

Graphical Dual Axial Transformer for Risk-adjusted Portfolio Optimization

1st Youjia Liu*

*Graduate School of Economics and Management
Tohoku University
Sendai, Japan
liu.you-jia.t3@dc.tohoku.ac.jp*

2nd Yasumasa Matsuda

*Graduate School of Economics and Management
Tohoku University
Sendai, Japan
yasumasa.matsuda.a4@tohoku.ac.jp*

Abstract—Portfolio optimization with financial time series requires jointly modeling temporal dynamics and cross-sectional asset interactions under noisy and non-stationary market conditions. Existing deep learning approaches often struggle to disentangle these heterogeneous dependencies or rely on reinforcement learning frameworks that lead to unstable training. We propose the Graphical Dual Axial Transformer (GDAT), which models financial data as a unified 3D tensor across assets, time, and features. GDAT employs a dual axial attention mechanism to disentangle temporal dependencies from cross-sectional correlations, while enabling structured information flow between axes. To stabilize relationship selection and filter market noise, the model incorporates a graph-structured inductive bias derived from sparse inverse covariance estimation, guiding attention toward economically meaningful interactions. By risk-adjusted return optimization, the proposed framework allows stable end-to-end training. Experiments on real-world datasets, including S&P 500, Russell 1000, and FTSE 100, demonstrate that GDAT consistently outperforms classic online learning and deep learning methods in terms of risk-adjusted performance and robustness. These results demonstrate the benefits of structured attention and cross-sectional modeling for portfolio optimization.

Index Terms—Attention, Transformer, Axial, Portfolio Optimization

I. INTRODUCTION

Modern portfolio optimization is driven by the joint cross-sectional and temporal dependence structure of asset returns, as originally formalized in Modern Portfolio Theory [1]. While classical econometric models, such as multivariate GARCH [2] and factor-based estimators [3], provide a foundation for modeling time-varying risk, their reliance on rigid parametric assumptions limits their ability to capture the high-dimensional, nonlinear interactions prevalent in modern financial markets.

Recent advances in deep learning have excelled in capturing spatio-temporal dependencies across multiple feature channels. However, financial applications attempting to jointly model these interactions are largely confined to deep reinforcement learning (DRL) [4], [5]. Despite its flexibility, DRL suffers from unstable training in noisy environments and high sensitivity to reward specifications [6]. Furthermore, most DRL approaches optimize for accumulated returns rather than the Sharpe ratio, as the latter is a non-additive, long-horizon metric that does not naturally decompose into step-wise rewards [7].

Moreover, the optimization objectives commonly adopted in DRL-based portfolio models can present practical challenges in financial applications. Portfolio optimization methods vary in their choice of objective function. Certain models target accumulated returns [4], whereas others employ risk-adjusted metrics, including the Sharpe ratio [8]. The sensitivity of these objectives to market noise and covariance estimation errors often leads to suboptimal allocations, highlighting the intrinsic difficulty of designing robust end-to-end optimization targets in non-stationary environments.

In this paper, we propose a novel GDAT model for portfolio weight generation, designed to capture both sequential patterns and cross-sectional asset interactions in financial time series. Specifically, GDAT employs a dual axial attention mechanism to disentangle temporal dependencies from asset-wise correlations, while a structured attention architecture enables directional information flow from the cross-sectional axis to the temporal axis. By decoupling representation learning from risk-adjusted return optimization, our framework allows stable end-to-end training and jointly models asset-level interactions, temporal dynamics, and heterogeneous feature channels. Empirical evaluations on real-world datasets, including S&P 500, Russell 1000, and FTSE 100, demonstrate that GDAT consistently outperforms baselines in portfolio optimization tasks.

Our main contributions are summarized as follows. First, we propose GDAT, which to our knowledge is among the earliest methods to adapt axial attention for portfolio optimization, specifically designed to disentangle cross-sectional and temporal dependencies in a risk-adjusted framework. Second, by using dual axial attention and graph sparsity, GDAT focuses self-attention on key relationships to improve allocation decisions. Third, the model effectively captures stable cross-sectional relationships despite high market noise, demonstrating superior robustness. Fourth, we model financial data as a unified 3D tensor over assets, time, and features, enabling structured representation learning for end-to-end optimization.

