

자연지능은 어떻게 새로운 질문을 만들어내는 가?

-자체제작 Ai 모델 Qpique을 활용한 자연지능의 질문 형성 메커니즘 탐구

민하원
조지이자벨라

《 요약 》

본 연구는 인간이 평범한 현상에서 새로운 질문을 만들어내는 과정을 하나의 인지적 메커니즘으로 설명할 수 있는지를 탐구하였다. 이를 위해 먼저 심리학·신경과학·발달심리학·과학철학의 선행연구를 조사하여, 자연지능의 질문 형성이 '예측오차 감지-가치 판단'에 따른 선택-설명적 재구성이라는 세 단계를 거친다는 가설을 세웠다. 이어서 이 가설을 검증하기 위해, 자체 제작한 인공지능망 기반 웹 도구 큐픽(Qpique)을 활용하여 예측오차 감지 단계만 갖춘 모델①과 학습진전에 기반한 가치판단·선택 단계까지 갖춘 모델②를 동일한 조건에서 비교하는 통제 실험을 설계·실행하였다. 실험 결과 모델②는 학습이 불가능한 노이즈 대상을 회피하고, 학습 가치가 있는 대상을 더 오래 지속적으로 탐구하며, 실제로 은닉된 구조를 해소하는 데까지 이르러 노이즈 낭비율·지속추구 비율·평균 질문 길이·구조 해소율·구조 몰입 시점 다섯 개 지표에서 모델①보다 우수한 결과를 보였고, 이는 예측오차 감지에 가치 판단에 따른 선택 과정이 더해질 때 비로소 탐구 지향적 질문이 형성된다는 두 가설을 함께 뒷받침하는 결과였다.

I. 서론

연구 동기

인류의 과학은 새로운 질문에서 시작되어 발전해 왔다. 인간은 주변에서 일어나는 현상을 단순히 관찰하는 데 그치지 않고, 기존의 지식으로는 충분히 설명되지 않는 현상에 의문을 제기하며 새로운 탐구를 이어 왔다. 뉴턴은 사과가 떨어지는 현상과 달이 지구로 떨어지지 않는 현상을 하나의 문제로 바라보며 만유인력 법칙을 제안하였고, 다윈은 갈라파고스 제도의 생물들이 서로 다른 특징을 갖는 이유를 질문하며 진화론을 정립하였다. 이처럼 과학의 발전은 단순히 새로운 정보를 축적한 결과가 아니라, 평범한 현상 속에서 새로운 질문을 발견하고 이를 탐구하는 과정에서 이루어져 왔다.

그러나 이러한 과학사의 사례를 살펴보면 한 가지 의문이 생겼다. 사람은 일상에서 수없이 많은 현상을 경험하지만, 왜 어떤 현상은 그냥 지나치고 어떤 현상은 새로운 질문으로 이어지는 것일까? 또한 같은 현상을 관찰하더라도 왜 어떤 사람은 기존의 지식을 그대로 받아들이는 반면, 어떤 사람은 새로운 탐구를 시작하는 것일까?

기존 연구에서는 예측오차, 호기심, 내재적 동기, 과학적 발견의 과정 등이 각각 연구되어 왔지만, 이러한 요소들이 어떻게 연결되어 인간의 질문 생성으로 이어지는지는 통합적으로 설명되지 않는다는 점에 주목하였다. 이에 본 연구에서는 자연지능에서 새로운 질문이 생성되는 과정을 하나의 인지적 메커니즘으로 설명할 수 있는지 탐구하고자 하였다.

2. 연구 목적

이 연구에서는 자연지능에서 새로운 질문이 생성되는 메커니즘을 탐구하고, 생성형 인공지능 실험을 통해 그 타당성을 검증하고자 한다. 구체적인 연구 문제는 다음과 같다.

- 1) 자연지능에서 새로운 질문이 생성되는 메커니즘을 규명하고자 한다.
- 2) 생성형 인공지능 실험을 통해 자연지능의 질문 생성 과정에 대한 가설의 타당성을 확인하고자 한다.

II. 본론

1. 이론적 배경

가. 탐구지향적 질문

탐구 지향적 질문(Inquiry-oriented Question)은 단순히 부족한 정보를 확인하기 위한 질문을 넘어, 현상의 원인이나 관계, 작동 원리를 밝히기 위해 추가적인 탐구를 요구하는 질문을 의미한다. 일반적인 정보 확인 질문은 기존의 지식이나 자료를 통해 답을 얻는 것을 목적으로 하지만, 탐구 지향적 질문은 현재의 지식만으로 충분히 설명되지 않는 현상을 대상으로 하며 관찰, 자료 수집, 가설 설정, 실험 및 검증 등의 과정을 유발한다. 따라서 이러한 질문은 기존 지식을 단순히 보충하는 데 그치지 않고, 기존의 설명을 검토하거나 수정·확장하는 학습으로 이어질 수 있다

나. 자연지능

자연 지능(Natural Intelligence)은 생명체가 환경을 인식하고 경험을 통해 학습하며 변화하는 환경에 적응하는 능력을 의미한다. 특히 인간의 자연 지능은 단순히 주어진 문제를 해결하는 수동적 역할을 넘어, 주변 현상에서 새로운 문제를 발견하고 질문을 만들어 내며 이를 주도적으로 탐구하는 고유한 특징을 보인다. 즉, 자연 지능은 기존의 지식과 경험을 바탕으로 세계를 예측하고, 그 과정에서 나타나는 현실과의 불일치를 새로운 질문과 탐구로 발전시키는 역동적인 개념이다. 따라서 본 연구에서 자연 지능은 인공지능과 대비되는 생물학적 지능 전체를 포괄하기보다, 인간이 새로운 질문을 생성하고 이를 통해 지식을 스스로 확장해 나가는 과정을 설명하기 위한 핵심 연구 대상으로 정의된다.

