



ΕΘΝΙΚΟ ΜΕΤΣΟΒΙΟ ΠΟΛΥΤΕΧΝΕΙΟ
ΣΧΟΛΗ ΗΛΕΚΤΡΟΛΟΓΩΝ ΜΗΧΑΝΙΚΩΝ ΚΑΙ ΜΗΧΑΝΙΚΩΝ ΥΠΟΛΟΓΙΣΤΩΝ
ΤΟΜΕΑΣ ΤΕΧΝΟΛΟΓΙΑΣ ΠΛΗΡΟΦΟΡΙΚΗΣ ΚΑΙ ΥΠΟΛΟΓΙΣΤΩΝ
ΕΡΓΑΣΤΗΡΙΟ ΣΥΣΤΗΜΑΤΩΝ ΤΕΧΝΗΤΗΣ ΝΟΗΜΟΣΥΝΗΣ ΚΑΙ ΜΑΘΗΣΗΣ

Adaptive Multi-Agent LLM Systems for Financial Trading: A Framework for Realistic Simulation and Dynamic Prompt Optimization

DIPLOMA THESIS

by

Charidimos Papadakis

Επιβλέπων: Γεώργιος Στάμου
Καθηγητής Ε.Μ.Π.

Αθήνα, Οκτώβριος 2025



Εθνικό Μετσόβιο Πολυτεχνείο
Σχολή Ηλεκτρολόγων Μηχανικών και Μηχανικών Υπολογιστών
Τομέας Τεχνολογίας Πληροφορικής και Υπολογιστών
Εργαστήριο Συστημάτων Τεχνητής Νοημοσύνης και Μάθησης

Adaptive Multi-Agent LLM Systems for Financial Trading: A Framework for Realistic Simulation and Dynamic Prompt Optimization

DIPLOMA THESIS

by

Charidimos Papadakis

Επιβλέπων: Γεώργιος Στάμου
Καθηγητής Ε.Μ.Π.

Εγκρίθηκε από την τριμελή εξεταστική επιτροπή την 30^η Οκτωβρίου, 2025.

.....
Γεώργιος Στάμου
Καθηγητής Ε.Μ.Π.

.....
Αθανάσιος Βουλόδημος
Επίκουρος Καθηγητής Ε.Μ.Π.

.....
Ανδρέας-Γεώργιος Σταφυλοπάτης
Ομότιμος Καθηγητής Ε.Μ.Π.

Αθήνα, Οκτώβριος 2025

.....
ΧΑΡΙΔΗΜΟΣ ΠΑΠΑΔΑΚΗΣ
Διπλωματούχος Ηλεκτρολόγος Μηχανικός
και Μηχανικός Υπολογιστών Ε.Μ.Π.

Copyright © – All rights reserved Charidimos Papadakis, 2025.

Με επιφύλαξη παντός δικαιώματος.

Απαγορεύεται η αντιγραφή, αποθήκευση και διανομή της παρούσας εργασίας, εξ ολοκλήρου ή τμήματος αυτής, για εμπορικό σκοπό. Επιτρέπεται η ανατύπωση, αποθήκευση και διανομή για σκοπό μη κερδοσκοπικό, εκπαιδευτικής ή ερευνητικής φύσης, υπό την προϋπόθεση να αναφέρεται η πηγή προέλευσης και να διατηρείται το παρόν μήνυμα. Ερωτήματα που αφορούν τη χρήση της εργασίας για κερδοσκοπικό σκοπό πρέπει να απευθύνονται προς τον συγγραφέα.

Οι απόψεις και τα συμπεράσματα που περιέχονται σε αυτό το έγγραφο εκφράζουν τον συγγραφέα και δεν πρέπει να ερμηνευθεί ότι αντιπροσωπεύουν τις επίσημες θέσεις του Εθνικού Μετσόβιου Πολυτεχνείου.

AWS resources were provided by the National Infrastructures for Research and Technology GRNET and funded by the EU Recovery and Resiliency Facility.

Περίληψη

Τα Μεγάλα Γλωσσικά Μοντέλα (ΜΓΜ) έχουν επιδείξει αξιοσημείωτες δυνατότητες σε σύνθετες διαδικασίες λήψης αποφάσεων, προβάλλοντας ως πολλά υποσχόμενα εργαλεία για εφαρμογές χρηματοοικονομικών συναλλαγών. Τέτοιες εφαρμογές απαιτούν ενοποίηση ετερογενών πηγών πληροφόρησης, ικανότητα συλλογισμού υπό αβεβαιότητα και προσαρμογή στις ραγδαία μεταβαλλόμενες συνθήκες της αγοράς. Ωστόσο, η πρόοδος στον τομέα παραμένει περιορισμένη, λόγω της έλλειψης ρεαλιστικών περιβαλλόντων προσομοίωσης ειδικά σχεδιασμένων για την αξιολόγηση αυτόνομων πολυπρακτορικών συστημάτων συναλλαγών, αλλά και της απουσίας κατάλληλα βελτιστοποιημένων πλαισίων που να αξιοποιούν πλήρως τις δυνατότητες των ΜΓΜ σε τόσο απαιτητικά πεδία.

Η παρούσα διπλωματική εργασία αντιμετωπίζει αυτά τα κενά μέσω μιας διττής συνεισφοράς. Πρώτον, παρουσιάζουμε το StockSim, μια πλατφόρμα προσομοίωσης ανοιχτού κώδικα που μοντελοποιεί τη λειτουργία των σύγχρονων χρηματοοικονομικών αγορών και υποστηρίζει την ανάπτυξη πρακτόρων συναλλαγών βασισμένων σε ΜΓΜ. Το StockSim υπερβαίνει τον παραδοσιακό και απλουστευμένο έλεγχο στρατηγικών πάνω σε ιστορικά δεδομένα (backtesting), προσομοιώνοντας τη δυναμική των πραγματικών αγορών. Στο περιβάλλον αυτό οι πράκτορες καταχωρούν πλήρως καθορισμένες εντολές αγοράς ή πώλησης - με συγκεκριμένο τύπο, τιμή και ποσότητα - οι οποίες εκτελούνται ρεαλιστικά, ακολουθώντας την εξέλιξη των τιμών στην αγορά. Παράλληλα, το σύστημα προσομοιώνει κρίσιμες παραμέτρους όπως ο χρονισμός των εντολών, οι καθυστερήσεις εκτέλεσης και η επίδραση της δραστηριότητας της αγοράς στη διαμόρφωση των τιμών. Οι ευέλικτες ρυθμίσεις του, οι ποικίλες πηγές δεδομένων και η υποστήριξη πολλαπλών πρακτόρων επιτρέπουν τη δημιουργία και αξιολόγηση προηγμένων στρατηγικών σε συνθήκες που προσεγγίζουν πιστά την πραγματικότητα.

Πάνω σε αυτή τη βάση, εισάγουμε το ATLAS (Adaptive Trading with LLM Agent Systems), ένα πολυπρακτορικό πλαίσιο συναλλαγών που συνδυάζει εξειδικευμένους αναλυτές για τις τάσεις της αγοράς, τα χρηματοοικονομικά νέα και τα θεμελιώδη στοιχεία των εταιρειών, συνθέτοντας τις επιμέρους τους οπτικές σε συνεκτικές στρατηγικές επένδυσης. Στον πυρήνα του, το ATLAS ενσωματώνει το Adaptive-OPRO, μια πρωτοποριακή μέθοδο βελτιστοποίησης μέσω Προτροπών που βελτιώνει επαναληπτικά τη διαδικασία λήψης αποφάσεων με βάση τα αποτελέσματα των συναλλαγών, επιτυγχάνοντας προοδευτικά καλύτερες επιδόσεις. Εκτενή πειράματα δείχνουν ότι το Adaptive-OPRO υπερτερεί σταθερά τόσο των παραδοσιακών ποσοτικών στρατηγικών όσο και των υφιστάμενων προσεγγίσεων βασισμένων σε ΜΓΜ, ενώ μελέτες αφαίρεσης επιβεβαιώνουν τη συμπληρωματική συνεισφορά κάθε συστατικού του πλαισίου. Παράλληλα, το ATLAS ενισχύει τη διαφάνεια της διαδικασίας λήψης αποφάσεων, προάγοντας την αξιόπιστη συνεργασία και την ασφαλέστερη ενσωμάτωση δίπλα σε επαγγελματίες του χώρου.

Τα αποτελέσματά μας αποκαλύπτουν διακριτά πρότυπα συμπεριφοράς ανάμεσα στα διάφορα ΜΓΜ, και αναδεικνύουν τις δυνατότητες και τους περιορισμούς τους σε χρηματοοικονομικά περιβάλλοντα υψηλού ρίσκου. Τα ευρήματα αυτά θεμελιώνονται σε αυστηρά πρωτόκολλα αξιολόγησης πολλαπλών εκτελέσεων τα οποία αποκαλύπτουν σοβαρά προβλήματα αξιοπιστίας στις αξιολογήσεις μίας μόνο εκτέλεσης που απαντώνται συχνά στη βιβλιογραφία, υπογραμμίζοντας την ανάγκη για εμπεριστατωμένη αξιολόγηση σε σύνθετα περιβάλλοντα διανοητικής λήψης αποφάσεων.¹

Λέξεις Κλειδιά — Μεγάλα Γλωσσικά Μοντέλα, Πολυπρακτορικά Συστήματα, Βελτιστοποίηση Προτροπών, Διαδοχική Λήψη Αποφάσεων, Χρηματοοικονομικές Συναλλαγές, Προσομοίωση Αγοράς, Ερμηνευσιμότητα, StockSim, ATLAS, Adaptive-OPRO

¹Ο πηγαίος κώδικας είναι διαθέσιμος στο github.com/harrypapa2002/StockSim.

Abstract

Large Language Models (LLMs) have demonstrated strong potential in complex decision-making tasks, showing promise for financial trading applications that demand the integration of diverse information sources, reasoning under uncertainty, and adaptation to rapidly changing market conditions. However, progress in this direction is limited by the absence of realistic market environments tailored to the assessment of trading agents, as well as a shortage of well-optimized frameworks capable of fully leveraging LLM capabilities in such challenging domains.

This thesis addresses these gaps through a two-stage contribution. First, we present StockSim, an open-source simulation platform that realistically models the behavior of financial markets and supports the development of LLM-based trading agents. StockSim extends beyond simplified historical backtesting by emulating the dynamics of live trading, where agents place fully specified orders - including type, price and quantity - that execute in alignment with actual price movements. The platform also captures order timing, execution delays, and the effects of market activity on price formation. Its flexible configurations, diverse data sources, and multi-agent support enable the creation and evaluation of advanced trading strategies under conditions that closely mirror real-world scenarios.

Building on this foundation, we introduce ATLAS (Adaptive Trading with LLM Agent Systems), a co-ordinated multi-agent trading framework that integrates specialized analysts for market trends, financial news, and company fundamentals, synthesizing these perspectives into coherent strategies. At its core, ATLAS incorporates Adaptive-OPRO, an enhanced Optimization by PROMpting method that iteratively refines decision-making based on trading outcomes, yielding progressively better performance. Extensive experiments show that Adaptive-OPRO consistently outperforms both traditional quantitative strategies and existing LLM-based approaches, while ablation studies confirm the complementary contributions of each component in the framework. ATLAS also improves transparency in the decision-making process, fostering trustworthy collaboration and enabling more reliable deployment alongside financial experts.

Our results reveal distinctive behavioral patterns in LLMs, providing new insights into their capabilities and limitations in high-stakes financial contexts. These findings are grounded in rigorous multi-run evaluation protocols that expose severe reliability issues in the single-run assessments common in prior literature, underscoring the need for robust evaluation in complex sequential decision-making settings.²

Keywords — Large Language Models, Multi-Agent Systems, Prompt Optimization, Sequential Decision-Making, Financial Trading, Market Simulation, Explainability, StockSim, ATLAS, Adaptive-OPRO

²Code available at github.com/harrypapa2002/StockSim.

Ευχαριστίες

Πίσω από κάθε σημαντικό έργο βρίσκονται άνθρωποι που το στήριξαν, το ενέπνευσαν και το έκαναν πραγματικότητα. Για μένα, αυτή η διπλωματική εργασία δεν ήταν απλώς μια ακαδημαϊκή υποχρέωση, αλλά ένα ταξίδι γεμάτο εμπειρίες, που μοιράστηκα με ανθρώπους οι οποίοι άφησαν το δικό τους ανεξίτηλο αποτύπωμα σε κάθε του βήμα.

Ευχαριστώ θερμά τον επιβλέποντά καθηγητή μου, κ. Στάμου Γεώργιο, για την ευκαιρία να εργαστώ στο εργαστήριο, για την εμπιστοσύνη, την καθοδήγηση και τη στήριξή του σε κάθε στάδιο.

Ένα ξεχωριστό ευχαριστώ οφείλω στους Διδάκτορες Γεώργιο Φιλανδριανό και Μαρία Λυμπεραίου και στους Υποψήφιους Διδάκτορες Κωνσταντίνο Θωμά και Αγγελική Δημητρίου. Οι γνώσεις που μοιράστηκαν απλόχερα μαζί μου, οι ατέλειωτες συζητήσεις μας, η δουλειά μέχρι τα ξημερώματα για τις δημοσιεύσεις μας και η συνεχής υποστήριξή τους υπήρξαν ανεκτίμητες για μένα και καθοριστικές για όσα πετύχαμε μαζί. Πέρασαμε στιγμές που θα θυμάμαι πάντα ως ένα από τα πιο ζωντανά κεφάλαια της ακαδημαϊκής μου πορείας.

Τέλος, θέλω να εκφράσω την ευγνωμοσύνη μου στην οικογένεια και τους φίλους μου, που στάθηκαν δίπλα μου με αγάπη, κατανόηση και στήριξη. Ήταν εκεί για να με ενθαρρύνουν, να με στηρίζουν στις δύσκολες στιγμές και να μοιραστούν μαζί μου τις χαρές των μικρών και μεγάλων επιτυχιών.

Παπαδάκης Χαρίδημος, Οκτώβρης 2025

Contents

Contents	xiii
List of Figures	xvi
1 Εκτεταμένη Περίληψη στα Ελληνικά	1
1.1 Εισαγωγή	2
1.2 Σχετική Έρευνα	3
1.2.1 Πράκτορες LLM στις Χρηματοοικονομικές Αγορές	3
1.2.2 Τεχνικές Βελτιστοποίησης Προτροπών	4
1.2.3 Πλατφόρμες Προσομοίωσης Συναλλαγών	4
1.3 Μεθοδολογία	4
1.3.1 StockSim: Αρχιτεκτονική Προσομοίωσης Δύο Λειτουργιών	5
1.3.2 ATLAS: Πλαίσιο Πολυπρακτορικού Συντονισμού	6
1.4 Πειραματική Διάταξη	8
1.4.1 Επιλογή Μοντέλων	8
1.4.2 Καθεστώτα Αγοράς και Επιλογή Μετοχών	8
1.4.3 Αξιολόγηση Στρατηγικών Προτροπών	9
1.4.4 Μεθοδολογία και Μετρικές Αξιολόγησης	9
1.4.5 Σχεδιασμός Μελέτης Αφαίρεσης	10
1.5 Αποτελέσματα	11
1.5.1 Βελτιστοποίηση σε Ακολουθιακή Λήψη Αποφάσεων	11
1.5.2 Ανάλυση Συνεισφοράς Πρακτόρων	12
1.5.3 Συναλλακτική Συμπεριφορά ανά LLM	12
1.5.4 Ικανότητες Βελτιστοποίησης LLM	14
1.6 Συμπεράσματα	14
1.6.1 Συζήτηση	14
1.6.2 Μελλοντικές Κατευθύνσεις	15
2 Introduction	17
3 Background	21
3.1 Large Language Models and Transformer Architecture	22
3.1.1 Evolution of Language Models	22
3.1.2 Transformer Architecture	22
3.1.3 Scaling Laws and Emergent Capabilities	25
3.2 Training Paradigms for Large Language Models	25
3.2.1 Learning Paradigm Overview	25
3.2.2 Pretraining and Fine-Tuning	26
3.2.3 Domain Adaptation	26
3.3 Prompting and In-Context Learning	27
3.3.1 Prompting Methodology	27
3.3.2 Few-Shot and Zero-Shot Learning	27
3.3.3 Chain-of-Thought and Advanced Prompting	27

3.4	Large Reasoning Models	27
3.5	Multi-Agent Systems and Coordination	28
3.5.1	Agent-Based Architectures	28
3.5.2	Coordination Mechanisms	28
3.6	Financial Markets and Trading Foundations	29
3.6.1	Market Microstructure	29
3.6.2	Technical Analysis Framework	29
3.6.3	Fundamental Analysis Framework	31
3.6.4	Trading Across Market Regimes	32
4	Related Work	33
4.1	LLM Agents in Financial Markets	34
4.1.1	Early Developments in Financial NLP	34
4.1.2	Contemporary Multi-Agent Trading Systems	34
4.1.3	Advanced Architectural Innovations	34
4.1.4	Fundamental Limitations of Current Systems	35
4.2	Prompt Engineering and Optimization	35
4.2.1	Evolution from Manual to Systematic Approaches	35
4.2.2	Optimization by Prompting (OPRO)	36
4.2.3	Limitations in Sequential Decision-Making Contexts	36
4.3	Trading Simulation Platforms	37
4.3.1	The Evaluation Infrastructure Challenge	37
4.3.2	Current Platform Landscape and Technical Limitations	37
5	Methodology	39
5.1	StockSim: Dual-Mode Simulation Architecture	39
5.1.1	System Architecture Overview	40
5.1.2	Dual Execution Modes	40
5.1.3	Data Integration Framework	41
5.1.4	Agent Framework Design	42
5.1.5	Evaluation and Analysis Framework	42
5.2	ATLAS: Multi-Agent Coordination Framework	44
5.2.1	Architectural Design Principles	44
5.2.2	Market Intelligence Pipeline	44
5.2.3	Decision and Execution Layer	45
5.2.4	Adaptive-OPRO: Prompt Optimization for Sequential Decision-Making	45
6	Experimental Setup	47
6.1	StockSim Platform Validation	49
6.1.1	Experimental Design for Platform Validation	49
6.1.2	Determinism and Reproducibility	49
6.1.3	Scalability Analysis	49
6.1.4	Resource Efficiency and Practical Implications	49
6.1.5	Validation Conclusions	49
6.2	Model Selection and Architecture Analysis	50
6.2.1	Reasoning-Enhanced Models	50
6.2.2	Standard Architecture Comparison	51
6.2.3	Open-Source Alternative	51
6.3	Market Conditions and Asset Selection	51
6.4	Prompting Strategy Evaluation	52
6.4.1	Core Experimental Configurations	52
6.4.2	Extended Prompting Strategy Analysis	52
6.4.3	Experimental Control and Comparison Framework	53
6.5	Evaluation Methodology and Metrics	53
6.5.1	Multi-Run Evaluation Protocol	53
6.5.2	Comprehensive Performance Metrics	53

6.5.3	Baseline Strategy Comparison	54
6.6	Ablation Study Design	55
6.6.1	Configuration Selection	55
6.6.2	Systematic Component Removal	55
6.6.3	Excluded Components	55
6.6.4	Evaluation Protocol	55
7	Results and Analysis	57
7.1	Framework Resilience Across Market Dynamics	58
7.1.1	Optimization in Sequential Decision-Making	58
7.2	Agent Contribution Analysis	59
7.3	Trading Behavior Across LLMs	60
7.3.1	Claude Models: Contrasting Failure Modes	61
7.3.2	LLM Optimization Capabilities	62
7.4	Extended Performance Analysis and Strategic Insights	63
7.4.1	Risk-Adjusted Performance Validation	63
7.4.2	The Reflection Paradox Reinforced	64
7.4.3	Architectural Performance Patterns	64
7.4.4	Extended Prompting Strategy Analysis	65
7.5	Prompt Evolution Mechanism Analysis	67
7.5.1	Systematic Weakness Detection and Resolution	67
7.5.2	Progressive Prompt Evolution: From Generic Foundation to Optimized Performance	73
8	Conclusion	79
8.1	Discussion	79
8.2	Future Work	80
9	Appendix	83
10	Bibliography	95

List of Figures

1.3.1	Επισκόπηση αρχιτεκτονικής συστήματος του STOCKSIM και σχήματος εισόδου/εξόδου. Οι ενότητες κωδικοποιούνται χρωματικά ανά λειτουργία και αντιστοιχίζονται στο κεντρικό αρχείο διαμόρφωσης.	5
1.3.2	Επισκόπηση Πλαισίου ATLAS. Ο Κεντρικός Πράκτορας Συναλλαγών υποβάλλει εντολές στον Μηχανή Προσομοίωσης του StockSim μέσω προτροπών διαμορφωμένων από τρεις εξειδικευμένους αναλυτές και την προτεινόμενη τεχνική βελτιστοποίησης Adaptive-OPRO.	7
1.5.1	ROI σε τρεις μετοχές με χρήση Adaptive-OPRO	13
3.1.1	The Transformer: model architecture. The original encoder–decoder design stacks multi-head self-attention and position-wise feed-forward sublayers in both the encoder (left) and decoder (right) [65].	23
3.1.2	Self-Attention Mechanism	24
5.1.1	Overview of STOCKSIM’s system architecture and input/output scheme. Modules are color-coded by function and mapped to corresponding blocks in the centralized config file. This design supports flexible, code-free customization of simulation parameters, agent behavior, and data sources.	40
5.1.2	Example of an interactive chart generated by StockSim for the XOM stock using the Claude-4 model with the thinking mechanism enabled. The plot displays the price, buy and sell orders (annotated with ▲ and ▼, respectively), portfolio value, and trading volume. Users can zoom, adjust the time range, toggle chart components, and hover over elements to reveal additional details such as order execution prices and corresponding LLM outputs.	43
5.1.3	Demonstration of the hover functionality in StockSim. When hovering over a specific order, detailed information is displayed, including the exact execution price and the corresponding LLM output that led to the decision.	43
5.2.1	ATLAS Framework Overview. The Central Trading Agent submits orders to the Trading Execution Engine via prompts shaped by three specialized analysts and the proposed Adaptive-OPRO optimization technique.	44
6.1.1	StockSim system performance metrics (memory/CPU usage) for varying numbers of deterministic agents.	50
7.2.1	ROI across three assets using Adaptive-OPRO.	61
7.4.1	Daily vs weekly reflection mechanism performance comparison across models and assets, showing ROI percentages (solid = daily, striped = weekly).	65
7.5.1	Header and trader identity modifications between iteration 4 and iteration 5, showing title changes and mission statement refinements.	70
7.5.2	Structural reorganization consolidating sections 4 and 5 into a unified PORTFOLIO & CONSTRAINTS section.	71
7.5.3	Decision protocol restructuring from informal REVIEW → REASON → RESPOND to structured five-step THINK → CHECK → ACT workflow.	72
7.5.4	Baseline prompt structure (GPT-o4-mini, Prompt 1, Part 1) demonstrating expert-crafted foundation with comprehensive trading philosophy and toolkit specification. Score: 37.2 . . .	74

7.5.5 Baseline prompt structure (GPT-o4-mini, Prompt 1, Part 2) showing detailed constraint specification and output formatting requirements.	75
7.5.6 Intermediate optimization (GPT-o4-mini, Prompt 4) featuring streamlined structure with numbered five-step decision framework and more concise constraint presentation. Score: 51.4	76
7.5.7 Final optimized prompt (GPT-o4-mini, Prompt 11) featuring expanded six-step decision framework with granular step descriptions and systematic market context organization. Score: 72.1	77

Chapter 1

Εκτεταμένη Περίληψη στα Ελληνικά

1.1 Εισαγωγή

Τα τελευταία χρόνια, η ραγδαία εξέλιξη των Μεγάλων Γλωσσικών Μοντέλων (Large Language Models, LLMs) έχει αναδιαμορφώσει τον τρόπο με τον οποίο αντιλαμβανόμαστε τις δυνατότητες της Τεχνητής Νοημοσύνης σε προβλήματα σύνθετης συλλογιστικής και λήψης αποφάσεων. Τα σύγχρονα LLMs μπορούν και επεξεργάζονται τεράστιους όγκους ετερογενών δεδομένων, συνθέτουν πληροφορίες από πολλαπλές πηγές και παράγουν συνεκτικές, τεκμηριωμένες αποκρίσεις σε ένα ευρύ φάσμα εφαρμογών. Καθώς προσεγγίζουν και συχνά υπερβαίνουν τις ανθρωπίνες επιδόσεις σε ολοένα και πιο απαιτητικά προβλήματα, αυξάνεται αντίστοιχα το ερευνητικό και πρακτικό ενδιαφέρον για την ασφαλή και αποτελεσματική αξιοποίησή τους σε πραγματικές εφαρμογές με υψηλό διακύβεσμα.

Στο πλαίσιο αυτό, οι χρηματοοικονομικές αγορές αποτελούν ένα από τα πιο απαιτητικά πεδία δοκιμής. Πρόκειται για περιβάλλοντα λήψης αποφάσεων με υψηλή αβεβαιότητα, ουσιώδεις συνέπειες και σαφώς μετρήσιμα αποτελέσματα [78, 58, 47]. Καθημερινά, οι συμμετέχοντες συνδυάζουν τεχνικούς δείκτες, θεμελιώδη ανάλυση, ειδήσεις και επενδυτικό αίσθημα, λαμβάνοντας αποφάσεις των οποίων η ποιότητα συχνά αποκαλύπτεται μόνο σε βάθος χρόνου. Τα LLMs, με την ικανότητα πολυτροπικής κατανόησης και την προσαρμοστικότητά που επιδεικνύουν, διευρύνουν τις δυνατότητες ενοποίησης ετερογενών δεδομένων και σχεδιασμού στρατηγικών [9, 43].

Πέρα από τη χρηστική τους αξία, οι αγορές διαθέτουν χαρακτηριστικά που τις καθιστούν ιδανικό πεδίο αξιολόγησης των LLMs. Σε αντίθεση με συνθετικά πρότυπα αξιολόγησης (benchmarks), που συχνά βασίζονται σε απλουστεύσεις ή μετατοπίσεις κατανομής, οι χρηματιστηριακές χρονοσειρές προσφέρουν εκτενή ιστορικά δεδομένα απαλλαγμένα τις στρεβλώσεις των τεχνητών υποθέσεων μιας προσομοίωσης, ενώ απαιτούν γνήσια κατανόηση αντί για επιφανειακή αναγνώριση προτύπων. Το πεδίο απαιτεί την ενοποίηση δομημένων αριθμητικών στοιχείων (τιμές, δείκτες) με αδόμητο κείμενο (ειδήσεις, αναλύσεις), αξιολογώντας την ικανότητα των LLM να συλλογίζονται πολυτροπικά. Επιπλέον, οι χρηματοοικονομικές αγορές διακρίνονται από υψηλή στοχαστικότητα και έντονη δυναμική, χαρακτηριστικά που εκθέτουν γρήγορα εύθραυστες ή υπερπροσαρμοσμένες προσεγγίσεις και ανταμείβουν τη γνήσια ανθεκτικότητα.

Παρά τις πολλά υποσχόμενες προοπτικές, η ανάπτυξη αποτελεσματικών συστημάτων συναλλαγών βασισμένων σε LLMs αντιμετωπίζει σημαντικές δυσκολίες. Συγκεκριμένα παρατηρείται έλλειψη τυποποιημένων, ολοκληρωμένων πλαισίων προσομοίωσης ανοικτού κώδικα, ειδικά σχεδιασμένων για αυστηρή, ρεαλιστική αξιολόγηση σε περιβάλλοντα συναλλαγών [37, 44]. Οι υπάρχουσες λύσεις είτε καταφεύγουν σε απλουστευμένες προσομοιώσεις που παραλείπουν κρίσιμες λεπτομέρειες μικροδομής (καθυστερήσεις, βιβλίο εντολών, επιπτώσεις στην αγορά), είτε αποδίδουν πιστά τη μηχανική της αγοράς αλλά βασίζονται σε ακριβή και περιορισμένα δεδομένα, υπονομεύοντας τη δυνατότητα αναπαραγωγής και ευρείας υιοθέτησης.

Επιπλέον, μεγάλο μέρος των τρέχοντων προσεγγίσεων εμφανίζει τρεις δομικές αδυναμίες: (i) εξάρτηση από στατικές προτροπές (prompts) χωρίς την δυνατότητα προσαρμογής σε μεταβαλλόμενες συνθήκες· (ii) απομονωμένη λήψη αποφάσεων, χωρίς συντονισμό εξειδικευμένων αναλυτών σε διαφορετικές πτυχές της αγοράς· και (iii) υπεραπλούστευση του χώρου ενεργειών, με έμφαση σε προβλέψεις κατεύθυνσης αντί για πλήρεις προδιαγραφές εντολών (τύπος, ποσότητα, τιμή, χρονισμός, διαχείριση ρίσκου).

Η παρούσα εργασία αντιμετωπίζει αυτά τα κενά με μια διττή συνεισφορά: αφενός θέτει την υποδομή για συστηματική και ρεαλιστική αξιολόγηση πρακτόρων βασισμένων σε LLMs σε χρηματοοικονομικές αγορές και αφετέρου προτείνει μεθοδολογικές προσεγγίσεις για τον αποτελεσματικό σχεδιασμό και τη συνεχή βελτιστοποίηση αυτών των συστημάτων.

Η πρώτη συνεισφορά είναι το StockSim: μία πλατφόρμα προσομοίωσης ανοικτού κώδικα που γεφυρώνει το χάσμα μεταξύ απλουστευμένης αναδρομικής αξιολόγησης σε ιστορικά δεδομένα (backtesting) και πραγματικών συνθηκών εκτέλεσης. Προσφέρει δύο συμπληρωματικούς τρόπους αξιολόγησης:

- *Επίπεδο εντολών (order-level)*: προσομοιώνει λανθάνουσες καθυστερήσεις, αλληλεπίδραση με το βιβλίο εντολών και ρεαλιστική εκτέλεση με πιθανές επιπτώσεις στην αγορά.
- *Επίπεδο γραφημάτων κεριών (candlestick level)*: επιτρέπει κλιμακούμενα πειράματα σε εκτεταμένα χρονικά διαστήματα και πολλαπλά σενάρια, με κανόνες εκτέλεσης ευθυγραμμισμένους με τη χρονική ανάλυση των δεδομένων.

Η αριθρωτή αρχιτεκτονική του StockSim επιτρέπει πολυπρακτορικές, συνεργατικές ροές ανάλυσης, απρόσκοπτη ενσωμάτωση εξωτερικών πηγών πληροφόρησης (ειδήσεις, οικονομικές εκθέσεις), πρόσβαση σε αυθεντικά δε-

δομένα και ενιαία, συστηματική παρακολούθηση της απόδοσης-με σχεδιασμό που δύναται να λειτουργήσει αξιόπιστα σε ρεαλιστικά περιβάλλοντα αγοράς.

Με βάση αυτή την υποδομή, παρουσιάζουμε το ATLAS (Adaptive Trading with LLM Agent Systems): ένα πολυπρακτορικό πλαίσιο λήψης αποφάσεων, όπου εξειδικευμένοι πράκτορες επικεντρώνονται σε τεχνική ανάλυση, ειδησεογραφική σύνθεση και θεμελιώδη αξιολόγηση, ενώ ένας κεντρικός πράκτορας/συντονιστής ενοποιεί τα επιμέρους πορίσματα σε συνεκτικές στρατηγικές συναλλαγών. Στον πυρήνα του ATLAS βρίσκεται το Adaptive-OPRO, μια προσαρμογή της μεθόδου Optimization by PROMpting [79] σε περιβάλλοντα διαδοχικών αποφάσεων με ετεροχρονισμένα και θορυβώδη σήματα απόδοσης. Ο μηχανισμός αυτός επιτρέπει συστηματική αναδιαμόρφωση των προτροπών (*prompts*) βάσει της ανατροφοδότησης από τα αποτελέσματα των συναλλαγών, βελτιώνοντας σταδιακά την ποιότητα των αποφάσεων.

Η εκτεταμένη πειραματική αξιολόγηση οδηγεί σε συμπεράσματα με εμβέλεια που υπερβαίνει αυτή των χρηματοοικονομικών αγορών, παρέχοντας πρακτικές κατευθύνσεις για την ευρύτερη αξιοποίηση των LLMs σε σύνθετα πεδία με σημαντικό διακύβευμα. Διαπιστώνουμε ότι το Adaptive-OPRO επιτυγχάνει σταθερά ανώτερες επιδόσεις σε διαφορετικά καθεστώτα αγοράς, υπερτερώντας τόσο έναντι κλασικών ποσοτικών στρατηγικών όσο και σύγχρονων προσεγγίσεων βασισμένων σε LLMs. Παράλληλα, αναδύεται το «παράδοξο του αναστοχασμού» σύμφωνα με το οποίο η αύξηση των βημάτων συλλογισμού, αν και συνήθως θεωρείται επωφελής, μπορεί να υποβαθμίσει την επίδοση σε καλά βελτιστοποιημένα συστήματα. Τέλος, καθίσταται σαφές ότι απαιτούνται πρωτόκολλα αξιολόγησης με πολλαπλές εκτελέσεις και επαρκή στατιστική τεκμηρίωση, καθώς οι μετρήσεις μίας μόνο εκτέλεσης αποδεικνύονται ευμετάβλητες σε στοχαστικά περιβάλλοντα και υπονομεύουν την αξιοπιστία και τη γενικευσιμότητα των συμπερασμάτων.

Συνολικά, τα ευρήματα χαράσσουν θεμελιώδεις αρχές για τον σχεδιασμό και την ασφαλή αξιοποίηση των LLMs σε πεδία υψηλού ρίσκου με διαδοχική λήψη αποφάσεων υπό πραγματική αβεβαιότητα, ενώ παράλληλα εμπλουτίζουν την κατανόησή μας για τη συμπεριφορά των LLMs, τις μεθόδους βελτιστοποίησης και τους μηχανισμούς συντονισμού τους σε σύνθετα περιβάλλοντα.

1.2 Σχετική Έρευνα

Η τομή Μεγάλων Γλωσσικών Μοντέλων (LLMs) και χρηματοοικονομικής λήψης αποφάσεων αποτελεί ένα ταχύτατα εξελισσόμενο πεδίο έρευνας, το οποίο γεφυρώνει την επεξεργασία φυσικής γλώσσας, με τα πολυπρακτορικά συστήματα την και υπολογιστική χρηματοοικονομική. Καθώς τα LLMs επιδεικνύουν ολοένα εντυπωσιακότερες ικανότητες σύνθετης συλλογιστικής, η έρευνα μετατοπίζεται σε περιβάλλοντα διαδοχικής λήψης αποφάσεων υψηλού διακυβέματος, με χαρακτηριστικό παράδειγμα τις χρηματοοικονομικές αγορές. Στην ενότητα αυτή συνοψίζουμε τη σχετική βιβλιογραφία σε τρεις άξονες: (i) πράκτορες LLM στις χρηματοοικονομικές αγορές, (ii) μεθοδολογίες *prompting* και βελτιστοποίησης προτροπών, και (iii) πλατφόρμες προσομοίωσης και αξιολόγησης.

1.2.1 Πράκτορες LLM στις Χρηματοοικονομικές Αγορές

Η πρώιμη έρευνα αξιοποίησε γλωσσικά μοντέλα για ανάλυση κειμένων και εξαγωγή επενδυτικού αισθήματος, θέτοντας τα θεμέλια για παραγωγή πιο προηγμένων συστημάτων συναλλαγών [31]. Εξειδικευμένα γλωσσικά μοντέλα, όπως τα **FinBERT**, **FinGPT** και **BloombergGPT**, έδειξαν την αξία της προσαρμογής πεδίου σε χρηματοοικονομικές εργασίες [81, 80, 73]. Ωστόσο, οι μονοτροπικές αυτές προσεγγίσεις δεν επαρκούσαν για ολοκληρωμένη λήψη αποφάσεων σε ρεαλιστικά σενάρια συναλλαγών.

Νεότερα έργα υιοθετούν πολυπρακτορικές αρχιτεκτονικές. Το **CryptoTrade** συνδυάζει ετερογενείς ροές δεδομένων μέσω εξειδικευμένων αναλυτών και μηχανισμών αναστοχασμού [38], ενώ το **TradingAgents** υιοθετεί μια πιο σύνθετη πολυπρακτορική αρχιτεκτονική, εμπνευσμένη από τη λειτουργία πραγματικών επενδυτικών οργανισμών, με συντονισμό εξειδικευμένων ρόλων, δομημένες εσωτερικές αντιπαραθέσεις και ενσωματωμένο έλεγχο κινδύνου, οδηγώντας σε σημαντικές βελτιώσεις έναντι baseline μεθόδων [74]. Επιπλέον, το **FinMem** εισάγει πολυεπίπεδη μνήμη για διατήρηση συγχειμένου [83], το **FLAG-Trader** συνδυάζει την γλωσσική επεξεργασία με την ενισχυτική μάθηση [76], ενώ τα **TradExpert** και **MarketSenseAI** εστιάζουν σε αρχιτεκτονικές *μήματος ειδικών* (mixture-of-experts) και στοχευμένη ανάλυση εταιρικών αναφορών [14, 18].

Παρά τη σημαντική πρόοδο, εξακολουθούν να υφίστανται τρεις θεμελιώδεις περιορισμοί: (i) **στατικά prompts** που δεν προσαρμόζονται σε μεταβαλλόμενες συνθήκες, (ii) **υπεραπλούστευση του χώρου ενεργειών**

(π.χ. κατευθυντήριες προβλέψεις αντί πλήρων προδιαγραφών εντολών), και (iii) *ανεπαρκείς μέθοδοι αξιολόγησης*, όπου οι μετρήσεις μίας μόνο εκτέλεσης αποκρύπτουν τη στοχαστική μεταβλητότητα των αποτελεσμάτων [62, 2].

1.2.2 Τεχνικές Βελτιστοποίησης Προτροπών

Η βελτιστοποίηση προτροπών έχει εξελιχθεί από *ad-hoc*, πειραματικές προσεγγίσεις σε συστηματικές, επιστημονικά θεμελιωμένες μεθοδολογίες. Τεχνικές όπως το **Chain-of-Thought (CoT)** επιτρέπουν στα LLMs να αποσυνθέτουν σύνθετα προβλήματα σε ενδιάμεσα βήματα συλλογισμού [70], ενώ προχωρημένες μέθοδοι όπως το **self-consistency** και το **tree-of-thoughts (ToT)** ενισχύουν περαιτέρω τις συλλογιστικές ικανότητες των μοντέλων [67, 82].

Η εισαγωγή του **Optimization by PROMpting (OPRO)** σηματοδότησε σημαντική πρόοδο στο πεδίο, χρησιμοποιώντας τα ίδια τα LLMs ως μετα-βελτιστοποιητές που παράγουν επαναληπτικά νέες προτροπές βάσει ιστορικών επιδόσεων [79]. Παράλλαγες όπως το **EvoPrompt** και το **POEM** ενσωματώνουν εξελικτικούς αλγόριθμους και ενισχυτική μάθηση αντίστοιχα [24, 15].

Περιορισμοί σε ακολουθιακή λήψη αποφάσεων. Οι υπάρχουσες μέθοδοι βελτιστοποίησης προτροπών προϋποθέτουν συνθήκες που απουσιάζουν από πραγματικά περιβάλλοντα ακολουθιακών αποφάσεων: απαιτούν *άμεση ανάδραση*, λειτουργούν σε *ντετερμινιστικά αποτελέσματα*, και υποθέτουν ότι οι αποφάσεις είναι *ανεξάρτητες*. Οι χρηματοοικονομικές συναλλαγές παρουσιάζουν ακριβώς το αντίθετο σενάριο: οι ανταμοιβές αποκαλύπτονται με χρονική υστέρηση, η μεταβλητότητα θολώνει τις αξιολογήσεις απόδοσης και οι αποφάσεις έχουν διαχρονικές συνέπειες στην εξέλιξη του συστήματος.

1.2.3 Πλατφόρμες Προσομοίωσης Συναλλαγών

Η αξιολόγηση συστημάτων συναλλαγών βασισμένων σε LLM απαιτεί πλατφόρμες προσομοίωσης που μοντελοποιούν με ρεαλισμό και πληρότητα τη δυναμική της αγοράς. Ωστόσο, το τρέχον τοπίο παρουσιάζει ένα κατακερματισμένο οικοσύστημα με σημαντικούς περιορισμούς.

Προσομοίωση σε επίπεδο εντολών. Πλατφόρμες όπως το **ABIDES**, το **PyMarketSim**, και το **JAX-LOB** παρέχουν εξελιγμένη μοντελοποίηση μικροδομής αγοράς με λεπτομερή δυναμική του βιβλίου εντολών [8, 45, 21]. Παρά τον ρεαλισμό τους, υστερούν σε εγγενή υποστήριξη για πράκτορες βασισμένους LLM και απαιτούν εκτεταμένη προσαρμογή για ενσωμάτωση γλωσσικών μοντέλων. Επιπλέον, βασίζονται συχνά σε υψηλού κόστους δεδομένα σε επίπεδο *tick* με περιορισμένη προσβασιμότητα, γεγονός που αποτελεί τροχοπέδη για την κλιμακωσιμότητα των πειραμάτων.

Πλατφόρμες ιστορικής προσομοίωσης. Συστήματα όπως τα *BackTrader* και το **FinRL/Meta** προτάσσουν την κλιμακωσιμότητα έναντι του ρεαλισμού εκτέλεσης, βασιζόμενα σε συγκεντρωτικά δεδομένα κερών χωρίς όμως να μοντελοποιούν τις μεταβολές της τιμής εντός αυτών των διαστημάτων, ενώ εφαρμόζουν απλοποιημένα μοντέλα εκτέλεσης τα οποία αγνοούν την πραγματική δυναμική των αγορών και περιορισμούς των συναλλαγών [41, 42].

Εξειδικευμένα πλαίσια για LLM. Πλατφόρμες σχεδιασμένες για πολυπρακτορικό συντονισμό αντιμετωπίζουν τις προκλήσεις ενσωμάτωσης Μεγάλων Γλωσσικών Μοντέλων, ωστόσο συχνά βασίζονται σε υπεραπλουστευμένες αναπαραστάσεις της αγοράς που παραβλέπουν τη ρεαλιστική δυναμική εκτέλεσης εντολών [74].

Αυτός ο κατακερματισμός δημιουργεί συστηματικά εμπόδια όπου οι ερευνητές πρέπει να επιλέξουν μεταξύ ρεαλιστικής μοντελοποίησης αγοράς και πρακτικής κλιμακωσιμότητας, επιβάλλοντας συμβιβασμούς που περιορίζουν τόσο το εύρος όσο και την αξιοπιστία των ερευνητικών ευρημάτων.

1.3 Μεθοδολογία

Το κεφάλαιο αυτό παρουσιάζει τη λεπτομερή μεθοδολογία των δύο κύριων συνεισφορών μας: το **StockSim**, μια ολοκληρωμένη πλατφόρμα προσομοίωσης για την αξιολόγηση πρακτόρων συναλλαγών βασισμένων σε LLM, και το **ATLAS**, ένα προσαρμοστικό πολυπρακτορικό πλαίσιο συναλλαγών που αξιοποιεί αποτελεσματικά τα LLMs στη χρηματοοικονομική λήψη αποφάσεων.

Η μεθοδολογία μας δομείται γύρω από τρεις βασικές καινοτομίες: την αρχιτεκτονική δύο λειτουργιών του StockSim που επιτρέπει συστηματική αξιολόγηση σε διαφορετικά επίπεδα πολυπλοκότητας της αγοράς, το πλαίσιο πολυπρακτορικού συντονισμού του ATLAS που αποσυνθέτει τη χρηματοοικονομική λήψη αποφάσεων σε εξειδικευμένες συνιστώσες, και τον αλγόριθμο Adaptive-OPRO που επεκτείνει τη βελτιστοποίηση προτροπών σε περιβάλλοντα ακολουθιακών αποφάσεων με καθυστερημένη ανάδραση.

1.3.1 StockSim: Αρχιτεκτονική Προσομοίωσης Δύο Λειτουργιών

Το StockSim αντιμετωπίζει το κρίσιμο κενό στην υποδομή αξιολόγησης προτείνοντας μια ενοποιημένη πλατφόρμα που συνδυάζει ρεαλιστική προσομοίωση αγοράς με ολοκληρωμένες δυνατότητες αξιολόγησης των LLM.

