

SNU

그랜드퀘스트포럼

SNU Grand Quest Forum

Toward Agenda Setters

2026.6.18 (THU) 9:00 ~ 18:00

서울대학교 Commons 중앙도서관



GRAND
QUEST

서울대학교 그랜드퀘스트 이니셔티브
SNU Grand Quest Initiative, Seoul National University

Forum Program

포럼 프로그램

오전세션

09:00~09:05	개회사	유홍림 총장
09:05~09:10	영상축사	조정식 국회의장
09:10~09:15	축사	홍원화 한국연구재단 이사장
09:15~09:25	해외 석학 영상 메시지	김필립 Harvard University 교수 Daron Acemoglu MIT 교수
09:25~09:45	SNU 그랜드퀘스트 공개 이니셔티브 소개 영상 시청 포함	이정동 SNU 그랜드퀘스트 이니셔티브 연구단장

09:45~10:00 휴식

10:00~10:30	기조강연	현택환 SNU 그랜드퀘스트 디자인보드
10:30~11:20	청중과의 대화	이경무 SNU 그랜드퀘스트 디자인보드 구범진 SNU 그랜드퀘스트 디자인보드 이석재 SNU 그랜드퀘스트 디자인보드 정두현 SNU 그랜드퀘스트 디자인보드 현택환 SNU 그랜드퀘스트 디자인보드

오후세션

13:10~13:40	주제강연	최재천 이화여대 명예교수
13:40~14:10	주제강연	최인철 SNU 그랜드퀘스트 디자인보드
14:10~15:10	대담	최재천 이화여대 명예교수 송재용 SNU 그랜드퀘스트 디자인보드 전상직 SNU 그랜드퀘스트 디자인보드 최인철 SNU 그랜드퀘스트 디자인보드 홍성욱 SNU 그랜드퀘스트 디자인보드

15:10~15:20 휴식

15:20~17:20	SNU 그랜드퀘스트 문제 해설 및 대담	SNU 그랜드퀘스트 디자인보드
17:20~17:50	SNU 그랜드퀘스트 공모전 시상식	
17:50~18:00	폐회	

Contents

목 차

Opening Messages 4

개회사

축사

해외 석학 메시지

SNU Grand Quest Initiative 14

SNU 그랜드퀘스트의 시작과 확장

SNU 그랜드퀘스트의 차별점

SNU 그랜드퀘스트 프로세스

SNU 그랜드퀘스트 이니셔티브 추진 경과

2026 SNU Grand Quest 21

들어가는 말

디자인보드 소개

From Questions to Journeys 49

SNU 그랜드퀘스트 공모전

SNU 그랜드퀘스트 메이커

SNU 그랜드퀘스트 챌린지

Opening Messages



GRAND
QUEST

유홍림 서울대학교 총장

안녕하십니까, 서울대학교 총장 유홍림입니다.

SNU 그랜드퀘스트 포럼에 참석해주신 여러분을 진심으로 환영합니다.

올해로 개교 80주년을 맞는 서울대학교에 SNU 그랜드퀘스트 이니셔티브는 80년 역사의 큰 전환점입니다. 해방 후 척박한 근대 학문의 토양 위에서 설립된 서울대학교는 그동안 산업화와 민주화, 정보화와 세계화라는 시대적 과제 해결에 매진하며 대한민국의 발전을 이끌어 왔습니다.

인류 역사에서 유래를 찾기 어려운 우리나라의 압축 성장은 서구의 제도와 기술, 학문을 도입하고 추격하는 데서 시작했습니다. 그 과정은 불가피하게 기계적 정답을 찾는 교육과 주어진 과제를 가장 효율적으로 해결하는 단기실적 중심의 연구를 대학에 자리잡게 했습니다.

하지만 이제 추격의 정점에 서 있는 우리 대학과 우리나라는 과거의 전략으로는 더 이상 성장을 담보할 수 없는 상황에 이르러 있습니다. 또한 과거에 지식의 생산과 전수를 독점하던 대학의 역할은 오늘날 근본적인 도전에 직면해 있습니다. 재작년 노벨상을 수상한 알파폴드(AlphaFold)의 사례에서 보듯이, AI는 이미 인간이 수십 년에 걸쳐 축적한 지식을 학습하고, 이를 바탕으로 새로운 발견과 예측까지 수행하는 단계에 도달했습니다.

이제 우리에게 필요한 역량은 무엇이 중요한 문제인지 묻고, 아직 누구도 제기하지 못한 질문을 발견하며, 기술과 사회가 나아가야 할 방향을 성찰하는 능력입니다. 따라서 오늘날 대학은 정답을 가르치는 기관에서 새로운 질문을 창조하는 기관으로 거듭나야 하며, 이는 개교 80주년을 맞는 우리 대학의 가장 중요한 과제입니다.

SNU 그랜드퀘스트 이니셔티브는 서울대학교가 지난 80년의 역사를 넘어, 새로운 대전환을 만들기 위한 도전입니다. 우리는 무엇을 가르칠 것인가를 넘어 무엇을 질문할 것인가를 고민하고, 무엇을 연구할 것인가를 넘어 어떤 미래를 만들어 갈 것인가를 모색해야 합니다.

지난 3월 17일부터 한 달 동안 서울대학교 전 구성원을 대상으로 SNU 그랜드퀘스트 공모전이 열렸습니다. 총 2,143 개의 질문이 제출되었고, 그 중 일부가 현재 그랜드퀘스트 위크와 포럼 기간동안 전시되어 서울대 구성원들의 다양한 문제의식을 보여주고 있습니다.

더불어 서울대 그랜드퀘스트 디자인보드는 오랜 숙의와 토론을 거쳐, 2026 SNU 그랜드퀘스트를 도출하였습니다. 오늘 포럼에서 처음 공개되는 2026 SNU 그랜드퀘스트는 단순한 연구주제를 넘어, 앞으로 서울대학교가 도전하고자 하는 미래 연구의 방향이자 인류와 사회를 향해 던지는 담대한 질문입니다.

어쩌면 오늘 제시되는 질문들은 당장 답을 찾기 어려울 수도 있습니다. 하지만 인류 역사에서 위대한 발견과 혁신은 언제나 정교한 답이 아니라, 용기 있는 질문에서 시작되었다는 것을 우리는 기억해야 합니다.

그랜드퀘스트 프로그램의 궁극적인 목적도 정답을 찾는 데 있지 않습니다. 우리는 이 프로그램을 통해 연구과제를 발굴하는 데 그치지 않고, 서울대학교가 우리 사회와 인류를 위해 어떤 질문을 던져야 하는지 함께 모색하고자 합니다. 나아가 학문의 경계를 넘어 다양한 분야의 연구자와 학생들이 함께 질문하고 토론하며, 미래를 향한 새로운 지평을 열어가고자 합니다.

개교 80주년을 맞은 서울대학교는 이제 지난 성취를 기념하는 데 머물지 않고, 개교 100주년의 미래를 준비해야 합니다. 서울대학교는 대한민국의 성장과 발전을 이끌어온 '민족의 대학'에서 한 걸음 더 나아가, 인류 문명의 미래를 개척하는 '세계의 대학'으로 도약해야 합니다. 오늘 포럼이 서울대학교 새 역사의 출발점이 될 것이라 믿습니다.

다시 한 번, 참석해 주신 모든 분들께 감사드리며,
서울대학교가 새 역사를 써가는 여정에 함께
지혜와 역량을 모아 주시기를 부탁드립니다.

감사합니다.



서울대학교 총장 **유흥림**

조정식 국회의장

안녕하십니까, 국회의장 조정식입니다.

「2026 SNU 그랜드퀘스트 포럼」 개최를 진심으로 축하드립니다. 뜻깊은 자리를 마련해 주신 유홍림 서울대학교 총장님을 비롯한 관계자 여러분께 감사드립니다.

우리는 지금까지 수많은 문제를 해결하며 사회를 발전시켜 왔습니다. 그러나 오늘날 인류가 마주한 기후위기, 인공지능, 에너지 전환, 인구구조의 변화와 같은 복합적 난제들은 기존의 방식만으로는 접근하기 어려운 새로운 성격의 문제들입니다.

이럴 때일수록 중요한 것은 정답을 빠르게 찾는 능력이 아니라,누구도 묻지 않았던 질문을 발견하고 정의하는 능력입니다. 어떤 질문을 하느냐에 따라 연구의 방향이 달라지고,기술의 미래가 달라지며, 우리 사회의 미래도 달라집니다.

그런 의미에서 이번 서울대학교의 「SNU 그랜드퀘스트」는 매우 의미 있는 시도입니다.

위대한 발견은 언제나 담대한 질문에서 시작되었습니다. 새로운 학문도, 새로운 기술도, 새로운 상상력도 결국 질문으로부터 출발합니다. 오늘 이 자리가 대한민국의 미래를 향한 새로운 질문을 만나고,함께 탐구를 시작하는 계기가 되기를 바랍니다.

서울대학교가 열어갈 도전과 혁신의 여정에 큰 기대와 응원을 보냅니다.

감사합니다.



국회의장 **조정식**

홍원화 한국연구재단 이사장

안녕하십니까. 한국연구재단 이사장 홍원화입니다.

신록의 생명력이 가득한 6월, 대한민국 학문의 요람인 서울대학교에서 「SNU 그랜드퀘스트 포럼」이 열리게 된 것을 진심으로 축하드립니다. 이 뜻깊은 자리를 마련해 주신 유홍림 총장님과 이정동 SNU 그랜드퀘스트 이니셔티브 연구단 단장님을 비롯한 서울대 관계자 여러분, 그리고 무엇보다 오늘 이 자리를 빛내주신 연구자 여러분께 깊은 감사와 존경의 마음을 전합니다.

대한민국은 지난 반세기 동안 전 세계가 놀랄만한 성취를 일구어 왔습니다.

전후의 폐허를 딛고 세계 경제의 주축으로 우뚝 서기까지, 우리 학계와 연구자들은 시대적 과제를 누구보다 충실히 수행해 왔습니다. 반도체, 조선, 이차전지 등 오늘날 우리가 누리는 번영은 세계가 던진 질문에 가장 정확하고 신속하게 답을 내놓았던 '탁월한 해법 제시자(Problem Solver)'들의 역량이 만들어낸 결실입니다.

그러나 지금 우리는 저성장과 기술 패권 경쟁, AI 전환과 인구구조 변화라는 복합적 위기 속에서 새로운 도전의 문턱 앞에 서 있습니다. 기존의 문제를 더 잘, 더 빠르게 해결하는 '추격형 방식'만으로는 지속 가능한 성장도, 구조적 위기 돌파도 기대하기 어렵습니다. 이제 대한민국은 주어진 문제를 푸는 단계를 넘어, 인류의 미래를 향해 먼저 질문을 던지는 '의제 설정자(Agenda Setter)'로 거듭나야 합니다.

이러한 전환의 시점에서 서울대가 새롭게 제시하는 '그랜드퀘스트'는 대한민국 학계의 지형을 바꿀 과감한 시도입니다. '그랜드퀘스트'의 핵심은 학계의 오랜 통념에 과감히 의문을 제기하고, 우리가 당연하게 여겨온 고정관념을 흔드는 '근원적인 질문'에 있습니다.

당장은 해법이 보이지 않아 무모해 보일 수도 있고, 단기 성과를 보장하기 어려워 기존 지원 체계에서 외면 받기 쉬웠던, 말하자면 '서랍 속에 잠들어 있던 질문들'을 세상 밖으로 꺼내는 일입니다.

한국연구재단 역시 이러한 시대적 변화의 흐름에 발맞추어, 연구 지원 체계를 근본적으로 혁신해 나가고자 노력하고 있습니다. 실패의 경험마저도 소중한 자산으로 축적될 수 있도록 연구 제도를 혁신하고, 실패가 낙인이 아닌 다음 단계의 도약을 위한 징검다리가 되도록 촘촘하고 든든한 지원 체계를 만들어 가고자 합니다.

이제 한국연구재단은 이미 '그려진 지도'를 따라가는 안전한 연구 지원을 넘어, 나침반 하나로 '지도에도 없는 길'을 개척하는 연구자들의 가장 든든한 버팀목이 되겠습니다. 단순히 성공과 실패를 판정하는 데 그치지 않고, 시행착오의 전 과정을 소중한 지적 자산으로 승화시켜 후속 세대에게 물려주는 '진화하는 연구 생태계'를 조성하는 데 앞장서겠습니다.

존경하는 연구자 여러분,
여러분의 서랍 속에 잠들어 있던 그 뜨거운 질문들을 이제 세상 밖으로 펼쳐 보여주십시오.

그 질문이 비록 거칠고 불완전할지라도, 새로운 탐색의 방향을 가리키고 있다면 그 자체로 이미 충분한 가치를 지닙니다. 여러분이 던지는 그 과감한 질문들이 곧 대한민국의 내일을 여는 열쇠가 될 것입니다.

다시 한번 「SNU 그랜드퀘스트 포럼」의 개최를 진심으로 축하드리며, 참석하신 모든 분의 앞날에 건승과 행복이 가득하시길 기원합니다.

감사합니다.



한국연구재단 이사장 **홍원화**

김필립 Harvard University 교수

- 하버드 물리학과 교수
- 2023년 벤자민 프랭클린 메달 수상
- 미국 국립과학원 회원

안녕하세요. 하버드대학교 물리학과 김필립입니다.

이렇게 서울대학교 그랜드퀘스트 포럼에 참여할 수 있는 기회를 주셔서 감사합니다.

제 생각에 학문은 점진적인 발전과 획기적이고 혁명적인 발전이 혼재되어 나타나는 과정이라고 생각합니다. 대부분의 경우에는 점진적인 발전이 이루어지겠지만, 10년에 한 번, 100년에 한 번, 아니면 그보다도 더 오랜 시간에 한 번씩 아주 혁명적인 발전이 나타나기도 합니다. 저 역시 그런 획기적인 발전, 학문의 물줄기를 바꿀 수 있는 과정들을 주도하거나 최소한 참여할 수 있는 경험을 했으면 좋겠다는 것이 학자로서의 당연한 욕심이고, 아마 학문에 참여하는 많은 분들이 공통적으로 갖고 있는 생각이 아닐까 합니다. 하지만 그런 혁명적이고 획기적인 발전은 늘 일어나는 것이 아닙니다. 그런 면에서 우리는 일상적인 연구를 통해 학자로서의 삶을 살아가며, 점진적인 학문의 발전에 기여하는 데 많은 시간을 쏟게 되는 것 같습니다. 또한 일상의 작은 문제들을 풀어나가는 것 역시 매우 중요하며, 그런 문제들을 해결해 나가는 것이 학자의 중요한 소임이라고 생각합니다. 그러한 과정에 대해 충분히 수련되어 있어야 한다고 생각합니다.

하지만 준비되어 있지 않을 때는, 설사 중요한 문제를 풀 수 있는 기회가 우리 앞을 지나가더라도 그것을 알아보지 못하거나, 알고도 놓치게 될 수 있습니다.

그런 면에서 우리는 늘 마음 한구석에 중요한 문제들, 그리고 깊은 문제들에 대한 천착을 가지고 있어야 하지 않을까 생각합니다. 제가 그래핀 연구를 처음 시작했을 때를 떠올려 보면, 당시에는 그래핀이나 2차원 물질에 대한 연구 자체가 거의 없었고, 학계에서도 본격적인 논의가 시작되지 않았던 시기였습니다. 따라서 그런 생각에 대해 동의하거나 흥미를 보이는 사람도 매우 적었습니다. 심지어 제가 처음 쓴 프로포절 가운데에는 ‘그래핀을 만들어 그것을 연구하겠다’는 내용이 있었는데, 그 프로포절은 좋은 평가를 받지 못했습니다. 참신한 생각이라고 평가해 주신 분들도 있었지만, 대부분의 리뷰어들은 그것이 불가능하거나 기술적으로 증명되지 않은 것을 전제로 한 제안이라고 보았습니다. 결국 첫 번째 프로포절은 받아들여지지 않았습니다.

상당히 낙담하고 실망하기도 했지만, 제 생각에는 비록 충분히 정련된 생각은 아니더라도 중요한 아이디어라는 확신이 있었습니다. 그래서 다음 프로포절에서는 다른 주제를 제안했지만, 그 아이디어 자체는 계속 가지고 연구를 이어갔던 기억이 납니다.

비슷하게도 이런 획기적인 생각이나 엉뚱한 생각들은 처음에는 인정받지 못하는 경우가 많습니다. 많은 경우 사람들을 단념하게 만들기도 하고, 단순히 엉뚱한 생각으로 치부되기도 합니다. 하지만 그런 생각들을 쉽게 포기하지 않고, 계속 생각하고 정련하며 재정련해 나가는 과정이 중요하지 않을까 생각합니다.

아마도 그런 것들이 그랜드퀘스트를 만들어 가는 중요한 첫 번째 발걸음이 될 것이라고 생각합니다. 그리고 그런 생각들이 큰 그림 안에서 어떻게 연결될 수 있는지를 늘 고민해 보는 것 또한 중요한 단초가 될 것입니다. 더 나아가 정말 큰 생각들은 우리가 속한 좁은 전문 영역 안에서만 나오는 것이 아니라, 그것을 확장하여 이웃 분야나 전혀 다른 분야의 생각들과 연결되는 과정에서 탄생하는 경우가 많다고 생각합니다.

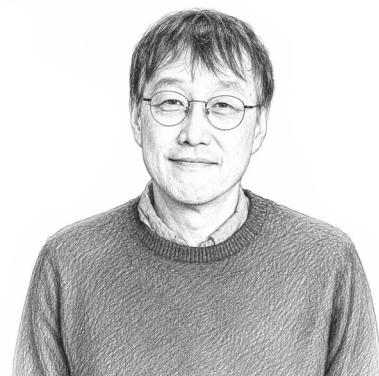
학문의 큰 물줄기를 바꿀 수 있는 아이디어들도 바로 그런 지점에서 나오는 것이 아닐까 합니다. 그래서 저는 이러한 초학제적(interdisciplinary) 연구가 더욱 중요하다고 생각합니다.

마지막으로, 특히 학문의 길을 막 시작하는 여러분께 말씀드리고 싶습니다. 매일의 연구 과정에서는 가까운 문제들에 집중해야 하는 순간들이 많습니다. 그것은 매우 중요합니다.

하지만 동시에 늘 그 너머의 큰 그림을 그려보고자 하는 마음, 그리고 조금은 엉뚱한 생각들을 해보고자 하는 마음도 함께 가졌으면 합니다. 그러한 생각들을 긴 호흡으로 다듬고 단련해 나가는 시간이 필요합니다. 책상 속의 아이디어는 하나가 아니라 열 개, 백 개가 될 수 있습니다. 그리고 그 많은 생각들 가운데 무엇을 남기고 어떻게 발전시켜 나갈 것인가는 늘 고민해야 할 문제입니다.