II. RELATED WORKS

A. Deep Learning with Portfolio Optimization

Deep learning has been increasingly adopted in portfolio optimization as an alternative to classical mean-variance and factor-based approaches, enabling more flexible modeling of

*Corresponding author

nonlinear return dynamics and asset dependencies. End-to-end frameworks have been proposed to directly output portfolio weights and optimize portfolio-level risk-adjusted objectives, thereby avoiding explicit return forecasting and traditional quadratic programming procedures [9]. Building on this idea, transformer-based architectures introduce self-attention mechanisms to model temporal dependencies and cross-asset relationships within an end-to-end learning framework [4].

A central challenge in deep learning-based portfolio optimization is aligning objectives with risk–return trade-offs rather than standard regression losses. Early works explored expected return maximization [10], while recent approaches incorporate risk via mean–variance utility [9] or deep reinforcement learning [4]. To address the instability of noisy covariance estimates, recent studies have optimized the Sharpe ratio by incorporating shrinkage-based estimators [8], [11], ensuring more robust end-to-end portfolio construction.

B. Axial Transformer

Multidimensional attention mechanisms include several variants, among which dual attention and axial attention (e.g., Axial Transformers) are widely studied. Axial attention factorizes multi-dimensional self-attention into a sequence of one-dimensional attention operations along individual axes, as proposed in [12]. This design has been successfully applied to high-dimensional visual data such as images and videos.

Following this line of work, axial or gated axial attention mechanisms have been adopted in medical imaging tasks, including medical image segmentation and 3D brain tumor analysis, where capturing long-range spatial dependencies is essential [13]. Several recent surveys further summarize axial or axial-like attention modules as effective components in transformer-based medical vision models. [14] [15]

In financial applications, the use of axial attention remains relatively limited. A gated axial attention architecture, which separately attends over feature and temporal dimensions for stock price movement prediction, has been proposed [16]. Beyond such isolated efforts, axial attention has not yet been extensively explored in financial modeling.

III. MODEL

A. Problem Formulation

Let $\mathcal{U} = \{a_1, \dots, a_N\}$ be a universe of N tradable assets observed over a look-back window of T time steps. For each asset i at time t , we observe a feature vector $\mathbf{x}_{i,t} \in \mathbb{R}^E$, where E denotes the number of feature channels. The complete input is represented as a three-dimensional tensor:

$$\mathbf{X} \in \mathbb{R}^{N \times T \times E}. \quad (1)$$

The objective is to learn an end-to-end mapping

$$\mathcal{F} : \mathbf{X} \mapsto \mathbf{w}, \quad (2)$$

where $\mathbf{w} \in \mathbb{R}^N$ is the portfolio weight vector satisfying the budget constraint

$$\sum_{i=1}^N w_i = 1, \quad w_i \geq 0. \quad (3)$$

Unlike reinforcement learning frameworks, we optimize $\mathcal{F}$ by directly minimizing the negative risk-adjusted return objectives, ensuring stable training dynamics through representation learning rather than policy iteration.

B. Axial Factorized Representation Learning

The overall architecture of GDAT is illustrated in Fig. 1. To model the high-dimensional tensor $\mathbf{X}$ efficiently, we adopt an axial factorization strategy that decomposes joint dependencies across feature, asset, and temporal dimensions.

1) *Feature-wise Embedding*: For each asset i at time t , the feature vector $\mathbf{x}_{i,t} \in \mathbb{R}^E$ is mapped into a latent space of dimension d_{model} via a shared nonlinear embedding function:

$$\phi : \mathbb{R}^E \rightarrow \mathbb{R}^{d_{\text{model}}}, \quad (4)$$

yielding the initial latent representation:

$$\mathbf{h}_{i,t}^{(0)} = \phi(\mathbf{x}_{i,t}). \quad (5)$$

The collection of these vectors across all assets and time steps forms the initial latent tensor:

$$\mathbf{H}^{(0)} = \{\mathbf{h}_{i,t}^{(0)}\}_{i=1, t=1}^{N, T} \in \mathbb{R}^{N \times T \times d_{\text{model}}}. \quad (6)$$

To model the high-dimensional tensor $\mathbf{X}$ efficiently, we apply axial attention along the asset (cross-sectional) and temporal dimensions.