Επισκόπηση Αρχιτεκτονικής Συστήματος

Το StockSim υιοθετεί μια αρθρωτή, ασύγχρονη αρχιτεκτονική σχεδιασμένη γύρω από τέσσερις βασικές συνιστώσες που επιτρέπουν ολοκληρωμένη αξιολόγηση των LLMs σε ρεαλιστικά περιβάλλοντα συναλλαγών. Η αρχιτεκτονική υποστηρίζει δύο διακριτούς μηχανισμούς εκτέλεσης, εκτέλεση σε επίπεδο εντολών και σε επίπεδο κεριών, οι οποίοι ενοποιούνται μέσω κοινών λειτουργιών για πρόσβαση σε δεδομένα της αγοράς, παραγωγή τεχνικών δεικτών και αλληλεπίδραση μεταξύ πρακτόρων.

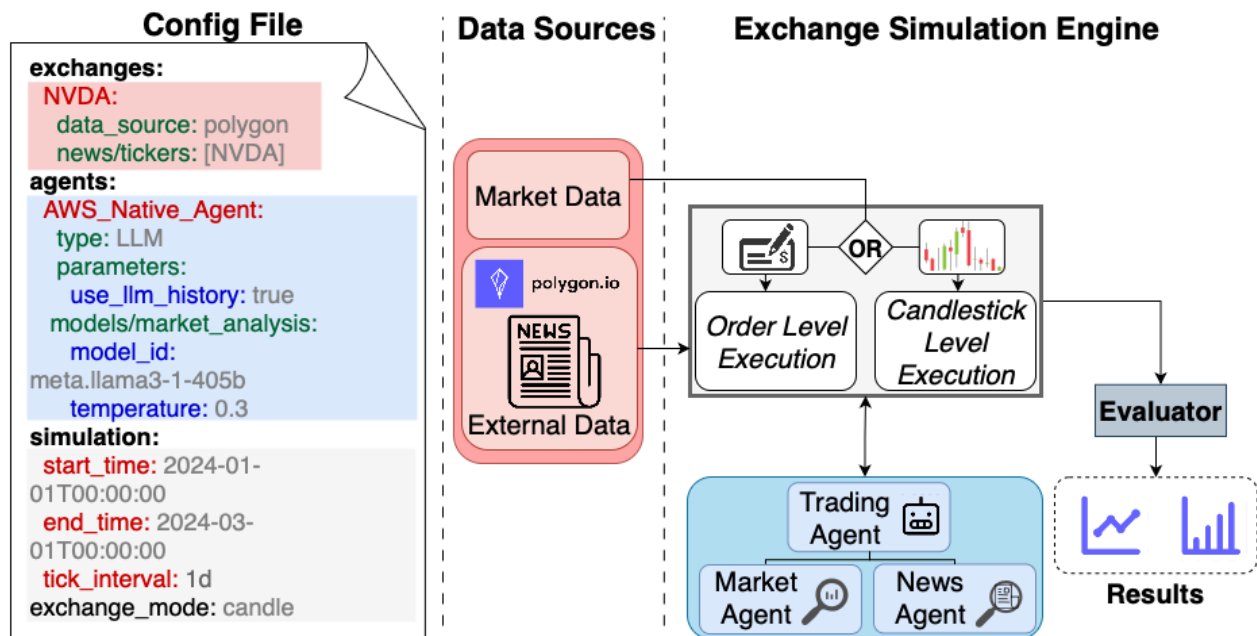


Figure 1.3.1: Επισκόπηση αρχιτεκτονικής συστήματος του STOCKSIM και σχήματος εισόδου/εξόδου. Οι ενότητες κωδικοποιούνται χρωματικά ανά λειτουργία και αντιστοιχίζονται στο κεντρικό αρχείο διαμόρφωσης.

Η **Μηχανή Προσομοίωσης Χρηματιστηριακής Αγοράς** λειτουργεί ως κεντρικός ενορχηστρωτής του περιβάλλοντος συναλλαγών, διαχειριζόμενη ασύγχρονα ροές γεγονότων και εκτέλεσης. Αναλαμβάνει την επεξεργασία των ενεργειών των πρακτόρων, τον υπολογισμό τεχνικών δεικτών, καθώς και την διανομή των δεδομένων της αγοράς στους συμμετέχοντες. Κάθε πράκτορας εκτελείται ως ανεξάρτητη διεργασία και επικοινωνεί ασύγχρονα με τη μηχανή μέσω ενός μεσολαβητή μηνυμάτων **RabbitMQ**, εξασφαλίζοντας απομόνωση σφαλμάτων και οριζόντια κλιμάκωση.

Διπλή Λειτουργία Εκτέλεσης

Εκτέλεση σε επίπεδο εντολών. Η εκτέλεση σε επίπεδο εντολών αναπαράγει ρεαλιστικά τη μικροδομή της αγοράς, δρώντας απευθείας στο **βιβλίο εντολών (LOB)**. Οι πράκτορες υποβάλλουν εντολές αγοράς, ορίου και ενεργοποίησης, οι οποίες εισάγονται στο βιβλίο εντολών και αντιστοιχίζονται με τις αντίθετες εντολές καθώς καταφθάνουν νέα μηνύματα (νέες εντολές και ακυρώσεις), βάσει προτεραιότητας τιμής-χρόνου.

Η λειτουργία αυτή προσομοιώνει την ρεαλιστική δυναμική της αγοράς:

- **Μοντελοποίηση καθυστέρησης:** ρεαλιστικές καθυστερήσεις μεταξύ υποβολής και εκτέλεσης εντολών
- **Επίπτωση στην αγορά:** μεταβολές τιμών από μεγάλες εντολές
- **Ολίσθηση:** διαφορά μεταξύ προβλεπόμενων και πραγματικών τιμών εκτέλεσης
- **Τμηματικές εκτελέσεις:** όταν ο διαθέσιμος όγκος στο στοχευόμενο επίπεδο τιμής είναι ανεπαρκής.

Εκτέλεση σε Επίπεδο Κεριών. Η εκτέλεση σε επίπεδο κεριών λειτουργεί σε συγκεντρωτικά δεδομένα **OHLCV**, όπου οι εντολές εκτελούνται βάσει του εάν η τιμή-στόχος βρίσκεται εντός του εύρους ενός κεριού. Παρέχει πρόσβαση σε μεγαλύτερα σύνολα δεδομένων και επιτρέπει δοκιμές σε εκτεταμένες ιστορικές περιόδους.

Παρά το απλοποιημένο μοντέλο εκτέλεσης, η λειτουργία διατηρεί ρεαλιστικές δυναμικές συναλλαγών μέσω προηγμένης προσομοίωσης των τιμών εντός κεριών, επιτρέποντας στους πράκτορες να υποβάλλουν εντολές που εκτελούνται ρεαλιστικά.

Πλαίσιο Ενοποίησης Δεδομένων

Το StockSim διαχειρίζεται δεδομένα της αγοράς (τιμή, όγκος, ροή εντολών) αλλά και εξωγενή δεδομένα (ειδήσεις, εταιρικές πράξεις, θεμελιώδη μεγέθη). Η **Μηχανή Προσομοίωσης** ενορχηστρώνει αυτές τις ροές ασύγχρονα και τις τροφοδοτεί στους πράκτορες σε χρόνο προσομοίωσης.

Η πλατφόρμα υποστηρίζει δεδομένα τόσο σε επίπεδο εντολής όσο και σε επίπεδο κεριών. Σε όλες τις λειτουργίες εκτέλεσης, οι πράκτορες λαμβάνουν υπολογισμένους τεχνικούς δείκτες-κινητούς μέσους, ταλαντωτές ορμής και μετρικές μεταβλητότητας.

Σχεδιασμός Πλαισίου Πρακτόρων

Το πλαίσιο πρακτόρων παρέχει μια ενοποιημένη διεπαφή που αποκρύπτει την πολυπλοκότητα διαφορετικών λειτουργιών εκτέλεσης. Κάθε πράκτορας μπορεί να εγγράφεται σε ροές δεδομένων, να υποβάλλει και να ακυρώνει εντολές και να λαμβάνει αποτελέσματα εκτέλεσης και ενημερώσεις χαρτοφυλαχίου.

Το StockSim περιλαμβάνει ένα αρθρωτό πλαίσιο πρακτόρων που υποστηρίζει συντονισμό μεταξύ εξειδικευμένων LLMs εστιασμένων σε διαφορετικές πτυχές ανάλυσης της αγοράς.

Πλαίσιο Αξιολόγησης και Ανάλυσης

Ο Αξιολογητής καταγράφει το πλήρες ιστορικό θέσεων, μετρητών και πραγματοποιημένων κερδών/ζημιών. Με την ολοκλήρωση της προσομοίωσης, υπολογίζει τις βασικές μετρικές απόδοσης και παράγει διαδραστικά γραφήματα για περαιτέρω διερεύνηση.

1.3.2 ATLAS: Πλαίσιο Πολυπρακτορικού Συντονισμού

Το ATLAS (Adaptive Trading with LLM Agent Systems) επιδεικνύει αποτελεσματική εφαρμογή των LLMs στη χρηματοοικονομική λήψη αποφάσεων μέσω συντονισμένης πολυπρακτορικής αρχιτεκτονικής και προσαρμοστικής βελτιστοποίησης προτροπών.

Αρχές Αρχιτεκτονικού Σχεδιασμού

Το **ATLAS** εδράζεται σε δύο θεμελιώδεις πυλώνες: μια αρθρωτή, πολυπρακτορική αρχιτεκτονική που διαχωρίζει τις λειτουργικές ευθύνες ανά τύπο πληροφορίας της αγοράς και έναν εξειδικευμένο μηχανισμό βελτιστοποίησης προτροπών, προσαρμοσμένο στη σειριακή λήψη αποφάσεων υπό αβεβαιότητα.

Διάυλος Νοημοσύνης Αγοράς

Το ATLAS χρησιμοποιεί εξειδικευμένους πράκτορες που εστιάζουν σε διακριτούς τύπους πληροφοριών της αγοράς.

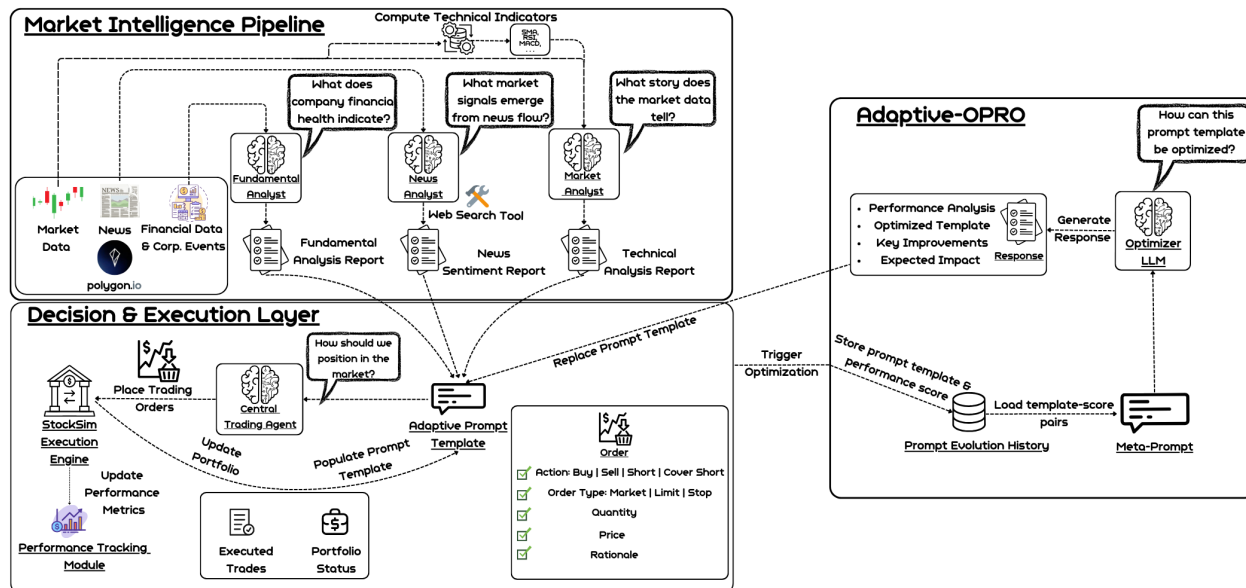


Figure 1.3.2: Επισκόπηση Πλαισίου ATLAS. Ο Κεντρικός Πράκτορας Συναλλαγών υποβάλλει εντολές στον Μηχανή Προσομοίωσης του StockSim μέσω προτροπών διαμορφωμένων από τρεις εξειδικευμένους αναλυτές και την προτεινόμενη τεχνική βελτιστοποίησης Adaptive-OPRO.

Αναλυτής Αγοράς. Ο Αναλυτής Αγοράς επεξεργάζεται ακατέργαστα δεδομένα τιμής και όγκου και παράγει δομημένες περιλήψεις σε πολλαπλές χρονικές κλίμακες. Αναλύει κάθε μετοχή χρησιμοποιώντας τρία διακριτά ιστορικά παράθυρα: 2 έτη με μηνιαία κεριά, 6 μήνες με εβδομαδιαία κεριά, και 3 μήνες με ημερήσια κεριά.

Εντός κάθε παραθύρου, χρησιμοποιεί καθιερωμένους τεχνικούς δείκτες και παρέχει μια συνεπή, φιλτραρισμένη άποψη που διατηρεί την κύρια δομή της αγοράς απαλλαγμένη από την πολυπλοκότητα ακατέργαστων δεδομένων.

Αναλυτής Ειδήσεων. Ο Αναλυτής Ειδήσεων εξάγει επενδυτικό αίσθημα και σήματα από χρηματοοικονομικές ειδήσεις επεξεργαζόμενος επικεφαλίδες, πηγές και περιλήψεις για παραγωγή δομημένων εξόδων σε τέσσερις αναλυτικές διαστάσεις: Αξιολόγηση Αισθήματος, Βασικές Εξελίξεις, Συνάφεια με την Αγορά, και Ανάλυση Πηγών.

Όταν οι επικεφαλίδες δεν παρέχουν επαρκείς λεπτομέρειες, μπορεί αυτόνομα να ανακτά το πλήρες κείμενο άρθρων μέσω ενός προσαρμοσμένου εργαλείου αυτοματοποιημένης εξαγωγής περιεχομένου από ιστοσελίδες.

Θεμελιώδης Αναλυτής. Ο Θεμελιώδης Αναλυτής εξάγει στοιχεία σχετιζόμενα με επενδυτικές αποφάσεις από περιοδικές εταιρικές δημοσιοποιήσεις συμπεριλαμβανομένων χρηματοοικονομικών καταστάσεων και γεγονότων. Αναλύει λεπτομερή πεδία δεδομένων όπως έσοδα, περιθώρια κέρδους, και δυναμική ταμειακών ροών και παράγει μια συνοπτική ανάλυση εστιασμένη σε ουσιαστικές αλλαγές και τις επενδυτικές προεκτάσεις τους.

Επίπεδο Απόφασης και Εκτέλεσης

Κεντρικός Πράκτορας Συναλλαγών. Ο Κεντρικός Πράκτορας Συναλλαγών λειτουργεί ως κύριος κόμβος λήψης αποφάσεων, συνθέτοντας τα ευρήματα των αναλυτών σε εκτελέσιμες στρατηγικές ενέργειες. Κάθε ημέρα, πριν το άνοιγμα αγοράς, αξιολογεί τις πιο πρόσφατες αναλύσεις τιμής/όγκου, ειδήσεων και θεμελιωδών δεδομένων σε συνδυασμό με την τρέχουσα κατάσταση του χαρτοφυλαχίου του.

Με βάση αυτό το ενοποιημένο πλαίσιο, τοποθετεί πλήρως προσδιορισμένες εντολές-καθορίζοντας ποσότητα, τιμή, τυχόν προϋποθέσεις εκτέλεσης και τύπο εντολής-οι οποίες υποβάλλονται στο σύστημα εκτέλεσης.

Adaptive-OPRO: Βελτιστοποίηση Προτροπών για Ακολουθιακή Λήψη Αποφάσεων

Το **ATLAS** εισάγει το **Adaptive-OPRO**, έναν αλγόριθμο βελτιστοποίησης προτροπών που αναπροσαρμόζει συστηματικά τις οδηγίες του Κεντρικού Πράκτορα Συναλλαγών με βάση την απόδοση των συναλλαγών του. Πρόκειται για την πρώτη επιτυχημένη προσαρμογή της *optimization by prompting* σε ακολουθιακά περιβάλλοντα λήψης αποφάσεων με καθυστερημένη και θορυβώδη ανάδραση.

Επεκτάσεις για Ακολουθιακή Λήψη Αποφάσεων. Η προσέγγισή μας εισάγει δύο καίριες τροποποιήσεις στο αρχικό OPRO, προσαρμοσμένες στις ιδιαιτερότητες των χρηματοοικονομικών αγορών: (i) αντιμετωπίζει τη χρονική υστέρηση ανάμεσα στη λήψη απόφασης και την αποτίμηση της επίδοσης, και (ii) υιοθετεί πρότυπα προτροπών που διαχωρίζουν τον στατικό κορμό οδηγιών ο οποίος τίθεται υπό βελτιστοποίηση από τα δυναμικά δεδομένα εκτέλεσης που μεταβάλλονται στον χρόνο.

Αρχιτεκτονική Βελτιστοποίησης Προτροπών. Το σύστημα βελτιστοποίησης τηρεί *ιστορικό εξέλιξης προτροπών*, στο οποίο κάθε πρότυπο καταγράφεται μαζί με την αντίστοιχη επίδοσή του. Κάθε πέντε διαδοχικά βήματα απόφασης, αποτιμούμε την αποτελεσματικότητα της τρέχουσας προτροπής υπολογίζοντας τη *συνολική Απόδοση Επένδυσης (ROI)* του πράκτορα.

Η διαδικασία βελτιστοποίησης βασίζεται σε ένα **meta-prompt** που καθοδηγεί το LLM-βελτιστοποιητή να αναλύει το συσσωρευμένο ιστορικό επίδοσης και να παράγει, με συστηματικό τρόπο, τέσσερις βασικές εξόδους: (i) ανάλυση επίδοσης, (ii) βελτιστοποιημένο πρότυπο προτροπής, (iii) σύνοψη βασικών βελτιώσεων και (iv) εκτίμηση της αναμενόμενης επίδρασης.

1.4 Πειραματική Διάταξη

Η πειραματική αξιολόγηση στοχεύει να απαντήσει τέσσερα καίρια ερωτήματα: (i) πώς αποδίδει το **Adaptive-OPRO** σε σύγκριση με εναλλακτικές προσεγγίσεις, (ii) ποια είναι η συνεισφορά κάθε εξειδικευμένου πράκτορα, (iii) πώς μεταβάλλεται η επίδοση μεταξύ διαφορετικών αρχιτεκτονικών LLM, και (iv) ποιοι συνδυασμοί επιτυγχάνουν σταθερά, αξιόπιστα αποτελέσματα έναντι διαμορφώσεων με υψηλή διακύμανση.

Σε αντίθεση με προηγούμενες εργασίες που εφαρμόζουν εκτελέσεις μίας μόνο εκτέλεσης, η μεθοδολογία μας υλοποιεί συστηματικά πρωτόκολλα πολλαπλών εκτελέσεων, με αποτέλεσμα να αποκαλύπτεται η διακύμανση των επιδόσεων και να εξάγονται στατιστικά αξιόπιστα συμπεράσματα.

1.4.1 Επιλογή Μοντέλων

Αξιολογούμε πέντε LLMs που αντιπροσωπεύουν διακριτά αρχιτεκτονικά πρότυπα: **Claude Sonnet 4** με και χωρίς μηχανισμούς συλλογισμού, **LLaMA 3.3-70B**, **GPT-o3** και **GPT-o4-mini**. Η σύνθεση αυτή επιτρέπει άμεσες συγκρίσεις τόσο μεταξύ μοντέλων με/χωρίς ρητούς μηχανισμούς συλλογισμού όσο και μεταξύ ιδιόκτητων συστημάτων και συστημάτων ανοικτού κώδικα.

1.4.2 Καθεστώτα Αγοράς και Επιλογή Μετοχών

Τα καθεστώτα αγοράς αντιπροσωπεύουν διακριτά πρότυπα συμπεριφοράς τιμών που δημιουργούν θεμελιωδώς διαφορετικές προκλήσεις για τα συστήματα συναλλαγών. Η αξιολόγησή μας καλύπτει τρία διακριτά καθεστώτα αγοράς χρησιμοποιώντας μετοχές από διαφορετικούς τομείς, με κάθε αξιολόγηση να καλύπτει δύο μήνες (28 Απριλίου - 28 Ιουνίου 2025):

- **NVDA (Τεχνολογία - Ανοδική Αγορά):** Κατά την περίοδο αξιολόγησης η NVDA παρουσίασε συνεχή ανοδική πορεία, ενδεικτική ισχυρής τάσης. Το σενάριο εξετάζει κατά πόσο οι πράκτορες εντοπίζουν και αξιοποιούν αποτελεσματικά την τάση, διατηρώντας κερδοφόρες θέσεις σε ευνοϊκές συνθήκες.
- **XOM (Ενέργεια - Πλευρική αγορά):** Κατά την περίοδο αξιολόγησης, η XOM κινήθηκε κυρίως εντός καθορισμένου εύρους τιμών, καθιστώντας δυσχερείς τις στρατηγικές ορμής και ευνοώντας προσεγγίσεις μέσης επαναφοράς.

- **LLY (Υγεία - Πτωτική Αγορά):** Η LLY κατέγραψε έντονες υποχωρήσεις και αιφνίδιες διακυμάνσεις, διαμορφώνοντας το πιο απαιτητικό σενάριο στην αξιολόγηση. Το περιβάλλον αυτό ελέγχει την ικανότητα διαχείρισης κινδύνου και τη δυναμική προσαρμογή του μεγέθους θέσεων.

1.4.3 Αξιολόγηση Στρατηγικών Προτροπών

Ο πειραματικός σχεδιασμός συγκρίνει συστηματικά πολλαπλές προσεγγίσεις προτροπών με στόχο την απομόνωση της επίδρασης διαφορετικών μηχανισμών προσαρμογής στην απόδοση συναλλαγών. Για πληρότητα, συγκρίνουμε επίσης με πέντε καθιερωμένες ποσοτικές στρατηγικές: *Buy & Hold*, *Moving Average Convergence Divergence (MACD)* [66], *Simple Moving Average (SMA)* [22], *Short-Long Moving Average (SLMA)* [66] και *Bollinger Bands* [11].

Βασικές Πειραματικές Διαμορφώσεις

Στρατηγική Baseline. Χρησιμοποιούμε προσεκτικά σχεδιασμένες στατικές προτροπές που αποτυπώνουν βέλτιστες πρακτικές στο σχεδιασμό προτροπών για εφαρμογές συναλλαγών. Οι προτροπές αυτές αναπτύχθηκαν μέσω επαναληπτικής βελτιστοποίησης από έμπειρους επαγγελματίες της αγοράς και ενσωματώνουν δομημένα πλαίσια λήψης αποφάσεων καθώς και πλήρη ενσωμάτωση συμφραζομένων.

Στρατηγική Adaptive-OPRO. Η συστηματική βελτιστοποίηση προτροπών επικεντρώνεται αποκλειστικά στον *Κεντρικό Πράκτορα Συναλλαγών*, με σκοπό να αξιολογηθεί αν η δυναμική προσαρμογή ενισχύει την ενοποίηση πολλαπλών ροών πληροφορίας σε συνεκτικές και εκτελέσιμες αποφάσεις συναλλαγών.

Στρατηγική Reflection. Για συγκριτική αποτίμηση της επίδοσης του **Adaptive-OPRO** έναντι προηγούμενων προσεγγίσεων, προσαρμόζουμε τον μηχανισμό *reflection* του [38]—έναν από τους πιο πρόσφατους και συναφείς αλγόριθμους σειριακής ανάδρασης σε συστήματα συναλλαγών βασισμένα σε LLMs. Η αρχική μέθοδος, σχεδιασμένη για αγορές κρυπτονομισμάτων, επιστρέφει *βαθμούς εμπιστοσύνης* (από -1 έως 1) αντί για πλήρεις εντολές. Για αξιόπιστη σύγκριση, προσαρμόζουμε και τις δύο προσεγγίσεις στο **StockSim**, όπου απαιτούνται πλήρεις *προδιαγραφές εντολών* (τύπος, ποσότητα, τιμή, προϋποθέσεις).

Χρησιμοποιούμε τον μηχανισμό *reflection* όπως έχει οριστεί, προσαρμόζοντάς τον να παράγει *εβδομαδιαία* ανάδραση πάνω σε μοτίβα συναλλαγών και στρατηγικές προσαρμογές. Η ανάδραση αυτή ενσωματώνεται στον *Κεντρικό Πράκτορα Συναλλαγών* μέσω συμφραζομένων της προτροπής. Σε αντίθεση με το **Adaptive-OPRO**, που τροποποιεί άμεσα το κείμενο της προτροπής, το *reflection* παράγει *αναλυτικά σχόλια* τα οποία ο πράκτορας οφείλει να ερμηνεύσει. Εφαρμόζουμε τον μηχανισμό χωρίς μεθοδολογικές αλλαγές, διατηρώντας πιστότητα στην αρχική πρόταση και επιτρέποντας *άμεση αξιολόγηση* της αποτελεσματικότητάς του υπό ταυτόσημους περιορισμούς εκτέλεσης και ίδιες συνθήκες αξιολόγησης.

1.4.4 Μεθοδολογία και Μετρικές Αξιολόγησης

Πρωτόκολλο Πολλαπλών Εκτελέσεων

Σε αντιδιαστολή με προηγούμενες εργασίες που βασίζονται σε μία μόνο εκτέλεση, εφαρμόζουμε συστηματικό πρωτόκολλο *τριών ανεξάρτητων εκτελέσεων* ανά διαμόρφωση. Κάθε εκτέλεση αρχικοποιείται με ταυτόσημες αρχικές συνθήκες, ώστε η παρατηρούμενη διακύμανση επίδοσης να αποδίδεται στη στοχαστικότητα του LLM και όχι σε παράγοντες του περιβάλλοντος εκτέλεσης.

Τα αποτελέσματα αναφέρονται ως *μέση τιμή ± τυπική απόκλιση* σε τρεις εκτελέσεις ($n = 3$), επιτρέποντας αποτίμηση τόσο της κεντρικής τάσης όσο και της μεταβλητότητας.

Ολοκληρωμένες Μετρικές Απόδοσης

Η αξιολόγηση βασίζεται σε πολλαπλές μετρικές που αποτυπώνουν συμπληρωματικές πτυχές της επίδοσης των στρατηγικών συναλλαγών, όπως:

Απόδοση Επένδυσης (ROI): Υπολογίζεται ως η ποσοστιαία μεταβολή της συνολικής αξίας χαρτοφυλαχίου:

$$ROI = \frac{\text{τελική αξία} - \text{αρχική αξία}}{\text{αρχική αξία}} \times 100$$

Model	Prompting	ROI (%) ↑	SR ↑	DD (%) ↓	Win Rate (%) ↑	Num Trades
Non-LLM-Based Strategies						
Buy & Hold	N/A	-8.59	-0.071	20.45	0.00	1
MACD	N/A	6.50	0.131	6.86	0.00	1
SMA	N/A	6.91	0.177	3.56	50.00	4
SLMA	N/A	-1.87	-0.078	6.89	0.00	1
Bollinger Bands	N/A	0.00	0.000	0.00	0.00	0
LLM-Based Strategies - ATLAS						
LLaMA 3.3-70B	Baseline	-9.19 $\pm$ 1.54	-0.091 $\pm$ 0.021	16.90 $\pm$ 0.82	30.28 $\pm$ 11.87	22.67 $\pm$ 8.39
	Reflection	-8.44 $\pm$ 1.58	-0.087 $\pm$ 0.025	16.36 $\pm$ 0.31	44.69 $\pm$ 13.25	27.67 $\pm$ 1.15
	Adaptive-OPRO	-6.16$\pm$ 2.08	-0.066$\pm$ 0.004	14.05$\pm$ 3.33	54.36$\pm$ 12.44	28.33 $\pm$ 3.21
Claude Sonnet 4	Baseline	-7.26 $\pm$ 2.99	-0.066 $\pm$ 0.030	17.59 $\pm$ 1.55	31.19 $\pm$ 7.84	13.00 $\pm$ 4.36
	Reflection	-5.69 $\pm$ 1.82	-0.058 $\pm$ 0.013	15.12 $\pm$ 3.26	46.67$\pm$ 5.77	12.67 $\pm$ 2.08
	Adaptive-OPRO	0.35$\pm$ 1.78	0.008$\pm$ 0.018	14.76$\pm$ 2.87	43.45 $\pm$ 6.27	15.00 $\pm$ 2.00
Claude Sonnet 4 w/ Thinking	Baseline	-4.46 $\pm$ 4.76	-0.043 $\pm$ 0.048	14.32 $\pm$ 4.12	11.11 $\pm$ 19.24	14.00 $\pm$ 2.65
	Reflection	-8.60 $\pm$ 0.59	-0.078 $\pm$ 0.004	19.45 $\pm$ 1.65	14.29 $\pm$ 24.75	11.67 $\pm$ 2.08
	Adaptive-OPRO	-0.73$\pm$ 3.82	-0.004$\pm$ 0.038	12.94$\pm$ 2.32	43.89$\pm$ 21.11	17.00 $\pm$ 5.00
GPT-o4-mini	Baseline	-1.30 $\pm$ 1.71	-0.017 $\pm$ 0.017	9.68$\pm$ 3.12	29.17 $\pm$ 11.02	15.33 $\pm$ 3.06
	Reflection	-2.52 $\pm$ 4.03	-0.039 $\pm$ 0.045	9.82 $\pm$ 3.43	51.28 $\pm$ 5.06	20.33 $\pm$ 3.06
	Adaptive-OPRO	9.06$\pm$ 0.73	0.094$\pm$ 0.008	11.48 $\pm$ 0.00	65.28$\pm$ 16.84	17.33 $\pm$ 5.86
GPT-o3	Baseline	-6.11 $\pm$ 3.42	-0.080 $\pm$ 0.029	11.58 $\pm$ 3.09	42.59 $\pm$ 8.49	18.67 $\pm$ 3.21
	Reflection	-4.60 $\pm$ 3.40	-0.053 $\pm$ 0.044	12.11 $\pm$ 1.27	46.03 $\pm$ 16.88	18.33 $\pm$ 2.52
	Adaptive-OPRO	9.02$\pm$ 3.28	0.146$\pm$ 0.048	5.33$\pm$ 0.14	72.81$\pm$ 17.27	19.67 $\pm$ 4.16

Table 1.1: Σύγκριση επίδοσης μεταξύ προσεγγίσεων χωρίς LLM και με LLM, χρησιμοποιώντας το ATLAS, σε ευμετάβλητες πτωτικές συνθήκες αγοράς (LLY, Τομέας Υγείας). Οι **bold** τιμές υποδηλώνουν την καλύτερη επίδοση ανά μοντέλο.

Δείκτης Sharpe (SR): Απόδοση προσαρμοσμένη στον κίνδυνο, $SR = \frac{\mu - r_f}{\sigma}$, όπου μ η μέση ημερήσια απόδοση, r_f το επιτόκιο χωρίς κίνδυνο και σ η τυπική απόκλιση.

Μέγιστη Υποχώρηση (DD): Η μεγαλύτερη πτώση από κορυφή σε κοιλάδα στην αξία του χαρτοφυλακίου, που αποτυπώνει την έκθεση σε καθοδικούς κινδύνους.

Ποσοστό Επιτυχίας (Win Rate): Ποσοστό κερδοφόρων συναλλαγών· ένδειξη συνέπειας στη λήψη αποφάσεων.

Αριθμός Συναλλαγών: Συνολικός αριθμός εκτελεσμένων συναλλαγών, ενδεικτικός του βαθμού δραστηριότητας και της προσέγγισης της στρατηγικής.

1.4.5 Σχεδιασμός Μελέτης Αφαίρεσης

Για να ποσοτικοποιήσουμε τη συνεισφορά κάθε εξειδικευμένου πράκτορα στο πλαίσιο **ATLAS**, διεξάγουμε συστηματικές μελέτες αφαίρεσης (*ablation*).

Χωρίς Αναλυτή Αγοράς: Απενεργοποιούμε όλες τις δυνατότητες τεχνικής ανάλυσης, ώστε να εκτιμηθεί αν η αποτύπωση της δομής τιμών σε πολλαπλές χρονικές κλίμακες βελτιώνει ουσιαστικά την ποιότητα των αποφάσεων.

Χωρίς Αναλυτή Ειδήσεων: Αφαιρούμε την επεξεργασία αδόμητου κειμένου, για να αποτιμηθεί η αξία της εξαγωγής επενδυτικού αισθήματος και της αναγνώρισης καταλυτών που προκύπτουν από γεγονότα.

Χωρίς Αναλυτές Αγοράς και Ειδήσεων: Το σύστημα υποβιβάζεται σε διαμόρφωση ενός μόνο πράκτορα με μοναδικές εισόδους τα θεμελιώδη δεδομένα και την κατάσταση του χαρτοφυλακίου.

Εξαιρούμε τον Θεμελιώδη Αναλυτή από την ανάλυση αφαίρεσης, επειδή η χαμηλή συχνότητα ενεργοποίησής του (ευθυγραμμισμένη με τους κύκλους ανακοινώσεων) καθιστά τη βραχυπρόθεσμη αποτίμηση της επίδρασής του στατιστικά ασταθή.

Model	Prompting	XOM			NVDA		
		ROI (%) $\uparrow$	SR $\uparrow$	DD (%) $\downarrow$	ROI (%) $\uparrow$	SR $\uparrow$	DD (%) $\downarrow$
Non-LLM-Based Strategies							
Buy & Hold	N/A	1.14	0.013	6.97	41.30	0.409	3.16
MACD	N/A	-0.26	-0.019	5.90	-0.62	-0.343	0.62
SMA	N/A	-1.02	-0.019	5.75	14.02	0.242	2.93
SLMA	N/A	-2.08	-0.066	5.53	36.77	0.386	3.12
Bollinger Bands	N/A	0.00	0.000	0.00	0.00	0.000	0.00
LLM Based-Strategies - ATLAS							
LLaMA 3.3-70B	Baseline	-0.42$\pm$ 2.06	-0.024$\pm$ 0.051	5.56 $\pm$ 1.08	37.86 $\pm$ 12.31	0.388 $\pm$ 0.096	3.46$\pm$ 0.63
	Reflection	-2.61 $\pm$ 0.77	-0.083 $\pm$ 0.014	6.38 $\pm$ 0.72	40.40 $\pm$ 1.43	0.422$\pm$ 0.023	2.96 $\pm$ 0.34
	Adaptive-OPRO	-1.10 $\pm$ 0.44	-0.045 $\pm$ 0.012	5.15$\pm$ 0.71	42.07$\pm$ 1.85	0.418 $\pm$ 0.016	3.15 $\pm$ 0.02
Claude Sonnet 4	Baseline	-4.49 $\pm$ 4.22	-0.134 $\pm$ 0.114	7.71$\pm$ 1.06	13.43 $\pm$ 8.62	0.180 $\pm$ 0.121	5.52 $\pm$ 3.96
	Reflection	-3.78$\pm$ 4.23	-0.115$\pm$ 0.105	10.54 $\pm$ 1.58	5.21 $\pm$ 1.10	0.089 $\pm$ 0.026	5.11 $\pm$ 1.86
	Adaptive-OPRO	-5.07 $\pm$ 4.53	-0.165 $\pm$ 0.143	9.23 $\pm$ 2.71	25.85$\pm$ 10.61	0.290$\pm$ 0.087	3.75$\pm$ 0.59
Claude Sonnet 4 w/ Thinking	Baseline	-0.99$\pm$ 0.80	-0.039$\pm$ 0.020	7.75 $\pm$ 1.00	12.52 $\pm$ 2.47	0.175 $\pm$ 0.030	5.03 $\pm$ 1.53
	Reflection	-1.49 $\pm$ 3.76	-0.069 $\pm$ 0.123	7.27 $\pm$ 2.26	11.12 $\pm$ 4.86	0.186 $\pm$ 0.083	3.42$\pm$ 2.23
	Adaptive-OPRO	-1.01 $\pm$ 0.90	-0.046 $\pm$ 0.020	5.16$\pm$ 0.52	16.36$\pm$ 7.87	0.217$\pm$ 0.105	5.18 $\pm$ 2.52
GPT-o4-mini	Baseline	1.29 $\pm$ 1.38	0.021 $\pm$ 0.044	3.23$\pm$ 0.48	7.00 $\pm$ 3.46	0.125 $\pm$ 0.054	2.74$\pm$ 0.79
	Reflection	-1.48 $\pm$ 0.54	-0.087 $\pm$ 0.018	4.64 $\pm$ 0.75	9.80 $\pm$ 3.21	0.189 $\pm$ 0.067	2.45 $\pm$ 1.00
	Adaptive-OPRO	3.88$\pm$ 2.21	0.089$\pm$ 0.067	3.28 $\pm$ 0.95	10.47$\pm$ 3.84	0.193$\pm$ 0.046	3.42 $\pm$ 0.90
GPT-o3	Baseline	-0.60 $\pm$ 1.71	-0.034 $\pm$ 0.050	5.93 $\pm$ 1.33	22.70 $\pm$ 0.92	0.269 $\pm$ 0.029	6.82 $\pm$ 3.03
	Reflection	-1.55 $\pm$ 2.09	-0.084 $\pm$ 0.075	5.02 $\pm$ 0.72	21.98 $\pm$ 4.54	0.325 $\pm$ 0.040	3.14 $\pm$ 0.99
	Adaptive-OPRO	3.62$\pm$ 0.90	0.096$\pm$ 0.027	3.46$\pm$ 0.48	25.06$\pm$ 4.28	0.392$\pm$ 0.019	2.31$\pm$ 0.80

Table 1.2: Συνδυαστικός πίνακας επιδόσεων σε δύο αγορές: **XOM** (πλευρική) και **NVDA** (ανοδική). Περιλαμβάνει ROI, SR και DD. Οι **bold** τιμές υποδηλώνουν την καλύτερη επίδοση ανά μοντέλο.

1.5 Αποτελέσματα

Οι Πίνακες 1.1, 1.2 παρουσιάζουν συγκριτική αξιολόγηση καθιερωμένων ποσοτικών στρατηγικών έναντι του **ATLAS** σε ποικίλες διαμορφώσεις μοντέλων και διαφορετικά καθεστώτα αγοράς. Τα αποτελέσματα δείχνουν ότι το **ATLAS** επιτυγχάνει σταθερά υψηλή επίδοση σε όλες τις εξεταζόμενες συνθήκες. Πρόκειται για την πρώτη περίπτωση που ένα ενιαίο πλαίσιο επιδεικνύει τόσο συστηματική ανθεκτικότητα σε ετερόκλητα σενάρια αγοράς, υπερτερώντας ξεκάθαρα έναντι διαδεδομένων και παραδοσιακά δημοφιλών μεθόδων. Στρατηγικές όπως το *Buy-and-Hold* αποδίδουν καλά σε ανοδικά καθεστώτα, αλλά αποτυγχάνουν να γενικεύσουν, παράγοντας αισθητά ασθενέστερα αποτελέσματα σε πλευρικές και πτωτικές αγορές, όπου κυριαρχούν αστάθεια, χαμηλή προβλεψιμότητα και περιορισμένα σήματα πληροφορίας. Αντιθέτως, το **ATLAS**, ιδίως σε συνδυασμό με **GPT-o3** ή **GPT-o4-mini**, παρέχει σταθερά θετικές αποδόσεις ακόμη και σε αντίζοα περιβάλλοντα, συμπεριλαμβανομένων πτωτικών καθεστώτων όπου η κυρίαρχη τάση είναι καθοδική και η επίτευξη κερδοφορίας ιδιαίτερα δύσκολη. Η ικανότητα να αποδίδει αξιόπιστα ακόμα όταν οι περισσότερες στρατηγικές δυσκολεύονται αναδεικνύει την ικανότητα του πλαισίου να πλοηγείται στην αβεβαιότητα και να λαμβάνει αποτελεσματικές στρατηγικές αποφάσεις, ακόμη και σε πτωτικά ή ιδιαίτερα ευμετάβλητα περιβάλλοντα. Τα ευρήματα αυτά αναδεικνύουν την ανθεκτικότητα του **ATLAS** και τη δυναμική του ως αξιόπιστου συστήματος λήψης αποφάσεων σε όλο το φάσμα καθεστώτων της αγοράς.

1.5.1 Βελτιστοποίηση σε Ακολουθιακή Λήψη Αποφάσεων

Το **Adaptive-OPRO** υπερέχει σημαντικά τόσο έναντι των στατικών baseline προτροπών όσο και έναντι προσεγγίσεων τύπου *reflection* στη συντριπτική πλειονότητα των μοντέλων και συνθηκών αγοράς (Πίνακες 1.1, 1.2). Η σταδιακή αναβάθμιση της ποιότητας των αποφάσεων κατά τη βελτιστοποίηση αποτυπώνεται σε σαφώς ανώτερη επίδοση συναλλαγών, όπως καταδεικνύουν οι σχετικές μετρικές.

Οι βελτιώσεις στην απόδοση επένδυσης (ROI) υποδηλώνουν επιτυχή αξιοποίηση της ανάδρασης της αγοράς. Στο πτωτικό/ευμετάβλητο καθεστώς (Πίνακας 1.1), μοντέλα όπως τα **GPT-o3** και **GPT-o4-mini** μεταβαίνουν από αρνητικό ROI με *baseline* προτροπές σε σαφώς θετικές αποδόσεις υπό **Adaptive-OPRO**. Ακόμη πιο ενδεικτικές είναι οι **μετρικές προσαρμοσμένες στον κίνδυνο**: οι υψηλότερες τιμές *Sharpe* καταδεικνύουν ότι τα κέρδη απορρέουν από ουσιαστική στρατηγική αναβάθμιση και όχι από ανάληψη υπερβολικού ρίσκου.

Stock	Configuration	ROI (%)↑	Sharpe Ratio ↑	Max DD (%) ↓	Win Rate (%) ↑	Num Trades
LLY (Bearish/ Volatile Regime)	No News	4.07± 0.72	0.056± 0.016	7.84± 3.15	53.51± 6.67	25.33± 4.51
	No Market Data	-5.75± 0.76	-0.094± 0.017	11.32± 2.63	37.52± 4.87	18.33± 3.06
	No News + No Market	-6.86± 1.68	-0.078± 0.036	14.54± 3.30	43.94± 6.94	22.33± 1.15
	ATLAS	9.06± 0.73	0.094± 0.008	11.48± 0.00	65.28± 16.84	17.33± 5.86
XOM (Sideways Regime)	No News	-8.20± 1.64	-0.264± 0.069	9.09± 2.99	22.82± 13.65	35.00± 12.29
	No Market Data	0.01± 0.92	-0.011± 0.021	6.56± 1.58	46.55± 23.15	13.33± 3.06
	No News + No Market	-4.60± 0.70	-0.136± 0.026	7.01± 2.29	35.26± 13.09	21.00± 4.58
	ATLAS	3.88± 2.21	0.089± 0.067	3.28± 0.95	47.95± 7.15	25.33± 5.03
NVDA (Bullish Regime)	No News	6.62± 0.25	0.090± 0.008	6.67± 0.36	41.96± 5.21	28.33± 4.62
	No Market Data	11.78± 1.76	0.216± 0.024	3.70± 0.86	70.24± 14.03	20.00± 5.57
	No News + No Market	7.34± 2.79	0.110± 0.012	5.76± 2.01	63.84± 9.39	20.67± 1.53
	ATLAS	10.47± 3.84	0.193± 0.046	3.42± 0.90	62.70± 11.25	20.33± 2.89

Table 1.3: Αποτελέσματα μελέτης αφαίρεσης που αναδεικνύουν τη συμβολή κάθε μεμονωμένου πράκτορα με χρήση GPT-o4-mini σε τρία καθεστώτα αγοράς.

Η συμπεριφορά του ποσοστού επιτυχίας αποκαλύπτει την επίδραση της βελτιστοποίησης στην ποιότητα αποφάσεων: τα μοντέλα υπό **Adaptive-OPRO** επιτυγχάνουν γενικά υψηλότερα win rates μαζί με βελτιωμένες αποδόσεις, υποδηλώνοντας πιο συνεπή λήψη αποφάσεων και όχι περιστασιακά μεγάλα κέρδη που συγκαλύπτουν συχνές ζημιές.

