

AUTOCALIBRATION AND TWEEDIE-DOMINANCE FOR INSURANCE PRICING WITH MACHINE LEARNING

Denuit, M., Charpentier, A. and J. Trufin

LIDAM Discussion Paper ISBA/2021/13

AUTOCALIBRATION AND TWEEDIE-DOMINANCE FOR INSURANCE PRICING WITH MACHINE LEARNING

Michel Denuit

Institute of Statistics, Biostatistics and Actuarial Science
UCLouvain

Louvain-la-Neuve, Belgium

Arthur Charpentier

Université du Québec à Montréal (UQAM)
Montreal, Quebec, Canada

Julien Trufin

Department of Mathematics
Université Libre de Bruxelles (ULB)
Brussels, Belgium

March 4, 2021

Abstract

Boosting techniques and neural networks are particularly effective machine learning methods for insurance pricing. Often in practice, there are nevertheless endless debates about the choice of the right loss function to be used to train the machine learning model, as well as about the appropriate metric to assess the performances of competing models. Also, the sum of fitted values can depart from the observed totals to a large extent and this often confuses actuarial analysts. The lack of balance inherent to training models by minimizing deviance outside the familiar GLM with canonical link setting has been empirically documented in Wüthrich (2019, 2020) who attributes it to the early stopping rule in gradient descent methods for model fitting. The present paper aims to further study this phenomenon when learning proceeds by minimizing Tweedie deviance. It is shown that minimizing deviance involves a trade-off between the integral of weighted differences of lower partial moments and the bias measured on a specific scale. Autocalibration is then proposed as a remedy. This new method to correct for bias adds an extra local GLM step to the analysis. Theoretically, it is shown that it implements the autocalibration concept in pure premium calculation and ensures that balance also holds on a local scale, not only at portfolio level as with existing bias-correction techniques. The convex order appears to be the natural tool to compare competing models, putting a new light on the diagnostic graphs and associated metrics proposed by Denuit et al. (2019).

Keywords: Risk classification, Tweedie distribution family, Concentration curve, Bregman loss, Convex order.

1 Introduction and motivation

In the 1960s, North-American actuaries pioneered risk classification with the help of minimum bias methods, after Bailey and Simon (1960) and Bailey (1963). The central idea is that an acceptable set of premiums should reproduce the experience within sub-portfolios corresponding to each level of meaningful risk factors (like gender or age, for instance) and also the overall experience, i.e. be balanced for each level and in total (leading to the method of marginal totals, or MMT in short).

In the late 1980s, it turned out that for any GLM with canonical link and a score containing an intercept, there is an exact balance between fitted and observed aggregated responses over the whole data set and for any level of the categorical features. This formally related GLMs to minimum bias and MMT, as documented in Mildenhall (1999). This connection greatly facilitated the wide acceptance of GLMs in actuarial practice that has been particularly fast after the development of powerful computer tools.

Nowadays, actuaries resort to advanced statistical methods to be able to accurately assess policyholders' expected claim costs, or pure premium. Modern risk classification tools can deal with highly segmented problems resulting from the massive amount of information now available to insurers. Actuarial pricing models are generally calibrated so that a statistical goodness-of-fit measure is optimized on the training set whereas model tuning (like early-stopping decisions) is performed by cross-validation. Competing models are compared on a separate test or validation set. The objective function used for model training (often call loss function in machine learning) generally corresponds to deviance or log-likelihood, in a model accounting for the nature of data under consideration: typically, Poisson for claim counts, Gamma for average claim severities and compound Poisson sums with Gamma-distributed terms for claim totals, all belonging to the Tweedie family with power variance function. The metric used to assess model performances sometimes differs from the loss function used for training. Lift diagrams developed after Meyers and Cummings (2009) appear to be useful in that respect.

Actuarial risk classification remained bridged to statistical regression models and naturally followed their evolution from GLMs to GAMs, trees and random forests, gradient and statistical boosting, projection pursuit and neural networks, to name just a few. This evolution took place gradually, following the availability of computational resources in statistical software, without questioning the relevance of algorithms regarding minimum bias and interpretability of the objective function: the deviance established itself as the only objective function, even when the interpretability of the likelihood equations in terms of balance was lost. The underlying MMT caution to GLMs has been progressively forgotten and actuaries massively adopted every new tool extending the GLM machinery.

However, it turned out that tree-based boosting models and neural networks trained to minimize deviance generally end up with a significant under-estimation of total claims: the total balance is generally broken, even on the training data set. This seems to be an intrinsic result of using such an objective function even if no formal explanation has been provided so far beyond empirical findings. In this paper, we question the relevance of deviance as objective function, without the global balance constraint. It is shown that deviance may lead to dubious candidate premiums because the latter can deviate a lot from observed losses when totals are not kept unchanged. This is formally established with the help of

analytical results and demonstrated on empirical motor third-party liability insurance data.

A natural remedy consists in restoring global balance at each step of the iterative procedure used to optimize the loss function. This is easily implemented by revising the intercept (under canonical link function) or by constraining the optimization so that the observed total matches its fitted counterpart. This simple solution restores global balance but does not ensure that financial equilibrium also holds in meaningful sub-portfolios (remember that GLMs with canonical link not only imposes global balance, at portfolio level when the score comprises an intercept, but also at the level of risk classes determined by binary features).

For this reason, we propose a new strategy based on the concept of autocalibration (see, e.g., Kruger and Ziegel, 2020). This approach guarantees global balance as well as local equilibrium in the spirit of the original MMT. This simple and effective solution to the problem is implemented by adding an extra step implementing MMT within a local GLM analysis. Specifically, after the analysis has been performed with a method that does not necessarily respect marginal totals, a local constant GLM fit is achieved in order to restore the connection with MMT. Thus, as advocated by Wüthrich (2019, 2020), we also combine GLMs with advanced statistical learning tools but in a different way. Wüthrich (2019, 2020) used Neural Networks to produce new features in the last hidden layer, to be used in a GLM replacing the output layer. In this approach, Neural Networks allow actuaries to perform feature-engineering to feed the GLM score and marginal totals are respected by the use of GLMs in the last step. Wüthrich (2019, 2020) then explains how to interpret the working features generated by the last hidden layer of Neural Networks. This approach can also be related to the polishing procedure proposed by Zumel (2019) where random forests predictions are used to train a linear model in a second step. As pointed out by this author, the extra step should not be performed on the same data in order to avoid a potential source of overfitting (called nested-model bias). The local GLM approach proposed in this paper applies to any statistical learning model, not specifically to Neural Networks or random forests. It consists in using the candidate premium itself as a new feature in the last step. This ensures that the feature space reduces to the real line. By using a local constant, or intercept-only GLM, these predictors define optimal neighborhoods to perform local averaging of observed losses.

Lack of balance is sometimes imposed on purpose by the analyst, with the hope to stabilize the results. This is typically the case when the loss function comprises a penalty term (Lasso, Ridge or Elastic-Net, for instance), in addition to a goodness-of-fit measure. The local GLM step restoring balance advocated in the present paper is thus relevant in that setting, too.

The remainder of the paper is structured as follows. Section 2 provides the reader with formal definitions and notation used throughout the text. Section 3 recalls the form of loss functions that are consistent for the mean, and hence useful in pure premium determination. In Section 4, a mixture representation of Tweedie deviance is derived. Precisely, Tweedie deviance is decomposed into the sum of the integral of weighted differences of lower partial moments and the bias measured on a specific scale. This decomposition is used to understand the consequences of training the model by minimizing Tweedie deviance. Special attention is devoted to the Poisson regression case. The concept of Tweedie dominance is also introduced there, as a natural counterpart to Bregman dominance, or forecast dominance discussed in Kruger and Ziegel (2020). In Section 5, we propose a simple and powerful method to restore

balance at both global and local levels, based on the concept of autocalibration. Tweedie dominance between autocalibrated predictors reduces to the well-known convex order, or stop-loss order with equal means that has been proposed to compare predictors by Denuit et al. (2019) and Kruger and Ziegel (2020). This allows us to derive new results for the diagnostic tools and performance metrics proposed by Denuit et al. (2019). Precisely, it is shown that the concentration curve reduces to the Lorenz curve for autocalibrated predictors, so that the corresponding ABC metric is zero. A non-zero estimated ABC thus suggests that the predictor lacks of balance. This also provides a sound methodological argument for using Lorenz curves and lift diagrams to assess model performances in insurance pricing. From a practical point of view, actuaries could also replace Lorenz and lift curves with the simpler stop-loss transforms for comparing the predictors under consideration. A numerical study is provided in Section 6. The final Section 7 discusses the results and concludes.

2 Context, definition and notation

Let us now describe the notation used in this paper. We consider a response Y and a set of features $X_1, \dots, X_p$ gathered in the vector $\mathbf{X} \in \mathcal{X}$ (classically, $\mathcal{X} \subset \mathbb{R}^p$). In this paper, the response is typically the number of claims reported to the insurance company by a given policyholder, the average claim severity or the total claim amount in relation with this contract. The dependence structure inside the random vector $(Y, X_1, \dots, X_p)$ is exploited to extract the information contained in $\mathbf{X}$ about Y . In actuarial pricing, the aim is to evaluate the pure premium as accurately as possible. This means that the target is the conditional expectation $\mu(\mathbf{X}) = \mathbb{E}[Y|\mathbf{X}]$ of the response Y (claim number or claim amount) given the available information $\mathbf{X}$. Henceforth, $\mu(\mathbf{X})$ is referred to as the true (pure) premium. Notice that in some applications, $\mu(\mathbf{X})$ only refers to one component of the pure premium. For instance, working in the frequency-severity decomposition of insurance losses, $\mu(\mathbf{X})$ can be either the expected number of insured events or the expected claim size, or severity.

Example 2.1. To illustrate various concepts introduced in this paper, we use simulated Poisson counts, conditional on a single covariate $X \sim \text{Uniform}([0, 10])$, with mean $\mu(x) = 8 - x + 3(x - 5)_+$, which is continuous and piecewise linear. Simulated data can be visualized on the left of Figure 2.1.

Example 2.2. In addition to the Poisson counts described in Example 2.1, we also use a second set of simulated Poisson counts, conditional on two independent covariates $X_1, X_2 \sim \text{Uniform}([0, 10])$, with mean

$$\begin{aligned} \mu(x_1, x_2) = & 3 \left(1 - \frac{x_1 - 5}{3}\right)^2 \exp \left(- \left(\frac{x_1 - 5}{3}\right)^2 - \left(\frac{x_2 - 5}{3} + 1\right)^2 \right) \\ & - 10 \left(\frac{x_1 - 5}{15} - \left(\frac{x_1 - 5}{3}\right)^3 - \left(\frac{x_2 - 5}{3}\right)^5 \right) \exp \left(- \left(\frac{x_1 - 5}{3}\right)^2 - \left(\frac{x_2 - 5}{3}\right)^2 \right) \\ & - \frac{1}{3} \exp \left(- \left(\frac{x_1 - 5}{3} + 1\right)^2 - \left(\frac{x_2 - 5}{3}\right)^2 \right) + 8 \end{aligned}$$

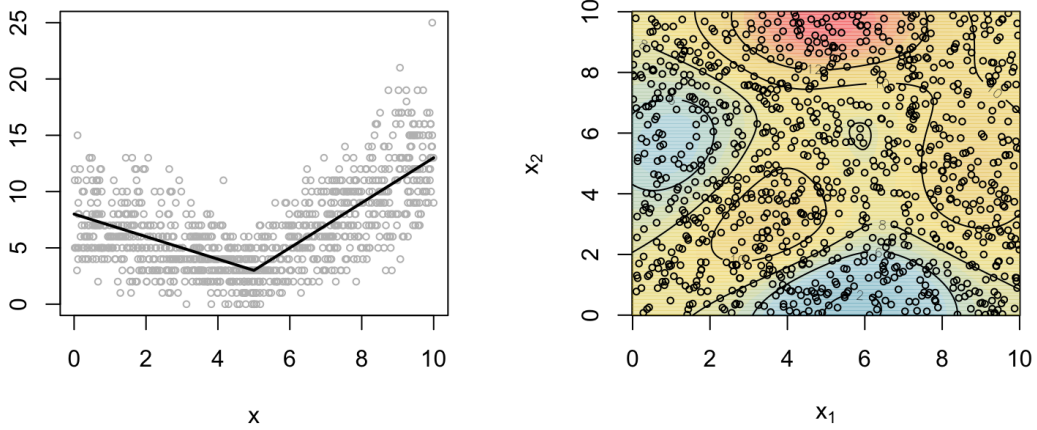


Figure 2.1: Generated data, with one and two covariates.

which is continuous and not separable. Simulated data are displayed in the right panel of Figure 2.1.

The function $\mathbf{x} \mapsto \mu(\mathbf{x}) = \mathbb{E}[Y|\mathbf{X} = \mathbf{x}]$ is unknown to the actuary, and may exhibit a complex behavior in $\mathbf{x}$. This is why this function is approximated by a (working, or actual) premium $\mathbf{x} \mapsto \pi(\mathbf{x})$ with a simpler structure. When the analyst is working in the frequency-severity decomposition of insurance losses, $\pi(\mathbf{x})$ targets the expected number of insured events or the mean claim severity, separately. Once fitted on the training data set using an appropriate learning procedure, this produces estimates $\hat{\pi}(\mathbf{x})$ for $\mu(\mathbf{x})$, or fitted values.

The developments in this paper apply in any setting where a global balance is desirable, that is, where it is important that the sum of estimates does not deviate too much from the sum of actual observations at both the entire portfolio level and also more locally, in meaningful classes of policyholders. The reason is obvious: the sum of the pure premiums must match the claim total as accurately as possible so that the insurance company is able to indemnify all third-parties and beneficiaries in execution of the contracts, without excess nor deficit, by the very definition of pure premium (expense loadings and cost-of-capital charges are added into the calculation at a later stage, when moving to commercial premiums). This naturally translates into a global balance constraint: considering that the total claim figures are representative of next-year's experience, it is important that the sum of fitted premiums $\hat{\pi}(\mathbf{x})$ match the sum of responses Y taken as proxy for the total premium income (i.e., the sum of $\mu(\mathbf{x})$), as closely as possible. But local equilibrium is also essential to guarantee a competitive pricing.

The merits of a given pricing tool can be assessed using the pair $(\mu(\mathbf{X}), \hat{\pi}(\mathbf{X}))$ so that we are back to the bivariate case even if there were thousands of features comprised in $\mathbf{X}$. What really matters is the correlation between $\hat{\pi}(\mathbf{X})$ and $\mu(\mathbf{X})$ but as $\mu(\mathbf{X})$ is unobserved the actuary can only use its noisy version Y to reveal the agreement of the true premium

$\mu(\mathbf{X})$ with its working counterpart $\hat{\pi}(\mathbf{X})$. In insurance applications, $\hat{\pi}(\mathbf{X})$ is supposed to be used as a premium so that correlation is important, but it is also essential that the sum of predictions $\hat{\pi}(\mathbf{X})$ matches the sum of actual losses as closely as possible, as explained before. This is expressed by the global balance condition and its local version.

To ease the exposition, we assume that predictor $\hat{\pi}(\mathbf{X})$ under consideration, as well as the conditional expectation $\mu(\mathbf{X})$ are continuous random variables admitting probability density functions. This is generally the case when there is at least one continuous feature contained in the available information $\mathbf{X}$ and the function $\hat{\pi}$ is a continuously increasing function of a real score built from $\mathbf{X}$. However, this rules out predictions based on discrete features only, as well as piecewise constant predictors, e.g., a single tree. Indeed, then $\hat{\pi}(\mathbf{X})$ takes only a limited number of values. As actuarial pricing is nowadays based on more sophisticated models (trees being combined into random forests, for instance), this continuity assumption does not really restrict the generality of the approach.

3 Consistent loss functions for pure premium

3.1 Bregman loss functions

Pure premiums corresponding to conditional expectations, they can be consistently estimated only if the expected loss is minimum for the mean response. This corresponds to the concept of consistent loss functions. Precisely, a loss function is consistent for the mean if no other quantity leads to a lower expected loss than the mean. Bregman loss functions are of the form

$$L(y, m) = \ell(y) - \ell(m) - \ell'(m)(y - m) \text{ for a convex function } \ell. \quad (3.1)$$

The loss function $L(\cdot, \cdot)$ in (3.1) is also referred to as q -loss function after Efron (1986). Bregman loss function measures the discrepancy between y and m on two distinct scales: on the ℓ -scale with the difference $\ell(y) - \ell(m)$ and on the original scale with the difference $y - m$. Both differences are then combined using the relative weight $\ell'(m)$. Sometimes, (3.1) is defined in terms of the concave function $-\ell$.

Savage (1971) showed, subject to weak regularity conditions, that for any loss function (3.1) and response Y ,

$$\mathbb{E}[L(Y, \mathbb{E}[Y])] \leq \mathbb{E}[L(Y, m)] \text{ for any } m \text{ in the support of } Y.$$

The class of Bregman loss functions is then said to be consistent for the (conditional) mean functional. It is worth to stress that Savage (1971) refers to insurance to motivate his proposed approach, stating in Section 2 that “Money payable subject to a contingency, such as the accidental burning of a house or the outcome of a race, can be regarded as a commodity. Such commodities are explicitly dealt in by insurance companies and bookmakers”. Savage (1971) then justifies the use of loss function (3.1) by the revelation of individual rate of substitution when presented a single, relatively simple, economic choice.

Bregman loss functions can take a wide variety of shapes: these loss functions can be asymmetric, with either under-predictions or over-predictions being more heavily penalized, and they can be strictly convex or have concave segments. Thus, restricting attention to loss

	Type	Name
$\xi < 0$	Continuous	-
$\xi = 0$	Continuous	Normal
$0 < \xi < 1$	Non existing	-
$\xi = 1$	Discrete	Poisson
$1 < \xi < 2$	Mixed, non-negative	Compound Poisson sum with Gamma-distributed terms
$\xi = 2$	Continuous, positive	Gamma
$2 < \xi < 3$	Continuous, positive	-
$\xi = 3$	Continuous, positive	Inverse Gaussian
$\xi > 3$	Continuous, positive	-

Table 3.1: Tweedie distributions and corresponding power parameters.

functions that generate the mean as the optimum value does not require imposing symmetry or other assumptions on the loss function. In fact, out of the infinite number of Bregman loss functions, only one is symmetric: the quadratic loss function $L(y, m) = (y - m)^2$ obtained with $\ell(m) = m^2$.

In actuarial studies targeting the pure premium, every Bregman loss function would be eligible. But a clever choice of loss function reflecting the nature of the response under consideration or its mean-variance relationship yields more accurate results. Notice that the term $\ell(y)$ is sometimes subtracted in (3.1). This does not modify the predictions and ensures that the loss function is defined over the whole range of the response. Let us also mention that there are classes of loss functions that are consistent with other indicators, such as quantiles or expectiles for instance. We refer the interested reader to Gneiting (2011) for a general overview. Loss functions are sometimes called scoring functions in the statistical literature but we do not adopt this terminology here in order to avoid any confusion with the score involved in $\pi(\mathbf{x})$.

3.2 Tweedie deviance

In this paper, we assume that the response obeys a probability distribution belonging to the Tweedie subclass of the Exponential Dispersion family. Precisely, this means that we assume that the logarithm of the probability mass function for a discrete response, or of the probability density function for a continuous response, is of the form $(y\theta - a(\theta))/\phi$ up to a constant term, for some known dispersion parameter ϕ and non-decreasing and convex cumulant function $a(\cdot)$ corresponding to a variance function of the form $V(\mu) = \mu^\xi$ for some power parameter ξ . For appropriate choice of ℓ in (3.1), we recover Tweedie deviance. Precisely, the function $\ell(m) = 2(m\theta_m - a(\theta_m))$ with $a'(\theta_m) = m$ results in the deviance loss function

$$L(y, m) = 2(y(\theta_y - \theta_m) - a(\theta_y) + a(\theta_m))$$

of the Exponential Dispersion distribution, where $a'(\theta_y) = y$.

