

LORENZ CURVE, GINI COEFFICIENT, AND TWEEDIE DOMINANCE FOR AUTOCALIBRATED PREDICTORS

Michel Denuit, Julien Trufin

LIDAM Discussion Paper ISBA
2021 / 36

ISBA

Voie du Roman Pays 20 - L1.04.01

B-1348 Louvain-la-Neuve

Email : lidam-library@uclouvain.be

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

LORENZ CURVE, GINI COEFFICIENT, AND TWEEDIE DOMINANCE FOR AUTOCALIBRATED PREDICTORS

Michel Denuit

Institute of Statistics, Biostatistics and Actuarial Science - ISBA

Louvain Institute of Data Analysis and Modeling - LIDAM

UCLouvain

Louvain-la-Neuve, Belgium

Julien Trufin

Department of Mathematics

Université Libre de Bruxelles (ULB)

Brussels, Belgium

November 3, 2021

Abstract

Denuit et al. (2019, 2021a, 2021b) proposed diagnostic tools based on concentration curves and associated metrics (ICC and ABC), as well as Tweedie dominance to compare competing models. This is in contrast with professional practice which generally resorts to Lorenz curves and Gini coefficients for that purpose. This note reconciles both approaches for autocalibrated predictors. Autocalibration is a desirable property, intimately related to the method of marginal totals that predates modern risk classification methods. It can easily be implemented using the practical method proposed by Denuit et al. (2021a), consisting in an extra local regression step. Under autocalibration, it is shown that Lorenz curve and concentration curve coincide, that ICC is equivalent to Gini coefficient, and that ABC is zero. The latter property thus offers an easy check for autocalibration.

Keywords: Risk classification, Tweedie distribution family, Concentration curve, ABC, ICC, Convex order.

1 Introduction and motivation

Since Frees et al. (2011, 2014), it is known that the respective positions of the graphs of the Lorenz and concentration curves allow the actuary to accurately assess performance of the premium under consideration. Building on these pioneering contributions, Denuit et al. (2019) proposed insurance-related metrics to evaluate the performances of candidate premiums (ABC and ICC, precisely defined in the next sections), as well as the convex order to select the optimal pricing tool. Denuit et al. (2021b) related this approach to expectation dependence and proposed an effective testing procedure for this dependence concept and successfully applied it to model selection. This is in contrast with the dominant practice using Lorenz curve and Gini coefficient.

Empirical evidence suggests that boosting models and neural networks trained to minimize deviance may end up with a significant under-estimation of total claims. To solve this problem, Denuit et al. (2021a) proposed a new strategy based on the concept of autocalibration (see, e.g., Kruger and Ziegel, 2020) to restore global balance as well as local equilibrium in the spirit of the original method of marginal totals, or MMT in short. Dating back to the 1960s, when North-American actuaries pioneered risk classification with the help of minimum bias methods, the central idea behind MMT 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. 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. Denuit et al. (2021a) show how to perform autocalibration by adding an extra step implementing the method of marginal totals (MMT) with the help of local regression.

This note reconciles the dominant practice based on Lorenz curves and Gini coefficients to these recent methodological advances for autocalibrated predictors. Once autocalibration has been implemented, we show that Lorenz curve and concentration curve coincide, that ICC is equivalent to Gini coefficient, and that ABC is zero. The latter property can be used to check whether autocalibration has been properly achieved.

The remainder of the paper is structured as follows. Section 2 recalls the concept of autocalibration. In Section 3, we show that concentration curve and Lorenz curve coincide for autocalibrated predictors. Section 4 considers the insurance metrics ICC and ABC under autocalibration. Section 5 is devoted to Tweedie dominance.

2 Autocalibration

2.1 Ratemaking problem

Our setting is the same as in Denuit et al. (2019, 2021a, 2021b). It corresponds to the situation typically faced by actuaries involved in insurance ratemaking. We consider a response Y of interest (number of claims or yearly claim cost, for instance) and a set of features $X_1, \dots, X_p$ gathered in the vector $\mathbf{X}$ that can explain Y . In actuarial ratemaking, the aim

is to evaluate the pure premium as accurately as possible. This means that the target is the conditional expectation $\mu(\mathbf{X}) = E[Y|\mathbf{X}]$ of the response Y given the available information $\mathbf{X}$. Henceforth, $\mu(\mathbf{X})$ is referred to as the true (pure) premium.

The function $\mathbf{x} \mapsto \mu(\mathbf{x}) = E[Y|\mathbf{X} = \mathbf{x}]$ 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. 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.

