

# TWEEDIE DOMINANCE FOR AUTOCALIBRATED PREDICTORS AND LAPLACE TRANSFORM ORDER

Michel Denuit, Julien Trufin

LIDAM Discussion Paper ISBA  
2022 / 40

## **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.html>

# TWEEDIE DOMINANCE FOR AUTOCALIBRATED PREDICTORS AND LAPLACE TRANSFORM ORDER

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 26, 2022

## Abstract

While Krüger and Ziegel (2021) defined forecast dominance, or Bregman dominance as dominance for every Bregman loss function, this short note explores Tweedie dominance proposed by Denuit et al. (2021) to compare competing models. A necessary and sufficient condition is established under autocalibration. Moreover, Laplace transform order turns out to be a sufficient condition for Tweedie dominance between autocalibrated predictors. This shows that Tweedie dominance is a rather weak concept compared to Bregman dominance that reduces to the well-known convex order among autocalibrated predictors.

**Keywords:** Tweedie distribution family, Autocalibrated models, Tweedie dominance, Laplace transform order.

# 1 Introduction and motivation

## 1.1 Regression for the mean

Consider a response  $Y$  and a set of features  $X_1, \dots, X_p$  gathered in the vector  $\mathbf{X} \in \mathcal{X}$  (classically,  $\mathcal{X} \subset \mathbb{R}^p$ ). In a regression setting, the dependence structure inside the random vector  $(Y, X_1, \dots, X_p)$  is exploited to extract the information contained in  $\mathbf{X}$  about  $Y$ . Here, we assume that the target is the conditional expectation  $\mu(\mathbf{X}) = \mathbb{E}[Y|\mathbf{X}]$  of the response  $Y$  given the available information  $\mathbf{X}$ .

The function  $\mathbf{x} \mapsto \mu(\mathbf{x}) = \mathbb{E}[Y|\mathbf{X} = \mathbf{x}]$  is approximated by a predictor  $\mathbf{x} \mapsto \pi(\mathbf{x})$  with a simpler structure. Once fitted by minimizing an empirical loss function on the training data set, this produces estimates  $\hat{\pi}(\mathbf{x})$  for  $\mu(\mathbf{x})$ , or fitted values.

## 1.2 Bregman loss functions

Conditional expectations can be consistently estimated only if the expected loss is minimum for the mean response. This corresponds to the concept of consistent loss function. Precisely, a loss function is consistent for the mean if no other quantity leads to a lower expected loss than the mean.

Bregman loss functions are of the form

$$L(y, m) = \ell(y) - \ell(m) - \ell'(m)(y - m) \text{ for a convex function } \ell. \quad (1.1)$$

Savage (1971) showed, subject to weak regularity conditions, that for any loss function (1.1) and response  $Y$ ,

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

The class of Bregman loss functions is thus consistent for the (conditional) mean functional.

## 1.3 Tweedie loss functions

Generalized Linear Models (GLMs) with power variance function have been successfully used in many applications. Poisson, Gamma and Inverse Gaussian distributions belong to this family, as well as compound Poisson distributions with Gamma-distributed summands. Within the Tweedie subclass of the Exponential Dispersion family, 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  $\phi$  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  $\xi \geq 1$ .

For an appropriate choice of  $\ell$  in (1.1), we recover Tweedie deviance. Precisely, the

function  $\ell(m) = 2(m\theta_m - a(\theta_m))$  with  $a'(\theta_m) = m$  results in the loss function

$$\begin{aligned}
L(y, m) &= 2(y(\theta_y - \theta_m) - a(\theta_y) + a(\theta_m)) \text{ where } a'(\theta_y) = y \\
&= \begin{cases} 2(y \ln \frac{y}{m} - (y - m)) & \text{for } \xi = 1 \\ 2(-\ln \frac{y}{m} + \frac{y}{m} - 1) & \text{for } \xi = 2 \\ 2\left(\frac{y^{2-\xi}}{(1-\xi)(2-\xi)} - \frac{ym^{1-\xi}}{1-\xi} + \frac{m^{2-\xi}}{2-\xi}\right) & \text{for } \xi > 1 \text{ and } \xi \neq 2. \end{cases} \quad (1.2)
\end{aligned}$$

The deviance loss function (1.2) is thus a Bregman loss function

## 1.4 Performance evaluation

The developments in this paper hold the training set as fixed. Precisely,  $\hat{\pi}$  is assumed to have been obtained by minimizing the empirical loss on 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  $\{(Y_i, \mathbf{X}_i), i = 1, 2, \dots, n\}$ , 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.

