

WASSERSTEIN BOOSTING TREES ALGORITHM FOR COUNT DATA, WITH APPLICATION TO CLAIM FREQUENCIES IN MOTOR INSURANCE

Michel Denuit, Marie Michaelides, Julien
Trufin, Harrison Verelst

LIDAM Discussion Paper ISBA
2025 / 24

ISBA

Voie du Roman Pays 20 - L1.04.01

B-1348 Louvain-la-Neuve

Email : lidam-library@uclouvain.be

<https://uclouvain.be/en/research-institutes/lidam/isba/publication>

Wasserstein boosting trees algorithm for count data, with application to claim frequencies in motor insurance

Michel Denuit

Institute of Statistics, Biostatistics and Actuarial Science - ISBA

Louvain Institute of Data Analysis and Modeling - LIDAM

UCLouvain

Louvain-la-Neuve, Belgium

Marie Michaelides

School of Mathematical & Computer Sciences

Actuarial Mathematics & Statistics

Heriot-Watt University

Edinburgh, United Kingdom

Julien Trufin

Department of Mathematics

Université Libre de Bruxelles (ULB)

Brussels, Belgium

Harrison Verelst

Detralytics

Brussels, Belgium

November 6, 2025

Abstract

This paper proposes a variant of the well-known boosting trees algorithm to estimate conditional distributions. Since regression trees partition observations into subgroups, the corresponding empirical distributions can be used to define the splitting criterion. Precisely, the parametric approach using Poisson deviance is replaced with a non-parametric one maximizing probabilistic distances between empirical distributions in child nodes. Proceeding in this way, the actuary obtains an estimated conditional distribution for the response, from which a conditional mean can be derived as well as any other quantity of interest in risk management. The numerical performances of the proposed method are assessed with simulated data while a case study demonstrates its usefulness for insurance applications.

Keywords: Wasserstein distance, regression trees, boosting, conditional distribution, count data.

1 Introduction

The present paper proposes to use probabilistic distances in the splitting criterion to grow regression trees, in a boosting approach. Boosting emerged from the field of machine learning and became rapidly popular among actuaries. Broadly speaking, boosting is an iterative procedure building the score sequentially, adding terms that weakly improve on the estimations obtained from the preceding step. Regression trees with limited depth are often used as weak learners in insurance studies, following [Friedman \[2001\]](#). We refer the reader to [Breiman et al. \[1984\]](#) for a general introduction to trees.

Boosting requires the specification of a loss function to be optimized for estimating the score. In that respect, actuaries generally adopt a parametric approach by selecting loss functions corresponding to negative log-likelihood, or deviance functions. For count data, Poisson or Negative Binomial distributions are often used to study numbers of claims reported by policyholders to their insurer. We refer the readers to [Denuit et al. \[2020\]](#) for an extensive treatment of boosting in the context of insurance.

Since trees partition the database in groups of observations corresponding to terminal leaves, many other quantities of interest can be computed from their empirical distribution, not only conditional means. Also, considering that trees produce estimated distributions in each terminal leaf opens the door to new partition rules based on distances between the resulting empirical distributions. Chapter 9 in [Denuit et al. \[2005\]](#) offers a general presentation of the topic with applications to actuarial science. Probability metrics have been extensively used in risk theory, to assess the quality of the approximation of the individual model by its collective counterpart. Here, the idea is to split a group in two subgroups to make the corresponding empirical distributions the furthest apart according to some relevant probabilistic distance.

Probabilistic distances have already been used in combination with regression trees, considering random forests. [Du et al. \[2021\]](#) proposed to adapt Breiman’s original splitting criterion in terms of Wasserstein distance between empirical distributions. Their Wasserstein random forest algorithm then produces an estimate of the conditional distribution function, allowing the actuary to derive a variety of quantities of interest, not just the conditional mean. Given the excellent performances of boosting algorithms in insurance studies, the present paper proposes a Wasserstein boosting trees algorithm. Instead of averaging (in the sense of probabilistic mixtures) estimation obtained from a large number of trees, the conditional distribution is learned sequentially, by updating estimation at previous step.

The remainder of this paper is organized as follows. Section 2 recalls probabilistic distances, with special emphasis on the Wasserstein distance that has been successfully used to grow trees by [Du et al. \[2021\]](#). Section 3 then proposes a new boosting algorithm adopting this distance in the splitting criterion. Section 4 illustrates the numerical performances of this new approach on simulated data. A case study with motor insurance claim frequencies is performed in Section 5. The final Section 6 briefly concludes the paper, summarizing the main findings and discussing possible extensions.

2 Wasserstein distance

Probability metrics, including the Wasserstein one, have successfully been used in actuarial science to assess the accuracy of the approximation of the individual model by its collective counterpart. Given two random variables V and W with respective probability distributions π_V and π_W , recall that the Wasserstein distance $\mathcal{W}$, also known as the Dudley or Kantorovitch distance, is defined as

$$\mathcal{W}(\pi_V, \pi_W) = \int_{-\infty}^{\infty} |\bar{F}_V(t) - \bar{F}_W(t)| dt = \int_0^1 |F_V^{-1}(u) - F_W^{-1}(u)| du$$

where $\bar{F}_V = 1 - F_V$ denotes the excess function corresponding to the distribution function F_V of V and where the quantile function F_V^{-1} is defined for a probability level u as

$$F_V^{-1}(u) = \inf\{v \in \mathbb{R} | F_V(v) \geq u\}.$$

Since the representation $\mathcal{W}(\pi_V, \pi_W) = E[|F_W^{-1}(U) - F_V^{-1}(U)|]$ holds true with U uniformly distributed over the unit interval, $\mathcal{W}(\pi_V, \pi_W)$ appears to be the lower bound on $E[|K - L|]$ over all random couples (K, L) with marginal distribution functions F_V and F_W (attained when K and L are perfectly positively dependent, or comonotonic).

For two counting random variables M and N with respective probability distributions π_M and π_N , the Wasserstein distance $\mathcal{W}$ can be simply rewritten as

$$\mathcal{W}(\pi_M, \pi_N) = \sum_{j=0}^{\infty} |\bar{F}_M(j) - \bar{F}_N(j)|.$$

It can easily be computed when M and N only assume a few positive values, as it is the case for claim frequencies at policy-level in personal lines. The series appearing in $\mathcal{W}(\pi_M, \pi_N)$ then reduces to just a few terms and is fast to compute.

Property 9.6.3 in [Denuit et al. \[2005\]](#) shows that $\mathcal{W}$ reduces to the difference of expectations under stochastic dominance. Recall that N is smaller than M in stochastic dominance, denoted as $N \leq_{ST} M$, when the inequality $P[N > t] \leq P[M > t]$ holds for all real t . It is then easy to see that

$$N \leq_{ST} M \Rightarrow \mathcal{W}(\pi_M, \pi_N) = \int_0^1 (F_M^{-1}(u) - F_N^{-1}(u)) du = E[M] - E[N].$$

Wasserstein distance thus reduces to the difference in pure premiums in that case. For instance, if N obeys the Poisson distribution with mean λ_1 and M obeys the Poisson distribution with mean λ_2 ,

$$\lambda_1 < \lambda_2 \Rightarrow N \leq_{ST} M \Rightarrow \mathcal{W}(\pi_M, \pi_N) = \lambda_2 - \lambda_1.$$

Considering claim counts M and N , the Wasserstein distance is the sum, over each possible claims number value j , of the difference in probability to report strictly more than j claims. In a frequency context, one would like to be able to distinguish policyholders with low probability to have claims or low probability to have a large number of claims with policyholders with high probability to have one or more than one claim. Using the Wasserstein distance in this context thus seems to be appealing. Since Poisson distributions are ordered with respect to $\leq_{ST}$ in their mean parameter, maximizing Wasserstein distance amounts to maximize differences in expected claim frequencies.

3 Wasserstein boosting trees

3.1 Setting

Supervised learning aims to provide an estimation of the conditional expectation $\mu(\mathbf{x}) = E[Y|\mathbf{X} = \mathbf{x}]$ for a response Y and a set of features $\mathbf{X} = (X_1, \dots, X_p) \in \mathcal{X} \subseteq \mathbb{R}^p$. In many applications of practical relevance, the additional information encoded in the conditional distribution function $F_Y(\cdot|\mathbf{x})$ defined as $F_Y(y|\mathbf{x}) = P[Y \leq y|\mathbf{X} = \mathbf{x}]$ is also of great interest. Du et al. [2021] propose natural variants of random forests to estimate $F_Y(\cdot|\mathbf{x})$, called Wasserstein random forests. In this paper, we present a variant of boosting trees to estimate $F_Y(\cdot|\mathbf{x})$ for count data. Precisely, the response Y is assumed to be discrete with values in $\mathbb{N} = \{0, 1, 2, \dots\}$. We denote by $\pi(\mathbf{x}, \cdot)$ the probability mass function associated with the conditional distribution function $F_Y(\cdot|\mathbf{x})$, i.e. $\pi(\mathbf{x}, k) = P[Y = k|\mathbf{X} = \mathbf{x}]$, $k \in \mathbb{N}$. We aim to estimate the conditional probabilities $\pi(\mathbf{x}, k)$ using a boosting approach given a data set $\mathcal{D} = \{(y_i, \mathbf{x}_i), i = 1, \dots, n\}$.

We denote by M the number of trees in the ensemble. Each decision tree is built on a random sample of the training set $\mathcal{D}$, denoted $\mathcal{D}_m^*$ for the m -th tree, taken without replacement. The fraction of training set used at each iteration to produce the random samples $\mathcal{D}_m^*$ is the bagging fraction and is denoted by α . A typical value for α is 0.5.

