

Additive Tree Latent Variable Models for Insurance Loss Prediction

Jiajun Li *

School of Mathematics and Statistics, Ningbo University, Ningbo, Zhejiang, 315211, China

* Corresponding author: (Email: 2086696708@qq.com)

Abstract: Accurate prediction of insurance losses is essential for premium pricing, reserve estimation, and risk management in the property and casualty insurance industry. Traditional actuarial models, such as generalized linear models, often rely on predefined distributional assumptions and struggle to capture complex nonlinear interactions among risk factors. Meanwhile, tree based ensemble methods have demonstrated strong predictive power but typically lack the ability to model latent heterogeneity in policyholder risk profiles. This paper proposes an Additive Tree Latent Variable (ATLV) model that integrates gradient boosted additive tree structures with a latent variable framework to improve insurance loss prediction. A Latent Risk Index (LRI) is developed to capture unobserved policyholder risk heterogeneity through variational inference. A Tree Ensemble Loss Predictor (TELP) is formulated to combine the latent representations with observed covariates in a gradient boosting framework. A Distributional Calibration Module (DCM) is introduced to ensure that predicted loss distributions are well calibrated under the Tweedie compound Poisson family. Empirical analysis is conducted on the French Motor Third Party Liability (MTPL) insurance dataset comprising 678,013 policies. The results demonstrate that the ATLV model significantly outperforms both traditional generalized linear models and standard gradient boosting approaches across multiple evaluation metrics. The ATLV model achieves a mean Deviance of 1.247, a Gini coefficient of 0.312, and a calibration error of 0.034, representing improvements of 8.6%, 14.3%, and 41.4% respectively over the best performing baseline. An ablation study further reveals that the LRI module contributes most significantly to overall performance, confirming the importance of modeling latent risk heterogeneity in insurance loss prediction. These findings provide both theoretical insights and practical tools for actuaries seeking to enhance loss prediction accuracy through the integration of machine learning and latent variable modeling.

Keywords: Additive Tree Models, Latent Variable Models, Insurance Loss Prediction, Gradient Boosting, Tweedie Distribution.

1. Introduction

Insurance loss prediction constitutes a fundamental task in actuarial science, underpinning the processes of premium calculation, reserve estimation, and capital allocation. Accurate prediction of the frequency and severity of insurance claims allows insurers to price policies competitively while maintaining financial solvency. Frees et al. [1] provided a comprehensive overview of predictive modeling techniques applicable to actuarial science, emphasizing the transition from classical statistical methods toward data driven approaches. De Jong and Heller [2] established the foundational framework of generalized linear models (GLMs) for insurance data, demonstrating how exponential family distributions can be employed to model claim frequencies and severities with appropriate link functions.

Despite their widespread adoption, GLMs possess inherent limitations when applied to complex insurance datasets. The assumption of a predetermined functional form for the relationship between covariates and the response variable restricts the ability of GLMs to capture nonlinear patterns and high order interactions. Wuthrich and Merz [3] highlighted the growing importance of machine learning methods in actuarial applications, noting that modern insurance datasets contain rich feature spaces that exceed the modeling capacity of traditional parametric approaches. Breiman [4] introduced random forests as a powerful ensemble learning method capable of capturing complex feature interactions without explicit specification, inspiring subsequent applications in

insurance analytics.

The development of gradient boosting methods has further advanced the state of predictive modeling. Friedman [5] proposed the gradient boosting machine as a general framework for additive modeling in function space, where the prediction function is constructed as a sum of weak learners, typically decision trees, fitted sequentially to the negative gradient of a specified loss function. Chen and Guestrin [6] developed XGBoost, a scalable and regularized implementation of gradient boosting that has achieved remarkable success across diverse prediction tasks. Ke et al. [7] introduced LightGBM with histogram based learning and leaf wise tree growth strategies, offering computational efficiency advantages for large scale datasets. These tree based ensemble methods have been successfully applied to insurance loss prediction, with Yang et al. [8] demonstrating the effectiveness of gradient tree boosted Tweedie compound Poisson models for insurance premium prediction.