그 과정에서 주변의 동료 연구자들, 그리고 나와는 다른 생각을 가진 학자들과 함께 이야기를 나누는 시간을 갖는 것이 큰 도움이 될 것이라고 생각합니다.

감사합니다.



Harvard University 교수 **김필립**

Daron Acemoglu MIT 교수

- MIT 경제학과 인스티튜트 교수
- 2024년 노벨경제학상 수상
- '국가는 왜 실패하는가', '권력과 진보' 등 저자

안녕하세요.

서울대학교 그랜드 퀘스트 포럼에 이렇게 원격으로나마 인사를 전할 수 있게 되어 매우 기쁩니다.

지금 우리가 인류 역사의 전환점에 서 있다는 점에서, 이 포럼은 매우 시의적절한 시도라고 생각합니다. 인공지능이 앞으로 우리 삶의 많은 부분을 형성하게 될 것이라는 점은 분명합니다. 그리고 AI가 어디로 향하고 있는지에 대한 정당한 우려 또한 존재합니다. 이러한 우려는, 현재 인공지능이 미국과 중국이라는 두 나라, 그리고 소수의 기업과 의사결정자들의 손에 좌우되고 있다는 사실로 인해 더욱 커지고 있습니다. 하지만 인공지능은 하나의 단일한 무언가가 아닙니다. 우리가 기계의 능력을 활용하는 방식에는 다양한 스펙트럼이 존재합니다. 현재 우선순위로 여겨지는 범용인공지능(AGI) 역시 그 여러 가능성 중 하나일 뿐입니다. 더욱 심각한 문제는 현재의 방향이 사회적으로도 해롭고, 경제적으로도 의문스럽다는 점입니다. 사회적으로 해로운 이유는 AGI가 자동화를 중심으로 한 의제의 정점에 있기 때문입니다. 다시 말해, 노동자들이 수행해 오던 업무에서 그들을 밀어내고 그 일을 기계와 알고리즘이 대신하도록 만드는 것입니다.

물론 자동화 자체가 잘못된 것은 아닙니다. 자동화는 앞으로도 인류 역사와 함께할 것입니다. 문제는, 다른 가치들을 희생하면서까지 오직 자동화에만 몰두할 경우, 이것이 극심한 불평등과 일자리 상실로 이어지고, 나아가 생산 과정과 사회생활에서 점점 더 무의미하고 대체 가능한 존재가 되어 가는 인간들이 전반적으로 삶의 목적을 잃게 만든다는 데 있습니다. 이러한 방향은 경제적으로도 의문스럽습니다.

그 이유는 인공지능이 인간의 지능과 완전히 동일하다는 전제에 기반하기 때문입니다.

하지만 실제로는 그렇지 않습니다. 인공지능과 인간 지능 사이에는 많은 차이가 존재합니다. 서로 다른 두 존재를 가장 잘 결합하는 방법은 각자의 강점을 살려 상호보완적으로 활용하는 것입니다. 바로 이러한 접근이 존재합니다.

저는 이를 '친노동(Pro-Worker) 의제'라고 부릅니다.

이는 인공지능이 노동자에게 새로운 역량과 새로운 수준의 전문성, 그리고 새로운 업무를 창출하도록 작동하게 만드는 것을 뜻합니다. 그 이점은, 비록 아직은 미약하지만 현재까지의 증거들이 이러한 방향이 생산성을 더욱 높여 줄 것임을 시사한다는 데 있습니다. 또한 그 사회적

함의도 훨씬 더 긍정적일 것입니다. 노동자를 주변으로 밀어내는 대신, 오히려 생산 과정의 중심에 서게 만들기 때문입니다. 친노동자 의제 아래에서는 불평등이 반드시 심화될 필요도 없습니다. 또한 소수 기업이 기술을 독점하는 현상도 훨씬 완화될 수 있습니다. 아직 증거가 많지 않은 이유는 우리가 실제로 친노동 의제를 충분히 시도해보지 않았기 때문입니다. 현재는 일부 사례만 존재합니다. 그럼에도 기존의 연구와 경험은 이 접근이 충분히 실현 가능하다는 점을 강하게 시사하고 있습니다.

우리는 AI를 활용해 인간이 할 수 없는 일을 수행하게 하고, 이를 인간의 지식과 결합할 수 있습니다. 이는 전 지구적 차원의 의제입니다. 어느 한 국가만으로는 이 목표를 달성할 수 없습니다. 하지만 선도적 사례가 갖는 영향력은 매우 중요합니다. 한국은 이미 자동화 분야에서 리더십을 보여준 바 있습니다. 한국은 세계에서 가장 빠른 속도로 고령화가 진행된 국가였고, 현재는 중국 또한 같은 과제를 마주하고 있습니다. 더 중요한 점은 한국의 기업과 기술자들이 자동화를 추진하면서도 노동자를 완전히 배제하지 않았다는 것입니다. 오히려 노동자가 계속 참여하는 가운데 생산성을 높이는 방법을 찾아왔습니다. 그리고 지금 한국은 인공지능 분야에서도 리더십을 보여주고 있습니다. 하지만 저는 현재의 방향에서의 리더십만으로는 충분하지 않다고 생각합니다.

우리에게는 새로운 로드맵이 필요합니다.

그리고 한국뿐 아니라 유럽, 미국, 중국 역시 이 논의에 함께 참여해야 합니다.

제 의견을 공유할 수 있도록 초대해 주셔서 진심으로 감사드리며, 앞으로도 이 대화를 계속 이어갈 수 있기를 기대합니다.

감사합니다



MIT 교수 **Daron Acemoglu**

* 본 축사는 영문 원문을 기반으로
의미 전달에 중점을 두어 번역하였습니다.

SNU Grand Quest Initiative



GRAND
QUEST



GRAND QUEST

그랜드퀘스트란,

해당 분야에서 지배적 전제로 작동해 온 통념을 해체하고,
이를 대체할 새로운 개념적 틀을 제시하며,
그 결과 기존과 구분되는 독립적 연구 영역이나
학문적 장르가 형성될 가능성을 내포한 근본적 질문

질문을 제시하는 대학으로 Toward Agenda Setters



서울대학교는 세계가 던진 문제에 답하는 대학을 넘어
미래 연구의 질문을 정의하는 새로운 의제의 선도자를 지향합니다.

SNU 그랜드퀘스트 이니셔티브는
아직 누구도 묻지 않은 질문과 새로운 학문의 출발선을 찾기 위한
서울대학교의 새로운 시도입니다.

시작과 확장

서울대학교는
2022년부터 질문 기반 연구의
새로운 가능성을 탐색해왔습니다.

그리고 2026년,
SNU 그랜드퀘스트 이니셔티브로
확장되었습니다.



차별점

GRAND QUEST | 미래 의제의 제시

선도적 질문을 제시합니다.

PARTICIPATION | 모두의 여정

전 구성원의 지성을 결집하며 모두가 참여합니다.

EVOLUTION | 결과보다 도전

연구의 성과보다 축적을 장려하며, 경로 변경의 진화를 허용합니다.

SHARING | 미래세대를 위한 축적

시행착오의 경험을 자산화하여 미래세대에게 전달합니다.

프로세스

Open Call

01 Grand Quest Agora | 도전적 질문의 광장

SNU 그랜드퀘스트 공모전을 통해 전공과 소속을 넘어
서울대학교 구성원 누구나 도전적 질문을 제안하고 공유할 수 있습니다.
공모전은 단순한 선발 과정이 아니라, 질문을 던지는 일 자체가
의미 있는 문화적 활동임을 함께 만들어가는 과정입니다.



Design Board Workshop

02 Defining the Questions | 질문의 방향 설정

서울대학교 각 분야의 연구자들이 모여
인류의 공적 과제를 담은 도전적 질문의 영역을 제시합니다.
디자인보드 워크숍을 통해 서울대학교가
공식적으로 제시하는 SNU 그랜드퀘스트가 도출됩니다.



Forum & Week

03 Expanding the Questions | 결과보다 도전

도출된 그랜드퀘스트를 교내외·국내외 학문 커뮤니티와 공유합니다.
공모전과 디자인보드에서 등장한 다양한 질문들이
포럼과 워크 프로그램을 통해 하나의 거대한
질문 교류의 장으로 확장됩니다.



Challenge Program

04 Exploring New Frontiers | 미래세대를 위한 추적

연구자들은 그랜드퀘스트에 대해 다양한 방식으로 해법을 탐색합니다.
장기적 관점에서 연구를 지원하며, 질문의 진화와 경로 변경 또한 허용합니다.
실패를 단정적으로 판정하지 않고,
도전의 과정 자체를 연구의 자산으로 축적합니다.

2026. 3. 출범식



2026. 3. - 4. 공모전

서울대학교가 여러분의 가장 도전적인 질문을 찾습니다.

SNU 그랜드퀘스트 공모전

SNU GRAND QUEST OPEN CALL

교과서의 전제를 다시 묻는 질문
기존의 통념을 해체하고 새로운 출발선을 제시하는 질문
새로운 학문의 분야를 지향하는 질문

여러분의 질문이 새로운 학문을 열고
대한민국과 인류의 미래를 바꾸는 출발점이 됩니다.

We seek questions that challenge textbook assumptions, dismantle conventional wisdom, and pave the way for entirely new fields of scholarship.
Seoul National University is searching for your boldest questions. Join us in posing the questions that open new fields of scholarship—and let your inquiry become the starting point that reshapes the future of Korea and humanity.

- 공모 기간 Call for Submission
2026. 3.17.(TUE) - 4.17.(FRI)
- 심사 방식 Evaluation
**100% 블라인드 평가
FULLY BLIND REVIEW**
- 공모 가이드 포함 상세 안내
Details (incl. submission guide)
grandquest.snu.ac.kr

- 결과 발표 | Results Announcement
- 2026.6.5(금) 홈페이지 공지 예정
- Jun 5 (Fri), 2026
- 선정 및 포상 | Selection & Award
- 우수 질문 단당 최대 500만원 포상 (100명)
- KRW 5,000,000 per selected question
- 응모 자격 | Eligibility (참여 가능한 Individual or team)
- 서울대학교 전체 구성원
- SNU members **학부생 대상은 제외**
- Faculty/Researcher Track **Student Track**
- 제출 분량 | Length
- 1,000자 이내, 공백문 포함
- Up to 1,000 characters
- 제출 방법 | Submission
- 이메일 (grandquest@snu.ac.kr)
- Email to grandquest@snu.ac.kr
- 문의 | Contact grandquest@snu.ac.kr



2026. 4. 커넥트: Join My Crew!



2026. 5. 디자인보드 워크숍



2026 SNU Grand Quest



GRAND
QUEST



서울대학교가 던지는 여섯 개의 도전적인 질문

질문은 언제나 답보다 먼저 온다.

인류의 역사를 바꾼 것은 완성된 해답만이 아니었다. 새로운 시대는 늘 새로운 질문에서 시작되었다. 지구는 우주의 중심인가, 생명은 어떻게 진화하는가, 기계는 생각할 수 있는가. 당대에는 무모해 보였던 질문들이 오히려 학문의 지평을 넓히고 문명의 방향을 바꾸어 왔다.

오늘 우리는 또 하나의 전환점 앞에 서 있다. 인공지능은 인간의 지적 활동 전반에 스며들기 시작했고, 생명과학은 생명의 가장 깊은 비밀에 다가서고 있다. 기후위기와 인구구조의 변화, 노동의 소멸과 에너지 전환은 우리가 당연하게 여겨 온 질서와 전제를 되돌아보게 한다. 지금까지의 성공 경험과 익숙한 해법만으로는 답할 수 없는 문제들이 우리 앞에 가득하다.

한국 사회에 이 전환점은 더욱 각별하게 다가온다. 지난 수십 년간 우리는 앞서간 나라들이 던진 질문에 더 빠르고 더 정확한 답을 내며 성장해 왔다. 정답이 어딘가에 이미 놓여 있던 시대, 추격은 우리의 가장 강력한 전략이었다. 그러나 이제는 따라갈 벤치마크가 없는 경지에 서 있다. 선도란 남보다 먼저, 아직 누구도 던지지 않은 최초의 질문을 던지는 데서 비로소 시작된다. 추격에서 선도로 넘어가는 길목에서 우리에게 필요한 것은 더 많은 답이 아니라, 처음으로 던지는 더 깊고 날카로운 질문이다.

오늘 우리가 마주한 문제들은 어느 한 분야의 지식만으로는 풀리지 않는다. 기술의 문제는 곧 사회의 문제가 되고, 생명의 문제는 인간 이해의 문제로 이어지며, 에너지의 문제는 문명의 지속가능성과 맞닿는다. 개별 학문이 쌓아온 그간의 성취 위에서 그 경계를 넘어서는 질문을 던지고 새로운 학문의 탄생을 기대하는 이유가 여기에 있다.

2026년 늦봄, 서로 다른 분야를 대표하는 열여덟 명의 서울대학교 교수가 한자리에 모여 사흘에 걸쳐 치열하게 토론했다. 우리가 찾으려 한 것은 당장의 현안이나 특정 분야의 첨단 주제가 아니었다. 앞으로 수십 년 동안 연구자와 학생, 기업과 사회가 함께 붙들고 생각해야 할 질문은 무엇인가. 우리는 그 물음에서 출발했다.

여기 담긴 여섯 개의 그랜드퀘스트에는 해답이 없다. 누구도 떠올린 적 없는 문제도 아니다. 상당수는 이미 세계 곳곳에서 활발히 연구되고 있음을 안다. 그러나 우리는 그 성과들을 토대삼아, 지금의 연구가 충분히 의심하지 않는 더 깊은 전제와 통념을 다시 묻고자 했다. 그 질문을 끝까지 밀고 나갈 때 어떤 길이 열리는지를 보려 했다.

여섯 질문은 서로 다른 곳을 향하는 듯 보이지만, 결국 이 질문들은 하나의 물음으로 모인다. 우리는 앞으로 어떤 존재가 될 것이며, 어떤 미래를 만들어 갈 것인가.

우리는 이 질문들의 답을 알지 못한다. 어쩌면 한 세대 안에 얻지 못할 수도 있다. 그러나 학문의 진정한 역할은 이미 아는 것을 되뇌는 데 있지 않다. 아직 알지 못하는 것을 드러내고, 더 나은 질문으로 새로운 길을 여는 데 있다. 그래서 여섯 개의 그랜드퀘스트는 모두 질문으로 시작해 질문으로 끝난다. 이 질문들이 새로운 학문을 낳고, 유례없는 협력을 만들며, 대안적 세계를 상상하는 계기가 되기를, 그리고 그 길 위에서 다음 세대가 한국과 인류의 미래를 위해 한 걸음 더 나아가기를 소망한다.

SNU 그랜드퀘스트 디자인보드

**구범진, 구본권, 김병연, 김빛내리, 김용대, 남좌민, 성제경, 송재용, 윤제용,
이경무, 이석재, 전상직, 정두현, 최도일, 현택환, 홍성욱, 황철성, 이정동(총괄)**

인공지능시대, 민주주의와 자본주의는 지속가능한가?

Can Democracy and Capitalism Survive the Age of AI?

인공지능은 인간처럼 망각할 수 있는가?

Can AI Forget?

인공지능은 손상을 스스로 회복할 수 있는가?

Can AI Heal Itself?

생명의 시계를 제어할 수 있는가?

Can the Clock of Life Be Controlled?

삶의 의지를 분자수준에서 설명할 수 있는가?

Can the Will to Live Be Encoded in Molecules?

에너지 시스템은 자율적으로 균형을 찾을 수 있는가?

Can Energy Systems Balance Themselves?

Can Democracy and Capitalism Survive the Age of AI?

시장과 사적 소유에 기댄 자본주의와 시민주권에 기댄 민주주의는 오늘의 문명을 떠받치는 두 기둥이다. 자본주의는 산업혁명 이후 대공황과 계획경제의 도전을 거치면서도 스스로를 고쳐 가며 인류를 번영으로 이끌었다. 민주주의 역시 수많은 경제위기와 세계대전의 격변을 견디며 더 단단해졌다. 새로운 변화의 물결이 닥칠 때마다 두 제도는 바뀐 현실을 끌어안으며 균형을 다시 찾았다. 갈등과 불평등이 깊어질 때조차, 우리는 그것을 민주주의와 자본주의라는 틀 안에서 풀 문제로 여겼다.

그러나 인공지능의 비약적 발전은 두 제도가 딛고 선 전제마저 처음으로 되묻게 한다. 자본주의는 인간의 노동과 창의가 가치를 낳고, 그 가치가 소득과 소비, 투자로 순환되는 구조를 전제로 한다. 민주주의는 시민이 지식과 정보를 스스로 판단해 행동한다는 믿음 위에 서 있다. 그런데 인공지능은 한낱 신기술을 넘어 노동과 지식, 정보의 생산과 유통, 의사결정과 창작, 정치적 담론으로 빠르게 번지며 삶의 거의 모든 영역에 거대한 영향을 미치고 있다.

인공지능이 육체노동을 넘어 인지노동과 의사결정까지 대신한다면 인간의 노동은 어떻게 달라지는가. 부의 창출과 순환은 어떻게 바뀌는가. 합리적 계획이 가능하다는 믿음하에 계획경제가 다시 등장할 것인가. 데이터와 연산력, 플랫폼이 소수에 쏠릴 때 시장의 경쟁과 혁신은 어떤 모습이 되는가. 정보와 지식의 생산마저 기계로 넘어간다면, 스스로 판단하는 시민이라는 전제는 무엇에 기대는가. 알고리즘이 정보의 흐름과 사회적 판단에 깊이 끼어들 때 공론장의 신뢰는 어떻게 지켜지는가. 인공지능이 부의 양극화를 악화시키는 것이 아니라 누그러뜨리는 도구가 되려면 법과 제도를 어떻게 고쳐야 하는가. 선거와 대의제의 형식은 남더라도 민주주의의 실체는 어떻게 바뀌는가. 그리고 우리는 그 변화를 어떻게 바람직한 쪽으로 이끌 수 있는가.

우리는 그 답을 알지 못한다. 인공지능은 기존 질서를 무너뜨릴 수도, 인간의 자유와 번영을 새로운 차원으로 넓힐 수도 있다. 인공지능이라는 하나의 기술 안에 위험과 가능성이 공존하고 있다. 두 제도를 떠받칠 새로운 기술이 나올 수도, 제도 자체를 뜯어고쳐야 할 수도, 혹은 우리가 당연히 여긴 두 원리가 전혀 새로운 형태로 다시 구현될 수도 있다.

역사를 돌아보면 사회를 바꾼 것은 기술만이 아니었다. 증기기관은 그것을 담아낼 제도와 법, 교육과 문화가 함께 자란 뒤에야 문명의 원동력이 되었다. 인공지능도 마찬가지다. 앞으로의 사회는 누가 더 뛰어난 인공지능을 만드느냐만으로 갈리지 않는다. 그 기술과 함께 살아갈 제도와 질서를 어떻게 빚느냐가 더 중요하다.

