

# COMPARISON OF PREDICTORS' PERFORMANCE IN INSURANCE PRICING: TESTING FOR BREGMAN DOMINANCE BASED ON MURPHY DIAGRAMS

Michel Denuit, Julien Trufin, Thomas Verdebout

REPRINT | 2025 / 15

## **ISBA**

Voie du Roman Pays 20 - L1.04.01

B-1348 Louvain-la-Neuve

Email : [lidam-library@uclouvain.be](mailto:lidam-library@uclouvain.be)

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



# Comparison of predictors' performance in insurance pricing: testing for Bregman dominance based on Murphy diagrams

Michel Denuit<sup>1</sup> · Julien Trufin<sup>2</sup> · Thomas Verdebout<sup>3</sup>

Received: 16 December 2024 / Revised: 11 April 2025 / Accepted: 30 June 2025 /

Published online: 4 August 2025

© The Author(s), under exclusive licence to European Actuarial Journal Association 2025

## Abstract

Ehm et al. (2016) defined forecast dominance, or Bregman dominance as dominance for every Bregman loss function. This letter explores Bregman dominance to compare competing candidate pure premiums. An actuarial interpretation is given for the expectation of extremal consistent loss functions characterizing Bregman dominance, showing their relevance in the context of insurance pricing. An effective testing procedure for Bregman dominance is proposed based on Murphy diagrams and its performance is evaluated through a simulation study. Compared to the existing literature, the proposed testing procedure is easier because it exploits the independence and homogeneity assumptions made in insurance studies and focuses on the conditional mean. An application to a Swiss motor insurance data set demonstrates the potential of the proposed procedure.

**Keywords** Insurance pricing · Model comparison · Bregman dominance · Murphy diagram · Statistical test

## 1 Setting

Given a response  $Y$  (claim number or claim amount) and a set of features  $X_1, \dots, X_p$  gathered in the vector  $X$ , actuarial pricing aims to evaluate the pure premium  $\mu(X) = E[Y|X]$  as accurately as possible. To this end, the function  $\mathbf{x} \mapsto \mu(\mathbf{x}) = E[Y|X = \mathbf{x}]$  is approximated by a (working, or actual) pure premium  $\mathbf{x} \mapsto \pi(\mathbf{x})$  based on an additive or piece-wise constant score. Once fitted on the training data set using an

---

✉ Julien Trufin  
julien.trufin@ulb.ac.be

<sup>1</sup> Institute of Statistics, Biostatistics and Actuarial Science - ISBA, Louvain Institute of Data Analysis and Modeling - LIDAM, UCLouvain, Louvain-la-Neuve, Belgium

<sup>2</sup> Department of Mathematics, Université Libre de Bruxelles (ULB), Bruxelles, Belgium

<sup>3</sup> ECARES and Department of Mathematics, Université Libre de Bruxelles (ULB), Brussels, Belgium

appropriate learning procedure, this produces estimates, or predicted values  $\hat{\pi}(x)$  for  $\mu(x)$ .

The developments in this paper hold the training set as fixed. Precisely,  $\hat{\pi}$  is assumed to have been obtained by minimizing the average empirical loss given some training set and all formulas in the text are meant given this training set. The respective performances of competing models can then be assessed on the basis of a validation set  $\{(Y_i, X_i), i = 1, 2, \dots, n\}$ , that has not been used to obtain  $\hat{\pi}$ . Provided the validation set comprises a large number of independent and identically distributed random vectors  $(Y_i, X_i)$ , this allows us to perform calculations with a generic random vector  $(Y, X)$ , independent of, and distributed as the observations contained in the training set. This assumption is retained throughout the paper.

## 2 Bregman dominance

Patton [9] established that Bregman loss functions are those attaining their minimum at the expected value of the response, subject to weak regularity conditions. The class of Bregman loss functions is then said to be strictly consistent for the (conditional) mean functional. We refer the interested reader to [5] for a general overview of consistency in statistics.

Since pure premiums correspond to conditional expectations, every Bregman loss function is thus eligible to fit  $\hat{\pi}$ . See e.g. Table 1 in [6] as well as [8] for examples of Bregman loss functions. Interestingly, all deviance functions associated to Exponential Dispersion distributions are Bregman loss functions.

We refer the reader to [4] for a general account of model comparison in an insurance context. As pointed out in 2015 by Patton (the reference is now updated to Patton, 2020), rankings of two predictors may depend on the specific consistent loss function used for out-of-sample evaluation. That is why [2, 3, 14, 16] and [1] proposed graphical tools and hypothesis tests to assess the robustness of empirical forecast rankings. In their terminology, one predictor dominates another if it performs better with respect to every consistent loss function. This approach leads to Bregman dominance when considering pure premiums. Stated precisely,  $\hat{\pi}_2$  outperforms  $\hat{\pi}_1$  in terms of Bregman dominance if the inequality  $E[L(Y, \hat{\pi}_2(X))] \leq E[L(Y, \hat{\pi}_1(X))]$  holds true for every Bregman loss function  $L$ . This guarantees that the better performance of  $\hat{\pi}_2$  compared to  $\hat{\pi}_1$  is robust across all Bregman loss functions.