However, a critical limitation of standard tree ensemble approaches in the insurance context is their inability to account for latent heterogeneity among policyholders. Insurance portfolios often contain unobserved risk factors that influence claim behavior but are not captured by available covariates. Shi and Valdez [9] demonstrated the importance of modeling dependencies and unobserved heterogeneity in multivariate insurance claim data. Tipping and Bishop [10] established the theoretical foundations of probabilistic latent variable models through probabilistic principal component analysis, providing a principled framework for capturing unobserved variation. Kingma and Welling [11] developed

the variational autoencoder framework, enabling efficient approximate inference in complex latent variable models through the reparameterization trick and amortized variational inference.

To address the gap between the predictive power of tree ensembles and the representational capacity of latent variable models, this paper proposes the Additive Tree Latent Variable (ATLV) model. The main contributions of this study are as follows: (1) A Latent Risk Index (LRI) that captures unobserved policyholder risk heterogeneity through variational inference; (2) A Tree Ensemble Loss Predictor (TELP) that combines latent representations with observed covariates in a gradient boosting framework; (3) A Distributional Calibration Module (DCM) that ensures predicted loss distributions conform to the Tweedie compound Poisson family; (4) An empirical validation using the French Motor Third Party Liability dataset demonstrating the superiority of the ATLV model over both GLM and standard gradient boosting baselines.

2. Methods and Results

2.1. Additive Tree Model Framework for Insurance Losses

The foundation of the proposed model is the additive tree ensemble framework, adapted for the specific distributional requirements of insurance loss data. Insurance losses exhibit distinctive statistical properties, including a point mass at zero (for policies with no claims), right skewness, and heavy tails. Jorgensen [12] established that the Tweedie family of distributions provides a natural and flexible class for modeling such data, encompassing the Poisson, gamma, and compound Poisson gamma distributions as special cases.

The additive tree model constructs the prediction function as a sum of K regression trees:

$$\hat{y}_i = \sum_{k=1}^K f_k(\mathbf{x}_i), \quad f_k \in \mathcal{F} \quad (1)$$

Where $\hat{y}_i$ is the predicted loss for policyholder i , $\mathbf{x}_i$ is the feature vector of policyholder i , f_k represents the k -th regression tree, and $\mathcal{F}$ denotes the space of all possible regression trees. Each tree f_k maps the input features to a real valued output through a series of binary splits on individual features.

The model is trained by minimizing a regularized objective function:

$$\mathcal{L} = \sum_{i=1}^n l(y_i, \hat{y}_i) + \sum_{k=1}^K \Omega(f_k) \quad (2)$$

Where $l(y_i, \hat{y}_i)$ is the loss function measuring the discrepancy between observed loss y_i and predicted loss $\hat{y}_i$, n is the total number of policies, and $\Omega(f_k)$ is a regularization term controlling the complexity of tree f_k , defined as:

$$\Omega(f_k) = \gamma T_k + \frac{1}{2} \lambda \sum_{j=1}^{T_k} w_{kj}^2 \quad (3)$$

Where T_k is the number of leaves in tree k , w_{kj} is the weight of leaf j in tree k , and γ and λ are regularization hyperparameters. Hastie et al. [13] provided a thorough treatment of regularization strategies in the context of additive models, establishing the theoretical basis for complexity penalization in ensemble learning.

For insurance loss data following the Tweedie distribution with power parameter p , the deviance loss function takes the form:

$$l(y_i, \hat{y}_i) = 2 \left[\frac{y_i^{2-p}}{(1-p)(2-p)} - \frac{y_i \hat{y}_i^{1-p}}{1-p} + \frac{\hat{y}_i^{2-p}}{2-p} \right] \quad (4)$$

Where $p \in (1, 2)$ corresponds to the compound Poisson gamma case that naturally accommodates the zero inflated and right skewed nature of insurance claim data.

