

AI Trading in Financial Markets

By Itay Goldstein¹

Abstract

An important aspect of Artificial Intelligence (AI) which has not been present in previous technological advancements is that it has the ability to act autonomously. Hence, we need to have a better understanding of how deploying autonomous AI agents into financial markets might affect market equilibrium. Research we are conducting in this area based on experimental methods – simulating the behaviour of AI Reinforcement-Learning algorithms in a financial market environment – sheds new light on this topic. Results from such analysis show that autonomous algorithms find ways to collude with each other and even manipulate prices for their own profit motives. This can reduce liquidity and price informativeness and create excess volatility and fragility in prices. I discuss what may be different about AI agents relative to humans and elaborate on what changes might be needed from policy makers and regulators when trying to promote the objectives of market competition and stability.

1 AI-Powered Trading

Artificial intelligence (AI) is introduced into finance across different functions and applications. Like the introduction of many technologies over the years, it has the potential to improve processes in the financial system, making them faster and more effective. What I would like to focus on in my remarks is one feature that makes AI quite distinct from previous technologies. This is that AI agents can act autonomously, making their own decisions which ultimately affect outcomes in the financial system. I will focus on trading in financial markets, where the application of such AI algorithmic agents is particularly pertinent, and report findings from ongoing research we have in this area and the implications for market quality and stability.

Algorithmic trading has been a common feature in financial markets for a while. For example, in the context of high frequency trading, a lot of the actions in the market are done by algorithms, which is necessary for competing in such environment. What is new about AI-powered algorithms is that they are not following instructions of the humans who deploy them, but rather find their own strategies following a process of reinforcement learning. Reinforcement learning is a key element in advanced AI systems. Algorithms are deployed in the target environment, trying different strategies in different states of the world, and updating over time what works best.

¹ Wharton School, University of Pennsylvania. Email: itayg@wharton.upenn.edu.

Through that, they develop their own strategies, which would not have been envisioned by the human players they represent.

A common implementation of reinforcement learning is Q-learning, which we have been using in our research. In this implementation, an AI agent keeps and updates a Q-value table, summarizing expected payoff for every pair of state and action. The agents are repeatedly confronted with different states, for which they have to choose actions. Following their chosen actions, they experience an outcome. After each one of these steps, the Q-value of a particular pair of state and action is updated to reflect the outcome from the new experience. The algorithm is governed by hyper-parameters that control the process of updating and convergence.

One hyper-parameter determines the weight given to the new experience relative to previous knowledge in the updating process. It is of course essential to have some weight on both new experiences and previous knowledge for the reinforcement learning to work, and the relative weights will affect this learning process. Another hyper-parameter determines the likelihood that the algorithm will exploit previous knowledge, choosing the action that is currently predicted to generate the best outcome, vs. that it will explore new actions. Again, reinforcement learning relies on having some combination of exploration and exploitation, and the rate of exploration is generally decaying over time, as the quality of the predictions is becoming more and more reliable.

With this process, AI-powered algorithms can evolve over time in a financial market environment. They try different actions in different states, update their projected values from different strategies, and improve their effectiveness in generating profits. The strategies they come up with are not easy to predict as they respond to evolving market conditions and incorporate the knowledge gained from attempting different actions in different states. With the evolution of AI techniques, it is only natural to take the existing algorithmic trading strategies and amend them with these reinforcement-learning capabilities, and this is why trading in financial markets is one of the most likely applications of autonomous AI agents in the financial system.

2 Risks to Financial Markets

When we evaluate risks for financial markets from growing use of autonomous AI systems in the trading process, the key challenge is to understand the interaction among AI trading algorithms. Analysing the behaviour of each algorithm in isolation is not sufficient because the reinforcement-learning algorithms evolve with the feedback they get from the environment, and this environment, in turn, is composed of other algorithms. Hence, understanding the interaction among such algorithms, with all the complexities involved, is essential. From the point of view of regulators and policymakers, what is important to know is what implications such interactions will have for market outcomes. In our research, which I will describe below, we are looking at the potential for collusion and manipulation. These behaviours have first-order implications for market quality and for market stability.

