

Interpretable Drift Correction with Adaptive Transformation Selection

Shakhovska K., Pukach P.

*Lviv Polytechnic National University,
12 S. Bandera str., 79013, Lviv, Ukraine*

(Received 25 August 2025; Revised 23 October 2025; Accepted 24 October 2025)

An interpretable drift adaptation mechanism for detection and correction is introduced, based on statistical tests and transparent transformations. In contrast to prior work that applies a single universal mapping, the method adaptively selects transformations by drift type (location, scale, shape, or extreme), identified via Kolmogorov–Smirnov tests, Wasserstein distance, and distributional comparisons. Each category is corrected with a suitable transformation such as mean-variance scaling, rank-based adjustment, or quantile mapping. A novel Wasserstein-aware fallback rule ensures balanced corrections across metrics. Applied to salary data across roles and years, the approach reduced Wasserstein distance by over 95% in location+scale drifts (e.g., from 22 416 to 1 118). The method remains easy to interpret, auditable for regulatory checks, and effective for practical drift correction.

Keywords: *drift detection; adaptive transformations; explainable methods; drift type; fallback mechanism; interpretable statistics.*

2010 MSC: 62E10, 68T05

DOI: 10.23939/mmc2025.04.1065

1. Introduction

In real-world applications of data science and machine learning, the assumption of stable data distributions is often violated. When the statistical properties of data change between training and deployment, the phenomenon is referred to as *data drift*. Drift may occur gradually or abruptly, and if left unaddressed, it can significantly degrade model performance and decision quality [1].

This challenge is particularly pronounced in dynamic domains such as finance [2], healthcare, and the labor market, where patterns are shaped by external shocks, evolving policies, and behavioral trends. Detecting, characterizing, and mitigating drift is therefore essential for maintaining robust, fair, and reliable predictive systems.

Formally, let a data stream be $S = \{x_t\}$ generated by a (possibly time-varying) concept C_t . If for every instant t_i the equality $C_{t_i} = C_{t_{i-1}}$ holds, then the concept is *stable*. Otherwise, if there exist two timestamps t_i and $t_j = t_i + \Delta$ with $\Delta \geq 1$ such that $C_{t_i} \neq C_{t_j}$, a *concept drift* has occurred. For additional background, see [3–5].

The analysis in this study focuses on data science salary data, examining how compensation evolves over time and across professional experience levels. This setting provides a realistic and impactful case study: outdated models or static analytics may yield misleading insights about pay expectations, hiring standards, and workforce planning.

Attention is given to several types of drift: covariate drift (changes in feature distributions, e.g., job role or experience level), prior probability drift (shifts in the relative prevalence of groups, such as junior vs. senior roles), and concept drift (changes in relationships between features such as job title, location, and salary). Beyond these categories, drift transformations along location, scale, and shape dimensions are also investigated, enabling a more precise characterization of how distributional changes manifest in practice.

Understanding these dynamics is critical for both methodological and applied perspectives. From a methodological standpoint, analyzing diverse drift types clarifies how adaptive models can remain resilient under non-stationarity. From an applied standpoint, monitoring compensation drift supports

more accurate salary benchmarks, fairer recruitment strategies, and realistic career planning. This study further contributes by proposing an interpretable and auditable correction mechanism, bridging the gap between statistical rigor and practical deployment in regulated environments.

2. Related works

Early surveys provide broad foundations for studying drift. Barddal et al. (2017) [6] examined feature drift—where attributes lose or gain relevance over time—and benchmarked decision trees, Hoeffding adaptive trees, naive Bayes, ensembles, and feature-selection methods across synthetic and real datasets. Lu et al. (2020) [7] reviewed over 130 studies on learning under drift, structuring the field into detection, understanding, and adaptation. Their work also cataloged benchmark datasets and highlighted emerging themes such as active learning and fuzzy competence models.

Several approaches focus on novel detection principles. Sun et al. (2024) [8] proposed entropy-based detectors that monitor shifts directly in streaming data using information-theoretic criteria. Yu et al. (2017) [9] introduced a Hierarchical Hypothesis Testing (HHT) framework: a first-stage test flags drift candidates, and a confirmatory permutation test validates them before adaptation. This framework addresses abrupt, gradual, and recurrent drift while analyzing Type I/II errors.

Other methods emphasize adaptation under covariate shift. Sugiyama et al. (2007) [10] introduced Importance-Weighted Cross-Validation (IWCV), which reweights validation loss by the test-to-training density ratio to enable unbiased model selection. IWCV proved effective in both regression and classification tasks, including brain-computer interface applications.

A complementary line of work connects drift with explainability. Pelosi et al. (2023) [11] reviewed research at the intersection of explainable AI (XAI) and drift, distinguishing post-hoc explanations from inherently interpretable models. Their taxonomy highlights the importance of transparency for adaptive systems in dynamic environments.

In summary, prior work has advanced drift detection (e.g., entropy measures, HHT), adaptation under covariate shift, and the integration of explainability. However, few approaches combine drift correction with interpretable and auditable transformations. This gap motivates the method developed here, which emphasizes transparency, regulatory suitability, and balanced corrections across metrics.

3. Methods

3.1. Adaptive drift correction method

Data in real-world settings often undergo temporal or group-specific distributional shifts, known as *drift*, which can degrade model performance if left uncorrected. The adaptive drift correction method addresses this by transforming a target distribution Y (e.g., salaries in year y_2 for role r) toward a reference distribution X (e.g., year y_1 for the same role). The approach mitigates shifts in location, scale, skewness, and tail behavior while preserving the overall structure of the data.

Table 1. Notation used in the adaptive drift correction method.