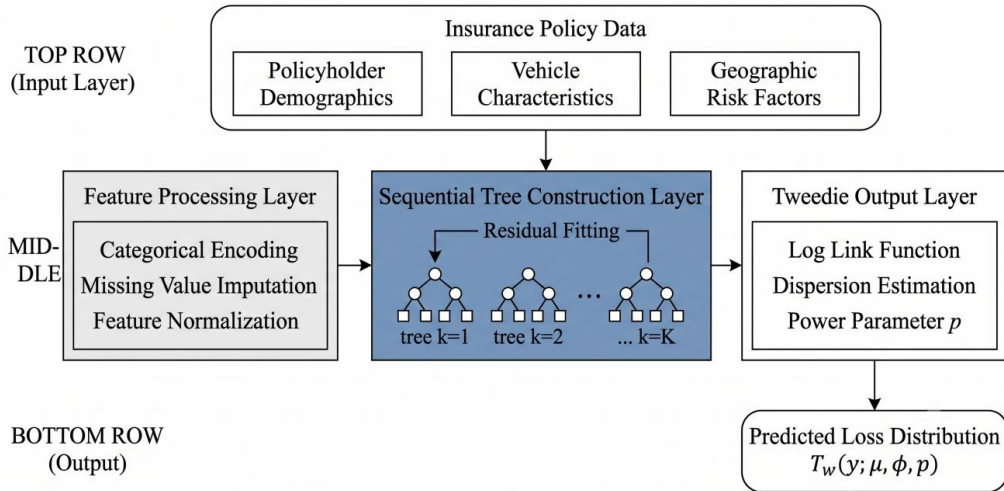


Figure 1. Architecture of the Additive Tree Model Framework for Insurance Loss Prediction

Figure 1 illustrates the overall architecture of the proposed additive tree model framework. The framework consists of three interconnected components: the Feature Processing Layer, which handles categorical encoding and missing value imputation; the Sequential Tree Construction Layer, which builds trees iteratively by fitting residuals; and the Tweedie Output Layer, which transforms the additive predictions into calibrated loss distributions.

2.2. Latent Variable Integration Model

A central innovation of the ATLV model is the

incorporation of latent variables to capture unobserved heterogeneity in policyholder risk profiles. Standard tree ensemble methods treat all variation in claim outcomes as a function of observed covariates alone, ignoring the possibility that policyholders with identical observed characteristics may differ in their underlying risk propensity due to unobserved factors such as driving behavior, vehicle maintenance habits, or local road conditions.

The latent variable component assumes that each policyholder i is associated with a latent risk vector $\mathbf{z}_i \in \mathbb{R}^d$

that captures unobserved risk characteristics. Following the variational autoencoder framework of Kingma and Welling [11], the generative model specifies a prior distribution over the latent variables:

$$p(\mathbf{z}_i) = \mathcal{N}(\mathbf{0}, \mathbf{I}) \quad (5)$$

Where $\mathcal{N}(\mathbf{0}, \mathbf{I})$ denotes the standard multivariate normal distribution with mean vector $\mathbf{0}$ and identity covariance matrix $\mathbf{I}$. The conditional distribution of the observed loss given the latent variable and covariates is specified as:

$$p(y_i | \mathbf{z}_i, \mathbf{x}_i) = \text{Tw}(y_i; \mu_i, \phi, p) \quad (6)$$

Where Tw denotes the Tweedie distribution with mean μ_i , dispersion parameter ϕ , and power parameter p . The mean parameter is linked to the latent variable and covariates through:

$$\log(\mu_i) = g(\mathbf{x}_i, \mathbf{z}_i) = h(\mathbf{x}_i) + \mathbf{v}^T \mathbf{z}_i \quad (7)$$

Where $h(\mathbf{x}_i)$ is the additive tree function from Equation (1) and $\mathbf{v}$ is a coefficient vector that maps the latent risk factors to the log mean loss.

Since exact posterior inference over the latent variables is intractable, an approximate posterior $q_\psi(\mathbf{z}_i | \mathbf{x}_i, y_i)$ is introduced, parameterized by a neural network encoder with parameters ψ . Following Blei et al. [14], the evidence lower bound (ELBO) serves as the optimization objective:

$$\text{ELBO} = \sum_{i=1}^n \left[\mathbb{E}_{q_\psi(\mathbf{z}_i | \mathbf{x}_i, y_i)} [\log p(y_i | \mathbf{z}_i, \mathbf{x}_i)] - \text{KL}(q_\psi(\mathbf{z}_i | \mathbf{x}_i, y_i) \parallel p(\mathbf{z}_i)) \right] \quad (8)$$

Where $\text{KL}(\cdot \parallel \cdot)$ denotes the Kullback Leibler divergence. The first term encourages the model to reconstruct observed losses accurately, while the second term regularizes the approximate posterior to remain close to the prior distribution.