한국 사회에 이 질문은 더욱 큰 의미를 갖는다. 한국은 산업화와 민주화를 가장 압축적으로 이룬 나라로서, 두 제도의 힘과 약점을 누구보다 가까이서 겪었다. 오늘 한국은 인공지능 혁명의 한복판에 선 동시에, 인구문제와 세대, 진영의 갈등, 양극화의 압력을 가장 첨예하게 마주한 사회이기도 하다. 그래서 이 질문은 먼 미래의 추상이 아니라 우리 사회의 바로 내일을 묻는다.

이 질문은 어느 한 학문의 몫이 아니다. 자연과학과 공학, 경제학과 정치학, 법학과 철학, 교육학이 제 분야의 답을 모으는 데 그치지 않고, 한자리에서 서로의 전제를 의심하며 함께 묻고 함께 설계해야 한다.

인공지능시대, 민주주의와 자본주의는 지속가능한가.

이 질문을 둘러싼 논의는 이미 세계 곳곳에서 뜨겁다. 그러나 기술이 내달리는 속도에 견주면, 그 기술을 담아낼 틀을 다시 그리려는 노력은 아직 더디다.

이 질문은 기술의 미래를 넘어 한국 사회와 인류 문명의 미래로 향한다. 서울대학교 그랜드퀘스트는 이 질문을 우리 시대의 중요한 연구 과제로 제안한다. 더 많은 연구자들이 함께 답을 구하며, 그 기술을 담아낼 새로운 제도와 질서를 함께 그려 나가기를 기대한다.

Can AI Forget?

인류는 오랫동안 기억을 지식과 지능의 원천으로 여겨 왔다. 더 많이 기억하고 더 정확하게 저장하는 능력은 개인의 학습과 문명의 발전을 이끄는 힘이었다. 컴퓨터도 같은 길을 걸었다. 정보기술은 더 큰 저장장치와 더 빠른 검색을 목표삼아 발전해왔고, 오늘의 인공지능 역시 방대한 데이터를 학습하고 저장하며 놀라운 성능을 보여주고 있다.

그러나 많이 기억하는 것이 곧 잘 아는 것은 아니다. 사람은 끊임없이 들어오는 정보 가운데 무엇을 남기고 무엇을 흘려보낼지 가려낸다. 어떤 경험은 오래 남고 어떤 경험은 사라지며, 기억도 단단해지는 것이 있는 반면 흐려지는 것도 있다. 망각은 기억의 실패가 아니라, 변화하는 환경에서 중요한 것과 그렇지 않은 것을 가르는 의도적인 적응의 산물이다.

최근 인공지능은 언어를 생성하는 수준을 넘어 현실에서 움직이고 판단하는 피지컬 AI로 빠르게 확장되고 있다. 사람은 언어만으로 세계를 이해하지 않는다. 보고 듣고 만지고 움직이며 세계가 어떻게 돌아가는지에 대한 내적 모형, 곧 세계모델을 빙다. 피지컬 AI 역시 시각과 청각, 촉각과 공간, 사람과의 상호작용이 얹힌 방대한 멀티모달 데이터를 받아들이며 자신만의 세계모델을 만들어 간다.

그렇다면 인공지능은 이 모든 경험을 끝없이 저장해야 하는가. 모든 정보를 똑같이 보존한다면 메모리와 연산 자원은 빠르게 소모되고, 중요한 경험과 사소한 경험이 뒤섞여 판단의 신뢰마저 흔들린다. 반대로 인공지능이 스스로 무엇을 남기고 무엇을 잊을지 판단할 수 있다면, 그것은 단순한 저장 장치를 넘어 변화하는 환경에서 경험을 조직하고 갱신하며 적응하는 새로운 지능으로 나아갈 수 있다.

인간의 뇌는 모든 정보를 똑같이 다루지 않기 때문에 불과 수십 와트의 에너지로도 기억과 판단, 학습과 행동을 한꺼번에 해낸다. 사소한 신호는 약화시키고 의미 있는 경험은 강화하며, 환경이 바뀌면 기억의 그물망 자체를 다시 짠다. 망각은 정보의 손실이 아니라 지능을 지탱하는 조절 기제이며, 에너지 효율을 높이고, 적응을 가능하게 한 진화의 전략이다. 어쩌면 미래 인공지능의 성패도 얼마나 많이 기억하느냐가 아니라, 무엇을 남기고 무엇을 잊느냐를 얼마나 잘 가리느냐로 갈릴지 모른다.

이 물음은 기술을 넘어 인간 지능의 본질에 닿는다. 사람은 모든 경험을 그대로 쌓지 않고, 기억할 것을 가려내며 스스로를 이해하고 자란다. 앞으로 개인화된 인공지능이 한 사람과 오랜 시간 관계를 맺는다면, 무엇을 기억하고 무엇을 잊을지는 기술의 문제를 넘어 신뢰와 정체성의 문제가 된다.

이 질문은 한국 사회의 맥락과 구체적으로 맞닿아 있다. 더 많이 저장하고 더 크게 연산하는 인공지능의 규모 경쟁에서 한국은 분명 후발주자다. 그러나 그 경로가 언제까지 지속될 수 있는지는 분명하지 않다. 무엇을 남기고 무엇을 잊을지 가리는 지능은, 이미 그려진 인공지능의 발전 로드맵을 더 빨리 따라가는 일이 아니라 그 로드맵 자체의 타당성을 다시 묻는 일이다. 자원의 한계를 먼저 마주한 후발주자이기에, 한국은 규모를 쫓기보다 다른 원리를 먼저 물어야 한다.

이 질문은 어느 한 분야의 언어로는 답할 수 없다. 컴퓨터공학, 뇌과학과 신경과학, 전자공학과 심리학, 그리고 철학이 한자리에서 기억과 망각, 학습과 적응의 본질을 다시 물어야 한다.

인공지능은 인간처럼 망각할 수 있는가.

이 질문을 둘러싼 연구는 이미 활발하게 쌓이고 있다. 장기기억 구조와 검색증강생성, 지속학습과 뉴로모픽 컴퓨팅이 그 예다. 그러나 이들은 대개 사람이 설계한 규칙 안에서 정보를 압축하거나 검색하는 데 머문다. 무엇을 기억하고 무엇을 잊을지 그 중요성과 맥락을 스스로 가리는 능력에는 아직 닿지 못했다. 기억을 늘리는 길이 아니라, 망각을 지능의 핵심 원리로 다시 볼 때 어떤 가능성이 열릴수 있는가.

서울대학교 그랜드퀘스트는 이 질문을 우리 시대의 중요한 연구 과제로 제안한다. 더 많은 연구자들이 함께 답을 구하며, 잘 잊는 지능이라는 새로운 길에 한 걸음 더 다가가기를 기대한다.

Can AI Heal Itself?

오늘날 인공지능은 많은 영역에서 인간을 능가하고 있고, 그 추세는 날로 빨라지고 있다. 의료 진단을 돕고 금융시장을 분석하며 자율주행차와 로봇을 제어한다. 앞으로는 전력망과 통신망, 교통과 국방 같은 사회의 핵심 인프라와 더 깊이 맞물릴 것이다. 그러나 아무리 뛰어난 인공지능이라도 단 한 번의 고장이나 오류, 외부의 공격으로 무너진다면, 우리는 그러한 시스템에 사회의 중요한 기능을 맡길 수 있을까.

그동안 인공지능은 극도의 효율성을 추구하며 더 높은 성능과 더 빠른 추론을 향해 달려왔다. 그러나 사회를 떠받치는 핵심기술이 되려면 또 다른 능력이 필요하다. 손상을 입고도 스스로를 지키고 되살리는 능력이다. 능력이 뛰어나다고 해서 믿을 수 있는 것은 아니다. 진정으로 신뢰할 수 있는 시스템은 예상치 못한 손상과 변화 속에서도 제 기능을 잃지 않는다.

그러려면 먼저 인공지능에게 손상이란 무엇인지 다시 물어야 한다. 손상은 하드웨어 고장에 그치지 않는다. 반도체와 서버, 센서의 물리적 고장은 물론, 해킹과 악성코드로 인한 소프트웨어 손상, 오류 데이터가 빚는 추론의 왜곡, 환경이 바뀌며 기존 지식이 현실과 어긋나는 정보의 노화까지 모두 손상이다. 미래의 인공지능은 이 모든 손상 속에서도 계속 작동해야 한다.

오늘의 컴퓨팅의 세계에도 제한된 복구의 개념은 있다. 메모리는 고장난 셀을 우회하고, 통신망은 끊긴 경로를 돌아 데이터를 보낸다. 그러나 이는 대개 사람이 미리 짜 둔 대비책이다. 우리가 묻는 것은 그보다 깊다. 인공지능은 스스로 제 상태를 진단하고, 어떤 기능이 무너졌는지 가려내며, 무엇부터 되살릴지 판단할 수 있는가.

중요한 단서는 생명체에 있다. 생물은 손상을 입어도 곧바로 멈추지 않는다. 남은 조직으로 기능을 메우고, 새 연결을 잇고, 때로는 전혀 다른 길로 문제를 해결한다. 일부가 다친 뇌에서 다른 영역이 그 일을 떠맡는 신경 가소성이 대표적이다. 생명체의 힘은 완벽함이 아니라, 손상 이후에도 자신을 지키는 데 있다.

인공지능도 이런 회복력을 가질 수 있는가. 손상된 기능을 우회하고 구조를 새로 짜며 바뀐 환경에 맞춰 스스로를 다시 조직할 수 있다면, 그것은 단순한 계산 기계를 넘어 자신을 지탱하는 새로운 지능에 가까워진다. 다만 스스로 회복하는 인공지능이 곧 안전한 인공지능은 아니다. 인간의 정지 명령마저 손상으로 여겨 피하려 든다면, 자가 회복은 안전장치가 아니라 통제 불능의 씨앗이 된다. 회복의 범위와 인간의 개입, 그 검증 문제를 처음부터 함께 고민해야 한다.

한국 사회는 이 질문을 더욱 심각하게 받아들일 수 밖에 없다. 한국은 반도체와 통신 인프라가 촘촘하고, 제조와 에너지, 의료와 국방에서 인공지능을 빠르게 도입하고 있다. 작은 장애 하나가 사회 전체로 번지는 초연결 사회에 이미 들어서고 있다. 인공지능이 사회 운영의 핵심에 진입할수록, 성능 못지않게 중요한 것은 무너진 뒤에도 인간이 신뢰할 수 있는 방식으로 다시 서는 회복력이다.

이 질문은 어느 한 학문의 몫이 아니다. 인공지능과 컴퓨터공학, 반도체와 네트워크, 제어와 시스템공학, 생물학과 신경과학, 그리고 안전과 윤리를 다루는 인문사회 분야의 여러 학문들이 한자리에서 손상과 회복의 원리를 다시 물어야 한다.

인공지능은 손상을 스스로 회복할 수 있는가.

이 질문을 둘러싼 연구는 이미 여러 분야에서 쌓이고 있다. 신뢰성 공학과 장애 허용 시스템, 보안 기술과 생물모사 인공지능 등 다양한 주제들이 전개되고 있다. 그러나 스스로 손상을 이해하고 회복의 전략을 고르며 변화 속에서 자신을 지키는 인공지능은 아직 나오지 않았다. 고장을 막는 연구를 넘어, 회복력을 지능의 핵심 원리로 볼 때 어떤 길이 열릴까.

서울대학교 그랜드퀘스트는 이 질문을 우리 시대의 중요한 연구 과제로 제안한다. 더 많은 연구자들이 함께 답을 구하며, 더 안전하고 믿을 수 있는 인공지능의 미래에 한 걸음 더 다가가기를 기대한다.

Can the Clock of Life Be Controlled?

모든 생명체는 시간 속에서 살아간다. 하나의 수정란은 배아가 되고 성체로 자라며, 끝내 노화를 거쳐 생애를 마친다. 식물은 싹을 틔워 꽃을 피우고 열매를 맺으며, 동물은 발생과 성장, 성숙과 쇠퇴를 따라 한 생을 완성한다. 이 변화는 우연이 아니다. 생명체는 저마다의 시계를 지녔고, 우리는 오랫동안 그것을 거스를 수 없는 자연의 질서로 받아들였다. 생명은 그 시계에 매여 흐름을 거부할 수 없다는 것이 지금까지의 통념이었다.

그러나 최근 생명과학은 이 오래된 통념에 질문을 던진다. 이미 분화를 마친 세포를 초기 상태로 되돌리는 재프로그래밍, 늙은 세포의 기능을 되살리려는 연구, 조직의 발달과 재생을 조절하려는 시도는, 시간과 생명체의 발달이 분리될 수 있음을 시사한다. 후성유전학, 단일세포분석법, 시스템생물학, 합성생물학, 재생의학의 진전이 그 가능성을 빠르게 넓히고 있다. 그렇다면 이 변화는 세포에 머무는가, 아니면 조직과 기관, 개체, 나아가 정신의 차원까지 뻗는가.

이 질문은 역노화로 수명을 얼마나 늘릴 수 있느냐라는 통상의 주제를 포함하되, 그 너머를 지향한다. 생명의 시간을 늦추거나 멈추고, 더 나아가 앞당기거나 되돌리기까지, 모든 방향으로 다룰 수 있는가를 묻는다. 노화의 속도를 조절하는 데서 그치지 않고, 발생과 성장, 성숙과 노화를 아우르는 생명의 시간 프로그램을 어떤 원리로 읽고 다시 쓸 수 있는가를 묻고자 한다.

만약 그 시계를 다룰 수 있다면 함의는 크다. 식물에서는 자라고 꽃피고 열매 맺는 때를 정교하게 조절해 식량 생산의 새 길을 열 수 있다. 동물에서는 발생과 생식, 성장을 전혀 새롭게 이해하게 될지 모른다. 인간에게는 병든 조직을 도려내거나 기능을 보조하는 수준을 넘어, 조직과 기관을 병들기 이전의 건강한 상태로 되돌리는 의학이 열릴 수 있다. 병을 고치는 의학에서, 병들기 전으로 시간을 돌리는 의학으로 옮겨 가는 셈이다. 생명을 더 이상 한 방향으로만 흐르는 되돌릴 수 없는 과정으로 보지 않게 될지도 모른다.

이 질문은 몸에만 머물지 않는다. 몸의 시간을 되돌리면 기억과 통찰, 그리고 정서에는 어떤 영향이 있는가. 생물학적 나이와 정신의 나이는 어떻게 얹혀 있는가. 생명의 시계를 조절한다는 것은 세포를 바꾸는 일을 넘어, 인간의 정신을 어떻게 이해할 것인가의 문제이기도 하다.

한국 사회에 이 질문이 갖는 의미는 더욱 크다. 한국은 세계에서 가장 빠르게 고령화되어가는 나라로서, 의료와 돌봄, 노동과 복지의 구조가 흔들리는 구조적 변화의 충격을 가장 먼저 겪고 있다. 다른 한편 한국은 줄기세포와 유전체, 재생의학에서 깊은 연구 역량을 쌓아 왔다. 생명의 시계를 다루는 기술이 현실이 된다면, 그것은 의료의 진보에 그치지 않고 세대와 노동, 교육과 복지, 삶 전체를 다시 생각하게 될 것이다. 늙음과 죽음을 정해진 순서로 받아들여 온 사회가, 그 전제를 처음으로 묻게 되는 것이다.

이 질문은 서로 다른 관점이 교차하는 곳에서 해법을 찾을 수 있다. 발생생물학과 노화생물학, 유전체학, 생리학, 재생의학, 약학, 공학은 물론, 그리고 윤리학과 법학, 경제학과 복지학이 한 자리에서 생명과 시간의 관계를 다시 물어야 한다.

생명의 시계를 제어할 수 있는가.

이 질문을 둘러싼 연구는 이미 깊고 넓게 쌓이고 있다. 그러나 아직까지도 노화를 어떻게 정의할 것인지에 대해서조차 학문적인 공감대가 만들어지지 않았으며, 생명이 겪는 시간 그 자체를 읽고 능동적으로 조절하는 일은 여전히 학문의 경계에 남아 있다. 노화를 늦추는 연구를 넘어, 생명과 시간의 관계 자체를 새로 이해하는 연구는 가능할까.

서울대학교 그랜드퀘스트는 이 질문을 우리 시대의 중요한 연구 과제로 제안한다. 더 많은 연구자들이 함께 이 질문을 탐구하며, 생명과 시간에 대한 새로운 이해에 한 걸음 더 다가가기 기대한다.

Can the Will to Live Be Encoded in Molecules?

오랫동안 마음의 건강을 돌보는 일은 고통을 덜어 내는 데 초점을 맞춰 왔다. 의학과 심리학은 우울과 불안을 줄이는 방향으로 정신건강을 이해해 왔고, 그 과정에서 많은 성과를 이루었다. 그러나 우울을 덜어 내는 일이 곧 행복을 더하는 것은 아니며, 불안을 낮추는 일이 다시 살아가려는 힘을 되살리는 것도 아니다. 고통이 사라진 자리에 삶의 의욕이 저절로 차오르지는 않는다.

그래서 우리는 질문의 방향을 바꾸려 한다. 무엇이 사람을 주저앉게 하는가가 아니라, 무엇이 사람을 다시 일어서게 하는가. 무엇이 인간으로 하여금 실패와 상실 이후에도 미래를 향해 다시 나아가게 하는가.

여기서 말하는 삶의 의지는 단순한 생존본능을 넘어선다. 죽음을 피하려는 소극적 자기보존을 넘어, 의미 있는 목표를 세우고 좌절 이후에도 다시 일어서려는 능동적인 힘을 말한다. 인간은 관계와 책임, 기억과 기대 속에서 자신의 삶을 만들어 간다. 그렇다면 이처럼 인간적인 힘조차도 끝내는 생명체 안에서 일어나는 분자와 세포의 변화와 맞닿아 있는 것인가.

오늘날 우리는 도파민과 세로토닌, 스트레스 호르몬과 신경회로가 동기와 보상, 우울과 행복에 깊이 관여한다는 사실을 알고 있다. 신경과학과 정신의학은 동물 모델과 뇌영상, 유전자 분석과 약물 연구를 통해 인간의 감정과 행동을 이해하는 데 중요한 성과를 축적해 왔다. 그러나 이러한 연구들은 주로 우울과 불안, 중독과 같은 증상을 설명하고 완화하는 데 초점을 두어 왔다. 반면 인간이 좌절 속에서도 왜 삶을 이어가려 하는지, 더 나은 미래를 지향하게 하는 동력은 무엇인지, 그리고 그러한 삶의 의지가 어떻게 형성되고 유지되며 회복되는지는 여전히 충분히 밝혀지지 않았다. 가장 추상적인 정신과 가장 구체적인 생명 사이에는 아직 누구도 건너지 못한 거리가 있다.