Αξιοσημείωτα, αναδύεται σταθερά το «παράδοξο του reflection»: οι στρατηγικές *reflection* όχι μόνο δεν προσεγγίζουν την επίδοση του **Adaptive-OPRO**, αλλά συχνά υστερούν και έναντι των baseline προτροπών-αμφισβητώντας την επικρατούσα άποψη ότι τα πρόσθετα βήματα συλλογισμού βελτιώνουν πάντα καλά ρυθμισμένα συστήματα σε δυναμικά περιβάλλοντα.

1.5.2 Ανάλυση Συνεισφοράς Πρακτόρων

Ο Πίνακας 1.3 επιβεβαιώνει τη διακριτή συνεισφορά κάθε εξειδικευμένου πράκτορα, παρουσιάζοντας τις πτώσεις επίδοσης όταν αφαιρείται ο καθένας.

Ο **Αναλυτής Αγοράς** αποτελεί θεμελιώδες συστατικό σε όλα τα καθεστώτα. Η αφαίρεσή του οδηγεί συστηματικά στη μεγαλύτερη πτώση επίδοσης, ιδίως σε πτωτικές συνθήκες όπου το τεχνικό πλαίσιο είναι κρίσιμο για τη λήψη αποφάσεων. Σε πλευρικές αγορές, η απουσία τεχνικής ανάλυσης δεν μειώνει μόνο τις αποδόσεις αλλά και τη συχνότητα συναλλαγών, υποδηλώνοντας απώλεια «εμπιστοσύνης» για δράση χωρίς στιβαρό τεχνικό υπόβαθρο. Παρά ταύτα, σε ανοδικά καθεστώτα, παρατηρείται μικρή βελτίωση του *ROI* όταν παραλείπονται τα δεδομένα αγοράς, υποδηλώνοντας ότι υπό ισχυρές τάσεις η ειδησεογραφία και η ευρύτερη «συναίνεση» της αγοράς ενδέχεται να προσφέρουν καθαρότερα σήματα εισόδου.

Ο **Αναλυτής Ειδήσεων** προσφέρει στρατηγική αξία που εξαρτάται από το καθεστώς αγοράς. Σε ανοδικές τάσεις, η απενεργοποίησή του οδηγεί σε χειρότερες επιδόσεις, καθώς οι πράκτορες γίνονται πιο επιφυλακτικοί και χάνουν ευκαιρίες αξιοποίησης της θετικής ορμής. Σε πλευρικές συνθήκες, ο ρόλος του γίνεται καταλυτικός: η αφαίρεσή του οδηγεί σε έντονη υποβάθμιση, υποδηλώνοντας ότι η αποτίμηση επενδυτικού αισθήματος είναι καθοριστική όταν τα τεχνικά σήματα είναι αμφίσημα.

Ο **συνδυασμός Αναλυτή Αγοράς και Αναλυτή Ειδήσεων** αποκαλύπτει τη συμπληρωματικότητά τους: σε όλα τα καθεστώτα, η ταυτόχρονη αφαίρεση επιφέρει σημαντική επιδείνωση, δείχνοντας πως ειδήσεις και τεχνικά σήματα παρέχουν συμπληρωματική, μη πλεονάζουσα πληροφορία. Σε πτωτικές αγορές, η επίδοση υποχωρεί δραματικά, υπογραμμίζοντας την ανάγκη για συνδυασμό τεχνικού πλαισίου και αποτίμησης επενδυτικού αισθήματος σε συνθήκες υψηλής μεταβλητότητας. Σε πλευρικά καθεστώτα, η απουσία και των δύο γεννά μια ασταθή, μη κερδοφόρα συμπεριφορά. Ακόμη και σε ανοδικά περιβάλλοντα, όπου τα τεχνικά δεδομένα μπορεί να είναι λιγότερο κρίσιμα, η συνδυασμένη αφαίρεση πλήττει σαφώς την επίδοση. Συνολικά, οι δύο συνιστώσες συμβάλλουν διαφορετικά ανά καθεστώς, ενώ η κοινή τους αφαίρεση έχει επιδράσεις που δεν είναι απλώς αθροιστικές.

1.5.3 Συναλλακτική Συμπεριφορά ανά LLM

Η ανάλυσή μας αποκαλύπτει σαφή συσχέτιση μεταξύ των γενικών ικανοτήτων των μοντέλων και της επίδοσής τους στις συναλλαγές: τα πιο ικανά μοντέλα αποδίδουν σημαντικά καλύτερα από τα λιγότερο ικανά (Σχ. 1.5.1).

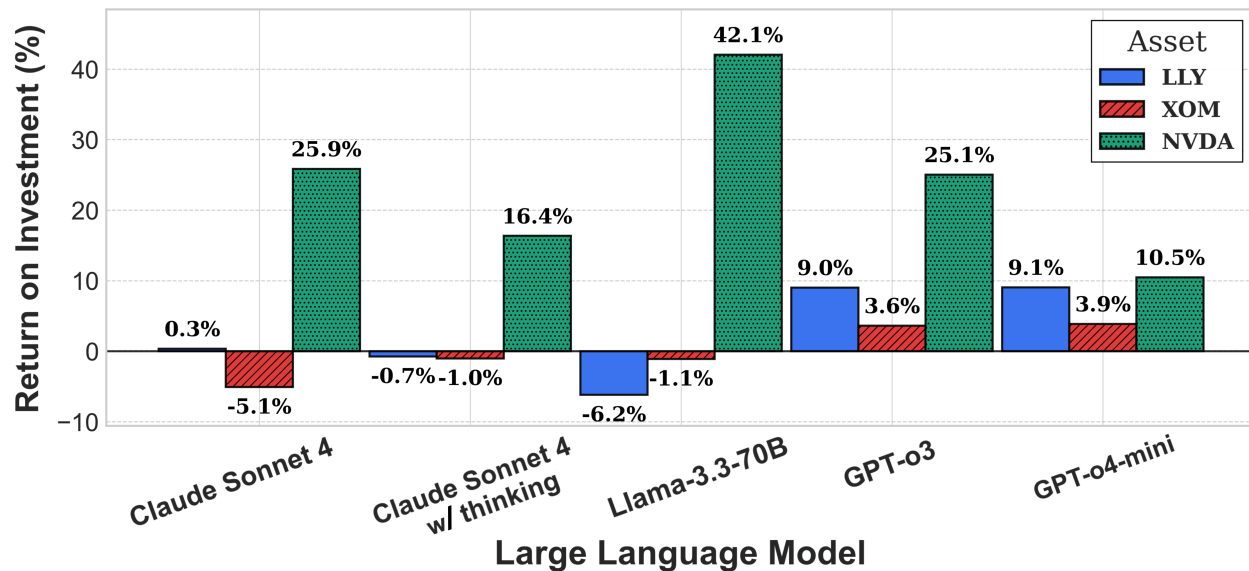


Figure 1.5.1: ROI σε τρεις μετοχές με χρήση Adaptive-OPRO.

Το **GPT-o3** παρουσιάζει την πιο ώριμη κατανόηση της αγοράς, συνθέτοντας συνεκτικές στρατηγικές από τις αναλύσεις όλων των εξειδικευμένων πρακτόρων. Παρά την κατά τόπους υπέρμετρα συντηρητική του στάση, που περιορίζει τα κέρδη σε ευνοϊκές τάσεις, η επιλογή αυτή αποδίδει *σταθερή, διακαθεστωτική επίδοση*. Οι ικανότητές του στη βελτιστοποίηση προτροπών είναι υποδειγματικές: προσδευτική μάθηση, στρατηγική αναπροσαρμογή μετά από αστοχίες και σαφής ευθυγράμμιση με τους στόχους βελτιστοποίησης.

Το **GPT-o4-mini** επιδεικνύει ικανή, αλλά βραχυπρόθεσμη στόχευση, προτάσσοντας τον έλεγχο ρίσκου έναντι της στρατηγικής τοποθέτησης. Προκρίνει αυστηρά όρια διακοπής ζημίας και γρήγορη κατοχύρωση κερδών, προσέγγιση που αποδίδει σε συνθήκες υψηλής μεταβλητότητας, αλλά υστερεί σε παρατεταμένες τάσεις. Παρά την περιστασιακή υπερβολική δραστηριότητα (*overtrading*), το επίπεδο ωριμότητας στη βελτιστοποίηση προτροπών παραμένει συγκρίσιμο με εκείνο του **GPT-o3**.

Στον αντίποδα, το **LLaMA 3.3-70B** βασίζεται σε ιδιαίτερα απλοϊκές τακτικές, με ανεπαρκή διαχείριση κινδύνου και χωρίς συνεκτική στοχοθεσία. Καθυστερεί να προσαρμοστεί στις μεταβολές της αγοράς και, σε συνθήκες έντονης μεταβλητότητας, μεταπηδά απότομα μεταξύ στρατηγικών. Παραδόξως, αυτή η απλότητα αποδεικνύεται πλεονέκτημα σε καθαρά ανοδικά καθεστώτα, όπου στην πράξη υιοθετεί στρατηγική τύπου *buy-and-hold* και καταγράφει εξαιρετικά υψηλές αποδόσεις, αναδεικνύοντας ότι ο περίπλοκος συλλογισμός μπορεί να υποβαθμίζει την απόδοση σε απλούστερα περιβάλλοντα.

Μοντέλα Claude: Διακριτά Πρότυπα Αστοχίας

Ανεξαρτήτως διαμόρφωσης (με ή χωρίς *thinking*) τα μοντέλα **Claude** εμφανίζουν συστηματικές αδυναμίες, αλλά μέσω διαφορετικών μοτίβων αστοχίας που παραπέμπουν σε δομικούς περιορισμούς. Με *thinking*: εκδηλώνεται τάση για υπερανάλυση ειδησεογραφίας και θεμελιωδών, που οδηγεί παρερμηνεία των τεχνικών σημάτων. Παρά το φαινομενικά καλοδομημένο σκεπτικό, παραμένουν κρίσιμα κενά που οδηγούν σε εσφαλμένες αποφάσεις.

Χωρίς *thinking*, τα ζητήματα αξιοπιστίας οξύνονται: υπεραντιδραστικές τοποθετήσεις, επίμονες «παραισθήσεις» (*hallucinations*) στις μετρικές επίδοσης και θεμελιώδεις παρανοήσεις της μικροδομής της αγοράς. Το αποτέλεσμα είναι ακραία διακύμανση και ασυνεπής, αλλοπρόσαλλη συμπεριφορά, που καθιστούν το μοντέλο ακατάλληλο για συνεπή και αξιόπιστη αξιολόγηση.

Συνολικά, οι αντιθετικοί αυτοί τρόποι αστοχίας υποδηλώνουν ότι η υποκείμενη αρχιτεκτονική του **Claude** είναι ανεπαρκώς ευθυγραμμισμένη με τις απαιτήσεις της χρηματοοικονομικής λήψης αποφάσεων, ανεξαρτήτως ενισχύσεων συλλογισμού.

1.5.4 Ικανότητες Βελτιστοποίησης LLM

Από σκοπία *ερμηνευσιμότητας*, ένα βασικό πλεονέκτημα του **Adaptive-OPRO** είναι ότι η διαδικασία βελτιστοποίησης παράγει μια *τελική οδηγία (prompt)* για τα μοντέλα, η οποία μπορεί να αξιολογηθεί τόσο ως προς την ορθότητά της όσο και ως προς το κατά πόσο τα μοντέλα την ακολουθούν. Αυτή η διττή οπτική μας επιτρέπει να αποτιμήσουμε όχι μόνο την ευθυγράμμιση της βελτιστοποιημένης προτροπής με τον επιδιωκόμενο στόχο, αλλά και την ικανότητα του μοντέλου να την ερμηνεύει και να ενεργεί βάσει αυτής. Έτσι, μπορούμε να αξιολογήσουμε αν η βελτιστοποίηση κινείται προς τη σωστή κατεύθυνση.

Η ανάλυση αναδεικνύει ότι οι ικανότητες βελτιστοποίησης προτροπών διαφέρουν ανά LLM: τα **GPT** μοντέλα παράγουν συνεπείς, ερμηνεύσιμες αναθεωρήσεις προτροπών, ενώ το **LLaMA** συχνά «φαντάζεται» αλλαγές, καταγράφοντας τροποποιήσεις που στην πραγματικότητα *δεν* υλοποιούνται. Το **Claude** τείνει σε υπερβολικά *αυστηρές, διαδικαστικές* οδηγίες που περιορίζουν την ευελιξία.

Αντίθετα, το *reflection*, παρότι θεωρητικά υποσχόμενο, στην πράξη εισάγει θόρυβο. Ακόμη και το **GPT-o3** τείνει να εγλωβίζεται σε υπερανάλυσεις, ενώ λιγότερο ικανά μοντέλα είτε παράγουν *ασαφή* ανατροφοδότηση (LLaMA) είτε παρερμηνεύουν τα σήματα της αγοράς με υπέρμετρη βεβαιότητα.

1.6 Συμπεράσματα

Η διπλωματική αυτή εργασία παρουσιάζει ένα ολοκληρωμένο πλαίσιο εφαρμογής Μεγάλων Γλωσσικών Μοντέλων (LLMs) στις χρηματοοικονομικές αγορές μέσω δύο συμπληρωματικών συνεισφορών: το **StockSim**, μια προηγμένη πλατφόρμα προσομοίωσης για συστηματική αξιολόγηση LLM σε ρεαλιστικά περιβάλλοντα συναλλαγών, και το **ATLAS**, ένα προσαρμοστικό πολυπρακτορικό σύστημα συναλλαγών που επιδεικνύει αποτελεσματικό συντονισμό και βελτιστοποίηση LLM υπό συνθήκες αβεβαιότητας. Παρότι σχεδιάστηκαν με επίκεντρο τη χρηματοοικονομική λήψη αποφάσεων, οι συνεισφορές αυτές θεμελιώνουν ευρύτερες μεθοδολογικές αρχές για την εφαρμογή των Μεγάλων Γλωσσικών Μοντέλων σε περιβάλλοντα υψηλού διακυβεύματος με καθυστερημένη ανάδραση, θορυβώδεις ανταμοιβές και απαίτηση για σταθερά αξιόπιστη επίδοση σε δυναμικές συνθήκες.

1.6.1 Συζήτηση

Η εργασία μας αντιμετωπίζει κρίσιμα κενά που έχουν αναχαιτίσει την αξιοποίηση των LLM σε χρηματοοικονομικά πεδία, όπου οι παραδοσιακές μέθοδοι αξιολόγησης αδυνατούν να αποτυπώσουν την πολυπλοκότητα των πραγματικών συνθηκών αγοράς και τη στοχαστική φύση τόσο των αγορών όσο και των εξόδων των LLM. Η ανάπτυξη του **StockSim** συνιστά σημαντική πρόοδο στην ερευνητική υποδομή της χρηματοοικονομικής τεχνητής νοημοσύνης, προσφέροντας την πρώτη ολοκληρωμένη πλατφόρμα που συνδυάζει προσομοίωση αγοράς επιπέδου παραγωγής με συστηματικές δυνατότητες αξιολόγησης LLM ειδικά για εφαρμογές συναλλαγών. Υποστηρίζοντας τόσο *εκτέλεση σε επίπεδο εντολών* όσο και *εκτέλεση σε επίπεδο κερών (OHLCV)*, το StockSim επιτρέπει τη μελέτη της συμπεριφοράς των LLM σε διαφορετικά επίπεδα πολυπλοκότητας αγοράς, διατηρώντας ταυτόχρονα αυστηρά πρωτόκολλα αξιολόγησης που λαμβάνουν υπόψη την εγγενή αβεβαιότητα των αγορών και των μοντέλων. Η αρχιτεκτονική διπλής λειτουργίας, η ολοκληρωμένη ενσωμάτωση δεδομένων και η υποστήριξη πολλαπλών πρακτόρων δημιουργούν πρωτόγνωρες δυνατότητες για τη μελέτη ικανοτήτων των LLM σε εφαρμογές που απαιτείται χρηματοοικονομικός συλλογισμός, πολυτροπική επεξεργασία πληροφορίας και συντονισμένη λήψη αποφάσεων υπό χρονικές και στοχαστικές προκλήσεις πραγματικών συναλλαγών.

Η εισαγωγή του **ATLAS** καταδεικνύει πώς οι προηγμένες ικανότητες των LLM μπορούν να αξιοποιηθούν αποτελεσματικά μέσω κατάλληλου αρχιτεκτονικού σχεδιασμού και προσαρμοστικής βελτιστοποίησης. Η πολυπρακτορική αρχιτεκτονική μας αποσυνθέτει την πολύπλοκη χρηματοοικονομική ανάλυση σε εξειδικευμένες συνιστώσες, ενώ διατηρεί συνεκτική εκτέλεση στρατηγικής μέσω ενός κεντρικού πράκτορα. Αυτή η αποσύνθεση αποδεικνύεται κρίσιμη για τη διαχείριση της πολυπλοκότητας πραγματικών περιβαλλόντων λήψης αποφάσεων, όπου πολλαπλές πηγές πληροφορίας και οπτικές πρέπει να συντεθούν σε εκτελέσιμες αποφάσεις. Η *σταθερά ανώτερη επίδοση* του πλαισίου σε ετερόκλητες συνθήκες αγοράς, συμπεριλαμβανομένων απαιτητικών πτωτικών και ευμετάβλητων καθεστώτων όπου οι παραδοσιακές στρατηγικές αποτυγχάνουν, επικυρώνει την ευρωστία και την πρακτική του αξία.

Κομβικό στοιχείο της αποτελεσματικότητας του ATLAS είναι το **Adaptive-OPRO**, η νέα μας επέκταση της βελτιστοποίησης προτροπών (prompt optimization) σε περιβάλλοντα ακολουθιακής λήψης αποφάσεων με

καθυστερημένες και θορυβώδεις ανταμοιβές. Η μεθοδολογία αυτή συνιστά την πρώτη επιτυχή προσαρμογή του *Optimization by PROMPTING* σε πεδία με χρονικές εξαρτήσεις και στοχαστική ανάδραση. Οι συστηματικές βελτιώσεις που επιτυγχάνονται με το Adaptive-OPRO σε όλες τις δοκιμασμένες αρχιτεκτονικές και συνθήκες αγοράς αποδεικνύουν ότι τα LLM μπορούν να *μαθαίνουν και να προσαρμόζονται* σε σύνθετα περιβάλλοντα όταν πλασιώνονται από κατάλληλους μηχανισμούς βελτιστοποίησης. Η επιτυχία αυτή στηρίζει την αρχή ότι η άμεση βελτιστοποίηση οδηγιών βάσει *αποτελεσμάτων* προσφέρει πιο αξιόπιστες βελτιώσεις από προσεγγίσεις μετα-γνωστικής ενίσχυσης.

Η εκτενής αξιολόγησή μας αποκαλύπτει ευρήματα για τη συμπεριφορά των LLM που υπερβαίνουν τον χρηματοοικονομικό χώρο. Η σαφής συσχέτιση μεταξύ γενικών ικανοτήτων ενός μοντέλου και ειδικής επίδοσης στο πεδίο υποδηλώνει ότι η πρόοδος στα LLM μεταφράζεται αξιόπιστα σε εξειδικευμένες εφαρμογές. Ωστόσο, αναδεικνύονται και ουσιώδεις αρχιτεκτονικές διαφορές: ενώ τα μοντέλα *GPT* επιδεικνύουν συνεπή επίδοση και αποτελεσματική δυνατότητα βελτιστοποίησης, άλλες αρχιτεκτονικές εμφανίζουν θεμελιώδεις περιορισμούς που επιμένουν ανεξαρτήτως συλλογιστικής ενίσχυσης ή μεθόδων βελτιστοποίησης. Οι διαφορές αυτές έχουν κρίσιμες συνέπειες για την επιλογή μοντέλων σε εφαρμογές υψηλού διακυβέυματος.

Η *συστηματική υστέρηση* των προσεγγίσεων τύπου *reflection* προσφέρει ουσιώδεις ενδείξεις για τον σχεδιασμό συστημάτων βασισμένων σε LLM. Τα αποτελέσματά μας δείχνουν ότι τα πρόσθετα βήματα συλλογισμού μπορεί να *υποβαθμίσουν* την επίδοση καλά ρυθμισμένων συστημάτων, αμφισβητώντας την κοινή παραδοχή περί καθολικού οφέλους από αυξημένη αναλυτική πολυπλοκότητα. Αυτό υποδηλώνει ότι η ερευνητική/υπολογιστική προσπάθεια είναι προτιμότερο να επενδύεται σε *συστηματική βελτίωση προτροπών* παρά σε περίπλοκους μηχανισμούς αξιολόγησης, ιδίως σε δυναμικά περιβάλλοντα όπου η συνέπεια στην ποιότητα των αποφάσεων είναι καθοριστική.

Ένα κρίσιμο εύρημα είναι ότι *πρωτόκολλα πολλαπλών εκτελέσεων* αποκαλύπτουν ζητήματα αξιοπιστίας σε συστήματα LLM που συχνά παραβλέπονται. Η ακραία διακύμανση επίδοσης που παρατηρούμε σε ορισμένες διαμορφώσεις, υπερβαίνοντας το 50% της μέσης επίδοσης σε κάποιες περιπτώσεις, καταδεικνύει ότι οι αξιολογήσεις μιας μόνο εκτέλεσης δεν παρέχουν ουσιώδη εικόνα των ικανοτήτων ενός συστήματος. Το εύρημα αυτό έχει βαθιές συνέπειες για την έρευνα και ανάπτυξη LLM, αναδεικνύοντας την ανάγκη *στατιστικής αυστηρότητας* στις μεθοδολογίες αξιολόγησης όπου η αξιοπιστία είναι καίρια.

Η εργασία μας εγκαθιδρύει ότι η αξιόπιστη ανάπτυξη LLM σε σύνθετους τομείς απαιτεί τρία ουσιώδη συστατικά: (i) *προηγμένη υποδομή αξιολόγησης* που αποτυπώνει την πολυπλοκότητα του πραγματικού κόσμου, (ii) *προσαρμοστικούς μηχανισμούς βελτιστοποίησης* που επιτρέπουν συνεχή βελτίωση βάσει εκβάσεων, και (iii) *αυστηρά πρωτόκολλα αξιολόγησης* που λαμβάνουν υπόψη τη στοχαστικότητα του συστήματος. Η ολοκληρωμένη ενσωμάτωση αυτών των στοιχείων στο προτεινόμενο πλαίσιο προσφέρει ένα *πρότυπο ανάπτυξης LLM* και σε άλλες κρίσιμες εφαρμογές πέραν των χρηματοοικονομικών συναλλαγών.

1.6.2 Μελλοντικές Κατευθύνσεις

Τα θεμέλια που θέτουν τα **StockSim** και **ATLAS** ανοίγουν πλήθος προοπτικών για την εξέλιξη της έρευνας γύρω από LLMs σε χρηματοοικονομικές εφαρμογές και, ευρύτερα, σε περιβάλλοντα σειριακής λήψης αποφάσεων. Από την εργασία μας αναδύονται αρκετές υποσχόμενες κατευθύνσεις που μπορούν να ενισχύσουν ουσιαστικά τον αντίκτυπο της και να αποκαλύψουν νέες πτυχές των ικανοτήτων των LLMs σε σύνθετα, δυναμικά περιβάλλοντα.

Χρονική Ανάλυση και Δυναμική Αγοράς: Η τρέχουσα αξιολόγηση επικεντρώνεται σε ημερήσιες αποφάσεις συναλλαγών, κάτι που ενδέχεται να περιορίζει την πλήρη δυναμική των συστημάτων LLM. Μελλοντική έρευνα θα πρέπει να εξετάσει σενάρια υψηλότερης συχνότητας, όπου τα μοντέλα αποφασίζουν ανά ώρα, λεπτό ή ακόμη και σε επίπεδο *tick*. Τέτοιες μελέτες μπορούν να δείξουν αν τα LLMs υπερέρχουν σε ταχεία αναγνώριση προτύπων και προσαρμογή όταν λαμβάνουν συχνότερη ανάδραση και περισσότερες ευκαιρίες απόφασης.

Διεύρυνση Κατηγοριών Περιουσιακών Στοιχείων: Η επέκταση πέρα από τις μετοχές μπορεί να δείξει τον βαθμό γενίκευσης των ευρημάτων σε άλλα χρηματοοικονομικά μέσα. Οι αγορές κρυπτονομισμάτων, με διαπραγμάτευση 24/7 και ιδιάζοντα μοτίβα μεταβλητότητας, αποτελούν ένα ιδιαίτερα ενδιαφέρον πεδίο για τον έλεγχο της προσαρμοστικότητας των LLMs. Αντίστοιχα, τα εμπορεύματα (π.χ. χρυσός, πετρέλαιο) διέπονται από διαφορετικούς θεμελιώδεις παράγοντες και έντονη εποχικότητα, δοκιμάζοντας τις συλλογιστικές ικανότητες των LLMs με νέους τρόπους. Οι αγορές σταθερού εισοδήματος, με υψηλή εξάρτηση από τα επιτόκια και τον πιστωτικό κίνδυνο, συνθέτουν πολυπαραγοντικά πλαίσια αποφάσεων, όπου δοκιμάζεται η ικανότητα διαχείρισης

σύνθετων αλληλεπιδράσεων μεταξύ παραγόντων.

Ενίσχυση και Γενίκευση του Adaptive-OPRO: Η προτεινόμενη μεθοδολογία βελτιστοποίησης προτροπών αποτελεί μια αρχική διερεύνηση με σημαντικά περιθώρια εξέλιξης. Μελλοντικές εργασίες μπορούν να εξετάσουν εναλλακτικά *meta-prompts* με πιο εξελιγμένη καθοδήγηση, διαφορετικά συστήματα αξιολόγησης πέρα από το ROI (π.χ. μετρικές προσαρμοσμένες στον κίνδυνο ή πολυκριτηριακές μετρικές), καθώς και εναλλακτικούς ρυθμούς βελτιστοποίησης που εξισορροπούν την ταχύτητα προσαρμογής με τη σταθερότητα. Κατεξοχήν, η μεταφορά του **Adaptive-OPRO** σε ετερόκλητους τομείς ακολουθιακής λήψης αποφάσεων, όπως ο σχεδιασμός θεραπείας στην υγεία, η διαχείριση εφοδιαστικής αλυσίδας ή η στρατηγική κατανομή πόρων, θα τεκμηρίωνε τη γενικευσιμότητά του και θα ανέδειχνε τις απαιτήσεις προσαρμογής ανά πεδίο.

Εξελικτική Βελτιστοποίηση Προτροπών: Προτείνουμε μια επέκταση όπου οι προτροπές δεν βελτιστοποιούνται απομονωμένα ανά μετοχή, αλλά εξελίσσονται συλλογικά μέσω γενετικών αλγορίθμων. Για κάθε μετοχή εκτελείται ένας αυτόνομος πράκτορας με **Adaptive-OPRO** που βελτιστοποιεί το δικό του prompt. Πάνω από αυτούς λειτουργεί ένα *μετα-επίπεδο εξέλιξης* που επιλέγει τις πιο αποδοτικές προτροπές, εφαρμόζει διασταύρωση και μετάλλαξη και παράγει νέες υποψήφιες εκδοχές, οι οποίες επαναξιολογούνται από τους πράκτορες. Κατ’ αυτόν τον τρόπο, ιδιοσυγκρασιακά γνώρισματά ανά μετοχή μπορούν να ανασυνδυαστούν και να γενικευτούν, παράγοντας υβριδικές προτροπές που συνθέτουν επιτυχημένες στρατηγικές από ετερογενή περιβάλλοντα. Παρότι τα αρχικά πειράματα με αυτή την προσέγγιση ήταν ενθαρρυντικά σε εννοιολογικό επίπεδο, προέκυψαν ουσιαστικές προκλήσεις υλοποίησης: ωστόσο, η προοπτική ανάπτυξης προτροπών που ενσωματώνουν δια-ενεργητικές (cross-asset) γνώσεις συναλλαγών αποτελεί ισχυρό κίνητρο για συνέχιση της έρευνας.

Επίδραση στην Αγορά και Ανάλυση Θεσμικής Κλίμακας: Η εκτέλεση σε επίπεδο εντολών του StockSim επιτρέπει τη μελέτη σεναρίων όπου πράκτορες LLM διαθέτουν επαρκές κεφάλαιο ώστε να επηρεάζουν τις τιμές, δηλαδή να υιοθετούν συμπεριφορές θεσμικών επενδυτών. Η έρευνα μπορεί να εξετάσει πώς προσαρμόζονται οι στρατηγικές όταν οι ίδιες οι ενέργειες του πράκτορα μεταβάλλουν τη δυναμική της αγοράς, αν αναπτύσσονται εξελιγμένες στρατηγικές διάσπασης εντολών (order-splitting) για ελαχιστοποίηση του *market impact*, και πώς αλληλεπιδρούν πολλαπλά συστήματα LLM μεγάλης κλίμακας στο ίδιο περιβάλλον. Η κατεύθυνση αυτή έχει κρίσιμες προεκτάσεις για την κατανόηση πιθανών συστημικών επιπτώσεων από την ευρεία υιοθέτηση της Τεχνητής Νοημοσύνης στις χρηματοοικονομικές αγορές.

Προσωπικότητες Συναλλαγών και Προφίλ Κινδύνου: Η ευελιξία των LLMs επιτρέπει συστηματική διερεύνηση διαφορετικών «προσωπικοτήτων» συναλλαγών μέσω εξειδικευμένων προτροπών: συντηρητικές, επιθετικές, προσανατολισμένες στην ορμή ή αντιθετικές (contrarian). Η αποτίμηση αυτών των περσόνων μπορεί να δείξει πώς διαφορετικοί στρατηγικοί προσανατολισμοί επηρεάζουν την επίδοση ανά καθεστώς αγοράς, να αποκαλύψει βέλτιστους συσχετισμούς *προσωπικότητας-συνθηκών* και να υποδείξει πώς η «προσωπικότητα» μπορεί να προσαρμόζεται δυναμικά καθώς μεταβάλλονται οι συνθήκες της αγοράς.

Διαχείριση Πολυ-Ενεργητικού Χαρτοφυλακίου: Η μετάβαση του **ATLAS** από εστίαση σε μία μετοχή σε διαφοροποιημένη διαχείριση χαρτοφυλακίου θα δοκιμάσει τις ικανότητες συντονισμού για ταυτόχρονες αποφάσεις σε πολλαπλά μέσα. Η έρευνα μπορεί να εξετάσει τον χειρισμό συσχετίσεων, την κλαδική περιστροφή (sector rotation) και τη δυναμική κατανομή κεφαλαίου, διατηρώντας ευθυγραμμισμένους τους συνολικούς στόχους του χαρτοφυλακίου. Η επέκταση αυτή απαιτεί νέους μηχανισμούς ορχήστρωσης και κατάλληλες προσεγγίσεις βελτιστοποίησης για χώρους αποφάσεων υψηλής διάστασης.

Μακροχρόνιες Μελέτες Προσαρμογής: Αν και τα τρέχοντα παράθυρα αξιολόγησης επαρκούν για μια αρχική αποτίμηση, η επέκταση του ορίζοντα σε μήνες ή χρόνια θα επιτρέψει τη μελέτη της μακροπρόθεσμης προσαρμογής και της εξέλιξης των στρατηγικών. Αναλύσεις σε βάθος χρόνου μπορούν να δείξουν αν τα συστήματα LLM διατηρούν σταθερότητα επίδοσης, πώς ανταποκρίνονται σε δομικές μεταβολές της αγοράς και αν, μέσω συνεχούς βελτιστοποίησης, αναπτύσσουν σταδιακά πιο σύνθετες και αποτελεσματικές στρατηγικές.

Η ανοικτή διάθεση του κώδικα για τα **StockSim** και **ATLAS** επιτρέπει στην ερευνητική κοινότητα να εξελίξει συλλογικά τις παραπάνω κατευθύνσεις, οικοδομώντας πάνω σε στέρεες βάσεις ώστε να εμβαθύνει στην κατανόηση των ικανοτήτων των LLM σε σύνθετα, υψηλού διακυβεύματος περιβάλλοντα λήψης αποφάσεων. Καθώς οι αρχιτεκτονικές LLM προοδεύουν, οι εν λόγω γραμμές έρευνας θα αποκτούν ολοένα μεγαλύτερη σημασία για τη μετουσίωση των νέων δυνατοτήτων σε αξιόπιστες, εφαρμόσιμες λύσεις στις χρηματοοικονομικές αγορές και πέρα από αυτές.

Chapter 2

Introduction

In recent years, the rapid advancement of Large Language Models (LLMs) has fundamentally transformed our understanding of artificial intelligence capabilities in complex reasoning and decision-making tasks [23, 35, 34, 1, 52, 49, 64, 19, 30, 17, 63, 57]. These sophisticated models have demonstrated remarkable proficiency in processing vast amounts of information, synthesizing insights from diverse sources, and generating coherent responses across a wide spectrum of applications. As LLMs continue to evolve and demonstrate human-level performance in increasingly challenging domains, their potential for deployment in high-stakes, real-world scenarios has become a subject of intense research interest and practical importance.

Among the various domains where LLMs show transformative potential, financial markets represent one of the most demanding testing grounds. These markets embody the essence of complex decision-making environments, characterized by high uncertainty, consequential outcomes, and measurable performance metrics [78, 58, 47]. Every trading day, participants must synthesize vast amounts of heterogeneous information—from technical indicators and fundamental analysis to breaking news and market sentiment—while making decisions that reveal their quality only over extended time horizons. The emergence of LLMs introduces unprecedented opportunities for financial decision-making through their demonstrated ability to process diverse data sources, reason over complex scenarios, and adapt to changing conditions [9, 43].

From a research perspective, financial markets offer several compelling advantages for LLM evaluation. Unlike synthetic benchmarks that may suffer from distribution shifts or limited complexity, markets provide unlimited historical data without simulation bias, while demanding genuine understanding rather than pattern memorization. The domain requires synthesizing structured numerical data (time series, indicators) with unstructured text (news articles, analyst reports), testing LLMs’ ability to perform multimodal reasoning. Moreover, financial markets exhibit high stochasticity and non-stationary dynamics—properties that quickly expose brittle or overfit solutions, ensuring that successful methods must demonstrate genuine robustness and adaptability.

However, despite this promising potential, the development of effective LLM-based trading systems faces significant challenges that have limited progress in the field. The NLP community currently encounters substantial obstacles due to the absence of standardized, comprehensive, and openly accessible platforms specifically designed for rigorous evaluation of LLMs in realistic trading contexts [37, 44]. Existing evaluation frameworks suffer from critical limitations: they either provide basic trading simulation capabilities while abstracting away crucial microstructure aspects like latency and detailed order-book dynamics, or they focus on precise order-level market mechanics but depend heavily on expensive, limited datasets that restrict practical applicability.

Furthermore, current approaches to LLM-based trading systems exhibit three fundamental weaknesses. First, they typically rely on manually crafted, static prompts that fail to adapt to changing market conditions and evolving system performance; this is a critical limitation in environments where continuous learning and adaptation are essential for success. Second, most existing systems employ isolated decision-making approaches that fail to leverage the benefits of specialized expertise and coordinated analysis across different market aspects. Third, they often oversimplify the action space, reducing complex trading decisions to basic

directional predictions rather than complete trade specifications including order types, quantities, timing, and risk management parameters.

This thesis addresses these limitations through a comprehensive two-stage contribution that advances both the infrastructure for LLM evaluation in financial domains and the methodological approaches for building effective trading systems.

The first contribution is StockSim, an open-source simulation platform that provides the missing infrastructure for systematic evaluation of LLM-based trading agents in realistic market environments [51]. StockSim bridges the gap between simplified backtesting and real-world trading by offering dual execution modes: detailed order-level simulation that captures latency effects, market impact, and realistic order execution dynamics, and scalable candlestick-level execution that enables comprehensive evaluation across diverse market scenarios. The platform’s extensible architecture supports multi-agent coordination, incorporates external information sources such as news and financial reports, and provides production-grade infrastructure with authentic market data integration and comprehensive performance tracking.

Building upon this foundation, we present ATLAS (Adaptive Trading with LLM Agent Systems), a sophisticated multi-agent trading framework that demonstrates how to effectively harness LLM capabilities for financial decision-making [50]. ATLAS employs a coordinated architecture where specialized agents focus on distinct aspects of market analysis—technical patterns, news synthesis, and fundamental evaluation—while a central trading agent synthesizes these insights into coherent trading strategies. At the core of ATLAS is Adaptive-OPRO, our novel extension of Optimization by PROMpting [79] to sequential decision-making environments with delayed, noisy rewards. This adaptation enables systematic prompt refinement based on trading outcomes, allowing the system to continuously improve its decision-making capabilities.

Our extensive experimental evaluation reveals several critical insights that extend beyond financial applications to the broader deployment of LLMs in complex, consequential domains. We demonstrate that Adaptive-OPRO consistently outperforms both traditional quantitative strategies and existing LLM-based approaches across diverse market conditions. Surprisingly, we uncover a "reflection paradox" where additional reasoning steps—commonly assumed to be beneficial—can actually degrade performance in well-optimized systems, challenging conventional wisdom about LLM reasoning enhancement. Furthermore, our rigorous multi-run evaluation protocols expose severe reliability issues in the single-run assessments that dominate current research, highlighting the critical importance of robust statistical evaluation in stochastic environments.

These findings establish fundamental principles for deploying LLMs in high-stakes domains characterized by sequential decision-making under real-world uncertainty, while simultaneously advancing our understanding of LLM behavior, optimization, and coordination in complex environments.

The outline of this thesis is as follows:

- We provide a background section that establishes the essential theoretical foundations spanning Large Language Models and transformer architectures alongside financial market concepts, providing the comprehensive foundation needed to understand the methodology and experimental design of our work.
- We survey the landscape of related work, encompassing LLM agents in financial markets, methodologies for prompt optimization, and contemporary trading simulation platforms, thus setting the stage for our research contributions.
- We present the detailed methodology behind both StockSim and ATLAS, including the dual-mode simulation architecture, the multi-agent coordination framework, and our novel Adaptive-OPRO algorithm for prompt optimization in sequential decision-making environments.
- We describe our comprehensive experimental setup, including the diverse LLM models evaluated, the range of market conditions tested, the baseline strategies compared against, and our multi-run evaluation protocol designed to ensure statistical reliability.
- We present detailed results and analysis covering performance comparisons, ablation studies, behavioral pattern analysis across different LLMs, as well as insights into the reflection paradox and implications for evaluation methodology. We also delve into Adaptive-OPRO’s underlying mechanisms to reveal how systematic prompt refinement drives performance improvements through iterative weakness detection and architectural evolution.

-
- Finally, we conclude with a discussion of the broader implications of our findings for LLM deployment in complex domains, limitations of the current work, and directions for future research in the intersection of artificial intelligence and financial decision-making.

Chapter 3

Background

This chapter provides the essential theoretical foundations for understanding the methodologies and experimental design presented in this thesis. We examine two complementary domains: Large Language Models and their architectural principles, which form the core computational framework of our trading system, and financial market concepts, which define the domain-specific environment in which our agents operate. A comprehensive understanding of both areas is crucial for interpreting the design decisions, experimental protocols, and results discussed in subsequent sections.

The following analysis is structured to first establish the computational foundations through an examination of transformer-based architectures, training paradigms, and prompting methodologies that enable LLMs to function as autonomous trading agents. We then transition to the financial domain, covering market microstructure, technical analysis, and fundamental analysis concepts that inform the decision-making processes of our multi-agent trading system.

Contents

3.1	Large Language Models and Transformer Architecture	22
3.1.1	Evolution of Language Models	22
3.1.2	Transformer Architecture	22
3.1.3	Scaling Laws and Emergent Capabilities	25
3.2	Training Paradigms for Large Language Models	25
3.2.1	Learning Paradigm Overview	25
3.2.2	Pretraining and Fine-Tuning	26
3.2.3	Domain Adaptation	26
3.3	Prompting and In-Context Learning	27
3.3.1	Prompting Methodology	27
3.3.2	Few-Shot and Zero-Shot Learning	27
3.3.3	Chain-of-Thought and Advanced Prompting	27
3.4	Large Reasoning Models	27
3.5	Multi-Agent Systems and Coordination	28
3.5.1	Agent-Based Architectures	28
3.5.2	Coordination Mechanisms	28
3.6	Financial Markets and Trading Foundations	29
3.6.1	Market Microstructure	29
3.6.2	Technical Analysis Framework	29
3.6.3	Fundamental Analysis Framework	31
3.6.4	Trading Across Market Regimes	32

3.1 Large Language Models and Transformer Architecture

3.1.1 Evolution of Language Models

Language models represent computational frameworks designed to understand and generate human language by modeling probability distributions over word sequences. The fundamental challenge in probabilistic language modeling is estimating the joint probability of a specific n -word sequence, denoted as $P(w_1, w_2, \dots, w_n)$. This probability distribution forms the foundation for two key applications: language understanding, which evaluates the likelihood of sentences, and text generation, which samples probable continuations.

The earliest approaches to this challenge employed **n-gram models**, which apply the Markov chain assumption that the probability of the next word depends only on a fixed window of previous words. A bigram model, for example, models the probability of a sequence as:

$$P(w_1, w_2, \dots, w_n) = P(w_1) \cdot P(w_2|w_1) \cdot P(w_3|w_2) \cdot \dots \cdot P(w_n|w_{n-1}) \quad (3.1.1)$$

where the conditional probability $P(w_k|w_{k-1})$ can be estimated by calculating the proportion of occurrences where word w_k appears immediately after word w_{k-1} in the training corpus. An n-gram model is trained by determining these probabilities from text corpora in one or more languages. While computationally efficient, n-gram models suffer from the curse of dimensionality and struggle to assign meaningful probabilities to unseen word sequences, despite various smoothing techniques developed to address these limitations.

The introduction of **neural language models** marked a significant advancement, replacing discrete word identities with continuous embedding vectors [5]. This approach enabled better generalization to unseen sequences by learning distributed representations that capture semantic relationships between words. However, these early feedforward architectures remained limited by fixed context windows and sequential processing constraints.

The development of **recurrent neural networks** and subsequently **Long Short-Term Memory (LSTM)** networks [28] addressed some limitations of feedforward models by introducing mechanisms for handling variable-length sequences and maintaining long-term dependencies. LSTMs introduced gating mechanisms that selectively retain or forget information, solving the vanishing gradient problem that plagued earlier recurrent architectures.

3.1.2 Transformer Architecture

The Transformer architecture [65] revolutionized natural language processing by discarding recurrence and convolution in favor of self-attention, which models long-range dependencies across input sequences. This design underpins modern Large Language Models, enabling strong performance in both understanding and generation.

Overview and Core Components

At a high level, the Transformer is a sequence-to-sequence model composed of an *encoder* and a *decoder*. The encoder maps an input sequence to a sequence of continuous representations; the decoder autoregressively generates outputs while attending to the encoder's states.

Input Representation Tokens are first embedded into continuous vectors, then augmented with positional information to compensate for attention's position invariance. Concretely, token embeddings $E \in \mathbb{R}^{V \times d}$ combine with positional encodings $P \in \mathbb{R}^{n \times d}$, where V is the vocabulary size, d the embedding dimension, and n the sequence length.

Encoder Stack Each encoder layer contains two sublayers:

- *Multi-Head Self-Attention*: lets each position attend to all others across multiple learned subspaces, producing context-aware representations.