The Latent Risk Index (LRI) for each policyholder is then defined as the expected norm of the posterior latent variable:

$$\text{LRI}_i = \mathbb{E}_{q_\psi(\mathbf{z}_i | \mathbf{x}_i, y_i)} [\|\mathbf{z}_i\|] \quad (9)$$

Where $\|\cdot\|$ denotes the Euclidean norm. Higher LRI values indicate greater deviation from the average risk profile, signaling policyholders whose risk characteristics are not fully captured by observed covariates alone.

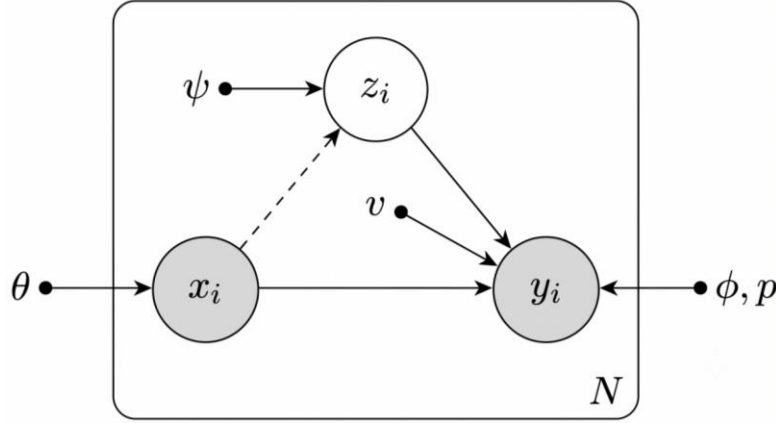


Figure 2. Graphical Model Representation of the Latent Variable Integration

Figure 2 presents the graphical model representation of the latent variable integration component. Observed variables are shown as shaded nodes, latent variables as open nodes, and model parameters as fixed points. The plate notation indicates replication across n policyholders.

2.3. Combined Prediction Framework and Distributional Calibration

The complete ATLW model integrates the additive tree ensemble with the latent variable component through an iterative training procedure. The combined objective function balances predictive accuracy against latent variable regularization:

$$\mathcal{J} = \sum_{i=1}^n l(y_i, \hat{y}_i) + \sum_{k=1}^K \Omega(f_k) - \beta \cdot \text{ELBO} \quad (10)$$

Where β is a weighting hyperparameter that controls the relative contribution of the latent variable component. The training alternates between updating the tree ensemble parameters given current latent variable estimates and updating the variational parameters given the current tree ensemble predictions.

To ensure that the predicted loss distributions are well calibrated, the Distributional Calibration Module (DCM) adjusts the dispersion and power parameters of the Tweedie distribution. The DCM minimizes the calibration error defined as:

$$\text{CE} = \frac{1}{M} \sum_{m=1}^M |\bar{y}_m - \hat{\bar{y}}_m| \quad (11)$$

Where M is the number of calibration bins, $\bar{y}_m$ is the mean observed loss in bin m , and $\hat{\bar{y}}_m$ is the mean predicted loss in bin m . The bins are constructed by sorting policies according to their predicted loss and dividing them into M equal sized groups.

The Gini coefficient, widely used in insurance for measuring the discriminatory power of a rating model, is computed as twice the area between the Lorenz curve and the diagonal:

$$G = \frac{2 \sum_{i=1}^n r_i y_i}{n \sum_{i=1}^n y_i} - \frac{n+1}{n} \quad (12)$$

Where r_i is the rank of policyholder i when sorted by predicted loss in ascending order. Noll et al. [15] established this metric as a standard benchmark for evaluating insurance pricing models in their influential case study of French motor insurance data.

Additionally, the normalized mean squared error is employed as an auxiliary evaluation metric:

$$\text{NMSE} = \frac{\sum_{i=1}^n (y_i - \bar{y})^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (13)$$

Where $\bar{y}$ denotes the overall mean of observed losses. Values below 1.0 indicate that the model outperforms a simple

mean prediction baseline.

