

國立臺灣大學理學院統計與數據科學研究所

碩士論文

Institute of Statistics and Data Science

College of Science

National Taiwan University

Master's Thesis



使用自助法技術對 LassoNet 模型去偏誤

Bias Reduction in LassoNet Models Using Bootstrap
Techniques

劉冠毅

Guan-Yi Liu

指導教授：楊鈞濤 博士

Advisor: Chun-Hao Yang Ph.D.

中華民國 114 年 2 月

February, 2025



摘要

LassoNet 是一種 Lasso 的擴展模型，其結合了前饋神經網絡以捕捉非線性關係，並同時保留其稀疏解。然而，它跟 Lasso 一樣有偏誤問題，特別是在高維廣義線性模型中。本文中將 LassoNet 擴展到高維廣義線性模型中，使其能夠在維持稀疏性的同時建模複雜的數據結構。然而因為該擴展模型的估計結果仍然與 Lasso 一樣存在偏誤問題，所以我們引入了 [van de Geer et al. \(2014\)](#) 提出的去偏誤架構，並通過基於自助法的校正方法作進一步改進。我們的改進提供了一個具有良好預測性和解釋性的去偏 LassoNet 估計式。我們提出的方法拓展 LassoNet 的應用範圍，並為分析高維數據中預測變量與響應變量之間的非線性和稀疏關係提供了一個可靠的框架。

關鍵字：神經網路、去偏誤、高維度模型、變數選擇、LassoNet





Abstract

LassoNet, an extension of Lasso, incorporates feed-forward neural networks (FFNs) to capture nonlinear relationships while retaining sparse solutions. However, it inherits Lasso's bias issues, particularly in high-dimensional GLMs. In this thesis, we extend LassoNet to GLMs, enabling it to model complex data structures while maintaining sparsity. As the estimations provided by this extended model still exhibit bias similar to Lasso, we address this issue by incorporating the debiasing framework introduced by [van de Geer et al. \(2014\)](#) and further enhancing it with a bootstrap-based correction. These refinements yield a debiased LassoNet estimator that is both predictive and interpretable. The proposed method broadens the applicability of LassoNet, providing a reliable framework for analyzing high-dimensional data with nonlinear and sparse relationships between predictors and response.

Keywords: Neural Networks, Bias Reduction, High-dimensional Models, Variable Selection, LassoNet

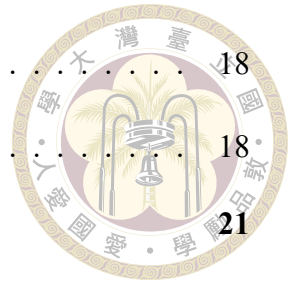




Contents

	Page
摘要	i
Abstract	iii
Contents	v
List of Figures	vii
Chapter 1 Introduction	1
1.1 Background	1
1.2 Proposed method	4
1.3 Organization of the thesis	4
Chapter 2 Preliminary	5
2.1 Lasso	5
2.2 Fully-Connected Neural Network	6
Chapter 3 Method	13
3.1 Lassonet	13
3.2 Debiasing Lasso Estimators	14
3.2.1 Node-wise Lasso regression	14
3.2.2 Extensions of node-wise Lasso regression	16
3.3 Debiasing the LassoNet	17

3.3.1	Merge LassoNet and GLM	18
3.3.2	Bootstrapping the Bias of LassoNet	18
Chapter 4	Simulation	21
4.1	Evaluation of estimation bias	22
4.2	Evaluating variable selection performance	24
Chapter 5	Conclusion	27
References		29





List of Figures

4.1	Comparison of the overall loss, measured as $\ \hat{\beta} - \beta\ _2$ across Lasso, the LassoNet, the debiased LassoNet estimator (DB-Net), the un-debiased bootstrap LassoNet (BS-Net), and the debiased bootstrap LassoNet (BS-DB-Net).	22
4.2	Comparison of the signal loss, measured as $\ \hat{B} - B\ _2$ across Lasso, the LassoNet, the debiased LassoNet estimator (DB-Net), the un-debiased bootstrap LassoNet (BS-Net), and the debiased bootstrap LassoNet (BS-DB-Net).	23
4.3	Comparison of the number of selections among Lasso, the debiased LassoNet estimator (DB-Net), the un-debiased bootstrap LassoNet estimator (BS-Net), and the debiased bootstrap LassoNet estimator (BS-DB-Net).	24
4.4	False negatives for the debiased LassoNet estimator (DB-Net), the un-debiased bootstrap LassoNet estimator (BS-Net), and the debiased bootstrap LassoNet estimator (BS-DB-Net).	25
4.5	False positives for the debiased LassoNet estimator (DB-Net), the un-debiased bootstrap LassoNet estimator (BS-Net), and the debiased bootstrap LassoNet estimator (BS-DB-Net).	25





Chapter 1 Introduction

1.1 Background

In recent years, high-dimensional regression models have been widely applied across various domains, particularly in fields where the number of predictors far exceeds the number of observations, such as genomics ([Feng and Simon, 2019](#)), finance ([Andrews and Lu, 2001](#)), and image analysis ([He et al., 2019](#); [Liu et al., 2024](#)). In these scenarios, only a few predictors are informative, while including too many irrelevant predictors can overly complicate the model, leading to overfitting. This has spurred the development of advanced techniques for variable selection, aiming to enhance model interpretability and improve predictive accuracy. Until now, numerous methods for variable selection in linear models have been developed, including stepwise selection, screening, and penalty-based approaches. Among these, the most classical penalty-based method is Lasso ([Tibshirani, 1996](#)), which has become a widely recognized approach, particularly in large p smaller n scenarios. The primary advantage of Lasso lies in its ability to efficiently achieve feature sparsity while simultaneously estimating coefficients by incorporating an ℓ_1 -regularization term into the loss function. Since the success of variable selection via Lasso, numerous Lasso-based methods and related studies have been developed, such as the *Adaptive Lasso* ([Zou, 2006](#)), the *Dantzig Selector* ([Candes and Tao, 2007](#)), and the *Elastic Net* ([Zou and](#)

[Hastie, 2005](#)), among others. It is well known that, under irrepresentable condition, the nonzero coefficients selected by Lasso are consistent with the true predictors with high probability. However, this condition primarily relies on the sample covariance matrix of the predictors satisfying certain structural assumptions, which are often difficult to verify in practice ([Zhao and Yu, 2006](#)). Despite these challenges, Lasso has provided a solid foundation for subsequent research and improvements in variable selection techniques.

In the aforementioned Lasso-based variable selection methods, the focus has been primarily on standard linear models, which assume a linear relationship between the response and predictors. [Lemhadri et al. \(2021\)](#) proposed LassoNet, which additionally incorporates a feed-forward neural network (FFN) to model the nonlinear effects. Although a nonlinear term is introduced into the model, they successfully mimic the Lasso framework to achieve sparse solutions and simultaneous estimation. LassoNet shares the same primary limitation as Lasso. Its inherent estimation bias is caused by the ℓ_1 -regularization term, which shrinks all coefficients, including the true predictors, toward zero. This bias not only reduces the accuracy of prediction performance but also complicates the interpretation of those nonzero coefficients. Luckily, many debiasing methods have been developed for the Lasso estimator. Since the distribution of the Lasso estimator is not tractable, a direct and standard way to debias the regression estimator is bootstrap ([Efron, 1992](#)). However, [Lahiri \(2010\)](#) demonstrated that the residual bootstrap approximation for Lasso is inconsistent, making it difficult to draw valid inferences directly from the Lasso estimates. To address this issue, [Chatterjee and Lahiri \(2011\)](#) proposed a modified bootstrap Lasso estimator, which mitigates the inconsistency problem by employing a thresholding technique. On the other hand, in the large p small n scenario, [Zhang and Zhang \(2014\)](#) proposed a low-dimension projection estimator (LDP), also known as "de-biased Lasso".