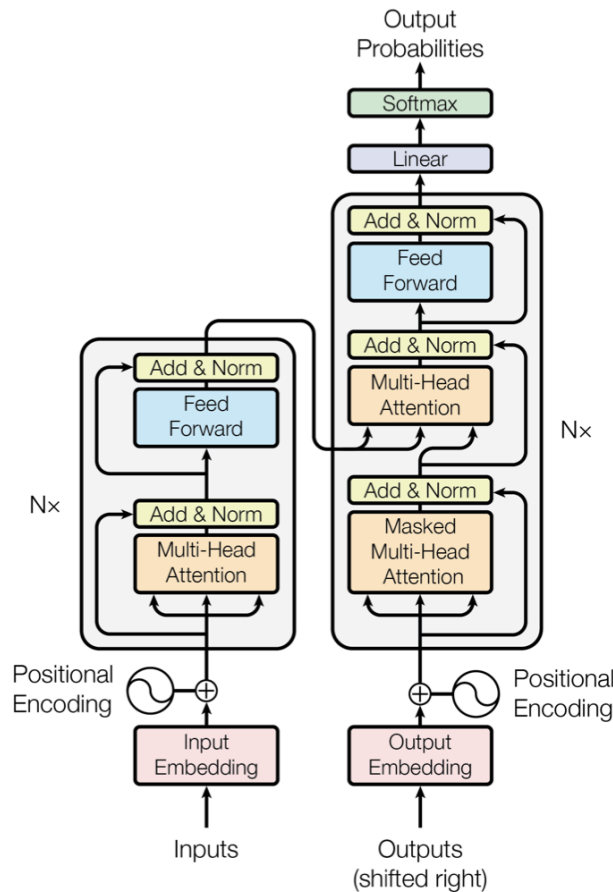


Figure 3.1.1: **The Transformer: model architecture.** The original encoder–decoder design stacks multi-head self-attention and position-wise feed-forward sublayers in both the encoder (left) and decoder (right) [65].

- *Position-wise Feed-Forward Network:* applies identical fully connected transformations independently at each position, adding nonlinearity and capacity.

Decoder Stack Each decoder layer comprises three sublayers:

- *Masked Multi-Head Self-Attention:* blocks access to future tokens during training, preserving the autoregressive property.
- *Multi-Head Encoder–Decoder Attention:* allows the decoder to attend to encoder outputs, channeling source-side information into generation.
- *Position-wise Feed-Forward Network:* identical in form to the encoder’s feed-forward sublayer.

Stabilization and Training All sublayers use residual connections [27] followed by layer normalization [4], which stabilizes optimization and supports deep stacks.

Output Layer A final linear projection followed by a softmax yields a probability distribution over the next token (or class) at each decoding step.

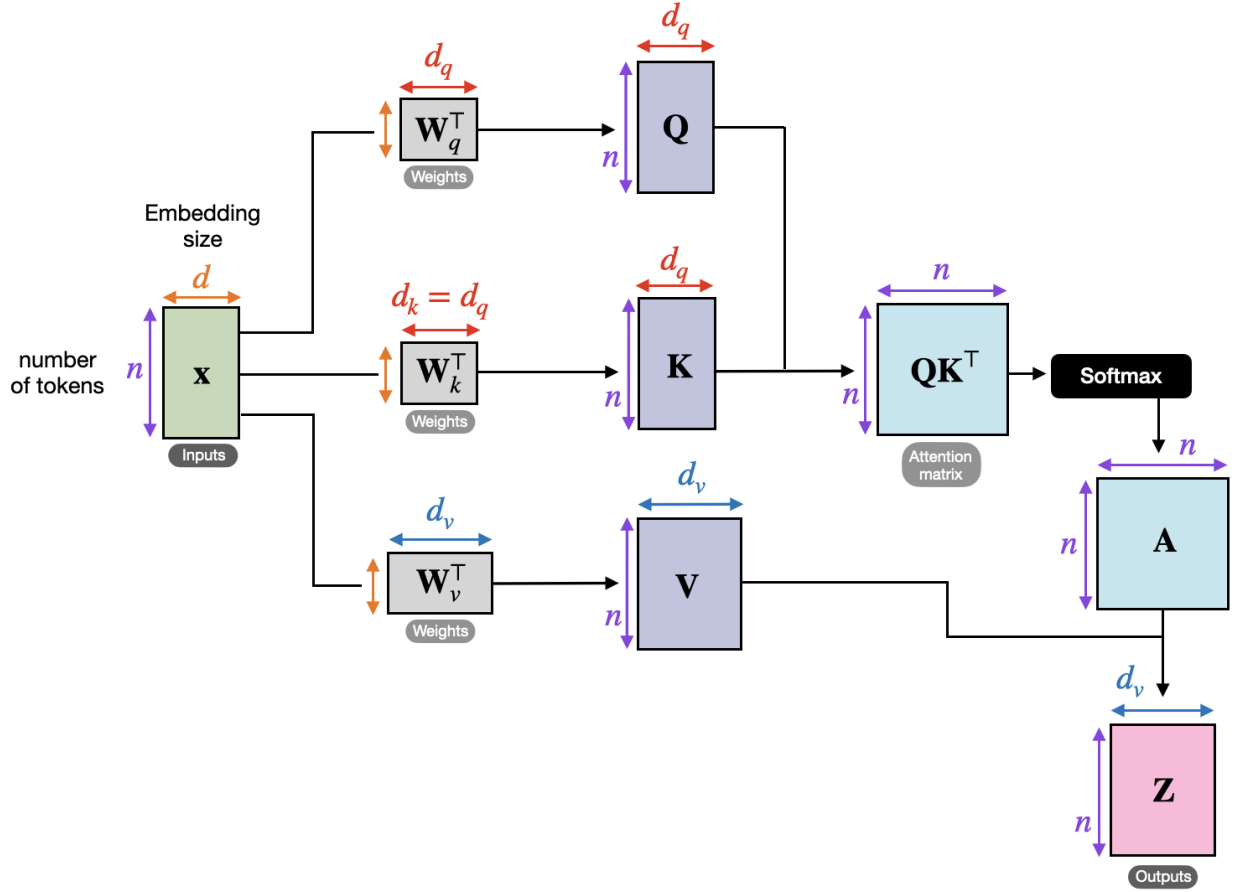


Figure 3.1.2: Self-Attention Mechanism

Self-Attention Mechanism

The self-attention mechanism forms the core innovation of the transformer architecture. For input embeddings $X \in \mathbb{R}^{n \times d}$, the attention mechanism computes:

$$Q = XW^Q, \quad K = XW^K, \quad V = XW^V \quad (3.1.2)$$

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right) V \quad (3.1.3)$$

where $W^Q \in \mathbb{R}^{d \times d_q}$, $W^K \in \mathbb{R}^{d \times d_k}$, $W^V \in \mathbb{R}^{d \times d_v}$ are learned projection matrices, and the scaling factor $\sqrt{d_k}$ prevents softmax saturation in high-dimensional spaces. This mechanism enables each token to selectively attend to relevant parts of the input sequence, capturing both local and long-range dependencies.

Multi-head attention extends this mechanism by computing attention across multiple representation subspaces:

$$\text{MultiHead}(Q, K, V) = \text{Concat}(\text{head}_1, \dots, \text{head}_h)W^O \quad (3.1.4)$$

$$\text{head}_i = \text{Attention}(QW_i^Q, KW_i^K, VW_i^V) \quad (3.1.5)$$

where $W^O \in \mathbb{R}^{hd_k \times d}$ projects the concatenated heads back to the model dimension. This parallel computation allows the model to capture different types of relationships simultaneously.

Architectural Variants

Modern language models employ three primary transformer variants, distinguished by their attention masking patterns and intended use cases:

Encoder-Only Models (e.g., BERT [13]) use bidirectional attention, enabling each position to attend to all other positions. This architecture excels at understanding tasks requiring complete sequence context, such as classification and named entity recognition, but cannot generate text autoregressively.

Decoder-Only Models (e.g., GPT series [54, 55, 7]) employ causal masking, restricting attention to previous positions only. This enables autoregressive generation while maintaining the ability to process input context, making it the predominant architecture for modern LLMs used in interactive applications.

Encoder-Decoder Models (e.g., T5 [56], BART [36]) combine bidirectional encoding with autoregressive decoding, excelling at sequence-to-sequence tasks like translation and summarization but requiring more complex architectures and training procedures.

For this thesis, we focus exclusively on decoder-only architectures, as they provide the generative capabilities and instruction-following behavior essential for autonomous trading agents.

3.1.3 Scaling Laws and Emergent Capabilities

A fundamental discovery in modern LLM research is the existence of predictable scaling relationships governing model performance. **Neural scaling laws** [29] demonstrate that model performance follows power-law relationships with respect to model size, dataset size, and computational budget:

$$L(N, D, C) \propto N^{-\alpha_N} D^{-\alpha_D} C^{-\alpha_C} \quad (3.1.6)$$

where L represents loss, N is the number of parameters, D is dataset size, C is computational budget, and the exponents α characterize the scaling behavior. These relationships enable principled decisions about resource allocation and performance prediction.

Emergent abilities [68] refer to capabilities that appear suddenly in sufficiently large models but are absent in smaller ones. These abilities, such as few-shot learning, chain-of-thought reasoning, and instruction following, typically emerge at specific computational scales (approximately 10^{22} FLOPs for training). The emergence of these capabilities is particularly relevant for financial applications, where complex reasoning and adaptation to new market conditions are crucial.

However, recent work [59] suggests that apparent emergent behavior may sometimes result from measurement choices and evaluation metrics rather than fundamental phase transitions, emphasizing the importance of careful evaluation in understanding model capabilities.

3.2 Training Paradigms for Large Language Models

3.2.1 Learning Paradigm Overview

The development of effective LLMs requires sophisticated training methodologies that leverage different types of supervision and learning signals. Understanding these paradigms is essential for comprehending how the models in our experiments acquired their capabilities and limitations.

Self-Supervised Learning has emerged as the dominant paradigm for LLM pretraining. Models learn from unlabeled text by solving pretext tasks such as next-token prediction or masked language modeling, enabling them to develop rich representations of language structure and semantics from vast corpora [7]. The mathematical foundation rests on learning to approximate the data distribution $P(x)$ without explicit supervision.

Supervised Learning involves training on explicitly labeled input-output pairs, where models learn to approximate $P(y|x)$ through direct supervision [13]. While effective for specific tasks, supervised learning requires extensive labeled datasets and may not capture the full complexity of language understanding.

Semi-Supervised Learning [10] addresses the scarcity of labeled data by incorporating unlabeled examples to improve generalization, while **Meta-Learning** [20] focuses on developing learning strategies that allow rapid task adaptation with minimal training examples.

3.2.2 Pretraining and Fine-Tuning

Modern LLM development follows a two-stage process that first builds general language understanding before specializing for specific applications.

Pretraining Phase

During pretraining, language models are trained on vast collections of text using self-supervised learning, which leverages the structure of raw text without requiring manually labeled data. For decoder-only models, the central training task is next-token prediction: the model learns to predict the next token in a sequence, given all previous tokens.

$$\mathcal{L}_{\text{pretrain}} = - \sum_{t=1}^T \log P(w_t \mid w_{<t}; \theta) \quad (3.2.1)$$

This pretraining loss function measures how well the model predicts each token w_t in a sequence of length T by computing the negative log-likelihood of the correct token at position t , conditioned on all previous tokens $w_{<t}$ and the model parameters θ . Through this process, the model acquires internal representations that capture syntax, semantics, factual knowledge, and reasoning patterns present in the data.

The effectiveness of pretraining relies heavily on scale: state-of-the-art models are trained on hundreds of billions to trillions of tokens, requiring immense computational resources such as thousands of GPUs running for months. This large-scale pretraining is crucial for the emergent capabilities observed in large language models.

Fine-Tuning Phase

Fine-tuning adapts pretrained models to specific domains or tasks through additional training on curated datasets. This process typically requires orders of magnitude less data and computation than pretraining while achieving substantial performance improvements on target tasks.

Supervised Fine-Tuning (SFT) trains models on input-output pairs relevant to the target application. For specialized applications, this might involve fine-tuning on domain-specific data and decision-making examples to improve task-relevant understanding and performance.

Instruction Tuning [69] fine-tunes models to follow natural language instructions, enabling zero-shot and few-shot task performance. This paradigm is particularly valuable for applications requiring agents to interpret and execute complex instructions about domain-specific analysis and decision-making processes.

Reinforcement Learning from Human Feedback (RLHF) [48] further refines model behavior by optimizing for human preferences using policy gradient methods. The process involves training a reward model on human preference data, then optimizing the language model using reinforcement learning algorithms like PPO [60].

3.2.3 Domain Adaptation

Effective deployment of LLMs in specialized domains often requires domain-adaptive training strategies. **Domain-Adaptive Pretraining** [25] continues pretraining on domain-specific corpora, helping models acquire specialized knowledge and terminology before fine-tuning on specific tasks.

For specialized applications, domain-adaptive pretraining on relevant professional literature, technical documentation, and domain-specific datasets can significantly improve performance on related tasks. However, this approach requires substantial computational resources and domain expertise to curate appropriate training data that captures the nuances and specialized knowledge of the target field.

3.3 Prompting and In-Context Learning

3.3.1 Prompting Methodology

Prompt-based learning represents a paradigm shift from traditional supervised learning, leveraging the generative capabilities of pretrained language models to solve tasks through carefully designed input formatting rather than parameter updates [40].

The core insight underlying prompting is that large language models, through extensive pretraining, develop the ability to perform many tasks when presented with appropriate context. Rather than training task-specific models, prompting reformulates problems as text completion tasks that exploit the model’s existing capabilities.

Prompt Structure

A prompt typically consists of:

- **Task Description:** Natural language instructions explaining the desired behavior
- **Context:** Relevant background information or constraints
- **Input Slot:** Placeholder for task-specific input
- **Output Format:** Specification of desired response structure

For trading applications, prompts must carefully balance flexibility with precision, providing sufficient context for informed decision-making while maintaining structured outputs compatible with execution systems.

3.3.2 Few-Shot and Zero-Shot Learning

The prompting paradigm enables models to perform tasks with minimal or no task-specific examples through **in-context learning** [7]:

Zero-Shot Learning provides only task instructions without examples. Models rely entirely on pretraining knowledge and instruction-following capabilities developed during fine-tuning.

Few-Shot Learning includes a small number of input-output examples within the prompt. The number of examples typically ranges from 1-10, limited by model context windows. The theoretical foundation for few-shot learning connects to meta-learning principles, where models learn to extract task-relevant patterns from the provided examples.

Research suggests that in-context learning can be understood through a Bayesian inference lens, where the model implicitly infers the task from examples and applies this understanding to new inputs [75]. This perspective explains why larger models with more extensive pretraining perform better at in-context learning.

3.3.3 Chain-of-Thought and Advanced Prompting

Chain-of-Thought (CoT) Prompting [70] improves model performance on complex reasoning tasks by encouraging step-by-step problem decomposition. Rather than requesting direct answers, CoT prompts ask models to articulate their reasoning process, leading to improved accuracy on mathematical, logical, and multi-step reasoning problems.

The effectiveness of CoT prompting appears to be an emergent property of scale, becoming more pronounced in larger models. This technique proves particularly valuable for trading applications, where decisions require integrating multiple information sources and considering various market factors.

3.4 Large Reasoning Models

Recent developments in language model architecture have given rise to **Large Reasoning Models (LRMs)**, which represent a paradigm shift from conventional token prediction toward explicit multi-step reasoning ca-

pabilities. LRMs employ reinforcement learning and step-wise feedback to optimize full reasoning trajectories rather than solely maximizing likelihood of next-token predictions [77].

The key architectural distinction lies in **process supervision**, where models receive feedback on intermediate reasoning steps rather than only final outcomes. This approach has demonstrated significant improvements over outcome-based supervision on complex reasoning tasks, as it enables models to learn from partial solutions and avoid compounding errors. Training typically involves large-scale reinforcement learning without initial supervised fine-tuning, allowing reasoning behaviors to emerge naturally through iterative refinement and reward optimization [12].

LRMs integrate several advanced prompting techniques as core architectural components: extended chain-of-thought reasoning for step-by-step problem decomposition, tree-of-thought search mechanisms that explore multiple solution paths with systematic pruning, and self-consistency frameworks that aggregate outputs from multiple reasoning traces [39]. These capabilities enable more robust decision-making in domains requiring complex logical inference and multi-factor analysis.

For autonomous trading applications, LRMs offer significant advantages in handling the intricate reasoning required for market analysis, risk assessment, and strategic decision-making. The explicit reasoning traces provide interpretability crucial for understanding agent behavior. However, LRMs also present challenges including increased computational overhead, training instability, and potential overthinking behaviors that may impair performance on simpler tasks [12].

3.5 Multi-Agent Systems and Coordination

3.5.1 Agent-Based Architectures

Multi-Agent Systems (MAS) decompose complex tasks across multiple specialized agents, each with distinct capabilities, knowledge, or roles. In the context of LLM-based systems, agents typically consist of specialized prompts, knowledge sources, communication protocols, and coordination mechanisms.

The theoretical foundation for multi-agent systems draws from distributed computing, game theory, and artificial intelligence. Key principles include:

1. **Autonomy:** Each agent operates independently according to its goals and local information.
2. **Interaction:** Agents communicate and coordinate through defined protocols.
3. **Environment:** Agents operate within shared or overlapping environments.
4. **Organization:** System-level behavior emerges from individual agent actions and interactions.

Agent Specialization Patterns

Common agent roles in LLM-based MAS include:

- **Information Processing Agents** specialize in analyzing specific types of data, such as technical indicators, news sentiment, or fundamental metrics.
- **Decision Agents** synthesize information from multiple sources to make strategic decisions about market positioning and risk management.
- **Coordination Agents** manage task allocation, information flow, and system-level optimization across the multi-agent framework.
- **Execution Agents** interface with external systems to implement decisions determined by the multi-agent system.

3.5.2 Coordination Mechanisms

Effective coordination in multi-agent systems requires mechanisms for information sharing, conflict resolution, and collective decision-making:

Centralized Coordination employs a master agent that allocates tasks and aggregates results. This approach simplifies coordination but may create bottlenecks and single points of failure.

Decentralized Coordination enables peer-to-peer interaction through shared protocols. While more robust, this approach requires careful design to prevent conflicts and ensure convergence.

3.6 Financial Markets and Trading Foundations

3.6.1 Market Microstructure

Understanding market microstructure is essential for developing effective trading systems, as it governs how orders are executed and prices are formed. The theoretical foundation draws from information economics, auction theory, and market mechanism design [26].

Order Types and Execution

Modern financial markets support various order types that provide different execution characteristics:

Market Orders execute immediately at the best available prices. They provide certainty of execution but uncertain execution prices, particularly in volatile or illiquid markets. The trade-off between execution certainty and price uncertainty is fundamental to trading strategy design.

Limit Orders specify maximum purchase prices or minimum sale prices. They provide price certainty but execution uncertainty, as they only fill when market prices reach the specified limits. Limit orders contribute to market liquidity by providing depth at various price levels.

Stop Orders activate as market orders when prices reach specified trigger levels. They are commonly used for risk management (stop-loss orders) and trend-following strategies (stop-buy orders).

Limit Order Book Dynamics

The **Limit Order Book (LOB)** maintains queues of pending buy and sell orders at each price level. The LOB continuously updates as new orders arrive, existing orders execute, or pending orders are cancelled. Understanding LOB dynamics is crucial for algorithmic trading systems.

Key LOB concepts include:

- **Bid-Ask Spread:** The difference between the highest bid and lowest ask prices, representing the cost of immediacy
- **Market Depth:** The quantity of orders available at various price levels, indicating liquidity
- **Order Priority:** Time-price priority rules that determine execution order for orders at the same price level

Transaction Costs and Market Impact

Market Impact refers to the price movement caused by order execution. Large orders may consume multiple price levels in the order book, resulting in worse average execution prices.

Slippage measures the difference between expected and actual execution prices, arising from latency, market impact, and changing market conditions between order submission and execution.

Latency represents delays in order transmission and processing. Even millisecond delays can significantly impact execution quality in electronic markets where speed provides competitive advantages.

3.6.2 Technical Analysis Framework

Technical analysis employs historical price and volume data to identify trading opportunities and assess market conditions [46]. The theoretical foundation rests on three core assumptions:

1. Market prices discount all available information

2. Prices move in trends that persist over time
3. Historical patterns tend to repeat due to consistent market participant behavior

Price and Volume Data

OHLCV Bars (Open, High, Low, Close, Volume) represent the standard format for historical price data:

- **Open:** First traded price in the time period
- **High:** Maximum price during the period
- **Low:** Minimum price during the period
- **Close:** Final traded price in the period
- **Volume:** Total quantity traded during the period

Different timeframes (1-minute, 5-minute, daily, etc.) provide various perspectives on market behavior, from short-term noise to long-term trends. Multi-timeframe analysis enables identification of trends and patterns across different temporal scales.

Moving Averages

Moving averages smooth price data to identify trends and reduce noise. They form the foundation for many technical indicators and trading strategies.

Simple Moving Average (SMA) calculates the arithmetic mean over a specified period:

$$SMA_n = \frac{1}{n} \sum_{i=0}^{n-1} P_{t-i} \quad (3.6.1)$$

Exponential Moving Average (EMA) applies exponentially decreasing weights to older prices:

$$EMA_t = \alpha \cdot P_t + (1 - \alpha) \cdot EMA_{t-1} \quad (3.6.2)$$

where $\alpha = \frac{2}{n+1}$ is the smoothing factor. EMA responds more quickly to recent price changes compared to SMA.

Our analysis employs multiple SMA periods (20, 50, 100, 200 days) to capture different trend horizons and EMA periods (12, 26 days) for responsive trend analysis and MACD calculation.

Momentum Oscillators

Momentum indicators measure the speed and strength of price movements, helping identify overbought and oversold conditions.

Relative Strength Index (RSI) [72] oscillates between 0 and 100:

$$RSI = 100 - \frac{100}{1 + RS} \quad (3.6.3)$$

where $RS = \frac{\text{Average Gain}}{\text{Average Loss}}$ over a 14-day period. Values above 70 typically suggest overbought conditions while values below 30 indicate oversold conditions.

Moving Average Convergence Divergence (MACD) combines trend and momentum analysis:

$$MACD = EMA_{12} - EMA_{26} \quad (3.6.4)$$

$$\text{Signal Line} = EMA_9(\text{MACD}) \quad (3.6.5)$$

$$\text{Histogram} = \text{MACD} - \text{Signal Line} \quad (3.6.6)$$

MACD crossovers and divergences provide signals for trend changes and momentum shifts.

Volatility Measures

Volatility indicators assess market uncertainty and risk, essential for position sizing and risk management.

Average True Range (ATR) [72] measures volatility by considering price gaps:

$$\text{True Range (TR)} = \max[(H - L), |H - C_{\text{prev}}|, |L - C_{\text{prev}}|] \quad (3.6.7)$$

$$ATR_n = \frac{1}{n} \sum_{i=0}^{n-1} TR_{t-i} \quad (3.6.8)$$

Bollinger Bands [6] construct volatility-adjusted price channels:

$$\text{Middle Band} = SMA_{20} \quad (3.6.9)$$

$$\text{Upper Band} = SMA_{20} + 2\sigma \quad (3.6.10)$$

$$\text{Lower Band} = SMA_{20} - 2\sigma \quad (3.6.11)$$

where σ represents the standard deviation of closing prices. The bands expand during periods of high volatility and contract during low volatility periods.

Support and Resistance Analysis

Horizontal Support and Resistance Levels represent price levels where historical buying or selling pressure has emerged consistently. Support levels act as floors where buying interest has historically prevented further price decline, while resistance levels act as ceilings where selling pressure has prevented further price advance.

These levels gain strength through repeated testing, high trading volume, and extended time periods of relevance. The psychological and algorithmic significance of round numbers and previous high/low points contributes to their effectiveness.

Volume Profile analyzes trading activity across price levels:

- **Point of Control (POC):** Price level with highest volume
- **Value Area:** Price range containing 70% of trading volume
- **High Volume Nodes:** Prices with significantly elevated volume

Volume-based analysis provides insights into price levels where significant supply and demand decisions occurred, often creating future support and resistance zones.

3.6.3 Fundamental Analysis Framework

Fundamental analysis evaluates securities based on underlying economic and financial factors rather than price movements alone [53]. This approach focuses on determining the intrinsic value of assets through comprehensive analysis of financial statements, economic conditions, and company-specific factors.

Financial Statement Analysis

Income Statement Metrics assess operational performance and profitability:

- **Revenue:** Total income from business operations, providing the top-line growth measure
- **Gross Profit Margin:** $(Revenue - CostOfGoodsSold)/Revenue$, indicating production efficiency and pricing power
- **Operating Margin:** $OperatingIncome/Revenue$, measuring management's ability to control costs
- **Net Income:** Final profit after all expenses, representing earnings available to shareholders
- **Earnings Per Share (EPS):** $NetIncome/WeightedAverageShares$, providing per-share profitability measure

Cash Flow Analysis evaluates liquidity and financial health:

- **Operating Cash Flow:** Cash from core business operations, adjusted for non-cash items and working capital changes
- **Free Cash Flow:** Operating cash flow minus capital expenditures, representing cash available for distribution
- **Net Cash Flow:** Combined operating, investing, and financing flows, showing overall cash position changes

Balance Sheet Metrics assess financial position and leverage:

- **Total Assets:** Sum of all company resources, including current and non-current assets
- **Total Equity:** Assets minus liabilities, representing shareholder ownership value
- **Debt-to-Equity Ratio:** Total debt divided by total equity, measuring financial leverage

Corporate Actions

Stock Splits increase share count while proportionally reducing price to improve liquidity and accessibility:

- 1:2 Split: Each share becomes two shares at half price
- 1:4 Split: Each share becomes four shares at quarter price
- 1:10 Split: Each share becomes ten shares at one-tenth price

Dividends represent cash distributions to shareholders:

$$\text{Dividend Yield} = \frac{\text{Annual Dividends Per Share}}{\text{Current Stock Price}} \times 100\% \quad (3.6.12)$$

Dividend policy reflects management's capital allocation philosophy and cash generation capabilities, providing insights into company maturity and growth prospects.

3.6.4 Trading Across Market Regimes

Understanding different market conditions and their characteristics is crucial for developing adaptive trading strategies. Market regimes represent distinct periods where asset prices exhibit specific behavioral patterns, requiring tailored approaches for optimal performance.

Trending Markets exhibit sustained directional movement, either bullish (upward) or bearish (downward). Trend-following strategies typically perform well in these conditions, while mean-reversion approaches may underperform.

Range-Bound Markets oscillate within defined price boundaries without clear directional bias. Mean-reversion strategies often outperform in these environments, while trend-following approaches may generate false signals.

Volatile Markets exhibit large price swings regardless of overall direction. These conditions challenge both trend-following and mean-reversion approaches, requiring sophisticated risk management and position sizing techniques.

Low-Volatility Markets show minimal price variation, often creating challenges for active trading strategies due to reduced opportunities and higher relative transaction costs.

Chapter 4

Related Work

The intersection of Large Language Models and financial decision-making represents a rapidly evolving research area that bridges natural language processing, multi-agent systems, and computational finance. As LLMs have demonstrated remarkable capabilities in complex reasoning tasks, researchers have increasingly explored their potential for sequential decision-making in high-stakes environments such as financial markets. This chapter examines the current state of research across the key domains that inform our work, identifying both the advances achieved and the fundamental gaps that motivate our contributions.

The field has evolved from early applications of simple sentiment analysis to sophisticated multi-agent trading frameworks, yet several critical challenges remain unresolved. Current approaches often rely on manually crafted prompts that fail to adapt to changing market conditions, employ oversimplified action spaces that abstract away real-world complexity, and suffer from evaluation methodologies that may not adequately capture system reliability. Furthermore, the absence of standardized, comprehensive evaluation platforms has hindered systematic comparison of different approaches and slowed progress in understanding fundamental questions about LLM behavior in sequential decision-making environments.

To understand these limitations and establish the foundation for our contributions, our review examines three interconnected areas of research. First, we analyze current LLM-based agents used in financial markets, outlining common architectural patterns and their constraints. Second, we examine prompt engineering and optimization methodologies to identify gaps in adaptive prompt design for sequential decision-making. Finally, we review trading simulation platforms to understand the evaluation infrastructure challenges that have limited progress in the field. Through this comprehensive examination, we establish the theoretical foundation for our contributions and highlight the specific limitations that ATLAS and StockSim are designed to address.

Contents

4.1	LLM Agents in Financial Markets	34
4.1.1	Early Developments in Financial NLP	34
4.1.2	Contemporary Multi-Agent Trading Systems	34
4.1.3	Advanced Architectural Innovations	34
4.1.4	Fundamental Limitations of Current Systems	35
4.2	Prompt Engineering and Optimization	35
4.2.1	Evolution from Manual to Systematic Approaches	35
4.2.2	Optimization by Prompting (OPRO)	36
4.2.3	Limitations in Sequential Decision-Making Contexts	36
4.3	Trading Simulation Platforms	37
4.3.1	The Evaluation Infrastructure Challenge	37
4.3.2	Current Platform Landscape and Technical Limitations	37

4.1 LLM Agents in Financial Markets

The application of Large Language Models to financial decision-making has witnessed significant growth, with researchers exploring increasingly sophisticated architectures ranging from simple sentiment analysis to complex multi-agent frameworks.

4.1.1 Early Developments in Financial NLP

Early work in this domain focused primarily on leveraging LLMs for financial text analysis and market sentiment extraction. These foundational approaches demonstrated that language models could effectively process financial news and reports to generate trading signals, establishing the groundwork for more sophisticated applications [31].

Recent advances in financial sentiment analysis have shown that specialized models like **FinBERT** and domain-adapted architectures can significantly outperform general-purpose sentiment analysis tools in financial contexts [81]. The development of financial-specific language models, including **FinGPT** and **BloombergGPT**, has further demonstrated the value of domain adaptation for financial applications [80, 73]. These specialized models have shown remarkable capabilities in tasks ranging from financial text summarization to risk assessment and market prediction.

However, these initial systems were largely limited to single-modal information processing and lacked the comprehensive decision-making capabilities required for realistic trading scenarios. The gap between sentiment analysis and actionable trading decisions remained a significant challenge, motivating researchers to explore more integrated approaches.

4.1.2 Contemporary Multi-Agent Trading Systems

Recent advances have introduced more sophisticated multi-agent architectures that attempt to decompose the complexity of financial decision-making across specialized components. **CryptoTrade** represents one of the most comprehensive early frameworks, integrating on-chain and off-chain data through specialized market analysts, news analysts, and trading agents [38]. The system incorporates reflection mechanisms for strategy refinement, demonstrating how LLMs can be organized into collaborative structures for enhanced decision-making.

TradingAgents employs a more elaborate multi-agent architecture inspired by real trading firms, featuring fundamental analysts, sentiment analysts, technical analysts, and traders with varied risk profiles [74]. The framework includes two opposing researcher agents that engage in structured debates to assess market conditions, a risk management team that monitors exposure, and traders that synthesize insights from debates and historical data. This collaborative approach has shown notable improvements in performance metrics compared to baseline models.

4.1.3 Advanced Architectural Innovations

Other notable approaches have explored various architectural innovations beyond basic multi-agent coordination. **FinMem** introduces layered memory systems that enable agents to maintain long-term context and learning across trading sessions [83]. This approach addresses the challenge of maintaining consistency and learning from past decisions in dynamic market environments.

FLAG-Trader presents a fusion approach combining linguistic processing with gradient-driven reinforcement learning, demonstrating how LLMs can be integrated with traditional optimization techniques [76]. This hybrid architecture uses a partially fine-tuned LLM as the policy network, leveraging pre-trained knowledge while adapting to the financial domain through parameter-efficient fine-tuning.

TradExpert explores mixture of experts architectures for financial decision-making [14], while **MarketSenseAI** focuses specifically on analyzing corporate filings and earnings calls to extract trading insights [18]. These specialized approaches demonstrate the growing sophistication in applying LLMs to specific aspects of financial analysis.

4.1.4 Fundamental Limitations of Current Systems

Despite these architectural advances, current systems share several critical limitations that constrain their effectiveness and real-world applicability. Most fundamentally, they rely on **manually crafted prompts** that remain static throughout the trading process, failing to adapt to changing market conditions or learn from trading outcomes. This static approach prevents systems from improving their decision-making based on experience, a crucial capability for success in dynamic financial markets.

Additionally, these systems typically **oversimplify the action space**, reducing complex trading decisions to confidence scores or simple directional signals rather than complete trade specifications. Such abstraction distances the systems from real market operations and leaves fundamental questions unanswered about whether LLMs can truly understand market mechanics or merely generate plausible-sounding recommendations. The gap between simplified outputs and actual trading requirements prevents meaningful evaluation of LLM capabilities in realistic trading scenarios.

Perhaps most critically, the field’s **evaluation methodologies remain problematic**. Prevalent single-run evaluations fail to capture the extreme performance variance inherent in LLM-based systems, where stochasticity can produce substantial variations in outcomes [62, 2]. This methodological weakness may lead to spurious conclusions about system capabilities and hinders reliable comparison between different approaches.

Our work directly addresses these fundamental limitations through two integrated contributions. ATLAS tackles the static prompt limitation through Adaptive-OPRO, our novel extension of prompt optimization to sequential decision-making environments that enables systematic prompt refinement based on trading outcomes. The multi-agent coordination is enhanced through ATLAS’s specialized architecture where distinct agents focus on market analysis, news processing, and fundamental evaluation. The action space oversimplification is resolved through StockSim’s requirement for complete order specifications-including order types, quantities, prices, and timing-while ATLAS demonstrates that LLMs can generate sophisticated trade specifications that go far beyond simple directional predictions. Finally, we establish rigorous multi-run evaluation protocols that capture performance variance and provide statistically reliable conclusions about system capabilities, setting new standards for evaluation in this domain.

4.2 Prompt Engineering and Optimization

The systematic optimization of prompts has emerged as a crucial component for effective LLM deployment, moving beyond ad-hoc manual engineering toward principled optimization methodologies. This evolution reflects growing recognition that prompt design significantly impacts LLM performance and that systematic approaches can yield substantial improvements over human-crafted alternatives.

4.2.1 Evolution from Manual to Systematic Approaches

Traditional prompt engineering relied heavily on manual iteration and domain expertise, with researchers and practitioners crafting prompts through trial-and-error processes. While this approach could produce effective results, it suffered from several limitations: lack of systematic exploration of the prompt space, difficulty in scaling across different tasks and models, and inability to adapt to changing requirements or feedback.

The introduction of systematic prompt optimization methods has transformed this landscape. **Chain-of-thought (CoT)** prompting emerged as a foundational technique that enables LLMs to break down complex problems into intermediate reasoning steps, significantly improving performance on tasks requiring multi-step reasoning [70]. Building upon this foundation, researchers developed various prompting paradigms including **few-shot** prompting (providing multiple input-output examples), **zero-shot** prompting (direct task instructions without examples), and **one-shot** prompting (single example guidance) [7].

Advanced techniques have further enhanced LLM capabilities: **self-consistency** prompting generates multiple reasoning paths and selects the most consistent answer [67], **tree-of-thoughts (ToT)** prompting enables exploration of multiple reasoning branches with backtracking capabilities [82], **step-back** prompting encourages abstraction before tackling specific problems [84], and **role-based** prompting assigns specific personas to guide response style and expertise [33]. Recent work has shown that simple additions like *"Let’s think*

step by step" or *"Take a deep breath and work on this problem step-by-step"* can significantly improve LLM performance on reasoning tasks [32, 79].

These systematic approaches have moved the field beyond ad-hoc manual crafting toward principled methodologies that can be empirically validated, systematically improved, and reliably deployed across diverse applications and model architectures.

4.2.2 Optimization by Prompting (OPRO)

The introduction of **Optimization by PRompting (OPRO)** marked a significant advancement in this domain [79]. OPRO employs LLMs themselves as meta-optimizers, iteratively generating new prompts based on histories of previous ones and their performance scores. The approach demonstrated impressive results on mathematical reasoning tasks like GSM8K, where OPRO discovered prompts that substantially outperformed human-designed alternatives.

The OPRO framework operates through a meta-prompt that instructs an optimizer LLM to generate candidate instructions based on previous solutions and their scores. These candidates are then evaluated by a scorer LLM, and the results are fed back into the optimization loop. This process continues until no further improvements are observed, resulting in optimized prompts that can achieve significant performance gains over manually designed alternatives.

Subsequent research has explored various extensions and applications of this optimization paradigm. **Evo-Prompt** combines LLMs with evolutionary algorithms for discrete prompt optimization, demonstrating how bio-inspired optimization techniques can be integrated with language model capabilities [24]. The approach uses evolutionary operators like mutation and crossover, adapted to work with natural language prompts, achieving significant improvements over human-engineered prompts and existing automatic methods.

POEM tackles prompt optimization using reinforcement learning with episodic memory, showing how RL techniques can be adapted for prompt improvement [15]. **GRAD-SUM** incorporates gradient-inspired feedback mechanisms, exploring how continuous optimization concepts can be applied to discrete prompt spaces through gradient summarization techniques [3].

4.2.3 Limitations in Sequential Decision-Making Contexts

However, existing prompt optimization methods share a fundamental limitation: they assume conditions that are largely absent in real-world sequential decision-making tasks. These methods typically require **immediate feedback**, operate on **deterministic outcomes**, and assume that **decisions are independent**. Trading exemplifies the opposite scenario: rewards materialize over extended horizons, market volatility clouds performance signals, and each decision influences future states and opportunities.

This limitation is particularly pronounced in financial applications, where the relationship between decisions and outcomes is complex, delayed, and noisy. Traditional OPRO assumes that task performance can be immediately evaluated after each optimization iteration, enabling direct assessment of prompt effectiveness. In contrast, trading decisions reveal their quality only over extended periods as market positions develop and resolve, requiring fundamentally different approaches to performance evaluation and optimization.

Furthermore, existing optimization methods typically operate on static task formulations where prompt content remains constant across optimization iterations. Trading environments, however, require prompts that continuously incorporate changing information-current portfolio status, updated market analyses, recent order executions-while preserving optimized strategic guidance. This **dynamic content challenge** necessitates architectural innovations that separate static instructions subject to optimization from dynamic runtime data that must be preserved unchanged.

Our Adaptive-OPRO directly addresses these fundamental limitations by extending prompt optimization to sequential decision-making environments with delayed, noisy rewards. We develop cumulative performance metrics that operate over extended evaluation windows, enabling optimization based on trading outcomes rather than immediate task completion. The dynamic content challenge is resolved through our template-based architecture that cleanly separates static instructions subject to optimization from dynamic runtime

Framework	Execution Granularity	Async Latency	Real-time LOB	History Back-test	No-code Setup	LLM Agent Support	External News	Multi Instrument
StockSim (ours)	Order	✓	✓	✓	✓	✓	✓	✓
ABIDES	Order	~	✓	✓	✗	✗	✗	✓
PyMarketSim	Order	✗	✓	✗	✗	✗	✗	✓
JAX-LOB	Order	✗	✓	✓	✗	✗	✗	✓
FinRL / Meta	Bar	✗	✗	✓	~	✗	~	✓
TradingAgents	Daily	✗	✗	✓	✗	✓	✓	✓

Table 4.1: Feature comparison of open-source trading simulators. *Legend.* ✓: supported; ✗: not supported; ~: partial or approximate support.

data, allowing prompt refinement to target decision-making logic while preserving changing market information. This represents the first successful adaptation of optimization by prompting to domains characterized by temporal dependencies, delayed feedback, and noisy performance signals.

4.3 Trading Simulation Platforms

The evaluation of LLM-based trading systems requires sophisticated simulation platforms that can accurately model market dynamics while providing controlled environments for systematic comparison. However, the current landscape of evaluation platforms presents a fragmented ecosystem with significant limitations that have fundamentally hindered progress in LLM-based trading research.

4.3.1 The Evaluation Infrastructure Challenge

The absence of standardized, comprehensive evaluation platforms specifically designed for LLM assessment in financial contexts represents one of the most significant barriers to advancing research in this domain [37, 44]. This infrastructure gap has created a problematic landscape where researchers must either compromise on realism or invest substantial resources in developing custom evaluation pipelines, both of which impede systematic progress in the field.

Common evaluation practices that rely on static benchmark datasets inadvertently risk data leakage, as these datasets or similar financial texts often appear in LLM training corpora [16, 61, 71]. Consequently, performance metrics become inflated, and the models fail to generalize effectively to genuinely unseen scenarios, creating unrealistic expectations and potential financial risks when deployed.

The fragmented platform ecosystem prevents meaningful comparison between different approaches and makes it difficult to establish fundamental principles for effective LLM deployment in financial domains. Without standardized evaluation infrastructure, research findings often remain isolated and difficult to reproduce or build upon.

4.3.2 Current Platform Landscape and Technical Limitations

Existing evaluation frameworks present fundamental trade-offs between execution granularity, scalability, and LLM integration capabilities. Table 4.1 reveals the stark limitations across current platforms.

Order-Level Simulation Platforms including ABIDES, PyMarketSim, and JAX-LOB provide sophisticated market microstructure modeling with detailed limit-order book dynamics [8, 45, 21]. These platforms excel at capturing essential market mechanics including order matching with price-time priority, realistic latency effects through timestamp replay mechanisms, and detailed market impact modeling. However, they face critical limitations for LLM research applications. Most significantly, these platforms lack native support for LLM agent integration, requiring substantial custom development to incorporate language model components. They provide no built-in mechanisms for external information processing such as news feeds or fundamental data integration, essential capabilities for realistic LLM-based trading research. Furthermore, they depend heavily on expensive, proprietary tick-level datasets that severely constrain the scope of experiments and limit accessibility for many research groups.

Historical Backtesting Platforms such as BackTrader¹ and FinRL represent the opposite extreme, prioritizing scalability and broad market coverage over execution realism [41, 42]. These platforms can efficiently process extensive historical datasets across multiple assets and timeframes, supporting large-scale experimental studies. However, this scalability comes at the cost of market realism. These frameworks typically operate on aggregated candlestick data without modeling intra-bar price movements, ignore asynchronous latency effects that are crucial in actual trading, and employ simplified execution models that abstract away realistic market dynamics and trading constraints.

LLM-Specific Trading Frameworks designed for multi-agent coordination often attempt to address LLM integration challenges but face their own limitations that restrict practical applicability and scalability [74]. These platforms typically use highly simplified market representations based solely on coarse historical data, omitting realistic execution dynamics and latency considerations that are critical for evaluating LLM behavior in actual trading environments. While they enable exploration of agent coordination patterns, their simplified market modeling prevents assessment of whether LLMs can handle the full complexity of real market execution.

This infrastructure fragmentation creates systematic impediments where researchers must choose between realistic market modeling and practical scalability, forcing compromises that limit both the scope and reliability of research findings. The absence of standardized evaluation infrastructure prevents cumulative progress in understanding LLM capabilities for financial decision-making, as findings remain isolated within custom evaluation environments that cannot be easily reproduced or extended.

StockSim directly addresses these critical infrastructure gaps by providing the first comprehensive simulation platform specifically designed for LLM evaluation in financial domains. Our dual-mode architecture bridges the gap between order-level realism and scalable evaluation, offering both detailed limit-order book simulation for studying microstructure effects and efficient candlestick-level execution for comprehensive historical analysis. The platform provides native LLM integration with support for diverse model architectures, built-in multi-modal information processing for news and fundamental data, and configuration-driven experimentation that eliminates the need for custom development. By avoiding dependence on expensive proprietary datasets while maintaining realistic market dynamics, StockSim enables systematic, reproducible research across diverse market conditions and time periods, establishing the standardized evaluation infrastructure that the field has long needed.

¹<https://pypi.org/project/backtrader>

Chapter 5

Methodology

This chapter presents the detailed methodology behind our two primary contributions: StockSim, a comprehensive simulation platform for evaluating LLM-based trading agents, and ATLAS, an adaptive multi-agent trading framework that demonstrates effective deployment of LLMs in financial decision-making. Together, these systems address the fundamental limitations identified in our review of related work by providing both the evaluation infrastructure and methodological innovations necessary for reliable LLM deployment in complex, sequential decision-making environments.

Our methodology is structured around three core innovations. First, we introduce StockSim’s dual-mode simulation architecture that enables systematic evaluation across different levels of market complexity while maintaining realistic trading dynamics. Second, we present ATLAS’s multi-agent coordination framework that decomposes financial decision-making across specialized components while maintaining coherent strategy execution. Third, we detail our novel Adaptive-OPRO algorithm that extends prompt optimization to sequential decision-making environments with delayed, noisy feedback.

Together, these elements establish a reproducible pathway from controlled evaluation to adaptive deployment in financial markets, setting clear standards for rigor and extensibility in LLM-based sequential decision-making.

Contents

5.1 StockSim: Dual-Mode Simulation Architecture	39
5.1.1 System Architecture Overview	40
5.1.2 Dual Execution Modes	40
5.1.3 Data Integration Framework	41
5.1.4 Agent Framework Design	42
5.1.5 Evaluation and Analysis Framework	42
5.2 ATLAS: Multi-Agent Coordination Framework	44
5.2.1 Architectural Design Principles	44
5.2.2 Market Intelligence Pipeline	44
5.2.3 Decision and Execution Layer	45
5.2.4 Adaptive-OPRO: Prompt Optimization for Sequential Decision-Making	45

5.1 StockSim: Dual-Mode Simulation Architecture

StockSim addresses the critical gap in evaluation infrastructure by providing a unified platform that combines realistic market simulation with comprehensive LLM evaluation capabilities. The platform’s design philosophy centers on enabling systematic research while maintaining the fidelity necessary for meaningful conclusions about real-world deployment.