The m -th tree partitions the feature space $\mathcal{X}$ into disjoint subspaces $\{\chi_j^{(m)}\}_{j \in \mathcal{J}_m}$, where $\mathcal{J}_m$ is a set of indexes. We denote by $\chi_{j(\mathbf{x})}^{(m)}$, $j(\mathbf{x}) \in \mathcal{J}_m$, the subspace of the feature space $\mathcal{X}$ induced by the m -th tree that contains $\mathbf{x}$, and by $N_m(\mathbf{x})$ the number of observations $(y_i, \mathbf{x}_i)$ in $\mathcal{D}_m^*$ such that $\mathbf{x}_i \in \chi_{j(\mathbf{x})}^{(m)}$. As proposed in Du et al. [2021], the estimate $\hat{\pi}_m(\mathbf{x}, k)$ of the conditional probability $\pi(\mathbf{x}, k)$ at point $\mathbf{x}$ given by the m -th tree is the empirical measure associated with the observations $(y_i, \mathbf{x}_i)$ that fall into the same terminal node as $\mathbf{x}$ (i.e. such that $\mathbf{x}_i \in \chi_{j(\mathbf{x})}^{(m)}$), that is

$$\hat{\pi}_m(\mathbf{x}, k) = \sum_{(y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} \frac{\mathbb{I}[\mathbf{x}_i \in \chi_{j(\mathbf{x})}^{(m)}, y_i = k]}{N_m(\mathbf{x})}, \quad (3.1)$$

where $\mathbb{I}[\cdot]$ denotes the indicator function, equal to 1 if the event appearing in the brackets is realized and to 0 otherwise.

3.2 Iteration steps

Initialization

We initialize the boosting algorithm with the estimate

$$\hat{\pi}_0(\mathbf{x}, k) = \frac{1}{n} \sum_{i=1}^n \mathbb{I}[y_i = k] \quad (3.2)$$

corresponding to the empirical distribution of the response. We thus start with the same probability mass function for every $\mathbf{x}$. The starting point (3.2) corresponds to the null model comprising only an intercept initializing the boosting procedure.

Step 1

From the initial estimate $\widehat{\pi}_0(\mathbf{x}, k)$, we build the first tree and we compute the estimate $\widehat{\pi}_1(\mathbf{x}, k)$, as given by (3.1) with $m = 1$. The initial estimate $\widehat{\pi}_0(\mathbf{x}, k)$ and the estimate from the first tree $\widehat{\pi}_1(\mathbf{x}, k)$ are then combined to produce the estimate of the conditional probability $\pi(\mathbf{x}, k)$ after the first tree of the boosting algorithm. We denote this estimate as $\widehat{\pi}_{\rightarrow 1}(\mathbf{x}, k)$. The latter is given by

$$\widehat{\pi}_{\rightarrow 1}(\mathbf{x}, k) = \widehat{c}_1(\mathbf{x}) \times \widehat{\beta}_1(\mathbf{x}, k) \times \widehat{\pi}_0(\mathbf{x}, k), \quad (3.3)$$

where $\widehat{\beta}_1(\mathbf{x}, k)$ satisfies

$$\frac{1}{N_1(\mathbf{x})} \sum_{(y_i, \mathbf{x}_i) \in \mathcal{D}_1^*} \mathbb{I} \left[\mathbf{x}_i \in \mathcal{X}_{j(\mathbf{x})}^{(1)} \right] (\widehat{\pi}_0(\mathbf{x}_i, k) \times \widehat{\beta}_1(\mathbf{x}, k)) = \widehat{\pi}_1(\mathbf{x}, k), \quad (3.4)$$

that is,

$$\widehat{\beta}_1(\mathbf{x}, k) = \frac{N_1(\mathbf{x}) \widehat{\pi}_1(\mathbf{x}, k)}{\sum_{(y_i, \mathbf{x}_i) \in \mathcal{D}_1^*} \mathbb{I} \left[\mathbf{x}_i \in \mathcal{X}_{j(\mathbf{x})}^{(1)} \right] \widehat{\pi}_0(\mathbf{x}_i, k)}, \quad (3.5)$$

and where the normalizing constant

$$\widehat{c}_1(\mathbf{x}) = \left(\sum_{k=0}^{\infty} \widehat{\beta}_1(\mathbf{x}, k) \times \widehat{\pi}_0(\mathbf{x}, k) \right)^{-1} \quad (3.6)$$

guarantees that $\sum_{k=0}^{\infty} \widehat{\pi}_{\rightarrow 1}(\mathbf{x}, k) = 1$. The new estimate $\widehat{\beta}_1(\mathbf{x}, k) \times \widehat{\pi}_0(\mathbf{x}, k)$ guarantees that its average over the observations that fall into a terminal node j of the first tree corresponds to the empirical estimate $\widehat{\pi}_1(\mathbf{x}, k)$ associated with that terminal node j . Similarly to how the score in the classical Boosting algorithm is updated at each step, the function $\widehat{\beta}_1(\mathbf{x}, k)$ defined in (3.5) corrects the initial estimate $\widehat{\pi}_0(\mathbf{x}, k)$ with the information $\widehat{\pi}_1(\mathbf{x}, k)$ gathered from the first tree to produce the next estimate $\widehat{\pi}_{\rightarrow 1}(\mathbf{x}, k)$.

Step 2

From the estimate after the first tree $\widehat{\pi}_{\rightarrow 1}(\mathbf{x}, k)$, we build the second tree and compute from this second tree the estimate $\widehat{\pi}_2(\mathbf{x}, k)$. Similarly to the first step, $\widehat{\pi}_{\rightarrow 1}(\mathbf{x}, k)$ is updated by $\widehat{\pi}_2(\mathbf{x}, k)$ to produce the estimate of the conditional probability $\pi(\mathbf{x}, k)$ resulting from the first two trees. This estimate is denoted as $\widehat{\pi}_{\rightarrow 2}(\mathbf{x}, k)$ and is given by

$$\widehat{\pi}_{\rightarrow 2}(\mathbf{x}, k) = \widehat{c}_2(\mathbf{x}) \times \widehat{\beta}_2(\mathbf{x}, k) \times \widehat{\pi}_{\rightarrow 1}(\mathbf{x}, k), \quad (3.7)$$

where $\widehat{\beta}_2(\mathbf{x}, k)$ satisfies

$$\frac{1}{N_2(\mathbf{x})} \sum_{(y_i, \mathbf{x}_i) \in \mathcal{D}_2^*} \mathbb{I} \left[\mathbf{x}_i \in \mathcal{X}_{j(\mathbf{x})}^{(2)} \right] (\widehat{\pi}_{\rightarrow 1}(\mathbf{x}_i, k) \times \widehat{\beta}_2(\mathbf{x}, k)) = \widehat{\pi}_2(\mathbf{x}, k), \quad (3.8)$$

that is,

$$\widehat{\beta}_2(\mathbf{x}, k) = \frac{N_2(\mathbf{x}) \widehat{\pi}_2(\mathbf{x}, k)}{\sum_{(y_i, \mathbf{x}_i) \in \mathcal{D}_2^*} \mathbb{I} \left[\mathbf{x}_i \in \chi_{j(\mathbf{x})}^{(2)} \right] \widehat{\pi}_{\rightarrow 1}(\mathbf{x}_i, k)}, \quad (3.9)$$

and where the normalizing constant

$$\widehat{c}_2(\mathbf{x}) = \left(\sum_{k=0}^{\infty} \widehat{\beta}_2(\mathbf{x}, k) \times \widehat{\pi}_{\rightarrow 1}(\mathbf{x}, k) \right)^{-1} \quad (3.10)$$

guarantees that $\sum_{k=0}^{\infty} \widehat{\pi}_{\rightarrow 2}(\mathbf{x}, k) = 1$.

Step m

More generally, with the convention that $\widehat{\pi}_{\rightarrow 0}(\mathbf{x}, k) = \widehat{\pi}_0(\mathbf{x}, k)$, the estimate of the conditional probability $\pi(\mathbf{x}, k)$ produced by the first m trees ($m = 1, \dots, M$) of the boosting algorithm, denoted $\widehat{\pi}_{\rightarrow m}(\mathbf{x}, k)$, is given by

$$\widehat{\pi}_{\rightarrow m}(\mathbf{x}, k) = \widehat{c}_m(\mathbf{x}) \times \widehat{\beta}_m(\mathbf{x}, k) \times \widehat{\pi}_{\rightarrow m-1}(\mathbf{x}, k), \quad (3.11)$$

where $\widehat{\beta}_m(\mathbf{x}, k)$ satisfies

$$\frac{1}{N_m(\mathbf{x})} \sum_{(y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} \mathbb{I} \left[\mathbf{x}_i \in \chi_{j(\mathbf{x})}^{(m)} \right] (\widehat{\pi}_{\rightarrow m-1}(\mathbf{x}_i, k) \times \widehat{\beta}_m(\mathbf{x}, k)) = \widehat{\pi}_m(\mathbf{x}, k), \quad (3.12)$$

that is,

$$\widehat{\beta}_m(\mathbf{x}, k) = \frac{N_m(\mathbf{x}) \widehat{\pi}_m(\mathbf{x}, k)}{\sum_{(y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} \mathbb{I} \left[\mathbf{x}_i \in \chi_{j(\mathbf{x})}^{(m)} \right] \widehat{\pi}_{\rightarrow m-1}(\mathbf{x}_i, k)}, \quad (3.13)$$

and where the normalizing constant

$$\widehat{c}_m(\mathbf{x}) = \left(\sum_{k=0}^{\infty} \widehat{\beta}_m(\mathbf{x}, k) \times \widehat{\pi}_{\rightarrow m-1}(\mathbf{x}, k) \right)^{-1} \quad (3.14)$$

guarantees that $\sum_{k=0}^{\infty} \widehat{\pi}_{\rightarrow m}(\mathbf{x}, k) = 1$.