This estimator is an asymptotically Gaussian-distributed estimator of low-dimension parameters in the high-dimensional linear regression model. The LDP approach introduces a correction score designed to address the bias in the Lasso estimator. Additionally, this approach can be used for variable selection, to estimate the entire regression coefficient vector, and to construct confidence intervals to quantify the uncertainty of the estimator.

Following the research line of [Zhang and Zhang \(2014\)](#), [van de Geer et al. \(2014\)](#) approached the problem from a different perspective by starting with the Karush–Kuhn–Tucker (KKT) conditions ([Boyd and Vandenberghe, 2004](#)) for the Lasso estimator. They reformulated these conditions into a form of the Lasso estimator plus a bias term. In this framework, approximating the inverse of the sample covariance matrix became necessary, and the authors successfully addressed this using node-wise Lasso regression ([Meinshausen and Bühlmann, 2006](#)). Furthermore, they extended their debiasing approach to generalized linear models (GLMs) with convex loss functions, providing a broader framework for bias correction. Building on this work, [Xia et al. \(2023\)](#) further addressed the limitations of the sparsity assumption on the inverse of the information matrix in GLM settings. They proposed a refined debiased Lasso estimator designed for the “large n , diverging p ” scenario. On the other hand, [Li \(2020\)](#) took an approach more closely aligned with [Zhang and Zhang \(2014\)](#), proposing a modified correction score to reduce the bias of the modified debiased Lasso estimator through a bootstrap method. A notable advantage of their correction is that it does not rely on the irrepresentable condition, making it particularly useful in high-dimensional settings.

1.2 Proposed method



In the GLM framework, the link function primarily captures the linear relationship between predictor x and response y , neglecting potential nonlinear effects. To address this limitation, this thesis extends LassoNet to GLMs to model more complex structures, while maintaining the ability to obtain sparse solutions in high-dimensional settings, similar to Lasso in GLMs. However, like Lasso in GLMs, this extension also results in a biased solution. To overcome this, we incorporate the debiasing method proposed by [van de Geer et al. \(2014\)](#) into this extension, achieving a debiased LassoNet estimator. Furthermore, we further reduced the bias of this debiased estimator through a bootstrap approach. With these refinements, we obtained a model that is not only predictable but also interpretable and reliable.

1.3 Organization of the thesis

The rest of the thesis is organized as follows. In Chapter [2](#), we introduce some preliminary concepts necessary for understanding the LassoNet model, including an overview of Lasso, the framework of fully connected neural networks, and the bootstrap procedure for estimation debiasing. In Chapter [3](#), we provide a detailed description of the LassoNet model, its extension to GLMs, and the exact debiasing procedure. In Chapter [4](#), we present the empirical results and compare our approach with other existing methods. Finally, we discuss our findings and summarize the main contributions of this work in Chapter [5](#).



Chapter 2 Preliminary

In this chapter, we introduce some preliminary knowledge relevant to the LassoNet model, including Lasso, fully-connected neural networks, and the framework of estimation debiasing with bootstrap. In Section 2.1, we present the ℓ_1 -regularization function, provide intuition to explain why the Lasso optimization leads to sparse solutions, and discuss the associated challenges and limitations of Lasso. In Section 2.2, we describe the framework of neural network models, the process of training these models, and explore some common network-based architectures, such as recurrent neural networks and convolutional neural networks.

2.1 Lasso

As mentioned in the introduction, Lasso is a regularization method designed to address overfitting in models. When considering a linear regression model:

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon}, \quad \boldsymbol{\epsilon} \sim N(\mathbf{0}, \sigma^2 \mathbb{I}), \quad (2.1)$$

where $\mathbf{Y} = (y_1, \dots, y_n) \in \mathbb{R}^n$ is the response vector, $\mathbf{X} = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p] \in \mathbb{R}^{n \times p}$ is the design matrix with column vectors $\mathbf{x}_i$, and $\boldsymbol{\beta} = (\beta_1, \beta_2, \dots, \beta_p) \in \mathbb{R}^p$ represents the unknown regression coefficients. Based on the residual sum of squares criterion, the

Lasso solution is written as:

$$\hat{\beta}^{Lasso} = \arg \min_{\beta} \left\{ \frac{1}{2n} (\mathbf{Y} - \mathbf{X}\beta)^T (\mathbf{Y} - \mathbf{X}\beta) + \lambda \|\beta\|_1 \right\} \quad (2.2)$$

where $\|\cdot\|_1$ denotes the ℓ_1 norm, and λ is a regularization parameter controlling the sparsity of the Lasso estimator. Specifically, when λ increases, more components in $\hat{\beta}^{Lasso}$ become zero; conversely, a smaller λ results in fewer zero components. In extreme cases, when $\lambda = 0$, the Lasso estimator reduces to the ordinary least squares (OLS) solution, whereas as $\lambda \rightarrow \infty$, it converges to a zero estimator. In fact, (2.2) is equivalent to the following Lagrangian form

$$\begin{aligned} \hat{\beta}^{Lasso} = \arg \min_{\beta} \left\{ \frac{1}{2n} (\mathbf{Y} - \mathbf{X}\beta)^T (\mathbf{Y} - \mathbf{X}\beta) \right\} \\ \text{subject to } \|\beta\|_1 \leq t \end{aligned} \quad (2.3)$$

where t is a non-negative tuning parameter. An intuitive way to understand why the Lasso can obtain a sparse solution is as follows. When considering the case $p = 2$ the constraint region for the Lasso is a diamond-shaped region defined by $|\beta_1| + |\beta_2| \leq t$, which has four corners. Meanwhile, the residual sum of squares has elliptical contours. When the solution lies at a corner of the diamond, it results in one of the parameters β_j being zero. As p increases, the diamond generalizes to a higher-dimensional rhomboid with more corners, providing more opportunities for additional β'_j s to become zero.

2.2 Fully-Connected Neural Network

In this section, we introduce one of the main structures in the LassoNet model: the fully-connected neural network (FCNN) (Hastie et al., 2009). A neural network, in



essence, is a regression model constructed by stacking multiple hidden layers. Assume the input, also called the covariate, is $\mathbf{x} = (x_1, \dots, x_p)^T \in \mathbb{R}^p$, and the response is $y \in \mathbb{R}$. Given an activation function $\sigma_d : \mathbb{R} \rightarrow \mathbb{R}$ and the number of neurons in the $(d-1)^{th}$ and d^{th} layers as q_{d-1} and q_d , respectively, the d th layer of a deep FCNN is defined by the mapping $\mathbf{z} : \mathbb{R}^{q_{d-1}} \rightarrow \mathbb{R}^{q_d}$:

$$\mathbf{r} \mapsto \mathbf{z}^{(d)} = (z_1^{(d)}(\mathbf{r}), \dots, z_{q_d}^{(d)}(\mathbf{r}))^T, \quad (2.4)$$

where $z_j^{(d)}(\mathbf{r})$, for $j = 1, \dots, q_d$, are the neurons in the d^{th} layer, and $\mathbf{r} = (r_1, \dots, r_{q_{d-1}})^T \in \mathbb{R}^{q_{d-1}}$ is the output variable from the $(d-1)^{th}$ layer:

$$z_j^{(d)}(\mathbf{r}) = \sigma_d \left(w_{0,j}^{(d)} + \langle \mathbf{w}_j^{(d)}, \mathbf{r} \rangle \right) = \sigma_d \left(w_{0,j}^{(d)} + \sum_{\ell=1}^{q_{d-1}} w_{\ell,j}^{(d)} r_\ell \right), \quad (2.5)$$

where $\mathbf{w}_j^{(d)} = (w_{1,j}^{(d)}, \dots, w_{q_{d-1},j}^{(d)})^T \in \mathbb{R}^{q_{d-1}}$ are the unknown network weights, and $w_{0,j}^{(d)} \in \mathbb{R}$ is the bias term. A depth- D , $D \in \mathbb{N}$, FCNN is obtained by composing D layers as in (2.5). When the input $\mathbf{x}$ is given, a depth- D FCNN can be expressed as:

$$g_{\mathbf{w}}^{(D)}(\mathbf{x}) \equiv (\mathbf{z}^{(D)} \circ \mathbf{z}^{(D-1)} \circ \dots \circ \mathbf{z}^{(1)}) (\mathbf{x}), \quad (2.6)$$

where the D th layer is the output layer, this network structure models the underlying relationship between y and $\mathbf{x}$. Common nonlinear activation functions include the sigmoid function, $\sigma_{\text{sigmoid}}(u) = \frac{1}{1+e^{-u}}$, and the Rectified Linear Unit (ReLU), $\sigma_{\text{ReLU}}(u) = \max(0, u)$, which enhance model complexity and help mitigate underfitting.

Specifically, if the response variable is continuous, there will be only one neuron in the output layer i.e. $q_d = 1$, and the identity function is typically chosen as the activation function for the output layer. On the other hand, when considering a classification problem

with K classes, there will be $q_d = K$ neurons in the output layer. In this case, the network often employs a softmax activation function at the output layer to normalize the outputs into probabilities:

$$\sigma_{\text{softmax}}(u_k) = \frac{e^{u_k}}{\sum_{j=1}^K e^{u_j}}, \quad k = 1, \dots, K, \quad (2.7)$$

where $u_k = w_{0,k}^{(D)} + \langle \mathbf{w}_k^D, \mathbf{z}^{D-1}(\mathbf{x}) \rangle$.

Typically, training the FCNN involves solving an optimization problem to minimize the loss function. A widely used algorithm for optimizing the loss function is gradient descent. For a training set $\{\mathbf{x}_i, y_i\}_{i=1}^n$, $\mathbf{x}_i \in \mathbb{R}^p$, $y_i \in \mathbb{R}$ following the empirical risk minimization criterion and considering a classification problem as an example, we minimize the empirical loss function:

$$L(\mathbf{w}) = -\frac{1}{n} \sum_{i=1}^n (y_i \log(g_{\mathbf{w}}^{(D)}(\mathbf{x}_i))) \quad (2.8)$$

using the gradient descent algorithm outlined in Algorithm 1.

Algorithm 1 Gradient Descent Algorithm

Require: Initial weight $\mathbf{w}^0$, learning rate $\eta > 0$

$k \leftarrow 0$

while $\mathbf{w}^{(k)}$ not covered **do**

$k \leftarrow k + 1$

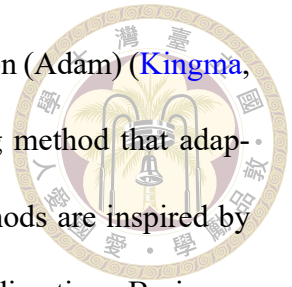
$\Delta \mathbf{w}^{(k)} = -\nabla L(\mathbf{w}^{(k)})$

$\mathbf{w}^{(k+1)} = \mathbf{w}^{(k)} + \eta \Delta \mathbf{w}^{(k)}$

end while

However, due to the high dimensionality of $L(\mathbf{w})$, the gradient may converge slowly or even get trapped in suboptimal regions, making optimization inefficient. To address this

issue, more advanced algorithms, such as Adaptive Moment Estimation (Adam) (Kingma, 2014), have been developed. Adam is a momentum-based learning method that adaptively adjusts the learning rate for each parameter. Momentum methods are inspired by physical analogies, where the velocity determines the momentum's direction. By incorporating a momentum term, the algorithm accumulates past gradients' effects, enabling faster convergence when the gradients consistently point in the same direction. Additionally, momentum helps the optimization escape shallow local minima or flat regions by counteracting oscillations when gradients point in opposing directions. Adam further improves on momentum-based methods by introducing adaptive learning rates for each parameter. It computes moving averages of the first and second moments of the gradient, enabling it to scale the updates dynamically. This adaptive behavior makes Adam particularly effective in high-dimensional and sparse settings, where the scale of gradients can vary significantly. For these reasons, Adam has become a standard choice for training neural networks, including LassoNet.



Algorithm 2 Adam Algorithm

Require: Global learning rate η , decay rates $\rho_1, \rho_2 \in [0, 1)$

Require: Small constant ϵ for numerical stable

Split the dataset into B batches with batch size m

$k \leftarrow 0$

Initialize $\mathbf{w}_B^{(k)}$, $\mathbf{s} = 0$, $\mathbf{r} = 0$

while $\mathbf{w}_B^{(k)}$ not converged **do**

$k \leftarrow k + 1$

for $b = 1, \dots, B$ **do**

$\mathbf{G} = \frac{1}{|B_b|} \nabla \sum_{j \in B_b} \left(y_j \log \left(g_{\mathbf{w}_b^{(k)}}^{(D)}(\mathbf{x}_j) \right) \right)$ ▷ Compute the gradient

$\mathbf{s} = \rho_1 \mathbf{s} + (1 - \rho_1) \mathbf{G}$ ▷ update biased first moment estimate

$\mathbf{r} = \rho_2 \mathbf{r} + (1 - \rho_2) \mathbf{G} \odot \mathbf{G}$ ▷ update biased second moment estimate

$\hat{\mathbf{s}} = \frac{\mathbf{s}}{1 - \rho_1}$ ▷ correct bias in first moment

$\hat{\mathbf{r}} = \frac{\mathbf{r}}{1 - \rho_2}$ ▷ correct bias in second moment

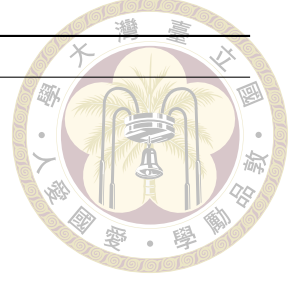
$\Delta \mathbf{w} = -\frac{\eta}{\sqrt{\hat{\mathbf{r}} + \epsilon}} \hat{\mathbf{s}}$

$\mathbf{w}_{b+1}^{(k)} = \mathbf{w}_b^{(k)} + \Delta \mathbf{w}$

end for

end while

Note: $\odot$ denotes element-wise multiplication



To conclude, while we have introduced gradient descent and Adam as two key optimization algorithms for training FCNNs, many other algorithms are available depending on the specific needs of the task and the user's familiarity with the methods. Readers are encouraged to select the algorithm that aligns best with their understanding and the requirements of their application. In addition to FCNNs, other neural network architectures, such as recurrent neural networks (RNNs) (Chen and Li, 2021) and convolutional neural networks (CNNs) (Alzubaidi et al., 2021), are widely used for tasks involving sequential data or image data, respectively. These architectures introduce unique structures and train-

ing procedures tailored to their specific applications. However, since our research does not involve RNNs or CNNs, we do not provide further discussion on these architectures here.