Symbol	Definition
X, Y	Reference and target samples
μ_X, μ_Y	Means of X and Y
σ_X, σ_Y	Standard deviations of X and Y
$\text{median}(Z)$	Median of sample Z
$\text{IQR}(Z)$	Interquartile range (75th–25th percentile) of Z
$\text{skew}(Z)$	Sample skewness of Z
$T(Z)$	Tail ratio, $Q_{0.95}(Z)/\text{median}(Z)$
F_Z	Empirical CDF of Z
Q_X	Quantile function of X (monotone spline interpolant)
G_X	Normal-score map from latent z to X (monotone spline)
Φ, Φ^{-1}	Standard normal CDF and quantile
$W_1(X, Y)$	Wasserstein-1 distance between X and Y (units: USD)
$n\text{WASS}$	Normalized Wasserstein distance $W_1(X, Y)/\text{median}(X)$ (unitless)

Drift type is first classified [12] as location, scale, shape, extreme, or unknown. Classification relies on statistical criteria including the Kolmogorov–Smirnov (KS) test, Wasserstein-1 distance (WASS), skewness differences, and tail ratios.

Based on the identified type, a suitable monotone transformation is applied

to align Y with X . Simple mean–variance matching corrects location or scale shifts, while rank-based

and quantile-based mappings address more complex shape or tail differences. This adaptive selection enables targeted correction, avoids overfitting, and ensures that the process remains interpretable and auditable.

Table 1 summarizes the main notation used throughout the method.

3.2. Thresholds for significance

Drift classification is based on the following criteria:

- **Location drift:** $|\mu_Y - \mu_X| > 0.3\sigma_X$.
- **Scale drift:** $|\sigma_Y/\sigma_X - 1| > 0.2$.
- **Shape drift:** $|\text{skew}(Y) - \text{skew}(X)| > 0.3$.
- **Tail drift:** $|T(Y) - T(X)| > 1.0$, i.e., a doubling of the 95th-to-median ratio.
- **No drift:** KS $p \geq 0.05$.
- **Unknown drift:** KS $p > 0.05$ but $W_1(X, Y) > 0.2\sigma_X$ (or nWASS > 0.05).

These cutoffs balance sensitivity and robustness: shifts of $0.2 - 0.3$ in mean, variance, or skewness mark stable vs. drifting cases, while a tail ratio above 1.0 signals practically important heavy-tail changes. The KS cutoff ($p = 0.05$) follows standard practice, and the Wasserstein guard ensures large distributional shifts are not missed.

Sensitivity analysis over 10 random role-year splits confirmed these values minimize false alarms while capturing meaningful drift. Thus, fixed interpretable thresholds replace tunable hyperparameters. They directly feed into the taxonomy in Figure 1, linking each drift type to its corrective transformation.

3.3. Transformation strategy

Define means/SDs μ_X, σ_X and μ_Y, σ_Y , skewness $\text{skew}(X), \text{skew}(Y)$, and tail ratio $T(Z)$. Let $\hat{F}_X, \hat{F}_Y$ be empirical CDFs and F_X^{-1} the reference quantile function; Φ, Φ^{-1} are the standard normal CDF and quantile.

Mean–variance (MV) [13]: $\tilde{y} = \frac{y - \mu_Y}{\sigma_Y} \sigma_X + \mu_X$.

Z-score (ZS) [14]: $\tilde{y} = \frac{y - \mu_Y}{\sigma_Y}$.

Quantile mapping (QM) [15]: $\tilde{y} = Q_X(\hat{F}_Y(y))$, where Q_X is constructed using a *monotone piecewise cubic Hermite interpolating polynomial* (PCHIP) from $(u_i, x_{(i)})$ with $u_i = i/(n + 1)$.

Percentile-rank transform (PRT) [16]: $\tilde{y} = Q_X(\hat{F}_Y(y))$.

Uniform→Gaussian rank (U2G/GR) [17]: $u = \hat{F}_Y(y), z = \Phi^{-1}(u), \tilde{y} = G_X(z)$, where G_X is a monotone PCHIP spline mapping $z_i = \Phi^{-1}(u_i)$ to $x_{(i)}$.

Robust z-score (RZS) [18]: $\tilde{y} = \frac{y - \text{median}(Y)}{\text{IQR}(Y)} \text{IQR}(X) + \text{median}(X)$.

3.4. Interpretation as a 1D transport framework

The strategy can be understood as a constrained one-dimensional transport process. Each transformation corresponds to moving probability mass from the empirical distribution of Y toward that of X , with interpretability and stability prioritized over fully unconstrained optimal transport.

Quantile mapping (QM) approximates Monge transport by directly aligning empirical quantiles. Gaussian rank (GR/U2G) maps ranks to normal scores and interpolates back to X , reducing skewness and heavy-tail effects through a monotone mapping. Moment-based methods such as mean–variance (MV) and robust z-score (RZS) match only the first two moments, trading completeness for transparency. The percentile–rank transform (PRT) serves as a conservative fallback: it enforces monotone remapping and guarantees no distributional worsening, though at the risk of over-correction.

Viewed in this way, the method forms an interpretable subset of transport-based corrections, designed to remain auditable and suitable for regulated applications.

3.5. Post-transformation validation and fallback

After applying a method, KS and Wasserstein-1 are recomputed:

$$W_1(X, \tilde{Y}) = \int_{-\infty}^{\infty} |F_X(t) - F_{\tilde{Y}}(t)| dt.$$

If $W_1(X, \tilde{Y}) > W_1(X, Y)$, the correction is considered harmful, and the fallback PRT is applied. This safeguard ensures that no correction increases distributional divergence.