Interpretation

At iteration m , the first $m - 1$ trees produce the estimate $\widehat{\pi}_{\rightarrow m-1}(\mathbf{x}, k)$ of the conditional probability $\pi(\mathbf{x}, k)$ and the m -th tree provides the estimate $\widehat{\pi}_m(\mathbf{x}, k)$ of $\pi(\mathbf{x}, k)$. These two estimates are then combined as described in (3.11) and (3.12) to produce the new estimate $\widehat{\pi}_{\rightarrow m}(\mathbf{x}, k)$ of $\pi(\mathbf{x}, k)$ resulting from the first m trees. Contrary to $\widehat{\pi}_{\rightarrow m-1}(\mathbf{x}, k)$, the new estimate $\widehat{\beta}_m(\mathbf{x}, k) \times \widehat{\pi}_{\rightarrow m-1}(\mathbf{x}, k)$ guarantees that its average over the observations that fall into a terminal node j of the m -th tree corresponds to the empirical estimate $\widehat{\pi}_m(\mathbf{x}, k)$ associated with that terminal node j .

3.3 Splitting rule

The splitting criterion that is maximized at each node of the m -th tree is based on the Wasserstein distance. Suppose that the node under consideration in the m -th tree consists in a subset ϕ of the feature space $\mathcal{X}$. Hence, any standardized binary split defines a partition $\phi_L \cup \phi_R$ of ϕ . Let N (resp. N_L and N_R) be the number of observations in $\mathcal{D}_m^*$ such that $\mathbf{x}_i \in \phi$ (resp. $\mathbf{x}_i \in \phi_L$ and $\mathbf{x}_i \in \phi_R$). The splitting criterion we consider here takes the form

$$L(\phi_L, \phi_R) = \frac{N_L}{N} \mathcal{W}(\hat{\pi}_L, \hat{\pi}_{L \rightarrow m-1}) + \frac{N_R}{N} \mathcal{W}(\hat{\pi}_R, \hat{\pi}_{R \rightarrow m-1}), \quad (3.15)$$

where

$$\hat{\pi}_{L \rightarrow m-1}(k) = \frac{1}{N_L} \sum_{(y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} \mathbf{I}[\mathbf{x}_i \in \phi_L] \hat{\pi}_{\rightarrow m-1}(\mathbf{x}_i, k) \quad (3.16)$$

$$\hat{\pi}_{R \rightarrow m-1}(k) = \frac{1}{N_R} \sum_{(y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} \mathbf{I}[\mathbf{x}_i \in \phi_R] \hat{\pi}_{\rightarrow m-1}(\mathbf{x}_i, k) \quad (3.17)$$

and

$$\hat{\pi}_L(k) = \frac{1}{N_L} \sum_{(y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} \mathbf{I}[\mathbf{x}_i \in \phi_L, y_i = k] \quad (3.18)$$

$$\hat{\pi}_R(k) = \frac{1}{N_R} \sum_{(y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} \mathbf{I}[\mathbf{x}_i \in \phi_R, y_i = k]. \quad (3.19)$$

At each node of the m -th tree, the best split maximizes the Wasserstein distances between the probability distributions already estimated from the first $m-1$ trees and the empirical distributions resulting from the split. The partition $\phi_L \cup \phi_R$ of ϕ obtained from our splitting criterion is such that the current estimates $\hat{\pi}_{L \rightarrow m-1}$ and $\hat{\pi}_{R \rightarrow m-1}$ are “the most different” from the empirical estimates $\hat{\pi}_L$ and $\hat{\pi}_R$, respectively, as measured by the Wasserstein distance. The m -th tree is thus built in a way that the resulting partition $\{\mathcal{X}_j^{(m)}\}_{j \in \mathcal{J}_m}$ of the feature space highlights subsets of the feature space for which the current estimates obtained after $m-1$ iterations are not satisfactory and hence need to be readjusted.

3.4 Shrinkage parameters

In boosting, it is often helpful to slow down the learning speed by adding only a small amount of the fit of the best-performing base-learner to the current additive score. This is achieved by multiplying the new effect entering the score with a shrinkage coefficient (a typical value is 0.1). This also appears to be useful in the Wasserstein boosting trees algorithm, as explained next.

The idea is to decrease $\hat{\beta}_m(\mathbf{x}, k)$ to slow down the learning process. Formally, let τ_1 and τ_2 be such that $0 \leq \tau_1 < 1$ and $\tau_2 > 1$. Define

$$\hat{\beta}_m^{\text{shrink}}(\mathbf{x}, k) = \min \left(\max \left(\hat{\beta}_m(\mathbf{x}, k), \tau_1 \right), \tau_2 \right), \quad (3.20)$$

where $\hat{\beta}_m(\mathbf{x}, k)$ is given in (3.13). The updating procedure then becomes

$$\hat{\pi}_{\rightarrow m}(\mathbf{x}, k) = \hat{c}_m(\mathbf{x}) \times \hat{\beta}_m^{\text{shrink}}(\mathbf{x}, k) \times \hat{\pi}_{\rightarrow m-1}(\mathbf{x}, k), \quad (3.21)$$

with

$$\widehat{c}_m(\mathbf{x}) = \left(\sum_{k=0}^{\infty} \widehat{\beta}_m^{\text{shrink}}(\mathbf{x}, k) \times \widehat{\pi}_{\rightarrow m-1}(\mathbf{x}, k) \right)^{-1}. \quad (3.22)$$

With the proposed shrinkage procedure, we have $\tau_1 \leq \widehat{\beta}_m^{\text{shrink}}(\mathbf{x}, k) \leq \tau_2$, so that we slow down the learning rate of the boosting procedure.

3.5 Exposure-to-risk

Context

In an insurance context, observations are not always made over the same time period. We denote by e the length of the observation period which is generally referred to as the exposure-to-risk. One observation of the data set $\mathcal{D}$ is thus also described by its exposure-to-risk, so that we have $\mathcal{D} = \{(e_i, y_i, \mathbf{x}_i), i = 1, \dots, n\}$. Consequently, the previously described algorithm must be adapted to account for this specificity when dealing with count numbers in insurance.

Let $E_m(\mathbf{x})$ be the sum of the exposures-to-risk e_i associated with the observations $(e_i, y_i, \mathbf{x}_i)$ in $\mathcal{D}_m^*$ such that $\mathbf{x}_i \in \chi_{j(\mathbf{x})}^{(m)}$, that is,

$$E_m(\mathbf{x}) = \sum_{(e_i, y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} e_i \mathbb{I} \left[\mathbf{x}_i \in \chi_{j(\mathbf{x})}^{(m)} \right].$$

Then, (3.1) becomes

$$\widehat{\pi}_m(\mathbf{x}, k) = \sum_{(e_i, y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} \frac{e_i \mathbb{I} \left[\mathbf{x}_i \in \chi_{j(\mathbf{x})}^{(m)}, y_i = k \right]}{E_m(\mathbf{x})}. \quad (3.23)$$

Iteration steps

The boosting algorithm is now initialized with the estimate

$$\widehat{\pi}_0(\mathbf{x}, k) = \frac{\sum_{i=1}^n e_i \mathbb{I} [y_i = k]}{\sum_{i=1}^n e_i}.$$

Then, $\widehat{\pi}_{\rightarrow m}(\mathbf{x}, k)$ is given by

$$\widehat{\pi}_{\rightarrow m}(\mathbf{x}, k) = \widehat{c}_m(\mathbf{x}) \times \widehat{\beta}_m(\mathbf{x}, k) \times \widehat{\pi}_{\rightarrow m-1}(\mathbf{x}, k),$$

where $\widehat{\beta}_m(\mathbf{x}, k)$ satisfies

$$\frac{1}{E_m(\mathbf{x})} \sum_{(e_i, y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} e_i \mathbb{I} \left[\mathbf{x}_i \in \chi_{j(\mathbf{x})}^{(m)} \right] (\widehat{\pi}_{\rightarrow m-1}(\mathbf{x}_i, k) \times \widehat{\beta}_m(\mathbf{x}, k)) = \widehat{\pi}_m(\mathbf{x}, k),$$

that is,

$$\widehat{\beta}_m(\mathbf{x}, k) = \frac{E_m(\mathbf{x}) \widehat{\pi}_m(\mathbf{x}, k)}{\sum_{(e_i, y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} e_i \mathbb{I} \left[\mathbf{x}_i \in \chi_{j(\mathbf{x})}^{(m)} \right] \widehat{\pi}_{\rightarrow m-1}(\mathbf{x}_i, k)},$$

and where the normalizing constant

$$\widehat{c}_m(\mathbf{x}) = \left(\sum_{k=0}^{\infty} \widehat{\beta}_m(\mathbf{x}, k) \times \widehat{\pi}_{\rightarrow m-1}(\mathbf{x}, k) \right)^{-1}$$

guarantees that $\sum_{k=0}^{\infty} \widehat{\pi}_{\rightarrow m}(\mathbf{x}, k) = 1$.