Chapter 3 Method

In Section 3.1, we introduce the structure of the LassoNet model and provide an explanation of how it successfully achieves sparse solutions. In Section 3.2 we provide the debiasing procedure [van de Geer et al. \(2014\)](#) proposed for high-dimensional GLMs. In Section 3.3 we provide the procedure of debiasing the LassoNet model.

3.1 LassoNet

First, we consider a residual feed-forward neural network (FFN):

$$\mathcal{H} = \{h : h_{\beta,w}(\mathbf{x}) = \beta' \mathbf{x} + g_w^{(D)}(\mathbf{x})\},$$

where $g_w^{(D)}$ was defined in Section 2.2, $\beta \in \mathbb{R}^p$ and $\mathbf{x} \in \mathbb{R}^p$. This structure models both the linear and non-linear effects between the predictor $\mathbf{x}$ and the response $y \in \mathbb{R}$ simultaneously. The linear component typically enhances interpretability, while the non-linear part captures complex relationships. In sparse high-dimensional linear models, Lasso ([Tibshirani, 1996](#)) is a commonly used tool for selecting predictive regressors. Similarly, within the FFN framework described above, applying LassoNet ([Lemhadri et al., 2021](#)) enables variable selection and estimation simultaneously. Given n training observations and an input design matrix $\mathbf{X} = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p] \in \mathbb{R}^{n \times p}$, with K neurons in the first hidden

layer, the LassoNet optimization function is defined by

$$\begin{aligned} \min_{\beta, w} \quad & L(\beta, w) + \lambda \|\beta\|_1 \\ \text{subject to} \quad & \|w_k^{(1)}\|_\infty \leq M \|\beta_k\|, \quad k = 1, \dots, p \end{aligned} \quad (3.1)$$



where λ and $M > 0$ are hyper parameters, $L(\beta, w)$ is the loss function, and $w_k^{(1)}$ denotes the weights corresponding to the k -th predictor in the first hidden layer. In fact, based on the constraint above, we can observe that traditional Lasso is a special case of this model. When $M = 0$, the k -th regressor does not pass through the network, reducing the model to standard Lasso. Conversely, as M approaches infinity, the model behaves exactly as a FNN. Thus, M serves as a balancing parameter, adjusting the weights between linearity and nonlinearity within the model.

3.2 Debiasing Lasso Estimators

3.2.1 Node-wise Lasso regression

We present the procedure for debiasing the Lasso estimator proposed by [van de Geer et al. \(2014\)](#). Here, we start with the case of linear regression. Considering the linear regression model in (2.1), the Lasso estimator in (2.2) satisfies the KKT stationarity condition:

$$\frac{-1}{n} \mathbf{X}^T (\mathbf{Y} - \mathbf{X} \hat{\beta}^{Lasso}) + \lambda \hat{\mathbf{s}} = 0, \quad (3.2)$$

where $\|\hat{\mathbf{s}}\|_\infty \leq 1$ and $\hat{s}_j = \text{sign}(\hat{\beta}_j^{Lasso})$ if $\hat{\beta}_j^{Lasso} \neq 0$. We can also represent (3.2) with the notation $\hat{\Sigma} = \mathbf{X}^T \mathbf{X} / n$ as:

$$\hat{\Sigma} (\hat{\beta}^{Lasso} - \beta) + \lambda \hat{\mathbf{s}} = \mathbf{X}^T \epsilon / n \quad (3.3)$$

The main idea in their procedure is to approximate the inverse of $\widehat{\Sigma}$ using node-wise Lasso, denoted as $\widehat{\Omega}$. Incorporating this approximation into (3.3), they derive:

$$\widehat{\beta}^{Lasso} + \widehat{\Omega}\lambda\widehat{s} - \beta = \widehat{\Omega}\mathbf{X}^T\epsilon/n - \delta/\sqrt{n}, \quad (3.4)$$

where $\delta = \sqrt{n} \left(\widehat{\Omega}\widehat{\Sigma} - \mathbb{I} \right) \left(\widehat{\beta} - \beta \right)$. In the same paper, it is proven that δ is asymptotically negligible under certain conditions. Thus, using (3.4) and (3.2) the following debiased estimator is suggested:

$$\widehat{\beta}^{DB-Lasso} = \widehat{\beta}^{Lasso} + \widehat{\Omega}\lambda\widehat{s} = \widehat{\beta}^{Lasso} + \widehat{\Omega}\mathbf{X}^T \left(\mathbf{Y} - \mathbf{X}\widehat{\beta} \right) / n \quad (3.5)$$

To construct $\widehat{\Omega}$ for the debiasing procedure, we rely on node-wise Lasso. This involves treating each covariate as a response variable and regressing it on the remaining covariates using Lasso below, we outline the detailed steps for calculating $\widehat{\Omega}$: For the design matrix $\mathbf{X}$, let $\mathbf{x}_j$ denote the j^{th} column and let $\mathbf{X}_{-j}$ denote the submatrix excluding the j^{th} column.

- Calculate the Lasso regression coefficients for each covariate:

$$\widehat{\alpha}_j \equiv \arg \min_{\alpha \in \mathbb{R}^{p-1}} \left\{ \frac{1}{n} (\mathbf{x}_j - \mathbf{X}_{-j}\alpha)^T (\mathbf{x}_j - \mathbf{X}_{-j}\alpha) + 2\lambda_j \|\alpha\|_1 \right\}, \text{ for } j = 1, \dots, p$$

with components of $\widehat{\alpha}_j = \{\widehat{\alpha}_{j,k} ; k = 1, \dots, p, k \neq j\}$

- $\widehat{\gamma}_j^2 \equiv (\mathbf{x}_j - \mathbf{X}_{-j}\widehat{\alpha}_j)^T (\mathbf{x}_j - \mathbf{X}_{-j}\widehat{\alpha}_j) / n + \lambda_j \|\widehat{\alpha}_j\|_1$

We then define $\widehat{\Omega}$ as:

$$\widehat{\Omega} = \widehat{\Gamma}^{-2} \widehat{A}$$

where

$$\widehat{\Gamma}^2 \equiv \text{diag} (\widehat{\gamma}_1^2, \dots, \widehat{\gamma}_p^2) \quad (3.6)$$

and

$$\hat{A} \equiv \begin{bmatrix} 1 & -\hat{\alpha}_{1,2} & \cdots & -\hat{\alpha}_{1,p} \\ -\hat{\alpha}_{2,1} & 1 & \cdots & -\hat{\alpha}_{2,p} \\ \vdots & \vdots & \ddots & \vdots \\ -\hat{\alpha}_{p,1} & -\hat{\alpha}_{p,2} & \cdots & 1 \end{bmatrix} \quad (3.7)$$



Based on these steps, they successfully construct the debiased Lasso estimator in the case of linear model.