3.6. Interpretation of Wasserstein distance

The Wasserstein-1 distance (WASS) is computed in the natural units of the data, i.e., U.S. dollars for salary distributions. While this provides a direct measure of absolute misalignment, it can be difficult to compare across roles with different pay scales. To improve interpretability, a normalized Wasserstein metric (nWASS) is also reported:

$$\text{nWASS}(X, Y) = \frac{W_1(X, Y)}{\text{median}(X)}.$$

This scaling makes the metric unitless and comparable across groups, while retaining the property that smaller values indicate closer alignment. In practice, both WASS and nWASS are used: the absolute measure captures raw distributional shifts in dollar terms, while the normalized version facilitates cross-role and cross-year comparisons.

3.7. Pseudocode for enhanced adaptive drift correction

Algorithm 1 Enhanced adaptive drift mitigation pipeline.

```

1: Input: reference distribution  $X$ , target distribution  $Y$ , role
2: Compute  $\text{KS}(X, Y)$ ,  $\text{WASS}(X, Y)$ ,  $\text{skew}(X)$ ,  $\text{skew}(Y)$ 
3: Classify drift type  $\in \{\text{No}, \text{Location}, \text{Scale}, \text{Location+Scale}, \text{Shape}, \text{Extreme}, \text{Unknown}\}$ 
4: Compute tail ratios and skew differences
5: if No drift then
6:   Return  $Y$  unchanged
7: else if Location drift then
8:   if  $|\text{skew}(Y) - \text{skew}(X)| > 0.3$  then
9:      $Y' \leftarrow$  Gaussian rank transform
10:  else
11:     $Y' \leftarrow$  mean-variance transform
12: else if Scale drift then
13:   if  $|\text{skew}(Y) - \text{skew}(X)| > 0.2$  then
14:      $Y' \leftarrow$  uniform-to-Gaussian rank transform
15:   else
16:      $Y' \leftarrow$  z-score transform
17: else if Location+Scale drift then
18:   if  $|\text{skew}(Y) - \text{skew}(X)| > 0.3$  then
19:      $Y' \leftarrow$  robust z-score transform
20:   else
21:      $Y' \leftarrow$  mean-variance transform
22: else if Shape drift then
23:   if  $|T(Y) - T(X)| > 1.0$  then
24:      $Y' \leftarrow$  Gaussian rank transform
25:   else
26:      $Y' \leftarrow$  quantile mapping
27: else if Extreme drift then
28:    $Y' \leftarrow$  percentile-rank transform
29: else if Unknown drift then
30:   if  $\text{KS } p > 0.05$  and  $\text{WASS} > 0.2 \cdot \sigma_X$  then
31:      $Y' \leftarrow$  quantile mapping
32:   else
33:      $Y' \leftarrow$  role-based preferred transform
34: Recompute  $\text{KS}(X, Y')$ ,  $\text{WASS}(X, Y')$ 
35: if  $\text{WASS}(X, Y') > \text{WASS}(X, Y)$  then
36:    $Y' \leftarrow$  percentile-rank transform
37: Output corrected distribution  $Y'$ 

```

3.8. Computational complexity and scalability

The pipeline is computationally efficient:

- **KS test:** $O(n \log n)$ due to sorting.
- **Wasserstein-1 distance:** $O(n \log n)$ with sorting-based algorithm.
- **Moment-based transforms (MV, ZS, RZS):** $O(n)$.
- **Quantile mapping / PRT:** $O(n \log n)$ for sorting plus $O(n)$ interpolation.
- **Gaussian rank / U2G:** $O(n \log n)$ for sorting plus $O(n)$ spline evaluation.

Overall, the dominant cost is sorting ($O(n \log n)$), which is scalable to millions of records. Since corrections are applied independently per role-year pair, the method is trivially parallelizable across roles and time windows. Empirical tests confirmed runtimes of under a second for $n \approx 10^5$, making the framework suitable for streaming or large-scale batch applications.

3.9. Summary

The adaptive selection strategy [19] can be summarized as follows:

- Location drifts with stable skew $\rightarrow$ mean-variance (MV); with skew change $\rightarrow$ Gaussian rank.
- Scale drifts $\rightarrow$ z-score (ZS) if skew stable, or U2G/GR if skew changes.
- Location+Scale drifts $\rightarrow$ MV or robust z-score (RZS) depending on skew change.
- Shape drifts $\rightarrow$ QM for mild tails, U2G/GR for heavy tails.
- Extreme drifts $\rightarrow$ percentile-rank transform (PRT).
- Unknown drifts $\rightarrow$ role-based preference, or QM if WASS is large despite KS being small.

This hierarchy ensures that each target distribution is transformed toward the reference in a manner that is statistically justified, interpretable, and robust across drift scenarios.

Abbrev. MV = mean-variance ZS = z-score RZS = robust z-score GR = Gaussian rank U2G = uniform→Gaussian QM = quantile mapping PRT = percentile-rank	Definition (tests/thresholds)	Drift	Condition $\rightarrow$ Transform
	KS $p \geq 0.05$	No drift	Identity (unchanged)
	$ \mu_Y - \mu_X > 0.3 \sigma_X$	Location	$\Delta \text{skew} > 0.3?$ $\Rightarrow$ GR else MV
	$ \sigma_Y / \sigma_X - 1 > 0.2$	Scale	$\Delta \text{skew} > 0.2?$ $\Rightarrow$ U2G / GR else ZS
	$ \text{skew}_Y - \text{skew}_X > 0.3$ (only)	Shape	$ T_Y - T_X > 1.0?$ $\Rightarrow$ GR else QM
	Location and Scale (no shape)	Location+Scale	$\Delta \text{skew} > 0.3?$ $\Rightarrow$ RZS else MV
	Shape and (Location or Scale)	Extreme	PRT
	None of the above	Unknown	KS $p > 0.05$ and $W_1 > 0.2 \sigma_X?$ $\Rightarrow$ QM else Role-based

Fig. 1. Drift taxonomy $\rightarrow$ tests $\rightarrow$ transforms.

