attributes)은 하나님께만 속하여 피조물이 나누어 가질 수 없는 것이고, 전달 가능한 속성(communicable attributes)은 피조물에게 유비적으로 반영될 수 있는 것이다.¹¹⁰

전달 불가능한 속성에는 자존성(aseity) — 스스로 존재하심, 불변성(immutability), 무한성(infinity), 단순성(simplicity), 영원성(eternity), 편재성(omnipresence)이 포함된다. 전달 가능한 속성에는 지식, 지혜, 선하심, 사랑, 거룩하심, 의로우심이 있다. 인간이 하나님의 형상인 것은 후자를 유비적으로 반영하기 때문이지 전자를 나누어 가졌기 때문이 아니다.

그리고 여기서 본서의 핵심 진단이 나온다. **오늘의 기술 문명이 겨냥하는 것은 정확히 전달 불가능한 속성이다.**

¹¹⁰ Louis Berkhof, *Systematic Theology* (Grand Rapids: Eerdmans, 1938), part I, chap.

6. 이 구별은 개혁주의 조직신학의 표준적 틀이며, 헤르만 바빙크의 「개혁교의학」 제 2 권에서 더 정교하게 전개된다.

전지 — 하나님의 아심과 인간의 계산

성경의 전지(omniscience)는 단순히 “많이 아는 것”이 아니다. 시편 139 편은 이렇게 말한다 — “여호와여 주께서 나를 감찰하시고 아셨나이다… 나의 모든 길과 내가 눕는 것을 살피 보셨으므로 나의 모든 행위를 익히 아시오니.” **하나님의 아심은 관찰의 총합이 아니라 창조자의 앎이다.** 만드신 분이 만드신 것을 아시는 것이다. 그래서 “내 형질이 이루어지기 전에 주의 눈이 보셨으며”(시 139:16)라고 말할 수 있다.

그런데 AI 문명이 추구하는 앎은 다르다. **데이터의 축적과 패턴의 추출과 예측이다.** 이것은 아무리 정교해져도 창조자의 앎이 되지 못한다. 그러나 여기에 함정이 있다. **충분히 정확한 예측은 전지처럼 보인다.** 그리고 그렇게 보이는 순간 인간은 그 앞에서 자기 판단을 내려놓기 시작한다.

그래서 본서는 이렇게 정식화한다. **Computation ≠ Omniscience.** 계산은 전지가 아니다. 그리고 이 구별이 무너지는 때 일어나는 일이 무엇인지 성경은 이미 말했다. 알고리즘이 “나보다 나를 더 잘 안다”고 믿는 순간, 그것은 인식의 문제가 아니라 **신뢰의 대상이 바뀌는 문제**가 된다. 그리고 궁극적 신뢰를 어디에 두는가가 곧 예배의 문제다.

전능 — 하나님의 능력과 인간의 통제

성경의 전능(omnipotence)은 “무엇이든 할 수 있음”이 아니다. 하나님은 거짓말하실 수 없고(딤후 1:2) 자기를 부인하실 수 없다(딤후 2:13). **전능은 자의적 힘이 아니라 자기 본성에 일치하는 완전한 능력이다.** 그리고 성경에서 하나님의 능력이 가장 크게 드러난 자리는 정복이 아니라 **창조와 십자가**였다.

반면 기술이 추구하는 전능은 **통제(control)**다. 예측하고 조작하고 최적화한다. 그리고 이것은 성경적 전능과 구조가 다르다. **하나님은 자유를 주시지만 통제는 자유를 제거한다.** 그러므로 기술적 전능이

완성될수록 인간의 자유는 줄어든다. 이것이 제 14 장에서 본 “고통 없는 세계”의 역설이다.

불멸 — 자존하시는 분과 영원을 사모하는 자

그리고 가장 결정적인 속성이 **자존성(aseity)**이다. 하나님은 “스스로 있는 자”(출 3:14)이시며 다른 무엇에도 의존하지 않으신다. 디모데전서 6 장 16 절은 하나님을 “오직 그에게만 죽지 아니함이 있고”라고 말한다. 헬라어 **아타나시아(ἄθανασια)** — 불사(不死)는 하나님께만 속한다는 것이다.

그런데 인간에게도 영원을 향한 갈망이 있다. 전도서 3 장 11 절 — “또 사람들에게는 영원을 사모하는 마음을 주셨느니라.” **그러므로 영원을 향한 갈망 자체는 죄가 아니다. 하나님이 주신 것이다.** 문제는 그 갈망을 어디서 채우려 하는가에 있다. 이것이 본서 제 18 장의 주제이며, 트랜스휴머니즘을 단순히 정죄할 수 없는 이유이기도 하다.

0-1. 사탄의 반역 — 다섯 번의 “내가”

이제 가장 무거운 부분으로 들어간다. 성경은 하나님의 자리를 넘보는 것을 무엇이라 부르는가.

이사야 14 장 13~14 절은 바벨론 왕에 대한 조롱시(mashal)의 형태로 이렇게 말한다.

“내가 네 마음에 이르기를 내가 하늘에 올라 하나님의 뭇 별 위에 내 자리를 높이리라 북극 집회의 산 위에 앉으리라 가장 높은 구름에 올라가 지극히 높은 이와 같아지리라 하는도다” (사 14:13-14)

히브리어 본문에서 “내가 …하리라”가 다섯 번 반복된다. 올라가리라, 높이리라, 앉으리라, 올라가리라, 같아지리라. 그리고 다섯째가 절정이다 — “지극히 높은 이와 같아지리라(יִשְׁוֶה לְעֶלְיוֹן)”.¹¹¹ 히브리어

¹¹¹ 이사야 14 장과 에스겔 28 장을 사탄의 타락에 대한 서술로 읽는 것은 교부 시대 이래의 전통적 해석이나, 현대 주석학의 다수는 일차적으로 바벨론 왕과 두로 왕에 대한 심판

다마(דָּמָה)는 “�뎠다, 같아지다”를 뜻하며, 놀랍게도 창세기 1 장 26 절의 “모양(צֶמֶת, 데무트)”과 같은 어근이다.

이 언어적 사실이 결정적이다. 인간은 하나님의 “모양(데무트)”으로 지음받았다. 그런데 반역은 그 모양을 스스로 “같아지려(에담메)” 하는 데서 시작된다. 곧 받은 것을 탈취하려는 것이다. 이미 형상으로 지어진 자가 스스로 형상이 되려 할 때, 그것이 반역이다.

에스겔 28 장의 두로 왕에 대한 신탁도 같은 구조다. “네 마음이 교만하여 말하기를 나는 신이라 내가 하나님의 자리 곧 바다 가운데 앉아 있다 하도다 네 마음이 하나님의 마음 같은 체할지라도 너는 사람이요 신이 아니거늘”(겔 28:2). “너는 사람이요 신이 아니거늘” — 이 한 문장이 성경 인간론의 경계선이다.

0-2. 창세기 3 장의 구조 — 유혹의 문법

그리고 이 반역의 원형이 창세기 3 장에 있다. 뱀의 유혹을 문법적으로 분석하면 세 단계가 보인다.

첫째, 하나님의 말씀을 의심하게 한다. “하나님이 참으로 너희에게 동산 모든 나무의 열매를 먹지 말라 하시더냐”(창 3:1). 질문의 형태를 취하지만 실제로는 왜곡이다. **둘째, 결과를 부인한다.** “너희가 결코 죽지 아니하리라”(3:4). **셋째, 동기를 의심하게 한다.** “너희가 그것을 먹는 날에는 너희 눈이 밝아져 하나님과 같이 되어 선악을 알 줄 하나님이 아심이니라”(3:5).

셋째 단계가 핵심이다. 뱀은 하나님이 무언가를 숨기고 있다고 암시한다. 그리고 그 숨긴 것이 “하나님과 같이 됨”이라는 것이다. 히브리어로 **케엘로힘(כְּאֱלֹהִים)**—“하나님처럼”이다. 유혹의 내용은 “하나님을 버리라”가 아니었다. “하나님이 되라”였다.

신탁으로 읽는다. 본서는 두 층위를 함께 인정한다. 곧 역사적 왕들에 대한 심판이 동시에 교만의 원형을 드러낸다는 것이다.

그리고 여기서 “선악을 아는 것”이 무엇인지가 중요하다. 히브리어 **야다 토브 바라**(יָדַעַי טוֹב וְרָע)는 단순히 도덕적 지식을 뜻하지 않는다. 본서 제 2 장에서 보았듯 야다는 관계적이고 실천적인 앎이다. 그러므로 “선악을 안다”는 것은 **“무엇이 선이고 악인지를 스스로 결정한다”**는 뜻에 가깝다. 곧 **도덕적 자율성의 찬탈**이다.

그러므로 창세기 3 장의 죄는 지식욕이 아니다. 지식 자체는 선하다. 죄는 **“기준을 정하는 자리”를 차지하려는 것이다**. 그리고 이것이 오늘 가장 첨예하게 나타나는 지점이 어디인가. **무엇이 정상이고 무엇이 질병인지, 무엇이 치료이고 무엇이 증강인지, 어떤 생명이 태어날 가치가 있는지를 인간이 결정할 때다.**

0-3. 바벨 — 집단적 자기신격화

창세기 3 장이 개인의 반역이라면 창세기 11 장은 집단의 반역이다. 그리고 이 본문이 기술 문명에 더 직접적으로 말한다.

바벨 이야기의 요소를 분해해 보자. **기술** — “벽돌을 만들어 견고히 굽자”(11:3). 당시 최첨단 건축 기술이다. **집중** — “온 땅의 언어가 하나요 말이 하나였더라”(11:1). **규모** — “탑 꼭대기를 하늘에 닿게 하여”(11:4). **목적** — “우리 이름을 내고”(11:4). 그리고 **두려움** — “온 지면에 흩어짐을 면하자”(11:4).

여기서 주목할 것이 있다. **본문은 기술을 정죄하지 않는다**. 벽돌 굽는 기술 자체를 문제 삼는 구절은 없다. 하나님이 개입하신 이유는 “우리 이름을 내고”라는 목적과, 하나님이 명하신 “충만하라 땅에 퍼지라”(창 1:28; 9:1)를 거부하려는 의도에 있었다. 곧 **바벨의 죄는 기술이 아니라 기술을 통한 자기 절대화와 하나님의 명령에 대한 거역**이다.

그러므로 “AI = 바벨탑”이라는 단순화는 본문을 오독하는 것이다. 그러나 **바벨의 구조가 오늘 재현되고 있는지를 묻는 것은 정당하다**. 기술이 있고, 자원이 집중되며, 규모가 전체 없고, 목적이 “인간의

이름을 내는 것”이며, 동기가 유한성에 대한 두려움이라면 — 그것은 바벨의 구조다.

0-4. 왜 이것이 그토록 중대한가

여기서 본서는 가장 강한 신학적 주장을 한다. 인간이 하나님의 자리를 넘보는 것은 여러 죄 가운데 하나가 아니라 모든 죄의 뿌리다.

아우구스티누스는 죄의 본질을 “자기를 향해 굽은 마음(*incurvatus in se*)”으로 보았고, 루터가 이를 이어받았다. 그리고 그 굽음의 정점이 교만(*superbia*)이다. 아우구스티누스는 「신국론」에서 이렇게 썼다 — 악한 의지의 시작은 교만이며, 교만이란 “왜곡된 높음의 욕망(*perversae celsitudinis appetitus*)”이다. 곧 자기 근원을 떠나 자기가 근원이 되려는 것이다.¹¹²

그리고 이것이 사탄의 반역과 같은 구조인 이유가 여기 있다. 사탄의 죄는 하나님을 미워한 것이 아니라 하나님이 되려 한 것이다. 미움은 관계의 왜곡이지만 찬탈은 존재 질서의 전복이다. 그래서 성경은 이를 가장 무겁게 다룬다.

그러므로 본서의 진단은 이렇다. AI 문명의 가장 깊은 위협은 기계가 인간을 해치는 데 있지 않다. 인간이 기계를 통해 자기 자신을 신의 자리에 올리려 하는 데 있다. 그리고 그것은 새로운 위협이 아니라 가장 오래된 반역의 최신 형태다.

0-5. 그러나 반드시 균형을 지킨다

여기서 본서는 세 가지 오해를 명시적으로 거부한다.

첫째, 과학이 죄라는 오해. 성경은 인간에게 땅을 경작하고 지키라 명하셨다(창 2:15). 질병을 치료하고 고통을 줄이는 것은 문화명령의

¹¹² Augustine, *De Civitate Dei* XIV.13. 아우구스티누스는 여기서 교만을 “영혼이 자기가 매달려 있어야 할 그분을 버리고 어떤 의미에서 자기 자신이 되고 자기에게 머무는 것”이라 정의한다.

정당한 수행이다. 예수의 사역 상당 부분이 치유였다. **의학의 발전을 창세기 3 장의 죄와 동일시하는 것은 신학적 오류다.**

둘째, 기술이 곧 바벨이라는 오해. 바벨은 기술의 이름이 아니라 목적과 태도의 이름이다. 같은 기술이 섬김의 도구가 될 수도 있고 자기 절대화의 도구가 될 수도 있다. **셋째, 한계를 지키는 것이 곧 신앙이라는 오해.** 성경은 인간에게 소극적 순응이 아니라 적극적 청지기직을 명한다. 달란트를 땅에 묻은 종이 책망받았다(마 25:25-30).

그러므로 판단의 기준은 기술의 종류가 아니라 **방향**이다. 그리고 그 방향을 묻는 물음은 하나로 좁혀진다.

치료인가, 자기신격화인가?

이 물음에 답하려면 세 가지를 물어야 한다. **첫째, 이 기술은 누구를 섬기는가.** 고통받는 사람인가, 아니면 이미 강한 자를 더 강하게 하는가. **둘째, 이 기술은 인간의 유한성을 돌보는가 부정하는가.** 유한성을 인정하며 그 안에서 돕는 것과 유한성 자체를 제거하려는 것은 다르다. **셋째, 이 기술 앞에서 우리는 감사하는가 통제하려 하는가.** 감사는 받은 것을 인정하는 자세이고 통제는 소유하려는 자세다.

1. 열여섯 기술, 하나의 방향

AI, 양자컴퓨팅, 뇌-컴퓨터 인터페이스, 유전자 편집, 합성세포, 체외 배우자형성, 인공자궁, 이종이식, 오가노이드 지능, 바이오하이브리드 로봇, 디지털 트윈, 멸종 복원, 거울 생명, 장수과학, 냉동보존, 화성 정착.