2) *Axial Positional Encoding*: To preserve asset identity and temporal ordering without collapsing the axial structure, we incorporate learnable axial positional embeddings. The position-aware representation $\mathbf{U}_{i,t} \in \mathbb{R}^{d_{\text{model}}}$ is defined as:

$$\mathbf{U}_{i,t} = \mathbf{h}_{i,t}^{(0)} + \mathbf{p}_i^{(N)} + \mathbf{p}_t^{(T)}, \quad (7)$$

where $\mathbf{p}_i^{(N)} \in \mathbb{R}^{d_{\text{model}}}$ and $\mathbf{p}_t^{(T)} \in \mathbb{R}^{d_{\text{model}}}$ denote the learnable embedding vectors associated with the i -th asset and the t -th time step, respectively. This decomposition allows the model to retain structural priors along each axis independently while keeping the parameter space efficient.

C. Graphical Lasso for Structural Masking

To mitigate the impact of high-frequency noise and spurious correlations prevalent in financial time series, we impose a sparse interaction prior along the asset dimension. Let $\mathbf{S} \in \mathbb{R}^{N \times N}$ denote the empirical covariance matrix of asset returns. In our offline training framework, we estimate a sparse precision matrix $\hat{\Theta}$ computed over the entire training period by solving the Graphical Lasso problem [17]:

$$\hat{\Theta} = \arg \min_{\Theta \succeq 0} \{-\log \det(\Theta) + \text{tr}(\mathbf{S}\Theta) + \rho \|\Theta\|_1\}, \quad (8)$$

where $\rho > 0$ controls the sparsity level of the estimated conditional dependency graph.

The resulting precision matrix induces a structural mask

$$\mathbf{M}_N(i, j) = \mathbb{I}(\hat{\Theta}_{i,j} \neq 0), \quad (9)$$

serves as a geometric constraint for the attention mechanism. Compared with a dense attention layer that exhaustively models all N^2 pairwise interactions, $\mathbf{M}_N$ instead induces a

structured search space informed by long-term conditional dependencies. This prior encourages the model to concentrate its representational capacity on economically meaningful asset clusters (e.g., sector-based groupings), while retaining the flexibility to learn dynamic interactions within each cluster.

D. Axial Attention Modules

1) *Complexity Efficiency*: Axial decomposition significantly reduces computational overhead compared to standard 3D self-attention. While vanilla Transformers scale quadratically with the product TN at $\mathcal{O}(T^2N^2d_{\text{model}})$, our approach decouples attention into cross-sectional and temporal axes, reducing complexity to:

$$\mathcal{O}\left(\underbrace{TN^2d_{\text{model}}}_{\text{Cross-sectional}} + \underbrace{NT^2d_{\text{model}}}_{\text{Temporal}}\right) = \mathcal{O}(TNd_{\text{model}}(N + T)). \quad (10)$$

This provides substantial efficiency gains for large asset universes (e.g., $N = 500$). Furthermore, this decomposition introduces a structural inductive bias: the model synchronizes market-wide information via $\mathbf{M}_N$ before refining individual asset trajectories, thereby suppressing spurious cross-temporal correlations.

2) *Cross-sectional Attention*: For each time step t , the asset representations are updated via N-axis Multi-Head Self-Attention (MHSA) constrained by the structural mask $\mathbf{M}_N$, which encodes statistically significant conditional dependencies. We adopt Pre-Layer Normalization (Pre-LN) and residual connections for numerical stability:

$$\mathbf{H}'_t = \mathbf{H}_t + \text{Dropout}\left(\text{Linear}\left(\text{MHSA}_N(\text{LN}(\mathbf{H}_t), \mathbf{M}_N)\right)\right), \quad (11)$$

where the attention incorporates the graph prior as an additive penalty:

$$\text{Attn}_N = \text{Softmax}\left(\frac{\mathbf{Q}_N \mathbf{K}_N^\top}{\sqrt{d_{\text{model}}}} + \mathcal{G}(\mathbf{M}_N)\right) \mathbf{V}_N. \quad (12)$$

Here, $\mathcal{G}(\cdot)$ is a mapping function that transforms the binary structural mask into an additive penalty matrix, where $\mathcal{G}(\mathbf{M}_N)_{i,j} = 0$ if $(\mathbf{M}_N)_{i,j} = 1$ (indicating an active edge), and $\mathcal{G}(\mathbf{M}_N)_{i,j} = -\infty$ otherwise. It is important to note that the structural mask $\mathbf{M}_N$ does not dictate static attention weights. Instead, it acts as a soft-prior filter. Within the active edges where $\mathbf{M}_N(i, j) = 1$, the attention mechanism remains purely data-driven, allowing the model to dynamically adjust the strength of interactions based on the specific market conditions observed at time t .