Table 3.1 lists all Tweedie distributions. Negative values of ξ gives continuous distributions on the whole real axis. For $0 < \xi < 1$, there is no Exponential Dispersion distribution

with such variance function. Only the cases $\xi \geq 1$ are thus interesting for applications in insurance. In the remainder of this paper, we thus restrict our analysis to $\xi \geq 1$, and the corresponding deviance loss function is

$$L(y, m) = \begin{cases} 2 \left(y \ln \frac{y}{m} - (y - m) \right) & \text{for } \xi = 1 \\ 2 \left(-\ln \frac{y}{m} + \frac{y}{m} - 1 \right) & \text{for } \xi = 2 \\ 2 \left(\frac{y^{2-\xi}}{(1-\xi)(2-\xi)} - \frac{ym^{1-\xi}}{1-\xi} + \frac{m^{2-\xi}}{2-\xi} \right) & \text{for } \xi > 1 \text{ and } \xi \neq 2. \end{cases} \quad (3.2)$$

3.3 Comparison criterion

Let $\hat{\pi}$ be the estimated mean response Y built from some training set (all formulas in this paper are meant given this training set). The respective performances of competing models can then be assessed on the basis of a validation set $\{(Y_i, \mathbf{X}_i), i = 1, 2, \dots, n\}$, that has not been used to obtain $\hat{\pi}$. Performances of $\hat{\pi}$ are generally assessed with the help of out-of-(training) sample deviance, also called predictive deviance and given by

$$D_n(\xi, \hat{\pi}) = \frac{1}{n} \sum_{i=1}^n L(Y_i, \hat{\pi}(\mathbf{X}_i)).$$

If n is large enough then we can resort to the limiting value

$$D_n(\xi, \hat{\pi}) \rightarrow E[L(Y^{\text{new}}, \hat{\pi}(\mathbf{X}^{\text{new}}))] \text{ as } n \rightarrow \infty,$$

where $(Y^{\text{new}}, \mathbf{X}^{\text{new}})$ is a new observation, independent of, and distributed as those $(Y_i, \mathbf{X}_i)$ contained in the training set. In this paper, we compare models on the basis of the large-sample version of the predictive deviance. This approach is meaningful in insurance applications where the analyst is typically in a data-rich situation.

Henceforth, we thus compare models on the basis of the predictive Tweedie deviance

$$E[L(Y^{\text{new}}, \hat{\pi}(\mathbf{X}^{\text{new}}))] = \begin{cases} 2 \left(E[\hat{\pi}(\mathbf{X}^{\text{new}}) - Y^{\text{new}} \ln \hat{\pi}(\mathbf{X}^{\text{new}})] + E[Y^{\text{new}} \ln Y^{\text{new}} - Y^{\text{new}}] \right) & \text{for } \xi = 1 \\ 2 \left(E \left[\ln \hat{\pi}(\mathbf{X}^{\text{new}}) + \frac{Y^{\text{new}}}{\hat{\pi}(\mathbf{X}^{\text{new}})} \right] - E[\ln Y^{\text{new}} - 1] \right) & \text{for } \xi = 2 \\ 2 \left(E \left[\frac{\hat{\pi}(\mathbf{X}^{\text{new}})^{2-\xi}}{2-\xi} - \frac{Y^{\text{new}} \hat{\pi}(\mathbf{X}^{\text{new}})^{1-\xi}}{1-\xi} \right] + E \left[\frac{Y^{\text{new} \, 2-\xi}}{(1-\xi)(2-\xi)} \right] \right) & \text{for } \xi > 1 \text{ and } \xi \neq 2 \end{cases} \quad (3.3)$$

which depends on the power parameter ξ and the predictor $\hat{\pi}$ under consideration. Notice that for any value of ξ , the last term in (3.3) does not depend on $\hat{\pi}$ and is thus irrelevant for model comparison, so that in the remainder of the paper, we actually use

$$D(\xi, \hat{\pi}) = \begin{cases} E[\hat{\pi}(\mathbf{X}^{\text{new}}) - Y^{\text{new}} \ln \hat{\pi}(\mathbf{X}^{\text{new}})] & \text{for } \xi = 1 \\ E \left[\ln \hat{\pi}(\mathbf{X}^{\text{new}}) + \frac{Y^{\text{new}}}{\hat{\pi}(\mathbf{X}^{\text{new}})} \right] & \text{for } \xi = 2 \\ E \left[\frac{\hat{\pi}(\mathbf{X}^{\text{new}})^{2-\xi}}{2-\xi} - \frac{Y^{\text{new}} \hat{\pi}(\mathbf{X}^{\text{new}})^{1-\xi}}{1-\xi} \right] & \text{for } \xi > 1 \text{ and } \xi \neq 2 \end{cases} \quad (3.4)$$

to assess the performances of a given predictor $\hat{\pi}$.

4 Performances assessment

4.1 Bregman dominance

Bregman dominance, also called forecast dominance is defined as dominance for every Bregman loss function (3.1). Precisely, $\hat{\pi}_2$ outperforms $\hat{\pi}_1$ in terms of Bregman dominance if the inequality $E[L(Y, \hat{\pi}_2)] \leq E[L(Y, \hat{\pi}_1)]$ holds true for every loss function L of the form given in (3.1). We refer the interested reader to Kruger and Ziegel (2020) and the references therein for an extensive presentation of this concept. Bregman dominance is thus a particular stochastic order relation used to compare the performances of two estimators $\hat{\pi}_1$ and $\hat{\pi}_2$ for the conditional means. We refer the reader to Shaked and Shanthikumar (2007) for a general presentation of stochastic order relations and to Denuit et al. (2005) for applications to insurance.

It is well known that any convex function $y \mapsto \ell(y)$ can be represented as a mixture of stop-loss functions $y \mapsto (y - t)_+$. Under weak regularity conditions, Ehm et al. (2016) established that every loss function (3.1) admits a mixture representation of the form

$$L(y, m) = \int L_t(y, m) dH(t)$$

where H is a non-negative measure, and

$$L_t(y, m) = (y - t)_+ - (m - t)_+ - (y - m)I[m > t].$$

Thus, every scoring function consistent for the mean can be written as a weighted average over elementary or extremal loss functions L_t . As an important consequence, an estimate that is preferable in terms of each extremal L_t is preferable in terms of any Bregman loss function and is thus superior in terms of Bregman dominance.

Based on the mixture representation derived by Ehm et al. (2016), Kruger and Ziegel (2020) established in their Theorem 2.1 that $\hat{\pi}_2$ outperforms $\hat{\pi}_1$ in terms of Bregman dominance if, and only if, the inequality

$$E[(\hat{\pi}_2 - t)_+] + E[(Y - \hat{\pi}_2)I[\hat{\pi}_2 > t]] \geq E[(\hat{\pi}_1 - t)_+] + E[(Y - \hat{\pi}_1)I[\hat{\pi}_1 > t]] \quad (4.1)$$

holds for all t . Both sides of the latter inequality is the expectation of the elementary loss function $L_t(y, m)$ once the stop-loss premiums $E[(Y - t)_+]$ appearing on both sides have been canceled. It is interesting to notice that (4.1) comprises one term involving the predictor, only (its stop-loss transform with respect to the threshold t) while the joint distribution of the pair $(Y, \hat{\pi}_2)$ only enters the second term. This echoes to the optimality conditions introduced in Denuit et al. (2019) which also combined both aspects: marginal distribution of the predictor and joint distribution with the response.

4.2 Tweedie dominance

In insurance pricing, it seems natural to restrict Bregman loss functions to Tweedie deviances (3.4). This leads to the definition of Tweedie dominance: $\hat{\pi}_2$ outperforms $\hat{\pi}_1$ in

terms of Tweedie dominance if the inequality $D(\xi, \hat{\pi}_2) \leq D(\xi, \hat{\pi}_1)$ holds true for every power parameter $\xi \geq 1$.

The next result provides the actuary with a sufficient condition for a model to outperform a competitor in terms of Tweedie dominance.

Proposition 4.1. *Define*

$$\psi_\xi(\pi) = \begin{cases} \ln \pi & \text{for } \xi = 2 \\ \frac{\pi^{2-\xi}}{2-\xi} & \text{else.} \end{cases} \quad (4.2)$$

Then, $\hat{\pi}_2$ outperforms $\hat{\pi}_1$ in terms of Tweedie dominance if

$$\mathbb{E}[\psi_\xi(\hat{\pi}_1(\mathbf{X}^{\text{new}}))] \geq \mathbb{E}[\psi_\xi(\hat{\pi}_2(\mathbf{X}^{\text{new}}))] \text{ for all } \xi \geq 1 \quad (4.3)$$

and

$$\mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}_1(\mathbf{X}^{\text{new}}) \leq t]] \geq \mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}_2(\mathbf{X}^{\text{new}}) \leq t]] \text{ for all } t \geq 0. \quad (4.4)$$

Proof. For $\xi = 1$, $\hat{\pi}_2$ is superior to $\hat{\pi}_1$ if

$$\begin{aligned} & \mathbb{E}[\hat{\pi}_1(\mathbf{X}^{\text{new}})] - \mathbb{E}[Y^{\text{new}} \ln \hat{\pi}_1(\mathbf{X}^{\text{new}})] \geq \mathbb{E}[\hat{\pi}_2(\mathbf{X}^{\text{new}})] - \mathbb{E}[Y^{\text{new}} \ln \hat{\pi}_2(\mathbf{X}^{\text{new}})] \\ \Leftrightarrow & \mathbb{E}[\hat{\pi}_1(\mathbf{X}^{\text{new}})] - \mathbb{E}[\hat{\pi}_2(\mathbf{X}^{\text{new}})] + \mathbb{E}[Y^{\text{new}} \ln \hat{\pi}_2(\mathbf{X}^{\text{new}})] - \mathbb{E}[Y^{\text{new}} \ln \hat{\pi}_1(\mathbf{X}^{\text{new}})] \geq 0. \end{aligned} \quad (4.5)$$

Since the identity

$$\begin{aligned} \mathbb{E}[Y^{\text{new}} \ln \hat{\pi}(\mathbf{X}^{\text{new}})] &= \int_0^\infty \mathbb{E}[Y^{\text{new}} \mathbb{I}[\ln \hat{\pi}(\mathbf{X}^{\text{new}}) > t]] dt - \int_{-\infty}^0 \mathbb{E}[Y^{\text{new}} \mathbb{I}[\ln \hat{\pi}(\mathbf{X}^{\text{new}}) \leq t]] dt \\ &= \int_1^\infty \mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}(\mathbf{X}^{\text{new}}) > s]] \frac{1}{s} ds - \int_0^1 \mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}(\mathbf{X}^{\text{new}}) \leq s]] \frac{1}{s} ds \end{aligned}$$

holds true for any predictor $\hat{\pi}$, we have

$$\begin{aligned} & \mathbb{E}[Y^{\text{new}} \ln \hat{\pi}_2(\mathbf{X}^{\text{new}})] - \mathbb{E}[Y^{\text{new}} \ln \hat{\pi}_1(\mathbf{X}^{\text{new}})] \\ &= \int_1^\infty (\mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}_2(\mathbf{X}^{\text{new}}) > s]] - \mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}_1(\mathbf{X}^{\text{new}}) > s]]) \frac{1}{s} ds \\ & \quad - \int_0^1 (\mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}_2(\mathbf{X}^{\text{new}}) \leq s]] - \mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}_1(\mathbf{X}^{\text{new}}) \leq s]]) \frac{1}{s} ds \\ &= \int_1^\infty (\mathbb{E}[Y^{\text{new}} (1 - \mathbb{I}[\hat{\pi}_2(\mathbf{X}^{\text{new}}) \leq s])] - \mathbb{E}[Y^{\text{new}} (1 - \mathbb{I}[\hat{\pi}_1(\mathbf{X}^{\text{new}}) \leq s])]) \frac{1}{s} ds \\ & \quad + \int_0^1 (\mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}_1(\mathbf{X}^{\text{new}}) \leq s]] - \mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}_2(\mathbf{X}^{\text{new}}) \leq s]]) \frac{1}{s} ds \\ &= \int_0^\infty (\mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}_1(\mathbf{X}^{\text{new}}) \leq s]] - \mathbb{E}[Y^{\text{new}} \mathbb{I}[\hat{\pi}_2(\mathbf{X}^{\text{new}}) \leq s]]) \frac{1}{s} ds \\ &\geq 0 \end{aligned}$$

by (4.4). Hence, both conditions (4.3) and (4.4) ensure inequality (4.5).

Turning to the case $\xi = 2$, $\hat{\pi}_2$ is superior to $\hat{\pi}_1$ if

$$\begin{aligned} & \mathbb{E} [\ln \hat{\pi}_1(\mathbf{X}^{\text{new}})] + \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_1(\mathbf{X}^{\text{new}})} \right] \geq \mathbb{E} [\ln \hat{\pi}_2(\mathbf{X}^{\text{new}})] + \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_2(\mathbf{X}^{\text{new}})} \right] \\ \Leftrightarrow & \mathbb{E} [\ln \hat{\pi}_1(\mathbf{X}^{\text{new}})] - \mathbb{E} [\ln \hat{\pi}_2(\mathbf{X}^{\text{new}})] + \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_1(\mathbf{X}^{\text{new}})} \right] - \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_2(\mathbf{X}^{\text{new}})} \right] \geq 0. \end{aligned} \quad (4.6)$$

Hence, it suffices to notice that (4.4) implies $\mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_1(\mathbf{X}^{\text{new}})} \right] - \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_2(\mathbf{X}^{\text{new}})} \right] \geq 0$ since the identity

$$\begin{aligned} \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}(\mathbf{X}^{\text{new}})} \right] &= \int_0^\infty \mathbb{E} \left[Y^{\text{new}} \mathbb{I} \left[\frac{1}{\hat{\pi}(\mathbf{X}^{\text{new}})} \geq t \right] \right] dt \\ &= \int_0^\infty \mathbb{E} \left[Y^{\text{new}} \mathbb{I} \left[\hat{\pi}(\mathbf{X}^{\text{new}}) \leq \frac{1}{t} \right] \right] dt \\ &= \int_0^\infty \mathbb{E} [Y^{\text{new}} \mathbb{I} [\hat{\pi}(\mathbf{X}^{\text{new}}) \leq s]] \frac{1}{s^2} ds. \end{aligned}$$

is valid for every predictor $\hat{\pi}$.

Finally, in the remaining cases, i.e. $\xi > 1$ and $\xi \neq 2$, $\hat{\pi}_2$ is superior to $\hat{\pi}_1$ if

$$\begin{aligned} & \mathbb{E} \left[\frac{\hat{\pi}_1(\mathbf{X}^{\text{new}})^{2-\xi}}{2-\xi} \right] + \frac{1}{\xi-1} \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_1(\mathbf{X}^{\text{new}})^{\xi-1}} \right] \geq \mathbb{E} \left[\frac{\hat{\pi}_2(\mathbf{X}^{\text{new}})^{2-\xi}}{2-\xi} \right] + \frac{1}{\xi-1} \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_2(\mathbf{X}^{\text{new}})^{\xi-1}} \right] \\ \Leftrightarrow & \mathbb{E} \left[\frac{\hat{\pi}_1(\mathbf{X}^{\text{new}})^{2-\xi}}{2-\xi} \right] - \mathbb{E} \left[\frac{\hat{\pi}_2(\mathbf{X}^{\text{new}})^{2-\xi}}{2-\xi} \right] + \frac{1}{\xi-1} \left(\mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_1(\mathbf{X}^{\text{new}})^{\xi-1}} \right] - \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_2(\mathbf{X}^{\text{new}})^{\xi-1}} \right] \right) \geq 0. \end{aligned} \quad (4.7)$$

Whatever the predictor $\hat{\pi}$, we can write

$$\begin{aligned} \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}(\mathbf{X}^{\text{new}})^{\xi-1}} \right] &= \int_0^\infty \mathbb{E} [Y^{\text{new}} \mathbb{I} [\hat{\pi}(\mathbf{X}^{\text{new}})^{1-\xi} \geq t]] dt \\ &= \int_0^\infty \mathbb{E} \left[Y^{\text{new}} \mathbb{I} \left[\hat{\pi}(\mathbf{X}^{\text{new}}) \leq \frac{1}{t^{\frac{1}{\xi-1}}} \right] \right] dt \\ &= \int_0^\infty \mathbb{E} [Y^{\text{new}} \mathbb{I} [\hat{\pi}(\mathbf{X}^{\text{new}}) \leq s]] \frac{\xi-1}{s^\xi} ds, \end{aligned}$$

The announced result then follows from

$$\begin{aligned} & \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_1(\mathbf{X}^{\text{new}})^{\xi-1}} \right] - \mathbb{E} \left[\frac{Y^{\text{new}}}{\hat{\pi}_2(\mathbf{X}^{\text{new}})^{\xi-1}} \right] \\ &= \int_0^\infty \left(\mathbb{E} [Y^{\text{new}} \mathbb{I} [\hat{\pi}_1(\mathbf{X}^{\text{new}}) \leq s]] - \mathbb{E} [Y^{\text{new}} \mathbb{I} [\hat{\pi}_2(\mathbf{X}^{\text{new}}) \leq s]] \right) \frac{\xi-1}{s^\xi} ds. \end{aligned}$$

This ends the proof. $\square$

Notice that from the proof of Proposition 4.1, one can easily deduce a mixture representation of the predictive deviance (3.4) used for model comparison, that is,

$$D(\xi, \hat{\pi}) = \int_0^\infty \left(\frac{1}{t^{\xi-1}} P[\hat{\pi} \leq t] - \frac{1}{t^\xi} E[YI[\hat{\pi} \leq t]] \right) dt. \quad (4.8)$$

This representation of $D(\xi, \hat{\pi})$ enables to relate Bregman dominance to Tweedie dominance. Recall that if $\hat{\pi}_2$ outperforms $\hat{\pi}_1$ in terms of Bregman dominance then the expected loss for every loss function L of the form given in (3.1) is smaller for $\hat{\pi}_2$, which is in particular the case for predictive Tweedie deviances. From

$$\begin{aligned} E[(\hat{\pi}_1 - t)I[\hat{\pi}_1 > t]] + E[(Y - \hat{\pi}_1)I[\hat{\pi}_1 > t]] &= E[(Y - t)I[\hat{\pi}_1 > t]] \\ &= E[Y] - t - E[YI[\hat{\pi}_1 \leq t]] + tP[\hat{\pi}_1 \leq t], \end{aligned}$$

we can see that inequality (4.1) can be rewritten as

$$E[Y] - t - E[YI[\hat{\pi}_2 \leq t]] + tP[\hat{\pi}_2 \leq t] \geq E[Y] - t - E[YI[\hat{\pi}_1 \leq t]] + tP[\hat{\pi}_1 \leq t]$$

or equivalently as

$$tP[\hat{\pi}_2 \leq t] - E[YI[\hat{\pi}_2 \leq t]] \geq tP[\hat{\pi}_1 \leq t] - E[YI[\hat{\pi}_1 \leq t]]. \quad (4.9)$$

Then, dividing both sides of (4.9) by t^ξ leads to

$$\frac{1}{t^{\xi-1}} P[\hat{\pi}_2 \leq t] - \frac{1}{t^\xi} E[YI[\hat{\pi}_2 \leq t]] \geq \frac{1}{t^{\xi-1}} P[\hat{\pi}_1 \leq t] - \frac{1}{t^\xi} E[YI[\hat{\pi}_1 \leq t]], \quad (4.10)$$

so that (4.8) yields $D(\xi, \hat{\pi}_2) \leq D(\xi, \hat{\pi}_1)$, as expected.

Instead of Tweedie dominance, the actuary could select a specific parameter ξ , only. This is typically the case with $\xi = 1$ (Poisson regression for counts), $\xi = 2$ or $\xi = 3$ (Gamma or Inverse Gaussian regression for claim severities). In this case, condition (4.3) in Proposition 4.1 is imposed only for that specific valued of ξ .

