

# ANOTHER LOOK AT THE ZERO INTEGRAL DIFFERENCE BETWEEN LORENZ AND CONCENTRATION CURVES IN SUPERVISED LEARNING

Michel Denuit, Julien Trufin

REPRINT | 2026 / 24

## **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>

# ANOTHER LOOK AT THE ZERO INTEGRAL DIFFERENCE BETWEEN LORENZ AND CONCENTRATION CURVES IN SUPERVISED LEARNING

Michel Denuit

Institute of Statistics, Biostatistics and Actuarial Science  
LIDAM, UCLouvain  
Louvain-la-Neuve, Belgium

Julien Trufin

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

December 11, 2025

## Abstract

We revisit the area between the concentration and Lorenz curves (ABC) criterion for model assessment in supervised learning. Insurance pricing is considered throughout the paper to illustrate the concepts but the results apply to any other setting where the mean of a response must be estimated from data. Building on the characterization of these curves, we provide new equivalent formulations for the case where the ABC vanishes. First, we characterize a vanishing ABC as the absence of correlation between pricing error and the ranks induced by the candidate premiums, making the link with Gini and Co-Gini coefficients. In both the discrete and continuous cases, we then show that a vanishing ABC corresponds to global balance in a modified portfolio that overweights lower premium classes. These results complement existing work on auto-calibration and contribute to a better understanding of ABC as a diagnostic tool in insurance pricing and related applications.

**Keywords:** Insurance pricing, Lorenz curve, concentration curve, area between the curves.

# 1 Introduction and motivation

Consider a response  $Y$  of interest. It can be a discrete (e.g. count) or continuous random variable. To explain the mean response, the analyst has at his or her disposal a set of features  $X_1, \dots, X_p \in \mathcal{X} \subseteq \mathbb{R}^p$  gathered in the vector  $\mathbf{X}$ . The aim is to evaluate the expected response as accurately as possible, so that the target is the conditional expectation  $\mu(\mathbf{X}) = \mathbb{E}[Y|\mathbf{X}]$  of the response  $Y$  given the available information  $\mathbf{X}$ . Here,  $(Y, \mathbf{X})$  is a generic observation, taken at random from the population of interest.

The function  $\mathbf{x} \mapsto \mu(\mathbf{x}) = \mathbb{E}[Y|\mathbf{X} = \mathbf{x}]$  is unknown to the analyst, and may exhibit a complex behavior in  $\mathbf{x}$ . This is why this function is approximated by a function  $\mathbf{x} \mapsto m(\mathbf{x})$  with a simpler structure. For instance,  $m(\mathbf{X})$  may be a monotonic function of a linear combination of the components of  $\mathbf{X}$  with Generalized Linear Models (GLMs). Typically,  $m(\mathbf{x}) = \exp(\boldsymbol{\beta}^\top \mathbf{x})$ . The linear combination  $\boldsymbol{\beta}^\top \mathbf{x}$  is generally referred to as the score in insurance. Once fitted on the training data, this produces estimates  $\hat{m}(\mathbf{x})$  for  $\mu(\mathbf{x})$ , or fitted values. With GLMs, we simply have  $\hat{m}(\mathbf{x}) = \exp(\hat{\boldsymbol{\beta}}^\top \mathbf{x})$  where  $\hat{\boldsymbol{\beta}}$  is the maximum-likelihood estimate of  $\boldsymbol{\beta}$ .

In general,  $\hat{m}(\cdot)$  can also be obtained with machine learning algorithms like random forests or neural networks, for instance. The performance of  $\hat{m}(\cdot)$  that has been obtained by minimizing the empirical loss on some training set can then be assessed on the basis of an independent random sample  $\{(Y_i, \mathbf{X}_i), i = 1, 2, \dots, n\}$  following the same law as  $(Y, \mathbf{X})$ . This sample is called validation set. All formulas in the following are meant given the training set, which means that the function  $\mathbf{x} \mapsto \hat{m}(\mathbf{x})$  is considered as known. In the case of GLMs,  $\hat{\boldsymbol{\beta}}$  is frozen at its estimation on the training set and treated as a known constant when  $\hat{m}(\cdot)$  is applied to the validation set.

Provided the validation set comprises a large number of independent and identically distributed pairs  $(Y_i, \mathbf{X}_i)$ , this allows us to perform calculations with a generic random vector  $(Y, \mathbf{X})$  distributed as the observations contained in the training set. This approach is meaningful in applications where the analyst is in a data-rich situation.

As an illustration, we consider the problem of setting pure premiums in general insurance. In that context,  $Y$  represents insurance losses (number or cost of claims, for instance) and  $\mu(\mathbf{X})$  is the true pure premium that should be charged to a policyholder with individual characteristics summarized in the vector  $\mathbf{X}$  while  $\hat{m}(\mathbf{X})$  is the actual pure premium charged by the insurer as an approximation of  $\mu(\mathbf{X})$ . By pure premium, actuaries mean the expected amount of benefits claimed by a given policyholder. Commercial premiums charged to customers supplement pure premiums with a safety loading (corresponding to the cost of capital ensuring solvency) as well as running expenses, intermediaries' commissions and taxes.

Throughout the paper, we assume that  $\hat{m}(\cdot)$  is unbiased in the sense

$$\mathbb{E}[\hat{m}(\mathbf{X})] = \mathbb{E}[\mu(\mathbf{X})] = \mathbb{E}[Y] = \mu \quad (1.1)$$

where  $\mu$  denotes the unconditional mean response. Assumption (1.1) is referred to as global balance in insurance. This is a minimal requirement as it requires that the premiums collected from policyholders balance claim amounts at portfolio level, on average. The pricing error arises when the insurer charges  $\hat{m}(\mathbf{X})$  instead of  $\mu(\mathbf{X})$ . Global balance (1.1) ensures that

the average pricing error is zero. At portfolio level, replacing  $\mu(\mathbf{X})$  with  $\widehat{m}(\mathbf{X})$  has no effect, on average, but this induces possible undesirable premium transfers between policyholders depending on their risk profile  $\mathbf{X}$ .

Lorenz and concentration curves are widely used as diagnostic tools in supervised learning, in particular in insurance pricing and risk modeling. For a detailed discussion of their properties in the context of insurance, we refer the reader to Denuit et al. (2019) and the references therein. For an exhaustive review of the properties of concentration and Lorenz curves, as well as related notions like Gini dispersion and dependence measures, we refer the interested reader to the book by Yitzhaki and Schechtman (2013). The area between the two curves (ABC), introduced in Denuit et al. (2019), has recently emerged as a quantitative diagnostic, which vanishes in the special case of unbiased and auto-calibrated predictors (see, e.g., Denuit et al., 2024). The purpose of this paper is to investigate more thoroughly the interpretation of a vanishing ABC for unbiased predictors, considering both discrete and continuous cases.

