

QuantaCo Markets

Intelligence Layer: Technical Overview

Aaditya Kumar

Quantaco Markets · quantacomarkets.com

Abstract

This document provides a technical overview of the QuantaCo Markets Meta-Intelligence Layer — the system that monitors investment policy statement mandate compliance for boutique discretionary wealth managers and delivers a calibrated breach probability score before markets open each day. We describe the conceptual architecture, the anomaly detection framework, the causal validation layer, and the explainability mechanism that makes every prediction auditable. Implementation details and proprietary hyperparameters are not disclosed.

Contents

1 Problem Statement	2
2 System Overview	2
3 Belief Propagation Anomaly Layer	2
3.1 Factor Graph Structure	2
3.2 Anomaly Detection via BP Residuals	2
4 Granger Causality: Factor Edge Validation	3
4.1 Motivation	3
4.2 Regression and Test	3
4.3 Edge Classification	3
5 Breach Probability Model	4
5.1 Meta-Model Architecture	4
5.2 Training on Real Data	4
6 Explainability via SHAP Attribution	4
7 Mandate Status Decision Rule	4
8 Validation Methodology	4
9 Output: PM Intelligence Brief	5

1 Problem Statement

Independent boutique wealth managers managing discretionary client portfolios are legally bound by Investment Policy Statement (IPS) mandates — documents specifying asset allocation bands, beta limits, sector concentration caps, and risk thresholds for each individual client. Breaching these mandates creates regulatory liability, client relationship damage, and potential legal exposure.

Most boutique firms monitor IPS compliance manually using spreadsheets and basic market alerts, spending 2–3 hours each morning checking positions before markets open. This process is slow, error-prone, and reactive — by the time a breach is detected, it has already occurred.

The QuantaCo intelligence layer is designed to change this. Rather than alerting a portfolio manager when a breach has happened, the system estimates the probability that a breach will occur within a forward horizon, enabling proactive intervention before mandate limits are crossed.

2 System Overview

The Meta-Intelligence Layer aggregates outputs from all pipeline factor models into a single calibrated breach-probability estimate:

$$\hat{p}_t = P(\text{IPS breach within } N \text{ trading days} \mid \mathcal{F}_t) \quad (1)$$

where $\mathcal{F}_t$ denotes the full information set at time t , comprising portfolio metrics, Monte Carlo VaR outputs, belief propagation salience residuals, Granger causal signals, sentiment scores, rolling beta exposures, and historical return series.

The pipeline runs in four sequential stages:

$$\text{load_factor_outputs}() \rightarrow \text{compute_anomaly_signals}() \rightarrow \text{build_feature_vector}() \rightarrow \hat{p}_t = \sigma(F(\mathbf{f}_t)) \quad (2)$$

The output $\hat{p}_t \in (0, 1)$ is a calibrated probability. A score of 0.68, for example, means the system estimates a 68% probability of an IPS mandate breach within the next 10 trading days given current portfolio conditions.

3 Belief Propagation Anomaly Layer

3.1 Factor Graph Structure

The system models portfolio risk factors as nodes in a graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ where:

$$\mathcal{V} = \{\text{rates, sentiment, betas, metrics, weights}\} \quad (3)$$

and edges $\mathcal{E}$ encode conditional dependence between factor nodes derived from the portfolio’s historical correlation regime. Belief propagation over this graph allows the system to model how information and risk flow across correlated factors simultaneously, rather than treating each factor in isolation.

3.2 Anomaly Detection via BP Residuals

For each factor node $i \in \mathcal{V}$, the system computes the difference between its historically expected salience and its actual observed salience in the current period. This residual is the core anomaly signal:

$$r_i = \hat{s}_i^{\text{hist}} - s_i^{\text{actual}} \quad (4)$$

The interpretation is direct:

$$r_i > 0 \implies \text{factor less salient than history predicts (settling)} \quad (5)$$

$$r_i < 0 \implies \text{factor **more** salient than history predicts} \implies \text{breach precursor} \quad (6)$$

When a factor is behaving more extremely than its entire historical record would predict, the system flags it as a potential breach precursor. This signal is captured before any mandate limit is formally crossed, providing genuine early warning rather than reactive alerting.

The most anomalous factor at time t is identified as:

$$i^* = \arg \min_{i \in \mathcal{V}} r_i \quad (7)$$

These anomaly residuals feed directly into the downstream breach probability model as named, interpretable features — `bp_res_rates`, `bp_res_sentiment`, `bp_res_betas`, `bp_res_metrics`, `bp_res_weights` — ensuring that the anomaly signal is fully traceable in every prediction.

4 Granger Causality: Factor Edge Validation

4.1 Motivation

The belief propagation layer assumes a graph topology connecting the five factor nodes. Not all edges in this graph reflect genuine statistical lead-lag relationships. The Granger causality layer validates each edge independently using historical factor time series, producing a transparent verdict for each directed relationship.

4.2 Regression and Test

For each directed edge ($i \rightarrow j$) in the factor graph, the Granger causality test fits a vector autoregression and tests whether past values of factor i provide statistically significant predictive power over factor j beyond j 's own history:

$$H_0 : \gamma_1 = \gamma_2 = \dots = \gamma_p = 0 \quad (8)$$

where γ_k are the cross-lag coefficients representing the predictive contribution of factor i to factor j .