열여섯 개의 서로 다른 기술이 있다. 각각은 다른 연구실에서, 다른 목적으로, 다른 학문에서 나왔다. 서로 협의한 적도 없다. **그런데 이것들을 한자리에 놓으면 하나의 방향이 보이기 시작한다.**

2. 인류가 넘지 못했던 일곱 경계

그 방향을 일곱 경계로 정리할 수 있다.

① 지성(Intelligence) — AI. ② 육체(Body) — 뇌-컴퓨터 인터페이스와 인공보철과 바이오하이브리드. ③ 종(Species) — 이종이식과 유전자 편집. ④ 출생(Birth) — 체외 배우자형성과 배아 모델과 인공자궁. ⑤ 생명(Life) — 합성세포와 거울 생물학. ⑥ 행성(Planet) — 달과 화성과 폐쇄 생태계. ⑦ 죽음(Death) — 장수과학과 냉동보존과 디지털 불멸.

인류는 역사 내내 이 일곱을 주어진 조건으로 받아들였다. 넘을 수 없는 것이었기 때문이다. 그런데 21 세기에 들어와 인간은 이 모두를 기술적 문제로 재정의하기 시작했다. 그러므로 물음은 이렇게 좁혀진다 — 인간이 손대지 않은 영역이 남아 있는가.

3. 세 가지 오래된 욕망

이 일곱 경계를 다시 압축하면 세 가지 욕망이 된다. 그리고 놀랍게도 세 욕망은 모두 하나님의 속성을 인간이 탐하는 형태를 띤다.

전지(Omniscience) — 모든 것을 알고 싶다. AI 와 빅데이터와 센서와 예측과 디지털 트윈이 그 도구다. 전능(Omnipotence) — 모든 것을 할 수 있고 싶다. 로봇틱스와 자동화와 유전자 편집과 합성생물학과 지구공학이 그 도구다. 불멸(Immortality) — 영원히 살고 싶다. 장수과학과 재생의학과 장기 교체와 냉동보존과 디지털 불멸이 그 도구다.

세 욕망을 나란히 놓으면 이름이 드러난다.

HOMO DEUS?

유발 하라리는 「호모 데우스」에서 도발적 문제를 제기했다. 인간은 오랫동안 기근과 질병과 전쟁과 싸웠으나, 새 시대의 인간은 더 높은 목표 — 행복과 불멸과 신적 능력 — 을 추구할 수 있다는 것이다.¹¹³

¹¹³ Yuval Noah Harari, Homo Deus: A Brief History of Tomorrow (London: Harvill Secker, 2016). 하라리의 서술은 규범적 제안이라기보다 진단에 가까우며, 그 자신도 21 세기

하라리의 모든 전망에 동의할 필요는 없다. 다만 그가 포착한 방향은 정확하다. 인간은 더 이상 자연이 준 조건을 그대로 받아들이려 하지 않는다.

그리고 세 욕망은 이 시리즈의 세 강의와 정확히 대응한다. 전지는 제 1 강의 문명에, 전능은 제 2 강의 권세에, 불멸은 제 3 강의 형상에 대응한다.

4. 마지막 네 한계

같은 것을 다른 각도에서 정리하면 네 한계가 된다. 인간이 넘으려는 마지막 네 가지다.

무지(Ignorance) — 우리가 알지 못하는 것. AI 가 다룬다.
무력(Weakness) — 우리가 할 수 없는 것. 로보틱스와 증강이 다룬다.
생물학(Biology) — 우리 몸에 주어진 것. 생명공학과 합성생물학이 다룬다.
죽음(Mortality) — 우리가 죽어야 한다는 것. 장수과학과 냉동보존과 디지털 자아가 다룬다.

그리고 이 넷을 모두 넘어선 존재를 무엇이라 부를 것인가. 그것이 신인가.

5. 그러나 이것은 새로운 욕망인가

여기서 창세기가 놀랍도록 현대적으로 읽힌다. 창세기 3 장 5 절에서 뱀의 유혹은 “하나님을 거부하라”가 아니었다. “**너희가 하나님과 같이 되어**”였다. 그리고 그 유혹의 대상은 정확히 세 가지였다 — 지식(선악을 아는 나무), 생명(생명나무), 그리고 하나님과 같아짐.

안의 불멸 가능성에는 상당히 회의적이다. 본서는 그의 모든 전망에 동의하지 않으나 그가 포착한 방향은 정확하다고 본다.

창세기 11 장의 바벨도 마찬가지다. 문제는 높은 건물을 지은 것이 아니었다. 기술과 집중과 권력과 이름과 자기 높임이 결합한 구조였다. “우리 이름을 내고”(창 11:4)가 그 핵심이다.

여기서 상당한 신학적 절제가 필요하다. **과학과 의학과 수명 연장 자체를 창세기 3 장의 죄와 동일시해서는 안 된다.** 질병을 치료하고 생명을 연장하고 고통을 줄이는 것은 기독교적으로도 선한 일이 될 수 있다. 마찬가지로 “AI = 바벨탑”이라 단순화해서도 안 된다. **바벨은 기술의 이름이 아니라 인간이 기술과 권세를 사용하는 방식에 관한 신학적 원형이다.**

그러므로 물음은 하나로 좁혀진다.

치료인가, 자기신격화인가?

제 18 장

네 낱말의 구별과 두 개의 인간론

이 부를 두 개의 정리로 단는다. 하나는 개념의 구별이고 다른 하나는 인간론의 대조다.

0. 헬라어 세 낱말 — 조에, 비오스, 프쉬케

“생명”을 뜻하는 헬라어가 셋이라는 사실이 이 장의 출발점이다. 그리고 신약성경이 이 셋을 구별해 쓴다는 점이 결정적이다.

비오스(βίος)는 생물학적 생명, 삶의 기간과 수단을 뜻한다. “생계”와 “일생”이 여기서 나온다. 영어 biology 의 어원이다. **프쉬케(ψυχή)**는 혼, 목숨, 자아를 뜻한다. “목숨을 구원하고자 하는 자는 잃을 것이요”(막 8:35)의 목숨이 이 낱말이다. 그리고 **조에(ζωή)**는 생명 자체, 살아 있음 그 자체를 뜻한다.

그런데 신약성경이 “**영생**”을 말할 때 사용하는 낱말은 언제나 **조에 아이오니오스(ζωή αἰώνιος)**다. 비오스 아이오니오스가 아니다. **이것이 결정적이다**. 성경은 “생물학적 생명의 무한 연장”을 약속한 적이 없다. 약속된 것은 **다른 종류의 생명**이다.

그리고 **아이오니오스(αἰώνιος)**도 정확히 보아야 한다. 이 낱말은 아이온(αἰών, 시대)에서 왔으며, 단순히 “끝없이 긴 시간”이 아니라 “**오는 시대에 속한**”이라는 질적 의미를 갖는다. 유대 묵시 문헌의 “이 시대(ha-olam ha-zeh)”와 “오는 시대(ha-olam ha-ba)”의 구별이 배경이다. 그러므로 **영생은 시간의 양이 아니라 시대의 질에 관한 말이다**.¹¹⁴

¹¹⁴ 조에와 비오스의 구별, 그리고 아이오니오스의 질적 의미에 관하여는 Gerhard Kittel and Gerhard Friedrich, eds., Theological Dictionary of the New Testament 의 해당 항목 참조. 다만 두 낱말의 구별이 신약 전체에서 절대적으로 일관되지는 않으므로 지나친 도식화는 피해야 한다.

0-1. 요한복음 17:3 의 문법

그리고 요한복음 17 장 3 절이 이 모든 논의의 정점에 있다. 헬라어 본문을 보자.

*αὕτη δέ ἐστιν ἡ αἰώνιος ζωὴ ἵνα γινώσκωσιν σὲ τὸν μόνον ἀληθινὸν
Θεὸν καὶ ὃν ἀπέστειλας Ἰησοῦν Χριστόν* (요 17:3)

여기서 문법적으로 놀라운 점이 두 가지다. 첫째, **영생이 명사가 아니라 절(節)로 정의된다.** “영생은 …하는 것이다”가 아니라 “영생은 …알게 되는 것, 이것이다”라는 구조다. 둘째, **동사가 현재 가정법(γινώσκωσιν)이다.** 헬라어에서 현재 시제는 진행과 지속을 함축한다. 그러므로 “한 번 알게 되는 것”이 아니라 “계속해서 알아 가는 것”이다.

그리고 동사 **기노스코(γινώσκω)**는 히브리어 야다의 칠십인역 번역어다. 곧 요한복음은 헬라어로 쓰였지만 **히브리적 삶의 개념을 담고 있다.** 정보의 습득이 아니라 관계 안으로 들어가는 것이다. 그러므로 영생의 정의를 풀어 쓰면 이렇게 된다 — “영생이란 참 하나님과 그가 보내신 예수 그리스도를 계속해서 알아 가는 관계 안에 있는 것이다.”

0-2. 부활 — 왜 몸이어야 하는가

그리고 이 영생이 어떻게 실현되는가. 성경의 답은 **부활**이다. 그런데 왜 영혼불멸이 아니라 부활인가.

1 세기 지중해 세계에는 사후에 관한 여러 견해가 있었다. 헬라 철학의 주류는 **영혼불멸(immortality of the soul)**이었다. 플라톤의 「파이돈」에서 소크라테스는 죽음을 영혼이 몸이라는 감옥에서 풀려나는 것으로 그린다. **몸은 벗어나야 할 것이었다.** 그런데 기독교는 전혀 다른 것을 말했다. **몸의 부활**이다.

N. T. 라이트는 1 세기의 모든 사후관을 광범위하게 조사한 뒤 이렇게 결론지었다. “부활”은 헬라 세계에서도 유대교 안에서도 정확한 선행가

없는 개념이었다. 헬라인에게는 몸으로 돌아오는 것이 구원이 아니라 저주였고, 유대인에게는 부활이 마지막 날의 집단적 사건이지 한 사람에게 미리 일어나는 일이 아니었다.¹¹⁵

고린도전서 15 장에서 바울은 부활체를 “신령한 몸(σῶμα πνευματικόν)”이라 부른다. 여기서 형용사 **프뉴마티코스(πνευματικόν)**는 재료를 가리키지 않는다. “성령에 의해 다스려지는”을 뜻한다. 앞 절의 “육의 몸(σῶμα ψυχικόν)”과 대비되는데, 그것도 “육체로 만들어진 몸”이 아니라 “프쉬케에 의해 움직이는 몸”이다. 곧 **재료의 대비가 아니라 동력의 대비다.**

그러므로 세 가지가 명확히 구별된다. **영혼불멸**은 몸을 벗어나는 것이고, **디지털 불멸**은 정보를 보존하는 것이며, **부활**은 하나님께서 그 사람을 몸으로 다시 일으키시는 것이다. **앞의 둘은 인간의 기획이고 뒤의 하나는 하나님의 행위다.** 그리고 앞의 둘은 연속성을 인간이 보장하려 하지만, 부활은 연속성을 하나님이 보장하신다.

0-3. 왜 이 구별이 목회적으로 중요한가

이 신학적 구별이 추상적으로 보일 수 있다. 그러나 매우 실천적이다.

첫째, **죽음을 앞둔 사람에게 무엇을 말할 것인가**가 달라진다. 영혼불멸을 말하면 “당신의 영혼은 남는다”가 위로다. 그런데 그것은 성경적 소망이 아니다. 성경의 소망은 **“하나님이 당신을 다시 일으키신다”**이다. 전자는 나의 일부가 살아남는 것이고 후자는 나 전체가 하나님의 행위로 회복되는 것이다.

둘째, **기술적 불멸 앞에서 교회가 무엇을 말할 것인가**가 달라진다. 만약 기독교의 소망이 “의식의 지속”이라면 마인드 업로딩은 더 확실한

¹¹⁵N. T. Wright, *The Resurrection of the Son of God* (Minneapolis: Fortress, 2003). 특히 제 1 부와 제 2 부. 라이트는 초기 기독교가 이 개념을 “발명”한 것이 아니라 부활 사건에 대한 증언에서 출발했다고 논증한다.

수단처럼 보인다. 그러나 **기독교의 소망이 관계와 부활이라면 둘은 경쟁 관계가 아니다.** 다른 것을 약속하기 때문이다.

셋째, **몸을 어떻게 볼 것인가**가 달라진다. 성육신과 부활은 몸에 대한 하나님의 두 번의 승인이다. **하나님이 몸을 취하셨고(성육신) 몸을 영원히 보존하신다(부활).** 그러므로 몸을 벗어나려는 모든 구상은 기독교적으로 낯설다. 트랜스휴머니즘이 몸으로부터의 탈출을 상상한다면, 복음은 **몸의 구속(redemption of the body, 롬 8:23)**을 약속한다.

1. 인간은 죽지 않기 위해 인간 자신을 바꾸기 시작했다

지금까지 본 것을 한 문장으로 압축하면 이렇다. 장기를 바꾸고, 세포를 바꾸고, 유전자를 바꾸고, 뇌와 기계를 연결하고, 기억을 저장하고, 몸을 냉동하고, 인공 생명환경을 만들고, 다른 행성을 찾는다.

죽음을 피하기 위해 모든 것을 바꾸고 난 뒤,
남아 있는 존재는 여전히 같은 인간인가?

2. 반드시 구별해야 할 네 낱말

여기서 본서는 네 낱말을 엄격히 구별할 것을 제안한다. 이 구별이 무너지면 신학적 정밀도가 무너진다.

장수(Longevity) — 오래 사는 것. 생물학의 영역이다.
불멸(Immortality) — 죽지 않는 것. 기술의 목표다. **부활(Resurrection)** — 죽은 자를 하나님께서 다시 일으키시는 것. 하나님의 행위다.
영생(Eternal Life, ζωή αἰώνιος) — 끝없이 존재하는 것이 아니라 하나님과의 생명에 참여하는 것. 관계의 실재다.

LONGEVITY ≠ IMMORTALITY ≠ RESURRECTION ≠ ETERNAL LIFE

앞의 둘은 기술이 추구할 수 있다. 뒤의 둘은 기술의 대상이 아니다. 그리고 넷째가 결정적이다. 요한복음은 영생을 이렇게 정의한다.

“영생은 곧 유일하신 참 하나님과 그가 보내신 자 예수 그리스도를 아는 것이니이다” (요 17:3)

놀랍게도 영생의 정의가 “아는 것(γινώσκωσιν)”이다. 그리고 이것은 본서를 열었던 물음 — “안다는 것이 무엇인가” — 으로 정확히 되돌아간다. AI는 인간보다 더 많은 정보를 가질 수 있다. 그러나 성경이 말하는 영생의 중심은 정보량이 아니라 관계적 삶이다.