The remainder of the paper is organized as follows. In Section 2, we recall the definition of the Lorenz and concentration curves, as well as of the ABC criterion. Section 3 gives a first characterization of a vanishing ABC in terms of pricing error. Section 4 shows that this condition can also be interpreted as global balance (1.1) within a modified portfolio where risk profiles are appropriately re-weighted.

## 2 The concentration curve, the Lorenz curve, and the area between them

### 2.1 Concentration curve

The concentration curve of  $\mu(\mathbf{X})$  with respect to  $\widehat{m}(\cdot)$  based on the information contained in the vector  $\mathbf{X}$  is defined as the function

$$\alpha \mapsto \text{CC}[\mu(\mathbf{X}), \widehat{m}(\mathbf{X}); \alpha] = \frac{1}{\mu} \mathbb{E}[\mu(\mathbf{X}) \mathbb{I}[\widehat{m}(\mathbf{X}) \leq F_{\widehat{m}(\mathbf{X})}^{-1}(\alpha)]], \quad \alpha \in (0, 1),$$

where  $\mathbb{I}[\cdot]$  denotes the indicator function (equal to 1 if the event appearing within brackets is realized and to 0 otherwise) and  $F_{\widehat{m}(\mathbf{X})}^{-1}$  is the quantile function, or generalized inverse of the distribution function  $F_{\widehat{m}(\mathbf{X})}$  of  $\widehat{m}(\mathbf{X})$ .

Considering an insurance context,  $\text{CC}[\mu(\mathbf{X}), \widehat{m}(\mathbf{X}); \alpha]$  represents the proportion of the total true pure premium that should be collected from policyholders who are actually charged the  $\alpha\%$  lower premiums computed with  $\widehat{m}(\cdot)$ . Notice that the concentration curve does not depend on the scale of  $\widehat{m}(\cdot)$  in the sense that replacing it with any increasing transformation (for instance, doubling  $\widehat{m}(\cdot)$ ) does not modify it. Only the ranking induced by  $\widehat{m}(\cdot)$  matters for the concentration curve.

In practice, we can replace the unknown  $\mu(\mathbf{X})$  with the response  $Y$  in the concentration curve. Precisely, the following identity holds

$$\text{CC}[\mu(\mathbf{X}), \widehat{m}(\mathbf{X}); \alpha] = \text{CC}[Y, \widehat{m}(\mathbf{X}); \alpha] = \frac{1}{\mu} \mathbb{E}[Y \mathbb{I}[\widehat{m}(\mathbf{X}) \leq F_{\widehat{m}(\mathbf{X})}^{-1}(\alpha)]] \quad (2.1)$$

for every probability level  $\alpha$ . Formula (2.1) shows that we can equivalently replace the unknown  $\mu(\mathbf{X})$  with the observed  $Y$ , which opens the door to estimation. Performances of competing models are then evaluated with the help of the empirical counterpart to (2.1).

## 2.2 Lorenz curve

The Lorenz curve associated with  $\hat{m}(\mathbf{X})$  is defined as the function

$$\begin{aligned}\alpha \mapsto \text{LC}[\hat{m}(\mathbf{X}); \alpha] &= \text{CC}[\hat{m}(\mathbf{X}), \hat{m}(\mathbf{X}); \alpha] \\ &= \frac{\mathbb{E}[\hat{m}(\mathbf{X}) \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]]}{\mathbb{E}[\hat{m}(\mathbf{X})]}, \quad \alpha \in (0, 1).\end{aligned}\quad (2.2)$$

Notice that the Lorenz curve only depends on  $\hat{m}(\mathbf{X})$  and not on the response  $Y$  (beyond its estimation on the validation set). In an insurance context,  $\text{LC}[\hat{m}(\mathbf{X}); \alpha]$  gives the percentage of the total pure premium contributed by policyholders who are charged the lower  $\alpha\%$  actual premiums computed with  $\hat{m}(\cdot)$ .

Ideally, we should have  $\text{LC}[\hat{m}(\mathbf{X}); \alpha] = \text{CC}[\mu(\mathbf{X}), \hat{m}(\mathbf{X}); \alpha]$  for all  $\alpha$ . This implies that  $\hat{m}(\cdot)$  is informative and measured on the correct scale, that is that pure premiums effectively balance insurance losses. By (1.1), the identity can indeed be rewritten as

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

$$\Leftrightarrow \mathbb{E}[(\mu(\mathbf{X}) - \hat{m}(\mathbf{X})) \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]] = 0 \text{ for all } \alpha, \quad (2.3)$$

which means that the pricing error  $\mu(\mathbf{X}) - \hat{m}(\mathbf{X})$  is on average 0 when considering every sub-portfolio corresponding to the condition  $\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)$ , that is, all those policyholders who are charged the  $\alpha\%$  lower premiums with  $\hat{m}(\cdot)$ .

Condition (2.3) is precisely related to the concept of auto-calibration. Recall that an estimator  $\hat{m}$  is said to be auto-calibrated if  $\mathbb{E}[Y | \hat{m}(\mathbf{X})] = \hat{m}(\mathbf{X})$  almost surely (Krüger and Ziegel, 2021). Under the unbiasedness assumption (1.1), this definition is equivalent to the equality between the Lorenz and concentration curves; see Denuit et al. (2024).

## 2.3 Area between the curves

The area between the curves, ABC in short, introduced in Denuit et al. (2019), is defined as

$$\text{ABC}[\hat{m}(\mathbf{X})] = \int_0^1 \left( \text{CC}[\mu(\mathbf{X}), \hat{m}(\mathbf{X}); \alpha] - \text{LC}[\hat{m}(\mathbf{X}); \alpha] \right) d\alpha. \quad (2.4)$$

The idea is to make (2.3) operational through a single numerical value. In Denuit et al. (2019) and Denuit et al. (2024), the concentration curve uniformly dominates the Lorenz curve in the numerical illustrations. However, this dominance is not guaranteed to hold in general. In fact, the ABC can be equal to zero even though the two curves do not coincide. Local violations may occur, as shown next.