3.2.2 Extensions of node-wise Lasso regression

Having established the Node-wise Lasso procedure for linear regression, we now explore its extension to GLMs. Specifically, the procedure is extended to strictly convex loss functions. For the data set $\{\mathbf{x}_i, y_i\}_{i=1}^n$, $\mathbf{x}_i \in \mathbb{R}^p$, $y_i \in \mathbb{R}$ we let $\rho_{\beta}(\mathbf{x}_i, y_i) = \rho(\beta^T \mathbf{x}_i, y_i)$ be a convex loss function in $\beta \in \mathbb{R}^p$ and define. The first derivative and the Hessian matrix of ρ_{β} are defined as:

$$\rho_{\beta}^{(1)} = \frac{\partial}{\partial \beta} \rho_{\beta}, \quad \rho_{\beta}^{(2)} = \frac{\partial}{\partial \beta \partial \beta^T} \rho_{\beta}$$

Moreover, for a given function g , the empirical loss is defined as:

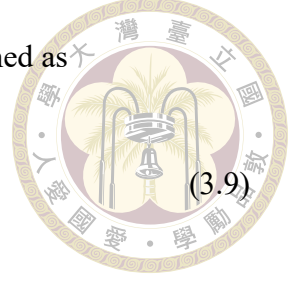
$$L_g(\beta) = \frac{1}{n} \sum_{i=1}^n g(\mathbf{x}_i, y_i)$$

and the corresponding ℓ_1 -regularized estimator is written as:

$$\hat{\beta}^{Lasso} = \arg \min_{\beta} \{L_g(\beta) + \lambda \|\beta\|_1\} \quad (3.8)$$

The general debiased estimator for this extended version is then defined as

$$\hat{\beta}^{DB-Lasso} = \hat{\beta}^{Lasso} - \hat{\Omega} L_{\rho_{\hat{\beta}^{Lasso}}^{(1)}}^{(1)} \left(\hat{\beta}^{Lasso} \right) \quad (3.9)$$



To compute $\hat{\Omega}$ we first define the input matrix as

$$\hat{\Sigma} \equiv L_{\rho_{\hat{\beta}^{Lasso}}^{(2)}}^{(2)} \left(\hat{\beta}^{Lasso} \right)$$

Given this input matrix, the node-wise Lasso procedure is applied as follows. For each $j = 1, \dots, p$ we solve:

$$\hat{\alpha}_j = \arg \min_{\alpha \in \mathbb{R}^{p-1}} \{ \hat{\Sigma}_{j,j} - 2\hat{\Sigma}_{j,\setminus j} \alpha + \alpha^T \hat{\Sigma}_{\setminus j,\setminus j} \alpha + 2\lambda \|\alpha\|_1 \} \quad (3.10)$$

where $\hat{\Sigma}_{j,\setminus j}$ denotes j^{th} row of $\hat{\Sigma}$ excluding the j^{th} element, and $\hat{\Sigma}_{\setminus j,\setminus j}$ is the submatrix of $\hat{\Sigma}$ obtained by removing the j^{th} row and column. Next, we calculate:

$$\hat{\gamma}_j = \hat{\Sigma}_{j,j} - \hat{\Sigma}_{j,\setminus j} \hat{\alpha}_j, \quad j = 1, \dots, p \quad (3.11)$$

to construct $\hat{\Omega} = \hat{\Gamma}^{-2} \hat{A}$, where $\hat{\Gamma}$ and $\hat{A}$ are defined as in (3.6) and (3.7) respectively.

3.3 Debiasing the LassoNet

In this section, we provide further details on how we improve the statistical inference of LassoNet. The debiasing methods discussed above primarily consider GLM, where the link function connects only the linear predictor to the response. However, in many cases, the relationship between predictors and responses is more complex.



3.3.1 Merge LassoNet and GLM

Here we extend the LassoNet to GLM. Specifically, for any GLM with a link function $\ell(\cdot)$ we incorporate a feed-forward network (FFN) g_w to model the mean of the response, such that $E(\mathbf{y}|\mathbf{x}) = \ell^{-1}(\beta'\mathbf{x} + g_w(\mathbf{x}))$. For example, consider logistic regression with the canonical logit link, which can be formulated as:

$$y_i \sim \text{Bin}(1, \text{logit}^{-1}(\beta'\mathbf{x}_i + g_w(\mathbf{x}_i))), \quad i = 1, \dots, n \quad (3.12)$$

In this setting, we can use the optimization function in (3.1) with the negative log-likelihood of (3.12) as the loss function to obtain a sparse model. This approach allows us to retain the interpretability benefits of traditional GLMs while adapting to more complex data structures.

3.3.2 Bootstrapping the Bias of LassoNet

Although a sparse model can be obtained via (3.1), like most regularization methods, suffers from estimation bias. In other words, the estimator produced by LassoNet is not unbiased. Our simulation experiments (4.1) also demonstrate that the LassoNet estimator exhibits a larger bias compared to Lasso. To correct the bias, we adopt a bootstrap strategy inspired by Li (2020), who proposed a method for bootstrapping the bias of a noiseless Lasso estimator. While Li (2020) utilized a noiseless estimator for a standard linear model, we extend this idea to LassoNet. Specifically, we first obtain a noiseless estimator for LassoNet using node-wise regression and then bootstrap the bias of this noiseless estimator. The detailed debiasing procedure is as follows:

Algorithm 3 Procedure for Bootstrap Debiased LassoNet

- 1: Obtain the initial estimator $\hat{\beta}$ by training LassoNet using the optimization function in (3.1).
- 2: Let $\hat{\beta}^{DB-Net}$ denote the noiseless estimator, obtained by following the procedure outlined in Section 3.2.
- 3: Define $\mathbf{X}$ as the design matrix. Resample the rows of $\mathbf{X}$ with replacement B times to create bootstrap design matrices $\mathbf{X}_b^*$, $b = 1, \dots, B$.
- 4: Generate the bootstrap response $\mathbf{Y}_b^*$ from an exponential family distribution denoted as $f_{\mathbf{Y}_b^*|\mathbf{X}_b^*}(\mathbf{y}_b^*|\mathbf{X}_b^*, \hat{\beta}^{DB-Net})$ with canonical link function $\ell(\cdot)$ satisfying $E(\mathbf{Y}_b^*|\mathbf{X}_b^*) = \ell^{-1}(\mathbf{X}_b^* \hat{\beta}^{DB-Net})$.
- 5: Obtain the bootstrap Lasso estimator $\hat{\beta}_b^*$ by minimizing the negative log-likelihood with ℓ_1 penalty:

$$\hat{\beta}_b^* = \arg \min_{\beta} \{-\log\{f_{\mathbf{Y}_b^*|\mathbf{X}_b^*}(\mathbf{y}_b^*|\mathbf{X}_b^*, \beta)\} + \lambda \|\beta\|_1\},$$

where λ is a tuning parameter.

- 6: Estimate the bias $\hat{b}$ of $\hat{\beta}^{DB-Net}$ as

$$\hat{b} = \frac{1}{B} \sum_{b=1}^B (\hat{\beta}_b^* - \hat{\beta}^{DB-Net}),$$

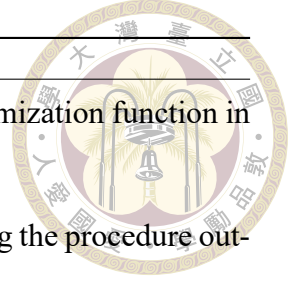
and define the bootstrap debiased LassoNet estimator as

$$\hat{\beta}^{BS-DB-Net} = \hat{\beta}^{DB-Net} - \hat{b}.$$

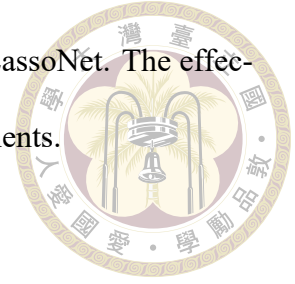
- 7: Construct a two-sided $100 \times (1 - \alpha)\%$ confidence interval for β_j as

$$(q_{\alpha/2}(\hat{\beta}_{(j)}^*), q_{1-\alpha/2}(\hat{\beta}_{(j)}^*)),$$

where $q_{\alpha}(\cdot)$ is the α -quantile, $\hat{\beta}^* = \{\hat{\beta}_1^*, \hat{\beta}_2^*, \dots, \hat{\beta}_B^*\}$, and $\hat{\beta}_{(j)}^*$ represents all values of the j -th component across the elements in $\hat{\beta}^*$.