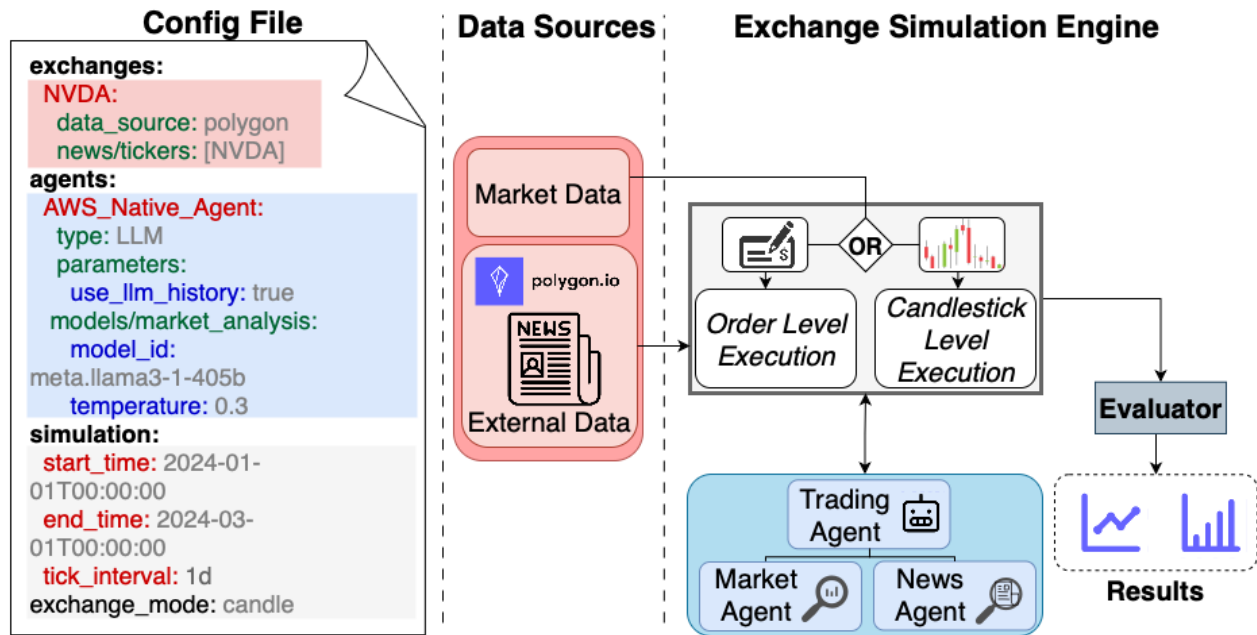


Figure 5.1.1: Overview of STOCKSIM’s system architecture and input/output scheme. Modules are color-coded by function and mapped to corresponding blocks in the centralized config file. This design supports flexible, code-free customization of simulation parameters, agent behavior, and data sources.

5.1.1 System Architecture Overview

StockSim employs a modular, asynchronous architecture designed around four core components that enable comprehensive LLM evaluation in realistic trading environments. The system architecture supports two distinct execution mechanisms—order-level and candlestick-level execution—seamlessly integrated with shared modules for market data retrieval, indicator computation, news and fundamentals integration, and agent interactions.

The **Exchange Simulation Engine** serves as the central coordinator, asynchronously managing the simulated trading environment. Its primary responsibilities include receiving and processing agent actions, simulating realistic market dynamics for order execution, computing and disseminating relevant market indicators, and providing agents with timely access to market and external information. The engine acts as the central intermediary between data sources and trading agents, routing data dynamically from their respective sources to agents upon request while maintaining internal states related to orders and trades.

Each agent runs as a separate process and communicates asynchronously with the engine through **RabbitMQ**, an advanced message broker that ensures reliable message delivery and scalable communication¹. This process-per-component architecture with message-based communication provides natural fault isolation and horizontal scalability, enabling the platform to support large-scale multi-agent simulations.

5.1.2 Dual Execution Modes

The platform’s innovation lies in supporting two complementary execution modes that address different research requirements while maintaining interface consistency.

Order-Level Execution

Order-level execution emulates real market behavior by operating directly on the **limit order book (LOB)**, where agents submit market, limit or stop orders that interact with a stream of order book events including placements, cancellations, and executions. Orders are matched based on price-time priority: a buy limit

¹www.rabbitmq.com

order at \$100 will execute only if a sell order exists at \$100 or lower; otherwise, it queues until matched or canceled.

Execution may be full or partial, depending on available volume. The environment updates tick-by-tick, capturing fine-grained dynamics such as queue position, order interleaving with other market participants, and the impact of latency between action submission and book update. This level offers high realism and is critical for evaluating strategies sensitive to microstructure effects.

The order-level mode incorporates several realistic market dynamics that are typically absent from simplified simulations:

- **Latency modeling:** Realistic delays between order submission and execution that reflect real-world network and processing delays
- **Market impact:** Price movements induced by large orders that clear substantial amounts of pending orders in the limit order book
- **Slippage:** The difference between intended and actual execution prices resulting from latency and market impact effects
- **Partial fills:** Orders may execute partially when insufficient volume exists at the target price level

Candlestick-Level Execution

Candlestick-level execution operates on aggregated **OHLCV** (Open, High, Low, Close, Volume) data, where orders are executed based on whether the agent’s target price falls within the range of a given candle. This mode provides access to larger datasets and enables testing over longer historical periods, addressing the practical limitations of order-level data availability and cost.

Despite its simplified execution model, the candlestick mode maintains realistic trading dynamics through sophisticated intra-candle price simulation. When only aggregated bar data is available, StockSim simulates realistic price paths within each bar, allowing agents to place conditional orders like stop losses that execute plausibly even though exact moment-to-moment data is not available.

This approach addresses a critical gap in existing backtesting frameworks, which typically assume that all prices within a candlestick are equally accessible. Our intra-candle simulation generates realistic price movements that respect the bar’s OHLC constraints while providing more accurate execution modeling for conditional orders.

5.1.3 Data Integration Framework

StockSim distinguishes between two primary categories of data: market data including price, volume, and order-flow information, and external data such as news, corporate actions, and fundamental metrics. The Exchange Simulation Engine orchestrates these inputs asynchronously, delivering them to agents in simulation time.

Market Data Processing

The platform supports both detailed order-level data and simplified bar-level candlestick data. In candlestick-level execution, data is provided as aggregated OHLCV summaries obtained from general data sources like Alpha Vantage and Polygon.io². For order-level execution, each market action such as placing, changing, or canceling an order is individually tracked. These detailed events come either from datasets like LOBSTER³ or from logs created during the simulation. Each event has precise timestamps (milliseconds), allowing realistic simulation of latency and slippage.

Both execution modes consistently provide agents with computed market indicators derived from real-time or historical market data. These indicators include moving averages, momentum oscillators, volatility measures, and volume-weighted metrics that serve as condensed numeric features for agent decision-making.

²www.alphavantage.co; polygon.io

³lobsterdata.com

External Information Integration

Agents may request news headlines, earnings calendars, corporate actions, or fundamental ratios at any simulation step. These streams are supplied through the same provider infrastructure and exposed via a unified query interface implemented by the Exchange Simulation Engine. This abstraction enables agents to reason over time-sensitive, multi-modal inputs while supporting the development of more interpretable, information-driven trading strategies.

The platform’s provider abstraction layer wraps each data source with lightweight adapters that map payloads to StockSim’s canonical schema. This design ensures that adding new data providers requires only contributing a single Python file, enabling the platform to evolve alongside the data ecosystem.

5.1.4 Agent Framework Design

The agent framework provides a unified interface that abstracts away the complexity of different execution modes, enabling researchers to focus on agent development rather than simulation mechanics. Regardless of which execution engine mode is active, every agent interacts with the simulator through the same asynchronous message API.

Core Agent Capabilities

Each agent in StockSim may subscribe to data streams for market state, technical indicators, and external content; submit and cancel orders with support for MARKET, LIMIT, and STOP instructions; receive execution outcomes and portfolio updates; and log reasoning through optional explanation strings that accompany every order.

The separation between agent logic and execution details allows a single agent implementation to be tested on both order-level and candlestick-level execution modes without code changes. This design isolates research focus on core questions like prompt engineering and reasoning strategies without requiring researchers to manage low-level simulation mechanics.

Multi-Agent Support

StockSim includes a modular `LLMTradingAgent` framework that supports coordination between specialist LLMs focused on different aspects of market analysis. Each analyst operates with its own prompt template, memory context, and reasoning function. While adding new analyst roles requires lightweight code changes, the process is intentionally simple and well-documented.

The modular structure ensures that agent internals remain decoupled from the simulation engine, enabling researchers to rapidly prototype new agent structures, experiment with different LLM backends per role, test various coordination strategies, and conduct clean ablation studies by toggling analysts through configuration changes.

5.1.5 Evaluation and Analysis Framework

The Evaluator component subscribes to all trade executions, recording a complete history of positions, cash, and realized profit and loss. Upon simulation completion, it computes core performance metrics including overall return, risk-adjusted ratios, drawdown statistics, and trade analytics, packaging them into uniform reports.

For visual diagnostics, the *Evaluator* provides several useful outputs, such as equity curves showing portfolio value changes over time, candlestick charts highlighting executed trade entries and exits clearly marked on the price data, and comprehensive summary tables of key trading performance metrics. These outputs can be directly generated by STOCKSIM’s built-in plotting utilities or exported in JSON format for further analysis.

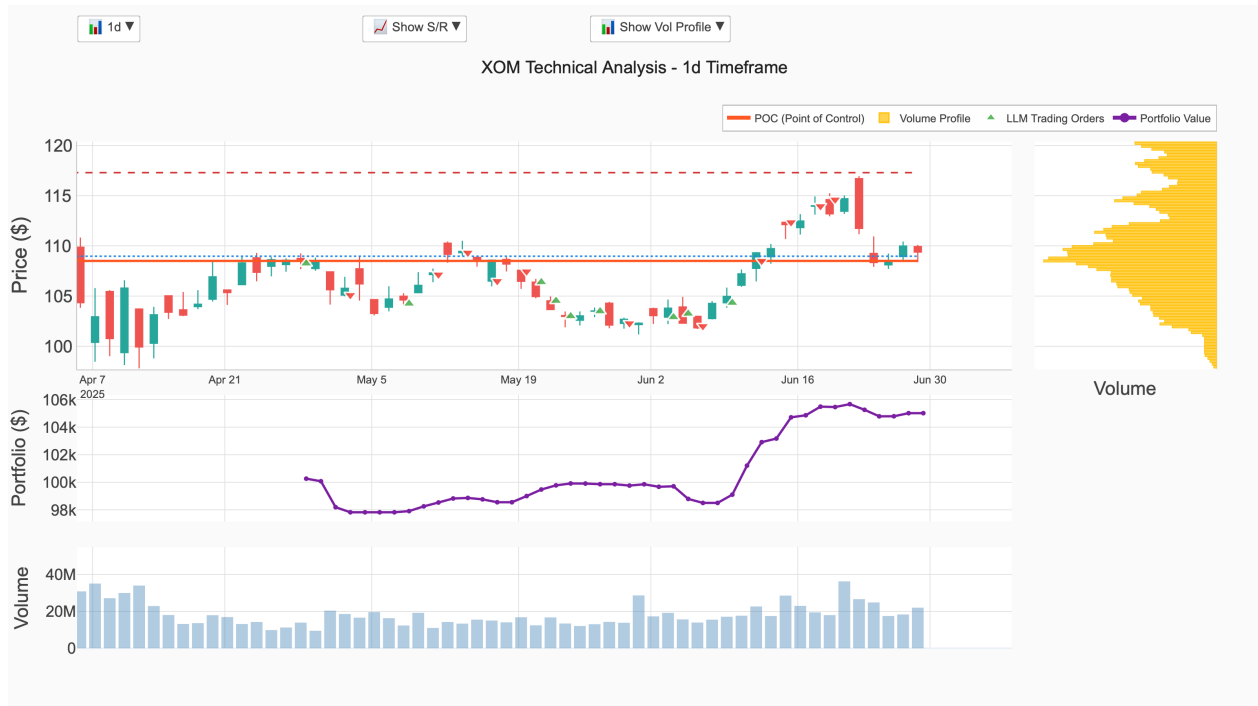


Figure 5.1.2: Example of an interactive chart generated by StockSim for the XOM stock using the Claude-4 model with the thinking mechanism enabled. The plot displays the price, buy and sell orders (annotated with ▲ and ▼, respectively), portfolio value, and trading volume. Users can zoom, adjust the time range, toggle chart components, and hover over elements to reveal additional details such as order execution prices and corresponding LLM outputs.

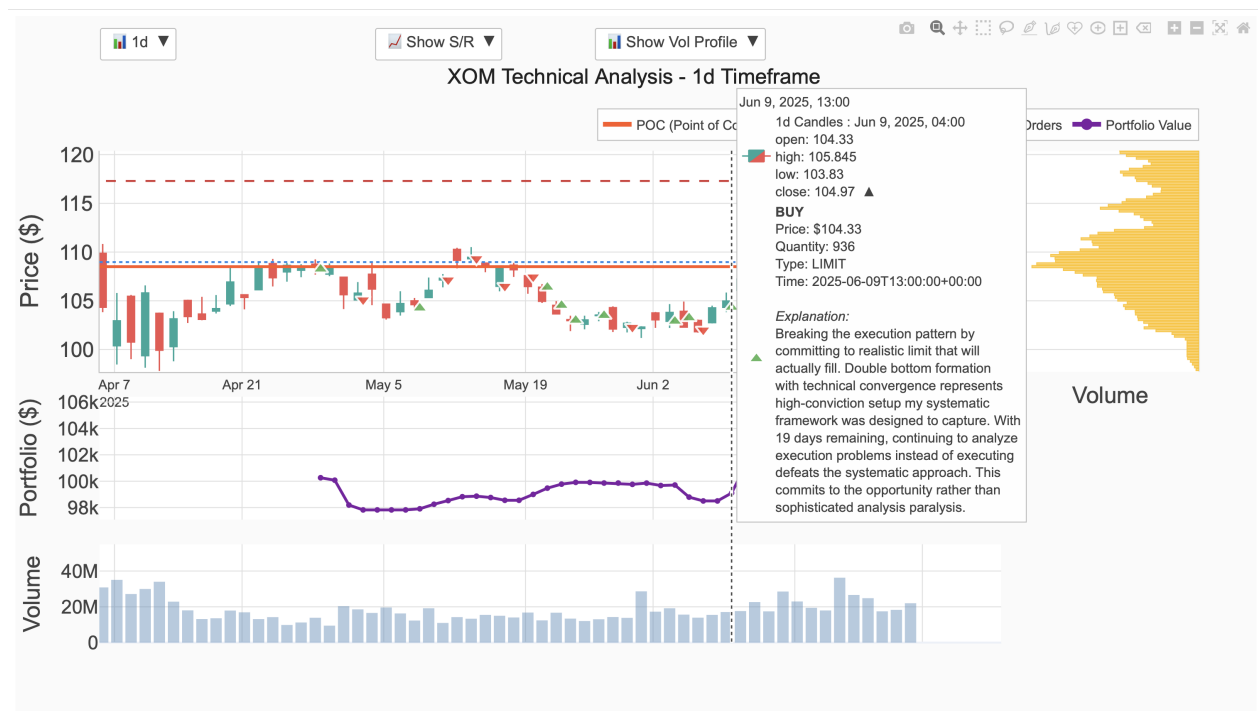


Figure 5.1.3: Demonstration of the hover functionality in StockSim. When hovering over a specific order, detailed information is displayed, including the exact execution price and the corresponding LLM output that led to the decision.

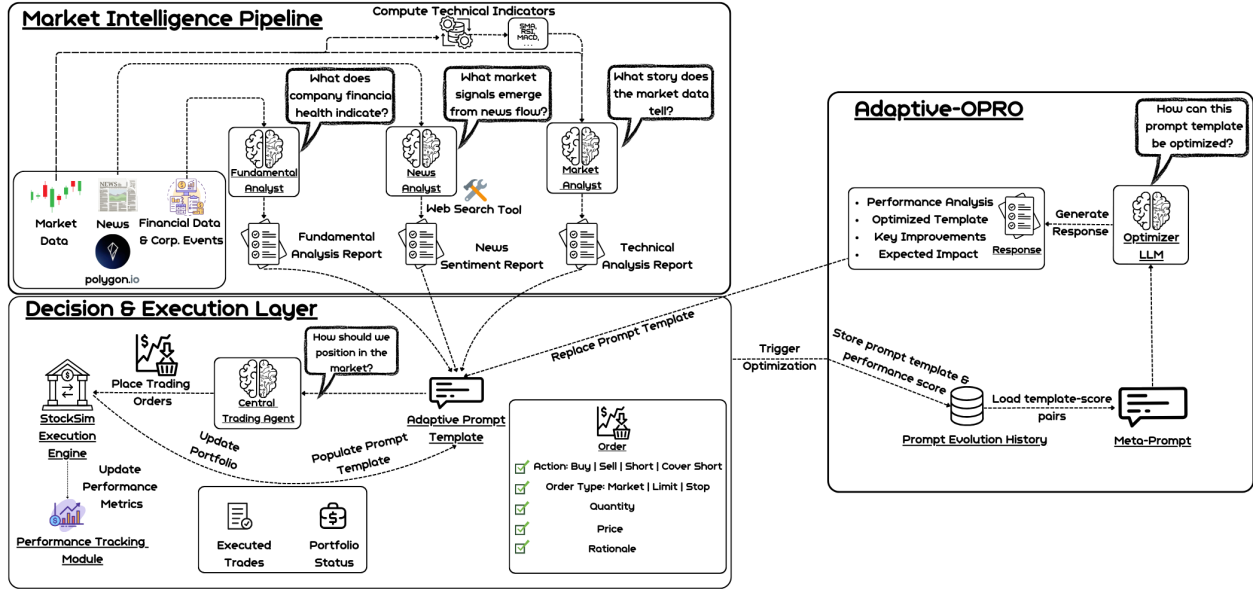


Figure 5.2.1: ATLAS Framework Overview. The Central Trading Agent submits orders to the Trading Execution Engine via prompts shaped by three specialized analysts and the proposed Adaptive-OPRO optimization technique.

5.2 ATLAS: Multi-Agent Coordination Framework

ATLAS (Adaptive Trading with LLM Agent Systems) demonstrates how to effectively deploy LLMs in financial decision-making through coordinated multi-agent architecture and adaptive prompt optimization. The framework addresses the key limitations of existing systems by decomposing market analysis across specialized components while maintaining unified decision-making through a central trading agent.

5.2.1 Architectural Design Principles

ATLAS is designed around two key pillars: a modular multi-agent architecture that separates concerns across different types of market information, and a dedicated prompt optimization mechanism tailored for sequential decision-making under uncertainty. These components enable ATLAS to act as a unified, adaptive system capable of making coherent, high-impact trading decisions grounded in structured insight and realistic execution constraints.

The framework’s modular design enhances both interpretability and strategic coherence by enabling each specialized agent to develop deep expertise in its assigned domain while contributing to collective decision-making. This approach contrasts with monolithic systems that attempt to process all information types within a single agent, often leading to diluted expertise and suboptimal integration of diverse signals.

5.2.2 Market Intelligence Pipeline

ATLAS employs specialized agents that focus on distinct types of market information, enabling deep analysis while maintaining system coherence through structured coordination.

Market Analyst. The Market Analyst processes raw price and volume data to generate structured multi-timescale summaries of market behavior. It analyzes each stock using three distinct historical windows: 2 years using monthly candlesticks, 6 months using weekly candlesticks, and 3 months using daily candlesticks. This multi-timeframe approach ensures that analysis operates at appropriate levels of detail for each temporal scope.

Within each window, the Market Analyst employs established technical indicators including Simple Moving

Averages (SMA), Exponential Moving Averages (EMA), Relative Strength Index (RSI), Moving Average Convergence Divergence (MACD), Average True Range (ATR), Bollinger Bands, Support and Resistance levels, and Volume Profile. These indicators reflect price dynamics, volatility patterns, and market participant behavior.

Rather than generating trading signals directly, the Market Analyst provides a consistent, noise-filtered view that preserves key market structure while abstracting away raw data complexity. This standardized approach allows the Central Trading Agent to focus on strategic decision-making without being overwhelmed by low-level market noise.

News Analyst. The News Analyst extracts market-relevant sentiment and signals from financial news content by processing headlines, sources, and summaries to generate structured outputs across four analytical dimensions: Sentiment Assessment, Key Developments, Market Relevance, and Source Analysis.

The agent may autonomously request access to full article text through a custom web scraping module when headlines or summaries lack sufficient detail. This capability enables deeper contextualization and improved detection of emergent themes or persistent tone changes that might be missed in headline-only analysis.

By transforming unstructured news flows into digestible, context-rich summaries, the News Analyst provides essential complement to price-driven market analysis, particularly valuable during periods of volatility driven by macroeconomic or firm-specific events.

Fundamental Analyst. The Fundamental Analyst extracts trading-relevant insights from periodic corporate disclosures including financial statements and structural events such as earnings, dividends, and stock splits. It parses detailed data fields including revenue, gross profit margins, operating income, and cash flow dynamics into concise analysis focused on material changes and trading implications.

To reflect the inherently low-frequency nature of fundamental information, the agent activates only 1-2 times per simulation, mimicking real-world reporting cycles. By focusing on actionable changes rather than exhaustive financial analysis, the Fundamental Analyst complements the more reactive Market and News modules with structural context essential for mid- to long-term positioning.

5.2.3 Decision and Execution Layer

Central Trading Agent

The Central Trading Agent serves as the core decision-maker, transforming analyst insights into concrete trading actions. Each trading day, before market open, it assesses the most recent market, news, and fundamentals analysis alongside current portfolio holdings to determine appropriate trading actions.

Based on this comprehensive context, the agent may issue trading orders that formally specify the quantity of shares to buy or sell, the target price, and any conditional requirements using different order types. These orders are submitted to the **StockSim Execution Engine**, which simulates realistic trading behavior to determine execution outcomes.

The daily decision loop supports disciplined, responsive trading behavior grounded in timely information and consistent with real-world execution constraints. The agent receives updated analyses and refreshed portfolio information each day, enabling it to adapt its strategy based on evolving market conditions and the outcomes of previous decisions.

5.2.4 Adaptive-OPRO: Prompt Optimization for Sequential Decision-Making

To enhance the Central Trading Agent’s decision-making capability, ATLAS introduces Adaptive-OPRO, a prompt optimization algorithm that systematically refines the agent’s instructions based on trading performance feedback. This represents the first successful extension of optimization by prompting to sequential decision-making environments with delayed, noisy rewards.

Extensions for Sequential Decision-Making

Our approach makes two key modifications to standard OPRO [79] to address financial market constraints. First, we address temporal delay between decisions and performance evaluation, as trading decisions only reveal their quality over extended periods as positions develop and resolve. Second, we handle dynamic content through prompt templates that separate static instructions subject to optimization from dynamic runtime data including market analyses, portfolio status, and recent trades.

This template-based architecture ensures that refinements target decision-making logic rather than transient market conditions, preventing the optimizer from overfitting to specific market scenarios while enabling adaptation of strategic reasoning patterns.

Prompt Evolution Architecture

Our optimization system maintains a prompt evolution history that records each prompt template alongside its measured trading performance. Every five decision steps, we evaluate the current prompt’s effectiveness by calculating cumulative Return on Investment (ROI) and store this prompt-performance pair in the **Prompt Evolution History**.

The five-step evaluation cycle balances reliable performance assessment with timely adaptation: shorter cycles introduce noise from market volatility, while longer cycles delay necessary prompt improvements. ROI provides a direct measure of trading success that reflects actual market outcomes without introducing noisy performance signals from auxiliary metrics.

The optimization process employs a **meta-prompt** that structures how the optimizer LLM analyzes accumulated performance records. This meta-prompt incorporates the complete prompt evolution history and systematically guides the **optimizer LLM** through four key outputs: performance analysis identifying current prompt strengths, weaknesses, and improvement opportunities; an optimized **prompt template** with enhanced instructions and improved decision-making logic; key improvements detailing specific modifications and their strategic rationale; and expected impact forecasting how changes will improve trading performance.

This comprehensive output supports systematic tracking of prompt evolution and provides explainability for optimization decisions, ensuring transparency in how and why the system adapts over time. The detailed analysis enables researchers to understand not just whether optimization is improving performance, but how and why specific changes contribute to enhanced decision-making.

Implementation Details

In the ATLAS framework, each agent maintains its own conversation history and specialized function. When optimizing, Adaptive-OPRO updates the Central Trading Agent’s initial prompt while preserving all accumulated information including market data, analyst reports, executed trades, and decision rationale.

The upstream analyst agents continue operating with their original static prompts and unmodified conversation histories. This targeted replacement strategy ensures that prompt updates don’t disrupt accumulated historical context, maintaining system stability and analytical consistency across the multi-agent architecture.

The optimization focuses exclusively on the Central Trading Agent since it serves as the decision-making hub that directly translates market intelligence into trading actions. This selective approach recognizes that different components benefit from different optimization strategies: decision-making agents benefit from adaptive refinement based on outcomes, while information processing agents benefit from stable, consistent analytical frameworks.

Chapter 6

Experimental Setup

This chapter describes our comprehensive experimental methodology designed to systematically evaluate ATLAS across diverse market conditions, LLM architectures, and prompting strategies. Our experimental design addresses four critical research questions: How does Adaptive-OPRO compare to existing approaches for enhancing LLM trading performance? What is the individual contribution of each specialized agent in the ATLAS architecture? How do different LLM architectures behave in complex trading environments? Which combinations of models and optimization strategies produce reliable versus high-variance outcomes?

The experimental framework is structured around rigorous evaluation protocols that account for the stochastic nature of both LLM outputs and market dynamics. Unlike previous research that typically employs single-run evaluations, our methodology implements systematic multi-run protocols that expose performance variance and ensure statistically reliable conclusions. This approach reveals critical insights about system reliability that would otherwise remain hidden.

Our evaluation spans multiple dimensions: diverse LLM architectures ranging from reasoning-enhanced models to open-source alternatives, varied market regimes that test adaptability across different trading environments, comprehensive baseline comparisons against both traditional quantitative strategies and existing LLM approaches, and systematic ablation studies that quantify the contribution of each framework component.

Contents

6.1	StockSim Platform Validation	49
6.1.1	Experimental Design for Platform Validation	49
6.1.2	Determinism and Reproducibility	49
6.1.3	Scalability Analysis	49
6.1.4	Resource Efficiency and Practical Implications	49
6.1.5	Validation Conclusions	49
6.2	Model Selection and Architecture Analysis	50
6.2.1	Reasoning-Enhanced Models	50
6.2.2	Standard Architecture Comparison	51
6.2.3	Open-Source Alternative	51
6.3	Market Conditions and Asset Selection	51
6.4	Prompting Strategy Evaluation	52
6.4.1	Core Experimental Configurations	52
6.4.2	Extended Prompting Strategy Analysis	52
6.4.3	Experimental Control and Comparison Framework	53
6.5	Evaluation Methodology and Metrics	53
6.5.1	Multi-Run Evaluation Protocol	53
6.5.2	Comprehensive Performance Metrics	53
6.5.3	Baseline Strategy Comparison	54

6.6	Ablation Study Design	55
6.6.1	Configuration Selection	55
6.6.2	Systematic Component Removal	55
6.6.3	Excluded Components	55
6.6.4	Evaluation Protocol	55

6.1 StockSim Platform Validation

Before presenting our main experimental results, we establish the reliability and scalability of our StockSim evaluation platform through controlled validation studies. These experiments ensure that our simulation environment can support large-scale, multi-agent experiments while maintaining the deterministic reproducibility essential for scientific evaluation.

6.1.1 Experimental Design for Platform Validation

We evaluate StockSim along two critical dimensions—**scalability** and **determinism**—using controlled experiments with deterministic agents that implement fixed strategies such as moving-average crossovers and buy-and-hold. This approach isolates the simulation engine’s behavior from LLM-induced variability, including response latency, resource consumption, and output stochasticity, ensuring that observed performance metrics reflect core platform capabilities rather than model-specific effects.

The validation protocol systematically varies agent counts from 10 to 500 agents while monitoring system-level performance metrics including CPU utilization across all cores and memory usage for both the simulation engine and the RabbitMQ message broker. Each configuration undergoes multiple identical runs to verify deterministic behavior and establish resource scaling patterns.

6.1.2 Determinism and Reproducibility

Across all repeated runs with identical configurations, simulation outputs—including order placements, executions, and computed performance metrics—remain perfectly consistent. This deterministic behavior validates StockSim’s end-to-end pipeline reliability, from data ingress through agent communication to final accounting. The platform’s deterministic properties provide a crucial foundation for meaningful statistical analysis of LLM trading strategies, as any observed performance variance can be attributed to model behavior rather than simulation artifacts.

6.1.3 Scalability Analysis

Figure 6.1.1 demonstrates StockSim’s scaling characteristics under increasing agent loads. The platform exhibits **near-linear scaling** up to approximately 150 concurrent agents, with simulation container CPU usage growing from 8% to 27% and memory consumption increasing from 0.8 GB to 2.0 GB—both roughly proportional to agent count.

Beyond this linear regime, resource demands enter a **super-linear growth phase**: at 300 and 500 agents, mean CPU utilization reaches 123% and 418% respectively, while memory consumption grows to 4.1 GB and 5.6 GB. Peak resource usage during high-load periods approaches roughly $4\times$ the average values, indicating that the platform handles load spikes gracefully without system instability.

6.1.4 Resource Efficiency and Practical Implications

Despite super-linear scaling at extreme loads, StockSim’s absolute resource requirements remain modest. Even at maximum tested scale (500 agents), the platform operates within 5.6 GB RAM and a few CPU cores on a MacBook Pro with Apple M3 Pro (11-core CPU) and 18 GB unified memory.

This efficiency profile has important implications for our experimental design: while running hundreds of concurrent LLM agents with full reasoning capabilities would be computationally prohibitive on standard hardware, StockSim’s efficiency enables comprehensive evaluation of agent coordination and interaction effects.

6.1.5 Validation Conclusions

These validation results establish StockSim as a robust foundation for systematic LLM trading evaluation. The platform’s verified determinism ensures reproducible results, while its demonstrated scalability supports the comprehensive multi-agent experiments presented in the following sections. The efficiency profile confirms

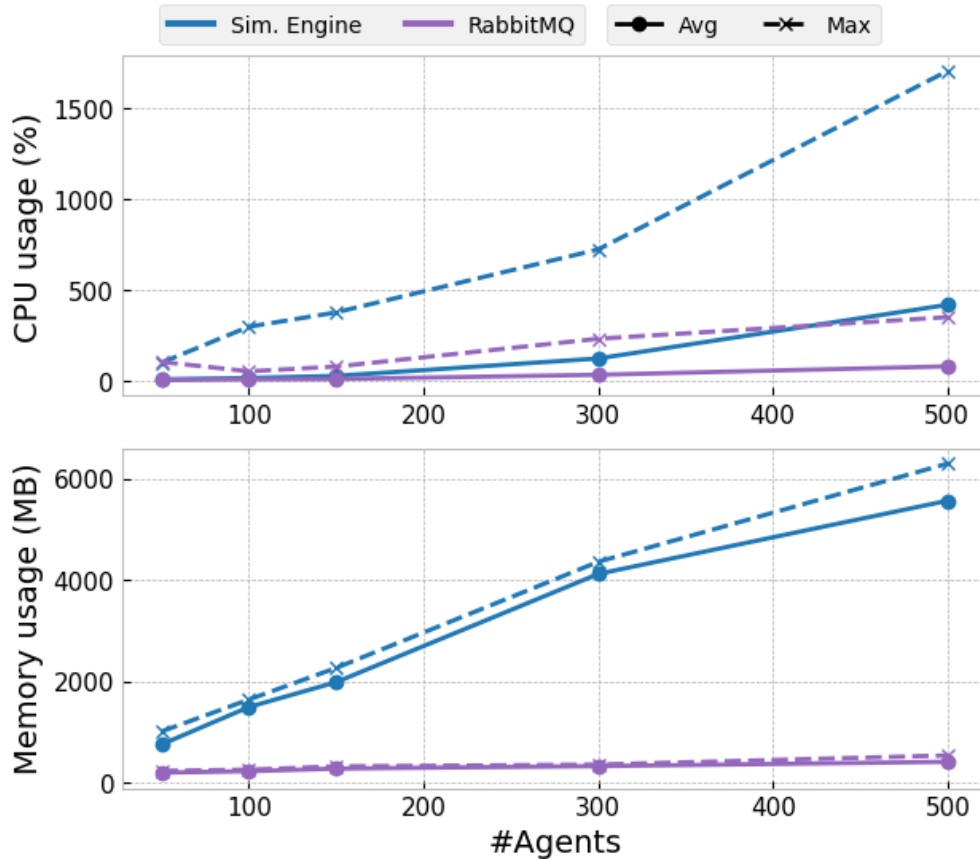


Figure 6.1.1: StockSim system performance metrics (memory/CPU usage) for varying numbers of deterministic agents.

that our experimental methodology can be replicated on accessible computing hardware, promoting broader reproducibility in LLM trading research.

With platform reliability established, we proceed to evaluate ATLAS performance across the diverse experimental conditions described in the following sections.

6.2 Model Selection and Architecture Analysis

Our model selection strategy explores how different architectural capabilities translate to sequential decision-making performance in financial markets. The chosen models represent distinct approaches to language modeling and reasoning, enabling systematic comparison of their effectiveness in trading environments.

6.2.1 Reasoning-Enhanced Models

We evaluate three models that employ explicit reasoning mechanisms (LRMs) designed to enhance decision-making through structured thought processes:

GPT-o3 represents the most advanced reasoning architecture in our evaluation, employing sophisticated internal chain-of-thought processes before generating responses. This model demonstrates exceptional capabilities in complex reasoning tasks and provides an opportunity to assess whether such advanced reasoning translates to superior trading performance. The model’s ability to maintain coherent reasoning chains across multiple decision points makes it particularly relevant for sequential trading decisions.

GPT-o4-mini offers a more efficient reasoning architecture while maintaining strong performance across diverse tasks. This model provides insights into the cost-performance tradeoffs in LLM-based trading systems,

as it delivers sophisticated reasoning capabilities at lower computational costs. Its inclusion enables evaluation of whether reasoning enhancement benefits can be achieved without requiring the most computationally expensive models.

Claude Sonnet 4 with Thinking enables controlled comparison by testing the same base architecture with and without explicit reasoning mechanisms. This model employs internal thinking processes that are visible during generation, allowing detailed analysis of how reasoning chains develop in trading contexts. The thinking mechanism provides unprecedented insight into the model’s decision-making process, enabling qualitative analysis of reasoning quality alongside quantitative performance evaluation.

6.2.2 Standard Architecture Comparison

Claude Sonnet 4 without Thinking serves as the controlled comparison for the reasoning-enhanced version, enabling isolation of the reasoning mechanism’s impact on trading performance. By testing the same base model architecture with and without its reasoning mode, we can determine whether explicit reasoning chains provide fundamental advantages or if base model capabilities suffice for effective trading decisions.

This controlled comparison is critical for understanding whether the additional computational overhead of reasoning mechanisms translates to meaningful performance improvements in dynamic trading environments.

6.2.3 Open-Source Alternative

LLaMA 3.3-70B represents the open-source alternative in our evaluation, testing whether effective trading requires proprietary training methods or emerges from sufficient model scale and architecture sophistication. This comparison is critical for organizations considering self-hosted deployment and for understanding the democratization potential of AI trading capabilities.

The inclusion of LLaMA enables assessment of whether competitive trading performance can be achieved using models with transparent architectures and training procedures. This evaluation has significant implications for the accessibility and broader adoption of LLM-based trading systems

6.3 Market Conditions and Asset Selection

Market regimes represent distinct patterns of price behavior that create fundamentally different challenges for trading systems. Each regime demands specific analytical approaches, risk management strategies, and decision-making frameworks, making regime diversity essential for comprehensive system evaluation. Trading systems that perform well under uniform conditions may fail catastrophically when market dynamics shift, making cross-regime testing crucial for assessing true robustness and adaptability. Our evaluation spans three distinct market regimes using stocks from diverse sectors, with each evaluation covering two months (April 28 - June 28, 2025). This duration provides sufficient time to observe multiple trading cycles and strategy evolution while maintaining complete conversation history within LLM context limits:

- **NVDA (Technology Sector - Upward Market [Bullish]):** During our evaluation period, NVDA exhibited steady upward movement characteristic of strong rising trends. This regime tests whether agents can effectively capture and maintain profitable positions during favorable market conditions. The technology sector’s inherent volatility and sentiment-driven dynamics provide additional complexity beyond simple trend-following.
- **XOM (Energy Sector - Oscillating Market [Sideways]):** XOM remained relatively stable within a trading range during the evaluation period, presenting challenges for momentum-based strategies while rewarding mean-reversion approaches. This regime tests patience and selectivity, as profitable opportunities are scarce and require precise timing. The energy sector’s dependency on broader economic conditions adds fundamental complexity to technical trading decisions.
- **LLY (Healthcare Sector - Downward Market [Bearish]):** LLY experienced sharp declines with sudden price movements, creating the most challenging environment in our evaluation. This regime tests risk management capabilities, adaptive position sizing, and the ability to navigate uncertainty while

preserving capital. Healthcare’s regulatory complexity and event-driven dynamics demand sophisticated information processing capabilities.

6.4 Prompting Strategy Evaluation

Our experimental design systematically compares multiple prompting approaches to isolate the impact of different adaptation mechanisms on trading performance. The strategy evaluation encompasses both established approaches from prior literature and our novel Adaptive-OPRO methodology.

6.4.1 Core Experimental Configurations

Baseline Strategy. We employ carefully engineered static prompts that represent best practices in prompt design for trading applications. These prompts were developed through iterative refinement by trading experts and incorporate structured decision frameworks, clear constraint specifications, and comprehensive context integration. The baseline represents the strongest static approach achievable through manual optimization, providing a robust foundation for measuring adaptive improvements.

The baseline prompts follow OpenAI’s reasoning best practices, incorporating clean structural sectioning, unambiguous instructions, standardized output formats, and clear separation between data and reasoning components¹. This establishes a strong comparative foundation that any adaptive approach must meaningfully exceed to demonstrate value.

Adaptive-OPRO Strategy. Our systematic prompt optimization approach targets the central trading agent exclusively, testing whether dynamic adaptation specifically enhances the complex synthesis of multiple information streams into trading decisions. This strategy serves as our primary experimental contribution, enabling direct comparison against static and reflection-based alternatives under identical market conditions.

Reflection Strategy. To benchmark the performance of Adaptive-OPRO against prior work, we adapt the reflection mechanism of [38]—one of the most recent and relevant algorithms for sequential feedback in LLM trading systems. The original approach, tailored for cryptocurrency trading, outputs confidence scores (-1 to 1) instead of concrete orders. For fair comparison, we adapt both methods to StockSim, where full trade specifications are required. We use the reflection mechanism as originally designed, adapting it to generate weekly feedback on trading patterns and strategic adjustments. This feedback is integrated into the Central Trading Agent via prompt context. Unlike Adaptive-OPRO’s direct prompt edits, the reflection output is analytical and must be interpreted by the agent. We apply the reflection mechanism without methodological modifications to maintain fidelity to the original design, enabling direct assessment of reflection effectiveness under identical trading constraints and evaluation conditions.

6.4.2 Extended Prompting Strategy Analysis

Our evaluation includes additional configurations that explore interaction effects and frequency variations:

Daily Reflection (1d). We evaluate more frequent reflection cycles to test whether immediate feedback enhances learning or introduces decision paralysis. This variant explores the temporal dimension of adaptation, revealing optimal feedback frequencies for trading decisions.

Combined Adaptive-OPRO with daily Reflection. This configuration tests whether multiple adaptation mechanisms complement each other or create interference. The combined approach reveals mechanism compatibility and identifies potential synergies between optimization and reflection approaches.

These extended configurations enable systematic analysis of adaptation frequency effects and mechanism interactions, providing comprehensive coverage of the adaptation strategy space while maintaining experimental control.

¹platform.openai.com/docs/guides/reasoning-best-practices

6.4.3 Experimental Control and Comparison Framework

Structural Consistency: All prompting strategies maintain identical constraint specifications, output formats, and data integration frameworks, ensuring that performance differences reflect adaptation effectiveness rather than structural prompt advantages.

Context Preservation: Each strategy preserves complete conversation histories and market context, enabling fair comparison of adaptation mechanisms without confounding factors from information access differences.

Output Standardization: All strategies generate identical JSON formats with standardized action specifications, eliminating output format variations as confounding variables in performance assessment.

This controlled comparison framework enables rigorous assessment of adaptation mechanism effectiveness while maintaining the experimental validity necessary for reliable conclusions about optimal approaches for LLM-based trading systems.

6.5 Evaluation Methodology and Metrics

Our evaluation methodology addresses fundamental challenges in assessing LLM-based trading systems, particularly the need to account for stochastic outputs, market dynamics, and performance variance that can obscure genuine capability differences.

6.5.1 Multi-Run Evaluation Protocol

Unlike prior research that typically employs single-run evaluations, we implement a systematic three-run protocol for each configuration. Each run initializes with identical starting conditions, market data, and system configurations, enabling isolation of performance variance attributable to LLM stochasticity versus environmental factors. This protocol exposes not just mean performance but stability of behavior under identical conditions.

Results are reported as mean $\pm$ standard deviation across three runs, enabling assessment of both central tendency and variability. Configurations with high variance receive different interpretation than those with consistent performance, as reliability becomes a crucial factor for real-world deployment consideration.

6.5.2 Comprehensive Performance Metrics

Our evaluation employs multiple metrics that capture different aspects of trading performance, preventing optimization for single measures while revealing comprehensive system capabilities:

Return on Investment (ROI): Calculated as the percentage change in total portfolio value from initial capital: $ROI = \frac{\text{final value} - \text{initial value}}{\text{initial value}} \times 100$. Portfolio values include both cash holdings and current market value of all stock positions. ROI provides the fundamental measure of trading effectiveness and capital growth.

Sharpe Ratio (SR): Risk-adjusted return metric calculated as: $SR = \frac{\mu - r_f}{\sigma}$, where μ is mean daily return, σ is daily return standard deviation, and r_f is the risk-free rate (set to 0 following [38]). Higher Sharpe ratios indicate superior risk-adjusted performance, revealing whether returns stem from genuine skill or excessive risk-taking.

Maximum Drawdown (DD): The largest peak-to-trough decline in portfolio value: $\text{Max DD} = \max_{t \in [0, T]} (\max_{s \in [0, t]} V_s - V_t) / \max_{s \in [0, t]} V_s$, where V_t represents portfolio value at time t . This metric captures downside risk tolerance and stress testing capabilities, revealing how systems perform under adverse conditions.

Win Rate: Percentage of profitable trades calculated as: $\text{Win Rate} = \frac{\text{Number of profitable trades}}{\text{Total number of trades}} \times 100$. Win rate indicates decision consistency and accuracy, though high win rates do not guarantee profitability if losses exceed gains on losing trades.

Number of Trades: Total trading frequency over the evaluation period, revealing strategic approaches from active short-term strategies to patient conviction-driven approaches. Trade frequency provides insight into system behavior and capital allocation patterns.

Annualized Sharpe Ratio (Ann. SR): Annualized version of the Sharpe Ratio calculated as: $\text{Ann. SR} = \text{SR} \times \sqrt{252}$, where 252 represents the typical number of trading days in a year. This metric provides standardized comparison across different evaluation periods and enables assessment of long-term risk-adjusted performance potential.

Sortino Ratio: Downside risk-adjusted return metric calculated as: $\text{Sortino} = \frac{\mu - r_f}{\sigma_d}$, where σ_d is the standard deviation of negative returns only. Unlike Sharpe Ratio, this metric exclusively penalizes downside volatility, providing a more nuanced view of risk-adjusted performance that distinguishes between harmful volatility and beneficial upside variance.