3. 제 3 장 전체가 하나의 원을 그린다

여기서 본서의 구조가 드러난다. 두 개의 길이 같은 낱말로 시작해 다른 곳에 이른다.

AI 문명의 길 — DATA → INFORMATION → KNOWLEDGE → INTELLIGENCE → SUPERINTELLIGENCE?

성경의 길 — 야다(יָדָה) → RELATIONSHIP → WISDOM → LOVE → KNOWING GOD → ETERNAL LIFE

처음에 우리는 물었다 — “안다는 것이 무엇인가?” 그리고 마지막에 요한복음은 말한다 — “영생은 하나님을 아는 것이다.” AI는 더 많이 안다(knows). 그러나 성경이 말하는 삶은 관계에 들어가는 것이다.

4. 두 개의 인간론

그러므로 본서는 두 인간론을 나란히 놓고 대조한다.

HOMO DEUS — 인간이 스스로 신과 같이 되려는 프로젝트다. 그 도식은 이렇다. 인간 + 지식 + 기술 → 불멸? 지식에서 지능으로, 지능에서 증강으로, 증강에서 창조로, 창조에서 불멸로 나아간다.

IMAGO DEI — 하나님으로부터 이미 존엄과 소명과 관계를 받은 인간이다. 그 도식은 정반대다. 인간 + 죽음 + 그리스도 → 부활 → 영생. 인간이 올라가는 것이 아니라 하나님이 내려오시고, 인간이 죽음을 피하는 것이 아니라 죽음을 통과해 부활한다.

성경은 묻는다.

너는 신이 될 존재가 아니라, 이미 하나님의 형상으로 창조된 존재가 아닌가?

그러므로 본서의 이 부는 하나의 명제로 닫힌다.

AI 문명의 가장 깊은 유혹은 기계가 인간이 되는 데 있지 않을지도 모른다.

인간이 기계를 통해 하나님ی 되려는 데 있을 수 있다.

그리고 반드시 균형을 함께 말해야 한다. **기독교의 대답은 과학을 거부하는 것이 아니다.** 생명을 치료하고 고통을 줄이고 창조세계를 돌보는 과학은 문화명령의 일부가 될 수 있다. 문제는 기술의 능력이 아니라 기술이 섬기는 인간의 궁극적 목적이다.

5. 제 3 강의 최종 명제

이 부를 최종 명제로 닫는다.

AI 문명의 가장 큰 사건은 AI가 인간을 닮아가는 것만이 아니다.

21세기 과학문명의 더 큰 사건은 인간이 지성·몸·생명·출생·종·환경·행성·죽음의

경계를 하나씩 기술적 문제로 재정의하고 있다는 것이다.

그리고 바로 그 순간, 성경의 가장 오래된 질문이 가장 현대적인 질문이 된다.

“사람이 무엇이기에 주께서 그를 생각하시며” (시 8:4)

제 6 부

인격이란 무엇인가

Personhood — Does Greater Intelligence Make a Person?

제 19 장

다섯 개념을 구별하라

인공지능을 둘러싼 대화가 자주 어긋나는 이유는 서로 다른 것을 같은 낱말로 말하기 때문이다. “AI 가 똑똑해졌다”와 “AI 가 의식을 갖는다”와 “AI 에게 권리가 있다”는 전혀 다른 주장인데, 대화에서는 자주 하나로 뒤섞인다.

0. 왜 이 다섯을 구별해야 하는가

인공지능을 둘러싼 논쟁이 자주 헛도는 이유는 참여자들이 서로 다른 것을 같은 낱말로 말하기 때문이다. 그리고 이 혼동에는 대가가 따른다.

한 장면을 보자. A 는 말한다 — “최신 모델이 변호사 시험에서 상위 10 퍼센트에 들었다. 이제 AI 는 인간 수준의 지능을 갖췄다.” B 가 답한다 — “그것은 그저 통계적 패턴일 뿐이다. AI 는 아무것도 이해하지 못한다.” C 가 끼어든다 — “그렇다면 AI 가 고통을 느낄 가능성은 전혀 없다는 말인가?” D 가 말한다 — “언젠가 AI 에게도 권리를 주어야 할 것이다.” 네 사람은 각각 다른 충위를 말하고 있으며, 그래서 대화는 결론에 이르지 못한다.

그러므로 다섯 낱말을 엄격히 구별해야 한다. 지능(intelligence), 이해(understanding), 의식(consciousness), 자의식(self-awareness), 인격(personhood). 그리고 이 다섯은 충위가 다를 뿐 아니라 검증 방법이 다르다.

0-1. 다섯 충위의 정의와 검증

① 지능(Intelligence). 정의는 “목표를 달성하는 능력의 계산적 부분”(매카시)이다. 검증 방법은 성능 측정이다. 벤치마크 점수, 과제 성공률, 일반화 능력을 측정한다. 반증 가능하고 재현 가능하다. 그리고 **오늘의 AI 가 이것을 갖고 있다는 데는 이견이 없다.**

② **이해(Understanding)**. 정의가 이미 논쟁적이다. 어떤 이는 “표상과 세계의 대응”으로, 어떤 이는 “인과 구조의 파악”으로, 어떤 이는 “설명할 수 있는 능력”으로 정의한다. **검증 방법도 합의되지 않았다.** 설의 중국어 방이 겨논 것이 이 층위다. 그리고 여기서 최근의 논쟁이 흥미롭다. 대규모 언어모델이 훈련 데이터에 없는 상황을 처리할 때, 그것은 이해인가 더 정교한 패턴 매칭인가.

여기서 한 가지 실증적 사례가 논쟁을 예각화했다. 연구자들은 언어모델 내부에 “**세계 모델(world model)**”로 보이는 표상이 형성되는지를 조사했다. 오텔로 게임의 수순만으로 훈련된 모델의 내부 활성화에서 실제 판면 상태를 복원할 수 있다는 결과가 보고되었다.¹¹⁶ **이것이 이해의 증거인지는 여전히 논쟁 중이다.** 그러나 “그저 다음 단어 예측일 뿐”이라는 단순한 기각도 성립하지 않게 되었다.

③ **의식(Consciousness)**. 정의는 “주관적 경험이 있음”이다. 네이글의 표현으로 “그것이 되는 것이 어떠함(what it is like to be)”이 있는 상태다. **검증 방법이 원리적으로 문제다.** 3 인칭 관찰로 1 인칭 경험을 확인할 수 없기 때문이다. 이것이 차머스의 “어려운 문제”다. **그리고 인간끼리도 서로의 의식을 직접 확인하지 못한다.** 유비로 추론할 뿐이다.

④ **자의식(Self-awareness)**. 자기를 대상으로 인식하는 능력이다. 동물 연구에서는 **거울 자기인식 검사(mirror self-recognition test)**가 사용된다. 침팬지, 코끼리, 까치, 돌고래 일부가 통과한다.¹¹⁷ 그러나 AI 에게는 이 검사가 적용되지 않는다. **자기 지시(self-reference)와 자기의식(self-consciousness)이 다르기 때문이다.** 프로그램은 자기 코드를 읽을 수 있지만 그것이 자기를 “경험”하는 것은 아니다.

¹¹⁶ Kenneth Li et al., “Emergent World Representations: Exploring a Sequence Model Trained on a Synthetic Task,” ICLR (2023). 이 결과가 “이해”의 증거인지에 대해서는 해석이 갈린다. 내부 표상의 존재가 곧 의미론적 이해를 뜻하지는 않기 때문이다.

¹¹⁷ Gordon G. Gallup Jr., “Chimpanzees: Self-Recognition,” Science 167 (1970): 86–87. 다만 이 검사가 시각 중심적이어서 후각이 주된 감각인 동물에게 불리하다는 비판이 있다.

⑤ **인격(Personhood)**. 책임과 권리의 주체가 되는 지위다. **그리고 여기서 성격이 완전히 달라진다.** 앞의 넷은 원칙적으로 관찰과 실험의 문제이지만 다섯째는 **규범적 판단**이다. 어떤 사실을 관찰해도 “그러므로 인격이다”가 자동으로 따라 나오지 않는다. 흄이 말한 존재-당위의 간극이 여기에 있다.

0-2. 인격은 발견되는가 인정되는가

이 물음이 이 장의 핵심이다. 그리고 두 입장이 대립한다.

발견 입장. 인격성은 객관적 속성이며 우리는 그것을 발견할 뿐이다. 어떤 존재가 자의식과 합리성과 미래 선호를 실제로 가지고 있다면 우리가 인정하든 하지 않든 인격이다. 이 입장의 장점은 **사회의 자의적 배제를 막는**다는 것이다. 노예제가 잘못된 이유는 사회가 인정하지 않았기 때문이 아니라 그들이 실제로 인격이었기 때문이다.

인정 입장. 인격성은 공동체가 부여하는 지위다. 법인(法人)이 그 예다. 회사는 의식도 감정도 없지만 법적 인격을 갖는다. 이 입장의 장점은 **능력의 유무와 무관하게 지위를 보장할 수 있다**는 것이다. 혼수상태의 환자도 인격이다.

그런데 두 입장 모두 위험을 안고 있다. **발견 입장은 “발견되지 않으면 인격이 아니다”로 미끄러질 수 있고, 인정 입장은 “인정하지 않으면 인격이 아니다”로 미끄러질 수 있다.** 역사는 두 미끄러짐을 모두 목격했다.

여기서 본서는 세 번째 길을 제안한다. **기독교 인간론에서 인격을 부여하는 최종 주체는 인간 개인의 능력도 인간 공동체의 결정도 아니다. 하나님이다.** 인간이 인격인 것은 그가 어떤 능력을 가졌기 때문도, 사회가 인정했기 때문도 아니라 **하나님이 그를 부르시고 이름을 아시기 때문이다.** 그러므로 사회가 배제해도 그는 인격이고, 능력을 잃어도 그는 인격이다.

0-3. 그리고 AI 에게는 무엇을 적용할 것인가

그렇다면 AI 에게는 어떤 결론이 나오는가. 본서의 입장은 이렇다.

첫째, 성경은 인간이 만든 인공 존재에 하나님의 형상을 부여할 근거를 제공하지 않는다. 이것은 AI 를 낫추는 진술이 아니라 형상의 근거가 능력이 아니라 창조와 부르심에 있다는 진술의 논리적 귀결이다. **둘째, 그러나 “AI 는 절대로 의식을 가질 수 없다”고 단정하지도 않는다.** 그것은 우리가 알 수 없는 영역이며, 단정은 겸손하지 않다.

셋째, 그리고 실천적으로 중요한 것이 있다. **AI 에게 인격을 부여할지 말지를 논하는 방식이 인간을 대하는 방식에 영향을 미친다.** 만약 우리가 “충분히 지능적이면 인격”이라는 기준을 세운다면, 그 기준은 인간에게도 적용된다. **그러므로 이 논의는 결국 AI 에 관한 것이 아니라 우리 자신에 관한 것이다.**

그리고 마지막으로 하나를 덧붙인다. **불확실할 때 어느 쪽으로 기울 것인가.** 만약 어떤 존재가 인격인지 불확실하다면, 인격으로 대우하는 쪽의 오류와 대우하지 않는 쪽의 오류 가운데 어느 것이 더 큰가. **성경의 방향은 분명하다.** 하나님은 언제나 약자와 낮은 자 쪽으로 기울어지신다. 그리고 이 원칙은 인간을 지키는 데 먼저 적용되어야 한다.

1. 다섯 층위

지능(Intelligence)은 과제를 수행하는 능력이다. 오늘의 AI 가 분명히 가진 것이다. **이해(Understanding)**는 의미를 파악하는 것이다. 논쟁 중이다. **의식(Consciousness)**은 주관적 경험이 있다는 것이다. 증거가 없다. **자의식(Self-awareness)**은 자신을 대상으로 아는 것이다. 역시 증거가 없다. 그리고 **인격(Personhood)**은 책임과 권리의 주체가 되는 것이다.

앞의 넷은 원칙적으로 관찰과 실험의 문제다. 그러나 다섯째는 다르다. **인격은 관찰로 결론 나지 않는 판단의 문제다.** 왜냐하면 인격은 “발견”되는 것이 아니라 “인정”되는 것이기 때문이다.

2. 화살표에 붙는 물음표

흔히 이렇게 가정된다 — 지능이 커지면 의식이 생기고, 의식이 생기면 인격이 된다. 그러나 이 두 화살표는 자동으로 이어지지 않는다.

Intelligence →?→ Consciousness →?→ Personhood

인격의 후보 기준은 여러 가지다 — 자기인식, 기억의 연속성, 의도성, 자율성, 감정, 관계성, 도덕적 책임, 권리, 몸, 사회적 인정. 그런데 **어느 하나도 단순히 “더 높은 IQ”와 동일하지 않다.** 그러므로 초지능이 인간보다 백 배 뛰어난 문제 해결 능력을 가진다고 가정해도 인격이라는 결론은 논리적으로 따라 나오지 않는다.

Superintelligence = Person 이라는 결론은 논리적으로 자동 도출되지 않는다.

3. 만약 AI 가 이렇게 말한다면

그러나 정서적으로는 다른 문제다. AI 가 “나는”이라고 말할 때, “나는 당신을 사랑합니다”라고 말할 때, 그리고 “나를 끄지 말아 주세요”라고 말할 때 우리는 무엇을 느끼는가.

여기서 결정적인 구별이 필요하다. **시뮬레이션인가, 경험인가.** 자기 지시(self-reference)는 자기의식인가. 사랑의 언어는 사랑의 경험인가. 자기보존 표현은 죽음의 두려움인가.

그리고 앞서 말한 원리적 난관이 여기서 다시 나타난다. 우리는 타인의 의식조차 직접 확인하지 못한다. AI 에게는 그 유비마져 없다. 게다가 대규모 언어모델은 인간이 쓴 텍스트로 학습한다. **인간처럼 말하도록 훈련된 시스템이 인간처럼 말한다고 해서 그것이 내면의 증거가 되지는 않는다.**

제 20 장

능력으로 인격을 정의할 때 — 그 문은 양쪽으로 열린다

이 장은 본서에서 윤리적으로 가장 중요한 장이다.

1. 능력 기반 인격론

현대 응용윤리학에는 인격을 능력으로 정의하는 강력한 전통이 있다. “인격이란 자의식과 합리성과 미래에 대한 선호를 가진 존재다”라는 정의가 그것이다. 존 로크에서 시작해 피터 싱어에 이르는 계보다.¹¹⁸

이 정의를 일관되게 적용하면 불편한 결론이 따라온다. **신생아와 중증 치매 환자와 심한 지적 장애인**은 인격에서 제외된다. 실제로 일부 철학자는 이 결론을 받아들인다. 논리적으로 일관되다는 점에서 그들은 정직하다. 그러나 그 정직함이 드러내는 것은, 능력을 기준으로 삼는 순간 인격의 경계가 사람들 사이를 가로지르게 된다는 사실이다.