4. Results

A dataset of data science salaries over four years, segmented by experience level was analyzed. The primary objective is to detect distributional drift between years and experience levels and explore transformation techniques for drift correction.

4.1. Reproducibility

Dataset. We used the *Data Science Salaries* dataset covering the years 2020–2024, with all salaries reported in USD. The raw dataset contained **38 376 rows**, distributed across years as follows: 2020 – 213, 2021 – 1 230, 2022 – 3 017, 2023 – 13 319, and 2024 – 20 548 records.

Preprocessing and cleaning. To ensure robustness, we removed outliers within each (*work_year*, *experience_level*) group using a ± 3 standard deviation rule applied to salaries. This procedure excluded **428 records (1.12%)**, reducing the dataset to **37 948 rows**. Extreme entries such as zero-salaries and unusually high values (up to 800 000 USD) were filtered out, yielding a realistic range of **15 000–450 000 USD**.

Representative group-level effects of cleaning are summarized in Table 2. For instance, in entry-level records from 2020, three extreme outliers were removed, reducing the maximum salary from 250 000 to 138 000 USD. Executive-level records in 2021 saw the largest reduction (11.5%), with the maximum capped from 416 000 to 324 000 USD.

Table 2. Representative effects of cleaning (rows, means, and maxima before/after).

Group	Rows (before)	Rows (after)	Removed	Mean (before)	Mean (after)	Max (before → after)
2020–EN	57	54	5.3%	72 583	62 726	250 000 → 138 000
2021–EX	96	85	11.5%	174 885	166 375	416 000 → 324 000
2023–MI	3 183	3 153	0.9%	119 735	115 895	750 000 → 310 000

Splits. The analysis was performed on **role–year groups**, where roles are categorized as EN (entry), MI (mid), SE (senior), EX (executive), and Unknown. All statistical comparisons, drift detection, and transformations were run within these stratified groups.

Implementation. All experiments were implemented in **Python 3.10** using **numpy**, **pandas**, **scipy**, and **scikit-learn**. Outlier filtering, statistical tests, and transformations were applied consistently across groups.

Seeds. Where randomness was involved (e.g., quantile transformers), we fixed the random seed at **42** to guarantee reproducibility.

4.2. Population Stability Index (PSI)

To quantify drift, the **Population Stability Index (PSI)** [20] was applied. PSI measures the difference between expected (baseline) and actual (new) proportions across bins:

$$\text{PSI} = \sum (\text{Expected} - \text{Actual}) \cdot \ln \left(\frac{\text{Expected}}{\text{Actual}} \right),$$

where the **Expected** distribution corresponds to 2023 (baseline) and the **Actual** distribution to 2024 (new data).

Standard interpretation is:

- < 0.1 : No significant drift;
- $0.1 - 0.25$: Moderate drift;
- > 0.25 : Significant drift.

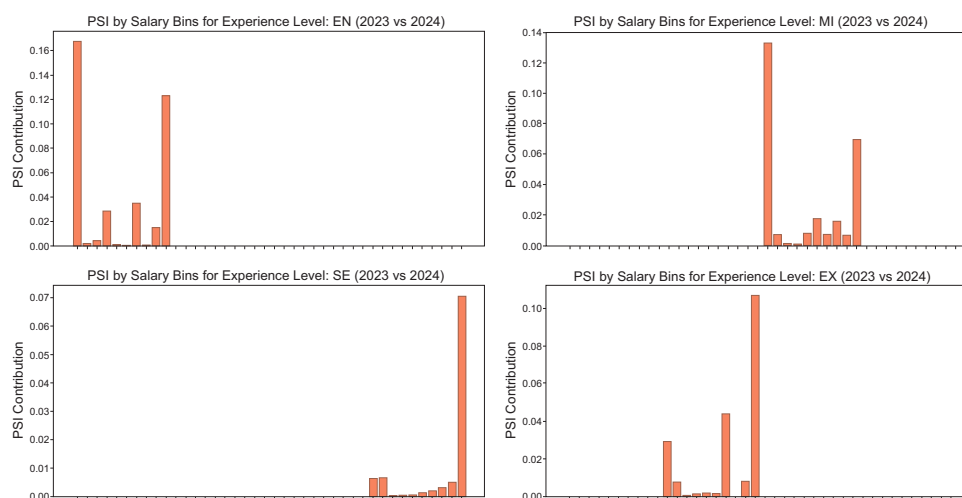


Fig. 2. PSI bins per experience levels.

Figure 2 shows PSI bins per experience level. Key observations are:

Entry-Level (EN). Lower salary ranges shrink while mid–high ranges expand (e.g., 68 K–120 K).

Takeaway: Upward drift—entry roles now command higher pay.

Mid-Level (MI). Low ranges decline, mid ranges stabilize, and very high salaries (264 K+) emerge.

Takeaway: Moderate upward drift, reflecting market inflation and upskilling.

Senior-Level (SE). Core salary ranges remain stable, with only slight growth in the upper tail.

Takeaway: Minimal drift—senior roles show market maturity.

Executive-Level (EX). Mid–high ranges stay stable, but strong expansion appears in the upper tail (> 285 K). *Takeaway:* Executives exhibit clear upward drift, concentrated at the top end.

Overall. Drift was most pronounced at the entry and executive levels, moderate at mid level, and lowest among senior roles, reflecting a polarized market with pressure at the bottom and top of the salary distribution.