**Example 2.1.** The difference between the concentration and Lorenz curves can be written as

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

If there exists  $\alpha^* \in (0, 1)$  such that  $\delta(m) = \mathbb{E}[Y | \hat{m}(\mathbf{X}) = m] - m$  is non-positive and non-decreasing in  $m$  for  $m \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*)$ , then  $\delta(\hat{m}(\mathbf{X})) \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*)]$  and  $\mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]$  are negatively correlated for all  $\alpha \in (0, 1)$ , since  $m \mapsto \delta(m) \mathbb{I}[m \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*)]$  is non-decreasing in  $m$  and  $m \mapsto \mathbb{I}[m \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]$  is non-increasing in  $m$ . Consequently, for all  $\alpha \in (0, 1)$ , we have

$$\begin{aligned} &\text{Cov}[\delta(\hat{m}(\mathbf{X})) \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*)], \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]] \\ &= \mathbb{E}[\delta(\hat{m}(\mathbf{X})) \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*)] \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]] \\ &\quad - \mathbb{E}[\delta(\hat{m}(\mathbf{X})) \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*)]] \mathbb{E}[\mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]] \\ &\leq 0, \end{aligned}$$

so that we obtain

$$\begin{aligned} &\mathbb{E}[\delta(\hat{m}(\mathbf{X})) \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*)] \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]] \\ &\leq \mathbb{E}[\delta(\hat{m}(\mathbf{X})) \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*)]] \mathbb{E}[\mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]] \\ &\leq 0, \end{aligned}$$

since  $\delta(\hat{m}(\mathbf{X}))$  is non-positive for  $\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*)$ . Therefore, we have

$$\text{CC}[\mu(\mathbf{X}), \hat{m}(\mathbf{X}); \alpha] \leq \text{LC}[\hat{m}(\mathbf{X}); \alpha], \quad \text{for all } \alpha \leq \alpha^*.$$

The latter inequality may invert for some  $\alpha > \alpha^*$ , because  $\mathbb{E}[\delta(\hat{m}(\mathbf{X}))] = 0$  implies that  $\delta(\hat{m}(\mathbf{X}))$  must be non-negative for certain values of  $\hat{m}(\mathbf{X})$  exceeding  $F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*)$ , and

$$\begin{aligned} &\mathbb{E}[\delta(\hat{m}(\mathbf{X})) \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]] \\ &= \mathbb{E}[\delta(\hat{m}(\mathbf{X})) \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*)]] + \mathbb{E}[\delta(\hat{m}(\mathbf{X})) \mathbb{I}[F_{\hat{m}(\mathbf{X})}^{-1}(\alpha^*) < \hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]], \end{aligned}$$

for  $\alpha > \alpha^*$ .

Under (1.1), we get

$$\text{ABC}[\hat{m}(\mathbf{X})] = \frac{1}{\mu} \int_0^1 \mathbb{E}[(Y - \hat{m}(\mathbf{X})) \mathbb{I}[\hat{m}(\mathbf{X}) \leq F_{\hat{m}(\mathbf{X})}^{-1}(\alpha)]] d\alpha. \quad (2.5)$$

This expression allows us to estimate the area between the curves, by substituting the empirical counterparts to the concentration and Lorenz curves in the integral.

### 3 Vanishing area between the curves and pricing error

#### 3.1 Continuous predictors

In this section, we assume that  $\widehat{m}(\mathbf{X})$  is a continuous random variable. This is typically the case when the score has been produced by a Generalized Additive Model including continuous features or by machine learning algorithms like random forests or neural networks.

The next result revisits the developments in Denuit et al. (2019), offering a re-statement in terms of pricing error.

**Proposition 3.1.** *Assume that  $\widehat{m}(\mathbf{X})$  is a continuous random variable and that (1.1) holds true. Then,  $ABC[\widehat{m}(\mathbf{X})] = 0$  is equivalent to require that the pricing error  $\mu(\mathbf{X}) - \widehat{m}(\mathbf{X})$  is uncorrelated with the rank induced by the actual premium  $\widehat{m}(\cdot)$ , that is,*

$$ABC[\widehat{m}(\mathbf{X})] = 0 \Leftrightarrow \text{Cov}[\mu(\mathbf{X}) - \widehat{m}(\mathbf{X}), F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))] = 0. \quad (3.1)$$

*Proof.* The integrated Lorenz curve is equal to

$$\begin{aligned} \int_0^1 LC[\widehat{m}(\mathbf{X}); \alpha] d\alpha &= \frac{1}{\mu} \mathbb{E} \left[ \widehat{m}(\mathbf{X}) \int_0^1 \mathbb{I}[\widehat{m}(\mathbf{X}) \leq F_{\widehat{m}(\mathbf{X})}^{-1}(\alpha)] d\alpha \right] \\ &= \frac{1}{\mu} \mathbb{E} \left[ \widehat{m}(\mathbf{X}) \int_0^1 \mathbb{I}[F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X})) \leq \alpha] d\alpha \right] \\ &= \frac{1}{\mu} \mathbb{E} [\widehat{m}(\mathbf{X}) (1 - F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X})))] \\ &= \frac{1}{2} - \frac{1}{\mu} \text{Cov}[\widehat{m}(\mathbf{X}), F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))]. \end{aligned}$$

The same reasoning applied to the concentration curve gives

$$\begin{aligned} \int_0^1 CC[\mu(\mathbf{X}), \widehat{m}(\mathbf{X}); \alpha] d\alpha &= \frac{1}{\mu} \mathbb{E} [\mu(\mathbf{X}) (1 - F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X})))] \\ &= \frac{1}{2} - \frac{1}{\mu} \text{Cov}[\mu(\mathbf{X}), F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))]. \end{aligned}$$

This ends the proof. □

Proposition 3.1 shows that a vanishing ABC is equivalent to the absence of correlation between the pricing error and the ranks induced by the candidate premium. Stated otherwise, there is no more information contained in these ranks about the pricing error.