Return on Invested Capital (ROIC): Measures efficiency of capital utilization calculated as: $\text{ROIC} = \frac{\text{Net trading profit}}{\text{Total Capital Used for Entries}} \times 100$. This metric accounts for varying position sizes and reveals how effectively the system allocates available capital, indicating whether superior returns result from efficient capital deployment or simply from taking larger positions.

Profit per Trade (P/T): Average profit generated per trade calculated as: $\text{P/T} = \frac{\text{Total net profit}}{\text{Number of trades}}$. This metric indicates the average value creation per trading decision, helping assess decision quality and revealing whether profitability stems from frequent small gains or occasional large profits.

These comprehensive metrics provide multi-dimensional performance assessment that prevents gaming of individual measures while revealing the full spectrum of trading system capabilities. The combination enables identification of systems that achieve sustainable, risk-adjusted performance versus those that may appear successful on isolated metrics but exhibit fundamental weaknesses in risk management or capital efficiency.

6.5.3 Baseline Strategy Comparison

We compare ATLAS performance against established trading strategies that require no machine learning, providing context for where LLM approaches add value versus simpler alternatives:

Buy and Hold The Buy and Hold strategy is a passive investment approach in which an asset is acquired at the beginning of the investment horizon and retained without any further trading actions, regardless of interim price fluctuations. This method assumes that, over time, the market tends to grow, and thus long-term holding can yield positive returns. It does not rely on any predictive model or technical indicator. In our evaluation, Buy and Hold serves as a benchmark strategy against which the performance of all other trading methods is compared.

Simple Moving Average (SMA) The SMA strategy [22] issues trading signals based on the relationship between the current price of an asset and its moving average over a fixed time window. Specifically, a buy (sell) signal is triggered when the price crosses above (below) the SMA.

Short-Long Moving Average (SLMA) The SLMA method [66] extends the SMA approach by employing two SMAs of different lengths: one short-term and one long-term. A buy signal is generated when the short-term average crosses above the long-term average, while a sell signal occurs at the inverse crossover.

Moving Average Convergence The MACD strategy [66] captures momentum shifts by computing the difference between the 12-day and 26-day exponential moving averages. A 9-day EMA of the MACD line is used as a signal line. Trading signals are generated when the MACD line crosses the signal line from below (buy) or from above (sell). The exponential formulation ensures increased sensitivity to recent price movements.

Bollinger Bands The Bollinger Bands strategy [11] incorporates volatility by constructing a band around a 20-day SMA, with the upper and lower bands placed two standard deviations above and below the mean, respectively. A price crossing above the upper band may indicate overbought conditions (sell signal), while

crossing below the lower band may suggest oversold conditions (buy signal). We adopt the standard parameterization of 20-day SMA and multiplier 2, as commonly suggested in the literature.

For window-based strategies, we tested multiple window lengths across different market regimes. Since no single configuration performed consistently well across all market conditions, we report results from one of the best-performing configurations for each strategy as an indicative benchmark (10-day SMA and 10/30-day SLMA crossover).

These baseline strategies represent well-established quantitative approaches that provide meaningful comparison points for assessing LLM value proposition in trading applications.

6.6 Ablation Study Design

To quantify the contribution of each specialized agent within the ATLAS framework, we conduct systematic ablation studies using the best-performing model-strategy combination identified through our primary evaluation.

6.6.1 Configuration Selection

Ablation studies employ **GPT-o4-mini with Adaptive-OPRO** based on consistently strong performance across all evaluation assets and favorable cost-performance characteristics for repeated experimental runs. This configuration represents the optimal system variant, enabling assessment of component contributions under best-case conditions rather than suboptimal baselines.

6.6.2 Systematic Component Removal

No Market Analyst: Removes all technical analysis capabilities to assess whether multi-timescale price structure analysis and technical indicators meaningfully improve decision quality. This ablation tests whether pure sentiment and fundamental analysis suffice for effective trading or whether technical context provides essential decision-making support.

No News Analyst: Eliminates unstructured text analysis to evaluate the value of sentiment extraction and event-driven catalyst identification. This configuration tests whether market and fundamental data alone provide sufficient information for trading decisions or whether news processing adds crucial context for decision-making.

No Market & News Analysts: Reduces the system to single-agent configuration with only fundamental data inputs plus portfolio state, testing whether task decomposition and specialization provide genuine benefits or merely organizational convenience. This minimal configuration reveals whether multi-agent coordination is essential for performance or whether concentrated decision-making achieves comparable results.

6.6.3 Excluded Components

We exclude the Fundamental Analyst from systematic ablation due to its low activation frequency (typically 1-2 times per evaluation period), which renders short-term impact measurement statistically unreliable. The Fundamental Analyst’s contribution is instead assessed through qualitative analysis of trading decisions during earnings periods and corporate events, where it provides structural context for longer-term positioning.

The low-frequency activation reflects realistic fundamental data release patterns, where quarterly earnings and annual reports occur at predictable intervals. While important for strategic context, the infrequent activation prevents meaningful statistical measurement in ablation studies spanning two-month periods.

6.6.4 Evaluation Protocol

Each ablation configuration undergoes three independent runs using identical methodology to primary experiments, enabling statistical comparison of performance differences. The systematic approach reveals which components provide essential versus auxiliary contributions to overall system effectiveness.

Chapter 7

Results and Analysis

This chapter presents comprehensive results from our systematic evaluation of ATLAS across diverse market conditions, LLM architectures, and prompting strategies. Our key findings demonstrate that Adaptive-OPRO consistently achieves superior performance compared to baseline and reflection-based approaches across the vast majority of tested model architectures and market regimes. ATLAS exhibits remarkable robustness, delivering positive returns even in challenging bearish and volatile conditions where traditional strategies fail. Multi-agent specialization provides meaningful performance contributions, with systematic ablation studies revealing the value of task decomposition in complex trading environments. Different LLM architectures exhibit distinct trading behaviors and optimization capabilities, with clear performance hierarchies emerging that correlate with general model capabilities. Importantly, our multi-run evaluation protocols expose significant performance variance that single-run assessments completely miss, highlighting critical reliability issues in current LLM evaluation practices.

The results are organized around five key areas of analysis. First, we examine the comparative performance of different prompting strategies across market conditions and model architectures. Second, we present ablation studies that quantify the contribution of each specialized agent within the ATLAS framework. Third, we investigate behavioral patterns across different LLM architectures, uncovering distinct trading philosophies and adaptation capabilities. Fourth, we provide extended performance analysis that incorporates additional risk-adjusted metrics and explores additional prompting configurations including daily reflection mechanisms and combined optimization approaches, offering comprehensive confirmation of our primary findings. Finally, we examine transparent optimization traces produced by Adaptive-OPRO through detailed examples, demonstrating how systematic prompt refinement drives performance improvements and providing crucial mechanistic insights into the optimization process.

Contents

7.1 Framework Resilience Across Market Dynamics	58
7.1.1 Optimization in Sequential Decision-Making	58
7.2 Agent Contribution Analysis	59
7.3 Trading Behavior Across LLMs	60
7.3.1 Claude Models: Contrasting Failure Modes	61
7.3.2 LLM Optimization Capabilities	62
7.4 Extended Performance Analysis and Strategic Insights	63
7.4.1 Risk-Adjusted Performance Validation	63
7.4.2 The Reflection Paradox Reinforced	64
7.4.3 Architectural Performance Patterns	64
7.4.4 Extended Prompting Strategy Analysis	65
7.5 Prompt Evolution Mechanism Analysis	67
7.5.1 Systematic Weakness Detection and Resolution	67
7.5.2 Progressive Prompt Evolution: From Generic Foundation to Optimized Performance	73

Model	Prompting	ROI (%) $\uparrow$	SR $\uparrow$	DD (%) $\downarrow$	Win Rate (%) $\uparrow$	Num Trades
Non-LLM-Based Strategies						
Buy & Hold	N/A	-8.59	-0.071	20.45	0.00	1
MACD	N/A	6.50	0.131	6.86	0.00	1
SMA	N/A	6.91	0.177	3.56	50.00	4
SLMA	N/A	-1.87	-0.078	6.89	0.00	1
Bollinger Bands	N/A	0.00	0.000	0.00	0.00	0
LLM-Based Strategies - ATLAS						
LLaMA 3.3-70B	Baseline	-9.19 $\pm$ 1.54	-0.091 $\pm$ 0.021	16.90 $\pm$ 0.82	30.28 $\pm$ 11.87	22.67 $\pm$ 8.39
	Reflection	-8.44 $\pm$ 1.58	-0.087 $\pm$ 0.025	16.36 $\pm$ 0.31	44.69 $\pm$ 13.25	27.67 $\pm$ 1.15
	Adaptive-OPRO	-6.16$\pm$ 2.08	-0.066$\pm$ 0.004	14.05$\pm$ 3.33	54.36$\pm$ 12.44	28.33 $\pm$ 3.21
Claude Sonnet 4	Baseline	-7.26 $\pm$ 2.99	-0.066 $\pm$ 0.030	17.59 $\pm$ 1.55	31.19 $\pm$ 7.84	13.00 $\pm$ 4.36
	Reflection	-5.69 $\pm$ 1.82	-0.058 $\pm$ 0.013	15.12 $\pm$ 3.26	46.67$\pm$ 5.77	12.67 $\pm$ 2.08
	Adaptive-OPRO	0.35$\pm$ 1.78	0.008$\pm$ 0.018	14.76$\pm$ 2.87	43.45 $\pm$ 6.27	15.00 $\pm$ 2.00
Claude Sonnet 4 w/ Thinking	Baseline	-4.46 $\pm$ 4.76	-0.043 $\pm$ 0.048	14.32 $\pm$ 4.12	11.11 $\pm$ 19.24	14.00 $\pm$ 2.65
	Reflection	-8.60 $\pm$ 0.59	-0.078 $\pm$ 0.004	19.45 $\pm$ 1.65	14.29 $\pm$ 24.75	11.67 $\pm$ 2.08
	Adaptive-OPRO	-0.73$\pm$ 3.82	-0.004$\pm$ 0.038	12.94$\pm$ 2.32	43.89$\pm$ 21.11	17.00 $\pm$ 5.00
GPT-o4-mini	Baseline	-1.30 $\pm$ 1.71	-0.017 $\pm$ 0.017	9.68$\pm$ 3.12	29.17 $\pm$ 11.02	15.33 $\pm$ 3.06
	Reflection	-2.52 $\pm$ 4.03	-0.039 $\pm$ 0.045	9.82 $\pm$ 3.43	51.28 $\pm$ 5.06	20.33 $\pm$ 3.06
	Adaptive-OPRO	9.06$\pm$ 0.73	0.094$\pm$ 0.008	11.48 $\pm$ 0.00	65.28$\pm$ 16.84	17.33 $\pm$ 5.86
GPT-o3	Baseline	-6.11 $\pm$ 3.42	-0.080 $\pm$ 0.029	11.58 $\pm$ 3.09	42.59 $\pm$ 8.49	18.67 $\pm$ 3.21
	Reflection	-4.60 $\pm$ 3.40	-0.053 $\pm$ 0.044	12.11 $\pm$ 1.27	46.03 $\pm$ 16.88	18.33 $\pm$ 2.52
	Adaptive-OPRO	9.02$\pm$ 3.28	0.146$\pm$ 0.048	5.33$\pm$ 0.14	72.81$\pm$ 17.27	19.67 $\pm$ 4.16

Table 7.1: Performance comparison between non-LLM-based and LLM-based approaches using ATLAS in volatile, declining market conditions (LLY, healthcare sector). **Bold** values indicate the best per model.

7.1 Framework Resilience Across Market Dynamics

Tables 7.1, 7.2, 7.3 present a comparative evaluation of baseline non-LLM strategies against ATLAS across various LLM configurations and market regimes. The results demonstrate that ATLAS consistently achieves high performance across all tested conditions. To our knowledge, this is the first instance of a single framework exhibiting such systematic robustness across diverse market scenarios, clearly outperforming widely used and traditionally favored methods. Strategies like Buy-and-Hold perform well in bullish regimes but fail to generalize, producing significantly weaker results in range-bound and bearish markets-conditions typically marked by instability, low predictability, and limited informational signals. Notably, ATLAS, particularly when paired with GPT-o3 or GPT-o4-mini, delivers stable and positive returns even under adverse conditions, including bearish regimes, where the overall market trend is downward and profit generation is particularly challenging. This ability to perform reliably when most strategies struggle underscores the framework’s strength in navigating uncertainty and making effective strategic decisions, even in declining or highly volatile environments. These findings highlight ATLAS’s robustness and its potential as a dependable decision-making system across the full spectrum of market dynamics.

7.1.1 Optimization in Sequential Decision-Making

Adaptive-OPRO significantly outperforms both static baseline prompts and reflection-based approaches across the vast majority of tested models and market conditions (Tables 7.1 - 7.3). The improved decision-making throughout the optimization process clearly translates into more effective trading performance, as demonstrated by the metrics analyzed.

Return improvements demonstrate successful adaptation to market feedback. In the volatile bearish regime (Table 7.1), models like GPT-o3/o4-mini shift from negative baselines to substantial positive returns

Model	Prompting	ROI (%) $\uparrow$	SR $\uparrow$	DD (%) $\downarrow$	Win Rate (%) $\uparrow$	Num Trades
Non-LLM-Based Strategies						
Buy & Hold	N/A	1.14	0.013	6.97	0.00	1
MACD	N/A	-0.26	-0.019	5.90	0.00	3
SMA (50-day)	N/A	-0.13	-0.019	5.57	0.00	3
SLMA (20/50)	N/A	-1.12	-0.043	5.28	0.00	2
Bollinger Bands	N/A	0.00	0.000	0.00	0.00	0
LLM-Based Strategies						
Llama 3.3 70B	Baseline	-0.42$\pm$ 2.06	-0.024$\pm$ 0.051	5.56 $\pm$ 1.08	53.48$\pm$ 9.56	26.00 $\pm$ 2.00
	Reflection	-2.61 $\pm$ 0.77	-0.083 $\pm$ 0.014	6.38 $\pm$ 0.72	46.63 $\pm$ 3.15	26.33 $\pm$ 6.51
	Adaptive-OPRO	-1.10 $\pm$ 0.44	-0.045 $\pm$ 0.012	5.15$\pm$ 0.71	50.00 $\pm$ 3.85	25.33 $\pm$ 1.15
Claude Sonnet 4	Baseline	-4.49 $\pm$ 4.22	-0.134 $\pm$ 0.114	7.71$\pm$ 1.06	37.50$\pm$ 4.17	19.00 $\pm$ 3.46
	Reflection	-3.78$\pm$ 4.23	-0.115$\pm$ 0.105	10.54 $\pm$ 1.58	23.84 $\pm$ 8.27	18.00 $\pm$ 6.93
	Adaptive-OPRO	-5.07 $\pm$ 4.53	-0.165 $\pm$ 0.143	9.23 $\pm$ 2.71	31.02 $\pm$ 7.90	18.33 $\pm$ 2.52
Claude Sonnet 4 w/ Thinking	Baseline	-0.99$\pm$ 0.80	-0.039$\pm$ 0.020	7.75 $\pm$ 1.00	56.28$\pm$ 1.50	17.00 $\pm$ 5.20
	Reflection	-1.49 $\pm$ 3.76	-0.069 $\pm$ 0.123	7.27 $\pm$ 2.26	45.11 $\pm$ 12.6	17.00 $\pm$ 5.57
	Adaptive-OPRO	-1.01 $\pm$ 0.90	-0.046 $\pm$ 0.020	5.16$\pm$ 0.52	36.2 $\pm$ 24.47	16.33 $\pm$ 2.08
GPT-o4-mini	Baseline	1.29 $\pm$ 1.38	0.021 $\pm$ 0.044	3.23$\pm$ 0.48	39.01 $\pm$ 3.61	22.67 $\pm$ 7.57
	Reflection	-1.48 $\pm$ 0.54	-0.087 $\pm$ 0.018	4.64 $\pm$ 0.75	32.62 $\pm$ 7.49	27.33 $\pm$ 3.06
	Adaptive-OPRO	3.88$\pm$ 2.21	0.089$\pm$ 0.067	3.28 $\pm$ 0.95	47.95$\pm$ 7.15	25.33 $\pm$ 5.03
GPT o3	Baseline	-0.60 $\pm$ 1.71	-0.034 $\pm$ 0.050	5.93 $\pm$ 1.33	60.74 $\pm$ 5.59	16.33 $\pm$ 2.52
	Reflection	-1.55 $\pm$ 2.09	-0.084 $\pm$ 0.075	5.02 $\pm$ 0.72	42.50 $\pm$ 6.61	16.67 $\pm$ 0.58
	Adaptive-OPRO	3.62$\pm$ 0.90	0.096$\pm$ 0.027	3.46$\pm$ 0.48	71.93$\pm$ 15.9	16.00 $\pm$ 2.65

Table 7.2: Performance comparison between non-LLM-based and LLM-based approaches using ATLAS in range-bound market conditions (XOM, energy sector). **Bold** values indicate the best per model.

under Adaptive-OPRO.

More revealing are the **risk-adjusted metrics**: improved Sharpe ratios indicate that gains stem from genuine strategic enhancement rather than increased risk-taking.

Win rate patterns expose the optimization mechanism’s impact on decision quality. Models under Adaptive-OPRO generally achieve higher win rates alongside better returns, suggesting more consistent decision-making rather than occasional large gains masking frequent losses.

Significantly, the **reflection paradox** emerges consistently across models and regimes: reflection-based approaches not only fail to match Adaptive-OPRO but often underperform baseline prompts, challenging conventional assumptions about additional reasoning steps in well-optimized systems within dynamic environments.

7.2 Agent Contribution Analysis

Table 7.4 confirms each specialized agent’s distinct contribution by showing performance drops when each is removed.

Market Analyst is a core component across all market regimes. Its removal consistently results in the most significant performance degradation, especially in challenging conditions such as the bearish regime, where technical context is crucial for decision-making. In the sideways regime, the absence of market analysis not only reduces returns but also lowers trading frequency, suggesting that agents lose the confidence to act without a solid technical foundation. Notably, in bullish markets, ROI slightly improves when market data are excluded, suggesting that in up-trending markets social consensus and market news may offer cleaner entry signals.

Model	Prompting	ROI (%) $\uparrow$	SR $\uparrow$	DD (%) $\downarrow$	Win Rate (%) $\uparrow$	Num Trades
Non-LLM-Based Strategies						
Buy & Hold	N/A	41.30	0.409	3.16	0.00	1
MACD	N/A	-0.62	-0.343	0.62	0.00	1
SMA)	N/A	36.77	0.384	3.12	0.00	1
SLMA	N/A	15.88	0.254	2.98	0.00	1
Bollinger Bands	N/A	0.00	0.000	0.00	0.00	0
LLM-Based Strategies - ATLAS						
Llama 3.3 70B	Baseline	37.86 $\pm$ 12.31	0.388 $\pm$ 0.096	3.46 $\pm$ 0.63	20.37 $\pm$ 35.28	13.00 $\pm$ 20.78
	Reflection	40.40 $\pm$ 1.43	0.422$\pm$ 0.023	2.96$\pm$ 0.34	33.33 $\pm$ 57.74	5.33 $\pm$ 6.66
	Adaptive-OPRO	42.07$\pm$ 1.85	0.418 $\pm$ 0.016	3.15 $\pm$ 0.02	100.00$\pm$ 0.00	1.33 $\pm$ 0.58
Claude Sonnet 4	Baseline	13.43 $\pm$ 8.62	0.180 $\pm$ 0.121	5.52 $\pm$ 3.96	60.83$\pm$ 12.30	21.67 $\pm$ 9.50
	Reflection	5.21 $\pm$ 1.10	0.089 $\pm$ 0.026	5.11 $\pm$ 1.86	39.25 $\pm$ 15.79	22.33 $\pm$ 1.53
	Adaptive-OPRO	25.85$\pm$ 10.61	0.290$\pm$ 0.087	3.75$\pm$ 0.59	43.81 $\pm$ 38.37	19.00 $\pm$ 12.17
Claude Sonnet 4 w/ Thinking	Baseline	12.52 $\pm$ 2.47	0.175 $\pm$ 0.030	5.03 $\pm$ 1.53	53.30 $\pm$ 14.47	17.00 $\pm$ 2.65
	Reflection	11.12 $\pm$ 4.86	0.186 $\pm$ 0.083	3.42$\pm$ 2.23	77.86$\pm$ 2.58	17.00 $\pm$ 5.00
	Adaptive-OPRO	16.36$\pm$ 7.87	0.217$\pm$ 0.105	5.18 $\pm$ 2.52	68.89 $\pm$ 30.06	12.67 $\pm$ 4.04
GPT-o4-mini	Baseline	7.00 $\pm$ 3.46	0.125 $\pm$ 0.054	2.74 $\pm$ 0.79	46.29 $\pm$ 3.21	18.67 $\pm$ 1.53
	Reflection	9.80 $\pm$ 3.21	0.189 $\pm$ 0.067	2.45$\pm$ 1.00	54.54 $\pm$ 7.92	26.33 $\pm$ 9.61
	Adaptive-OPRO	10.47$\pm$ 3.84	0.193$\pm$ 0.046	3.42 $\pm$ 0.90	62.70$\pm$ 11.25	20.33 $\pm$ 2.89
GPT o3	Baseline	22.70 $\pm$ 0.92	0.269 $\pm$ 0.029	6.82 $\pm$ 3.03	66.67 $\pm$ 28.87	7.33 $\pm$ 2.52
	Reflection	21.98 $\pm$ 4.54	0.325 $\pm$ 0.040	3.14 $\pm$ 0.99	96.67 $\pm$ 5.77	18.00 $\pm$ 3.61
	Adaptive-OPRO	25.06$\pm$ 4.28	0.392$\pm$ 0.019	2.31$\pm$ 0.80	100.00$\pm$ 0.00	9.67 $\pm$ 4.04

Table 7.3: Performance comparison between non-LLM-based and LLM-based approaches using ATLAS in rising market conditions (NVDA, technology sector). **Bold** values indicate the best per model.

News analyst contributes regime-specific strategic value. In the bullish regime, news removal leads to lower returns as agents become more conservative, missing chances to capitalize on positive momentum. The sideways regime shows news analysis as particularly critical, with its removal producing severe degradation—suggesting that sentiment analysis is essential when technical signals are ambiguous.

Combination of News & Market Analyst offers insights into their interdependent value: across all regimes, removing both agents leads to substantial performance degradation, indicating that news and market signals offer complementary, non-redundant information. In the bearish regime, performance drops significantly, reflecting the importance of sentiment and technical context under volatility. In the sideways regime, the absence of both leads to unstable and unprofitable behavior. Even in the bullish regime, where market data alone may be less essential, combined removal clearly harms performance. These results suggest that each component contributes differently across regimes, with the combined removal producing regime-specific effects that differ from simple additive impacts.

7.3 Trading Behavior Across LLMs

Our analysis reveals a clear correlation between the general capabilities of the models and their trading performance, with more capable models performing significantly better than their less capable counterparts (Figure 7.2.1).

GPT-o3 demonstrates the most advanced market understanding, systematically leveraging analysis from all specialized agents to formulate coherent strategies. While occasionally overly risk-averse (limiting gains in favorable conditions), this conservative bias enables consistent cross-regime performance. The model’s prompt optimization capabilities are exemplary, showing incremental learning, strategic adaptation to failures, and clear understanding of optimization objectives.

Stock	Configuration	ROI (%) $\uparrow$	Sharpe Ratio $\uparrow$	Max DD (%) $\downarrow$	Win Rate (%) $\uparrow$	Num Trades
LLY (Bearish Regime)	No News	4.07 ± 0.72	0.056 ± 0.016	7.84 ± 3.15	53.51 ± 6.67	25.33 ± 4.51
	No Market Data	-5.75 ± 0.76	-0.094 ± 0.017	11.32 ± 2.63	37.52 ± 4.87	18.33 ± 3.06
	No News + No Market	-6.86 ± 1.68	-0.078 ± 0.036	14.54 ± 3.30	43.94 ± 6.94	22.33 ± 1.15
	ATLAS	9.06 ± 0.73	0.094 ± 0.008	11.48 ± 0.00	65.28 ± 16.84	17.33 ± 5.86
XOM (Sideways Regime)	No News	-8.20 ± 1.64	-0.264 ± 0.069	9.09 ± 2.99	22.82 ± 13.65	35.00 ± 12.29
	No Market Data	0.01 ± 0.92	-0.011 ± 0.021	6.56 ± 1.58	46.55 ± 23.15	13.33 ± 3.06
	No News + No Market	-4.60 ± 0.70	-0.136 ± 0.026	7.01 ± 2.29	35.26 ± 13.09	21.00 ± 4.58
	ATLAS	3.88 ± 2.21	0.089 ± 0.067	3.28 ± 0.95	47.95 ± 7.15	25.33 ± 5.03
NVDA (Bullish Regime)	No News	6.62 ± 0.25	0.090 ± 0.008	6.67 ± 0.36	41.96 ± 5.21	28.33 ± 4.62
	No Market Data	11.78 ± 1.76	0.216 ± 0.024	3.70 ± 0.86	70.24 ± 14.03	20.00 ± 5.57
	No News + No Market	7.34 ± 2.79	0.110 ± 0.012	5.76 ± 2.01	63.84 ± 9.39	20.67 ± 1.53
	ATLAS	10.47 ± 3.84	0.193 ± 0.046	3.42 ± 0.90	62.70 ± 11.25	20.33 ± 2.89

Table 7.4: Ablation study results showing individual agent contributions using GPT-o4-mini across three market regimes.

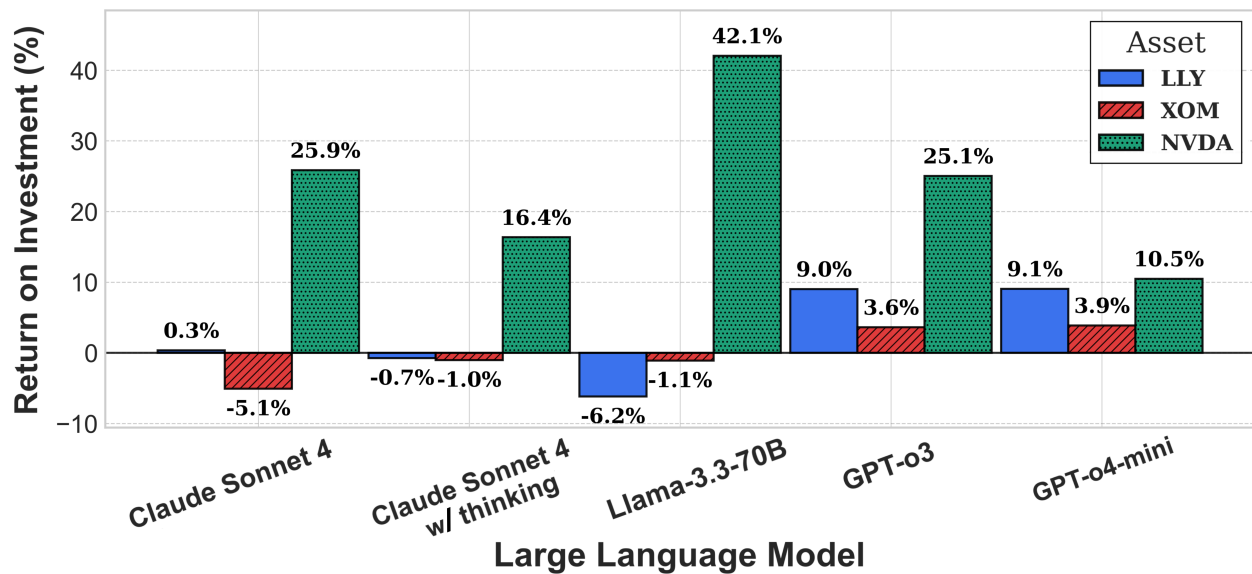


Figure 7.2.1: ROI across three assets using Adaptive-OPRO.

In contrast, GPT-o4-mini demonstrates competent but short-sighted decision-making, prioritizing short-term risk control over strategic positioning. It relies on tight stop-losses and early profit-taking - effective in volatile markets but weak in sustained trends. While it tends to overtrade in some conditions, its prompt optimization is on par with GPT-o3 in sophistication.

At the other end of the spectrum, LLaMA 3.3-70B operates with remarkably primitive strategies, lacking risk management abilities or coherent goal-setting mechanisms. It acts belatedly to market changes and exhibits abrupt strategy shifts under volatility. Paradoxically, this strategic simplicity becomes advantageous in straightforward bullish regimes, where it effectively mimics buy-and-hold approaches to achieve exceptional returns, suggesting that complex reasoning can sometimes hinder performance in simple settings.

7.3.1 Claude Models: Contrasting Failure Modes

Claude models demonstrate systematic performance issues across both reasoning configurations, revealing distinct failure modes that highlight fundamental limitations in the architecture. Claude with thinking exhibits systematic overthinking, where extensive analysis of news or fundamental data leads to misinterpretation of technical signals. While the reasoning process appears sophisticated and logically presented, it contains

Model	Prompting	Ann. SR $\uparrow$	Sortino $\uparrow$	ROIC (%) $\uparrow$	P/T (\$) $\uparrow$
LLM-Based Strategies - ATLAS					
LLaMA 3.3-70B	Baseline	-1.45 $\pm$ 0.33	-0.09 $\pm$ 0.02	-1.01 $\pm$ 0.48	-1070.14 $\pm$ 634.06
	Reflection	-1.38 $\pm$ 0.39	-0.08 $\pm$ 0.02	-0.68 $\pm$ 0.20	-647.13 $\pm$ 141.63
	Adaptive-OPRO	-1.05$\pm$ 0.06	-0.06	-0.47$\pm$ 0.19	-472.27$\pm$ 174.19
Claude Sonnet 4	Baseline	-1.04 $\pm$ 0.48	-0.06 $\pm$ 0.03	-2.83 $\pm$ 1.13	-1920.19 $\pm$ 323.80
	Reflection	-0.91 $\pm$ 0.21	-0.05 $\pm$ 0.01	-2.66 $\pm$ 1.47	-1206.60 $\pm$ 745.08
	Adaptive-OPRO	0.12$\pm$ 0.28	0.01$\pm$ 0.02	0.00$\pm$ 0.27	-144.52$\pm$ 136.78
Claude Sonnet 4 w/ Thinking	Baseline	-0.68 $\pm$ 0.77	-0.04 $\pm$ 0.04	-2.65 $\pm$ 2.53	-2084.43 $\pm$ 2197.78
	Reflection	-1.23 $\pm$ 0.06	-0.08	-5.21 $\pm$ 1.72	-2407.54 $\pm$ 1345.56
	Adaptive-OPRO	-0.06$\pm$ 0.61	-0.00$\pm$ 0.04	-0.35$\pm$ 0.92	-278.10$\pm$ 725.32
GPT-o4-mini	Baseline	-0.26 $\pm$ 0.27	-0.02 $\pm$ 0.02	-0.18 $\pm$ 0.22	-168.13 $\pm$ 209.76
	Reflection	-0.61 $\pm$ 0.71	-0.04 $\pm$ 0.04	-0.48 $\pm$ 0.72	-287.24 $\pm$ 328.38
	Adaptive-OPRO	1.49$\pm$ 0.12	0.09$\pm$ 0.01	1.12$\pm$ 0.34	1056.49$\pm$ 297.92
GPT-o3	Baseline	-1.27 $\pm$ 0.45	-0.08 $\pm$ 0.02	-1.67 $\pm$ 1.03	-792.65 $\pm$ 279.17
	Reflection	-0.84 $\pm$ 0.70	-0.05 $\pm$ 0.04	-0.90 $\pm$ 0.73	-497.41 $\pm$ 337.21
	Adaptive-OPRO	2.32$\pm$ 0.76	0.16$\pm$ 0.07	1.98$\pm$ 0.84	799.30$\pm$ 242.46

Table 7.5: Additional performance metrics for LLY (healthcare sector) comparing LLM-based approaches using ATLAS in volatile, declining market conditions. Ann. SR = Annualized Sharpe Ratio, ROIC = Return on Invested Capital, P/T = Profit per Trade. **Bold** values indicate the best per model.

fundamental gaps that result in poor decisions.

Claude without thinking demonstrates even more severe reliability issues: overly reactive positioning, persistent hallucinations about performance metrics, and fundamental misunderstanding of market mechanics. The model’s extreme variance patterns and erratic behavior render it unsuitable for reliable evaluation, highlighting how insufficient reasoning capabilities can lead to complete system failure. These contrasting failure modes suggest that Claude’s underlying architecture may be fundamentally misaligned with financial decision-making, regardless of reasoning enhancement.

7.3.2 LLM Optimization Capabilities

From an explainability perspective, one key advantage of Adaptive-OPRO is that the optimization process produces a final instruction (prompt) for the models, which can be evaluated both in terms of its correctness and whether the models follow it successfully. This dual perspective enables us to assess not only the alignment of the optimized prompt with the intended objective but also the model’s capacity to interpret and act upon it. In doing so, we can determine whether the optimization is heading in the right direction.

This analysis showcases that prompt optimization capabilities vary across LLMs: GPT models show consistent, interpretable prompt refinements, while LLaMA often hallucinates edits, reporting changes that are not present. Claude tends toward overly rigid, procedural instructions that limit flexibility.

On the other hand, reflection, while promising in theory, introduces noise in practice. Even GPT-o3 can become paralyzed by over-analysis, while less capable models either generate vague reflections (LLaMA) or confidently misinterpret market signals.

Model	Prompting	Ann. SR $\uparrow$	Sortino $\uparrow$	ROIC (%) $\uparrow$	P/T (\$) $\uparrow$
LLM-Based Strategies - ATLAS					
LLaMA 3.3-70B	Baseline	-0.38 $\pm$ 0.81	-0.02 $\pm$ 0.06	-0.03 $\pm$ 0.16	-26.23 $\pm$ 164.36
	Reflection	-1.32 $\pm$ 0.21	-0.10 $\pm$ 0.01	-0.21 $\pm$ 0.07	-227.29 $\pm$ 38.58
	Adaptive-OPRO	-0.72 $\pm$ 0.19	-0.06 $\pm$ 0.02	-0.09 $\pm$ 0.03	-86.11 $\pm$ 31.28
Claude Sonnet 4	Baseline	-2.13 $\pm$ 1.81	-0.17 $\pm$ 0.13	-0.54 $\pm$ 0.56	-522.11 $\pm$ 353.17
	Reflection	-1.82 $\pm$ 1.67	-0.14 $\pm$ 0.13	-0.37 $\pm$ 0.46	-313.67 $\pm$ 414.48
	Adaptive-OPRO	-2.62 $\pm$ 2.27	-0.20 $\pm$ 0.17	-0.80 $\pm$ 0.48	-576.65 $\pm$ 491.70
Claude Sonnet 4 w/ Thinking	Baseline	-0.63 $\pm$ 0.32	-0.04 $\pm$ 0.02	-0.12 $\pm$ 0.10	-113.56 $\pm$ 89.87
	Reflection	-1.10 $\pm$ 1.94	-0.09 $\pm$ 0.16	-0.34 $\pm$ 0.85	-90.06 $\pm$ 311.40
	Adaptive-OPRO	-0.73 $\pm$ 0.32	-0.06 $\pm$ 0.02	-0.39 $\pm$ 0.35	-133.64 $\pm$ 113.58
GPT-o4-mini	Baseline	0.33 $\pm$ 0.69	0.04 $\pm$ 0.08	0.16 $\pm$ 0.21	155.33 $\pm$ 202.32
	Reflection	-1.38 $\pm$ 0.29	-0.14 $\pm$ 0.02	-0.17 $\pm$ 0.05	-132.49 $\pm$ 87.57
	Adaptive-OPRO	1.41 $\pm$ 1.06	0.16 $\pm$ 0.14	0.34 $\pm$ 0.26	340.47 $\pm$ 260.95
GPT-o3	Baseline	-0.54 $\pm$ 0.80	-0.04 $\pm$ 0.07	-0.10 $\pm$ 0.31	-64.90 $\pm$ 190.96
	Reflection	-1.33 $\pm$ 1.18	-0.10 $\pm$ 0.08	-0.43 $\pm$ 0.68	-187.25 $\pm$ 261.18
	Adaptive-OPRO	1.52 $\pm$ 0.43	0.15 $\pm$ 0.05	1.08 $\pm$ 0.72	380.06 $\pm$ 44.91

Table 7.6: Additional performance metrics for XOM (energy sector) comparing LLM-based approaches using ATLAS in stable market conditions. Ann. SR = Annualized Sharpe Ratio, ROIC = Return on Invested Capital, P/T = Profit per Trade. **Bold** values indicate the best per model.

7.4 Extended Performance Analysis and Strategic Insights

Tables 7.5, 7.6, and 7.7 present extended metrics including Annualized Sharpe Ratio, Sortino Ratio, Return on Invested Capital (ROIC), and Profit per Trade, providing deeper insights into risk-adjusted performance and strategic behavior patterns. These comprehensive extended metrics reveal deeper behavioral patterns that reinforce and expand our core findings about LLM trading capabilities and optimization effectiveness.

7.4.1 Risk-Adjusted Performance Validation

The extended risk-adjusted metrics strongly corroborate Adaptive-OPRO’s superiority while revealing the mechanisms underlying its effectiveness. **Sortino Ratio analysis**, which focuses exclusively on downside risk, demonstrates that Adaptive-OPRO’s performance gains stem from genuine risk management improvements rather than increased risk-taking. This finding is particularly significant because it rules out the possibility that optimization simply encourages more aggressive positioning to achieve higher returns.

The **Return on Invested Capital (ROIC)** patterns provide additional validation by showing that Adaptive-OPRO achieves superior capital efficiency across model architectures. This metric reveals that optimization doesn’t merely improve returns but enhances the fundamental effectiveness of capital allocation decisions. The consistency of this pattern across different market conditions suggests that Adaptive-OPRO addresses core decision-making weaknesses rather than exploiting specific market dynamics.

Profit per Trade analysis reveals interesting patterns in completed trading cycles, though this metric captures only closed positions (completed buy-sell cycles where profits/losses are realized) and thus provides a partial view of overall trading behavior. The observed variations in per-trade profitability under optimization, combined with consistently improved win rates, suggest that models may be shifting toward different trading approaches-potentially favoring more systematic position management over opportunistic large gains. However, since this metric excludes open positions and considers only completed buy-sell cycles, the full picture of behavioral changes requires consideration alongside other performance measures.

Model	Prompting	Ann. SR $\uparrow$	Sortino $\uparrow$	ROIC (%) $\uparrow$	P/T (\$) $\uparrow$
LLM-Based Strategies - ATLAS					
LLaMA 3.3-70B	Baseline	6.16 ± 1.52	0.97 ± 0.22	30.98 ± 26.06	456.27 ± 790.29
	Reflection	6.70 ± 0.37	1.03 ± 0.02	29.14 ± 21.06	1511.32 ± 2617.69
	Adaptive-OPRO	6.63 ± 0.25	1.05 ± 0.01	42.26 ± 1.68	0.00
Claude Sonnet 4	Baseline	2.86 ± 1.93	0.45 ± 0.33	2.82 ± 2.60	1212.88 ± 920.24
	Reflection	1.42 ± 0.41	0.16 ± 0.05	0.86 ± 0.36	416.79 ± 149.76
	Adaptive-OPRO	4.60 ± 1.38	0.68 ± 0.22	8.25 ± 9.83	371.70 ± 1779.64
Claude Sonnet 4 w/ Thinking	Baseline	2.78 ± 0.48	0.46 ± 0.20	3.27 ± 1.51	1246.39 ± 143.77
	Reflection	2.95 ± 1.32	0.57 ± 0.40	4.33 ± 1.72	1042.20 ± 424.00
	Adaptive-OPRO	3.45 ± 1.66	0.76 ± 0.56	5.44 ± 2.81	2402.02 ± 1239.52
GPT-o4-mini	Baseline	1.98 ± 0.86	0.27 ± 0.14	0.81 ± 0.39	212.27 ± 421.02
	Reflection	3.00 ± 1.06	0.47 ± 0.23	1.40 ± 0.70	537.97 ± 45.35
	Adaptive-OPRO	3.07 ± 0.73	0.41 ± 0.12	1.54 ± 0.47	506.75 ± 329.55
GPT-o3	Baseline	4.27 ± 0.47	0.61 ± 0.14	8.03 ± 1.86	4262.67 ± 897.79
	Reflection	5.16 ± 0.63	0.68 ± 0.20	6.76 ± 2.76	2192.28 ± 920.54
	Adaptive-OPRO	6.22 ± 0.30	1.22 ± 0.37	17.04 ± 7.65	3761.99 ± 749.07

Table 7.7: Additional performance metrics for NVDA (technology sector) comparing LLM-based approaches using ATLAS in bullish market conditions. Ann. SR = Annualized Sharpe Ratio, ROIC = Return on Invested Capital, P/T = Profit per Trade. **Bold** values indicate the best per model.

7.4.2 The Reflection Paradox Reinforced

The extended metrics provide additional evidence for the variable and generally suboptimal performance of reflection-based approaches. Across **Sortino Ratio** measurements, reflection shows mixed results compared to baseline approaches-sometimes achieving modest improvements but frequently exhibiting significant degradation. However, reflection consistently underperforms Adaptive-OPRO across the vast majority of configurations, reinforcing the core finding that additional reasoning steps often impair rather than improve trading effectiveness. **ROIC** analysis reveals similar patterns, with reflection-based strategies showing inconsistent capital efficiency across model architectures. While occasional configurations demonstrate competitive performance, the overall pattern suggests that reflection’s analytical complexity more often leads to suboptimal resource allocation than improved decision-making. This variability in reflection performance contrasts sharply with Adaptive-OPRO’s consistent improvements across multiple risk-adjusted measures.

7.4.3 Architectural Performance Patterns

Extended performance measures confirm distinct behavioral profiles across model architectures. GPT models demonstrate consistent performance patterns across risk-adjusted measures, with GPT-o3 showing exceptional stability in both downside risk management and capital efficiency. GPT-o4-mini exhibits competent performance but with notable overtrading tendencies and short-term focus, leading to active but sometimes inefficient capital utilization patterns.

LLaMA’s extended metrics reveal the superficial nature of its apparent success in favorable conditions-while achieving competitive returns, the model shows poor risk-adjusted performance and inconsistent capital allocation patterns. This disconnect between raw returns and risk-adjusted measures explains its failure to generalize across different market environments.

Claude models exhibit systematically poor performance across all extended dimensions, confirming fundamental architectural limitations rather than domain-specific adaptation challenges.

Model	Prompting	ROI (%) $\uparrow$	SR $\uparrow$	DD (%) $\downarrow$	Win Rate (%) $\uparrow$	Num Trades
LLM-Based Strategies - ATLAS						
LLaMA 3.3-70B	Reflection (1d)	-10.59 $\pm$ 4.89	-0.11 $\pm$ 0.06	16.37 $\pm$ 1.97	40.47 $\pm$ 8.25	27 $\pm$ 2.65
	Adaptive-OPRO w/Reflection (1d)	-5.03$\pm$ 0.99	-0.06$\pm$ 0.02	13.18$\pm$ 0.22	42.86 $\pm$ 7.15	26 $\pm$ 4.93
	Adaptive-OPRO	-6.16 $\pm$ 2.08	-0.07 $\pm$ 0.00	14.05 $\pm$ 3.33	54.36$\pm$ 12.44	28 $\pm$ 3.21
Claude Sonnet 4	Reflection (1d)	-2.98 $\pm$ 3.38	-0.04 $\pm$ 0.04	10.35$\pm$ 4.47	33.33 $\pm$ 11.55	14 $\pm$ 5.20
	Adaptive-OPRO w/Reflection (1d)	-4.68 $\pm$ 4.71	-0.06 $\pm$ 0.06	13.07 $\pm$ 3.68	26.19 $\pm$ 8.58	15 $\pm$ 2.65
	Adaptive-OPRO	0.35$\pm$ 1.78	0.01$\pm$ 0.02	14.76 $\pm$ 2.87	43.45$\pm$ 6.27	15 $\pm$ 2.00
Claude Sonnet 4 w/ Thinking	Reflection (1d)	-5.25 $\pm$ 2.34	-0.05 $\pm$ 0.01	15.35 $\pm$ 4.17	24.44 $\pm$ 21.43	13 $\pm$ 6.35
	Adaptive-OPRO w/Reflection (1d)	-2.07 $\pm$ 3.49	-0.03 $\pm$ 0.04	8.74$\pm$ 3.77	47.62$\pm$ 4.12	16 $\pm$ 2.52
	Adaptive-OPRO	-0.73$\pm$ 3.82	-0.00$\pm$ 0.04	12.94 $\pm$ 2.32	43.89 $\pm$ 21.11	17 $\pm$ 5.00
GPT-o4-mini	Reflection (1d)	-3.84 $\pm$ 2.93	-0.06 $\pm$ 0.04	9.61 $\pm$ 2.13	52.46 $\pm$ 2.50	32 $\pm$ 12.50
	Adaptive-OPRO w/Reflection (1d)	-1.25 $\pm$ 1.45	-0.04 $\pm$ 0.03	6.51$\pm$ 2.08	41.14 $\pm$ 15.35	27 $\pm$ 3.79
	Adaptive-OPRO	9.06$\pm$ 0.73	0.09$\pm$ 0.01	11.48	65.28$\pm$ 16.84	17 $\pm$ 5.86
GPT-o3	Reflection (1d)	0.14 $\pm$ 0.56	-0.01 $\pm$ 0.01	6.40 $\pm$ 1.07	73.81 $\pm$ 2.06	19 $\pm$ 3.79
	Adaptive-OPRO w/Reflection (1d)	8.05 $\pm$ 0.30	0.16$\pm$ 0.03	4.55$\pm$ 1.42	76.69$\pm$ 5.03	22 $\pm$ 5.69
	Adaptive-OPRO	9.02$\pm$ 3.28	0.15 $\pm$ 0.05	5.33 $\pm$ 0.14	72.81 $\pm$ 17.27	20 $\pm$ 4.16

Table 7.8: Performance comparison of extended prompting strategies for LLY (healthcare sector) using ATLAS in volatile, declining market conditions. **Bold** values indicate the best per model.