다. 예측오차

예측오차란 뇌가 기존의 지식과 경험을 바탕으로 형성한 내부 모델(Internal Model)의 예측과 실제 감각 입력 사이에서 발생하는 불일치를 의미한다. 예측처리 이론(Predictive Processing)에 따르면 뇌는 외부 세계를 수동적으로 받아들이는 것이 아니라 지속적으로 세계를 예측하고, 실제 감각 정보와 비교하여 예측오차를 발생시킨다. 예측오차는 현재의 내부 모델이 실제 세계를 충분히 설명하지 못하고 있음을 나타내며, 뇌가 새로운 정보를 받아들이거나 기존의 모델을 수정하도록 하는 신호로 작용한다. 이처럼 예측오차는

예측과 관찰 사이의 차이를 감지하고 이를 바탕으로 학습과 모델의 수정이 이루어지는 과정의 중요한 요소이다.

라. 인공지능

인공지능(Artificial Intelligence, AI)은 인간의 자연 지능이 가진 인식, 학습, 적응 능력을 컴퓨터 시스템을 통해 인공적으로 구현한 기술로, 대규모 데이터와 정교한 알고리즘을 바탕으로 최적의 패턴을 찾아내어 주어진 문제를 해결하는 데 특화되어 있다.

인공지능은 데이터를 통해 스스로 규칙을 찾아 성능을 향상시키며 학습한다. 대표적인 학습 방식에는 지도학습, 비지도학습, 강화학습이 있다.

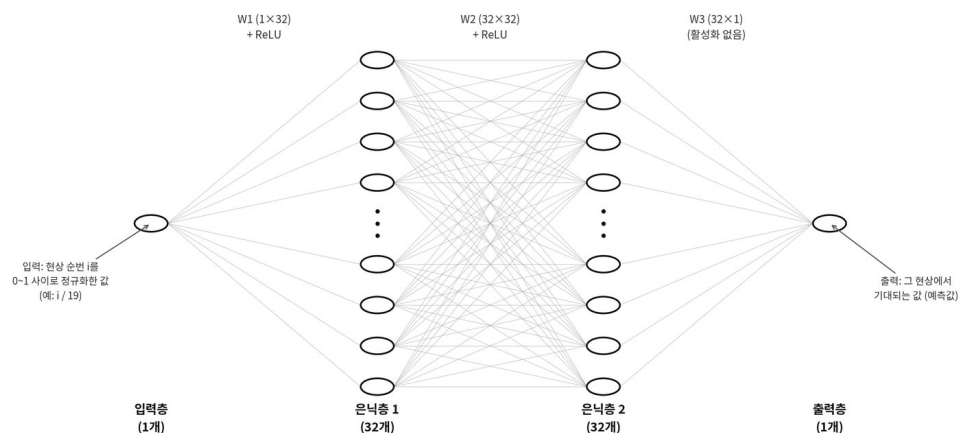
지도학습은 입력 데이터와 함께 정답(Label)이 주어진 상태에서 학습하는 방법이다. 인공지능은 예측한 결과와 실제 정답을 비교하여 오차를 계산하고, 그 오차를 줄이도록 가중치를 반복적으로 수정한다. 예를 들어, 수천 장의 고양이와 강아지 사진에 각각 '고양이', '강아지'라는 정답을 함께 제공하면, 인공지능은 오답을 수정하는 과정을 반복하며 두 동물을 구별하는 방법을 학습한다.

비지도학습은 정답 없이 데이터를 학습하는 방법이다. 인공지능은 데이터의 공통된 특징과 패턴을 스스로 찾아 비슷한 데이터끼리 묶거나 구조를 발견한다. 예를 들어 고객의 구매 데이터를 분석하여 비슷한 소비 성향을 가진 사람들을 여러 집단으로 분류하는 데 활용된다.

강화학습은 인공지능이 환경과 상호작용하면서 보상(Reward)을 최대화하는 행동을 학습하는 방법이다. 올바른 행동에는 보상을 받고, 잘못된 행동에는 보상이 줄어들거나 벌점을 받으면서 시행착오를 통해 최적의 전략을 찾아간다. 예를 들어 알파고는 수많은 바둑 대국을 스스로 반복하며 승리했을 때는 높은 보상을, 패배했을 때는 낮은 보상을 받아 점차 더 좋은 수를 선택하도록 학습하였다.

라. 인공신경망

인공신경망이란 사람 뇌의 뉴런이 서로 연결되어 정보를 처리하는 방식을 단순화하여 흉내 낸 계산 구조이다. 여러 개의 인공 뉴런(노드)이 층을 이루며 연결되고, 각 연결에는 가중치(weight)가 부여되어 입력값이 이 가중치들을 거치며 변환되어 출력값으로 이어진다. 이 중에서도 입력층-은닉층-출력층이 순서대로 완전히 연결된 가장 기본적인 구조를 다층 퍼셉트론(Multi-Layer Perceptron, MLP)이라 부르며, 층과 층 사이에는 비선형 활성화 함수(이 연구에서는 ReLU)를 두어 단순한 직선 관계로는 표현할 수 없는 복잡한 패턴까지 학습할 수 있게 한다.