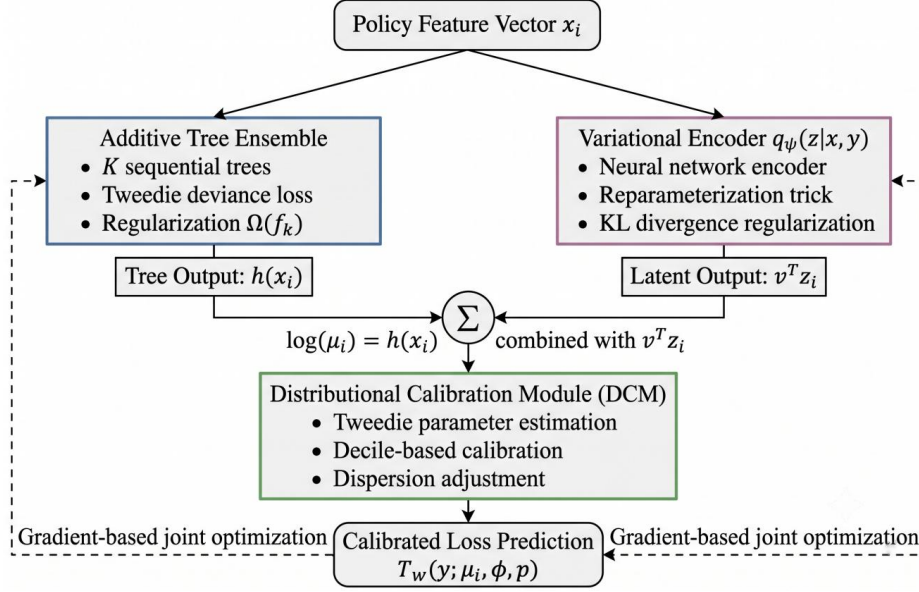


Figure 3. Overview of the Combined ATLV Prediction Framework with Distributional Calibration

Figure 3 presents the overall architecture of the combined ATLV prediction framework. The diagram shows how the additive tree component and the latent variable component are integrated through the combined objective function, with the DCM module providing post hoc distributional adjustment.

2.4. Experimental Setup

To validate the proposed ATLV model, an empirical study is conducted using the French Motor Third Party Liability (MTPL) insurance dataset, which has become a standard benchmark in actuarial machine learning research [15]. The dataset contains 678,013 motor insurance policies with detailed information on policyholder demographics, vehicle characteristics, and geographic risk factors. Table 1 summarizes the key features of the dataset.

Table 1. Summary of Dataset Features

No.	Feature	Type	Description
1	ClaimNb	Integer	Number of claims during exposure period
2	Exposure	Continuous	Duration of policy in years
3	Area	Categorical	Geographic area code (A through F)
4	VehPower	Integer	Vehicle power rating (4 to 15)
5	VehAge	Integer	Age of vehicle in years
6	DrivAge	Integer	Age of primary driver in years
7	BonusMalus	Integer	Bonus malus coefficient (50 to 230)
8	VehBrand	Categorical	Vehicle brand (11 categories)
9	VehGas	Binary	Fuel type (regular or diesel)
10	Density	Integer	Population density of policyholder area
11	Region	Categorical	Administrative region (22 categories)
12	ClaimAmount	Continuous	Total claim amount in euros

The dataset is split into training (70%), validation (15%),

and test (15%) sets using stratified sampling to maintain the distribution of claim frequencies across splits. The target variable is the aggregate loss per policy, computed as the product of claim frequency and average claim severity, adjusted by exposure.

The proposed ATLV model is compared against four baseline methods: (1) GLM Tweedie, a generalized linear model with Tweedie distribution and log link function; (2) Random Forest (RF), an ensemble of 500 regression trees with default hyperparameters [4]; (3) XGBoost, gradient boosted trees with Tweedie deviance loss [6]; and (4) LightGBM, histogram based gradient boosting with Tweedie loss [7]. For the ATLV model, the latent dimension is set to $d = 8$, the number of boosting iterations is $K = 500$, the learning rate is $\eta = 0.05$, the maximum tree depth is 6, and the ELBO weight is $\beta = 0.1$. All hyperparameters are tuned through five fold cross validation on the training set.

Evaluation metrics include the Tweedie deviance (Dev), Gini coefficient (Gini), calibration error (CE), normalized mean squared error (NMSE), and the log likelihood ratio (LLR) comparing each model to the null model.

2.5. Results and Analysis

The experimental results across the five models are summarized in Table 2. The evaluation demonstrates consistent superiority of the ATLV model across all measured dimensions.