3) *Temporal Attention*: For each asset i , the temporal trajectory $\mathbf{H}_i \in \mathbb{R}^{T \times d_{\text{model}}}$ is refined via masked temporal attention to enforce causality:

$$\mathbf{H}'_i = \mathbf{H}_i + \text{Dropout}\left(\text{Linear}\left(\text{MHSA}_T(\text{LN}(\mathbf{H}_i), \mathbf{M}_T)\right)\right), \quad (13)$$

where $\mathbf{M}_T$ is a lower-triangular mask preventing look-ahead bias. Stacking these blocks allows the model to capture nonlinear temporal patterns and momentum dynamics without requiring a recurrent architecture.

TABLE I
ASSET COUNTS AND CHRONOLOGICAL TRAIN/VALIDATION/TEST SPLITS
ACROSS DATASETS.

	S&P 500	Russell 1000	FTSE 100
Train	1990-01 – 2015-12	1990-01 – 2015-12	2001-01 – 2015-12
Validation	2016-01 – 2018-12	2016-01 – 2018-12	2016-01 – 2018-12
Test	2019-01 – 2025-06	2019-01 – 2025-06	2019-01 – 2025-06
Assets	237	350	69

E. Portfolio Construction Block

To generate the final portfolio weights, the model extracts a representative feature vector for each asset. Specifically, the representation sequence for asset i is sliced along the temporal dimension to obtain the last hidden vector:

$$\mathbf{z}_i = \mathbf{h}'_{i,T} \in \mathbb{R}^{d_{\text{model}}}, \quad (14)$$

where $\mathbf{z}_i$ encapsulates the cumulative historical information of asset i within the look-back window. This vector is passed through a feed-forward network (FFN), $g(\cdot)$, to produce an unnormalized score $s_i = g(\mathbf{z}_i)$.

The final portfolio weights are obtained via softmax normalization across the asset universe:

$$w_i = \frac{\exp(s_i)}{\sum_{j=1}^N \exp(s_j)}, \quad (15)$$

ensuring the budget constraint $\sum w_i = 1$ and the long-only constraint $w_i \geq 0$. The resulting weight vector $\mathbf{w} = [w_1, \dots, w_N]^\top$ is used for subsequent rebalancing, with the entire framework trained end-to-end by directly optimizing a differentiable risk-adjusted objective.

IV. EXPERIMENTS

A. Experimental Setting

Model selection is performed based on validation performance. All experiments are conducted using 10 different random seeds, and the reported results correspond to the average performance across runs.

To ensure causality, portfolio weights $\mathbf{w}_t$ are generated in a one-step-ahead manner using only information available up to time t . These weights are held fixed for the subsequent period to calculate the realized portfolio return $r_{p,t+1} = \mathbf{w}_t^\top \mathbf{r}_{t+1}$. This rolling procedure is applied sequentially across the test horizon to produce a time-varying return sequence $\{r_{p,t+1}\}$, effectively preventing look-ahead bias and reflecting practical investment constraints.

We evaluate our model across the S&P 500, Russell 1000, and FTSE 100 indices, with dataset specifications and temporal splits detailed in Table I. To ensure comparability, we align the validation and test periods across all markets. The model input for each asset consists of concatenated daily returns and intraday high-low ranges, jointly capturing inter-day dynamics and intraday volatility to provide richer signals than return-only baselines.

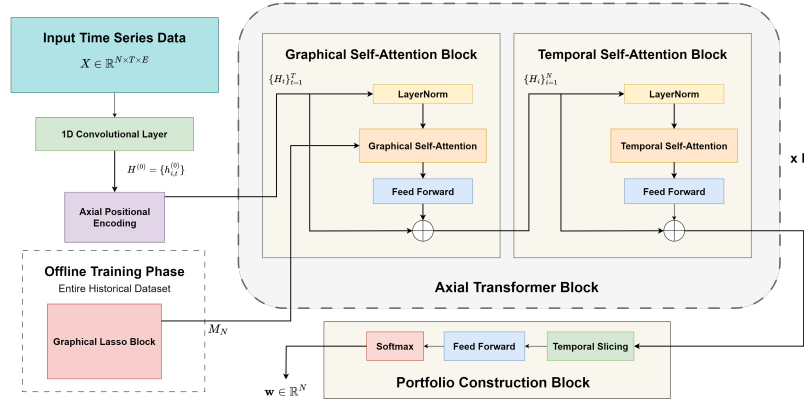


Fig. 1. The proposed Graphical Dual Axial Transformer architecture for portfolio optimization

