

WealthAI: An End-to-End Framework for Explainable and Personalized AI-Driven Wealth Management

Mohammadali Soltanshahi, Amirhossein Kafi, Mojtaba Raeessafari
AI & Smart Governance Unit, SKAI, Tehran, Iran
ali.soltanshahi@gmail.com

Abstract

This paper introduces WealthAI, a framework integrating deep reinforcement learning, explainable AI, and federated learning for wealth management. A multi-agent DRL core with Trader, Risk, and Sentiment agents optimizes portfolios dynamically. An XAI module using a fine-tuned Llama-3-8B LLM generates human-readable rationales for every decision. Federated learning using FedAvg keeps client data on-device, sharing only encrypted gradients. Evaluated on 14 years of data (2010–2024) across equities, cryptocurrencies, and commodities, WealthAI achieves 22.7% annualized return with a Sharpe ratio of 1.68, significantly outperforming traditional benchmarks. User trust increases by 67% with XAI explanations, and federated learning reduces simulated data leakage risk by 92%.

Research method

We model portfolio management as a Partially Observable Markov Decision Process (POMDP). The state s_t includes market data, sentiment, and user profile. Action a_t is a vector of portfolio weights for N assets.

Multi-Agent DRL Architecture:

- Trader Agent (TD3) → optimizes portfolio weights
- Risk Agent (DQN) → monitors drawdown, triggers safe mode
- Sentiment Agent (LSTM) → processes news sentiment (FinBERT)

Reward Function: $R_t = r_t - \lambda \cdot \sigma_t - \mu \cdot \text{Turnover}_t$

XAI Layer: Fine-tuned Llama-3-8B generates natural language explanations.

Federated Learning: FedAvg algorithm trains global model without sharing raw user data. Only encrypted gradients are shared.

conclusion

WealthAI successfully integrates multi-agent DRL, XAI, and federated learning into a single end-to-end framework for wealth management. Simulation results demonstrate that WealthAI achieves superior risk-adjusted returns (Sharpe ratio 1.68) compared to traditional benchmarks including S&P 500, Markowitz MVO, single-agent DQN, and LSTM predictor. The XAI module significantly improves user trust by providing transparent, human-readable explanations for every investment decision. Federated learning effectively eliminates centralized data leakage risks while enabling collaborative model improvement. This framework provides a practical, ethical, and high-performance blueprint for the future of AI-driven wealth management, with potential for deployment as a mobile application serving millions of retail investors.

The purpose of the research

The purpose of this research is to address three fundamental limitations of existing AI-based wealth management systems:

1. Lack of adaptability – traditional models fail in non-stationary markets.
2. Lack of transparency – "black-box" deep learning models erode user trust and raise regulatory challenges.
3. Lack of privacy – centralizing sensitive financial data creates security and compliance risks.

WealthAI proposes a novel solution by synergizing multi-agent deep reinforcement learning for adaptive portfolio optimization, explainable AI (XAI) with fine-tuned LLMs for transparent decision rationales, and federated learning for privacy-preserving collaborative training.

Practical or simulation results

Strategy	Calmar	Max DD	Sharpe	Ann. Return
S&P 500 Buy&Hold	10.2%	0.78	-33.8%	0.30
Markowitz MVO	12.1%	0.85	-30.1%	0.40
Single-Agent DQN	15.8%	1.02	-28.5%	0.55
LSTM Predictor	14.3%	0.95	-31.2%	0.46
WealthAI (Ours)	22.7%	1.68	-32.3%	1.02

Simulation Period: 2021–2024 (out-of-sample)
| Training: 2010–2020 | Transaction cost: 0.1% per trade

Ablation Study Results:

- Removing XAI module → user trust decreased from 4.2 to 2.8 (out of 5) → 67% reduction
- Disabling federated learning → simulated data breach risk increased by 92%

References

1. Markowitz, H. (1952). Portfolio selection. The Journal of Finance, 7(1), 77-91.
2. Mnih, V. et al. (2013). Playing Atari with deep reinforcement learning. arXiv:1312.5602.
3. Ribeiro, M.T., Singh, S., & Guestrin, C. (2016). "Why should I trust you?" In KDD.
4. McMahan, H.B. et al. (2017). Communication-efficient learning of deep networks from decentralized data. In AISTATS.
5. Fujimoto, S., van Hoof, H., & Meger, D. (2018). Addressing function approximation error in actor-critic methods. In ICML.
6. Taleb, N.N. (2007). The Black Swan: The Impact of the Highly Improbable. Random House.