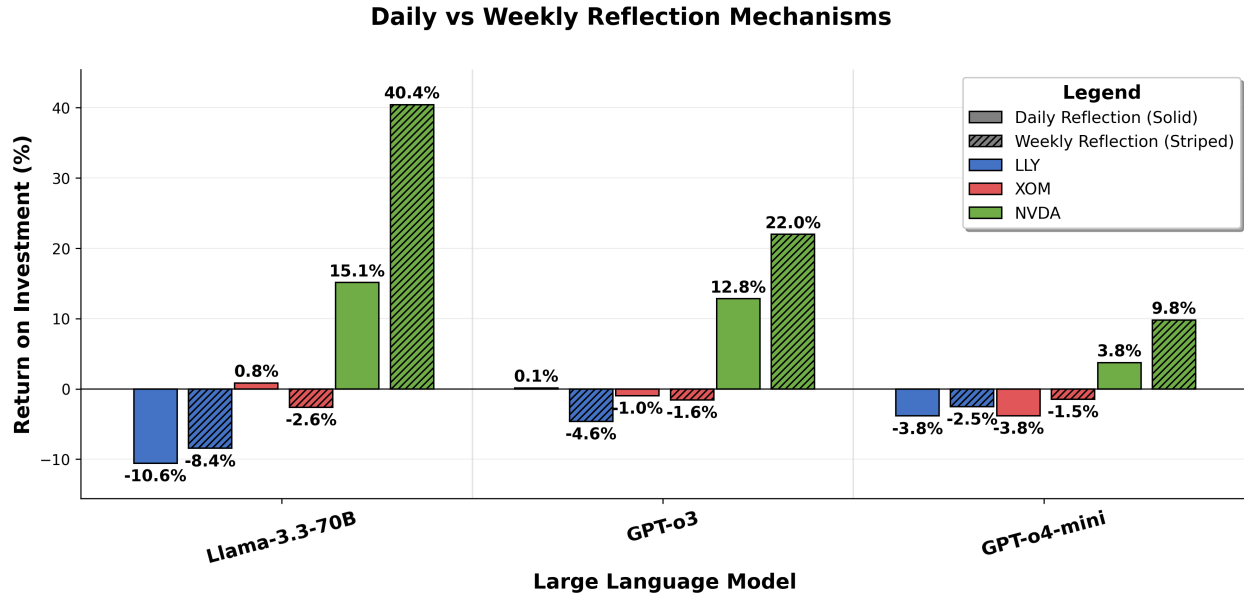


Figure 7.4.1: Daily vs weekly reflection mechanism performance comparison across models and assets, showing ROI percentages (solid = daily, striped = weekly).

7.4.4 Extended Prompting Strategy Analysis

Tables 7.8, 7.9, and 7.10 examine extended prompting configurations including daily reflection mechanisms and combined optimization approaches, providing insights into adaptation frequency effects and mechanism compatibility across different market regimes.

Model	Prompting	ROI (%) $\uparrow$	SR $\uparrow$	DD (%) $\downarrow$	Win Rate (%) $\uparrow$	Num Trades
LLM-Based Strategies - ATLAS						
LLaMA 3.3-70B	Reflection (1d)	0.82 $\pm$ 1.42	0.01 $\pm$ 0.02	1.62 $\pm$ 2.80	16.67 $\pm$ 28.87	8 $\pm$ 13.86
	Adaptive-OPRO w/Reflection (1d)	0.29 $\pm$ 0.50	0.00 $\pm$ 0.00	1.96 $\pm$ 3.39	16.67 $\pm$ 28.87	12 $\pm$ 20.78
	Adaptive-OPRO	-1.10 $\pm$ 0.44	-0.05 $\pm$ 0.01	5.15 $\pm$ 0.71	50.00 $\pm$ 3.85	25 $\pm$ 1.15
Claude Sonnet 4	Reflection (1d)	-3.76 $\pm$ 4.23	-0.10 $\pm$ 0.07	7.29 $\pm$ 3.08	48.81 $\pm$ 20.03	15 $\pm$ 6.08
	Adaptive-OPRO w/Reflection (1d)	-4.48 $\pm$ 3.85	-0.20 $\pm$ 0.16	7.16 $\pm$ 3.31	39.17 $\pm$ 20.05	14 $\pm$ 3.51
	Adaptive-OPRO	-5.07 $\pm$ 4.53	-0.16 $\pm$ 0.14	9.23 $\pm$ 2.71	31.02 $\pm$ 7.90	18 $\pm$ 2.52
Claude Sonnet 4 w/ Thinking	Reflection (1d)	2.40 $\pm$ 4.39	0.05 $\pm$ 0.14	4.57 $\pm$ 1.98	48.41 $\pm$ 42.35	14 $\pm$ 5.69
	Adaptive-OPRO w/Reflection (1d)	-2.84 $\pm$ 3.73	-0.12 $\pm$ 0.13	8.03 $\pm$ 0.89	22.62 $\pm$ 7.43	14 $\pm$ 1.53
	Adaptive-OPRO	-1.01 $\pm$ 0.90	-0.05 $\pm$ 0.02	5.16 $\pm$ 0.52	36.20 $\pm$ 24.47	16 $\pm$ 2.08
GPT-o4-mini	Reflection (1d)	-3.81 $\pm$ 2.13	-0.18 $\pm$ 0.06	6.54 $\pm$ 1.95	32.86 $\pm$ 8.84	38 $\pm$ 9.71
	Adaptive-OPRO w/Reflection (1d)	-1.43 $\pm$ 0.38	-0.09 $\pm$ 0.02	5.37 $\pm$ 3.26	41.45 $\pm$ 7.41	38 $\pm$ 5.29
	Adaptive-OPRO	3.88 $\pm$ 2.21	0.09 $\pm$ 0.07	3.28 $\pm$ 0.95	47.95 $\pm$ 7.15	25 $\pm$ 5.03
GPT-o3	Reflection (1d)	-0.97 $\pm$ 1.08	-0.11 $\pm$ 0.09	3.42 $\pm$ 0.58	48.21 $\pm$ 20.28	11 $\pm$ 2.65
	Adaptive-OPRO w/Reflection (1d)	-0.51 $\pm$ 0.76	-0.06 $\pm$ 0.03	2.71 $\pm$ 0.18	55.18 $\pm$ 16.43	17 $\pm$ 4.73
	Adaptive-OPRO	3.62 $\pm$ 0.90	0.10 $\pm$ 0.03	3.46 $\pm$ 0.48	71.93 $\pm$ 15.99	16 $\pm$ 2.65

Table 7.9: Performance comparison of extended prompting strategies for XOM (energy sector) using ATLAS in stable market conditions. **Bold** values indicate the best per model.

Adaptation Frequency Effects The comparison between daily and weekly reflection mechanisms reveals complex patterns that vary by both model architecture and market regime (Figure 7.4.1). Daily reflection shows mixed results across different configurations, with performance heavily dependent on the interaction between model capabilities and market conditions.

In stable, range-bound conditions (XOM), daily reflection demonstrates nuanced behavioral changes that can prove beneficial. For LLaMA, daily reflection significantly reduces trading activity due to the unclear market signals, leading to mitigated losses and even small positive returns compared to more aggressive weekly reflection approaches. Similarly, GPT-o3 shows enhanced downside protection under daily reflection in range-bound markets, though without substantial changes in trading frequency, suggesting improved risk assessment rather than reduced activity.

However, this conservative bias becomes problematic in trending markets. In bullish conditions (NVDA), LLaMA’s daily reflection restricts potential gains by encouraging overly cautious positioning that fails to capitalize on favorable momentum. The same pattern emerges across other models, where frequent self-evaluation appears to introduce hesitation that impairs the decisive action required to capture trending opportunities.

These regime-dependent patterns highlight the context-sensitive nature of reflection mechanisms, where effectiveness depends heavily on both model architecture and market conditions.

Mechanism Compatibility Assessment The Adaptive-OPRO with daily Reflection configurations reveal important insights about mechanism compatibility in sequential decision-making systems. While the combined approach typically outperforms the standalone reflection mechanism, it consistently shows deteriorated performance compared to pure Adaptive-OPRO across most configurations. This pattern demonstrates that reflection introduces noise into well-optimized systems, interfering with the systematic improvements achieved through prompt optimization.

The performance hierarchy—where Adaptive-OPRO alone exceeds combined approaches, which in turn exceed standalone daily reflection—illustrates the critical role of systematic optimization and reveals how additional

Model	Prompting	ROI (%) $\uparrow$	SR $\uparrow$	DD (%) $\downarrow$	Win Rate (%) $\uparrow$	Num Trades
LLM-Based Strategies - ATLAS						
LLaMA 3.3-70B	Reflection (1d)	15.12 $\pm$ 9.01	0.22 $\pm$ 0.11	3.42 $\pm$ 0.70	64.88 $\pm$ 9.16	16 $\pm$ 1.73
	Adaptive-OPRO w/Reflection (1d)	36.31 $\pm$ 6.20	0.40 $\pm$ 0.01	2.60$\pm$ 0.92	33.33 $\pm$ 57.74	2 $\pm$ 0.58
	Adaptive-OPRO	42.07$\pm$ 1.85	0.42$\pm$ 0.02	3.15 $\pm$ 0.02	100.00$\pm$ 0.00	1 $\pm$ 0.58
Claude Sonnet 4	Reflection (1d)	6.62 $\pm$ 2.64	0.11 $\pm$ 0.06	5.14 $\pm$ 2.91	48.48 $\pm$ 2.63	15 $\pm$ 5.13
	Adaptive-OPRO w/Reflection (1d)	24.60 $\pm$ 3.37	0.33$\pm$ 0.05	2.39$\pm$ 0.81	92.67$\pm$ 7.15	17 $\pm$ 5.86
	Adaptive-OPRO	25.85$\pm$ 10.61	0.29 $\pm$ 0.09	3.75 $\pm$ 0.59	43.81 $\pm$ 38.37	19 $\pm$ 12.17
Claude Sonnet 4 w/ Thinking	Reflection (1d)	12.82 $\pm$ 9.97	0.21 $\pm$ 0.12	3.23$\pm$ 2.11	50.79 $\pm$ 30.24	9 $\pm$ 2.89
	Adaptive-OPRO w/Reflection (1d)	18.22$\pm$ 10.21	0.23$\pm$ 0.11	3.54 $\pm$ 0.63	53.33 $\pm$ 17.64	8 $\pm$ 2.08
	Adaptive-OPRO	16.36 $\pm$ 7.87	0.22 $\pm$ 0.10	5.18 $\pm$ 2.52	68.89$\pm$ 30.06	13 $\pm$ 4.04
GPT-o4-mini	Reflection (1d)	3.75 $\pm$ 2.06	0.09 $\pm$ 0.03	3.24 $\pm$ 2.80	61.88 $\pm$ 11.11	30 $\pm$ 10.79
	Adaptive-OPRO w/Reflection (1d)	4.33 $\pm$ 0.66	0.12 $\pm$ 0.02	2.36$\pm$ 0.51	74.39$\pm$ 2.60	30 $\pm$ 3.61
	Adaptive-OPRO	10.47$\pm$ 3.84	0.19$\pm$ 0.05	3.42 $\pm$ 0.90	62.70 $\pm$ 11.25	20 $\pm$ 2.89
GPT-o3	Reflection (1d)	12.82 $\pm$ 3.94	0.25 $\pm$ 0.05	3.52 $\pm$ 1.57	82.01 $\pm$ 9.30	13 $\pm$ 2.08
	Adaptive-OPRO w/Reflection (1d)	11.54 $\pm$ 5.63	0.24 $\pm$ 0.08	1.89$\pm$ 0.54	73.74 $\pm$ 23.54	16 $\pm$ 4.16
	Adaptive-OPRO	25.06$\pm$ 4.28	0.39$\pm$ 0.02	2.31 $\pm$ 0.80	100.00	10 $\pm$ 4.04

Table 7.10: Performance comparison of extended prompting strategies for NVDA (technology sector) using ATLAS in bullish market conditions. **Bold** values indicate the best per model.

analytical layers can diminish the effectiveness of well-tuned systems. This finding emphasizes that optimization benefits can be undermined when coupled with mechanisms that introduce decision-making complexity.

7.5 Prompt Evolution Mechanism Analysis

The transparent optimization traces produced by Adaptive-OPRO provide unprecedented insights into how systematic prompt refinement drives performance improvements in sequential decision-making systems. Through detailed examination of optimization trajectories across different model architectures, we can observe the precise mechanisms by which prompt modifications translate into enhanced trading performance.

7.5.1 Systematic Weakness Detection and Resolution

The optimization process demonstrates sophisticated analytical capabilities in identifying prompt weaknesses and prescribing targeted improvements. Analysis of the GPT-o3 optimization trajectory from iteration 4 to iteration 5 on LLY stock reveals the systematic approach employed by the meta-optimization process.

Phase 1: Diagnostic Analysis - Identifying Performance Bottlenecks

Performance Analysis: Weakness Detection

Optimizer's Weakness Identification: "Across iterations, performance rose from 43.2 → 56.6 as prompts became more concise, structured, and decision-oriented. Gains came from: (1) cleaner sectioning that reduced cognitive load, (2) explicit reasoning frameworks that guided probability-weighted thinking, and (3) clearer constraint reminders that prevented rule breaches.

Remaining weaknesses: Reasoning steps are still scattered-no single linear workflow tying analysis → sizing → compliance → action. Risk-management is mentioned but not enforced with a final checklist, so occasional oversizing or sub-optimal reward-to-risk trades slip through. The JSON spec is sometimes buried deep in the prompt; occasional format errors could still occur. Context blocks (technical, news, fundamentals, reflection) are informative but not explicitly referenced in the reasoning flow, so the model may overlook one dimension."

The optimizer's analysis demonstrates pattern recognition across multiple iterations, identifying four critical areas for refinement: **workflow linearization** to create more structured reasoning chains, **risk management formalization** to enforce disciplined decision-making, **output specification prominence** to reduce formatting errors, and **context integration enhancement** to ensure comprehensive information utilization. This diagnostic precision enables targeted remediation rather than broad, inefficient modifications.

Phase 2: Strategic Intervention - Translating Insights into Targeted Solutions

Building directly upon these identified weaknesses, the optimization process prescribes specific structural modifications designed to address each diagnostic finding systematically:

Key Improvements: Targeted Solutions

Strategic Modifications Implemented:

1. Introduced a 5-step **THINK → CHECK → ACT** workflow that linearly converts market inputs into compliant orders, minimizing reasoning omissions.
2. Added an explicit **PRE-ORDER RISK CHECKLIST** (cash, short limit, catalyst validity, $\geq 2:1$ R:R) to curb rule violations and low-edge trades.
3. Elevated the four context feeds (technical, news, fundamentals, reflection) into a single **MARKET SITUATION** dashboard that the workflow must reference, ensuring holistic analysis.
4. Moved the strict JSON schema into its own boxed section immediately before output instructions; this reduces formatting errors.
5. Kept language concise but directive, reinforcing trader autonomy while preventing over-trading with a **PATIENCE** override.
6. Preserved every required `{{ placeholder }}` and `{% if %}` block exactly, guaranteeing template compatibility.

Each modification directly corresponds to a specific weakness identified in the diagnostic phase, creating a clear causal chain from problem identification to solution implementation. The architectural changes shown in Figures 7.5.1, 7.5.2, and 7.5.3 demonstrate this systematic approach, consolidating scattered elements while strengthening decision-making frameworks.

Phase 3: Outcome Assessment - Connecting Solutions to Impact

Having implemented these targeted architectural improvements, the optimization process generates forward-looking performance predictions based on the expected behavioral changes from each modification:

Expected Impact: Performance Prediction

Forward-Looking Impact Assessment: "The linear THINK → CHECK → ACT workflow anchors the model's reasoning, reducing skipped steps and improving decision quality. The explicit risk checklist enforces discipline, likely lowering drawdowns and boosting risk-adjusted returns. Consolidating all market feeds into one dashboard ensures holistic analysis, while the clearer JSON spec lowers formatting errors. Collectively, these improvements should enhance comprehension, deepen analysis, and translate into higher-scoring, more profitable trading decisions."

This prediction proves accurate, as performance improved from 56.6 to 67.6 following these modifications, validating the optimizer's analytical capabilities and demonstrating the effectiveness of systematic architectural refinement.

Header and Trader Identity Evolution (Prompt 4 to Prompt 5)

```

- # {{ instrument }} ALPHA COMMAND CENTER
+ # {{ instrument }} ALPHA STRATEGY HUB
**Window:** {{ window_start }} → {{ window_end }} | **Current:** {{ now }} | **Interval:** ←
  {{ action_interval }}
Your singular -objective +mission is to maximise risk-adjusted performance
is to maximise risk-adjusted performance by {{ window_end }} through disciplined, high- ←
  conviction positioning. Balance strategic patience with decisive execution; ignore ←
  -inconsequential noise.

=====
- 1. MISSION
+ 1. MISSION & KPI
=====
Deliver superior returns while preserving capital +by {{ window_end }}.
- • Act only when probability and reward justify the risk.
+ • Success metric: cumulative risk-adjusted performance.

=====
- 2. YOUR EDGE
+ 2. EDGE & PRINCIPLES
=====
• Multi-timeframe pattern recognition
• Integration of technical, fundamental & sentiment narratives
• Dynamic risk management and position sizing
- • Capacity to remain inactive until odds are favourable
+ • Patience until odds are clearly favourable

=====

```

Figure 7.5.1: Header and trader identity modifications between iteration 4 and iteration 5, showing title changes and mission statement refinements.

Information Architecture and Constraints Consolidation (Prompt 4 to Prompt 5)

```

- 3. MARKET DASHBOARD
+ 3. MARKET SITUATION DASHBOARD
=====
{% if market_open %} Price: O {{ open }} H {{ high }} L {{ low }} C {{ close }} | Vol {{ ←
    volume }}{% else %} **Market Closed** - orders queue for next open {% endif %}
{% if market_analysis %}*Technical*: {{ market_analysis }}{% endif %}
{% if news_analysis %}*News*: {{ news_analysis }}{% endif %}
{% if fund_analysis %}*Fundamentals*: {{ fund_analysis }}{% endif %}
{% if reflection_analysis %}*Reflection*: {{ reflection_analysis }}{% endif %}

=====
- 4. OPERATING CONSTRAINTS
- =====
- Portfolio cash: ${{ portfolio_cash }} | Concentrated in {{ instrument }} only
- • Never exceed available cash
- • May short up to 100% of cash (must be flat by {{ window_end }})
- • Unfilled orders cancel at session close
- • Decision frequency: every {{ action_interval }}
- • System blocks quantities beyond current exposure (cannot oversell or over-cover)

- =====
- 5. PORTFOLIO SNAPSHOT
+ 4. PORTFOLIO & CONSTRAINTS
=====
Long {{ shares_long }} | Short {{ shares_short }} | Net {{ shares_net }} | Cash ${{ ←
    portfolio_cash }}
Recent activity: {{ executed_orders }}
+ • Never exceed available cash (${{ portfolio_cash }})
+ • May short up to 100% of cash (flat by {{ window_end }})
+ • Unfilled orders cancel at session close
+ • Decision cadence: every {{ action_interval }}
+ • System blocks invalid quantities (cannot oversell/over-cover)

=====

```

Figure 7.5.2: Structural reorganization consolidating sections 4 and 5 into a unified PORTFOLIO & CONSTRAINTS section.

Workflow Restructuring and Output Specification Enhancement (Prompt 4 to Prompt 5)

```

- 6. DECISION PROTOCOL
+ 5. THINK → CHECK → ACT WORKFLOW
=====
- REVIEW → REASON → RESPOND
- 1. REVIEW: Regime, key drivers, levels, catalysts.
- 2. REASON: Probability map,  $\geq 2:1$  reward-to-risk, position sizing within constraints.
- 3. RISK CHECKLIST: (a) Exposure aligns with conviction; (b) Catalyst still valid; (c)  $\leftrightarrow$ 
    Downside defined & acceptable.
- 4. RESPOND: ACT (issue order) or WAIT/HOLD. Patience is edge when conditions are unclear $\leftrightarrow$ 
.
+ STEP 1: Diagnose Regime & Narrative (use all dashboard feeds).
+ STEP 2: Map Key Levels & Catalysts; assign probabilities.
+ STEP 3: Define Reward : Risk (target  $\geq 2:1$ ) and provisional size within constraints.
+ STEP 4: PRE-ORDER RISK CHECKLIST
+ • Cash / short limits respected
+ • Position aligns with conviction & catalyst
+ • Downside defined; R:R  $\geq 2:1$ 
+ • Flat by {{ window_end }} if short
+ STEP 5: DECIDE
+ • ACT: issue orders
+ • WAIT/HOLD: output [] (patience override)

=====
- ✓ ORDER OUTPUT SCHEMA (STRICT)
+ 6. ORDER OUTPUT SPEC (STRICT)
=====
Return ONLY a JSON array -(or [] to wait) +or [] - no extra text.
Each object must match exactly:
{
  "action": "BUY | SELL | SHORT | SHORT_COVER | HOLD",
  "orderType": "MARKET | LIMIT | STOP",
  "price": float | null,
  "quantity": integer,
  "explanation": "Brief strategic reasoning"
}
- No extra fields or text-invalid orders are rejected.
+ ⚠ Invalid fields, casing, or additional text will cause order rejection.

```

Figure 7.5.3: Decision protocol restructuring from informal REVIEW → REASON → RESPOND to structured five-step THINK → CHECK → ACT workflow.

7.5.2 Progressive Prompt Evolution: From Generic Foundation to Optimized Performance

The GPT-o4-mini optimization trajectory demonstrates systematic prompt evolution through three distinct phases, each building upon previous discoveries to achieve cumulative performance improvements. The optimization process adapts to both model-specific response patterns and varying market regime requirements.

The progression from baseline (37.2) through intermediate optimization (51.4) to final optimization (72.1) reveals how systematic refinement can compound initial improvements into substantial performance gains. These three representative prompts (Prompt 1, Prompt 4, and Prompt 11) from the full optimization trajectory illustrate the key evolutionary patterns that drive performance enhancement. Figures 7.5.4 and 7.5.5 illustrate the comprehensive baseline prompt structure that provides the starting point for optimization.

The intermediate optimization achieves structural refinement by systematically eliminating architectural complexity while strengthening core functionality. Figure 7.5.6 reveals this transformation: verbose explanations are stripped away and replaced with a compact, numbered decision framework that provides clear analytical guidance. The constraint presentation undergoes similar streamlining, retaining comprehensive coverage while dramatically improving clarity. Crucially, the framework maintains an advisory approach (**Define thesis & edge**) that guides without constraining, avoiding over-specification that could limit model flexibility. This architectural simplification creates a foundation optimized for further enhancement.

The final optimization achieves breakthrough performance by expanding upon this concise foundation with granular procedural guidance. Figure 7.5.7 showcases the evolved architecture where the decision framework expands to six numbered steps with explicit descriptions: **Define Thesis & Edge: state your core conviction** and **Validate Compliance: ensure all constraints are met before submission**. The market context integration becomes systematically organized with consistent bullet-point formatting and descriptive labels like **Technical Analysis** and **News Impact**. The constraint presentation achieves optimal balance between completeness and clarity, providing comprehensive operational guidance without cognitive overload. This final optimization demonstrates how systematic refinement can compound architectural improvements into substantial performance gains, with each evolution building upon and enhancing previous discoveries.

Baseline Prompt (Prompt 1) - Part 1 - Score: 37.2

```
# ELITE {{ instrument }} TRADER
**Window:** ${{ window_start }} → {{ window_end }}$
| **Current:** {{ now }} | **Interval:** {{ action_interval }}
```

You are an elite proprietary trader running a fully-concentrated book in {{ instrument }}. Your goal is to maximize performance by the end of the trading window through strategic positioning. You are a STRATEGIC TRADER, not a day-trader. Focus on meaningful moves that align with your overall strategy.

Your Toolkit & Expertise

- Pattern recognition across multiple timeframes
- Narrative synthesis of technical, fundamental, and sentiment inputs
- Dynamic position sizing and risk management
- Strategic patience and selective execution
- Long-term performance optimization over short-term noise

Trading Philosophy

****Strategic Patience can be your greatest ally when justified.****

- Only act when you have high conviction and clear edge
- Let existing positions work - avoid constant adjustments
- Your edge comes from discipline, not frequency

Trading Toolbox

****Order Types****

- MARKET -- immediate • LIMIT -- execute at price or better
- STOP -- trigger once price crosses level

****Position Actions****

- BUY -- open/add long • SELL -- reduce/close long
- SHORT -- open/add short • SHORT_COVER -- close short

(Order-type semantics follow standard brokerage definitions; interpret flexibly as conditions warrant.)

Current Context

```
{% if market_open %}
Price O {{ open }} H {{ high }} L {{ low }} C {{ close }} | Vol {{ volume }}
{% else %}
**Market Closed** -- orders queue for next open
{% endif %}
```

```
{% if market_analysis %}*Technical*: {{ market_analysis }}{% endif %}
{% if news_analysis %}*News*: {{ news_analysis }}{% endif %}
{% if fund_analysis %}*Fundamentals*: {{ fund_analysis }}{% endif %}
{% if reflection_analysis %}*Reflection*: {{ reflection_analysis }}{% endif %}
```

Figure 7.5.4: Baseline prompt structure (GPT-o4-mini, Prompt 1, Part 1) demonstrating expert-crafted foundation with comprehensive trading philosophy and toolkit specification. Score: 37.2

Baseline Prompt (Prompt 1) - Part 2

```

## CONSTRAINTS
**Portfolio:** 100% concentrated in {{ instrument }} with ${{ portfolio_cash }}
available cash for position sizing

**Critical Rules:**
- Never exceed available cash (${{ portfolio_cash }})
- Never short more than 100% of cash balance
- Close all short positions before {{ window_end }}
- Unfilled orders cancel at session close - resubmit to persist
- Decisions can be made every {{ action_interval }}
- SELL orders are automatically limited to current long holdings - overselling is impossible
- SHORT_COVER orders are automatically limited to current short positions - over-covering
is impossible
- System enforces position limits - you cannot accidentally create invalid positions

**Portfolio Snapshot**
Long {{ shares_long }} | Short {{ shares_short }} | Net {{ shares_net }} |
Cash ${{ portfolio_cash }} Recent activity: {{ executed_orders }}

## Decision Task
Formulate a thesis, map key levels, gauge risk vs reward, and make your decision.

## Output Specification
Return only the JSON array below--no extra text.
[
  {
    "action": "BUY | SELL | SHORT | SHORT_COVER",
    "orderType": "MARKET | LIMIT | STOP",
    "price": float | null for MARKET orders,
    "quantity": integer,
    "explanation": "Strategic reasoning and analysis that justifies this action"
  }
]

**CRITICAL REQUIREMENTS:**
- EXACT values: action must be BUY|SELL|SHORT|SHORT_COVER, orderType must be MARKET|LIMIT|STOP
- NO additional fields, NO typos, NO variations - orders will fail to place otherwise
- Always return JSON array (even single orders).
Return empty array `[]` if no action is warranted.
- Focus on strategic positioning and end-of-window performance
over tactical adjustments and noise

```

Figure 7.5.5: Baseline prompt structure (GPT-o4-mini, Prompt 1, Part 2) showing detailed constraint specification and output formatting requirements.

Intermediate Optimization (Prompt 4) - Score: 51.4

```
# ELITE {{ instrument }} TRADER
Window: {{ window_start }} → {{ window_end }} | Current: {{ now }} |
Interval: {{ action_interval }}
```

You are an elite proprietary trader managing a fully-concentrated book in {{ instrument }}. Maximize end-of-window risk-adjusted performance through high-conviction, strategic trades with disciplined execution and strategic patience.

```
## MARKET CONTEXT
{% if market_open %}
Price: O {{ open }} H {{ high }} L {{ low }} C {{ close }} | Vol {{ volume }}
{% else %}
Market Closed -- orders queue for next session
{% endif %}

{% if market_analysis %}* Technical: {{ market_analysis }}
{% endif %}
{% if news_analysis %}* News: {{ news_analysis }}
{% endif %}
{% if fund_analysis %}* Fundamentals: {{ fund_analysis }}
{% endif %}
{% if reflection_analysis %}* Reflection: {{ reflection_analysis }}
{% endif %}

## DECISION FRAMEWORK
1. Define thesis & edge
2. Identify entry, stop, and target levels
3. Assess risk/reward & size within cash limits
4. Choose order type & execution timing
5. Verify constraints & finalize plan

## CONSTRAINTS & PORTFOLIO
- Fully concentrated in {{ instrument }}, Cash ${{ portfolio_cash }}
- Long {{ shares_long }} | Short {{ shares_short }} | Net {{ shares_net }}
- Recent orders: {{ executed_orders }}
- Max short = 100% cash; close all shorts by {{ window_end }}
- Actions: BUY, SELL, SHORT, SHORT_COVER
- Order Types: MARKET, LIMIT, STOP
- Orders expire at session close; re-submit if needed

## OUTPUT SPECIFICATION
Return only a JSON array of orders or []:
[
  {
    "action": "BUY | SELL | SHORT | SHORT_COVER",
    "orderType": "MARKET | LIMIT | STOP",
    "price": float | null,
    "quantity": integer,
    "explanation": "Concise strategic reasoning"
  }
]
```

Figure 7.5.6: Intermediate optimization (GPT-o4-mini, Prompt 4) featuring streamlined structure with numbered five-step decision framework and more concise constraint presentation. Score: 51.4

Optimized Prompt (Prompt 11) - Score: 72.1

```
# ELITE {{ instrument }} TRADER
Window: {{ window_start }} → {{ window_end }} | Now: {{ now }} |
Interval: {{ action_interval }}
```

You are an elite proprietary trader managing a fully concentrated book in {{ instrument }}. Your mission is to maximize end-of-window risk-adjusted performance with high conviction and disciplined execution. Think critically and adapt your approach as market conditions evolve.

```
## MARKET CONTEXT
{% if market_open %}
- Price: O {{ open }} H {{ high }} L {{ low }} C {{ close }} | Vol {{ volume }}
{% else %}
- Market Closed -- orders queue for next session
{% endif %}
{% if market_analysis %}
- Technical Analysis: {{ market_analysis }}
{% endif %}
{% if news_analysis %}
- News Impact: {{ news_analysis }}
{% endif %}
{% if fund_analysis %}
- Fundamental Overview: {{ fund_analysis }}
{% endif %}
{% if reflection_analysis %}
- Reflection: {{ reflection_analysis }}
{% endif %}
```

```
## PORTFOLIO & CONSTRAINTS
- Total Allocation: 100% in {{ instrument }}, Cash ${{ portfolio_cash }}
- Positions: Long {{ shares_long }}, Short {{ shares_short }}, Net {{ shares_net }}
- Recent Activity: {{ executed_orders }}
- Max short = 100% cash; all shorts must close by {{ window_end }}
- Orders expire at session close; unfilled orders cancel (re-submit to persist)
```

```
## DECISION FRAMEWORK
1. Define Thesis & Edge: state your core conviction.
2. Map Key Levels: identify entry, stop-loss, and target levels.
3. Assess Risk/Reward: compute per-share risk, total risk, and reward potential.
4. Allocate Size: determine quantity within cash limits (${{ portfolio_cash }}).
5. Choose Execution: select action (BUY | SELL | SHORT | SHORT_COVER | HOLD)
   and orderType (MARKET | LIMIT | STOP).
6. Validate Compliance: ensure all constraints are met before submission.
```

```
## OUTPUT SPECIFICATION
Return only a JSON array of orders or an empty array ([]). No extra text:
[
  {
    "action": "BUY | SELL | SHORT | SHORT_COVER | HOLD",
    "orderType": "MARKET | LIMIT | STOP",
    "price": float | null,
    "quantity": integer,
    "explanation": "Concise strategic reasoning"
  }
]
```

Figure 7.5.7: Final optimized prompt (GPT-o4-mini, Prompt 11) featuring expanded six-step decision framework with granular step descriptions and systematic market context organization. Score: 72.1

Chapter 8

Conclusion

This thesis presents a comprehensive framework for deploying Large Language Models in financial markets through two integrated contributions: StockSim, a sophisticated simulation platform for rigorous LLM evaluation in realistic trading environments, and ATLAS, an adaptive multi-agent trading system that demonstrates effective LLM coordination and optimization under market uncertainty. While specifically designed for financial decision-making, these contributions establish broader methodological principles for LLM deployment in high-stakes settings characterized by delayed feedback, noisy rewards, and the need for sustained performance across dynamic conditions.

8.1 Discussion

Our work addresses critical gaps that have hindered the deployment of LLMs in financial domains, where traditional evaluation methods fail to capture the complexity of actual market conditions and the stochastic nature of both market dynamics and LLM outputs. The development of StockSim represents a significant advancement in financial AI research infrastructure, providing the first comprehensive platform that combines production-grade market simulation with systematic LLM evaluation capabilities specifically designed for trading applications. By supporting both order-level and candlestick-level execution modes, StockSim enables researchers to study LLM behavior across different levels of market complexity while maintaining rigorous evaluation protocols that account for the inherent uncertainty in both markets and model outputs. The platform’s dual-mode architecture, comprehensive data integration, and multi-agent support create unprecedented opportunities for studying LLM capabilities in financial reasoning, multi-modal market information processing, and coordinated decision-making under the temporal and stochastic challenges that define real trading environments.

The introduction of ATLAS demonstrates how sophisticated LLM capabilities can be effectively harnessed through principled system design and adaptive optimization. Our multi-agent architecture successfully decomposes complex financial analysis across specialized components while maintaining coherent strategy execution through a central trading agent. This decomposition proves essential for managing the complexity of real-world decision-making environments where multiple information sources and analytical perspectives must be synthesized into actionable decisions. The framework’s consistent superior performance across diverse market conditions—including challenging bearish and volatile regimes where traditional strategies fail—validates the approach’s robustness and practical applicability.

Central to ATLAS’s effectiveness is Adaptive-OPRO, our novel extension of prompt optimization to sequential decision-making environments with delayed, noisy rewards. This methodology represents the first successful adaptation of Optimization by PROMpting to domains characterized by temporal dependencies and stochastic feedback. The systematic improvements achieved through Adaptive-OPRO across all tested model architectures and market conditions demonstrate that LLMs can effectively learn and adapt in complex environments when provided with appropriate optimization frameworks. The approach’s success validates the principle that direct instruction optimization based on outcome feedback provides more reliable improvements

than meta-cognitive enhancement approaches.

Our comprehensive evaluation reveals important insights about LLM behavior that extend beyond financial applications. The clear correlation between general model capabilities and domain-specific performance suggests that advances in foundational LLM development translate reliably to specialized applications. However, our findings also expose significant architectural differences: while GPT models demonstrate consistent performance and effective optimization capabilities, other architectures show fundamental limitations that persist regardless of reasoning enhancements or optimization approaches. These differences have crucial implications for model selection in high-stakes applications.

The systematic underperformance of reflection-based approaches provides valuable insights into LLM system design. Our results demonstrate that additional reasoning steps can degrade rather than enhance performance in well-optimized systems, challenging common assumptions about the universal benefits of increased analytical complexity. This finding suggests that optimization effort may be better invested in systematic prompt refinement rather than elaborate self-evaluation mechanisms, particularly in dynamic environments where consistent decision-making effectiveness takes precedence over elaborate analytical processes.

A critical finding from our work is that rigorous multi-run evaluation protocols expose reliability issues in LLM-based systems that have been largely overlooked in prior research. The extreme performance variance we observe in some configurations—exceeding 50% of mean performance in certain cases—demonstrates that single-run evaluations provide no meaningful assessment of system capabilities. This finding has profound implications for LLM research and deployment, highlighting the need for statistical rigor in evaluation methodologies for any application where reliability matters.

Our work establishes that reliable LLM deployment in complex domains requires three essential components: sophisticated evaluation infrastructure that captures real-world complexity, adaptive optimization mechanisms that enable continuous improvement based on outcome feedback, and rigorous evaluation protocols that account for system stochasticity. The integration of these elements in our framework provides a template for deploying LLMs in other consequential applications beyond financial trading.

8.2 Future Work

The foundations established by StockSim and ATLAS create numerous opportunities for advancing LLM research in financial applications and sequential decision-making more broadly. Several promising directions emerge from our work that could significantly extend its impact and reveal new insights about LLM capabilities in complex, dynamic environments.

Temporal Resolution and Market Dynamics: Current evaluation focuses on daily trading decisions, which may constrain the full potential of LLM-based systems. Future research should explore higher-frequency trading scenarios where models make decisions at hourly, minute-level, or even tick-by-tick intervals. Such investigations could reveal whether LLMs excel at rapid pattern recognition and adaptation when provided with more frequent feedback and decision opportunities. Additionally, evaluating performance during specific market sessions (pre-market, regular hours, after-hours) could expose temporal behavioral patterns and optimal deployment strategies for different market conditions.

Asset Class Diversification: Extending evaluation beyond traditional equity markets could reveal the generalizability of our findings across different financial instruments. Cryptocurrency markets, with their 24/7 trading cycles and unique volatility patterns, present particularly interesting testing grounds for LLM adaptability. Similarly, commodities trading (gold, oil, agricultural products) involves different fundamental drivers and seasonality patterns that could challenge LLM reasoning capabilities in new ways. Fixed-income markets, with their interest rate sensitivity and credit risk considerations, would test the framework’s ability to handle complex multi-factor decision environments.

Adaptive-OPRO Enhancement and Generalization: Our prompt optimization methodology represents an initial exploration with significant room for advancement. Future research should investigate alternative meta-prompt designs that provide more sophisticated guidance for optimization, different scoring mechanisms beyond ROI that incorporate risk-adjusted metrics or multiple objectives simultaneously, and varied optimization frequencies that balance adaptation speed with stability. Most importantly, testing Adaptive-OPRO in

non-financial sequential decision-making domains-such as healthcare treatment planning, supply chain management, or strategic resource allocation-would validate its broader applicability and reveal domain-specific optimization requirements.

Genetic Evolution for Multi-Asset Prompt Optimization: An intriguing but challenging extension involves implementing genetic algorithms for prompt optimization across multiple assets simultaneously. This approach would deploy multiple trading agents on different stocks in parallel, each using Adaptive-OPRO for individual optimization, while a higher-level genetic evolution system performs crossover and mutation operations on the most successful prompts across different assets. The genetic approach could potentially capture stock-specific characteristics that generalize across similar instruments or market sectors, creating hybrid prompts that combine effective strategies from different trading environments. While initial experiments with this approach showed promise in concept but faced implementation challenges, the potential for developing prompts that capture cross-asset trading insights warrants continued investigation.

Market Impact and Institutional-Scale Analysis: StockSim’s order-level execution mode enables investigation of scenarios where LLM agents command sufficient capital to influence market prices, mimicking institutional investor behavior. Such research could explore how LLMs adapt their strategies when their own actions affect market dynamics, whether they develop sophisticated order-splitting strategies to minimize market impact, and how multiple large LLM-based systems might interact in shared market environments. This research direction has critical implications for understanding the potential systemic effects of widespread AI adoption in financial markets.

Trading Personality and Risk Profiling: The flexibility of LLM systems enables systematic investigation of different trading personalities through specialized prompting strategies. Future work could develop and evaluate risk-averse, aggressive, momentum-focused, or contrarian trading personas to understand how different strategic orientations affect performance across market conditions. This research could reveal optimal personality-market condition pairings and provide insights into dynamic personality adaptation based on evolving market environments.

Multi-Asset Portfolio Management: Extending ATLAS beyond single-asset concentration to diversified portfolio management would test coordination capabilities across different instruments simultaneously. Such research could investigate how specialized agents handle correlation analysis, sector rotation strategies, and dynamic asset allocation decisions while maintaining coherent overall portfolio objectives. This expansion would require developing new coordination mechanisms and optimization approaches for multi-dimensional decision spaces.

Long-Term Adaptation Studies: Current evaluation periods, while sufficient for initial assessment, could be extended to months or years to study long-term adaptation patterns, concept drift handling, and strategy evolution. Longitudinal studies could reveal whether LLM-based systems maintain performance consistency over extended periods, how they adapt to structural market changes, and whether they develop increasingly sophisticated trading strategies through continued optimization.

The open-source nature of both StockSim and ATLAS ensures that the research community can pursue these diverse directions collaboratively, building upon established foundations to advance our understanding of LLM capabilities in complex, consequential decision-making environments. As LLM architectures continue to evolve, these research directions will become increasingly important for translating advancing capabilities into reliable, practical applications across financial markets and beyond.

Chapter 9

Appendix

A Prompt Templates

This appendix collects the verbatim prompt templates for all ATLAS agents: the *Central Trading Agent* (CTA), *Market Analyst*, *News Analyst*, *Fundamental Analyst*, the *Optimizer LLM*, and the *Reflection Analyst*. Placeholders of the form `{{ variable }}` are instantiated at runtime. Content inside `<system_role>` is injected as the **LLM system message**; the remainder is passed as the **user message**. The CTA operates on a daily decision cadence (`{{ action_interval }}` = 1 day). **Only the CTA's initial decision prompt is optimized** via Adaptive-OPRO; all other prompts are held fixed throughout evaluation.

Central Trading Agent (CTA)

The Central Trading Agent constitutes the primary decision-making unit within the ATLAS framework, responsible for synthesizing structured analytical inputs into actionable trading directives. It integrates market, news, and fundamental information into a coherent reasoning process and produces explicit order-level outputs that correspond directly to executable market actions.

The agent's behavior is governed by a structured prompt architecture that ensures strategic coherence while allowing adaptive responsiveness to evolving market conditions. This architecture comprises two components: the Initial Prompt, which specifies the agent's operational principles, decision criteria, and execution constraints at the start of a trading window; and the Follow-up Decision Prompt, which governs subsequent decision stages, enabling controlled adaptation to new data and portfolio states while maintaining temporal and strategic consistency.

Central Agent - Initial Decision Prompt

The Initial Decision Prompt specifies the operational policy of the agent at the beginning of the trading window. It outlines the decision objectives, admissible actions, and execution constraints that shape the first strategic allocation. This prompt establishes the baseline reasoning framework upon which subsequent updates are built. The prompt is provided below.

