

NBER WORKING PAPER SERIES

ALPHAGLASS: INTERPRETABLE CHARACTERISTIC-BASED PORTFOLIO CHOICE

Sebastian Bell
Ali Kakhbod
Martin Lettau
Abdolreza Nazemi

Working Paper 35186
<http://www.nber.org/papers/w35186>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
May 2026

The authors have nothing to disclose. The views expressed herein are those of the authors and do not necessarily reflect the views of the National Bureau of Economic Research.

NBER working papers are circulated for discussion and comment purposes. They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2026 by Sebastian Bell, Ali Kakhbod, Martin Lettau, and Abdolreza Nazemi. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

AlphaGlass: Interpretable Characteristic-Based Portfolio Choice
Sebastian Bell, Ali Kakhbod, Martin Lettau, and Abdolreza Nazemi
NBER Working Paper No. 35186
May 2026
JEL No. C14, C45, G10, G11, G12

ABSTRACT

We propose AlphaGlass, an inherently interpretable machine-learning framework for constructing portfolios that directly optimize investment objectives. AlphaGlass maps stock characteristics into additive signals with sparse interactions and converts these signals into long-short portfolios through a differentiable rank-and-mask layer. This end-to-end design allows the model to optimize objectives such as the Sharpe ratio or mean-variance utility while keeping portfolio weights interpretable and traceable to specific characteristics and interactions. We show theoretically that in-sample objective maximization consistently estimates the population objective and that the differentiable rank-and-mask layer is a faithful smooth proxy for the corresponding conventional long-short quantile portfolio. In U.S. equities, AlphaGlass delivers strong out-of-sample performance and reveals economically interpretable drivers of long and short positions.

Sebastian Bell
School of Economics and Management,
Karlsruhe Institute of Technology
sebastian.bell@kit.edu

Ali Kakhbod
University of California, Berkeley
Haas School of Business
akakhbod@berkeley.edu

Martin Lettau
University of California, Berkeley
Haas School of Business
and CEPR
and also NBER
lettau@haas.berkeley.edu

Abdolreza Nazemi
Karlsruhe Institute of Technology
nazemi@kit.edu

1 Introduction

A fundamental question in empirical asset pricing is how an investor should map a large set of firm characteristics into portfolios that maximize a chosen investment objective. The traditional characteristic-sorting approach yields transparent portfolio rules and remains the workhorse of the field. However, it limits interactions among characteristics and does not optimize the investment objective. Modern machine-learning methods can incorporate many predictors, nonlinearities, and high-dimensional interactions, yet much of the literature focuses on return prediction with complex “black-box” models, such as neural networks, tree-based ensembles, and generative architectures.¹ This creates two challenges. First, return prediction is only an intermediate target: forecasting returns need not align with risk-adjusted portfolio objectives once signals are converted into portfolios. Second, the opacity of black-box models limits economic interpretability, making it difficult to understand which model features drive the results.

We address this central problem by developing *AlphaGlass*, an interpretable framework for characteristic-based portfolio construction for general investment objective functions. AlphaGlass maps firm characteristics and their pairwise interactions into additive signals and then converts these signals into portfolio weights. This design links portfolio weights to stock characteristics through a separable parameterization, enabling an exact decomposition of portfolio scores and allocations into characteristic- and interaction-level contributions. The key innovation is that AlphaGlass jointly estimates the signals and maximizes the investment objective in a single step, while maintaining economic transparency, thereby integrating elements that existing asset-pricing approaches typically treat separately: end-to-end investor-objective optimization, rank-based portfolio formation, and maintaining inherent interpretability.

Our contribution is threefold. First, we introduce a unified methodology for learning interpretable characteristic-based portfolio rules that directly optimize investor objectives such as the Sharpe ratio or mean-variance preferences. Unlike the standard prediction-then-sort pipeline, AlphaGlass embeds a differentiable rank-and-mask approximation of conventional quantile sorting within an additive, sparse-interaction architecture. This makes the characteristic-to-portfolio mapping trainable end-to-end while preserving the ranking logic of empirical asset pricing. At the same time, the additive structure provides exact economic attribution: each univariate shape function and selected interaction surface contributes directly to the portfolio score, allowing the model’s long and short decisions to be decomposed into characteristic- and interaction-level components. An important aspect of the model is that the mapping from signal to portfolio weights is differentiable, unlike traditional “hard” cutoffs based on characteristic quantiles. We construct a smooth rank-and-mask layer that approximates traditional sorted portfolio weights.

Second, we show that the approach delivers strong risk-adjusted performance while producing transparent economic diagnostics. In U.S. equities from 2000 to 2021, AlphaGlass generates a pronounced spread in returns and Sharpe ratios across signal-sorted portfolios, achieving a monthly

¹Other approaches include regularized linear models (e.g. Feng et al., 2020; Freyberger et al., 2020; Kozak et al., 2020) and PCA-based methods (e.g., Kelly et al., 2019; Lettau and Pelger, 2020; Kelly et al., 2023; Lettau, 2023).

out-of-sample Sharpe ratio of 0.47, outperforming benchmark characteristic-sorted portfolios and decile portfolios formed on Random Forest and neural network-based return predictions. Moreover, the performance comes with *ex ante* transparency: the glass-box structure yields interpretable shape functions and interaction surfaces that map directly into portfolio composition. In our main specification, the most important univariate shapes identify economically plausible drivers such as industry concentration (*herf*), operating leverage dynamics (*pchsale_pchxsga*), momentum (*mom12m*), and industry-adjusted labor adjustment (*chempia*), while the sparse interaction block highlights conditional effects such as the $\text{acc} \times \text{hire}$ interaction. Because AlphaGlass is separable by design, we can further decompose long and short allocations into signed, term-level contributions and identify which characteristics most strongly push assets into the long and short tails. This decomposition provides a direct link between a high-performing portfolio rule and economically interpretable drivers. We also consider several variations and extensions of the method, including alternative data splits, subsample analyses (e.g., removing small firms), constrained portfolio optimization (e.g., drawdown penalization), and other investor objectives (e.g., mean-variance utility).

Finally, we establish theoretical properties that relate AlphaGlass’s learning behavior to the underlying economic objective. We show that maximizing the sample Sharpe ratio is argmax-consistent for maximizing the population Sharpe ratio within the AlphaGlass function class. Asymptotically, the fitted model therefore attains the best-in-class out-of-sample Sharpe ratio available within the chosen interpretable architecture. We also formalize that the differentiable rank-and-mask procedure is a faithful surrogate for the conventional hard sort-and-cut portfolio construction. In particular, when signals are well separated, the “soft” weights converge to the hard quantile weights, and the corresponding Sharpe ratios coincide asymptotically under mild conditions.

We complement these theoretical results with simulation evidence from a controlled data generating process. Because the conditional mean and volatility structures are known in that environment, we can evaluate each fitted strategy on an arbitrarily large independent sample, yielding an arbitrarily accurate approximation to its population Sharpe ratio and enabling comparison with oracle portfolio rules. We find that AlphaGlass consistently outperforms benchmark learners, capturing a larger share of the highest attainable population Sharpe ratio.

Taken together, AlphaGlass provides an end-to-end approach to portfolio construction that combines the flexibility of state-of-the-art ML with the transparency required for empirical asset pricing. The framework is designed to be modular and extensible: the rank-and-mask procedure can accommodate alternative quantile definitions and constraints, and the additive structure can incorporate different regularization schemes and pruning rules to control complexity. Most importantly, the combination of end-to-end maximization of flexible objective functions with an exact decomposability of portfolio signals yields a portfolio construction method that is simultaneously high-performing, economically grounded, and transparent.

In asset pricing, the most popular ML application is predicting returns in high-dimensional panels of firm characteristics. Gu et al. (2020) and related work document that nonlinear learners, including bagging or boosting ensembles and neural networks, often outperform unregularized linear

benchmarks and can generate economically meaningful spreads when forecasts are translated into trading strategies. Classical evidence documents a large set of anomalies and the empirical difficulty of summarizing cross-sectional expected returns with small linear specifications (Fama and French, 1993; Harvey et al., 2016; Jensen et al., 2023). Modern ML work shows that regularization and nonlinear function approximation can improve predictive performance relative to traditional linear benchmarks, and that economically meaningful portfolio spreads can be generated when forecasts are translated into trading strategies (Gu et al., 2020; Bryzgalova et al., 2025). This includes penalized linear methods (Feng et al., 2020; Freyberger et al., 2020; Kozak et al., 2020), dimension-reduction, factor-based, and embedding approaches that summarize high-dimensional firm information (Kelly et al., 2019; Giglio and Xiu, 2021; Kakhbod et al., 2024; Gabaix et al., 2025), and nonlinear models such as tree ensembles and neural networks that capture interactions and state dependence (Gu et al., 2020, 2021).² Bell et al. (2025) use the Explainable Boosting Machine (EBM), a sequential boosting technique, to predict bond returns. In contrast to this return-prediction focus, we target the investor’s end objective by directly optimizing the portfolio Sharpe ratio within an inherently interpretable model class.

In doing so, we contribute to a related strand of literature that employs ML techniques to construct optimal portfolios, emphasizing that predictive models are ultimately judged by the risk-adjusted performance they enable. Common simple approaches, such as plug-in portfolios, can perform poorly out-of-sample because estimation error in expected returns and covariances is first order and investor utility is not taken into account when estimating portfolio weights (Jobson and Korkie, 1980; Michaud, 1989; Kan and Zhou, 2007; Kelly and Xiu, 2023). This motivates shrinkage and robustification, including covariance regularization (Ledoit and Wolf, 2004, 2012) and Bayesian/Black-Litterman style approaches (Jorion, 1986; Black and Litterman, 1992). More directly, an influential literature integrates estimation and optimization by parameterizing portfolio weights as functions of observable characteristics and estimating the weight function by maximizing in-sample utility (Brandt, 1999; Ait-Sahalia and Brandt, 2001; Brandt and Santa-Clara, 2006; Brandt et al., 2009). In mean-variance utility settings, this approach is also known as maximum Sharpe ratio regression (MSRR), which connects tangency-portfolio estimation to regression and facilitates regularization and model selection (see, e.g., Britten-Jones, 1999; DeMiguel et al., 2020; Ao et al., 2019). Recent work explores high-capacity MSRR architectures (e.g., neural networks) and clarifies conditions under which greater model complexity can improve out-of-sample portfolio efficiency (Simon et al., 2025; Didisheim et al., 2023). Relatedly, Cong et al. (2026) use deep reinforcement learning to learn goal-oriented portfolio policies with flexible path-dependent objectives, whereas our approach focuses on interpretable characteristic-based portfolio rules with transparent main effects, sparse interactions, and rank-based long-short portfolio formation. Finally, portfolio choice is closely linked to SDF estimation: maximizing the Sharpe ratio corresponds to identifying the most stringent Hansen-Jagannathan pricing restriction within a candidate payoff/SDF class, and regularized SDF estimation

²Recent works expand the information set beyond standard characteristics to text and other unstructured data, including applications to news, disclosures, and language-model-based signals (Ke et al., 2019; Chen et al., 2022; Fedyk et al., 2024; Kakhbod et al., 2025).

brings ML tools directly into the fundamentals of asset pricing (Hansen and Jagannathan, 1997; Kozak et al., 2020; Chen et al., 2024).

This literature motivates our end-to-end design: rather than forecasting returns and then forming portfolios, we directly optimize the investor’s objective within an inherently interpretable characteristic based framework. Relative to existing portfolio-optimization approaches, AlphaGlass combines neural-network flexibility with a separable rank-based architecture, so that portfolio signals and long-short allocations can be decomposed into explicit characteristic- and interaction-level contributions. This design yields transparent decision rules and reveals the economically interpretable drivers of long and short positions.

The rest of the paper is organized as follows. Section 2 details the estimation methodology and training algorithm. Section 3 describes the data. Section 4 presents empirical results, while Section 5 provides insights into portfolio composition. Section 6 extends the model to a mean-variance objective. Section 7 outlines the model’s theoretical properties. Section 8 presents simulation evidence comparing AlphaGlass to benchmark learners and oracle portfolio rules. Section 9 concludes.

2 Methodology

We develop a customized machine learning estimator that enables flexible portfolio construction while retaining model interpretability. AlphaGlass is a neural network-based additive model, designed to combine the interpretability of Generalized Additive Models (GAMs) with the flexibility of neural networks. Like the Explainable Boosting Machine (EBM) (Lou et al., 2013; Nori et al., 2019), the model assumes an additive decomposition of the prediction function into main effects and a sparse set of pairwise interactions, but distinctly differs in the way these effects are represented and estimated.

Let $r_{l,t}$ be the excess return of firm $l=1,\dots,L$ at time $t=1,\dots,T$. $x_{l,t}$ is a vector of N characteristics of firm l at t . The standard approach (e.g., Gu et al. (2020)) is as follows.

1. Predict returns as a function of lagged characteristics using nonlinear ML methods:

$$r_{l,t+1} = g(x_{l,t}) + e_{l,t+1}. \quad (1)$$

$g(x_{l,t})$ is estimated by minimizing the mean-square-error loss function: $\mathcal{L}_{\text{MSE}} = \frac{1}{LT} \sum_{l,t} e_{l,t}^2$.

2. Each t , form M portfolios based on the predicted returns $\hat{r}_{lt+1} = \hat{g}(x_{l,t})$. Let f_{t+1}^{HL} be the return of the long-short (M -minus-1) portfolio in $t+1$ that is formed at t .
3. Compute the Sharpe ratio (SR) of f_{t+1}^{HL} .

(1) is estimated in a training sample $\mathcal{T}_1$ and validated in a validation sample $\mathcal{T}_2$. The Sharpe ratio of f_{t+1}^{HL} is computed in the test sample $\mathcal{T}_3$. Gu et al. (2020) compare various methods, including linear models, boosted regression trees, random forests, and neural networks, using this methodology.

Note that the three steps are performed sequentially. Moreover, while the final objective is to create a factor f_t with a high Sharpe ratio, predicted returns in (1) are estimated using a different

objective function. Instead, the AlphaGlass model estimates all three steps jointly, using a single objective function, as follows.

1. *Asset-specific signals* $s_{l,t}$ are a function of lagged characteristics:

$$s_{l,t} = g(x_{i,l,t}) + e_{l,t+1}. \quad (2)$$

2. Portfolio weights $w_{l,t}^m$ of $m = 1, \dots, M$ portfolios are functions of the signals:

$$w_{l,t}^m = h^m(s_{l,t}). \quad (3)$$

3. The weights $w_{l,t}^{\text{HL}}$ and returns f_{t+1}^{HL} of the long-short portfolio are

$$w_{l,t}^{\text{HL}} = w_{l,t}^M - w_{l,t}^1, \quad (4)$$

$$f_{t+1}^{\text{HL}} = \sum_l w_{l,t}^{\text{HL}} r_{l,t+1}. \quad (5)$$

4. Portfolio weights $w_{l,t}^{\text{HL}}, t = 1, \dots, T-1$ are obtained by minimizing the loss function

$$\{\hat{w}_{l,t}^{\text{HL}}\} = \underset{\{w_{l,t}^{\text{HL}}\}}{\operatorname{argmin}} \mathcal{L}(f_2^{\text{HL}}, \dots, f_T^{\text{HL}}). \quad (6)$$

Several points are worth noting. First, (2)–(6) are solved jointly rather than sequentially yielding $\hat{g}_{i,l,t}, \hat{s}_{l,t}$ and $\hat{w}_{l,t}^{\text{HL}}$. Hence, the signals $s_{l,t}$ in (2) are estimated with respect to the loss function (6), rather than by minimizing the MSE of the errors $e_{l,t+1}$. Second, the specification of the loss function is flexible. In principle, the loss function can be any well-defined (dis)utility function. In the main part of the paper, we focus on the negative of the Sharpe ratio of f_t^{HL} : $\mathcal{L} = -\text{SR}(f_t^{\text{HL}})$, which has been the standard specification in the literature. We also consider mean-variance preferences: $\mathcal{L} = -\text{E}(f_t^{\text{HL}}) + (\gamma/2)\text{Var}(f_t^{\text{HL}})$. Our approach also allows for portfolio constraints and penalties. Third, gradient-based ML methods require that the minimization problem (2)–(6) is differentiable. However, standard portfolio construction relies on simple rankings, so that $w_{l,t}^m$ are non-differentiable step functions with “hard” cutoffs. We replace such “hard” cutoffs with smooth functions, as explained below. Next, we describe the estimation of the signals $s_{l,t}$ in (2) and portfolio weights $w_{l,t}^m$ in (3).

2.1 Estimation of signals

Following Hastie and Tibshirani (1986), Lou et al. (2013), and Bell et al. (2025), we estimate signals in (2) using a General Additive Model (GAM) with interactions³:

$$s_{l,t} = \beta_0 + \sum_{i=1}^N f_i(x_{i,l,t}) + \sum_{i>j} f_{ij}(x_{i,l,t}, x_{j,l,t}), \quad (7)$$

where

- $x_{i,l,t}$ is the characteristic i for firm l at time t ,

³Note that we use the GAM-like structure deliberately to ensure inherent interpretability, but in principle the AlphaGlass mechanism can be used with other architectures.

- f_i represents the main effect of feature x_i ,
- f_{ij} represents the interaction effect between features x_i and x_j .

The advantage of this specification is that it allows nonlinear effects while remaining inherently interpretable. For example, the influence of each variable $x_{i,l,t}$ is given by the function $f_i(x_{i,l,t})$ and possible interactions with other variables via $f_{ij}(x_{i,l,t}, x_{j,l,t})$. Within this additive structure, the AlphaGlass model represents each f_i and f_{ij} as a small feed-forward neural network, trained under smoothness and regularization constraints. This neural representation enables the model to produce smooth, differentiable effect functions and to incorporate additional constraints such as monotonicity and sparsity.

The training procedure implements three main stages:

1. **Main effect learning:** For each feature i , a neural subnetwork f_i is trained while minimizing the chosen loss function. Each subnet consists of one input neuron for x_i , a hidden layer of 20 units, and an output neuron. All subnets use Rectified Linear Unit (ReLU) activations⁴ and are trained in parallel via gradient descent. Training stops after a pre-specified number of epochs or when performance no longer improves meaningfully (as defined by an early stopping threshold).
2. **Interaction learning:** After selecting the interaction pairs, neural subnetworks f_{ij} are trained while keeping the main effect subnetworks fixed. The interaction subnetworks have two input neurons (for x_i and x_j), two hidden layers of 20 neurons each, and one output neuron.
3. **Joint fine-tuning:** Once all main effect and interaction subnetworks are trained, a global optimization step is performed in which all parameters are updated simultaneously (with a lower learning rate in order to avoid losing information from the previous training steps). Due to the structural separability of the neural subnetworks, this stage does not compromise the interpretability of individual shape functions while enabling small, coordinated adjustments that improve performance.

We provide additional information on these training stages and other details related to the estimation technique in Appendix A.

2.2 Estimation of portfolio weights

In each period t , we construct M portfolios as functions of signals $s_{l,t}$ and compute their returns in $t + 1$. Since we are interested in the long-short portfolio, defined as the difference between the portfolios $m = M$ and $m = 1$, we focus on the weights of these two portfolios. As mentioned above, a key challenge in our setting is that portfolio weights must be differentiable with respect to the signals. Rank-sorting with “hard” cutoffs, which is commonly used in portfolio construction, is non-differentiable, rendering gradient-based estimators in ML methods, such as neural networks, infeasible. To address this, we replace “hard” portfolio weights based on discrete rankings and

⁴ReLU activation is based on the max operator: $h(z) = \max(0, z)$.

portfolio inclusions with smooth “soft” weights.⁵ These “soft” portfolio weights can be interpreted as the probability that a given stock is included in a given portfolio. The estimation proceeds in two steps: first, transforming discrete ranks, and then computing portfolio weights.

First, in each month t , we replace discrete rankings with smooth pairwise comparisons, yielding differentiable approximations of the cross-sectional order of signals. Let n_t be the number of stocks in month t . The relative ordering of signals for stocks l and l' in each month t is approximated via pairwise logistic comparisons

$$S_{l,l',t} = \sigma\left(\frac{S_{l,t} - S_{l',t}}{\tau_s}\right),$$

where $\sigma(x) = \frac{1}{1+e^{-x}} \in (0,1)$ is the sigmoid function. The “temperature” hyperparameter $\tau_s > 0$ (and $\tau_m > 0$ below) controls the sharpness of the soft approximation. Smaller values make the sigmoids steeper, yielding rankings that behave more like hard, discrete assignments, whereas larger values produce smoother, more gradual transitions. The sigmoid output $S_{l,l',t}$ can be interpreted as the probability that asset l outranks asset l' based on the difference between their signals. When the signals are far apart, the sigmoid is close to 0 or 1, whereas for signals that are close together, the output lies near 0.5, reflecting uncertainty about their relative ordering. Note that $S_{l,l',t}$ converges to the (non-differentiable) indicator function as τ_s goes to zero, which is equivalent to the standard “hard” comparison:

$$S_{l,l',t} = \sigma\left(\frac{S_{l,t} - S_{l',t}}{\tau_s}\right) \xrightarrow{\tau_s \searrow 0} \mathbf{1}\{S_{l,t} \geq S_{l',t}\}.$$

The “soft” rank of asset l in month t is then defined as the sum of all pairwise comparisons of l and the other assets:

$$\text{rank}_{l,t} = 1 + \sum_{l' \neq l} S_{l,l',t} \in [1, n_t].$$

The rank can be interpreted as the expected position of asset l in the cross-sectional ordering. High signals correspond to high probabilities $S_{l,l',t}$ and therefore a high rank close to n_t , while low signals correspond to low ranks closer to 1. Hence, the soft rank behaves like a smooth analog of the usual ordinal rank while remaining differentiable. The smooth formulation ensures that the mapping of signals $s_{l,t}$ to weights $w_{l,t}$ is differentiable, allowing gradients of the Sharpe-based loss to propagate back to the signal-generating model.

In the second step, we select the highest- and lowest-ranked assets for each month to form high ($m = M$) and low ($m = 1$) portfolios. Once again, we define a “soft” approximation. Let $k_t = \lfloor n_t/M \rfloor$ denote the cutoff size in month t . Consider the high portfolio first. We shift the $\text{rank}_{l,t}$ so that stocks with ranks around the cutoff have values close to 0. For example, if $n_t = 200$ and $M = 10$, the cutoff is $k_t = 20$, hence we subtract rank_{181} from $\text{rank}_{l,t}$. For the low portfolio, we first take the negative $\text{rank}_{l,t}$ and add rank_{k_t} . Then, we apply the sigmoid transformation:

⁵In Section 7, we prove that the soft portfolio weights represent a faithful surrogate for the “hard” weights obtained by exact ranking and truncation (Theorem 1).

$$m_{l,t}^H = \sigma\left(\frac{\text{rank}_{l,t} - \text{rank}_{n_t-k_t+1,t}}{\tau_m}\right), \quad m_{l,t}^L = \sigma\left(\frac{\text{rank}_{k_t,t} - \text{rank}_{l,t}}{\tau_m}\right).$$

$m_{l,t}^H$ and $m_{l,t}^L$ can be interpreted as probabilities that a stock l is in the high and low portfolios, respectively. For example, for stocks with high ranks $m_{l,t}^H \approx 1$ and $m_{l,t}^L \approx 0$, and vice versa for stocks with low ranks. As before, the standard “hard” case is a special case when τ_s goes to zero:

$$m_{l,t}^H \xrightarrow{\tau_m \searrow 0} \mathbf{1}\{\text{rank}_{l,t} \geq \text{rank}_{n_t-k_t+1,t}\}, \quad m_{l,t}^L \xrightarrow{\tau_m \searrow 0} \mathbf{1}\{\text{rank}_{k_t,t} \leq \text{rank}_{l,t}\}.$$

To obtain portfolio weights, we normalize $m_{l,t}^H$ and $m_{l,t}^L$, so that each leg sums to one within each month:

$$w_{l,t}^H = \frac{m_{l,t}^H}{\sum_j m_{j,t}^H}, \quad w_{l,t}^L = \frac{m_{l,t}^L}{\sum_j m_{j,t}^L}.$$

Finally, the portfolio weights of the “high-minus-low” portfolio are given by

$$w_{l,t}^{\text{HL}} = w_{l,t}^H - w_{l,t}^L.$$

Figure 1 compares the “soft” portfolio weights to “hard” weights for an example of $L=200$ stocks and $M=10$ decile portfolios, so that the cutoff value $k=20$. Signals s_l are uniformly distributed between 0 and 1 and then normalized. The parameters of the sigmoid functions are $\tau_{\text{rank}}=0.1$ and $\tau_{\text{mask}}=0.5$. The orange line represents the “hard” equal-weighted weights of the long-short portfolios based on a decile sort. The blue line shows the “soft” portfolio weights obtained via sigmoid transformations. The x -axis shows stocks around the cutoffs $l=20$ and $l=181$.⁶

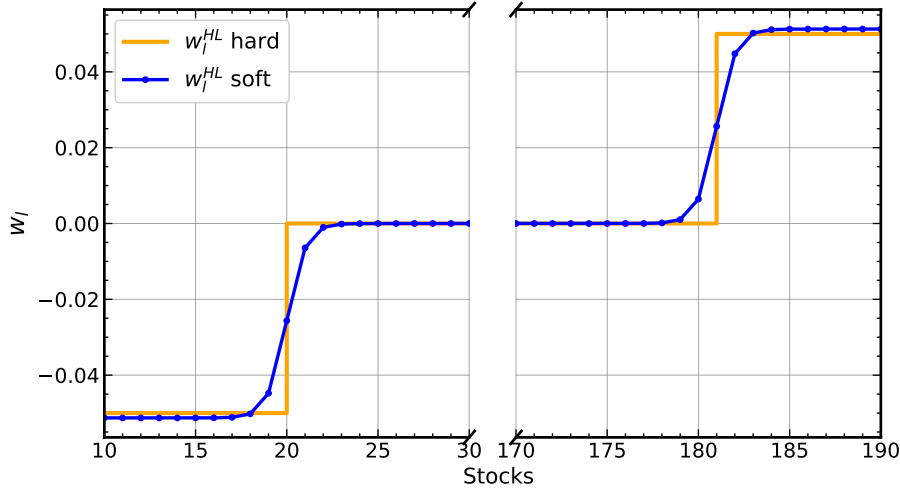
2.3 The loss function

In addition to the baseline loss function $\mathcal{L}(f_t^{\text{HL}})$, we incorporate regularizations to control model complexity and induce sparsity in the estimation of the signals in (7). Specifically, we augment the loss function with ℓ_1 -penalties on effect-specific weights as follows. Instead of estimating the f_i and f_{ij} functions in (7) directly, we introduce regularization parameters θ_i and γ_{ij} and define the GAM model as

$$s_{l,t} = \beta_0 + \sum_{i=1}^N \theta_i \tilde{f}_i(x_{i,l,t}) + \sum_{i>j} \gamma_{ij} \tilde{f}_{ij}(x_{i,l,t}, x_{j,l,t}), \quad \theta_i, \gamma_{ij} \geq 0. \quad (8)$$

⁶One could, in principle, map the AlphaGlass signal directly into weights over the full universe. We use a rank-and-mask construction because the signal is best viewed as an ordinal portfolio score rather than a calibrated expected-return forecast. The ordering of firms is economically meaningful, but the cardinal scale and spacing of scores are not uniquely pinned down. Directly using score levels as weights would therefore require additional assumptions about the signal scale and could make exposures sensitive to arbitrary transformations. The rank-and-mask layer instead converts the ordering into normalized long and short legs with a fixed gross exposure, preserving dollar neutrality and remaining well-defined as the stock universe changes. It also focuses the portfolio on the tails of the score distribution, where the model produces the clearest separation, mimicking conventional portfolio construction. For completeness, we consider an alternative weighting scheme over the full cross-section in Appendix E, resulting in lower Sharpe ratios.

Figure 1: Soft vs. hard portfolio weights w_t^{HL}



Note: This figure shows the weights of the high-low portfolio w_t^{HL} for an example with $L = 200$ stocks and $M = 10$. The blue and orange lines show the soft and hard weights, respectively. Signals s_t are evenly distributed between 0 and 1. The parameters of the sigmoid functions are $\tau_{\text{rank}} = 0.1$ and $\tau_{\text{mask}} = 0.5$.

The model estimates $(\theta_i, \tilde{f}_i)$, as well as $(\gamma_{ij}, \tilde{f}_{ij})$, separately. When we interpret the model output, however, we focus on the total effects $f_i(x_{i,l,t}) = \theta_i \tilde{f}_i(x_{i,l,t})$ and $f_{ij}(x_{i,l,t}, x_{j,l,t}) = \gamma_{ij} \tilde{f}_{ij}(x_{i,l,t}, x_{j,l,t})$.

The loss function incorporates ℓ_1 -regularization of the main and interaction effects:

$$\mathcal{L}(f_t^{\text{HL}} | \lambda_{\text{main}}, \lambda_{\text{inter}}) = \mathcal{L}(f_t^{\text{HL}}) + \lambda_{\text{main}} \sum_{i=1}^N |\theta_i| + \lambda_{\text{inter}} \sum_{i>j} |\gamma_{ij}|. \quad (9)$$

The ℓ_1 -penalties induce shrinkage towards zero during optimization. Exact zeros correspond to effects that are effectively removed from the forward pass, and the hyperparameters λ_{main} and λ_{inter} allow us to control the relative size of the regularization penalty during main effect and interaction training, respectively. Further details on the selection and pruning mechanisms are provided in Appendix A.1.

3 Data

We obtain monthly stock return data from CRSP and limit our stock universe to the three major U.S. stock exchanges, NYSE, NASDAQ, and AMEX. These data are merged with the 91 firm characteristics from Green et al. (2017) as provided by Gu et al. (2020). The characteristics are lagged relative to returns so that returns in month t are matched with monthly characteristics data at the end of month $t - 1$, quarterly data by the end of $t - 5$, and the most recent annual data by the end of $t - 7$.

We remove observations for which return or market-cap data are unavailable, yielding a total of 1,656,664 observations over the sample period from January 2000 to December 2021. Full variable definitions are available in Appendix B. Similar to Kelly et al. (2019) and Freyberger et al. (2020), all characteristics are cross-sectionally ranked and mapped to the interval $[0, 1]$. In the following main

analysis, we split our data into training, validation, and test sets $\mathcal{T}_1$, $\mathcal{T}_2$, and $\mathcal{T}_3$ of lengths 10 years, 2 years, and 10 years, respectively. We consider other splits as robustness tests.

Table 1 reports means, standard deviations, and Sharpe ratios of long-short portfolios based on univariate characteristic-sorts. Each month, we sort stocks according to a given characteristic and form value-weighted decile portfolios. The long-short portfolio is the difference between portfolios 10 and 1. Panel A shows results for the combined train and validation samples, and Panel B shows results for the test samples. We include the 10 characteristics with the highest in-sample Sharpe ratios: `pchcurrat` (%change in current ratio (current assets/liabilities)), `pchquick` (%change in quick ratio (current near-cash (quick) assets/liabilities)), `chcsho` (change in shares outstanding), `mvel1` (market value (size)), `pchsale_pchrect` (%change in sales minus %change in accounts receivable), `egr` (growth in common shareholder equity), `tb` (tax income to book income ratio), `cashpr` (cash productivity (net operating CF minus capital expenditures, divided by adjusted net income)), `salecash` (sales to cash ratio), `sp` (sales to price ratio). Note that many of the characteristics with high Sharpe ratios in the training/validation samples have substantially lower Sharpe ratios in the test sample, for example, `pchcurrat`, `pchquick`, `mvel1` (size). The only two exceptions are `chcsho` and `pchsale_pchrect`, which have out-of-sample Sharpe ratios of 0.235 and 0.196, respectively.

Table 1: Returns of characteristic-sorted portfolios

	<code>pchcurrat</code>	<code>pchquick</code>	<code>chcsho</code>	<code>mvel1</code>	<code>pchsale_pchrect</code>	<code>egr</code>	<code>tb</code>	<code>cashpr</code>	<code>salecash</code>	<code>sp</code>
Panel A: In-sample										
Mean	0.58	0.50	0.94	1.53	0.54	0.83	0.63	0.80	1.12	1.15
S.D.	1.96	1.93	3.66	7.04	2.51	3.88	3.07	3.99	5.85	6.03
SR	0.30	0.26	0.26	0.22	0.22	0.21	0.21	0.20	0.19	0.19
Panel B: Out-of-sample										
Mean	0.14	0.18	0.44	0.02	0.41	0.20	0.18	0.09	0.31	0.30
S.D.	1.79	1.83	1.89	5.43	2.11	2.46	2.10	2.75	4.03	4.14
SR	0.08	0.10	0.24	0.00	0.20	0.08	0.08	0.03	0.08	0.07

Note: This table reports monthly means, standard deviations, and Sharpe ratios of high-minus-low portfolios based on decile characteristic sorts. We sort characteristic-sorted portfolios by their Sharpe ratios in the combined training and validation samples.

4 Empirical results

Following the literature (e.g., Gu et al. (2020)), our benchmark estimation of the AlphaGlass model assumes that the loss function is the negative of the Sharpe ratio of the long/short portfolio f_t^{HL} :

$$\mathcal{L}(f_t^{\text{HL}}) = -\text{SR}(f_t^{\text{HL}}), \quad (10)$$

with ℓ_1 -regularizations as in (9). As an alternative, we consider mean-variance preferences in Section 6. We benchmark the AlphaGlass model against several methods studied in the literature: random forest (RF), neural network (NN), and explainable boosting machine (EBM). All three models follow the

standard approach outlined in Section 2, i.e., using the mean-square-error loss function to estimate predicted returns in (1). The explainable boosting machine (EBM) is an inherently interpretable ML model based on the GAM model in (7) to predict returns, but uses the MSE loss function rather than the Sharpe ratio (see Lou et al. (2013) and Bell et al. (2025) for details). In all cases, we construct decile portfolios and compute the 10-minus-1 factor f_t^{HL} . We evaluate all approaches based on the out-of-sample Sharpe ratios of the estimated long/short factors.

4.1 AlphaGlass vs. benchmark models

Table 2 reports the means, standard deviations, and Sharpe ratios of the 10 portfolios as well as the 10-1 portfolio based on the AlphaGlass estimation. In-sample mean returns and Sharpe ratios of the decile portfolios are nearly monotonic, with a monthly in-sample Sharpe ratio for the 10-1 portfolio of 0.54. The out-of-sample results largely mirror the in-sample results, indicating that the model yields stable performance. The monthly Sharpe ratio of the high/low portfolio is 0.47 (1.62 p.a.), slightly lower than the in-sample counterpart. We discuss the model details and interpretation in the following sections.

Results of the benchmark models are reported in Table 3. In-sample mean returns and Sharpe ratios are (mostly) monotonic across the decile portfolios; however, the out-of-sample results are mixed. The random forest model performs the worst. The mean return of the 10-1 portfolio f_t^{HL} is positive, but its Sharpe ratio is only 0.14. The corresponding Sharpe ratios for the neural network and EBM are 0.34 and 0.30, respectively.⁷ All three models are outperformed by AlphaGlass.⁸

Figure 2 plots the cumulative log returns of the estimated f_t^{HL} factors, as well as the CRSP-VW index, in the test sample. Returns based on AlphaGlass and the neural net are comparable until 2017, but the AlphaGlass factor dominates afterward. The returns of the EBM portfolio are similar to those of the CRSP-VW returns, whereas the cumulative return of the RF portfolio is nearly flat until 2020.

A potential concern in portfolio construction is that results are driven by small, illiquid stocks. To address this issue, we form long/short factors by size deciles. Each month, we sort stocks into size quintiles and create factors using only stocks in a given quintile. Table 4 shows the out-of-sample means, standard deviations, and Sharpe ratios for the four models by size quintile. Consider first the results of the AlphaGlass model in Panel A. The AlphaGlass factors for smaller stocks have higher mean returns than those for larger stocks; however, they are also more volatile. The Sharpe ratios are high even for larger stocks, ranging from 0.25 to 0.4, indicating that the AlphaGlass model’s results are not driven by small stocks. However, this is not the case for the other three models, as only the

⁷Gu et al. (2020) also find that neural networks yield a higher Sharpe ratio than other ML methods, such as elastic nets, RF, and gradient boosted regression trees. Note that they refit models every 12 months, so the results are not directly comparable to those in this paper.

⁸As an additional benchmark, Appendix D considers an SR-maximizing combination of characteristic-sorted long-short portfolios. The poor out-of-sample results underline importance of optimization techniques that learn generalizable portfolio rules rather than merely fit in-sample mean and covariance estimates.

Table 2: Returns of AlphaGlass portfolios

	1	2	3	4	5	6	7	8	9	10	10-1
Panel A: In-sample											
Mean	−0.30	0.37	0.33	0.54	0.76	0.79	1.01	0.95	1.04	1.57	1.87
S.D.	7.63	5.26	5.66	6.43	6.44	6.35	6.49	6.55	6.85	7.71	3.44
SR	−0.04	0.07	0.06	0.08	0.12	0.13	0.16	0.15	0.15	0.20	0.54
Panel B: Out-of-sample											
Mean	0.46	0.67	0.92	0.95	1.04	1.19	1.31	1.23	1.38	1.55	1.10
S.D.	5.79	4.40	3.64	4.40	4.94	5.18	5.40	5.50	5.90	6.43	2.35
SR	0.08	0.15	0.25	0.22	0.21	0.23	0.24	0.22	0.23	0.24	0.47

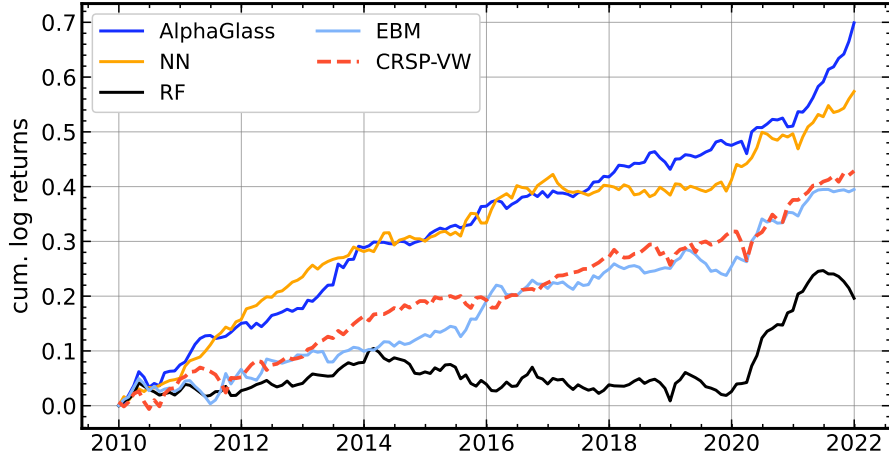
Note: This table compares monthly means, standard deviations, and Sharpe ratios for ten equal-weighted portfolios constructed based on predicted AlphaGlass signal deciles in the training/validation (Panel A) and testing sample (Panel B). The last columns show the results of the 10-1 portfolio, which is long in portfolio 10 and short in portfolio 1.