## 3 Murphy diagram

Krüger and Ziegel [6] established that  $\hat{\pi}_2$  outperforms  $\hat{\pi}_1$  in terms of Bregman dominance if, and only if,

$$E[(t - Y)(I[\hat{\pi}_2(X) > t] - I[\hat{\pi}_1(X) > t])] \leq 0 \quad \text{for all } t \in \mathbb{R}, \quad (1)$$

where  $I[\cdot]$  is the indicator function (equal to 1 if the event appearing within the brackets is realized and to 0 otherwise). Note that in applications to insurance pricing,  $\hat{\pi}_k$  is

positive and finite for  $k \in \{1, 2\}$  so that condition (1) is automatically fulfilled for every  $t < 0$  and we can restrict our attention to finite values of  $t \geq 0$ . Therefore, we will assume below that the set on which (1) must hold to get dominance is a compact subset— $\Omega$  say—of the positive real line.

The aim of insurance ratemaking, or actuarial pricing is to evaluate the pure premium as accurately as possible. This means that the target is the conditional expectation  $\mu(X)$  of the response  $Y$ . Thus, the goal is to estimate the unknown  $\mu(X)$  and not to predict response values. This point is crucial for proper understanding of insurance ratemaking: predicting the actual losses  $Y$  is meaningless since the very nature of insurance is to compensate losses within large groups. If actuaries were able to predict the response  $Y$  accurately, insurance would even become useless because every policyholder would be charged his or her actual losses. Instead, actuaries are only interested in the true pure premium  $\mu(X)$  because insurance exploits the fact that individual departures  $Y - \mu(X)$  just cancel each other when they are summed over a sufficiently large portfolio. We refer the reader to [7] for more details about the very aim of insurance ratemaking, which is not to predict the actual losses  $Y$  but to create accurate estimates of  $\mu(X)$ , which is unobserved. This point of view differs from the literature devoted to predictive accuracy and forecast dominance.

Inequality (1) possesses a nice interpretation in terms of performances of the candidate pure premiums  $\hat{\pi}_1$  and  $\hat{\pi}_2$  as it can be rewritten as

$$E[(\mu(X) - t)I[\hat{\pi}_1(X) > t]] \leq E[(\mu(X) - t)I[\hat{\pi}_2(X) > t]] \quad \text{for all } t \in \Omega. \quad (2)$$

Inequality (2), which is equivalent to (1), compares  $\hat{\pi}_1$  and  $\hat{\pi}_2$  with the help of the true pure premium  $\mu(X)$ . Now, (2) explains why the actuary should consider  $\hat{\pi}_2$  as being superior to  $\hat{\pi}_1$  when (1) holds true. Inequality (2) requires that whatever the threshold  $t$ , the average excess of true pure premium  $\mu(X)$  over threshold  $t$  when the candidate premium exceeds this threshold is larger for  $\hat{\pi}_2$  compared to  $\hat{\pi}_1$ . When this holds for every  $t$ , it means that  $\hat{\pi}_2$  better captures the behavior of the true pure premium above any selected threshold  $t$  than  $\hat{\pi}_1$ . The same interpretation holds with shortfalls since (1) can be rewritten as

$$E[(t - \mu(X))I[\hat{\pi}_1(X) \leq t]] \leq E[(t - \mu(X))I[\hat{\pi}_2(X) \leq t]] \quad \text{for all } t \in \Omega.$$

Of course, since the true pure premium  $\mu(\cdot)$  is unknown in applications, the inequalities with  $\mu(X)$  remain conceptual and only serve to provide actuaries with an intuitive meaning for (1) and its relevance in insurance studies. In practice, the inequalities with the observed response  $Y$  replacing the true pure premium  $\mu(X)$  can be estimated from the available data and can be used to assess the performances of the candidate pure premiums  $\hat{\pi}_1$  and  $\hat{\pi}_2$  under consideration.

Subtracting  $E[(\mu(X) - t)_+]$  to both sides of inequality (2), we recover the external, or elementary scores identified by Ehm et al. ([2]), but where  $Y$  is replaced with  $\mu(X)$ . However, the interpretation given by these authors in the context of forecasting does not translate to insurance ratemaking.

Denoting as

$$L_t(\widehat{\pi}_k(X), Y) = \frac{1}{2}(t - Y)\mathbb{I}[\widehat{\pi}_k(X) > t] \quad (3)$$

the extremal consistent loss function of [2] for the mean indexed by the parameter  $t \in \Omega$  (up to a term that does not depend on  $\widehat{\pi}_k(X)$ ),  $k = 1, 2$ , we see that (1) is equivalent to  $\mathbb{E}[L_t(\widehat{\pi}_2(X), Y)] \leq \mathbb{E}[L_t(\widehat{\pi}_1(X), Y)]$  for all  $t \in \Omega$ .

