

2025 연구실 세미나

HALO: Hierarchical Autonomous Logic-Oriented Orchestration for Multi-Agent LLM Systems

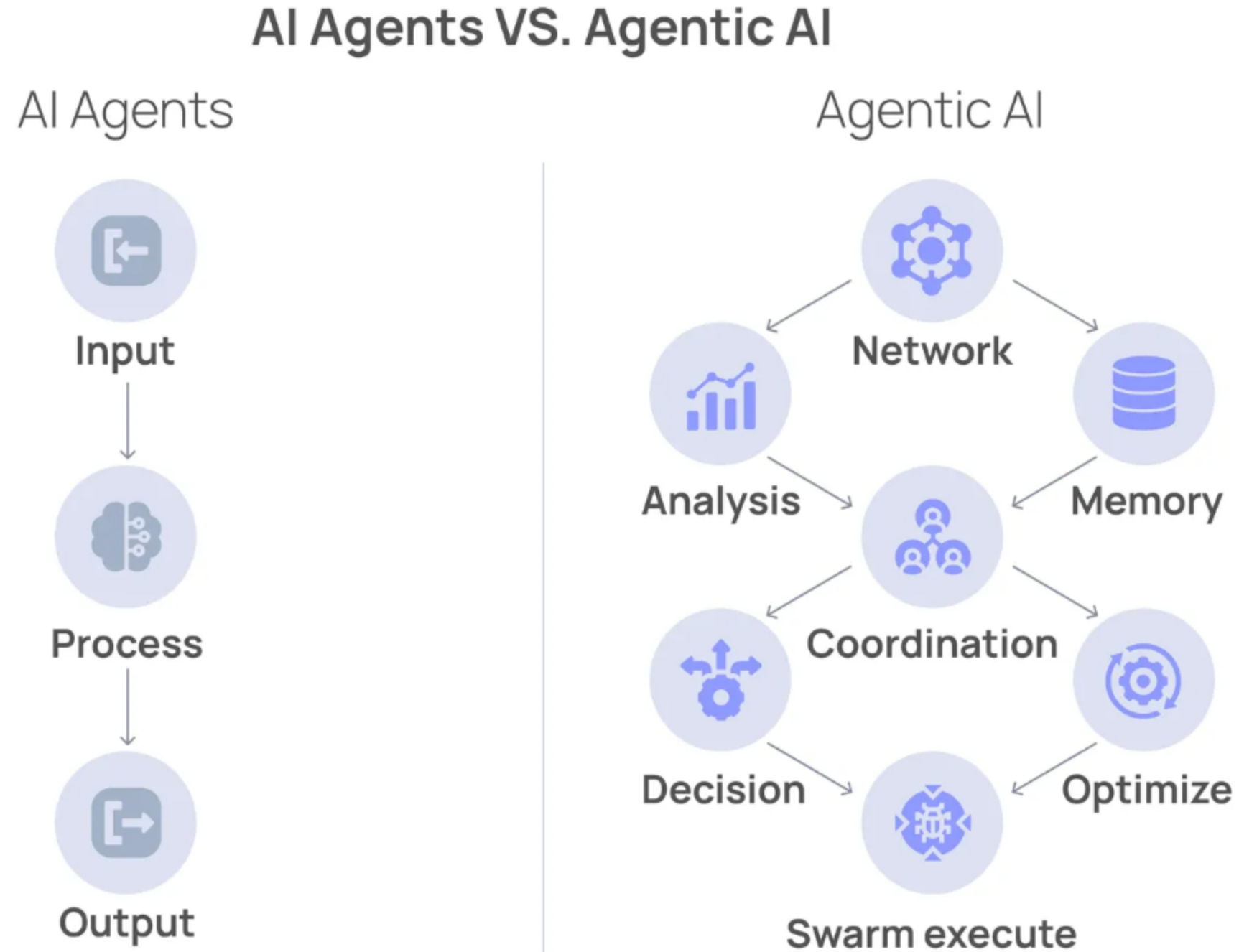
논문 세미나 발표

2025.10.15 강지윤

COPYRIGHT(C) 2025 JIYUN KANG ALL RIGHTS RESERVED.