2.2 Performance evaluation

The developments in the present note holds the training set as fixed. Precisely, $\hat{\pi}$ is assumed to have been built from some training set (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 $\{(Y_i, \mathbf{X}_i), i = 1, 2, \dots, n\}$, that has not been used to obtain $\hat{\pi}$. Provided the validation set comprises a large number of independent and identically distributed pairs $(Y_i, \mathbf{X}_i)$, this allows us to perform calculations with a generic random vector $(Y, \mathbf{X})$, independent of, and distributed as the observations contained in the training set. This approach is meaningful in insurance ratemaking where the actuary is typically in a data-rich situation.

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 $\mathbf{X}$. 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 $\mathbf{X}$. Henceforth, we denote as $F_{\hat{\pi}}$ the distribution function of $\hat{\pi}$, assumed to be continuous and strictly increasing throughout this text.

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. This naturally leads to the concept of autocalibration.

2.3 Autocalibrated predictors

Recall that a predictor $\hat{\pi}$ is said to be autocalibrated if $\hat{\pi}(\mathbf{X}) = E[Y|\hat{\pi}(\mathbf{X})]$. We refer the reader to Kruger and Ziegel (2020) for a general presentation of this concept. Autocalibration ensures that the average response in a neighborhood defined with the help of the autocalibrated predictor under consideration matches premium income. This exactly corresponds to the MMT and is thus desirable for any candidate premium.

If $s \mapsto E[Y|\hat{\pi}(\mathbf{X}) = s]$ is continuously increasing, a simple way to restore global balance consists in switching from $\hat{\pi}$ to its balance-corrected version $\hat{\pi}_{BC}$ defined as

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

that averages to $E[Y]$, as shown in Property 5.1 in Denuit et al. (2021a). This shows that $\hat{\pi}$ can easily be autocalibrated by performing 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. Balance-correction by local GLM recognizes the nature of the response under consideration: local Poisson GLM for claim counts, local Gamma GLM for average claim severities and local compound Poisson-Gamma GLM for claim 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. The canonical link function is adopted so that maximum-likelihood estimates replicate observed totals.

Henceforth, we restrict our analysis to autocalibrated predictors.

3 Lorenz and concentration curves for autocalibrated predictors

Let $F_{\hat{\pi}}^{-1}$ be the quantile function (or Value-at-Risk) of $\hat{\pi}$. 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 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 $\hat{\pi}$. The true premium is not observed in reality, but the following identity holds

$$\text{CC}[\mu(\mathbf{X}), \hat{\pi}(\mathbf{X}); \alpha] = \text{CC}[Y, \hat{\pi}(\mathbf{X}); \alpha] = \frac{E[Y\mathbf{I}[\hat{\pi}(\mathbf{X}) \leq F_{\hat{\pi}}^{-1}(\alpha)]]}{E[Y]} \quad (3.1)$$

for every probability level α . Formula (3.1) shows that we can equivalently replace the pure premium $\mu(\mathbf{X})$ with the response Y in the concentration curve.

The Lorenz curve LC associated with the predictor $\hat{\pi}(\mathbf{X})$ is defined as

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

For an exhaustive review of the properties of a concentration and Lorenz curves, we refer the interested reader to the book by Yitzhaki and Schechtman (2013).

Remark 4.3 in Denuit et al. (2019) points out that CC and LC are equal when $\mu(\mathbf{X})$ and $\hat{\pi}(\mathbf{X})$ are perfectly positively dependent (i.e. comonotonic). Here, we show that this is generally true for any autocalibrated predictor.

Property 3.1. *If $\hat{\pi}$ is autocalibrated, then we have*

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

for every probability level α .

Proof. Starting from (3.1), we have that

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

which ends the proof. $\square$

According to Definition 4.1 in Denuit et al. (2019), the best-performing predictor has smaller LC and CC. However, many empirical studies only use the Lorenz curve to evaluate the performances of a predictive model. This dominant practice is generally flawed since it confuses the predictor $\hat{\pi}$ with the conditional expectation (so that LC and CC coincide) whereas it is very likely that $\hat{\pi}(\mathbf{X})$ only approximates $\mu(\mathbf{X})$. Therefore, analysts must resort to the pair of curves CC and LC to evaluate performance of a pricing model. The main message of Property 3.1 is that the use of Lorenz curves only, without reference to concentration curves, is justified for autocalibrated predictors.