The main ingredient of the proof of Proposition 4.1 is the integral of the difference of functions in (4.4) weighted by

$$\frac{\xi - 1 + I[\xi = 1]}{s^\xi} = \frac{\xi - 1 + I[\xi = 1]}{V(s)}$$

where $V(\cdot)$ is the variance function associated with the response distribution. Weights thus appear to be inversely proportional to the variability. Under the Tweedie variance function, $V(s)$ is a power of s so that smaller weights are assigned to the differences at larger values of the response.

Proposition 4.1 shows that improving a predictor $\hat{\pi}$ can be achieved by

- (i) increasing the overall bias measured on a modified scale induced by the auxiliary function ψ_ξ defined in (4.2). This may act against the conservation of observed totals.

(ii) increasing the dependence between $\hat{\pi}$ and the response, so to decrease

$$\begin{aligned} \mathbb{E}\left[Y^{\text{new}}\mathbb{I}[\hat{\pi}(\mathbf{X}^{\text{new}}) \leq t]\right] &= \text{Cov}\left[Y^{\text{new}}, \mathbb{I}[\hat{\pi}(\mathbf{X}^{\text{new}}) \leq t]\right] + \mathbb{E}\left[Y^{\text{new}}\right]\mathbb{P}\left[\hat{\pi}(\mathbf{X}^{\text{new}}) \leq t\right] \\ &= \mathbb{E}\left[Y^{\text{new}}\middle|\hat{\pi}(\mathbf{X}^{\text{new}}) \leq t\right]\mathbb{P}\left[\hat{\pi}(\mathbf{X}^{\text{new}}) \leq t\right]. \end{aligned}$$

The latter lower partial moment essentially depends on the correlation structure of the pair $(Y^{\text{new}}, \hat{\pi}(\mathbf{X}^{\text{new}}))$. Notice that the quantities appearing in the decomposition above are closely related to expectation dependence as defined by Wright (1987).

In an insurance ratemaking context, the lower partial moment can be interpreted as best-profile premium income

$$\begin{aligned} \mathbb{E}\left[Y^{\text{new}}\mathbb{I}[\hat{\pi}(\mathbf{X}^{\text{new}}) \leq t]\right] &= \mathbb{E}\left[\mu(\mathbf{X}^{\text{new}})\mathbb{I}[\hat{\pi}(\mathbf{X}^{\text{new}}) \leq t]\right] \\ &\approx \frac{1}{n} \sum_{i|\hat{\pi}(\mathbf{X}_i) \leq t} \mu(\mathbf{X}_i). \end{aligned}$$

Here, $\sum_{i|\hat{\pi}(\mathbf{X}_i) \leq t} \mu(\mathbf{X}_i)$ is the true premium income for the sub-portfolio formed by gathering all policyholders with predicted premium at most equal to t . These policyholders exhibit the best risk profiles according to the candidate premium $\hat{\pi}$.

Of course, there is a trade-off between these two goals when computing $D(\xi, \hat{\pi})$. If the model is very flexible then it can produce a predictor that correlates a lot with the response and the lower partial moment can be decreased to a large extent compared to models imposing a rigid form for the predictor (like GLMs). More bias can then be allowed.

Remark 4.2. The discussion in the present paper holds the training set as fixed. This is different from the discussion devoted to model stability, or bias-variance trade-off. Indeed, variance is meant as a property varying the training set, and the same for bias. Expectations and variances are thus taken with respect to varying training sets when bias-variance trade-off is discussed whereas here, the training set is kept fixed and we consider that a large enough validation set is available. We believe that this approach is more in line with insurance practice.

4.3 Application to Poisson regression

Poisson deviance is by far the most widely used one in insurance applications. It applies for instance to claim counts in property and casualty insurance, death counts in life insurance, and numbers of transitions in health insurance. We assume that we deal with a response Y obeying the Poisson distribution. A vector $\mathbf{X}$ of features is available to predict the mean response $\mu(\mathbf{X})$. To ease the presentation, we assume unit exposures.

Considering (4.2)-(4.3) with $\xi = 1$, the bias is thus measured on the response scale. Condition (4.3) then favors $\hat{\pi}_2$ if

$$\mathbb{E}[\mu(\mathbf{X}^{\text{new}})] - \mathbb{E}[\hat{\pi}_2(\mathbf{X}^{\text{new}})] \geq \mathbb{E}[\mu(\mathbf{X}^{\text{new}})] - \mathbb{E}[\hat{\pi}_1(\mathbf{X}^{\text{new}})]$$

or, equivalently,

$$\mathbb{E}[Y^{\text{new}}] - \mathbb{E}[\hat{\pi}_2(\mathbf{X}^{\text{new}})] \geq \mathbb{E}[Y^{\text{new}}] - \mathbb{E}[\hat{\pi}_1(\mathbf{X}^{\text{new}})].$$

A larger bias is thus advantageous. As a consequence, adopting Poisson deviance as loss function outside GLMs may create a total premium income gap

$$\begin{aligned} \mathbb{E}[Y^{\text{new}}] - \mathbb{E}[\hat{\pi}(\mathbf{X}^{\text{new}})] &= \mathbb{E}[\mu(\mathbf{X}^{\text{new}})] - \mathbb{E}[\hat{\pi}(\mathbf{X}^{\text{new}})] \\ &\approx \frac{1}{n} \sum_{i=1}^n \mu(\mathbf{X}_i) - \frac{1}{n} \sum_{i=1}^n \hat{\pi}(\mathbf{X}_i) \end{aligned}$$

where $\sum_{i=1}^n \mu(\mathbf{X}_i)$ is the true premium income for the validation set whereas $\sum_{i=1}^n \hat{\pi}(\mathbf{X}_i)$ is the one obtained by adopting predictor $\hat{\pi}$ for premium calculation. This gap may become quite large with highly flexible models such as Neural Networks or boosting.

5 Restoring balance at global and local scales

In order to give properties of a predictor $\hat{\pi}$, let $F_{\hat{\pi}}^{-1}$ be the quantile function (or Value-at-Risk) of $\hat{\pi}$, defined as the generalized inverse of its distribution function $F_{\hat{\pi}}$, that is,

$$F_{\hat{\pi}}^{-1}(\alpha) = \inf\{t | F_{\hat{\pi}}(t) \geq \alpha\} \text{ for a probability level } \alpha.$$

One can also consider quantile based region as sets of $\mathcal{X}$, associated with $\hat{\pi}$, defined either as

$$\mathcal{X}_{\hat{\pi}, \alpha\%} = \{\mathbf{x} \in \mathcal{X} : \hat{\pi}(\mathbf{x}) \leq F_{\hat{\pi}}^{-1}(\alpha)\}$$

or

$$\mathcal{X}_{\hat{\pi}, \overline{\alpha\%}} = \{\mathbf{x} \in \mathcal{X} : \hat{\pi}(\mathbf{x}) > F_{\hat{\pi}}^{-1}(\alpha)\}$$

Heuristically, $\mathcal{X}_{\hat{\pi}, \alpha\%}$ is the set of insured characteristics related to the smallest risk, with proportion $\alpha\%$, while $\mathcal{X}_{\hat{\pi}, \overline{\alpha\%}}$ is related to the largest risks, with proportion $1 - \alpha\%$ (where small and large risks are related to predictor $\hat{\pi}$). Of course, those sets can also be defined for the true premium μ .

On Figure 5.1, we can visualize set $\mathcal{X}_{\mu, 20\%}$ and $\mathcal{X}_{\mu, \overline{80\%}}$ of bottom-20% and top-20% true risks, according to μ , in the univariate case described in Example 2.1 and the bivariate case described in Example 2.2. On Figure 5.2, sets $\mathcal{X}_{\hat{\pi}, 20\%}$ and $\mathcal{X}_{\hat{\pi}, \overline{80\%}}$ can be visualized, with two GAM models, with a low degree of freedom on the left, and a high degree of freedom on the right, for the univariate case described in Example 2.1. On Figure 5.3, sets $\mathcal{X}_{\hat{\pi}, 20\%}$ and $\mathcal{X}_{\hat{\pi}, \overline{80\%}}$ can be visualized (respectively on the left and on the right), with an additive GAM model on top ($\hat{\pi}(x_1, x_2) = \hat{s}_1(x_1) + \hat{s}_2(x_2)$), and a bivariate spline based Poisson model below ($\hat{\pi}(x_1, x_2) = \hat{s}(x_1, x_2)$), for the bivariate case described in Example 2.2. Plain areas are $\mathcal{X}_{\hat{\pi}, \alpha}$ sets, while dashed ones are $\mathcal{X}_{\mu, \alpha}$ sets.

5.1 Autocalibration

Recall that a predictor $\hat{\pi}$ is said to be autocalibrated if $\hat{\pi}(\mathbf{X}) = \mathbb{E}[Y | \hat{\pi}(\mathbf{X})]$. We refer the reader to Kruger and Ziegel (2020) for a general presentation of this concept. By Jensen inequality, autocalibration thus ensures that

$$\mathbb{E}[g(\hat{\pi}(\mathbf{X}))] \leq \mathbb{E}[g(Y)] \text{ for every convex function } g,$$

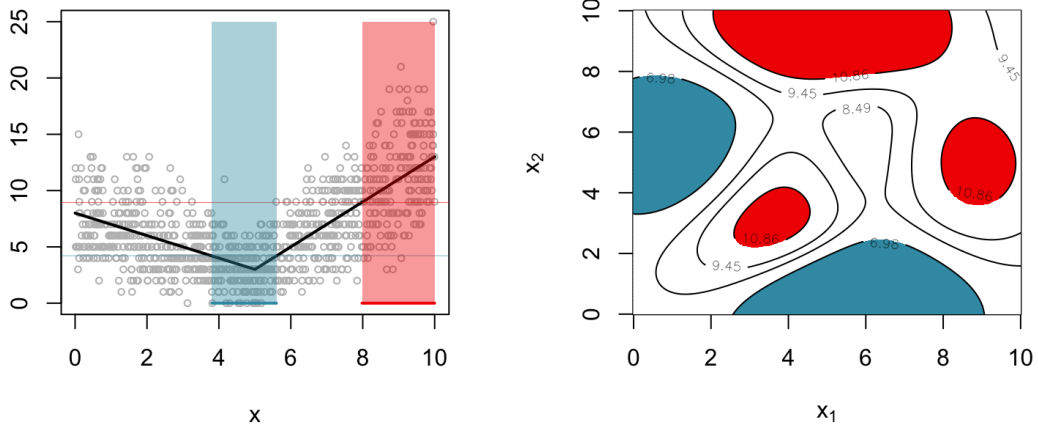


Figure 5.1: Theoretical quantiles for μ , $\mathcal{X}_{\mu,20\%}$ and $\mathcal{X}_{\mu,80\%}$ for the Poisson counts described in Example 2.1 in the left panel and in Example 2.2 in the right panel.

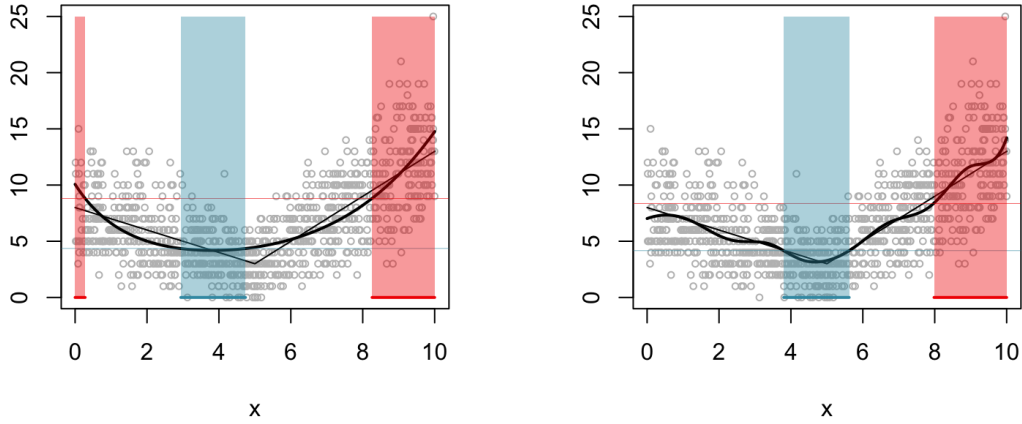


Figure 5.2: Model based quantiles for simulated Poisson counts described in Example 2.1, according to $\hat{\pi}_{\text{low-splines}}$ (Poisson regression with splines and low degree of freedom, on the left) and $\hat{\pi}_{\text{high-splines}}$ (Poisson regression with splines and high degree of freedom, on the right), $\mathcal{X}_{\hat{\pi},20\%}$ and $\mathcal{X}_{\hat{\pi},80\%}$.

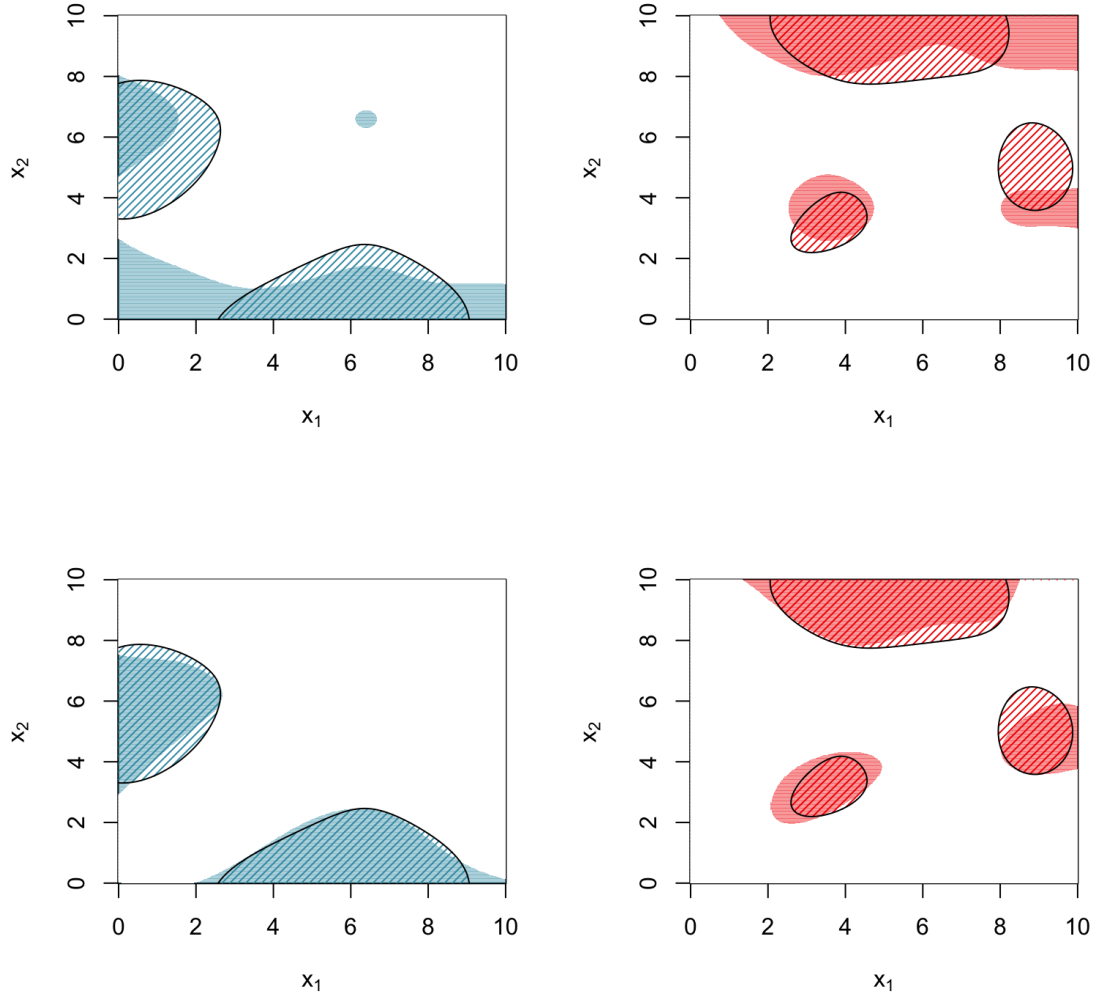


Figure 5.3: Model based quantiles for simulated Poisson counts described in Example 2.2, according to $\hat{\pi}_{\text{add-splines}}$ (Poisson regression with additive splines, on top) and $\hat{\pi}_{\text{biv-splines}}$ (Poisson regression with bivariate splines, below), in plain colors, $\mathcal{X}_{\hat{\pi}, 20\%}$ and $\mathcal{X}_{\hat{\pi}, 80\%}$. Dashed regions are theoretical quantiles.

or equivalently, that

$$\mathbb{E}[\hat{\pi}(\mathbf{X})] = \mathbb{E}[Y] \text{ and } \mathbb{E}[(\hat{\pi}(\mathbf{X}) - t)_+] \leq \mathbb{E}[(Y - t)_+] \text{ for all } t \geq 0.$$

These inequalities correspond to the convex order between $\hat{\pi}(\mathbf{X})$ and Y . Thus, autocalibration implies that the predictor is less variable than the response, in the sense of the convex order.

Notice that Bregman dominance reduces to the convex order for autocalibrated predictors, as pointed out by Kruger and Ziegel (2020). This is easily deduced from (4.1) since

$$\mathbb{E}[(Y - \hat{\pi})\mathbb{I}[\hat{\pi} > t]] = \mathbb{E}[(Y - \mathbb{E}[Y|\hat{\pi}])\mathbb{I}[\hat{\pi} > t]] = 0$$

for every autocalibrated predictor.

Under mild technical requirement, a simple way to restore global balance consists in switching from $\hat{\pi}$ to its balance-corrected version $\hat{\pi}_{\text{BC}}$ defined as

$$\hat{\pi}_{\text{BC}}(\mathbf{X}) = \mathbb{E}[Y|\hat{\pi}(\mathbf{X})]$$

that averages to $\mathbb{E}[Y]$, as shown in the next result.

Property 5.1. *If $s \mapsto \mathbb{E}[Y|\hat{\pi}(\mathbf{X}) = s]$ is continuously increasing then the balance-corrected version $\hat{\pi}_{\text{BC}}$ of the candidate premium $\hat{\pi}$ satisfies the autocalibration property.*

Proof. If $s \mapsto \mathbb{E}[Y|\hat{\pi}(\mathbf{X}) = s]$ is continuously increasing, that is, if $\hat{\pi}_{\text{BC}}(\mathbf{X})$ is continuously increasing in $\hat{\pi}(\mathbf{X})$, then Lemma 2.2 in Shaked et al. (2012) allows us to write

$$\mathbb{E}[Y|\hat{\pi}(\mathbf{X})] = \mathbb{E}[Y|\hat{\pi}_{\text{BC}}(\mathbf{X})] = \hat{\pi}_{\text{BC}}(\mathbf{X})$$

so that the resulting $\hat{\pi}_{\text{BC}}(\mathbf{X})$ is indeed autocalibrated. This ends the proof. $\square$

In insurance applications, autocalibration induces local balance and imposes financial equilibrium not only at portfolio level but also in any sufficiently large sub-portfolio. This concept thus appears to be particularly appealing in a ratemaking context.