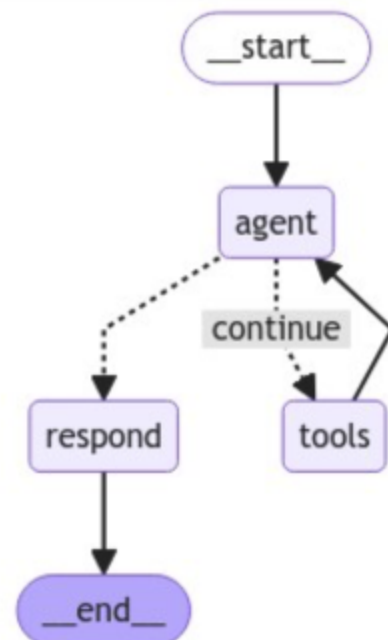
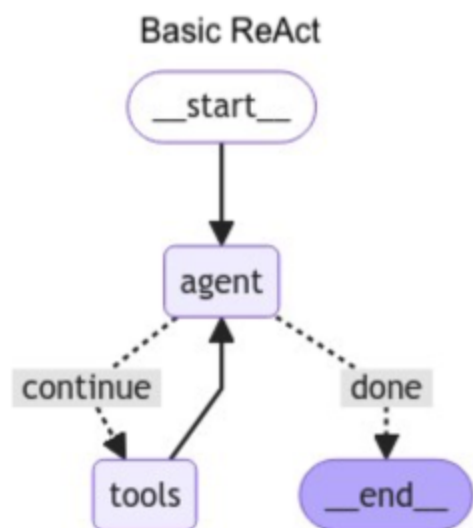
AI AGENT VS AGENTIC AI