Based on this procedure, we successfully corrected the bias in LassoNet. The effectiveness of this correction is demonstrated by our simulation experiments.





Chapter 4 Simulation

In this chapter, we demonstrate the empirical performance of our debiasing procedure for the LassoNet model. For this experiment, we consider a logistic regression model as an example to evaluate the effectiveness of bias correction. The sample size is set to $n = 200$, and the number of covariates $p = 50$. We generate y_i 's following the model in 3.12, where each x_i is sampled from $\mathcal{N}(0, 1)$ for $i = 1, \dots, n$, and g_w is a FCNN with 10 layers and 50 nodes per layer. Let β_s denote the non-zero coefficients. We set the true coefficient vector as $\beta = (B, \mathbf{0}_{p-6})^T$ where $B = (1.5, 0.5, 1.5, 0.5, 1.5, 0.5)^T$ represents the true signal, and $\mathbf{0}_k$ is a zero vector of length k . This setting demonstrates a weak sparsity level, as well as a smaller contrast between strong signals (1.5) and weak signals (0.5). For this synthetic data, we select the regularization parameter λ by training the Lasso model using 10-fold cross-validation, rather than tuning it in every repetition. This decision is based on findings from Lemhadri et al. (2021), which demonstrated in their experiments that the regularization parameter λ in the LassoNet model exhibits the same property as in Lasso; specifically, as λ increases, the shrinkage effect becomes stronger. However, the effect of M on bias remains unclear. Thus, in this experiment, we focus solely on investigating the effect of different values of M on the debiasing performance. Additionally, to prevent overfitting, we train the model using a single layer with 10 nodes.

To evaluate the validity of the debiasing performance with bootstrapping, we com-

pare the performances of the following estimators: the original Lasso estimator, the LassoNet estimator, the debiased LassoNet estimator (DB-Net), the un-debiased bootstrap LassoNet estimator (BS-Net), where the LassoNet estimator is bootstrapped directly, and the debiased bootstrap LassoNet estimator (BS-DB-Net). The following reports are based on 50 repetitions.

4.1 Evaluation of estimation bias

First, we present the overall loss for all estimators, measured as $\|\hat{\beta} - \beta\|_2$. As shown in Figure 4.1

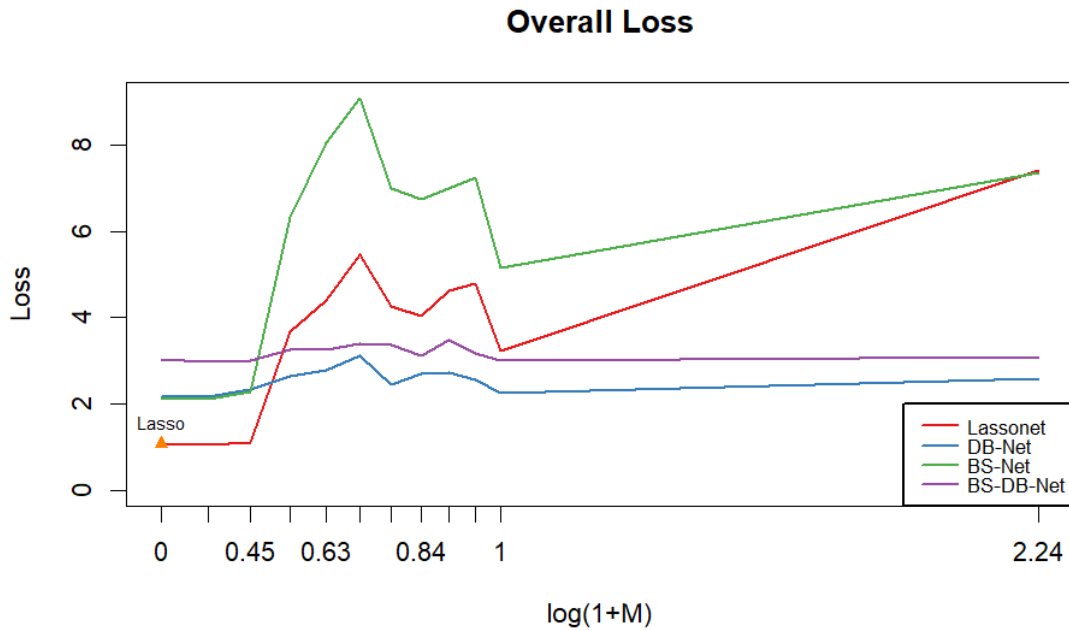


Figure 4.1: Comparison of the overall loss, measured as $\|\hat{\beta} - \beta\|_2$ across Lasso, the LassoNet, the debiased LassoNet estimator (DB-Net), the un-debiased bootstrap LassoNet (BS-Net), and the debiased bootstrap LassoNet (BS-DB-Net).

it is evident that as M increases, the bias of the LassoNet model becomes larger, especially in the weak non-linearity region ($M < 1$), where the bias increases dramatically.

Meanwhile, the BS-Net also exhibits a large bias, as it is bootstrapped from an estimator with significant bias. This observation highlights that directly bootstrapping the LassoNet estimator cannot achieve good debiasing performance. Instead, our two-step debiasing procedure is necessary to address this issue effectively. Moreover, we observe that BS-Net and BS-DB-Net are competitive in terms of overall loss; however, their bias remains larger than that of the Lasso estimator. Nonetheless, as shown in Figure 4.2

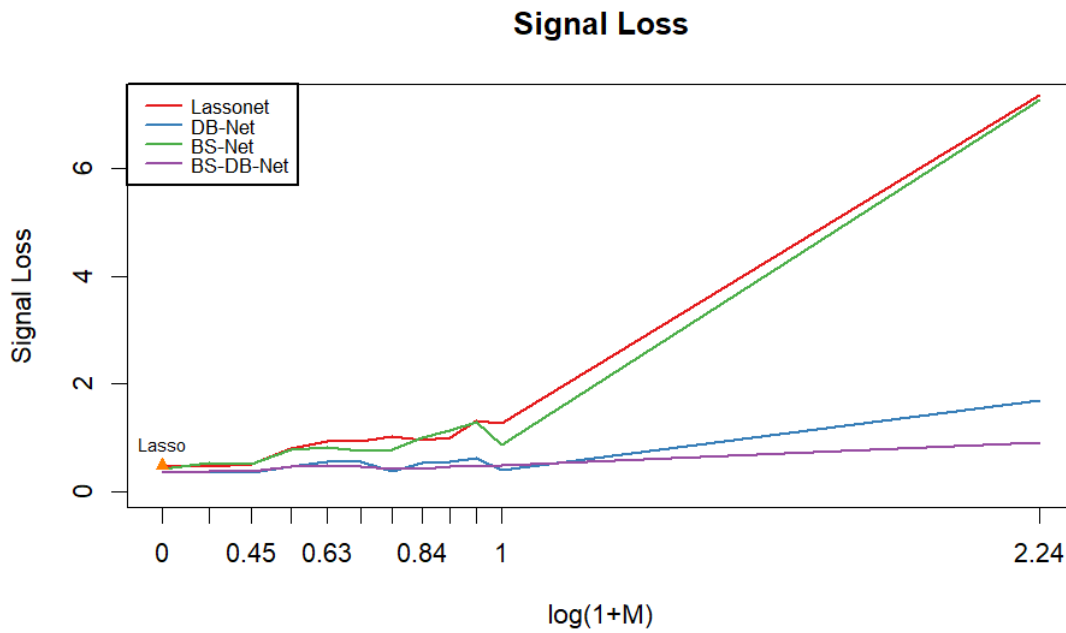
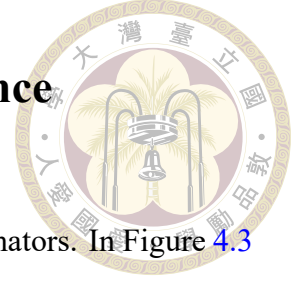