Ehm et al. [2] used the extremal consistent loss functions (3) to graphically compare performances of two predictors. Those graphs are known as Murphy diagrams; see, e.g., [4] for applications to insurance. Coming back to the interpretation of (1) in terms of performances of candidate pure premiums, we see that Murphy diagrams actually assess the accuracy of the pricing structure. The actuary must then favor candidate premiums leading to smaller expected loss (3) because this means that the corresponding premium better capture the behavior of the true, unknown pure premium.

While the Murphy diagram is a useful tool, it only provides graphical evidence of the performance differences but gives no formal statistical justification. [3, 6] and [14] already developed statistical tests for Bregman dominance based on the extremal consistent loss functions of [2]. However, these authors consider forecasts based on time series whereas the present paper considers predictions made from independent data as commonly found in insurance. This allows us to propose a simple testing procedure of Kolmogorov-Smirnov type.

Every Bregman loss function can be recovered as a weighted average of the extremal consistent loss functions (3). Let us illustrate this result with the classical mean squared error (MSE):

$$\begin{aligned} \mathbb{E} \left[ \int_0^\infty L_t(\widehat{\pi}_k(X), Y) dt \right] &= \frac{1}{2} \mathbb{E} \left[ \int_0^{\widehat{\pi}_k(X)} (t - Y) dt \right] \\ &= \frac{1}{4} \mathbb{E}[(\widehat{\pi}_k(X) - Y)^2 - Y^2] \end{aligned} \quad (4)$$

which reduces to the MSE since the second term does not involve the predictor  $\widehat{\pi}_k(X)$ . This shows that MSE is an aggregation of the finer information provided by Murphy diagram.

## 4 Testing for Bregman dominance

Consider two predictors  $\widehat{\pi}_1(X)$  and  $\widehat{\pi}_2(X)$  valued in the sets  $\Omega_1 \subseteq \mathbb{R}$  and  $\Omega_2 \subseteq \mathbb{R}$ , respectively, and let  $\Omega = [\min(\Omega_1 \cup \Omega_2), \max(\Omega_1 \cup \Omega_2)] \subseteq \mathbb{R}$  be a compact interval. From (1), we want to test the null hypothesis

$$\mathcal{H}_0 : \mathbb{E}[(t - Y)(\mathbb{I}[\widehat{\pi}_2(X) > t] - \mathbb{I}[\widehat{\pi}_1(X) > t])] \leq 0 \quad \text{for all } t \in \Omega$$

against the alternative

$$\mathcal{H}_1 : \mathbb{E}[(t - Y) (\mathbb{I}[\widehat{\pi}_2(X) > t] - \mathbb{I}[\widehat{\pi}_1(X) > t])] > 0 \quad \text{for some } t \in \Omega.$$

Let

$$D(t) := \mathbb{E}[(t - Y) (\mathbb{I}[\widehat{\pi}_2(X) > t] - \mathbb{I}[\widehat{\pi}_1(X) > t])], \quad t \in \Omega, \quad \text{and } \overline{D} := \sup_{t \in \Omega} D(t).$$

The null hypothesis  $\mathcal{H}_0$  is then equivalent to  $\overline{D} \leq 0$ . A natural estimator  $D_n(t)$  of  $D(t)$  is

$$D_n(t) = \frac{1}{n} \sum_{i=1}^n (t - Y_i) (\mathbb{I}[\widehat{\pi}_2(X_i) > t] - \mathbb{I}[\widehat{\pi}_1(X_i) > t]).$$

A test statistic of Kolmogorov-Smirnov type can then be defined as

$$T_n = \sup_{t \in \Omega} \sqrt{n} D_n(t).$$

The test will reject the null hypothesis  $\mathcal{H}_0$  at level  $\beta \in (0, 1)$  if the value of  $T_n$  exceeds some critical value  $c_\beta$ . The critical values  $c_\beta$  can be obtained by studying the limiting behavior of the empirical process  $\sqrt{n}(D_n(t) - D(t))$  under the null hypothesis. We have the following result.