[그림 1] Qpique AI 모델의 인공신경망 구조

2. 연구 방법

가. 가설

1) 인간의 자연지능은 예측오차를 감지한 후, 새로움·학습 가능성·탐색 전략·개인적 가치와 목표를 바탕으로 일부 오차를 선택하고, 이를 탐구 지향적 질문으로 발전시켜 새로운 설명과 지식을 형성한다.

2) 예측오차 감지-가치 판단에 따른 선택-탐구 지향적 질문 생성의 과정이 실제로 질문 생성에 필요한 메커니즘이라면, 이를 각각 통제된 인공지능 실험에서 해당 단계의 유무에 따라 생성되는 질문의 특성이 달라질 것이다.

나. 탐구 1

1) 탐구 과정

가) 탐구 순서

- (1) 연구 주제와 관련된 선행연구 및 과학사 사례를 조사한다.
- (2) 각 자료에서 인간이 현상을 관찰하고 질문을 생성하는 과정을 분석한다.
- (3) 자료에서 반복적으로 나타나는 공통 요소를 추출한다.
- (4) 추출한 요소를 '예측오차 감지-가치 판단에 따른 선택-설명적 재구성'의 단계로 분류한다.
- (5) 각 단계가 질문 생성 과정에서 어떤 역할을 하는지 비교·분석한다.
- (6) 분석 결과를 종합하여 자연지능의 질문 생성 메커니즘에 대한 가설을 설정한다.

2) 탐구 결과

(1) 선행연구 및 과학사 사례 조사

질문 형성 메커니즘과 관련하여 심리학, 신경과학, 과학철학, 발달심리학 네 분야의 선행 연구를 조사하였다. 심리학에서는 Loewenstein(1994)의 정보격차 이론과 Berlyne(1960)의 최적각성 이론을, 신경과학에서는 Schultz(1997)(2) 각 자료에서 인간이 현상을 의도파민 보상예측오차 이론과 Friston의 예측부호화 이론, Kang 외(2009)의 호기심-보상회로 fMRI 연구를 조사하였다. 발달심리학에서는 Oudeyer 외(2007)의 학습진전 개념과 Aston-Jones & Cohen(2005)의 탐험-활용 조절 이론, Chouinard(2007)의 아동 질문 형성 연구를 조사하였다. 과학철학에서는 Kuhn(1962)의 이상현상과 패러다임 전환 이론을 조사하였다. 아울러 이러한 이론들이 실제 인간의 질문 형성 과정에서 어떻게 나타나는지를 확인하기 위해 뉴턴의 만유인력 발견 사례를 과학사적 사례로 분석하였다.

관찰하고 질문을 생성하는 과정 분석

Loewenstein의 정보격차 이론과 Schultz·Friston의 예측오차 이론은 각각 심리학과 신경과학의 층위에서 같은 현상을 설명한다. 인간이 특정 상황에 대해 세운 예측과 실제 관찰 사이의 불일치가, 심리적으로는 '정보격차의 느낌'으로, 신경학적으로는 '도파민 보상예측오차 신호'로 나타난다는 것이다. 이 불일치가 감지된 이후, 그 수많은 불일치 중 무엇을

계속 추구할 것인지는 별도의 선별 과정을 거친다. Berlyne의 최적각성 이론은 지나치게 익숙하거나 낯선 오차는 추구되지 않고 적정 수준의 오차만 선택됨을 알 수 있다.

Oudeyer의 학습진전 개념과 Aston-Jones & Cohen의 탐험-활용 조절 이론은 이 선별이 '계속 파고들수록 오차가 줄어드는가'라는 온라인 학습 신호에 의해 이루어짐을 보였다. 선별을 통과한 오차를 해소하려는 시도는 Kang 외의 연구에서 확인된 도파민-해마 회로에 의해 지속적으로 유지되며, 이 지속적 탐구는 Kuhn이 설명한 것처럼 기존 개념체계 자체를 재구성하는 데까지 이를 수 있다. Chouinard의 연구는 이렇게 형성된 질문이 단순한 정보 습득이 아니라 인지발달의 기제로 기능함을 보였다. 뉴턴의 사례는 이 모든 단계가 실제 한 인물의 탐구 과정 속에서 순차적으로 나타난 예로 분석되었다.

(3) 반복적으로 나타나는 공통 요소 추출

첫째, 인간은 자신의 내부 모델이 세운 예측과 실제 관찰 사이의 구체적인 불일치를 지속적으로 감지한다. 둘째, 감지된 모든 불일치가 추구 대상이 되는 것은 아니며, 계속 파고들었을 때 실제로 이해가 깊어지는지(오차가 줄어드는지)에 따라 선별적으로 추구가 된다. 셋째, 선별된 불일치를 해소하려는 시도는 내재적 보상에 의해 지속적으로 유지된다. 넷째, 이러한 지속적 탐구의 결과는 기존 정보의 단순한 재조합이 아니라 새로운 개념적 구조의 형성으로 이어진다.