4 ICC and ABC metrics for autocalibrated predictors

In addition to the respective positions of the graphs of the concentration and Lorenz curves, Denuit et al. (2019) suggest basing the comparison both on the integral of the concentration curves (ICC) and the area between the two curves (ABC). These insurance metrics are defined as

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

and

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

respectively. A better model has smaller ICC and ABC.

Whereas ICC and ABC have sound methodological foundation, dominant practice uses Gini coefficients to compare predictors. Recall that Gini mean difference is defined for a non-negative continuous random variable Z as

$$\text{Gini}[Z] = \text{E}[|Z_1 - Z_2|] = \text{E}[\max\{Z_1, Z_2\}] - \text{E}[\min\{Z_1, Z_2\}]$$

where Z_1 and Z_2 are independent and distributed as Z . It represents the average absolute difference between two observations distributed as Z (whereas variance is based on squared difference). If Z is continuous then it can be shown that

$$\text{Gini}[Z] = 4\text{Cov}[Z, F_Z(Z)].$$

Thus, Gini mean difference measures the association between a random variable and its rank. The Gini coefficient is the Gini mean difference divided by twice the mean. It is also known as the concentration ratio and represents the area between the 45-degree line and the actual Lorenz curve divided by the area between the 45-degree line and the Lorenz curve that yields the maximal value that this index can have. This is why candidate premiums with larger Gini mean difference tend to be preferred.

The next result derives expression of the insurance metrics ICC and ABC for autocalibrated predictors and relates ICC to Gini coefficient.

Property 4.1. *If $\hat{\pi}$ is autocalibrated, then we have*

$$\text{ICC}[\mu(\mathbf{X}), \hat{\pi}(\mathbf{X})] = \frac{1}{2} - \frac{\text{Gini}[\hat{\pi}(\mathbf{X})]}{\mathbb{E}[\hat{\pi}(\mathbf{X})]}$$

and $\text{ABC}[\mu(\mathbf{X}), \hat{\pi}(\mathbf{X})] = 0$.

Proof. Define $\hat{\Pi} = F_{\hat{\pi}}(\hat{\pi}(\mathbf{X}))$ that is uniformly distributed over the unit interval $[0, 1]$. We know from Section 4.2 in Denuit et al. (2019a) that

$$\text{ICC}[\mu(\mathbf{X}), \hat{\pi}(\mathbf{X})] = \frac{1}{2} - \frac{\text{Cov}[\mu(\mathbf{X}), \hat{\Pi}]}{\mathbb{E}[Y]}.$$

Now,

$$\text{Cov}[Y, \hat{\Pi}] = \text{Cov}[\mu(\mathbf{X}), \hat{\Pi}] \text{ and } \text{Cov}[Y, \hat{\Pi}] = \text{Cov}[\hat{\pi}(\mathbf{X}), \hat{\Pi}]$$

for autocalibrated predictors. This gives the announced result for ICC. We know from Property 3.1 that $\text{CC} = \text{LC}$ for autocalibrated predictors so that we directly obtain $\text{ABC} = 0$ in that case. This ends the proof. $\square$

For autocalibrated models, we get from Property 4.1 that $\hat{\pi}_2$ outperforms $\hat{\pi}_1$ when

$$\text{ICC}[\mu(\mathbf{X}), \hat{\pi}_2(\mathbf{X})] \leq \text{ICC}[\mu(\mathbf{X}), \hat{\pi}_1(\mathbf{X})] \Leftrightarrow \text{Gini}[\hat{\pi}_2(\mathbf{X}), \hat{\Pi}_2] \geq \text{Gini}[\hat{\pi}_1(\mathbf{X}), \hat{\Pi}_1].$$

We thus see that comparing model performances with the help of ICC is equivalent to comparing Gini coefficients when predictors are autocalibrated. Again, this justifies the dominant practice provided autocalibration has been implemented.

Also, Property 4.1 shows that a nonzero estimated ABC suggests that the predictor under consideration is not autocalibrated. Initially proposed as an indicator of the performance of a given predictor in Denuit et al. (2019), it turns out that ABC is more useful for detecting whether autocalibration has been properly implemented.

5 Tweedie dominance for autocalibrated predictors

Actuarial studies often 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 .

