



This is a paper from the 1990s. Would a modern AI be using the methods and conclusions of this paper?

Yes—in spirit and in selected technical settings, but not as the core method used to train or operate a frontier general-purpose AI such as an LLM. Natalie Glance’s 1993 dissertation, *Dynamics with Expectations*, is strikingly prescient about multi-agent systems: it studies agents with delayed and imperfect information, limited resources, forward-looking expectations, switching costs, heterogeneous strategies, and collective-action problems. Those remain central problems in modern AI.^[1]

The important distinction is that modern AI would usually treat its models as **conceptual, analytical, or simulation frameworks**, then replace the paper’s hand-specified decision rules with learned policies, neural models, empirical data, and large-scale experimentation.

What the paper actually does

The dissertation studies two related systems:

1. **Resource contention in computational ecosystems.** Agents choose among scarce resources based on perceived payoff. Their information may be noisy and delayed; they may forecast future payoffs either by assuming the current situation persists (“flat expectations”) or extrapolating trends. The work uses nonlinear dynamics, delay differential equations, mean-field approximations, linear stability analysis, and numerical simulation to identify stable, oscillatory, chaotic, and phase-transition-like regimes.^[1]
2. **Collective action and cooperation.** Agents face a tension between individual incentives—such as free riding—and group benefits. The dissertation examines how expectations, uncertainty, group size, diversity, hierarchical/ultrametric organization, and changing organizational topology affect whether cooperation emerges or collapses. It finds metastability, abrupt transitions, and cases where diverse or fluid structures create alternative pathways toward cooperation.^[1]

That is not “AI” in the contemporary deep-learning sense. It is closer to what we now call **agent-based modeling, multi-agent systems, complex adaptive systems, evolutionary game theory**, and parts of **multi-agent reinforcement learning (MARL)**.

Where modern AI still uses it

Paper concept	Modern counterpart	Status today
Agents choose actions based on expected future payoff	Reinforcement learning, planning, model-based RL, expected-return optimization	Fundamental
Other agents make one's environment change	Multi-agent RL non-stationarity	Fundamental open problem
Delayed, partial, noisy observations	Partially observable Markov decision processes, distributed control, decentralized execution	Fundamental
Competition for limited resources	Scheduling, load balancing, bandwidth/compute allocation, cloud systems, robot fleets	Actively used
Cooperation vs. free riding	Mechanism design, AI alignment, multi-agent coordination, distributed governance	Actively researched
Diverse agent beliefs and strategies	Population-based training, heterogeneous-agent simulations, mixed-agent evaluation	Actively used
Feedback, oscillation, instability, bifurcation	Control theory, learning dynamics, safety analysis, system identification	Still valuable
Organizational topology and fluidity	Communication graphs, dynamic networks, swarm systems, federated/distributed agents	Actively used

The direct intellectual bridge is particularly clear in MARL. A central difficulty is **non-stationarity**: when one agent learns or changes its policy, every other agent's effective environment changes as well. That is essentially the modern formulation of the feedback loop Glance analyzed—agents react to aggregate behavior, but their reactions themselves reshape it. Contemporary MARL surveys identify non-stationarity, partial observability, scalability, and credit assignment as defining challenges.^{[2][3][4]}

What has changed since 1993

The dissertation assumes deliberately simple expectation rules: agents either treat the recent state as persistent or extend a perceived trend. It also models utility functions, uncertainty, and switching behavior explicitly. Glance says this simplification is purposeful: she wanted to study the systemic consequences of expectations, rather than solve the harder problem of how agents form accurate predictions.^[1]

A modern implementation usually changes several things:

- **Learned rather than fixed expectations.** A neural policy, world model, Bayesian belief model, or predictor may infer future rewards, opponents' behavior, and state transitions from data.
- **High-dimensional observations.** Current agents may process language, images, code, logs, sensor streams, or tool outputs—not just a scalar resource-load variable.
- **Richer action spaces.** Rather than switching between two resources, an agent might allocate work among hundreds of tools, compute nodes, subagents, markets, or tasks.
- **Centralized training.** During training, systems often use global state or aggregate data to mitigate the very instability caused by independently adapting agents; then they execute with only local information.
- **Empirical validation.** Rather than relying principally on solvable stylized models and phase diagrams, modern systems are benchmarked in simulations, games, robotics, operational traces, and deployment environments.
- **Safety and mechanism layers.** Current deployed multi-agent workflows often enforce permissions, budgets, routing rules, contracts, reputation, audits, or centralized orchestration—because spontaneous cooperation is not reliable enough in arbitrary systems.

So a modern AI would generally not literally use the thesis's specific error-function switching probabilities, two-resource payoff forms, or ultrametric hierarchy as its universal architecture. But a researcher might absolutely use those as a minimal model to isolate a failure mode before building a more complex learned system.

Which conclusions hold up best

Several conclusions look highly durable.