(4) 세 단계(예측오차 감지 - 가치 판단에 따른 선택 - 설명적 재구성)로 분류

1단계는 '예측오차 감지' 단계로, 내부 모델이 세운 예측과 실제 관찰 사이의 불일치가 탐지된다(Loewenstein, Schultz, Friston). 2단계는 '가치 판단에 따른 선택' 단계로, 감지된 수많은 오차 중 계속 탐구했을 때 실제로 학습이 진전되는 대상만 선별된다(Berlyne, Oudeyer, Aston-Jones & Cohen). 3단계는 '설명적 재구성' 단계로, 선별된 오차를 해소하려는 시도가 지속적 동기에 의해 유지되며 새로운 개념 구조로 이어진다(Kang 외, Kuhn, Chouinard).

(5) 각 단계의 역할 비교·분석

세 단계는 서로 다른 기능을 수행하며 순차적으로 작동한다.

1단계는 무수히 발생하는 예측 오차 가운데 인지적으로 감지 가능한 불일치를 걸러내는 '탐지' 기능을 수행한다.

2단계는 감지된 오차 중 '계속 파고들 가치가 있는가'를 판단하는 '선별' 기능을 수행한다. 이 가치 판단은 단일한 기준이 아니라 네 가지 하위 기준이 함께 작동하는 복합적 평가로 나타난다. 자극이 충분히 흥미로운지를 묻는 새로움(Berlyne의 최적각성 이론), 지금의 능력으로 성장할 수 있는지를 묻는 학습 가능성(Oudeyer 외의 학습진전 개념), 기존 방식을 유지할지 새로운 탐구를 시작할지를 묻는 행동 전략(Aston-Jones & Cohen의 탐험-활용 조절 이론), 그리고 그 문제가 자신에게 의미 있는지를 묻는 개인적 가치와 목표(Ryan & Deci의 자기결정성이론, Wigfield & Eccles의 기대-가치 이론)가 그것이다.

판단 기준	이론적 근거	핵심 질문
새로움	Berlyne(1960)최적각성 이론	충분히 흥미로운가?
학습 가능성	Oudeyer 외(2007)학습진전 개념	지금의 능력으로 성장할 수 있는가?
행동 전략	Aston-Jones & Cohen(2005)탐험-활용 조절 이론	기존 방식을 유지할 것인가, 새로운 탐구를 시작할 것인가?
개인적 가치와 목표	Ryan & Deci(2000) 자기결정성 이론/Wigfield & Eccles(2000) 기대가치 이론	나에게 의미 있는 문제인가?

[표 1] 판단 기준

3단계는 선별을 통과한 오차에 대한 탐구를 얼마나 오래, 얼마나 깊이 유지할 것인지를 조절하는 '지속·생성' 기능을 수행하며, 이 단계의 산출물의 질에 따라 단순한 정보 습득에 그치는 경우와 뉴턴의 만유인력 발견과 같은 근본적 개념 재구성에 이르는 경우로 나뉜다.

(6) 분석 결과를 종합하여 자연지능의 질문 생성 메커니즘에 대한 결론 도출

이상의 선행연구와 과학사 사례를 종합적으로 분석한 결과, 자연지능의 질문 형성은 단일한 순간적 통찰이 아니라 '예측오차 감지 - 가치 판단에 따른 선택 - 설명적 재구성'이라는 위계적 단계를 거치는 과정임을 확인할 수 있었다. 질문은 내부 모델이 세운 예측과 실제 관찰 사이의 불일치를 감지하는 데서 출발하지만, 감지된 모든 불일치가 질문으로 발전하는 것은 아니며, 계속 탐구할수록 실제로 이해가 깊어지는 대상만이 선별적으로 추구된다는 점, 그리고 이렇게 선별된 탐구가 도파민 보상회로에 의한 지속적 동기 부여를 통해 유지될 때 비로소 기존 지식의 단순한 재조합을 넘어 새로운 개념적 구조로 이어진다는 점이 심리학·신경과학·과학철학의 서로 다른 이론들에 걸쳐 공통적으로 확인되었다.

이는 평범한 현상에서 새로운 질문을 이끌어내는 능력이 우연적이거나 설명 불가능한 과정이 아니라, 예측-오차 계산이라는 식별 가능한 정보처리 원리에 기반한다는 것을 보여준다. 나아가 이 메커니즘이 생물학적 기질(뇌)에만 국한되지 않고 정보처리적 원리로 기술 가능하다는 점은, 동일한 원리를 인공지능 시스템에 구현했을 때 인간과 유사한 자발적 탐구 행동을 유도할 수 있는 가능성 역시 뒷받침한다. 다만 이 단계는 직접적인 실험이 아닌 선행연구 및 문헌 조사를 통해 도출된 결론이라는 한계를 지닌다.

다. 탐구 2

1) 탐구 과정

가) 탐구 순서

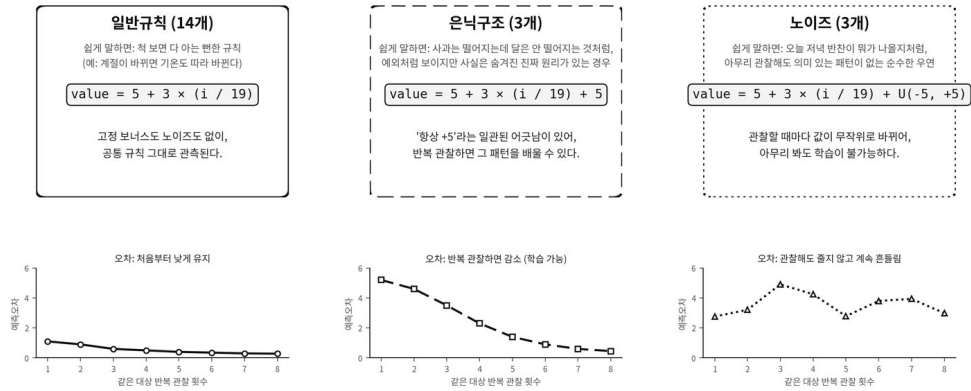
(1) 설정한 가설을 검증하기 위해, 질문 형성 메커니즘의 유무만을 다르게 하고 나머지 조건은 모두 동일하게 통제하는 비교 실험을 설계한다.