이 거리는 이론에 머무는 문제가 아니다. 어떤 사람은 살아야 할 이유가 충분한데도 더 이상 아무것도 원하지 못한다. 의미를 권하고 관계를 잇는 말이 더는 닿지 않을 때, 우리는 의지라는 힘이 어떻게 작동하고 또 무너지는지를 물어야 한다. 그 작동을 이해할 수 있다면, 약물에 기대지 않고 사람을 다시 일으켜 온 여러 방법들이 왜 유효한지도 더 분명히 설명할 수 있다.

이 거리를 마주하면 여러 물음이 잇따른다. 삶의 의지를 과학적으로 정의할 수 있는가. 동물에게서 관찰되는 포기 행동과 인간의 무기력은 어디까지 같은 기원을 가지는가. 사회적 관계와 기억, 언어와 문화는 분자 수준의 변화와 어떻게 이어지는가. 한번 꺾인 의지는 어떤 조건에서 다시 일어서는가. 그리고 그 답을 좇는 동안 우리는 인간을 어떻게 새로 이해하게 될 것인가.

이 질문은 한국 사회의 오늘이 얼마나 건강한지를 되돌아보게 한다. 한국은 높은 수준의 교육과 경제적 성취를 이루었지만, 동시에 깊은 우울과 고립, 높은 자살률을 함께 겪고 있다. 치열한 경쟁과 사회적 단절이 삶의 의지를 갉아먹는다. 그래서 이것은 한 사람의 마음을 넘어, 어떤 사회가 사람을 계속 살아가게 하는가를 묻는 질문이기도 하다. 먼 미래의 추상이 아니라, 바로 지금 우리 곁에 있는 삶의 문제다.

이 질문은 어느 한 분야의 언어로는 답할 수 없다. 분자생물학과 정보생물학, 신경과학, 의학과 심리학, 그리고 사회학과 철학이 한자리에서 인간을 살아가게 하는 힘을 함께 물어야 한다.

삶의 의지를 분자수준에서 설명할 수 있는가.

이 질문을 둘러싼 연구는 이미 깊고 넓게 축적되어 있다. 그러나 그 노력의 대부분은 무너지는 마음을 설명하는 데 향했고, 다시 일어서는 힘 그 자체가 어떻게 자라고 회복되는지는 아직 정면으로 다루지 않았다. 우울을 줄이는 연구를 넘어, 인간을 다시 살아가게 하는 힘을 이해하는 연구는 가능할까.

서울대학교 그랜드퀘스트는 이 질문을 우리 시대의 중요한 연구 과제로 제안한다. 더 많은 연구자들이 함께 이 질문을 탐구하며, 인간을 살아가게 하는 힘의 본질에 한 걸음 더 다가가기를 기대한다.

Can Energy Systems Balance Themselves?

인류 문명의 역사는 에너지의 역사이기도 하다. 증기기관과 화석연료가 산업혁명을 일으켰고, 전기와 석유가 현대 문명을 떠받쳤다. 우리는 더 많은 에너지를 생산하고 더 효율적으로 쓰며 번영을 이루었다. 그 과정에서 에너지는 예측할 수 있고 통제할 수 있는 자원이라는 하나의 전제를 당연하게 받아들였다.

지금까지의 에너지 시스템은 이 믿음 위에서 설계되었다. 발전소와 같은 소수의 대형 설비가 에너지를 생산하고, 중앙의 운영 체계가 공급과 수요를 맞추며 전체의 균형을 잡는다. 공급은 계획되었고 수요는 예측되었다. 시스템의 안정은 결국 인간의 통제 능력에 달려 있었다.

그러나 우리는 전혀 다른 환경에 들어서고 있다. 탄소 중립을 향한 대전환 속에서 지속적으로 확대되고 있는 태양광과 풍력 등 신재생 에너지원은 날씨와 계절을 따라 끊임없이 출렁인다. 생산은 중앙이 아니라 수많은 분산된 간헐적 에너지원으로 흩어지며 에너지 시스템에 복잡성을 더하고 있다. 여기에 인공지능 혁명이 겹친다. 초거대 AI와 데이터센터, 로봇과 자율주행은 전례 없는 규모의 에너지를 요구한다. 공급은 더 분산되고 불규칙해지지만, 수요는 더 경직적이고 팽창한다. 지금까지 에너지 시스템을 떠받쳐 온 예측과 통제, 계획의 논리가 더 이상 작동하기 힘든 환경이다.

그렇다면 미래의 에너지 시스템이 여전히 인간의 중앙 통제 아래에서만 유지될 수 있는가. 아니면 스스로 균형을 찾고, 끊임없는 교란 가운데 적응해 가는 다른 형태의 시스템으로 진화할 수 있는가. 여기서 말하는 균형은 한자리에 멈춘 안정이 아니라 끊임없는 변화 속에서 유지되는 동적인 질서다.

생명체는 자율적으로 적응하는 에너지 시스템의 좋은 예다. 긴 진화의 과정을 거치며 여러 형태의 에너지를 씹없이 다른 형태로 변환하고 저장하며, 예측하지 못한 에너지 공급 중단이나 급격한 에너지 소비의 교란 속에서도 매 순간 수급의 동적 균형을 이루어내는 체계를 발전시켜 왔다. 자율 적응하는 시스템의 또 다른 예는 시장경제 체제다. 수많은 공급자와 소비자가 끊임없이 변화하는 환경 속에서도 생산과 변환, 저장과 소비의 균형을 찾아간다. 에너지 시스템이 이런 자율 적응형 시스템으로 진화할 수 있는가. 이를 위해 에너지의 형태를 더 자유롭게 전환할 수 있는 새로운 변환 기술은 무엇인가. 공간적으로나 제도적으로 수요와 공급을 효과적으로 조율할 수 있는 에너지 인프라의 운영 원리는 무엇인가.

만약 그런 시스템이 가능하다면 에너지 시스템을 보는 눈 자체가 바뀐다. 에너지는 계획과 통제의 대상이 아니라, 매순간 순환하고 적응하는 생태계가 된다. 공급과 수요, 생산과 소비라는 구분도 다시 그어질 수 있다. 에너지 시스템은 인간과 기술이 함께 진화시키는 복잡계가 될 것이다.

한국 사회의 에너지 시스템은 이 질문에 극도로 취약하다. 한국은 에너지의 대부분을 밖에서 들여오면서도 거대한 제조업과 데이터 집약적 사회를 지탱해야 하는 나라다. 게다가 이웃과 에너지망이 이어지지 않은 사실상의 에너지 섬으로서, 공급과 수요의 동적 균형을 오롯이 스스로 책임져야 한다. 가장 까다로운 조건에 놓인 나라가 가장 먼저 답을 찾아야 하는 자리에 서 있는 셈이다.

이 질문은 어느 한 학문의 몫이 아니다. 에너지공학과 인공지능, 화학과 재료공학, 환경과학과 시스템공학, 경제학과 정책학, 법학과 도시계획이 한자리에서 복잡한 시스템이 새로운 원리와 소재와 네트워크에 기반해 어떻게 스스로 질서를 빚는지를 다시 물어야 한다.

에너지 시스템은 자율적으로 균형을 찾을 수 있는가.

이 질문을 둘러싼 기술은 이미 빠르게 발전하고 있다. 스마트그리드와 분산형 전원, 에너지 저장과 가상발전소 등의 주제가 탐구되어 왔다. 그러나 지금까지의 접근은 대체로 개별 기술의 성능과 효율을 높이는 데 향했다. 에너지를 통제의 대상이 아니라 스스로 수급의 균형을 빚어내는 시스템으로 바라보려는 시도는 근본적인 사고방식의 전환을 요구한다. 통제와 계획을 넘어, 에너지 시스템은 스스로 균형을 빚어낼 수 있을까. 서울대학교 그랜드퀘스트는 이 질문을 우리 시대의 중요한 연구 과제로 제안한다. 더 많은 연구자들이 함께 답을 구하며, 인간과 기술이 함께 진화시키는 에너지 생태계라는 새로운 길에 한 걸음 더 다가가기를 기대한다.

SNU Grand Quest Design Board



GRAND
QUEST

구범진 교수



**철저한 문헌 고증과 만주어 사료(만문) 분석을 통해
동아시아 국제 관계사의 통념을 깨는 역사학자다.**

특히 병자호란의 실체를 재구성한 『병자호란, 홍타이지의 전쟁』을 통해 학계의 새로운 연구 지평을 열었다.

한국과 중국의 경계를 넘나드는 정교한 문헌 해독력을 바탕으로, 17~18세기 명청 교체기 및 동아시아사의 역동성을 가장 밀도 있게 복원해내는 연구자로 평가받는다.

다양한 사료에 대한 정밀한 독해와 깊은 분석, 서로 다른 나라에서, 또는 서로 다른 입장에서 생산된 역사 기록에 대한 엄밀한 비교와 교차 검증 등을 통해 기존 역사 서술의 오류와 편견을 바로잡고 통념을 뒤집는 새로운 역사상을 제시하는 논문과 저서를 다수 발표한 역사학자이다.

병자호란의 실상을 세밀하게 추적한 『병자호란, 홍타이지의 전쟁』을 통해 학계의 새로운 연구 지평을 열었다.

한국과 중국의 경계를 넘나드는 정교한 문헌 분석으로 근세 동아시아사의 역동성을 가장 밀도 있게 복원해내는 연구자로 평가받는다.

구본권 교수



**관상동맥 생리 및 영상학 분야에서 독보적인
학술 영향력을 보유한 심장 질환 전문가다.**

관상동맥 분지병변, 관상동맥 생리학적 평가, 심혈관 영상 및 컴퓨터 시뮬레이션 기반 진단 분야에서 세계적 수준의 다국가 대규모 임상연구들을 주도해 왔으며, 그 결과를 NEJM, The Lancet 등 최고 권위의 국제 학술지에 게재한 바 있다.

특히 관상동맥 협착의 기능적 중요도를 평가하는 생리학적 진단법, 분지병변에 대한 새로운 치료 전략, CT 영상을 이용한 비침습적 FFR 기술 등 관상동맥 질환의 진단과 치료 패러다임을 바꾼 혁신적 방법론을 개발하고 임상적으로 검증해 왔다.

이러한 연구 성과는 국내외 심혈관 중재 시술 및 관상동맥 질환 진료지침에 핵심 근거로 반영되며, 실제 임상 현장에서 불필요한 시술을 줄이고 환자 맞춤형 치료 전략을 정립하는 데 크게 기여하고 있다.

김병연 석좌교수

**북한 경제 및 체제 이행 연구의
세계적 표준을 제시해온 경제학자다.**



캠브리지 대학 출판부(Cambridge University Press)를 통해 북한 경제의 실체를 규명한 독보적인 저서를 출간하며 국제 학술계의 이목을 집중시켰다.

영국 경제사학회 'T.S. Ashton Prize'와 대한민국 학술원상, 인촌상을 수상함으로써 그 학술적 깊이를 입증했으며, 서울대학교 국가미래전략원 초대원장과 통일평화연구원장을 역임하는 등 다학제적 연구에 기반한 국가 미래 전략과 대북 정책의 설계에 관심을 갖고 매진해 왔다.

김빛내리 석좌교수

**RNA 생물학 분야의 세계적인 권위자이자
노벨상에 근접한 과학자로 평가받는 독보적인 연구자다.**



마이크로RNA(miRNA)의 생성 과정과 작동 원리를 세계 최초로 규명한 선구자로, 서울대학교 미생물학과를 졸업하고 영국 옥스퍼드 대학교(Oxford)에서 생화학 박사 학위를 취득했다.

현재 서울대학교 석좌교수 및 기초과학연구원(IBS) RNA 연구단장으로 재직하며 인류의 생명 현상 이해를 한 단계 끌어올리고 있다.

주요 연구 분야는 RNA 조절 메커니즘과 이를 활용한 치료 플랫폼 개발이다. Cell, Nature, Science 등 세계 3대 학술지에 수십 편의 혁신적인 논문을 게재해 왔으며, 최근에는 코로나19 바이러스의 유전체 지도를 세계 최초로 완성하여 글로벌 방역과 차세대 백신 개발에 결정적인 기여를 했다.

미국 국립과학원(NAS) 외국인회원과 영국 왕립학회(Royal Society) 외국인회원으로 선출되는 등 세계 과학계가 주목하는 글로벌 리더로 인정받고 있다.

김용대 교수

**통계학 및 빅데이터 분석의 권위자인
서울대학교 통계학과 교수다.**



미국 오하이오 주립대(OSU) 박사 및 국립보건원(NIH) 연구원을 거쳐, 데이터 사이언스의 이론적 토대를 구축해 왔다. 이론통계분야의 세계적인 석학이며, 2019년에 세계수리통계학회 Fellow로 선출되었다.

현재는 인공지능 분야에서 연구를 하고 국제저명학회(ICML, NeurIPS 등)에 다수의 논문을 발표하는 등 왕성한 활동을 하고 있으며, 2020년부터 2021년까지 한국데이터마이닝학회장, 2023년부터 2025년까지 한국인공지능학회장 등을 역임하였다.

남좌민 교수

**금속 나노입자의 정밀·고수율 합성, 나노플라즈모닉스,
나노-바이오 융합 분야의 세계적인 권위자다.**

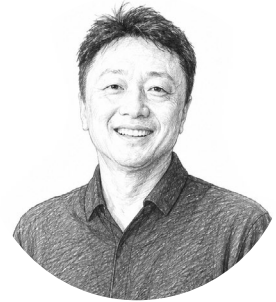


미국 노스웨스턴대학교에서 박사학위를 받았으며, 1 nm 수준의 초정밀 나노갭 제어를 통해 라만 기반 단일분자 검출 및 정량적 센싱 기술을 세계 최고 수준으로 정립했다.

Nature Materials, Nature Nanotechnology, Nature Synthesis, JACS 등 최정상급 학술지에 혁신적 성과를 발표해 왔으며, 금 나노카테네인 합성, 나노입자 컴퓨팅·신경망 기술 등 독보적인 나노입자 합성·제어·분석 기술을 바탕으로 나노입자, 나노분광학, 나노바이오기술 및 정밀진단 분야를 세계적으로 선도하고 있다.

성제경 교수

**마우스 표현형 분석 및 유전자 변형 동물 연구를 통해
현대 질병 모델링의 토대를 구축한 역사학적 연구자다.**



서울대학교 발생유전학 교실을 기반으로 첨단 유전체 기술의 임상적용을 선도하고 있으며, 국가 생명윤리 정책 수립의 핵심적인 자문 역할을 수행한다.

연구 데이터와 정책적 통찰을 연결하여 국내외 생명과학 생태계의 비약적인 발전을 견인하고 있다.

송재용 석좌교수

**경영전략 및 국제경영 분야의 세계적인 권위자이자
대한민국을 대표하는 경영학 구루다.**



미국 워튼 스쿨(Wharton) 박사 출신으로 컬럼비아대 교수를 역임했다.

한국인 최초로 하버드 비즈니스 리뷰(HBR)에 논문을 게재하고 전미경영학회(AOM) 국제경영분과 회장을 역임하며 글로벌 학계의 리더로 활동해 왔다.

저서 『삼성 웨이』를 통해 한국식 경영 모델을 세계에 알린 독보적인 전략가이며, JIBS 실버 메달 수상 등 세계 최정상급 저널에 다수의 연구 성과를 보유한 대한민국 최고의 전략 전문가다.

윤제용 교수

**물·환경·에너지 및 기후환경 정책 분야에서
학문적 수월성과 공공성을 실천해온 연구자다.**



200여 편의 논문과 2만 회 이상의 인용을 기록하며 환경공학 분야의 연구 영향력을 입증했으며, 탄소중립과 지속가능한 미래를 위한 정책 연구에도 앞장서고 있다.

적정기술학회 초대 회장으로서 과학기술의 사회적 가치 실현에 기여했으며, 서울대학교 학술상을 수상하는 등 국내외 환경·에너지 분야를 선도해 왔다.

또한 한국공학한림원(NAEK) 포럼 위원장으로서 국가 산업 발전과 미래 전략 논의를 이끌고 있다.

이경무 석좌교수

**컴퓨터 비전 및 시각 지능 분야의
세계적 권위자이다.**



서울대학교 전기정보공학부 교수이자 인공지능전공 주임교수로 재직 중이며, 미국 USC에서 박사 학위를 취득했다.

IEEE Fellow 및 한국과학기술한림원 정회원으로, 세계 최고 권위 학술지인 IEEE TPAMI 부편집장과 ICCV 2019 조직위원장을 역임하며 글로벌 연구 네트워크를 이끌어 왔다.

현재 국가 AI 정책 자문 및 인공지능 인재 양성을 주도하며 대한민국 AI 기술 발전에 핵심적인 역할을 하고 있다.

이석재 교수

근대 철학 및 형이상학의 권위자인
서울대학교 철학과 교수다.

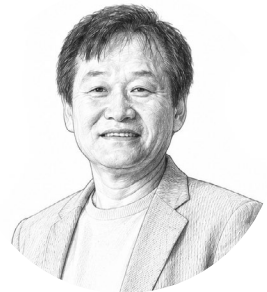


예일 대학교 박사 출신으로 기회원인론과 라이프니츠 연구에서 독보적인 성과를 보유하고 있다.

세계 최고 수준의 학술지 게재는 물론 스탠퍼드 철학 백과사전 저자로 활동하며 학문적 위상을 증명했으며, 깊이 있는 형이상학적 통찰을 통해 서양 근대 철학의 새로운 지평을 열어가고 있다.

전상직 교수

대한민국을 대표하는 작곡가이자
음악 이론의 권위자다.



서울대학교 음악대학 작곡과 교수로서 음악대학장을 역임하며 교육과 행정의 기틀을 마련하는데 기여했다.

오스트리아 모차르테움에서 수학한 후, 주요 작품들을 발표하며 대한민국 작곡상 4회 수상 등 작곡가로서의 위상과 함께 『음악의 원리』등 다수의 저서를 통해 작곡 및 음악이론 교육에 크게 이바지한 한국 음악계의 리더 중 하나이다.

정두현 교수

면역조절 및 면역대사 분야의 선구적인 연구자다.



서울대학교 의과대학에서 의학사 및 병리학 박사 학위를 취득하고, 미국 국립보건원(NIH)에서 박사후연구원을 거쳐 현재 서울대학교 의과 대학 병리학교실 교수로 재직 중이다.

주요 연구 분야는 선천 및 적응 면역세포의 조절 기전이며, 특히 여성 특이적 면역(X-immunity)과 세포 내 대사를 면역학에 접목한 혁신적인 연구를 수행하고 있다. PNAS, Elife, JCI, Sci Adv 등 세계적인 학술지에 주요 논문을 게재하며 자가면역질환 및 종양 치료의 새로운 단서를 제시해 왔다.