**Proposition 4.1** *Assume that the second moment of  $Y$  exists, i.e.  $\mathbb{E}[Y^2] < \infty$ . Then, when  $D(t) = 0$ , which is at the boundary between  $\mathcal{H}_0$  and  $\mathcal{H}_1$ , the empirical process  $\sqrt{n}D_n(t)$  converges weakly to a Gaussian process with mean zero and covariance function*

$$\begin{aligned} \Sigma(t_1, t_2) := & \mathbb{E} \left[ (t_1 - Y) (\mathbb{I}[\widehat{\pi}_2(X) > t_1] - \mathbb{I}[\widehat{\pi}_1(X) > t_1]) \right. \\ & \left. (t_2 - Y) (\mathbb{I}[\widehat{\pi}_2(X) > t_2] - \mathbb{I}[\widehat{\pi}_1(X) > t_2]) \right]. \end{aligned} \quad (5)$$

The proof of Proposition 4.1 is given in Appendix A. From Proposition 4.1 and the continuous mapping theorem, it follows that a Kolmogorov-Smirnov type test can be derived by rejecting the null hypothesis  $\mathcal{H}_0$  at the asymptotic level  $\beta \in (0, 1)$  when

$$T_n = \sup_{t \in \Omega} \sqrt{n} D_n(t) > c_\beta, \quad (6)$$

where the critical value  $c_\beta$  can, in principle, be obtained from Proposition 4.1. Although Proposition 4.1 shows that  $T_n$  in (6) is a natural test statistic for the problem considered, computing the critical value  $c_\beta$  is challenging as it requires the computation of  $\Sigma(t_1, t_2)$ , which is impractical due to the unknown distribution of  $(Y, \widehat{\pi}_1(X), \widehat{\pi}_2(X))$ . Hence, we propose in Appendix B a Monte Carlo procedure to perform the test, following [15]. A simulation exercise is performed in Appendix C, demonstrating that the procedure described in Appendix B achieves the correct nominal  $\beta$ -level constraint and maintains power under the alternative hypothesis.

## 5 Case study

### 5.1 Data set

We consider the Swiss motor insurance data set used in Wüthrich [12] and in Wüthrich and Buser [13] consisting of 500 000 insurance policies. This is a simulated data set which as often been used in the literature. In the present setting, it has the advantage to contain the true pure premium  $\mu(\mathbf{X})$  so that we can illustrate the intuitive meaning of Murphy diagram in terms of the assessment of the superiority of some candidate pure premium. Of course, the analysis can be conducted without the knowledge of  $\mu(\mathbf{X})$ , as it is the case with a real data set.

For each policy  $i$  ( $i = 1, \dots, 500\,000$ ), the data set records the numbers of claims  $Y_i$  filed by the policyholder  $i$ , the corresponding exposure-to-risk  $e_i \leq 1$  and the features  $\mathbf{X}_i = (X_{i1}, \dots, X_{i8})$  where  $X_{i1}$  is the age of the driver (age),  $X_{i2}$  is the age of the car (ac),  $X_{i3}$  is the power of the car (power),  $X_{i4}$  is the fuel type of the car (gas),  $X_{i5}$  is the vehicle brand (brand),  $X_{i6}$  is the area code of the living place of the driver (area),  $X_{i7}$  is the population density at the living place of the driver (dens), and  $X_{i8}$  is the Swiss canton of the license plate of the car (ct).

This data set is unique because the values  $e_i \mu(\mathbf{X}_i)$  of the true model are also provided. In fact, the number of claims  $Y_i$  has been generated from a Poisson distribution with mean  $e_i \mu(\mathbf{X}_i)$ . We partition the dataset into a training set  $\mathcal{D}$  and a validation set  $\bar{\mathcal{D}}$ . The training set  $\mathcal{D}$  comprises 80% of the observations from the entire dataset, while the validation set  $\bar{\mathcal{D}}$  contains the remaining 20% of observations.

### 5.2 Models under consideration

The observations  $(Y_i, \mathbf{X}_i)$  are assumed to be independent. Given  $\mathbf{X}_i = \mathbf{x}_i$  and the exposure-to-risk  $e_i$ , the response  $Y_i$  is assumed to follow a Poisson distribution with mean  $e_i \mu(\mathbf{x}_i)$ . Thus,  $\mu(\mathbf{x}_i)$  represents the expected annual claim frequency for policyholder  $i$ . Our objective is to estimate the function  $\mathbf{x} \mapsto \mu(\mathbf{x})$  using the observations  $(Y_i, \mathbf{X}_i)$  in the training set  $\mathcal{D}$ .

We fit generalized additive models (GAMs) on  $\mathcal{D}$  with Poisson deviance loss and log-link function using the R package `gam`. More precisely, we fit two GAMs producing the following predictors:

- $\hat{\pi}^{\text{GAM1}}$ , using one feature, namely  $X_1$  (age);
- $\hat{\pi}^{\text{GAM2}}$ , using two features, namely  $X_1$  (age) and  $X_6$  (area).

In both models, the effect of the covariate age is captured by splines, and we do not consider interaction terms.