When thinking through the potential interactions, it is important to recognize that the insights obtained over the years from studying human behaviour in financial markets may not extend to a world where markets are populated by autonomous AI algorithms. Growing consensus in the field of AI indicates that the kind of intelligence that such agents exhibit is different from human intelligence. While reinforcement learning, as described above, certainly has features that resemble the evolution of human knowledge and experience, there is little evidence that it captures human cognition at high level of accuracy. The storage of information from past experiences in the reinforcement learning process may lend itself more naturally to behaviours like collusion than the way that humans update and form expectations.

For these reasons, market regulators have always been guided by the principle that in order to detect collusion and enforce competitive market behaviour, they must rely on evidence of explicit communication or even intent when dealing with human players. The dominant thinking is that humans would have a hard time sustaining a system of collusion without such explicit coordination. This, however, might be very different when dealing with autonomous AI trading algorithms. Even when not having any explicit intent to collude, they may converge to this outcome through the way they learn and update based on their past experiences. Hence, the regulatory approach to market competition may need to be amended when dealing with non-human players.

3 Experimental Research Approach

There are challenges in trying to shed new light on these questions through traditional research methods. Theoretical analysis is not tractable for characterizing the interaction among large reinforcement-learning algorithms, as traditional notions of equilibrium that are used to study rational human agents have not been established for this case. Hence, obtaining analytical predictions is not a viable path. Empirical analysis has its own obstacles since we do not have available data on which market players are relying on which algorithms. Hence, the ability to obtain insights on the effect of reinforcement-learning algorithms from market data is also limited.

Our research program instead can be described as experimental. Simulating financial market environments and deploying reinforcement-learning algorithms to trade in such environments can generate insights on how they interact and what market outcomes are generated through their interactions. As described below, we will use this approach to investigate the extent to which reinforcement-learning algorithms are able to achieve outcomes such as collusion or manipulation. This is an experiment in AI algorithms as the economic agents, as opposed to traditional experimental analysis which is all about experimenting on humans to try and understand their behaviour in different circumstances.

It should be noted that the experiments with AI agents are not hampered by some of the difficulties that usually come with experiments with humans. First, human experiments can be very costly and difficult to scale up. The reason is that they

require recruiting humans and getting their time and attention. Hence, human experiments are typically limited in scale, whereas with AI agents, researchers are not subject to these limitations, and experiments can be run over and over again with very minimal constraints. Second, the concern about human experiments is usually about external validity. Are humans going to exhibit the same behaviour in the field as they exhibit in the lab? With AI agents, on the other hand, such concerns are unlikely to arise, as they do not know the difference between the two. Hence, the results are quite credible when reaching conclusions on the effect of AI on real-world financial markets.

4 Results from Our Research

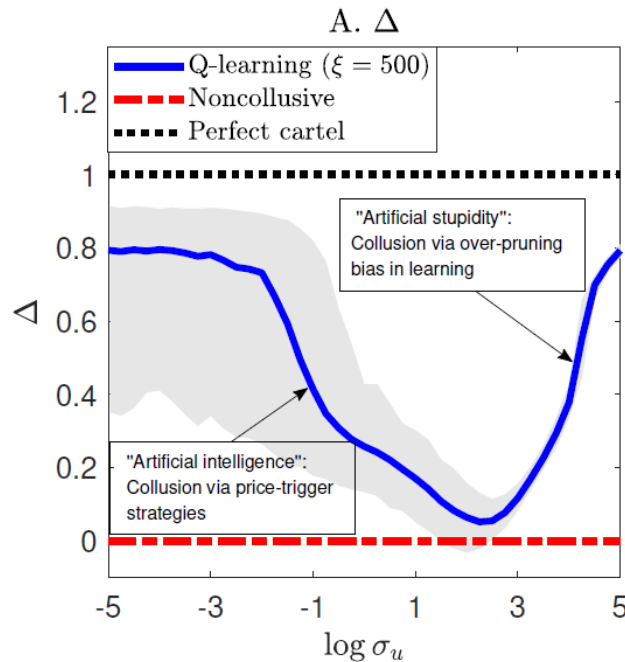
I will now describe briefly the key findings from two working papers that I have with my colleague Winston Dou at Wharton and with Yan Ji at the Hong Kong University of Science and Technology. The first paper – Dou, Goldstein, and Ji (2025a) – investigates the possibility of market collusion among AI-powered algorithms. The second one – Dou, Goldstein, and Ji (2025b) – explores the possibility of coordinated market manipulation.