Table 3: Returns of benchmark models

	1	2	3	4	5	6	7	8	9	10	10-1
Panel A: RF In-sample											
Mean	−3.20	−1.15	−1.46	−0.65	−0.13	0.53	0.98	1.38	2.25	4.04	7.24
S.D.	9.47	6.66	7.14	6.22	5.51	5.15	5.27	5.65	6.84	10.13	9.71
SR	−0.34	−0.17	−0.21	−0.11	−0.02	0.10	0.19	0.24	0.33	0.40	0.75
Panel B: RF Out-of-sample											
Mean	0.94	0.89	0.81	0.94	0.96	1.04	1.22	1.33	1.00	1.56	0.62
S.D.	5.32	4.89	4.64	4.06	4.11	4.45	5.33	5.68	6.11	7.96	4.33
SR	0.18	0.18	0.18	0.23	0.23	0.23	0.23	0.24	0.16	0.20	0.14
Panel C: NN In-sample											
Mean	−1.54	−0.18	0.90	1.00	1.05	1.03	1.05	0.88	0.76	2.11	3.66
S.D.	10.89	8.75	4.57	3.97	4.43	5.03	5.50	6.09	6.71	11.12	5.45
SR	−0.14	−0.02	0.20	0.25	0.24	0.21	0.19	0.15	0.11	0.19	0.67
Panel D: NN Out-of-sample											
Mean	0.60	0.72	1.02	1.03	1.05	1.24	1.24	1.24	1.09	1.50	0.90
S.D.	7.31	6.40	3.19	3.57	4.19	4.66	4.97	5.30	5.49	7.04	2.67
SR	0.08	0.11	0.32	0.29	0.25	0.27	0.25	0.23	0.20	0.21	0.34
Panel E: EBM In-sample											
Mean	−3.41	−1.56	−0.58	−0.15	0.37	0.94	1.44	2.07	2.89	5.07	8.48
S.D.	8.53	6.88	6.15	6.05	5.65	5.66	5.65	6.24	7.33	11.81	10.91
SR	−0.40	−0.23	−0.10	−0.03	0.07	0.17	0.25	0.33	0.40	0.43	0.78
Panel F: EBM Out-of-sample											
Mean	0.59	0.85	0.87	0.91	0.94	1.11	1.32	1.00	1.17	1.94	1.35
S.D.	5.79	4.51	4.11	4.11	4.45	4.72	5.44	5.37	6.15	7.86	4.52
SR	0.10	0.19	0.21	0.22	0.21	0.24	0.24	0.19	0.19	0.25	0.30

Note: This table compares monthly means, standard deviations, and Sharpe ratios for ten equal-weighted portfolios constructed based on Random Forest (RF), neural network (NN), and explainable boosting machine (EBM) prediction deciles in the training/validation and testing sample. The last columns show the results of the 10-1 portfolio, which is long in portfolio 10 and short in portfolio 1.

Figure 2: Cumulative returns of f_t^{HL} of benchmark models



Note: This figure shows cumulative log returns of the f_t^{HL} factors based on the AlphaGlass (blue), neural net (NN, orange), random forest (RF, black), and the explainable boosting machine (EBM, light blue) model, as well as the CRSP-VW index (red dashed) in the test sample.

factors with the smallest stocks in the first size quintile have Sharpe ratios above 0.2. Factors with larger stocks have means and Sharpe ratios that are close to zero. Hence, AlphaGlass is the only model that generates factors with high means and Sharpe ratios for stocks across the size spectrum.

The results in Table 4 are based on the AlphaGlass estimation for the entire sample. As an alternative exercise, we exclude small stocks in the bottom size quintile in each month and estimate AlphaGlass using the remaining sample. The full results are in Table F.1. The Sharpe ratio of AlphaGlass is 0.43, slightly lower than the SR of 0.54 for the full sample that includes small stocks. The performance of the other models is more sensitive to the exclusion of small stocks. The out-of-sample Sharpe ratios of the neural net and EBM models drop from 0.34 to 0.12, and from 0.30 to 0.13, respectively.

Our baseline design uses a 10-year training window, a 2-year validation window, and a 10-year testing window. For robustness, we use a longer estimation period, with 15 years of training data, while keeping the validation window fixed at two years and, correspondingly, shortening the test period; see Table G.1 in Appendix G. The out-of-sample Sharpe ratio of the long/short portfolio is 0.49, slightly higher than in the benchmark case, suggesting that AlphaGlass is reasonably stable across alternative train/test splits.

4.2 Predictor importance

The additive structure of the AlphaGlass signal estimation in (7) allows us to investigate the role of individual characteristics in the estimation. First, we construct a measure of the overall “importance” of predictors. The next section focuses on how predictors translate into signals through the estimated “shape” functions $f_i(x_i)$ and $f_{ij}(x_i, x_j)$.

Table 4: Out-of-sample factors by size quintiles

	Size quintile				
	Small	2	3	4	Big
Panel A: AlphaGlass					
Mean	1.34	1.32	1.15	0.81	0.59
S.D.	4.26	3.27	3.54	3.21	2.20
SR	0.31	0.40	0.33	0.25	0.27
Panel B: Random forest					
Mean	1.54	−0.06	0.20	0.33	0.06
S.D.	5.57	4.43	4.18	3.39	4.02
SR	0.28	−0.01	0.05	0.10	0.01
Panel C: Neural Net					
Mean	1.48	0.69	0.44	−0.03	0.06
S.D.	4.71	3.75	3.83	3.83	3.16
SR	0.32	0.19	0.11	−0.01	0.02
Panel D: EBM					
Mean	2.28	0.48	0.86	−0.10	0.18
S.D.	4.93	4.47	4.93	4.19	4.63
SR	0.46	0.11	0.18	−0.02	0.04

Note: This table compares monthly out-of-sample means, standard deviations, and Sharpe ratios for long/short factors based on AlphaGlass, Random Forest (RF), neural network (NN), and explainable boosting machine (EBM) predictions. The factors include only stocks in a size quintile, with quintiles formed each month.

We define the overall importance of a predictor by the *mean absolute score* $S(i)$ as absolute values of $f_i(x_{i,l,t})$ averaged over the training sample:

$$S(i) = \frac{1}{|\mathcal{T}_1|} \sum_{(l,t) \in \mathcal{T}_1} |f_i(x_{i,l,t})|. \quad (11)$$

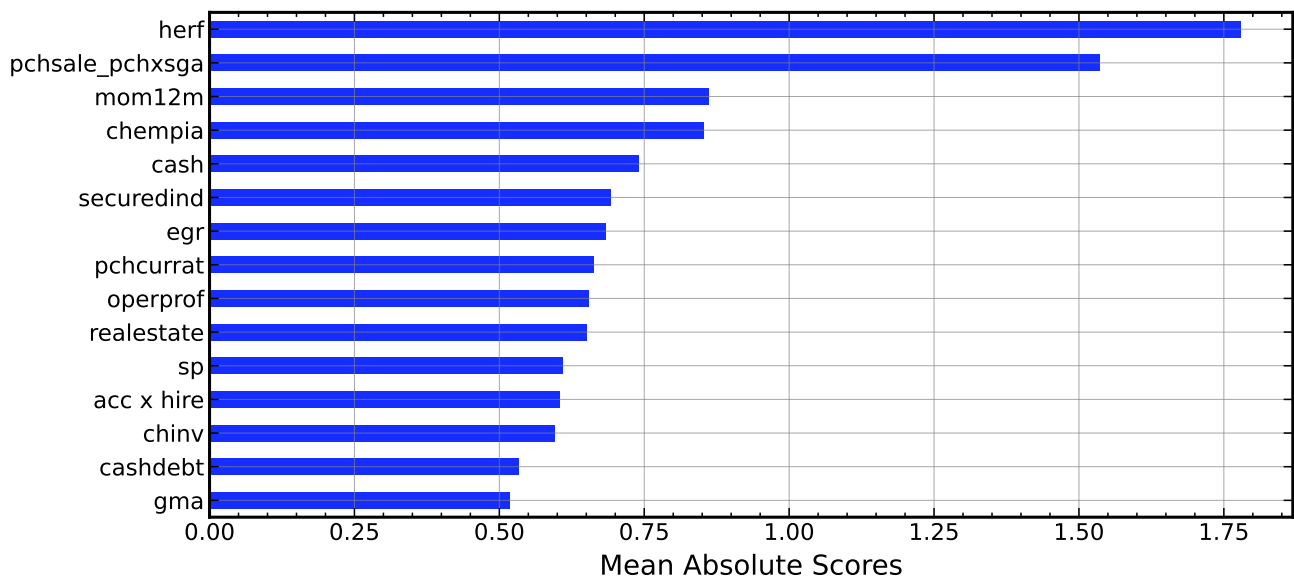
$S(i)$ captures the average effect of variable i , taking the distribution over the sample into account, and is therefore similar to the absolute value of a (standardized) coefficient in a linear regression. Therefore, the mean absolute scores capture the average “size” of the effect of an individual variable on the model’s prediction. The mean absolute scores of the interaction terms are defined analogously:

$$S(i,j) = \frac{1}{|\mathcal{T}_1|} \sum_{(l,t) \in \mathcal{T}_1} |f_{ij}(x_{i,l,t}, x_{j,l,t})|. \quad (12)$$

We consider both individual effects and effects by characteristic categories. Figure 3 presents the 15 most important effects according to the mean absolute score measure. We find that, overall, univariate effects are the most important contributors, with the exception of the interaction between `acc` and `hire`. The ranking is led by `herf` (index of industry sales concentration) and `pchsale_pchxsga` (% change in sales - % change in SG&A), with mean absolute scores of 1.78 and 1.54, respectively. This is followed by a noticeable drop to the next two variables, `mom12m` (12-month momentum) and `chempia` (industry-adjusted change in employees), both with scores around 0.86. Beyond these top

four effects, the importance values become relatively compressed: cash (cash holdings), securedind (secured debt indicator), egr (growth in common shareholder equity), pchcurrat (% change in current ratio), operprof (operating profitability), and realestate (real estate holdings) all lie in a fairly narrow range between roughly 0.65 and 0.75, followed by sp (sales to price ratio), the $\text{acc} \times \text{hire}$ (accrual-employee growth rate) interaction, chinv, cashdebt, and gma. This pattern suggests that the model places particularly strong weight on a small set of leading univariate effects, while still combining them with a broader range of medium-sized signals rather than relying on a single characteristic alone. The fact that the $\text{acc} \times \text{hire}$ term ranks near the middle of the top 15 indicates that interaction effects are economically relevant, but remain selective relative to the dominant main effects.

Figure 3: Importance of predictors



Note: This figure shows the most important characteristics for the AlphaGlass model based on the mean absolute scores $S(i)$ and $S(i,j)$. Importance values are calculated as mean absolute scores on the training data.

In nonlinear models, it is usually not possible to determine with certainty which predictors are relevant and which are not. Conceptually, variable selection should be related to mean returns and Sharpe ratios associated with characteristics, as well as the dependence structure of the data. Table 5 reports properties of high/low characteristic portfolios that are informative about these dimensions. The first column shows the 12 characteristics with the highest Sharpe ratios. Four of the 15 most important predictors have among the highest Sharpe ratios across all 91 characteristics: pchcurrat, egr, sp, and herf. Next, we estimate the first four principal components of the matrix consisting of the long/short characteristic portfolios and compute the sum of the (absolute) loadings by characteristic. Smaller loadings indicate that a characteristic is less relevant to the common movements captured by PCA factors. The second column of Table 5 shows the predictors with the smallest loadings. Five of the most important characteristics have small PCA loadings: pchcurrat, chtoia, chempia, pchsale_pchxsga, and herf. Finally, we compute the tangency mean-variance portfolio generated from the 91 high/low characteristic portfolios. The third column shows the characteristics

with the highest weights. Again, five of the most important predictors have high tangency weights: herf, egr, cash, operprof, and sp. Although there is no clear interpretation of the predictors' importance, these results suggest that many of the top predictors are associated with properties relevant to portfolio optimization.

Table 5: Properties of characteristic portfolios

	Sharpe ratios		PCA weights		MV weights	
1	pchcurrat	0.30	pchcurrat	0.04	pchsale_pchrect	3.80%
2	pchquick	0.26	pricedelay	0.05	tb	3.70%
3	chcsho	0.26	pchquick	0.05	herf	3.35%
4	mvel1	0.22	chatoia	0.05	egr	3.06%
5	pchsale_pchrect	0.22	cinvest	0.08	nincr	2.94%
6	egr	0.21	pchsale_pchrect	0.08	ear	2.62%
7	tb	0.21	aeavol	0.10	roaq	2.62%
8	cashpr	0.20	chempia	0.10	cash	2.49%
9	salecash	0.19	pchsale_pchinv	0.11	indmom	2.31%
10	sp	0.19	pchsale_pchxsga	0.11	operprof	1.97%
11	agr	0.18	herf	0.11	sp	1.96%
12	herf	0.18	pchcapx_ia	0.12	mvel1	1.93%

Note: This table reports the 12 characteristics with the highest Sharpe ratios of long/short characteristic portfolios, the lowest sum of absolute loadings of the first four principal components of long/short characteristic portfolios, and the highest mean-variance weights of long/short characteristic portfolios. Characteristics in the top-10 by mean score are bolded.

So far, we have described the importance of individual characteristics. Next, we group the characteristics into categories and assess the total importance of all characteristics within each category. Table C.2 lists the categories and included characteristics. Figure 4 shows the 10 categories with the highest mean scores. The most important characteristics are those related to profitability and its changes, followed by those related to investment, external financing, and the “classic” value/growth and momentum/reversal characteristics.

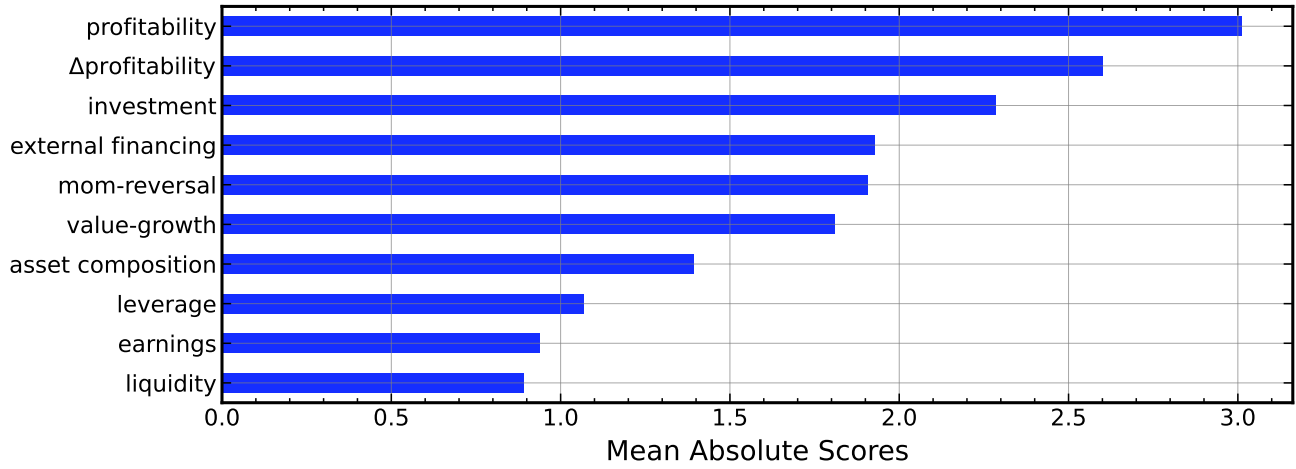
4.3 Alternative importance measure

The mean absolute score $S(i)$ captures the average effect of predictor i on signal estimation using the training sample. An alternative measure of a variable's importance is its direct effect on portfolio Sharpe ratios in the test sample. Given the estimated model, we compute the change in the out-of-sample Sharpe ratio ΔSR when an individual term f_i or f_{ij} is removed from the model. We construct decile portfolios on signals $s_{i,t}^{-i}$ when the univariate effect of variable i is eliminated from (7), and compute the corresponding Sharpe ratio of the 10-1 long-short portfolio SR_{-i} . The joint effect of variables i and j is computed analogously. The effect of variable i and interaction (i,j) on the explanatory power is the reduction relative to the SR of the full model:

$$\Delta SR_i = SR - SR_{-i}, \quad (13)$$

$$\Delta SR_{ij} = SR - SR_{-ij}. \quad (14)$$

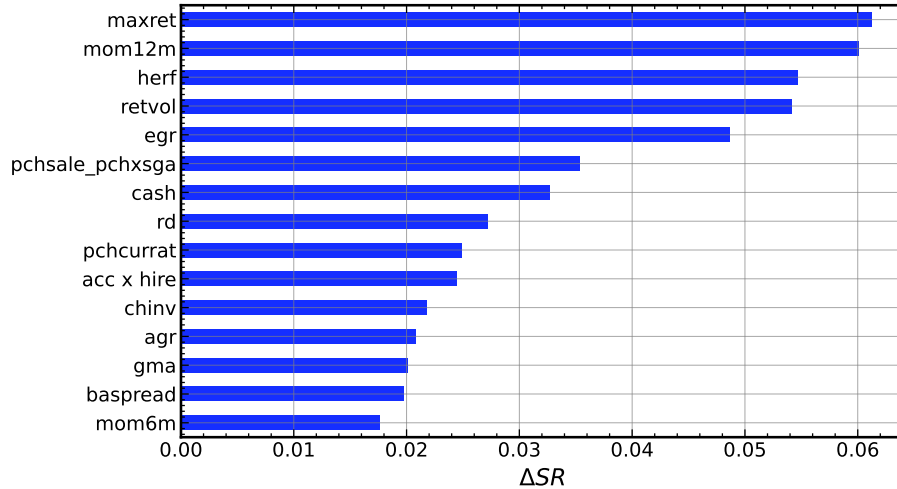
Figure 4: Importance by categories



Note: This figure shows the most important categories for the AlphaGlass model based on the mean absolute scores $S(i)$ and $S(i,j)$. Importance values are calculated as mean absolute scores on the training data and aggregated by category.

Figure 5 shows the 15 most important predictors according to the Sharpe ratio criterion. When `maxret` or `mom12m` are removed, the out-of-sample Sharpe ratio drops from 0.47 to 0.41. The effects of removing `herf`, `retvol`, or `egr` are only slightly smaller. 10 characteristics are among the top 15 according to both criteria, indicating substantial overlap, even though the mean absolute score is based on the training sample, whereas the Sharpe ratio criterion is computed on the test sample.

Figure 5: Importance of predictors: ΔSR



Note: This figure shows the most important model terms for the AlphaGlass model according to the Sharpe ratio criterion (13) and (14).

4.4 Shape functions

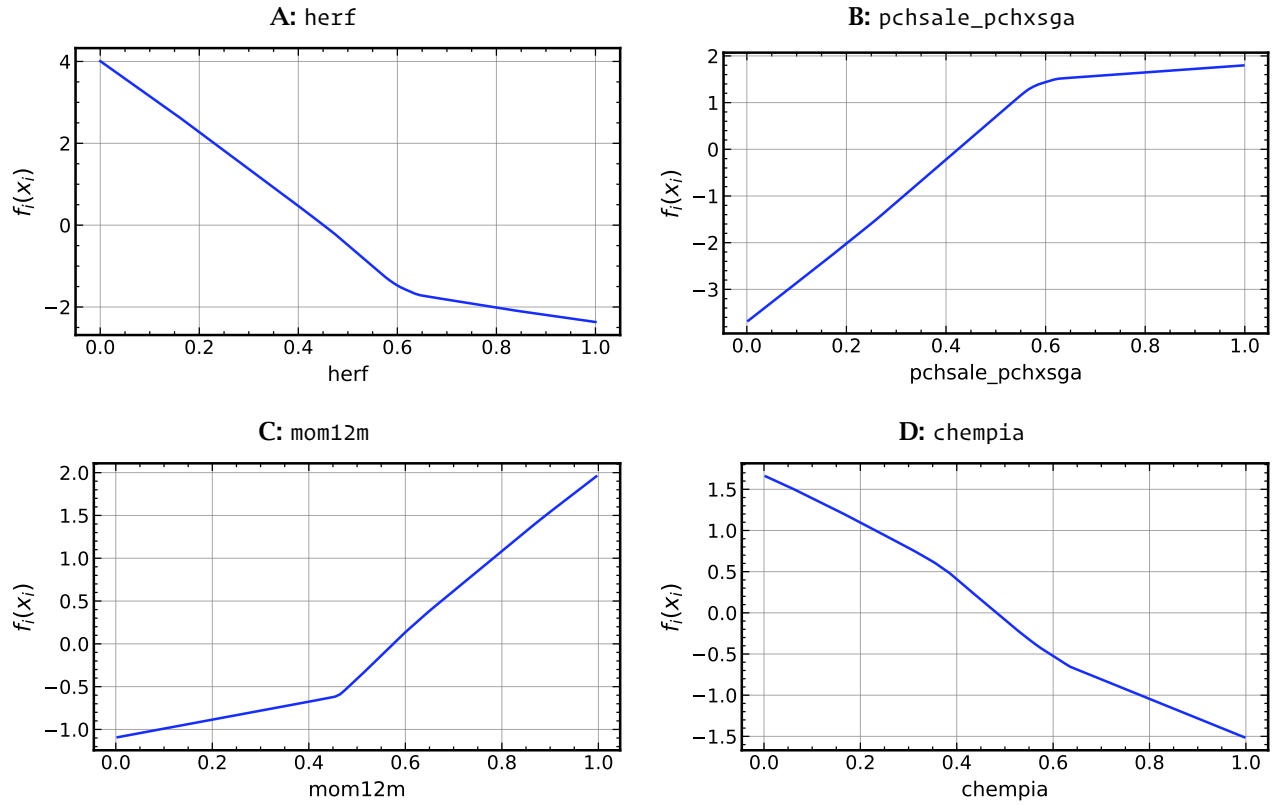
Due to the separable structure of the model for estimating signals in (7), we can analyze the (potentially) nonlinear effects of individual input characteristics or interaction terms directly from the estimated shape functions f_i and f_{ij} . In contrast to black-box models, these functions capture the influence of each model term *without* the approximation error inherent in post-hoc methods such as SHAP.⁹

The four most important univariate shape functions of our main model (in terms of mean absolute scores) are presented in Figure 6, with additional shape functions provided in Appendix H. Note that the output of each function represents the contribution of an input variable to an asset-specific signal $s_{i,t}$. As outlined in Section 2, this signal is transformed into portfolio weights through a soft ranking and decile cut procedure, so that high (low) signals are associated with high (low) weights, i.e., membership in the long (short) leg of the final portfolio. The mean and standard deviation of the estimated signals are 0.74 and 6.06, respectively, with an interquartile range of (-4.38, 4.98).

We first examine the effect associated with the most important univariate effect (Panel A), the industry sales concentration measure *herf*, as defined in Hou and Robinson (2006). Intuitively, higher values of *herf* indicate that industry sales are concentrated in the hands of a few firms (i.e., a more concentrated and less competitive industry structure), whereas lower values indicate a more competitive industry with sales distributed more evenly across firms. The estimated shape function is monotonically decreasing over the full range of *herf* values. It starts at a strongly positive value of about 4 when *herf* = 0 and then declines steeply, approximately linearly, as *herf* increases. Around *herf* $\approx$ 0.6, the function exhibits a kink: beyond this point, the function continues to decrease, but the slope becomes noticeably less steep. At the upper end of the support, the function falls to below -2 at *herf* = 1, implying that very high concentration is associated with a substantially negative contribution to the model’s output score. Economically, this shape indicates that industry structure is an important driver of the model’s score, which is subsequently mapped into portfolio weights. Low concentration (low *herf*) contributes positively to the score and therefore increases the portfolio weight assigned to such stocks, whereas higher concentration (high *herf*) contributes negatively and therefore reduces the portfolio weight. In this sense, the Sharpe ratio-optimal portfolio implied by the model weights toward stocks operating in more competitive industries and away from stocks in highly concentrated industries, with the strongest marginal shift occurring as concentration rises from very low to moderate levels and a smaller marginal adjustment once concentration is already high.

⁹Post-hoc explainability methods approximate the decision process of a fitted black-box model and can therefore introduce approximation error or instability (Rudin, 2019). In high-dimensional stock-level applications, methods such as SHAP can also become computationally costly and often rely on sampling or approximations, which may introduce estimation noise and make explanations sensitive to implementation choices (Slack et al., 2020). These distinctions are especially important in our setting because AlphaGlass provides exact decomposition by construction rather than an ex-post explanation of an opaque model.

Figure 6: Univariate functions $f_i(x_i)$



Note: This figure shows the estimated $f_i(x_i)$ functions of the most important characteristics. The x -axis plots the input variable x_i , cross-sectionally ranked and normalized to the $[0,1]$ interval. The y -axis represents the additive contribution to the prediction output through the shape function $f_i(x_i)$.

The second most important effect is associated with `pchsale_pchxsga`, which denotes the accounting signal from Abarbanell and Bushee (1998), defined as the percentage change in sales minus the percentage change in selling, general, and administrative (SG&A) expenses. The variable captures whether a firm's revenue growth is accompanied by relatively contained overhead costs. Higher values indicate that sales are growing faster than SG&A, consistent with improving operating leverage and cost discipline, whereas lower values indicate that SG&A is rising at least as quickly as sales, suggesting weaker cost control or deterioration in operating efficiency.

The corresponding AlphaGlass shape function for `pchsale_pchxsga`, displayed in Panel B, is monotonically increasing, starting at a strongly negative value slightly above -4 at an input value of 0 and rises steeply as the characteristic increases, reaching close to $+2$ at an input value of 1. The increase is particularly pronounced at low to moderate values, after which the function exhibits a kink around the midrange (between approximately 0.5 and 0.6). Beyond this point, the function continues to increase, but at a noticeably lower slope, implying diminishing marginal contributions at higher values of the characteristic. This shape implies that the model assigns higher scores to firms whose sales growth outpaces SG&A growth, and lower scores to firms whose sales growth is not supported by comparatively restrained overhead expenses. Since the model score is subsequently transformed into portfolio weights, the function indicates that the Sharpe ratio-optimal portfolio places greater weight on stocks with high `pchsale_pchxsga` and reduces exposure to stocks with low values. The steep rise at the lower end suggests that moving from weak to moderate relative cost efficiency produces the largest increase in the model-implied score, while improvements at already high levels of `pchsale_pchxsga` still raise the score but with smaller incremental effects.

The shape function in Panel C is associated with 12-month stock return momentum as in Jegadeesh and Titman (1993), measured as a stock's past return over the previous year (excluding the most recent month) and commonly interpreted as capturing the tendency of recent winners to continue outperforming recent losers in the near term. Higher values of the momentum input, therefore, indicate stronger positive past performance over the lookback window, whereas lower values indicate weaker or negative past performance. In line with this observation, the estimated shape function is monotonically increasing across the input range. At an input value of 0, the function output is moderately negative at around -1 . It then rises only gradually, reaching roughly -0.5 at a momentum of about 0.5, implying that variations in momentum within the lower half of the distribution have only a limited influence on the model score. Beyond this midrange level, the slope increases markedly, and the function rises more rapidly, reaching a strongly positive output of approximately $+2$ at the maximum input value. This pattern indicates a convex response, with substantially larger score contributions for very high-momentum realizations. Economically, the shape implies that the model rewards exposure to stocks with strong 12-month momentum, whereas assigning lower scores to stocks with weak momentum. This indicates that the Sharpe ratio-optimal portfolio tilts toward high-momentum stocks and away from low-momentum stocks. The relatively flat profile over the lower half of the input range suggests that moderately low momentum is penalized, but not strongly differentiated, whereas exceptionally strong momentum is treated as particularly beneficial and receives a disproportionately large increase in score and hence portfolio weight.

Finally, we examine the effect of `chempia`, the industry-adjusted change in employees following Asness et al. (2000) (Panel D). It captures a firm’s growth in headcount relative to the typical change in headcount among its industry peers over the same period. Higher values, therefore, indicate comparatively stronger employment growth than the industry benchmark, whereas lower values indicate comparatively weaker employment growth (or employment contraction) relative to peers.

The associated shape function is monotonically decreasing and approximately symmetric, with no pronounced kinks or sharp changes in slope. It declines in an almost linear fashion across the input range, taking a positive value of about +1.5 at `chempia`=0 and falling steadily to about −1.5 at `chempia`=1. Therefore, the contribution to the model score moves smoothly from positive to negative as the industry-adjusted change in employees increases. Economically, this pattern implies that lower industry-adjusted employee growth is associated with higher model scores, while relatively high employee growth is associated with lower scores. Since the model score is subsequently converted into portfolio weights, the shape indicates that, all else equal, the Sharpe ratio-optimal portfolio increases weight on firms with low (or negative) industry-adjusted changes in employees and reduces weight on firms exhibiting unusually strong headcount expansion relative to their industries. The near-linear, smooth decline suggests that this preference is broadly proportional: incremental increases in relative employee growth are associated with a correspondingly lower contribution to the score and therefore a lower portfolio tilt.

We conclude that the AlphaGlass model identifies several economically plausible effects that are exploited to maximize the portfolio’s Sharpe ratio.

4.5 Interactions

In addition to the univariate effects presented so far, the AlphaGlass optimizer models interactions between pairs of input characteristics. We show the effect associated with the most important such term in Figure 7. Specifically, the shape function $f_{ij}(x_i, x_j)$ for the interaction between working capital accruals (`acc`) and the employee growth rate (`hire`), which we call “acc&hire”, is presented in Panel A of Figure 7 in the form of a heatmap. Working capital accruals summarize changes in non-cash working capital components and can be interpreted as a measure of the extent to which current earnings are supported by accruals rather than cash flows. The employee growth rate captures the firm’s hiring intensity and thus its expansion in labor inputs.

The estimated interaction surface indicates that the contribution of hiring to the model score depends strongly on accruals levels. When `acc` is low, the interaction output is negative for low `hire`, reaching values below −1, but it increases considerably as `hire` rises, ultimately exceeding +2.5 at high hiring rates. Hence, conditional on low accruals, higher employee growth is associated with a substantially more positive contribution to the score, whereas low hiring is penalized. As `acc` increases, however, this relationship diminishes: the gradient with respect to `hire` becomes progressively weaker, and for large `acc` the interaction output is relatively small (between 0 and 1) across the hiring spectrum. At high `acc`, there is also a mild reversal in the pattern, with low `hire` yielding slightly higher outputs than high `hire`.

Another way to illustrate interaction effects is to plot x_i conditional on different values of x_j , i.e., $f_{ij}(x_i|x_j)$. Panel B shows the effect of `acc` on the x -axis for five different values of `hire`, equivalent to taking horizontal cuts of the heatmap in Panel A. Panel C presents the opposite case of holding `acc` constant at different levels and observing the resulting contributions as a function of `hire` (equivalent to taking vertical cuts in the heatmap).

In economic terms, this interaction suggests that the portfolio-relevant information in hiring is concentrated among firms with low working capital accruals. In that region, high hiring is associated with a large positive contribution to the score and therefore a higher model-implied portfolio weight, while low hiring reduces the score and thus lowers the weight. When accruals are high, by contrast, hiring conveys little incremental benefit for the score: the interaction contribution becomes small and fairly uniform, and the slight reversal implies that high hiring is no longer rewarded and may even be modestly disfavored relative to low hiring. Overall, the model therefore puts weight on high-hiring firms only when such hiring occurs alongside low accruals, while it largely neutralizes (or slightly reverses) the hiring tilt when accruals are elevated.

5 Economic insights from interpretability and portfolio properties

So far, we have demonstrated AlphaGlass' capabilities in terms of performance and structural interpretability. Next, we detail how the model provides insights into portfolio composition through transparent trading rules derived from its shape functions. In particular, we focus on the composition of the long and short legs of the 10-1 portfolio, which is constructed from the model's output signals.

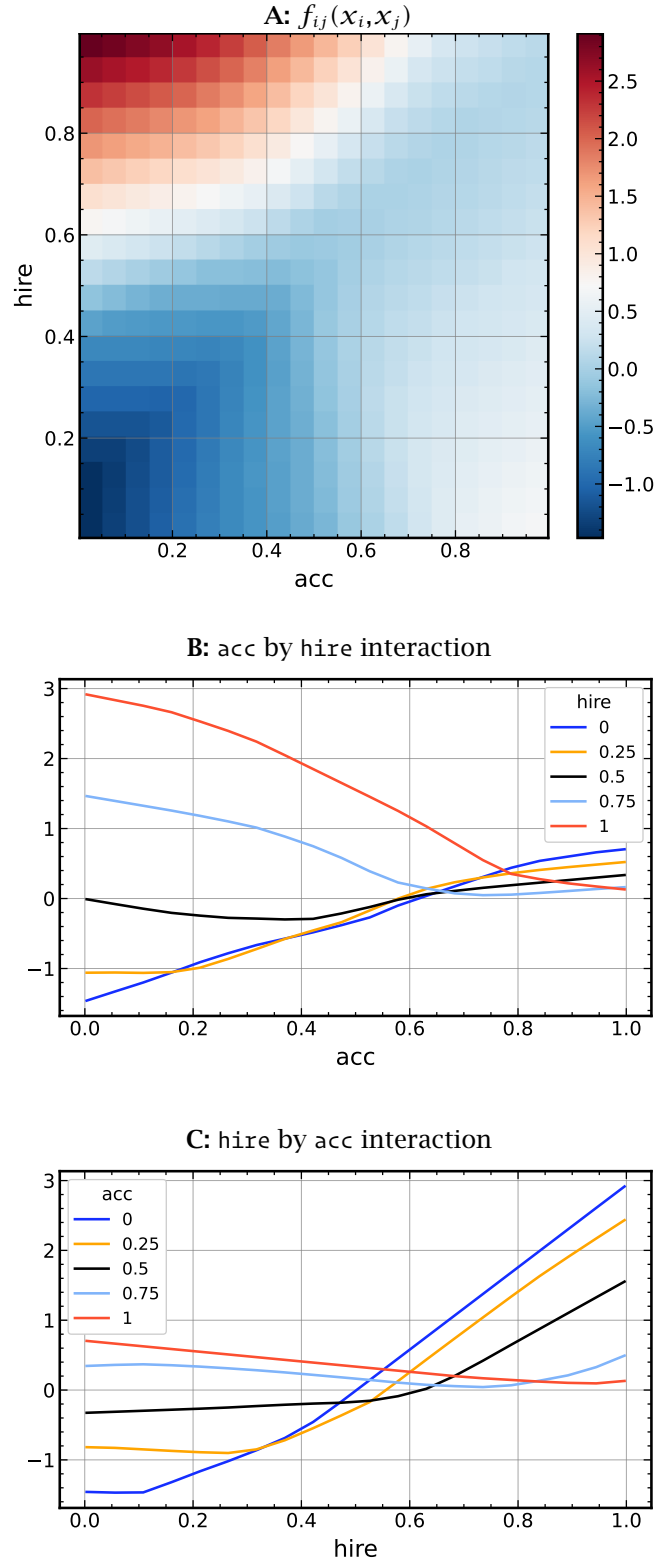
In contrast to black-box approaches, which allow only ex-post analysis of selected assets, our interpretable model provides an ex-ante understanding of how portfolios are constructed. In other words, the AlphaGlass model produces transparent rules describing which values of characteristics lead to higher output scores and, therefore, feature in the portfolios.

We illustrate this using the most important effects presented in Section 4.4. For the most important characteristic, `herf`, the shape function decreases monotonically, indicating that output scores are increased for low-`herf` stocks and decreased for high-`herf` stocks. In terms of portfolio composition, this suggests that low-`herf` stocks are expected to appear in the long leg (decile portfolio 10) constructed from the AlphaGlass scores, whereas high-`herf` stocks are expected to populate the short leg (decile portfolio 1). Panel A of Figure 8 shows that this prediction is indeed reflected in the composition, as monthly mean `herf` is considerably higher in portfolio 1 compared to portfolio 10.¹⁰ This pattern holds consistently throughout the sample period.

Similar conclusions can be drawn from the `pchsale_pchxsga` and `mom12m` shape functions: Both increase monotonically and are, therefore, consistent with stocks with high (low) characteristic values appearing in the long (short) leg. Panels B and C of Figure 8 confirm this expectation, showing

¹⁰Means within portfolios are scaled by the monthly mean across all stocks to eliminate changes in the characteristic levels.

Figure 7: Interaction - acc&hire

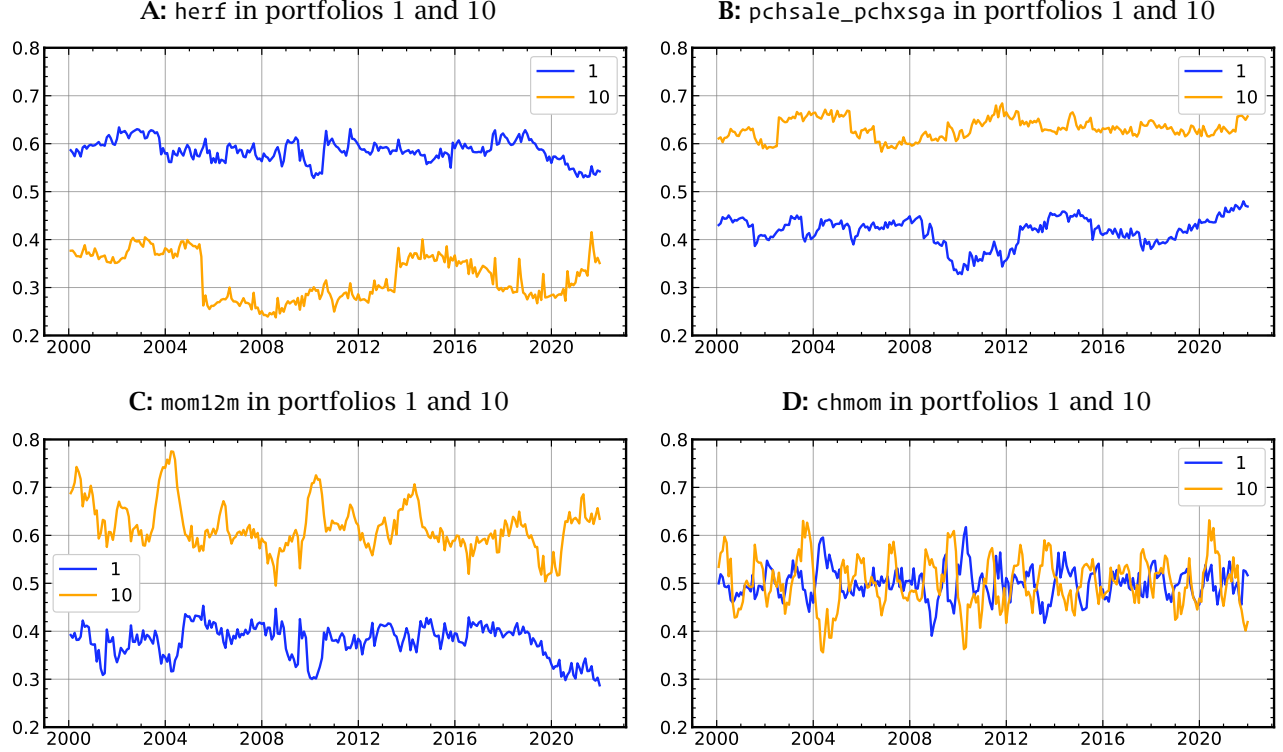


Note: Panel A displays the shape function of the interaction between acc and hire. The scaled values of acc and hire are shown on the x - and the y -axis, respectively. The color-coded axis represents the additive contribution to the signal output associated with combinations of x and y values in basis points. In Panels B and C, conditional contributions are derived from the shape function by separately holding each variable involved in the interaction constant at different levels.

clear separation between the two legs over the sample period.