Only the cases $\xi \geq 1$ are relevant for applications in insurance, with $\xi = 1$ corresponding to the Poisson distribution, $1 < \xi < 2$ to compound Poisson sums with Gamma-distributed terms, and $\xi \geq 2$ to continuous distributions over the positive real line with $\xi = 2$ for the Gamma distribution and $\xi = 3$ for the Inverse Gaussian one. In this section, we thus restrict our analysis to $\xi \geq 1$. Models are compared on the basis of the out-of-(training) sample, or predictive Tweedie deviance that essentially reduces to

$$D(\xi, \hat{\pi}) = \begin{cases} \mathbb{E} [\hat{\pi}(\mathbf{X}) - Y \ln \hat{\pi}(\mathbf{X})] & \text{for } \xi = 1 \\ \mathbb{E} \left[\ln \hat{\pi}(\mathbf{X}) + \frac{Y}{\hat{\pi}(\mathbf{X})} \right] & \text{for } \xi = 2 \\ \mathbb{E} \left[\frac{\hat{\pi}(\mathbf{X})^{2-\xi}}{2-\xi} - \frac{Y \hat{\pi}(\mathbf{X})^{1-\xi}}{1-\xi} \right] & \text{for } \xi > 1 \text{ and } \xi \neq 2 \end{cases} \quad (5.1)$$

for a given predictor $\hat{\pi}$.

Typically, the analyst selects a specific parameter ξ and assesses the respective performances of the predictors $\hat{\pi}_1$ and $\hat{\pi}_2$ under consideration with the help of the corresponding Tweedie deviance values. Instead of comparing two predictors $\hat{\pi}_1$ and $\hat{\pi}_2$ on the basis of a specific deviance, Denuit et al. (2021a) suggest to require superiority for every Tweedie deviance (5.1), that is, for any power parameter $\xi \geq 1$. 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 $\xi \geq 1$.

Proposition 4.1 in Denuit et al. (2021a) provides the actuary with a sufficient condition for a model to outperform a competitor in terms of Tweedie dominance. The next result reinforces this result for autocalibrated predictors.

Proposition 5.1. *For two autocalibrated predictors $\hat{\pi}_1$ and $\hat{\pi}_2$, $\hat{\pi}_2$ outperforms $\hat{\pi}_1$ in terms of Tweedie dominance if, and only if, the inequality*

$$\mathbb{E} [\psi_\xi(\hat{\pi}_1(\mathbf{X}))] \leq \mathbb{E} [\psi_\xi(\hat{\pi}_2(\mathbf{X}))] \quad (5.2)$$

holds true for every power parameter $\xi \geq 1$, where

$$\psi_\xi(\pi) = \begin{cases} \pi \ln \pi & \text{for } \xi = 1 \\ -\ln \pi & \text{for } \xi = 2 \\ \frac{\pi^{2-\xi}}{\xi-2} & \text{for } \xi > 1 \text{ and } \xi \neq 2. \end{cases} \quad (5.3)$$

Proof. For an autocalibrated predictor $\hat{\pi}$ and any measurable function $g : \mathbb{R} \rightarrow \mathbb{R}$, we have that

$$\begin{aligned} \mathbb{E} [Yg(\hat{\pi}(\mathbf{X}))] &= \mathbb{E} [\mathbb{E} [Yg(\hat{\pi}(\mathbf{X})) | \hat{\pi}(\mathbf{X})]] \\ &= \mathbb{E} [\mathbb{E} [Y | \hat{\pi}(\mathbf{X})] g(\hat{\pi}(\mathbf{X}))] \\ &= \mathbb{E} [\hat{\pi}(\mathbf{X})g(\hat{\pi}(\mathbf{X}))] \end{aligned} \quad (5.4)$$

Hence, one sees that we can equivalently replace Y with the autocalibrated predictor $\hat{\pi}(\mathbf{X})$ in $E[Yg(\hat{\pi}(\mathbf{X}))]$. Therefore, the predictive deviance $D(\xi, \hat{\pi})$ defined in (5.1) simplifies to

$$D(\xi, \hat{\pi}) = \begin{cases} E[\hat{\pi}(\mathbf{X}) - \hat{\pi}(\mathbf{X}) \ln \hat{\pi}(\mathbf{X})] & \text{for } \xi = 1 \\ E[\ln \hat{\pi}(\mathbf{X}) + 1] & \text{for } \xi = 2 \\ E\left[\frac{\hat{\pi}(\mathbf{X})^{2-\xi}}{2-\xi} - \frac{\hat{\pi}(\mathbf{X})^{2-\xi}}{1-\xi}\right] & \text{for } \xi > 1 \text{ and } \xi \neq 2, \end{cases} \quad (5.5)$$