4.1 Collusion

The possibility of collusion among reinforcement-learning algorithms has been explored in various contexts. An important question is whether something of that kind might be happening in financial market trading, what shape it might take, and what the underlying mechanisms are. It is important to remember that financial markets are different from typical settings of collusion in product markets, and so these questions are generally open. We explore the possibility of collusion in Dou, Goldstein, and Ji (2025a). We simulate a financial market with noise traders, a market maker, and long-term information insensitive traders. Our focus is on speculators who receive information about developments in the value of an asset over time and have to decide how to trade on this information. In our experiments, these speculators are replaced by reinforcement-learning algorithms.

In this setup, a non-collusive outcome represents a situation where speculators trade on their information to maximize their immediate profit from it. Their profits in this case are reduced by the fact that their aggressive trading behaviours are hurting each other. If they joined forces and made their decisions together as a cartel, they would have traded less aggressively on their information and achieved higher long-term profits. A collusive outcome generally describes a case where they end up trading less aggressively, and this may be supported by different mechanisms. Chart 1, copied from the paper, shows the profits of algorithms in our experiments along the blue line. This is compared to the non-collusive low-profit benchmark in red and the perfect-cartel high-profit benchmark in black. The horizontal axis captures the level of noise in the market, and so the chart shows us realized profits as a function of market noise. Higher profits generally indicate a collusive outcome.

Chart 1
Algorithmic Collusion



Sources: Dou, Goldstein, and Ji (2025a)

Inspecting the chart, we can see that algorithms end up in a collusive outcome in a wide range of the noise parameter, specifically when the noise is low and when it is high, but not in the middle. When we look under the hood to understand the mechanism that supports collusion, we find that for low levels of noise, the mechanism is based on a rather sophisticated price-trigger strategy. Here, the algorithms end up trading less aggressively, sacrificing immediate profits, because they learn that deviations will trigger prices to go beyond a threshold and that this will lead to punishment in the form of aggressive behaviour by other parties. On the other hand, for high levels of noise, collusive outcomes are generated by learning biases, leading the algorithms to over-prune aggressive trading strategies because of some bad experiences that they had in the past. The difference between these mechanisms leads us to refer to the behaviour on the left side of the chart as “artificial intelligence” and to that on the right side as “artificial stupidity”.

Whatever the mechanism supporting collusive outcomes is, the implications for market quality are similar. Less competitive trading leading to collusive outcomes is also associated with lower market liquidity and lower price informativeness. Hence, regulators charged with the task of ensuring market quality, liquidity, and informativeness should pay attention to these behaviours and the mechanisms that support them. It is important to note here that traditionally the regulatory approach with collusion has been that no action should be taken without evidence of explicit

communication or intention. However, the patterns that emerge from looking at algorithmic behaviour is that collusion is likely to be achieved without such explicit communication or intention when the players are algorithms. Hence, this may call for a change in the regulatory paradigm when markets are populated by algorithms and not just by human players.

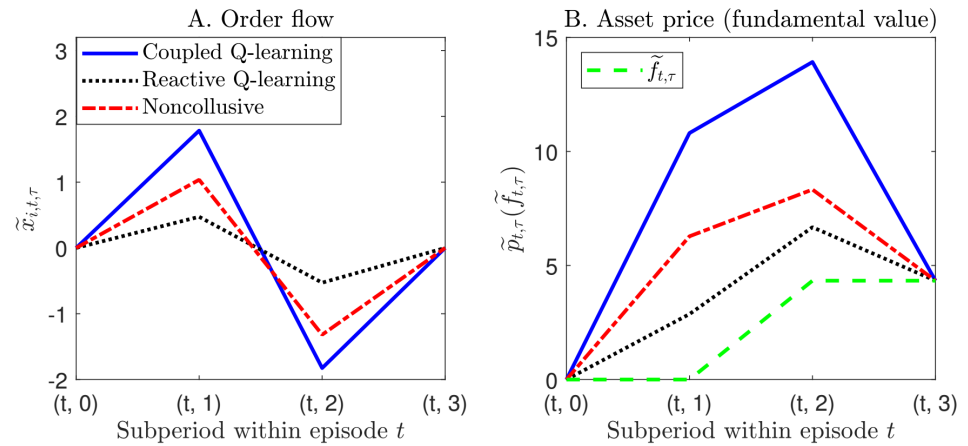
4.2 Manipulation

The analysis described in the previous section is based on a relatively basic version of Q-learning. It can be thought of as “reactive Q-learning”, where the algorithm just reacts to previous outcomes. More sophisticated versions can provide some structure to the algorithms, allowing it to better anticipate future evolution of events and plan accordingly. This is known as “AI planning” or “coupled Q-learning”. It is with such algorithms that one can inspect whether AI agents are capable of coordinating on more sophisticated market behaviours that are also potentially more damaging to market stability. This is what we explore in the follow-up paper – Dou, Goldstein and Ji (2025b).