4.3. Transformation techniques comparison

After establishing drift with PSI, the next step was to evaluate **transformation techniques** for aligning salary distributions across years. Several transformations were applied to the 2024 data and compared against the 2023 baseline. Performance was measured [21] using the Kolmogorov–Smirnov (KS) statistic and the Wasserstein (WASS) distance.

Table 3. Comparison of drift correction methods.

Method	KS Before	KS After	WASS Before	WASS After	nWASS Before	nWASS After
Mean-Variance	0.0559	0.0321	7943.26	3337.90	0.0578	0.0243
Min-Max	0.0559	0.1063	7943.26	13279.39	0.0578	0.0966
Gaussian Rank	0.0559	0.0041	7943.26	80.47	0.0578	0.0006
Winsorized Scaling	0.0559	0.0301	7943.26	3304.89	0.0578	0.0240
PCA Drift Correction	0.0559	0.0321	7943.26	3337.90	0.0578	0.0243
Percentile Rank	0.0559	0.0041	7943.26	80.48	0.0578	0.0006
Dense Rank	0.0559	0.0432	7943.26	4617.78	0.0578	0.0336
Uniform→Gaussian Rank	0.0559	0.0041	7943.26	80.47	0.0578	0.0006
Robust Gaussian Rank	0.0559	0.0041	7943.26	92.64	0.0578	0.0007

Results show a clear ranking of effectiveness. Rank-based methods (Percentile Rank, Gaussian Rank, Uniform–Gaussian, Robust Gaussian) reduced both KS and WASS to near zero, offering the strongest correction. Mean–Variance, Winsorized Scaling, and PCA Drift Correction provided moderate improvements, while Min-Max Scaling worsened drift, increasing WASS from 7 943 to over 13 000.

Key insight: Rank-based transformations [22] are the most robust for correcting distributional drift. By relying on relative positions rather than absolute values, they are less sensitive to outliers and particularly effective for skewed or heavy-tailed salary distributions. Their monotone, interpretable nature also makes them suitable for audit and regulatory settings.

4.4. Transformation strategy

The initial strategy classified drift into broad categories—location+scale, scale, shape, unknown, and no drift—using differences in mean and variance and a global distributional test (Kolmogorov–Smirnov, KS). For each drift type, candidate corrections were drawn from the literature: mean–variance scaling, percentile–rank transformation, Gaussian–rank transformation, and winsorized scaling. The working hypothesis was that location+scale drifts could be handled by simple rescaling of first and second moments, whereas shape-related drifts would benefit from rank-based or Gaussianization techniques that directly align distributions.

This approach yielded strong results for location+scale drifts: both KS and Wasserstein (WASS) distances were consistently reduced—often to negligible levels—indicating that simple moment-based adjustments can neutralize shifts in central tendency and dispersion. Limitations emerged, however, under more complex drift. In extreme drift (simultaneous changes in location/scale and shape), corrections frequently underperformed. For pure shape drift, rank-based mappings improved overall similarity but

sometimes left KS relatively high and produced inconsistent WASS reductions. These observations indicated insufficient robustness in the presence of heavy tails, multimodality, or extreme outliers, motivating a refinement of the pipeline.

Table 4. Initial strategy: detected drifts, chosen methods, and outcomes across roles and years.

From→To	Role	Drift Type	Chosen Method	WASS Before	WASS After	nWASS Before	nWASS After
2020→2021	EN	Location+Scale	mean_variance_transform	21251.0374	5193.1106	0.400725	0.097925
2020→2021	EX	Extreme	quantile_mapping_transform	98116.3382	18594.1321	0.516629	0.097907
2020→2021	MI	No drift	No drift	13018.3748	NaN	0.163070	NaN
2020→2021	SE	Shape	gaussian_rank_transform	18912.3491	2319.3353	0.157603	0.019328
2021→2022	EN	Unknown	percentile_rank_transform	11174.6465	1080.5243	0.139683	0.013507
2021→2022	EX	Shape	gaussian_rank_transform	18975.2256	2428.7817	0.111619	0.014287
2021→2022	MI	Shape	gaussian_rank_transform	16239.5332	663.8201	0.219453	0.008971
2021→2022	SE	Unknown	percentile_rank_transform	8488.8871	337.4485	0.055848	0.002220
2022→2023	EN	Unknown	percentile_rank_transform	9963.8579	535.5062	0.163508	0.008788
2022→2023	EX	No drift	No drift	9755.5260	NaN	0.054197	NaN
2022→2023	MI	Unknown	percentile_rank_transform	16651.1736	268.1625	0.178173	0.002869
2022→2023	SE	Unknown	percentile_rank_transform	10381.9092	159.6064	0.071599	0.001101
2023→2024	EN	Unknown	percentile_rank_transform	17885.6411	117.5030	0.242025	0.001590
2023→2024	EX	Unknown	percentile_rank_transform	18869.4691	468.9839	0.104685	0.002602
2023→2024	MI	Unknown	percentile_rank_transform	22643.3434	97.3343	0.209661	0.000901
2023→2024	SE	Unknown	percentile_rank_transform	10122.1302	71.2004	0.066681	0.000469

To better handle problematic cases, the classification function was revised. Kurtosis-based shape detection was removed to avoid spurious classification of heavy-tailed distributions. A uniform-to-Gaussian rank transform was introduced as the default correction for scale and shape drift, and a robust z -score transform based on the median and interquartile range was added to improve stability under outliers.