We also train gradient boosting trees (GBTs) on  $\mathcal{D}$  with Poisson deviance loss and log-link function with the help of the R package `gbm`. The bagging fraction  $\gamma$  is set at 0.5. The shrinkage parameter  $\kappa$  is set at the low value 0.01. The size of the trees is controlled by the interaction depth `ID`. We fit two GBTs on 80% of the observations in  $\mathcal{D}$  and the optimal values for the number of trees  $T$  are determined by minimizing the out-of-sample Poisson deviance loss computed on the remaining 20% of observations in  $\mathcal{D}$ , producing the following estimators:



**Table 1** Out-of-sample estimates (on  $\overline{\mathcal{D}}$ ) of the generalization error

|                           |                        |
|---------------------------|------------------------|
| $\mu$                     | $279.23 \cdot 10^{-3}$ |
| $\hat{\pi}^{\text{GAM1}}$ | $313.26 \cdot 10^{-3}$ |
| $\hat{\pi}^{\text{GAM2}}$ | $312.74 \cdot 10^{-3}$ |
| $\hat{\pi}^{\text{GBT1}}$ | $291.74 \cdot 10^{-3}$ |
| $\hat{\pi}^{\text{GBT2}}$ | $279.75 \cdot 10^{-3}$ |

**Table 2** Values of  $\hat{p}$  for  $B = 500$ . The null hypothesis  $\mathcal{H}_0$ :  $E[(t - Y)(I[\hat{\pi}_2(X) > t] - I[\hat{\pi}_1(X) > t])] \leq 0$  for all  $t \in \Omega$  is rejected at the level 0.05 when  $\hat{p} < 0.05$ 

|                             | $\mu$ | $\hat{\pi}^{\text{GAM1}}$ | $\hat{\pi}_2^{\text{GAM2}}$ | $\hat{\pi}^{\text{GBT1}}$ | $\hat{\pi}^{\text{GBT2}}$ |
|-----------------------------|-------|---------------------------|-----------------------------|---------------------------|---------------------------|
| $\mu$                       | /     | 0.000                     | 0.000                       | 0.000                     | 0.272                     |
| $\hat{\pi}^{\text{GAM1}}$   | 1.000 | /                         | 0.000                       | 0.908                     | 0.960                     |
| $\hat{\pi}_1^{\text{GAM2}}$ | 1.000 | 0.000                     | /                           | 0.924                     | 0.962                     |
| $\hat{\pi}^{\text{GBT1}}$   | 1.000 | 0.000                     | 0.000                       | /                         | 0.976                     |
| $\hat{\pi}^{\text{GBT2}}$   | 0.980 | 0.000                     | 0.000                       | 0.000                     | /                         |

- $\hat{\pi}^{\text{GBT1}}$ , using features  $X_1$  (age) and  $X_6$  (area) and  $\text{ID} = 2$ ;
- $\hat{\pi}^{\text{GBT2}}$ , using all 8 available features and  $\text{ID} = 2$ .

We expect  $\hat{\pi}^{\text{GBT1}}$  to produce better results than with the GAMs because we allow for interaction effects between feature components age and area. Moreover,  $\hat{\pi}^{\text{GBT2}}$  is expected to perform better than the others since it uses all features, unlike the other models considered.

In Table 1, we compute out-of-sample estimates (on  $\overline{\mathcal{D}}$ ) of the generalization errors (with the Poisson loss function) for the four models under consideration as well as for the true model. As expected, we notice that  $\hat{\pi}^{\text{GAM2}}$  slightly outperforms  $\hat{\pi}^{\text{GAM1}}$  and that  $\hat{\pi}^{\text{GBT1}}$  outperforms  $\hat{\pi}^{\text{GAM2}}$ . Moreover,  $\hat{\pi}^{\text{GBT2}}$  has the lowest error among the four models considered. It is interesting to notice that the error of  $\hat{\pi}^{\text{GBT2}}$  is very closed to the one of the true model  $\mu$ , so that  $\hat{\pi}^{\text{GBT2}}$  seems to be a good candidate to approximate  $\mu$ .

### 5.3 Testing for Bregman dominance

Table 2 contains the values of  $\hat{p}$  of our testing procedure computed on the validation set  $\overline{\mathcal{D}}$  with  $B = 500$ . We observe that the null hypothesis  $\mathcal{H}_0$  is always rejected when  $\hat{\pi}_1 = \mu$ , as expected, except for  $\hat{\pi}_2 = \hat{\pi}^{\text{GBT2}}$ . Similarly,  $\mathcal{H}_0$  is always rejected when  $\hat{\pi}_1 = \hat{\pi}^{\text{GBT2}}$  except for  $\hat{\pi}_2 = \mu$ , as expected. Hence, for  $\hat{\pi}^{\text{GBT2}}$  and  $\mu$ , our test cannot determine if either of the two models outperforms the other in terms of Bregman dominance. Moreover, for  $\hat{\pi}_1 = \hat{\pi}^{\text{GAM1}}$  (resp.  $\hat{\pi}_1 = \hat{\pi}^{\text{GAM2}}$ ), we do not reject  $\mathcal{H}_0$ , except for  $\hat{\pi}_2 = \hat{\pi}^{\text{GAM2}}$  (resp.  $\hat{\pi}_2 = \hat{\pi}^{\text{GAM1}}$ ). Therefore, for  $\hat{\pi}^{\text{GAM1}}$  and  $\hat{\pi}^{\text{GAM2}}$ , one can say that neither model outperforms the other in terms of Bregman dominance. Finally, for  $\hat{\pi}_1 = \hat{\pi}^{\text{GBT1}}$ ,  $\mathcal{H}_0$  is rejected for  $\hat{\pi}_2 = \hat{\pi}^{\text{GAM1}}$  and  $\hat{\pi}_2 = \hat{\pi}^{\text{GAM2}}$  and is not rejected for  $\hat{\pi}_2 = \hat{\pi}^{\text{GBT2}}$  and  $\hat{\pi}_2 = \mu$ , which is not particularly surprising here.