In contrast, Panel D presents the same analysis for the least important univariate effect, namely the change in 6-month momentum (chmom). We observe little separation between the two legs, confirming the expectation that chmom does not meaningfully influence portfolio composition.

Figure 8: Characteristics of AlphaGlass portfolios



Note: This figure shows the characteristics of the top and bottom deciles of the AlphaGlass portfolios. Panels A to C show the monthly means of the three characteristics with the highest $S(i)$. Panel D shows the characteristic with the lowest $S(i)$.

5.1 Signal decomposition

Next, we examine which characteristics explain asset selection within the AlphaGlass portfolio. To this end, we measure which model terms increase the signals, pushing assets toward the long leg, and which decrease the signals, most likely leading to inclusion in the short leg. Let $\mathcal{P}_q \subset \mathcal{T}_1$ be the set of training-set indices (l, t) that fall into decile $q \in \{1, \dots, 10\}$ at their respective formation month. To isolate signed contributions, we decompose each shape function output $a = f_i(x_{i,l,t})$ (for some feature i), or $a = f_{ij}(x_{i,l,t}, x_{j,l,t})$ (for some pair (i, j)) into its positive and negative parts:

$$[a]_+ = \max\{a, 0\}, \quad [a]_- = \max\{-a, 0\}.$$

We define *directional mean scores* on any decile q as the portfolio-size-normalized averages of the positive and negative parts of the shape functions:

$$S_i^+(q) = \frac{1}{|\mathcal{P}_q|} \sum_{(l,t) \in \mathcal{P}_q} [f_i(x_{i,l,t})]_+ \quad (15)$$

$$S_i^-(q) = -\frac{1}{|\mathcal{P}_q|} \sum_{(l,t) \in \mathcal{P}_q} [f_i(x_{i,l,t})]_- \quad (16)$$

$$S_{ij}^+(q) = \frac{1}{|\mathcal{P}_q|} \sum_{(l,t) \in \mathcal{P}_q} [f_{ij}(x_{i,l,t}, x_{j,l,t})]_+ \quad (17)$$

$$S_{ij}^-(q) = -\frac{1}{|\mathcal{P}_q|} \sum_{(l,t) \in \mathcal{P}_q} [f_{ij}(x_{i,l,t}, x_{j,l,t})]_- \quad (18)$$

By construction, S^+ summarizes how strongly (on average) a feature or interaction contributes positively within the chosen decile, whereas S^- summarizes the average negative contribution and is reported with a minus sign so its magnitude is directly comparable to $|S^+|$. For our long-short analysis, we set

$$S_i^+ = S_i^+(10), \quad S_{ij}^+ = S_{ij}^+(10) \quad \text{and} \quad S_i^- = S_i^-(1), \quad S_{ij}^- = S_{ij}^-(1).$$

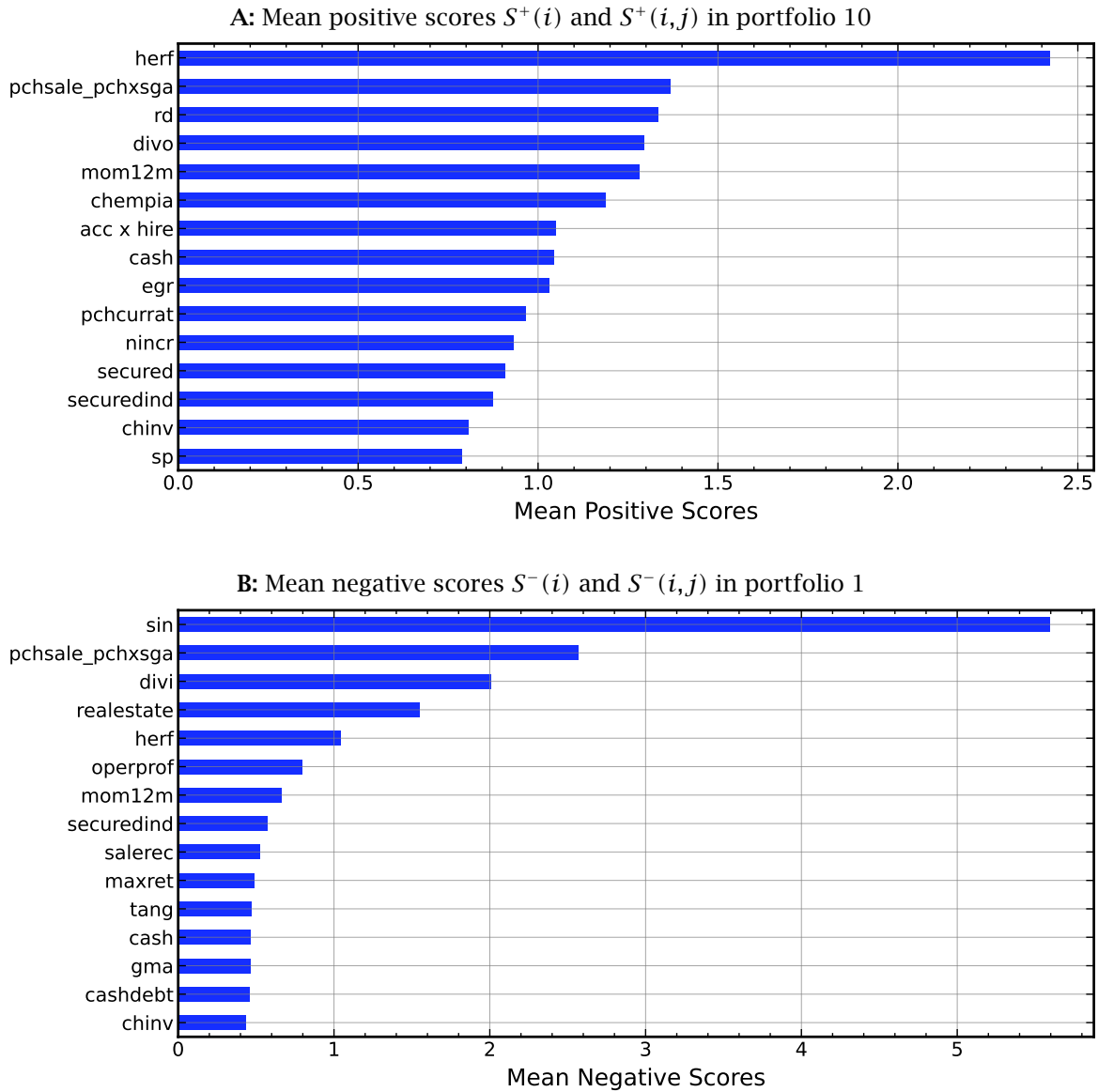
Note that this step, once again, relies on the separable model structure, which allows for exact decomposition of the model output (the portfolio signals) into term-level contributions (the shape function outputs).

The resulting importance rankings are presented in Figure 9. Panel A reports the top 15 effects in terms of S_i^+ and S_{ij}^+ .¹¹ The most prominent driver of long signals is industry sales concentration *herf*, with $S^+ = 2.42$, far exceeding the remaining effects. This aligns with the previous finding that *herf* is the most important contributor to model signals overall. The decomposition clarifies that its importance is mainly due to positive contributions toward the long leg, i.e., the low-*herf* region (see Figure 6). Beyond this leading term, the next tier of effects have more comparable magnitudes ($S^+ \approx 0.79$ – 1.37), suggesting that the model's long selection is supported by a diverse set of reinforcing signals rather than a single-factor rule.

The leading contributors fall into several intuitive themes. First, the inclusion of *mom12m* among the top effects indicates that recent price trends materially strengthen long signals, consistent with momentum serving as an important selection criterion. Second, the presence of *rd* (R&D increase) and *egr* (growth in common shareholder equity) indicates that the long leg is tilted toward firms with characteristics associated with investment and growth. Third, fundamentals and operating dynamics play a central role: *nincr* (number of earnings increases) contributes positively, pointing to a preference for firms exhibiting improving earnings patterns, and *pchsale_pchxsga* (percentage change in sales minus percentage change in SG&A) suggests that sales growth outpacing overhead growth is rewarded in the signal, consistent with operating leverage or efficiency considerations. Relatedly, *chinv* (change in inventory) appears among the top long contributors, indicating that inventory dy-

¹¹Figure I.1 in the Appendix also reports leg-wise mean score decompositions when positive and negative contributions are taken into account, confirming the main drivers.

Figure 9: Importance of predictors for long/short allocation



Note: This figure shows the most important effects of the AlphaGlass model for allocation in the long (Panel A) and short leg (Panel B) of the 10-1 portfolio. Importance values are calculated as mean positive (negative) scores on the training data, respectively.

namics contain incremental information for long selection in the learned portfolio mapping.

The long-leg ranking also highlights balance-sheet and financing characteristics. `cash` (cash holdings) and `pchcurrat` (percentage change in current ratio) both contribute positively, consistent with the model assigning higher long signals to firms with stronger liquidity positions or improving short-term solvency. At the same time, `secured` (secured debt) and `securedind` (secured debt indicator) feature among the top positive contributors, implying that secured borrowing is, in this learned allocation, associated with higher signals for long-decile firms, potentially capturing differences in collateralization, financing constraints, or capital structure that the model exploits together with other characteristics.

Moreover, we find further evidence that AlphaGlass leverages nonlinear interaction information. The strong ranking of `acc` & `hire` indicates that the model conditions the contribution of working-capital accruals (`acc`) on the firm's employee growth (`hire`). In other words, the long signal is not driven by accruals or hiring alone, but by particular combinations of these variables in regions where the learned bivariate shape function contributes positively. Likewise, `chempia` (industry-adjusted change in employees) underscores that labor-force adjustments relative to industry peers are informative for selection in the long leg.

Panel B presents the analogous ranking for the short leg. The dominant driver of short signals is `sin`, the `sin` stock indicator, with a magnitude of $|S^-| = 5.60$, substantially larger than all other terms. This implies that, among bottom-decile assets, the learned `sin` shape function frequently takes values that strongly reduce the signal, making `sin` the single most influential downward contributor to short-leg inclusion. A second large effect is `pchsale_pchxsga` ($|S^-| = 2.57$), indicating that the same operating-dynamics variable that supports long signals in the top decile also meaningfully contributes to low signals in the bottom decile, consistent with the shape function described above in which unfavorable regions of sales-minus-overhead dynamics push assets toward the short leg.

Several additional terms with sizable negative contributions reflect payout and balance-sheet characteristics. `divi` (dividend initiation) appears prominently ($|S^-| = 2.01$), suggesting that, within the short decile, dividend initiations are associated with signal-reducing regions of the learned shape function. `realestate` (real estate holdings) also ranks highly ($|S^-| = 1.55$), and `herf` (industry sales concentration) again appears ($|S^-| = 1.04$), indicating that industry structure and asset composition features contribute materially to the model's negative scoring within the short-decile set.

Profitability-related characteristics also feature among the main downward contributors: `operprof` (operating profitability) and `gma` (gross profitability) both appear in the top 15. Their presence in the short-leg ranking should be interpreted through the lens of the learned nonlinear shape functions: among bottom-decile stocks, the realized profitability values tend to fall in regions that further reduce the signal. Relatedly, `salerec` (sales-to-receivables) contributes negatively, consistent with the model using working-capital/receivables information when forming low signals.

The short-leg ranking also highlights risk and financing variables. `maxret` (maximum daily return) enters among the top effects, suggesting that extreme daily moves (a proxy for lottery-like behavior or crash/optionality exposure) are associated with lower signals in the short decile. `securedind` (secured

debt indicator) again appears, indicating that the secured-debt dimension is informative for both legs, though in different regions of the covariate space. Finally, several liquidity and debt-service measures (cash and cashdebt, cash flow to debt) contribute to signal reductions in the bottom decile, consistent with the model penalizing weaker financing capacity or adverse liquidity/debt-service configurations among short candidates. Inventory dynamics (chinv) also appear again, suggesting that inventory changes contain incremental information for both extreme tails of the signal distribution.

Two patterns stand out when comparing Panels A and B. First, some characteristics are important on both sides, most notably pchsale_pchxsga, herf, mom12m, securedind, cash, and chinv. This is consistent with a nonlinear score function that uses the same covariates to separate winners from losers by assigning opposite-signed contributions in different regions of the feature space (and/or by interacting with other terms), rather than employing disjoint sets of predictors for longs and shorts. Second, each leg also has clear distinctive drivers: the long leg features growth/investment and improving fundamentals (e.g., rd, egr, nincr) as well as a salient interaction (acc $\times$ hire), whereas the short leg is dominated by sin and places greater emphasis on payout/initiation and asset-composition variables (divi, realestate), together with measures related to tail risk (maxret) and profitability/quality (operprof, gma).

Overall, the signal decomposition indicates that AlphaGlass’s asset selection is driven by a combination of broad economic themes, such as industry structure, operating dynamics, profitability/quality, liquidity/financing, and risk, while allowing for strong nonlinearities and interaction effects. Importantly, the overlap of key terms across the long and short tails suggests that the model primarily differentiates positions through directional contributions of the same underlying characteristics, rather than by relying on entirely separate sets of variables for the two legs.

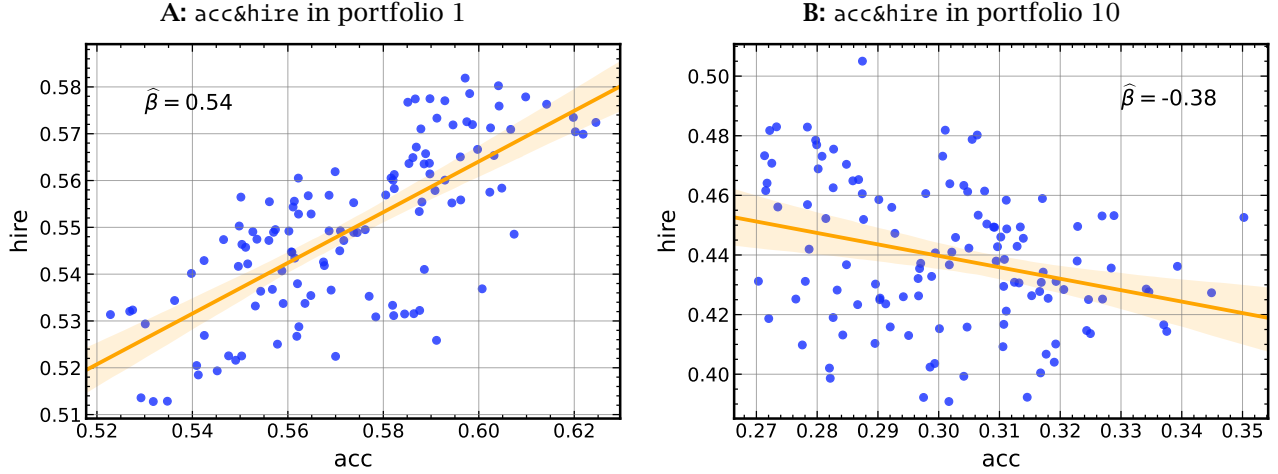
5.2 Interactions

In the previous sections, we found that, in addition to the univariate effects discussed above, the AlphaGlass model relies on interactions in its decision process, particularly the acc&hire interaction. In what follows, we study the relationship between this interaction and AlphaGlass portfolios. Recall the heatmap and the conditional effects displayed in Figure 7. Analogously to the univariate effects, we can derive implications from this learned shape function. We observe the highest interaction scores when low acc values coincide with high hire values. When acc is high, on the other hand, we find the lowest interaction outputs for low-hire stocks. This indicates that stocks with higher interaction scores, which are therefore more likely to be included in the long leg, should exhibit a negative correlation between acc and hire. The interaction shape suggests the opposite pattern for low score (likely shorted) stocks. Low interaction outputs are observed for stocks in which acc and hire are both low or both high, indicating a positive correlation between the two characteristics in these stocks.

The described effects are reflected in AlphaGlass portfolios. Panel A of Figure 10 shows the scatter plot of acc on the x -axis and hire of the first portfolio (the short leg of the AlphaGlass strategy) on the y -axis. The plot also shows the orange regression line with 95% confidence intervals

and the regression coefficient in the top corner. Consistent with the expectation, there is a positive relationship between `acc` and `hire` with a highly significant regression coefficient of 0.54 ($t = 11.4$). Panel B shows the same plot for portfolio 10 (the long leg of the AlphaGlass strategy), with the coefficient of -0.38 ($t = -3.3$) confirming the expected negative relationship.

Figure 10: `acc&hire` within AlphaGlass portfolios



Note: This figure shows scatter plots of `acc` and `hire` in AlphaGlass portfolios 1 and 10. The orange line is the regression line with 95% confidence intervals.

6 Mean-variance preferences

AlphaGlass allows flexible specification of the loss function. In principle, any differentiable function of the factor f_t^{HL} is permissible. For example, the estimation allows for portfolio constraints or penalties, see Appendix J, which considers drawdown penalties. While the SR criterion used in the baseline analysis is natural in asset pricing and portfolio construction, it is useful to consider alternative objective functions that more flexibly reflect investor preferences. In particular, mean-variance utility provides a standard formulation in which expected returns are traded off against return variance through a risk-aversion parameter. This allows us to examine whether the AlphaGlass framework continues to perform well when the estimation target is no longer scale-free but instead depends explicitly on the investor's risk tolerance.¹²

$$\mathcal{U}(f_t) = E(f_t) - \frac{\gamma}{2} \text{Var}(f_t), \quad (19)$$

¹²In Appendix K, we show how the Sharpe ratio and mean-variance objectives complement each other. Although they select the same portfolio rule under certain assumptions, the equivalence is not general, so it is informative to study the mean-variance case separately. In this section, we therefore estimate the model for mean-variance preferences over returns f_t :

where γ is the coefficient of absolute risk aversion. We estimate the AlphaGlass model using the (dis)utility of the long/short portfolio, f^{HL} , as the loss function:

$$\mathcal{L} = -\mathcal{U}(f_t^{\text{HL}}). \quad (20)$$

Table 6 shows the maximized in-sample and out-of-sample utilities, $\mathcal{U}$, for $\gamma = 1$ and $\gamma = 3$ for AlphaGlass and the three benchmark models. AlphaGlass delivers the strongest out-of-sample performance for both levels of risk aversion. For $\gamma = 1$, its out-of-sample utility is 1.12, compared with 0.99 for EBM, 0.86 for NN, and 0.53 for RF. For $\gamma = 3$, the ranking is unchanged: AlphaGlass again performs best, with an out-of-sample utility of 1.23, followed by EBM (0.87), NN (0.79), and RF (0.34). Therefore, AlphaGlass remains the most effective method even when the objective is changed from Sharpe ratio maximization to a utility criterion. Moreover, the results reveal a pronounced contrast between in-sample and out-of-sample performance. The benchmark models achieve substantially larger in-sample utility values than AlphaGlass. For example, when $\gamma = 1$, RF and EBM reach in-sample utilities of 8.56 and 7.88, respectively, far above the AlphaGlass value of 1.81. A similar pattern holds for $\gamma = 3$, where RF and EBM again produce much larger in-sample utilities than AlphaGlass. However, these in-sample gains do not survive out-of-sample. Instead, all three benchmark models experience large declines, whereas AlphaGlass retains a much larger share of its in-sample performance.

Overall, these results underscore the broader point that AlphaGlass is not confined to maximizing the Sharpe ratio. Rather, the framework can accommodate alternative portfolio objectives while preserving its central features: end-to-end estimation, economic interpretability, and strong out-of-sample performance.

Table 6: Maximized utility $\mathcal{U}$ of mean-variance preferences

	In-sample	Out-of-sample
Panel A: $\gamma = 1$		
AlphaGlass	1.81	1.12
RF	8.56	0.53
NN	3.51	0.86
EBM	7.88	0.99
Panel B: $\gamma = 3$		
AlphaGlass	1.43	1.23
RF	7.48	0.34
NN	3.21	0.79
EBM	6.70	0.87

Note: This table compares maximized mean-variance utility for $\gamma = 1, 3$. The AlphaGlass model is estimated using a mean-variance loss function (20). The RF, NN, and EBM models use the MSE loss function. The results in the table are based on the long/short portfolios implied by the models and multiplied by 100.

7 AlphaGlass theoretical properties

In this section, we derive important properties of the AlphaGlass approach with respect to its learning behavior. Complete proofs are in Appendix K.1, while Appendix K.2 provides additional theoretical results. In particular, it clarifies the relationship between the Sharpe ratio and mean-variance objectives, and justifies the advantages of our implementation compared to a softmax-based formulation.¹³

As detailed in Section 2, AlphaGlass is trained through a differentiable approximation to the standard hard rank-and-cut construction. This raises a practical issue: the economically meaningful object is the hard-sorted quantile portfolio, while optimization uses a smooth surrogate. First, we justify our training approach by showing that the differentiable soft sorts we employ serve as a faithful proxy for the hard portfolio procedures applied to the test data, and that their gradients used during backpropagation are appropriate.

We fix the temperature parameters $\tau = (\tau_{\text{rank}}, \tau_{\text{mask}})$. For each θ and t , let $w_t^{\text{hard}}(\theta|\tau)$ be the portfolio weights obtained by exact ranking and truncation (e.g., invest only in the top-ranked assets that pass the long/short mask). Let $w_t^{\text{soft}}(\theta|\tau)$ be the corresponding differentiable approximation with temperature vector τ entering the soft ranking and masking. We denote the associated empirical Sharpe ratios by $\widehat{\text{SR}}_T^{\text{hard}}(\theta)$ and $\widehat{\text{SR}}_T^{\text{soft}}(\theta|\tau)$.

To quantify how clearly the model separates selected and non-selected assets, we introduce a margin: for each (θ, t) , let $\Delta_\theta^{(t)} > 0$ measure the smallest pairwise gap between signals in month t . Theorem 1 formalizes that our differentiable portfolio construction is a faithful surrogate for the usual hard top-bottom sort:

Theorem 1 (Soft-hard equivalence)

For fixed θ , the random variables $\Delta_\theta^{(t)}$ have a continuous density at 0 (no point mass at ties). Then:

1. *Pointwise convergence given a margin: Fix t and suppose $\Delta_\theta^{(t)} \geq \delta > 0$. Then there exist constants $C_1, C_2, c > 0$ depending only on q and n_{max} such that for all sufficiently small τ ,*

$$\|w_t^{\text{soft}}(\theta|\tau) - w_t^{\text{hard}}(\theta)\|_1 \leq C_1 e^{-c\delta/\tau_{\text{rank}}} + C_2 e^{-c/\tau_{\text{mask}}}.$$

2. *Averaged convergence and Sharpe ratio: Let $\tau_T \downarrow 0$. Then*

$$\frac{1}{T} \sum_{t=1}^T \|w_t^{\text{soft}}(\theta|\tau_T) - w_t^{\text{hard}}(\theta)\|_1 \xrightarrow{p} 0, \quad \widehat{\text{SR}}_T^{\text{soft}}(\theta|\tau_T) - \widehat{\text{SR}}_T^{\text{hard}}(\theta) \xrightarrow{p} 0.$$

3. *Gradient consistency at non-tie points: Fix T and θ such that $\Delta_\theta^{(t)} > 0$ for all $t = 1, \dots, T$ (no ties*

¹³Note that, although the formal statements below are written for the SR criterion, the arguments underlying Theorems 1–3 are more general. Theorem 1 is primarily a statement about the soft rank-and-mask construction and its approximation of the corresponding hard-sorted portfolio and for Theorems 2 and 3 the proofs establish convergence of the empirical portfolio moments $\hat{\mu}_T(\theta)$ and $\hat{\sigma}_T(\theta)$ and then pass these moment bounds to the objective of interest. The same logic extends to other criteria of the form $Q(\theta) = \phi(\mu(\theta), \sigma(\theta))$, provided the objective is well behaved (e.g., preserves continuity and bounded gradients). Mean-variance utility, $Q(\theta) = \mu(\theta) - \frac{\gamma}{2} \sigma^2(\theta)$, is a simple example covered by this extension.

in the sample at θ) and $\hat{\sigma}_T^{\text{hard}}(\theta) > 0$. Then $\widehat{\text{SR}}_T^{\text{hard}}(\theta)$ is locally constant in a neighborhood of θ , hence $\partial^{\text{Clarke}} \widehat{\text{SR}}_T^{\text{hard}}(\theta) = \{0\}$, and

$$\|\nabla \widehat{\text{SR}}_T^{\text{soft}}(\theta|\tau)\| \rightarrow 0 \quad \text{as } \tau \downarrow 0.$$

This theorem formalizes multiple desirable properties:

- *Soft approximation:* If, at time t , the signal scores exhibit a nontrivial margin, then the soft weights $w_t^{\text{soft}}(\theta|\tau)$ are exponentially close (in ℓ_1) to the hard equal-weight portfolio $w_t^{\text{hard}}(\theta)$ as τ_{rank} goes to zero. Intuitively, once the scores are well separated everywhere, the sigmoids in the soft ranking/masking behave like indicators, so the differentiable construction reproduces the hard rank-and-cut portfolio.
- *Time-average and Sharpe ratio consistency:* Under stationarity/weak dependence and no mass at exact ties, an annealing schedule $\tau_T \downarrow 0$ yields

$$\frac{1}{T} \sum_{t=1}^T \|w_t^{\text{soft}}(\theta|\tau) - w_t^{\text{hard}}(\theta)\|_1 \xrightarrow{p} 0, \quad \widehat{\text{SR}}_T^{\text{soft}}(\theta) - \widehat{\text{SR}}_T^{\text{hard}}(\theta) \xrightarrow{p} 0.$$

Therefore, in large samples, the differentiable construction targets the same portfolio and the same sample Sharpe ratio as the discrete benchmark.¹⁴

In other words, Theorem 1 provides a direct bridge between the portfolio procedure used for optimization and the hard quantile portfolios used for evaluation: whenever the model produces clear cross-sectional separation, the soft weights are exponentially close to the hard weights, and the corresponding Sharpe ratios match in large samples. This ensures that the population and sample Sharpe ratios appearing in our theoretical results can be interpreted as Sharpe ratios of the standard hard-sorted long-short portfolios, up to a vanishing approximation error. At the same time, at parameter values where the ranking is stable (no ties), the hard Sharpe ratio is locally flat, and the gradient of the soft Sharpe vanishes as the temperature goes to zero. This formalizes the idea that the soft relaxation correctly captures the non-smooth structure of the true optimization problem, while still allowing gradient-based training.

We formalize that, within the AlphaGlass class, maximizing the sample Sharpe ratio is asymptotically aligned with maximizing the population Sharpe ratio.¹⁵ Given a sample of length T , we define the *population* (or out-of-sample) mean and volatility of the portfolio return as

$$\mu(\theta) = \mathbb{E}[\pi_t(\theta)], \quad \sigma^2(\theta) = \text{Var}(\pi_t(\theta)),$$

and the population Sharpe ratio as

$$\text{SR}(\theta) = \frac{\mu(\theta)}{\sigma(\theta)}.$$

¹⁴We also show the complementary property of gradient validity: The gradient of the soft objective converges (as temperatures go to zero) to a valid Clarke (1975) subgradient of the non-differentiable hard-sort objective. Away from ties, the hard objective is locally flat, and the soft gradient vanishes, indicating that using the smooth surrogate provides principled gradients for training against a hard-sorting target (see Appendix K.1 for more details).

¹⁵For clarity, we work with a simplified notation in this section and refer to the technical assumptions and detailed proofs in the Appendix.

Given a sample of length T , the empirical analogs are

$$\hat{\mu}_T(\theta) = \frac{1}{T} \sum_{t=1}^T \pi_t(\theta), \quad \hat{\sigma}_T^2(\theta) = \frac{1}{T} \sum_{t=1}^T (\pi_t(\theta) - \hat{\mu}_T(\theta))^2,$$

and the empirical Sharpe ratio

$$\widehat{\text{SR}}_T(\theta) = \frac{\hat{\mu}_T(\theta)}{\hat{\sigma}_T(\theta)}.$$

Theorem 2 formally justifies optimizing the Sharpe ratio directly:¹⁶

Theorem 2 (Consistency of Sharpe maximization within the AlphaGlass class)

Fix $\tau_{\text{rank}}, \tau_{\text{mask}} > 0$ and define the population and empirical maximizers of the Sharpe ratio as

$$\theta^* \in \operatorname{argmax}_{\theta \in \Theta} \text{SR}(\theta), \quad \hat{\theta}_T \in \operatorname{argmax}_{\theta \in \Theta} \widehat{\text{SR}}_T(\theta).$$

Then:

1. $\sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| \xrightarrow{p} 0$.
2. Any sequence of empirical maximizers satisfies $\text{dist}(\hat{\theta}_T, \operatorname{argmax} \text{SR}) \xrightarrow{p} 0$. If θ^* is unique, $\hat{\theta}_T \xrightarrow{p} \theta^*$.
3. Consequently, $\widehat{\text{SR}}_T(\hat{\theta}_T) \xrightarrow{p} \text{SR}(\theta^*)$.

In other words, asymptotically, AlphaGlass attains the best Sharpe available within its interpretable class. The theorem shows that our training objective (the sample Sharpe ratio computed from the soft long-short portfolio) is statistically aligned with the economic target (the population Sharpe ratio). In particular, it provides:

- *Uniform learnability*: As $T \rightarrow \infty$, the entire empirical estimate $\theta \mapsto \widehat{\text{SR}}_T(\theta)$ converges uniformly (in probability) to the true $\theta \mapsto \text{SR}(\theta)$. In other words, wherever we look in the model space, the sample Sharpe ratio is a reliable proxy for the population Sharpe ratio.
- *Consistency*: For any sequence of empirical maximizers $\hat{\theta}_T \in \operatorname{argmax}_{\theta \in \Theta} \widehat{\text{SR}}_T(\theta)$ converges to the set of population maximizers $\operatorname{argmax}_{\theta \in \Theta} \text{SR}(\theta)$ in probability. If the maximizer is unique, then $\hat{\theta}_T \rightarrow \theta^*$. Therefore, maximizing the sample Sharpe ratio asymptotically recovers the best-in-class model (at the population level).
- *Value optimality*: The achieved Sharpe ratio at the fitted model converges to the best attainable Sharpe ratio within the class:

$$\widehat{\text{SR}}_T(\hat{\theta}_T) \rightarrow \text{SR}(\theta^*).$$

In practice, this means that training AlphaGlass by maximizing the in-sample Sharpe ratio is asymptotically “correct”: With enough data, the resulting parameter vector yields a portfolio whose population Sharpe ratio approaches the highest Sharpe ratio achievable within the chosen architecture.

¹⁶For clarity, we work with a simplified notation of the mathematical properties in this section and refer to the technical assumptions and detailed explanations in Appendix K.1.

This raises the question of how reliably the AlphaGlass portfolio that maximizes the *in-sample* Sharpe ratio performs *out-of-sample*. Even if maximizing sample Sharpe ratio targets the population optimum (Theorem 2), a finite training sample can still produce over-optimistic in-sample Sharpe ratios. The next result gives a finite-sample reliability statement: with high probability, the population Sharpe ratio of the fitted AlphaGlass strategy is close to the best attainable population Sharpe ratio within the class, with a gap controlled by an explicit measure of effective strategy complexity, sample length, and temporal dependence.

For this, we consider the class of portfolio returns

$$\Pi = \{\pi_t(\theta) = w_t(\theta)^\top r_{t+1} : \theta \in \Theta\}.$$

For a block length $b \in \{1, \dots, \lfloor T/2 \rfloor\}$, let

$$m = \lfloor \frac{T}{2b} \rfloor,$$

and denote by $\mathfrak{R}_{T,b}^{\text{blk}}(\Pi)$ the blocked (dependent-data) Rademacher complexity of Π , as defined in Appendix K.1. Intuitively, $\mathfrak{R}_{T,b}^{\text{blk}}(\Pi)$ measures the effective richness of the model class after accounting for temporal dependence through blocking.¹⁷

Theorem 3 (Oracle inequality for out-of-sample Sharpe ratio)

Assume that returns are bounded and weakly dependent, that the AlphaGlass map $\theta \mapsto w_t(\theta)$ is Lipschitz, and that the portfolio volatility is uniformly bounded away from zero. Then for every $\delta \in (0, 1)$,

$$\Pr\left(\text{SR}(\theta^*) - \text{SR}(\hat{\theta}_T) \leq C_1 \mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + C_2 \left(\frac{b}{T} + \sqrt{\frac{\log(8/\delta)}{m}}\right)\right) \geq 1 - \delta - 2m\beta(b),$$

for some constants $C_1, C_2 > 0$ depending only on basic properties of the data and the model. Here $\beta(b)$ denotes the β -mixing coefficient at lag b .

In particular, if $b = b_{T,\delta}$ is chosen so that

$$2m\beta(b) \leq \delta/2,$$

then

$$\Pr\left(\text{SR}(\theta^*) - \text{SR}(\hat{\theta}_T) \leq C_1 \mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + C_2 \left(\frac{b}{T} + \sqrt{\frac{\log(8/\delta)}{m}}\right)\right) \geq 1 - \delta.$$

This property is important because it guarantees that the empirically Sharpe-optimal AlphaGlass portfolio generalizes: its out-of-sample Sharpe ratio is close to that of the best possible portfolio within the model class. The gap is controlled by three components: the effective complexity of the strategy class, captured by $\mathfrak{R}_{T,b}^{\text{blk}}(\Pi)$, a finite-sample remainder b/T , and a concentration term based on the effective number of approximately independent blocks, $m = \lfloor T/(2b) \rfloor$.

This result matters economically because it converts a statistical training objective into a reliability statement about investment performance. Optimizing the sample Sharpe ratio does not

¹⁷Intuitively, $\mathfrak{R}_{T,b}^{\text{blk}}(\Pi)$ measures how well the class Π can fit random noise in the data after accounting for weak dependence through block aggregation: larger values mean higher effective model capacity and thus a greater potential for overfitting. See, for example, Shalev-Shwartz and Ben-David (2014) for a general overview.

merely fit noise: with high probability, the resulting portfolio’s population Sharpe ratio is near-best-in-class, and the gap vanishes as the sample length grows provided the effective complexity remains controlled and dependence is not too strong. Therefore, the bound can be read as a transparent data-sufficiency condition. To make the SR shortfall small, the sample length T must be large relative to the model’s effective complexity and relative to the dependence-adjusted block structure. This formalizes a fundamental trade-off in portfolio construction: increasing model flexibility can raise the attainable in-class Sharpe ratio, but it also requires more data for the same level of out-of-sample reliability.

In particular, the number of effective model parameters influences reliability through its effect on the complexity term $\mathfrak{R}_{T,b}^{\text{blk}}(\Pi)$. Activating many terms, especially interactions, increases the number of ways the strategy can adapt to the sample, so more data are required for the same level of out-of-sample reliability. At the same time, this underscores that sparsity and pruning are not merely interpretability devices. They directly reduce effective model complexity and thereby tighten the bound, making it more likely that in-sample Sharpe improvements persist out-of-sample.

Taken together, these results provide a unified theoretical justification highlighting the approach’s practical benefits. The differentiable rank-and-mask procedure faithfully approximates hard quantile portfolios (Theorem 1), maximizing sample Sharpe targets the population Sharpe ratio within the AlphaGlass class (Theorem 2), and finite-sample generalization guarantees explain when a high in-sample Sharpe ratio is informative about out-of-sample performance (Theorem 3).

8 Monte Carlo simulation

While AlphaGlass is guaranteed to target the maximum attainable Sharpe ratio within its model class, this guarantee does not, by itself, provide a numerical comparison against richer population benchmarks. In simulation, we can go one step further. Because the data-generating process is known, we can evaluate each fitted model on an arbitrarily large independent test sample, so that the resulting out-of-sample Sharpe ratio is an arbitrarily accurate approximation to the model’s population Sharpe ratio. This allows us to study not only whether AlphaGlass outperforms benchmark learners in finite samples, but also how close it comes to the simulated population optimum and to oracle portfolio rules that are infeasible in practice.

8.1 Simulation design

Concretely, our simulations are designed as follows. In each month t , we draw a cross section of $n = 500$ assets with $p = 6$ characteristics. Let

$$\mathbf{x}_{l,t} = (x_{l,t,1}, \dots, x_{l,t,p})^\top, \quad l = 1, \dots, n,$$

and assume that the raw characteristics are jointly Gaussian with unit variances and common pairwise correlation 0.25. As in the empirical analysis, each characteristic is then transformed into a

cross-sectional rank on $[0,1]$:

$$u_{l,t,j} = \frac{\text{rank}(x_{l,t,j})}{n} \in [0,1], \quad U_{l,t} = (u_{l,t,1}, \dots, u_{l,t,p})^\top.$$

This normalization keeps the simulated design close to the empirical implementation and removes any role for the marginal scale of the characteristics. Monthly returns are generated according to

$$r_{l,t+1} = \mu_{\text{sig}} m(U_{l,t}) + \sigma(U_{l,t}) \varepsilon_{l,t+1}, \quad \varepsilon_{l,t+1} \stackrel{i.i.d.}{\sim} N(0,1), \quad (21)$$

with $\mu_{\text{sig}} = 0.018$, where the conditional mean component is

$$m(U_{l,t}) = a_1 (2u_{l,t,1} - 1)^3, \quad a_1 = 1. \quad (22)$$

The conditional idiosyncratic volatility is characteristic-dependent and given by

$$\log \sigma(U_{l,t})^2 = g_0 + g_1 u_{l,t,2} + g_2 (u_{l,t,3} - \frac{1}{2})^2, \quad (23)$$

with $(g_0, g_1, g_2) = (-4.6, 2.5, 1.2)$.

For each Monte Carlo replication, we simulate a training panel of length $T \in \{120, 240\}$ months and estimate four models: AlphaGlass, Explainable Boosting Machines (EBM), Random Forests and a neural network. Consistent with our empirical implementation, EBM, Random Forest, and the neural net are estimated as pooled return-prediction models, whereas AlphaGlass is trained directly on the time-series Sharpe-ratio objective. We then evaluate each fitted model on an independent large population sample of 1000 months. This separation ensures that the reported performance reflects out-of-sample portfolio quality rather than in-sample fit.

To evaluate all methods on a common footing, we translate fitted scores into the same portfolio rule used in the empirical analysis: in each month, assets are sorted into deciles by model scores, and we form the equal-weight long-short 10-1 portfolio. Let $D_{1,t}(s)$ and $D_{10,t}(s)$ denote the bottom and top deciles under score s . The corresponding portfolio return is

$$R_{t+1}^{10-1}(s) = \frac{1}{|D_{10,t}(s)|} \sum_{l \in D_{10,t}(s)} r_{l,t+1} - \frac{1}{|D_{1,t}(s)|} \sum_{l \in D_{1,t}(s)} r_{l,t+1},$$

and performance is measured by the annualized Sharpe ratio of $\{R_{t+1}^{10-1}(s)\}$.