Notice that the area below the Lorenz curve is closely related to an indicator of variability, referred to as Gini coefficient in the literature. To see this, consider two risk profiles  $\mathbf{X}$  and  $\mathbf{X}'$  taken at random from the population of insured individuals, so that  $\mathbf{X}$  and  $\mathbf{X}'$  are independent and identically distributed. Recall that the Gini coefficient of dispersion of  $\widehat{m}(\mathbf{X})$  is defined as

$$\text{Gini}[\widehat{m}(\mathbf{X})] = \frac{1}{2} \mathbb{E}[|\widehat{m}(\mathbf{X}) - \widehat{m}(\mathbf{X}')|].$$

The larger the expected absolute difference appearing in the Gini coefficient, the finer the risk classification induced by  $\widehat{m}(\cdot)$  in the sense that the absolute difference in the premiums  $\widehat{m}(\mathbf{X})$  and  $\widehat{m}(\mathbf{X}')$  charged to two policyholders taken at random from the insurance portfolio gets larger on average. Stated otherwise, premiums charged to different policyholders differ on average to a larger extent. Now,

$$\begin{aligned} \mathbb{E}[|\widehat{m}(\mathbf{X}) - \widehat{m}(\mathbf{X}')|] &= \mathbb{E}[\max\{\widehat{m}(\mathbf{X}), \widehat{m}(\mathbf{X}')\}] - \mathbb{E}[\min\{\widehat{m}(\mathbf{X}), \widehat{m}(\mathbf{X}')\}] \\ &= \int_0^\infty \left(1 - (F_{\widehat{m}(\mathbf{X})}(m))^2\right) dm - \int_0^\infty \left(1 - F_{\widehat{m}(\mathbf{X})}(m)\right)^2 dm \\ &= 2 \int_0^\infty F_{\widehat{m}(\mathbf{X})}(m)(1 - F_{\widehat{m}(\mathbf{X})}(m)) dm. \end{aligned}$$

Integration by parts then gives

$$\begin{aligned} \mathbb{E}[|\widehat{m}(\mathbf{X}) - \widehat{m}(\mathbf{X}')|] &= 4 \left( \mathbb{E}[\widehat{m}(\mathbf{X}) F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))] - \frac{1}{2} \mathbb{E}[\widehat{m}(\mathbf{X})] \right) \\ &= 4 \text{Cov}[\widehat{m}(\mathbf{X}), F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))] \end{aligned}$$

so that we recover the covariance appearing in the integral of the Lorenz curve as well as the classical identities

$$\text{Gini}[\widehat{m}(\mathbf{X})] = 2 \text{Cov}[\widehat{m}(\mathbf{X}), F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))] = \mu \left( 1 - 2 \int_0^1 \text{LC}[\widehat{m}(\mathbf{X}); \alpha] d\alpha \right).$$

This covariance thus appears to be an indicator of dispersion of the premiums obtained with  $\widehat{m}(\cdot)$ .

The covariance representation of the Gini coefficient of dispersion opens the door to its extension into a dependence measure. This leads to the Co-Gini coefficient defined as

$$\text{CoGini}[\mu(\mathbf{X}), \widehat{m}(\mathbf{X})] = 2 \text{Cov}[\mu(\mathbf{X}), F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))].$$

The Co-Gini coefficient is related to the integral of the concentration curve, as it can be seen from the proof of Proposition 3.1. Co-Gini coefficient measures the information contained in the ranks induced by the candidate premium about the true pure premium. The equivalent condition (3.1) can then be re-stated as the equality between Gini and Co-Gini coefficients.

### 3.2 Discrete predictors

Assume that  $\widehat{m}(\mathbf{X})$  takes only  $K < \infty$  distinct finite values  $\widehat{m}_1 < \dots < \widehat{m}_K$  such that  $\Pr[\widehat{m}(\mathbf{X}) = \widehat{m}_k] = p_k > 0$  for all  $k = 1, \dots, K$ . This is typically the case when the score has been produced by a GLM integrating discrete features, or by a single regression tree. Let

$$T^{(k)} = \mathbb{E}[(Y - \widehat{m}(\mathbf{X})) \mathbb{I}[\widehat{m}(\mathbf{X}) \leq \widehat{m}_k]], \quad k = 1, \dots, K, \quad \text{with } T^{(K)} = 0.$$

Then, (2.5) can be expressed as

$$\text{ABC}[\widehat{m}(\mathbf{X})] = \frac{1}{\mathbb{E}[Y]} \sum_{k=1}^K p_k T^{(k)}. \quad (3.2)$$

*Remark 3.2.* Notice that under unbiasedness  $E[Y] = E[\hat{m}(\mathbf{X})]$ ,  $E[Y]ABC[\hat{m}(\mathbf{X})]$  corresponds to the unscaled version  $ABC^\circ[\hat{m}(\mathbf{X})]$  defined in Wüthrich (2025), while our expression for  $ABC^\circ[\hat{m}(\mathbf{X})]$  deduced from (3.2) does not exactly correspond to the expression for  $ABC^\circ[\hat{m}(\mathbf{X})]$  given in Wüthrich (2025).

## 4 Vanishing area between the curves and portfolio re-weighting

### 4.1 Continuous predictors

We know from (3.1) that  $ABC[\hat{m}(\mathbf{X})] = 0$  is equivalent to the absence of correlation between the pricing error  $\mu(\mathbf{X}) - \hat{m}(\mathbf{X})$  and the rank  $F_{\hat{m}(\mathbf{X})}(\hat{m}(\mathbf{X}))$ . In this section, we establish a second equivalence with global balance in an hypothetical portfolio where every risk profile receives a new relative weight. This is precisely stated next.

**Proposition 4.1.** *Assume that  $\hat{m}(\mathbf{X})$  is a continuous random variable and that (1.1) holds true. Let  $(Y^*, \mathbf{X}^*)$  be such that*

$$f_{\hat{m}(\mathbf{X}^*)}(m) = 2f_{\hat{m}(\mathbf{X})}(m)(1 - F_{\hat{m}(\mathbf{X})}(m)), \quad m \geq 0,$$

and  $E[Y|\hat{m}(\mathbf{X}) = m] = E[Y^*|\hat{m}(\mathbf{X}^*) = m]$ . Then,

$$ABC[\hat{m}(\mathbf{X})] = 0 \Leftrightarrow E[Y^*] = E[\hat{m}(\mathbf{X}^*)]. \quad (4.1)$$

*Proof.* First,

$$\begin{aligned} E[\hat{m}(\mathbf{X}^*)] &= 2 \int_0^\infty m f_{\hat{m}(\mathbf{X})}(m)(1 - F_{\hat{m}(\mathbf{X})}(m)) dm \\ &= 2E[\hat{m}(\mathbf{X})(1 - F_{\hat{m}(\mathbf{X})}(\hat{m}(\mathbf{X})))]. \end{aligned}$$

Moreover

$$\begin{aligned} E[Y^*] &= \int_0^\infty E[Y^*|\hat{m}(\mathbf{X}^*) = m] f_{\hat{m}(\mathbf{X}^*)}(m) dm \\ &= 2 \int_0^\infty E[Y^*|\hat{m}(\mathbf{X}^*) = m] f_{\hat{m}(\mathbf{X})}(m)(1 - F_{\hat{m}(\mathbf{X})}(m)) dm \\ &= 2 \int_0^\infty E[Y|\hat{m}(\mathbf{X}) = m] f_{\hat{m}(\mathbf{X})}(m)(1 - F_{\hat{m}(\mathbf{X})}(m)) dm \\ &= 2E[E[Y|\hat{m}(\mathbf{X})](1 - F_{\hat{m}(\mathbf{X})}(\hat{m}(\mathbf{X})))]. \end{aligned}$$