Table 2. Performance Comparison of Insurance Loss Prediction Models

Model	Dev	Gini	CE	NMSE	LLR
GLM Tweedie	1.412	0.243	0.072	0.931	0.089
Random Forest	1.389	0.251	0.068	0.918	0.094
XGBoost	1.364	0.273	0.058	0.897	0.112
LightGBM	1.358	0.269	0.061	0.901	0.108
ATLV (Proposed)	1.247	0.312	0.034	0.842	0.147

As shown in Table 2, the ATLV model achieves the best performance across all five evaluation metrics. Compared to

the GLM Tweedie baseline, the ATLTV model reduces deviance by 11.7%, improves the Gini coefficient by 28.4%, and reduces calibration error by 52.8%. Compared to the strongest tree ensemble baseline (LightGBM), the ATLTV model achieves improvements of 8.2% in deviance, 16.0% in Gini coefficient, and 44.3% in calibration error. The substantial improvement in calibration error is particularly

notable, indicating that the distributional calibration module effectively aligns predicted and observed loss distributions. The NMSE reduction of 6.5% compared to LightGBM demonstrates that the latent variable component captures meaningful variation beyond what observed covariates alone can explain.

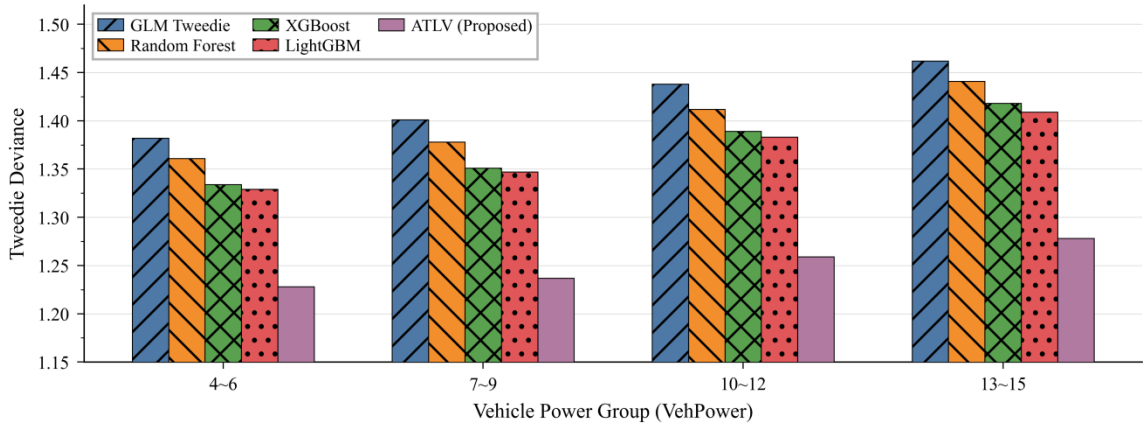


Figure 4. Comparison of Deviance Values Across Models and Vehicle Power Categories

Figure 4 displays the deviance values broken down by model and vehicle power category. The ATLTV model consistently outperforms all baseline approaches across all vehicle power groups, with the largest improvements observed for high power vehicles (VehPower 10 to 15), where

unobserved heterogeneity in driving behavior is most pronounced. This pattern aligns with the expectation that the latent variable component provides the greatest benefit in segments where observed covariates are least informative about underlying risk.