Figure 1 displays Murphy diagrams for the extremal consistent loss functions  $L_t(\hat{\pi}(X), Y)$  defined in (3). More precisely, it shows the estimates of  $E[L_t(\hat{\pi}(X), Y)]$  on the validation set  $\overline{D}$  for the models  $\hat{\pi}$  under consideration as well as for the true model  $\mu$ . It is interesting to notice that Murphy diagrams almost coincide for  $\hat{\pi} = \hat{\pi}^{\text{GBT2}}$  and  $\hat{\pi} = \mu$ , which confirms the difficulty of our test in distinguishing between these two models. Also, one sees that Murphy diagrams intersect for  $\hat{\pi} = \hat{\pi}^{\text{GAM1}}$  and  $\hat{\pi} = \hat{\pi}^{\text{GAM2}}$ , which suggests that neither model dominates the other, as confirmed by our testing procedure. Notice that the difference in MSE can be obtained (up to a constant factor) as the area between the curves in Figure 1 as shown by (4).

## Appendix

### A Proof of Proposition 4.1

Let

$$g_n(t) := \sqrt{n}D_n(t) = \frac{1}{\sqrt{n}} \sum_{i=1}^n (t - Y_i) (\mathbb{I}[\hat{\pi}_2(X_i) > t] - \mathbb{I}[\hat{\pi}_1(X_i) > t]).$$

For  $t_1 < t_2$ , since  $D(t_1) = D(t_2) = 0$  and the random vectors  $\{(Y_i, X_i), i = 1, \dots, n\}$  are i.i.d., it is easy to see that

$$E[g_n(t_1)g_n(t_2)] = E\left[(t_1 - Y) (\mathbb{I}[\hat{\pi}_2(X) > t_1] - \mathbb{I}[\hat{\pi}_1(X) > t_1])\right. \\ \left. (t_2 - Y) (\mathbb{I}[\hat{\pi}_2(X) > t_2] - \mathbb{I}[\hat{\pi}_1(X) > t_2])\right].$$

Note that

$$g_n(t) = t\sqrt{n}(F_2^{(n)}(t) - F_1^{(n)}(t)) - \frac{1}{\sqrt{n}} \sum_{i=1}^n Y_i (\mathbb{I}[\hat{\pi}_2(X_i) > t] - \mathbb{I}[\hat{\pi}_1(X_i) > t])$$