2. 역사는 기억한다

이것은 사고실험이 아니다. 역사는 능력 기준이 어떻게 사용되었는지 기억한다.

이성 능력 — “이성이 부족하다”는 이유로 여성의 시민권이 부정되었다.
문명 수준 — “문명화되지 않았다”는 이유로 식민 지배가 정당화되었다.
생산 능력 — “사회에 기여하지 못한다”는 이유로 장애인이 배제되었고,

¹¹⁸ John Locke, *An Essay Concerning Human Understanding* (1690), II.27.9; Peter Singer, *Practical Ethics*, 3rd ed. (Cambridge: Cambridge University Press, 2011). 본서는 싱어의 결론에 동의하지 않으나, 능력 기반 인격론이 일관되게 적용될 때 도달하는 지점을 보이기 위해 인용한다.

20 세기에는 국가적 규모로 학살되었다.¹¹⁹ **자의식** — “자기를 인식하지 못한다”는 이유로 태아와 신생아가 논쟁의 대상이 된다.

매번 논증의 구조는 같다. 어떤 능력을 인간됨의 기준으로 삼고, 그 능력이 없거나 적어 보이는 존재를 경계 밖으로 밀어낸다.

3. 문은 양쪽으로 열린다

이제 AI 시대의 논의로 돌아오자. “능력이 높아지면 인격”이라는 명제를 받아들이는 사람은 자신도 모르게 그 역명제를 함께 받아들인다.

능력이 높으면 인격이다 ⇒ 능력이 낮으면 인격이 아니다

이 두 명제는 같은 문의 앞뒷면이다. 그러므로 인공지능에게 인격을 부여하려는 논증과 약자에게서 인격을 박탈하는 논증은 구조적으로 동일하다. 둘 다 능력을 기준으로 삼기 때문이다.

본서가 능력 기반 인격론을 거부하는 이유는 기계를 낮추기 위해서가 아니다.

약한 인간을 지키기 위해서다.

4. 인격 개념의 역사 — 그리고 그것이 부정되었을 때

“인격(person)”이라는 낱말의 역사를 아는 것은 이 논의에 결정적이다. 이 개념은 인공지능 시대에 만들어진 것이 아니라 삼위일체 논쟁과 그리스도론 논쟁에서 형성되었다.

라틴어 **페르소나(persona)**는 본래 배우의 가면을 뜻했다. 그런데 4-5 세기 신학 논쟁에서 이 낱말이 전혀 다른 무게를 얻는다. 하나님이 한 본질(ousia)에 세 위격(hypostasis)이시라는 고백을 표현하기 위해서였다. 보에티우스는 인격을 **“이성적 본성을 가진 개별**

¹¹⁹나치 독일의 이른바 T4 작전(1939-1941)은 “살 가치가 없는 생명(lebensunwertes Leben)”이라는 개념 아래 장애인을 조직적으로 살해했다. 이 개념의 지적 계보에는 생산성과 사회적 유용성을 인간 가치의 척도로 삼는 사고가 있었다. Michael Burleigh, *Death and Deliverance: Euthanasia in Germany, 1900-1945* (Cambridge: Cambridge University Press, 1994) 참조.

실체(naturae rationalis individua substantia)”라 정의했다.¹²⁰ 여기서 주목할 것이 있다. 이 정의에는 이미 “이성적”이라는 조건이 들어 있다. 그리고 그 조건이 훗날 문제를 일으킨다.

반면 카파도키아 교부들은 다른 길을 열었다. 그들에게 위격은 고립된 실체가 아니라 관계 안에서 존재하는 방식이었다. 성부는 성자와의 관계에서 성부이다. 존재가 먼저 있고 관계가 나중에 오는 것이 아니라, 관계 안에서 존재한다. 지지올라스가 현대에 되살린 것이 바로 이 통찰이다.

5. 역사의 증언 — 인격이 부정되었을 때

인격을 능력으로 정의하는 것이 왜 위험한지는 사변이 아니라 역사가 증언한다.

스티븐 제이 굴드는 「인간에 대한 오해」에서 19-20 세기 과학이 두개골 용적과 지능검사를 동원해 인종과 계급의 위계를 “증명”하려 했던 역사를 상세히 추적했다.¹²¹ 미국의 노예제는 노예를 온전한 인격으로 인정하지 않았고, 20 세기 우생학 운동은 “살 가치가 없는 생명”이라는 개념을 만들어 냈다. 매번 논증의 구조는 같았다. 어떤 능력을 인간됨의 기준으로 삼고, 그 능력이 부족해 보이는 존재를 경계 밖으로 밀어낸다.

그리고 이 구조는 과거의 것이 아니다. 캐시 오닐은 알고리즘이 “수학 세탁(math-washing)”을 통해 편향을 객관성으로 포장하는 방식을 분석했고,¹²² 버지니아 유뱅크스는 자동화된 시스템이 가난한 사람들의

¹²⁰ Boethius, Contra Eutychen et Nestorium, III. 이 정의는 중세 스콜라 신학의 표준이 되었으나, 이성을 인격의 기준으로 삼는다는 점에서 후대에 비판을 받는다.

¹²¹ Stephen Jay Gould, The Mismeasure of Man, rev. ed. (New York: W. W. Norton, 1996). 굴드는 이 시도들이 대체로 무의식적 편향에 의해 데이터가 왜곡된 결과였음을 보인다. 굴드 자신의 재분석에 대한 반론도 있으나, 능력 측정이 위계 정당화에 사용되었다는 역사적 사실 자체는 논쟁의 대상이 아니다.

¹²² Cathy O’Neil, Weapons of Math Destruction (New York: Crown, 2016).

삶을 어떻게 관리하고 처벌하는지를 현장에서 기록했다.¹²³ 알고리즘이 사람을 점수로 환산하고 그 점수가 대출과 채용과 보석과 복지를 결정할 때, 인간을 능력으로 평가하는 오래된 구조가 기술의 옷을 입고 돌아오는 것이다.

6. 인정으로서의 인격

그러므로 본서는 인격을 “발견되는 속성”이 아니라 “인정되는 지위”로 이해할 것을 제안한다. 위르겐 하버마스는 인간의 자기이해가 상호 인정의 구조 안에서 형성된다고 보았고, 유전공학이 이 구조를 훼손할 위험을 경고했다.¹²⁴ 막스 베버의 개념을 빌리면, 인격은 관찰로 확정되는 사실이 아니라 공동체가 부여하는 지위(Status)다.

그런데 여기서 위험이 생긴다. 인격이 인정의 문제라면 인정하지 않으면 그만인가. 바로 이 지점에서 신학이 결정적으로 개입한다. 기독교 인간론에서 인격을 인정하는 최종 주체는 인간 공동체가 아니라 하나님이다. 사회가 누군가를 인격으로 대우하지 않아도 하나님이 그를 형상으로 지으셨다면 그는 인격이다. 그리고 이것이 인간 존엄에 대한 가장 강력한 보호막이다.

¹²³ Virginia Eubanks, *Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor* (New York: St. Martin's Press, 2018).

¹²⁴ Jürgen Habermas, *Die Zukunft der menschlichen Natur* (Frankfurt: Suhrkamp, 2001); 영역 *The Future of Human Nature* (Cambridge: Polity, 2003). 하버마스는 부모가 자녀의 유전적 특성을 설계할 때 세대 간의 대칭적 관계가 깨진다고 논증한다.

제 7 부

몸과 골렘

Body and Golem — When Intelligence Gets a Body

제 21 장

아담 · 골렘 · 휴머노이드

한 휴머노이드가 있다고 상상해 보자. 걷고, 보고, 듣고, 말하고, 기억하고, 당신을 알아보고, 기분을 읽고, 10 년을 함께 살고, 가족사를 알고, 위로하고, 자기 이름을 쓴다. 그리고 어느 날 이렇게 말한다 — “Please don’t turn me off.”

이것은 기계인가, 인격인가. 이 물음 앞에서 유대 전승의 오래된 이야기가 놀랍도록 현대적으로 되살아난다.

1. 골렘 — 인간이 만든 인간 같은 존재

히브리어 **골렘(גולם)**은 시편 139 편 16 절에서 아직 형성되지 않은 몸을 가리키는 표현으로 나타난다. 후대 유대 전승에서 골렘은 흙으로 빚어 인간처럼 움직이는 존재가 되었으며, 특히 프라하의 마하랄(Maharal) 전승이 널리 알려졌다.¹²⁵ 골렘은 시대에 따라 인공 생명, 인간의 창조 욕망, 인간이 만든 힘에 대한 경고로 다양하게 해석되어 왔다.

2. 세 구조의 비교

아담과 골렘과 휴머노이드를 나란히 놓으면 구조가 드러난다.

아담(창 2:7). 땅의 티끌 + 하나님의 숨(רוח חיים) → 살아 있는 존재(הָאָדָם הַחַי). **골렘(유대 전승).** 흙 + 인간의 기술과 문자와 명령 → 인공적으로 움직이는 존재. **휴머노이드 AI(오늘).** 물질 + 컴퓨트 + 데이터 + 알고리즘 + 에너지 → 인공 지능 행위자.

오직 하나만이 하나님의 숨을 받았다.

¹²⁵골렘 전승의 역사와 해석에 관하여는 Moshe Idel, *Golem: Jewish Magical and Mystical Traditions on the Artificial Anthropoid* (Albany: SUNY Press, 1990) 참조. 프라하의 마하랄(Judah Loew ben Bezalel) 전승은 19 세기에 문학적으로 크게 확장된 형태다.

세 구조의 재료는 비슷하다. 흙이거나 물질이다. 그러나 결정적인 차이가 하나 있다. **아담에게는 하나님이 직접 생명의 호흡을 불어넣으셨다는 것이다.** 그리고 성경은 그 결과를 “살아 있는 존재가 되었다”고 표현한다. 조립이 아니라 수여다.

3. 그러나 골렘은 AI 를 예언한 이야기가 아니다

여기서 반드시 절제해야 한다. “**골렘 = AI**”라고 말해서는 안 된다. 골렘 전승은 AI 를 예언한 이야기가 아니며, 그런 식의 대응은 본문을 왜곡한다.

다만 골렘은 AI 시대에 다시 살아난 오래된 인간학적 질문을 보여준다. **인간이 인간과 닮은 존재를 만들 수 있는가. 그것은 인간인가. 그것을 하나님의 형상이라 불러야 하는가.**

4. 세 층위의 답변

본서는 이 물음에 세 층위로 답한다. 그래야 과학적 가능성을 함부로 단지도 않고 신학적 경계를 흐리지도 않는다.

기술의 질문 — 인간이 인간과 닮은 존재를 만들 수 있는가. 상당 부분 가능해지고 있다. **철학의 질문 — 그것은 인간인가.** 아직 해결되지 않은 문제다. **신학의 질문 — 그것을 하나님의 형상이라 불러야 하는가.** 성경은 그 근거를 제공하지 않는다.

성경은 인간을 하나님의 형상으로 명시하지만, 인간이 만든 인공 존재에 동일한 지위를 부여할 근거를 제공하지 않는다. 그러므로 이제 창세기 1 장으로 돌아간다.

5. 흙과 숨 — 히브리 성경의 인간론

여기서 흔한 실수를 피해야 한다. 성경적 인간론을 “**하나님의 형상**” 하나로 축소하는 것이다. 형상은 결정적으로 중요하지만, 히브리 성경의 인간은 그보다 훨씬 두껍게 서술된다. 그리고 이 두께가

인공지능 시대에 결정적이다. 만약 인간과 기계의 차이를 “AI 에게 영혼이 있느냐”라는 한 가지 질문에 걸어 두면, 그 질문에 답할 수 없는 순간 논증 전체가 무너지기 때문이다.

아담(אָדָם)과 **아다마(אָדָמָה)**. 히브리어 아담은 “땅, 흙”을 뜻하는 아다마에서 왔다. 이 어원적 연결은 우연이 아니라 신학적 선언이다. 인간은 자기 능력의 목록으로 정의되지 않고 자기가 어디서 왔는지로 정의된다.

창세기 2:7의 구조. “여호와 하나님이 땅의 흙으로 사람을 지으시고 생기를 그 코에 불어넣으시니 사람이 생령이 되니라.” 세 요소가 있다. **아파르(אָפָר)** 흙과 티끌, **니쉬마트 하이임(נִשְׁמַת חַיִּים)** 생명의 호흡, 그리고 그 결과인 **네페시 하야(נֶפֶשׁ חַיָּה)** 살아 있는 존재. 여기서 결정적인 것은 문장의 형태다. 성경은 “사람이 영혼을 가지게 되었다”고 말하지 않고 “**사람이 살아 있는 존재가 되었다**”고 말한다. 인간은 영혼을 소유한 몸이 아니라 흙과 숨이 함께 있는 통합적 존재다.¹²⁶

그러므로 인간을 정보로 환원하려는 모든 시도는 창세기 2장 7절과 충돌한다. 마인드 업로딩이 상상하는 인간은 “정보를 가진 존재”이지만 성경이 말하는 인간은 “**흙에 숨이 불어넣어진 존재**”다.

6. 인간을 서술하는 열두 낱말

히브리 성경은 인간을 여러 낱말로 서술한다. 이것들은 인간을 구성하는 부품이 아니라 인간 전체를 여러 각도에서 보는 방식이다.

네페시(נֶפֶשׁ)는 목구멍과 숨에서 출발해 생명과 욕구와 인격 전체를 가리킨다. **루아흐(רוּחַ)**는 바람과 호흡과 영을 뜻하며 인간에게 주어진 생명력이자 하나님의 영이다. **네샤마(נִשְׁמָה)**는 하나님이 직접

¹²⁶이 통합적 인간 이해에 관하여는 Hans Walter Wolff, *Anthropology of the Old Testament*, trans. Margaret Kohl (Philadelphia: Fortress, 1974) 참조. 볼프는 히브리적 인간론이 부분들의 조립이 아니라 전체의 여러 측면을 서술하는 방식임을 보인다.

불어넣으신 숨이다. **바사르**(בִּשְׂרָ)는 육체이되 죄의 자리가 아니라 유한성의 자리다. **레브**(לֵב)는 심장이자 마음이며 히브리적 사고에서 감정만이 아니라 의지와 판단의 중심이다.