Heuristically, function $s \mapsto \mathbb{E}[Y|\hat{\pi}(\mathbf{X}) = s]$ is the average value of points $\mu(\hat{\pi}^{-1}(s))$, where $\hat{\pi}^{-1}(s)$ simply denotes the set of $\mathbf{x}$ fulfilling this condition. To visualize such a function, consider the univariate case discussed in Example 2.1. Consider first the case where $\hat{\pi}$ is obtained from a standard Poisson GLM regression with log-link, so that $\hat{\pi}$ is strictly monotonic, and therefore convertible. On Figure 5.4, on the left $\hat{\pi}$ is the strong curve. The horizontal line is at level s . There, $x = \hat{\pi}^{-1}(s)$ is the axis of the two points, the one below being the one of interest, $\mu(\hat{\pi}^{-1}(s))$ that we wish to compute. With several values s , we get a collection of points $\mu(\hat{\pi}^{-1}(s))$ that are on the right of Figure 5.4. Here, $s \mapsto \mathbb{E}[Y|\hat{\pi}(X) = s]$ is not monotonic, but $\hat{\pi}$ is clearly underfitting. On Figure 5.5, two GAM models are considered, and $\hat{\pi}$ is no-longer monotonic. On top, $\hat{\pi}$ is a convex function, and $\hat{\pi}^{-1}(s)$ is generally a collection of two points. Function $s \mapsto \mathbb{E}[Y|\hat{\pi}(X) = s]$ is then the average of those two points, since in this example, we assume that X is uniformly distributed (otherwise, weights should be proportional to the density). This gives the black strong line that we can observe on the right, with a simple GAM model, with a low degree of freedom, on top, and a GAM model

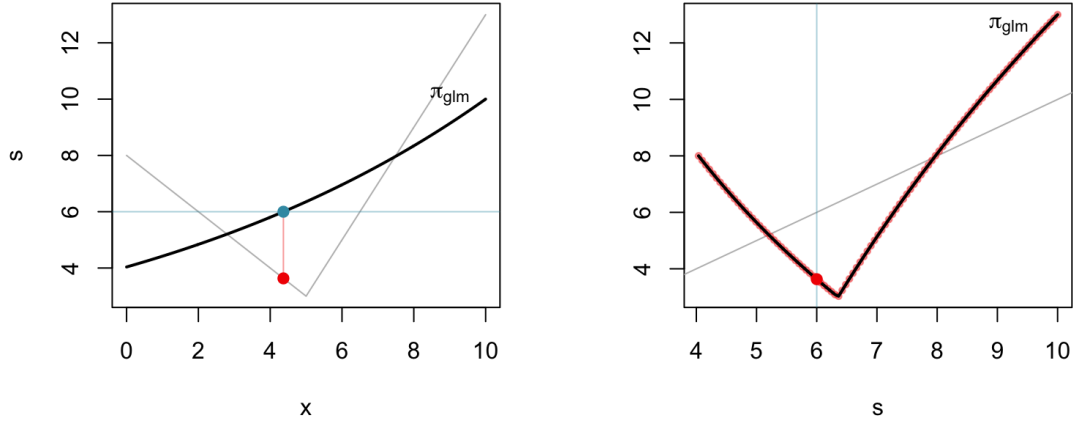


Figure 5.4: Evolution of $s \mapsto E[Y|\hat{\pi}(X) = s] = \mu(\hat{\pi}^{-1}(s))$ for a simple Poisson regression (glm), where $\hat{\pi}$ is monotonic, on the right. On the left, we can visualize μ and $\hat{\pi}$, for all x 's. s is on the y -axis, $(\hat{\pi}^{-1}(s), s)$ is the blue point, $(\hat{\pi}^{-1}(s), \mu(\hat{\pi}^{-1}(s)))$ is the red point. Points $s \mapsto \mu(\hat{\pi}^{-1}(s))$ can be visualized on the right side.

with more degrees of freedom, below. Note that on those two examples, $s \mapsto E[Y|\hat{\pi}(X) = s]$ is an increasing function.

In those examples, we computed the true function $s \mapsto E[Y|\hat{\pi}(X) = s]$, since we used simulated data, with true premium μ , and a known distribution for X . In practice, it is not possible, but one can approximate that function using a local regression on the pairs $(\hat{\pi}(x_i), y_i)$. This has been performed on Figure 5.6, with the GLM on the left (the true function is the thin line in the back) and the GAM model on the right. Note that here also, if $\hat{\pi}$ does not underfit, the approximation of $s \mapsto E[Y|\hat{\pi}(X) = s]$ is an increasing function.

In the bivariate case described in Example 2.2, the set of $\mathbf{X}$'s such that $\hat{\pi}(\mathbf{X}) = s$ is now a collection of curves if $\hat{\pi}$ is continuous, as we can visualize on Figure 5.7, with the same model, but two levels s , on top and below. Here $\hat{\pi}$ is the standard additive GAM model. On the right, we can visualize the distribution of the random variable $\mu(\hat{\pi}^{-1}(s))$, that is also Y given $\hat{\pi}(\mathbf{X}) = s$. The expected value is represented by the vertical lines. Varying s , we obtain the curve on the left of Figure 5.8, which is $s \mapsto E[\mu(\mathbf{X})|\hat{\pi}(\mathbf{X}) = s]$. The two points are the two values of Figure 5.7. Note here that the function is not increasing, for low values of s , but μ was absolutely not an additive function, so again, underfitting seems to compromise the increasing property of $s \mapsto E[\mu(\mathbf{X})|\hat{\pi}(\mathbf{X}) = s]$ (that we will have with bivariate splines). The approximated version (without the knowledge of μ) is drawn on the right.

On Figure 5.9, π is now the bivariate spline based Poisson regression model. On Figure 5.10, we can visualize the theoretical function $s \mapsto E[\mu(\mathbf{X})|\hat{\pi}(\mathbf{X}) = s]$ on the left, and its approximated version on the right. Increasingness is clearly visible there.

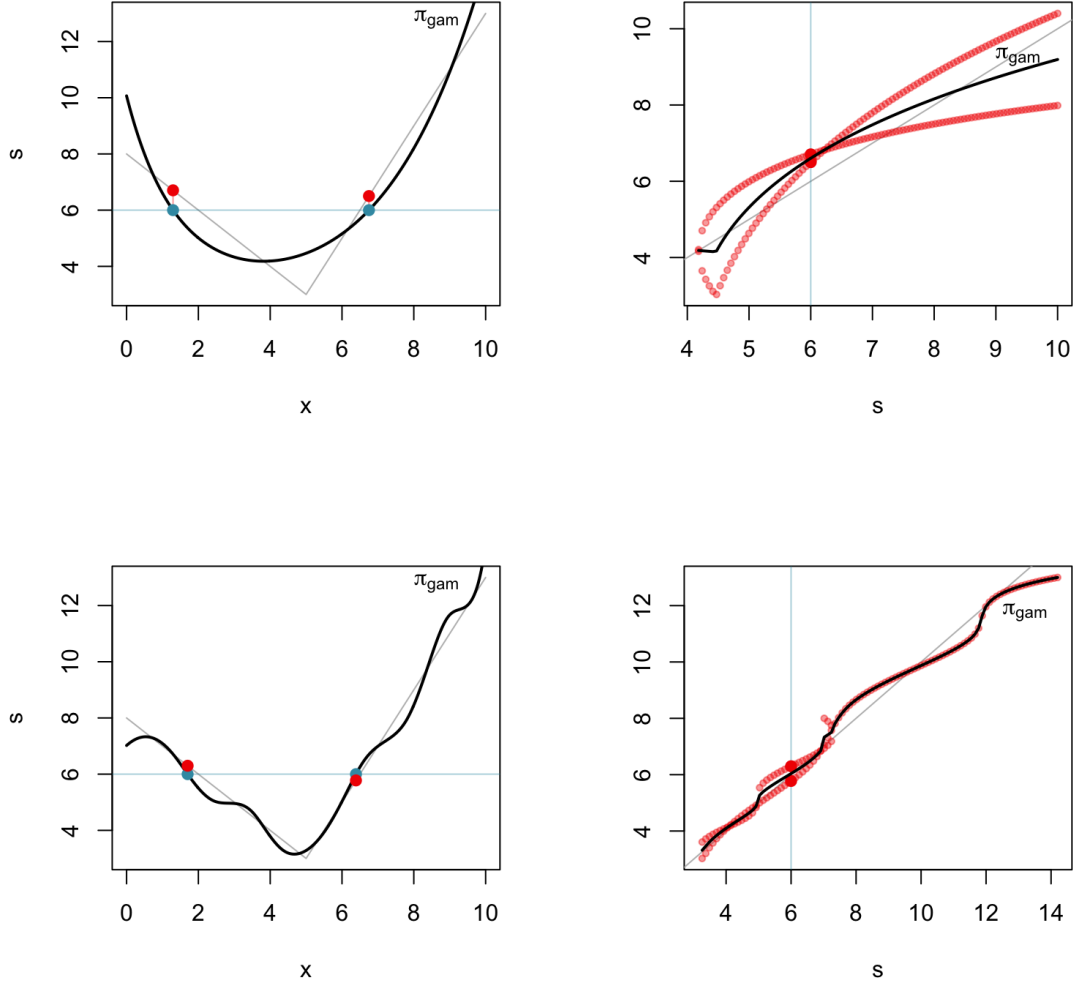


Figure 5.5: Evolution of $s \mapsto \mathbb{E}[Y|\hat{\pi}(X) = s] = \mathbb{E}[\mu(\hat{\pi}^{-1}(s))]$ for two Poisson regression, with splines (gam, with more degrees of freedom below) on the right. On the left, we can visualize μ and $\hat{\pi}$, for all x 's. s is on the y -axis, $(\hat{\pi}^{-1}(s), s)$ are the blue points. Since $\hat{\pi}$ is no longer monotonic there can be several x 's such that $\hat{\pi}(x) = s$. $(\hat{\pi}^{-1}(s), \mu(\hat{\pi}^{-1}(s)))$ are the red points. The average of points $\mu(\hat{\pi}^{-1}(s))$'s is function $s \mapsto \mathbb{E}[\mu(\hat{\pi}^{-1}(s))]$ that can be visualized on the right side, with the plain black line.

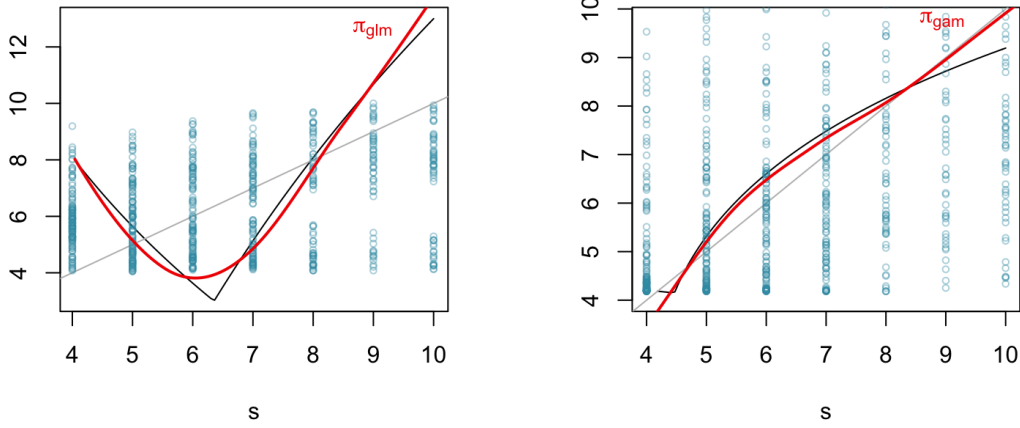


Figure 5.6: Evolution of $s \mapsto E[Y|\hat{\pi}(X) = s] = E[\mu(\hat{\pi}^{-1}(s))]$ for the Poisson GLM (on the left) and GAM (on the right), with the plain thin dark line. The approximated version of that function is drawn, from the local regression on the pairs $(\hat{\pi}(x_i), y_i)$, in red.

5.2 Autocalibrating a given predictor

We can restore global balance, or unbiased the predictor by reconciling the predicted and observed total on the training set. In this way, we recover the global balance property imposed after the seminal work by Bailey and Simon (1960). However, this does not ensure that balance holds locally. Indeed, global balance is only one of the GLM likelihood equations under canonical link, corresponding to the intercept. In order to extend the other likelihood equations imposed in the GLM setting to general machine learning procedures, we need to mimic the way local GLM proceeds for fitting, by defining meaningful neighborhoods for statistical learning.

To this end, we work under canonical link function. The only feature entering the analysis is the logarithm of the predictor, that is, $\ln \hat{\pi}(\mathbf{x}_i)$ so that the local GLM is thus fitted to $(Y_i, e_i, \hat{\pi}(\mathbf{x}_i))$ for some relevant exposure e_i . An intuitively acceptable solution would consist in imposing marginal constraints on local neighborhoods defined by mean of $\hat{\pi}$. This allows for some local transfers of claims and premiums from neighboring policyholders and so implements local balance conditions in sub-portfolios corresponding to these neighborhoods. This is in essence the local GLM approach (see Loader, 1999, for a detailed account) that allows the actuary to maintain the relationship with MMT.

In order to obtain an autocalibrated version $\hat{\pi}_{BC}$ of $\hat{\pi}$, let us consider a specific risk profile $\mathbf{x}$. A weight $\nu_i(\hat{\pi}(\mathbf{x}))$ is assigned to each $(Y_i, e_i, \hat{\pi}(\mathbf{x}_i))$, $i = 1, \dots, n$, computed from some weight function $\nu(\cdot)$ chosen to be continuous, symmetric, peaked at 0 and defined on $[-1, 1]$. These weights depend on the relative distance of $\hat{\pi}(\mathbf{x}_i)$ with respect to $\hat{\pi}(\mathbf{x})$. A common choice for $\nu(\cdot)$ is the tricube weight function but several alternatives are available, including rectangular (or uniform), Gaussian or Epanechnikov, for instance. Here, $\nu_i(\hat{\pi}(\mathbf{x}))$ is larger for policyholders i such that $\hat{\pi}(\mathbf{x}_i)$ is close to $\hat{\pi}(\mathbf{x})$ and decreases when $\hat{\pi}(\mathbf{x}_i)$ gets far away

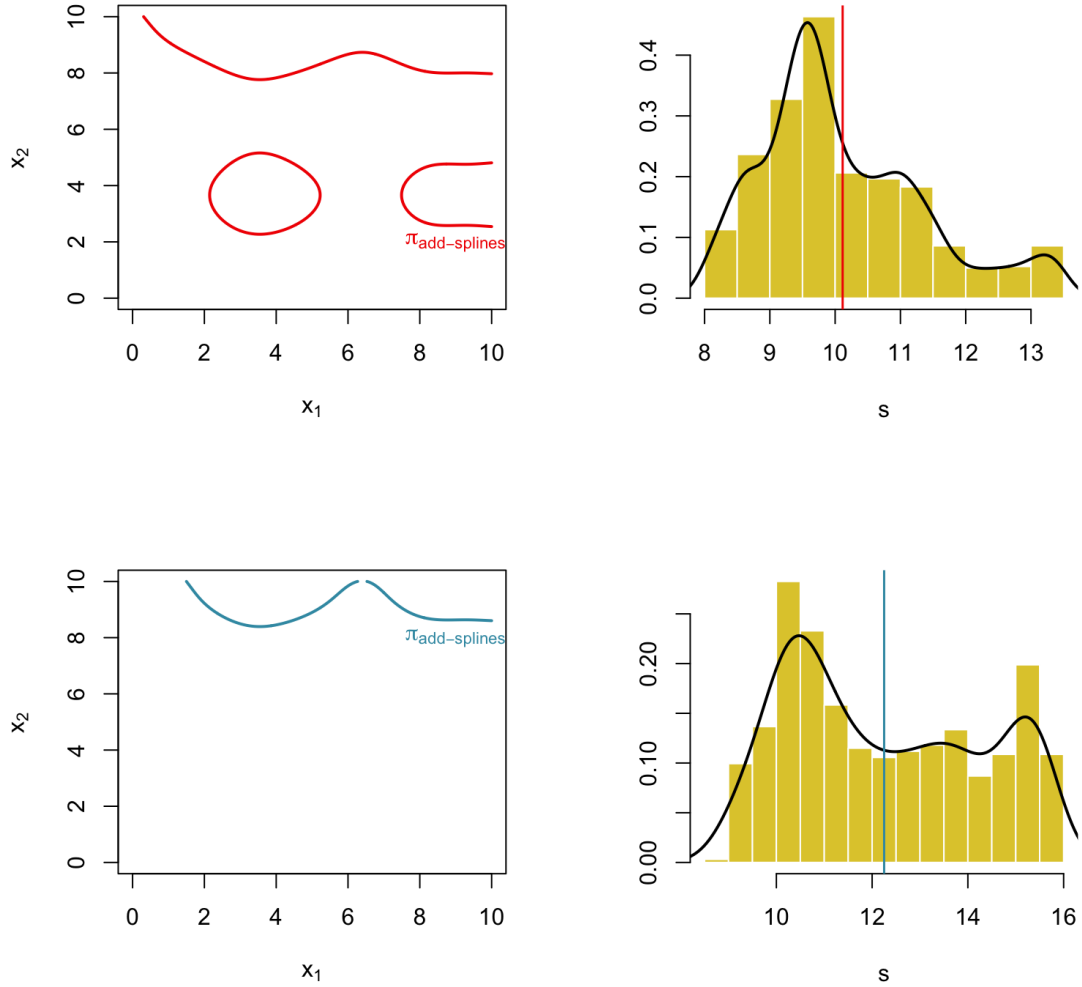


Figure 5.7: On the left, sets of points $\mathbf{x}$ such that $\hat{\pi}(\mathbf{x}) = s$ for $s = 10$ on top and $s = 11.5$ below, when $\hat{\pi}$ is the prediction from the classical GAM model, with additive components $s_1(x_1) + s_2(x_2)$. On the right, density of variable $\mu(\mathbf{X})$ conditional on $\hat{\pi}(\mathbf{X}) = s$. Vertical lines are $E[\mu(\mathbf{X}) | \hat{\pi}(\mathbf{X}) = s]$ for the two values of s .

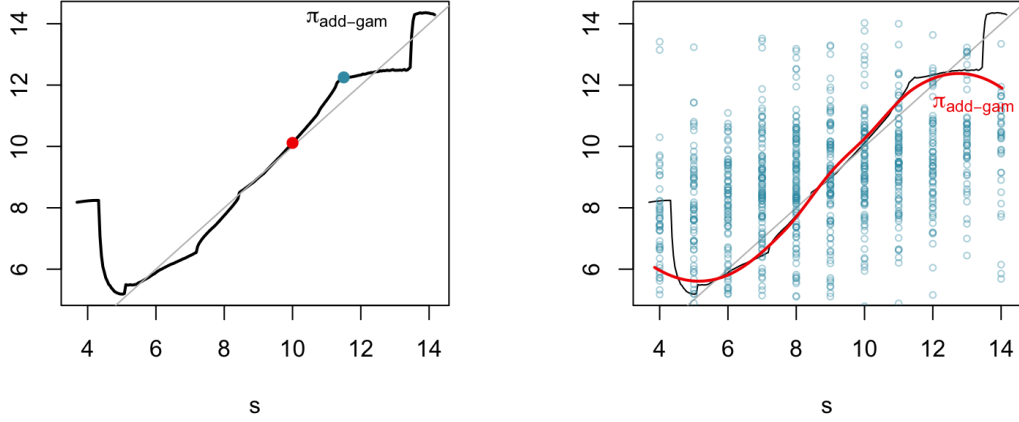


Figure 5.8: Evolution of $s \mapsto E[Y|\hat{\pi}(\mathbf{X}) = s] = E[\mu(\hat{\pi}^{-1}(s))]$ for the bivariate spline Poisson model, with the theoretical curve on the left. Points red and blue correspond to the levels discussed on Figure 5.9. The approximated version of that function is drawn, from the local regression on the pairs $(\hat{\pi}(\mathbf{x}_i), y_i)$, in red.

from $\hat{\pi}(\mathbf{x})$.

The local GLM likelihood equation

$$\sum_{i=1}^n \nu_i(\hat{\pi}(\mathbf{x})) y_i = \sum_{i=1}^n \nu_i(\hat{\pi}(\mathbf{x})) e_i \hat{\pi}_{BC}(\mathbf{x})$$

thus matches MMT constraints: smoothing is ensured by transferring part of the experience at neighboring $\hat{\pi}$ values to obtain $\hat{\pi}_{BC}$. A local constant, or intercept-only GLM thus provides the appropriate fitting procedure when balance must be respected. Opting for a rectangular weight function complies with MMT: the weights ν_i are constant within the smoothing window and zero otherwise so that the sums reduce to observations comprised within this window and the uniform weights factor out.