4.3 Edge Classification

Each factor graph edge receives a transparent verdict based on the minimum p-value obtained across tested lag orders:

$$\text{Verdict}(i \rightarrow j) = \begin{cases} \text{SUPPORTED} & \text{if } p^* < 0.05 \\ \text{WEAK} & \text{if } 0.05 \leq p^* < 0.15 \\ \text{UNSUPPORTED} & \text{if } p^* \geq 0.15 \end{cases} \quad (9)$$

This layer is deliberately honest about the limits of the model. Rather than assuming all factor relationships are causally valid, the system reports which edges are statistically supported and which are coincidental. A portfolio manager or compliance officer reviewing the intelligence output can see exactly which causal relationships the system trusts and which it does not.

5 Breach Probability Model

5.1 Meta-Model Architecture

The breach probability model fuses the outputs of all pipeline factor models — including volatility estimates, tail risk measures, sequence forecasts, and belief propagation anomaly residuals — into a single fixed-dimensional feature vector $\mathbf{f}_t$. This vector is scored by a gradient-boosted ensemble classifier trained on real historical portfolio data.

The raw logit z_t produced by the ensemble is mapped to a probability via sigmoid activation:

$$\hat{p}_t = \sigma(z_t) = \frac{1}{1 + e^{-z_t}}, \quad \hat{p}_t \in (0, 1) \quad (10)$$

The model is trained using binary cross-entropy loss directly against real breach labels, which produces calibrated probabilities without requiring a separate post-hoc calibration step.

5.2 Training on Real Data

No synthetic data is used in training. The model learns from two real data sources:

Backtest bootstrap: Historical portfolio return series generate labeled training rows by checking whether realized future returns breach the VaR mandate limit within the forecast horizon. Labels are derived from realized returns, not simulated values.

Audit trail: Each recorded pipeline run contributes a real portfolio state vector. Breach labels are assigned based on whether the VaR mandate was subsequently exceeded in consecutive audit runs.

The model retrains on every pipeline run, automatically incorporating the latest audit history. No separate retraining step is required.

6 Explainability via SHAP Attribution

Every breach probability score is accompanied by a SHAP (SHapley Additive exPlanations) attribution that decomposes $\hat{p}_t$ into named, signed contributions from each input feature. The top five drivers by absolute contribution are surfaced in the PM brief as plain-English explanations, making every prediction fully auditable and every rebalancing decision documentable for compliance purposes.

7 Mandate Status Decision Rule

The final mandate status is a deterministic function of the breach probability score $\hat{p}_t$ and the current VaR breach flag:

$$\text{Status}_t = \begin{cases} \text{BREACH} & \text{if current VaR exceeds mandate limit} \vee \hat{p}_t \text{ above high threshold} \\ \text{AT_RISK} & \text{if } \hat{p}_t \text{ above moderate threshold} \vee \text{criticalalerts} \text{ present} \\ \text{CLEAR} & \text{otherwise} \end{cases} \quad (11)$$

The specific thresholds are calibrated to the empirical breach rate distribution and are not disclosed publicly.

8 Validation Methodology

The breach probability model is validated using time series cross-validation with an expanding window walk-forward scheme to prevent look-ahead bias. At each fold, the model is trained on all data up to time t and evaluated on the subsequent forward window:

$$\text{Train: } \{1, \dots, t\} \quad (12)$$

$$\text{Test: } \{t + 1, \dots, t + H\} \quad (13)$$

Evaluation metrics include:

Brier Score — calibration of probability estimates:

$$\text{BS} = \frac{1}{T} \sum_{t=1}^T (\hat{p}_t - y_t)^2 \quad (14)$$

Brier Skill Score — improvement over the base rate baseline:

$$\text{BSS} = 1 - \frac{\text{BS}}{\bar{y}(1 - \bar{y})} \quad (15)$$

ROC-AUC — discriminative power across all thresholds:

$$\text{AUC} = \int_0^1 \text{TPR}(\text{FPR}^{-1}(u)) \, du \quad (16)$$

Average Precision — area under the precision-recall curve, appropriate for the imbalanced breach label distribution:

$$\text{AP} = \sum_k (R_k - R_{k-1}) P_k \quad (17)$$

The full model is benchmarked against two baselines: a VaR-only model using standard risk metrics alone, and a naive model predicting the historical breach rate for all periods. Incremental value of the belief propagation anomaly layer is measured by comparing performance between the full model and an ablated version with BP residuals removed.

9 Output: PM Intelligence Brief

The pipeline culminates in two deliverables generated automatically before market open:

Dashboard: A live unified view showing the breach probability gauge, mandate status, top SHAP feature drivers, belief propagation field salience, and Granger causality edge verdicts for the current portfolio.

White-labeled PDF investment memo: An auto-generated report branded to the client firm containing the executive summary, priority flags with specific actions, position-level recommendations with urgency ratings, mandate compliance section, and tax efficiency guidance. Every recommendation is traceable to the specific quantitative inputs that generated it.

The combination of a calibrated forward-looking breach probability score, a transparent causal validation layer, full SHAP attribution on every prediction, and an audit-ready PDF report represents a genuinely different approach from rule-based compliance alerting tools. The system tells a portfolio manager not just that a breach has occurred, but which factors are causing it, how likely a formal breach is within the next ten trading days, and exactly what to do about it.

QuantaCo Markets · quantacomarkets.com · Toronto, Canada

*This document describes the conceptual architecture of the QuantaCo intelligence layer for informational purposes.
Proprietary implementation details, hyperparameters, and training specifications are not disclosed. All rights reserved.*