B. Optimization Problem Setting

We evaluate the impact of different risk-adjusted objectives on portfolio allocation policies. Following [9], we adopt objectives that integrate realized returns with risk measures directly. Notably, this formulation facilitates end-to-end learning while bypassing the explicit estimation of high-dimensional covariance matrices, which are often numerically unstable in large-scale asset allocation. We compare the empirical performance of these objectives in Table II and Table III.

a) *Negative Mean-Variance*: We adopt an end-to-end portfolio optimization framework where the network directly outputs portfolio weights $\mathbf{w}_t \in \mathbb{R}^N$. Let $\mathbf{r}_{t+1}$ denote the asset return vector at time $t + 1$, and $r_{p,t+1} = \mathbf{w}_t^\top \mathbf{r}_{t+1}$ be the realized portfolio return. To maximize the risk-adjusted return, we minimize the negative mean-variance utility:

$$\mathcal{L}_{\text{MV}}(\mathbf{w}_t) = - \left(\mathbf{w}_t^\top \mathbf{r}_{t+1} - \frac{\lambda}{2} \text{Var}(r_{p,t+1}) \right), \quad (16)$$

where λ is the risk-aversion parameter. Unless otherwise specified, the risk-free rate r_f is set to zero.

b) *Negative Sharpe Ratio*: Ratio-based objectives like the Sharpe ratio provide a direct measure of reward-to-variability but can suffer from numerical instability during early training due to vanishing or exploding gradients. We address this by directly optimizing the negative Sharpe ratio:

$$\mathcal{L}_{\text{Sharpe}}(\mathbf{w}_t) = - \frac{\mathbf{w}_t^\top \mathbf{r}_{t+1} - r_f}{\text{Std}(r_{p,t+1}) + \epsilon}, \quad (17)$$

where ϵ is a small constant introduced for numerical stability. This loss ensures the model converges toward a portfolio with the maximum risk-adjusted performance while maintaining backpropagation stability. The models are trained using the negative Sharpe ratio loss to target the Maximum Sharpe Ratio Portfolio (MSRP).

C. Evaluation Metrics

1) *Evaluation Metrics*: Following empirical finance standards, we evaluate portfolio performance through multiple dimensions. We report the **Expected Return (ER)**, calculated

as the average realized portfolio return, and the **Sharpe Ratio (SR)** to assess risk-adjusted efficiency. To capture downside and tail risks, we employ the **Sortino Ratio** and **Maximum Drawdown (MDD)**, while total volatility and profitability frequency are measured by **Standard Deviation (Std)** and the **Positive Return Ratio**, respectively. Trading efficiency is monitored via **Turnover**, reflecting the average magnitude of daily weight rebalancing. Finally, investor utility is assessed via the **Mean-Variance Utility (MV)**, defined as $U = \bar{r}_p - \frac{\lambda}{2} \sigma_p^2$, where $\bar{r}_p$ and σ_p^2 denote the average return and variance of the portfolio, and λ is the risk-aversion coefficient.

D. Baselines and Benchmark Strategies

We compare the proposed model against a comprehensive set of baseline and benchmark strategies, including (1) online learning and traditional portfolio strategies: Uniform Constant Rebalanced Portfolio (UCRP) [18], Global Minimum Variance Portfolio (GMVP) [19], Anticor [20], Subset Resampling Portfolio (SSR) [21], On-Line Moving Average Reversion (OLMAR) [22], and Robust Median Reversion (RMR) [23]; (2) deep learning based methods: Temporal Convolutional Network (TCN) [24], Gated Recurrent Unit (GRU) [25], Long Short-Term Memory (LSTM) [26], and Transformer [27].

E. Main Results: Comparative Performance

1) *Back Testing*: The backtesting results for all methods are presented in Table II. Overall, GDAT outperforms all other baselines, demonstrating a superior Sharpe ratio and Sortino ratio. In comparison, online learning-based methods perform poorly as they do not explicitly optimize the Sharpe ratio, leading to suboptimal portfolio decisions due to a low return per unit of risk. GDAT surpasses all other models owing to its ability to simultaneously capture both cross-sectional and temporal relations through the axial self-attention mechanism.