We adopt this approach in this note and we assess performances of  $\hat{\pi}$  with the help of the predictive Tweedie deviance corresponding to (1.2) given by

$$D(\xi, \hat{\pi}) = \begin{cases} E[\hat{\pi}(\mathbf{X}) - Y \ln \hat{\pi}(\mathbf{X})] & \text{for } \xi = 1 \\ E\left[\ln \hat{\pi}(\mathbf{X}) + \frac{Y}{\hat{\pi}(\mathbf{X})}\right] & \text{for } \xi = 2 \\ E\left[\frac{\hat{\pi}(\mathbf{X})^{2-\xi}}{2-\xi} - \frac{Y\hat{\pi}(\mathbf{X})^{1-\xi}}{1-\xi}\right] & \text{for } \xi > 1 \text{ and } \xi \neq 2 \end{cases} \quad (1.3)$$

where the terms not involving  $\hat{\pi}(\mathbf{X})$  have been discarded.

## 2 Tweedie dominance and autocalibration

Henceforth, we consider  $Y \geq 0$  as well as  $\hat{\pi}(\mathbf{X}) \geq 0$ . We assume that  $\hat{\pi}(\mathbf{X})$  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  $\mathbf{X}$ .

Bregman dominance, also called forecast dominance is defined as dominance for every Bregman loss function (1.1). Precisely,  $\hat{\pi}_2$  outperforms  $\hat{\pi}_1$  in terms of Bregman dominance if the inequality  $E[L(Y, \hat{\pi}_2(\mathbf{X}))] \leq E[L(Y, \hat{\pi}_1(\mathbf{X}))]$  holds true for every loss function  $L$  of

the form given in (1.1). We refer the interested reader to Krüger and Ziegel (2021) and the references therein for an extensive presentation of this concept.

Tweedie dominance introduced in Denuit et al. (2021) restricts Bregman loss functions to Tweedie deviances (1.3). Precisely,  $\hat{\pi}_2$  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  $\xi \geq 1$ . In their Proposition 4.1, Denuit et al. (2021) established sufficient conditions for Tweedie dominance in terms of simple test functions.

Recall that a predictor  $\hat{\pi}$  is said to be autocalibrated if  $\hat{\pi}(\mathbf{X}) = \mathbb{E}[Y|\hat{\pi}(\mathbf{X})]$ . Autocalibration thus ensures that the average response in a neighborhood defined with the help of the autocalibrated predictor under consideration matches the latter. See, e.g., Krüger and Ziegel, (2021).

The next result shows that Tweedie dominance can be characterized by the comparison of simple test functions for autocalibrated predictors.

**Proposition 2.1.** *Consider two autocalibrated predictors  $\hat{\pi}_1$  and  $\hat{\pi}_2$ . Then,  $\hat{\pi}_2$  outperforms  $\hat{\pi}_1$  in terms of Tweedie dominance if, and only if, the inequality*

$$\mathbb{E}[\psi_\xi(\hat{\pi}_1(\mathbf{X}))] \leq \mathbb{E}[\psi_\xi(\hat{\pi}_2(\mathbf{X}))] \quad (2.1)$$

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

$$\psi_\xi(\pi) = \begin{cases} \pi \ln \pi & \text{for } \xi = 1 \\ -\ln \pi & \text{for } \xi = 2 \\ \frac{\pi^{2-\xi}}{\xi-2} & \text{for } \xi > 1 \text{ and } \xi \neq 2. \end{cases} \quad (2.2)$$

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

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

Therefore, the predictive deviance  $D(\xi, \hat{\pi})$  defined in (1.3) simplifies to

$$D(\xi, \hat{\pi}) = \begin{cases} \mathbb{E}[\hat{\pi}(\mathbf{X}) - \hat{\pi}(\mathbf{X}) \ln \hat{\pi}(\mathbf{X})] & \text{for } \xi = 1 \\ \mathbb{E}[\ln \hat{\pi}(\mathbf{X}) + 1] & \text{for } \xi = 2 \\ \mathbb{E}\left[\frac{(\hat{\pi}(\mathbf{X}))^{2-\xi}}{2-\xi} - \frac{(\hat{\pi}(\mathbf{X}))^{2-\xi}}{1-\xi}\right] & \text{for } \xi > 1 \text{ and } \xi \neq 2 \end{cases}. \quad (2.3)$$