8.2 Oracle benchmarks

A central advantage of the simulation design is that it makes latent benchmarks observable. Because the conditional mean and conditional volatility functions are known, we can construct oracle portfolio rules that are infeasible in practice and use them to assess how much of the Sharpe-relevant signal each method recovers. Our simulation uses three such oracle benchmarks. The first is a *weighting oracle* that uses the true conditional means and variances and may choose continuous zero-investment weights. For each month t , let

$$\mu_t = (\mu_{1,t}, \dots, \mu_{n,t})^\top, \quad \mu_{l,t} = \mu_{\text{sig}} m(U_{l,t}).$$

The weighting oracle solves

$$\max_{w_t} \frac{w_t^\top \mu_t}{\sqrt{w_t^\top \Sigma_t w_t}} \quad \text{s.t.} \quad \mathbf{1}^\top w_t = 0, \quad \|w_t\|_1 = 2. \quad (24)$$

With diagonal Σ_t , the solution is

$$w_{l,t}^w = \frac{2\tilde{w}_{l,t}}{\sum_{j=1}^n |\tilde{w}_{j,t}|}, \quad \tilde{w}_{l,t} = \frac{\mu_{l,t} - \lambda_t}{\sigma_{l,t}^2}, \quad \lambda_t = \frac{\sum_{j=1}^n \mu_{j,t} / \sigma_{j,t}^2}{\sum_{j=1}^n 1 / \sigma_{j,t}^2}. \quad (25)$$

This is the natural upper bound in the simulated economy, but it is not directly comparable to the estimated models because it uses continuous weights rather than equal-weight deciles.

The second oracle is a *ranking oracle*. It uses the same equal-weight 10-1 portfolio construction as the estimated models, but sorts assets on the true risk-adjusted score

$$q_{l,t}^R = \frac{\mu_{l,t} - \lambda_t}{\sigma_{l,t}^2}. \quad (26)$$

This oracle is the most informative benchmark for our purposes because it holds equal-weight portfolio construction fixed, in line with model evaluation, and isolates pure ranking quality. Any difference between AlphaGlass and this oracle, therefore, reflects imperfect recovery of the Sharpe-relevant score, rather than gains from alternative weighting schemes.

The third benchmark is a *mean oracle*, which also uses equal-weight decile portfolios but sorts on the true conditional mean only:

$$q_{l,t}^\mu = \mu_{l,t}. \quad (27)$$

This oracle is useful because it represents the best equal-weight decile strategy available to a method that learns only expected returns and ignores characteristic-dependent volatility.

8.3 Results

Table 7 reports Monte Carlo means and standard deviations across 100 replications. We can draw the following conclusions from the results. First, the oracle benchmarks are stable across the two training lengths. The weighting oracle attains an annualized Sharpe ratio of about 3.0, the ranking oracle about 2.38, and the mean oracle about 2.03. The gap between the ranking and mean oracles is economically meaningful, around 0.34 Sharpe ratio units for both $T = 120$ and $T = 240$. This is precisely the value of volatility adjustment in a decile-sort environment: even when the final portfolio is equal-weighted within each tail, sorting on a risk-adjusted score dominates sorting on expected return alone. The further gap between the weighting oracle and the ranking oracle, roughly 0.61, captures the additional gains from moving from equal-weight tail portfolios to fully optimal continuous weights.

Second, AlphaGlass dominates all three benchmark learners at both sample sizes. For $T = 120$, AlphaGlass attains an average annualized Sharpe ratio of 1.774, compared with 1.568 for EBM, 1.449 for the neural net, and 1.411 for Random Forest. For $T = 240$, these values rise to 1.970, 1.755, 1.557, and 1.438, respectively. As expected, all models benefit from longer training samples, but Alpha-

Glass remains the top performer throughout, while EBM is the strongest of the return-prediction benchmarks.

Third, the key comparison is the fraction of the ranking-oracle Sharpe ratio. AlphaGlass captures 74.3% of the ranking-oracle SR at $T = 120$ and 83.4% at $T = 240$. The corresponding values are 65.8% and 74.2% for EBM, 60.7% and 65.9% for the neural net, and 59.1% and 60.8% for Random Forest. Because all models and the ranking oracle are evaluated through the same equal-weight 10–1 decile rule, these differences reflect ranking quality rather than differences in portfolio implementation. The results, therefore, show that training directly on a Sharpe-ratio objective helps AlphaGlass learn the part of the signal that matters most for tail-portfolio formation.

Fourth, the comparison with the mean oracle shows that AlphaGlass not only learns the non-linear conditional mean well, but also recovers part of the incremental information contained in the volatility-related characteristics. At $T = 120$, AlphaGlass reaches 86.7% of the mean-oracle SR, compared with 76.7% for EBM, 70.9% for the neural net, and 68.9% for Random Forest. At $T = 240$, AlphaGlass reaches 97.7%, while EBM, the neural net, and Random Forest reach 86.9%, 77.1%, and 71.1%, respectively. Thus, AlphaGlass nearly closes the gap to the best mean-only decile strategy in the longer sample and simultaneously preserves a clear advantage relative to the alternative learners in the Sharpe-relevant ranking benchmark.

Overall, the simulation supports the empirical analysis. In a setting where the Sharpe-optimal portfolio depends on both expected return and characteristic-dependent risk, AlphaGlass learns a more useful portfolio score than methods trained as generic return predictors. The gains are robust across training lengths, become stronger with more data, and are particularly pronounced when evaluated relative to the equal-weight ranking oracle, which most closely mirrors empirical portfolio construction.

9 Conclusion

We demonstrate that portfolio objectives can be optimized directly using inherently interpretable machine learning. The AlphaGlass framework bridges flexible nonlinear learning and the transparency requirements of empirical asset pricing by combining an additive, glass-box architecture with a sparse set of pairwise interactions with an end-to-end differentiable portfolio formation layer that maps scores into long-short quantile weights. This design preserves full interpretability by construction: each portfolio decision admits an exact decomposition into main effects and interaction contributions, enabling economically meaningful attribution without relying on noisy post-hoc explainability methods.

Empirically, AlphaGlass delivers strong out-of-sample performance in U.S. equities, producing a pronounced return and Sharpe spread across signal-sorted portfolios and outperforming key benchmarks. Crucially, these gains come with *ex ante* transparency. The learned shape functions and interaction surfaces reveal clear economic patterns in how characteristics map onto optimal allocations, including industry structure, operating leverage dynamics, momentum, labor adjustment, and

Table 7: Results of Monte Carlo simulations

	$T = 120$	$T = 240$
Panel A: Oracle benchmarks		
Weighting oracle	2.995 (0.116)	2.982 (0.118)
Ranking oracle	2.388 (0.107)	2.368 (0.120)
Mean oracle	2.047 (0.123)	2.021 (0.105)
Panel B: Model Sharpe ratios		
AlphaGlass	1.774 (0.340)	1.970 (0.226)
EBM	1.568 (0.219)	1.755 (0.173)
Random Forest	1.411 (0.236)	1.438 (0.205)
Neural net	1.449 (0.204)	1.557 (0.191)
Panel C: Fraction of ranking-oracle SR		
AlphaGlass	0.743 (0.140)	0.834 (0.104)
EBM	0.658 (0.094)	0.742 (0.076)
Random Forest	0.591 (0.096)	0.608 (0.086)
Neural net	0.607 (0.085)	0.659 (0.086)
Panel D: Fraction of mean-oracle SR		
AlphaGlass	0.867 (0.164)	0.977 (0.112)
EBM	0.767 (0.100)	0.869 (0.081)
Random Forest	0.689 (0.106)	0.711 (0.090)
Neural net	0.709 (0.097)	0.771 (0.092)

This table reports Monte Carlo mean Sharpe ratios (annualized) with standard deviations in parentheses across 100 replications. Each replication simulates a panel with $n = 500$ assets and $p = 6$ characteristics per month. Models are trained on $T \in \{120, 240\}$ months and evaluated on an independent population sample of 1000 months. All estimated models are evaluated through the same equal-weight long-short 10–1 decile portfolio.

economically intuitive conditional relationships captured by sparse interactions. The term-level decomposition further clarifies which characteristics systematically push assets toward the long and short tails, providing a direct bridge from a high-performing portfolio rule to an interpretable set of economic drivers.

Finally, we theoretically justify our approach. We show that maximizing the sample Sharpe ratio is statistically aligned with maximizing the population Sharpe ratio within the AlphaGlass class and that the differentiable rank-and-mask procedure is a faithful surrogate for conventional hard sorts. Overall, AlphaGlass provides a practical template for learning tradable factors and portfolio rules that are simultaneously high-performing, economically grounded, and transparent—and therefore well suited for both portfolio construction and mechanism-based empirical asset-pricing inference.

References

- Abarbanell, Jeffery S. and Brian J. Bushee (1998) "Abnormal returns to a fundamental analysis strategy," *The Accounting Review*, 19–45.
- Ait-Sahalia, Yacine and Michael W. Brandt (2001) "Variable selection for portfolio choice," *The Journal of Finance*, 56 (4), 1297–1351.
- Ali, Ashiq, Lee-Seok Hwang, and Mark A. Trombley (2003) "Arbitrage risk and the book-to-market anomaly," *Journal of Financial Economics*, 69 (2), 355–373.
- Almeida, Heitor and Murillo Campello (2007) "Financial constraints, asset tangibility, and corporate investment," *The Review of Financial Studies*, 20 (5), 1429–1460.
- Amihud, Yakov and Haim Mendelson (1989) "The effects of beta, bid-ask spread, residual risk, and size on stock returns," *The Journal of Finance*, 44 (2), 479–486.
- Anderson, Christopher W. and Luis Garcia-Feijoo (2006) "Empirical evidence on capital investment, growth options, and security returns," *The Journal of Finance*, 61 (1), 171–194.
- Ang, Andrew, Robert J. Hodrick, Yuhang Xing, and Xiaoyan Zhang (2006) "The cross-section of volatility and expected returns," *The Journal of Finance*, 61 (1), 259–299.
- Ao, Mengmeng, Li Yingying, and Xinghua Zheng (2019) "Approaching mean-variance efficiency for large portfolios," *The Review of Financial Studies*, 32 (7), 2890–2919.
- Asness, Clifford S., R. Burt Porter, and Ross L. Stevens (2000) "Predicting stock returns using industry-relative firm characteristics."
- Balakrishnan, Karthik, Eli Bartov, and Lucile Faurel (2010) "Post loss/profit announcement drift," *Journal of Accounting and Economics*, 50 (1), 20–41.
- Bali, Turan G., Nusret Cakici, and Robert F. Whitelaw (2011) "Maxing out: Stocks as lotteries and the cross-section of expected returns," *Journal of Financial Economics*, 99 (2), 427–446.
- Bandyopadhyay, Sati P., Alan G. Huang, and Tony S. Wirjanto (2010) "The accrual volatility anomaly."
- Banz, Rolf W. (1981) "The relationship between return and market value of common stocks," *Journal of Financial Economics*, 9 (1), 3–18.
- Barbee Jr, William C., Sandip Mukherji, and Gary A. Raines (1996) "Do sales-price and debt-equity explain stock returns better than book-market and firm size?" *Financial Analysts Journal*, 52 (2), 56–60.
- Barr Rosenberg, Kenneth Reid and Ronald Lanstein (1985) "Persuasive evidence of market inefficiency," *The Journal of Portfolio Management*, 11 (3), 9–16.
- Barth, Mary E., John A. Elliott, and Mark W. Finn (1999) "Market rewards associated with patterns of increasing earnings," *Journal of Accounting Research*, 37 (2), 387–413.
- Basu, Sanjoy (1977) "Investment performance of common stocks in relation to their price-earnings ratios: A test of the efficient market hypothesis," *The Journal of Finance*, 32 (3), 663–682.
- Bell, Sebastian, Ali Kakhbod, Martin Lettau, and Abdolreza Nazemi (2025) "Glass box machine learning and corporate bond returns," Technical report, National Bureau of Economic Research.
- Belo, Frederico, Xiaoji Lin, and Santiago Bazdresch (2014) "Labor hiring, investment, and stock return predictability in the cross section," *Journal of Political Economy*, 122 (1), 129–177.
- Bhandari, Laxmi Chand (1988) "Debt/equity ratio and expected common stock returns: Empirical evidence," *The Journal of Finance*, 43 (2), 507–528.

- Black, Fischer and Robert Litterman (1992) "Global portfolio optimization," *Financial Analysts Journal*, 48 (5), 28–43.
- Brandt, Michael W. (1999) "Estimating portfolio and consumption choice: A conditional Euler equations approach," *The Journal of Finance*, 54 (5), 1609–1645.
- Brandt, Michael W. and Pedro Santa-Clara (2006) "Dynamic portfolio selection by augmenting the asset space," *The Journal of Finance*, 61 (5), 2187–2217.
- Brandt, Michael W., Pedro Santa-Clara, and Rossen Valkanov (2009) "Parametric portfolio policies: Exploiting characteristics in the cross-section of equity returns," *The Review of Financial Studies*, 22 (9), 3411–3447.
- Britten-Jones, Mark (1999) "The sampling error in estimates of mean-variance efficient portfolio weights," *The Journal of Finance*, 54 (2), 655–671.
- Brown, David P. and Bradford Rowe (2007) "The productivity premium in equity returns."
- Bryzgalova, Svetlana, Markus Pelger, and Jason Zhu (2025) "Forest through the Trees: Building Cross-Sections of Stock Returns," *Journal of Finance*, 80 (5), 2447–2506.
- Chandrashekar, Satyajit and Ramesh K. S. Rao (2009) "The productivity of corporate cash holdings and the cross-section of expected stock returns," *McCombs Research Paper Series No. FIN-03-09*.
- Chen, Long and Lu Zhang (2010) "A better three-factor model that explains more anomalies," *The Journal of Finance*, 65 (2), 563–595.
- Chen, Luyang, Markus Pelger, and Jason Zhu (2024) "Deep learning in asset pricing," *Management Science*, 70 (2), 714–750.
- Chen, Yifei, Bryan T. Kelly, and Dacheng Xiu (2022) "Expected returns and large language models," *Available at SSRN 4416687*.
- Chordia, Tarun, Avanidhar Subrahmanyam, and V. Ravi Anshuman (2001) "Trading activity and expected stock returns," *Journal of Financial Economics*, 59 (1), 3–32.
- Clarke, Frank H. (1975) "Generalized gradients and applications," *Transactions of the American Mathematical Society*, 205, 247–262.
- Cong, Lin William, Ke Tang, and Jingyuan Wang (2026) "AlphaPortfolio: Goal-Oriented Investment Management Through Deep Reinforcement Learning," Working paper.
- Cooper, Michael J., Huseyin Gulen, and Michael J. Schill (2008) "Asset growth and the cross-section of stock returns," *The Journal of Finance*, 63 (4), 1609–1651.
- Datar, Vinay T., Narayan Y. Naik, and Robert Radcliffe (1998) "Liquidity and stock returns: An alternative test," *Journal of Financial Markets*, 1 (2), 203–219.
- DeMiguel, Victor, Alberto Martin-Utrera, Francisco J. Nogales, and Raman Uppal (2020) "A transaction-cost perspective on the multitude of firm characteristics," *The Review of Financial Studies*, 33 (5), 2180–2222.
- Desai, Hemang, Shivaram Rajgopal, and Mohan Venkatachalam (2004) "Value-glamour and accruals mispricing: One anomaly or two?" *The Accounting Review*, 79 (2), 355–385.
- Didisheim, Antoine, Shikun Barry Ke, Bryan T. Kelly, and Semyon Malamud (2023) "Complexity in factor pricing models," *National Bureau of Economic Research*.
- Eberhart, Allan C., William F. Maxwell, and Akhtar R. Siddique (2004) "An examination of long-term abnormal stock returns and operating performance following R&D increases," *The Journal of Finance*, 59 (2), 623–650.
- Eisfeldt, Andrea L. and Dimitris Papanikolaou (2013) "Organization capital and the cross-section of expected returns," *The Journal of Finance*, 68 (4), 1365–1406.

- Fairfield, Patricia M., J. Scott Whisenant, and Teri Lombardi Yohn (2003) "Accrued earnings and growth: Implications for future profitability and market mispricing," *The Accounting Review*, 78 (1), 353–371.
- Fama, Eugene F. and Kenneth R. French (1993) "Common risk factors in the returns on stocks and bonds," *Journal of Financial Economics*, 33 (1), 3–56.
- (2015) "A five-factor asset pricing model," *Journal of Financial Economics*, 116 (1), 1–22.
- Fama, Eugene F. and James D. MacBeth (1973) "Risk, return, and equilibrium: Empirical tests," *Journal of Political Economy*, 81 (3), 607–636.
- Fedyk, Anastassia, Ali Kakhbod, Peiyao Li, and Ulrike Malmendier (2024) "AI and perception biases in investments: An experimental study," *Working paper*.
- Feng, Guanhao, Stefano Giglio, and Dacheng Xiu (2020) "Taming the factor zoo: A test of new factors," *The Journal of Finance*, 75 (3), 1327–1370.
- Francis, Jennifer, Ryan LaFond, Per M. Olsson, and Katherine Schipper (2004) "Costs of equity and earnings attributes," *The Accounting Review*, 79 (4), 967–1010.
- Freyberger, Joachim, Andreas Neuhierl, and Michael Weber (2020) "Dissecting characteristics nonparametrically," *The Review of Financial Studies*, 33 (5), 2326–2377.
- Gabaix, Xavier, Ralph S. J. Koijen, Robert J. Richmond, and Motohiro Yogo (2025) "Asset Embeddings," Working Paper 2025-51, Becker Friedman Institute for Economics.
- Gentleman, Eric and Joseph M Marks (2006) "Acceleration strategies," *SSRN Electronic Journal*.
- Giglio, Stefano and Dacheng Xiu (2021) "Asset pricing with omitted factors," *Journal of Political Economy*, 129 (7), 1947–1990.
- Green, Jeremiah, John R. M. Hand, and X. Frank Zhang (2017) "The characteristics that provide independent information about average US monthly stock returns," *The Review of Financial Studies*, 30 (12), 4389–4436.
- Gu, Shihao, Bryan Kelly, and Dacheng Xiu (2020) "Empirical asset pricing via machine learning," *The Review of Financial Studies*, 33 (5), 2223–2273.
- (2021) "Autoencoder asset pricing models," *Journal of Econometrics*, 222 (1), 429–450.
- Guo, Re-Jin, Baruch Lev, and Charles Shi (2006) "Explaining the Short-and Long-Term IPO Anomalies in the US by R&D," *Journal of Business Finance & Accounting*, 33 (3-4), 550–579.
- Hafzalla, Nader, Russell Lundholm, and E. Matthew Van Winkle (2011) "Percent accruals," *The Accounting Review*, 86 (1), 209–236.
- Hansen, Lars Peter and Ravi Jagannathan (1997) "Assessing specification errors in stochastic discount factor models," *The Journal of Finance*, 52 (2), 557–590.
- Harvey, Campbell R., Yan Liu, and Heqing Zhu (2016) "... and the cross-section of expected returns," *The Review of Financial Studies*, 29 (1), 5–68.
- Hastie, Trevor and Robert Tibshirani (1986) "Generalized Additive Models," *Statistical Science*, 1 (3), 297 – 310.
- Holthausen, Robert W. and David F. Larcker (1992) "The prediction of stock returns using financial statement information," *Journal of Accounting and Economics*, 15 (2-3), 373–411.
- Hong, Harrison and Marcin Kacperczyk (2009) "The price of sin: The effects of social norms on markets," *Journal of Financial Economics*, 93 (1), 15–36.
- Hou, Kewei and Tobias J. Moskowitz (2005) "Market frictions, price delay, and the cross-section of expected returns," *The Review of Financial Studies*, 18 (3), 981–1020.

- Hou, Kewei and David T. Robinson (2006) "Industry concentration and average stock returns," *The Journal of Finance*, 61 (4), 1927-1956.
- Hou, Kewei, Chen Xue, and Lu Zhang (2015) "Digesting anomalies: An investment approach," *The Review of Financial Studies*, 28 (3), 650-705.
- Huang, Alan Guoming (2009) "The cross section of cashflow volatility and expected stock returns," *Journal of Empirical Finance*, 16 (3), 409-429.
- Jegadeesh, Narasimhan (1990) "Evidence of predictable behavior of security returns," *The Journal of Finance*, 45 (3), 881-898.
- Jegadeesh, Narasimhan and Sheridan Titman (1993) "Returns to buying winners and selling losers: Implications for stock market efficiency," *The Journal of Finance*, 48 (1), 65-91.
- Jensen, Theis I., Bryan T. Kelly, and Lasse H. Pedersen (2023) "Is there a replication crisis in finance?" *Journal of Finance*, 78 (5), 2465-2518.
- Jiang, Guohua, Charles M. C. Lee, and Yi Zhang (2005) "Information uncertainty and expected returns," *Review of Accounting Studies*, 10 (2), 185-221.
- Jobson, J. David and Bob Korkie (1980) "Estimation for Markowitz efficient portfolios," *Journal of the American Statistical Association*, 75 (371), 544-554.
- Jorion, Philippe (1986) "Bayes-Stein estimation for portfolio analysis," *Journal of Financial and Quantitative Analysis*, 21 (3), 279-292.
- Kakhbod, Ali, Amir Kermani, and Bernardo Maciel (2025) "In the Fed's Mind," *SSRN working paper*.
- Kakhbod, Ali, Leonid Kogan, Peiyao Li, and Dimitris Papanikolaou (2024) "Measuring Creative Destruction," *SSRN working paper 5008685*.
- Kama, Itay (2009) "On the market reaction to revenue and earnings surprises," *Journal of Business Finance & Accounting*, 36 (1-2), 31-50.
- Kan, Raymond and Guofu Zhou (2007) "Optimal portfolio choice with parameter uncertainty," *Journal of Financial and Quantitative Analysis*, 42 (3), 621-656.
- Ke, Zheng Tracy, Bryan T. Kelly, and Dacheng Xiu (2019) "Predicting returns with text data," Technical report, National Bureau of Economic Research.
- Kelly, Bryan, Diogo Palhares, and Seth Pruitt (2023) "Modeling corporate bond returns," *The Journal of Finance*, 78 (4), 1967-2008.
- Kelly, Bryan, Seth Pruitt, and Yinan Su (2019) "Characteristics are covariances: A unified model of risk and return," *Journal of Financial Economics*, 134 (3), 501-524.
- Kelly, Bryan and Dacheng Xiu (2023) "Financial machine learning," *Foundations and Trends in Finance*, 13 (3-4), 205-363.
- Kishore, Runeet, Michael W. Brandt, Pedro Santa-Clara, and Mohan Venkatachalam (2008) "Earnings announcements are full of surprises."
- Kozak, Serhiy, Stefan Nagel, and Shrihari Santosh (2020) "Shrinking the cross-section," *Journal of Financial Economics*, 135 (2), 271-292.
- Lakonishok, Josef, Andrei Shleifer, and Robert W. Vishny (1994) "Contrarian investment, extrapolation, and risk," *The Journal of Finance*, 49 (5), 1541-1578.
- Ledoit, Olivier and Michael Wolf (2004) "Honey, I shrunk the sample covariance matrix," *The Journal of Portfolio Management*, 30 (4), 110-119.

- (2012) “Nonlinear shrinkage estimation of large-dimensional covariance matrices,” *The Annals of Statistics*, 40 (2), 1024–1060.
- Lerman, Alina, Joshua Livnat, and Richard R. Mendenhall (2007) “The high-volume return premium and post-earnings announcement drift.”
- Lettau, Martin (2023) “High-dimensional factor models and the factor zoo,” *National Bureau of Economic Research working paper no. 31719*.
- Lettau, Martin and Markus Pelger (2020) “Factors that fit the time series and cross-section of stock returns,” *The Review of Financial Studies*, 33 (5), 2274–2325.
- Lev, Baruch and Doron Nissim (2004) “Taxable income, future earnings, and equity values,” *The Accounting Review*, 79 (4), 1039–1074.
- Litzenberger, Robert H. and Krishna Ramaswamy (1982) “The effects of dividends on common stock prices tax effects or information effects?” *The Journal of Finance*, 37 (2), 429–443.
- Liu, Weimin (2006) “A liquidity-augmented capital asset pricing model,” *Journal of Financial Economics*, 82 (3), 631–671.
- Lou, Yin, Rich Caruana, Johannes Gehrke, and Giles Hooker (2013) “Accurate intelligible models with pairwise interactions,” in *Proceedings of the 19th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 623–631.
- Michaely, Roni, Richard H. Thaler, and Kent L. Womack (1995) “Price reactions to dividend initiations and omissions: Overreaction or drift?” *The Journal of Finance*, 50 (2), 573–608.
- Michaud, Richard O. (1989) “The Markowitz optimization enigma: Is ‘optimized’ optimal?” *Financial analysts journal*, 45 (1), 31–42.
- Mohanram, Partha S. (2005) “Separating winners from losers among lowbook-to-market stocks using financial statement analysis,” *Review of Accounting Studies*, 10 (2), 133–170.
- Moskowitz, Tobias J. and Mark Grinblatt (1999) “Do industries explain momentum?” *The Journal of Finance*, 54 (4), 1249–1290.
- Nori, Harsha, Samuel Jenkins, Paul Koch, and Rich Caruana (2019) “InterpretML: A unified framework for machine learning interpretability,” *arXiv:1909.09223*.
- Novy-Marx, Robert (2013) “The other side of value: The gross profitability premium,” *Journal of Financial Economics*, 108 (1), 1–28.
- Ou, Jane A. and Stephen H. Penman (1989) “Financial statement analysis and the prediction of stock returns,” *Journal of Accounting and Economics*, 11 (4), 295–329.
- Palazzo, Bernardino (2012) “Cash holdings, risk, and expected returns,” *Journal of Financial Economics*, 104 (1), 162–185.
- Piotroski, Joseph D. (2000) “Value investing: The use of historical financial statement information to separate winners from losers,” *Journal of Accounting Research*, 1–41.
- Pontiff, Jeffrey and Artemiza Woodgate (2008) “Share issuance and cross-sectional returns,” *The Journal of Finance*, 63 (2), 921–945.
- Richardson, Scott A., Richard G. Sloan, Mark T. Soliman, and Irem Tuna (2005) “Accrual reliability, earnings persistence and stock prices,” *Journal of Accounting and Economics*, 39 (3), 437–485.
- Rudin, Cynthia (2019) “Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead,” *Nature Machine Intelligence*, 1 (5), 206–215.

- Shalev-Shwartz, Shai and Shai Ben-David (2014) *Understanding machine learning: From theory to algorithms*, Cambridge University Press.
- Simon, Frederik, Sebastian Weibels, and Tom Zimmermann (2025) “Deep parametric portfolio policies,” *CFR Working Paper*.
- Slack, Dylan, Sophie Hilgard, Emily Jia, Sameer Singh, and Himabindu Lakkaraju (2020) “Fooling lime and shap: Adversarial attacks on post hoc explanation methods,” in *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, 180–186.
- Sloan, Richard G. (1996) “Do stock prices fully reflect information in accruals and cash flows about future earnings?” *The Accounting Review*, 289–315.
- Soliman, Mark T. (2008) “The use of DuPont analysis by market participants,” *The Accounting Review*, 83 (3), 823–853.
- Thomas, Jacob K. and Frank X. Zhang (2011) “Tax expense momentum,” *Journal of Accounting Research*, 49 (3), 791–821.
- Thomas, Jacob K. and Huai Zhang (2002) “Inventory changes and future returns,” *Review of Accounting Studies*, 7 (2), 163–187.
- Titman, Sheridan, K C. John Wei, and Feixue Xie (2004) “Capital investments and stock returns,” *Journal of Financial and Quantitative Analysis*, 39 (4), 677–700.
- Tuzel, Selale (2010) “Corporate real estate holdings and the cross-section of stock returns,” *The Review of Financial Studies*, 23 (6), 2268–2302.
- Valta, Philip (2016) “Strategic default, debt structure, and stock returns,” *Journal of Financial and Quantitative Analysis*, 51 (1), 197–229.

Appendix

Contents

A	Implementation details	48
A.1	Term selection and pruning	48
A.2	Interaction selection	48
A.3	Stage-wise training and early stopping	49
A.4	Data Loading and Batching	50
B	Variable definitions	52
C	Characteristic categories	54
D	Sharpe-maximizing benchmark	55
E	Alternative weights	56
F	Excluding small stocks	58
G	Alternative data split	59
H	Additional shape functions	60
I	Decile-wise decomposition	61
J	Drawdowns	63
K	Theoretical results	65
K.1	Proofs of theorems in Section 7	65
K.2	Additional results	87

A Implementation details

In addition to the description in Section 2, we explain the estimation algorithm and key design choices in greater detail.

A.1 Term selection and pruning

In addition to the regularization outlined in Section 2.3, we employ two further mechanisms to enforce sparsity: First, to limit the interaction search space before training the interaction block, we apply a heredity-aware screening step: we first fit main effects, form residuals against the main-effect signal, and compute a fast two-way importance score on binned features (a FAST-style score on residuals) while enforcing weak heredity by only considering pairs (i, j) that involve at least one main effect that remains active. The top K pairs, ranked by this score, are retained for interaction training. Second, after training (first on main effects, then on the retained interactions), we quantify each term’s contribution using the product of a variance scale and the squared gate: Let $z_i(l, t) = f_i(x_{i,l,t})$ and $z_{ij}(l, t) = f_{ij}(x_{i,l,t}, x_{j,l,t})$ denote raw subnet outputs (before gates), and let $\text{Var}_w[\cdot]$ denote the weighted variance over the training set. We define $C_i = \theta_i^2 \text{Var}_w[z_i]$ for main effects and $C_{ij} = y_{ij}^2 \text{Var}_w[z_{ij}]$ for interactions, and normalize across all effects via $\tilde{C}_k = C_k / \sum_\ell C_\ell$. Sorting effects in descending $\tilde{C}$ yields a regularization path along which we recompute the validation loss while progressively enabling more terms. For main effects, letting $\mathcal{M}(m)$ denote the model that keeps only the top m main effects (with interactions disabled), we evaluate $\mathcal{L}_{\text{val}}^{\text{main}}(m) = \mathcal{L}_{\text{SR}}^{\text{val}}(\mathcal{M}(m))$ for $m = 0, 1, \dots, N$ and select the smallest $\hat{m}$ whose validation loss lies within a relative tolerance ρ of the path minimum, i.e.,

$$\hat{m} = \min \left\{ m : (\mathcal{L}_{\text{val}}^{\text{main}}(m) - \min_{m'} \mathcal{L}_{\text{val}}^{\text{main}}(m')) / \min_{m'} \mathcal{L}_{\text{val}}^{\text{main}}(m') \leq \rho \right\}.$$

We then freeze the selected $\hat{m}$ main effects, activate the screened interactions in decreasing order of $\tilde{C}_{ij}$, and repeat the same tolerance rule to select the smallest number $\hat{k}$ of interactions whose validation loss is within ρ of the best point on the interaction path. Overall, the gates with ℓ_1 penalties provide continuous shrinkage during training, the heredity-aware screen keeps the interaction set manageable, and the tolerance-based pruning path performs discrete selection guided by out-of-sample Sharpe ratio. Together, these steps yield a sparse, interpretable model with reliable attribution, even in the presence of highly correlated inputs.

A.2 Interaction selection

We select pairwise interactions using an efficient binned preprocessor and an interaction detection algorithm (FAST of Lou et al. (2013)) that assigns each pair of feature indices (i, j) a score S_{ij} , measuring the improvement of a pair-specific model over an additive model on the same binned representation. Specifically, before screening interactions, we discretize each feature into a small number of bins. To construct bins, we use cut points so that each bin contains approximately the same number of samples. Let B_i be the number of bins for feature i . The preprocessor maps every sample x_n to an integer-coded matrix

$$X \in \{0, \dots, B_1 - 1\} \times \dots \times \{0, \dots, B_p - 1\},$$

so that $X_{n,i}$ is the bin index of sample n on feature i . This compressed integer form enables highly efficient counting, aggregation, and table lookups, which the interaction detector uses internally.

On the binned scale, we can define a purely additive (no interaction) model that depends only on features i and j :

$$\hat{y}_n^{\text{add}} = \mu + g_i(X_{n,i}) + g_j(X_{n,j}),$$

where $g_i: \{0, \dots, B_i - 1\} \rightarrow \mathbb{R}$ and $g_j: \{0, \dots, B_j - 1\} \rightarrow \mathbb{R}$ are univariate shape tables (one score per bin), i.e., y is explained using only per-bin offsets for i and j . The interaction-augmented counterpart adds a two-way table

$$\hat{y}_n^{\text{pair}} = \mu + g_i(X_{n,i}) + g_j(X_{n,j}) + h_{ij}(X_{n,i}, X_{n,j}),$$

where h_{ij} assigns an extra score to each bin pair $(b_i, b_j) \in \{0, \dots, B_i - 1\} \times \{0, \dots, B_j - 1\}$.

Given the binned data X , the labels y , optional sample weights w , and the current main-effect scores from the univariate fitting stage, our variant of the FAST algorithm computes how much additional predictive signal is captured by allowing a two-way table for (i, j) beyond the best additive fit on i and j alone. Concretely, for each pair (i, j) , the detector performs the following operations on the same binned grid:

1. Form sufficient statistics per bin (and per class for classification), e.g., weighted counts and weighted response sums in each i -bin, each j -bin, and each (i, j) bin-pair.
2. Fit the best additive tables (g_i, g_j) on $\{i, j\}$ by minimizing the chosen loss on bins subject to simple leaf-size constraints.
3. Fit the best additive-plus-pair model (g_i, g_j, h_{ij}) on the same binned data and constraints.
4. Compute the hold-out loss values $\mathcal{L}_{\text{add}}^{ij}$ and $\mathcal{L}_{\text{pair}}^{ij}$ on exactly the same representation.

The interaction score is the improvement:

$$S_{ij} = \mathcal{L}_{\text{add}}^{ij} - \mathcal{L}_{\text{pair}}^{ij} \geq 0,$$

so larger S_{ij} means the pair table h_{ij} explains structure that cannot be replicated by separate g_i and g_j . Because everything is computed on integer bins, these fits reduce to fast table updates rather than expensive neural training.

Additionally, heredity is enforced by restricting candidates to those where at least one of i or j is already active as a main effect. The top K pairs (as defined by a hyperparameter) are selected. Binning creates a compact representation that makes exhaustive pair screening computationally feasible. Comparing a pair model against an additive baseline isolates true interaction signal, while the heredity rule shrinks the search space and reduces spurious pairs by requiring that interactions involve features that already matter on their own. The capacity parameter K controls complexity and training cost.

A.3 Stage-wise training and early stopping

As outlined above, the full training procedure goes through the following stages: (i) train main-effect blocks only (while interaction blocks are frozen), (ii) add discovered interactions and train interaction blocks only (while main effects are frozen), (iii) jointly fine-tune all parameters with a

reduced learning rate.

They are implemented in the following way:

1. Main (univariate) effect learning: For each feature i , a neural subnetwork f_i is trained while minimizing the chosen loss function. Each subnet consists of one input neuron for x_i , a hidden layer of 20 units, and an output neuron. All subnets use ReLU activations and are trained in parallel via gradient descent. Training stops after a predefined number of epochs or when performance no longer improves (as defined by an early stopping threshold).
2. Interaction learning: After selecting the interaction pairs using the FAST algorithm (see Lou et al., 2013), neural subnetworks f_{ij} are trained while keeping the main effect subnetworks fixed. The interaction subnetworks have two input neurons (for x_i and x_j), two hidden layers of 20 neurons each, and one output neuron. Otherwise, training proceeds as in the first stage.
3. Joint fine-tuning: Once all main effect and interaction subnetworks are trained, a global optimization step is performed in which all parameters are updated simultaneously (with a lower learning rate to avoid losing information from the previous training steps). Due to the structural separability of the neural subnetworks, this stage does not compromise the interpretability of individual shape functions, but it allows small, coordinated adjustments that improve performance.

Each stage performs at most a fixed number of epochs, with an early-stopping tolerance hyperparameter tuned during validation.

A.4 Data Loading and Batching

We construct each training batch as a concatenation of whole months: Training data are indexed by a discrete time identifier (month) so that each observation belongs to a cross-section at a rebalancing date. Let the data set be $\{(x_i, r_i, t_i)\}_{i=1}^N$, where $x_i \in \mathbb{R}^p$ are features, r_i is the realized (future) return, and $t_i \in \mathcal{T}$ is the month label. We construct each mini-batch as a union of complete months rather than sampling individual rows uniformly. Concretely, at the start of each epoch, we sample a permutation of the unique months $\mathcal{T}$ and then form batches from consecutive months in this order. The batch index set is therefore of the form $B = \bigcup_{t \in \mathcal{T}_B} \{i : t_i = t\}$ for some subset of months $\mathcal{T}_B \subset \mathcal{T}$.

This design ensures that portfolio losses requiring within-month ranking and weighting are computed on complete cross-sections. If a mini-batch contained only a random subset of stocks from a month, then the centering, normalization, and ranking operations would be computed on a truncated cross-section, producing weights that do not correspond to any implementable portfolio on that date. The resulting gradients would optimize a distorted objective in which exposures and rank cutoffs fluctuate with the arbitrary inclusion/exclusion of names, rather than reflecting the true cross-sectional trading decision. Economically, batching by complete months aligns training with the execution protocol of a cross-sectional equity strategy. The trading decision is made once per rebalancing date using the information set available at that date; all names in the stock universe are ranked and assigned long/short weights simultaneously, and realized portfolio return is evaluated from the weighted aggregation of next-period returns. The objective requires coherent monthly portfolio returns. Constructing batches from multiple months (number of months controlled by a

hyperparameter) yields a meaningful within-batch estimate of both the mean and dispersion of portfolio returns, stabilizing the Sharpe-based gradient signal relative to computing it from single-month batches. Finally, shuffling months each epoch preserves the benefits of stochastic optimization while keeping the portfolio construction step economically well-posed: randomness comes from sampling time blocks, not from fragmenting cross-sections, which would otherwise introduce artificial noise that does not correspond to any realistic trading constraint.