Therefore,  $E[Y^*] = E[\hat{m}(\mathbf{X}^*)]$  if, and only if,

$$\begin{aligned} &E[\hat{m}(\mathbf{X})(1 - F_{\hat{m}(\mathbf{X})}(\hat{m}(\mathbf{X})))] = E[E[Y|\hat{m}(\mathbf{X})](1 - F_{\hat{m}(\mathbf{X})}(\hat{m}(\mathbf{X})))] \\ \Leftrightarrow &E[\hat{m}(\mathbf{X})(1 - F_{\hat{m}(\mathbf{X})}(\hat{m}(\mathbf{X})))] = E[Y(1 - F_{\hat{m}(\mathbf{X})}(\hat{m}(\mathbf{X})))] \\ \Leftrightarrow &E[\hat{m}(\mathbf{X})F_{\hat{m}(\mathbf{X})}(\hat{m}(\mathbf{X}))] = E[\mu(\mathbf{X})F_{\hat{m}(\mathbf{X})}(\hat{m}(\mathbf{X}))] \\ \Leftrightarrow &\text{Cov}[\hat{m}(\mathbf{X}), F_{\hat{m}(\mathbf{X})}(\hat{m}(\mathbf{X}))] = \text{Cov}[\mu(\mathbf{X}), F_{\hat{m}(\mathbf{X})}(\hat{m}(\mathbf{X}))] \end{aligned}$$

since  $E[\hat{m}(\mathbf{X})] = E[Y]$  by (1.1). This ends the proof.  $\square$

The proof of the equivalence (4.1) provides us with another interpretation of the condition  $\text{ABC}[\widehat{m}(\mathbf{X})] = 0$  in terms of global balance under a quota-share policy condition. Consider a portfolio where the indemnity paid by the insurer for a loss amount  $Y$  is reduced to  $Y F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))$ . Then, the premium is also proportionally reduced by the same factor. In this setting, a policyholder with a larger premium  $\widehat{m}(\mathbf{X})$  is indemnified to a larger extent, in relative terms. From the proof of Proposition 4.1, we see that

$$\text{ABC}[\widehat{m}(\mathbf{X})] = 0 \Leftrightarrow \mathbb{E}[\widehat{m}(\mathbf{X}) F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))] = \mathbb{E}[Y F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))]$$

which corresponds to global balance in this quota-share adjusted portfolio. This ensures that expected premiums  $\mathbb{E}[\widehat{m}(\mathbf{X}) F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))]$  match expected claims  $\mathbb{E}[Y F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))]$ . Notice that the same interpretation holds by replacing  $F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))$  with  $1 - F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))$  thanks to (1.1), meaning that policyholders with larger  $\widehat{m}(\mathbf{X})$  are indemnified to a lesser extent.

The next result reveals the way the portfolio is re-weighted.

**Property 4.2.** (i)  $\widehat{m}(\mathbf{X}^*)$  is smaller than  $\widehat{m}(\mathbf{X})$  in the likelihood ratio order, that is, the ratio  $f_{\widehat{m}(\mathbf{X})}(m)/f_{\widehat{m}(\mathbf{X}^*)}(m)$  of their respective probability density functions is non-decreasing in  $m$ .

(ii)  $\widehat{m}(\mathbf{X}^*)$  is distributed as  $\min\{\widehat{m}(\mathbf{X}'), \widehat{m}(\mathbf{X})\}$  where  $\mathbf{X}$  and  $\mathbf{X}'$  are independent and identically distributed.

*Proof.* The result stated under (i) directly follows from

$$\frac{f_{\widehat{m}(\mathbf{X})}(m)}{f_{\widehat{m}(\mathbf{X}^*)}(m)} = \frac{1}{2(1 - F_{\widehat{m}(\mathbf{X})}(m))}$$

which is non-decreasing in  $m$ . To get (ii), notice that the probability density function of  $\min\{\widehat{m}(\mathbf{X}'), \widehat{m}(\mathbf{X})\}$  is obtained from

$$-\frac{d}{dm} \Pr[\min\{\widehat{m}(\mathbf{X}'), \widehat{m}(\mathbf{X})\} > m] = -\frac{d}{dm} (1 - F_{\widehat{m}(\mathbf{X})}(m))^2 = 2(1 - F_{\widehat{m}(\mathbf{X})}(m)) f_{\widehat{m}(\mathbf{X})}(m)$$

which is indeed  $f_{\widehat{m}(\mathbf{X}^*)}(m)$ . This ends the proof.  $\square$

Property 4.2 shows that the hypothetical portfolio where global balance must hold puts more weight on the risk profiles with smaller premiums  $\widehat{m}(\cdot)$ .

## 4.2 Discrete predictors

It turns out that the condition  $\text{ABC}[\widehat{m}(\mathbf{X})] = 0$  is also equivalent to require global balance in an hypothetical portfolio where the relative weight of each risk profile has been modified. This is precisely stated next.

**Proposition 4.3.** Define  $\alpha_0 = 0$ ,  $\alpha_j = \sum_{k=1}^j p_k$  for  $j = 1, \dots, K-1$ ,

$$p^* = \sum_{k=1}^K (1 - \alpha_{k-1}) p_k \text{ and } p_k^* = \frac{(1 - \alpha_{k-1}) p_k}{p^*}, \quad k = 1, \dots, K.$$

Then,

$$ABC[\widehat{m}(\mathbf{X})] = 0 \Leftrightarrow \sum_{k=1}^K p_k^* \mathbb{E}[Y | \widehat{m}(\mathbf{X}) = \widehat{m}_k] = \sum_{k=1}^K p_k^* \widehat{m}_k. \quad (4.2)$$

*Proof.* Let

$$S^{(k)} = \mathbb{E}[(Y - \widehat{m}(\mathbf{X})) \mathbb{I}[\widehat{m}(\mathbf{X}) = \widehat{m}_k]], \quad k = 1, \dots, K.$$

We have  $T^{(k)} = \sum_{j=1}^k S^{(j)}$ . Then

$$\begin{aligned} ABC[\widehat{m}(\mathbf{X})] &= \frac{1}{\mathbb{E}[Y]} \sum_{k=1}^K p_k \sum_{j=1}^k S^{(j)} \\ &= \frac{1}{\mathbb{E}[Y]} \sum_{j=1}^K \sum_{k=j}^K p_k S^{(j)} \\ &= \frac{1}{\mathbb{E}[Y]} \sum_{j=1}^K (1 - \alpha_{j-1}) S^{(j)}. \end{aligned} \quad (4.3)$$