Since  $\hat{\pi}_1$  and  $\hat{\pi}_2$  are autocalibrated, we have  $\mathbb{E}[\hat{\pi}_1(\mathbf{X})] = \mathbb{E}[\hat{\pi}_2(\mathbf{X})]$ . Hence,

$$D(1, \hat{\pi}_2) \leq D(1, \hat{\pi}_1) \Leftrightarrow \mathbb{E}[\hat{\pi}_1(\mathbf{X}) \ln \hat{\pi}_1(\mathbf{X})] \leq \mathbb{E}[\hat{\pi}_2(\mathbf{X}) \ln \hat{\pi}_2(\mathbf{X})]$$

which corresponds to the stated inequality (2.1) for  $\psi_1$  defined according to (2.2). Similarly,

$$D(2, \hat{\pi}_2) \leq D(2, \hat{\pi}_1) \Leftrightarrow \mathbb{E}[\ln \hat{\pi}_2(\mathbf{X})] \leq \mathbb{E}[\ln \hat{\pi}_1(\mathbf{X})]$$

which corresponds to (2.1) with  $\psi_2$  as in (2.2). Finally, for  $\xi > 1$  and  $\xi \neq 2$ ,

$$D(\xi, \hat{\pi}_2) \leq D(\xi, \hat{\pi}_1) \Leftrightarrow \mathbb{E} \left[ \frac{(\hat{\pi}_2(\mathbf{X}))^{2-\xi}}{(\xi-2)(1-\xi)} \right] \leq \mathbb{E} \left[ \frac{(\hat{\pi}_1(\mathbf{X}))^{2-\xi}}{(\xi-2)(1-\xi)} \right]$$

which is indeed (2.1) with  $\psi_\xi$  in (2.2), noticing that  $1 - \xi < 0$ . This ends the proof.  $\square$

For any  $\xi \geq 1$ , the function  $\psi_\xi$  defined in (2.2) is convex. Therefore, if  $\hat{\pi}_1(\mathbf{X})$  and  $\hat{\pi}_2(\mathbf{X})$  satisfy  $\hat{\pi}_1(\mathbf{X}) \preceq_{\text{CX}} \hat{\pi}_2(\mathbf{X})$  where  $\preceq_{\text{CX}}$  is the convex order defined as  $\mathbb{E}[g(\hat{\pi}_1(\mathbf{X}))] = \mathbb{E}[g(\hat{\pi}_2(\mathbf{X}))]$  for every convex function  $g$  such that the expectations exist, then  $\hat{\pi}_2$  outperforms  $\hat{\pi}_1$  in Tweedie dominance. This was expected since we know from Theorem 3.1 in Krüger and Ziegel (2021) that for two autocalibrated predictors  $\hat{\pi}_1$  and  $\hat{\pi}_2$ ,  $\hat{\pi}_2$  outperforms  $\hat{\pi}_1$  in terms of Bregman dominance if, and only if,  $\hat{\pi}_2(\mathbf{X})$  dominates  $\hat{\pi}_1(\mathbf{X})$  in the convex order. The convex order is thus a sufficient condition for Tweedie dominance in this case.

The test functions  $\psi_\xi$  in (2.2) are actually not only convex but have second, third, fourth, etc. derivatives alternating in sign. Precisely, the  $m$ th derivative of  $\psi_\xi$  is positive when  $m$  is even and negative when  $m$  is odd for any  $m \geq 2$ . This echoes the characterization of the Laplace transform order with the help of functions with a completely monotone first derivative (Theorem 5.A.4 in Shaked and Shanthikumar, 2007). It turns out that Laplace transform order is a sufficient condition for Tweedie dominance under autocalibration, as shown next.

**Proposition 2.2.** *Let  $L_{\hat{\pi}}(\cdot)$  denote the Laplace transform of  $\hat{\pi}$ , that is,  $L_{\hat{\pi}}(s) = \mathbb{E}[\exp(-s\hat{\pi})]$ . 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.* The second derivative of the test functions  $\psi_\xi$  in (2.2) is given by  $\psi_\xi''(t) = t^{-\xi}$  for all  $\xi \geq 1$ . We thus see that  $\psi_\xi''$  is non-negative and possesses negative odd derivatives and positive even derivatives. With the representation of completely monotone functions given e.g. in Lemma 2.1 of Merkle (2014), we obtain

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

for some measure  $\vartheta$  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  $\psi_\xi$ :

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

which ends the proof.  $\square$

Proposition 2.2 shows that Tweedie dominance for autocalibrated predictors is a rather weak dominance concept compared to Bregman dominance that reduces to the convex order. The testing procedure proposed by Bhattacharyya et al. (2021) can be applied to check whether two autocalibrated predictors are ranked according to Laplace order. Contrary to Bregman dominance that is equivalent to the ordering of the Lorenz curves associated to the autocalibrated predictors under consideration, Tweedie dominance thus allows multiple crossings between the Lorenz curves. We refer the interested reader to Chiu (2007, 2021) for more information about intersecting Lorenz curves.

It is tempting to conclude that Laplace transform order is in fact equivalent to Tweedie dominance. This is indeed the case if  $\ln \hat{\pi}_1 \geq 0$  and  $\ln \hat{\pi}_2 \geq 0$ , that is, if  $\hat{\pi}_1 \geq 1$  and  $\hat{\pi}_2 \geq 1$ , as shown in the next result.