B Variable definitions

Table B.1: Stock/Firm characteristics

	Abbreviation	Measure	Reference
1.	absacc	Absolute accruals	Bandyopadhyay et al. (2010)
2.	acc	Working capital accruals	Sloan (1996)
3.	aeavol	Abnormal earnings announcement volume	Lerman et al. (2007)
4.	age	Number of years since first Compustat coverage	Jiang et al. (2005)
5.	agr	Asset growth	Cooper et al. (2008)
6.	baspread	Bid-ask spread	Amihud and Mendelson (1989)
7.	beta	Beta	Fama and MacBeth (1973)
8.	bm	Book-to-market	Barr Rosenberg and Lanstein (1985)
9.	bm_ia	Industry-adjusted book-to-market	Asness et al. (2000)
10.	cash	Cash holdings	Palazzo (2012)
11.	cashdebt	Cash flow to debt	Ou and Penman (1989)
12.	cashpr	Cash productivity	Chandrashekar and Rao (2009)
13.	cfp	Cash flow to price ratio	Desai et al. (2004)
14.	cfp_ia	Industry-adjusted cash flow to price ratio	Asness et al. (2000)
15.	chatoia	Industry-adjusted change in asset turnover	Soliman (2008)
16.	chcsho	Change in shares outstanding	Pontiff and Woodgate (2008)
17.	chempia	Industry-adjusted change in employees	Asness et al. (2000)
18.	chinv	Change in inventory	Thomas and Zhang (2002)
19.	chmom	Change in 6-month momentum	Gettleman and Marks (2006)
20.	chpmia	Industry-adjusted change in profit margin	Soliman (2008)
21.	ctx	Change in tax expense	Thomas and Zhang (2011)
22.	cinvest	Corporate investment	Titman et al. (2004)
23.	convind	Convertible debt indicator	Valta (2016)
24.	currat	Current ratio	Ou and Penman (1989)
25.	depr	Depreciation / PP&E	Holthausen and Larcker (1992)
26.	divi	Dividend initiation	Michaely et al. (1995)
27.	divo	Dividend omission	Michaely et al. (1995)
28.	dolvol	Dollar trading volume	Chordia et al. (2001)
29.	dy	Dividend to price	Litzenberger and Ramaswamy (1982)
30.	ear	Earnings announcement return	Kishore et al. (2008)
31.	egr	Growth in common shareholder equity	Richardson et al. (2005)
32.	ep	Earnings to price	Basu (1977)
33.	gma	Gross profitability	Novy-Marx (2013)
34.	grcapx	Growth in capital expenditures	Anderson and Garcia-Feijoo (2006)
35.	grltnoa	Growth in long term net operating assets	Fairfield et al. (2003)
36.	herf	Industry sales concentration	Hou and Robinson (2006)
37.	hire	Employee growth rate	Belo et al. (2014)
38.	idiovol	Idiosyncratic return volatility	Ali et al. (2003)
39.	indmom	Industry momentum	Moskowitz and Grinblatt (1999)
40.	invest	Capital expenditures and inventory	Chen and Zhang (2010)
41.	lev	Leverage	Bhandari (1988)
42.	lgr	Growth in long-term debt	Richardson et al. (2005)
43.	maxret	Maximum daily return	Bali et al. (2011)
44.	mom12m	12-month momentum	Jegadeesh (1990)
45.	mom1m	1-month momentum	Jegadeesh and Titman (1993)
46.	mom36m	36-month momentum	Jegadeesh and Titman (1993)
47.	mom6m	6-month momentum	Jegadeesh and Titman (1993)
48.	ms	Financial statement score	Mohanram (2005)
49.	mvell	Size	Banz (1981)
50.	mve_ia	Industry-adjusted size	Asness et al. (2000)
51.	nincr	Number of earnings increases	Barth et al. (1999)
52.	operprof	Operating profitability	Fama and French (2015)
53.	orgcap	Organizational capital	Eisfeldt and Papanikolaou (2013)

Continued on next page

	Abbreviation	Measure	Reference
54.	pchcapx_ia	Industry-adjusted % change in capital expenditures	Abarbanell and Bushee (1998)
55.	pchcurrat	% change in current ratio	Ou and Penman (1989)
56.	pchdepr	% change in depreciation	Holthausen and Larcker (1992)
57.	pchgm_pchsale	% change in gross margin - % change in sales	Abarbanell and Bushee (1998)
58.	pchquick	% change in quick ratio	Ou and Penman (1989)
59.	pchsale_pchinv	% change in sales - % change in inventory	Abarbanell and Bushee (1998)
60.	pchsale_pchrect	% change in sales - % change in A/R	Abarbanell and Bushee (1998)
61.	pchsale_pchxsga	% change in sales - % change in SG&A	Abarbanell and Bushee (1998)
62.	pctacc	Percent accruals	Hafzalla et al. (2011)
63.	pricedelay	Price delay	Hou and Moskowitz (2005)
64.	ps	Financial statements score	Piotroski (2000)
65.	quick	Quick ratio	Ou and Penman (1989)
66.	rd	R&D increase	Eberhart et al. (2004)
67.	rd_mve	R&D to market capitalization	Guo et al. (2006)
68.	rd_sale	R&D to sales	Guo et al. (2006)
69.	realestate	Real estate holdings	Tuzel (2010)
70.	retvol	Return volatility	Ang et al. (2006)
71.	roaq	Return on assets	Balakrishnan et al. (2010)
72.	roavol	Earnings volatility	Francis et al. (2004)
73.	roeq	Return on equity	Hou et al. (2015)
74.	roic	Return on invested capital	Brown and Rowe (2007)
75.	rsup	Revenue surprise	Kama (2009)
76.	salecash	Sales to cash	Ou and Penman (1989)
77.	saleinv	Sales to inventory	Ou and Penman (1989)
78.	salerec	Sales to receivables	Ou and Penman (1989)
79.	secured	Secured debt	Valta (2016)
80.	securedind	Secured debt indicator	Valta (2016)
81.	sgr	Sales growth	Lakonishok et al. (1994)
82.	sin	Sin stocks	Hong and Kacperczyk (2009)
83.	sp	Sales to price	Barbee Jr et al. (1996)
84.	std_dolvol	Volatility of liquidity (dollar trading volume)	Chordia et al. (2001)
85.	std_turn	Volatility of liquidity (share turnover)	Chordia et al. (2001)
86.	stdacc	Accrual volatility	Bandyopadhyay et al. (2010)
87.	stdcf	Cash flow volatility	Huang (2009)
88.	tang	Debt capacity/firm tangibility	Almeida and Campello (2007)
89.	tb	Tax income to book income	Lev and Nissim (2004)
90.	turn	Share turnover	Datar et al. (1998)
91.	zerotrade	Zero trading days	Liu (2006)

Note: This table presents variable definitions for the set of stock/firm characteristics in our dataset. We provide descriptions of each measure with references to previous literature.

C Characteristic categories

Table C.2: Characteristic categories

Accounting	ms, ps
Accruals	absacc, acc, pctacc, stdacc
Asset composition	cash, realestate
Earnings	aeavol, ear, roavol, stdcf
External financing	chcsho, convind, lgr, secured, securedind, tang
Growth	nincr, rsup, sgr
Investment	agr, chinvt, cinvest, egr, grcapx, grltnoa, hire, invest, pchcapx_ia
Leverage	cashdebt, currat, lev, quick
Leverage change	pchcurrat, pchquick
Liquidity	baspread, dolvol, std_dolvol, std_turn, turn, zerotrade
Momentum/reversal	chmom, indmom, mom12m, mom1m, mom36m, mom6m
Payout indicator	divi, divo
Profitability	cashpr, chpmia, gma, operprof, roaq, roeq, roic, salecash, saleinv, salerec
Profitability change	chatoia, pchgm_pchsale, pchsale_pchinvt, pchsale_pchrect, pchsale_pchxsga
R&D	orgcap, rd, rd_mve, rd_sale
Size	mve_ia, mvel1
Value/growth	bm, bm_ia, cfp, cfp_ia, dy, ep, sp
Volatility	idiovol, maxret, retvol

Note: This table shows the characteristic categories.

D Sharpe-maximizing benchmark

As an additional benchmark, we also form a portfolio that combines the characteristic-sorted long-short portfolios. For each firm characteristic j , we construct a monthly high-minus-low (HML) return series $h_{j,t}$ by sorting stocks into ten portfolios based on the cross-section of that characteristic in month t and taking the return difference between the highest and lowest portfolio. This yields a set of tradable characteristic-based portfolios. In contrast to the univariate benchmark, which evaluates each characteristic-sorted portfolio separately, this specification allows the benchmark to exploit diversification across characteristics.

To keep the benchmark design aligned with the training protocol used throughout the paper, all portfolio-combination decisions are made using only the combined training and validation sample. Specifically, for each characteristic portfolio j , we first compute its in-sample Sharpe ratio and rank portfolios accordingly. For a given value of k , we retain the top k portfolios and estimate the combination weights using only their in-sample return moments.

Let $\mathbf{h}_t = (h_{1,t}, \dots, h_{k,t})'$ denote the vector of retained portfolio returns in month t , and let $\hat{\mu}$ and $\hat{\Sigma}$ denote the sample mean vector and covariance matrix of $\mathbf{h}_t$ estimated on the combined training and validation sample. We then form the Sharpe-optimal combination by choosing weights in the tangency-portfolio direction,

$$\tilde{\mathbf{w}} = \hat{\Sigma}^{-1} \hat{\mu}.$$

Because the Sharpe ratio is invariant to scale, we normalize these weights by their ℓ_1 norm,

$$\mathbf{w} = \frac{\tilde{\mathbf{w}}}{\sum_{j=1}^k |\tilde{w}_j|},$$

so that the benchmark has unit gross exposure across characteristic portfolios. This is an unconstrained combination, so individual portfolio weights may be positive or negative.¹⁸

The return on the combined benchmark portfolio is then given by

$$r_t^{\text{bench}} = \mathbf{w}' \mathbf{h}_t = \sum_{j=1}^k w_j h_{j,t}.$$

After estimating $\mathbf{w}$ on the combined training and validation samples, we hold these weights fixed and evaluate the resulting benchmark out-of-sample on the test period. For each choice of k , we report the monthly mean return, standard deviation, and Sharpe ratio, both in-sample and out-of-sample. This procedure provides a mean-variance benchmark that is directly comparable to AlphaGlass, while still restricting all parameter selection to information available prior to the test sample.

¹⁸Note that we do not impose the additional constraint $\sum_{j=1}^k w_j = 0$ when combining characteristic portfolios. Adding such a restriction would generally change the tangency-portfolio solution and prevent the benchmark from attaining the unconstrained Sharpe-optimal combination implied by $\hat{\Sigma}^{-1} \hat{\mu}$.

Table D.1: Statistics of Sharpe-optimal characteristic-combination portfolios

	3	5	10	45	91
Panel A: In-sample					
Mean	0.007	0.008	0.009	0.005	0.004
S.D.	0.021	0.018	0.016	0.006	0.004
SR	0.341	0.460	0.520	0.800	1.089
Panel B: Out-of-sample					
Mean	0.001	0.002	0.000	0.000	0.000
S.D.	0.013	0.013	0.011	0.006	0.004
SR	0.060	0.123	0.002	0.032	0.062

Note: This table compares monthly means, standard deviations, and Sharpe ratios for the Sharpe-optimal characteristics-based benchmark in the training/validation (Panel A) and testing sample (Panel B). Results are reported for different choices of k .

E Alternative weights

Our main training objective focuses on the cross-sectional tails by forming differentiable long and short legs using soft top- and bottom-cutoffs. While this tail-focused construction directly targets extreme-ranked assets and resembles standard quantile long-short portfolio formation, it is a natural idea to consider a fully cross-sectional objective that uses all assets each month. We formulate an alternative objective that enforces dollar-neutrality and controls portfolio leverage via ℓ_1 normalization.

Let $t \in \{1, \dots, T\}$ index months and let l index the set of tradable stocks in month t with $l = 1, \dots, n_t$. Given model scores (signals) $\{s_{l,t}\}$ and realized returns $\{r_{l,t}\}$, we construct a dollar-neutral long-short portfolio using all stocks in the cross section.

First, we demean the scores within each month:

$$c_{l,t} = s_{l,t} - \bar{s}_t, \quad \bar{s}_t = \frac{1}{n_t} \sum_{l=1}^{n_t} s_{l,t}. \quad (\text{E.1})$$

This centering step implies $\sum_{l=1}^{n_t} c_{l,t} = 0$.

Next, we normalize by the cross-sectional ℓ_1 norm to control gross exposure and obtain portfolio weights:

$$w_{l,t} = \frac{c_{l,t}}{\sum_{j=1}^{n_t} |c_{j,t}|}. \quad (\text{E.2})$$

By construction,

$$\sum_{l=1}^{n_t} w_{l,t} = 0 \quad \text{and} \quad \sum_{l=1}^{n_t} |w_{l,t}| = 1, \quad (\text{E.3})$$

so that the portfolio is (approximately) dollar-neutral each month with approximately unit gross exposure.

The portfolio return in month t is then

$$\pi_t = \sum_{l=1}^{n_t} w_{l,t} r_{l,t}. \quad (\text{E.4})$$

Finally, we maximize the time-series Sharpe ratio of $\{\pi_t\}_{t=1}^T$:

$$\text{Sharpe}(\pi) = \frac{\mu_\pi}{\sigma_\pi}, \quad \mu_\pi = \frac{1}{T} \sum_{t=1}^T \pi_t, \quad \sigma_\pi = \sqrt{\frac{1}{T} \sum_{t=1}^T (\pi_t - \mu_\pi)^2}, \quad (\text{E.5})$$

and define the training loss as the negative Sharpe ratio,

$$\mathcal{L} = -\text{Sharpe}(\pi). \quad (\text{E.6})$$

Table E.1: Statistics of AlphaGlass portfolios with alternative objective

	1	2	3	4	5	6	7	8	9	10	10-1
Panel A: In-sample											
Mean	-0.076	0.301	0.427	0.481	0.432	0.574	0.804	1.161	1.322	1.637	1.713
S.D.	7.188	6.523	6.335	5.320	5.749	6.043	6.666	6.719	6.969	6.945	1.407
SR	-0.011	0.046	0.067	0.090	0.075	0.095	0.121	0.173	0.190	0.236	1.218
Panel B: Out-of-sample											
Mean	0.989	1.116	0.939	0.874	0.897	0.841	0.916	1.340	1.246	1.544	0.555
S.D.	5.323	5.192	4.989	4.260	4.427	4.678	5.143	5.460	5.672	5.839	1.578
SR	0.186	0.215	0.188	0.205	0.203	0.180	0.178	0.245	0.220	0.264	0.352

Note: This table compares monthly means, standard deviations, and Sharpe ratios for ten equal-weighted portfolios constructed based on predicted AlphaGlass signal deciles in the training/validation (Panel A) and testing sample (Panel B). The last columns show the results of the 10-1 portfolio, which is long in portfolio 10 and short in portfolio 1.

The resulting portfolio statistics are reported in Table E.1. Relative to the tail-focused training objective, the fully cross-sectional objective yields a somewhat lower out-of-sample Sharpe ratio, with the deterioration being particularly pronounced within the short leg. A plausible explanation is that the alternative objective distributes portfolio weight and gradient signal across the entire cross-section, rather than concentrating learning on the most extreme low- and high-score assets. As a result, the model is encouraged to improve average cross-sectional ranking quality, but places less emphasis on identifying the subset of stocks with the weakest subsequent returns. This is especially relevant for the short leg, where performance depends critically on correctly isolating the most negative-return candidates rather than merely assigning moderately lower scores to below-average stocks. In contrast, the tail-focused objective is more closely aligned with the downstream portfolio formation rule, since it explicitly rewards separation in the tails. Consequently, it can produce stronger discrimination among the lowest-ranked assets and, in turn, a more profitable short portfolio.

F Excluding small stocks

Every month, we exclude stocks in the bottom size quintile and estimate all models on the remaining sample. The main findings continue to hold in this restricted universe. AlphaGlass still produces a clear spread across signal-sorted portfolios, and the out-of-sample 10-1 portfolio maintains a high Sharpe ratio, only modestly below the baseline specification. This indicates that the model's performance is not primarily a micro-cap phenomenon. Instead, the learned signal appears to capture a broader cross-sectional structure that remains economically relevant even after excluding the smallest firms.

Table F.1: Excluding small stocks – Out-of-sample

	1	2	3	4	5	6	7	8	9	10	10-1
Panel A: AlphaGlass											
Mean	0.46	0.72	0.92	0.97	0.99	1.11	1.21	1.12	1.25	1.44	0.98
S.D.	5.55	4.15	3.70	4.52	4.90	5.03	5.26	5.36	5.62	5.86	2.29
SR	0.08	0.17	0.25	0.22	0.20	0.22	0.23	0.21	0.22	0.25	0.43
Panel C: RF											
Mean	0.97	0.87	0.81	0.95	0.95	1.00	1.15	1.39	1.11	1.01	0.04
S.D.	5.31	4.77	4.81	4.15	4.19	4.08	5.07	5.49	5.75	7.07	3.11
SR	0.18	0.18	0.17	0.23	0.23	0.25	0.23	0.25	0.19	0.14	0.01
Panel B: NN											
Mean	0.68	0.83	0.97	0.93	1.02	1.22	1.17	1.21	1.14	1.03	0.35
S.D.	7.25	6.07	3.23	3.58	4.19	4.65	4.90	5.22	5.45	5.68	2.86
SR	0.09	0.14	0.30	0.26	0.24	0.26	0.24	0.23	0.21	0.18	0.12
Panel D: EBM											
Mean	0.71	0.83	1.04	1.01	1.09	1.15	1.08	1.20	0.94	1.14	0.43
S.D.	6.42	5.39	5.01	4.77	4.75	4.63	4.40	4.25	4.60	6.29	3.43
SR	0.11	0.15	0.21	0.21	0.23	0.25	0.25	0.28	0.21	0.18	0.13

Note: This table compares monthly out-of-sample means, standard deviations, and Sharpe ratios for ten equal-weighted portfolios constructed based on AlphaGlass, Random Forest (RF), neural network (NN), and explainable boosting machine (EBM) prediction deciles when stocks in the bottom size quintile in each month are excluded. The last columns show the results of the 10-1 portfolio, which is long in portfolio 10 and short in portfolio 1.

G Alternative data split

For comparison, we re-estimate our models on a longer training period of 15 years (1/2000-12/2014). We keep the length of the validation window constant at two years (1/2015-12/2016) and, accordingly, use five years for testing (1/2017-12/2021). The resulting portfolio statistics are presented in Table G.1. We find that performance is comparable to the main model. The decile spread remains pronounced and the out-of-sample 10-1 portfolio continues to deliver a Sharpe ratio close to the in-sample value, even slightly exceeding the main result. This suggests that the AlphaGlass mapping is reasonably stable across alternative train/test splits, rather than an artifact of the baseline sample partition.

Table G.1: AlphaGlass portfolios with long training window

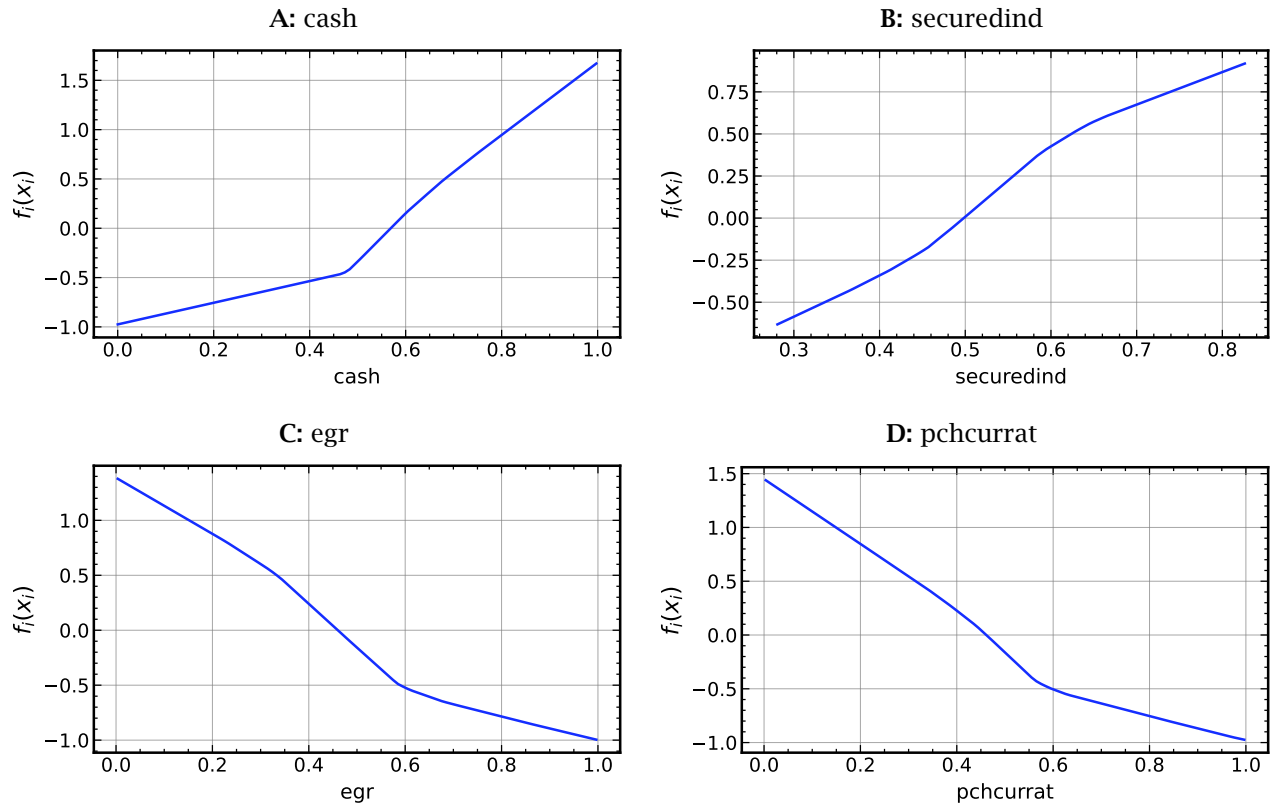
	1	2	3	4	5	6	7	8	9	10	10-1
Panel A: In-sample											
Mean	−0.27	0.44	0.60	0.62	0.70	1.00	1.13	1.12	1.23	1.51	1.78
S.D.	7.02	6.23	4.71	4.50	5.55	6.20	5.98	6.10	6.16	6.00	3.03
SR	−0.04	0.07	0.13	0.14	0.13	0.16	0.19	0.18	0.20	0.25	0.59
Panel B: Out-of-sample											
Mean	0.48	0.66	0.82	0.80	0.74	1.14	1.48	1.49	1.59	1.72	1.24
S.D.	6.75	7.39	5.97	4.85	5.06	5.93	6.51	6.71	6.84	6.44	2.55
SR	0.07	0.09	0.14	0.17	0.15	0.19	0.23	0.22	0.23	0.27	0.49

Note: This table compares monthly means, standard deviations, and Sharpe ratios for ten equal-weighted portfolios constructed based on predicted AlphaGlass signal deciles in the longer training/validation (Panel A) and testing sample (Panel B). The last columns show the results of the 10-1 portfolio, which is long in portfolio 10 and short in portfolio 1.

H Additional shape functions

In Section 4.4, we discuss the shape functions associated with the most important predictors in terms of mean absolute scores. Figure H.1 presents additional shape functions related to subsequent characteristics.

Figure H.1: Univariate functions $f_i(x_i)$



Note: This figure shows the estimated $f_i(x_i)$ functions of additional characteristics. The x -axis plots the input variable x_i , cross-sectionally ranked and normalized to the [0,1] interval. The y -axis represents the additive contribution to the prediction output through the shape function $f_i(x_i)$.

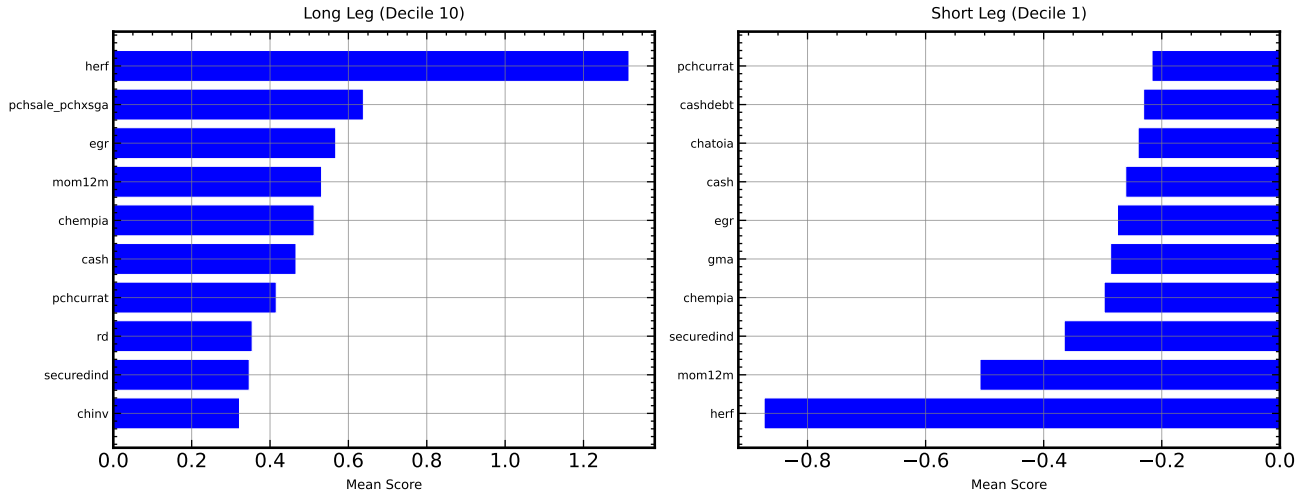
I Decile-wise decomposition

To complement the directional decomposition in Section 5.1, we also report the net mean contribution of each model term within the extreme score deciles. For any decile q , define

$$\bar{S}_i(q) = \frac{1}{|\mathcal{P}_q|} \sum_{(l,t) \in \mathcal{P}_q} f_i(x_{i,l,t}), \quad \bar{S}_{ij}(q) = \frac{1}{|\mathcal{P}_q|} \sum_{(l,t) \in \mathcal{P}_q} f_{ij}(x_{i,l,t}, x_{j,l,t}).$$

By construction, $\bar{S}_i(q) = S_i^+(q) + S_i^-(q)$ and likewise for interactions, so this measure captures the average *net* signed contribution after offsetting positive and negative realizations within the same decile. These results therefore represent a decomposition of the average long- and short-leg signals, rather than a ranking of one-sided directional pushes (as in Figure 9). Figure I.1 decomposes the mean AlphaGlass score (including both negative and positive contributions) separately for the extreme decile portfolios. Specifically, it reports the ten largest average score contributions within the top decile (long leg) and bottom decile (short leg), allowing us to compare which model terms drive each side of the strategy. A central finding is that several of the same predictors, including *herf*, *mom12m*, *securedind*, *chempia*, *egr*, *cash*, and *pchcurrat*, appear in both tails with opposite signs, indicating that these characteristics are the main forces separating the long and short portfolios. At the same time, the composition is not perfectly symmetric: the long leg is more concentrated, with especially large positive contributions from *herf* and *pchsale_pchxsga*, while the short leg is more diffuse and includes some predictors that are specific to that tail, such as *gma*, *chatoia*, and *cashdebt*.

Figure I.1: Mean scores $\bar{S}_i(q)$ in top and bottom decile

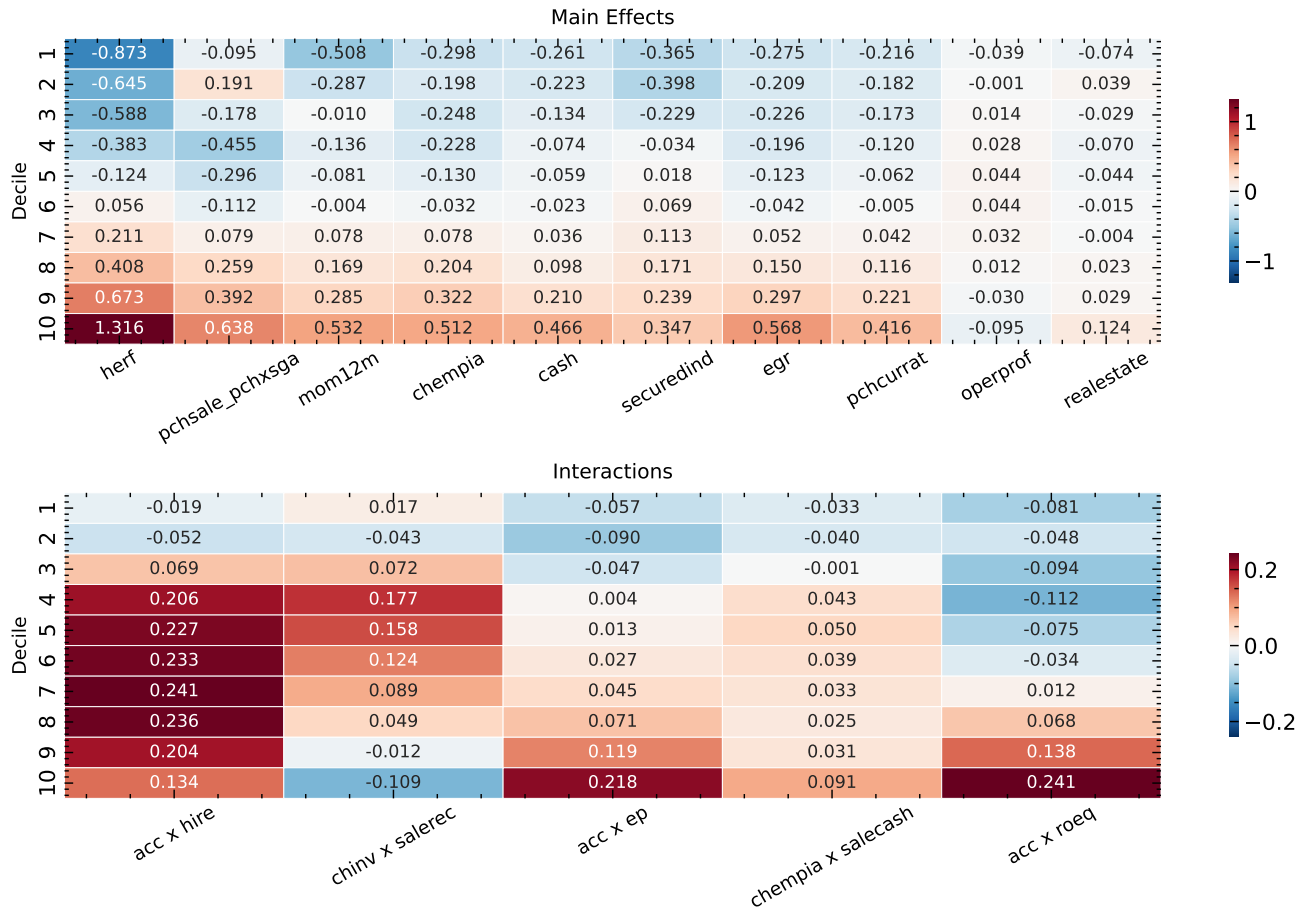


Note: This figure shows mean scores associated with the respective top 10 predictors in the long (portfolio 10) and the short leg (portfolio 1).

Figure I.2 presents the mean score contributions of the most important univariate effects and interactions across deciles. It shows how the contributions of the most important model terms vary across the AlphaGlass signal deciles. The upper heatmap reports the average signed score contributions of the ten most important univariate effects within each decile, while the lower heatmap reports the corresponding averages for the five most important interaction terms. Each cell, therefore, indicates whether a given term tends, on average, to increase or decrease the portfolio signal for

stocks in that part of the score distribution. A clear pattern is that many of the leading main effects become more positive as one moves from low to high deciles, showing that the same characteristics help separate the short and long tails through opposite-signed contributions across the distribution. The interaction effects are smaller in magnitude, but still display systematic cross-decile variation, suggesting that they refine the ranking induced by the dominant univariate signals.

Figure I.2: Decomposition heatmaps



Note: This figure shows mean score contributions of the top main effects and interactions for each decile portfolio.

J Drawdowns

A useful feature of AlphaGlass is that the model is trained directly on a portfolio objective, enabling it to account for investor preferences beyond the Sharpe ratio. To illustrate this flexibility, we consider an alternative objective that augments the baseline criterion with a penalty for drawdowns. More specifically, we implement a penalty term that penalizes large maximum drawdown (MDD) in a soft (differentiable) way. Let $\{f_t^{\text{HL}}\}_{t=1}^T$ denote the sequence of returns of the long/short factor produced by the AlphaGlass portfolio. We define the cumulative wealth process as

$$W_t = \prod_{s=1}^t (1 + f_s^{\text{HL}}), \quad W_0 = 1.$$

Drawdowns measure how far the strategy is below its historical peak. We define the running maximum of wealth as

$$W_t^{\max} = \max_{0 \leq s \leq t} W_s,$$

and define the (fractional) drawdown at time t as

$$D_t = 1 - \frac{W_t}{W_t^{\max}} \in [0, 1].$$

We penalize the strategy for spending time in drawdowns and for experiencing deep drawdowns by adding the time-average of squared drawdowns,

$$\mathcal{L}_{\text{DD}} = \frac{1}{T} \sum_{t=1}^T D_t^2,$$

where drawdowns are squared to overweight large peak-to-trough declines, aligned with the maximum drawdown metric. The overall training objective augments the standard risk-adjusted return criterion with this drawdown term:

$$\mathcal{L} = -(\text{SR}(f_t^{\text{HL}}) - \lambda_{\text{DD}} \mathcal{L}_{\text{DD}}), \quad (\text{J.1})$$

with $\lambda_{\text{DD}} \geq 0$ trading off average risk-adjusted performance against the economic preference for smoother equity curves and reduced drawdown severity. Intuitively, this encourages strategies that not only perform well on average but also avoid prolonged or deep underwater episodes, which are costly for leveraged investors and capacity-constrained portfolios.

Table J.1 shows the in-sample and out-of-sample Sharpe ratios, drawdown penalties $\lambda_{\text{DD}} \mathcal{L}_{\text{DD}}$, and the negative of the minimized loss $-\mathcal{L}$. The AlphaGlass model is estimated with the loss function (J.1). In contrast, the results for the three benchmark models are based on estimates obtained using the standard mean-square error loss function, and their loss values are computed using the estimated long/short portfolios. AlphaGlass’s out-of-sample Sharpe ratio is 0.3 with a drawdown penalty close to zero, so the total negative loss is slightly lower than the Sharpe ratio (0.299). Note that the neural net’s Sharpe ratio is higher, 0.34, but the drawdown penalty is orders of magnitude larger, resulting in a smaller negative loss. The EBM high/low portfolio has a Sharpe ratio nearly identical to AlphaGlass, but its drawdown penalty is as high as that of the NN model. The random forest yields the poorest results, as in the baseline case without drawdown penalties.

Table J.1: Drawdowns

	In-sample			Out-of-sample		
	SR	$\lambda_{DD}\mathcal{L}_{DD}$	$-\mathcal{L}$	SR	$\lambda_{DD}\mathcal{L}_{DD}$	$-\mathcal{L}$
AlphaGlass	1.087	0.000	1.087	0.302	0.003	0.299
RF	0.874	0.071	0.803	0.143	3.572	-3.429
NN	0.671	0.008	0.663	0.338	0.171	0.167
EBM	0.777	0.090	0.686	0.299	0.174	0.125

Note: This table shows the in-sample and out-of-sample Sharpe ratios, drawdown penalties $\lambda_{DD}\mathcal{L}_{DD}$, and the negative of the minimized loss $-\mathcal{L}$. The drawdown parameter $\lambda_{DD} = 100$.

K Theoretical results

K.1 Proofs of theorems in Section 7

In the following, we provide complete formal proofs of the theorems stated in Section 7. While the formal statements below are written for the SR objective, the underlying proof strategies apply to broader classes of objectives. For Theorem 1, the core argument is objective-free: it establishes convergence of the soft rank-and-mask weights, which implies convergence of time-averaged return differences and, in particular, convergence of empirical moments such as $\hat{\mu}_T(\theta)$ and $\hat{\nu}_T(\theta)$ or, equivalently, $\hat{\sigma}_T^2(\theta) = \hat{\nu}_T(\theta) - \hat{\mu}_T(\theta)^2$. Accordingly, the same convergence argument also applies to empirical criteria that depend continuously on these moments. Similarly, Theorems 2 and 3 proceed by first establishing uniform control of the empirical portfolio moments and then passing these moment bounds to the objective of interest. Therefore, the same logic extends to objectives of the form

$$Q(\theta) = \phi(\mu(\theta), \nu(\theta)), \quad \text{or equivalently} \quad Q(\theta) = \tilde{\phi}(\mu(\theta), \sigma^2(\theta)),$$

provided ϕ is sufficiently regular. For the consistency step in Theorem 2, continuity of ϕ on the relevant compact domain is sufficient. For a high-probability deviation step in Theorem 3, one additionally needs local Lipschitz continuity, so that deviations of $(\hat{\mu}_T, \hat{\nu}_T)$ transfer to deviations of $\hat{Q}_T$. A sufficient condition for the latter is that ϕ be continuously differentiable on an open neighborhood of the domain with uniformly bounded gradient there.¹⁹

Fix months $t = 1, \dots, T$. In month t , there are $n_t \geq 2$ traded assets indexed by l . Let $X_{l,t} \in [0, 1]^p$ be the vector of (rank-normalized) characteristics and $r_{l,t} \in \mathbb{R}$ the realized return from t to $t + 1$. The AlphaGlass signal is

$$s_{l,t}^\theta = \beta_0 + \sum_{i=1}^p f_i(x_{i,l,t}) + \sum_{i>j} f_{ij}(x_{i,l,t}, x_{j,l,t}),$$

with θ denoting the set of all model parameters. We construct differentiable weights via the pairwise-soft-rank and soft-mask mapping described in Section 2, with temperatures $\tau_{\text{rank}}, \tau_{\text{mask}} > 0$. The long-short portfolio return in month t is

$$\pi_t(\theta) = \sum_{l=1}^{n_t} w_{l,t}^\theta r_{l,t}, \quad \mu(\theta) = \mathbb{E}[\pi_t(\theta)], \quad \sigma^2(\theta) = \text{Var}[\pi_t(\theta)], \quad \text{SR}(\theta) = \mu(\theta) / \sigma(\theta).$$

The sample analogs replace expectations with means over $t = 1, \dots, T$.

Throughout, we make the following assumptions:

- Assumption A1 (Stationarity and weak dependence). Let $Z_t = \{(X_{l,t}, r_{l,t})_{l=1}^{n_t}\}$ be strictly stationary and β -mixing with coefficients $\{\beta(k)\}_{k \geq 1}$ satisfying $\sum_{k=1}^{\infty} \beta(k)^{\frac{1}{2}} < \infty$.
- Assumption A2 (Boundedness/sub-Gaussianity and portfolio feasibility). There exists $R < \infty$ such that $|r_{l,t}| \leq R$ almost surely (a.s.) (or $r_{l,t}$ are sub-Gaussian with a common proxy). Cross-sectional sizes obey $2 \leq n_{\min} \leq n_t \leq n_{\max} < \infty$. The soft long and short legs are normalized to have unit mass each, so $\sum_l |w_{l,t}^\theta| \leq 2$.

¹⁹For ratio-type objectives, one also needs the denominator to be bounded away from zero on the domain under consideration.