Our main message here is thus that a local, intercept-only GLM with rectangular weights implements local balance, or MMT in sub-portfolios gathering policyholders with about the same predicted value. The rectangular weight function involved in the statistical procedure optimally transfers part of claim experience between neighboring policyholders. This approach implements smoothness from a statistical point of view while remaining fully transparent and understandable. Indeed, local averaging can just be seen as an application of the mutuality principle at the heart of insurance.

5.3 Link with Lorenz and concentration curves

Denuit et al. (2019) suggested to use concentration and Lorenz curves to evaluate the performances of a candidate predictor or to compare two competing alternatives. As defined

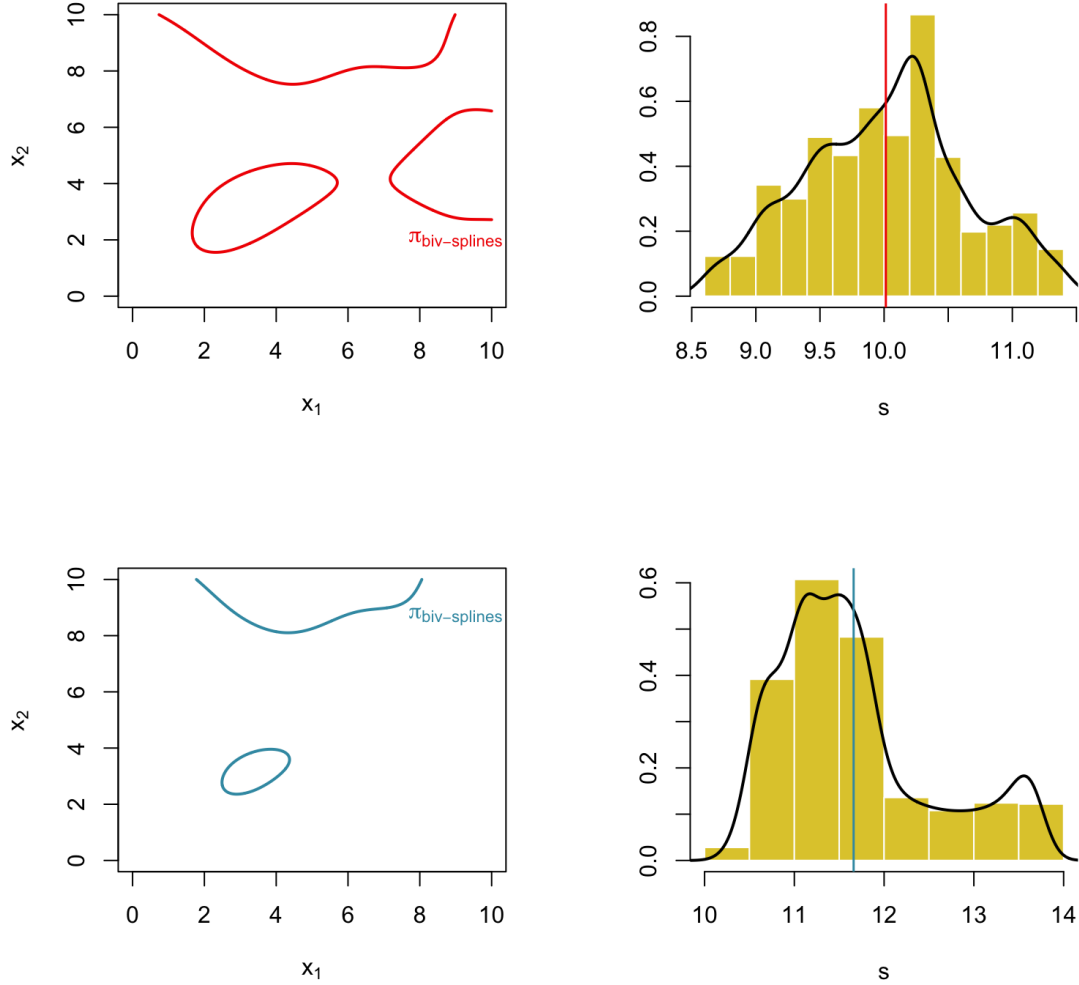


Figure 5.9: On the left, sets of points $\mathbf{x}$ such that $\hat{\pi}(\mathbf{x}) = s$ for $s = 10$ on top and $s = 11.5$ below, when $\hat{\pi}$ is the prediction from the Poisson regression with bivariate spline $s(x_1, x_2)$. On the right, density of variable $\mu(\mathbf{X})$ conditional on $\hat{\pi}(\mathbf{X}) = s$. Vertical lines are $E[\mu(\mathbf{X}) | \hat{\pi}(\mathbf{X}) = s]$ for the two values of s .

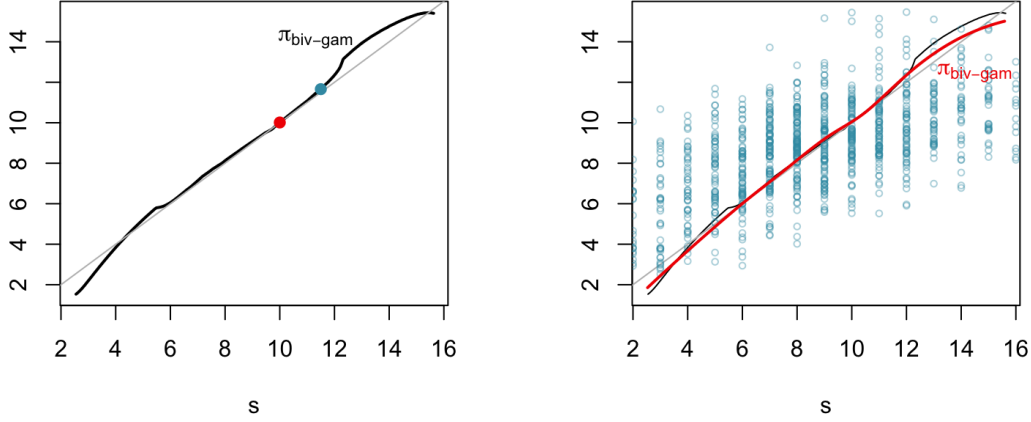


Figure 5.10: Evolution of $s \mapsto E[Y|\hat{\pi}(\mathbf{X}) = s] = E[\mu(\hat{\pi}^{-1}(s))]$ for the bivariate spline Poisson model, with the theoretical curve on the left. Points red and blue correspond to the levels discussed on Figure 5.9. The approximated version of that function is drawn, from the local regression on the pairs $(\pi(\mathbf{x}_i), y_i)$, in red.

previously, $F_{\hat{\pi}}^{-1}$ be the quantile function (or Value-at-Risk) of $\hat{\pi}$. On Figure 5.11, we can visualize the evolution of $F_{\hat{\pi}}$ (and actually also F_{μ}) in the univariate case discussed in Example 2.1, and the bivariate case of Example 2.2.

The concentration curve of the true premium $\mu(\mathbf{X})$ with respect to the working premium $\hat{\pi}$ based on the information contained in the vector $\mathbf{X}$ is defined there as

$$\alpha \mapsto \text{CC}[\mu(\mathbf{X}), \hat{\pi}(\mathbf{X}); \alpha] = \frac{E[\mu(\mathbf{X})\mathbf{I}[\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)]]}{E[\mu(\mathbf{X})]}.$$

We see that $\text{CC}[\mu(\mathbf{X}), \hat{\pi}(\mathbf{X}); \alpha]$ represents the percentage of the total premium income corresponding to the sub-portfolio $\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)$, i.e. to the $\alpha\%$ of contracts with the smallest premium $\hat{\pi}$. For an exhaustive review of the properties of a concentration curve we refer the interested reader to the book by Yitzhaki and Schechtman (2013). The Lorenz curve LC associated with the predictor $\hat{\pi}(\mathbf{X})$ is defined as

$$\begin{aligned} \alpha \mapsto \text{LC}[\hat{\pi}(\mathbf{X}); \alpha] &= \text{CC}[\hat{\pi}(\mathbf{X}), \hat{\pi}(\mathbf{X}); \alpha] \\ &= \frac{E[\hat{\pi}(\mathbf{X})\mathbf{I}[\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)]]}{E[\hat{\pi}(\mathbf{X})]}. \end{aligned}$$

The next result shows that Lorenz and concentration curves coincide for autocalibrated predictors.

Property 5.2. *If $\hat{\pi}$ is autocalibrated then*

$$\text{LC}[\hat{\pi}(\mathbf{X}); \alpha] = \text{CC}[\mu(\mathbf{X}), \hat{\pi}(\mathbf{X}); \alpha]$$

for every probability level α .

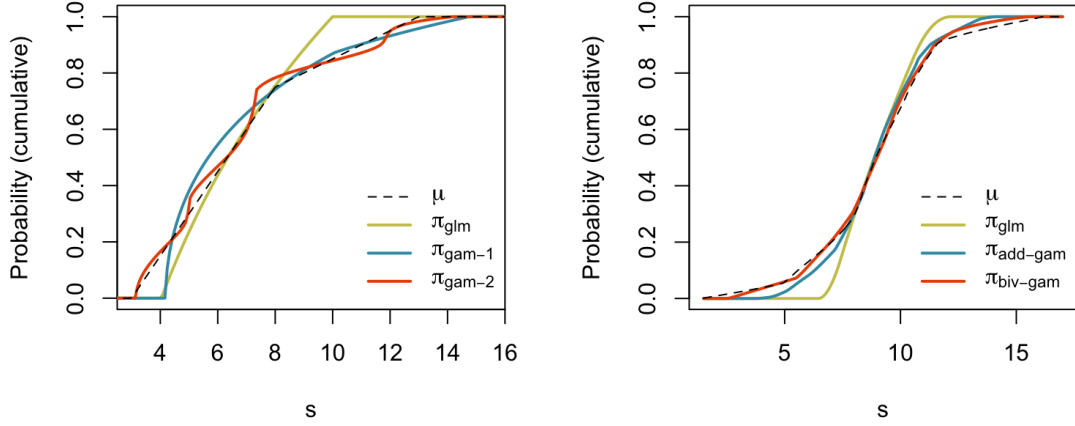


Figure 5.11: Evolution of $F_{\hat{\pi}}$, the distribution function of $\hat{\pi}$, i.e. $s \mapsto \mathbb{P}[\hat{\pi}(\mathbf{X}) \leq s]$, in the univariate case of Example 2.1 on the left, and the bivariate case of Example 2.2 on the right.

Proof. It suffices to write

$$\begin{aligned}
\text{LC}[\hat{\pi}(\mathbf{X}); \alpha] &= \frac{\mathbb{E}[\hat{\pi}(\mathbf{X}) \mathbb{I}[\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)]]}{\mathbb{E}[\hat{\pi}(\mathbf{X})]} \\
&= \frac{\mathbb{E}[\mathbb{E}[Y | \hat{\pi}(\mathbf{X})] \mathbb{I}[\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)]]}{\mathbb{E}[\hat{\pi}(\mathbf{X})]} \\
&= \frac{\mathbb{E}[\mathbb{E}[Y \mathbb{I}[\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)] | \hat{\pi}(\mathbf{X})]]}{\mathbb{E}[\hat{\pi}(\mathbf{X})]} \\
&= \frac{\mathbb{E}[Y \mathbb{I}[\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)]]}{\mathbb{E}[\hat{\pi}(\mathbf{X})]} \\
&= \frac{\mathbb{E}[\mathbb{E}[Y \mathbb{I}[\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)] | \mathbf{X}]]}{\mathbb{E}[\hat{\pi}(\mathbf{X})]} \\
&= \frac{\mathbb{E}[\mathbb{E}[Y | \mathbf{X}] \mathbb{I}[\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)]]}{\mathbb{E}[\hat{\pi}(\mathbf{X})]} \\
&= \frac{\mathbb{E}[\mu(\mathbf{X}) \mathbb{I}[\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)]]}{\mathbb{E}[\hat{\pi}(\mathbf{X})]} \\
&= \frac{\mathbb{E}[\mu(\mathbf{X}) \mathbb{I}[\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)]]}{\mathbb{E}[\mu(\mathbf{X})]} \\
&= \text{CC}[\mu(\mathbf{X}), \hat{\pi}(\mathbf{X}); \alpha],
\end{aligned}$$

which ends the proof. $\square$

The Area Between the Curves, or ABC metric is defined after Denuit et al. (2019) as the area between the two curves CC and LC. It is interesting to notice that $\text{ABC}[\hat{\pi}(\mathbf{X})] = 0$ for an

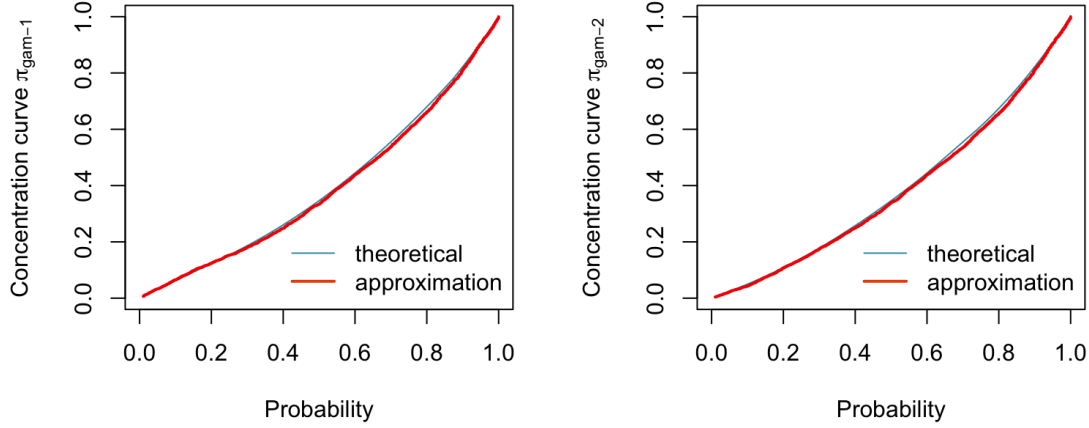


Figure 5.12: Evolution of $\text{CC}(\hat{\pi}, \mu; \alpha)$ and $\widetilde{\text{CC}}(\hat{\pi}, \mu; \alpha)$ in the univariate case.

autocalibrated predictor. A nonzero estimated ABC thus suggests that the predictor under consideration is not autocalibrated. Initially proposed as an indicator of the performance of a given predictor, it turns out that ABC is more useful for detecting the need to implement autocalibration.

Functions $\alpha \mapsto \text{CC}[\mu(\mathbf{X}), \hat{\pi}(\mathbf{X}); \alpha]$ can be visualized on the cases of Examples 2.1 and 2.2. On Figure 5.12 we can visualize the two GAM models in the univariate case of Example 2.1. The theoretical line is the plot of CC , while the approximated one, denote $\widetilde{\text{CC}}$ is

$$\alpha \mapsto \frac{1}{n\bar{y}} \sum_{i=1}^n y_i I[\hat{\pi}(\mathbf{x}_i) \leq q_\alpha]$$

where q_α is the empirical quantile of level α of $\hat{\pi}(\mathbf{x}_i)$'s. Note that computations of integrals for the theoretical lines were performed, here using quasi-Monte Carlo techniques.

On Figure 5.13 we can visualize not only the two GAM models of Example 2.1, in the univariate case, but also the GLM. On the right, we can visualize the derivative of that concentration curve. The analogous in the bivariate case of Example 2.2 can be visualized on Figure 5.14.

5.4 Bregman dominance under autocalibration

Let $\hat{\pi}_1$ and $\hat{\pi}_2$ be two autocalibrated predictors. Then, $\hat{\pi}_2$ outperforms $\hat{\pi}_1$ in terms of Bregman dominance if, and only if,

$$\begin{aligned} & \hat{\pi}_2 \text{ dominates } \hat{\pi}_1 \text{ in convex order} \\ \Leftrightarrow & \text{LC}[\hat{\pi}_2(\mathbf{X}); \alpha] \leq \text{LC}[\hat{\pi}_1(\mathbf{X}); \alpha] \text{ for all probability levels } \alpha \\ \Leftrightarrow & \text{CC}[\mu(\mathbf{X}), \hat{\pi}_2(\mathbf{X}); \alpha] \leq \text{CC}[\mu(\mathbf{X}), \hat{\pi}_1(\mathbf{X}); \alpha] \text{ for all probability levels } \alpha \end{aligned}$$

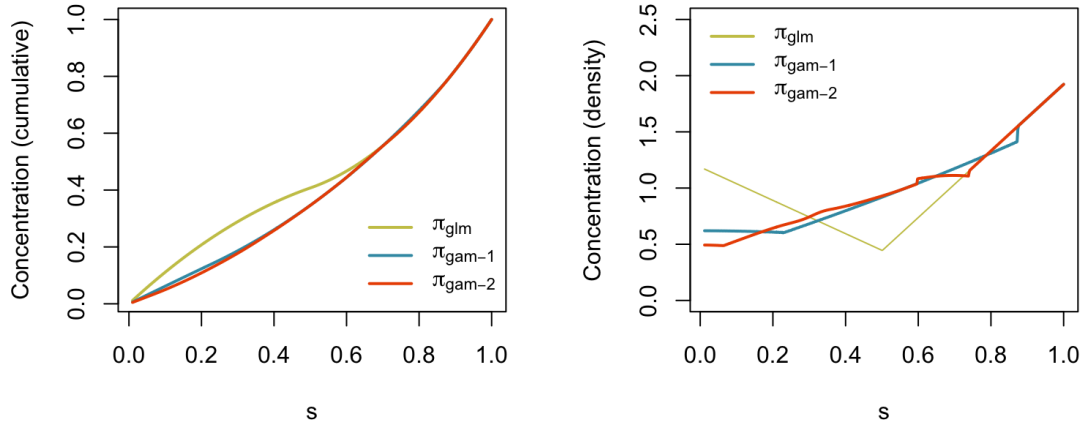


Figure 5.13: Evolution of $CC(\hat{\pi}, \mu; \alpha)$ and $CC'(\hat{\pi}, \mu; \alpha)$ is on the right, in the univariate case.

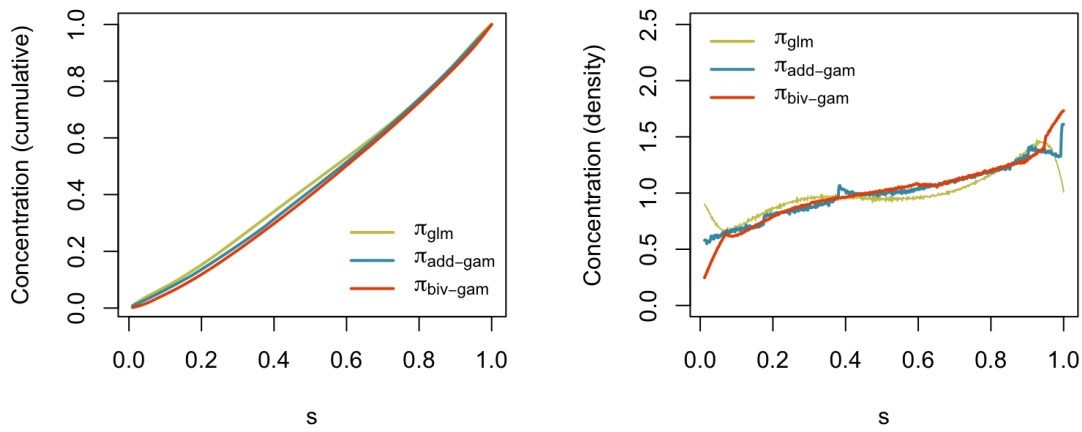


Figure 5.14: Evolution of $CC(\hat{\pi}, \mu; \alpha)$ and $CC'(\hat{\pi}, \mu; \alpha)$ is on the right, in the bivariate case.

since Lorenz and concentration curves coincide for autocalibrated predictors. It is interesting to notice that, for autocalibrated predictors, $\hat{\pi}_2$ outperforms $\hat{\pi}_1$ in terms of Bregman dominance if, and only if, $\hat{\pi}_2$ is more discriminatory than $\hat{\pi}_1$ in the sense defined in Denuit et al. (2019).