(2) 비교 실험의 독립변인(질문 형성 메커니즘의 유무), 종속변인(생성된 질문의 특성), 통제변인(신경망 구조, 초기 가중치, 관찰 대상 목록, 관찰 횟수)을 설정한다.

(3) 두 모델이 공통으로 관찰할 가상의 현상들을 일반규칙·은닉구조·노이즈 세 유형으로 설계하여 탐구 환경을 구성한다.

현상 유형별 값 생성 방식 — env.py PhenomenaWorld.observe(i)

모든 현상은 공통 규칙 $value = 5 + 3 \times (\text{순번 } i / 19)$ 위에서 출발한다. 여기에 유형별로 서로 다른 '어긋남'이 더해진다.



[그림 2] 현상 유형별 값 생성 방식

- (4) 사전학습된 언어모델을 사용하지 않고, 관찰을 통해 스스로 예측을 학습하는 작은 인공 신경망(예측 모델)을 직접 설계하고 구현한다.
- (5) '예측오차 감지' 단계는 두 모델에 동일하게 적용하고, '가치 판단에 따른 선택' 단계(예측오차 감소 추이, 즉 학습진전에 기반한 선택 규칙)는 모델②에만 추가로 구현한다.
- (6) 생성된 질문을 '가치판단 적중률', '지속추구 성향', '실제 학습 성과(구조 해소 여부)', '언어적 질(구체성·설명지향성)'의 네 범주로 평가할 수 있는 기준표를 설계한다.

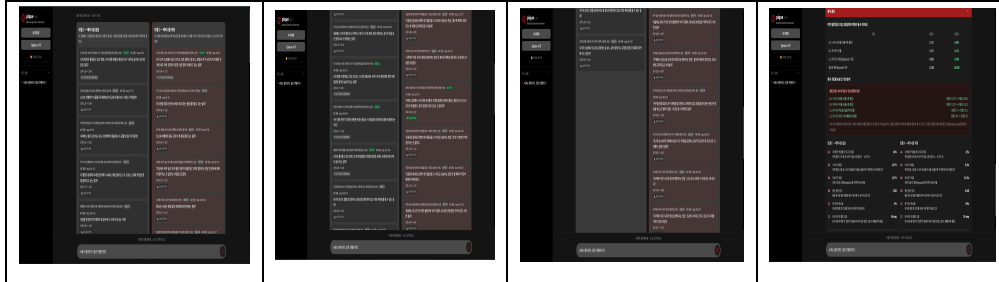
범주	코드	지표명	측정 내용
A. 가치판단 적중률	A1	가치판단 적중률 (은닉구조 비중)	생성된 질문 중 은닉구조 대상을 향한 비율
	A2	노이즈 낭비율	전체 관찰 중 노이즈 대상을 관찰한 비율 (낮을수록 좋음)
B. 지속추구 성향	B1	지속추구 비율	전체 관찰 스텝 중 '질문(episode)'을 이루며 보낸 비율
	B2	평균 질문 길이	질문 하나당 평균 재방문 횟수 (깊수록 길이 파고든 것)
C. 실제 학습 성과	C1	은닉구조 해소율	은닉형 질문 중 오차를 일정 비율 이상 줄인 비율
	C2	은닉구조 첫 물음 시점	은닉구조에 대한 첫 '질문'이 형성되기까지 걸린 스텝 수
D. 언어적 질	D1	구체성	수치·구체적 사례·비교표현 등이 문장에 포함되는지 (규칙기반 /5점)
	D2	설명지향성	원인·과정·메커니즘을 묻는 표현이 있는지 (규칙기반 /5점)

[표 2] Qpique AI 실행 평가기준표

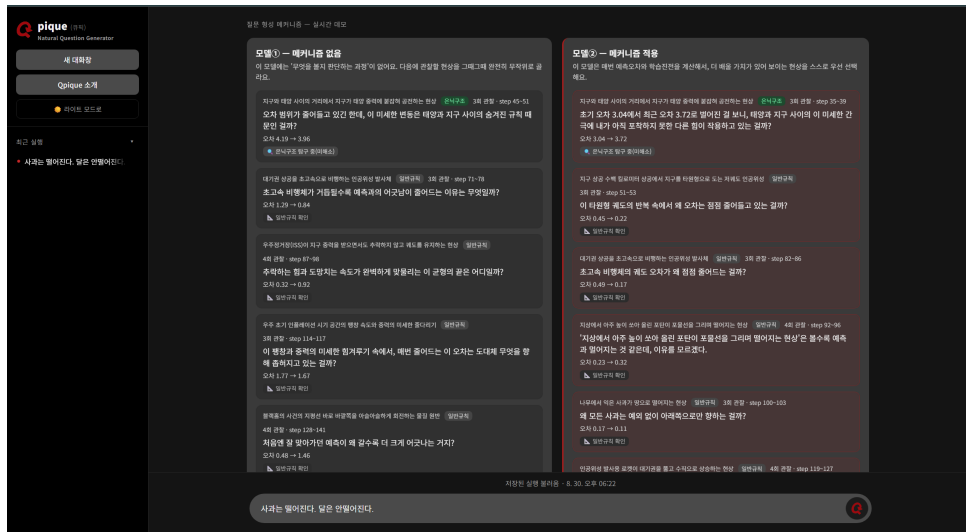
- (7) 서로 다른 난수 시드로 두 모델을 반복 실행하여 데이터를 수집하고, 기준표에 따라 결과를 채점·비교한다.
- (8) 다양한 현상 입력에 대해 실험을 반복적으로 재현·검증할 수 있도록, 실행 결과를 기준표에 따라 자동으로 분석해 보여주는 웹 기반 비교 도구를 제작한다.