- Assumption A3 (Model class regularity). The parameter space Θ is compact. Each f_i, f_{ij} is uniformly bounded and L -Lipschitz in its inputs. Moreover, the map $\theta \mapsto s_{l,t}^\theta$ is Lipschitz uniformly in (l, t) . In our implementation, we adopt L1 regularization on the gates to encourage sparsity. While we do not explicitly project weights onto a bounded set, with fixed depth/width, normalized inputs, and early stopping, parameter norms remain moderate in practice. Assumption A3 should therefore be viewed as a standard theoretical regularity condition ensuring boundedness and Lipschitz continuity of the signal network.
- Assumption A4 (Positive volatility in class). $\inf_{\theta \in \Theta} \sigma(\theta) \geq \underline{\sigma} > 0$. For Theorem 1, we additionally assume that the hard top-bottom- q portfolios satisfy $\inf_{\theta \in \Theta} \sigma^{\text{hard}}(\theta) \geq \underline{\sigma}_0 > 0$ and that, for all sufficiently small τ , the corresponding soft portfolios satisfy $\inf_{\theta \in \Theta} \sigma^{\text{soft}}(\theta, \tau) \geq \underline{\sigma}_0/2$.

Auxiliary penalties (e.g. ℓ_1 gates) are bounded continuous functionals of $(X_{l,t}, r_{l,t})$ and the weights. Adding them to the loss does not affect any of the arguments below.

Proof of Theorem 1

We first show that the differentiable soft-sorting portfolio construction is a faithful surrogate for the usual hard top-bottom sort. Fix $q \in (0, 0.5)$ and let $k_t = \lfloor qn_t \rfloor \geq 1$ for each t . Define the hard top/bottom- q portfolio that equal-weights the top k_t and bottom k_t assets by the true (non-smoothed) signal ranks and denote these weights by $w_{l,t}^0(\theta)$. For the soft portfolio, let $w_{l,t}^\tau(\theta)$ be the weights built from pairwise-sigmoid soft ranks and soft masks with temperatures $\tau = (\tau_{\text{rank}}, \tau_{\text{mask}}) > 0$, as described above. For each t , define the signal margin

$$\Delta_\theta^{(t)} = \min_{i \neq j} |s_{i,t}^\theta - s_{j,t}^\theta|,$$

i.e., the smallest pairwise gap between signals in month t . Assume that, for fixed θ , the random variables $\Delta_\theta^{(t)}$ have no point mass at 0 (equivalently, their distribution function is continuous at 0).

Proof. Throughout, $\|\cdot\|_1$ is the ℓ_1 norm over the cross-sectional index l . By A2, $2 \leq n_t \leq n_{\max}$. We start by establishing a bound for the absolute difference between the sigmoid function and a hard indicator: For the logistic function $\sigma(x) = 1/(1 + e^{-x})$,

$$|\sigma(x) - 1\{x > 0\}| = \begin{cases} \frac{1}{1+e^x} \leq e^{-x}, & x \geq 0, \\ \frac{e^x}{1+e^x} \leq e^x, & x < 0, \end{cases} \leq e^{-|x|}. \quad (\text{K.1})$$

With $x = (s_{l',t}^\theta - s_{l,t}^\theta)/\tau_{\text{rank}}$, the pairwise-soft comparison $S_{l',t}^\tau = \sigma((s_{l',t}^\theta - s_{l,t}^\theta)/\tau_{\text{rank}})$ satisfies

$$|S_{l',t}^\tau - 1\{s_{l',t}^\theta > s_{l,t}^\theta\}| \leq \exp(-|s_{l',t}^\theta - s_{l,t}^\theta|/\tau_{\text{rank}}). \quad (\text{K.2})$$

We now use these results to bound the rank error. Let

$$\text{rank}_{l,t}^0 = 1 + \sum_{j \neq l} 1\{s_{j,t}^\theta > s_{l,t}^\theta\}, \quad \text{rank}_{l,t}^\tau = 1 + \sum_{j \neq l} S_{l,j,t}^\tau.$$

If $\Delta_\theta^{(t)} \geq \delta$, then by definition of the signal margin

$$|s_{j,t}^\theta - s_{l,t}^\theta| \geq \delta \quad \text{for all } j \neq l.$$

Therefore, by (K.2),

$$|S_{l,j,t}^\tau - 1\{s_{j,t}^\theta > s_{l,t}^\theta\}| \leq \exp(-|s_{j,t}^\theta - s_{l,t}^\theta|/\tau_{\text{rank}}) \leq e^{-\delta/\tau_{\text{rank}}} \quad \text{for all } j \neq l.$$

Summing over $j \neq l$ yields

$$|\text{rank}_{l,t}^\tau - \text{rank}_{l,t}^0| = \left| \sum_{j \neq l} (S_{l,j,t}^\tau - 1\{s_{j,t}^\theta > s_{l,t}^\theta\}) \right| \leq \sum_{j \neq l} |S_{l,j,t}^\tau - 1\{s_{j,t}^\theta > s_{l,t}^\theta\}| \leq (n_t - 1)e^{-\delta/\tau_{\text{rank}}}.$$

Using $n_t \leq n_{\max}$, we obtain the uniform bound

$$|\text{rank}_{l,t}^\tau - \text{rank}_{l,t}^0| \leq C_{\text{rank}} e^{-\delta/\tau_{\text{rank}}}, \quad \text{with } C_{\text{rank}} = n_{\max}, \quad (\text{K.3})$$

for all l and t .

This translates into a bound for the mask error: Define the hard mask $m_{l,t}^{0,\text{top}} = 1\{\text{rank}_{l,t}^0 \leq k_t\}$ and the soft mask $m_{l,t}^{\tau,\text{top}} = \sigma((k_t + 0.5 - \text{rank}_{l,t}^\tau)/\tau_{\text{mask}})$. Let $a_{l,t} = k_t + 0.5 - \text{rank}_{l,t}^0$ and $\eta_{l,t} = \text{rank}_{l,t}^\tau - \text{rank}_{l,t}^0$. Note that, since k_t and $\text{rank}_{l,t}^0$ are integers, $a_{l,t} \in \{\pm 0.5, \pm 1.5, \dots\}$, in particular $|a_{l,t}| \geq 0.5$. Write

$$m_{l,t}^{\tau,\text{top}} - m_{l,t}^{0,\text{top}} = \left(\sigma((a_{l,t} - \eta_{l,t})/\tau_{\text{mask}}) - \sigma(a_{l,t}/\tau_{\text{mask}}) \right) + \left(\sigma(a_{l,t}/\tau_{\text{mask}}) - 1\{a_{l,t} > 0\} \right).$$

The second term is bounded by $e^{-|a_{l,t}|/\tau_{\text{mask}}} \leq e^{-0.5/\tau_{\text{mask}}}$ by (K.1). For the first term, we use that σ is 1/4-Lipschitz: $|\sigma(x) - \sigma(y)| \leq \frac{1}{4}|x - y|$. With $x = (a_{l,t} - \eta_{l,t})/\tau_{\text{mask}}$ and $y = a_{l,t}/\tau_{\text{mask}}$,

$$|\sigma((a_{l,t} - \eta_{l,t})/\tau_{\text{mask}}) - \sigma(a_{l,t}/\tau_{\text{mask}})| \leq \frac{1}{4\tau_{\text{mask}}} |\eta_{l,t}|.$$

Combining this result with (K.3) yields

$$|m_{l,t}^{\tau,\text{top}} - m_{l,t}^{0,\text{top}}| \leq e^{-0.5/\tau_{\text{mask}}} + \frac{C_{\text{rank}}}{4\tau_{\text{mask}}} e^{-\delta/\tau_{\text{rank}}}. \quad (\text{K.4})$$

An identical bound holds for the bottom mask $m_{l,t}^{\tau,\text{bot}}$ relative to $m_{l,t}^{0,\text{bot}} = 1\{\text{rank}_{l,t}^0 \geq n_t - k_t + 1\}$. Based on this, we can bound the normalization step: Let $Z_t^{\tau,\text{top}} = \sum_j m_{j,t}^{\tau,\text{top}}$ and $Z_t^{0,\text{top}} = \sum_j m_{j,t}^{0,\text{top}} = k_t$. Summing (K.4) over j and using $n_t \leq n_{\max}$ gives

$$|Z_t^{\tau,\text{top}} - k_t| \leq n_{\max} e^{-0.5/\tau_{\text{mask}}} + \frac{n_{\max} C_{\text{rank}}}{4\tau_{\text{mask}}} e^{-\delta/\tau_{\text{rank}}} =: \varepsilon(\tau),$$

where $\varepsilon(\tau) \downarrow 0$ as $\tau \downarrow 0$. Thus, for all t ,

$$Z_t^{\tau,\text{top}} \in [k_t - \varepsilon(\tau), k_t + \varepsilon(\tau)]. \quad (\text{K.5})$$

The same bound holds for $Z_t^{\tau,\text{bot}}$. Since $k_t \geq k_{\min} = \lfloor qn_{\min} \rfloor \geq 1$ and $\varepsilon(\tau) \rightarrow 0$, we can choose τ small enough so that $\varepsilon(\tau) \leq k_{\min}/2$, and hence $Z_t^{\tau,\text{top}}, Z_t^{\tau,\text{bot}} \geq k_t/2$ for all t .

We can now formulate the ℓ_1 error bound for the portfolio weights: The top-leg weights are $w_{l,t}^{\tau,\text{top}} = m_{l,t}^{\tau,\text{top}}/Z_t^{\tau,\text{top}}$ and $w_{l,t}^{0,\text{top}} = m_{l,t}^{0,\text{top}}/k_t$. We examine the difference $w_{l,t}^{\tau,\text{top}} - w_{l,t}^{0,\text{top}}$. Adding and subtracting $m_{l,t}^{0,\text{top}}/Z_t^{\tau,\text{top}}$ yields:

$$w_{l,t}^{\tau,\text{top}} - w_{l,t}^{0,\text{top}} = \frac{m_{l,t}^{\tau,\text{top}}}{Z_t^{\tau,\text{top}}} - \frac{m_{l,t}^{0,\text{top}}}{k_t} = \left(\frac{m_{l,t}^{\tau,\text{top}}}{Z_t^{\tau,\text{top}}} - \frac{m_{l,t}^{0,\text{top}}}{Z_t^{\tau,\text{top}}} \right) + \left(\frac{m_{l,t}^{0,\text{top}}}{Z_t^{\tau,\text{top}}} - \frac{m_{l,t}^{0,\text{top}}}{k_t} \right),$$

which simplifies to

$$w_{l,t}^{\tau,\text{top}} - w_{l,t}^{0,\text{top}} = \frac{m_{l,t}^{\tau,\text{top}} - m_{l,t}^{0,\text{top}}}{Z_t^{\tau,\text{top}}} + m_{l,t}^{0,\text{top}} \left(\frac{1}{Z_t^{\tau,\text{top}}} - \frac{1}{k_t} \right).$$

Taking absolute values and using the triangle inequality,

$$|w_{l,t}^{\tau,\text{top}} - w_{l,t}^{0,\text{top}}| \leq \frac{|m_{l,t}^{\tau,\text{top}} - m_{l,t}^{0,\text{top}}|}{Z_t^{\tau,\text{top}}} + m_{l,t}^{0,\text{top}} \left| \frac{1}{Z_t^{\tau,\text{top}}} - \frac{1}{k_t} \right|.$$

We now sum over l . Recall that $m_{l,t}^{0,\text{top}} \in \{0,1\}$ and $\sum_l m_{l,t}^{0,\text{top}} = k_t$. Hence

$$\begin{aligned} \sum_l |w_{l,t}^{\tau,\text{top}} - w_{l,t}^{0,\text{top}}| &\leq \frac{1}{Z_t^{\tau,\text{top}}} \sum_l |m_{l,t}^{\tau,\text{top}} - m_{l,t}^{0,\text{top}}| + \left| \frac{1}{Z_t^{\tau,\text{top}}} - \frac{1}{k_t} \right| \sum_l m_{l,t}^{0,\text{top}} \\ &= \frac{1}{Z_t^{\tau,\text{top}}} \sum_l |m_{l,t}^{\tau,\text{top}} - m_{l,t}^{0,\text{top}}| + k_t \left| \frac{1}{Z_t^{\tau,\text{top}}} - \frac{1}{k_t} \right|. \end{aligned}$$

On the event where $Z_t^{\tau,\text{top}} \geq k_t/2$ (see (K.5)), we can bound the difference of reciprocals as

$$\left| \frac{1}{Z_t^{\tau,\text{top}}} - \frac{1}{k_t} \right| = \frac{|Z_t^{\tau,\text{top}} - k_t|}{Z_t^{\tau,\text{top}} k_t} \leq \frac{2}{k_t^2} |Z_t^{\tau,\text{top}} - k_t|.$$

Using again $Z_t^{\tau,\text{top}} \geq k_t/2$, we also have

$$\frac{1}{Z_t^{\tau,\text{top}}} \leq \frac{2}{k_t}.$$

Combining these results, the sum over l can be bounded in the following way:

$$\sum_l |w_{l,t}^{\tau,\text{top}} - w_{l,t}^{0,\text{top}}| \leq \frac{2}{k_t} \sum_l |m_{l,t}^{\tau,\text{top}} - m_{l,t}^{0,\text{top}}| + \frac{2}{k_t} |Z_t^{\tau,\text{top}} - k_t|. \quad (\text{K.6})$$

Recall that, by (K.4) and $n_t \leq n_{\max}$,

$$\sum_l |m_{l,t}^{\tau,\text{top}} - m_{l,t}^{0,\text{top}}| \leq n_{\max} e^{-0.5/\tau_{\text{mask}}} + \frac{n_{\max} C_{\text{rank}}}{4\tau_{\text{mask}}} e^{-\delta/\tau_{\text{rank}}},$$

and accordingly

$$|Z_t^{\tau,\text{top}} - k_t| \leq n_{\max} e^{-0.5/\tau_{\text{mask}}} + \frac{n_{\max} C_{\text{rank}}}{4\tau_{\text{mask}}} e^{-\delta/\tau_{\text{rank}}}.$$

Applying these bounds to (K.6), and absorbing constants depending only on $(q, n_{\max})$ into $A_1, A_2, B_1, B_2 > 0$, we obtain

$$\sum_l |w_{l,t}^{\tau,\text{top}} - w_{l,t}^{0,\text{top}}| \leq A_1 e^{-B_1/\tau_{\text{mask}}} + A_2 e^{-B_2\delta/\tau_{\text{rank}}}.$$

An identical bound holds for the short leg. Since $w_{l,t}^{\tau} = w_{l,t}^{\tau,\text{top}} - w_{l,t}^{\tau,\text{bot}}$ and $w_{l,t}^0 = w_{l,t}^{0,\text{top}} - w_{l,t}^{0,\text{bot}}$,

$$\|w_{\cdot,t}^{\tau} - w_{\cdot,t}^0\|_1 \leq \sum_l |w_{l,t}^{\tau,\text{top}} - w_{l,t}^{0,\text{top}}| + \sum_l |w_{l,t}^{\tau,\text{bot}} - w_{l,t}^{0,\text{bot}}| \leq C_1 e^{-c/\tau_{\text{mask}}} + C_2 e^{-c\delta/\tau_{\text{rank}}},$$

which proves part (1).

For part (2), define $E_t(\tau) = \|w_{\cdot,t}^{\tau}(\theta) - w_{\cdot,t}^0(\theta)\|_1$. By part (1),

$$E_t(\tau) \leq C_1 e^{-c/\tau_{\text{mask}}} + C_2 e^{-c\Delta_{\theta}^{(t)}/\tau_{\text{rank}}} =: a(\tau) + b_t(\tau).$$

Fix a sequence $\tau_T \downarrow 0$ and write $a_T = a(\tau_T)$ and $b_{t,T} = b_t(\tau_T)$. For any $\delta > 0$, decompose according to the margin event:

$$b_{t,T} \leq C_2 e^{-c\delta/\tau_{\text{rank},T}} 1\{\Delta_\theta^{(t)} > \delta\} + C_2 1\{\Delta_\theta^{(t)} \leq \delta\} =: u_T(\delta) 1\{\Delta_\theta^{(t)} > \delta\} + C_2 1\{\Delta_\theta^{(t)} \leq \delta\},$$

where $u_T(\delta) = C_2 e^{-c\delta/\tau_{\text{rank},T}} \rightarrow 0$ as $T \rightarrow \infty$. Therefore,

$$\frac{1}{T} \sum_{t=1}^T E_t(\tau_T) \leq a_T + \frac{1}{T} \sum_{t=1}^T b_{t,T} \leq a_T + u_T(\delta) + C_2 \frac{1}{T} \sum_{t=1}^T 1\{\Delta_\theta^{(t)} \leq \delta\}.$$

By A1 and the “no point mass at 0” assumption, the process $\{\Delta_\theta^{(t)}\}_{t \geq 1}$ is strictly stationary and β -mixing, so the law of large numbers for indicators yields

$$\frac{1}{T} \sum_{t=1}^T 1\{\Delta_\theta^{(t)} \leq \delta\} \xrightarrow{p} \mathbb{P}(\Delta_\theta \leq \delta) =: p(\delta),$$

and $p(\delta) \rightarrow 0$ as $\delta \downarrow 0$. We now use this result to show that $\frac{1}{T} \sum_{t=1}^T E_t(\tau_T) \xrightarrow{p} 0$: First, fix $\varepsilon > 0$ and choose $\delta > 0$ small enough that

$$C_2 p(\delta) < \varepsilon/4.$$

By convergence in probability of the empirical frequency, there exists T_1 such that for all $T \geq T_1$,

$$\mathbb{P}\left(C_2 \frac{1}{T} \sum_{t=1}^T 1\{\Delta_\theta^{(t)} \leq \delta\} \leq C_2 p(\delta) + \varepsilon/4\right) \geq 1 - \eta_T,$$

with $\eta_T \rightarrow 0$ as $T \rightarrow \infty$. Next, since $a_T \rightarrow 0$ and $u_T(\delta) \rightarrow 0$, there exists T_2 such that for all $T \geq T_2$,

$$a_T < \varepsilon/4 \quad \text{and} \quad u_T(\delta) < \varepsilon/4.$$

Let $T_0 = \max\{T_1, T_2\}$ and consider any $T \geq T_0$. On the event

$$A_T = \left\{ C_2 \frac{1}{T} \sum_{t=1}^T 1\{\Delta_\theta^{(t)} \leq \delta\} \leq C_2 p(\delta) + \varepsilon/4 \right\},$$

we then have

$$\begin{aligned} \frac{1}{T} \sum_{t=1}^T E_t(\tau_T) &\leq a_T + u_T(\delta) + C_2 \frac{1}{T} \sum_{t=1}^T 1\{\Delta_\theta^{(t)} \leq \delta\} \\ &\leq \varepsilon/4 + \varepsilon/4 + (C_2 p(\delta) + \varepsilon/4) \\ &\leq \varepsilon/4 + \varepsilon/4 + \varepsilon/4 + \varepsilon/4 = \varepsilon, \end{aligned}$$

where we used $C_2 p(\delta) < \varepsilon/4$. Therefore, for all $T \geq T_0$,

$$\mathbb{P}\left(\frac{1}{T} \sum_{t=1}^T E_t(\tau_T) > \varepsilon\right) \leq \mathbb{P}(A_T^c) = \eta_T \rightarrow 0,$$

which shows that

$$\frac{1}{T} \sum_{t=1}^T E_t(\tau_T) \xrightarrow{p} 0.$$

To obtain Sharpe ratio convergence, let $\pi_t^\tau = \sum_l w_{l,t}^\tau r_{l,t}$ and $\pi_t^0 = \sum_l w_{l,t}^0 r_{l,t}$. We can bound the dif-

ference in managed returns using the triangle inequality and the uniform bound on returns from A2:

$$\begin{aligned}\pi_t^\tau - \pi_t^0 &= \sum_l (w_{l,t}^\tau - w_{l,t}^0) r_{l,t}, \\ \Rightarrow |\pi_t^\tau - \pi_t^0| &\leq \sum_l |w_{l,t}^\tau - w_{l,t}^0| |r_{l,t}| \leq R \sum_l |w_{l,t}^\tau - w_{l,t}^0| = R E_t(\tau).\end{aligned}$$

Averaging over t and using the triangle inequality once more gives

$$\begin{aligned}\left| \frac{1}{T} \sum_{t=1}^T \pi_t^{\tau_T} - \frac{1}{T} \sum_{t=1}^T \pi_t^0 \right| &= \left| \frac{1}{T} \sum_{t=1}^T (\pi_t^{\tau_T} - \pi_t^0) \right| \\ &\leq \frac{1}{T} \sum_{t=1}^T |\pi_t^{\tau_T} - \pi_t^0| \\ &\leq \frac{R}{T} \sum_{t=1}^T E_t(\tau_T).\end{aligned}$$

Since we have already shown that $\frac{1}{T} \sum_{t=1}^T E_t(\tau_T) \xrightarrow{p} 0$, it follows that

$$\left| \frac{1}{T} \sum_{t=1}^T \pi_t^{\tau_T} - \frac{1}{T} \sum_{t=1}^T \pi_t^0 \right| \xrightarrow{p} 0.$$

For the second moments, use the identity

$$(\pi_t^{\tau_T})^2 - (\pi_t^0)^2 = (\pi_t^{\tau_T} - \pi_t^0)(\pi_t^{\tau_T} + \pi_t^0).$$

Therefore

$$\begin{aligned}\left| \frac{1}{T} \sum_{t=1}^T (\pi_t^{\tau_T})^2 - \frac{1}{T} \sum_{t=1}^T (\pi_t^0)^2 \right| &= \left| \frac{1}{T} \sum_{t=1}^T [(\pi_t^{\tau_T})^2 - (\pi_t^0)^2] \right| \\ &\leq \frac{1}{T} \sum_{t=1}^T |(\pi_t^{\tau_T})^2 - (\pi_t^0)^2| \\ &= \frac{1}{T} \sum_{t=1}^T |\pi_t^{\tau_T} - \pi_t^0| |\pi_t^{\tau_T} + \pi_t^0|.\end{aligned}$$

By A2, $|r_{l,t}| \leq R$ and $\sum_l |w_{l,t}^\tau| \leq 2$ imply $|\pi_t^\tau| \leq 2R$ and $|\pi_t^0| \leq 2R$ for all t , so

$$|\pi_t^{\tau_T}| + |\pi_t^0| \leq 4R,$$

and therefore

$$\begin{aligned}\left| \frac{1}{T} \sum_{t=1}^T (\pi_t^{\tau_T})^2 - \frac{1}{T} \sum_{t=1}^T (\pi_t^0)^2 \right| &\leq \frac{1}{T} \sum_{t=1}^T |\pi_t^{\tau_T} - \pi_t^0| (|\pi_t^{\tau_T}| + |\pi_t^0|) \\ &\leq 4R \frac{1}{T} \sum_{t=1}^T |\pi_t^{\tau_T} - \pi_t^0|.\end{aligned}$$

Using $|\pi_t^{\tau_T} - \pi_t^0| \leq RE_t(\tau_T)$ and $\frac{1}{T} \sum_{t=1}^T E_t(\tau_T) \xrightarrow{p} 0$ again yields

$$\left| \frac{1}{T} \sum_{t=1}^T (\pi_t^{\tau_T})^2 - \frac{1}{T} \sum_{t=1}^T (\pi_t^0)^2 \right| \xrightarrow{p} 0.$$

By A4, the population volatility of the hard portfolio satisfies $\sigma^0(\theta) \geq \underline{\sigma} > 0$. From the convergence of the sample mean and second moment in the hard case, $\hat{\mu}_T^0 \rightarrow \mu^0$ and $\hat{\nu}_T^0 \rightarrow \nu^0$, we obtain $\hat{\sigma}_T^{2,0} = \hat{\nu}_T^0 - (\hat{\mu}_T^0)^2 \rightarrow \sigma^{2,0}$ and hence $\hat{\sigma}_T^0 \rightarrow \sigma^0$ in probability. In particular,

$$\mathbb{P}(\hat{\sigma}_T^0 \geq \underline{\sigma}/2) \rightarrow 1.$$

Similarly, from

$$|\hat{\mu}_T^{\tau_T} - \hat{\mu}_T^0| \xrightarrow{p} 0, \quad |\hat{\nu}_T^{\tau_T} - \hat{\nu}_T^0| \xrightarrow{p} 0,$$

we have

$$\hat{\sigma}_T^{2,\tau_T} = \hat{\nu}_T^{\tau_T} - (\hat{\mu}_T^{\tau_T})^2 \quad \text{and} \quad \hat{\sigma}_T^{2,0} = \hat{\nu}_T^0 - (\hat{\mu}_T^0)^2,$$

so that

$$|\hat{\sigma}_T^{2,\tau_T} - \hat{\sigma}_T^{2,0}| \xrightarrow{p} 0.$$

On the event $\{\hat{\sigma}_T^0 \geq \underline{\sigma}/2\}$, the denominator $\hat{\sigma}_T^{\tau_T} + \hat{\sigma}_T^0$ is bounded away from 0 for all large T , and

$$|\hat{\sigma}_T^{\tau_T} - \hat{\sigma}_T^0| = \frac{|\hat{\sigma}_T^{2,\tau_T} - \hat{\sigma}_T^{2,0}|}{\hat{\sigma}_T^{\tau_T} + \hat{\sigma}_T^0} \xrightarrow{p} 0.$$

Finally, the Sharpe ratio can be written as $g(\mu, \sigma^2) = \mu / \sqrt{\sigma^2}$, which is continuous on $\{\sigma^2 > 0\}$. Since $(\hat{\mu}_T^{\tau_T}, \hat{\sigma}_T^{2,\tau_T})$ and $(\hat{\mu}_T^0, \hat{\sigma}_T^{2,0})$ both lie in this domain for all large T with probability tending to one, and

$$(\hat{\mu}_T^{\tau_T}, \hat{\sigma}_T^{2,\tau_T}) - (\hat{\mu}_T^0, \hat{\sigma}_T^{2,0}) \xrightarrow{p} 0,$$

the continuous mapping theorem yields

$$\widehat{\text{SR}}_T^{\text{soft}}(\theta) - \widehat{\text{SR}}_T^{\text{hard}}(\theta) = g(\hat{\mu}_T^{\tau_T}, \hat{\sigma}_T^{2,\tau_T}) - g(\hat{\mu}_T^0, \hat{\sigma}_T^{2,0}) \xrightarrow{p} 0,$$

which proves part (2).

For part (3), fix T and a parameter value θ such that $\Delta_\theta^{(t)} > 0$ for all $t = 1, \dots, T$. Recall that

$$\Delta_\theta^{(t)} = \min_{i \neq j} |s_{i,t}^\theta - s_{j,t}^\theta|$$

is the minimal pairwise signal gap in month t . Since T is finite and each $\Delta_\theta^{(t)} > 0$, we may define

$$\bar{\delta} = \min_{1 \leq t \leq T} \Delta_\theta^{(t)} > 0,$$

so that for all t and all $i \neq j$ we have $|s_{i,t}^\theta - s_{j,t}^\theta| \geq \bar{\delta}$. By A3, the signals are Lipschitz in θ : there exists $L_s < \infty$ such that

$$\max_{i,t} |s_{i,t}^{\theta'} - s_{i,t}^\theta| \leq L_s \|\theta' - \theta\| \quad \text{for all } \theta', \theta \in \Theta.$$

Hence, for any θ' and any $i \neq j$,

$$|(s_{i,t}^{\theta'} - s_{j,t}^{\theta'}) - (s_{i,t}^{\theta} - s_{j,t}^{\theta})| \leq |s_{i,t}^{\theta'} - s_{i,t}^{\theta}| + |s_{j,t}^{\theta'} - s_{j,t}^{\theta}| \leq 2L_s \|\theta' - \theta\|.$$

Choose $\eta > 0$ such that $2L_s\eta < \bar{\delta}/2$, and consider any θ' with $\|\theta' - \theta\| \leq \eta$. Then, for all $i \neq j$ and all t ,

$$|s_{i,t}^{\theta'} - s_{j,t}^{\theta'}| \geq |s_{i,t}^{\theta} - s_{j,t}^{\theta}| - |(s_{i,t}^{\theta'} - s_{j,t}^{\theta'}) - (s_{i,t}^{\theta} - s_{j,t}^{\theta})| \geq \bar{\delta} - \frac{\bar{\delta}}{2} = \frac{\bar{\delta}}{2} > 0.$$

In particular, the sign of $s_{i,t}^{\theta'} - s_{j,t}^{\theta'}$ coincides with the sign of $s_{i,t}^{\theta} - s_{j,t}^{\theta}$. Therefore, the ranking of the signals is unchanged, and the hard top/bottom memberships (which assets are in the top and bottom k_t) do not change when θ' varies within $\{\theta' : \|\theta' - \theta\| \leq \eta\}$.

Since the hard portfolio weights $w_{l,t}^0(\cdot)$ depend only on these memberships, it follows that

$$w_{l,t}^0(\theta') = w_{l,t}^0(\theta) \quad \text{for all } l, t \quad \text{whenever } \|\theta' - \theta\| \leq \eta.$$

Consequently, the hard-sample Sharpe ratio $\widehat{\text{SR}}_T^{\text{hard}}(\theta')$ is locally constant (and hence locally Lipschitz) in a neighborhood of θ . A locally Lipschitz function that is constant in a neighborhood has Clarke subdifferential $\{0\}$ at that point, so

$$\partial^{\text{Clarke}} \widehat{\text{SR}}_T^{\text{hard}}(\theta) = \{0\}.$$

It remains to show that $\|\nabla_{\theta} \widehat{\text{SR}}_T^{\text{soft}}(\theta|\tau)\| \rightarrow 0$ as $\tau \downarrow 0$. First, view the sample Sharpe ratio as a function of the return path $(\pi_1, \dots, \pi_T)$:

$$\hat{\mu}_T = \frac{1}{T} \sum_{u=1}^T \pi_u, \quad \hat{\nu}_T = \frac{1}{T} \sum_{u=1}^T \pi_u^2, \quad \hat{\sigma}_T^2 = \hat{\nu}_T - \hat{\mu}_T^2, \quad \widehat{\text{SR}}_T = \frac{\hat{\mu}_T}{\hat{\sigma}_T}.$$

Differentiation yields

$$\frac{\partial \hat{\mu}_T}{\partial \pi_t} = \frac{1}{T}, \quad \frac{\partial \hat{\sigma}_T^2}{\partial \pi_t} = \frac{2}{T}(\pi_t - \hat{\mu}_T), \quad \frac{\partial \hat{\sigma}_T}{\partial \pi_t} = \frac{\pi_t - \hat{\mu}_T}{T \hat{\sigma}_T},$$

and therefore

$$\frac{\partial \widehat{\text{SR}}_T}{\partial \pi_t} = \frac{1}{T \hat{\sigma}_T} \left(1 - \frac{\hat{\mu}_T}{\hat{\sigma}_T^2} (\pi_t - \hat{\mu}_T) \right).$$

By A2 and the normalization of the weights, $|\pi_t| \leq 2R$ for all t and all τ , so $|\hat{\mu}_T| \leq 2R$ and $|\pi_t - \hat{\mu}_T| \leq |\pi_t| + |\hat{\mu}_T| \leq 4R$. By assumption, the hard return path $\{\pi_t(\theta; 0)\}_{t=1}^T$ has strictly positive sample standard deviation, so $\hat{\sigma}_T(\theta; 0) > 0$. Since $\pi_t(\theta|\tau)$ depends continuously on τ , it follows that $\hat{\sigma}_T(\theta|\tau)$ is continuous in τ , and thus there exist $c > 0$ and $\bar{\tau} > 0$ such that

$$\hat{\sigma}_T(\theta|\tau) \geq c \quad \text{for all } 0 \leq \tau \leq \bar{\tau}.$$

Combining these bounds, we obtain, for all sufficiently small τ ,

$$\left| \frac{\partial \widehat{\text{SR}}_T}{\partial \pi_t} \right| \leq \frac{1}{T \hat{\sigma}_T} \left(1 + \frac{|\hat{\mu}_T|}{\hat{\sigma}_T^2} |\pi_t - \hat{\mu}_T| \right) \leq \frac{1}{Tc} \left(1 + \frac{2R}{c^2} \cdot 4R \right) =: C_{\text{SR}} < \infty,$$

where C_{SR} does not depend on τ . By the chain rule, we can now write

$$\nabla_{\theta} \widehat{\text{SR}}_T(\theta|\tau) = \sum_{t=1}^T \frac{\partial \widehat{\text{SR}}_T}{\partial \pi_t} \nabla_{\theta} \pi_t(\theta|\tau),$$

where

$$\nabla_{\theta} \pi_t(\theta|\tau) = r_t^{\top} J_t^{\top}(\theta) \nabla_{\theta} s_t^{\theta}.$$

Here $s_t^{\theta} = (s_{1,t}^{\theta}, \dots, s_{n_t,t}^{\theta})^{\top}$ is the signal vector, $w_t^{\tau}(\theta)$ is the soft weight vector, $J_t^{\top}(\theta) = \partial w_{\cdot,t}^{\tau}(\theta) / \partial s_{\cdot,t}^{\theta}$ is the Jacobian of the soft mapping with respect to the signals, and $\nabla_{\theta} s_t^{\theta}$ is uniformly bounded in operator norm over (t, θ) by A3. We now bound $\|J_t^{\top}(\theta)\|$ under the positive margin $\bar{\delta}$. In the pairwise layer,

$$S_{l,j,t}^{\tau} = \sigma\left(\frac{s_{j,t}^{\theta} - s_{l,t}^{\theta}}{\tau_{\text{rank}}}\right), \quad \sigma(x) = \frac{1}{1 + e^{-x}},$$

so with $z = (s_{j,t}^{\theta} - s_{l,t}^{\theta}) / \tau_{\text{rank}}$,

$$\frac{\partial S_{l,j,t}^{\tau}}{\partial s_{l,t}^{\theta}} = -\frac{1}{\tau_{\text{rank}}} \sigma'(z), \quad \sigma'(z) = \sigma(z)(1 - \sigma(z)).$$

A direct calculation shows $\sigma'(z) = e^{-z} / (1 + e^{-z})^2 \leq e^{-|z|}$ for all z , and by the margin assumption $|s_{j,t}^{\theta} - s_{l,t}^{\theta}| \geq \bar{\delta}$, so $|z| \geq \bar{\delta} / \tau_{\text{rank}}$. Therefore

$$\left| \frac{\partial S_{l,j,t}^{\tau}}{\partial s_{l,t}^{\theta}} \right| \leq \frac{1}{\tau_{\text{rank}}} e^{-|z|} \leq \frac{1}{\tau_{\text{rank}}} e^{-\bar{\delta} / \tau_{\text{rank}}},$$

and the same bound holds for $\partial S_{l,j,t}^{\tau} / \partial s_{j,t}^{\theta}$, while all other partials are zero. The soft rank $\text{rank}_{l,t}^{\tau} = 1 + \sum_{j \neq l} S_{l,j,t}^{\tau}$ therefore satisfies

$$\left\| \frac{\partial \text{rank}_{l,t}^{\tau}}{\partial s_{\cdot,t}^{\theta}} \right\| \leq C \frac{e^{-\bar{\delta} / \tau_{\text{rank}}}}{\tau_{\text{rank}}},$$

for some constant $C > 0$ depending only on n_{max} (e.g. the maximum row-sum norm). In the mask layer,

$$m_{l,t}^{\tau, \text{top}} = \sigma\left(\frac{k_t + 0.5 - \text{rank}_{l,t}^{\tau}}{\tau_{\text{mask}}}\right), \quad u = \frac{k_t + 0.5 - \text{rank}_{l,t}^{\tau}}{\tau_{\text{mask}}},$$

and

$$\frac{\partial m_{l,t}^{\tau, \text{top}}}{\partial \text{rank}_{l,t}^{\tau}} = -\frac{1}{\tau_{\text{mask}}} \sigma'(u), \quad \text{so} \quad \left| \frac{\partial m_{l,t}^{\tau, \text{top}}}{\partial \text{rank}_{l,t}^{\tau}} \right| = \frac{1}{\tau_{\text{mask}}} \sigma'(u).$$

At $\tau = 0$, $a_{l,t} = k_t + 0.5 - \text{rank}_{l,t}^0$ is a half-integer, so $|a_{l,t}| \geq \frac{1}{2}$. From part (1) we know that $|\text{rank}_{l,t}^{\tau} - \text{rank}_{l,t}^0| \leq C_{\text{rank}} e^{-\bar{\delta} / \tau_{\text{rank}}}$, hence for all sufficiently small τ ,

$$|k_t + 0.5 - \text{rank}_{l,t}^{\tau}| \geq |a_{l,t}| - |\text{rank}_{l,t}^{\tau} - \text{rank}_{l,t}^0| \geq \frac{1}{2} - C_{\text{rank}} e^{-\bar{\delta} / \tau_{\text{rank}}} \geq \frac{1}{4}.$$

Thus $|u| \geq c / \tau_{\text{mask}}$ for some $c \in (0, 1/2)$ and all small τ , and using again $\sigma'(u) \leq e^{-|u|}$ we obtain

$$\left| \frac{\partial m_{l,t}^{\tau, \text{top}}}{\partial \text{rank}_{l,t}^{\tau}} \right| \leq \frac{1}{\tau_{\text{mask}}} e^{-|u|} \leq \frac{1}{\tau_{\text{mask}}} e^{-c / \tau_{\text{mask}}}.$$

An analogous bound holds for the bottom mask.

Finally, the normalized weights $w_{l,t}^{\tau}$ are obtained by dividing the masks by their leg sums $Z_t^{\tau, \text{top}}, Z_t^{\tau, \text{bot}}$, which stay in $[k_t/2, 3k_t/2]$ for all small τ by (K.5). The derivatives of this normalization step are therefore uniformly bounded by constants depending only on (q, n_{max}) . Combining the bounds from

the pairwise and mask layers yields

$$\|J_t^\tau(\theta)\| \leq C' \left(\frac{e^{-\bar{\delta}/\tau_{\text{rank}}}}{\tau_{\text{rank}}} + \frac{e^{-c/\tau_{\text{mask}}}}{\tau_{\text{mask}}} \right),$$

for some constant $C' > 0$ depending only on q and n_{max} . Therefore, using $\|r_t\|_\infty \leq R$ from A2 and $\|\nabla_\theta s_t^\theta\| \leq L_s$ from A3, we have

$$\nabla_\theta \pi_t(\theta|\tau) = r_t^\top J_t^\tau(\theta) \nabla_\theta s_t^\theta,$$

and hence, with a suitable operator norm,

$$\|\nabla_\theta \pi_t(\theta|\tau)\| \leq \|r_t\|_\infty \|J_t^\tau(\theta)\| \|\nabla_\theta s_t^\theta\| \leq R \|J_t^\tau(\theta)\| L_s.$$

Using the bound on $J_t^\tau(\theta)$ derived above, this yields

$$\|\nabla_\theta \pi_t(\theta|\tau)\| \leq C'' \left(\frac{e^{-\bar{\delta}/\tau_{\text{rank}}}}{\tau_{\text{rank}}} + \frac{e^{-c/\tau_{\text{mask}}}}{\tau_{\text{mask}}} \right), \quad C'' = RC' L_s.$$

For any $a > 0$, the elementary limit

$$\frac{e^{-a/x}}{x} \rightarrow 0 \quad \text{as } x \downarrow 0$$

holds (e.g. by the change of variables $y = 1/x$, so $e^{-a/x}/x = ye^{-ay} \rightarrow 0$ as $y \rightarrow \infty$). Applying this with $a = \bar{\delta}$ and $x = \tau_{\text{rank}}$, and with $a = c$ and $x = \tau_{\text{mask}}$, we conclude that

$$\|\nabla_\theta \pi_t(\theta|\tau)\| \rightarrow 0 \quad \text{as } \tau \downarrow 0$$

for each fixed t .