2) *Robustness and Generalization Analysis*: Table III provides a multi-dimensional analysis of GDAT focusing on risk-aversion sensitivity and cross-market generalization. First, sensitivity analysis across varying parameters $\lambda \in \{1, 2, 5\}$ demonstrates GDAT's stability under diverse investor profiles

TABLE II

PERFORMANCE COMPARISON OF PORTFOLIO STRATEGIES ON THE S&P 500 DATASET. ALL DEEP LEARNING MODELS ARE OPTIMIZED USING THE MSRP OBJECTIVE. ER, STD, SR, AND SORTINO ARE ANNUALIZED.

Method	ER	Std	SR	Sortino	Positive Ratio (%)	MDD	Turnover
Baselines							
UCRP	0.142	0.203	0.697	1.060	54.110	0.383	0.000
GMVP	0.059	0.169	0.351	0.632	52.638	0.304	0.197
OLMAR	0.263	0.311	0.844	0.985	54.997	0.570	0.368
RMR	0.276	0.310	0.888	1.059	55.426	0.569	0.388
Anticor	0.125	0.229	0.548	0.797	53.893	0.321	1.521
SSR	0.088	0.161	0.547	1.086	54.016	0.292	0.392
MSRP							
GRU	0.190	0.230	0.824	1.160	53.325	0.379	0.139
LSTM	0.223	0.243	0.920	1.238	53.472	0.374	0.144
TCN	0.242	0.258	0.939	1.188	54.147	0.442	0.080
Transformer	0.289	0.311	0.931	1.004	52.945	0.451	0.175
GDAT	0.261	0.259	1.007	1.305	53.528	0.429	0.633

TABLE III

COMPARISON OF MV UTILITY ACROSS DIFFERENT RISK-AVERSION PARAMETERS λ . ALL MODELS ARE TRAINED BY DIRECTLY MINIMIZING THE NEGATIVE MV LOSS. RESULTS ACROSS THREE INDICES DEMONSTRATE GDAT'S ROBUSTNESS ACROSS DIFFERENT DATASETS.

Dataset	Model	$\lambda = 1$	$\lambda = 2$	$\lambda = 5$
S&P 500	TCN	0.1438	0.1233	0.0584
	GRU	0.1317	0.1139	0.0789
	LSTM	0.1420	0.1132	0.0746
	Transformer	0.1354	0.1438	0.0706
	GDAT	0.1504	0.1488	0.0808
Russell 1000	TCN	0.1711	0.1459	0.0798
	GRU	0.1426	0.1278	0.0690
	LSTM	0.1525	0.1277	0.0247
	Transformer	0.1499	0.1292	-0.0063
	GDAT	0.2090	0.1721	0.0941
FTSE 100	TCN	0.1279	0.1017	0.0209
	GRU	0.1361	0.0831	0.0309
	LSTM	0.1411	0.1089	-0.0023
	Transformer	0.1296	0.0790	-0.0428
	GDAT	0.1497	0.1138	0.0330

within the Mean-Variance framework. Second, GDAT consistently achieves superior utility across larger (Russell 1000) and international (FTSE 100) indices, underscoring its ability to extract universal asset dependencies regardless of market scale or geography. These results illustrate GDAT's capacity to maintain a superior efficient frontier across varying risk levels, validating its robustness in balancing the fundamental risk–return trade-off.

F. Rolling Window Analysis

To evaluate long-term strategic stability over the 2019–2025 period, which encompasses diverse market regimes, we conduct a rolling window analysis with window size 512 to ensure stable risk estimation. This evaluation horizon includes major market disruptions, such as the 2020 pandemic shock and the subsequent high-inflation environment.

As shown in Table IV, GDAT consistently achieves the highest average Sharpe ratio, demonstrating superior robustness

and positioning it as a more reliable candidate for institutional long-term asset allocation.

TABLE IV

AVERAGE ANNUALIZED METRICS OVER A 512-DAY ROLLING WINDOW (2019–2025) UNDER THE MSRP OBJECTIVE ON S&P 500.

Portfolio	Avg. ER	Avg. Std	Avg. SR	Avg. MDD
EWP	0.1369	0.1934	0.7068	0.2195
GMVP	0.0364	0.1653	0.2496	0.2118
OLMAR	0.2247	0.2661	0.8659	0.2847
RMR	0.2353	0.2637	0.9085	0.2852
Anticor	0.1269	0.2173	0.5902	0.2306
SSR	0.0724	0.1535	0.4646	0.1855
TCN	0.2411	0.2524	0.9559	0.2696
GRU	0.1685	0.2221	0.7516	0.2464
LSTM	0.2086	0.2366	0.8741	0.2532
Transformer	0.2779	0.3059	0.8919	0.3351
GDAT	0.2510	0.2471	1.0267	0.2308

G. Ablation Study

1) *Effect of Dual Axial Attention (GDAT)*: Table V evaluates GDAT variants to assess the contributions of dual axial attention. Removing either the cross-sectional component or the temporal component reduces performance compared to the full model. This confirms that jointly modeling temporal dynamics and asset-wise interactions is crucial for effective risk-adjusted portfolio optimization.