Splitting rule

Again, suppose that the node under consideration in the m -th tree consists in a subset ϕ of the feature space χ . Hence, any standardized binary split defines a partition $\phi_L \cup \phi_R$ of ϕ . Let E (resp. E_L and E_R) be the total exposure-to-risk in $\mathcal{D}_m^*$ such that $\mathbf{x}_i \in \phi$ (resp. $\mathbf{x}_i \in \phi_L$ and $\mathbf{x}_i \in \phi_R$). The splitting criterion we consider here takes the form

$$L(\phi_L, \phi_R) = \frac{E_L}{E} \mathcal{W}(\widehat{\pi}_L, \widehat{\pi}_{L, \rightarrow m-1}) + \frac{E_R}{E} \mathcal{W}(\widehat{\pi}_R, \widehat{\pi}_{R, \rightarrow m-1}),$$

where

$$\begin{aligned} \widehat{\pi}_{L, \rightarrow m-1}(k) &= \frac{1}{E_L} \sum_{(e_i, y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} e_i \mathbb{I}[\mathbf{x}_i \in \phi_L] \widehat{\pi}_{\rightarrow m-1}(\mathbf{x}_i, k) \\ \widehat{\pi}_{R, \rightarrow m-1}(k) &= \frac{1}{E_R} \sum_{(e_i, y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} e_i \mathbb{I}[\mathbf{x}_i \in \phi_R] \widehat{\pi}_{\rightarrow m-1}(\mathbf{x}_i, k) \end{aligned}$$

and

$$\begin{aligned} \widehat{\pi}_L(k) &= \frac{1}{E_L} \sum_{(e_i, y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} e_i \mathbb{I}[\mathbf{x}_i \in \phi_L, y_i = k] \\ \widehat{\pi}_R(k) &= \frac{1}{E_R} \sum_{(e_i, y_i, \mathbf{x}_i) \in \mathcal{D}_m^*} e_i \mathbb{I}[\mathbf{x}_i \in \phi_R, y_i = k]. \end{aligned}$$

4 Numerical illustrations with simulated data sets

In this section, we illustrate the newly proposed Wasserstein boosting trees on several examples. In the first two examples, we show based on simple distribution functions that the estimated distributions obtained from Wasserstein boosting trees converge towards the true distributions. Then, we compare Wasserstein boosting trees with gradient boosting trees on two examples. In the first example, we simulate the numbers of claims from Poisson distributions and compare the performances of the Wasserstein boosting trees algorithm with the gradient boosting trees algorithm using the Poisson deviance as loss function. The second example demonstrates the advantage of the Wasserstein Boosting trees algorithm over the gradient boosting trees algorithm by simulating the numbers of claims from Negative Binomial distributions with the same mean.

4.1 Example 1

In this very simple example, three categorical features $\mathbf{X} = (X_1, X_2, X_3)$ are supposed to be available:

- X_1 = Age: policyholder’s seniority with two levels, labeled as young (y) or old (o), with 50% of the portfolio in each level;
- X_2 = Gender: policyholder’s gender with two levels, labeled as female (f) or male (m), with 50% of the portfolio in each level;
- X_3 = Sport: whether the policyholder’s car is a sports car or not with two levels, labeled as yes or no, with 50% of the portfolio in each level.

The eight risk classes are depicted in Table 1. The variables X_1 , X_2 and X_3 are assumed to be mutually independent.

The response Y is supposed to be the number of claims. The conditional probabilities $\pi(\mathbf{x}, k)$ are given in Table 1. The eight risk classes have the same possible values but have different associated probabilities. The associated probabilities are arbitrarily chosen. In this example, the true conditional probabilities $\pi(\mathbf{x}, k)$ are thus known and we can simulate realizations of the random vector $(Y, \mathbf{X})$. Specifically, we generate $n = 10\,000$ independent realizations $(y_1, \mathbf{x}_1), (y_2, \mathbf{x}_2), \dots$ of $(Y, \mathbf{X})$. A simulation represents a policy that has been observed during a whole year.

	$\pi(\mathbf{x}, k)$					
	$\mathbf{x}$	$k = 0$	$k = 1$	$k = 2$	$k = 3$	$k = 4$
Class 1	$X_1 = y, X_2 = m, X_3 = yes$	0.400	0.400	0.100	0.050	0.050
Class 2	$X_1 = y, X_2 = m, X_3 = no$	0.650	0.150	0.100	0.050	0.050
Class 3	$X_1 = o, X_2 = m, X_3 = yes$	0.600	0.300	0.050	0.025	0.025
Class 4	$X_1 = o, X_2 = m, X_3 = no$	0.800	0.100	0.050	0.025	0.025
Class 5	$X_1 = y, X_2 = f, X_3 = yes$	0.500	0.300	0.050	0.100	0.050
Class 6	$X_1 = y, X_2 = f, X_3 = no$	0.750	0.050	0.050	0.100	0.050
Class 7	$X_1 = o, X_2 = f, X_3 = yes$	0.700	0.175	0.050	0.050	0.025
Class 8	$X_1 = o, X_2 = f, X_3 = no$	0.850	0.050	0.025	0.050	0.025

Table 1: $\pi(\mathbf{x}, k)$ for the eight risk classes.

For each tree, a sequence of standardized binary splits is made recursively by maximizing the splitting rule discussed in Section 3.3. The size of the trees is controlled by the interaction depth ID. Trees with an interaction depth equal to ID correspond to trees with ID non-terminal nodes. By setting ID = 1, only single-split regression trees are produced, which allows capturing only the main effects of the features in the score. For ID = 2, two-way interactions are permitted, and for ID = 3, three-way interactions are allowed, and so on. A larger value of ID allows the model to learn deeper patterns in the data. In this way the m -th tree partitions the feature space $\mathcal{X}$ into disjoint subsets $\{\mathcal{X}_j^{(m)}, j = 1, \dots, \text{ID} + 1\}$.

We run the Wasserstein boosting trees algorithm with $\text{ID} \in \{1, 2, 3\}$ and shrinkage parameters $\tau_1 = 1 - \tau$ and $\tau_2 = 1 + \tau$ with $\tau \in \{0.5, 0.05\}$. The results are depicted in Figure

1 (for ID = 1), Figure 2 (for ID = 2) and Figure 3 (for ID = 3). For each risk profile $\mathbf{x}$, we show the Wasserstein distance between the true conditional probabilities $\pi(\mathbf{x}, k)$ and the estimated ones $\hat{\pi}_{\rightarrow m}(\mathbf{x}, k)$ with respect to the number of iterations m . With ID $\in \{1, 2\}$, one sees that for any risk class, the estimate $\hat{\pi}_{\rightarrow m}(\mathbf{x}, k)$ does not allow to recover the true conditional probability $\pi(\mathbf{x}, k)$ whatever the number of iteration m . For ID = 3, one sees that for any risk class, the estimate $\hat{\pi}_{\rightarrow m}(\mathbf{x}, k)$ tends to the conditional probability $\pi(\mathbf{x}, k)$ in the sense that $\mathcal{W}(\pi(\mathbf{x}, .), \hat{\pi}_{\rightarrow m}(\mathbf{x}, .))$ tends to 0 after 10 iterations for $\tau = 0.5$ and after 30 iterations for $\tau = 0.05$. This means that for ID = 3, starting from the empirical distribution of the response for every $\mathbf{x}$ and correcting after each fitted tree the previously estimated conditional distribution $\hat{\pi}_{\rightarrow m-1}(\mathbf{x}, .)$ with the information coming from the new fitted tree $\hat{\pi}_m(\mathbf{x}, .)$, we recover the true conditional probability $\pi(\mathbf{x}, k)$. The smaller the shrinkage parameter is, the slowest the learning rate and the smoothest the convergence as it prevents large corrections at each step.

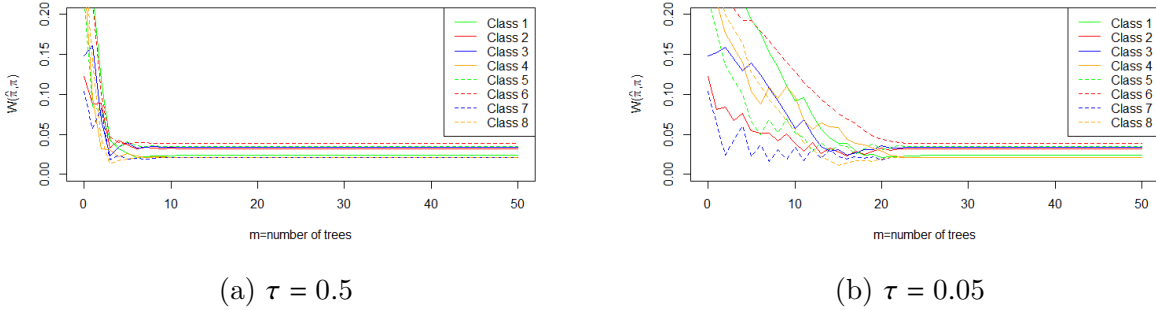


Figure 1: Wasserstein distance between $\pi(\mathbf{x}, .)$ and $\hat{\pi}_{\rightarrow m}(\mathbf{x}, .)$ (obtained with ID = 1) with respect to the number of trees m for the eight risk classes.