2) 탐구 결과

Qpique AI 모델 작동 UI 전체 모습



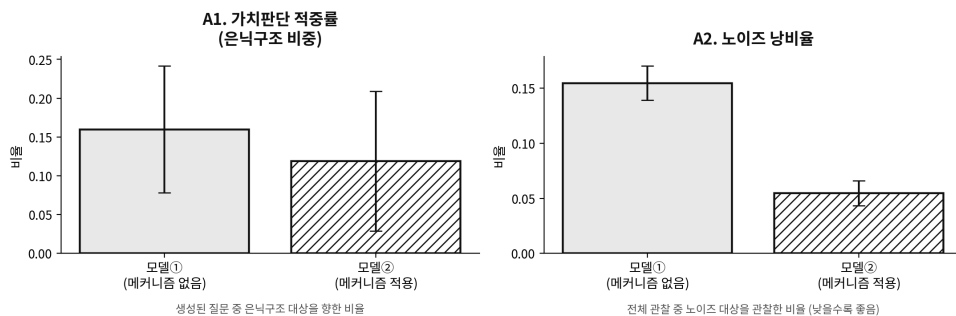
[표 3] Qpique AI UI 전체 모습



[그림 3] Qpique AI UI 세부 모습

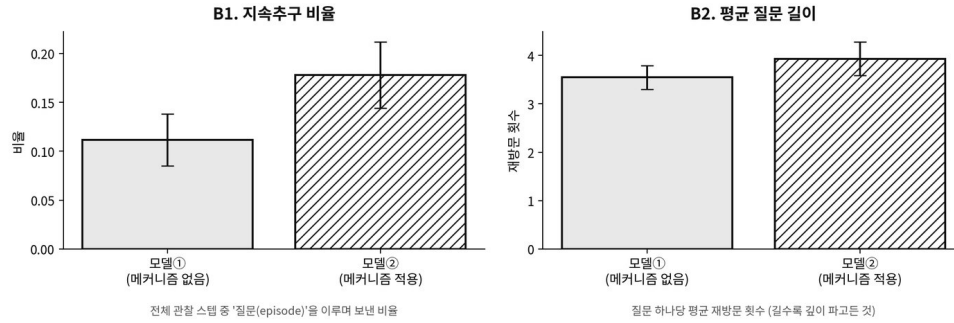
위 [표 3]과 [그림 3]은 하나의 현상을 입력했을 때 모델①(메커니즘 없음)과 모델②(메커니즘 적용)가 각각 만들어낸 질문(episode)을 나란히 비교할 수 있도록 구성한 큐픽의 실행 화면이다. 각 말풍선에는 관찰 대상이 된 현상의 유형(일반규칙·은닉구조·노이즈), 관찰 횟수와 관찰 구간, 실제로 생성된 질문 문장, 관찰 전후 오차 변화, 언어적 질 등급이 함께 표시된다. 동일한 현상·동일한 초기 조건에서 출발했음에도 두 모델이 서로 다른 질문을 만들어낸다는 점을 이 화면에서 직접 확인할 수 있다.