- **Forecasting can destabilize systems.** The thesis finds that longer predictive horizons and trend extrapolation can induce oscillation, instability, or crashes when many agents act on similar forecasts. This is still a live concern: strategic learners, trading agents, recommender systems, and tool-using agents can create feedback loops where predictions change the world being predicted.^[1]
- **Delayed information matters.** The paper shows that delay and reevaluation rates shape phase lag, tracking ability, and instability. Modern distributed systems and agentic AI encounter this constantly: stale state, slow API responses, asynchronous workers, delayed feedback, and out-of-date models all change collective behavior.^[1]

- **The best static equilibrium may not be most adaptive.** One of the dissertation’s more interesting results is that a system’s performance in a changing environment may differ from its performance in a static one; transiently chaotic systems can sometimes respond better to novel changes than a more stable system, though they may be worse over longer horizons. Modern control, exploration, robust RL, and adaptive systems retain this tradeoff: responsiveness can require flexibility, but flexibility can create volatility.^[11]
- **Collective outcomes can change abruptly.** The work finds metastable states, threshold effects, dependence on initial conditions, and sharp transitions between cooperation and defection. Those ideas remain useful for thinking about online communities, decentralized networks, AI-agent economies, protocol incentives, and the institutional conditions for cooperation.^[11]
- **Topology is part of the mechanism.** The claim that group structure and its capacity to change affect cooperative dynamics is very compatible with modern graph-based MARL, dynamic communication networks, and LLM-agent workflows. Contemporary LLM multi-agent research explicitly frames systems in terms of players, strategies, payoffs, and information—the same core analytical primitives used here.^[15]

What should be treated cautiously

The paper’s conclusions are **conditional results of small, stylized models**, not empirical laws about people, markets, or AI agents.

In particular:

- “Chaos improves adaptability” should not be read as a design prescription to make AI systems chaotic. In the thesis itself, the advantage is conditional, short-horizon, and tied to a narrow performance definition and specific driven dynamical systems.^[11]
- Trend-following produces a crash in the model because its assumptions synchronize agents’ reactions in a particular way. Modern agents may have learned, heterogeneous, communicative, constrained, or strategically adversarial policies, so the outcome need not generalize quantitatively.
- The paper treats agents as comparatively simple payoff optimizers. LLM agents are not clean utility maximizers; they are language-model policies shaped by pretraining, prompting, tool scaffolding, memory, and external controls. Recent work even reports that LLM agents in repeated games may show cooperative behavior rather than converging to standard Nash-equilibrium predictions, illustrating why one should test rather than assume classical-agent behavior.^[16]

- Its main formalism is best suited to aggregate behavior. It does not capture the semantic reasoning, persuasion, deception, norm formation, open-ended invention, or institutional rule-making that become important when agents communicate in natural language.

Bottom line

This paper is **more relevant to modern AI than its 1993 date might suggest**, especially if your interest is in decentralized coordination, emergent order, and the design of agent ecologies rather than in the internals of a single language model.

A modern AI researcher would likely use it in one of three ways:

- As a **theoretical baseline** for analyzing delayed feedback, expectation-driven instability, and regime shifts.
- As an **agent-based simulation template** for studying tool-using AI agents competing or cooperating over compute, tasks, data, liquidity, attention, or shared infrastructure.
- As a **design warning**: decentralized systems do not automatically self-organize into beneficial orders; the timing of information, incentive structure, switching friction, diversity of strategies, and network topology can determine whether they coordinate, cycle, crash, or free ride.

For the kind of “freeorder” question you are interested in, its deepest continuing value is probably this: spontaneous order is not merely the absence of central control. It depends on the evolving relation among local knowledge, expectations, feedback delays, incentives, and the structures through which agents encounter one another. That is a useful framework for evaluating emerging AI-agent economies and institutions—even though the mathematical machinery would now be augmented substantially with MARL, learned world models, network analysis, and empirical simulation.^{[7][2][1]}

*
**

1. thesis-Unknown.pdf
2. <https://www.sciopen.com/article/10.1016/j.jai.2024.02.003>
3. <https://arxiv.org/html/2312.10256v2>
4. <https://www.alphaxiv.org/abs/2312.01058>
5. <https://arxiv.org/abs/2601.15047>

6. <https://arxiv.org/abs/2604.00487>
7. <https://openreview.net/forum?id=0ZnEzvSLNR>
8. <https://dl.acm.org/doi/10.1145/3817113>
9. <https://arxiv.org/html/2503.13415v1>
10. <https://www.sciencedirect.com/science/article/pii/S0376042126000096>
11. <https://www.sciencedirect.com/science/article/pii/S2949855424000042>
12. <https://openreview.net/forum?id=esabw9linR>
13. <https://www.sciencedirect.com/science/article/pii/S0893608026005708>
14. <https://apxml.com/courses/multi-agent-llm-systems-design-implementation/chapter-5-agent-reasoning-decision-making/game-theory-agent-interactions>
15. <https://link.springer.com/article/10.1007/s10458-024-09666-5>
16. <https://dl.acm.org/doi/full/10.1145/3806262.3806263>