Now, let

$$R^{(k)} = \mathbb{E}[Y | \widehat{m}(\mathbf{X}) = \widehat{m}_k] - \widehat{m}_k, \quad k = 1, \dots, K.$$

Since  $S^{(k)} = R^{(k)} p_k$ , we get from (4.3) that

$$\begin{aligned} ABC[\widehat{m}(\mathbf{X})] &= \frac{1}{\mathbb{E}[Y]} \sum_{k=1}^K (1 - \alpha_{k-1}) p_k R^{(k)} \\ &= \frac{p^*}{\mathbb{E}[Y]} \sum_{k=1}^K p_k^* R^{(k)}. \end{aligned} \quad (4.4)$$

The announced result (4.2) then follows from (4.4).  $\square$

Equation (4.2) expresses global balance in an hypothetical portfolio with a modified composition, in the sense that the relative weight  $p_k$  of risk class  $k$  is replaced with  $p_k^*$ . The next result studies these modified weights. It can be seen as the discrete counterpart to Property 4.2.

**Property 4.4.** (i)  $(p_k^*)$  is smaller than  $(p_k)$  in the sense of the likelihood ratio order, that is, the ratio  $p_k/p_k^*$  is non-decreasing in  $k$ .

(ii) If  $\mathbf{X}'$  is an independent copy of  $\mathbf{X}$  then

$$p_k^* = \frac{\Pr[\min\{\widehat{m}(\mathbf{X}'), \widehat{m}(\mathbf{X})\} = \widehat{m}_k] + \Pr[\widehat{m}(\mathbf{X}') = \widehat{m}(\mathbf{X}) = \widehat{m}_k]}{1 + \Pr[\widehat{m}(\mathbf{X}') = \widehat{m}(\mathbf{X})]}.$$

*Proof.* Item (i) follows from the fact that the ratio

$$\frac{p_k}{p_k^*} = \frac{p^*}{1 - \alpha_{k-1}}$$

is non-decreasing in  $k$ . Turning to (ii),

$$\begin{aligned}
\Pr[\min\{\widehat{m}(\mathbf{X}'), \widehat{m}(\mathbf{X})\} = \widehat{m}_k] &= \Pr[\widehat{m}(\mathbf{X}') = \widehat{m}_k, \widehat{m}(\mathbf{X}) \geq \widehat{m}_k] \\
&\quad + \Pr[\widehat{m}(\mathbf{X}) = \widehat{m}_k, \widehat{m}(\mathbf{X}') > \widehat{m}_k] \\
&= 2 \Pr[\widehat{m}(\mathbf{X}') = \widehat{m}_k, \widehat{m}(\mathbf{X}) \geq \widehat{m}_k] \\
&\quad - \Pr[\widehat{m}(\mathbf{X}') = \widehat{m}(\mathbf{X}) = \widehat{m}_k],
\end{aligned}$$

so that we have

$$\begin{aligned}
p_k(1 - \alpha_{k-1}) &= \Pr[\widehat{m}(\mathbf{X}') = \widehat{m}_k, \widehat{m}(\mathbf{X}) \geq \widehat{m}_k] \\
&= \frac{\Pr[\min\{\widehat{m}(\mathbf{X}'), \widehat{m}(\mathbf{X})\} = \widehat{m}_k] + \Pr[\widehat{m}(\mathbf{X}') = \widehat{m}(\mathbf{X}) = \widehat{m}_k]}{2}.
\end{aligned}$$

Thus, we get

$$p^* = \frac{1 + \Pr[\widehat{m}(\mathbf{X}') = \widehat{m}(\mathbf{X})]}{2},$$

and hence

$$p_k^* = \frac{\Pr[\min\{\widehat{m}(\mathbf{X}'), \widehat{m}(\mathbf{X})\} = \widehat{m}_k] + \Pr[\widehat{m}(\mathbf{X}') = \widehat{m}(\mathbf{X}) = \widehat{m}_k]}{1 + \Pr[\widehat{m}(\mathbf{X}') = \widehat{m}(\mathbf{X})]},$$

which ends the proof.  $\square$

Property 4.4 shows that the hypothetical portfolio where (4.2) expresses global balance over-weights risk classes with smaller pure premiums compared to the initial portfolio.

## 5 Numerical illustrations

We present two numerical examples, one with a continuous predictor and one with a discrete predictor, illustrating how the equivalent characterizations of a vanishing ABC apply.

### 5.1 Continuous predictor

**Model specification.** Assume that the fitted pure premium

$$\widehat{m}(\mathbf{X}) \sim \text{Unif}(0, 1),$$

so that its distribution function satisfies

$$F_{\widehat{m}(\mathbf{X})}(m) = m, \quad m \in [0, 1].$$

We specify the conditional mean of the response given the fitted pure premium as

$$\mathbb{E}[Y|\widehat{m}(\mathbf{X})] = \widehat{m}(\mathbf{X}) + a(\widehat{m}(\mathbf{X}) - 0.5) + b(\widehat{m}(\mathbf{X}) - 0.5)^3, \quad (5.1)$$

where  $a$  and  $b$  are real constants. This specification allows for departures from auto-calibration while remaining analytically tractable.

**Unbiasedness.** Since  $\widehat{m}(\mathbf{X}) \sim \text{Unif}(0, 1)$ ,

$$\mathbb{E}[\widehat{m}(\mathbf{X}) - 0.5] = 0, \quad \mathbb{E}[(\widehat{m}(\mathbf{X}) - 0.5)^3] = 0.$$

Hence, from (5.1),

$$\mathbb{E}[Y] = \mathbb{E}[\widehat{m}(\mathbf{X})] + a \mathbb{E}[\widehat{m}(\mathbf{X}) - 0.5] + b \mathbb{E}[(\widehat{m}(\mathbf{X}) - 0.5)^3] = \mathbb{E}[\widehat{m}(\mathbf{X})] = 0.5.$$

Therefore, the unbiasedness (global balance) condition (1.1) holds for all values of  $a$  and  $b$ .