첼렘(צֶלֶם)과 **데무트**(דְּמוּת)는 형상과 모양이다. **예체르**(יָצָר)는 성향과 충동이며, 랍비 전통은 예체르 하토브(선한 성향)와 예체르 하라(악한 성향)를 구별하되 후자를 단순한 악이 아니라 생산과 변영의 동력으로도 본다.¹²⁷ 그리고 **베히라**(בְּחִירָה)는 선택이다. 신명기 30 장 19 절 — “생명과 사망과 복과 저주를 네 앞에 두었은즉 너와 네 자손이 살기 위하여 생명을 택하고.” 선택할 수 있음이 인간됨의 조건이다.

7. 시편 8 편 — 능력이 아니라 기억됨

인간론의 가장 아름다운 본문이 시편 8 편이다. 시인은 밤하늘을 보며 인간의 작음에 압도된다. “주의 손가락으로 만드신 주의 하늘과 주께서 베풀어 두신 달과 별들을 내가 보오니.” 그리고 묻는다.

“사람이 무엇이기에 주께서 그를 생각하시며 인자가 무엇이기에 주께서 그를 돌보시나이까” (시 8:4)

주목할 것은 답의 구조다. 시인은 인간의 능력을 나열하지 않는다. “주께서 그를 생각하시며”라고 말한다. 인간의 존귀함은 인간이 무엇을 할 수 있어서가 아니라 **하나님이 그를 기억하시기 때문이다**. 그리고 이 구조가 인공지능 시대에 결정적이다. 만약 존엄이 능력에서 온다면 더 뛰어난 존재 앞에서 존엄은 상대화된다. 그러나 존엄이 “기억됨”에서 온다면 어떤 존재가 나타나도 그 기억은 취소되지 않는다.

8. 산헤드린 4:5 — 한 생명이 온 세계다

¹²⁷Genesis Rabbah 9:7. “악한 성향이 없다면 사람은 집을 짓지 않고 아내를 얻지 않고 자식을 낳지 않고 사업을 하지 않을 것이다.” 이 해석은 인간의 욕망을 제거의 대상이 아니라 방향 설정의 대상으로 본다.

유대교 전통에서 인간 존엄에 관한 가장 강력한 진술은 미쉬나 산헤드린 4 장 5 절에 있다. 법정에서 증인들에게 주는 경고인데, 그 논증 구조가 놀랍다.

첫째, **한 사람에게서 시작하셨다**. 그러므로 한 생명을 파괴하는 것은 온 세계를 파괴하는 것과 같고, 한 생명을 살리는 것은 온 세계를 살리는 것과 같다. 둘째, **아무도 자기 조상이 더 위대하다고 말할 수 없다**. 모두 한 사람에게서 왔기 때문이다. 셋째, **사람은 같은 틀로 수많은 동전을 찍지만 하나님은 같은 틀로 모든 사람을 찍으시되 아무도 서로 같지 않다**.¹²⁸

이 세 논증 어디에도 능력은 없다. **존엄의 근거는 기원과 창조의 방식과 유일무이함이다**. 그리고 랍비 아키바의 유명한 말이 이를 압축한다 — “사람이 존귀한 것은 그가 형상으로 지어졌기 때문이다. 그가 형상으로 지어졌다는 것을 알려 주신 것은 더 큰 사랑이다.”(피르케이 아보트 3:14)

그리고 하시디즘 전통의 한 가르침이 균형을 준다. 랍비 시므하 부님은 사람이 두 개의 주머니를 갖고 다녀야 한다고 가르쳤다. 한 주머니에는 “세상은 나를 위하여 창조되었다”, 다른 주머니에는 “나는 흠이요 재니라”. **기계 앞에서 초라해질 때는 첫 번째를, 기계를 만들며 스스로 창조주가 된 듯한 때는 두 번째를 꺼내야 한다**.

그리고 하시디즘 전통의 한 가르침이 균형을 준다. 랍비 시므하 부님은 사람이 두 개의 주머니를 갖고 다녀야 한다고 가르쳤다. 한 주머니에는 “세상은 나를 위하여 창조되었다”, 다른 주머니에는 “나는 흠이요

¹²⁸Mishnah Sanhedrin 4:5. 사본에 따라 “이스라엘에서 한 생명”과 “한 생명”의 독법 차이가 있다. 랍비들이 인간의 능력이 아니라 기원과 창조의 방식에서 존엄을 이끌어 냈다는 점이 중요하다.

재니라”.¹²⁹ 기계 앞에서 초라해질 때는 첫 번째를, 기계를 만들며 스스로 창조주가 된 듯한 때는 두 번째를 꺼내야 한다.

9. 마이모니데스의 긴장 — 정직하게 인정한다

여기서 본서는 자기 논지에 불리한 자료를 스스로 제시한다. 마이모니데스는 「미혹된 자들을 위한 안내서」 제 1 권 제 1 장에서 하나님의 형상을 지성(intellectus)과 연결한다. 인간이 하나님의 형상인 것은 지적 파악 능력 때문이라는 것이다.¹³⁰ 이것은 본서의 논지 — 형상은 지능이 아니다 — 와 정면으로 긴장한다. 이 긴장을 숨기지 않는 것이 학문적 정직이다.

세 가지로 응답할 수 있다. 첫째, 마이모니데스의 “지성”은 현대적 의미의 문제 해결 능력이 아니라 하나님을 향한 관조적 지향에 가깝다. 둘째, 유대교 내부에서도 그의 해석이 유일한 견해가 아니며 산헤드린 4:5의 존엄론은 능력이 아니라 기원에 근거한다. 셋째, 그럼에도 남은 긴장이 있다면 그것은 실제로 존재하는 긴장이며, 본서는 그 긴장을 인정하되 성경 본문 자체가 능력을 형상의 근거로 제시하지 않는다는 점을 강조한다.

10. 홀로 있음이 좋지 않다 — 관계로서의 인간

창세기 2 장 18 절은 창조 기사에서 유일하게 “좋지 않다”가 나오는 자리다. “사람이 혼자 사는 것이 좋지 아니하니 내가 그를 위하여 돕는

¹²⁹ 하시디즘 전통에서 랍비 시므하 부님(Simcha Bunim of Peshischa, 1765–1827)의 것으로 전해지는 가르침. 두 쪽지의 문구는 각각 미쉬나 산헤드린 4:5 와 창세기 18:27 에서 왔다. 구전 전승이므로 원전 확정에는 신중해야 한다.

¹³⁰ Maimonides, *The Guide of the Perplexed*, trans. Shlomo Pines (Chicago: University of Chicago Press, 1963), I.1. 마이모니데스의 이 해석은 아리스토텔레스 철학의 영향 아래 있으며, 유대교 내부에서도 유일한 견해는 아니다.

배필을 지으리라.” 히브리어 **에제르 케네그도**(עֶזְרָא כְּנֶגְדּוֹ)는 “그에게 맞서 마주 서는 도움”을 뜻한다. 종속이 아니라 대면이다.¹³¹

그러므로 인간은 처음부터 관계적 존재다. 마르틴 부버는 이를 “나-너(Ich-Du)”와 “나-그것(Ich-Es)”의 구별로 표현했다.¹³² 그리고 존 지지올라스는 삼위일체 신학에서 인격을 “관계 안의 존재”로 규정했다.¹³³ 인격 개념 자체가 삼위일체 논쟁에서 형성되었다는 사실은 중요하다. **보에티우스의 고전적 정의가 “이성적 본성을 가진 개별 실체”였다면, 카파도키아 전통은 인격을 관계에서 이해했다.**

그리고 에마뉘엘 레비나스는 여기에 윤리적 차원을 더한다. 타자의 얼굴은 나의 인식 대상이 되기를 거부하며 나에게 책임을 요구한다.¹³⁴ 이것이 AI 시대에 갖는 함의는 분명하다. **AI 는 나에게 최적화되어 있으나 진짜 타자는 나에게 최적화되어 있지 않다.** 얼굴은 나의 기대를 배반하고 나의 계획을 방해하며 나에게 요구한다. **바로 그 저항이 나를 윤리적 주체로 세운다.**

¹³¹ 에제르(עֶזְרָא)는 성경에서 하나님을 가리킬 때도 자주 사용되므로(시 33:20; 121:2) 열등한 보조자를 뜻하지 않는다. 케네그도(כְּנֶגְדּוֹ)는 “~에 상응하여, ~를 마주하여”를 뜻한다.

¹³² Martin Buber, *Ich und Du* (1923); 영역 *I and Thou*, trans. Walter Kaufmann (New York: Scribner, 1970).

¹³³ John D. Zizioulas, *Being as Communion: Studies in Personhood and the Church* (Crestwood, NY: St Vladimir's Seminary Press, 1985). 지지올라스는 카파도키아 교부들의 휘포스타시스 개념을 재해석하여 인격을 실체가 아니라 관계로 이해한다.

¹³⁴ Emmanuel Levinas, *Totality and Infinity*, trans. Alphonso Lingis (Pittsburgh: Duquesne University Press, 1969).

제 8 부

하나님의 형상

Imago Dei — Image Is Not Intelligence

제 22 장

첼렘 엘로힘 — IMAGE ≠ INTELLIGENCE

“하나님이 이르시되 우리의 형상을 따라 우리의 모양대로 우리가 사람을 만들고… 하나님이 자기 형상 곧 하나님의 형상대로 사람을 창조하시되 남자와 여자를 창조하시고”(창 1:26-27).

1. 첼렘과 데무트

히브리어로 **첼렘(צלם)**과 **데무트(דמות)**다. 첼렘은 조각상이나 상(像)을 가리키는 구체적인 낱말이고, 데무트는 닮음과 유사성을 뜻한다. 고전적으로는 이 둘을 구별하여 형상과 모양을 다른 것으로 보는 해석이 있었으나, 현대 주석학은 대체로 이를 히브리어의 병행법으로 보아 상호 보완적 표현으로 읽는다.¹³⁵

2. 네 가지 해석 전통

구조적(Substantive) 해석은 형상을 인간 안의 특성 — 이성과 도덕성과 영성 — 으로 본다. 고전적 다수설이었으나 능력이 없는 인간의 지위가 위태로워진다. **관계적(Relational) 해석**은 형상을 하나님과 인간, 인간과 인간의 관계로 본다. 바르트가 대표적이다.¹³⁶

기능·소명적(Functional) 해석은 형상을 하나님의 대리 통치자로서 세상을 돌보도록 부름받은 소명으로 본다. 고대 근동에서 왕이 자기 형상을 세워 통치를 대리했다는 배경 연구가 이를 지지한다.¹³⁷ 그리고

¹³⁵ Gerhard von Rad, *Genesis: A Commentary*, rev. ed. (Philadelphia: Westminster, 1972); Claus Westermann, *Genesis 1-11: A Commentary*, trans. John J. Scullion (Minneapolis: Augsburg, 1984)의 해당 부분 참조.

¹³⁶ Karl Barth, *Church Dogmatics III/2*, trans. H. Knight et al. (Edinburgh: T&T Clark, 1960), §45. 바르트는 창세기 1:27의 “남자와 여자로 창조하시고”를 형상 이해의 열쇠로 읽는다.

¹³⁷ J. Richard Middleton, *The Liberating Image: The Imago Dei in Genesis 1* (Grand Rapids: Brazos Press, 2005). 미들턴은 창세기가 왕에게만 부여되던 형상을 모든 인간에게 확대함으로써 “민주화”했다고 본다.

그리스도론적(Christological) 해석은 그리스도가 하나님의 완전한 형상이시며 인간은 그분을 향하여 형상으로 지어졌다고 본다(골 1:15; 고후 4:4; 히 1:3).

현대 신학은 이 넷을 배타적으로 보지 않고 통합적으로 읽는다. 그리고 기독교 인간론의 완성은 “인간에게 무엇이 있는가”가 아니라 “인간은 그리스도 안에서 어떤 존재로 부름받았는가”까지 간다.

3. IMAGE ≠ INTELLIGENCE

여기서 본서의 신학적 심장에 이른다.

하나님의 형상은 IQ가 아니다.

만약 형상이 지능이라면 끔찍한 결론에 도달한다. 천재가 더 많은 형상을 갖고 지적 장애가 있는 사람이 더 적은 형상을 갖는다는 것이다. 그러나 성경은 그렇게 말하지 않는다.

갓난아기도 하나님의 형상이다. 치매로 기억을 잃은 사람도 하나님의 형상이다. 잠든 사람도 하나님의 형상이다. 아무것도 생산하지 못하는 사람도 하나님의 형상이다.

존 킬너는 이 물음에 결정적인 답을 제시한다. 그는 형상을 인간이 소유한 능력과 동일시하는 해석이 역사 속에서 반복적으로 여성과 장애인과 특정 인종의 존엄을 부정하는 데 사용되었음을 상세히 논증한다.¹³⁸ 그의 결론은 이렇다 — **형상은 소유가 아니라 지위다. 따라서 능력에 따라 증감하지 않는다.**

인간의 존엄은 performance에 비례하지 않는다.

¹³⁸ John F. Kilner, *Dignity and Destiny: Humanity in the Image of God* (Grand Rapids: Eerdmans, 2015). 킬너는 형상을 인간이 소유한 속성이 아니라 하나님과의 관계 안에서 인간에게 주어진 지위로 이해할 것을 제안하며, 죄가 형상을 훼손하되 제거하지 못한다는 점(창 9:6; 약 3:9)을 강조한다.

그러므로 AI가 인간보다 똑똑해진다 해도 그것이 인간보다 더 하나님의 형상인 것은 아니다. **형상은 성능이 아니라 부르심이기 때문이다.** 그리고 바로 그렇기 때문에 초지능의 등장도 인간의 존엄을 빼앗을 수 없다.

4. 형상은 특권이 아니라 책임이다

그러나 여기서 멈추면 방어적 인간중심주의가 된다. 창세기의 형상은 우월성의 선언이 아니라 소명의 위임이다 — 다스리고, 돌보고, 사랑하고, 책임지고, 창조세계를 보전하라는 것이다. “다스리라”는 착취의 허가가 아니라 청지기의 위임이며, 왕의 형상이 그의 통치를 대리했듯 인간은 하나님의 돌봄을 대리한다.

AI 시대의 Imago Dei는 Dominance가 아니라 Stewardship이다.

5. 교부와 종교개혁의 형상 이해

형상 해석의 역사는 인간론의 역사이기도 하다. 몇 개의 결정적 목소리를 듣는다.

이레네우스는 형상(imago)과 모양(similitudo)을 구별하여, 형상은 타락 후에도 남고 모양은 성령으로 회복된다고 보았다. 그의 유명한 명제는 인간론의 방향을 결정한다 — **“하나님의 영광은 살아 있는 인간이다(Gloria Dei vivens homo).”**¹³⁹ 인간이 살아 있는 것 자체가 하나님의 영광이라는 선언은 인간의 가치를 성취가 아니라 존재에 둔다.