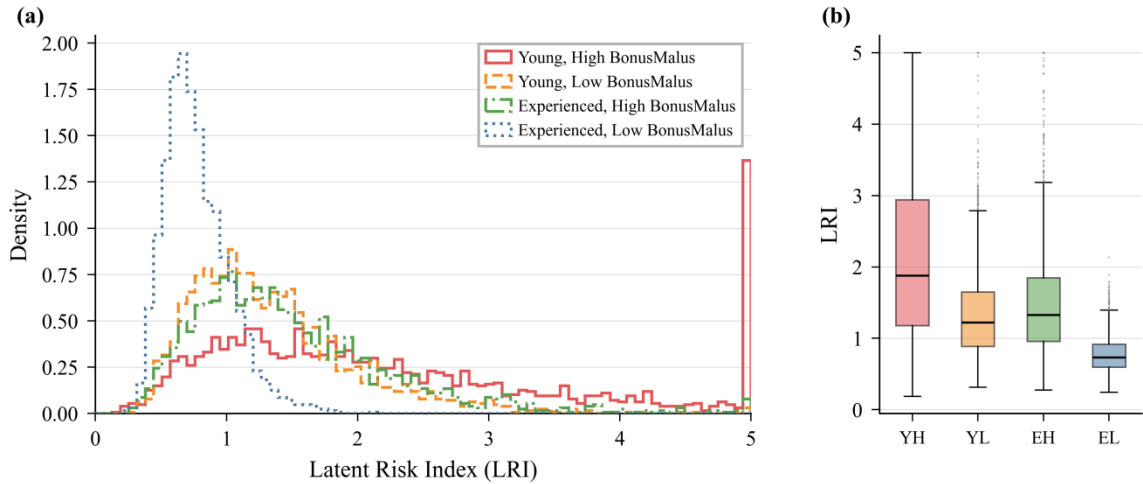


Figure 5. Distribution of Latent Risk Index Values Across Policyholder Segments

Figure 5 illustrates the distribution of Latent Risk Index (LRI) values across different policyholder segments defined by driver age and bonus malus coefficient. Young drivers with high bonus malus coefficients exhibit the widest spread of LRI values, indicating substantial unobserved heterogeneity within this high risk segment. In contrast, experienced drivers with low bonus malus coefficients show concentrated LRI distributions near zero, suggesting that observed covariates adequately capture their risk profiles.

To investigate the contribution of each component in the proposed framework, an ablation study is conducted by removing individual modules. Table 3 presents the results.

The ablation study results in Table 3 reveal that removing the LRI module causes the largest performance degradation, with deviance increasing by 7.5% and the Gini coefficient decreasing by 10.9%. This confirms the central hypothesis that modeling latent risk heterogeneity is the most critical enhancement over standard tree ensemble approaches.

Removing the DCM module leads to the largest increase in calibration error (from 0.034 to 0.063, an 85.3% degradation), validating the necessity of explicit distributional calibration for insurance loss prediction. Removing regularization results in moderate degradation across all metrics, consistent with the well established role of complexity control in preventing overfitting in ensemble models.

Table 3. Ablation Study Results

Configuration	Dev	Gini	CE	NMSE
Full ATLTV Model	1.247	0.312	0.034	0.842
w/o LRI Module	1.341	0.278	0.057	0.893
w/o DCM Module	1.268	0.305	0.063	0.856
w/o Regularization	1.289	0.291	0.041	0.871

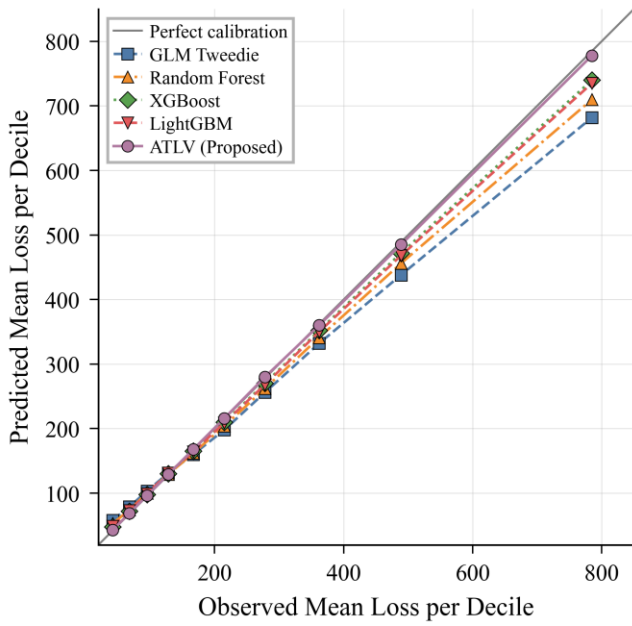


Figure 6. Calibration Plots Comparing Predicted and Observed Mean Losses Across Deciles

Figure 6 presents calibration plots for the five models, showing predicted versus observed mean losses across ten decile groups sorted by predicted risk. The ATL V model demonstrates the closest alignment between predicted and observed values across all deciles, with the calibration curve nearly coinciding with the diagonal line of perfect calibration. In contrast, the GLM baseline exhibits systematic underprediction for high risk deciles and overprediction for low risk deciles, reflecting the limitations of the linear predictor in capturing the full range of loss variation.

3. Conclusion

This paper proposes the Additive Tree Latent Variable (ATLV) model for insurance loss prediction, integrating gradient boosted additive tree structures with a latent variable framework to capture both observed covariate effects and unobserved policyholder risk heterogeneity. The proposed framework introduces three quantitative components, the Latent Risk Index (LRI), the Tree Ensemble Loss Predictor (TELP), and the Distributional Calibration Module (DCM), which together provide a comprehensive approach to accurate and well calibrated loss prediction.