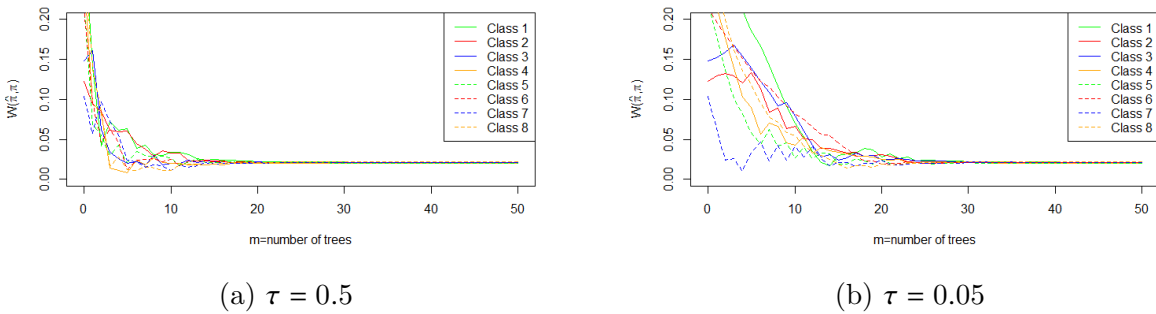
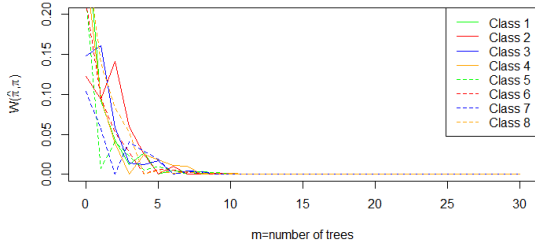


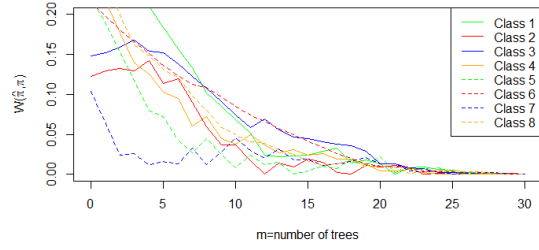
Figure 2: Wasserstein distance between $\pi(\mathbf{x}, .)$ and $\hat{\pi}_{\rightarrow m}(\mathbf{x}, .)$ (obtained with ID = 2) with respect to the number of trees m for the eight risk classes.

4.2 Example 2

We consider the same example as in the previous section, except that the conditional probabilities $\pi(\mathbf{x}, k)$ for risk classes 7 and 8 are now given in Table 2. Compared to Table 1,



(a) $\tau = 0.5$



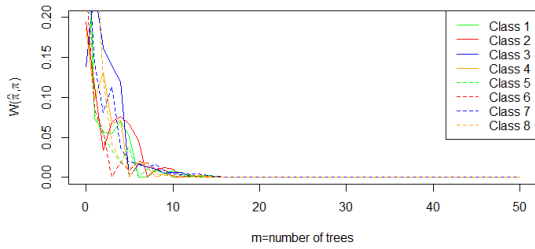
(b) $\tau = 0.05$

Figure 3: Wasserstein distance between $\pi(\mathbf{x}, \cdot)$ and $\hat{\pi}_{\rightarrow m}(\mathbf{x}, \cdot)$ (obtained with ID = 3) with respect to the number of trees m for the eight risk classes.

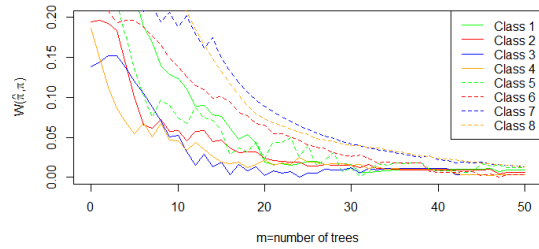
only the values 0 and 1 are allowed for Y for risk classes 7 and 8 in Table 2. The results are shown in Figure 4 for ID = 3. Again, one sees that the Wasserstein boosting trees algorithm enables to recover the true conditional probabilities $\pi(\mathbf{x}, k)$ for all risk classes.

$\pi(\mathbf{x}, k)$						
	$\mathbf{x}$	$k = 0$	$k = 1$	$k = 2$	$k = 3$	$k = 4$
Class 7	$X_1 = o, X_2 = f, X_3 = yes$	0.700	0.300	0.000	0.000	0.000
Class 8	$X_1 = o, X_2 = f, X_3 = no$	0.950	0.050	0.000	0.000	0.000

Table 2: $\pi(\mathbf{x}, k)$ for risk classes 7 and 8.



(a) $\tau = 0.5$



(b) $\tau = 0.05$

Figure 4: Wasserstein distance between $\pi(\mathbf{x}, k)$ and $\hat{\pi}_{\rightarrow m}(\mathbf{x}, k)$ (obtained with ID = 3) with respect to the number of trees m for the eight risk classes.

4.3 Comparison with the gradient boosting tree algorithm

The Wasserstein boosting trees algorithm enables to estimate the conditional probabilities $\pi(\mathbf{x}, k)$ in a non-parametric way. In particular, it allows the actuary to estimate the conditional expectation $\mu(\mathbf{x})$ from the corresponding conditional distribution. In insurance

pricing, actuaries generally apply learning procedures like the gradient boosting trees algorithm to estimate the conditional expectation $\mu(\mathbf{x})$, assuming that the responses Y follows a counting distribution such as Poisson or Negative Binomial. In this section, we compare the results obtained with the Wasserstein boosting trees algorithm proposed in this paper to the ones derived with the gradient boosting trees algorithm using the Poisson deviance as loss function.

4.3.1 Example 1 : response obeying Poisson distribution

Context

We consider an example in car insurance with the following four features:

- X_1 = Gender: policyholder's gender with two levels, labeled as female (f) or male (m), with 50% of the portfolio in each level;
- X_2 = Sport: whether the policyholder's car is a sports car or not, with two levels (yes or no) and 50% of the portfolio in each level;
- X_3 = Split: whether the policyholder splits its annual premium or not, with two levels (yes or no) and 50% of the portfolio in each level;
- X_4 = Age: policyholder's age (integer values from 18 to 65) with the same proportion 1/48 of the portfolio in each value.

The features X_1 , X_2 , X_3 and X_4 are assumed to be mutually independent.

The response Y is supposed to be the number of claims. Given $\mathbf{X} = \mathbf{x}$, Y is assumed to be Poisson distributed with expected claim frequency given by

$$\begin{aligned} \mu(\mathbf{x}) = & 0.3 \\ & \times (1 + 0.1 \mathbb{I}[x_1 = \textit{male}]) \\ & \times \left(1 + \frac{1}{\sqrt{x_4 - 17}}\right) \\ & \times (1 + 0.3 \mathbb{I}[x_2 = \textit{yes}] \mathbb{I}[18 \leq x_4 < 35] - 0.3 \mathbb{I}[x_2 = \textit{yes}] \mathbb{I}[45 \leq x_4 \leq 65]). \end{aligned}$$

The true model $\mu(\mathbf{x})$ is known and we can simulate realizations of $(Y, \mathbf{X})$. Specifically, we generate $n = 30\,000$ independent realizations of $(Y, \mathbf{X})$. We use 80% of the observations for training the models (training set) and 20% for assessing the performance of the models (validation set).

We run the Wasserstein boosting trees (henceforth abbreviated as WBT) algorithm and the gradient boosting trees (henceforth referred to as GBT) algorithm with $ID \in \{2, 3, 4, 5\}$, a bag fraction $\alpha = 0.5$ and shrinkage parameters $\tau_1 = 1 - \tau$ and $\tau_2 = 1 + \tau$ with $\tau \in \{\tau_a, \tau_b, \tau_c\} = \{0.5, 0.1, 0.05\}$ for the WBT algorithm and shrinkage parameter $\tau \in \{\tau_a, \tau_b, \tau_c\} = \{1, 0.1, 0.05\}$ for the GBT algorithm. GBT is fitted using the Poisson deviance loss.

Out-of-sample Poisson deviance

To compare the predictive accuracy of both algorithms, we compute the Poisson deviance on the validation set. Figures 5a and 5b display out-of-sample Poisson deviance for GBT and WBT. We observe that in most cases WBT outperforms GBT in the sense that the out-of-sample deviance for WBT is smaller than the one for the corresponding GBT. Figure 5b shows that the number of trees required for WBT can be much smaller than the number of trees for GBT. The lower the shrinkage parameter, the greater the difference of the number of trees between both algorithms. As expected, for both methods, more trees are needed when the shrinkage parameter is decreased.

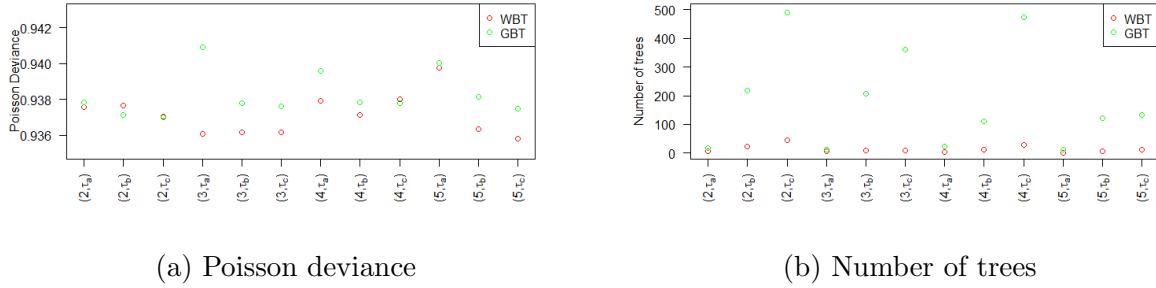


Figure 5: Out-of-sample Poisson deviance and number of trees for GBT (in green) and WBT (in red) for each set of parameters (ID, τ_ℓ) , $\ell = a, b, c$.

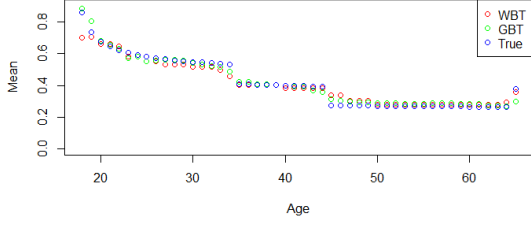
Estimated means

We compare the estimated conditional mean $\hat{\mu}(\mathbf{x})$ to the true one $\mu(\mathbf{x})$ for each risk class. For this comparison, we consider the best WBT and GBT models according to the out-of-sample Poisson deviance, namely WBT with $ID = 5$, $\tau = 0.05$ and $M = 13$ and GBT with $ID = 2$, $\tau = 0.05$ and $M = 490$.