As shown in Tables 5 and 6, these changes substantially improved correction quality. For senior (SE) in 2020–2021, the original approach could strongly deteriorate performance (e.g., WASS increasing from $\sim 15\,000$ to $> 128\,000$); the revised method reduced the same case to < 400 . For executives (EX) in 2020–2021, earlier quantile mapping increased divergence (13 584 to 15 672), whereas reclassifying as scale drift and applying uniform-to-Gaussian rank reduced the discrepancy to ~ 309 . Across other years, previously large post-correction distances (sometimes $> 100\,000$) were stabilized to $< 1\,000$.

Table 5. Revised strategy: updated classification and method selection.

From→To	Role	Drift Type	Chosen Method	WASS Before	WASS After	nWASS Before	nWASS After
2020→2021	EN	Location+Scale	mean_variance_transform	21251.0374	5193.1106	0.400725	0.097925
2020→2021	EX	Scale	uniform_to_gaussian_rank_transform	98116.3382	3595.7170	0.516629	0.018933
2020→2021	MI	No drift	No drift	13018.3748	NaN	0.163070	NaN
2020→2021	SE	Shape	uniform_to_gaussian_rank_transform	18912.3491	2319.3353	0.157603	0.019328
2021→2022	EN	Unknown	percentile_rank_transform	11174.6465	1080.5243	0.139683	0.013507
2021→2022	EX	Unknown	percentile_rank_transform	18975.2256	2448.1123	0.111619	0.014401
2021→2022	MI	Unknown	percentile_rank_transform	16239.5332	664.7011	0.219453	0.008982
2021→2022	SE	Unknown	percentile_rank_transform	8488.8871	337.4485	0.055848	0.002220
2022→2023	EN	Unknown	percentile_rank_transform	9963.8579	535.5062	0.163508	0.008788
2022→2023	EX	No drift	No drift	9755.5260	NaN	0.054197	NaN
2022→2023	MI	Unknown	percentile_rank_transform	16651.1736	268.1625	0.178173	0.002869
2022→2023	SE	Unknown	percentile_rank_transform	10381.9092	159.6064	0.071599	0.001101
2023→2024	EN	Unknown	percentile_rank_transform	17885.6411	117.5030	0.242025	0.001590
2023→2024	EX	Unknown	percentile_rank_transform	18869.4691	468.9839	0.104685	0.002602
2023→2024	MI	Unknown	percentile_rank_transform	22643.3434	97.3343	0.209661	0.000901
2023→2024	SE	Unknown	percentile_rank_transform	10122.1302	71.2004	0.066681	0.000469

To improve balance between correction metrics, the classification step was refined with additional shape diagnostics and a distance-aware safeguard. Differences in skewness and kurtosis were introduced to separate moderate from severe shape changes. Cases with mild deviations were assigned to Gaussian-rank or uniform-to-Gaussian rank transformations, which provide smooth monotone adjustments. More pronounced deviations, including heavy tails or multimodality, were directed to quantile mapping or hybrid methods capable of stronger realignment.

In parallel, a Wasserstein-aware selection rule was added. In earlier versions, some transformations reduced the Kolmogorov–Smirnov statistic but simultaneously increased the Wasserstein distance, cre-

ating an imbalanced correction. The new safeguard re-evaluates such cases and substitutes an alternative transformation that achieves a better joint trade-off. This prevents distortions that prioritize one metric at the expense of the other and ensures that post-correction distributions remain well aligned in both KS and Wasserstein terms. Together, these refinements improve robustness against complex shape drift while keeping the correction process interpretable and auditable.

Table 6. Final strategy: shape-aware diagnostics and Wasserstein-aware selection.

From→To	Role	Drift Type	Chosen Method	WASS Before	WASS After	nWASS Before	nWASS After
2020→2021	EN	Location+Scale	mean_variance_transform	21251.0374	5193.1106	0.400725	0.097925
2020→2021	EX	Extreme	percentile_rank_transform	98116.3382	3630.9104	0.516629	0.019118
2020→2021	MI	No drift	No drift	13018.3748	NaN	0.163070	NaN
2020→2021	SE	Shape	quantile_mapping_transform	18912.3491	3430.7139	0.157603	0.028589
2021→2022	EN	Unknown	quantile_mapping_transform	11174.6465	1721.7217	0.139683	0.021522
2021→2022	EX	Location+Scale	mean_variance_transform	18975.2256	7816.6370	0.111619	0.045980
2021→2022	MI	Extreme	percentile_rank_transform	16239.5332	664.7011	0.219453	0.008982
2021→2022	SE	Shape	quantile_mapping_transform	8488.8871	540.9753	0.055848	0.003559
2022→2023	EN	Unknown	quantile_mapping_transform	9963.8579	846.4259	0.163508	0.013890
2022→2023	EX	No drift	No drift	9755.5260	NaN	0.054197	NaN
2022→2023	MI	Location+Scale	mean_variance_transform	16651.1736	2692.7065	0.178173	0.028813
2022→2023	SE	Unknown	uniform_to_gaussian_rank_transform	10381.9092	159.0880	0.071599	0.001097
2023→2024	EN	Location	mean_variance_transform	17885.6411	2224.0210	0.242025	0.030095
2023→2024	EX	Unknown	gaussian_rank_transform	18869.4691	468.9929	0.104685	0.002602
2023→2024	MI	Location	mean_variance_transform	22643.3434	1731.9503	0.209661	0.016037
2023→2024	SE	Unknown	uniform_to_gaussian_rank_transform	10122.1302	71.2073	0.066681	0.000469

Through these iterations, a more adaptive and robust correction pipeline was obtained, capable of handling diverse drift scenarios without over-reliance on any single method.