**Proposition 2.3.** *Consider two autocalibrated predictors  $\hat{\pi}_1$  and  $\hat{\pi}_2$  such that  $\hat{\pi}_1 \geq 1$  and  $\hat{\pi}_2 \geq 1$ . Then,  $\hat{\pi}_2$  outperforms  $\hat{\pi}_1$  in terms of Tweedie dominance if, and only if,  $L_{\hat{\pi}_1}(s) \leq L_{\hat{\pi}_2}(s)$  for all  $s \geq 0$ .*

*Proof.* Tweedie dominance ensures that

$$\frac{E[\hat{\pi}_1(\mathbf{X})^{2-\xi}]}{\xi-2} \leq \frac{E[\hat{\pi}_2(\mathbf{X})^{2-\xi}]}{\xi-2} \quad \text{for all } \xi > 1 \text{ and } \xi \neq 2,$$

so that we have

$$E[\hat{\pi}_1(\mathbf{X})^{-s}] \leq E[\hat{\pi}_2(\mathbf{X})^{-s}] \quad \text{for all } s \geq 0,$$

which can be rewritten as

$$E[\exp(-s \ln \hat{\pi}_1(\mathbf{X}))] \leq E[\exp(-s \ln \hat{\pi}_2(\mathbf{X}))] \quad \text{for all } s \geq 0.$$

Since  $\ln \hat{\pi}_1 \geq 0$  and  $\ln \hat{\pi}_2 \geq 0$ , we can conclude that Tweedie dominance ensures that  $\ln \hat{\pi}_1$  and  $\ln \hat{\pi}_2$  are ordered in the Laplace transform order, that is,

$$\ln \hat{\pi}_2(\mathbf{X}) \preceq_{Lt} \ln \hat{\pi}_1(\mathbf{X}).$$

Now applying Theorem 5.A.7(a) in Shaked and Shanthikumar (2007) with  $g(t) = \exp(-t)$  gives

$$\exp(-\ln \hat{\pi}_2(\mathbf{X})) = \hat{\pi}_2(\mathbf{X})^{-1} \preceq_{Lt} \exp(-\ln \hat{\pi}_1(\mathbf{X})) = \hat{\pi}_1(\mathbf{X})^{-1}.$$

Then, applying once again the same theorem with  $g(t) = 1/t$  leads to

$$\hat{\pi}_2(\mathbf{X}) \preceq_{Lt} \hat{\pi}_1(\mathbf{X}),$$

so that  $\hat{\pi}_1$  and  $\hat{\pi}_2$  are also ordered in the Laplace transform order, which ends the proof.  $\square$

Proposition 2.3 applies when  $Y \geq 1$ . In the continuous case when  $Y$  obeys the Gamma distribution for instance, this essentially means that  $Y$  is at least equal to one measurement unit.

*Remark 2.4.* In general, i.e. when  $\hat{\pi}_1 \geq 0$  and  $\hat{\pi}_2 \geq 0$ , we can only get the following inequality. Define the function  $\chi(x) : x \in (0, \infty) \rightarrow x^{-s}$ ,  $s > 0$ . Since we have

$$\chi(x) = \int_0^\infty \exp(-ux) \frac{u^{s-1}}{\Gamma(s)} du,$$

then Tweedie dominance implies

$$\begin{aligned} & \mathbb{E}[\hat{\pi}_1(\mathbf{X})^{-s}] \leq \mathbb{E}[\hat{\pi}_2(\mathbf{X})^{-s}] \quad \text{for all } s > 0 \\ \Leftrightarrow & \mathbb{E} \left[ \int_0^\infty \exp(-u\hat{\pi}_1(\mathbf{X})) \frac{u^{s-1}}{\Gamma(s)} du \right] \leq \mathbb{E} \left[ \int_0^\infty \exp(-u\hat{\pi}_2(\mathbf{X})) \frac{u^{s-1}}{\Gamma(s)} du \right] \quad \text{for all } s > 0 \\ \Leftrightarrow & \int_0^\infty (L_{\hat{\pi}_1}(u) - L_{\hat{\pi}_2}(u)) u^{s-1} du \leq 0 \quad \text{for all } s > 0. \end{aligned}$$

## 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. (2021). Autocalibration and Tweedie dominance for insurance pricing with machine learning. *Insurance: Mathematics and Economics* 101, 485-497.
- Krüger, F., Ziegel, J.F. (2021). Generic conditions for forecast dominance. *Journal of Business & Economic Statistics* 39, 972-983.
- Merkle, M. (2014). Completely monotone functions: A digest. In “Analytic Number Theory, Approximation Theory, and Special Functions” (pp. 347-364). Springer, New York.
- Savage, L.J. (1971). Elicitation of personal probabilities and expectations. *Journal of the American Statistical Association* 66, 783-810.
- Shaked, M., Shanthikumar, J.G. (2007). *Stochastic Orders*. Springer, New York.