**Vanishing ABC via the covariance characterization.** Define

$$g(\widehat{m}(\mathbf{X})) = \mathbb{E}[Y | \widehat{m}(\mathbf{X})]$$

as the conditional mean given the fitted pure premium. Recall that  $\mu(\mathbf{X}) = \mathbb{E}[Y | \mathbf{X}]$  and that the pricing error is  $\mu(\mathbf{X}) - \widehat{m}(\mathbf{X})$ . By the tower property,

$$\mathbb{E}[\mu(\mathbf{X}) | \widehat{m}(\mathbf{X})] = \mathbb{E}[\mathbb{E}[Y | \mathbf{X}] | \widehat{m}(\mathbf{X})] = g(\widehat{m}(\mathbf{X})),$$

so that

$$\mathbb{E}[\mu(\mathbf{X}) - \widehat{m}(\mathbf{X}) | \widehat{m}(\mathbf{X})] = g(\widehat{m}(\mathbf{X})) - \widehat{m}(\mathbf{X}).$$

According to Proposition 3.1,

$$\text{ABC}[\widehat{m}(\mathbf{X})] = 0 \Leftrightarrow \text{Cov}[\mu(\mathbf{X}) - \widehat{m}(\mathbf{X}), F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X}))] = 0.$$

In the present setting,

$$F_{\widehat{m}(\mathbf{X})}(\widehat{m}(\mathbf{X})) = \widehat{m}(\mathbf{X}).$$

Using again the tower property together with unbiasedness,

$$\text{Cov}[\mu(\mathbf{X}) - \widehat{m}(\mathbf{X}), \widehat{m}(\mathbf{X})] = \text{Cov}[g(\widehat{m}(\mathbf{X})) - \widehat{m}(\mathbf{X}), \widehat{m}(\mathbf{X})].$$

Since  $\widehat{m}(\mathbf{X}) \sim \text{Unif}(0, 1)$ , a direct calculation gives

$$\text{Cov}[g(\widehat{m}(\mathbf{X})) - \widehat{m}(\mathbf{X}), \widehat{m}(\mathbf{X})] = \frac{a}{12} + \frac{b}{80}. \quad (5.2)$$

Therefore, according to Proposition 3.1, the ABC vanishes if and only if

$$\frac{a}{12} + \frac{b}{80} = 0. \quad (5.3)$$

**Choice of parameters.** To obtain a concrete numerical illustration, we choose

$$a = 0.03, \quad b = -0.2, \quad (5.4)$$

which satisfy the above condition (5.3). For this choice, the function  $m \mapsto g(m)$  is increasing on  $[0, 1]$  and remains strictly positive, with

$$g(0) = 0.01, \quad g(1) = 0.99.$$

**Lorenz and concentration curves do not coincide.** The Lorenz and concentration curves differ since  $g(m) \neq m$ .

The Lorenz curve associated with  $\widehat{m}(\mathbf{X})$  is

$$\text{LC}[\widehat{m}(\mathbf{X}); \alpha] = \frac{\mathbb{E}[\widehat{m}(\mathbf{X})\mathbb{I}[\widehat{m}(\mathbf{X}) \leq \alpha]]}{\mathbb{E}[\widehat{m}(\mathbf{X})]} = \alpha^2, \quad \alpha \in (0, 1).$$

By the tower property,

$$\mathbb{E}[\mu(\mathbf{X})\mathbb{I}[\widehat{m}(\mathbf{X}) \leq \alpha]] = \mathbb{E}[g(\widehat{m}(\mathbf{X}))\mathbb{I}[\widehat{m}(\mathbf{X}) \leq \alpha]].$$

Since  $\mathbb{E}[Y] = 0.5$ , the concentration curve is

$$\text{CC}[\mu(\mathbf{X}), \widehat{m}(\mathbf{X}); \alpha] = 2 \int_0^\alpha g(m) \, dm, \quad \alpha \in (0, 1).$$

Using (5.1), we obtain

$$\text{CC}[\mu(\mathbf{X}), \widehat{m}(\mathbf{X}); \alpha] = \frac{\alpha(-5\alpha^3 + 10\alpha^2 + 44\alpha + 1)}{50}.$$

Hence,

$$\text{CC}[\mu(\mathbf{X}), \widehat{m}(\mathbf{X}); \alpha] - \text{LC}[\widehat{m}(\mathbf{X}); \alpha] = -\frac{\alpha(\alpha - 1)(5\alpha^2 - 5\alpha + 1)}{50}.$$

The difference changes sign on  $(0, 1)$ , so the two curves cross, yet

$$\int_0^1 (\text{CC}[\mu(\mathbf{X}), \widehat{m}(\mathbf{X}); \alpha] - \text{LC}[\widehat{m}(\mathbf{X}); \alpha]) \, d\alpha = 0,$$

which shows that  $\text{ABC}[\widehat{m}(\mathbf{X})] = 0$ , as predicted by Proposition 3.1.

**Re-weighted portfolio.** The density of the re-weighted fitted pure premium  $\widehat{m}(\mathbf{X}^*)$  defined in Proposition 4.1 is

$$f_{\widehat{m}(\mathbf{X}^*)}(m) = 2(1 - m), \quad m \in [0, 1],$$

which coincides with the density of  $\min\{\widehat{m}(\mathbf{X}), \widehat{m}(\mathbf{X}')\}$  for an independent copy  $\mathbf{X}'$ . This re-weighting assigns relatively more mass to smaller fitted pure premiums.

Let  $\mu^*(\mathbf{X}^*) = \mathbb{E}[Y^*|\mathbf{X}^*]$  denote the conditional mean in the re-weighted portfolio. By construction,

$$\mathbb{E}[Y^*|\widehat{m}(\mathbf{X}^*) = m] = \mathbb{E}[Y|\widehat{m}(\mathbf{X}) = m] = g(m).$$

The tower property yields

$$\mathbb{E}[\mu^*(\mathbf{X}^*)] = \mathbb{E}[\mathbb{E}[Y^*|\widehat{m}(\mathbf{X}^*)]] = \mathbb{E}[g(\widehat{m}(\mathbf{X}^*))].$$

By definition, global balance in the re-weighted portfolio means

$$\mathbb{E}[\mu^*(\mathbf{X}^*)] = \mathbb{E}[\widehat{m}(\mathbf{X}^*)].$$

Combining these two relations, we see that global balance in the re-weighted portfolio is equivalent to

$$\mathbb{E}[g(\widehat{m}(\mathbf{X}^*))] = \mathbb{E}[\widehat{m}(\mathbf{X}^*)].$$

Finally,

$$\mathbb{E}[g(\widehat{m}(\mathbf{X}^*)) - \widehat{m}(\mathbf{X}^*)] = -2\text{Cov}[g(\widehat{m}(\mathbf{X})) - \widehat{m}(\mathbf{X}), \widehat{m}(\mathbf{X})] = 0$$

by (5.2) and (5.4), so that

$$\mathbb{E}[\mu^*(\mathbf{X}^*)] = \mathbb{E}[\widehat{m}(\mathbf{X}^*)].$$

Thus, global balance holds in this re-weighted portfolio, as predicted by Proposition 4.1.