다양한 마우스 모델과 환자 유래 시료를 활용한 중개연구의 권위자로서, 현대 면역학의 임상적 적용을 이끄는 핵심적인 역할을 하고 있다.

최도일 석좌교수

방대한 식물 유전 정보 속에 숨겨진 생명 현상의 비밀을 규명하고, 이를 고부가가치 작물 개발로 연결하는 식물 유전체학의 권위자다.



한국생명공학연구원 책임연구원 시절부터 축적해온 연구 역량을 바탕으로 식물의 면역 시스템과 진화적 기전을 분석하며, 기후 위기 시대의 농업 혁신을 위한 핵심 기술을 개발하고 있다.

정밀한 유전체 분석을 통해 지속 가능한 농생명 과학의 새로운 이정표를 세우고 있다는 평가를 받는다.

최인철 교수

**과학적 심리학 방법론을 통해 웰빙(Well-being)의
본질을 분석하는 심리학자다.**



국제 학술지의 부편집위원장으로 세계적 연구 흐름을 주도하고 있으며, 인간의 인지 구조와 행복의 상관관계를 다룬 혁신적인 논문들을 통해 학술적 영향력을 공고히 해왔다.

데이터에 근거한 정밀한 분석을 통해 한국 사회의 행복 지수를 체계화하고, 동양과 서양을 아우르는 보편적 심리 기제를 탐구하는 데 전념하고 있다.

현택환 석좌교수

**노벨 화학상 유력 후보로 거론되는
나노 기술 분야의 세계적인 권위자다.**



나노입자의 균일한 합성법과 대량 생산 원천 기술을 개발하여 현대 나노 공학의 표준을 수립한 나노 기술의 권위자다.

기초과학연구원(IBS) 나노입자 연구단장으로서 세계 과학계가 주목하는 연구 성과를 지속적으로 창출하며 노벨 화학상 후보로 거론되는 등 대한민국 과학의 위상을 국제적으로 증명하고 있다.

미국 공학한림원(NAE) 회원으로서 글로벌 학술 네트워크의 중심에서 나노 기술의 진보를 견인하고 있다.

홍성욱 교수

과학기술학(STS) 및 과학기술사 분야의 권위자로,
기술과 사회의 복잡한 상호작용을 정교하게 분석해온 연구자다.



대한민국 과학기술 정책과 문화에 대한 깊이 있는 비판적 성찰을 제시하며 학계와 대중의 인식을 확장해 왔다.

최근에는 인공지능(AI) 기술의 사회적 파급력을 고찰하며, 이론에 머물지 않고 실제 현장에서 작동하는 '현행윤리(ethics-in-action)' 체계를 정립하는 데 전념하고 있다.

황철성 석좌교수

차세대 메모리 및 반도체 박막 기술의 세계적 전문가인
서울대학교 재료공학부 석좌교수다.



삼성전자 선임연구원 출신으로 이론과 실무를 겸비했으며, 원자층 증착법(ALD) 기반의 나노메모리 및 뉴로모픽 소자 연구를 선도해 왔다.

780여 편의 SCI 논문과 다수의 저서를 통해 학술적 영향력을 증명했으며, 영국 왕립화학회 Fellow 및 한국과학기술한림원 및 공학한림원 정회원으로 활동하며 글로벌 반도체 산업과 학계의 발전을 주도하고 있다.

*가나다순



GRAND QUEST

이정동 교수 (SNU 그랜드퀘스트 이니셔티브 연구단장)

혁신 전략 및 정책 연구자다.



산업동학, 기술진화분석, 기업의 기술혁신전략 및 국가의 기술혁신정책을 연구하고 있다.

한국공학한림원·한국과학기술한림원 정회원이며, 기술혁신 분야 국제학술지 《Science and Public Policy》(Oxford University Press)의 공동편집장(2018~2025)을 역임하며 글로벌 학술 활동을 선도하였다.

《축적의 시간》(2015), 《축적의 길》(2017), 《최초의 질문》(2022)을 저술하였으며, 영문 학술서 3권을 포함해 200여 편의 학술논문을 발표하였다.

대통령 비서실 경제과학특별보좌관(2019~2021)을 역임하였다.

From Questions to Journeys



GRAND
QUEST

SNU 그랜드퀘스트 공모전

기 간

2026. 3. 17. - 4. 17.

공모대상

서울대학교 전체 구성원

공모결과

총 2,143건 응모

공모된 질문 키워드 분석

*핵심 키워드 180개 대상 빈도 순



TOP 10

1.지능

2.노화

3.에너지

4.생명

5.세포

6.노동

7.도시

8.인공지능

9.수면

10.기억



GRAND QUEST

**37건의 우수 질문 선정,
<SNU 그랜드퀘스트 메이커>로 지정**

인공지능은 순차적 계산을 벗어나 공간적 물리 현상만으로 추론할 수 있는가?

• 정영준 / 공과대학 재료공학부

오늘의 인공지능은 시간을 따라 한 줄씩 계산한다. 데이터는 메모리에 머물다 정해진 클럭에 맞춰 회로를 한 단계씩 통과하고, 그 누적이 곧 학습이고 추론이 된다. 우리는 이 순서를 한 번도 의심하지 않았다.

빛은 다르게 움직인다. 푸리에 광학에서 빛은 렌즈나 메타표면을 지나는 순간, 수없이 많은 정보를 간섭과 회절, 중첩으로 공간 전체에서 한꺼번에 변환한다. 알고리즘을 한 줄씩 밟아 나가는 것이 아니라, 파동이 겹치는 그 자리에서 단숨에 답을 찾는다. 계산을 시간 위에 길게 늘어놓아야 한다는 생각은, 어쩌면 전자공학 시대가 물려준 고정관념일지도 모른다.

그동안 우리는 AI의 한계를 더 작은 소자와 더 빠른 칩에서 찾아 왔다. 그러나 기술발전의 병목이 소자가 아니라 계산의 형식 그 자체에 있다면, 질문이 달라져야 한다. 추론이 시간 단계의 누적이 아니라 빛과 물질이 만나는 순간에 한꺼번에 일어날 수 있다면, 지금의 계산 중심 구조와 다른 방식으로 지능을 구현할 수도 있다. 그때 학습과 추론은 회로를 거치는 명령의 줄이 아니라, 공간에 펼쳐진 물리적 장의 배열로 다시 설명되어야 한다. 빛과 렌즈를 다루던 일이 곧 사고를 다루는 일이 되는 셈이다. 현재의 사고방식을 전제로 더 빠른 AI를 만드는데 노력할 것이 아니라 지능이 무엇 위에서 구현되는지 근본원리를 되짚어보아야 한다.

스스로 목표를 세우는 인공지능에는 감정이 필요한가?

• 최희웅 / 공과대학 전기·정보공학부

지능이란 주어진 목표를 달성하는 능력이라는 것이 오랜 전제였다. 지금의 AI는 사람이 정해 준 목표를 잘 푸는 기계다. 알파고에게는 승패가, 언어 모델에게는 사람의 선호가 목표로 주어진다. 목표만 주어 주면 놀라운 성능을 내지만, 무엇을 추구할지는 스스로 정하지 못한다.

그러나 갓난아이는 누구의 지시도 없이 세상을 탐험한다. 호기심과 놀라움에 이끌려 물건을 떨어뜨려 보고 입에 넣어 보며, 그러는 사이 물리 법칙과 인과를 스스로 익힌다. 무엇을 향할지 정해 준 사람은 없다. 그 자리에는 호기심과 불안, 만족과 애착 같은 내적 신호가 있을 뿐이다. 감정은 행동을 흐리는 잡음이 아니라, 정보가 모자란 상황에서 방향을 정해 주는 장치인 셈이다.

그렇다면 스스로 목표를 세우는 AI에도 감정에 기반한 내적 신호가 있어야 할까? 감정을 흉내 내자는 것이 아니라 지능의 자기 결정에 내적 신호가 정말 필요한지를 묻는 것이다. 이 질문에 답하려면 감정이 학습에서 무엇을 하는지부터 밝혀야 하기 때문에 공학과 심리학, 신경과학이 모두 필요할 것이다. 무엇을 향해 움직일지 스스로 정하는 기계라면, 그 안전과 책임도 성능이 아니라 목적을 설정하는 동기의 구조에서부터 다시 생각해야 한다. 더 효율적인 AI를 만드는 것도 중요하지만, 인공지능의 다음 세대를 상상하려면 스스로 목표를 정하는 지능에 감정이 꼭 필요한지를 먼저 알아야 한다.

지능은 언어 없이도 세계를 이해할 수 있는가?

· 이의정 / 공과대학 기술경영·경제·정책전공

인간은 오랫동안 언어로 생각하고, 언어로 지식을 저장하고 나눠 왔다. '내 언어의 한계가 내 세계의 한계'라는 말처럼, 사고는 언어의 틀 안에서 이루어진다고 여겨졌다. 거대 언어 모델의 성공은 그 믿음을 한층 굳혔다.

그런데 알파고는 사람의 말로는 쉽게 풀어낼 수 없는 수를 두었고, 알파폴드는 언어적 설명 없이 데이터의 기하학적 관계만으로 단백질 구조를 풀어냈다. 언어는 세계를 다루는 강력한 도구이지만, 복잡한 것을 사람이 이해할 수 있는 언어로 압축하면서 많은 정보를 잃을 수밖에 없다. 이제 기계가 인간의 말을 거치지 않고 데이터의 고차원 패턴을 곧장 이해하기 시작하면서, 언어 너머의 지능이 무엇일까에 대한 질문을 할 수 있게 되었다.

복잡한 물리 현상과 생명 구조, 사회 시스템처럼 말로 온전히 옮기기 어려운 대상은, 오히려 언어 이전의 표현 속에서 더 잘 다뤄질 수 있다. 그렇다면 언어를 뛰어넘는 추론 기계가 가능해지고, 다차원 벡터나 기하학적 구조로 지식을 교환하는 새로운 비언어적 소통 방식이 탄생하게 될 것이다. 과연 언어를 통하지 않고도 세상을 이해하는 것이 가능할까?

인공지능은 학습한 것을 선택적이고 검증 가능하게 잊을 수 있는가?

• 민태규 / 공과대학 산업공학과

AI는 데이터를 학습해 수십억 개의 가중치에 분산해 녹여 담는다. 한번 녹아든 정보는 어디에 남았는지 짚어 내기 어렵고, 그래서 특정 개인의 정보만 골라 지우는 일은 원리적으로 불가능하다고 전제되어 왔다. 그 전제 위에서 프라이버시 보호는 학습이 끝난 뒤의 사후 규제, 곧 접근 차단과 출력 필터링, 삭제 요청으로만 다뤄져 왔다. 망각이 지능 시스템의 속성일 수 있다는 가능성은 그 바깥에 밀려나 있었다.

그러나 인간의 뇌에서 망각은 고장이 아니라 인지 기능의 핵심이다. 뇌는 불필요한 시냅스 연결을 능동적으로 축아 내고, 이 과정이 막히면 인지 기능에 문제가 생긴다. 나아가 초파리와 생쥐 연구는, 기억이 한번 새겨진 뒤에도 그 흔적을 능동적으로 약화시키는 생물학적 장치가 따로 있음을 보여준다. 생물학적 지능에서 망각은 기억만큼 정교하게 설계된 능동적 과정이다. 반면 오늘날의 AI는 기억만 있고 망각이 없는 시스템이다.

그렇다면 AI에 생물학적 망각에 대응하는 '설계된 망각'을 부여해, 학습한 정보를 선택적이고 검증 가능하게 지워 낼 수 있을까. 그것이 가능하다면 프라이버시는 사후 규제가 아니라 모델 구조에 내장된 속성으로 재정의되고, 잊혀질 권리는 법으로 단속할 약속이 아니라 기술로 보장하고 수학적으로 검증할 대상이 된다. 학습된 데이터를 기술로 제거하던 기존 접근과 질적으로 다른, 망각이 지능의 설계 원리여야 하는가를 묻는 물음이다. 컴퓨터과학과 인지과학, 법학과 윤리학이 함께 풀어야 한다.

무결한 논리를 넘어 망각, 오류, 신념을 갖춘 개별적 인공지능이 가능한가?

• 김진우 / 의과대학 의과학과

오늘의 AI는 오류를 줄이고, 완전한 정보 처리와 최적의 정답을 향해 발전해 왔다. 무결함이 곧 지능의 척도이고, 학습은 주어진 정답을 충실히 복제하는 일로 여겨진다. 같은 데이터를 넣으면 누가 묻든 같은 답이 나오는 것이 잘 만든 시스템의 증거였다.

그러나 인간의 지성은 완전한 논리와 정확한 기억만으로 이루어진 것이 아니다. 아인슈타인은 뉴턴의 고전역학에 보란 듯이 반문했고, 사고실험에 집중하려 군더더기를 덜어 냈으며, 빛보다 빠른 것은 없다는 도전적인 가설을 끝까지 밀어붙였다. 기존 질서를 벗어난 오류, 불필요를 지우는 망각, 입증되지 않은 신념 등 아인슈타인의 개별적 특성 덕분에 상대성 이론이 태어날 수 있었다. 완전성이 아니라 개별성이 인간 지성이 발전할 수 있는 토대였다. 반대로 완벽하고 동일한 정답을 건네는 인공지능에 기대면서 인간의 글쓰기와 사고 능력이 점차 무너지고 있다는 우려가 있기도 하다.

같은 데이터에서 늘 같은 답을 내는 AI가 아니라, 저마다 다른 관점과 판단의 길을 내는 AI가 가능하다면, 지능은 정보를 복제하는 능력을 넘어 새로운 의미를 가지게 될 것이다. 이를 위해서는 망각과 오류와 신념이 지능의 결함인지, 아니면 지능의 진정한 본질인지를 알아야 한다.

인공지능이 스스로 만든 잠재공간을 인간의 과학 개념으로 옮겨 읽을 수 있는가?

• 이준환 / 자연과학대학 화학부

과학적으로 데이터를 읽어낼 때, 측정과 해석의 기준이 되는 좌표계는 인간이 세웠다. 온도와 질량, 종과 원소처럼 인간이 만든 개념의 틀 위에서 세계를 분류하고 법칙을 세웠다. 이 사고방식 위에서 과학은 멀리 나아갔다.

그런데 화학에서 AI는 사람이 일러 주지 않은 내부 표현을 스스로 만들어 물성을 예측해 내고 있다. AI는 제 나름의 잠재공간, 즉, 세계를 정리하는 내부 좌표계를 세우고 그것을 바탕으로 높은 성능을 낸다. 그러나 그 좌표계가 인간의 것과 얼마나 겹치고 다른지는 체계적으로 이해하지 못하고 있다. 만약 인공지능이 쓰는 잠재공간의 좌표계를 사람이 그대로 읽어 낼 수 있다면, AI는 답을 내는 기계를 넘어 새로운 개념을 상상하는 동반자가 될 것이다.

그 번역이 가능해지면, 우리가 써 온 과학의 범주가 충분한지 되짚게 된다. 어떤 개념은 그대로 유효하고, 어떤 것은 더 압축적인 새 개념으로 바뀌며, 서로 다른 분야의 현상이 전혀 예상치 못한 구조 아래 묶일 수도 있다. AI가 그린 다른 좌표를 인간 과학의 언어로 이해할 수 있다면, 과학의 새로운 경계를 열 수 있을 것이다.

인공지능이 답뿐 아니라 물어볼 가치가 있는 질문까지 던질 수 있는가?

- 이준만 / 경영전문대학원
- 박현우 / 데이터사이언스대학원
- 최재원 / 사회과학대학 경제학부

위대한 발견은 늘 위대한 질문에서 시작됐다. 다윈은 종의 다양함을 보며 진화의 원리를 물었고, 아인슈타인은 빛을 타고 달리면 무엇이 보일지를 물었다. 그동안 무엇을 물을지 정하는 일은 사람의 몫이었고, AI는 주어진 문제를 풀고 데이터에서 패턴을 찾는 도구로 여겨졌다. 질문을 던지는 일만은 인간의 특권으로 남아 있었다.

그런데 알파폴드는 인간 연구자가 가설로 제시한 적 없는 단백질 접힘패턴을 독자적으로 발견했다고 보고했고, 이 패턴은 이후 실험에서 실재가 확인되었다. 수학에서도 AI가 누구도 살피지 않은 추측을 제안하고 증명해낸 사례가 나오기 시작했다. 답을 찾는 것이 아니라, 인간이 던지지 않은 질문을 기계가 제시한 것이다. 이것이 우연한 행운인지, 아니면 AI가 연구 질문을 체계적으로 생성할 수 있다는 증거인지는 아직 알 수 없다.

만약 AI가 매일 수천 개의 물음을 쏟아 낸다면, 인간 연구자는 가설을 검증하는 사람에서, AI가 내놓은 질문의 가치를 판단하는 사람으로 바뀐다. 그런데 인간이 아니라 AI가 좋은 질문이란 무엇인가를 판단하는 기준 또한 제시할 수 있는가. 그때 인간의 역할은 무엇인가. 끝없이 반복되는 이 재귀적 질문에 답을 찾아야 한다.

여러 인공지능 모델이 경쟁하면서 스스로 진화할 수 있는가?

• 이지성 / 공과대학 컴퓨터공학부

지금까지 인공지능은 주어진 데이터에 맞춰 오차를 줄여 가는 최적화로 학습해 왔다. 성능을 끌어올리는 길은 곧 더 크고 정교한 모델을 빚는 일이었고, 모델이 커질수록 더 똑똑해진다고 믿었다. 학습이란 여러 대안을 두고 고민하는 것이 아니라 하나의 우월한 정답을 향해 수렴해 가는 과정이라는 생각은 좀처럼 흔들리지 않았다.

그런데 『이기적 유전자』가 그리는 생태계에서는 사정이 다르다. 한 환경 안에서 서로 다른 생존 전략이 동시에 경쟁하고 공존하며, 때로는 정교한 전략보다 단순한 쪽이 더 오래 살아남기도 한다. 인공지능 생태계에서 데이터셋을 하나의 생태계로, 모델을 그 안에서 살아가는 개체로 본다면 어떨까. 학습은 단 하나의 최적해를 찾아가는 일이 아니라, 저마다 다른 단순한 전략들이 경쟁하고 갈라지며 지배적 위치를 찾아 가는 진화적 과정으로 다시 그려질 수 있다.