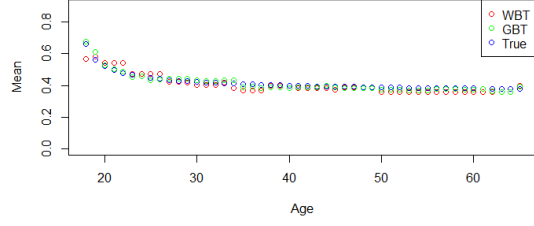
Since x_3 does not influence $\mu(\mathbf{x})$ in our example, we consider the following four groups:

- Group 1: Male driving a sports car;
- Group 2 : Male driving a regular car;
- Group 3 : Female driving a sports car;
- Group 4 : Female driving a regular car.

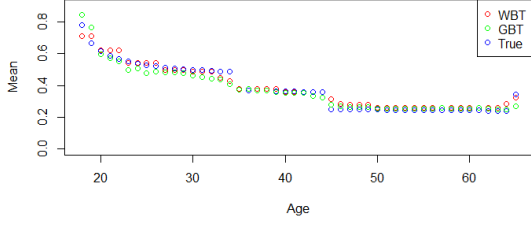
For each group, Figures 6a to 6d display the estimated means with respect to the age of the policyholder for the best WBT and GBT together with the true means. One sees that both models accurately estimate the conditional means $\mu(\mathbf{x})$. For most values of $\mathbf{x}$, both estimated conditional means are close to each other and close to the true conditional mean $\mu(\mathbf{x})$.



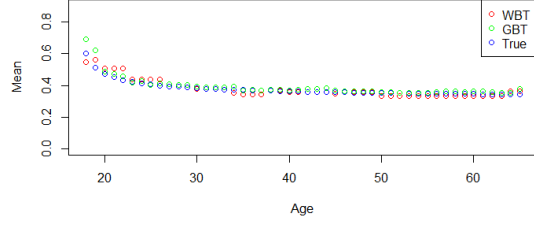
(a) Group 1



(b) Group 2



(c) Group 3



(d) Group 4

Figure 6: Estimated means with respect to the age for each group for WBT (in red) and GBT (in green) together with the true means (in blue).

Estimated conditional distributions

The GBT and WBT algorithms perform similarly in this example when it comes to estimating the conditional mean $\mu(\mathbf{x})$. However, WBT does not directly estimate the conditional mean $\mu(\mathbf{x})$, it rather focuses on the estimation of the conditional distribution $\pi(\mathbf{x}, \cdot)$ without imposing any rigid parametric assumption for $\pi(\mathbf{x}, \cdot)$ whereas GBT assumes that the response is Poisson distributed.

Let us now compare the estimated conditional distributions $\hat{\pi}_{\text{WBT}}(\mathbf{x}, \cdot)$ obtained from the WBT algorithm (with $\text{ID} = 5$, $\tau = 0.05$ and $M = 13$), the estimated conditional distribution $\hat{\pi}_{\text{GBT}}(\mathbf{x}, \cdot)$ obtained from the GBT algorithm (with $\text{ID} = 2$, $\tau = 0.05$ and $M = 490$) and the true Poisson distribution $\pi(\mathbf{x}, \cdot)$ for some risk classes. To do so, we consider the same four groups as before, and for each group, we examine the conditional distributions $\hat{\pi}_{\text{WBT}}(\mathbf{x}, \cdot)$, $\hat{\pi}_{\text{GBT}}(\mathbf{x}, \cdot)$ and $\pi(\mathbf{x}, \cdot)$ for certain values of X_4 (policyholder's age), specifically $X_4 = 20, 30, 40, 50, 60$.

Figure 7 shows the results of this comparison. For all groups and age values, both algorithms reproduce the main shape of the true distribution, with the highest probability at zero claims and rapidly decreasing mass for higher counts. At age 20, the WBT slightly underestimates the probability of observing zero claims and correspondingly overestimates the probability of one or more claims for all four groups, leading to a marginally heavier right tail. Starting from age 30, however, WBT aligns more closely with the true distribution than GBT, capturing more accurately the shift in probability mass as the expected number of claims decreases with age. The group-specific differences, namely higher claim frequencies for sports-car drivers (Groups 1 and 3) and lower ones for regular-car drivers (Groups 2 and

4), are consistently recovered by both models. Overall, WBT demonstrates slightly better stability and fidelity to the true conditional distributions, reflecting its capacity to learn the entire probability structure rather than only the conditional mean.

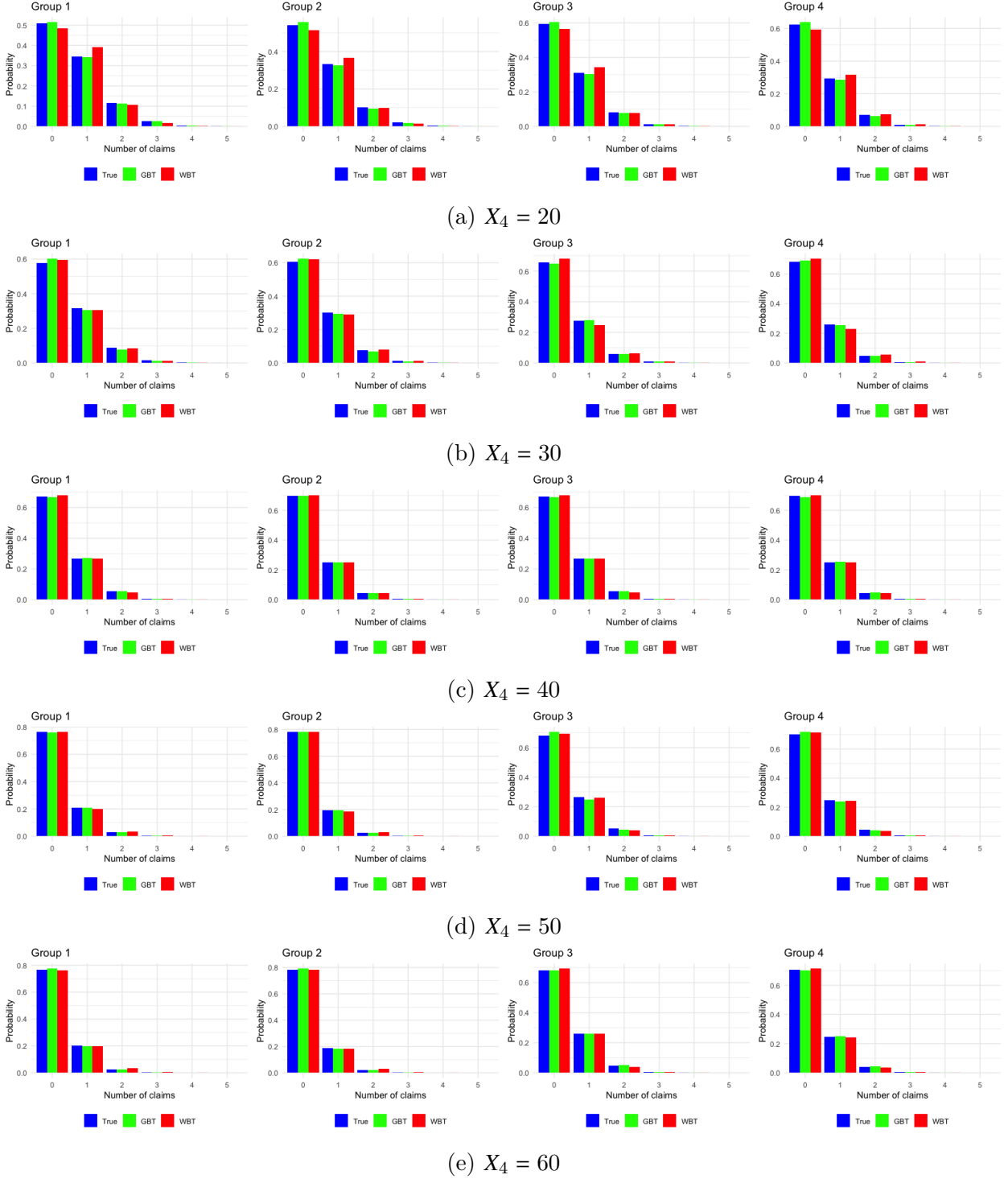


Figure 7: Estimated conditional distributions with respect to the age for each group for WBT (in red) and GBT (in green) together with the true means (in blue).

4.3.2 Example 2 : response obeying Negative Binomial distribution

Let us assume that conditionally on $\mathbf{X}$, the response is no more Poisson distributed but follows Negative Binomial distribution. This time, two different values $\mathbf{x}_1$ and $\mathbf{x}_2$ of $\mathbf{X}$ such that $\mu(\mathbf{x}_1) = \mu(\mathbf{x}_2)$ may have different associated distributions $\pi(\mathbf{x}_1, \cdot)$ and $\pi(\mathbf{x}_2, \cdot)$ since the Negative Binomial distribution is governed by two parameters (the mean and the dispersion parameter).