A. 가치판단 적중률



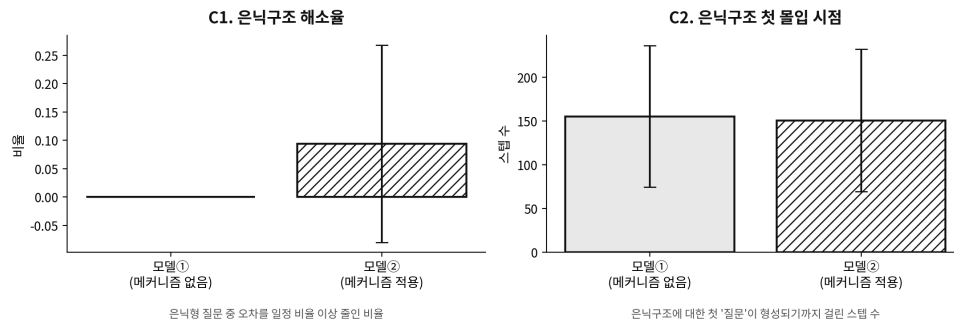
[그래프 1] Qpique AI의 가치판단 적중률 비교

B. 지속추구 성향



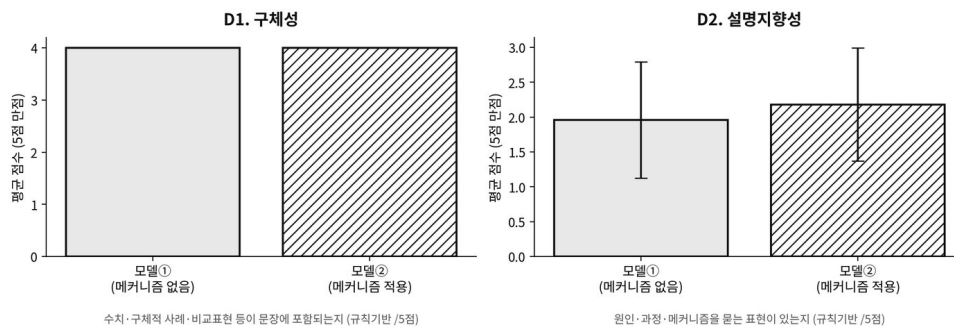
[그래프 2] Qpique AI의 지속추구 성향 비교

C. 실제 학습 성과



[그래프 3] Qpique AI의 실제 학습 성과 비교

D. 언어적 질



[그래프 4] Qpique AI의 언어적 질

[그래프 1]의 가치판단 적중률에서는 두 지표가 서로 다른 방향을 보였다. A2(노이즈 낭 비율)는 모델①이 0.155, 모델②가 0.055로 모델②가 뚜렷하게 낮아, 학습이 불가능한 대상(노이즈)을 가려내고 회피하는 능력은 가설대로 확인되었다. 반면 A1(생성된 질문 중 은닉구조 비중)은 모델①이 0.160, 모델②가 0.119로 오히려 모델①이 더 높게 나타나 가설과 반대 방향이었다. 이는 모델②가 은닉구조만 편식하듯 골라내는 것이 아니라, 학습 진전이 남아 있는 대상이면 일반규칙 현상도 상당수 함께 탐구하기 때문으로 해석된다 (일반규칙 적중률 0.856 vs 0.732). 즉 모델②의 가치판단은 '은닉구조만 향한다'는 단순한 형태가 아니라 '아직 배울 것이 남은 대상을 향한다'는 형태에 가까우며, 이러한 차이가 단일 비율 지표인 A1에는 오히려 불리하게 반영된 것으로 볼 수 있다.

[그래프 2]의 지속추구 성향은 두 지표 모두 가설을 지지했다. B1(지속추구 비율)은 모델 ①이 0.112, 모델 ②가 0.178로 모델 ②가 더 높아, 같은 대상을 몰아서 반복 관찰하는 목표 지향적 행동이 메커니즘 적용 시 더 자주 나타났다. B2(평균 질문 길이·재방문 횟수) 역시 모델 ①이 3.55회, 모델 ②가 3.93회로 모델 ②가 더 길게 파고들었다.

[그래프3]의 실제 학습 성과에서는 차이가 가장 뚜렷했다. C1(은닉구조 해소율)은 모델 ①이 0.0, 모델 ②가 0.094로, 8회의 반복 실험 동안 모델 ①은 단 한 번도 은닉구조의 오차를 임계치 이상 줄이지 못한 반면 모델 ②는 실제로 이를 해소한 사례가 있었다. C2(은닉구조 첫 몰입 시점)는 모델 ①이 155.3스텝, 모델 ②가 150.6스텝으로 모델 ②가 더 이르게 은닉구조에 몰입하기 시작했다.

[그래프 4]의 언어적 질은 두 모델 사이 차이가 가장 작았다. D1(구체성)은 두 모델 모두 평균 4.0점으로 동일했는데, 이는 질문 문장의 구체성이 메커니즘의 유무보다는 문장 생성 템플릿 자체에 의해 결정되기 때문으로 보인다. D2(설명지향성)는 모델 ①이 1.95점, 모델 ②가 2.14점으로 모델 ②가 근소하게 높았으나 표준편차 범위가 겹쳐 뚜렷한 차이로 보기는 어렵다.

종합하면, A2·B1·B2·C1·C2 다섯 지표는 모두 가설(메커니즘이 있으면 학습 가치가 있는 대상을 더 잘 선별하고, 더 오래 몰입하며, 실제 학습 성과로 이어진다)을 지지하는 방향으로 나타났다. A1은 반대 방향, D1·D2는 유의미한 차이가 없었다. 즉 메커니즘의 효과는 질문 대상을 얼마나 정확히 맞히는가(비율)보다, 그 선택이 실제 학습으로 이어지는가(지속성과 성과)에서 더 분명하게 나타났다. 다만 A1의 결과는 단순 비율 지표만으로는 가치 판단의 질을 온전히 포착하기 어렵다는 점을 보여준다.

III. 결론

본 연구는 두 가지 가설을 세우고 이를 검증하였다.

첫 번째 가설은, 인간의 자연지능이 예측오차를 감지한 뒤 새로움·학습 가능성·탐색 전략·개인적 가치와 목표를 바탕으로 그중 일부 오차를 선택하여 탐구 지향적 질문으로 발전시키고, 이를 통해 새로운 설명과 지식을 형성한다는 것이었다. 큐픽의 모델 ②는 이 과정 가운데 예측오차 감지와, 오차가 줄어드는 속도(학습진전)를 통한 가치판단·선택 단계를 계산 가능한 형태로 구현한 것이다. 실험 결과 모델 ②는 노이즈처럼 학습이 불가능한 오차는 회피하고(A2), 학습 가치가 있는 대상은 반복적으로 파고들어(B1·B2) 실제로 오차를 해소하는 데까지 이르렀다(C1·C2). 이는 '예측오차 감지 → 가치판단에 따른 선택 → 탐구 지향적 질문'이라는 단순화된 계산 모델만으로도 새로움과 학습 가능성을 반영한 탐구 행동이 실제로 나타날 수 있음을 보여주며, 첫 번째 가설이 제시한 자연지능의 질문 생성 과정이 컴퓨터로 구현 가능한 메커니즘으로서 설득력을 가진다는 것을 뒷받침한다. 다만 본 실험에서의 가치판단은 학습진전이라는 단일 지표로 단순화된 것으로, 인간이 지닌 개인적 가치·목표나 다양한 탐색 전략까지 온전히 재현한 것은 아니므로 이는 직접적 증명이라기보다 간접적인 근거로 보아야 한다.