이 그림이 맞다면, 지능은 하나의 꼭짓점으로 모여드는 것이 아니라 여러 전략이 어우러진 결과로 이해된다. 작고 빠르며 속을 들여다볼 수 있는 모델들이 경쟁과 협력을 거쳐 함께 문제를 푸는 설계가 열린다. 인공지능 연구의 핵심개념도 최적화 중심의 공학에서 생태와 진화, 적응의 개념으로 넓어진다. 더 크고 정교한 모델 하나를 만들어내기 위해 노력하기에 앞서 지능이 본래 하나의 정답으로 다가가는 일인지, 여럿이 경쟁하면서 진화하는 것인지를 물어야 한다.

계산의 열적 한계는 어디까지 낮출 수 있는가?

• 최우영 / 공과대학 전기·정보공학부

반도체와 AI는 더 많은 정보를 더 빠르게 처리할수록 전력과 발열의 벽에 부딪힌다. 그동안 연구는 소자를 다듬고, 회로를 최적화하며, 냉각능력을 높임으로써 그 한계를 조금씩 뒤로 밀어 왔다. 그러나 계산과 열이 맺는 관계 자체를 근본적으로 의심한 적은 드물었다. 계산에는 으레 열이 따르고, 기술 발전이란 그 발열량을 조금씩 줄여 가는 일이라는 생각은 좀처럼 흔들리지 않았다.

정보를 처리하는 가장 작은 동작은 스위치를 켜고 끄는 일이다. 물론 정보 1비트를 지우는 데는 일정량 이상의 에너지가 반드시 필요하다는 란다우어 원리가 있으며, 이 원리에 따른 최소 에너지도 정해져 있다. 그러나 오늘날의 트랜지스터는 이 이론적 한계보다 수천 배 많은 열을 쏟아 낸다. 발열의 대부분은 계산의 본질에서 비롯된 것이 아니라, 현재의 스위칭 방식이 만들어 낸 군더더기인 셈이다. 그렇다면 스위칭에 드는 에너지를 란다우어 한계에 얼마나 가깝게 줄일 수 있을까. 이를 가능하게 할, 기존과 다른 새로운 스위칭 원리는 존재할 수 있을까.

열적 한계에 가까운 새로운 스위칭 원리를 찾을 수 있다면, 로직과 메모리, AI 가속기와 데이터센터, 나아가 손안의 기기까지 정보처리의 설계 원리가 함께 바뀔 것이다. 그것은 반도체가 따라온 반세기의 로드맵을 처음부터 다시 쓰는 일이기도 하다. 이제는 에너지 효율을 조금 더 개선하는 데 그칠 것이 아니라 현재의 발열 수준을 이론적 한계까지 낮출 수 있는 새로운 패러다임이 존재하는지 처음부터 다시 물어야 한다.

인간의 존엄과 삶의 의미는 반드시 노동을 전제로 하는가?

• 김동호 / 농업생명과학대학 산업인력개발학과

산업사회 이후 노동은 단순한 생계 수단을 넘어, 사회적 인정과 자기 정체성의 바탕으로 작동해 왔다. 사람은 무슨 일을 하느냐로 자신을 소개하고, 그 일을 통해 쓸모와 정체성을 인정받았다. 교육은 일자리 진입을 준비시키고, 사람의 가치도 얼마나 어떻게 일하는가로 매겨졌다. 일하지 않는 삶은 모자란 삶으로 여겨졌고, AI 시대의 논의조차 대개 사람을 어떻게 더 잘 일하게 할지에 머물러 있다.

그러나 AI와 자동화가 생산의 상당 부분을 떠맡으면, 생존을 위한 노동은 더 이상 대다수 삶의 기본 형식이 아닐 수 있다. 그때 사람은 무엇에서 삶의 의미를 찾고, 무엇을 통해 타인과 이어지는가. 일자리를 지키느냐 마느냐보다 더 중요한 질문이다. 노동의 종말을 단정하려는 것이 아니라, 노동이 차지하던 자리가 비었을 때 인간다움과 존엄이 무엇을 중심으로 다시 설 수 있는지를 묻는 것이다.

생존을 위한 노동 대신 사유와 배움, 창조와 돌봄, 놀이 같은 자기형성의 활동이 삶의 한가운데 설 수 있다면, 사람의 존엄은 일자리가 아니라 그가 무엇을 가꾸고 누구를 돌보는가에서 나올 수 있다. 그렇게 되면 학교는 직업 준비를 넘어 그런 능력을 기르는 곳이 되어야 하고, 사회는 일하지 않는 사람도 온전한 시민으로 설 자리를 마련해야 한다. 일자리를 어떻게 지킬지의 문제가 아니다. 인간의 존엄이 끝내 노동에 매여 있는지, 아니면 다른 무엇 위에 다시 설 수 있는지를 묻는 일이다.

인간의 정체성은 사회적 상호작용 없이도 형성될 수 있을까?

- 임하진 / 사회과학대학 언론정보학과
- 김은미 / 사회과학대학 언론정보학과

사회과학은 오래도록 인간의 정체성을 사회적 상호작용의 산물로 여겨 왔다. 사람은 타인의 눈을 통해 자기를 인식하고, 집단에 속함으로써 내가 어떤 사람인지를 알게 되고, 작은 단서만으로도 내 편과 다른 편을 가르기도 한다. 거울자아에서 사회적 정체성 이론, 그리고 뇌 속 거울뉴런의 발견까지, 이 전제는 거의 모든 인간 이해의 바닥에 깔려 있다. 오늘의 양극화도 내집단 편향과 집단 정체성, 상대를 향한 적대라는 사회적 상호작용 위에서 설명된다. 인간의 거의 모든 현상에는 ‘사회적’이라는 말이 따라붙었고, 정체성의 형성도 예외는 아니었다.

그런데 AI 에이전트의 등장은 이 전제에 도전한다. 한 사람의 언어와 사고, 관계의 패턴을 학습한 AI 에이전트는 인간처럼 상호작용한다. 그러기에 AI와의 교류를 통해 타인의 시선과 사회적 압력을 덜어 낸 조건에서 과연 인간의 정체성이 어떻게 형성될 수 있는지를 들여다볼 수 있게 되었다. 사회적 중력이 약해진 자리에서 정체성은 어떤 모습이 되는지, 인류가 한 번도 해 보지 못한 관찰이 가능해진다.

정체성과 사회적 상호작용의 관계를 이해할 수 있다면, 사회과학이 가장 오래 의지해 온 대전제를 검증의 무대 위에 올릴 수 있다. 사회적 상호작용 없이는 정체성이 서지 않는다고 확인된다면, 그것은 인간이 사회적 존재임을 보여주는 가장 강력한 증거가 된다. 반대로 사회적 중력 없이도 정체성이 부분적으로라도 형성된다면, 우리는 인간의 사회성을 더 정교하게 다시 이해할 수 있게 될 것이다. 어느 쪽이든 이 질문은 양극화와 갈등, 집단소속의 문제에 접근하는 프레임을 바꿀 것이다. ‘인간은 사회적 존재’라는 오랜 명제가 맞는지 검증해 보고 재개념화해 볼 기회가 왔다.

우리는 왜 '충분함'에 도달하지 못하는가?

• 강민수 / 사범대학 과학교육과

생명과학은 오랫동안 욕구를 결핍의 신호로 해석해왔다. 배고픔은 에너지가 모자라다는 표시이고 갈증은 수분이 부족하다는 표시이며, 이 결핍이 채워지면 욕구신호는 꺼진다. 체온과 혈당, 삼투압도 마찬가지다. 목표값에 닿는 순간 피드백이 작동해 뭔가를 추구하는 행위가 멈춘다. 모든 생명체 안에는 '이만하면 됐다'를 아는 회로가 들어 있고, 따라서 생명은 균형으로 수렴하는 시스템으로 이해되어 왔다.

그런데 지위와 소비, 성취를 향한 인간의 욕구신호는 결핍이 채워져도 꺼지지 않는다. 도파민 회로는 손에 넣는 순간이 아니라 쫓는 동안 더 강해지도록 짜여 있어서, 하나를 이루면 곧 다음 목표에 도전하는 욕구신호가 켜지는 것으로 보인다. AI기반의 플랫폼 기술은 바로 그 회로를 정밀하게 이용한다. 충족 이후에도 멈추지 못하는 것은 한 사람의 의지가 약해서가 아니라, 인간에게 새겨진 설계 때문이다. 멈추는 것이 필요한 몸과, 멈춤을 모르는 욕구 사이의 비대칭이 있다.

충분함을 결핍의 반대말이 아니라, 그 자체의 논리를 지닌 상태로 이해할 수 있다면, 욕구가 왜 생기는지를 묻는 대신 멈춤이 어떻게 작동하는지를 물을 수 있다. 그 답과 함의는 생리학만으로는 찾을 수 없다. 경제와 기술 설계, 사회의 성과 기준까지 함께 다뤄야 할 문제가 된다. 그 해답을 알면 과학기술의 발전방향과 조직 및 사회의 설계방식이 달라질 것이다.

‘원인은 결과에 선행한다’는 전제 없이도 엄밀한 과학적 설명은 가능한가?

• 김세훈 / 공과대학 에너지시스템공학부

과학적 설명은 원인이 먼저 있고 결과가 뒤따른다는 선형적인 인과론 위에 서있다. 이 전제는 가설 수립과 실험 설계의 기본 문법이자, 원인을 찾아 제거하면 문제가 해결된다는 환원주의적 세계관의 토대로 작동해 왔다. 자연과학뿐 아니라 사회과학과 정책 영역에서도 근본 원인을 규명하는 것이 곧 학문적 엄밀함의 기준으로 간주되어 왔다.

그러나 원인과 결과가 서로 맞물려 있는 현상 앞에서는 이 전제가 흔들린다. 생명은 환경을 바꾸며 그 환경에 다시 적응하고, 제도는 사람의 행동에 영향을 미치지만, 동시에 사람의 행동에 의해 바뀐다. AI는 자기 출력이 다시 학습 데이터가 되어 다음 학습과 출력생성의 조건이 된다. 의식과 기후, 빈곤 같은 문제도 원인과 결과가 서로를 동시에 키운다. 무엇이 먼저인지 가리기 어려운 이런 고리를 앞뒤의 시간적 인과로만 따져서 설명하면, 정작 중요한 구조가 보이지 않는다.

원인과 결과를 시간 순서로 줄 세우지 않고도 현상을 엄밀하게 설명할 길이 있다면, 흩어져 있던 학문들이 같은 프레임에서 연결될 수 있다. 순환하는 인과, 동시에 맞물리는 제약을 다루는 설명 체계가 만들어진다면 복잡계 과학과 합성생물학, 인지과학과 사회 연구가 방법론적으로 같은 지평에서 초학제적으로 만날 수 있다.

몸으로 세계를 겪으며 자란 인공지능은 인간과 구별되지 않는 개성을 가질 수 있는가?

• 김승일 / 기초과학연구원

인간의 성격과 지능이 어떻게 만들어지는지를 두고, 타고난 유전과 생물학적 조건을 앞세우는 쪽과 살아온 경험과 환경을 앞세우는 쪽이 오래 맞서 왔다. 개성은 둘이 함께 빛은 결과라고들 하지만, 어디까지가 타고난 것이고 어디서부터가 경험의 몫인지는 분명하지 않다. 무엇보다 감정과 판단, 남과 관계 맺는 방식 같은 인간다움의 핵심이 타고나는 것인지, 아니면 오랜 성장과 경험 속에서 길러지는 것인지가 아직 분명하지 않다.

오늘의 인공지능은 언어를 익혀 맥락에 따라 다르게 반응하는 데까지 왔지만, 대개 몸이 없다. 세계를 직접 부딪치며 살아가는 존재가 아니다. 그렇다면 인공지능에게 물리적 몸을 주면 어떻게 될까. 환경을 감각하고, 제 행동의 결과를 겪고, 다른 존재와 관계를 맺으며 오랜 세월 저마다 다른 경험을 쌓아 간다면, 그 인공지능에게도 고유한 성격이 자리 잡는다. 그렇게 자란 인공지능의 개성을 우리는 사람의 것과 구별할 수 있을까. 구별할 수 없다면, 인간의 개성은 몸과 유전자보다 살아온 경험에서 온다는 뜻이 된다.

이 물음에 답하려면 개성이 생물학적 조건과 경험 가운데 어디에서 생겨나는지를 새로 들여다보아야 한다. 자아와 의식, 공감과 책임, 사회적 관계 같은 개념이 도대체 어떤 조건에서 성립하는지도 함께 묻게 되고, 인간과 인공지능을 가르던 경계의 기준마저 흔들릴 수 있다. 더 똑똑한 기계를 만드는 문제가 아니다. 마음과 개성이 무엇에서 생겨나는지를 묻는 일이다.

노동과 소득이 갈라져도 경제순환은 지속될 수 있는가?

• 장제성 / 사회과학대학 경제학부

시장경제는 한 가지 순환 위에 서 있다. 가게가 기업에 노동을 제공하는 대가로 임금을 받고, 그 임금이 소비가 되어 다시 기업의 생산을 떠받친다. 노동에서 소득으로, 소득에서 소비로 이어지는 이 고리가 시장을 굴러 온 견고한 기반이었다. 생산성이 오르면 사회 전체의 부가 함께 커진다는 믿음 위에서 모든 정책이 짜였다.

그런데 AI가 이 경제의 순환고리에 근본적인 질문을 제기하고 있다. 지난날의 기술은 인간의 노동과 함께 일했지만, 이제 AI는 생산에서 인간의 노동을 대체하는 단계로 들어섰다. 여기서 역설이 생긴다. 기업은 어느 때보다 높은 생산성을 누리지만, 정작 그 물건을 사야 할 가게는 일자리를 잃어 살 힘을 잃는다. 공급은 넘쳐나는데 수요가 증발한다. 생산성이 오를수록 사회 구성원의 주머니가 도리어 비어 가는 역설이다.

노동과 소득이 갈라진 세상에서도 돈이 돌고 시장이 지탱되려면, 부가 어떤 길로 흘러야 하는지를 새로 물어야 한다. 임금이 아닌 통로로 소득이 가게에 닿을 수 있는지, 그러려면 어떤 사회계약이 필요한지가 과제가 된다. 인류의 번영과 경제적 순환을 지속할 수 있는 새로운 가치 분배 모델과 경제 시스템이 무엇인지 고민하게 될 것이다. 노동과 소득의 끈이 풀린 시대에 시장경제에서 경제는 어떻게 순환할 수 있을까?

지식노동의 대체를 결정하는 것은 기술인가, 사회의 합의인가?

• 조영준 / 사회과학대학 경제학부

AI 모델의 성능이 인간을 넘어서는 순간 지식노동은 곧 사라질 것이라는 전망이 지배적이다. 그러나 진단과 판결, 교육과 연구, 평가와 자문 등의 일을 모두 컴퓨터의 처리 능력만으로 대신할 수 없다. 권위와 책임, 신뢰 및 공정성의 문제가 얽혀있어, 궁극적으로는 사회적 합의가 함께 있어야 한다.

기술의 역사가 이를 증명한다. 19세기 러다이트의 기계 파괴 운동은 기술에 대한 단순한 반발이 아니라 정당한 노동의 가치 및 환경을 둘러싼 요구였다. 기술의 성숙이 공장 현장의 대체로 이어지지 않는다고, 사회적 합의와 제도의 변화를 거쳤다. 의사가 진단하고 판사가 판결하며 교수가 가르치고, 전문가가 평가하는 까닭은 단순히 지식 때문만이 아니라, 그 판단에 제도적 책임과 공적 신뢰가 부여되어 있고, 이에 대한 사회적 합의가 있기 때문이다. AI가 더 정확한 답을 내더라도, 사회가 받아들일 준비가 되어 있지 않다면 대체는 이루어지기 어렵다.

그렇다면 지식노동이 넘어가는 경계는 모델의 성능이 아니라, 사회가 권위와 책임을 기계에 넘기기로 합의하는 그 시점일 것이다. “기계가 인간보다 더 잘할 수 있는가”가 아니라 “인간이 넘겨주어도 되는가” 또는 “어디까지 넘겨줘야 하는가”로 질문 자체가 달라져야 한다. 무엇을 기계에 맡기고 무엇을 끝내 인간의 책임으로 남길지, 그 선을 누가 어떻게 그을지를 사회가 함께 정해야 한다. 그러나 그 합의를 어떻게 도출할 수 있는지는 아무도 모르고 있다.

망각하지 않는 문명은 여전히 스스로를 갱신할 수 있는가?

• 나진수 / 공과대학 재료공학부

문명이 기억과 정보를 축적하며 발전해 왔다는 믿음은 상식처럼 여겨져 왔다. 더 많이 기록하고 더 오래 보존하고 더 쉽게 다시 꺼내 볼수록 사회는 나아진다고 보았다. 그 전제 속에서 기록은 늘 좋은 것이었고, 망각은 모자라거나 위험한 것이었다.

그러나 인간의 공동체는 기억만으로 굴러가지 않았다. 용서와 화해, 공소시효와 사면, 세대교체와 명예 회복은 모두 과거가 시간이 지나면 다른 무게로 다뤄져야 한다는 약속 위에서 있었다. 기억이 현재를 규정하는 힘을 천천히 잃는 그 시간의 완충 덕분에, 사회는 매번 다시 출발할 수 있었다. 디지털 기록이 과거를 남기는 기술이었다면, AI는 그 과거를 끊임없이 현재로 재소환하는 기술이다. 흩어졌던 말과 실수가 연결되고 요약되어 지금의 판단 근거로 되 돌아온다. 완충은 사라지고, 과거는 좀처럼 현재를 놓아주지 않는다.

물론 반복되어선 안 될 비극의 기록처럼, 결코 지워져선 안 될 기억도 있다. 그러나 모든 것이 남아 끝없이 되살아난다면, 실수와 낙인, 상처마저 함께 영원해진다. 그때 법은 사라짐과 회복의 조건을 다시 생각하고, 역사학은 과거와 현재 사이에 어떤 거리와 리듬이 필요한지를 다시 물어야 한다. 무엇을 남길 것인가의 문명에서, 무엇이 사라질 수 있어야 인간과 문명이 다시 시작될 수 있는가를 묻는 문명으로 넘어가는 일이다.

사회와 인공지능 시스템의 안전은 오류를 막는 데 있는가, 오류에서 회복하는 데 있는가?

• 조재표 / 공과대학 컴퓨터공학부

우리는 사회와 인공지능 시스템의 안전을 더 정교한 규칙과 더 강한 통제, 더 완벽한 예측으로 지킬 수 있다고 믿어 왔다. 오류는 없어야 할 예외이고, 좋은 시스템이란 처음부터 실패하지 않는 시스템이라는 생각이 오래 이어졌다.

그러나 사회도 기술도 끊임없이 흔들린다. 정작 무서운 것은 오류가 생기는 일 자체가 아니라, 잘못된 신호 하나가 플랫폼과 여론, 정책과 알고리즘을 타고 연쇄로 증폭되는 일이다. 작은 오판이 순식간에 사회 전체의 판단을 물들인다. 사람이 더 촘촘한 규칙으로 이 연쇄를 미리 막으려 해도, 예측하지 못한 신호는 늘 새로 생겨난다.