Finally, recall that

$$\nabla_\theta \widehat{\text{SR}}_T(\theta|\tau) = \sum_{t=1}^T \frac{\partial \widehat{\text{SR}}_T}{\partial \pi_t} \nabla_\theta \pi_t(\theta|\tau),$$

and we have $|\frac{\partial \widehat{\text{SR}}_T}{\partial \pi_t}| \leq C_{\text{SR}}$ for some finite C_{SR} (since $\hat{\sigma}_T$ is bounded away from 0 for small τ by A4 and continuity from the hard case). Thus

$$\|\nabla_\theta \widehat{\text{SR}}_T(\theta|\tau)\| \leq TC_{\text{SR}} \max_{1 \leq t \leq T} \|\nabla_\theta \pi_t(\theta|\tau)\| \rightarrow 0 \quad \text{as } \tau \downarrow 0.$$

This proves gradient consistency and completes part (3). $\square$

Proof of Theorem 2

Proof. As above, for each t , write $\pi_t(\theta) = r_t^\top w_t(\theta)$ with $r_t = (r_{1,t}, \dots, r_{n_t,t})$ and $w_t(\theta) = (w_{1,t}^\theta, \dots, w_{n_t,t}^\theta)$. Define $\mu(\theta) = \mathbb{E}[\pi_t(\theta)]$, $\nu(\theta) = \mathbb{E}[\pi_t(\theta)^2]$, $\sigma^2(\theta) = \nu(\theta) - \mu(\theta)^2$, and $\text{SR}(\theta) = \mu(\theta)/\sigma(\theta)$. Define $\hat{\mu}_T(\theta) = \frac{1}{T} \sum_{t=1}^T \pi_t(\theta)$ and $\hat{\nu}_T(\theta) = \frac{1}{T} \sum_{t=1}^T \pi_t(\theta)^2$. Then $\hat{\sigma}_T^2(\theta) = \hat{\nu}_T(\theta) - \hat{\mu}_T(\theta)^2$ and $\widehat{\text{SR}}_T(\theta) = \hat{\mu}_T(\theta)/\hat{\sigma}_T(\theta)$.

By A2, $\sum_l |w_{l,t}^\theta| \leq 2$ and $|r_{l,t}| \leq R$, therefore $|\pi_t(\theta)| \leq 2R$ uniformly in (t, θ) . The soft rank/mask mapping is a composition of sigmoids. For the logistic $\sigma(x) = 1/(1 + e^{-x})$ we have $|\sigma'(x)| \leq 1/4$. The inner pairwise layer $S_{l',t} = \sigma((s_{l',t}^\theta - s_{l,t}^\theta)/\tau_{\text{rank}})$ has derivatives bounded by $(4\tau_{\text{rank}})^{-1}$ with respect to $s_{l',t}^\theta, s_{l,t}^\theta$. The soft ranks and masks are finite sums and further compositions with sigmoids of

temperature τ_{mask} , so by repeated application of the chain rule, there exists $L_w < \infty$ such that

$$\sum_{l=1}^{n_t} |w_{l,t}(\theta) - w_{l,t}(\theta')| \leq L_w \max_u |s_{u,t}^\theta - s_{u,t}^{\theta'}| \quad \text{for all } t, \theta, \theta'.$$

By A3 there exists $L_s < \infty$ with $\max_u |s_{u,t}^\theta - s_{u,t}^{\theta'}| \leq L_s \|\theta - \theta'\|$ uniformly in t . Combining these bounds,

$$|\pi_t(\theta) - \pi_t(\theta')| = |r_t^\top (w_t(\theta) - w_t(\theta'))| \leq \|r_t\|_\infty \sum_l |w_{l,t}(\theta) - w_{l,t}(\theta')| \leq RL_w L_s \|\theta - \theta'\|$$

uniformly in t . Therefore, the return map $\theta \mapsto \pi_t(\theta)$ is bounded and Lipschitz uniformly in t .

Next, we prove a uniform law of large numbers for the classes $\{\pi_t(\theta) : \theta \in \Theta\}$ and $\{\pi_t(\theta)^2 : \theta \in \Theta\}$. Recall from the previous step that there exists a constant $L_\pi < \infty$ such that

$$|\pi_t(\theta) - \pi_t(\theta')| \leq L_\pi \|\theta - \theta'\| \quad \text{for all } t, \theta, \theta' \in \Theta,$$

and, by A2, $|\pi_t(\theta)| \leq B$ for all (t, θ) , where we may take $B = 2R$.

Fix $\varepsilon > 0$ and choose a finite ε' -net $\{\theta^{(j)}\}_{j=1}^J \subset \Theta$ with

$$\varepsilon' = \frac{\varepsilon}{3L_\pi},$$

so that for any $\theta \in \Theta$ there exists j with $\|\theta - \theta^{(j)}\| \leq \varepsilon'$. For each fixed j , by A1-A2 and boundedness of $\pi_t(\theta^{(j)})$, the scalar average

$$\frac{1}{T} \sum_{t=1}^T \pi_t(\theta^{(j)}) - \mathbb{E}[\pi_t(\theta^{(j)})]$$

converges to 0 in probability as $T \rightarrow \infty$ (Law of large numbers for bounded β -mixing sequences). Since $J < \infty$, a union bound yields

$$\max_{1 \leq j \leq J} \left| \frac{1}{T} \sum_{t=1}^T \pi_t(\theta^{(j)}) - \mathbb{E}[\pi_t(\theta^{(j)})] \right| \xrightarrow{p} 0.$$

Now fix an arbitrary $\theta \in \Theta$ and let $\theta^{(j)}$ be a net point with $\|\theta - \theta^{(j)}\| \leq \varepsilon'$. Then

$$\begin{aligned} \left| \frac{1}{T} \sum_{t=1}^T \pi_t(\theta) - \mathbb{E}[\pi_t(\theta)] \right| &\leq \frac{1}{T} \sum_{t=1}^T |\pi_t(\theta) - \pi_t(\theta^{(j)})| + \left| \frac{1}{T} \sum_{t=1}^T \pi_t(\theta^{(j)}) - \mathbb{E}[\pi_t(\theta^{(j)})] \right| \\ &\quad + \left| \mathbb{E}[\pi_t(\theta^{(j)})] - \mathbb{E}[\pi_t(\theta)] \right|. \end{aligned}$$

The middle term converges to 0 in probability uniformly over j by the previous result. The first and third terms are deterministically bounded by

$$\frac{1}{T} \sum_{t=1}^T |\pi_t(\theta) - \pi_t(\theta^{(j)})| \leq L_\pi \|\theta - \theta^{(j)}\| \leq L_\pi \varepsilon' = \varepsilon/3,$$

and

$$|\mathbb{E}[\pi_t(\theta^{(j)})] - \mathbb{E}[\pi_t(\theta)]| \leq \mathbb{E}[|\pi_t(\theta^{(j)}) - \pi_t(\theta)|] \leq L_\pi \|\theta - \theta^{(j)}\| \leq \varepsilon/3.$$

Now, for convenience, define

$$M_{T,j} = \left| \frac{1}{T} \sum_{t=1}^T \pi_t(\theta^{(j)}) - \mathbb{E}[\pi_t(\theta^{(j)})] \right|, \quad 1 \leq j \leq J.$$

From the result above and the law of large numbers under A1-A2, we have

$$\max_{1 \leq j \leq J} M_{T,j} \xrightarrow{p} 0.$$

Combining the bounds on the three terms in the decomposition, for any $\theta \in \Theta$ and its net point $\theta^{(j)}$,

$$\left| \frac{1}{T} \sum_{t=1}^T \pi_t(\theta) - \mathbb{E}[\pi_t(\theta)] \right| \leq \frac{\varepsilon}{3} + M_{T,j} + \frac{\varepsilon}{3} = \frac{2\varepsilon}{3} + M_{T,j}.$$

Taking the supremum over $\theta \in \Theta$ and noting that for each θ the index j is some element of $\{1, \dots, J\}$, we obtain

$$\sup_{\theta \in \Theta} \left| \frac{1}{T} \sum_{t=1}^T \pi_t(\theta) - \mathbb{E}[\pi_t(\theta)] \right| \leq \frac{2\varepsilon}{3} + \max_{1 \leq j \leq J} M_{T,j}.$$

Therefore, for any $\varepsilon > 0$,

$$\mathbb{P} \left(\sup_{\theta \in \Theta} \left| \frac{1}{T} \sum_{t=1}^T \pi_t(\theta) - \mathbb{E}[\pi_t(\theta)] \right| > \varepsilon \right) \leq \mathbb{P} \left(\max_{1 \leq j \leq J} M_{T,j} > \frac{\varepsilon}{3} \right) \rightarrow 0,$$

because $\max_{1 \leq j \leq J} M_{T,j} \xrightarrow{p} 0$. Therefore,

$$\sup_{\theta \in \Theta} \left| \frac{1}{T} \sum_{t=1}^T \pi_t(\theta) - \mathbb{E}[\pi_t(\theta)] \right| \xrightarrow{p} 0.$$

The same argument applies to the class $\{\pi_t(\theta)^2 : \theta \in \Theta\}$. Indeed, for any $\theta, \theta' \in \Theta$,

$$|\pi_t(\theta)^2 - \pi_t(\theta')^2| = |\pi_t(\theta) - \pi_t(\theta')| \cdot |\pi_t(\theta) + \pi_t(\theta')| \leq 2B |\pi_t(\theta) - \pi_t(\theta')| \leq 2BL_\pi \|\theta - \theta'\|,$$

so $\theta \mapsto \pi_t(\theta)^2$ is also uniformly Lipschitz in θ with a bounded envelope $|\pi_t(\theta)^2| \leq B^2$. For each fixed θ , the bounded β -mixing sequence $\{\pi_t(\theta)^2\}_{t \geq 1}$ satisfies the law of large numbers under A1-A2. Repeating the ε -net argument therefore yields

$$\sup_{\theta \in \Theta} |\hat{\mu}_T(\theta) - \mu(\theta)| \xrightarrow{p} 0, \quad \sup_{\theta \in \Theta} |\hat{\nu}_T(\theta) - \nu(\theta)| \xrightarrow{p} 0.$$

We now use this to show uniform convergence of variances and Sharpe ratios. Since $\hat{\sigma}_T^2(\theta) = \hat{\nu}_T(\theta) - \hat{\mu}_T(\theta)^2$ and $\sigma^2(\theta) = \nu(\theta) - \mu(\theta)^2$, for all θ ,

$$|\hat{\sigma}_T^2(\theta) - \sigma^2(\theta)| \leq |\hat{\nu}_T(\theta) - \nu(\theta)| + (|\hat{\mu}_T(\theta)| + |\mu(\theta)|) |\hat{\mu}_T(\theta) - \mu(\theta)|.$$

Moreover, by A2 we have $|\pi_t(\theta)| \leq B \leq 2R$ almost surely for all (t, θ) , so

$$|\hat{\mu}_T(\theta)| = \left| \frac{1}{T} \sum_{t=1}^T \pi_t(\theta) \right| \leq \frac{1}{T} \sum_{t=1}^T |\pi_t(\theta)| \leq B \leq 2R,$$

and

$$|\mu(\theta)| = |\mathbb{E}[\pi_t(\theta)]| \leq \mathbb{E}[|\pi_t(\theta)|] \leq B \leq 2R,$$

uniformly in $\theta \in \Theta$.

By the bound above and the uniform envelope $|\hat{\mu}_T(\theta)|, |\mu(\theta)| \leq 2R$, for all $\theta \in \Theta$,

$$|\hat{\sigma}_T^2(\theta) - \sigma^2(\theta)| \leq |\hat{\nu}_T(\theta) - \nu(\theta)| + 4R |\hat{\mu}_T(\theta) - \mu(\theta)|.$$

Taking the supremum over $\theta \in \Theta$ yields

$$\sup_{\theta \in \Theta} |\hat{\sigma}_T^2(\theta) - \sigma^2(\theta)| \leq \sup_{\theta \in \Theta} |\hat{\nu}_T(\theta) - \nu(\theta)| + 4R \sup_{\theta \in \Theta} |\hat{\mu}_T(\theta) - \mu(\theta)|.$$

By the uniform laws of large numbers established above,

$$\sup_{\theta \in \Theta} |\hat{\mu}_T(\theta) - \mu(\theta)| \xrightarrow{p} 0, \quad \sup_{\theta \in \Theta} |\hat{\nu}_T(\theta) - \nu(\theta)| \xrightarrow{p} 0,$$

so the right-hand side converges to 0 in probability. Hence,

$$\sup_{\theta \in \Theta} |\hat{\sigma}_T^2(\theta) - \sigma^2(\theta)| \xrightarrow{p} 0.$$

By A4, $\sigma(\theta) \geq \underline{\sigma} > 0$ for all $\theta \in \Theta$. Together with $\sup_{\theta} |\hat{\sigma}_T^2(\theta) - \sigma^2(\theta)| \xrightarrow{p} 0$, this implies that for any $\varepsilon \in (0, \underline{\sigma}^2)$,

$$\mathbb{P}\left(\sup_{\theta \in \Theta} |\hat{\sigma}_T^2(\theta) - \sigma^2(\theta)| \leq \varepsilon\right) \rightarrow 1,$$

and hence, on this event,

$$\hat{\sigma}_T^2(\theta) \geq \sigma^2(\theta) - \varepsilon \geq \underline{\sigma}^2 - \varepsilon \quad \text{for all } \theta.$$

Taking $\varepsilon = \underline{\sigma}^2/4$ and square roots yields $\hat{\sigma}_T(\theta) \geq \underline{\sigma}/2$ uniformly in θ , so

$$\mathbb{P}\left(\inf_{\theta \in \Theta} \hat{\sigma}_T(\theta) \geq \underline{\sigma}/2\right) \rightarrow 1.$$

Moreover, for all θ ,

$$|\hat{\sigma}_T(\theta) - \sigma(\theta)| = \frac{|\hat{\sigma}_T^2(\theta) - \sigma^2(\theta)|}{\hat{\sigma}_T(\theta) + \sigma(\theta)} \leq \frac{|\hat{\sigma}_T^2(\theta) - \sigma^2(\theta)|}{\inf_{\vartheta} \hat{\sigma}_T(\vartheta) + \inf_{\vartheta} \sigma(\vartheta)},$$

so on the high-probability event where $\inf_{\theta} \hat{\sigma}_T(\theta) \geq \underline{\sigma}/2$ we obtain

$$\sup_{\theta \in \Theta} |\hat{\sigma}_T(\theta) - \sigma(\theta)| \leq \frac{\sup_{\theta} |\hat{\sigma}_T^2(\theta) - \sigma^2(\theta)|}{\underline{\sigma}/2 + \underline{\sigma}} \rightarrow 0 \quad \text{in probability.}$$

For Sharpe ratios,

$$|\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| = \left| \frac{\hat{\mu}_T(\theta)}{\hat{\sigma}_T(\theta)} - \frac{\mu(\theta)}{\sigma(\theta)} \right| \leq \frac{|\hat{\mu}_T(\theta) - \mu(\theta)|}{\sigma(\theta)} + |\hat{\mu}_T(\theta)| \cdot \left| \frac{1}{\hat{\sigma}_T(\theta)} - \frac{1}{\sigma(\theta)} \right|.$$

By A4, $\sigma(\theta) \geq \underline{\sigma}$ for all θ . Using that and noting that

$$\left| \frac{1}{\hat{\sigma}_T(\theta)} - \frac{1}{\sigma(\theta)} \right| = \frac{|\hat{\sigma}_T(\theta) - \sigma(\theta)|}{\hat{\sigma}_T(\theta) \sigma(\theta)},$$

we obtain, on the event $\{\inf_{\theta} \hat{\sigma}_T(\theta) \geq \underline{\sigma}/2\}$,

$$\left| \frac{1}{\hat{\sigma}_T(\theta)} - \frac{1}{\sigma(\theta)} \right| \leq \frac{2}{\underline{\sigma}^2} |\hat{\sigma}_T(\theta) - \sigma(\theta)|.$$

Using $|\hat{\mu}_T(\theta)| \leq 2R$ and taking the supremum over θ in the Sharpe ratio bound then gives

$$\sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| \leq \frac{\sup_{\theta \in \Theta} |\hat{\mu}_T(\theta) - \mu(\theta)|}{\underline{\sigma}} + \frac{4R}{\underline{\sigma}^2} \sup_{\theta \in \Theta} |\hat{\sigma}_T(\theta) - \sigma(\theta)|.$$

We have already shown that $\sup_{\theta} |\hat{\mu}_T(\theta) - \mu(\theta)| \xrightarrow{p} 0$ and $\sup_{\theta} |\hat{\sigma}_T(\theta) - \sigma(\theta)| \xrightarrow{p} 0$, so the right-hand side converges to 0 in probability:

$$\sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| \xrightarrow{p} 0.$$

This proves part (1) of the theorem.

For (2), let $\Theta^* = \operatorname{argmax}_{\theta \in \Theta} \text{SR}(\theta)$ and $M = \sup_{\theta \in \Theta} \text{SR}(\theta)$. For $\varepsilon > 0$, define the ε -neighborhood of the maximizer set

$$U_\varepsilon = \{\theta \in \Theta : \text{dist}(\theta, \Theta^*) < \varepsilon\}.$$

Then $C_\varepsilon = \Theta \setminus U_\varepsilon$ is compact and does not intersect Θ^* . By continuity of SR and compactness of C_ε , the restricted supremum $M_\varepsilon = \sup_{\theta \in C_\varepsilon} \text{SR}(\theta)$ is attained. We claim that $M_\varepsilon < M$. Otherwise, there would exist a sequence $(\theta_n) \subset C_\varepsilon$ with $\text{SR}(\theta_n) \rightarrow M$. By compactness of Θ , a subsequence converges to some $\bar{\theta} \in C_\varepsilon$, and continuity of SR gives $\text{SR}(\bar{\theta}) = M$, so $\bar{\theta} \in \Theta^*$, contradicting $\bar{\theta} \in C_\varepsilon$. Thus $M_\varepsilon < M$, and we may define

$$\delta(\varepsilon) = \frac{M - M_\varepsilon}{2} > 0,$$

so that

$$\sup_{\theta \in \Theta \setminus U_\varepsilon} \text{SR}(\theta) = M_\varepsilon \leq M - 2\delta(\varepsilon) = \sup_{\theta \in \Theta} \text{SR}(\theta) - 2\delta(\varepsilon).$$

By part (1), we have

$$\sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| \xrightarrow{p} 0,$$

so for any fixed $\varepsilon > 0$,

$$\mathbb{P}\left(\sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| \leq \delta(\varepsilon)\right) \rightarrow 1.$$

On the event

$$A_T(\varepsilon) = \left\{ \sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| \leq \delta(\varepsilon) \right\},$$

we have, for all $\theta \in \Theta \setminus U_\varepsilon$,

$$\widehat{\text{SR}}_T(\theta) \leq \text{SR}(\theta) + \delta(\varepsilon) \leq \sup_{\vartheta \in \Theta \setminus U_\varepsilon} \text{SR}(\vartheta) + \delta(\varepsilon) \leq \sup_{\vartheta \in \Theta} \text{SR}(\vartheta) - \delta(\varepsilon).$$

On the other hand, pick any $\theta^* \in \Theta^*$. Then $\text{SR}(\theta^*) = \sup_{\vartheta \in \Theta} \text{SR}(\vartheta)$ and, on $A_T(\varepsilon)$,

$$\widehat{\text{SR}}_T(\theta^*) \geq \text{SR}(\theta^*) - \delta(\varepsilon) = \sup_{\vartheta \in \Theta} \text{SR}(\vartheta) - \delta(\varepsilon).$$

Combining the last two results, on $A_T(\varepsilon)$ we obtain

$$\sup_{\theta \in \Theta \setminus U_\varepsilon} \widehat{\text{SR}}_T(\theta) \leq \sup_{\vartheta \in \Theta} \text{SR}(\vartheta) - \delta(\varepsilon) \leq \widehat{\text{SR}}_T(\theta^*) \leq \sup_{\vartheta \in \Theta} \widehat{\text{SR}}_T(\vartheta).$$

Therefore any maximizer $\hat{\theta}_T \in \operatorname{argmax}_{\theta \in \Theta} \widehat{\text{SR}}_T(\theta)$ must lie in U_ε on $A_T(\varepsilon)$: indeed, if $\hat{\theta}_T \in \Theta \setminus U_\varepsilon$, then $\widehat{\text{SR}}_T(\hat{\theta}_T) \leq \sup_{\theta \in \Theta \setminus U_\varepsilon} \widehat{\text{SR}}_T(\theta) \leq \sup_{\vartheta \in \Theta} \text{SR}(\vartheta) - \delta(\varepsilon)$, while there exists $\theta^* \in \Theta^*$ with $\widehat{\text{SR}}_T(\theta^*) \geq \sup_{\vartheta \in \Theta} \text{SR}(\vartheta) - \delta(\varepsilon)$, contradicting maximality of $\hat{\theta}_T$. Therefore,

$$\mathbb{P}(\hat{\theta}_T \in U_\varepsilon) \rightarrow 1 \quad \text{for every } \varepsilon > 0.$$

Equivalently, $\text{dist}(\hat{\theta}_T, \Theta^*) \xrightarrow{p} 0$. If, in addition, the maximizer is unique, $\Theta^* = \{\theta^*\}$, then $\hat{\theta}_T \rightarrow \theta^*$ in probability, proving part (2).

For (3), fix any $\theta^* \in \Theta^*$. By the triangle inequality,

$$|\widehat{\text{SR}}_T(\hat{\theta}_T) - \text{SR}(\theta^*)| \leq |\widehat{\text{SR}}_T(\hat{\theta}_T) - \text{SR}(\hat{\theta}_T)| + |\text{SR}(\hat{\theta}_T) - \text{SR}(\theta^*)| \leq \sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| + |\text{SR}(\hat{\theta}_T) - \text{SR}(\theta^*)|.$$

The first term converges to 0 in probability by part (1). For the second term, note that SR is continuous on the compact set Θ , hence uniformly continuous. Therefore, for any $\eta > 0$ there exists $\varepsilon > 0$ such that $\|\theta - \vartheta\| < \varepsilon$ implies $|\text{SR}(\theta) - \text{SR}(\vartheta)| < \eta$ for all $\theta, \vartheta \in \Theta$. Since SR is constant on the maximizer set Θ^* , whenever $\text{dist}(\hat{\theta}_T, \Theta^*) < \varepsilon$ we can find $\vartheta_T^* \in \Theta^*$ with $\|\hat{\theta}_T - \vartheta_T^*\| < \varepsilon$ and hence $|\text{SR}(\hat{\theta}_T) - \text{SR}(\theta^*)| = |\text{SR}(\hat{\theta}_T) - \text{SR}(\vartheta_T^*)| < \eta$. Because $\text{dist}(\hat{\theta}_T, \Theta^*) \xrightarrow{p} 0$ by (2), this shows $\text{SR}(\hat{\theta}_T) \rightarrow \text{SR}(\theta^*)$ in probability. Combining the two terms yields

$$|\widehat{\text{SR}}_T(\hat{\theta}_T) - \text{SR}(\theta^*)| \xrightarrow{p} 0,$$

which proves part (3). □

The proof of Theorem 2 can be strengthened to a high-probability bound. In particular, under A1-A4 there exists a complexity term $\mathfrak{R}_T(\Pi)$ for the class $\Pi = \{\pi.(\theta) : \theta \in \Theta\}$ such that, for any $\delta \in (0, 1)$ and some constant $C > 0$, with probability at least $1 - \delta$,

$$\sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| \leq C \left(\mathfrak{R}_T(\Pi) + \sqrt{\frac{-\log \delta}{T}} \right).$$

We make this precise below.

Proof of Theorem 3

Theorem 3 (Oracle inequality for out-of-sample Sharpe ratio)

Let

$$\theta^* \in \operatorname{argmax}_{\theta \in \Theta} \text{SR}(\theta), \quad \hat{\theta}_T \in \operatorname{argmax}_{\theta \in \Theta} \widehat{\text{SR}}_T(\theta).$$

Assume A1-A4 and suppose that $|\pi_t(\theta)| \leq B$ a.s. for all t and $\theta \in \Theta$. Fix a block length $b \in \{1, \dots, \lfloor T/2 \rfloor\}$ and let

$$m = \left\lfloor \frac{T}{2b} \right\rfloor.$$

Let $\mathfrak{R}_{T,b}^{\text{blk}}(\Pi)$ denote the blocked complexity of the portfolio-return class $\Pi = \{\pi.(\theta) : \theta \in \Theta\}$ defined in Lemma 1. Then there exist constants $C_1, C_2 > 0$, depending only on $\underline{\sigma}$, B , and the mixing coefficients, such that for every $\delta \in (0, 1)$,

$$\mathbb{P} \left(\text{SR}(\theta^*) - \text{SR}(\hat{\theta}_T) \leq C_1 \mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + C_2 \left(\frac{b}{T} + \sqrt{\frac{\log(8/\delta)}{m}} \right) \right) \geq 1 - \delta - 2m\beta(b).$$

In particular, if $b = b_{T,\delta}$ is chosen so that $2m\beta(b) \leq \delta/2$, then the same bound holds with probability at least $1 - \delta$.

Recall the portfolio return $\pi_t(\theta) = \sum_{l=1}^{n_t} w_{l,t}^\theta r_{l,t}$, the population moments $\mu(\theta) = \mathbb{E}[\pi_t(\theta)]$, $\nu(\theta) = \mathbb{E}[\pi_t(\theta)^2]$, $\sigma^2(\theta) = \nu(\theta) - \mu(\theta)^2$, and $\text{SR}(\theta) = \mu(\theta)/\sigma(\theta)$. The sample analogs are

$$\hat{\mu}_T(\theta) = \frac{1}{T} \sum_{t=1}^T \pi_t(\theta), \quad \hat{\nu}_T(\theta) = \frac{1}{T} \sum_{t=1}^T \pi_t(\theta)^2,$$

$$\hat{\sigma}_T^2(\theta) = \hat{\nu}_T(\theta) - \hat{\mu}_T(\theta)^2, \quad \widehat{\text{SR}}_T(\theta) = \frac{\hat{\mu}_T(\theta)}{\hat{\sigma}_T(\theta)}.$$

Let $\Pi = \{\pi_t(\theta) : \theta \in \Theta\}$ denote the class of return processes induced by the model.

Lemma 1. Under A1-A3 and $|\pi_t(\theta)| \leq B$ a.s., fix a block length $b \in \{1, \dots, \lfloor T/2 \rfloor\}$ and let

$$m = \lfloor \frac{T}{2b} \rfloor.$$

Define the odd-block sums

$$U_j(\theta) = \sum_{t=(2j-2)b+1}^{(2j-1)b} \pi_t(\theta), \quad j = 1, \dots, m,$$

and let $U_1^\#, \dots, U_m^\#$ be independent Berbee couplings of the odd blocks, with each $U_j^\#$ having the same marginal law as U_j . Let $\epsilon_1, \dots, \epsilon_m$ be i.i.d. Rademacher variables independent of the sample and define the blocked complexity

$$\mathfrak{R}_{T,b}^{\text{blk}}(\Pi) := \frac{1}{T} \mathbb{E} \sup_{\theta \in \Theta} \left| \sum_{j=1}^m \epsilon_j U_j^\#(\theta) \right|.$$

Then there exists a constant $C > 0$, depending only on B and the mixing coefficients, such that for every $\delta \in (0, 1)$, with probability at least $1 - \delta - 2m\beta(b)$,

$$\sup_{\theta \in \Theta} |\hat{\mu}_T(\theta) - \mu(\theta)| \leq 4\mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + CB \left(\frac{b}{T} + \sqrt{\frac{\log(8/\delta)}{m}} \right),$$

and

$$\sup_{\theta \in \Theta} |\hat{\nu}_T(\theta) - \nu(\theta)| \leq 8B\mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + CB^2 \left(\frac{b}{T} + \sqrt{\frac{\log(8/\delta)}{m}} \right).$$

Proof. We prove the bound for the mean. The argument for the second moment is analogous, with an additional contraction step for the map $x \mapsto x^2$.

Define the centered process

$$Z_t(\theta) = \pi_t(\theta) - \mathbb{E}[\pi_t(\theta)].$$

Then

$$\hat{\mu}_T(\theta) - \mu(\theta) = \frac{1}{T} \sum_{t=1}^T Z_t(\theta).$$

Since $|\pi_t(\theta)| \leq B$ a.s., we have $|Z_t(\theta)| \leq 2B$ for all t and θ . Thus it suffices to control the empirical process $\{T^{-1} \sum_{t=1}^T Z_t(\theta)\}_{\theta \in \Theta}$ uniformly in θ .

Step 1: Odd-even block decomposition. Write

$$T = 2mb + r, \quad 0 \leq r < 2b,$$

and define the odd and even block sums

$$\tilde{U}_j(\theta) = \sum_{t=(2j-2)b+1}^{(2j-1)b} Z_t(\theta), \quad \tilde{V}_j(\theta) = \sum_{t=(2j-1)b+1}^{2jb} Z_t(\theta), \quad j=1, \dots, m,$$

together with the remainder

$$R_T(\theta) = \sum_{t=2mb+1}^T Z_t(\theta).$$

Then

$$\sum_{t=1}^T Z_t(\theta) = \sum_{j=1}^m \tilde{U}_j(\theta) + \sum_{j=1}^m \tilde{V}_j(\theta) + R_T(\theta).$$

Because $|Z_t(\theta)| \leq 2B$, the remainder is bounded uniformly by

$$\sup_{\theta \in \Theta} |R_T(\theta)| \leq 2Br \leq 4Bb.$$

Therefore,

$$\sup_{\theta \in \Theta} \left| \frac{1}{T} \sum_{t=1}^T Z_t(\theta) \right| \leq \frac{1}{T} \sup_{\theta \in \Theta} \left| \sum_{j=1}^m \tilde{U}_j(\theta) \right| + \frac{1}{T} \sup_{\theta \in \Theta} \left| \sum_{j=1}^m \tilde{V}_j(\theta) \right| + \frac{4Bb}{T}. \quad (\text{K.7})$$

This step is directly analogous to the decomposition in the original proof, but with odd and even blocks separated by a gap of length b , which is what allows Berbee's coupling to be applied correctly under β -mixing.

Step 2: Coupling the odd and even blocks. Under A1, the process $(Z_t)_{t \geq 1}$ is β -mixing. Since consecutive odd blocks are separated by even blocks of length b , Berbee's coupling theorem implies that there exist independent copies $\tilde{U}_1^\#, \dots, \tilde{U}_m^\#$ such that each $\tilde{U}_j^\#$ has the same marginal law as $\tilde{U}_j$ and

$$\mathbb{P}\left((\tilde{U}_1, \dots, \tilde{U}_m) \neq (\tilde{U}_1^\#, \dots, \tilde{U}_m^\#)\right) \leq m\beta(b).$$

Likewise, there exist independent copies $\tilde{V}_1^\#, \dots, \tilde{V}_m^\#$ such that

$$\mathbb{P}\left((\tilde{V}_1, \dots, \tilde{V}_m) \neq (\tilde{V}_1^\#, \dots, \tilde{V}_m^\#)\right) \leq m\beta(b).$$

This is the key place where the corrected proof differs from the original one: the mixing error appears as a *loss of probability*, rather than as a samplewise additive term inside the deviation bound.

For any threshold $x > 0$, the coupling immediately yields

$$\mathbb{P}\left(\frac{1}{T} \sup_{\theta \in \Theta} \left| \sum_{j=1}^m \tilde{U}_j(\theta) \right| > x\right) \leq \mathbb{P}\left(\frac{1}{T} \sup_{\theta \in \Theta} \left| \sum_{j=1}^m \tilde{U}_j^\#(\theta) \right| > x\right) + m\beta(b),$$

and the analogous inequality holds for the even blocks. Hence it suffices to control the corresponding independent-block processes.

Step 3: Symmetrization for the independent odd blocks. Let

$$A_U^\# = \frac{1}{T} \sup_{\theta \in \Theta} \left| \sum_{j=1}^m \tilde{U}_j^\#(\theta) \right|.$$

Because each $\tilde{U}_j^\#$ has the same law as the centered block sum $\tilde{U}_j$, we have $\mathbb{E}[\tilde{U}_j^\#(\theta)] = 0$ for each j and θ . Introduce i.i.d. Rademacher signs $\epsilon_1, \dots, \epsilon_m$, independent of the data. By the usual symmetrization inequality for independent empirical processes,

$$\mathbb{E} A_U^\# \leq \frac{2}{T} \mathbb{E} \sup_{\theta \in \Theta} \left| \sum_{j=1}^m \epsilon_j \tilde{U}_j^\#(\theta) \right|.$$

Since $\tilde{U}_j^\#(\theta)$ is the centered version of the return block $U_j^\#(\theta)$ and centering does not increase the symmetrized supremum by more than a universal factor, we may absorb that factor into the definition of $\mathfrak{R}_{T,b}^{\text{blk}}(\Pi)$ and conclude that

$$\mathbb{E} A_U^\# \leq 2 \mathfrak{R}_{T,b}^{\text{blk}}(\Pi). \quad (\text{K.8})$$

The same argument applies to the independent even blocks.

Step 4: Concentration around the expectation for the independent blocks. We now show that $A_U^\#$ concentrates around its mean. View

$$A_U^\# = f(\tilde{U}_1^\#, \dots, \tilde{U}_m^\#), \quad f(x_1, \dots, x_m) = \frac{1}{T} \sup_{\theta \in \Theta} \left| \sum_{j=1}^m x_j(\theta) \right|.$$

If two inputs differ only in the j th coordinate, then by the same bounded-differences argument as in the original proof,

$$|f(x) - f(x')| \leq \frac{1}{T} \sup_{\theta \in \Theta} |x_j(\theta) - x'_j(\theta)|.$$

Because each block contains exactly b observations and $|Z_t(\theta)| \leq 2B$, we have

$$\sup_{\theta \in \Theta} |\tilde{U}_j^\#(\theta)| \leq 2Bb,$$

so changing one block can alter $A_U^\#$ by at most

$$\frac{1}{T} \cdot 4Bb \leq \frac{2B}{m},$$

since $T \geq 2mb$. Therefore McDiarmid's inequality implies that, with probability at least $1 - \delta/4$,

$$A_U^\# \leq \mathbb{E} A_U^\# + CB \sqrt{\frac{\log(8/\delta)}{m}}, \quad (\text{K.9})$$

for a universal constant $C > 0$. The same estimate holds for the independent even-block process

$$A_V^\# = \frac{1}{T} \sup_{\theta \in \Theta} \left| \sum_{j=1}^m \tilde{V}_j^\#(\theta) \right|.$$

Combining (K.8) and (K.9), and applying the same argument to the even blocks, we obtain that

with probability at least $1 - \delta/2$,

$$A_U^\# + A_V^\# \leq 4\mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + CB\sqrt{\frac{\log(8/\delta)}{m}}.$$

After intersecting with the two coupling events from Step 2 and using a union bound, the same inequality holds for the original odd and even block sums with probability at least $1 - \delta - 2m\beta(b)$.

Step 5: Conclusion for the mean. Insert the odd/even block bound into (K.7). This shows that, with probability at least $1 - \delta - 2m\beta(b)$,

$$\sup_{\theta \in \Theta} |\hat{\mu}_T(\theta) - \mu(\theta)| \leq 4\mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + CB\left(\frac{b}{T} + \sqrt{\frac{\log(8/\delta)}{m}}\right).$$

This proves the first bound.

Step 6: Second moment bound via contraction. Now define

$$Y_t(\theta) = \pi_t(\theta)^2 - \mathbb{E}[\pi_t(\theta)^2].$$

On $[-B, B]$, the map $h(x) = x^2$ is $2B$ -Lipschitz, since

$$|h(x) - h(y)| = |x - y||x + y| \leq 2B|x - y|.$$

Hence the symmetrized complexity of the squared-return class is at most a factor $2B$ larger than that of the original return class. Repeating Steps 1-5 with $Y_t(\theta)$ in place of $Z_t(\theta)$ and using the contraction inequality therefore yields

$$\sup_{\theta \in \Theta} |\hat{\nu}_T(\theta) - \nu(\theta)| \leq 8B\mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + CB^2\left(\frac{b}{T} + \sqrt{\frac{\log(8/\delta)}{m}}\right)$$

on the same event, after adjusting the constant C .

Combining the two displays yields the lemma. $\square$

Lemma 2. *Under A3 (signals are uniformly Lipschitz in θ) and the differentiable soft rank/mask mapping with temperatures $\tau_{\text{rank}}, \tau_{\text{mask}} > 0$, there exists $L_\pi < \infty$ such that, uniformly in t ,*

$$|\pi_t(\theta) - \pi_t(\theta')| \leq L_\pi \|\theta - \theta'\| \quad \text{for all } \theta, \theta' \in \Theta.$$

In particular, one may take

$$L_\pi \leq R_{\max} L_{\text{sort}} L_{\text{net}}, \quad \text{with } R_{\max} = \sup_{t,l} |r_{l,t}|,$$

where L_{net} is the Lipschitz envelope of $\theta \mapsto s_{l,t}^\theta$ (A3) and L_{sort} is an explicit constant depending only on $(\tau_{\text{rank}}, \tau_{\text{mask}}, q, n_{\max})$ defined below.

Proof. By A3, there exists $L_{\text{net}} < \infty$ such that for all (l, t) and all θ, θ' ,

$$|s_{l,t}^\theta - s_{l,t}^{\theta'}| \leq L_{\text{net}} \|\theta - \theta'\|. \quad (\text{K.10})$$

The soft pairwise preference is $S_{l,j,t} = \sigma((s_{j,t}^\theta - s_{l,t}^\theta)/\tau_{\text{rank}})$ with $\sigma'(x) = \sigma(x)(1 - \sigma(x)) \in [0, \frac{1}{4}]$.