Figure 4.2: Comparison of the signal loss, measured as $\|\hat{\mathbf{B}} - \mathbf{B}\|_2$ across Lasso, the LassoNet, the debiased LassoNet estimator (DB-Net), the un-debiased bootstrap LassoNet (BS-Net), and the debiased bootstrap LassoNet (BS-DB-Net).

when focusing solely on the signal, we observe that as the non-linearity becomes stronger, the bias of LassoNet and BS-Net increases at a faster rate. Meanwhile, we also observe that BS-DB-Net outperforms DB-Net in cases of strong non-linearity.

4.2 Evaluating variable selection performance



Next, we evaluate the variable selection performance of all estimators. In Figure 4.3

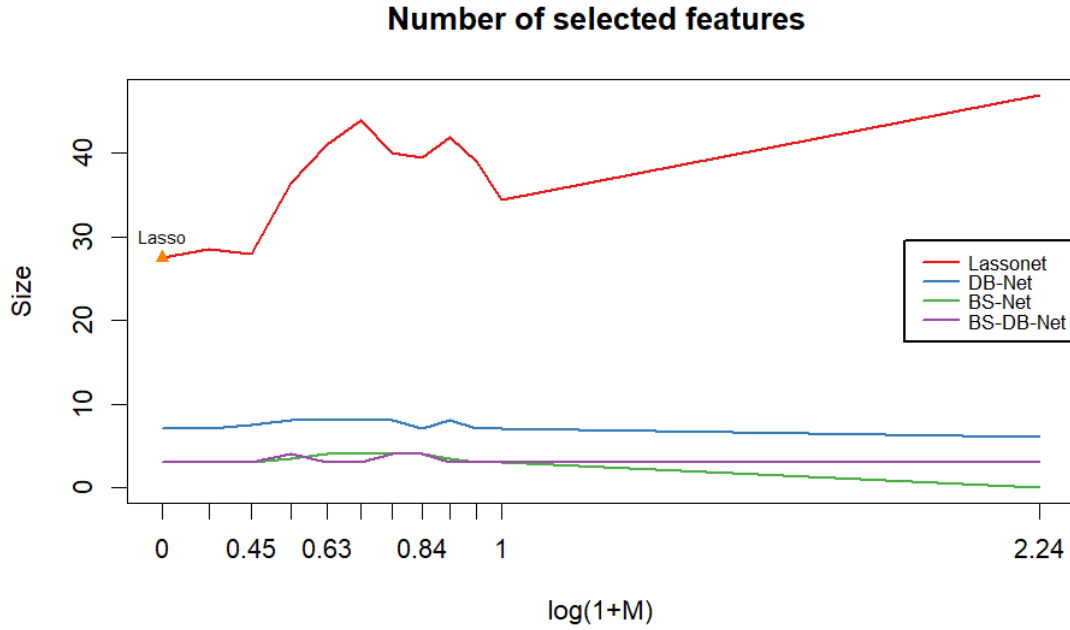


Figure 4.3: Comparison of the number of selections among Lasso, the debiased LassoNet estimator (DB-Net), the un-debiased bootstrap LassoNet estimator (BS-Net), and the debiased bootstrap LassoNet estimator (BS-DB-Net).

we compare the number of selected variables among the four estimators. Note that for DB-Net, BS-Net, and BS-DB-Net, the coefficients are not exactly zero. Variables are selected based on whether 0 falls within the 95% confidence interval defined in step 7 of Algorithm 3. Specifically, if 0 is within the interval, the variable is not selected; otherwise, it is selected. As shown in Figure 4.3, Lasso and LassoNet tend to select too many variables, especially under strong non-linearity, where LassoNet almost selects all variables. In contrast, BS-DB-Net, DB-Net, and BS-Net successfully achieve sparse solutions. However, Figures 4.4 and 4.5

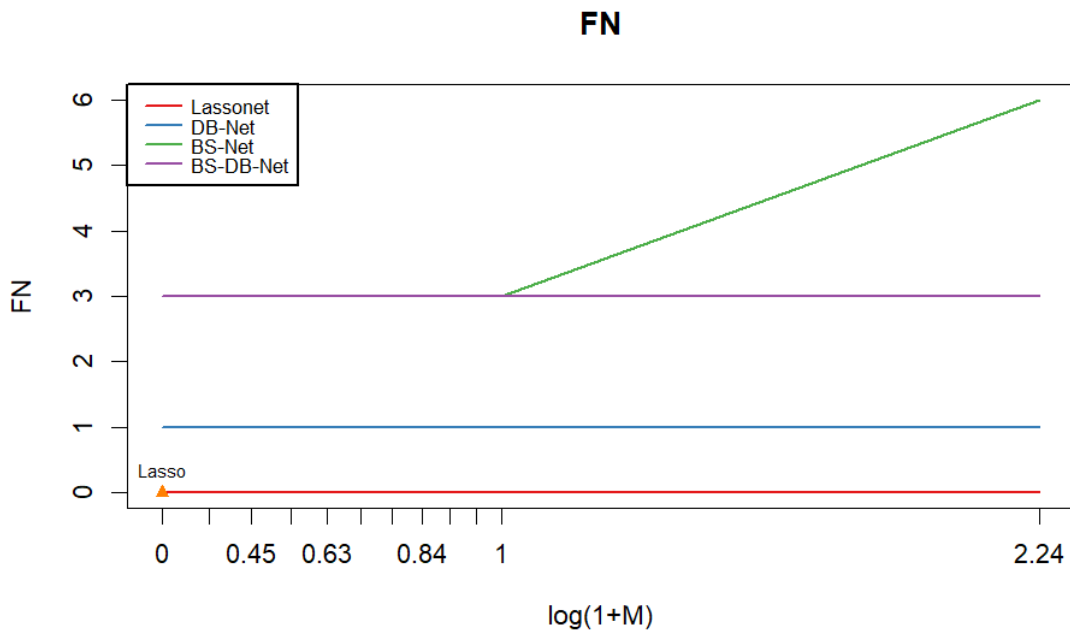


Figure 4.4: False negatives for the debiased LassoNet estimator (DB-Net), the un-debiased bootstrap LassoNet estimator (BS-Net), and the debiased bootstrap LassoNet estimator (BS-DB-Net).