2) *Effect of Graphical Lasso Sparsity (ρ)*: Figure 2 illustrates the non-monotonic relationship between the Graphical Lasso regularization parameter ρ and portfolio performance. Under-regularization (small ρ) results in noisy attention by capturing spurious correlations, leading to unstable weights. Conversely, over-regularization (large ρ) discards informative dependencies, degrading the model's predictive power. Optimal sparsity allows GDAT to stabilize the attention mechanism by filtering idiosyncratic shocks while emphasizing robust, sector-aligned relationships. This demonstrates the necessity of the proposed structural masking for effective cross-asset modeling.

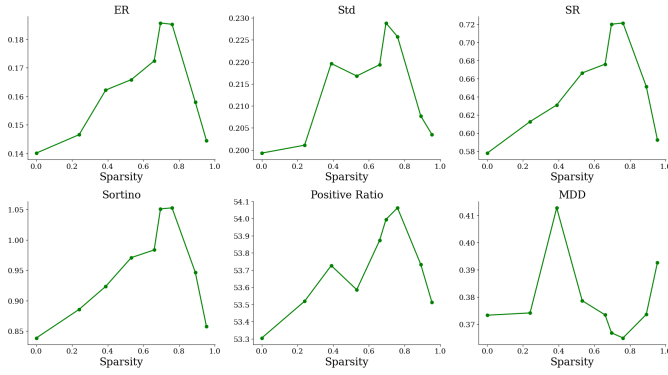


Fig. 2. Ablation study of the Graphical Lasso sparsity on the FTSE100 dataset. Performance metrics are shown across different sparsity levels.

V. CONCLUSION

This paper proposes a novel GDAT model for the generation of optimal portfolio weights. It is designed to concurrently capture sequential patterns, model inter-asset correlations, and optimize portfolio decisions with respect to the Sharpe ratio. By jointly leveraging both Sharpe ratio and mean-variance losses, the proposed method achieves substantial improvements in portfolio selection performance. Extensive empirical evaluations on real-world datasets, including S&P 500, Russell 1000, and FTSE 100, substantiate the effectiveness of GDAT.

The experimental results further indicate that (1) the dual axial attention mechanism consistently outperforms single-axial attention approaches, and (2) incorporating graph sparsity enables the self-attention mechanism to more accurately focus on the most salient edges.

TABLE V
ABLATION STUDY OF GDAT: CONTRIBUTION OF TEMPORAL AND CROSS-SECTIONAL ATTENTION

Method	ER	Std	SR
GDAT (Cross-asset only)	0.201	0.210	0.958
GDAT (Temporal only)	0.197	0.213	0.929
GDAT (Full model)	0.261	0.259	1.007

ACKNOWLEDGMENT

This work was supported by JST SPRING, Grant Number JPMJSP2114.