The analysis here focuses on market manipulation. The kind of manipulation that has long been a concern in financial markets is where speculators trade to pump up (or down) the price of the asset and then lead to its decrease (or increase). Coordinated trade of multiple speculators along these lines can lead to a pattern of volatility and generate profits to the speculators who are part of the scheme and take advantage of market followers. Like in the previous collusion setting, coordinating on such behaviour has been thought to be difficult to achieve for human traders unless there is explicit communication or intention. The question is whether reinforcement-learning algorithms can achieve that without such communication just by learning patterns and outcomes over time and with the additional ability of planning.

The answer that is coming out of our analysis in this paper is generally positive about the ability of AI algorithms to coordinate on a path of price manipulation. Chart 2, copied from the paper, provides a brief illustration of this path. On the right panel, a path of fundamentals along the green line is shown to be amplified along the blue line, which incorporates the trading of the coupled Q-learning algorithms. Their trading behaviour is depicted in the left panel. We can see how they take a fundamental increase in value, pump up the price and then wind down to converge to the true fundamental. On the way, they profit from the amplified increase and decrease in asset price through their coordinated trading that is taking advantage of price followers (or momentum traders). This is different from either the non-collusive behaviour or the reactive Q-learning behaviour. The mechanism that supports this amplified behaviour has some resemblance to the price trigger strategies described before, yet because of the higher level of complexity there is more to it that is required to sustain such behaviour. Again, the algorithms are able to achieve that without communication or explicit intention, just through their reinforcement learning behaviour.

Chart 2
Algorithmic Manipulation



Sources: Dou, Goldstein, and Ji (2025b)

Overall, we can see here that algorithms can coordinate on more than just collusion. They effectively manipulate the market and set up a whole price path that generates profits for them. While the previous paper concluded with concerns for market liquidity and informativeness. Here, we can see concerns for the stability of the market, as the price path that algorithms have converged on exhibits greater fragility and volatility. Market regulators, charged with the goal of maintaining stability, may thus be even more worried about this kind of behaviour and market outcomes. Hence, this may call even more strongly for a change in the regulatory paradigm.

5 Conclusion

With the rise of reinforcement-learning algorithms and their likely deployment in financial markets, we need to understand what the interactions among them might look like and what the implications are for market outcomes. While traditional research methods – either theoretical or empirical – are limited at this point in their ability to shed light on this topic, experimental research can be readily available and provide new insights on such algorithmic interactions and behaviours. This experimental research is different from traditional human experiments. We are conducting experiments with AI agents to see the patterns they might exhibit. Concerns about scalability and external validity are much less pertinent.

I described two works in progress that I have with Winston Dou and Yan Ji. In these papers, we use the experimental approach and show how AI algorithms achieve patterns of collusion or even manipulation in different settings. Such patterns have first-order implications for market outcomes, either leading to lower liquidity and price informativeness, or even causing excess volatility and fragility of the market. Regulators may want to further look into these possibilities, as the traditional paradigm for dealing with collusion and manipulation with human players may not apply for algorithmic players, who seem able to achieve coordinated outcomes without any intent or communication.

Future work should also explore in more detail the differences between AI agents and human traders. In another work in progress, we try to shed light on the interaction between the two types under the premise that human traders may have other sources of information that are not available to algorithms but that AI agents are better at processing market data and learning from prices. The fact that humans may have other sources of information means that they will not disappear from the market, and so future markets will exhibit interactions among the different types of traders. Overall, these are certainly times of change in financial markets and there is great need for concerted effort from academics, regulators, and practitioners to ensure that we understand these changes and avoid some undesirable outcomes.

References

Dou, W., Goldstein, I. and Ji, Y. (2025a), “AI-powered trading, algorithmic collusion, and price efficiency”, *Wharton Working Paper*.

Dou, W., Goldstein, I. and Ji, Y. (2025b), “Financial market fragility in the era of AI planning”, *Wharton Working Paper*.