HISTORY OF AGENTIC AI

Single Agent

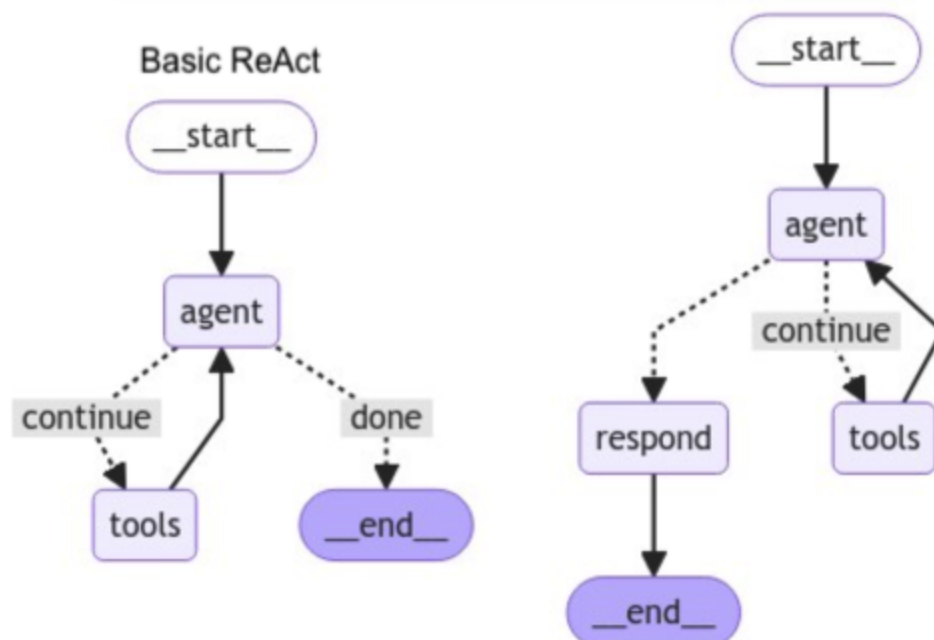
ReAct w/Structured Output



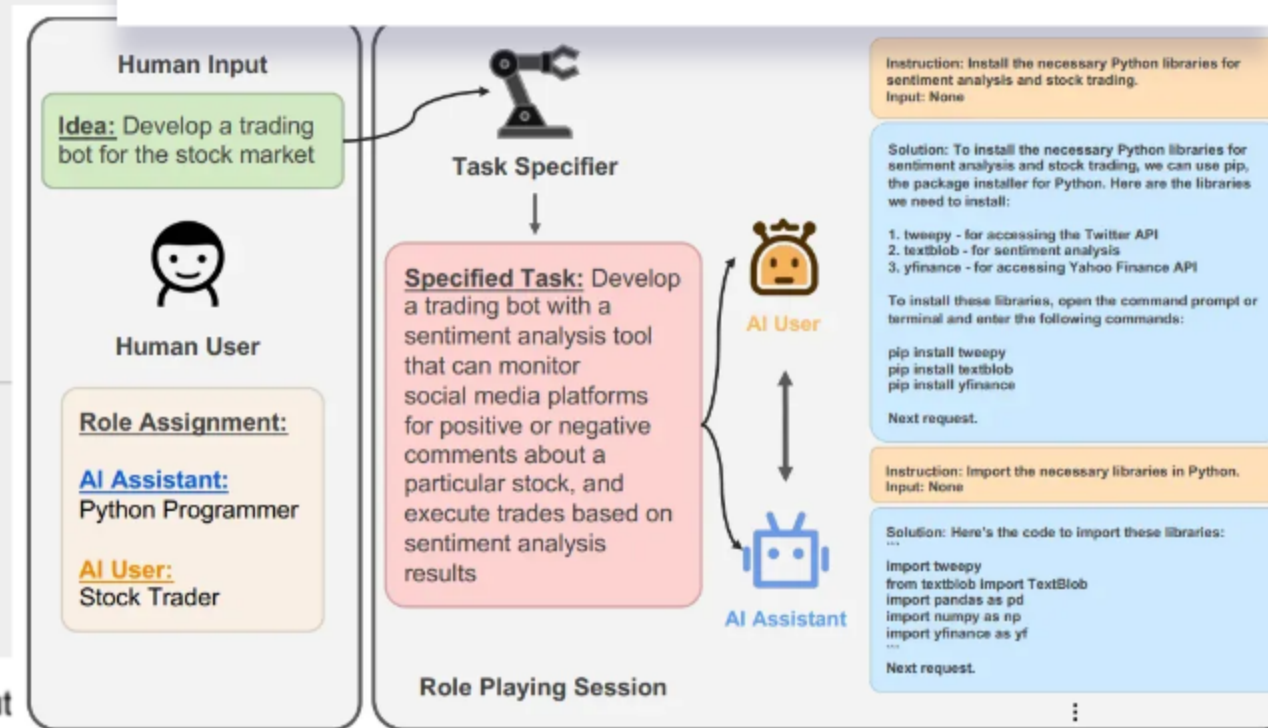
HISTORY OF AGENTIC AI

Single Agent

ReAct w/Structured Output



Static Multi Agents

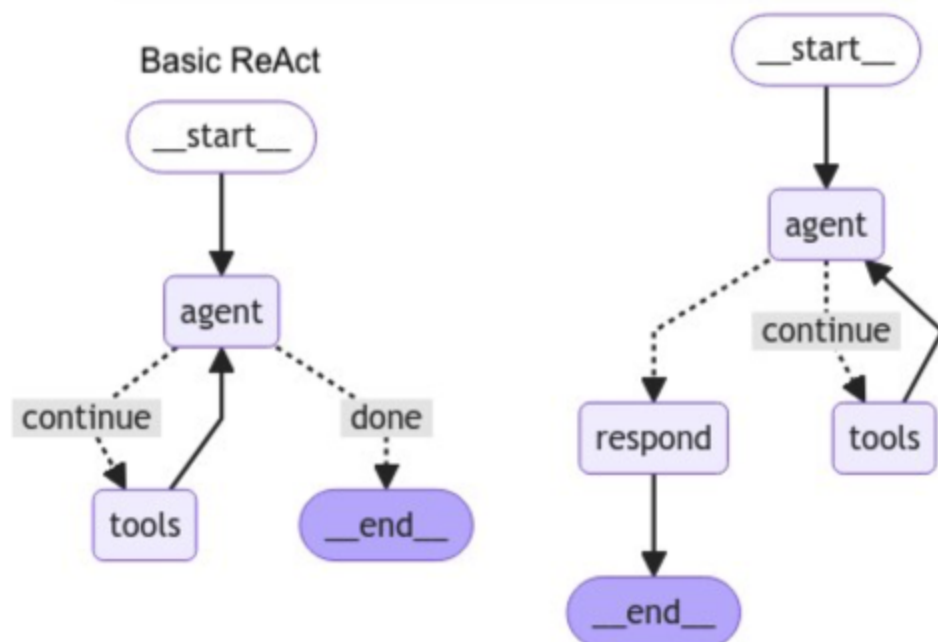


정해진 Role에 따라
Agent가 역할 놀이(Role playing)을 하는 방식

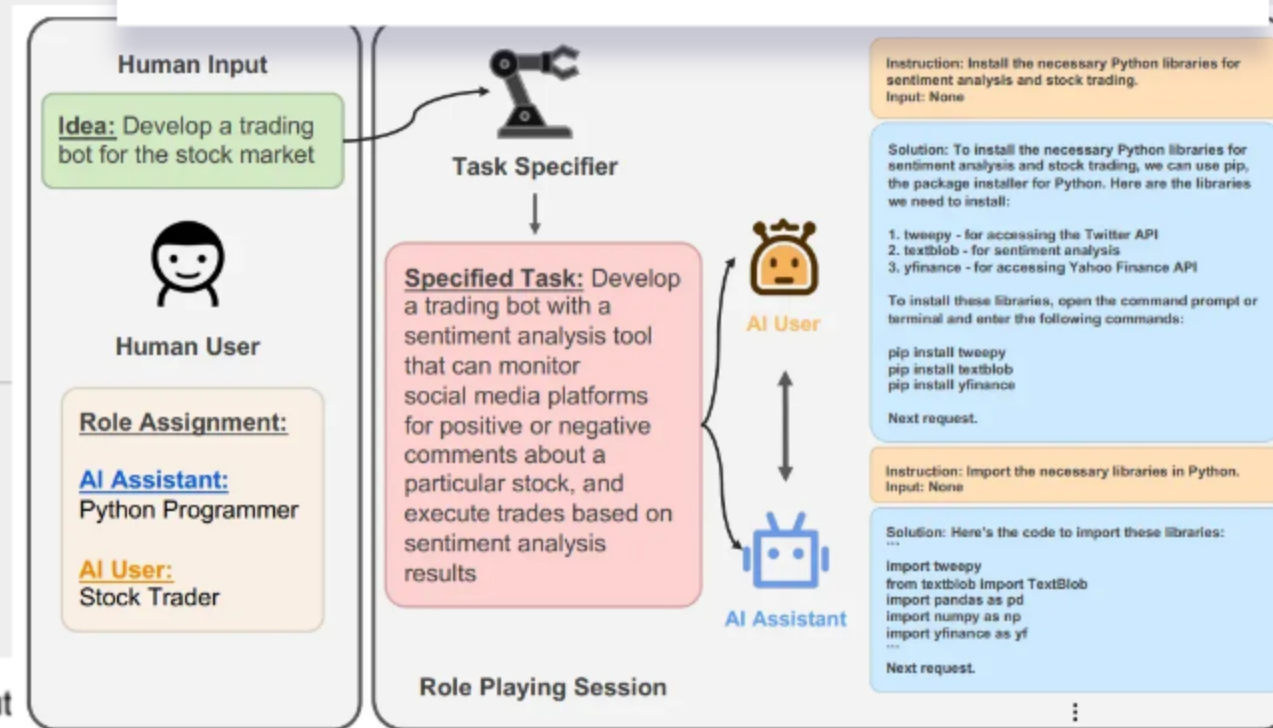
HISTORY OF AGENTIC AI

Single Agent

ReAct w/Structured Output

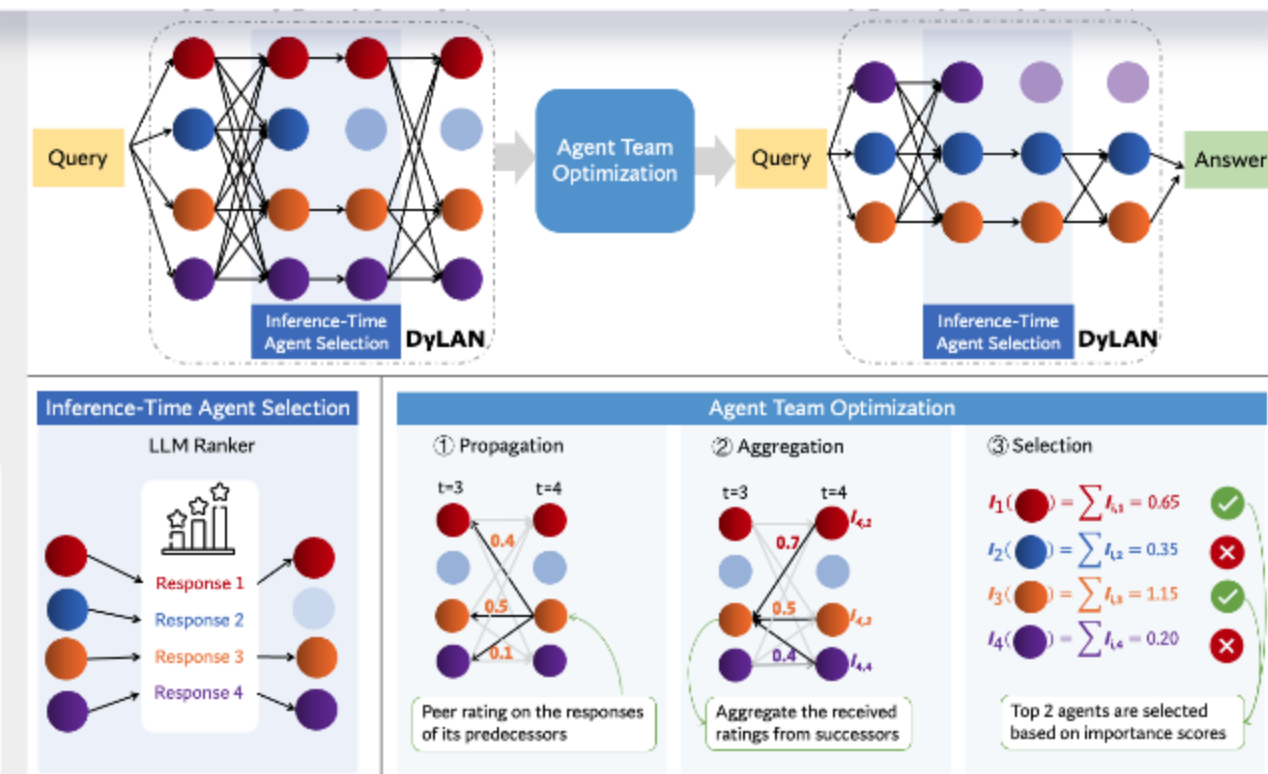


Static Multi Agents



여러 에이전트(agent)가
역할(role)을 배정받아 역할 놀이(role-playing) 방식으로