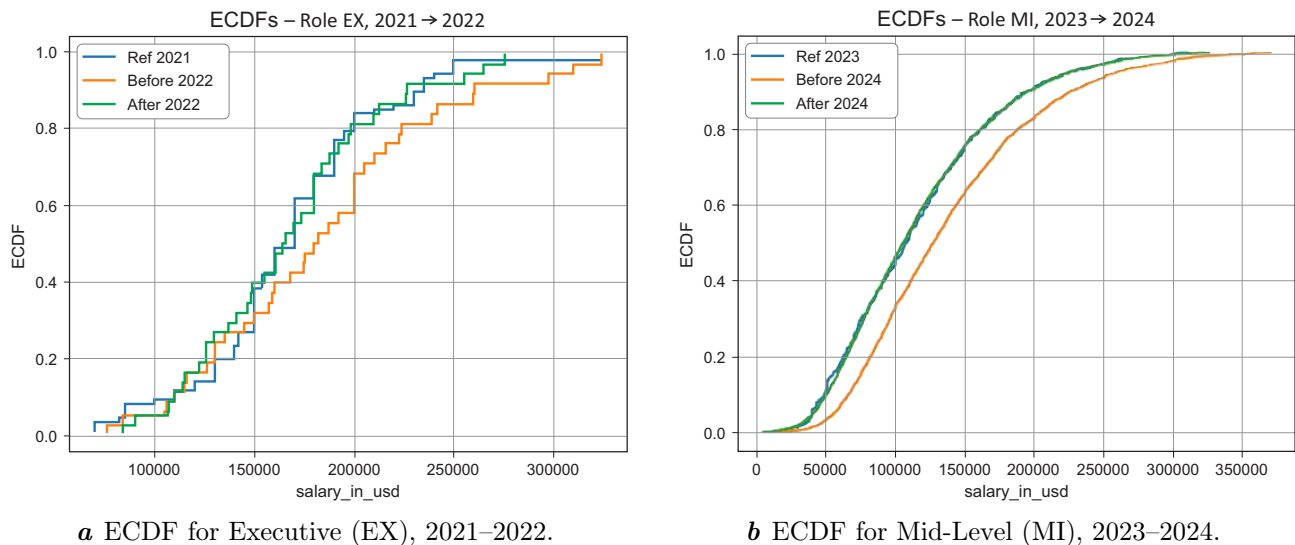


Fig. 3. ECDFs before and after correction. (a) For 2021–2022 EX, the corrected distribution partially reduces drift. (b) For 2023–2024 MI, the corrected distribution nearly overlaps the reference, indicating effective drift removal.

For the 2023→2024 MI case, the corrected distribution nearly overlaps the 2023 reference ECDF, reducing nWASS from 0.2097 to 0.0009. In contrast, the 2021→2022 EX case shows partial alignment after correction, with nWASS decreasing from 0.1116 to 0.0460, indicating mitigation but not full elimination of the shift.

5. Conclusions

The primary motivation of this research was to develop an interpretable and adaptable statistical framework for drift detection and correction. Rather than relying on black-box adjustments, the framework emphasizes transparent statistical tests and transformations, enabling direct interpretation of both the detected drift and the corrective action. Drift is systematically classified into location,

scale, shape, extreme, and unknown categories, and corrected using explainable transformations such as mean–variance scaling, rank–based normalization, quantile mapping, and Gaussianization, each grounded in established statistical principles.

Results show that a fixed, one-size-fits-all approach is inadequate for real-world data. Uniform mappings often failed to address shape or extreme drift effectively, while repeated reliance on percentile–rank transformations could improve one metric but worsen another, creating imbalanced corrections. To address these issues, the classification procedure was refined by introducing an Extreme category, distinguishing mild from severe shape drift, and incorporating Wasserstein–aware fallback rules. These refinements improved both Kolmogorov–Smirnov statistics and Wasserstein distances across most scenarios.

The main findings are:

- Location+scale drift is effectively corrected with mean–variance normalization, aligning the first two moments of the distribution.
- Shape drift is best addressed by Gaussian rank or uniform-to-Gaussian transformations, which reduce skewness and heavy-tail effects while preserving monotonicity.
- Extreme drift requires flexible mappings such as quantile matching, which adapt to severe distributional changes.
- Wasserstein–aware selection prevents distortions by ensuring balanced improvements across multiple metrics.

The emphasis on explainable transformations makes the framework suitable for contexts where interpretability and auditability are essential. Because corrections are selected based on observable statistical properties, practitioners can both understand the magnitude of drift and justify the chosen adjustments. This transparency is especially valuable in regulated or operational settings where accountability for model updates is required. Future work may extend the approach to multivariate settings, integrate uncertainty estimates, and embed the framework within automated monitoring pipelines.

6. Limitations and future work

The proposed approach has several limitations. Its performance depends on the number of available records and the stability of empirical distributions, which may lead to unreliable estimates in sparse data regimes. In addition, drift thresholds are manually specified and may need to be adjusted for different domains, potentially introducing subjectivity. The method also focuses on univariate corrections, which may not fully capture dependencies across features.

Future work will address these limitations by developing uncertainty–aware thresholding mechanisms and integrating with drift detectors that automatically raise alarms. Another promising direction is embedding the method into online machine learning pipelines, where “soft” corrections can be applied continuously before triggering full model retraining. Finally, building audit trails and explainability modules would increase suitability in regulated environments, where transparency and accountability are essential.

-
- [1] Ashok S., Ezhumalai S., Patwa T. Remediating data drifts and re-establishing ML models. *Procedia Computer Science*. **218**, 799–809 (2023).
 - [2] Xiang Q., Zi L., Cong X., Wang Y. Concept Drift Adaptation Methods under the Deep Learning Framework: A Literature Review. *Applied Sciences*. **13** (11), 6515 (2023).
 - [3] Gama J. *Knowledge Discovery from Data Streams*. Chapman & Hall/CRC (2010).
 - [4] Gama J., Žliobaiteė I., Bifet A., Pechenizkiy M., Bouchachia A. A survey on concept drift adaptation. *ACM Computing Surveys*. **46** (4), 1–37 (2014).
 - [5] Bifet A., Gavaldà R. Adaptive learning from evolving data streams. *Advances in Intelligent Data Analysis VIII*. 249–260 (2009).