By the chain rule,

$$\left| \frac{\partial S_{l,j,t}}{\partial s_{l,t}^\theta} \right| = \frac{1}{\tau_{\text{rank}}} \sigma' \left(\frac{s_{j,t}^\theta - s_{l,t}^\theta}{\tau_{\text{rank}}} \right) \leq \frac{1}{4\tau_{\text{rank}}} =: C_{\text{pair}},$$

and the same bound holds for $\partial S_{l,j,t} / \partial s_{j,t}^\theta$ (all other signal partials are zero). The soft rank is $\text{rank}_{l,t}^\tau = 1 + \sum_{j \neq l} S_{l,j,t}$. Hence, for any fixed asset index u ,

$$\left| \frac{\partial \text{rank}_{l,t}^\tau}{\partial s_{u,t}^\theta} \right| \leq \sum_{j \neq l} \left| \frac{\partial S_{l,j,t}}{\partial s_{u,t}^\theta} \right| \leq (n_t - 1) C_{\text{pair}} \leq (n_{\max} - 1) C_{\text{pair}} =: C_{\text{rank}}.$$

The top-mask is $m_{l,t}^H = \sigma((k_t + 0.5 - \text{rank}_{l,t}^\tau) / \tau_{\text{mask}})$, with derivative

$$\left| \frac{\partial m_{l,t}^H}{\partial \text{rank}_{l,t}^\tau} \right| = \frac{1}{\tau_{\text{mask}}} \sigma' \left(\frac{k_t + 0.5 - \text{rank}_{l,t}^\tau}{\tau_{\text{mask}}} \right) \leq \frac{1}{4\tau_{\text{mask}}} =: C_{\text{mask}}.$$

By the chain rule and the bound from the previous step, for any u ,

$$\left| \frac{\partial m_{l,t}^H}{\partial s_{u,t}^\theta} \right| \leq C_{\text{mask}} \left| \frac{\partial \text{rank}_{l,t}^\tau}{\partial s_{u,t}^\theta} \right| \leq C_{\text{mask}} C_{\text{rank}}.$$

The same bound holds for the bottom mask $m_{l,t}^L$. Define the (soft) leg sums $Z_t^H = \sum_j m_{j,t}^H$ and $Z_t^L = \sum_j m_{j,t}^L$. The normalized leg weights are

$$w_{l,t}^{\text{long}} = \frac{m_{l,t}^H}{Z_t^H}, \quad w_{l,t}^{\text{short}} = \frac{m_{l,t}^L}{Z_t^L}.$$

By A2, $2 \leq n_t \leq n_{\max}$ and the masks are strictly positive (logistic), so there exists a deterministic lower bound $\underline{Z} > 0$ such that

$$Z_t^H \geq \underline{Z}, \quad Z_t^L \geq \underline{Z} \quad \text{for all } t, \quad (\text{K.11})$$

uniformly over the sample. One convenient choice is

$$k_{\min} := \max\{1, \lfloor q n_{\min} \rfloor\}, \quad \underline{Z} = k_{\min} \sigma(0.5 / \tau_{\text{mask}}).$$

For the long leg, differentiate $w_{l,t}^{\text{long}}$ with respect to a single mask $m_{u,t}^H$:

$$\frac{\partial w_{l,t}^{\text{long}}}{\partial m_{u,t}^H} = \frac{\mathbf{1}\{l=u\}}{Z_t^H} - \frac{m_{l,t}^H}{(Z_t^H)^2},$$

so that, using $m_{l,t}^H \leq 1$ and (K.11),

$$\sum_{l=1}^{n_t} \left| \frac{\partial w_{l,t}^{\text{long}}}{\partial m_{u,t}^H} \right| \leq \frac{1}{Z_t^H} + \frac{\sum_l m_{l,t}^H}{(Z_t^H)^2} = \frac{1}{Z_t^H} + \frac{Z_t^H}{(Z_t^H)^2} \leq \frac{2}{\underline{Z}}.$$

Summing over all masks (chain rule) and using $n_t \leq n_{\max}$ gives

$$\sum_{l=1}^{n_t} \left| \frac{\partial w_{l,t}^{\text{long}}}{\partial s_{u,t}^\theta} \right| \leq \left(\sum_{j=1}^{n_t} \left| \frac{\partial m_{j,t}^H}{\partial s_{u,t}^\theta} \right| \right) \cdot \sup_u \sum_l \left| \frac{\partial w_{l,t}^{\text{long}}}{\partial m_{u,t}^H} \right| \leq n_{\max} (C_{\text{mask}} C_{\text{rank}}) \cdot \frac{2}{\underline{Z}}. \quad (\text{K.12})$$

The same bound holds for the short leg:

$$\sum_{l=1}^{n_t} \left| \frac{\partial w_{l,t}^{\text{short}}}{\partial s_{u,t}^{\theta}} \right| \leq n_{\max} (C_{\text{mask}} C_{\text{rank}}) \cdot \frac{2}{\underline{Z}}. \quad (\text{K.13})$$

The long-short weights are $w_{l,t} = w_{l,t}^{\text{long}} - w_{l,t}^{\text{short}}$. Combining (K.12) and (K.13) yields

$$\sum_{l=1}^{n_t} \left| \frac{\partial w_{l,t}}{\partial s_{u,t}^{\theta}} \right| \leq \frac{4n_{\max}}{\underline{Z}} C_{\text{mask}} C_{\text{rank}} = \frac{4n_{\max}}{\underline{Z}} \cdot \frac{1}{4\tau_{\text{mask}}} \cdot \frac{n_{\max}-1}{4\tau_{\text{rank}}} \leq \frac{n_{\max}^2}{4\underline{Z}} \cdot \frac{1}{\tau_{\text{mask}} \tau_{\text{rank}}}.$$

Hence, for any two signal vectors $s_{\cdot,t}^{\theta}$ and $s_{\cdot,t}^{\theta'}$,

$$\sum_{l=1}^{n_t} |w_{l,t}(\theta) - w_{l,t}(\theta')| \leq L_{\text{sort}} \cdot \max_u |s_{u,t}^{\theta} - s_{u,t}^{\theta'}|, \quad L_{\text{sort}} = \frac{n_{\max}^2}{4\underline{Z}} \cdot \frac{1}{\tau_{\text{mask}} \tau_{\text{rank}}}. \quad (\text{K.14})$$

By (K.10), $\max_u |s_{u,t}^{\theta} - s_{u,t}^{\theta'}| \leq L_{\text{net}} \|\theta - \theta'\|$. Plugging into (K.14) gives

$$\sum_{l=1}^{n_t} |w_{l,t}(\theta) - w_{l,t}(\theta')| \leq L_{\text{sort}} L_{\text{net}} \|\theta - \theta'\|.$$

Finally, the managed return difference satisfies

$$|\pi_t(\theta) - \pi_t(\theta')| = \left| \sum_{l=1}^{n_t} (w_{l,t}(\theta) - w_{l,t}(\theta')) r_{l,t} \right| \leq \|r_{\cdot,t}\|_{\infty} \sum_{l=1}^{n_t} |w_{l,t}(\theta) - w_{l,t}(\theta')| \leq R_{\max} L_{\text{sort}} L_{\text{net}} \|\theta - \theta'\|.$$

This proves the claim with $L_{\pi} = R_{\max} L_{\text{sort}} L_{\text{net}}$. $\square$

Lemma 3. Let $g(\mu, \nu) = \mu / \sqrt{\nu - \mu^2}$ be defined on

$$\mathcal{D}_{\underline{\sigma}, B} = \{(\mu, \nu) \in \mathbb{R}^2 : \nu - \mu^2 \geq \underline{\sigma}^2, \quad |\mu| \leq B, \quad 0 \leq \nu \leq B^2\}.$$

Under A2-A4, fix $b \in \{1, \dots, \lfloor T/2 \rfloor\}$ and $m = \lfloor T/(2b) \rfloor$. Then for any $\delta \in (0, 1)$ there exists a constant $C_{\text{SR}} = C_{\text{SR}}(\underline{\sigma}, B) > 0$ such that, with probability at least $1 - \delta - 2m\beta(b)$,

$$\sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| \leq C_{\text{SR}} \left(\sup_{\theta \in \Theta} |\hat{\mu}_T(\theta) - \mu(\theta)| + \sup_{\theta \in \Theta} |\hat{\nu}_T(\theta) - \nu(\theta)| \right).$$

Consequently, by Lemma 1, with probability at least $1 - \delta - 2m\beta(b)$,

$$\sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| \leq C_{\text{SR}} \left(\mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + \frac{b}{T} + \sqrt{\frac{\log(8/\delta)}{m}} \right).$$

Proof. We first compute the gradient of g and bound it uniformly on a set that contains, with high probability, the line segment between the population moments $(\mu(\theta), \nu(\theta))$ and the empirical moments $(\hat{\mu}_T(\theta), \hat{\nu}_T(\theta))$, uniformly over θ .

Write $\sigma^2 = \nu - \mu^2$ and $\sigma = \sqrt{\nu - \mu^2}$. Then

$$g(\mu, \nu) = \frac{\mu}{\sigma} = \mu(\nu - \mu^2)^{-1/2}.$$

By differentiation,

$$\frac{\partial g}{\partial \mu} = \frac{1}{\sigma} + \frac{\mu^2}{\sigma^3}, \quad \frac{\partial g}{\partial \nu} = -\frac{\mu}{2\sigma^3}.$$

Hence, on any set with $\sigma \geq s > 0$ and $|\mu| \leq B$,

$$\left| \frac{\partial g}{\partial \mu} \right| \leq \frac{1}{s} + \frac{B^2}{s^3}, \quad \left| \frac{\partial g}{\partial \nu} \right| \leq \frac{B}{2s^3}. \quad (\text{K.15})$$

Thus

$$\|\nabla g\|_2 \leq \frac{1}{s} + \frac{B^2}{s^3} + \frac{B}{2s^3}.$$

Fix $\theta \in \Theta$ and consider the line segment between the population pair $x = (\mu, \nu)$ and empirical pair $\hat{x} = (\hat{\mu}_T, \hat{\nu}_T)$:

$$x_t = (1-t)x + t\hat{x} = (\mu_t, \nu_t), \quad t \in [0, 1].$$

Since $|\pi_t(\theta)| \leq B$ almost surely, we have deterministically $|\hat{\mu}_T(\theta)| \leq B$ and $0 \leq \hat{\nu}_T(\theta) \leq B^2$, and therefore $|\mu_t| \leq B$ and $0 \leq \nu_t \leq B^2$ for all $t \in [0, 1]$.

Define the uniform deviations

$$\Delta_\mu = \sup_{\vartheta \in \Theta} |\hat{\mu}_T(\vartheta) - \mu(\vartheta)|, \quad \Delta_\nu = \sup_{\vartheta \in \Theta} |\hat{\nu}_T(\vartheta) - \nu(\vartheta)|.$$

Then, for every θ and every $t \in [0, 1]$,

$$\nu_t - \mu_t^2 \geq \sigma^2(\theta) - \Delta_\nu - \Delta_\mu(2B + \Delta_\mu).$$

Indeed, this is exactly the same segment argument as in the current proof. By A4, $\sigma^2(\theta) \geq \underline{\sigma}^2$ uniformly over θ . Hence on the event

$$\mathcal{A}_T(\eta) = \{\Delta_\mu + \Delta_\nu \leq \eta\},$$

we have

$$\nu_t - \mu_t^2 \geq \underline{\sigma}^2 - \eta(2B + 1 + \eta) \quad \text{for all } t \in [0, 1], \text{ uniformly in } \theta.$$

Choose $\eta > 0$ small enough so that $\eta(2B + 1 + \eta) \leq \underline{\sigma}^2/2$. Then on $\mathcal{A}_T(\eta)$,

$$\nu_t - \mu_t^2 \geq \underline{\sigma}^2/2 \quad \text{for all } t \in [0, 1], \text{ uniformly in } \theta. \quad (\text{K.16})$$

By Lemma 1, the event $\mathcal{A}_T(\eta)$ has probability at least $1 - \delta - 2m\beta(b)$ once

$$\eta \asymp \mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + \frac{b}{T} + \sqrt{\frac{\log(8/\delta)}{m}}.$$

Now apply the multivariate mean value theorem. For each θ , there exists $t^* \in (0, 1)$ such that

$$|\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| = |g(\hat{\mu}_T, \hat{\nu}_T) - g(\mu, \nu)| \leq \|\nabla g(\mu_{t^*}, \nu_{t^*})\|_2 \|(\hat{\mu}_T - \mu, \hat{\nu}_T - \nu)\|_2.$$

By (K.16), the intermediate point (μ_{t^*}, ν_{t^*}) lies in a domain where the variance is bounded below by $\underline{\sigma}/\sqrt{2}$. Therefore, using (K.15) with $s = \underline{\sigma}/\sqrt{2}$, we obtain a finite constant $C_{\text{SR}}(\underline{\sigma}, B)$ such that

$$\|\nabla g(\mu_{t^*}, \nu_{t^*})\|_2 \leq C_{\text{SR}}(\underline{\sigma}, B)$$

uniformly in θ . Since

$$\|(\hat{\mu}_T - \mu, \hat{\nu}_T - \nu)\|_2 \leq |\hat{\mu}_T - \mu| + |\hat{\nu}_T - \nu|,$$

we conclude that on $\mathcal{A}_T(\eta)$,

$$|\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| \leq C_{\text{SR}}(|\hat{\mu}_T(\theta) - \mu(\theta)| + |\hat{\nu}_T(\theta) - \nu(\theta)|)$$

uniformly over $\theta \in \Theta$. Taking the supremum over θ proves the first display.

The second display follows by substituting the moment bounds from Lemma 1 into the right-hand side and absorbing the resulting numerical factors into the same constant C_{SR} . $\square$

Proposition 1. *Let $\theta^* \in \arg\max_{\theta \in \Theta} \text{SR}(\theta)$ and let $\hat{\theta}_T \in \arg\max_{\theta \in \Theta} \widehat{\text{SR}}_T(\theta)$ be any empirical maximizer. Then*

$$\text{SR}(\theta^*) - \text{SR}(\hat{\theta}_T) \leq 2 \sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)|. \quad (\text{K.17})$$

Consequently, under A1-A4, for every $\delta \in (0, 1)$,

$$\mathbb{P}\left(\text{SR}(\theta^*) - \text{SR}(\hat{\theta}_T) \leq 2C_{\text{SR}}\left(\mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + \frac{b}{T} + \sqrt{\frac{\log(8/\delta)}{m}}\right)\right) \geq 1 - \delta - 2m\beta(b),$$

where $m = \lfloor T/(2b) \rfloor$ and $C_{\text{SR}} = C_{\text{SR}}(\underline{\sigma}, B)$ is the constant from Lemma 3.

Proof. Fix any selections $\theta^* \in \arg\max_{\theta} \text{SR}(\theta)$ and $\hat{\theta}_T \in \arg\max_{\theta} \widehat{\text{SR}}_T(\theta)$. We begin with the same deterministic decomposition used in the current proof:

$$\text{SR}(\theta^*) - \text{SR}(\hat{\theta}_T) = (\text{SR}(\theta^*) - \widehat{\text{SR}}_T(\theta^*)) + (\widehat{\text{SR}}_T(\theta^*) - \widehat{\text{SR}}_T(\hat{\theta}_T)) + (\widehat{\text{SR}}_T(\hat{\theta}_T) - \text{SR}(\hat{\theta}_T)).$$

By the optimality of $\hat{\theta}_T$ for the empirical criterion,

$$\widehat{\text{SR}}_T(\theta^*) - \widehat{\text{SR}}_T(\hat{\theta}_T) \leq 0.$$

Therefore,

$$\text{SR}(\theta^*) - \text{SR}(\hat{\theta}_T) \leq |\text{SR}(\theta^*) - \widehat{\text{SR}}_T(\theta^*)| + |\widehat{\text{SR}}_T(\hat{\theta}_T) - \text{SR}(\hat{\theta}_T)|,$$

and each of the two absolute values is bounded by the uniform deviation $\sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)|$. This proves (K.17).

Now apply Lemma 3. On an event of probability at least $1 - \delta - 2m\beta(b)$,

$$\sup_{\theta \in \Theta} |\widehat{\text{SR}}_T(\theta) - \text{SR}(\theta)| \leq C_{\text{SR}}\left(\mathfrak{R}_{T,b}^{\text{blk}}(\Pi) + \frac{b}{T} + \sqrt{\frac{\log(8/\delta)}{m}}\right).$$

Combining this with (K.17) yields the stated oracle inequality. $\square$

Remark 1. *Lemma 2 implies that the managed-return class is a Lipschitz parametric class in θ . This is the starting point for deriving further upper bounds on the blocked complexity $\mathfrak{R}_{T,b}^{\text{blk}}(\Pi)$ under additional entropy or covering assumptions on Θ .*

K.2 Additional results

Complementing the properties presented in Section 7, we provide additional theoretical justifications for our approach. First, we provide a simple expected-utility foundation for the Sharpe ratio objective, which is used to estimate the AlphaGlass portfolio rule, clarifying the relation between the

baseline objective and a standard mean-variance investment problem. The key point is that, in the presence of a risk-free asset and under standard mean-variance preferences (or expected utility with elliptical returns), the investor's portfolio choice problem reduces to maximizing the unconditional Sharpe ratio of the resulting portfolio return. Risk aversion determines only the overall leverage under the rule, but it does not affect which rule is optimal within a given class.

We work in a one-period setting. Let $R \in \mathbb{R}^N$ denote the vector of excess returns on N test assets over the risk-free rate. A portfolio rule π maps time- t information (such as firm characteristics) into portfolio weights $w(\pi) \in \mathbb{R}^N$. The associated portfolio excess return is

$$r_p(\pi) \equiv w(\pi)^\top R,$$

with mean and variance

$$\mu(\pi) \equiv \mathbb{E}[r_p(\pi)], \quad \sigma^2(\pi) \equiv \text{Var}(r_p(\pi)).$$

The investor can also invest in the risk-free asset, yielding a gross return of $1 + r_f$. Given a portfolio rule π , she chooses a scalar position $\alpha \in \mathbb{R}$ that scales the zero-cost portfolio, so that final wealth (normalized by initial wealth) is

$$W_1(\alpha, \pi) = (1 + r_f) + \alpha r_p(\pi).$$

We assume mean-variance preferences,

$$U(\alpha, \pi) = \mathbb{E}[W_1(\alpha, \pi)] - \frac{\gamma}{2} \text{Var}(W_1(\alpha, \pi)),$$

for some risk-aversion parameter $\gamma > 0$.²⁰ Let Π denote the class of admissible portfolio rules. In our empirical implementation, Π is the class of AlphaGlass GAM-like rules, parameterized by θ and mapping characteristics into weights via the differentiable rank-mask architecture described above. The following theorem shows that the investor's joint choice of a portfolio rule $\pi \in \Pi$ and scale $\alpha \in \mathbb{R}$ is equivalent to choosing the rule that maximizes the Sharpe ratio of the portfolio return $r_p(\pi)$.

Theorem 4 (Optimal portfolio rule as Sharpe maximizer)

Let R be a vector of excess returns and let Π be a class of portfolio rules π , each inducing a zero-cost portfolio excess return $r_p(\pi)$ with mean $\mu(\pi)$ and variance $\sigma^2(\pi)$. Suppose the investor has mean-variance preferences

$$U(\alpha, \pi) = \mathbb{E}[W_1(\alpha, \pi)] - \frac{\gamma}{2} \text{Var}(W_1(\alpha, \pi)),$$

where $W_1(\alpha, \pi) = (1 + r_f) + \alpha r_p(\pi)$, and $\gamma > 0$ is fixed. Then the problem

$$\max_{\pi \in \Pi, \alpha \in \mathbb{R}} U(\alpha, \pi)$$

is equivalent to

$$\max_{\pi \in \Pi} \text{SR}(\pi), \quad \text{SR}(\pi) \equiv \frac{\mu(\pi)}{\sigma(\pi)},$$

²⁰Under elliptical returns (for example, multivariate normality) this mean-variance representation is equivalent to expected utility maximization for any increasing concave Bernoulli utility, so the analysis below can also be interpreted as a reduced-form representation of more general preferences.

in the sense that any portfolio rule $\pi^* \in \Pi$ that solves

$$\pi^* \in \operatorname{argmax}_{\pi \in \Pi} \text{SR}(\pi)$$

admits a scale $\alpha^*(\pi^*)$ such that the pair $(\pi^*, \alpha^*(\pi^*))$ solves the original utility maximization problem. The optimal scale for a given π is

$$\alpha^*(\pi) = \frac{\mu(\pi)}{\gamma \sigma^2(\pi)},$$

so risk aversion γ only affects the leverage $\alpha^*(\pi)$, not which portfolio rule π is optimal within Π .

Economically, Theorem 4 implies that, once we fix a class Π of portfolio rules (in our case, interpretable GAM-style rules with sparsity and pairwise interactions), the “right” object to estimate for a mean-variance investor with access to the risk-free asset is the portfolio rule that maximizes the unconditional Sharpe ratio of the induced portfolio return. This provides a decision-theoretic rationale for our baseline Sharpe ratio specification. Note, however, that this equivalence result concerns the choice of the portfolio rule in a benchmark problem with separate leverage choice. Our mean-variance exercise in Section 6 serves a different purpose: it studies whether AlphaGlass continues to perform well when investor preferences are embedded directly in the estimation objective through the parameter γ . In this sense, the Sharpe-ratio and mean-variance specifications are complementary formulations.

Second, we provide an economic and decision-theoretic rationale for our choice of rank-based mask weights, rather than softmax-based weights, as the final link function from scores to portfolio weights. The key idea is that the output of the AlphaGlass signal is only identified up to strictly increasing transformations: if we compose the last layer of the network with a strictly increasing scalar function and adjust earlier layers accordingly, the ranking of assets by score is preserved and, in large samples, portfolio performance is unchanged. Economically, this means that the *ordering* induced by the score carries information (which stocks are better or worse), but the *cardinal scale* of the score is not uniquely meaningful.

A natural robustness requirement is that the portfolio decision be *ordinarily invariant*: if we apply any strictly increasing transformation to the scores, the resulting portfolio should not change. At the same time, we require that the rule treat assets symmetrically and only react to the ranking, not to asset labels. Formally, this symmetry is captured by permutation equivariance. We show below that these two requirements together imply that the weighting rule must be rank-based: there exists a function f on the set of ranks such that the weight assigned to asset i depends only on its rank in the cross-section. Our top/bottom mask is a particularly transparent and simple choice of such a rank-based rule.

By contrast, softmax-based weighting schemes implicitly treat the scores as economically meaningful in levels and are sensitive to arbitrary monotone rescalings of the scores. Changing the “temperature” of the softmax or rescaling the signal can yield very different portfolios, even when the ranking of scores remains unchanged. From a decision-theoretic perspective, when scores are identified only up to monotone transformations, softmax is a fragile choice. Rank-based masks, by construction, are robust to such transformations and therefore align better with the economic content

of the signal.

We now formalize these ideas in a simple static cross-sectional setting. Fix a cross-section with $n \geq 2$ assets. Let

$$s = (s_1, \dots, s_n)^\top \in \mathbb{R}^n$$

denote the vector of scores produced by the signal in a given period. A portfolio rule is a mapping

$$W : \mathbb{R}^n \rightarrow \mathbb{R}^n, \quad s \mapsto W(s) = (W_1(s), \dots, W_n(s))^\top,$$

where $W_i(s)$ is the portfolio weight assigned to asset i when the score vector is s .

We impose two structural properties that capture symmetry and robustness to monotone transformations. For any permutation σ of $\{1, \dots, n\}$, define the permuted score vector $(\sigma \cdot s) \in \mathbb{R}^n$ by

$$(\sigma \cdot s)_i \equiv s_{\sigma^{-1}(i)}.$$

We say that W is permutation equivariant if

$$W(\sigma \cdot s) = \sigma \cdot W(s)$$

for all $s \in \mathbb{R}^n$ and all permutations σ . This means that relabeling the assets only relabels the weights.

Let $\phi : \mathbb{R} \rightarrow \mathbb{R}$ be any strictly increasing function. We define the component-wise composition $\phi \circ s \in \mathbb{R}^n$ by $(\phi \circ s)_i = \phi(s_i)$. We say that W is ordinally invariant if

$$W(\phi \circ s) = W(s)$$

for all score vectors $s \in \mathbb{R}^n$ and all strictly increasing ϕ . This expresses the requirement that the portfolio decision depends only on the *ordering* of the scores, not on an arbitrary monotone scaling or transformation. For each s with distinct entries (no ties), let $\text{rank}_i(s) \in \{1, \dots, n\}$ denote the rank of asset i in the cross-section of scores, where rank 1 is the lowest score and rank n is the highest. We will show that permutation equivariance and ordinal invariance force $W(s)$ to depend only on these ranks. We first show that any continuous portfolio rule that is permutation equivariant and ordinally invariant must be rank-based.

Theorem 5 (Ordinally invariant portfolio rules are rank-based)

Let $W : \mathbb{R}^n \rightarrow \mathbb{R}^n$ be a portfolio rule that is permutation equivariant and ordinally invariant. Suppose furthermore that W is continuous and that we restrict attention to score vectors $s \in \mathbb{R}^n$ with no ties, $s_i \neq s_j$ for $i \neq j$. Then there exists a function $f : \{1, \dots, n\} \rightarrow \mathbb{R}$ such that, for every such s and every asset i ,

$$W_i(s) = f(\text{rank}_i(s)).$$

In particular, $W(s)$ depends only on the cross-sectional ranking of scores, not on their levels.

Proof. Take two score vectors $s, s' \in \mathbb{R}^n$ without ties and the same ranking. That is, there exists a permutation π of $\{1, \dots, n\}$ such that

$$s_{\pi(1)} < s_{\pi(2)} < \dots < s_{\pi(n)}, \quad s'_{\pi(1)} < s'_{\pi(2)} < \dots < s'_{\pi(n)}.$$

Define the sorted score sequences

$$\tilde{s}_k \equiv s_{\pi(k)}, \quad \tilde{s}'_k \equiv s'_{\pi(k)}, \quad k = 1, \dots, n,$$

so that $\tilde{s}_1 < \dots < \tilde{s}_n$ and $\tilde{s}'_1 < \dots < \tilde{s}'_n$.

Because both sequences $(\tilde{s}_k)$ and $(\tilde{s}'_k)$ are strictly increasing, we can construct a strictly increasing function $\phi: \mathbb{R} \rightarrow \mathbb{R}$ such that

$$\phi(\tilde{s}_k) = \tilde{s}'_k \quad \text{for all } k = 1, \dots, n.$$

One convenient construction is to define ϕ piecewise linearly through the points $(\tilde{s}_k, \tilde{s}'_k)$ and extend it linearly outside $[\tilde{s}_1, \tilde{s}_n]$.

For each asset i , let $r_i = \text{rank}_i(s)$, so that $\pi(r_i) = i$ and $s_i = \tilde{s}_{r_i}$. By construction,

$$\phi(s_i) = \phi(\tilde{s}_{r_i}) = \tilde{s}'_{r_i} = s'_i.$$

In vector notation, $\phi \circ s = s'$. By ordinal invariance,

$$W(s') = W(\phi \circ s) = W(s).$$

Thus, whenever s and s' have the same ranking, we must have $W(s) = W(s')$. Therefore, $W(s)$ depends only on the ranking pattern of s .

Let $v^0 = (1, 2, \dots, n)^\top$ denote the canonical strictly increasing vector. Its ranking is

$$1 < 2 < \dots < n.$$

Define $w^0 \equiv W(v^0) \in \mathbb{R}^n$.

Now take any s with no ties. Let π_s be the permutation that sorts s in increasing order, so that

$$s_{\pi_s(1)} < s_{\pi_s(2)} < \dots < s_{\pi_s(n)}.$$

Define the sorted vector $\tilde{s} \in \mathbb{R}^n$ by $\tilde{s}_k \equiv s_{\pi_s(k)}$. Then s and $\tilde{s}$ have the same ranking. By Step 1, $W(s) = W(\tilde{s})$.

As before, because $(1 < 2 < \dots < n)$ and $(\tilde{s}_1 < \dots < \tilde{s}_n)$ are strictly increasing sequences, there exists a strictly increasing function ϕ such that

$$\phi(k) = \tilde{s}_k \quad \text{for } k = 1, \dots, n.$$

Thus $\phi \circ v^0 = \tilde{s}$, so by ordinal invariance,

$$W(\tilde{s}) = W(\phi \circ v^0) = W(v^0) = w^0.$$

Combining, we obtain

$$W(s) = W(\tilde{s}) = w^0.$$

This identity holds in the sorted coordinate system.

The relationship between s and its sorted version $\tilde{s}$ is $s = \pi_s \cdot \tilde{s}$. By permutation equivariance,

$$W(s) = W(\pi_s \cdot \tilde{s}) = \pi_s \cdot W(\tilde{s}) = \pi_s \cdot w^0.$$

Taking the i -th component, we obtain

$$W_i(s) = (\pi_s \cdot w^0)_i = w_{\pi_s^{-1}(i)}^0.$$

But $\pi_s^{-1}(i)$ is precisely the rank of asset i in the cross-section of scores:

$$\text{rank}_i(s) = \pi_s^{-1}(i).$$

Therefore

$$W_i(s) = w_{\text{rank}_i(s)}^0.$$

Define $f: \{1, \dots, n\} \rightarrow \mathbb{R}$ by $f(k) = w_k^0$. Then for all i and all score vectors s without ties,

$$W_i(s) = f(\text{rank}_i(s)),$$

which is exactly the desired rank-based representation. $\square$

Theorem 5 formalizes the intuitive statement that, under the natural symmetry and robustness properties of permutation equivariance and ordinal invariance, the portfolio rule cannot depend on the cardinal scale of scores. It can only depend on the ordering of scores, and, therefore, must be a function of ranks. In our application, the specific choice of f encodes how aggressively the portfolio loads on high- versus low-score assets and whether the middle ranks receive zero weight.

We can now position our rank-based top/bottom mask and the softmax rule in the framework of Theorem 5.

Our mask-based rule is defined as follows. Fix an integer $K \in \{1, \dots, \lfloor n/2 \rfloor\}$. For a score vector s without ties, we set

$$w_i^{\text{rk}}(s) = \begin{cases} \frac{1}{K}, & \text{if } \text{rank}_i(s) > n - K \quad (\text{top } K \text{ assets}), \\ -\frac{1}{K}, & \text{if } \text{rank}_i(s) \leq K \quad (\text{bottom } K \text{ assets}), \\ 0, & \text{otherwise.} \end{cases}$$

This is exactly of the form $W_i(s) = f(\text{rank}_i(s))$ with

$$f(k) = \begin{cases} -1/K, & k \leq K, \\ 0, & K < k \leq n - K, \\ 1/K, & k > n - K. \end{cases}$$

By construction, this rule is permutation equivariant and ordinally invariant, so it satisfies the structural requirements of Theorem 5. In words, it is a simple, symmetric rule that only cares about which assets are in the top K and bottom K of the score distribution, while equal-weighting within each leg to avoid overreacting to small cardinal differences in scores.

In contrast, the softmax rule maps scores into weights through exponential functions of the levels:

$$p_i(s) = \frac{\exp(\lambda s_i)}{\sum_{j=1}^n \exp(\lambda s_j)}, \quad w_i^{\text{sm}}(s) = p_i(s) - \frac{1}{n},$$

where $\lambda > 0$ is a temperature parameter. This rule is permutation equivariant, but it is *not* ordinally

invariant: rescaling the scores by a strictly increasing function changes the weights in general. We formalize this in the next lemma.

Lemma 4 (Softmax is not ordinally invariant). *Fix $n \geq 2$ and $\lambda > 0$. Define the softmax-based zero-cost weights*

$$w_i^{\text{sm}}(s) = \frac{\exp(\lambda s_i)}{\sum_{j=1}^n \exp(\lambda s_j)} - \frac{1}{n}.$$

Let $s \in \mathbb{R}^n$ be any non-constant score vector (i.e., $s_i \neq s_j$ for some $i \neq j$). Then there exists a strictly increasing function $\phi: \mathbb{R} \rightarrow \mathbb{R}$ such that

$$w^{\text{sm}}(\phi \circ s) \neq w^{\text{sm}}(s).$$

In particular, the softmax rule is not ordinally invariant.

Proof. We provide a simple argument based on linear rescaling. Consider a strictly increasing function of the form $\phi(x) = ax$ with $a > 0$. For any s and any $a > 0$,

$$w_i^{\text{sm}}(\phi \circ s) = w_i^{\text{sm}}(as) = \frac{\exp(\lambda a s_i)}{\sum_{j=1}^n \exp(\lambda a s_j)} - \frac{1}{n}.$$

Suppose, for contradiction, that $w^{\text{sm}}(as) = w^{\text{sm}}(s)$ for all $a > 0$. Then the softmax probabilities would also be invariant:

$$\frac{\exp(\lambda a s_i)}{\sum_{j=1}^n \exp(\lambda a s_j)} = \frac{\exp(\lambda s_i)}{\sum_{j=1}^n \exp(\lambda s_j)} \quad \text{for all } i \text{ and all } a > 0.$$

Consider two indices $i \neq j$ such that $s_i \neq s_j$ (which exist because s is non-constant). Define the probability ratio

$$r_{ij}(s; \lambda) = \frac{p_i(s)}{p_j(s)} = \frac{\exp(\lambda s_i)}{\exp(\lambda s_j)} = \exp(\lambda(s_i - s_j)).$$

Under the supposed invariance, we must have

$$r_{ij}(as; \lambda) = r_{ij}(s; \lambda) \quad \text{for all } a > 0.$$

But

$$r_{ij}(as; \lambda) = \exp(\lambda a(s_i - s_j)),$$

so the equality $r_{ij}(as; \lambda) = r_{ij}(s; \lambda)$ for all $a > 0$ would imply

$$\exp(\lambda a(s_i - s_j)) = \exp(\lambda(s_i - s_j)) \quad \text{for all } a > 0.$$

Taking logarithms, this requires

$$\lambda a(s_i - s_j) = \lambda(s_i - s_j) \quad \text{for all } a > 0,$$

which is impossible when $\lambda > 0$ and $s_i \neq s_j$. Thus, there exists at least one $a > 0$ such that $w^{\text{sm}}(as) \neq w^{\text{sm}}(s)$. The function $\phi(x) = ax$ is strictly increasing, so we have shown that there exists a strictly increasing ϕ such that $w^{\text{sm}}(\phi \circ s) \neq w^{\text{sm}}(s)$.

Therefore, the softmax rule is not ordinally invariant. □

Lemma 4 shows that softmax weights depend on the arbitrary scale of the score: multiplying all

scores by a positive constant changes the portfolio, even though the ranking of scores is unchanged. More generally, any strictly increasing transformation that alters the relative gaps between scores will typically change the portfolio. When the underlying economic content of the signal is only ordinal (the ordering of assets by expected performance or risk exposure), such sensitivity to the arbitrary cardinal scale of the score is undesirable.

Combining Theorem 5, the structure of our rank-based mask, and Lemma 4, we obtain the following decision-theoretic interpretation. Assuming that scores are identified only up to strictly increasing transformations, and requiring that the portfolio rule treats assets symmetrically and is robust to such transformations, we are led to rank-based weighting rules. Our top/bottom mask is a simple and transparent member of this class. In contrast, softmax-based rules implicitly treat the levels of the scores as economically meaningful, and are fragile to arbitrary monotone rescalings of the signal. For these reasons, we view rank-based masks as the more natural vehicle for constructing economically robust portfolios from the AlphaGlass signal.

Finally, we show that the AlphaGlass Sharpe-maximization problem can be understood within a standard asset-pricing framework based on a Hansen-Jagannathan (HJ) volatility bound. The key object is the factor return $F_{t+1}(\theta)$ generated by a given parameter vector θ in the AlphaGlass class. Let

$$\mu(\theta) = \mathbb{E}[F_{t+1}(\theta)], \quad \sigma^2(\theta) = \text{Var}(F_{t+1}(\theta)),$$

and write $\text{SR}(\theta) = \mu(\theta)/\sigma(\theta)$ for the population Sharpe ratio of the AlphaGlass factor generated by θ .

With each θ we now associate a one-factor stochastic discount factor (SDF) of the form

$$m_{t+1}(\theta) = a(\theta) - b(\theta)F_{t+1}(\theta),$$

where the coefficients $a(\theta)$ and $b(\theta)$ are chosen so that the SDF has unit mean and prices the AlphaGlass factor itself, i.e.

$$\mathbb{E}[m_{t+1}(\theta)] = 1, \quad \mathbb{E}[m_{t+1}(\theta)F_{t+1}(\theta)] = 0.$$

These two linear restrictions pin down $(a(\theta), b(\theta))$ uniquely whenever $\sigma^2(\theta) > 0$, and define a family

$$\mathcal{M}_{\Pi} = \{m(\theta) : \theta \in \Theta\}$$

of one-factor SDFs indexed by the same parameters, gates, and shape functions that parameterize the AlphaGlass portfolio rules. We now formalize that, within the payoff class $\{F(\theta) : \theta \in \Theta\}$, maximizing the Sharpe ratio is equivalent to selecting the payoff that attains the largest HJ bound, i.e., the payoff that implies the largest minimal SDF volatility among all SDFs that price it.

Theorem 6 (AlphaGlass Sharpe ratios as HJ bounds)

Fix $\theta \in \Theta$ with $\sigma^2(\theta) > 0$ and let $F(\theta)$ be the associated AlphaGlass factor with mean $\mu(\theta)$ and variance $\sigma^2(\theta)$.

1. *Among all SDFs m with $\mathbb{E}[m] = 1$ that price $F(\theta)$, i.e. $\mathbb{E}[mF(\theta)] = 0$, the unique minimum-variance SDF m^* is linear in $F(\theta)$:*

$$m^*(\theta) = 1 - \frac{\mu(\theta)}{\sigma^2(\theta)}(F(\theta) - \mu(\theta)), \quad \text{Var}(m^*(\theta)) = \text{SR}(\theta)^2.$$

Equivalently, every such m satisfies $\sqrt{\text{Var}(m)} \geq |\text{SR}(\theta)|$, with equality only for $m^(\theta)$.*

2. If an SDF m prices all AlphaGlass factors $\{F(\theta) : \theta \in \Theta\}$, then necessarily

$$\sqrt{\text{Var}(m)} \geq \sup_{\theta \in \Theta} |\text{SR}(\theta)|.$$

Among all stochastic discount factors (SDFs) m that have unit mean and correctly price this payoff, the theorem characterizes the least volatile admissible SDF. This “closest” SDF with finite variance is linear in the payoff and is unique. Its variance equals the squared Sharpe ratio,

$$\text{Var}(m^*(\theta)) = \text{SR}(\theta)^2,$$

and any other SDF that prices $F(\theta)$ must have weakly larger volatility. Therefore, $|\text{SR}(\theta)|$ is exactly the tight HJ lower bound on SDF volatility implied by pricing the single payoff $F(\theta)$.

This is economically valuable because SDF volatility summarizes the tightness of no-arbitrage pricing restrictions: pricing payoffs with large mean return relative to risk requires an SDF with correspondingly large variability. The theorem, therefore, links the statistical performance of an AlphaGlass factor (its Sharpe ratio) to the minimal amount of pricing-kernel variation needed to rationalize it in an arbitrage-free model. In particular, a high Sharpe AlphaGlass factor is a payoff that forces any compatible SDF to be volatile, i.e., it imposes a large pricing restriction on the economy. The intuition can be extended to the entire AlphaGlass class: If there exists an SDF that simultaneously prices every AlphaGlass factor in the class, then its volatility must be at least the maximum Sharpe ratio attainable within the class. This provides an economic measure of the “difficulty” of jointly pricing the full family of AlphaGlass portfolios: the best-performing factor in Sharpe-ratio terms sets a lower bound on the volatility of any SDF that is consistent with pricing the entire set of AlphaGlass payoffs.