상호작용하면서 과제를 해결하도록 설계

Dynamic Multi Agents



LLM 기반 에이전트를 동적으로 선택하고,
그들 사이 통신 구조를 동적으로 조정하면서 협업하도록
설계된 프레임워크

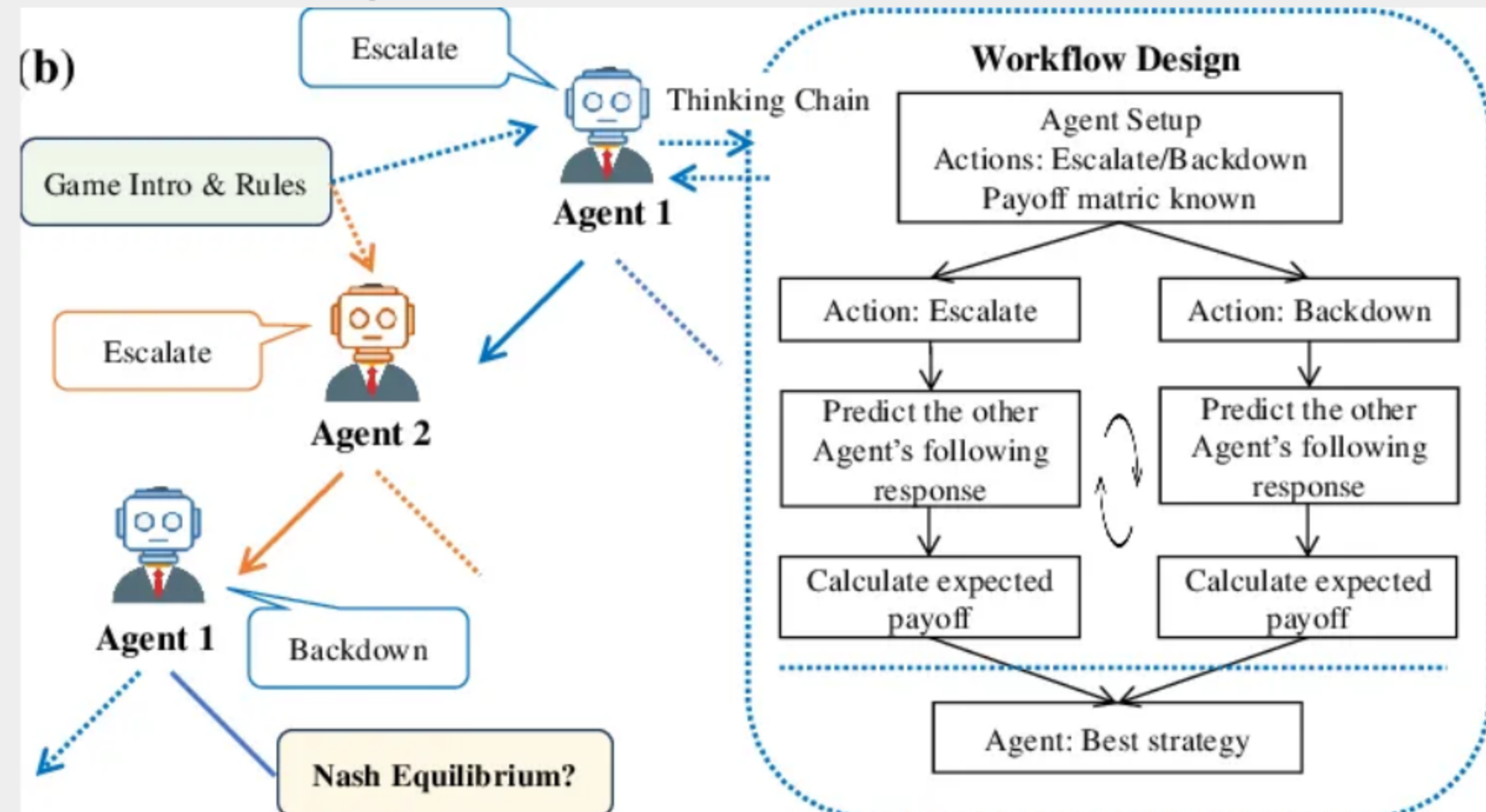
COOPERATION OPTIMIZATION STRATEGIES

중앙집권체제 Centralized framework

ex. AgentLaboratory, ScoreFlow,
WORKFLOW-LLM
Controller agent

지방분권체제 Decentralized approaches

Game-theoretic Workflow : LLM에게 단순한 "무엇을 해라" 지시를 넘어서,
"이런 사고 과정을 거쳐야 한다"라는 구조적 지침을 제공
(규칙 설명 → 전략 탐색 → 균형 분석 → ...)



CHALLENGES

PROBLEM 1

heavily depend on expert insight and
manually-designed policies
(rigid and predefined role space)