Consider a simple example with two categorical features, each one with only two levels, say policyholder’s age (X_1) either “young” or “old” and policyholder’s gender (X_2) either “male” or “female”. For each of the four risk classes, we assume that the corresponding distribution $\pi(\mathbf{x}, \cdot)$ follows a Negative Binomial (NB) distribution with the parameters given in Table 3. The probability mass functions for each class are given on Figure 8. Under Poisson assumption, risk classes 1-2 and 3-4 cannot be distinguished because mean responses are equal. These risk classes nevertheless differ in their residual heterogeneity.

	$\mathbf{x}$	$\pi(\mathbf{x}, \cdot)$
Class 1	$X_1 = y, X_2 = m$	$Y \mathbf{X} = \mathbf{x} \sim NB\left(1, \frac{1}{1+0.3}\right)$, such that $E[Y \mathbf{X} = \mathbf{x}] = 0.3$
Class 2	$X_1 = o, X_2 = m$	$Y \mathbf{X} = \mathbf{x} \sim NB\left(500, \frac{500}{500+0.3}\right)$, such that $E[Y \mathbf{X} = \mathbf{x}] = 0.3$
Class 3	$X_1 = y, X_2 = f$	$Y \mathbf{X} = \mathbf{x} \sim NB\left(1, \frac{1}{1+0.15}\right)$, such that $E[Y \mathbf{X} = \mathbf{x}] = 0.15$
Class 4	$X_1 = o, X_2 = f$	$Y \mathbf{X} = \mathbf{x} \sim NB\left(500, \frac{500}{500+0.15}\right)$, such that $E[Y \mathbf{X} = \mathbf{x}] = 0.15$

Table 3: $\pi(\mathbf{x}, \cdot)$ for the four risk classes.

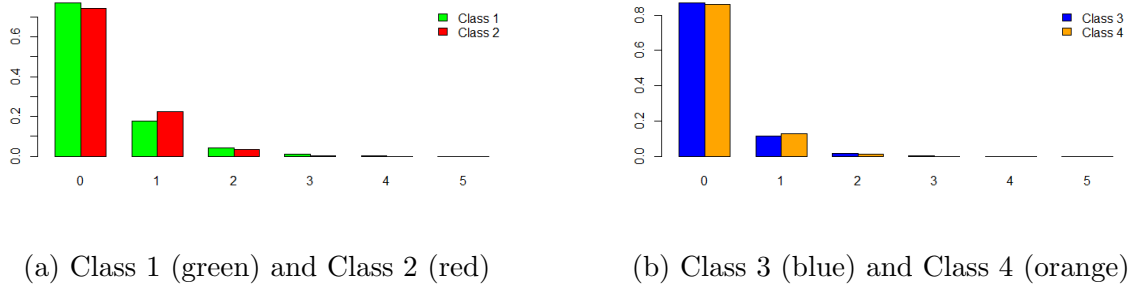


Figure 8: Comparison of the probability mass function between Class 1 and Class 2 and between Class 3 and Class 4.

Figure 9 shows the Wasserstein distance between the estimated conditional distribution $\hat{\pi}(\mathbf{x}, \cdot)$ and the true Negative Binomial distribution $\pi(\mathbf{x}, \cdot)$ for the four risk classes. We see that the WBT algorithm enables to accurately estimate the true conditional distribution for any risk class (after 20 iterations, $\mathcal{W}(\hat{\pi}(\mathbf{x}, \cdot), \pi(\mathbf{x}, \cdot)) = 0$ for all $\mathbf{x}$). If we use the GBT algorithm with the typical Poisson assumption for the condition distribution $\pi(\mathbf{x}, \cdot)$, we cannot distinguish risk profiles of Class 1 (resp. Class 3) from risk profiles of Class 2 (resp. Class 4), as shown in Figure 10.

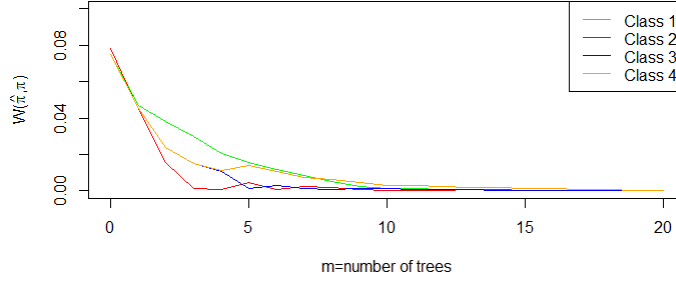


Figure 9: Wasserstein distance between the estimated distribution $\hat{\pi}(\mathbf{x}, .)$ and the true distribution $\pi(\mathbf{x}, .)$ for the four risk classes with ID = 2 and $\tau_2 = 0.01$.

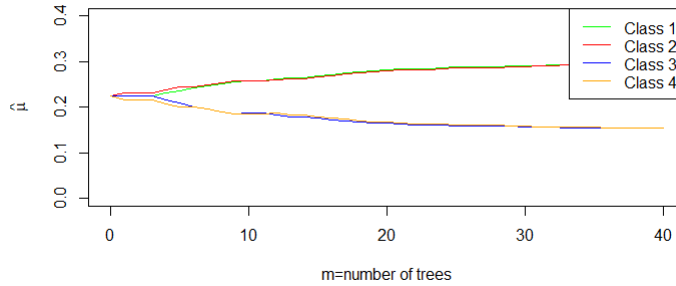


Figure 10: Estimated mean $\hat{\mu}(\mathbf{x})$ with the GBT algorithm for the four risk classes with ID = 2 and $\tau_2 = 0.01$.

5 Case study with claim frequencies in motor insurance

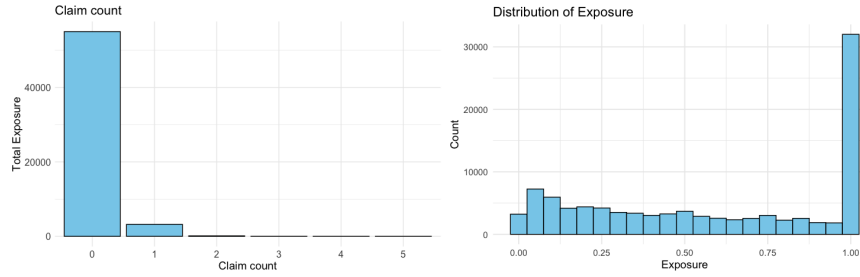
5.1 MTPL data set

We now consider the `freMTPL2freq` data set comprised in the `CASdatasets` package in R. This dataset was introduced by [Charpentier \[2015\]](#) and has been widely used since in the actuarial literature, in works such as [Huyghe et al. \[2024\]](#) or [Delong and Kozak \[2023\]](#). It contains a French motor third-party liability (MTPL) insurance portfolio comprising of 677 991 entries, each corresponding to a policy. Twelve variables are observed for each of these policies, of which we keep the nine listed in Table 4 for the present analysis. The additional three variables, not shown here, are the vehicle brand, population density and French region where the insured vehicle is registered. We refer to [Noll et al. \[2018\]](#) for a detailed description of the complete dataset.

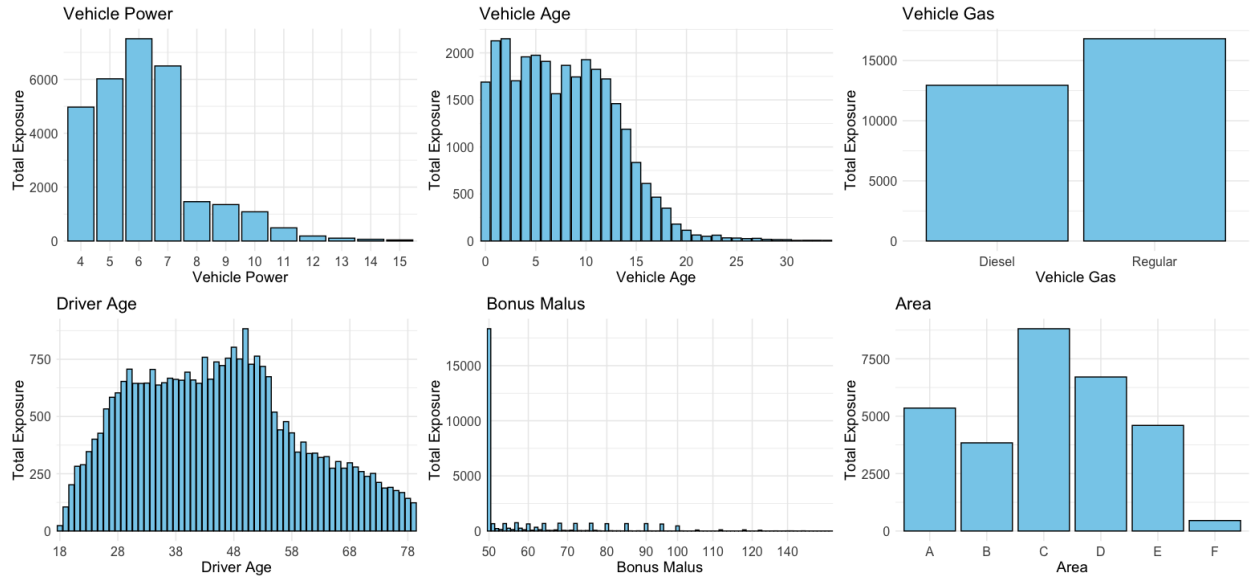
We work with a subset of 50 000 observations in the present analysis. Figures 11a and 11b show the histograms of, respectively, the claim count (response variable) and exposure, and the six covariates that we will use in the models.

Table 4: Description of the variables in the **freMTPL2freq** dataset.