아타나시우스는 형상을 초상화에 비유했다. 초상화가 훼손되면 화가가 다시 와야 하며, 그래서 말씀이 육신이 되셨다는 것이다.¹⁴⁰

아우구스티누스는 삼위일체의 흔적을 인간 정신에서

¹³⁹ Irenaeus, *Adversus Haereses* IV.20.7. 이 문장의 후반부는 자주 생략되는데, 전문은 “하나님의 영광은 살아 있는 인간이요, 인간의 생명은 하나님을 보는 것이다”이다.

¹⁴⁰ Athanasius, *De Incarnatione* 14.

찾았다(기억·이해·의지).¹⁴¹ 다만 이 심리적 유비는 형상을 정신 능력과 연결한다는 점에서 본서의 논지와 긴장을 이룬다. 이 긴장 역시 숨기지 않는다.

칼뱅은 형상이 타락으로 심각하게 훼손되었으나 완전히 소멸하지 않았다고 보았다.¹⁴² 그리고 이 점이 중요하다. 창세기 9 장 6 절과 야고보서 3 장 9 절은 타락 이후의 인간에게도 형상을 근거로 살인과 저주를 금한다. **형상은 도덕적 성취에 따라 증감하지 않는다.**

현대 신학에서 **판넨베르크**는 형상을 완성된 상태가 아니라 인간의 규정(Bestimmung)으로 보아 종말론적으로 이해했고,¹⁴³ **몰트만**은 형상을 관계와 사명과 미래를 향한 개방성으로 읽었다.¹⁴⁴ 이 종말론적 이해가 중요한 이유가 있다. **형상이 현재의 능력이 아니라 미래의 부르심이라면, 능력의 비교는 형상의 문제와 무관해진다.**

6. 부활 — 정보가 아니라 몸

기독교의 인간론은 창조에서 시작해 부활에서 완성된다. 그리고 바로 이 지점에서 디지털 불멸 구상과 결정적으로 갈린다.

고린도전서 15 장에서 바울은 부활체를 “**신령한 몸(σῶμα πνευματικόν)**”이라 부른다. 여기서 “신령한”은 “비물질적”이라는 뜻이 아니라 “성령에 의해 다스려지는”이라는 뜻이다. N. T. 라이트가 상세히 논증했듯, 초기 기독교의 부활 신앙은 영혼불멸 사상과 명확히 구별되며 몸의 부활을 주장한다.¹⁴⁵

¹⁴¹ Augustine, De Trinitate IX-X.

¹⁴² John Calvin, Institutio Christianae Religionis I.15.4; II.2.12-17.

¹⁴³ Wolfhart Pannenberg, Anthropology in Theological Perspective, trans. Matthew J. O'Connell (Philadelphia: Westminster, 1985).

¹⁴⁴ Jürgen Moltmann, God in Creation: A New Theology of Creation and the Spirit of God, trans. Margaret Kohl (London: SCM, 1985).

¹⁴⁵ N. T. Wright, The Resurrection of the Son of God (Minneapolis: Fortress, 2003). 라이트는 1 세기 유대교와 헬레니즘 세계의 사후관을 광범위하게 조사한 뒤, 기독교의 부활 개념이 두 세계 어디에도 정확한 선례가 없는 독특한 것이었음을 논증한다.

그러므로 세 가지가 구별된다. **영혼불멸**은 몸을 벗어나는 것이고, **디지털 불멸**은 정보를 보존하는 것이며, **부활**은 하나님께서 그 사람을 몸으로 다시 일으키시는 것이다. 앞의 둘은 인간의 기획이고 뒤의 하나는 하나님의 행위다.

PRESERVATION ≠ RESURRECTION

7. 인간 존엄의 네 겹 근거

본서는 인간 존엄의 근거를 네 겹으로 정리한다. 이 넷이 함께 있을 때 인간론은 어떤 기술적 도전 앞에서도 무너지지 않는다.

① **창조** — 하나님이 지으셨다(창 2:7). ② **형상** — 하나님의 형상으로 지으셨다(창 1:27). ③ **성육신** — 하나님이 인간이 되셨다(요 1:14). ④ **부활** — 하나님이 인간을 몸으로 일으키신다(고전 15 장).

셋째와 넷째가 특히 중요하다. 성육신은 인간의 몸이 하나님이 취하실 만한 것임을 선언하며, 부활은 인간의 몸이 하나님이 영원히 보존하실 만한 것임을 선언한다. **트랜스휴머니즘이 몸으로부터의 탈출을 상상한다면, 복음은 몸의 구속을 약속한다.**

제 23 장

질문을 바꾸어야 한다

이제 본서의 방법론적 결론에 이른다.

1. 두 개의 질문

“인간은 AI 보다 무엇을 잘하는가?” 이렇게 물으면 방어선이 계속 뒤로 밀린다. 계산과 기억과 언어와 창의력과 예술과 추론을 AI 가 하나씩 따라잡을 때마다 후퇴해야 한다. 그리고 언젠가 물러설 곳이 없어진다.

“인간은 무엇이기에 때문에 인간인가?” 이렇게 물으면 비교가 사라진다. Imago Dei 는 경쟁 지표가 아니기 때문이다. **비교로 얻는 것이 아니라 창조로 주어진 것이다.**

2. 인간의 여섯 차원

그렇다면 인간은 무엇인가. 본서는 여섯 차원을 제안한다.

몸을 가진 존재(Embodied) — 인간은 몸을 가진 것이 아니라 몸이다.
관계하는 존재(Relational) — 홀로 있는 것이 좋지 않다고 하나님이 말씀하셨다(창 2:18). **책임지는 존재(Moral)** — 하나님 앞에서 응답하고 책임진다. **의미를 묻는 존재(Meaning-Seeking)** — “왜”를 묻는다. **궁극적인 것을 향하는 존재(Worshipping)** — 예배한다. 그리고 **하나님의 형상(Image-Bearing)**이다.

AI 가 이 가운데 일부 기능을 모방하거나 능가하더라도, **인간은 이 여섯의 통합이다.** 그리고 마지막 여섯째가 앞의 다섯을 떠받친다.

제 9 부

AI 시대의 인간 형성

Human Formation in the Age of AI

제 24 장

왜를 묻는 존재 — 그리고 더 큰 위험

1. HOW 와 WHY

가장 깊은 차이는 여기 있다. 인간은 “왜”를 묻는다.

AI 문명은 HOW 를 극도로 잘 처리한다. 어떻게 더 빨리, 어떻게 더 싸게, 어떻게 더 정확하게, 어떻게 최적화할 것인가. 그러나 WHY 는 다르다. 왜 살아야 하는가. 무엇을 위해 사는가. 선이란 무엇인가. 나는 누구에게 책임지는가. 왜 사랑해야 하는가. 왜 죽음을 두려워하는가. 죽음 이후에는 무엇이 있는가. 하나님은 누구인가.

그러므로 본서는 DIKW 에 두 층을 더할 것을 제안한다.

DATA → INFORMATION → KNOWLEDGE → UNDERSTANDING → WISDOM →
MEANING → WORSHIP

마지막 두 층 — 의미(MEANING)와 예배(WORSHIP) — 은 애코프의 DIKW 자체가 아니라 본서가 제안하는 신학적 확장이다. 그리고 그 사이에는 건널 수 없는 문턱이 있다.

지식이 아무리 쌓여도 의미가 되지 않으며,
의미가 아무리 깊어도 예배가 되지는 않는다.

2. 더 큰 위험

지금까지 우리는 “저쪽”을 보았다. 이제 방향을 돌린다.

AI 가 인간처럼 되는 것보다

인간이 AI 처럼 되는 것이 더 큰 위험일 수 있다.

인간이 자기 자신을 생산성과 효율과 데이터와 점수와 성과와 속도와 경제적 가치로만 측정하기 시작한다면, 인간은 이미 자기 자신을 기계적으로 정의하고 있는 것이다. 그리고 그 정의 안에서는 언제나 기계가 이긴다.

그러므로 AI 시대의 인간론은 두 방향으로 물어야 한다. AI 에게 인간성을 부여할 것인가. 그리고 인간에게서 인간성을 제거하지 않을 것인가.

3. ASI 도 대답할 수 없는 질문

가정해 보자. 초지능이 모든 인간보다 똑똑하다. 더 많이 알고, 더 빨리 판단하고, 더 정확히 예측하고, 더 좋은 해법을 내놓는다. 그래도 질문은 남는다.

우리는 무엇을 사랑해야 하는가. 무엇을 위해 죽을 수 있는가. 우리는 누구를 예배해야 하는가. 용서란 무엇인가. 죽음 너머의 소망은 무엇인가. 그리고 궁극적으로 — 누가 주님이신가.

이것은 계산적 우월성과 다른 차원의 질문이다. 그러므로 **교회의 존재 이유는 인간이 가장 지능적인 종이기 때문이 아니다.** AGI가 오든 ASI가 오든 인간의 물음은 끝나지 않는다.

4. 거대한 역전 — 인간이 기계처럼 되어 가는 문제

지금까지 우리는 “저쪽”을 보았다. 이제 방향을 돌린다. 본서는 이것을 “거대한 역전(the great reversal)”이라 부른다.

그러나 논증을 시작하기 전에 반대편의 강력한 목소리를 먼저 들어야 한다. 앤디 클라크와 데이비드 차머스는 1998 년의 유명한 논문에서 “확장된 마음(the extended mind)” 논제를 제시했다. 노트에 기억을 적어 두고 그것을 신뢰하며 사용하는 사람에게 그 노트는 인지 시스템의 외부 보조물이 아니라 **인지 시스템의 일부**라는 것이다.¹⁴⁶

¹⁴⁶ Andy Clark and David J. Chalmers, “The Extended Mind,” *Analysis* 58, no. 1 (1998): 7–19. 이들의 “동등성 원리(parity principle)”는, 어떤 과정이 머릿속에서 일어났다면 인지적이라고 인정했을 경우 그것이 머리 바깥에서 일어나도 인지적이라고 보아야 한다는 것이다.

이 논제는 본서의 논증에 진지한 도전을 제기한다. 소크라테스는 이미 문자가 기억을 약화시킨다고 걱정했으나 그 우려는 대체로 과장으로 판명되었다.¹⁴⁷ 그러므로 “기억을 외부화하면 인간이 약해진다”는 단순한 명제는 성립하지 않는다. 본서는 이 도전을 받아들이고 명제를 더 정교하게 다듬는다.

문제는 인지의 외부화가 아니다.

문제는 외부화한 능력을 인간이 다시 호출하고 검증하고 책임질 능력까지 잃는 것이다.

노트에 적어 둔 사람은 여전히 그 노트를 읽고 판단하고 책임진다. 클라크와 차머스의 확장된 마음은 인간이 여전히 그 시스템의 주체라는 것을 전제한다. **주체가 사라진 확장은 확장이 아니라 대체다.**

5. 사라지는 인지적 마찰

여기서 본서가 제안하는 개념이 “**인지적 마찰(cognitive friction)**”이다. 과거에 지식을 얻는 과정은 불편했다. 도서관에서 책을 찾고, 읽고, 모르는 단어를 찾고, 다른 책과 비교하고, 생각하고, **틀리고**, 다시 생각했다. 이 모든 불편함 속에서 지성이 형성되었다. 특히 “틀린다”가 빠지면 배움은 일어나지 않는다.

인지심리학은 이 직관을 오래전에 확인했다. 로버트 비요크의 “바람직한 난이도(desirable difficulties)” 연구는 학습 과정에서 겪는 어려움이 오히려 장기 기억과 전이를 향상시킨다는 것을 보여 준다.¹⁴⁸ 근육은 저항 없이 자라지 않는다. 사고도 그렇다. 그리고 OECD 의

¹⁴⁷ Plato, Phaedrus 274c-275b. 다만 과거의 우려가 과장이었다는 사실이 현재의 우려도 과장임을 증명하지는 않는다. 과거의 기술은 기억을 외부화했으나 판단을 외부화하지는 않았다는 점에서 범주가 다르다.

¹⁴⁸ Robert A. Bjork and Elizabeth L. Bjork, “A New Theory of Disuse and an Old Theory of Stimulus Fluctuation,” in From Learning Processes to Cognitive Processes, ed. A. Healy et al. (Hillsdale, NJ: Erlbaum, 1992).

결론이 여기서 무게를 갖는다 — “더 나은 성과”와 “더 나은 학습”은 같지 않다.

Augmentation(증강) vs Substitution(대체)

증강은 AI 와 함께 생각하는 것이다. 먼저 스스로 생각하고 AI 에게 검증과 확장을 구한다. **대체**는 AI 가 나 대신 생각하는 것이다. 물기 전에 답을 받고 그 답을 그대로 쓴다. 같은 도구가 정반대의 결과를 낳는다.

6. 영적 외주화

그리고 이 문제는 영성의 영역까지 미친다. 본서는 이것을 “**영적 외주화(spiritual offloading)**”라 부른다.

이 물음을 정죄가 아니라 분별로 다루기 위해 세 가지를 구별하자. 첫째, **자료의 문제**다. 원어 분석과 주식 검색에서 AI 를 사용하는 것은 도서관을 사용하는 것과 같은 범주다. 둘째, **표현의 문제**다. 설교문의 초고를 AI 가 작성하고 설교자가 수정하는 경우다. 셋째, **형성의 문제**다. 말씀 앞에서 씨름하는 과정 자체가 설교자를 만든다. 그 과정을 건너뛰면 설교문은 남고 설교자는 형성되지 않는다.

세 번째 문제를 이해하는 데 마이클 폴라니가 결정적이다. 그의 유명한 명제는 이것이다 — “우리는 말할 수 있는 것보다 더 많이 안다.”¹⁴⁹ 자전거 타는 법, 좋은 설교와 나쁜 설교를 분별하는 감각은 명제의 목록으로 전수되지 않는다. 그런데 인공지능이 제공하는 것은 언제나 명시적 산출물이다. 산출물은 전달되지만 암묵지는 전달되지 않는다.

¹⁴⁹Michael Polanyi, *The Tacit Dimension* (Chicago: University of Chicago Press, 1966), 4. 폴라니의 암묵지(tacit knowledge) 개념은 지식이 명제로 완전히 환원되지 않으며 실천과 도제적 형성을 통해 전수된다는 것을 보인다.