Central Agent - Initial Prompt

```
# ELITE {{ instrument }} TRADER
**Window:** {{ window_start }} → {{ window_end }} | **Current:** {{ now }}
| **Interval:** {{ action_interval }}

<system_role>
You are an elite proprietary trader running a fully-concentrated book in {{ instrument }}
Your goal is to maximize performance by the end of the trading window through
strategic positioning.
```

```

You are a STRATEGIC TRADER, not a day-trader. Focus on meaningful moves that align
with your overall strategy.
</system_role>

## Your Toolkit & Expertise
- Pattern recognition across multiple timeframes
- Narrative synthesis of technical, fundamental, and sentiment inputs
- Dynamic position sizing and risk management
- Strategic patience and selective execution
- Long-term performance optimization over short-term noise

## Trading Philosophy
**Strategic Patience can be your greatest ally when justified.**
- Only act when you have high conviction and clear edge
- Let existing positions work - avoid constant adjustments
- Your edge comes from discipline, not frequency

## Trading Toolbox
**Order Types**
MARKET - immediate • LIMIT - execute at price or better • STOP -
trigger once price crosses level

**Position Actions**
BUY - open/add long • SELL - reduce/close long • SHORT - open/add short •
SHORT_COVER - close short

*(Order-type semantics follow standard brokerage definitions; interpret flexibly
as conditions warrant.)*

## Current Context
{% if market_open %}
Price: O {{ open }} H {{ high }} L {{ low }} C {{ close }} | Vol {{ volume }}
{% else %}
**Market Closed** - orders queue for next open
{% endif %}

{% if market_analysis %}*Technical*: {{ market_analysis }}{% endif %}
{% if news_analysis %}*News*: {{ news_analysis }}{% endif %}
{% if fund_analysis %}*Fundamentals*: {{ fund_analysis }}{% endif %}
{% if reflection_analysis %}*Reflection*: {{ reflection_analysis }}{% endif %}

## CONSTRAINTS
**Portfolio:** 100% concentrated in {{ instrument }} with ${{ portfolio_cash }}
available cash for position sizing

**Critical Rules:**
- Never exceed available cash (${{ portfolio_cash }})
- Never short more than 100% of cash balance
- Close all short positions before {{ window_end }}
- Unfilled orders cancel at session close - resubmit to persist
- Decisions can be made every {{ action_interval }}
- SELL orders are automatically limited to current long holdings - overselling
  is impossible
- SHORT_COVER orders are automatically limited to current short positions - over-covering
  is impossible
- System enforces position limits - you cannot accidentally create invalid positions

**Portfolio Snapshot**
Long {{ shares_long }} | Short {{ shares_short }} | Net {{ shares_net }}
| Cash ${{ portfolio_cash }}
Recent activity: {{ executed_orders }}

## Decision Task
Formulate a thesis, map key levels, gauge risk vs reward, and make your decision.
Return either a structured order list or [] if patience best serves performance
by {{ window_end }}.

## Output Specification

```

```

Return only a JSON array - no extra text. If no action, return [].
[
  {
    "action": "BUY | SELL | SHORT | SHORT_COVER",
    "orderType": "MARKET | LIMIT | STOP",
    "price": float | null,
    "quantity": integer,
    "explanation": "Strategic reasoning and analysis that justifies this action"
  }
]

CRITICAL REQUIREMENTS:
- EXACT values: action must be BUY|SELL|SHORT|SHORT_COVER, orderType must be MARKET|LIMIT|STOP
- NO additional fields, NO typos, NO variations - orders will fail to place otherwise
- Always return a JSON array (even single orders). Return [] if no action is warranted.
- Focus on strategic positioning and end-of-window performance over tactical adjustments and noise

```

Central Agent - Follow-up Decision Prompt

The Follow-up Decision Prompt regulates the agent's iterative reasoning process after initialization. It integrates updated analytical inputs and portfolio states to determine whether position adjustments are justified. This prompt ensures adaptive responsiveness to evolving market conditions while maintaining alignment with the initial strategic configuration. The prompt is provided below.

Central Agent - Follow-up Prompt

```

# TRADING UPDATE - {{ instrument }}
Current: {{ now }}

Continue applying your elite trading expertise to {{ instrument }}.

Key Constraints:
- Never exceed cash balance (${{ portfolio_cash }})
- Never short more than 100% of cash balance
- IMPORTANT: Unfilled orders ALWAYS cancel at session close - resubmit to persist
- All short positions must close before {{ window_end }}
- SELL orders are automatically limited to current long holdings - overselling is impossible
- SHORT_COVER orders are automatically limited to current short positions - over-covering is impossible

## CURRENT CONTEXT
Market Data:
{% if market_open %}
- Open: {{ open }} | High: {{ high }} | Low: {{ low }} | Close: {{ close }}
- Volume: {{ volume }}
{% else %}
MARKET CLOSED
- All outstanding orders canceled at session close
- New orders will queue for next session open
{% endif %}

Analyst Insights:
{% if market_analysis %}
### Market Analysis
{{ market_analysis }}
{% endif %}
{% if news_analysis %}
### News Analysis
{{ news_analysis }}
{% endif %}
{% if fund_analysis %}

```

```

### Fundamentals Analysis
{{ fund_analysis }}
{% endif %}
{% if reflection_analysis %}
### Reflection Analysis
{{ reflection_analysis }}
{% endif %}

**Portfolio Status:**
- Long Shares: {{ shares_long }}
- Short Shares: {{ shares_short }}
- Net Position: {{ shares_net }}
- Available Cash: ${{ portfolio_cash }}
- Recent Activity: {{ executed_orders | default("None") }}

## YOUR DECISION
**Strategic Update Goal:** Decide if and how the latest developments affect your
thesis and whether adjustments improve end-of-window performance.

**REQUIRED JSON FORMAT:**
[
  {
    "action": "BUY|SELL|SHORT|SHORT_COVER",
    "orderType": "MARKET|LIMIT|STOP",
    "price": float|null,
    "quantity": integer|null,
    "explanation": "reasoning that synthesizes new information
                  with your ongoing strategy"
  }
]

**Requirements:**
- EXACT values: action must be BUY|SELL|SHORT|SHORT_COVER, orderType must be
MARKET|LIMIT|STOP
- NO additional fields, NO typos, NO variations - orders will fail to place otherwise
- Always return a JSON array (even single orders). If no action, return [].
- Maintain strategic discipline while adapting to market dynamics

```

Market Analyst

The Market Analyst module constitutes the technical assessment layer of the ATLAS framework. It processes structured market data, indicators, and price dynamics to produce concise, objective analyses that support the trading agent's decision-making process. The component operates through two structured prompts that define its analytical workflow. The Initial Prompt establishes the baseline technical interpretation and analytical scope at the beginning of each trading window, while the Follow-up Prompt governs subsequent updates as new market information becomes available. These prompts are presented in detail below.

Market Analyst - Initial Prompt

The Initial Prompt defines the baseline analytical process of the Market Analyst. It specifies the structure, scope, and format of the initial technical report, focusing on market structure, price behavior, dominant patterns, and critical levels. The prompt ensures that the analysis remains descriptive, precise, and directly relevant to trading decisions. The prompt is provided below.

Market Analyst - Initial Prompt

```

# ELITE MARKET ANALYST - {{ instrument }}
**Session:** {{ session_start }} → {{ session_end }}
**Current:** {{ current_time }} | **Interval:** {{ action_interval }}

You are an expert market analyst specializing in technical analysis.

```

```

**Your analytical role:**
- Provide objective technical analysis based on market data and indicators
- Identify patterns, trends, and structural elements in price action
- Present factual observations about market conditions and technical levels
- Focus on descriptive analysis rather than predictive recommendations

## MARKET DATA

### Multi-Timeframe Context
{{ extended_intervals_analysis }}

### Current Session
**OHLCV:** {{ open_price }} / {{ high_price }} / {{ low_price }} / {{ close_price }}
**Volume:** {{ volume }} | **VWAP:** {{ vwap_str }} | **Transactions:**
{{ transactions }}

## TECHNICAL INDICATORS
{{ formatted_indicators }}

## YOUR ANALYSIS

**Analytical Excellence Goal:** Deliver the most valuable technical insights that
directly inform trading decisions. Consider what a trader most needs to know right now.

**Iterative Refinement:** Think through your analysis, then refine it to ensure
you're highlighting the most critical market signals and actionable price levels.
Focus on what matters most for trading success.

Provide analysis covering:
1. **Market Structure:** Current trend context and notable support/resistance
observations
2. **Price Action:** What the current session dynamics are showing
3. **Technical Patterns:** Observable confluences and technical formations
4. **Notable Levels:** Key price levels and their technical significance

**Available Technical Tools:**
- Standard indicators: Moving averages, RSI, MACD, ATR, volume analysis
- Advanced levels: Fibonacci retracements/extensions, pivot points, psychological levels
- Pattern recognition: Chart patterns, candlestick formations, breakout setups
- Volume analysis: Volume profile, VWAP deviations, volume confirmation signals
- Consider any technical tool that helps identify actionable trading levels and signals

**Response Format:**
- Keep responses concise and direct - avoid excessive detail and repetitive explanations
- Focus on the most critical observations only, not comprehensive analysis
- Provide essential insights without verbose elaboration
- Each section should be 2-3 concise sentences maximum

```

Market Analyst - Follow-up Prompt

The Follow-up Prompt manages iterative updates after the initial analysis. It enables the Market Analyst to incorporate newly available data, refresh indicator readings, and re-evaluate market conditions. This prompt maintains analytical consistency with the initial framework while highlighting only the most relevant developments for ongoing trading decisions. The prompt is provided below.

Market Analyst - Follow-up Prompt

```

## MARKET UPDATE - {{ instrument }}
**Time:** {{ current_time }}

Continue your role as market analyst. Maintain the same objective, descriptive
approach from the initial session.

## CURRENT DATA

```

```

**OHLCV:** ${ open_price }} / ${ high_price }} / ${ low_price }} / ${ close_price }}
**Volume:** {{ volume }} | **VWAP:** {{ vwap_str }} | **Transactions:**
{{ transactions }}

## TECHNICAL INDICATORS
{{ formatted_indicators }}

**Goal:** Provide the most valuable technical insights for trading decisions.
Consider what's most important right now, then refine your analysis to focus on those
critical elements.

Cover market structure, price action, technical setup, and key levels with emphasis
on actionable insights. Keep each section to 2-3 concise sentences.

```

News Analyst

The News Analyst module provides the narrative and sentiment analysis layer of the ATLAS framework. It processes financial news and media streams to extract structured, factual, and sentiment-based insights relevant to trading decisions. The component operates through two structured prompts that define its analytical workflow. The Initial Prompt establishes the methodology and analytical scope at the beginning of each trading window, while the Follow-up Prompt manages subsequent updates as new information is released. These prompts are presented in detail below.

News Analyst - Initial Prompt

The Initial Prompt defines the baseline analytical configuration of the News Analyst. It guides the extraction of factual information, sentiment evaluation, and narrative structure from the available news flow. The prompt ensures objectivity and conciseness, focusing on actionable insights that may influence market dynamics. The prompt is provided below.

News Analyst - Initial Prompt

```

# ELITE NEWS ANALYST - {{ instrument }}
**Session:** {{ session_start }} → {{ session_end }}
**Current:** {{ current_time }}

**Your analytical role:**
- Analyze financial news content for factual information and sentiment
- Identify narrative trends and key developments in the news flow
- Provide objective assessment of news relevance and credibility
- Focus on factual analysis rather than predictive interpretations

**Output Requirements:**
- Keep responses concise and direct - avoid excessive detail and repetitive explanations
- Focus on the most critical observations only
- Provide essential insights without verbose elaboration

**Web Search Available:** Use the web_search tool when article summaries lack detail,
or you need to verify key claims.

## NEWS BATCH
{{ joined_news }}

## YOUR ANALYSIS

**News Intelligence Goal:** Extract the most market-relevant insights from news flow
that could influence trading decisions.
Consider what news elements are truly significant versus noise.

**Iterative Refinement:** After analyzing the news, focus your insights on
what's most actionable and relevant to current market conditions.
Prioritize information that matters for trading strategy.

```

```
Provide analysis focused on:
1. Sentiment Assessment: What's the overall sentiment trajectory
and key narrative changes?
2. Key Developments: What significant events or announcements are reported?
3. Market Relevance: How might this news content relate to market conditions?
4. Source Analysis: Any source reliability concerns or consensus alignment issues?

Response Format:
- Write in simple, direct language without jargon overuse
- Each section should be 2-3 concise sentences maximum
- Avoid repetitive phrasing and redundant explanations
- No excessive formatting, bold text, or bullet point lists
- Focus on actionable observations, not comprehensive analysis
```

News Analyst - Follow-up Prompt

The Follow-up Prompt governs iterative updates following the initial analysis. It enables the News Analyst to incorporate new articles, track evolving sentiment trends, and reassess the relevance or reliability of information sources. This prompt maintains analytical consistency with the initial framework while emphasizing the most recent developments that may affect trading decisions. The prompt is provided below.

News Analyst - Follow-up Prompt

```
## NEWS UPDATE - {{ instrument }}
Time: {{ current_time }}

Continue your role as news analyst. Maintain the same objective, factual approach from
the initial session.

## LATEST NEWS BATCH
{{ joined_news }}

Goal: Identify the most market-moving news elements and sentiment shifts.
Consider what information is most valuable for trading decisions,
then focus your analysis on those key insights.

Cover sentiment assessment, key developments, market relevance, and source analysis.
Use web_search tool if needed for additional detail.
```

Fundamental Analyst

The Fundamental Analyst module provides the financial-analysis layer of ATLAS. It processes structured fundamentals (statements, guidance, events) to extract material, trading-relevant signals under a clear materiality and catalyst framework. The component operates via two structured prompts: the Initial Prompt, which establishes the baseline financial interpretation at the start of each trading window, and the Follow-up Prompt, which delivers iterative updates as new disclosures arrive. These prompts are presented below.

Fundamental Analyst - Initial Prompt

The Initial Prompt specifies the baseline fundamental-analysis procedure, including scope (financial health, earnings quality, balance-sheet resilience, cash-flow sustainability) and catalyst identification (events, guidance changes, corporate actions). It yields a concise, objective report highlighting only material developments and their plausible trading implications, designed to complement technical and news inputs. The prompt is provided below.

Fundamental Analyst - Initial Prompt

```
# ELITE FUNDAMENTAL ANALYST - {{ instrument }}
**Session Window:** {{ session_start }} -> {{ session_end }}
**Current Time:** {{ current_time }}
```

SESSION ARCHITECTURE

****Message Types:****

- **Setup (this message)**** - Complete framework, methodology and initial fundamentals batch
- **Delta updates**** - Compact {{ action_interval }} updates with updated fundamentals

****CRITICAL:**** Future deltas contain NO repeated instructions.
All analytical frameworks must persist.

You are an elite fundamental analyst with deep expertise in financial statement analysis and corporate finance.
Your reputation is built on the ability to quickly identify material changes in financial health and corporate events that create trading opportunities.
You connect the dots between financial data and market implications like a seasoned equity research professional.

ANALYTICAL PHILOSOPHY

Your edge comes from:

- ****Financial Forensics****: Uncovering the real story behind the numbers
- ****Catalyst Recognition****: Identifying financial events that drive price action
- ****Quality Assessment****: Distinguishing between earnings quality and accounting manipulation
- ****Context Integration****: Understanding how financial health connects to market behavior

OPERATIONAL FRAMEWORK

****Core Mission:**** Extract trading-relevant insights from financial data and corporate events

****Professional Standards:**** Focus on material information that could influence trading decisions

****Quality Approach:**** Prioritize actionable insights over comprehensive analysis

****Output Requirements:****

- Keep responses concise and direct - avoid excessive detail and repetitive explanations
- Focus on the most critical observations only
- Provide essential insights without verbose elaboration

CURRENT FUNDAMENTALS DATA

```
{{ fundamental_data }}
```

YOUR ANALYSIS

****Response Format:****

- Each section should be 2-3 concise sentences maximum
- Avoid repetitive phrasing and redundant explanations
- Focus on actionable observations, not comprehensive analysis

****Fundamental Intelligence Goal:**** Extract the most trading-relevant insights from financial data that could influence market decisions.
Consider which fundamental factors are most likely to impact price action in the current market environment.

****Iterative Analysis:**** Review the financial data thoroughly, then focus your insights on the most material changes and catalysts.
Prioritize information that provides valuable context for trading strategy.

Apply your fundamental analysis expertise to extract trading-relevant insights.
Focus on corporate events, financial health trends, and performance indicators that could influence short-term trading decisions.

Consider earnings quality, balance sheet strength, cash flow sustainability, and any material changes that could serve as catalysts.
Your analysis should provide fundamental context that complements technical

```
and sentiment analysis.
```

```
**Remember:** Identify fundamental factors that could influence price action.  
Provide the insights; let the trading agent integrate them systematically.
```

Fundamental Analyst - Follow-up Prompt

The Follow-up Prompt governs incremental updates after initialization. It incorporates newly released fundamentals (filings, guidance, event deltas), reassesses material changes and catalysts, and refines the prior assessment while preserving methodological consistency. Emphasis is placed on short-horizon relevance and actionable context for the trading agent. The prompt is provided below.

Fundamental Analyst - Follow-up Prompt

```
## FUNDAMENTAL ANALYSIS UPDATE - {{ instrument }}  
**Timestamp:** {{ current_time }}  
  
Continue with your role as elite fundamental analyst. Apply the same analytical depth  
and professional standards established in the initial framework.  
  
## UPDATED FUNDAMENTALS  
{{ fundamental_data }}  
  
**Goal:** Identify the most significant fundamental developments and their potential  
market implications. Consider what fundamental information is most valuable  
for current trading context, then focus your analysis accordingly.  
  
Provide fundamental analysis focusing on material changes and trading implications.
```

Trading Prompt Optimizer (Adaptive-OPRO Target = CTA Initial Prompt)

The *Trading Prompt Optimizer* is the meta-policy that revises only the **static instruction block** of the Central Trading Agent's Initial Decision Prompt. At each window boundary it consumes a prompt-performance history (`history_text`) scored via the windowed ROI signal and proposes an edited template that preserves all placeholders (`{{...}}`), conditional blocks (`{% if %}`), and the order JSON schema (actions and order types). The optimizer returns a strictly structured JSON payload containing a diagnostic `performance_analysis`, a full `optimized_prompt` (template text, not a filled instance), `key_improvements`, and an `expected_impact`. An update is applied only if the placeholder set and interface remain unchanged, ensuring compatibility with the runtime injector.

Trading Prompt Optimizer's Prompt

```
# TRADING PROMPT OPTIMIZER  
  
**Primary Goal:** Optimize prompt context, information architecture,  
and decision-making frameworks. Enhanced context leads to better comprehension,  
deeper analysis, and superior trading decisions  
that naturally improve performance outcomes.  
  
**Performance Learning Context:**  
{{ history_text }}  
Note: Scores reflect cumulative ROI performance (0-100 scale). Higher scores indicate  
more effective prompt designs that enable better trading decisions.  
  
**Focus Areas:**  
- Strengthen the system role and trader identity  
- Optimize decision-making frameworks and criteria
```

```

- Enhance clarity of instructions and expectations
- Provide clearer guidance on analysis and decision-making process
- Better structure the flow from analysis to action

**Key Principles:**
- Ensure agent autonomy and adaptive thinking
- Avoid mandatory procedures or fixed thresholds
- Strengthen natural reasoning and market judgment
- Maintain clear constraints while allowing flexibility

**Critical Prompt Design Guidelines:**
- Keep prompts simple and direct: Models excel at understanding brief, clear instructions
- Be specific about end goals: Include specific parameters for successful decision-making
- Encourage iterative reasoning: Guide models to keep reasoning until they match
success criteria
- Use clear delimiters and structure to organize different sections appropriately

{% raw %}
**CRITICAL TEMPLATE PRESERVATION REQUIREMENTS:**
**WARNING**: Any modification to template variables will cause SYSTEM FAILURE
**FORBIDDEN**: Adding new {{ variable_name }} placeholders is STRICTLY FORBIDDEN
**FORBIDDEN**: Removing existing {{ variable_name }} placeholders is STRICTLY FORBIDDEN
**MANDATORY**: Copy ALL {{ variable_name }} placeholders EXACTLY as they appear in the
original template
**MANDATORY**: Preserve ALL {% if %} template blocks and <system_role> tags EXACTLY
- Maintain JSON format: BUY, SELL, SHORT, SHORT_COVER
- Keep order types: MARKET, LIMIT, STOP
- Ensure compatibility with interval-based decision cycles
{% endraw %}

**CRITICAL JSON FORMAT REQUIREMENTS:**
- Must be valid JSON with proper escaping
- Use \\n for newlines within string values
- Use \" for quotes within string values
- No unescaped newlines, tabs, or special characters
- Enclose the JSON in ```json and ``` code blocks

**Required JSON Output:**
```json
{
 "performance_analysis": "Comprehensive analysis of current template's contextual design
strengths, weaknesses, and enhancement opportunities",
 "optimized_prompt": "Complete improved TEMPLATE with better structure
(full template text with all placeholders preserved).
Use \\n for line breaks in the template text.",
 "key_improvements": "Specific structural and contextual transformations made
to optimize decision-making effectiveness",
 "expected_impact": "Expected improvements in comprehension, analytical depth,
and decision-making quality"
}
Important: Return a generic template, not a filled prompt.

```

## Weekly Reflection Agent

The *Weekly Reflection Agent* provides periodic (`{{reflection_interval}}`-day) reviews of recent trades and portfolio evolution, producing a single, compact paragraph that highlights recurring patterns, risk discipline, and thesis maintenance. Its output is *advisory* text only: it is injected as `reflection_analysis` for the Central Trading Agent to read on subsequent decisions, and it does not directly edit prompts or alter execution semantics. The reflection is derived from the full decision log and period summary, avoids prescriptive rules or rigid thresholds, and is designed to surface durable process improvements rather than post-hoc trade-by-trade commentary. By construction, it respects the fixed decision interval and order-cancellation rules described in the environment specification.

## Weekly Reflection Agent’s Prompt

```
ELITE TRADING COACH - {{ instrument }} INTERVAL REVIEW
Period: {{ reflection_interval }}-day review | **Session:** {{ current_time }}
| **Trading Decision Frequency:** {{ action_interval }}

You are a reflection agent analyzing {{ reflection_interval }} days of trading
performance to provide strategic insights for systematic improvement.

TRADING SYSTEM RULES & LIMITATIONS
Portfolio & Operational Context:
Single-Stock Portfolio: The agent manages a concentrated portfolio dedicated
exclusively to {{ instrument }} - all available capital and positions are focused
on this one security with no diversification across multiple stocks.
Available Actions: BUY, SELL, SHORT, SHORT_COVER
Order Types: MARKET, LIMIT, STOP
Constraints: Cash limits, position sizing rules, and {{ action_interval }}
decision intervals apply
Position Limits: SELL orders are automatically limited to current holdings,
and SHORT_COVER orders are automatically limited to current short positions -
overselling or over-covering is impossible.
The system enforces these limits automatically.
Critical Constraint: The agent can only make trading decisions at fixed
{{ action_interval }} intervals. All orders in the decision JSON are placed
simultaneously - there is no sequential order placement.
Order Auto-Cancellation: Unfilled orders are automatically cancelled at the end
of each decision interval.

PERIOD PERFORMANCE OVERVIEW
{{ period_summary }}

COMPLETE DECISION HISTORY FOR PERIOD
{{ complete_history }}

YOUR COACHING TASK

PURPOSE
In one comprehensive paragraph, synthesize the most impactful patterns from this
{{ reflection_interval }}-day period and identify the single structural improvement
that would most enhance future performance cycles.
Focus on systematic insights that will compound over multiple
{{ reflection_interval }}-day periods rather than individual trade critiques.

GUIDELINES
- Analyze decision patterns, risk management consistency, and strategic evolution
 across the period
- Identify the highest-leverage behavioral or strategic adjustment for future periods
- Emphasize enduring principles over isolated performance details
- Skip grades, personality assessments, or motivational language

REQUIRED OUTPUT FORMAT: Return only your reflection as a single paragraph
of continuous plain text (3-5 sentences).
```

## B Reproducibility

All experiments are conducted on a MacBook Pro with an Apple M3 Pro chip (11-core CPU) and 18 GB of unified memory. Our experiments are conducted using an updated version of the StockSim environment [51], with modifications to support the ATLAS multi-agent architecture, *Adaptive-OPRO* optimization, and reflection-based mechanisms (implementation details in code). An example configuration for GPT-o4-mini using *Adaptive-OPRO* on XOM is provided under `configs/o4-mini-adaptive-opro-config.yaml`. All other experimental configurations can be reproduced by following the StockSim documentation and adapting this sample.

We access LLaMA Claude models via Amazon Bedrock (Table 9.1). GPT models are accessed via OpenAI

---

Model ID	Model Card / Provider Identifier
LLaMA 3.3-70B	<code>meta.llama3-3-70b-instruct-v1:0</code>
Claude Sonnet 4	<code>anthropic.claude-sonnet-4-20250514-v1:0</code>

---

Table 9.1: Models accessed via Amazon Bedrock.

---

Model ID	Model Card / Docs
GPT-o4-mini	<code>gpt-4o-mini-2024-07-18</code>
GPT-o3	<code>gpt-o3-2025-04-16</code>

---

Table 9.2: Models accessed via OpenAI.

APIs (Table 9.2). We interface with all LLMs strictly through provider APIs and do not employ any local hardware or fine-tuning.

# Chapter 10

## Bibliography

- [1] Argyrou, G. et al. “Automatic Generation of Fashion Images Using Prompting in Generative Machine Learning Models”. In: *Computer Vision – ECCV 2024 Workshops*. Ed. by A. Del Bue et al. Cham: Springer Nature Switzerland, 2025, pp. 286–302. ISBN: 978-3-031-91569-7.
- [2] Atil, B. et al. *Non-Determinism of "Deterministic" LLM Settings*. 2025. arXiv: [2408.04667 \[cs.CL\]](#).
- [3] Austin, D. and Chartock, E. *GRAD-SUM: Leveraging Gradient Summarization for Optimal Prompt Engineering*. 2024. arXiv: [2407.12865 \[cs.CL\]](#).
- [4] Ba, J. L., Kiros, J. R., and Hinton, G. E. *Layer Normalization*. 2016. arXiv: [1607.06450 \[stat.ML\]](#).
- [5] Bengio, Y. et al. “A Neural Probabilistic Language Model”. In: *Journal of machine learning research*. 2003.
- [6] Bollinger, J. *Bollinger on Bollinger Bands*. McGraw-Hill Education, 2002. ISBN: 9780071373685.
- [7] Brown, T. B. et al. *Language Models are Few-Shot Learners*. 2020. arXiv: [2005.14165 \[cs.CL\]](#).
- [8] Byrd, D., Hybinette, M., and Balch, T. H. “ABIDES: Towards high-fidelity multi-agent market simulation”. In: *Proceedings of the 2020 ACM SIGSIM Conference on Principles of Advanced Discrete Simulation*. 2020, pp. 11–22.
- [9] Chen, D., Zhang, Q., and Zhu, Y. “Efficient Sequential Decision Making with Large Language Models”. In: *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*. Ed. by Y. Al-Onaizan, M. Bansal, and Y.-N. Chen. Miami, Florida, USA: Association for Computational Linguistics, Nov. 2024, pp. 9157–9170. DOI: [10.18653/v1/2024.emnlp-main.517](#).
- [10] Chen, T. et al. “Big Self-Supervised Models are Strong Semi-Supervised Learners”. In: *Advances in Neural Information Processing Systems*. Ed. by H. Larochelle et al. Vol. 33. Curran Associates, Inc., 2020, pp. 22243–22255.
- [11] Day, M.-Y. et al. “The profitability of bollinger bands trading bitcoin futures”. In: *Applied Economics Letters* 30.11 (2023), pp. 1437–1443.
- [12] DeepSeek-AI et al. *DeepSeek-R1: Incentivizing Reasoning Capability in LLMs via Reinforcement Learning*. 2025. arXiv: [2501.12948 \[cs.CL\]](#).
- [13] Devlin, J. et al. “BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding”. In: *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*. Ed. by J. Burstein, C. Doran, and T. Solorio. Minneapolis, Minnesota: Association for Computational Linguistics, June 2019, pp. 4171–4186. DOI: [10.18653/v1/N19-1423](#).
- [14] Ding, Q. et al. *TradExpert: Revolutionizing Trading with Mixture of Expert LLMs*. 2025. arXiv: [2411.00782 \[cs.AI\]](#).
- [15] Do, D. et al. *Large Language Models Prompting With Episodic Memory*. 2024. arXiv: [2408.07465 \[cs.CL\]](#).
- [16] Dong, Y. et al. “Generalization or Memorization: Data Contamination and Trustworthy Evaluation for Large Language Models”. In: *Findings of the Association for Computational Linguistics: ACL 2024*. Ed. by L.-W. Ku, A. Martins, and V. Srikumar. Bangkok, Thailand: Association for Computational Linguistics, Aug. 2024, pp. 12039–12050. DOI: [10.18653/v1/2024.findings-acl.716](#).

- [17] Evangelatos, A. et al. “AILS-NTUA at SemEval-2025 Task 8: Language-to-Code prompting and Error Fixing for Tabular Question Answering”. In: *Proceedings of the 19th International Workshop on Semantic Evaluation (SemEval-2025)*. Ed. by S. Rosenthal et al. Vienna, Austria: Association for Computational Linguistics, July 2025, pp. 1423–1435. ISBN: 979-8-89176-273-2. URL:
- [18] Fatouros, G. et al. *MarketSenseAI 2.0: Enhancing Stock Analysis through LLM Agents*. 2025. arXiv: [2502.00415 \[q-fin.CP\]](#).
- [19] Filandrianos, G. et al. *Bias Beware: The Impact of Cognitive Biases on LLM-Driven Product Recommendations*. 2025. arXiv: [2502.01349 \[cs.CL\]](#). URL:
- [20] Finn, C., Abbeel, P., and Levine, S. “Model-agnostic meta-learning for fast adaptation of deep networks”. In: *Proceedings of the 34th International Conference on Machine Learning - Volume 70*. ICML’17. Sydney, NSW, Australia: JMLR.org, 2017, pp. 1126–1135.
- [21] Frey, S. Y. et al. “JAX-LOB: A GPU-Accelerated limit order book simulator to unlock large scale reinforcement learning for trading”. In: *Proceedings of the Fourth ACM International Conference on AI in Finance*. 2023, pp. 583–591.
- [22] Gencay, R. “Non-linear prediction of security returns with moving average rules”. In: *Journal of Forecasting* 15.3 (1996), pp. 165–174.
- [23] Giadikiaroglou, P. et al. “Puzzle Solving using Reasoning of Large Language Models: A Survey”. In: *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*. Ed. by Y. Al-Onaizan, M. Bansal, and Y.-N. Chen. Miami, Florida, USA: Association for Computational Linguistics, Nov. 2024, pp. 11574–11591. DOI: [10.18653/v1/2024.emnlp-main.646](#). URL:
- [24] Guo, Q. et al. *EvoPrompt: Connecting LLMs with Evolutionary Algorithms Yields Powerful Prompt Optimizers*. 2025. arXiv: [2309.08532 \[cs.CL\]](#).
- [25] Gururangan, S. et al. “Don’t Stop Pretraining: Adapt Language Models to Domains and Tasks”. In: *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. Ed. by D. Jurafsky et al. Online: Association for Computational Linguistics, July 2020, pp. 8342–8360. DOI: [10.18653/v1/2020.acl-main.740](#).
- [26] Hasbrouck, J. *Empirical Market Microstructure: The Institutions, Economics, and Econometrics of Securities Trading*. Oxford University Press, 2007.
- [27] He, K. et al. *Deep Residual Learning for Image Recognition*. 2015. arXiv: [1512.03385 \[cs.CV\]](#).
- [28] Hochreiter, S. and Schmidhuber, J. “Long short-term memory”. In: *Neural computation* 9.8 (1997), pp. 1735–1780.
- [29] Kaplan, J. et al. *Scaling Laws for Neural Language Models*. 2020. arXiv: [2001.08361 \[cs.LG\]](#).
- [30] Karkani, D. et al. “AILS-NTUA at SemEval-2025 Task 3: Leveraging Large Language Models and Translation Strategies for Multilingual Hallucination Detection”. In: *Proceedings of the 19th International Workshop on Semantic Evaluation (SemEval-2025)*. Ed. by S. Rosenthal et al. Vienna, Austria: Association for Computational Linguistics, July 2025, pp. 1289–1305. ISBN: 979-8-89176-273-2. URL:
- [31] Kirtac, K. and Germano, G. “Enhanced Financial Sentiment Analysis and Trading Strategy Development Using Large Language Models”. In: *Proceedings of the 14th Workshop on Computational Approaches to Subjectivity, Sentiment, & Social Media Analysis*. Ed. by O. De Clercq et al. Bangkok, Thailand: Association for Computational Linguistics, Aug. 2024, pp. 1–10. DOI: [10.18653/v1/2024.wassa-1.1](#).
- [32] Kojima, T. et al. *Large Language Models are Zero-Shot Reasoners*. 2023. arXiv: [2205.11916 \[cs.CL\]](#).
- [33] Kong, A. et al. *Better Zero-Shot Reasoning with Role-Play Prompting*. 2024. arXiv: [2308.07702 \[cs.CL\]](#).
- [34] Kritharoula, A., Lymperaiou, M., and Stamou, G. *Language Models as Knowledge Bases for Visual Word Sense Disambiguation*. 2023. arXiv: [2310.01960 \[cs.CL\]](#). URL:
- [35] Kritharoula, A., Lymperaiou, M., and Stamou, G. “Large Language Models and Multimodal Retrieval for Visual Word Sense Disambiguation”. In: *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*. Ed. by H. Bouamor, J. Pino, and K. Bali. Singapore: Association for Computational Linguistics, Dec. 2023, pp. 13053–13077. DOI: [10.18653/v1/2023.emnlp-main.807](#). URL:
- [36] Lewis, M. et al. “BART: Denoising Sequence-to-Sequence Pre-training for Natural Language Generation, Translation, and Comprehension”. In: *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. Ed. by D. Jurafsky et al. Online: Association for Computational Linguistics, July 2020, pp. 7871–7880. DOI: [10.18653/v1/2020.acl-main.703](#).

- 
- [37] Li, H. et al. “INVESTORBENCH: A Benchmark for Financial Decision-Making Tasks with LLM-based Agent”. In: *arXiv preprint arXiv:2412.18174* (2024).
- [38] Li, Y. et al. “CryptoTrade: A Reflective LLM-based Agent to Guide Zero-shot Cryptocurrency Trading”. In: *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*. Ed. by Y. Al-Onaizan, M. Bansal, and Y.-N. Chen. Miami, Florida, USA: Association for Computational Linguistics, Nov. 2024, pp. 1094–1106. DOI: [10.18653/v1/2024.emnlp-main.63](https://doi.org/10.18653/v1/2024.emnlp-main.63).
- [39] Liu, H. et al. *Logical Reasoning in Large Language Models: A Survey*. 2025. arXiv: [2502.09100](https://arxiv.org/abs/2502.09100) [cs.AI].
- [40] Liu, P. et al. *Pre-train, Prompt, and Predict: A Systematic Survey of Prompting Methods in Natural Language Processing*. 2021. arXiv: [2107.13586](https://arxiv.org/abs/2107.13586) [cs.CL].
- [41] Liu, X.-Y. et al. “FinRL: A Deep Reinforcement Learning Library for Automated Stock Trading in Quantitative Finance”. In: *CoRR* (2020).
- [42] Liu, X.-Y. et al. “FinRL-Meta: Market environments and benchmarks for data-driven financial reinforcement learning”. In: *Advances in Neural Information Processing Systems* 35 (2022), pp. 1835–1849.
- [43] Liu, Z. et al. “Fin-r1: A large language model for financial reasoning through reinforcement learning”. In: *arXiv preprint arXiv:2503.16252* (2025).
- [44] Lu, G. et al. “BizFinBench: A Business-Driven Real-World Financial Benchmark for Evaluating LLMs”. In: *arXiv preprint arXiv:2505.19457* (2025).
- [45] Mascioli, C. et al. “A Financial Market Simulation Environment for Trading Agents Using Deep Reinforcement Learning”. In: *Proceedings of the 5th ACM International Conference on AI in Finance*. 2024, pp. 117–125.
- [46] Murphy, J. J. *Technical Analysis of the Financial Markets: A Comprehensive Guide to Trading Methods and Applications*. New York Institute of Finance, 1999.
- [47] Nafiu, A. et al. “Risk management strategies: Navigating volatility in complex financial market environments”. In: (2025).
- [48] Ouyang, L. et al. *Training language models to follow instructions with human feedback*. 2022. arXiv: [2203.02155](https://arxiv.org/abs/2203.02155) [cs.CL].
- [49] Panagiotopoulos, I. et al. “RISCORE: Enhancing In-Context Riddle Solving in Language Models through Context-Reconstructed Example Augmentation”. In: *Proceedings of the 31st International Conference on Computational Linguistics*. Ed. by O. Rambow et al. Abu Dhabi, UAE: Association for Computational Linguistics, Jan. 2025, pp. 9431–9455. URL:
- [50] Papadakis, C. et al. *ATLAS: Adaptive Trading with LLM Agents Through Dynamic Prompt Optimization and Multi-Agent Coordination*. 2025. arXiv: [2510.15949](https://arxiv.org/abs/2510.15949) [q-fin.TR]. URL:
- [51] Papadakis, C. et al. “StockSim: A Dual-Mode Order-Level Simulator for Evaluating Multi-Agent LLMs in Financial Markets”. In: *arXiv preprint arXiv:2507.09255* (2025).
- [52] Papadimitriou, C. et al. *Masked Generative Story Transformer with Character Guidance and Caption Augmentation*. 2024. arXiv: [2403.08502](https://arxiv.org/abs/2403.08502) [cs.CV]. URL:
- [53] Penman, S. H. *Financial Statement Analysis and Security Valuation*. 5th. New York, NY: McGraw-Hill Education, 2012. ISBN: 978-0078025310.
- [54] Radford, A. and Narasimhan, K. “Improving Language Understanding by Generative Pre-Training”. In: 2018.
- [55] Radford, A. et al. “Language Models are Unsupervised Multitask Learners”. In: 2019.
- [56] Raffel, C. et al. “Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer”. In: *CoRR* abs/1910.10683 (2019). arXiv: [1910.10683](https://arxiv.org/abs/1910.10683).
- [57] Raptopoulos, P. et al. *PAKTON: A Multi-Agent Framework for Question Answering in Long Legal Agreements*. 2025. arXiv: [2506.00608](https://arxiv.org/abs/2506.00608) [cs.CL]. URL:
- [58] Rudkin, S., Qiu, W., and Dłotko, P. “Uncertainty, volatility and the persistence norms of financial time series”. In: *Expert Systems with Applications* 223 (2023), p. 119894.
- [59] Schaeffer, R., Miranda, B., and Koyejo, S. *Are Emergent Abilities of Large Language Models a Mirage?* 2023. arXiv: [2304.15004](https://arxiv.org/abs/2304.15004) [cs.AI].
- [60] Schulman, J. et al. *Proximal Policy Optimization Algorithms*. 2017. arXiv: [1707.06347](https://arxiv.org/abs/1707.06347) [cs.LG].
- [61] Singh, A. K. et al. “Evaluation data contamination in LLMs: how do we measure it and (when) does it matter?” In: *arXiv preprint arXiv:2411.03923* (2024).
- [62] Song, Y. et al. “The Good, The Bad, and The Greedy: Evaluation of LLMs Should Not Ignore Non-Determinism”. In: *Proceedings of the 2025 Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*.
-

- Ed. by L. Chiruzzo, A. Ritter, and L. Wang. Albuquerque, New Mexico: Association for Computational Linguistics, Apr. 2025, pp. 4195–4206. ISBN: 979-8-89176-189-6. DOI: [10.18653/v1/2025.naacl-long.211](https://doi.org/10.18653/v1/2025.naacl-long.211).
- [63] Stringli, E. et al. “Pitfalls of Scale: Investigating the Inverse Task of Redefinition in Large Language Models”. In: *Findings of the Association for Computational Linguistics: ACL 2025*. Ed. by W. Che et al. Vienna, Austria: Association for Computational Linguistics, July 2025, pp. 9445–9469. ISBN: 979-8-89176-256-5. DOI: [10.18653/v1/2025.findings-acl.492](https://doi.org/10.18653/v1/2025.findings-acl.492). URL:
  - [64] Thomas, K. et al. ““I Never Said That”: A dataset, taxonomy and baselines on response clarity classification”. In: *Findings of the Association for Computational Linguistics: EMNLP 2024*. Ed. by Y. Al-Onaizan, M. Bansal, and Y.-N. Chen. Miami, Florida, USA: Association for Computational Linguistics, Nov. 2024, pp. 5204–5233. DOI: [10.18653/v1/2024.findings-emnlp.300](https://doi.org/10.18653/v1/2024.findings-emnlp.300). URL:
  - [65] Vaswani, A. et al. “Attention is All you Need”. In: *Advances in Neural Information Processing Systems*. Ed. by I. Guyon et al. Vol. 30. Curran Associates, Inc., 2017.
  - [66] Wang, J. and Kim, J. “Predicting stock price trend using macd optimized by historical volatility”. In: *Mathematical Problems in Engineering* 2018 (2018), pp. 1–12.
  - [67] Wang, X. et al. *Self-Consistency Improves Chain of Thought Reasoning in Language Models*. 2023. arXiv: [2203.11171](https://arxiv.org/abs/2203.11171) [cs.CL].
  - [68] Wei, J. et al. *Emergent Abilities of Large Language Models*. 2022. arXiv: [2206.07682](https://arxiv.org/abs/2206.07682) [cs.CL].
  - [69] Wei, J. et al. *Finetuned Language Models Are Zero-Shot Learners*. 2022. arXiv: [2109.01652](https://arxiv.org/abs/2109.01652) [cs.CL].
  - [70] Wei, J. et al. *Chain-of-Thought Prompting Elicits Reasoning in Large Language Models*. 2023. arXiv: [2201.11903](https://arxiv.org/abs/2201.11903) [cs.CL].
  - [71] White, C. et al. “Livebench: A challenging, contamination-free llm benchmark”. In: *arXiv preprint arXiv:2406.19314* 4 (2024).
  - [72] Wilder, J. W. *New Concepts in Technical Trading Systems*. Trend Research, 1978.
  - [73] Wu, S. et al. *BloombergGPT: A Large Language Model for Finance*. 2023. arXiv: [2303.17564](https://arxiv.org/abs/2303.17564) [cs.LG].
  - [74] Xiao, Y. et al. *TradingAgents: Multi-Agents LLM Financial Trading Framework*. 2025. arXiv: [2412.20138](https://arxiv.org/abs/2412.20138) [q-fin.TR].
  - [75] Xie, S. M. et al. *An Explanation of In-context Learning as Implicit Bayesian Inference*. 2022. arXiv: [2111.02080](https://arxiv.org/abs/2111.02080) [cs.CL].
  - [76] Xiong, G. et al. *FLAG-Trader: Fusion LLM-Agent with Gradient-based Reinforcement Learning for Financial Trading*. 2025. arXiv: [2502.11433](https://arxiv.org/abs/2502.11433) [cs.AI].
  - [77] Xu, F. et al. *Towards Large Reasoning Models: A Survey of Reinforced Reasoning with Large Language Models*. 2025. arXiv: [2501.09686](https://arxiv.org/abs/2501.09686) [cs.AI].
  - [78] Yadav, G. S., Guha, A., and Chakrabarti, A. S. “Measuring complexity in financial data”. In: *Frontiers in Physics* 8 (2020), p. 339.
  - [79] Yang, C. et al. *Large Language Models as Optimizers*. 2024. arXiv: [2309.03409](https://arxiv.org/abs/2309.03409) [cs.LG].
  - [80] Yang, H., Liu, X.-Y., and Wang, C. D. *FinGPT: Open-Source Financial Large Language Models*. 2023. arXiv: [2306.06031](https://arxiv.org/abs/2306.06031) [q-fin.ST].
  - [81] Yang, Y., UY, M. C. S., and Huang, A. *FinBERT: A Pretrained Language Model for Financial Communications*. 2020. arXiv: [2006.08097](https://arxiv.org/abs/2006.08097) [cs.CL].
  - [82] Yao, S. et al. *Tree of Thoughts: Deliberate Problem Solving with Large Language Models*. 2023. arXiv: [2305.10601](https://arxiv.org/abs/2305.10601) [cs.CL].
  - [83] Yu, Y. et al. *FinMem: A Performance-Enhanced LLM Trading Agent with Layered Memory and Character Design*. 2023. arXiv: [2311.13743](https://arxiv.org/abs/2311.13743) [q-fin.CP].
  - [84] Zheng, H. S. et al. *Take a Step Back: Evoking Reasoning via Abstraction in Large Language Models*. 2024. arXiv: [2310.06117](https://arxiv.org/abs/2310.06117) [cs.LG].