## 5.2 Discrete predictor

We now consider a simple discrete example with a finite number of premium classes. It illustrates the discrete formulas derived in Section 3.2 and the re-weighted portfolio characterization in Section 4.2.

**Model specification.** Assume that the fitted pure premium  $\widehat{m}(\mathbf{X})$  can take only three distinct values

$$\widehat{m}_1 < \widehat{m}_2 < \widehat{m}_3,$$

with

$$(\widehat{m}_1, \widehat{m}_2, \widehat{m}_3) = (0.6, 1.0, 1.4),$$

and that

$$\Pr[\widehat{m}(\mathbf{X}) = \widehat{m}_k] = p_k, \quad k = 1, 2, 3,$$

with

$$(p_1, p_2, p_3) = (0.2, 0.5, 0.3).$$

The true class-wise conditional means are taken as

$$\mu_k = \mathbb{E}[Y | \widehat{m}(\mathbf{X}) = \widehat{m}_k], \quad k = 1, 2, 3,$$

with

$$(\mu_1, \mu_2, \mu_3) = \left(\frac{4}{5}, \frac{111}{125}, \frac{109}{75}\right) \approx (0.800, 0.888, 1.453).$$

The ordering of fitted premiums and true means is consistent:

$$\widehat{m}_1 < \widehat{m}_2 < \widehat{m}_3 \quad \text{and} \quad \mu_1 < \mu_2 < \mu_3.$$

The associated class-wise deviations are

$$R^{(k)} = \mu_k - \widehat{m}_k, \quad k = 1, 2, 3,$$

that is,

$$(R^{(1)}, R^{(2)}, R^{(3)}) = \left(\frac{1}{5}, -\frac{14}{125}, \frac{4}{75}\right) \approx (0.200, -0.112, 0.053).$$

**Unbiasedness and vanishing ABC.** Recall that the individual-level pricing error is  $\mu(\mathbf{X}) - \hat{m}(\mathbf{X})$ , with  $\mu(\mathbf{X}) = \mathbb{E}[Y|\mathbf{X}]$ . At the class level, the average deviation is  $R^{(k)} = \mu_k - \hat{m}_k$ .

First, global balance holds:

$$\mathbb{E}[\hat{m}(\mathbf{X})] = \sum_{k=1}^3 p_k \hat{m}_k = 1.04 = \sum_{k=1}^3 p_k \mu_k = \mathbb{E}[\mu(\mathbf{X})] = \mathbb{E}[Y].$$

Define, as in Section 4.2,

$$S^{(k)} = \mathbb{E}[(Y - \hat{m}(\mathbf{X}))\mathbb{I}[\hat{m}(\mathbf{X}) = \hat{m}_k]] = p_k R^{(k)}, \quad k = 1, 2, 3,$$

and

$$T^{(k)} = \mathbb{E}[(Y - \hat{m}(\mathbf{X}))\mathbb{I}[\hat{m}(\mathbf{X}) \leq \hat{m}_k]] = \sum_{j=1}^k S^{(j)}, \quad k = 1, 2, 3.$$

For the above specification,

$$(S^{(1)}, S^{(2)}, S^{(3)}) = \left(\frac{1}{25}, -\frac{7}{125}, \frac{2}{125}\right) \approx (0.040, -0.056, 0.016),$$

and

$$(T^{(1)}, T^{(2)}, T^{(3)}) = \left(\frac{1}{25}, -\frac{2}{125}, 0\right) \approx (0.040, -0.016, 0).$$

Using (3.2), we have

$$\text{ABC}[\hat{m}(\mathbf{X})] = \frac{1}{\mathbb{E}[Y]} \sum_{k=1}^3 p_k T^{(k)} = 0.$$

**Re-weighted portfolio.** Let

$$\alpha_0 = 0, \quad \alpha_1 = p_1 = 0.2, \quad \alpha_2 = p_1 + p_2 = 0.7,$$

and

$$p^* = \sum_{k=1}^3 (1 - \alpha_{k-1}) p_k = 0.690, \quad p_k^* = \frac{(1 - \alpha_{k-1}) p_k}{p^*}.$$

This yields

$$(p_1^*, p_2^*, p_3^*) = \left(\frac{20}{69}, \frac{40}{69}, \frac{3}{23}\right) \approx (0.29, 0.58, 0.13).$$

Compared to the original distribution (0.2, 0.5, 0.3), the re-weighted portfolio puts relatively more mass on lower premium classes, in agreement with Property 4.4.

Equation (4.2) states that vanishing ABC is equivalent to global balance under these modified weights:

$$\sum_{k=1}^3 p_k^* \mathbb{E}[Y | \hat{m}(\mathbf{X}) = \hat{m}_k] = \sum_{k=1}^3 p_k^* \hat{m}_k.$$

In the present example,

$$\sum_{k=1}^3 p_k^* \mu_k = \sum_{k=1}^3 p_k^* \hat{m}_k = \frac{323}{345} \approx 0.936.$$

This discrete example mirrors the conclusions of the continuous case. A vanishing ABC does not imply auto-calibration at the class level, since the deviations  $R^{(k)}$  are non-zero. Instead, it reflects exact global balance in a suitably re-weighted portfolio that over-emphasizes lower premium classes, as characterized by Property 4.4.

## Acknowledgements

Both authors wish to dedicate this paper to Professor Joseph Glaz, the founding editor of Methodology and Computing in Applied Probability (MCAP). They have had the privilege of serving as Associate Editors of MCAP (Michel from 2008 to 2013 and Julien from 2015 to 2024) and have greatly benefited, as both editors and authors, from this remarkable journal that brings together applied probabilists from diverse fields. 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

- Denuit, M., Huyghe, J., Trufin, J., Verdebout, T. (2024). Testing for auto-calibration with Lorenz and Concentration curves. Insurance: Mathematics and Economics 117, 130-139.
- 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.
- Krüger, F., Ziegel, J.F. (2021). Generic conditions for forecast dominance. Journal of Business & Economic Statistics 39, 972-983.
- Wüthrich, M.V. (2025). Auto-calibration tests for discrete finite regression functions. European Actuarial Journal 15, 335-341.
- Yitzhaki, S., Schechtman, E. (2013). The Gini Methodology: A Primer on Statistical Methodology. Springer.