Compared to the preceding sections, thresholds are now replaced with quantiles (and thus differ with the predictor under consideration), in line with what is done in practice when lift diagrams are used. There is thus a link between concentration curves used at validation stage (Denuit et al., 2019) and the adoption of deviance as objective function for calibration. Notice also that for two predictors $\hat{\pi}_1$ and $\hat{\pi}_2$ such that $s \mapsto E[Y|\hat{\pi}_k(\mathbf{X}) = s]$ is continuously increasing for $k \in \{1, 2\}$, $\hat{\pi}_{2,BC}$ outperforms $\hat{\pi}_{1,BC}$ in terms of Bregman dominance if, and only if, $CC[\mu(\mathbf{X}), \hat{\pi}_{2,BC}(\mathbf{X}); \alpha] \leq CC[\mu(\mathbf{X}), \hat{\pi}_{1,BC}(\mathbf{X}); \alpha]$ for all α , or equivalently, if, and only if, $CC[\mu(\mathbf{X}), \hat{\pi}_2(\mathbf{X}); \alpha] \leq CC[\mu(\mathbf{X}), \hat{\pi}_1(\mathbf{X}); \alpha]$ for all α . Therefore, one prefers $\hat{\pi}_2$ over $\hat{\pi}_1$ if $CC[\mu(\mathbf{X}), \hat{\pi}_2(\mathbf{X}); \alpha] \leq CC[\mu(\mathbf{X}), \hat{\pi}_1(\mathbf{X}); \alpha]$ for all α , so that $\hat{\pi}_{2,BC}$ outperforms $\hat{\pi}_{1,BC}$ in terms of Bregman dominance.

As noticed in Denuit et al. (2019), two predictors might well be incomparable because their respective concentration curves intersect. In such a case, we can base the comparison on the integral of the concentration curve defined as

$$ICC[\mu(\mathbf{X}), \hat{\pi}(\mathbf{X})] = \int_0^1 CC[\mu(\mathbf{X}), \hat{\pi}(\mathbf{X}); \alpha] d\alpha.$$

For two predictors $\hat{\pi}_1$ and $\hat{\pi}_2$ such that $s \mapsto E[Y|\hat{\pi}_k(\mathbf{X}) = s]$ is continuously increasing for $k \in \{1, 2\}$ and $ICC[\mu(\mathbf{X}), \hat{\pi}_2(\mathbf{X})] \leq ICC[\mu(\mathbf{X}), \hat{\pi}_1(\mathbf{X})]$, one thus have

$$ICC[\mu(\mathbf{X}), \hat{\pi}_{2,BC}(\mathbf{X})] \leq ICC[\mu(\mathbf{X}), \hat{\pi}_{1,BC}(\mathbf{X})]$$

and

$$ABC[\hat{\pi}_{2,BC}(\mathbf{X})] = ABC[\hat{\pi}_{1,BC}(\mathbf{X})] = 0,$$

so that $\hat{\pi}_{2,BC}(\mathbf{X})$ is preferred over $\hat{\pi}_{1,BC}(\mathbf{X})$ according to ICC (and ABC) metrics.

6 Application on real data

In order to illustrate the techniques discussed in the previous section, consider the **freMTPL2freq** dataset, from the **CASDataset** R package of Charpentier (2014). The variable of interest is here is the annual claim frequency, and the classical model will be a Poisson one, on variable **ClaimNb**, with exposure **Exposure** and various explanatory variables (continuous and categorical).

6.1 Construction of three models

Three models will be considered : a generalized linear model $\hat{\pi}^{glm}$, a generalized additive model $\hat{\pi}^{gam}$ where continuous features are transformed nonlinearly using spline functions, a boosting algorithm $\hat{\pi}^{bst}$. For the last models, we used the **h2o** package, with **h2o.gbm**. In order to model our counting variable, we set **distribution** = "poisson" and **offset_column** = "Exposure". For the boosting inference procedure, we used **ntrees** = 30 (the impact

of the number of iterations will be discussed later on) and `nfolds = 5`¹. From the initial dataset, with 678,013 rows, 70% are randomly chosen for our training dataset (474,609 rows) that is used to construct $\hat{\pi}$, and the remaining 30% (203,404 rows) are used as a validation dataset, to draw the following graphs.

On Figure 6.1, we can visualise the distribution of predictions $\hat{\pi}(\mathbf{x}_i)$. The average value $\bar{\pi}$ for the four models, on the training dataset, is given on the left of Table 6.1 (additional information is given, namely quantiles). As a benchmark, if we run a simple Poisson regression of the intercept only, we obtain $\hat{\beta}_0 = -2.2911$, corresponding to a baseline prediction $\bar{\pi} = 0.10115$.

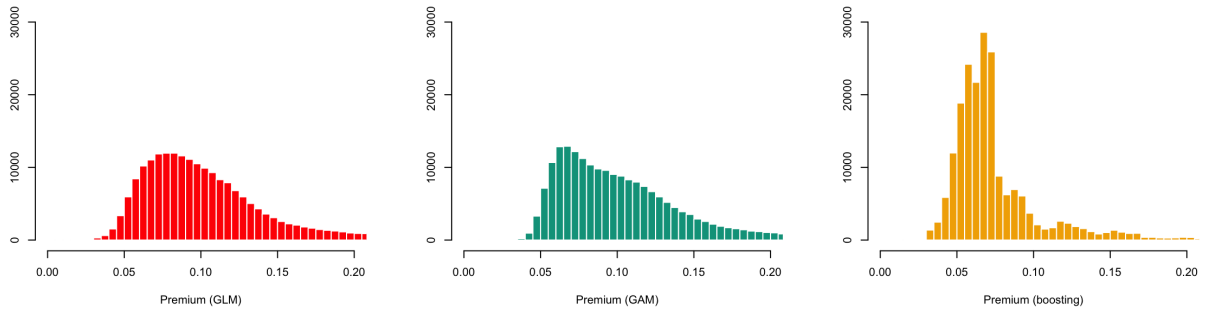


Figure 6.1: Histogram of $\{\hat{\pi}(\mathbf{x}_1), \dots, \hat{\pi}(\mathbf{x}_n)\}$ on the validation dataset. In this section, ■ corresponds to the standard Poisson regression (GLM), ■ corresponds to the smooth additive regression (GAM), and ■ is the boosting model.

	$\hat{\pi}^{\text{glm}}$	$\hat{\pi}^{\text{gam}}$	$\hat{\pi}^{\text{bst}}$		$\hat{\pi}_{BC}^{\text{glm}}$	$\hat{\pi}_{BC}^{\text{gam}}$	$\hat{\pi}_{BC}^{\text{bst}}$
average $\bar{\pi}$	0.1087	0.1092	0.0820	average $\bar{\pi}$	0.1063	0.1072	0.1037
10% quantile	0.0605	0.0598	0.0498	10% quantile	0.0606	0.0545	0.0500
90% quantile	0.1682	0.1713	0.1244	90% quantile	0.1706	0.1821	0.1747

Table 6.1: Summary statistics on $\{\hat{\pi}(\mathbf{x}_1), \dots, \hat{\pi}(\mathbf{x}_n)\}$, on the validation dataset (assuming an exposure of 1 to provide annualized predictions), for π on the left, and the corrected version π_{BC} on the right.

On Figure 6.2, we can visualize the evolution of $s \mapsto \mathbb{E}[Y|\hat{\pi}(\mathbf{X}) = s]$ when s varies on the range of possible premiums (here $s \in [0, 0.2]$). For the Generalized Linear Model $\mathbb{E}[Y|\hat{\pi}(\mathbf{X}) = s] \sim s$, while $\mathbb{E}[Y|\hat{\pi}(\mathbf{X}) = s] > s$ for the boosting model. The same information can be visualized on Figure 6.3, where the x -axis is not the premium, but the quantile version, i.e. $u \mapsto \mathbb{E}[Y|\hat{\pi}(\mathbf{X}) = F_{\hat{\pi}}^{-1}(u)]$. The positive bias we observe on $\hat{\pi}^{\text{bst}}$ means that this model *underestimates* the true price of the risk. As shown on Table 6.1, the bias correction will help to correct.

¹The codes used in this section can be found on the github repository <https://github.com/freakonometrics/autocalibration>

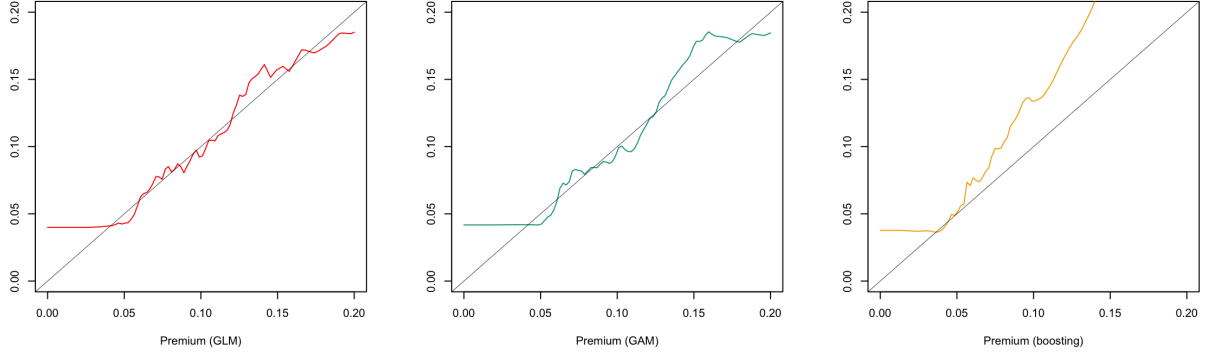


Figure 6.2: Evolution of $s \mapsto E[Y|\hat{\pi}(\mathbf{X}) = s]$ (the thin straight line corresponds to $\mu = \hat{\pi}$).

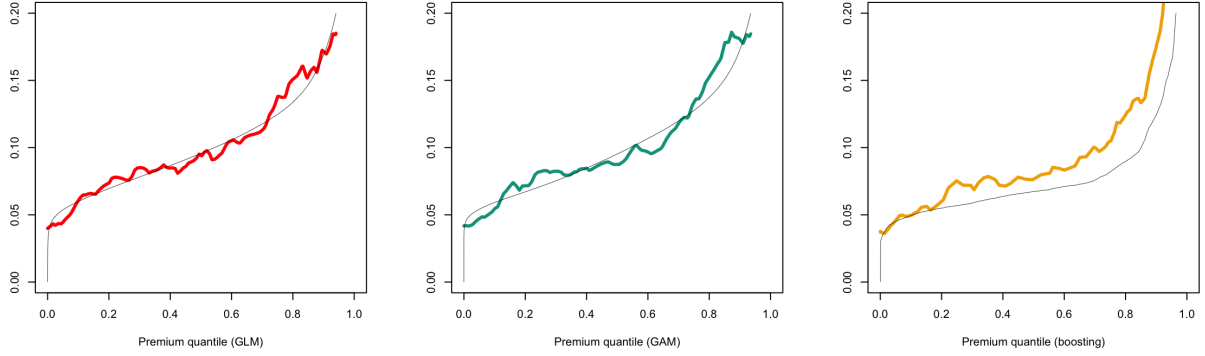


Figure 6.3: Evolution of $u \mapsto E[Y|\hat{\pi}(\mathbf{X}) = F_{\hat{\pi}}^{-1}(u)]$ (the thin straight line corresponds to $\mu = \hat{\pi}$).

On Figure 6.4, we can visualize the evolution of the difference between $E[Y|\hat{\pi}(\mathbf{X}) = F_{\hat{\pi}}^{-1}(u)]$ and $F_{\hat{\pi}}^{-1}(u)$, when u varies. This is interpreted as the additive local bias on model $\hat{\pi}$. A multiplicative version can be visualized on Figure 6.5, with $u \mapsto E[Y|\hat{\pi}(\mathbf{X}) = F_{\hat{\pi}}^{-1}(u)]/F_{\hat{\pi}}^{-1}(u)$, which is the (local) multiplier used to go from $\hat{\pi}$ to $\hat{\pi}_{BC}$.

6.2 From $\hat{\pi}$ to $\hat{\pi}_{BC}$

On Figure 6.6, we can compare π and $\hat{\pi}_{BC}$: on top, we have the QQ plot of $\hat{\pi}_{BC}$ against $\hat{\pi}$, which corresponds to the plot $u \mapsto \{F_{\hat{\pi}}^{-1}(u), F_{\hat{\pi}_{BC}}^{-1}(u)\}$, and on the bottom, a scatterplot of $\{\hat{\pi}(\mathbf{x}_i), \hat{\pi}_{BC}(\mathbf{x}_i)\}$ (for a subset of the validation database). The distribution of $\hat{\pi}_{BC}$ can be visualized on Figure 6.7.

On Figure 6.8, we can visualize the evolution of $u \mapsto E[Y|\hat{\pi}_{BC}(\mathbf{X}) = F_{\hat{\pi}_{BC}}^{-1}(u)]$, on the unit interval $[0, 1]$. On Figure 6.9, we can visualize the evolution of the difference between $E[Y|\hat{\pi}_{BC}(\mathbf{X}) = F_{\hat{\pi}_{BC}}^{-1}(u)]$ and $F_{\hat{\pi}}^{-1}(u)$, when u varies, which is the additive local bias on model $\hat{\pi}$. A multiplicative version can be visualized on Figure 6.10, with $u \mapsto E[Y|\hat{\pi}(\mathbf{X}) = F_{\hat{\pi}_{BC}}^{-1}(u)]/F_{\hat{\pi}_{BC}}^{-1}(u)$. On the boosting version, we see clearly the impact of the correction, and

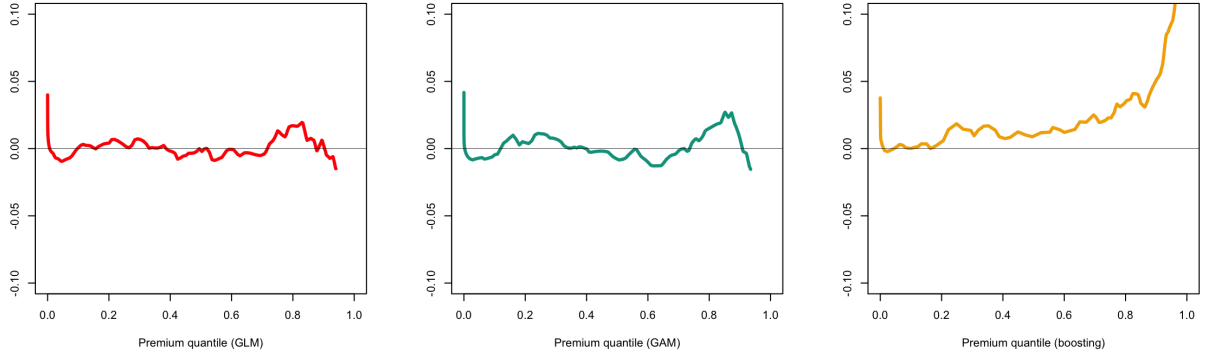


Figure 6.4: Evolution of $u \mapsto E[Y|\hat{\pi}(\mathbf{X}) = F_{\hat{\pi}}^{-1}(u)] - F_{\hat{\pi}}^{-1}(u)$ (the thin horizontal straight line $y = 0$ corresponds to $\mu = \hat{\pi}$).

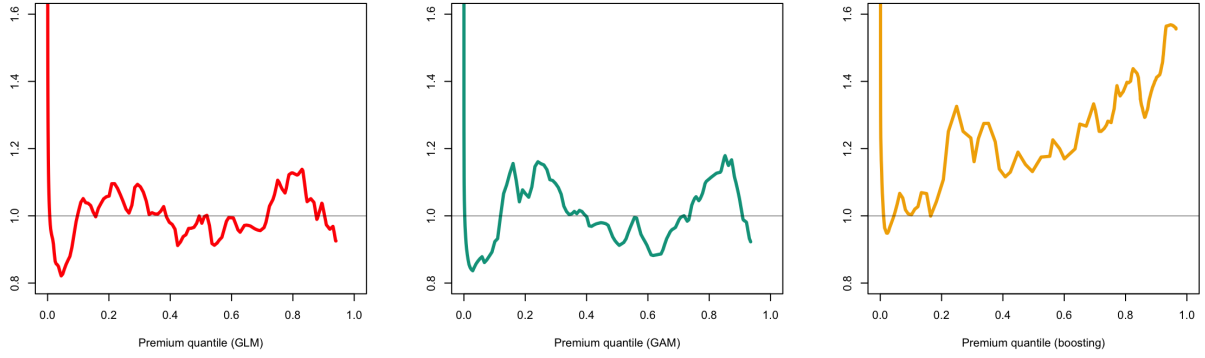


Figure 6.5: Evolution of $u \mapsto E[Y|\hat{\pi}(\mathbf{X}) = F_{\hat{\pi}}^{-1}(u)]/F_{\hat{\pi}}^{-1}(u)$ (the thin horizontal straight line $y = 1$ corresponds to $\mu = \hat{\pi}$).

that $\hat{\pi}_{\text{BC}}^{\text{bst}}$ is now locally unbiased.

6.3 Partial Dependence Plots

In order to compare the prices induced by this correction, we can compute partial dependence plots, i.e. for a variable of interest x_k out of $\mathbf{x}$, consider $x \mapsto E[\hat{\pi}(x, \mathbf{X}_{-k})]$ where $\hat{\pi}(x, \mathbf{x}_{-k})$ corresponds to vector $\mathbf{x}$ where x is substituted to x_k . The empirical version is

$$\text{pdp}_k(x) = \frac{1}{n} \sum_{i=1}^n \hat{\pi}(x_{1,i}, \dots, x_{k-1,i}, x, x_{k+1,i}, \dots, x_{p,i}).$$

On Figure 6.11, we can see the partial dependence plot for the age the driver, **DrivAge**, while on Figure 6.12, we can see the partial dependence plot for the population density in the city the driver lives in, **Density**.