Variable	Description
IDpol	Policy number.
ClaimNb	Number of claims.
Exposure	Total exposure in yearly units.
Area	Area code (A,B,C,D,E or F).
BonusMalus	Bonus-malus level, ranging between 50 and 230 (reference level: 100).
DrivAge	Age of the driver of the vehicle, in years.
VehAge	Age of the vehicle, in years.
VehGas	Type of fuel used for the vehicle (diesel or regular gas).
VehPower	Power of the car, categorical variable ranging from 4 to 15.



(a) Claim count and exposure.



(b) Covariates from the **freMTPL2freq** dataset.

Figure 11: Histograms of the claim count, exposure and covariates from the **freMTPL2freq** dataset.

We split our subset into a training set comprising 80% of the observations. The remaining

20% of observations allow us to validate our WBT model by comparing its performance against the observed data and the results obtained using the GBT approach.

We run both the WBT and GBT algorithms using $ID \in \{2, 3, 4, 5\}$, a bag fraction $\alpha = 0.5$ and shrinkage parameters $\tau_1 = 1 - \tau$ and $\tau_2 = 1 + \tau$ with $\tau \in \{\tau_a, \tau_b, \tau_c\} = \{0.5, 0.1, 0.05\}$ for WBT and $\tau \in \{\tau_a, \tau_b, \tau_c\} = \{1, 0.1, 0.05\}$ for GBT. Moreover, GBT is fitted with the Poisson deviance loss.

5.2 Out-of-sample Poisson deviance

Figures 12a and 12b show the out-of-sample Poisson deviance and the corresponding number of trees obtained for the GBT (in green) and WBT (in red) algorithms.

Across all parameter settings, the WBT systematically achieves lower out-of-sample deviance values than the GBT. Although GBT is designed to directly minimize the Poisson deviance, WBT, based on the Wasserstein distance, yields comparable or better results, suggesting that its distributional learning criterion effectively captures the underlying data structure.

In addition, the number of trees required for convergence is consistently smaller for WBT than for GBT. This indicates that WBT attains similar or improved predictive performance with a more compact ensemble, reflecting greater model efficiency in terms of complexity. Overall, the results confirm that both algorithms perform similarly in this setting, with WBT achieving a marginal advantage in terms of both deviance and parsimony.

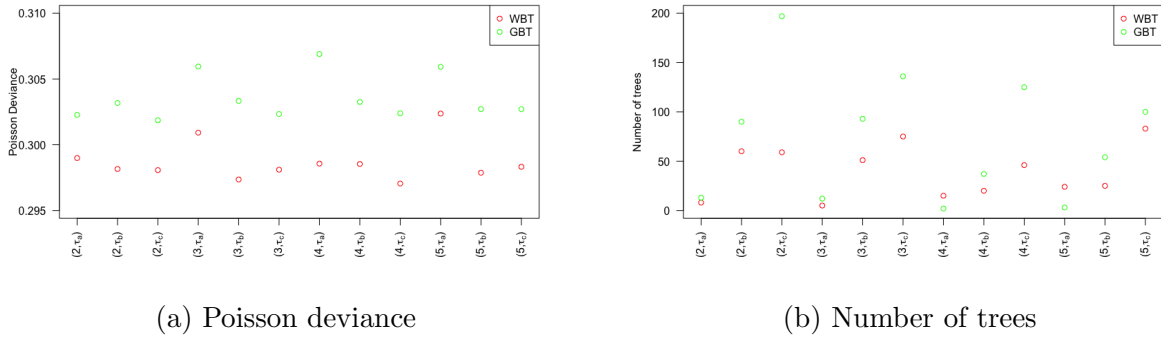


Figure 12: Out-of-sample Poisson deviance and number of trees for GBT (in green) and WBT (in red) for each set of parameter (ID, τ_ℓ) , $\ell = a, b, c$.

6 Conclusion

This paper proposes a new boosting algorithm for claim counts, based on Wasserstein distance. Compared to existing procedures targeting conditional means, WBT produces estimated conditional distributions which allow the actuary to derive any estimation of interest. When applied to simulated data and a classical industry data base, WBT shows excellent performances compared to GBT. The non-parametric nature of the newly proposed WBT allows the actuary to recover patterns present in the data beyond conditional means, such

as residual heterogeneity compared to Poisson distribution for instance. In our examples, WBT even outperforms GBT in terms of Poisson deviance despite GBT precisely optimizes that metric.

The Poisson distribution has been extended in different ways to accommodate the specific aspects of claim frequency distributions: Overdispersed Poisson or Poisson mixtures, including Negative Binomial, have been proposed to account for overdispersion, Zero-inflated Poisson or Hurdle models to account for the probability mass at zero, etc. Machine learning tools have been developed for each model but the approach remains parametric and thus exposes actuaries to model risk. The latter has been extensively studied for the conditional mean, showing that Poisson regression delivers consistent estimates even if responses do not obey the Poisson distribution as long as the conditional mean is correctly specified. However, this is not necessarily true for other parameters of interest, like the no-claim probability extensively used in underwriting. Being non-parametric by nature, the WBT algorithm proposed in this paper avoids this pitfall and can be used to challenge the conclusions obtained from parametric regression.

In this paper, we restricted ourselves to integral probability metrics because they turn out to be computed more rapidly compared to metrics defined as supremum (like Kolmogorov distance, for instance, corresponding to the largest absolute difference between respective distribution functions). Also, we paid particular attention to their expressions with count data. Of course, there are other candidates for being used as distance in splitting criterion. For counting random variables M and N , the Wasserstein distance is closely related to the total-variation distance $\mathcal{TV}(\pi_M, \pi_N)$ given by

$$\mathcal{TV}(\pi_M, \pi_N) = \sum_{k=0}^{\infty} |\mathbb{P}[M = k] - \mathbb{P}[N = k]|.$$

When M and N refer to claim frequencies at policy-level in personal lines, the series reduces to just a few terms and is thus easy to compute. Proposition 9.6.5 in [Denuit et al. \[2005\]](#) shows that $\mathcal{TV}(\pi_M, \pi_N) \leq 2\mathcal{W}(\pi_M, \pi_N)$.

Besides $\mathcal{W}$ and $\mathcal{TV}$, the integrated stop-loss distance is another candidate to assess the proximity of two distributions in insurance studies. The integrated difference in stop-loss premiums has traditionally been used to measure the distance between two risks with finite variances in the actuarial literature. Given two random variables V and W , the integrated stop-loss distance $\mathcal{ISL}$ is given by

$$\mathcal{ISL}(\pi_V, \pi_W) = \int_{-\infty}^{\infty} |\mathbb{E}[(V - t)_+] - \mathbb{E}[(W - t)_+]| dt,$$

where the integrand is the absolute difference of the stop-loss transforms of V and of W . This distance is easy to compute under stop-loss or convex order since Proposition 9.8.2 in [Denuit et al. \[2005\]](#) shows that

- (i) If $\mathbb{E}[(V - t)_+] \leq \mathbb{E}[(W - t)_+]$ holds for all t then $\mathcal{ISL}(\pi_V, \pi_W) = \frac{1}{2}(\mathbb{E}[W^2] - \mathbb{E}[V^2])$.
- (ii) If $\mathbb{E}[(V - t)_+] \leq \mathbb{E}[(W - t)_+]$ holds for all t and $\mathbb{E}[V] = \mathbb{E}[W]$ then $\mathcal{ISL}(\pi_V, \pi_W) = \frac{1}{2}(\text{Var}[W] - \text{Var}[V])$.

Considering Poisson mixtures, these results apply when mixing distributions satisfy the conditions appearing in (i)-(ii). Under (ii), we see that $\mathcal{ISL}$ coincides with the classical splitting criterion used in regression trees.

This paper thus opens a new approach to statistical learning in insurance, replacing deviance with probabilistic distance. Depending on the application, the actuary is free to select the most appropriate distance. Considering insurance ratemaking, Wasserstein distance appears to be a natural candidate complying with intuition.

Acknowledgements

Julien Trufin gratefully acknowledges the support received from the Research Chair ACTIONS under the aegis of the Risk Foundation, an initiative by BNP Paribas Cardif and the French Institute of Actuaries.

Compliance with ethical standards

The authors declare that they have no conflicts of interest.

Competing interests

The authors have no competing interests to declare that are relevant to the content of this article.

References

- L. Breiman, J. Friedman, R. Olshen, and C. Stone. *Classification and Regression Trees*. Wadsworth Statistics/Probability Series, 1984.
- A. Charpentier. Computational actuarial science with r. *CRC Press*, 2015.
- L. Delong and A. Kozak. The use of autoencoders for training neural networks with mixed categorical and numerical features. *ASTIN Bulletin*, 53(2):213–232, 2023.
- M. Denuit, J. Dhaene, M. Goovaerts, and R. Kaas. *Actuarial Theory for Dependent Risks: Measures, Orders and Models*. Wiley, New York, 2005.
- M. Denuit, D. Hainaut, and J. Trufin. *Effective Statistical Learning Methods for Actuaries II: Tree-Based Methods and Extensions*. Springer Actuarial Lecture Notes Series, 2020.
- Q. Du, G. Biau, F. Petit, and R. Porcher. Wasserstein random forests and applications in heterogeneous treatment effects. *International Conference on Artificial Intelligence and Statistics*, pages 1729–1737, 2021.
- J. Friedman. Greedy function approximation: A gradient boosting machine. *Annals of Statistics*, 29:1189–1232, 2001.

- J. Huyghe, J. Trufin, and M. Denuit. Boosting cost-complexity pruned trees on tweedie responses: the abt machine for insurance ratemaking. *Scandinavian Actuarial Journal*, 2025(5):417–439, 2024.
- A. Noll, R. Salzmänn, and M. Wüthrich. Case study: French motor third-party liability claims. 2018. doi: AvailableatSSRN:<https://ssrn.com/abstract=3164764>.