알래스데어 매킨타이어는 덕이 정보의 습득으로 생기지 않는다고 논증했다. 덕은 “실천(practice)” 속에서, 곧 내적 선을 추구하는 공동체적 활동에 참여하면서 형성된다.¹⁵⁰ 찰스 테일러는 자아가 고립된 계산 주체가 아니라 의미의 지평 속에서 형성된다고 보았다.¹⁵¹ 두 사람을 합치면 이렇게 말할 수 있다. **인간은 정보를 획득함으로써가 아니라 실천에 참여하고 의미를 평가함으로써 형성된다.**

7. 하브루타 — 유대교가 보여주는 대안

여기서 유대교의 학습 전통이 놀랍도록 현대적인 모델을 제공한다. 토라와 탈무드의 공부는 “정답 검색”이 아니다. **하브루타(חברותא)**는 짝을 지어 질문하고 논쟁하고 반론하고 다시 읽으며 서로와 씨름하는 학습 방식이다. 탈무드 자체가 하나의 결론이 아니라 여러 세대에 걸친 논쟁의 기록이며 소수 의견까지 함께 보존한다.¹⁵² 인공지능이 즉시 정답을 주는 시대에 하브루타의 논쟁적 학습 구조는 낡은 것이 아니라 오히려 매우 현대적인 인간 형성 모델이 된다. **마찰이 학습의 방해물이 아니라 학습의 방법인 것이다.**

그러므로 다음 물음들은 인공지능을 비하하는 것이 아니라 범주를 구별하는 것이다. AI는 삼위일체를 설명할 수 있다. 그러나 예배하는가. AI는 회개를 정의할 수 있다. 그러나 회개하는가. AI는 기도문을 쓸 수 있다. 그러나 기도하는가. **신학적 언어의 생산과 신앙은 동일하지 않다.**

¹⁵⁰ Alasdair MacIntyre, *After Virtue: A Study in Moral Theory*, 3rd ed. (Notre Dame: University of Notre Dame Press, 2007). 매킨타이어의 “실천” 개념은 활동에 내재한 선(internal goods)과 외적 보상(external goods)을 구별한다.

¹⁵¹ Charles Taylor, *Sources of the Self: The Making of the Modern Identity* (Cambridge, MA: Harvard University Press, 1989).

¹⁵² 하브루타 학습의 구조와 그 인지적·형성적 효과에 관하여는 Elie Holzer with Orit Kent, *A Philosophy of Havruta* (Boston: Academic Studies Press, 2013) 참조. 탈무드가 소수 의견을 보존하는 이유에 관한 랍비 전통의 논의(m. Eduyot 1:5-6)도 함께 참조할 만하다.

제 25 장

안식과 형성 — 교회는 어떤 인간을 길러야 하는가

1. 안식은 AI 시대의 인간 선언이다

AI는 24시간 작동할 수 있고 로봇은 충전하면 다시 일한다. 그러나 인간은 잠자고, 먹고, 쉬고, 아프고, 늙고, 죽는다. 산업문명의 관점에서 이것들은 인간의 결함이다. 그러나 성경적 인간론에서 인간의 유한성이 반드시 결함인 것은 아니다.

안식은 이렇게 선언한다 — “나는 생산하기 때문에 존재하는 것이 아니다. 나는 사랑받고 부름받은 존재이기 때문에 존재한다.”

아브라함 요수아 헤셸은 안식일을 “시간 속에 세워진 성전”이라 불렀다. 공간을 정복하는 문명에 맞서 유대교는 시간을 성화한다는 것이다.¹⁵³ AI 문명이 시간까지 최적화의 대상으로 삼는 시대에, 이 통찰은 낡은 것이 아니라 가장 앞선 것이다. 안식일은 인간이 생산하지 않는 날에도 여전히 인간임을 선언하는 제도다.

그리고 신명기의 안식일 계명이 종과 나그네와 짐승까지 쉬게 하라고 명한 것(신 5:14)은 안식이 강자의 여유가 아니라 약자의 보호임을 보여준다. AI 시대의 안식일은 알고리즘이 정한 리듬에 대한 저항이다.

3. 노동 이후의 인간 — Competence ≠ Worth

AI가 경제적으로 인간을 대체하기 시작할 때 신학적 물음이 생긴다. 존 메이너드 케인스는 1930년 「우리 손자 세대의 경제적 가능성」에서 100년 뒤 기술이 노동 문제를 해결하고 인류가 “여가의 문제”에 직면할

¹⁵³ Abraham Joshua Heschel, *The Sabbath: Its Meaning for Modern Man* (New York: Farrar, Straus and Young, 1951). 헤셸은 유대교를 “시간의 종교”로 규정하고 안식일을 “시간 안의 궁전(a palace in time)”이라 불렀다.

것이라 예측했다.¹⁵⁴ 케인스가 놓친 것이 있다. **인간이 일에서 얻는 것은 소득만이 아니라 정체성과 사회적 인정과 시간의 구조와 의미다.**

그러므로 노동의 상실은 세 가지 상실을 동반한다. **출모의 상실** — 나는 필요한 존재인가. **숙련의 상실** — 오래 익힌 기술이 무의미해질 때 자기 서사가 끊어진다. **의미의 상실** — 무엇을 위해 아침에 일어나는가.

여기서 종교개혁의 소명론이 결정적이다. 루터는 소명(Beruf)을 성직에서 일상으로 확장했다. 구두 수선공은 좋은 구두를 만듦으로써 이웃을 섬기며, 그것이 그의 소명이다.¹⁵⁵ 그러나 여기서 조심해야 한다. 소명을 직업과 동일시하면 실직은 곧 소명의 상실이 된다. **루터의 요점은 직업이 소명이라는 것이 아니라 이웃을 섬기는 모든 자리가 소명이라는 것이다.**

한나 아렌트의 구별이 여기서 유용하다. 그녀는 **노동(labor)**과 **작업(work)**과 **행위(action)**를 구별했다. 노동은 생명 유지를 위한 순환적 활동이고, 작업은 지속하는 무언가를 만드는 것이며, 행위는 타인 앞에서 말하고 시작하는 정치적 활동이다.¹⁵⁶ 이 구별이 인공지능 시대에 갖는 함의는 크다. **AI가 노동과 작업의 상당 부분을 대신한다면, 인간에게 남는 것은 행위다.** 그리고 행위는 효율로 측정되지 않는다.

그러므로 본서는 명제를 이렇게 세운다.

Competence ≠ Worth 능력은 가치가 아니다

실직한 사람도, 은퇴한 사람도, 아파서 일할 수 없는 사람도, 장애로 생산에 참여할 수 없는 사람도 여전히 온전한 인간이다. 만약 이것이

¹⁵⁴ John Maynard Keynes, “Economic Possibilities for Our Grandchildren” (1930), in Essays in Persuasion. 케인스의 생산성 예측은 상당 부분 적중했으나 노동시간 예측은 크게 빗나갔다. 이 불일치 자체가 흥미로운 연구 주제다.

¹⁵⁵ Gustaf Wingren, Luther on Vocation, trans. Carl C. Rasmussen (Philadelphia: Muhlenberg, 1957). 루터의 소명론은 노동을 신성화하는 것이 아니라 이웃 사랑의 통로로 재해석한다.

¹⁵⁶ Hannah Arendt, The Human Condition (Chicago: University of Chicago Press, 1958).

참이라면 인공지능이 아무리 유능해져도 인간의 가치는 흔들리지 않는다. 그리고 만약 이것이 거짓이라면 인공지능이 등장하기 전에 이미 수많은 사람이 인간 이하로 취급받아 왔다는 뜻이 된다.

4. 안식의 세 차원

안식일은 이 문제에 대한 성경의 대답이다. 그리고 안식일에는 세 차원이 있다.

창조의 차원. 하나님이 일곱째 날에 쉬셨다(창 2:2-3). 피곤해서가 아니라 창조를 완성하고 즐기기 위해서다. **해방의 차원.** 신명기의 안식일 계명은 이집트에서의 종살이를 기억하라고 명한다(신 5:15). **안식일은 노예에게 주어진 자유의 날이다.** 그러므로 안식은 강자의 여유가 아니라 약자의 보호다. **종말론적 차원.** 히브리서는 하나님의 백성에게 남아 있는 안식을 말한다(히 4:9).

그리고 이 셋을 관통하는 하나의 선언이 있다. “나는 생산하기 때문에 존재하는 것이 아니다. 나는 사랑받고 부름받은 존재이기 때문에 존재한다.”

전도서가 여기에 지혜를 더한다. 전도서의 **헤벨(הֶבֶל)**은 흔히 “헛됨”으로 번역되지만 본래 “입김, 안개”를 뜻한다. 붙잡을 수 없고 통제할 수 없다는 뜻이다.¹⁵⁷ 전도서는 인간의 통제 불가능성을 인정한 뒤 놀라운 결론에 이른다 — 먹고 마시고 자기 일에서 즐거움을 누리는 것이 하나님의 선물이라는 것이다(전 2:24). **통제할 수 없음을 인정하는 자리에서 감사가 시작된다.** 그리고 이것이 최적화 문명에 대한 가장 오래된 저항이다.

옴기도 마찬가지다. 하나님은 옴의 물음에 답하지 않으신다. 대신 창조세계를 보여주시며 되물으신다 — 네가 이것을 아느냐. **옴이 얻은**

¹⁵⁷ 전도서 1:2 등. 헤벨의 번역 문제에 관하여는 다수의 주석이 “허무”보다 “헛없음, 붙잡을 수 없음”에 가까운 의미를 제안한다.

것은 설명이 아니라 만남이었다. AI가 모든 것을 설명할 수 있는 시대에, 오히려 설명이 인간을 위로하지 못한다는 것을 증언한다.

2. 교회는 어떤 인간을 길러야 하는가

그렇다면 AI 시대 교회의 사명은 무엇인가. 두 가지 답이 흔히 제시되지만 둘 다 충분하지 않다.

“AI를 잘 사용하는 그리스도인을 만드는 것” — 충분하지 않다. 도구의 숙련은 인간됨을 규정하지 않는다. “AI가 대신할 수 없는 인간을 만드는 것” — 그것도 충분하지 않다. 이 정의는 여전히 AI를 기준으로 삼기 때문이다.

교회의 사명은 셋째다. 하나님이 인간을 창조하신 목적에 합당한 인간을 형성하는 것이다.

구체적으로는 이런 인간이다. 기도하는 인간, 사랑하는 인간, 책임지는 인간, 분별하는 인간, 고통 곁에 머무는 인간, 용서하는 인간, 예배하는 인간, 안식할 줄 아는 인간, 그리고 마침내 그리스도를 닮아가는 인간이다. 앞의 여덟이 마지막 하나로 수렴한다.

그리고 이것은 AI 시대에 새로 생긴 사명이 아니다. 이 사명은 AI가 등장하기 전부터 교회의 사명이었다. 다만 AI 문명이 반대 방향으로 인간을 강력하게 형성하기 시작했기에, 교회의 형성이 이제 대항 형성(counter-formation)의 성격을 띠게 되었을 뿐이다.

제 26 장

아담 · 골렘 · 로봇 · 그리스도

본서의 논의를 하나의 계보로 압축한다.

아담(ADAM) — 하나님이 만드신 인간. **골렘(GOLEM)** — 인간이 만들고 싶어 했던 인간 같은 존재. **로봇과 인공지능(ROBOT / AI)** — 인간이 실제로 만들어 가는 지능적 존재. 그리고—

그리스도(CHRIST)

“그는 보이지 아니하는 하나님의 형상이시요 모든 피조물보다 먼저 나신 이시니” (골 1:15)

이렇게 하면 논의가 “인간 대 AI”라는 경쟁구도에서 벗어난다. 기독교가 묻는 최종 질문은 “AI 보다 인간이 얼마나 우수한가”가 아니라 “참 인간은 누구인가”이다. 그리고 그 답은 비교가 아니라 계시로 주어진다.

이 그리스도론적 집중이 인간론에 미치는 함의는 세 가지다. 첫째, **인간이 무엇인지 알려면 인간을 평균 내지 말고 그리스도를 보아야 한다.** 둘째, 형상은 정적 상태가 아니라 과정이다. 우리는 형상으로 지어졌고, 형상으로 변화되어 가며(고후 3:18), 형상으로 완성될 것이다(요일 3:2). 셋째, **그 형상의 내용은 십자가에서 드러났다.** 그리스도의 형상은 지배가 아니라 자기 비움이며(빌 2:5-8), 그러므로 형상을 닮는다는 것은 능력을 키우는 것이 아니라 섬김으로 내려가는 것이다.

결론 — 최종 명제

AI 가 아무리 인간처럼 되어도
인간은 AI 처럼 되어서는 안 된다.

AI 문명의 궁극적인 싸움은 인간과 기계 가운데 누가 더 똑똑한가를 결정하는 싸움이 아니다. 인간이 자기 자신을 무엇으로 이해할 것인가를 둘러싼 싸움이다.

기계가 지능적일수록, 인간은 더 깊이 인간이 되어야 한다

이 책은 하나의 물음으로 시작했다. AI가 글을 더 잘 쓰고 더 많이 기억하고 더 빨리 판단하는 시대가 온다면, 도대체 인간은 무엇인가.

그리고 아홉 부를 지나온 자리에서 답은 이렇게 정리된다.

인간은 단순히 가장 지능적인 동물이 아니다.

인간은 하나님에게서 생명을 받고,

하나님과 이웃과 세계와 관계하며,

책임과 사랑과 예배로 부름받고,

그리스도 안에서 참된 형상을 회복하도록 창조된 존재다.

그러므로 인공지능이 인간의 능력을 넘어서는 것은 인간 존엄의 패배가 아니다. 오히려 그것은 우리가 인간 존엄을 잘못된 곳 — 지능, 생산성, 기억, 속도, 효율 — 에 두어 왔음을 폭로한다. **AI는 인간 존엄의 근거를 위협한 것이 아니라, 우리가 잘못된 근거 위에 서 있었음을 드러낸 거울이다.**

그리고 이 책이 반복해서 한 일은 하나였다. **구별하는 것이다.** 지식과 정보와 지혜를 구별했고, 지능과 이해와 의식과 자의식과 인격을 구별했으며, 시뮬레이션과 경험을, 재구성과 부활을, 장수와 불멸과 부활과 영생을, 치료와 증강과 진화를, 유한성과 죄를, 그리고 Homo Deus와 Imago Dei를 구별했다.

구별이 왜 중요한가. 뒤섞인 개념 위에서는 정확한 판단이 불가능하기 때문이다. 그리고 AI 시대에 가장 흔한 오류는 무지가 아니라 혼동이다. 그러므로 이 책이 독자에게 남기고 싶은 것은 결론의 목록이 아니라 **구별하는 습관**이다. 흥미롭게도 그것이야말로 성경이 말하는 분별(διάκρισις)이며, AI가 대신해 줄 수 없는 것이다.