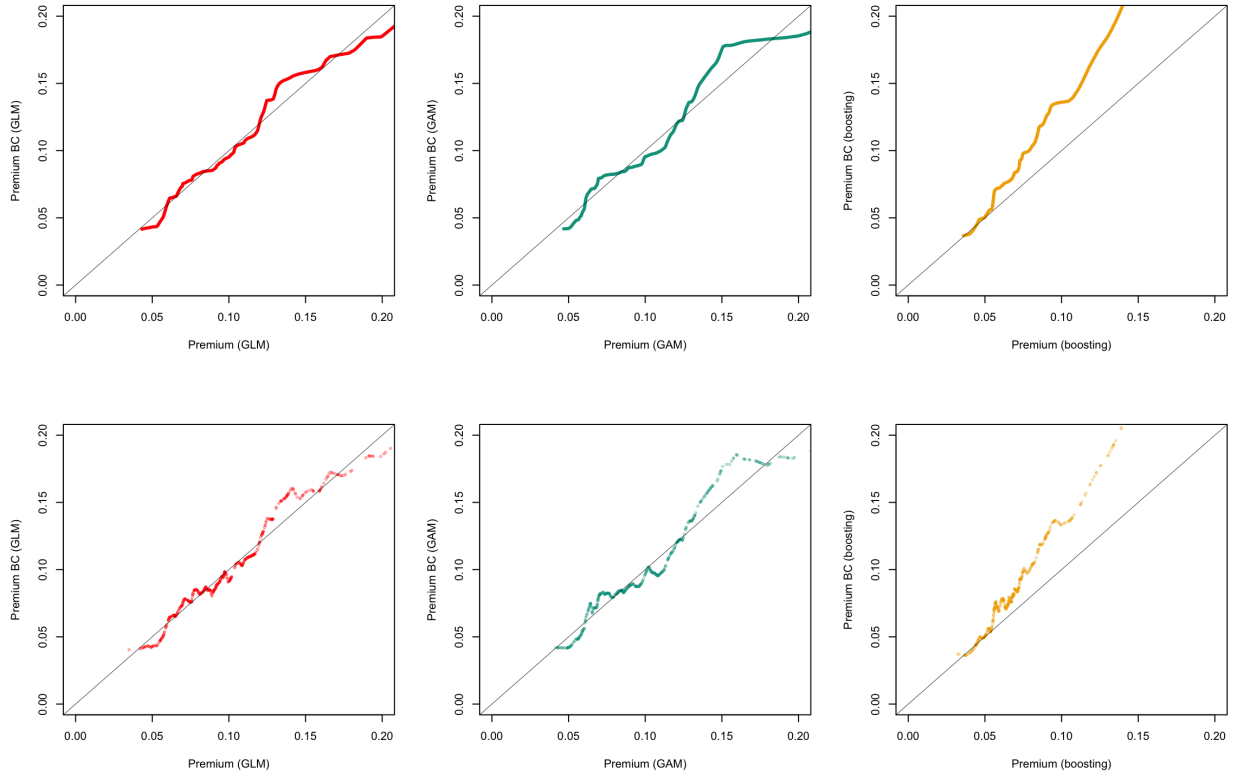


Figure 6.6: QQ plot of $\hat{\pi}_{BC}$ against $\hat{\pi}$ (plain line), on top, and scatterplot a subset of points $\{\hat{\pi}(\mathbf{x}_i), \hat{\pi}_{BC}(\mathbf{x}_i)\}$ from the validation dataset, below.

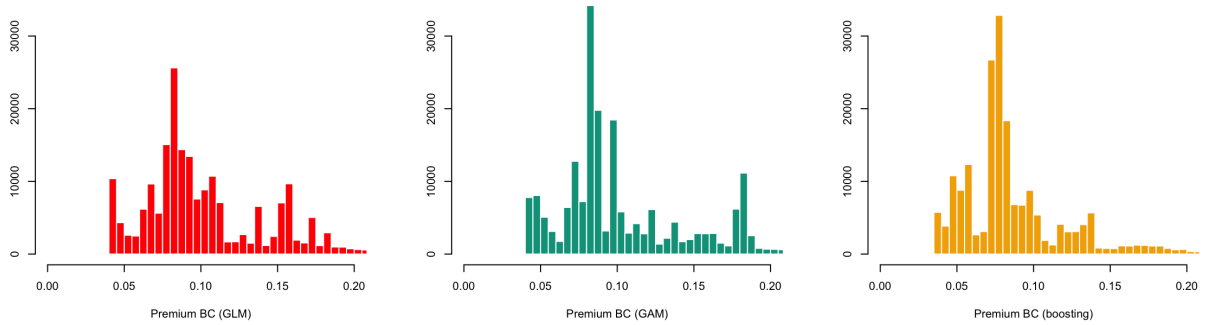


Figure 6.7: Histogram of $\{\hat{\pi}_{BC}(\mathbf{x}_1), \dots, \hat{\pi}_{BC}(\mathbf{x}_n)\}$ on the validation dataset.

6.4 Smoothing parameter α and off-sample validation

The specification of bandwidth $h(s)$ is discussed in Loader (1999, Section 2.2.1). For a nearest neighbor bandwidth, compute all distances $|s - s_i|$ between the fitting point s and the data points s_i , we choose $h(x)$ to be the k th smallest distance, where $k = \lceil n\alpha \rceil$. In the R function `locfit`, when `alpha` is given as a single number, it represents a nearest neighbor

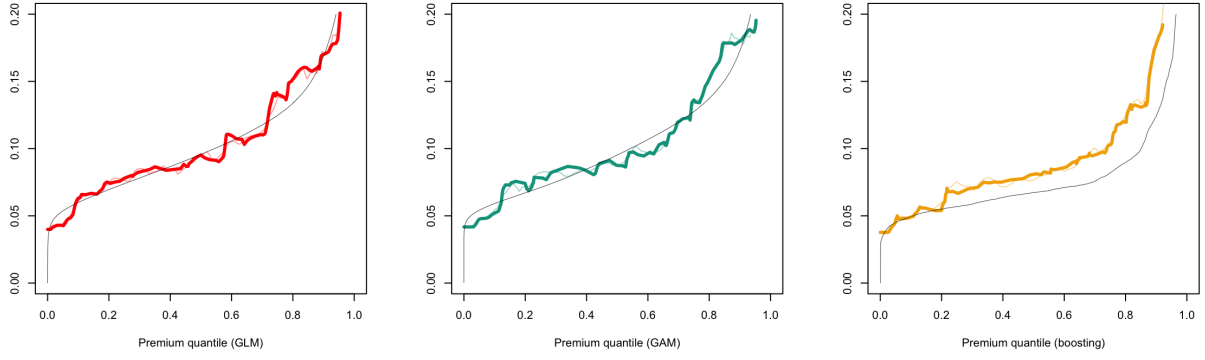


Figure 6.8: Evolution of $s \mapsto E[Y|\hat{\pi}_{BC}(\mathbf{X}) = F_{\hat{\pi}_{BC}}^{-1}(u)]$ (the thin colored line is the former version of $\hat{\pi}$).

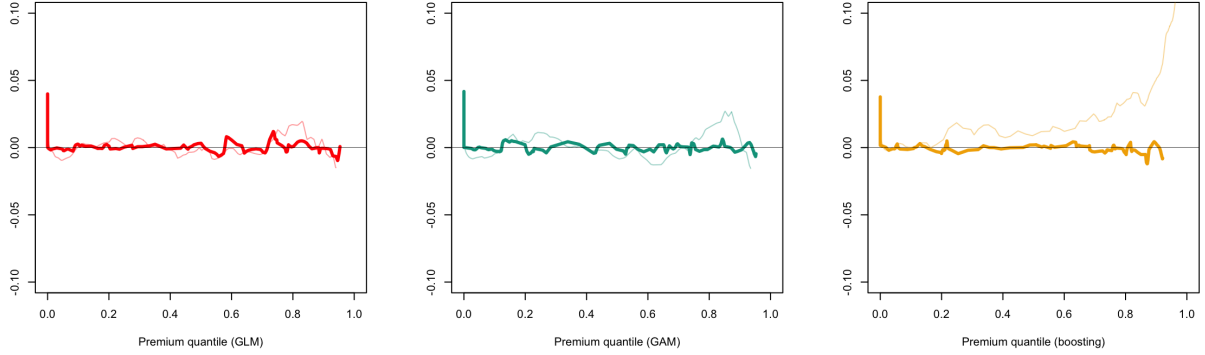


Figure 6.9: Evolution of $u \mapsto E[Y|\hat{\pi}_{BC}(\mathbf{X}) = F_{\hat{\pi}_{BC}}^{-1}(u)] - F_{\hat{\pi}_{BC}}^{-1}(u)$ (the thin colored line is the former version of $\hat{\pi}$).

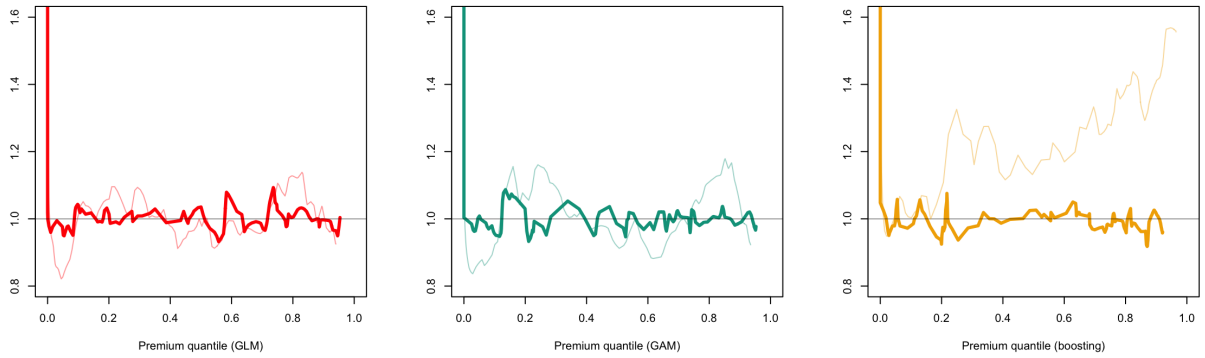


Figure 6.10: Evolution of the multiplicative correction, $u \mapsto E[Y|\hat{\pi}_{BC}(\mathbf{X}) = F_{\hat{\pi}_{BC}}^{-1}(u)]/F_{\hat{\pi}_{BC}}^{-1}(u)$ (the thin colored line is the former version of $\hat{\pi}$).

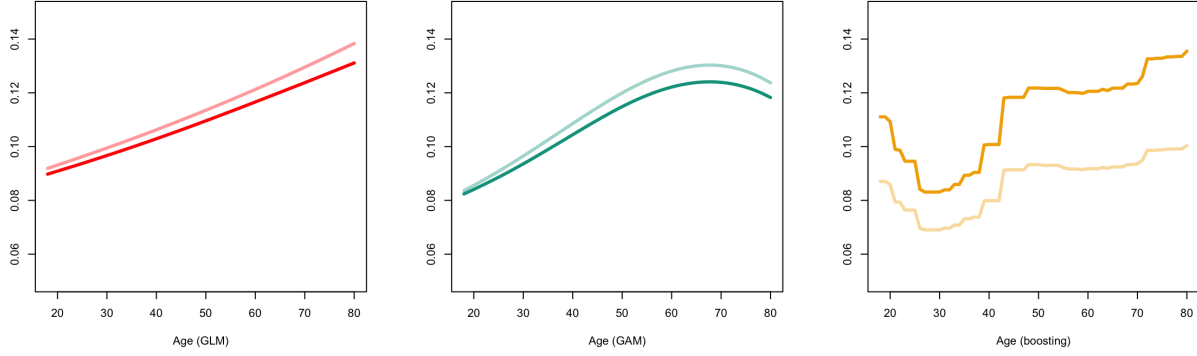


Figure 6.11: Partial Dependence Plot of the age of the driver, **DrvAge**, for the three models, with $\hat{\pi}_{BC}$ (strong line) and $\hat{\pi}$ (thin line).

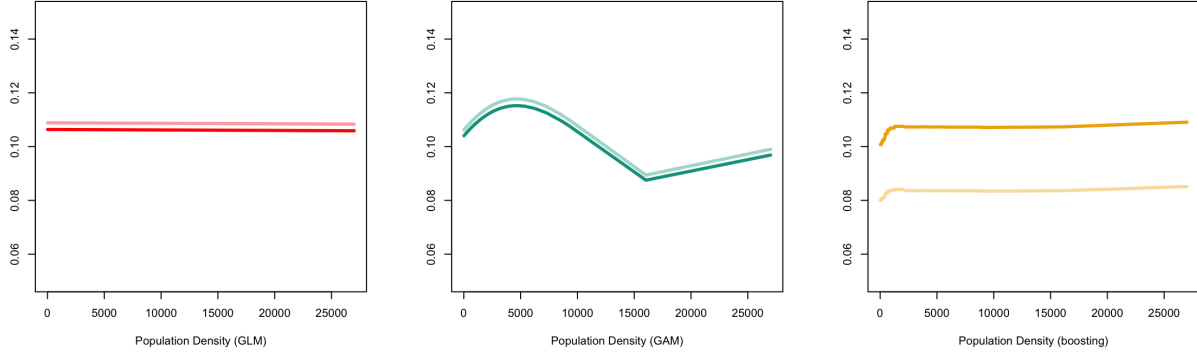


Figure 6.12: Partial Dependence Plot of the population density in the city the driver lives in, **Density**, for the three models, with $\hat{\pi}_{BC}$ (strong line) and $\hat{\pi}$ (thin line).

fraction (the default smoothing parameter is $\alpha = 70\%$. But a second component can be added, $\alpha = (\alpha_0, \alpha_1)$. That second component represents a constant bandwidth, and $h(s)$ will be computed as follows: as previously, $k = \lceil n\alpha_0 \rceil$, and if $d_{(i)}$ represents the ordered statistics of $d_i = |s - s_i|$, $h(s) = \max\{d_{(k)}, \alpha_1\}$. The default value in R is `alpha=c(0.7,0)`.

As we can see on Table 6.1, the boosting algorithm was globally underestimating the average premium. Using a local regression can actually lower the bias. On Figure 6.13, we used 60% of the original dataset to train our three models (training dataset), then the local regression was performed on 20% of the dataset (smoothing dataset) and finally various quantities (bias and empirical loss) were computed on the remaining 20% (validation dataset). On the

left, we can visualize the evolution of the bias $\frac{1}{n} \sum_{i=1}^n (\hat{\pi}_{BC}^{bst}(\mathbf{x}_i) - y_i)$ as a function of α . Here,

smoothing has no impact on the overall bias of GLM and GAM (here the two models have a (small) positive bias on the validation dataset). Smoothing has an impact on the correction of $\hat{\pi}^{bst}$. We can observe that a small α can lead to a much smaller bias (possibly null). Note that we used $\alpha = 5\%$ for previous graphs of this section. In the middle of Figure 6.13,

we can visualize the evolution of the empirical Poisson loss $\frac{1}{n} \sum_{i=1}^n (\hat{\pi}_{\text{BC}}^{\text{bst}}(\mathbf{x}_i) - y_i \ln \hat{\pi}_{\text{BC}}^{\text{bst}}(\mathbf{x}_i))$.

Note that again, a small α leads to a smaller loss. It is also interesting to see that the similar performances are obtained in case of overfitting. Thus, autocalibration also corrects for overfitting by locally averaging the initial predictors.

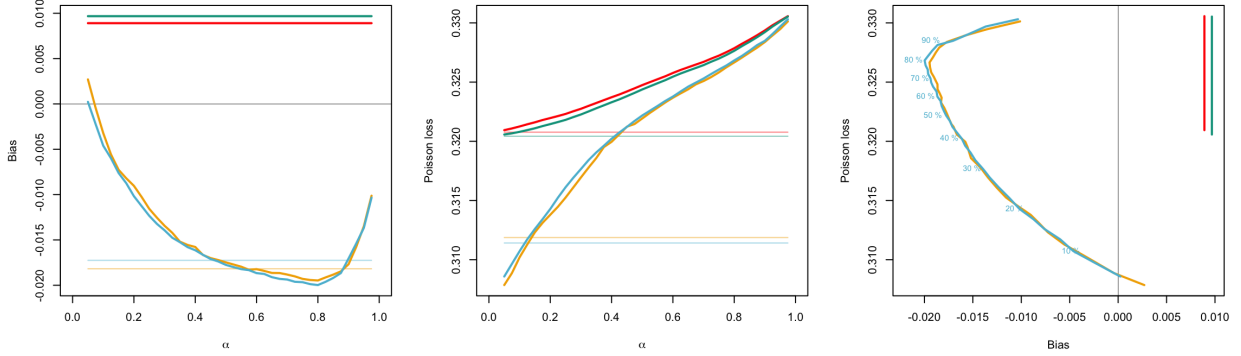


Figure 6.13: Evolution of the bias as a function a α , from 2.5% to 97.5%, on the left, of the Poisson loss in the center, and joint scatterplot of bias and empirical loss on the right. Horizontal light lines correspond to original estimates $\hat{\pi}$ while strong curves are $\hat{\pi}_{\text{BC}}$ for various smoothing parameter α . As previously $\blacksquare$ corresponds to the standard Poisson regression (GLM), $\blacksquare$ corresponds to the smooth additive regression (GAM), while $\blacksquare$ is the boosting model with 30 trees (as previously) and $\blacksquare$ is the boosting model with 1000 trees (overfitting).

To discuss further the impact of α , on Figures 6.14 to 6.16, we have the impact of various values for α (considered as a single value), with a 2% nearest neighbor fraction on Figure 6.14, 20% on Figure 6.15, and 80% on Figure 6.16. Note that $\alpha = 5\%$ was used in the previous section. The various graphs are those described in the previous section, but only to illustrate the construction of the construction of $\hat{\pi}_{\text{BC}}^{\text{gbm}}$ from $\hat{\pi}^{\text{gbm}}$.

On the top of Figures 6.14 to 6.16, we can visualize on the left $u \mapsto \mathbb{E}[Y|\hat{\pi}(\mathbf{X}) = F_{\hat{\pi}}^{-1}(u)]$. Then on the right, we can visualize the evolution of the difference between $\mathbb{E}[Y|\hat{\pi}(\mathbf{X}) = F_{\hat{\pi}}^{-1}(u)]$ and $F_{\hat{\pi}}^{-1}(u)$, when u varies, that is interpreted as the additive local bias on model $\hat{\pi}$, and the multiplicative version is on the right, with $u \mapsto \mathbb{E}[Y|\hat{\pi}(\mathbf{X}) = F_{\hat{\pi}}^{-1}(u)]/F_{\hat{\pi}}^{-1}(u)$, which is the (local) multiplier used to go from $\hat{\pi}$ to $\hat{\pi}_{\text{BC}}$. Then below, we can compare $\hat{\pi}$ and $\hat{\pi}_{\text{BC}}$: with the QQ plot of $\hat{\pi}_{\text{BC}}$ against π on the left – i.e. the plot $u \mapsto \{F_{\hat{\pi}}^{-1}(u), F_{\hat{\pi}_{\text{BC}}}^{-1}(u)\}$ – and in the center, a scatterplot of the pairs $(\hat{\pi}(\mathbf{x}_i), \hat{\pi}_{\text{BC}}(\mathbf{x}_i))$. The empirical distribution of $\hat{\pi}_{\text{BC}}(\mathbf{x}_i)$ is then given on the right, with an histogram. Then, on the third row, we can visualize $s \mapsto \mathbb{E}[Y|\hat{\pi}_{\text{BC}}(\mathbf{X}) = s]$ on the left and $u \mapsto \mathbb{E}[Y|\hat{\pi}_{\text{BC}}(\mathbf{X}) = F_{\hat{\pi}_{\text{BC}}}^{-1}(u)]$ in the center (as previously, the thin line in the background is the equivalent with the original $\hat{\pi}$). Then on the right, we can visualize the evolution of the difference between $\mathbb{E}[Y|\hat{\pi}(\mathbf{X}) = F_{\hat{\pi}}^{-1}(u)]$ and $F_{\hat{\pi}}^{-1}(u)$, the additive bias, while the multiplicative bias is on the left of the last row, on the bottom. If the correction worked well, we should have (almost) horizontal line for those two graphs, $y = 0$ for the additive bias, and $y = 1$ for the multiplicative one. And

finally, we can visualize partial dependence plots, with the partial dependence plot for the age the driver, `DrivAge`, in the center of that last row, and the partial dependence plot for the population density in the city the driver lives in, `Density`, on the right.

6.5 Comparing models

As we have seen when moving from $\hat{\pi}$ to $\hat{\pi}_{\text{BC}}$ the distribution of premiums changed substantially (see Figure 6.1 and 6.7) and the transformation was not (perfectly) monotonic (see Figure 6.6). Nevertheless, the rank correlation between $\hat{\pi}$ and $\hat{\pi}_{\text{BC}}$ is rather high (0.99 for the three models).

Finally, on Figure 6.18 we can visualize the concentration curves of the three models, against $\hat{\pi}_{\text{BC}}^{\text{gam}}$.

7 Discussion

The main message of this paper is as follows. Advanced learning models are able to produce scores that better correlate with the response, as well as with the true premium compared to classical GLMs. This comes from the additional freedom obtained by letting scores to depend in a flexible way of available features, not only linearly. But breaking the overall balance is the price to pay for this higher correlation. Because no constraint on the replication of the observed total, or global balance is imposed, machine learning tools are also able to substantially increase overall bias.