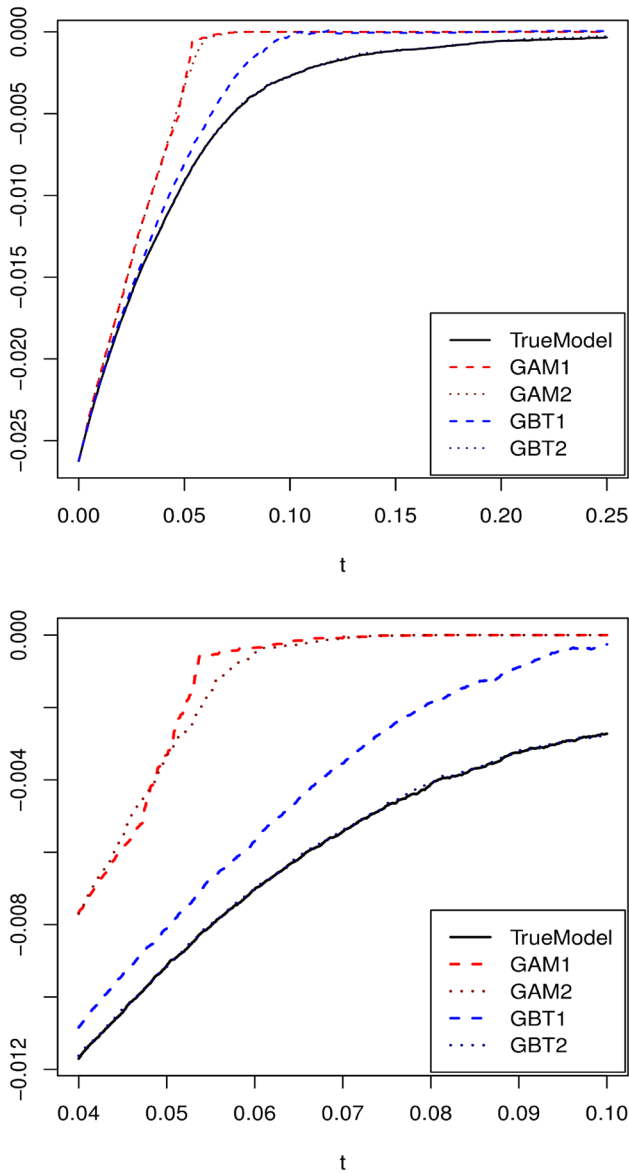
where

$$F_j^{(n)}(t) := n^{-1} \sum_{i=1}^n \mathbb{I}[\hat{\pi}_j(X_i) > t], \quad j = 1, 2.$$

Since  $t \rightarrow \frac{1}{n} \sum_{i=1}^n Y_i \mathbb{I}[\hat{\pi}_j(X_i) > t]$  is monotone, it is a VC-class of functions and it is therefore Donsker. Moreover, since  $\Omega$  is compact, there exists a  $C > 0$  such that

$$|tF_j^{(n)}(t)| \leq CF_j^{(n)}(t).$$

As the envelope of  $t \rightarrow tF_j^{(n)}(t)$  is square integrable, Theorem 19.14 of [11] directly entails that it is Donsker. Now, since the sum of Donsker classes is Donsker (see for instance [10]), the result follows.



**Fig. 1** Murphy diagrams for  $\hat{\pi}^{\text{GAM1}}$  and  $\hat{\pi}^{\text{GAM2}}$  in red,  $\hat{\pi}^{\text{GBT1}}$  and  $\hat{\pi}^{\text{GBT2}}$  in blue, and for  $\mu$  in black

## B Monte Carlo testing procedure

1. Generate  $B$  independent samples  $U_1, \dots, U_B$ , where the  $b$ th sample  $U_b = (U_{b1}, \dots, U_{bn})$  contains  $n$  independent standard Gaussian random variables.

## 2. Compute

$$\hat{p} := \frac{1}{B} \sum_{b=1}^B \mathbb{I} \left[ \max_{t \in \Omega} \Delta(t, D_n(t), U_b) > \max_{t \in \Omega} \sqrt{n} D_n(t) \right],$$

where

$$\Delta(t, D_n(t), U_b) := \frac{1}{\sqrt{n}} \sum_{i=1}^n \left( (t - Y_i) (\mathbb{I}[\hat{\pi}_2(X_i) > t] - \mathbb{I}[\hat{\pi}_1(X_i) > t]) - D_n(t) \right) U_{bi}.$$

## 3. Reject the null hypothesis at the level $\beta$ when $\hat{p} < \beta$ .

## C Simulated example

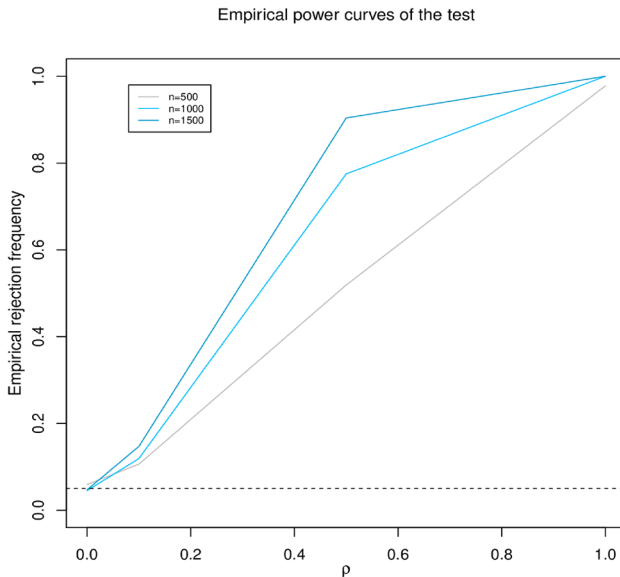
In this simulation exercise,  $\mu(X)$ ,  $\hat{\pi}_1(X)$  and  $\hat{\pi}_2(X)$  are assumed to be uniformly distributed over  $[a, b]$ , with  $b > a > 0$ ,  $\mu(X)$  and  $\hat{\pi}_2(X)$  are supposed to be mutually independent and  $\mu(X)$  and  $\hat{\pi}_1(X)$  are made dependent through the use of a Gaussian copula with a correlation coefficient  $\rho \geq 0$ . Therefore, since  $\mathbb{E}[(t - Y)\mathbb{I}[\hat{\pi}_k(X) > t]] = \mathbb{E}[(t - \mu(X))\mathbb{I}[\hat{\pi}_k(X) > t]]$  for  $k = 1, 2$ , we have