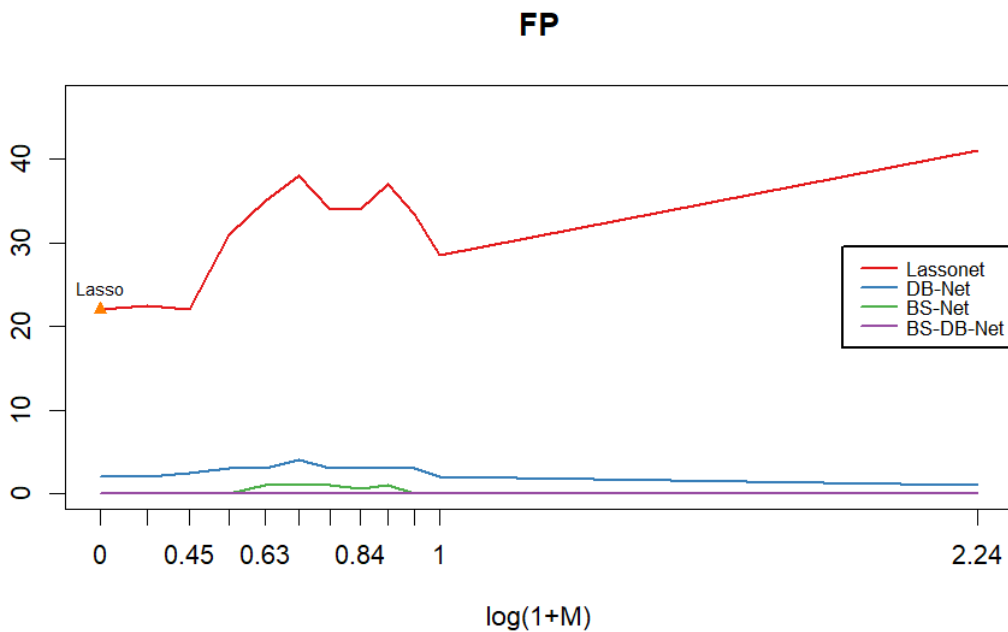
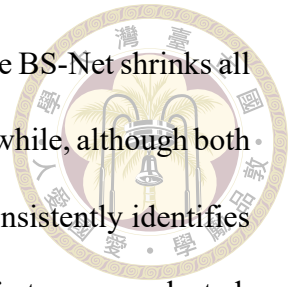


Figure 4.5: False positives for the debiased LassoNet estimator (DB-Net), the un-debiased bootstrap LassoNet estimator (BS-Net), and the debiased bootstrap LassoNet estimator (BS-DB-Net).

reveal that BS-Net exhibits a higher rate of false negatives compared to the other two

estimators, particularly under strong non-linearity. This occurs because BS-Net shrinks all coefficients toward zero, making it unable to detect true signals. Meanwhile, although both BS-DB-Net and DB-Net fail to detect all true signals, BS-DB-Net consistently identifies all strong signals with zero false positives, ensuring that no irrelevant features are selected. In contrast, DB-Net selects both incorrect signals and irrelevant features, compromising the model's interpretability.





Chapter 5 Conclusion

In this thesis, we extend the LassoNet model to GLMs, enabling it to capture both linear and nonlinear relationships between the response and covariates. Additionally, we address the inherent bias in Lasso-type models caused by the ℓ_1 -regularization penalty. To mitigate the bias in the extended version, we incorporate node-wise Lasso regression and bootstrap techniques into the debiasing process. This debiasing procedure allows the extended model not only to retain the advantages of Lasso-type models, namely, the ability to select relevant features but also to correct the shrinkage in the estimated coefficients. Our numerical studies demonstrate that the BS-DB-Net procedure achieves less signal loss compared to four alternative estimators when the relationship between the input and output is weakly linear. Furthermore, our procedure exhibits superior performance in variable selection by accurately identifying relatively strong signals and avoiding the selection of irrelevant features. While the current procedure performs effectively on synthetic data, reducing the computational cost of training the neural network and the bootstrap process could open avenues for applying this method to high-dimensional settings, including ultra-high-dimensional scenarios.





References

- Alzubaidi, L., Zhang, J., Humaidi, A. J., Al-Dujaili, A., Duan, Y., Al-Shamma, O., Santamaria, J., Fadhel, M. A., Al-Amidie, M., and Farhan, L. (2021). Review of deep learning: concepts, cnn architectures, challenges, applications, future directions. Journal of big Data, 8:1–74.
- Andrews, D. W. and Lu, B. (2001). Consistent model and moment selection procedures for gmm estimation with application to dynamic panel data models. Journal of Econometrics, 101(1):123–164.
- Boyd, S. and Vandenberghe, L. (2004). Convex optimization. Cambridge university press.
- Candes, E. and Tao, T. (2007). The dantzig selector: Statistical estimation when p is much larger than n . The Annals of Statistics, 35(6):2313–2351.
- Chatterjee, A. and Lahiri, S. N. (2011). Bootstrapping lasso estimators. Journal of the American Statistical Association, 106(494):608–625.
- Chen, Y. and Li, J. (2021). Recurrent neural networks algorithms and applications. In 2021 2nd International Conference on Big Data & Artificial Intelligence & Software Engineering (ICBASE), pages 38–43. IEEE.

Efron, B. (1992). Bootstrap methods: another look at the jackknife. In Breakthroughs in statistics: Methodology and distribution, pages 569–593. Springer.

Feng, J. and Simon, N. (2019). Sparse-input neural networks for high-dimensional non-parametric regression and classification.

Hastie, T., Tibshirani, R., and Friedman, J. (2009). The elements of statistical learning: data mining, inference and prediction. Springer, 2 edition.

He, K., Xu, H., and Kang, J. (2019). A selective overview of feature screening methods with applications to neuroimaging data. WIREs Computational Statistics, 11(2):e1454.

Kingma, D. P. (2014). Adam: A method for stochastic optimization. arXiv preprint arXiv:1412.6980.

Lahiri, S. (2010). Asymptotic properties of the residual bootstrap for lasso estimators. Proceedings of the American Mathematical Society, 138(12):4497–4509.

Lemhadri, I., Ruan, F., Abraham, L., and Tibshirani, R. (2021). Lassonet: A neural network with feature sparsity. Journal of Machine Learning Research, 22(127):1–29.

Li, S. (2020). Debiasing the debiased lasso with bootstrap. Electronic Journal of Statistics, 14(1):2298–2337.

Liu, B., Zhang, Q., Xue, L., Song, P. X. K., and Kang, J. (2024). Robust high-dimensional regression with coefficient thresholding and its application to imaging data analysis. Journal of the American Statistical Association, 119(545):715–729.

Meinshausen, N. and Bühlmann, P. (2006). High-dimensional graphs and variable selection with the Lasso. The Annals of Statistics, 34(3):1436 – 1462.

Tibshirani, R. (1996). Regression shrinkage and selection via the lasso. Journal of the Royal Statistical Society Series B: Statistical Methodology, 58(1):267–288.

van de Geer, S., Bühlmann, P., Ritov, Y., and Dezeure, R. (2014). On asymptotically optimal confidence regions and tests for high-dimensional models. The Annals of Statistics, 42(3):1166 – 1202.

Xia, L., Nan, B., and Li, Y. (2023). Debiased lasso for generalized linear models with a diverging number of covariates. Biometrics, 79(1):344–357.

Zhang, C.-H. and Zhang, S. S. (2014). Confidence intervals for low dimensional parameters in high dimensional linear models. Journal of the Royal Statistical Society Series B: Statistical Methodology, 76(1):217–242.

Zhao, P. and Yu, B. (2006). On model selection consistency of lasso. The Journal of Machine Learning Research, 7:2541–2563.

Zou, H. (2006). The adaptive lasso and its oracle properties. Journal of the American Statistical Association, 101(476):1418–1429.

Zou, H. and Hastie, T. (2005). Regularization and variable selection via the elastic net. Journal of the Royal Statistical Society Series B: Statistical Methodology, 67(2):301–320.