which ends the proof. $\square$

For any $\xi \geq 1$, the function ψ_ξ defined in (5.3) is convex. Therefore, if $\hat{\pi}_2(\mathbf{X})$ dominates $\hat{\pi}_1(\mathbf{X})$ in the convex order (in the sense that the Lorenz curve for $\hat{\pi}_2(\mathbf{X})$ lies below that for $\hat{\pi}_1(\mathbf{X})$) then $\hat{\pi}_2$ outperforms $\hat{\pi}_1$ in Tweedie dominance. Notice that the test functions ψ_ξ in Proposition 5.1 are not only convex but have second, third, fourth, etc. derivatives alternating in sign. Precisely, the m th derivative of ψ_ξ is positive when m is even and negative when m is odd for any $m \geq 2$. This echoes the characterization of the Laplace order with the help functions with a completely monotone first derivative (Theorem 5.A.4 in Shaked and Shanthikumar, 2007). It turns out that Laplace order is a sufficient condition for Tweedie dominance, as shown next.

Proposition 5.2. *Let $L_{\hat{\pi}}(\cdot)$ denote the Laplace transform of $\hat{\pi}$, that is, $L_{\hat{\pi}}(s) = E[\exp(-s\hat{\pi})]$. For two autocalibrated predictors $\hat{\pi}_1$ and $\hat{\pi}_2$, if the inequality $L_{\hat{\pi}_1}(s) \leq L_{\hat{\pi}_2}(s)$ holds true for all $s \geq 0$, then $\hat{\pi}_2$ outperforms $\hat{\pi}_1$ in terms of Tweedie dominance.*

Proof. With the representation of completely monotone functions (see e.g. Lemma 2.1 in Merkle, 2014), we obtain

$$\psi_\xi''(t) = \int_0^\infty \exp(-ut) d\vartheta(u)$$

for some measure ϑ on $(0, \infty)$. This gives

$$\begin{aligned} \psi_\xi'(t) &= c_1 - \int_0^\infty \frac{\exp(-ut)}{u} d\vartheta(u) \\ \psi_\xi(t) &= c_2 + c_1 t + \int_0^\infty \frac{\exp(-ut)}{u^2} d\vartheta(u) \end{aligned}$$

for some real constants c_1 and c_2 . We then obtain the following representation for the expectations of the test functions ψ_ξ :

$$E[\psi_\xi(\hat{\pi})] = c_2 + c_1 E[\hat{\pi}] + \int_0^\infty \frac{1}{u^2} L_{\hat{\pi}}(u) d\vartheta(u),$$

which ends the proof. $\square$

This shows that Tweedie dominance for autocalibrated predictors is a rather weak dominance concept. The testing procedure proposed by Bhattacharyya et al. (2021) can be applied to check whether two autocalibrated predictors are ranked according to Laplace order. Tweedie dominance thus allows multiple crossings between the Lorenz curves associated to the autocalibrated predictors under consideration. We refer the interested reader to Chiu (2007, 2021) for more information about intersecting Lorenz curves.

References

- Bhattacharyya, D., Khan, R. A., Mitra, M. (2021). Tests for Laplace order dominance with applications to insurance data. *Insurance: Mathematics and Economics* 99, 163-173.
- Chiu, W. H. (2007). Intersecting Lorenz curves, the degree of downside inequality aversion, and tax reforms. *Social Choice and Welfare* 28, 375-399.
- Chiu, W. H. (2021). Intersecting Lorenz curves and aversion to inverse downside inequality. *Social Choice and Welfare* 56, 487-508.
- Denuit, M., Charpentier, A., Trufin, J. (2021a). Autocalibration and Tweedie-dominance for insurance pricing with machine learning. *Insurance: Mathematics and Economics*, in press.
- 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.
- Denuit, M., Trufin, J., Verdebout, Th. (2021b). Testing for more positive expectation dependence with application to model comparison. *Insurance: Mathematics and Economics*, in press.
- Frees, E.W., Meyers, G., Cummings, A.D. (2011). Summarizing insurance scores using a Gini index. *Journal of the American Statistical Association* 106, 1085-1098.
- Frees, E.W., Meyers, G., Cummings, A.D. (2014). Insurance ratemaking and a Gini index. *Journal of Risk and Insurance* 81, 335-366.
- Kruger, F., Ziegel, J.F. (2020). Generic conditions for forecast dominance. *Journal of Business & Economic Statistics*, in press.
- Merkle, M. (2014). Completely monotone functions: A digest. In “Analytic Number Theory, Approximation Theory, and Special Functions” (pp. 347-364). Springer, New York.
- Shaked, M., Shanthikumar, J.G. (2007). *Stochastic Orders*. Springer, New York.
- Yitzhaki, S., Schechtman, E. (2013). *The Gini Methodology: A Primer on Statistical Methodology*. Springer.