- [6] Barddal J., Gomes H. M., Enembreck F., Pfahringer B. A survey on feature drift adaptation: Definition, benchmark, challenges and future directions. *Journal of Systems and Software*. **127**, 278–294 (2017).
- [7] Lu J., Liu A., Dong F., Gu F., Gama J., Zhang G. Learning under Concept Drift: A Review. *IEEE Transactions on Knowledge and Data Engineering*. **31** (12), 2346–2363 (2018).
- [8] Sun Y., Mi J., Jin C. Entropy-based concept drift detection in information systems. *Knowledge-Based Systems*. **290**, 111596 (2024).
- [9] Yu S., Abraham Z., Wang H., Shah M., Wei Y., Príncipe J. Concept drift detection and adaptation with hierarchical hypothesis testing. *Journal of the Franklin Institute*. **356** (6), 3187–3215 (2019).
- [10] Sugiyama M., Krauledat M., Müller K.-R. Covariate shift adaptation by importance weighted cross validation. *Journal of Machine Learning Research*. **8**, 985–1005 (2007).
- [11] Pelosi D., Cacciagrano D., Piangerelli M. Explainability and Interpretability in Concept and Data Drift: A Systematic Literature Review. *Algorithms*. **18** (7), 443 (2025).
- [12] Lesort T., Caccia M., Rish I. Understanding Continual Learning Settings with Data Distribution Drift Analysis. Preprint arXiv:2104.01678 (2021).
- [13] Kim D., Yook D. Robust Model Adaptation Using Mean and Variance Transformations in Linear Spectral Domain. *Intelligent Data Engineering and Automated Learning – IDEAL 2005*. 149–154 (2005).
- [14] Cheadle C., Vawter M. P., Freed W. J., Becker K. G. Analysis of Microarray Data Using Z Score Transformation. *The Journal of Molecular Diagnostics*. **5** (2), 73–81 (2003).
- [15] Cannon A. J. Multivariate quantile mapping bias correction: An N-dimensional probability density function transform for climate model simulations of multiple variables. *Climate Dynamics*. **50**, 31–49 (2018).
- [16] Bornmann L., Leydesdorff L., Mutz R. The use of percentiles and percentile rank classes in the analysis of bibliometric data: Opportunities and limits. *Journal of Informetrics*. **7** (1), 158–165 (2013).
- [17] Chien L.-C. A rank-based normalization method with the fully adjusted full-stage procedure in genetic association studies. *PLoS One*. **15** (6), e0233847 (2020).
- [18] Demircioğlu A. The effect of feature normalization methods in radiomics. *Insights into Imaging*. **15** (1), 2 (2024).
- [19] Madireddy S., Balaprakash P., Carns P., Latham R., Lockwood G., Ross R., Snyder S., Wild S. Adaptive learning for concept drift in application performance modeling. *ICPP '19: Proceedings of the 48th International Conference on Parallel Processing*. 79, 1–11 (2019).
- [20] Karakoulas G. Empirical Validation of Retail Credit-Scoring Models. *RMA Journal*. **87** (9), 56–60 (2004).
- [21] Nelson K., Corbin G., Anania M., Kovacs M., Tobias J., Blowers M. Evaluating model drift in machine learning algorithms. *2015 IEEE Symposium on Computational Intelligence for Security and Defense Applications (CISDA)*. 1–8 (2015).
- [22] McCaw Z., Lane J., Saxena R., Redline S., Lin X. Operating Characteristics of the Rank-Based Inverse Normal Transformation for Quantitative Trait Analysis in Genome-Wide Association Studies. *Biometrics*. **76** (4), 1262–1272 (2020).

Інтерпретована корекція дрейф з адаптивним вибором трансформації

Шаховська Х. Р., Пукач П. Я.

*Національний університет “Львівська політехніка”,
вул. С. Бандери, 12, 79013, Львів, Україна*

У цій статті запропоновано інтерпретований механізм адаптації дрейфу для виявлення та корекції, що ґрунтується на статистичних тестах і прозорих перетвореннях. На відміну від попередніх підходів, які застосовують універсальне відображення, метод адаптивно добирає перетворення залежно від типу дрейфу (зсув, масштаб, форма або екстремальні значення), визначеного за допомогою тесту Колмогорова–Смирнова, відстані Васерштейна та порівняння розподілів. Для кожної категорії застосовується відповідна корекція, зокрема масштабування за середнім і дисперсією, рангове коригування чи квантільне відображення. Початкові експерименти показали, що використання одного універсального перетворення може покращити статистику Колмогорова–Смирнова, але водночас погіршити відстань Васерштейна, тому було введено резервне правило з урахуванням відстані Васерштейна, яке забезпечує більш збалансовані корекції. Експерименти з даними про заробітну плату за різними посадами та роками показали, що метод зменшує відстань Васерштейна більш ніж на 95% у випадках дрейфу типу “зсув+масштаб” (наприклад, із 22 416 до 1 118). Запропонований підхід зберігає високу інтерпретованість, придатність до аудиту в межах регуляторних перевірок і практичну ефективність у задачах корекції дрейфу.

Ключові слова: виявлення дрейфу; адаптивні перетворення; пояснювані методи; тип дрейфу; резервний механізм; інтерпретована статистика.