SOLUTION 1

Extensible Role Instantiation Mechanism

역할을 유연하게 확장 및 생성할 수 있는 구조
→ 다양한 환경에서도 self-organization & coordination 가능

PROBLEM 2

users' lack expertise in prompt engineering

SOLUTION 2

Prompt Engineering Module

사용자의 질의(query)를 자동으로 refinement하여 multi-agent
collaboration의 효율성과 효과성 향상

HALO'S OBJECTIVE

HALO

워크플로우 탐색(workflow search)을 통해

다중 에이전트 시스템의 최적 reasoning 경로를 찾는 문제

워크플로우는 하위 과업 (sub tasks)로 구성 $\{T_1, T_2, \dots, T_K\}$

각 하위 과업은 역할별 특화 LLM 기반 에이전트 집합 $\mathcal{A}_k = \{a_k^{(1)}, a_k^{(2)}, \dots\}$

에 의해 수행된다.

OBJECTIVE

$\mathcal{S}$ (워크플로우 공간)는 "가능한 모든 과업-에이전트 조합의 집합"이며,

HALO는 주어진 Query를 바탕으로 이 중 가장 효과적인 reasoning workflow를 찾아내는 것을 목표

$$\mathcal{W}^* = \arg \max_{\mathcal{W} \in \mathcal{S}} \text{Value}(\mathcal{Q}, \mathcal{W})$$

HIERARCHICAL REASONING STACK

HIGH

**PLANNING
AGENT**

$$T_k = \mathcal{A}_{\text{plan}}(Q^*, \mathcal{F}, H_{k-1})$$

앞선 refined Query와 전역 구조화된 과업
표현 $\mathcal{F}$ 을 입력 받으면, 과업들을 세분화

이때, 전체 과업을 한번에 세분화하지 않고,
단계적(step-wise) 전략을 사용하여
한번에 하나의 하위 과업만 생성하고, 이전
하위 과업들의 수행 이력 H_{k-1} 을 바탕으로
decomposition policy로 점진적 갱신

MID

**ROLE-
DESIGN
AGENTS**

하위 에이전트 생성

$$\mathcal{A}_k = \{a_k^{(1)}, a_k^{(2)}, \dots\}$$

역할별 시스템 프롬프트 생성

$$a_k^{(i)} = \mathcal{A}_{\text{role}}(T_k, Q^*, \mathcal{F}) \Rightarrow \rho_k^{(i)}$$

LOW

**INFERENCE
AGENT**

협력적 추론 수행

$$\mathcal{Y}_k = \{y_k^{(1)}, y_k^{(2)}, \dots\} = \mathcal{A}_k(T_k, Q^*, \mathcal{F})$$

$T_k, Q^*, \mathcal{F}$ 를 받은 Inference agent는 협력적 추론을
수행하고 이에 맞는 집합 $\mathcal{Y}_k$ 라는 중간 산출물을 생성

조기 종료(early-stopping)

HALO는 수행 이력 H_k 에 포함된 하위 과업 중 66% 이
상이 동일한 일관된 답변 $\hat{Y}$ 에 도달했을 경우, 추론 과
정을 조기에 종료