To prevent this to occur, the balance-corrected version of any predictor can be obtained by local GLM, recognizing the nature of the response Y : local Poisson GLM for claim counts, local Gamma GLM for average claim severities and compound Poisson sums with Gamma-distributed terms for claim totals. The canonical link function is adopted so that maximum-likelihood estimates replicate observed totals. An intercept-only GLM is fitted locally, with rectangular weight function, on a set of observations close to the one of interest with proximity assessed with the help of the predictor to be corrected.

In this approach, the supervised learning model is used to produce a real-valued signal $\hat{\pi}(\mathbf{X})$, reducing the high-dimensional feature space to the real line. In fact, the score of the model is enough to compute $\hat{\pi}_{\text{BC}}$ so that the learning model is essentially used to assess proximity among individuals, before performing local averaging. In that respect, the proposed approach shares some similarities with k -NN, or k nearest-neighbors except that here, proximity is assessed with the help of the real-valued $\hat{\pi}$ produced by the learning model. And the proposed extra autocalibration step also counteracts possible overfitting by locally averaging the initial predictor.

The relationship established with consistent loss function and Bregman dominance also appears to be particularly instructive. EU General Data Protection Regulation (GDPR) gives consumers the right to ask for an explanation of an algorithmic decision concerning them. In an insurance context, this means that any provider should be able to explain in an understandable way the reasons for not granting coverage or why the premium charged is so expensive. Model transparency has been studied so far in terms of features interpretability, with the help of importance measures or partial dependence plots mainly. However, the role

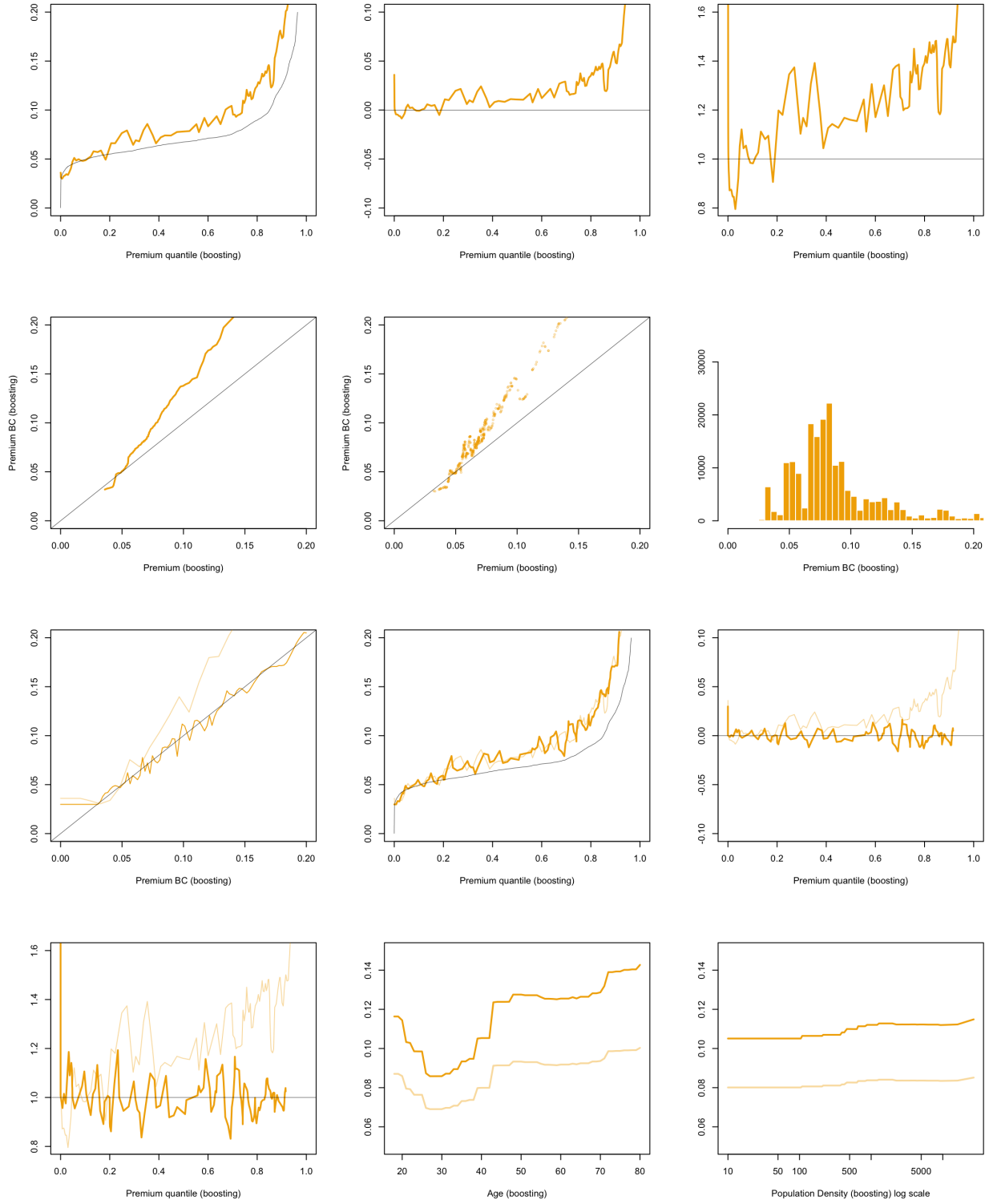


Figure 6.14: Various graphs of $\hat{\pi}^{\text{bst}}$ and $\hat{\pi}_{\text{BC}}^{\text{bst}}$ when $\alpha = 2\%$ for $\hat{\pi}^{\text{gbm}}$

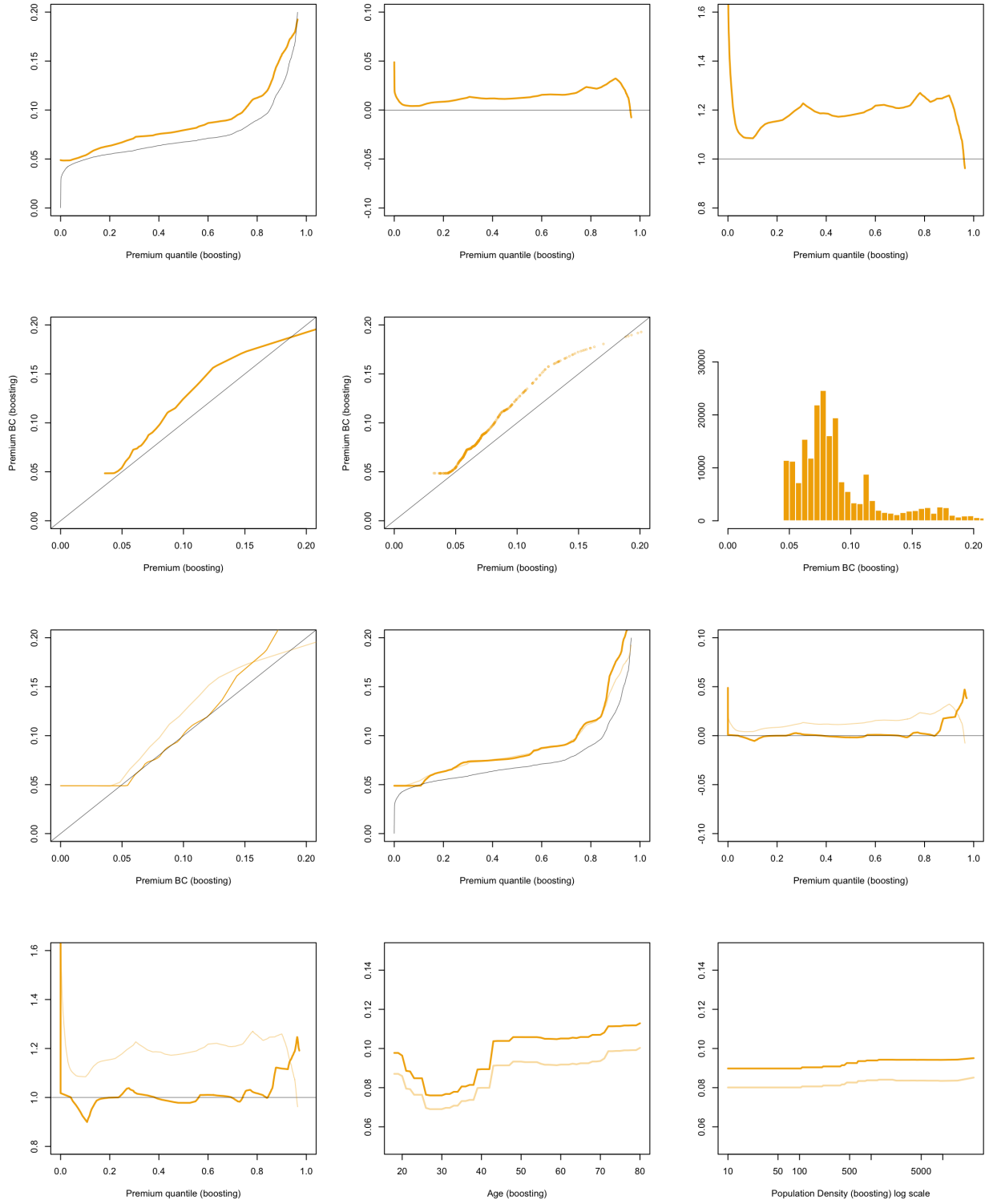


Figure 6.15: Various graphs of $\hat{\pi}^{bst}$ and $\hat{\pi}_{BC}^{bst}$ when $\alpha = 20\%$ for $\hat{\pi}^{gbm}$

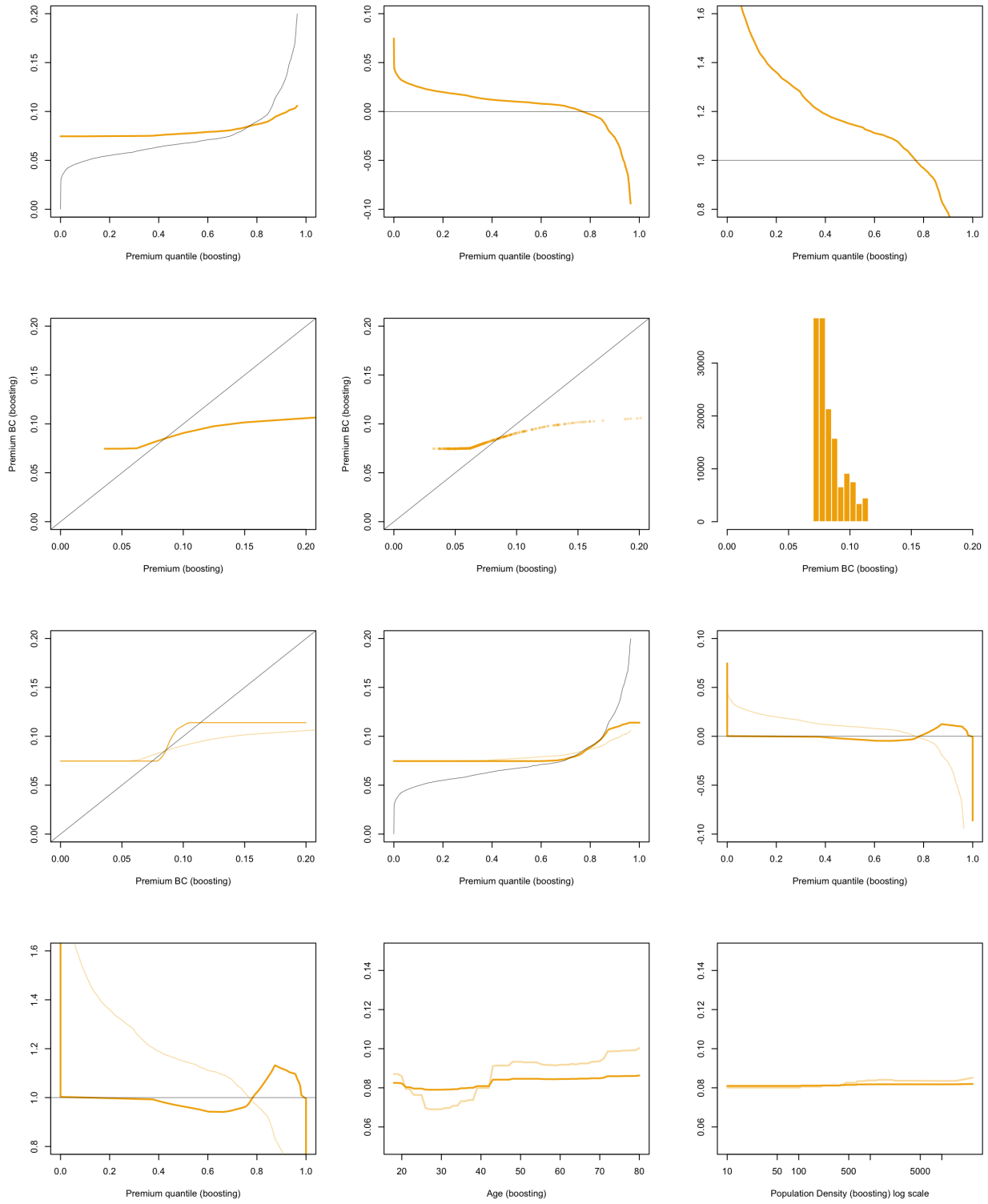


Figure 6.16: Various graphs of $\hat{\pi}^{\text{bst}}$ and $\hat{\pi}_{\text{BC}}^{\text{bst}}$ when $\alpha = 80\%$ for $\hat{\pi}^{\text{gbm}}$

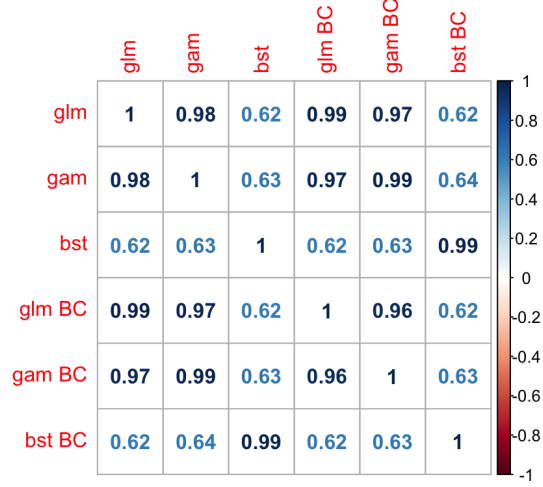


Figure 6.17: Spearman's rank correlation between $\hat{\pi}^{\text{glm}}$, $\hat{\pi}^{\text{gam}}$, $\hat{\pi}^{\text{bst}}$ and their autocalibrated counterparts, $\hat{\pi}_{\text{BC}}^{\text{glm}}$, $\hat{\pi}_{\text{BC}}^{\text{gam}}$ and $\hat{\pi}_{\text{BC}}^{\text{bst}}$, on the validation dataset

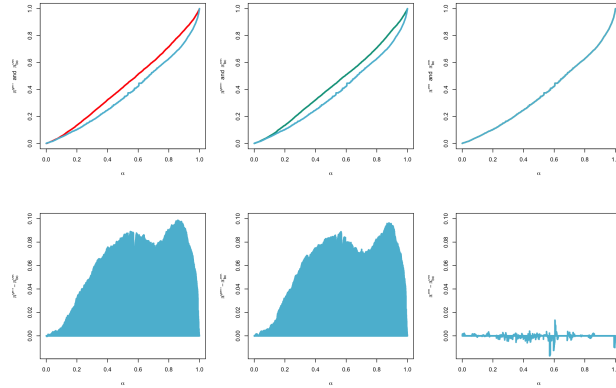


Figure 6.18: Concentration curves, with $\hat{\pi}^{\text{glm}}$, $\hat{\pi}^{\text{gam}}$, $\hat{\pi}^{\text{bst}}$, against $\hat{\pi}_{\text{BC}}^{\text{bst}}$, from the left to the right on top and $\hat{\pi}^{\text{glm}} - \hat{\pi}_{\text{BC}}^{\text{bst}}$, $\hat{\pi}^{\text{gam}} - \hat{\pi}_{\text{BC}}^{\text{bst}}$ and $\hat{\pi}^{\text{bst}} - \hat{\pi}_{\text{BC}}^{\text{bst}}$ from the left to the right, on the bottom.

of the objective function remains largely ignored and we believe that this component is also fundamental for transparency: end-users should be comfortable with the way the model is trained, which obviously impacts on the role granted to features to explain the response of interest and hence the premium. With minimum bias, the insurer can explain that the amount of premium charged has been computed in order to compensate insurance risks within meaningful sub-portfolios, maintaining an overall balance. The proposed balance correction mechanism implements the very same idea, using machine learning tools to define meaningful neighborhoods where collected premiums must match claims to be compensated. This is in accordance with the fundamental mutuality principle at the heart of insurance. But whether

it is transparent to explain policyholders that their premium minimizes Tweedie deviance is certainly debatable. The consistency of Bregman loss functions for pure premiums offers a convenient justification in that respect. By adopting such loss functions, actuaries can guarantee the policyholders, the consumers' organizations and the regulatory authorities that their approach ensures that the machine learning algorithm searches for the correct pure premium.

References

- Bailey, R.A. (1963). Insurance rates with minimum bias. *Proceedings of the Casualty Actuarial Society* 50, 4-11.
- Bailey, R.A., Simon, L.J. (1960). Two studies on automobile insurance ratemaking. *ASTIN Bulletin* 1, 192-217.
- Charpentier, A. (2014). *Computational Actuarial Science with R*. Chapman and Hall / CRC.
- Denuit, M., Dhaene, J., Goovaerts, M.J., Kaas, R. (2005). *Actuarial Theory for Dependent Risks: Measures, Orders and Models*. Wiley, New York.
- Denuit, M., Sznajder, D., Trufin, J. (2019). Model selection based on Lorenz and concentration curves, Gini indices and convex order. *Insurance: Mathematics and Economics* 89, 128-139.
- Efron, B. (1986). How biased is the apparent error rate of a prediction rule? *Journal of the American Statistical Association* 81, 461–470.
- Ehm, W., Gneiting, T., Jordan, A., Kruger, F. (2016) Of quantiles and expectiles: Consistent scoring functions, Choquet representations and forecast rankings. *Journal of the Royal Statistical Society – Series B* 78, 505-562.
- Gneiting, T. (2011). Making and evaluating point forecasts. *Journal of the American Statistical Association* 106, 746–762.
- Kruger, F., Ziegel, J.F. (2020). Generic conditions for forecast dominance. *Journal of Business & Economic Statistics*, in press.
- Loader, C. (1999). *Local Regression and Likelihood*. Springer, New York.
- Meyers, G., Cummings, A. D. (2009). “Goodness of Fit” vs. “Goodness of Lift”. *Actuarial Review* 36, 16-17.
- Mildenhall, S.J. (1999). A systematic relationship between minimum bias and generalized linear models. *Proceedings of the Casualty Actuarial Society* 86, 393-487.
- Savage, L.J. (1971). Elicitation of personal probabilities and expectations. *Journal of the American Statistical Association* 66, 783-810.

- Shaked, M., Shanthikumar, J.G. (2007). Stochastic Orders. Springer, New York.
- Shaked, M., Sordo, M. A., Suarez-Llorens, A. (2012). Global dependence stochastic orders. Methodology and Computing in Applied Probability 14, 617-648.
- Wright, R. (1987). Expectation dependence of random variables, with an application in portfolio theory. Theory and Decision 22, 111-124.
- Wüthrich, M.V. (2019). From Generalized Linear Models to Neural Networks, and back. Available at SSRN 3491790.
- Wüthrich, M.V. (2020). Bias regularization in neural network models for general insurance pricing. European Actuarial Journal 10, 179-202.
- Yitzhaki, S., Schechtman, E. (2013). The Gini Methodology: A Primer on Statistical Methodology. Springer.
- Zumel, N. (2019). An ad-hoc method for calibrating uncalibrated models. Win-Vector Blog, WordPress.