$$\begin{aligned} \mathbb{E}[(t - Y)\mathbb{I}[\hat{\pi}_2(X) > t]] &\leq \mathbb{E}[(t - Y)\mathbb{I}[\hat{\pi}_1(X) > t]] \quad \text{for all } t \in [a, b] \\ \Leftrightarrow \mathbb{E}[\mu(X)|\hat{\pi}_2(X) > t] &\geq \mathbb{E}[\mu(X)|\hat{\pi}_1(X) > t] \quad \text{for all } t \in [a, b] \\ \Leftrightarrow \mathbb{E}[\mu(X)] &\geq \mathbb{E}[\mu(X)|\hat{\pi}_1(X) > t] \quad \text{for all } t \in [a, b]. \end{aligned}$$

When  $\rho = 0$ , we get  $\mathbb{E}[\mu(X)|\hat{\pi}_1(X) > t] = \mathbb{E}[\mu(X)]$ , so that this case coincides with the boundary of the null hypothesis. Increasing the value of  $\rho$  increases  $\mathbb{E}[\mu(X)|\hat{\pi}_1(X) > t]$ , so that the larger  $\rho$ , the further the data generating process is from the null hypothesis.

For this simulation exercise, we assume that  $Y$  given  $X$  is Poisson distributed with mean  $\mu(X)$ , which is a typical case in insurance when the response represents the number of claims made by a policyholder over one year for instance, and we suppose that  $a = 0.05$  and  $b = 0.5$ . We consider  $\rho = 0, 0.1, 0.5, 1$ , and for each value of  $\rho$ , we generate  $R = 1000$  independent samples  $(Y_1, \hat{\pi}_1(X_1), \hat{\pi}_2(X_1)), \dots, (Y_n, \hat{\pi}_1(X_n), \hat{\pi}_2(X_n))$ .

Figure 2 displays empirical rejection frequencies of the test performed using steps 1-3 of the Monte-Carlo procedure described in the previous section with  $B = 500$  at the nominal level  $\beta = 5\%$  for sample sizes  $n = 500$ ,  $n = 1000$  and  $n = 1500$ . The test satisfies the level constraint since the rejection frequencies for  $\rho = 0$  are close



**Fig. 2** Empirical power curves of the test for various sample sizes

to the nominal level 5%. Moreover, the testing procedure shows higher power as  $n$  increases.

**Acknowledgements** The authors would like to express their gratitude to the anonymous Referee and Handling Editor whose comments have been extremely useful to revise a previous version of the present work. Julien Trufin gratefully acknowledges the support received from the Research Chair ACTIONS under the aegis of the Risk Foundation, an initiative by BNP Paribas Cardif and the French Institute of Actuaries.

## References

1. Barendse S, Patton AJ (2021) Comparing predictive accuracy in the presence of a loss function shape parameter. *Journal of Business & Economic Statistics* 40:1057–1069
2. 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
3. Ehm W, Kruger F (2018) Forecast dominance testing via sign randomization. *Electronic Journal of Statistics* 12:3758–3793
4. Fissler T, Lorentzen C, Mayer M (2022) Model comparison and calibration assessment: User guide for consistent scoring functions in machine learning and actuarial practice. *arXiv preprint arXiv:2202.12780*
5. Gneiting T (2011) Making and evaluating point forecasts. *J Am Stat Assoc* 106:746–762
6. Krüger F, Ziegel JF (2021) Generic conditions for forecast dominance. *Journal of Business & Economic Statistics* 39:972–983
7. Meyers G, Cummings AD (2009) “Goodness of Fit” vs. “Goodness of Lift”. *Actuarial Review* 36:16–17
8. Patton AJ (2020) Comparing possibly misspecified forecasts. *Journal of Business & Economic Statistics* 38:796–809
9. Savage LJ (1971) Elicitation of personal probabilities and expectations. *J Am Stat Assoc* 66:783–810
10. van der Vaart AW (1996) New Donsker classes. *Ann Probab* 24:2128–2140
11. van der Vaart AW (2000) *Asymptotic Statistics*. Cambridge University Press

12. Wüthrich MV (2020) Bias regularization in neural network models for general insurance pricing. *Eur Actuar J* 10:179–202
13. Wüthrich, M.V., Buser, C. (2023). Data analytics for non-life insurance pricing. SSRN Manuscript ID 2870308, Version of June 19, 2023
14. Yen T-J, Yen Y-M (2021) Testing forecast accuracy of expectiles and quantiles with the extremal consistent loss functions. *Int J Forecast* 37:733–758
15. Zhu X, Guo X, Lin L, Zhu L (2016) Testing for positive expectation dependence. *Ann Inst Stat Math* 68:135–153
16. Ziegel JF, Krüger F, Jordan A, Fasciati F (2020) Robust forecast evaluation of expected shortfall. *J Financ Economet* 18:95–120

**Publisher's Note** Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.