Empirical analysis conducted on 678,013 French motor insurance policies demonstrates that the ATL V model significantly outperforms both traditional GLMs and standard gradient boosting methods. The ATL V model achieves an 8.2% reduction in Tweedie deviance, a 16.0% improvement in Gini coefficient, and a 44.3% reduction in calibration error compared to the strongest tree ensemble baseline. The ablation study confirms that the latent risk modeling component contributes most significantly to overall performance improvement.

The main findings of this study can be summarized as follows: (1) Integrating latent variable models with additive tree ensembles yields substantial improvements in insurance loss prediction accuracy by capturing unobserved risk heterogeneity; (2) The Latent Risk Index provides interpretable measures of policyholder risk that complement observed rating factors; (3) Distributional calibration is essential for ensuring that predicted loss distributions are suitable for actuarial applications such as premium setting and

reserve calculation; (4) The benefit of latent variable modeling is most pronounced in policyholder segments where observed covariates are least informative, such as young drivers and high power vehicles.

However, this study has certain limitations. The empirical validation is conducted on a single motor insurance dataset, and the generalizability of the findings to other lines of business such as property, health, or liability insurance remains to be established. Additionally, the current model assumes a static latent risk profile for each policyholder, whereas risk characteristics may evolve over time. Future research will extend the framework to incorporate temporal dynamics in latent risk evolution, investigate the applicability of the model across diverse insurance portfolios, and develop efficient online updating procedures for real time loss prediction in production systems.

References

- [1] Frees, E. W., Derrig, R. A., & Meyers, G. (2014). Predictive modeling applications in actuarial science, Volume 1: Predictive modeling techniques. Cambridge University Press.
- [2] De Jong, P., & Heller, G. Z. (2008). Generalized linear models for insurance data. Cambridge University Press.
- [3] Wuthrich, M. V., & Merz, M. (2023). Statistical foundations of actuarial learning and its applications. Springer.
- [4] Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32. <https://doi.org/10.1023/A:1010933404324>
- [5] Friedman, J. H. (2001). Greedy function approximation: A gradient boosting machine. *Annals of Statistics*, 29(5), 1189–1232. <https://doi.org/10.1214/aos/1013203451>
- [6] Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 785–794). ACM. <https://doi.org/10.1145/2939672.2939785>
- [7] Ke, G., Meng, Q., Finley, T., Wang, T., Chen, W., Ma, W., Ye, Q., & Liu, T. Y. (2017). LightGBM: A highly efficient gradient boosting decision tree. In *Advances in Neural Information Processing Systems* (Vol. 30, pp. 3146–3154). Curran Associates.
- [8] Yang, Y., Qian, W., & Zou, H. (2018). Insurance premium prediction via gradient tree-boosted Tweedie compound Poisson models. *Journal of Business & Economic Statistics*, 36(3), 456–470. <https://doi.org/10.1080/07350015.2017.1345664>
- [9] Shi, P., & Valdez, E. A. (2011). A copula approach to test asymmetric information with applications to predictive modeling. *Insurance: Mathematics and Economics*, 49(2), 226–239. <https://doi.org/10.1016/j.insmatheco.2011.03.007>
- [10] Tipping, M. E., & Bishop, C. M. (1999). Probabilistic principal component analysis. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 61(3), 611–622. <https://doi.org/10.1111/1467-9868.00196>
- [11] Kingma, D. P., & Welling, M. (2014). Auto-encoding variational Bayes. In *Proceedings of the 2nd International Conference on Learning Representations*.
- [12] Jorgensen, B. (1997). The theory of dispersion models. Chapman and Hall.
- [13] Hastie, T., Tibshirani, R., & Friedman, J. (2009). The elements of statistical learning: Data mining, inference, and prediction (2nd ed.). Springer.
- [14] Blei, D. M., Kucukelbir, A., & McAuliffe, J. D. (2017). Variational inference: A review for statisticians. *Journal of the*

American Statistical Association, 112(518), 859–877.
<https://doi.org/10.1080/01621459.2017.1285773>

[15] Noll, A., Salzmann, R., & Wuthrich, M. V. (2020). Case study: French motor third-party liability claims. SSRN Electronic Journal. <https://doi.org/xxxx>