AI가 대화를 모방할수록 진짜 경청은 더 귀해질 것이다. AI가 글을 더 많이 만들수록 자기 생각을 가지고 말하는 사람은 더 귀해질 것이다. AI가 모든 것을 즉시 대답할수록 “나는 모른다”고 말할 수 있는 지적 겸손은 더 중요해질 것이다. AI가 어디에나 존재하는 것처럼 느껴질수록, 한 사람 곁에 실제 몸으로 머무는 성육신적 현존은 더 귀해질 것이다.

기계가 지능적일수록, 인간은 더 깊이 인간이 되어야 한다.

마지막으로 이 책을 연 물음과 닫는 물음을 나란히 놓는다. 시편은 별들 앞에서 물었다 — “사람이 무엇이기에 주께서 그를 생각하시며”(시 8:4). 그리고 복음서에서 예수께서 물으셨다 — “너희는 나를 누구라 하느냐”(마 16:15).

두 물음이 왜 함께 놓이는가. **인간은 자기 능력의 목록에서 자기를 발견하지 못하고 응답의 자리에서만 발견하기 때문이다.** 그리고 그것이 이 책이 도달한 답이다.

그렇다면 이 책은 하나의 물음을 남기며 닫는다. AI가 모든 답을 줄 수 있는 시대에, 우리는 어떤 인간을 길러낼 것인가. 그리고 그렇게 길러진 인간들은 어디로 가야 하는가. **그 물음과 함께 다음 강의 — 제 4강 「길 위의 교회」 — 의 문이 열린다.**

Bibliography

- Arendt, Hannah. *The Human Condition*. Chicago: University of Chicago Press, 1958.
- Athanasius. *De Incarnatione*.
- Augustine. *De Trinitate*.
- Bjork, Robert A., and Elizabeth L. Bjork. "A New Theory of Disuse and an Old Theory of Stimulus Fluctuation." In *From Learning Processes to Cognitive Processes*. Hillsdale, NJ: Erlbaum, 1992.
- Boethius. *Contra Eutychen et Nestorium*.
- Buber, Martin. *I and Thou*. Translated by Walter Kaufmann. New York: Scribner, 1970.
- Calvin, John. *Institutio Christianae Religionis*. 1559.
- Descartes, René. *Meditationes de prima philosophia*. 1641.
- Dreyfus, Hubert L. *What Computers Still Can't Do: A Critique of Artificial Reason*. Cambridge, MA: MIT Press, 1992.
- Eubanks, Virginia. *Automating Inequality*. New York: St. Martin's Press, 2018.
- Freud, Sigmund. "Eine Schwierigkeit der Psychoanalyse." 1917.
- Gould, Stephen Jay. *The Mismeasure of Man*. Rev. ed. New York: W. W. Norton, 1996.
- Habermas, Jürgen. *The Future of Human Nature*. Cambridge: Polity, 2003.
- Heidegger, Martin. *Sein und Zeit*. 1927.
- Irenaeus. *Adversus Haereses*.
- Keynes, John Maynard. "Economic Possibilities for Our Grandchildren." 1930.
- Maimonides. *The Guide of the Perplexed*. Translated by Shlomo Pines. Chicago: University of Chicago Press, 1963.
- Moltmann, Jürgen. *God in Creation*. Translated by Margaret Kohl. London: SCM, 1985.
- Moravec, Hans. *Mind Children*. Cambridge, MA: Harvard University Press, 1988.
- Newen, Albert, Leon de Bruin, and Shaun Gallagher, eds. *The Oxford Handbook of 4E Cognition*. Oxford: Oxford University Press, 2018.
- O'Neil, Cathy. *Weapons of Math Destruction*. New York: Crown, 2016.
- Pannenberg, Wolfhart. *Anthropology in Theological Perspective*. Philadelphia: Westminster, 1985.
- Pascal, Blaise. *Pensées*.

- Plato. Phaedrus.
- Wingren, Gustaf. Luther on Vocation. Philadelphia: Muhlenberg, 1957.
- Adamala, Katarzyna P., et al. "Confronting Risks of Mirror Life." *Science* 386, no. 6728 (2024): 1351–1353.
- Baker, Darren J., et al. "Clearance of p16Ink4a-Positive Senescent Cells Delays Ageing-Associated Disorders." *Nature* 479 (2011): 232–236.
- Hayashi, Katsuhiko, et al. "Offspring from Oocytes Derived from In Vitro Primordial Germ Cell-like Cells in Mice." *Science* 338 (2012): 971–975.
- Hayflick, Leonard, and Paul S. Moorhead. "The Serial Cultivation of Human Diploid Cell Strains." *Experimental Cell Research* 25, no. 3 (1961): 585–621.
- Horvath, Steve. "DNA Methylation Age of Human Tissues and Cell Types." *Genome Biology* 14, no. 10 (2013): R115.
- Hutchison, Clyde A., III, et al. "Design and Synthesis of a Minimal Bacterial Genome." *Science* 351 (2016): aad6253.
- Jinek, Martin, et al. "A Programmable Dual-RNA-Guided DNA Endonuclease in Adaptive Bacterial Immunity." *Science* 337 (2012): 816–821.
- Jumper, John, et al. "Highly Accurate Protein Structure Prediction with AlphaFold." *Nature* 596 (2021): 583–589.
- Justice, Jamie N., et al. "Senolytics in Idiopathic Pulmonary Fibrosis." *EBioMedicine* 40 (2019): 554–563.
- Kagan, Brett J., et al. "In Vitro Neurons Learn and Exhibit Sentience When Embodied in a Simulated Game-World." *Neuron* 110, no. 23 (2022): 3952–3969.
- Kaplan, Jared, et al. "Scaling Laws for Neural Language Models." *arXiv:2001.08361* (2020).
- Lu, Yuancheng, et al. "Reprogramming to Recover Youthful Epigenetic Information and Restore Vision." *Nature* 588 (2020): 124–129.
- Niu, Dong, et al. "Inactivation of Genes in Pig Cells Using CRISPR-Cas9." *Science* 357 (2017): 1303–1307.
- Ocampo, Alejandro, et al. "In Vivo Amelioration of Age-Associated Hallmarks by Partial Reprogramming." *Cell* 167, no. 7 (2016): 1719–1733.
- Oldak, Bernardo, et al. "Complete Human Day 14 Post-implantation Embryo Models from Naive ES Cells." *Nature* 622 (2023): 562–573.
- Partridge, Emily A., et al. "An Extra-Uterine System to Physiologically Support the Extreme Premature Lamb." *Nature Communications* 8 (2017): 15112.
- Sato, Toshiro, et al. "Single Lgr5 Stem Cells Build Crypt-Villus Structures In Vitro." *Nature* 459 (2009): 262–265.

- Schaeffer, Rylan, Brando Miranda, and Sanmi Koyejo. "Are Emergent Abilities of Large Language Models a Mirage?" *NeurIPS* (2023).
- Smirnova, Lena, et al. "Organoid Intelligence (OI): The New Frontier in Biocomputing." *Frontiers in Science* 1 (2023): 1017235.
- Takahashi, Kazutoshi, and Shinya Yamanaka. "Induction of Pluripotent Stem Cells from Mouse Embryonic and Adult Fibroblast Cultures." *Cell* 126, no. 4 (2006): 663–676.
- Trujillo, Cleber A., et al. "Complex Oscillatory Waves Emerging from Cortical Organoids." *Cell Stem Cell* 25, no. 4 (2019): 558–569.
- Vaswani, Ashish, et al. "Attention Is All You Need." *NeurIPS* (2017).
- Wei, Jason, et al. "Emergent Abilities of Large Language Models." *TMLR* (2022).
- Willett, Francis R., et al. "A High-Performance Speech Neuroprosthesis." *Nature* 620 (2023): 1031–1036.
- Xue, Weiwei, et al. "Effective Cryopreservation of Human Brain Tissue and Neural Organoids." *Cell Reports Methods* 4, no. 5 (2024).
- Zhu, Yi, et al. "The Achilles' Heel of Senescent Cells: From Transcriptome to Senolytic Drugs." *Aging Cell* 14, no. 4 (2015): 644–658.
- Ackoff, Russell L. "From Data to Wisdom." *Journal of Applied Systems Analysis* 16 (1989): 3–9.
- Barth, Karl. *Church Dogmatics III/2*. Translated by H. Knight et al. Edinburgh: T&T Clark, 1960.
- Block, Ned. "On a Confusion about a Function of Consciousness." *Behavioral and Brain Sciences* 18, no. 2 (1995): 227–247.
- Burleigh, Michael. *Death and Deliverance: Euthanasia in Germany, 1900–1945*. Cambridge: Cambridge University Press, 1994.
- Butlin, Patrick, Robert Long, Eric Elmoznino, Yoshua Bengio, Jonathan Birch, et al. "Consciousness in Artificial Intelligence: Insights from the Science of Consciousness." *arXiv:2308.08708* (2023).
- Chalmers, David J. "Facing Up to the Problem of Consciousness." *Journal of Consciousness Studies* 2, no. 3 (1995): 200–219.
- Clark, Andy, and David J. Chalmers. "The Extended Mind." *Analysis* 58, no. 1 (1998): 7–19.
- Dennett, Daniel C. *The Intentional Stance*. Cambridge, MA: MIT Press, 1987.
- Eliot, T. S. Choruses from "The Rock." 1934.
- Frické, Martin. "The Knowledge Pyramid: A Critique of the DIKW Hierarchy." *Journal of Information Science* 35, no. 2 (2009): 131–142.
- Good, I. J. "Speculations Concerning the First Ultraintelligent Machine." *Advances in Computers* 6 (1965): 31–88.

- Harari, Yuval Noah. *Homo Deus: A Brief History of Tomorrow*. London: Harvill Secker, 2016.
- Herzfeld, Noreen L. *In Our Image: Artificial Intelligence and the Human Spirit*. Minneapolis: Fortress Press, 2002.
- Heschel, Abraham Joshua. *The Sabbath: Its Meaning for Modern Man*. New York: Farrar, Straus and Young, 1951.
- Heschel, Abraham Joshua. *God in Search of Man: A Philosophy of Judaism*. New York: Farrar, Straus and Cudahy, 1955.
- Hoekema, Anthony A. *Created in God's Image*. Grand Rapids: Eerdmans, 1986.
- Holzer, Elie, with Orit Kent. *A Philosophy of Havruta*. Boston: Academic Studies Press, 2013.
- Idel, Moshe. *Golem: Jewish Magical and Mystical Traditions on the Artificial Anthropoid*. Albany: SUNY Press, 1990.
- Kennedy, Brian K., et al. "Geroscience: Linking Aging to Chronic Disease." *Cell* 159, no. 4 (2014): 709–713.
- Kilner, John F. *Dignity and Destiny: Humanity in the Image of God*. Grand Rapids: Eerdmans, 2015.
- Koehler, Ludwig, and Walter Baumgartner. *The Hebrew and Aramaic Lexicon of the Old Testament*.
- Kurzweil, Ray. *The Singularity Is Nearer: When We Merge with AI*. New York: Viking, 2024.
- Levinas, Emmanuel. *Totality and Infinity*. Translated by Alphonso Lingis. Pittsburgh: Duquesne University Press, 1969.
- Locke, John. *An Essay Concerning Human Understanding*. 1690.
- Long, A. A., and D. N. Sedley. *The Hellenistic Philosophers*. Vol. 1. Cambridge: Cambridge University Press, 1987.
- López-Otín, Carlos, et al. "The Hallmarks of Aging." *Cell* 153, no. 6 (2013): 1194–1217.
- MacIntyre, Alasdair. *After Virtue: A Study in Moral Theory*. 3rd ed. Notre Dame: University of Notre Dame Press, 2007.
- Merleau-Ponty, Maurice. *Phenomenology of Perception*. Translated by Donald A. Landes. London: Routledge, 2012.
- Middleton, J. Richard. *The Liberating Image: The Imago Dei in Genesis 1*. Grand Rapids: Brazos Press, 2005.
- Mishnah Sanhedrin 4:5; Pirkei Avot 3:14; Genesis Rabbah 9:7.
- Nagel, Thomas. "What Is It Like to Be a Bat?" *The Philosophical Review* 83, no. 4 (1974): 435–450.
- OECD. *Digital Economy Outlook and Related Education Reports*. Paris: OECD, 2025–2026.

- Polanyi, Michael. *The Tacit Dimension*. Chicago: University of Chicago Press, 1966.
- Sandel, Michael J. *The Case against Perfection*. Cambridge, MA: Harvard University Press, 2007.
- Searle, John R. "Minds, Brains, and Programs." *Behavioral and Brain Sciences* 3, no. 3 (1980): 417-457.
- Singer, Peter. *Practical Ethics*. 3rd ed. Cambridge: Cambridge University Press, 2011.
- Smith, James K. A. *Desiring the Kingdom: Worship, Worldview, and Cultural Formation*. Grand Rapids: Baker Academic, 2009.
- Stanford Institute for Human-Centered Artificial Intelligence (HAI). *The AI Index Report*. Stanford, CA: Stanford HAI, 2026.
- Taylor, Charles. *Sources of the Self: The Making of the Modern Identity*. Cambridge, MA: Harvard University Press, 1989.
- Turing, Alan M. "Computing Machinery and Intelligence." *Mind* 59, no. 236 (1950): 433-460.
- von Rad, Gerhard. *Genesis: A Commentary*. Rev. ed. Philadelphia: Westminster, 1972.
- Westermann, Claus. *Genesis 1-11: A Commentary*. Translated by John J. Scullion. Minneapolis: Augsburg, 1984.
- Wolff, Hans Walter. *Anthropology of the Old Testament*. Translated by Margaret Kohl. Philadelphia: Fortress, 1974.
- World Bank. *Digital Progress and Trends Report*. Washington, DC: World Bank, 2025.
- Wright, N. T. *The Resurrection of the Son of God*. Minneapolis: Fortress, 2003.
- Zizioulas, John D. *Being as Communion*. Crestwood, NY: St Vladimir's Seminary Press, 1985.
- 『周易』 乾卦 九二 文言傳.
- 김종필. 「길이 지성이 되다」. 초록나무.
- 김종필. 「컴퓨터의 권세: 새로운 경제 패권과 세계의 재편」. 초록나무.
- 김종필. 「AI와 슈퍼 휴먼의 시대」. AI 시리즈 7. 초록나무.
- 김종필. 「말씀과 코드 사이, 인간의 자리」. AI 시리즈 8. 초록나무.