

Discussion of
“Machine Learning Meets Markowitz”

by Wang, Gao, Harvey, Liu, and Tao.

Daniele Bianchi
Queen Mary, University of London

53rd EFA Annual Meeting

The paper, and where I want to push

The idea: stop training the pricing function on MSE. Embed the investor's *constrained* portfolio problem as a differentiable layer, and train $g(\mathbf{x}; \boldsymbol{\vartheta})$ on realized utility instead.

The setting: daily China A-shares, 2010–2023; long-only, position caps, homogeneous turnover costs, different risk-aversion levels.

My comments are about how much of the claim the empirics establish, and about the economics underneath:

#1: Standard ML *can* be made cost-aware.

#2: The gains come from the constraints, not the training.

#3: The output is a policy, not an expected return.

+ *minor comments at the end.*

Comment #1a: “Cannot train with cost awareness”

— but a weighted loss can

The motivation rests on this claim, made four times — pp. 2, 3, 5 and 36. On p. 36:

↪ “The prediction model is trained solely to minimize a statistical error like MSE, **without any awareness of the transaction costs** ... this inability to adapt is an important drawback.”

And again on p. 9, describing what the framework achieves:

“... different assets are **weighted differently in reducing the overall MSE**. This differential weighting is **absent from Bayesian or other traditional approaches**.”

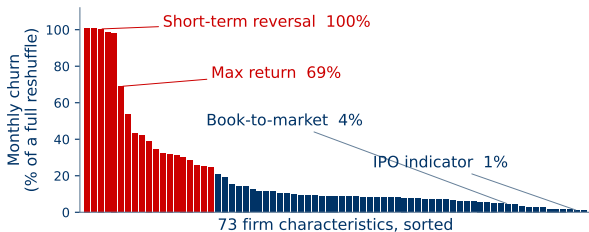
↪ But differential weighting of squared errors **is weighted least squares** — think of a value-weighted loss.

↪ *Objection: WLS weights are chosen, not endogenous. Are they arbitrary? ⇒ next slide.*

Comment #1b: The cost of a signal is measurable

— before any return model is fitted

Example: turnover of characteristic-managed portfolios.



100% = ranking fully reshuffled each month. US data, 73 characteristics — not your panel.

↪ A **loading** on short-term reversal means **replacing the position every month**; value and profitability churn **3–4%** as much — **a hundredfold spread**.

↪ **Hand-picked: YES** — so are the position cap, the architecture, and the smoothing constant your optimiser needs (§3.4, never reported).

↪ **Arbitrary: NO** — the cost is **measured**, not assumed.

↪ **Ask:** would a two-stage benchmark with a cost-weighted loss close the gap?

Comment #2a: Most of the gain is the constraint

Table 3, NN3, **annualized Sharpe**. Add a turnover penalty to the *optimiser* — no retraining:

	Transaction cost rate			
	10 bp	20 bp	30 bp	40 bp
Two-stage MSE, no cost term	1.23	0.51	−0.20	−0.91
+ turnover penalty	1.41	1.16	0.97	0.87
End-to-end	1.64	1.42	1.12	1.07
closed by the penalty alone	44%	71%	89%	90%

↪ Volatilities $\approx 20\%$ across all three — nothing hiding in risk.

↪ The advantage lives **only in the long-only top-decile portfolio E2E was trained to optimize**. On the long–short spread the two are indistinguishable. What improved is the **policy**, not the **forecast**.

Comment #2b: And the rest is just trading less

To concede: the gap over the cost-aware benchmark is real — 2.9 to 5.1 pp.

Two ways to save: pick **cheaper stocks (extensive)** or trade the **same stocks less (intensive)**. The paper tells the extensive story — p. 5, “*down-weighting high-turnover characteristics as transaction costs escalate*” — but...

Daily turnover, % of portfolio	0 bp	40 bp	saving
Before the penalty (extensive)	58.6	50.5	+8
After the penalty (intensive)	58.6	10.8	+40

↪ **80% of the saving is intensive** — the same stocks, traded less. The signal itself barely moves.

↪ **And the adapted signal still loses money:** 50.5%/day at 40 bp costs $\approx 49\%$ /yr vs $\approx 44\%$ gross.

↪ **Ask:** is a homogeneous cost sensible, and is the signal adjustment economically meaningful?

Comment #3: Is there such a thing as an investor-specific expected return?

Comment #3: Preferences move demand, not means

An investor's preferences move their **demand curve** and **optimal policy**.

They do not move the **conditional mean**.

↪ Two things actually differ across your investors: the **mandate** (long-only), and **where a limited model spends its accuracy** — p. 9, “we save the limited degree of freedom ... for assets that we likely long.”

↪ Both are facts about the **investor's problem**, not about returns.

↪ Same dividend yield, same present-value decomposition, **same forecast** — different risk aversion, different **willingness to act**. Net returns are investor-specific, but that is $E[r]$ **minus a cost**: the expected return is common, the cost is yours.

↪ Jensen et al. (2024) put the investor in Λ , not in μ . **The heterogeneity belongs in the cost function** — and yours has none: γ is uniform, constant, size-independent.

↪ **Ask**: if it is a forecast, **how good a forecast is it?** p. 2 already accepts “a lower cross-sectional R^2 ” — but no R^2 or MSE appears anywhere in the paper.

Minor comments — happy to take these offline

- ↪ **Significance.** NN3 is $t = 1.74$ equal-weighted and $t = 0.88$ value-weighted (Tables A.3, A.5) — yet NN3 alone carries Tables 3, 4 and 5. Stars appear to be one-sided; p. 35 claims 5% across all six models.
- ↪ **Table 5 pooling.** Nine η levels pooled into one t -stat across $\sim 21,900$ near-collinear daily observations; no clustering.
- ↪ **Dominance claims.** pp. 5 and 44 claim dominance “across the entire risk spectrum”; in Table 4, E2E has the best gross Sharpe in 0 of 4 rows at $\eta \geq 2.5$.
- ↪ **The cost model has no investor in it.** γ is uniform, constant, and independent of size — so no capacity dimension. AUM never appears; Amihud enters only as a *predictor*, never as a price. Price limits ($\pm 10\%$) never mentioned, though they bind hardest where the signal is strongest.
- ↪ **ζ and eq. (18).** §3.4 shows the LP has zero gradient without the $\zeta \|\mathbf{w}\|^2$ smoother — but eq. (18), the estimated equation, omits it, and ζ is never reported.
- ↪ **Reporting.** No alphas anywhere (CH-3 is available). MDD of 134% and 236% (Table 3) exceed the bound for a wealth process. Table A.1 $\times 243$ days implies $\approx 11.3\text{M}$ stock-days vs 5,962,768 reported (p. 32); no missing-data protocol. Table A.2 still lists CRSP, TAQ and GIM-governance characteristics for A-shares.

Nice Paper. A lot to think about.

Looking forward to reading the next version.