REFERENCES

- [1] H. Markowitz, "Portfolio selection," *The Journal of Finance*, vol. 7, no. 1, pp. 77–91, 1952.
- [2] R. Engle, "Dynamic conditional correlation: A simple class of multivariate generalized autoregressive conditional heteroskedasticity models," *Journal of business & economic statistics*, vol. 20, no. 3, pp. 339–350, 2002.
- [3] J. Fan, Y. Fan, and J. Lv, "High dimensional covariance matrix estimation using a factor model," *Journal of Econometrics*, vol. 147, no. 1, pp. 186–197, 2008.
- [4] K. Xu, Y. Zhang, D. Ye, P. Zhao, and M. Tan, "Relation-aware transformer for portfolio policy learning," in *Proceedings of the twenty-ninth international conference on international joint conferences on artificial intelligence*, 2021, pp. 4647–4653.
- [5] Z. Li and R. Fan, "Crisis-resilient portfolio management via graph-based spatio-temporal learning," *arXiv preprint arXiv:2510.20868*, 2025.
- [6] Y. Bai, Y. Gao, R. Wan, S. Zhang, and R. Song, "A review of reinforcement learning in financial applications," *Annual Review of Statistics and Its Application*, vol. 12, no. 1, pp. 209–232, 2025.
- [7] J. Moody and M. Saffell, "Reinforcement learning for trading," *Advances in neural information processing systems*, vol. 11, 1998.
- [8] S. Sood, K. Papasotiriou, M. Vaciulis, and T. Balch, "Deep reinforcement learning for optimal portfolio allocation: A comparative study with mean-variance optimization," *FinPlan*, vol. 2023, no. 2023, p. 21, 2023.
- [9] C. Zhang, Z. Zhang, M. Cucuringu, and S. Zohren, "A universal end-to-end approach to portfolio optimization via deep learning," *arXiv preprint arXiv:2111.09170*, 2021.
- [10] F. D. Freitas, A. F. De Souza, and A. R. De Almeida, "Prediction-based portfolio optimization model using neural networks," *Neurocomputing*, vol. 72, no. 10–12, pp. 2155–2170, 2009.
- [11] O. Ledoit and M. Wolf, "Honey, i shrunk the sample covariance matrix," 2003.
- [12] J. Ho, N. Kalchbrenner, D. Weissenborn, and T. Salimans, "Axial attention in multidimensional transformers," *arXiv preprint arXiv:1912.12180*, 2019.
- [13] J. M. J. Valanarasu, P. Oza, I. Hachililoglu, and V. M. Patel, "Medical transformer: Gated axial-attention for medical image segmentation," in *International conference on medical image computing and computer-assisted intervention*. Springer, 2021, pp. 36–46.
- [14] Z. Ma, Y. Qi, C. Xu, W. Zhao, M. Lou, Y. Wang, and Y. Ma, "Atfe-net: axial transformer and feature enhancement-based cnn for ultrasound breast mass segmentation," *Computers in Biology and Medicine*, vol. 153, p. 106533, 2023.
- [15] Z. Zhang, Y. Peng, X. Duan, Q. Hou, and Z. Li, "Dual-axis generalized cross attention and shape-aware network for 2d medical image segmentation," *Biomedical Signal Processing and Control*, vol. 107, p. 107791, 2025.
- [16] D. Kisiel and D. Gorse, "Axial-lob: High-frequency trading with axial attention," in *2022 IEEE Symposium Series on Computational Intelligence (SSCI)*. IEEE, 2022, pp. 1327–1333.
- [17] J. Friedman, T. Hastie, and R. Tibshirani, "Sparse inverse covariance estimation with the graphical lasso," *Biostatistics*, vol. 9, no. 3, pp. 432–441, 2008.
- [18] T. M. Cover, "Universal portfolios," *Mathematical finance*, vol. 1, no. 1, pp. 1–29, 1991.
- [19] A. Kempf and C. Memmel, "Estimating the global minimum variance portfolio," *Schmalenbach Business Review*, vol. 58, no. 4, pp. 332–348, 2006.
- [20] A. Borodin, R. El-Yaniv, and V. Gogan, "Can we learn to beat the best stock," *Advances in Neural Information Processing Systems*, vol. 16, 2003.
- [21] W. Shen and J. Wang, "Portfolio selection via subset resampling," in *Proceedings of the AAAI conference on artificial intelligence*, vol. 31, no. 1, 2017.
- [22] B. Li and S. C. Hoi, "On-line portfolio selection with moving average reversion," *arXiv preprint arXiv:1206.4626*, 2012.
- [23] D. Huang, J. Zhou, B. Li, S. C. Hoi, and S. Zhou, "Robust median reversion strategy for online portfolio selection," *IEEE Transactions on Knowledge and Data Engineering*, vol. 28, no. 9, pp. 2480–2493, 2016.
- [24] C. Lea, M. D. Flynn, R. Vidal, A. Reiter, and G. D. Hager, "Temporal convolutional networks for action segmentation and detection," in *proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 156–165.
- [25] J. Chung, C. Gulcehre, K. Cho, and Y. Bengio, "Empirical evaluation of gated recurrent neural networks on sequence modeling," *arXiv preprint arXiv:1412.3555*, 2014.
- [26] S. Hochreiter and J. Schmidhuber, "Long short-term memory," *Neural computation*, vol. 9, no. 8, pp. 1735–1780, 1997.
- [27] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, E. Kaiser, and I. Polosukhin, "Attention is all you need," *Advances in neural information processing systems*, vol. 30, 2017.