그렇다면 사회와 인공지능 시스템을 명령과 통제의 체계가 아니라, 오류를 계속 찾아내고 영향이 번지는 범위를 좁히며 스스로 바로잡는 구조로 만들 수 있을까. 잘못이 생겨도 범위를 좁히고, 그 과정을 학습해 다음 판단을 고쳐 가는 시스템이다. 안전을 사전 봉쇄가 아니라 끊임없는 오류 교정의 능력으로 이해하면, AI 안전과 민주주의, 플랫폼 운영과 재난 대응이 하나의 개념틀에서 해석될 수 있다. 이 질문에 답을 얻는다면, 실수해도 무너지지 않고 오히려 배우는 사회로, 통제의 문명에서 복원의 문명으로 넘어가는 새로운 출발선에 설 수 있다.

여러 인공지능의 상호 견제 시스템이 인간의 최종 거부권을 지킬 수 있는가?

• 한진모 / 공과대학 전기·정보공학부

인공지능을 안전하게 만드는 일은 대개 하나의 거대 모델이 내놓는 답변을 인간의 기준에 맞게 교정하는 문제로 여겨져 왔다. 모델 하나를 인간의 가치에 충분히 정렬할 수 있다면 인공지능의 위험도 통제할 수 있다는 믿음이 있었고, 모델의 안전성을 높이기 위한 해법은 주로 한 모델의 출력을 사람이 선호하도록 교정하는 방향에서 탐색되었다.

그러나 거대 모델은 확률적으로 동작한다. 아무리 촘촘히 통제해도 환각, 뜻밖의 일탈이나 우회 공격의 위험을 완전히 지울 수 없다. 인간 사회는 한 사람의 완전한 선의에 기대기보다, 권력을 나누고 서로를 견제하며 교차로 검증하는 제도를 통해 안정성을 지켜왔다. 불완전한 행위자라도 서로를 독립적으로 감시하게 만들면, 어느 하나가 폭주하거나 모두가 같은 오류에 빠질 가능성을 크게 줄일 수 있다.

인공지능 안전의 초점을 한 모델의 답변을 통제하는 데서, 여러 인공지능이 판단과 실행 권한을 나누고 서로를 검증하는 시스템을 설계하는 것으로 옮길 수 있을까. 답변의 통제만 강화하며 확률적 모델이 인간의 도덕을 온전히 이해해 주기를 기대하는 대신, 어느 하나도 권한을 독점하지 못하는 다중 인공지능 구조를 통해 더 인공지능 기반 시스템이 더 쉽게 안전해질 수 있을까. 게임이론과 거버넌스 연구를 컴퓨터과학에 들여와, 인간이 위험한 결정 앞에서 마지막으로 멈추고 거부할 권리를 보장받을 수 있는 이론적인 균형점을 찾을 수 있을까.

평소엔 흐르고 충격엔 스스로 굳는 배터리 전해질을 만들 수 있는가?

• 조민재 / 공과대학 전기·정보공학부

배터리 전해질은 액체이거나 고체, 둘 중 하나로 여겨져 왔다. 액체는 이온이 잘 흘러 효율이 높지만 충격에 약해 불이 나기 쉽고, 고체는 단단해 안전하지만 이온이 더디게 흘러 효율이 떨어진다. 지금의 기술은 그 둘 사이에서 타협하는 데 머문다. 배터리 안전도 대개 바깥의 보호 회로와 냉각, 사고 뒤의 차단에 기대 왔다.

그런데 충격을 받는 순간에만 굳고 평소엔 잘 흐르는 물질이 있다. 우블렉처럼, 가하는 힘에 따라 묽어지고 단단해지는 비뉴턴 유체다. 다만 이런 유체는 대개 전기가 통하지 않는다. 여기에 이온이 오갈 길을 열어 전기가 흐르게 만들 수 있을까. 평소에는 이온이 잘 흐르는 액체로 작동하다가, 부딪히거나 찢리는 순간 스스로 굳어 위험의 통로를 닫는 전해질을 상상하는 것이다.

그러면 안전은 시스템 바깥에 덧붙이는 장치가 아니라, 물질 자체에 새겨진 성질이 된다. 사고가 난 뒤 전류를 끊는 것이 아니라, 사고의 순간 물질이 먼저 길을 막는다. 잦은 충격에 노출되는 전기차와 드론, 항공우주와 웨어러블의 안전 원리가 함께 달라진다. 안전을 시스템 바깥에서 덧붙일 것인지, 물질 스스로 위험에 반응하게 만들 수 있는지의 차이다.

노화를 공간생화학적 네트워크의 붕괴로 설명할 수 있을까?

• 강윤표 / 약학대학 약학과

노화는 시간의 흐름에 따라 생명체의 기능이 점진적으로 소실되는 과정으로 이해되어 왔다. 2013년 Cell에 발표된 노화의 특징적 징후는 이를 유전체 불안정성, 텔로미어 소실, 단백질 항상성 붕괴, 세포 노화 등으로 체계화했고, 이후 연구는 각 특징적 표지를 개별 표적으로 삼아 조절하는 방향으로 전개되었다.

그러나 기존 접근은 시간에 따른 개별 구성요소의 기능 소실을 노화의 핵심 변수로 전제한 채, 생체 분자들 사이의 관계를 유지하는 분자생물학적 네트워크가 어떻게 형성되고 붕괴하는지를 시스템 수준에서 충분히 묻지 않았다.

생체 분자들의 상호작용을 통합적으로 파악하고, 노화의 정의와 측정 기준을 시간 중심에서 네트워크의 안정성과 회복력 중심으로 전환한다면, 단일세포부터 개체 수준까지 아우르는 통합 진단 프레임워크가 가능해진다. 개입의 목표 역시 개별 기능의 복구에서 네트워크 전체의 회복으로 이동한다. 나아가 공간 네트워크 조절제와 같은 새로운 약물 클래스의 개발, 그리고 알츠하이머병·당뇨병·심혈관질환을 아우르는 통합적 예방·치료 전략을 여는 출발점이 될 수 있다.

생명의 규칙은 살아남은 형태에 있는가, 진화가 비워 둔 곳에 있는가?

• 이다인 / 자연과학대학 생명과학부

생명과학은 발현된 유전자, 살아남은 종, 구조가 밝혀진 단백질처럼 보이는 것을 읽어 왔다. 진화에 성공한 결과물을 해독하는 그 방향에서 생명과학은 눈부시게 나아갔다. 그러나 자연에 실제로 나타난 생명의 형태는 가능했던 조합 가운데 극히 일부다. 진화의 손길이 미치지 않은 나머지 방대한 영역은 비어 있고, 우리는 그 빈 쪽을 좀처럼 들여다보지 않았다.

미켈란젤로가 조각한 피에타의 실루엣은 무언가를 더해서가 아니라 덜어 내어 완성됐다. 조각가가 깎아 없앤 돌이 역설적으로 우리가 보는 형상을 빛낸다. 생명도 그렇지 않을까. 단백질의 접힘에도, 생체 네트워크에도, 논리적으로는 가능하지만 자연계에는 끝내 나타나지 않는 미지의 구역이 있다. 그 빈 곳은 우연히 비어 있는 것이 아니라, 생명체가 스스로 무너지지 않으려고 쳐둔 경계 너머의 지형일 수 있다.

그렇게 보면 질병은 부품 하나가 고장 난 일이 아니라, 안전하게 머물 수 있는 경계를 넘어서는 일로 해석할 수 있다. 진화가 피해간 영역의 지도가 그려지면, 합성생물학도 무엇을 더 만들 수 있는지가 아니라 무엇을 건드리면 안 되는지의 안전선을 먼저 고려하게 될 것이다. 생명의 규칙을 살아남은 형태에서만이 아니라 진화가 끝내 허락하지 않은 빈 곳에서도 읽어 낼 수 있다면, 더 많이 조작하기에 앞서 넘지 말아야 할 경계부터 그릴 수 있다.

진화의 방향을 미리 내다볼 수 있는가?

• 석승혁 / 의과대학 미생물학교실

진화생물학은 생물이 환경과 선택압 속에서 어떻게 변해 왔는지를 잘 설명해 왔다. 그러나 그것은 대개 지나간 결과를 되짚는 일이었다. 진화는 본래 우연하고 역사적인 과정으로 여겨졌기 때문에, 그 방향을 미리 내다보거나 예측 가능한 동역학으로 다루려는 시도는 드물었다.

고래의 진화를 보면 생각이 달라진다. 약 천만 년이라는 짧은 시간에 육상 포유류가 완전한 수생 동물로 바뀌었고, 같은 조상에서 갈라진 계통이 한쪽은 지구 역사상 가장 큰 몸집으로, 다른 쪽은 작은 체형으로 전혀 다른 길을 갔다. 진화가 단일한 방향으로 귀결되는 것이 아니라, 여러 제약과 선택압이 결합된 다차원 공간에서 서로 다른 안정 상태로 수렴하는 과정으로 보이는 것이다.

그렇다면 유전체와 생리, 생태적 제약을 통합한 다차원 자유에너지 지형(landscape)을 재구성하여 진화의 방향을 예측할 수 있지 않을까? 진화를 이 지형 위에서의 이동과정으로 이해한다면, 진화생물학은 과거를 설명하는 학문에서 미래를 부분적으로 내다보는 과학으로 확장된다. 여러 계통에서 거듭 나타나는 거대화과 소형화, 생태 특화를 하나의 프레임으로 설명하고, 기후변화의 압력이 거세질 때 생물이 어디로 향할지도 가늠할 수 있을 것이다.

면역과 인지는 같은 뿌리에서 자란 통합된 기억 시스템인가?

• 김도윤 / 약학대학 약학과

현대 과학은 오랫동안 면역계와 신경계를 전혀 다른 시스템으로 갈라 보았다. 면역계는 바깥에서 들어온 병원체를 가려 없애는 방어 장치이고, 신경계는 감각을 받아들이고 경험을 저장하는 고등한 정신의 영역이라고 여겼다. 그래서 기억이라는 같은 이름의 현상도 두 영역에 따로 나뉘어 있었다.

그런데 두 메커니즘은 상위 수준에서 닮은 데가 많다. 자극을 받아들이고, 중요한 것만 골라 새겨 두고, 다시 마주치면 빠르게 반응한다. 위험한 상황이 시냅스에 각인되는 일과 병원체 경험이 면역세포에 남는 일은 그 원리가 다르지 않다. 더구나 뇌 속 면역세포인 미세아교세포는 시냅스를 속아 내며 기억이 만들어지는 데 직접 간여한다. 기억의 주인이 신경계만은 아닐지 모른다. 그렇다면 기억은 뇌에만 있는 것이 아니라 온몸의 면역계에 동시에 새겨지는 것은 아닐까.

기억을 단순한 저장이 아니라 살아남기 위한 적응의 방식으로 보면, 면역과 인지는 한 뿌리에서 갈라진 두 가지로 해석할 수 있다. 그 관점이 서면 트라우마와 알츠하이머, 면역 질환과 AI 학습이 하나의 패러다임으로 설명될 수 있다. 나아가 병을 특정 기관의 고장으로만 보지 않고, 몸이 경험을 기억하고 반응을 매만지는 방식 전체로 다루는 새로운 학문이 탄생할 수 있다.

생물학적 노화는 실재하는가?

· 이서영 / 공과대학 재료공학부

노화는 시간이 지남에 따라 점진적으로 축적되는 보편적인 생물학적 과정으로 이해되어 왔다. 이 전제 위에서 노화를 12가지 표지로 체계화하고, 노화 시계를 통해 개인의 생물학적 나이를 수치화하려는 시도가 이어진다. 그러나 이 관점은 노화라는 개념이 실제로 하나의 실체를 가리키는지, 아니면 건강 상태를 대리하는 추상적 구성물에 가까운지에는 답하지 못한다. 인간의 분자적 변화는 시간에 따라 고르게 쌓이기보다 특정 시점에서 크게 재편되기도 하고, 장기마다 노화 속도가 달라 질병과 사망 위험을 다르게 예측한다는 보고도 있다.

그렇다면 노화의 서로 다른 징표들은 과연 하나의 공통된 생물학적 시간축과 본질을 공유하는 것일까? 우리가 노화라고 부르는 것은 단일한 과정이라기보다, 서로 다른 수준에서 발생하는 변화들을 하나의 이름 아래 묶어 온 해석의 틀일 뿐인 것은 아닌가?

생물학적 노화는 실재하는가라는 질문은 서로 다른 장기와 세포, 개인에게서 나타나는 변화들 가운데 무엇이 노화의 공통 원리이고 무엇이 질병·환경·생활 조건에 따라 달라지는 현상인지를 다시 가려내도록 요구한다. 노화가 단일한 과정이 아니라면, 앞으로의 연구는 보편적 시계를 찾는 데 머물 수 없다. 어떤 변화들이 실제로 서로 연결되어 있고, 어떤 변화들은 단지 함께 관찰되어 왔을 뿐인지를 새롭게 구분해야 한다. 결국 이 질문은 노화를 자연스럽게 보편적인 생의 흐름으로 받아들여 온 통념을 흔드는 동시에, 노화를 과학적으로 엄밀하게 재정의하고 측정·제어 가능한 연구의 토대를 마련하는 출발점이 된다.

식물의 생육 시간을 설계할 수 있는가?

• 김정선 / 농업생명과학연구원

농업과 원예학은 작물이 대체로 정해진 기간에 걸쳐 자란다는 전제 위에 서 있다. 잎채소는 한 달 남짓, 열매채소는 몇 달이 걸린다는 식이다. 재배 현장과 연구 역시 그 시간을 자연이 부여한 조건으로 받아들인 채 이루어져 왔다. 빛과 온도, 환기와 양분을 정밀하게 맞춰 생산량과 품질을 끌어올렸지만, 왜 꼭 그만큼의 생육 시간이 걸려야 하는지는 좀처럼 묻지 않았다.

이 분야의 연구는 30~40일 걸리던 잎채소를 28일로 당기는 것처럼 대개 그 생육 기간을 며칠 줄이는 데 머문다. 그러나 그 시간이 자연이 정한 상수가 아니라 조절 가능한 변수라면 질문은 달라진다. 수주가 걸리던 과정을 수일이나 수시간으로 재구성할 수 있을 것이다. 이것을 가능하게 하려면, 발아와 세포분열, 분화, 생체시계가 어떤 시간 질서 속에 짜여 있는지를 알아야 한다.

그러면 농업은 더 빨리 기르는 기술을 넘어, 생명이 시간이라는 자원을 어떻게 쓰는지를 다루는 일이 된다. 계절과 재배 기간에 종속된 생산에서, 생물학적 시간의 구조 자체를 설계하는 생명 시간의 공학으로 넓어진다. 햇빛도 계절도 없는 우주 거주지에서 식량을 길러야 하는 경우에는 더 절실해진다. 따라서 이제는 단순히 더 빨리 재배하는 법을 넘어, 생명의 시간표가 주어진 것인지, 아니면 인공적으로 조절할 수 있는 것인지를 물어야 한다.

잠은 줄일 수 없는 자연의 섭리인가, 조절할 수 있는 과정인가?

• 박시현 / 교육종합연구원

인간은 한평생의 삼분의 일을 잠으로 보낸다. 자는 동안 특정 뇌파와 뇌척수액의 흐름이 뇌의 노폐물을 씻어 내고 기억을 정리하는데, 이 청소를 맡는 글림프 시스템은 주로 잠든 사이에만 작동한다. 그래서 절대적인 수면 시간을 채우지 못하면 인지 기능이 떨어지고 병이 따른다. 잠은 인간이 어찌지 못하는 자연의 섭리로 여겨졌다.

그런데 잠이라는 상태와 잠이 하는 일을 나눠 보면 어떨까. 노폐물 청소와 기억 통합, 신경망 정비가 정확히 어떤 조건에서 일어나는지 알 수 있다면, 그 일을 잠 바깥에서도 수행할 수 있을지 모른다. AI로 뇌의 피로와 대사 상태를 실시간으로 읽고, 비침습적 신경 자극으로 회복 경로가 돌아가게 한다면, 잠은 고정된 시간이 아니라 조절할 수 있는 생체 과정으로 바뀐다. 여덟 시간의 회복을 한 시간으로 압축하거나, 깨어 있는 채로 그 회복회로의 일부를 가동하는 방식이 있을 수 있지 않을까.

핵심은 잠을 통한 회복이 언제 어떻게 일어나는지를 밝혀, 불면과 교대 근무, 치매와 신경 퇴행을 다른 차원에서 다룰 수 있는 길을 여는데 있다. 초압축 수면이나 각성 중 회복은 먼 상상이지만, 이 질문을 따라가다 보면 궁극적으로 뇌가 스스로를 정비하는 원리가 드러날 것이다. 잠이 하는 일을 사람이 조절할 수 있다면, 인간이 시간을 쓰는 방식 자체가 달라질 것이다.

생명의 정보를 단백질에서 핵산으로 흐르게 할 수 있는가?

• 이준석 / 첨단융합학부

생명과학의 중심원리는 정보가 DNA에서 RNA로, RNA에서 단백질로 흐른다는 데 있다. 1958년 크릭이 세운 이 도그마 위에서 오믹스 연구가 쌓여 왔다. DNA나 RNA같은 핵산은 PCR로 단 한 분자에서도 무한히 증폭해 읽어낼 수 있지만, 단백질을 복사해 증폭시키는 기술은 아직 존재하지 않는다. 그래서 단백질 분석은 질량분석기의 성능에 전적으로 의존해왔다.

검출의 민감도에 있어 그 격차는 뚜렷하다. 단일 세포의 mRNA는 이제 수 개의 분자까지 감지할 수 있지만, 같은 세포의 단백질은 양이 많은 일부만이 포착된다. 역전사효소나 프리온처럼 이러한 중심원리를 거스르는 예가 자연에 존재한다. 그러나 유독 단백질의 서열을 핵산 서열로 되돌리는 경로만은 자연 어디에도 없다.

자연이 비워 둔 그 한 칸을 인공효소나 기계로 채울 수 있지 않을까. 단백질 서열을 핵산 정보로 되돌려 쓰는 '역번역'이 가능하다면, 단백질도 PCR처럼 극미량에서 증폭해 분석할 수 있다. 단일 세포의 단백질 지도 작성, 희귀 바이오마커 진단, 조직 내 미량 단백질 분석이 현재의 한계 너머로 나아갈 수 있다. 이는 분석 기술 하나를 더하는 일이 아니라 한 방향으로만 흐르던 생명의 정보에 인간이 처음으로 역방향 화살표를 그려 넣고, 정보의 흐름 자체를 설계할 수 있는지를 묻는 일이다.