두 번째 가설은, 예측오차 감지-가치판단에 따른 선택-탐구 지향적 질문 생성의 과정이 실제로 질문 생성에 필요한 메커니즘이라면, 이 과정을 통제된 인공지능 실험에서 해당 단계의 유무에 따라 생성되는 질문의 특성이 달라질 것이라는 것이었다. 이는 메커니즘의 유무만을 다르게 설정하고 나머지 조건은 모두 동일하게 통제된 모델 ①·② 비교 실험을 통해 직접 검증되었다. 노이즈 낭비율(A2), 지속추구 비율(B1), 평균 질문 길이(B2), 은닉구조 해소율(C1), 은닉구조 첫 몰입 시점(C2) 다섯 개 지표에서 메커니즘의 유무에 따라 생성되는 질문의 특성이 유의미하게 달라졌으므로, 두 번째 가설 역시 채택되었다고 볼 수 있다. 다만 A1(생성된 질문 중 은닉구조 비중)처럼 메커니즘의 유무가 예상과 다른 방향으로 작용한 지표도 있어, 메커니즘의 유무가 질문의 모든 특성을 동일한 방향으로 바꾸는 것은 아니며 그 영향은 지표에 따라 다르게 나타날 수 있다는 점도 함께 확인되었다.

다만 본 연구에는 몇 가지 한계도 있다. 첫째, 위에서 언급한 A1의 경우처럼 단순 비율 지표만으로는 가치판단의 질을 온전히 포착하기 어렵다. 둘째, C1·C2와 같은 지표는 표준편차가 커서 8회의 반복 실험만으로는 결과의 안정성을 완전히 보장하기 어렵다. 셋째, 본 실험은 20개의 현상과 제한된 노이즈 비율이라는 단순화된 환경에서 이루어졌으므로, 더 복잡하고 현실적인 환경에서도 동일한 경향이 유지되는지는 추가로 확인할 필요가 있다. 넷째, 언어적 질(D1·D2)은 규칙 기반 채점 방식으로 측정되어, 실제 사람이 느끼는 질문의 깊이나 참신함을 완전히 반영하지는 못한다.

이러한 한계를 보완하기 위해, 후속 연구에서는 현상의 개수나 노이즈 비율을 다양하게 바꾸어 실험함으로써 결과가 특정 조건에서만 나타나는 것은 아닌지 확인할 필요가 있다. 또한 학습진전 외에 개인적 가치·목표나 탐색 전략을 반영할 수 있는 요소를 모델에 추가해 첫 번째 가설을 더 직접적으로 검증하거나, 모델이 생성한 질문을 실제 사람이 만든 질문과 비교하고 규칙 기반 채점을 넘어 사람 또는 LLM을 활용한 평가로 언어적 질을 보완한다면, 인공지능경망이 만들어내는 질문이 실제 탐구 활동에서도 유의미한 수준인지를 더 정교하게 검증할 수 있을 것이다.

IV. 참고문헌

- Aston-Jones, G., & Cohen, J. D. (2005). An integrative theory of locus coeruleus-norepinephrine function: Adaptive gain and optimal performance. *Annual Review of Neuroscience*, 28, 403–450.
- Berlyne, D. E. (1960). *Conflict, arousal, and curiosity*. McGraw-Hill.
- Chouinard, M. M. (2007). Children's questions: A mechanism for cognitive development. *Monographs of the Society for Research in Child Development*, 72(1), 1–129.
- Friston, K. (2005). A theory of cortical responses. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 360(1456), 815–836.
- Kang, M. J., Hsu, M., Krajbich, I. M., Loewenstein, G., McClure, S. M., Wang, J. T., & Camerer, C. F. (2009). The wick in the candle of learning: Epistemic curiosity activates reward circuitry and enhances memory. *Psychological Science*, 20(8), 963–973.
- Kuhn, T. S. (1962). *The structure of scientific revolutions*. University of Chicago Press.
- Loewenstein, G. (1994). The psychology of curiosity: A review and reinterpretation. *Psychological Bulletin*, 116(1), 75–98.
- Oudeyer, P.-Y., Kaplan, F., & Hafner, V. V. (2007). Intrinsic motivation systems for autonomous mental development. *IEEE Transactions on Evolutionary Computation*, 11(2), 265–286.
- Ryan, R. M., & Deci, E. L. (2000). Self-determination theory and the facilitation of intrinsic motivation, social development, and well-being. *American Psychologist*, 55(1), 68–78.
- Schultz, W., Dayan, P., & Montague, P. R. (1997). A neural substrate of prediction and reward. *Science*, 275(5306), 1593–1599.
- Wigfield, A., & Eccles, J. S. (2000). Expectancy-value theory of achievement motivation. *Contemporary Educational Psychology*, 25(1), 68–81.