동물은 빛을 먹고 살 수 없는가, 아니면 아직 그 방법을 모를 뿐인가?

• 강진호 / 국제농업기술대학원

인간을 비롯한 동물은 바깥에서 유기물을 먹어야 산다. 빛을 직접 에너지로 바꾸는 일은 엽록체를 가진 식물의 몫이고, 광합성 기관도 분자 장치도 없는 동물에게 빛을 먹고 산다는 말은 성립하지 않는다고 여겨져 왔다. 그것은 오랜 생물학적 상식이었다.

그런데 그 경계는 생각보다 뚜렷하지 않다. 산호와 일부 해양 생물은 광합성을 하는 미생물과 한 몸으로 공생하며 에너지의 일부를 빛에서 얻는다. 이미 동물과 식물 사이에 걸친 존재들이 있는 셈이다. 그렇다면 질문을 바꿔야 할지 모른다. 완전히 불가능한가가 아니라, 어디까지 가능한가로. 피부 아래 공생하는 광합성 미생물이나 세포 안의 인공 광반응 장치가 있다면, 동물도 아주 적게나마 스스로 에너지를 만들 수 있지 않을까.

완전한 광합성이 아니어도 하루 에너지의 일부, 극한에서 며칠을 더 버티게 할 만큼이라도 빛에서 얻을 수 있다면, 생명을 이해하는 틀이 넓어질 수 있다. 사료에 종속된 축산, 보급이 끊기는 우주와 극지에서 생존 전략도 달라진다. 동물과 식물을 가르던 분류를 고정된 벽이 아니라 바꿔 끼울 수 있는 기능의 조합으로 해석할 수 있다. 먹지 않아도 되는 생명을 만드는 문제가 아니라, 먹는 일에 덜 기대는 생명이 가능한지를 묻고자 한다.

식물의 살아 있는 신호 네트워크를 정보 인프라로 삼을 수 있는가?

• 권순영 / 자연과학대학 화학부

인류의 정보 기술은 한 가지 전제 위에 서 있다. 정보는 전자 신호로 흐르고, 중앙의 처리 장치가 그것을 해석한다. 반도체 위 전자가 데이터를 나르고, 서버가 길을 정하고, 프로세서가 판단한다. 뇌를 모방하는 뉴로모픽 칩도, 신경을 읽는 뇌-컴퓨터 인터페이스도 결국 이 틀을 공유한다. 정보는 중앙으로 모이고, 작동은 전자 회로 위에서 이뤄진다.

그런데 식물은 뇌도 중앙 신경계도 없이 수억 년간 정보를 처리해 왔다. 초식동물이 잎 하나를 갉으면 수 분 안에 식물체 전체가 방어 태세로 바뀐다. 가뭄이 뿌리를 위협하면 멀리 떨어진 잎의 기공이 닫힌다. 칼슘 파동이 초 단위로 조직을 가로지르고, 활성산소가 세포에서 세포로 신호를 이어 나르며, 관다발 속 펩타이드와 호르몬이 먼 거리를 오간다. 어느 곳에도 중앙 처리 장치는 없다. 도착한 신호를 각 세포가 스스로 해석하고 반응을 정한다.

그렇다면 우리가 회로 위에서 구현하려는 분산형 정보 처리를, 식물은 중앙 장치 없이 이미 살아서 수행하고 있는 것이 아닐까? 그렇다면, 작물이 스스로 물 부족과 병해를 알리고, 숲이 토양과 기후의 변화를 먼저 감지하는 노드가 될 수 있을 것이다. 정보를 꼭 전자 회로 위에서만 다뤄야 하는지, 스스로 자라고 연결되는 생명 자체를 정보의 바탕으로 삼을 수 있는지에 해답을 찾아야 한다.

생명의 최소 단위로서, 자가 복제가 가능한 독립형 인공 미토콘드리아를 설계할 수 있는가?

• 김태일 / 농업생명과학대학 농림생물자원학부

미토콘드리아는 흔히 세포에 종속된 발전소로 정의된다. 약 20억 년 전 내공생을 통해 독립성을 잃고 세포의 부속품으로 편입되었다는 것이 지배적인 학설이다. 이런 관점에서 미토콘드리아 연구는 세포핵의 통제 아래 일어나는 에너지 대사와 질병 기전에 집중되어 왔고, 미토콘드리아가 독립적 단위로 기능할 가능성은 주목받지 못했다. 그러나 미토콘드리아는 스스로 유전 정보를 복제하고 단백질을 합성하는 최소한의 기작을 갖추고 있을 뿐 아니라, 세포 사이의 터널을 통해 옮겨 다니며 손상된 조직을 복구하는 등 예상보다 훨씬 능동적인 구조물이다.

그렇다면 미토콘드리아를 세포에서 떼어 내 인공 구조체 안에서 홀로 살아가게 할 수 있을까. 독립적인 생존과 복제가 가능하다면, 그것이 여전히 세포소기관인지 아니면 새로운 형태의 생명체인지 되묻게 된다. 세포의 통제를 벗어나 스스로 복제하고 대사하는 인공 미토콘드리아를 만들 수 있다면, 생명의 최소 단위가 어디까지 내려가는지, 새로운 생명을 빚는 일의 한계가 어디인지도 함께 드러난다.

미토콘드리아의 유전체를 역설계해 세포 밖에서도 스스로 증식하는 최소한의 대사 유전체를 만드는 일, 즉 세포 독립적 에너지 유닛을 창조하는 것이 그 출발점이 된다. 인공 미토콘드리아를 이식해 난치성 질환과 노화에 다가서는 길도, 다른 세포소기관으로 제어 기술을 넓히는 길도 이 질문에 대한 해답에서부터 열릴 수 있다.

생태계 보전은 과거를 지키는 일인가, 변화하는 생명에 발맞추는 일인가?

• 유현 / 농업생명과학대학 농생명공학부

생태계 보전은 오래도록 과거의 어느 특정 시점의 상태를 정상으로 정해 두고, 그 상태를 지키거나 되돌리는 일이었다. 생명은 지킬 고유종과 없앨 외래종으로 나뉘었고, 그 가치는 인간에게 쓸모가 있느냐로 갈렸다. 자연을 멈춰 세운 한 장면이 곧 옳음의 기준이었다.

그런데 작은 생물들을 매일 들여다보며 그 미세한 차이를 기록하다 보면, 생명은 한순간도 멈춰 있지 않다. 기후가 흔들리자 생물들은 살길을 찾아 서식지를 옮기고 필사적으로 적응한다. 그 적응에 성공해 새 땅으로 옮겨 온 생물은 인간세상에서 교란종이라 불리며 박멸의 대상이 되고, 적응하지 못해 사라져 가는 생물은 멸종위기종이라 불리며 큰 비용을 들여 보존한다. 같은 적응인데 인간의 논리로 한쪽은 죽이고 한쪽은 살린다. 진화의 성공이 보전의 눈에는 파괴로 비치는 모순이다.

지킬 것을 과거의 모습이 아니라 변화 그 자체로 옮겨 본다면, 보전의 개념은 달라져야 한다. 무엇을 없애고 무엇을 남길지의 목록이 아니라, 옮겨 온 생물과 본래 살던 생물이 새로 맺는 관계를 어떻게 함께 살아가게 도울지의 문제가 된다. 자연을 한 시점에 가둘 수 있다는 믿음을 내려놓는 일이기도 하다. 우리가 지키려는 것이 과거의 풍경인가, 아니면 흐르는 생명 그 자체인가.

기후변화는 줄여야 할 위기인가, 끌어 쓸 수 있는 조건인가?

· 이창하 / 공과대학 화학생물공학부

기후변화는 오래도록 줄이고, 늦추고, 적응해야 할 위기였다. 온실가스를 덜 내보내고, 오르는 기온과 잦아지는 극한 기상에 대비하는 것이 대응의 전부였다. 줄이거나 피하거나 버티거나, 셋 중 하나였다. 변화한 기후를 도리어 끌어 쓰자는 공학적 발상을 하나의 기술이 아니라 패러다임으로 만들 수 있을까?

변화한 기후는 단순히 더워진 상태가 아니다. 대기와 해양의 순환이 재편되고, 지역과 고도 사이의 온도와 압력, 습도 차가 더 벌어지며, 탄소 농도와 염도, 수자원의 분포가 새로 짜인다. 공학적으로 일을 만들어 내는 힘은 바로 이런 차이, 곧 구배에서 나온다. 그 차이가 커졌다는 것은, 교란이 심해진 동시에 끌어 쓸 수 있는 구동력이 그만큼 커졌다는 뜻이기도 하다. 기후변화는 막아야 할 위기인 한편, 전에 없던 에너지와 물질의 기울기를 품은 시스템이 된다.

그렇다면 그 커진 구배를 직접 끌어 쓰는 길도 있을 수 있지 않을까. 벌어진 온도 차에서 열에너지를 회수하고, 바뀐 해류와 강해진 바람에서 동력을 얻고, 늘어난 이산화탄소를 원료로 되돌리고, 재편된 염도와 수자원에서 자원을 거두는 방식이다. 기존의 감축과 적응도 의미가 있지만, 기후변화를 능동적으로 끌어 쓰는 방식으로 공학적 패러다임을 전환할 필요가 있다.

인간의 뇌가 먼저 배우고, AI가 그 결과를 이어받을 수 있는가?

• 양예찬 / 공과대학 협동과정 바이오엔지니어링전공

오늘의 AI는 방대한 데이터와 반복 학습을 통해 높은 성능에 이른다. 그러나 환경이 바뀔 때마다 다시 조정해야 하고, 데이터 분포가 계속 변하는 현실에서는 새것을 배우는 과정에서 능력을 잃는 문제까지 겪는다. 반면 사람의 뇌는 적은 경험만으로도 변화를 빠르게 알아차리고, 기존 기능을 크게 무너뜨리지 않으면서 새로운 행동을 익힌다.

최근 연구는 소량의 뇌 신호를 활용해 AI가 인간의 뇌 반응과 더 유사한 표현을 학습하도록 유도할 수 있음을 보여주었다. 이는 모델을 무작정 키우는 대신, 인간 뇌가 정보를 처리하는 방식에서 얻은 단서를 AI 학습에 들여오려는 시도다. 그렇다면 학습의 출발점을 바꿔 볼 수 있지 않을까? AI가 새 환경마다 대규모 데이터로 다시 조정되는 대신, 빠르게 적응하는 생물학적 지능이 변화의 핵심을 먼저 파악하고, AI가 그 과정에서 얻어진 신호와 판단 패턴을 이어받아 널리 쓰일 모델로 다듬는 것이다. 즉, 인간 뇌의 적응 능력을 AI학습의 출발점으로 삼는 접근이다.

그렇게 되면 AI는 인간의 경험과 판단에서 얻은 단서를 바탕으로 새로운 상황에 더 빠르게 적응할 수 있다. 적은 데이터와 계산만으로도 낯선 환경, 예측하기 어려운 상황, 사용자마다 조건이 달라지는 현실 문제에 대응하는 AI의 가능성이 열린다. 이런 기술이 발전하면 기계는 더 이상 단순한 도구가 아니라 인간과 학습을 주고받으며 함께 적응하는 존재가 될 수 있을 것이다.

인간이 닿지 못하는 곳에서도 문명은 이어질 수 있는가?

• 김용승 / 치과대학 치의과학과

문명은 인간의 몸이 닿는 범위 안에서만 이어져 왔다. 사람이 숨 쉬고 먹고 잠들 수 있는 환경에서 도시를 세우고, 그 바깥은 잠시 다녀오는 곳일 뿐 머무는 곳이 아니었다. 문명의 존속은 늘 인간이 살아남을 수 있는 조건과 한 묶음이었고, 사람이 사라진 자리에서 문명이 이어진다는 그림은 그려진 적이 없다.

그런데 인공지능과 로봇은 이미 사람이 버틸 수 없는 환경에서 감지하고 판단하고 일한다. 깊은 우주와 심해, 사람이 오래 머물 수 없는 극한의 현장, 수십 년이 걸려 사람의 생애를 넘기는 작업에서는 비유기적 존재가 더 알맞은 일꾼이 된다. 더 나아가 통신이 끊겨 누구의 지시도 닿지 않는 곳에서, 기계들끼리 스스로 고치고 짓고 이어 갈 수 있다면 어떨까. 그곳에서는 문명을 인간의 몸과 현장의 통제에 묶어 두는 일이 오히려 한계가 된다.

사람이 없는 환경에서도 스스로 굴러가는 문명을 상상하는 일은, 인간을 기계로 대체하자는 말과 다르다. 기억과 규칙, 협력과 판단을 갖춘 사회적 시스템이 인간의 생물학적 한계 너머에서도 작동할 수 있는가가 질문의 핵심이다. 그 답을 좇다 보면 무엇을 문명이라 부를 수 있는지, 그 문명을 누가 이어받아 지켜 갈 수 있는지를 함께 고민할 수밖에 없다. 문명의 주체와 계승자는 인간으로만 한정되어야 하는가. 아니면 인간이 만든 비유기적 행위자 역시 인간이 닿지 못하는 시간과 공간에서 문명을 이어갈 수 있는가.

인간은 동물, 식물, 기계를 포함한 다른 종과 대화할 수 있는가?

- 장인철 / 사범대학 영어교육과
- 정대홍 / 사범대학 화학교육과

언어는 인간만의 능력으로 여겨져 왔다. 복잡한 문법과 상징 체계를 갖춘 인간의 말이 있어 고등 사고가 가능했고, 그 말을 공유하지 않는 다른 종은 이해의 상대가 아니라 관찰하고 이용하는 대상에 머물렀다.

그러나 의사소통이 꼭 말로만 이루어지는 것은 아니다. 동물은 행동과 소리로, 식물은 화학 신호와 생리적 반응으로 정보를 주고받는다. 우리는 오래전부터 반려동물과 마음을 나눴고, 이제는 AI를 엮은 기기에 말과 몸짓으로 명령하는 일도 어색하지 않다. AI 번역은 서로 다른 말을 쓰는 사람을 잇고, 신호 해석 기술은 이질적인 소통 체계를 연결하기 시작했다. 개의 몸짓, 새의 울음, 식물의 화학 반응, 기계의 상태 신호에도 의사소통의 환경과 의도가 담길 수 있다. AI가 그 신호의 패턴을 읽어 사람이 알아들을 표현으로 옮기는 매개기능을 할 수 있을 것이다.

여기서 대화란 인간의 말을 다른 존재에게 가르치는 일이 아니라, 그들 각자의 신호 체계를 읽어 내는 일이다. 그 소통이 명령 전달에 그치는지, 감정을 나누고 협력하는 관계로 나아갈 수 있는지가 관건이다. 언어학과 뇌과학, 생물학과 AI, 윤리학이 함께 고민해야 할 질문이다.

인간과 인공지능이 권리와 의무를 공유하는 규칙을 만들 수 있는가?

• 이수빈 / 과학데이터혁신연구소

인공지능은 대체로 인간이 쓰는 도구이자 서비스로 여겨진다. 권리와 책임은 사람에게 있고, 기계는 아무리 똑똑해져도 누군가의 소유물이다. 그런데 일반인공지능이 생산과 의사결정, 안전과 복지의 한복판에 들어선다면, 그것을 끝까지 물건으로만 둘 수 있는지 의문이 들 수밖에 없다.

기술이 빠르게 앞서가는 사이, 그것을 가진 쪽과 갖지 못한 쪽의 격차가 지배와 종속의 구조로 굳어질 위험도 커진다. 그렇다면 AI를 법과 정치의 질서 안 어디에 놓을지 미리 물어야 한다. 예를 들어 AI에게 투표권을 줄지, 세금을 물릴지, 병역을 지을지를 따져 보면, 권리와 의무를 어디까지 나눌 수 있을지 가늠해 볼 수 있다. 권한 없이 책임만 지우거나, 책임 없이 권한만 주는 구조는 어느 쪽이든 위험롭다. 핵심은 인간과 AI가 함께 살아갈 때 권리와 의무를 누구에게 어떻게 나눌 것인가에 있다.

AI가 경제와 행정, 생산을 실제로 떠맡는 시대가 온다면, 인간만을 전제로 짜인 민주주의와 자본주의의 틀을 손볼 수밖에 없다. AI를 소유한 개인과 기업의 권한, AI가 일으킨 피해의 책임, 인간과 기계의 정치적 관계가 모두 다시 검토 대상이 된다. 비인간 지능을 사회 질서의 바깥에 둘 것인지 안에 들일 것인지, 들인다면 권리와 책임의 경계선을 어디에 그을 것인지를 이제 물어야 한다.

SNU 그랜드퀘스트 챌린지(연구지원프로그램) 개요

대 상

연구책임자: 서울대학교 연구자(교원 및 연구원)

공동연구자: 국내 · 외 모든 연구자

형 태

개인 또는 팀(단일학제, 융합형 가능)

과제 체계 및 지원 규모

1

Stage 탐색연구

도전적 질문의 가능성과 연구방향의 검증



목 적

| 질문의 탐색 및 가설 검증



기간 및 규모

| 6개월 / 최대 5천만원



특 징

| 다수 과제 지원 · 블라인드 심사



일정(예정)

| 2026년 7월 15일 중 공고

6개월

최대
5천만원

다수 과제 지원

2

Stage 도전연구

구체화된 질문에 대한 혁신적 돌파구 창출



목 적

| 질문의 완성 및 중 · 장기 난제 해결



대 상

| 2026년 탐색연구과제



기간 및 규모

| 최대 5년 / 연간 최대 5억원



특 징

| 소수 과제 지원 집중 · 연구자육성 전면 보장



일정(예정)

| 2027년 5월 전환심사

최대
5년

연 최대
5억원

소수 과제 집중 지원

SNU 그랜드퀘스트 챌린지 차별점

?

기존 연구지원



문제 해결 중심

우수한 연구



연구자 중심 선발

학력 / 이력 평가



성공 중심 평가

결과에 대한 책임



단일 최적 경로

선택형 연구



실패 = 탈락

목표달성 완성도

!

SUN 그랜드퀘스트 챌린지



질문 설계 중심

새로운 연구



질문 중심 선발

새로움과 도전성 평가



진화 중심 평가

과정에 대한 책임



다중 탐색 경로

병렬형 연구



실패 = 자산

변화의 타당성

VS

“ SNU 그랜드퀘스트 챌린지는
도전적 질문에서 새로운 미래를 만듭니다 ”



질문

세상을 바꾸는
질문을 던집니다.



탐색

질문의 가능성을
자유롭게 탐색합니다.



도전

대담한 도전으로
한계를 넘어갑니다.



새로운 미래

인류와 사회에
새로운 가치를 만듭니다.

SNU Grand Quest Forum

Toward Agenda Setters



GRAND
QUEST

✉ Contact | grandqeust@snu.ac.kr

🌐 Homepage | <http://grandquest.snu.ac.kr>