

High dimensional generalized linear models for temporal dependent data

YUEFENG HAN^{1,a}, RUEY S. TSAY^{2,b} and WEI BIAO WU^{3,c}

¹*Department of Statistics, Rutgers University, Piscataway, NJ 08854, USA.* yuefeng.han@rutgers.edu

²*Booth School of Business, University of Chicago, Chicago, IL 60637, USA.* ruey.tsay@chicagobooth.edu

³*Department of Statistics, University of Chicago, Chicago, IL 60637, USA.* wbwu@galton.uchicago.edu

High dimensional generalized linear models are widely applicable in many scientific fields with data having heavy tails. However, little is known about statistical guarantees on the estimates of such models in a time series setting. In this article, we establish statistical error bounds and support recovery guarantees of the classical ℓ_1 regularized procedure for generalized linear model with temporal dependent data. We also propose a new robust M -estimator for high dimensional time series. Properties of the proposed robust procedure are investigated both theoretically and numerically. As an extension, we introduce a robust estimator for linear regression and show that the proposed robust estimator achieves nearly the optimal rate as that for i.i.d sub-Gaussian data. Simulation results show that the proposed method performs well numerically in the presence of heavy-tailed and serially dependent covariates and/or errors, and it significantly outperforms the classical Lasso method. For applications, we demonstrate, in the supplementary material, the regularized robust procedure via analyzing high-frequency trading data in finance.

Keywords: High dimensional analysis; time series analysis; generalized linear model; robust estimation; support recovery

1. Introduction

In recent years, information technology has made high dimensional time series data increasingly common. The demand for modelling and forecasting such data arises naturally from communication engineering, environmental studies, market analysis in finance and panel studies in economics, among others. In many applications, we often face the challenge of dealing with a large number of complicated issues such as missing values or heavy tails. The Lasso regularized method, originally introduced by [66] and subsequently investigated by many others, is a popular technique for high dimensional linear regression models with sparse coefficients. As a matter of fact, the ℓ_1 -type penalty of the Lasso can also be applied to other models in high dimension, including, for example, logistic regression ([21,43,62,64], among others), multinomial logistic regression [39] or Cox regression [67] by replacing the ℓ_2 loss function by the corresponding negative log-likelihood function.

Generalized linear models (GLM, [47]) are a flexible generalization of the ordinary linear regression by allowing researchers to model the relationship between the predictors and a function of the mean of the response variable, which can follow a continuous or discrete distribution. In a variety of applications, the observed response consists of binary or count data for which GLM is especially useful. See, for example, social network analysis [5,40,58,65,85], biological neural networks [7,15,54], compressed sensing [38,56,57], power systems analysis [34], and seismology [53,74]. This paper deals with the Lasso penalty for GLM applied to high dimensional time series data. Under the independent and identically distributed (i.i.d.) setting, there exists substantial literature on the Lasso methods for high dimensional GLM. For instance, [73] showed non-asymptotic oracle inequalities for the empirical risk minimizer with Lasso penalty for high dimensional GLMs with a Lipschitz loss function. [32] studied Lasso estimator in the Cox proportional hazards regression when the covariates are time dependent,

and established oracle inequalities for prediction and estimation errors. A number of papers analyzed penalized methods beyond Lasso. [48] applied group Lasso to high dimensional logistic regression, proposed an efficient algorithm, and showed consistency of the estimator. [52] studied penalized M -estimators with a general class of regularization methods, including an ℓ_2 error bound for the Lasso in GLM. [33] studied weighted absolute penalty and its adaptive, multistage application in GLM. [20] investigated asymptotic equivalence of Lasso and other concave regularized methods in a thresholded parameter space. [37] studied adaptive Lasso and group-Lasso for the functional Poisson regression.

Despite the extensive research in GLMs for i.i.d data, very limited work focused on theoretical properties of the regularized estimates when the observations are dependent. [2] investigated theoretical properties of Lasso estimators with a random design for high dimensional Gaussian processes. [76] analyzed Lasso estimator with a fixed design matrix and Dantzig selector under random design. [27] extended the Lasso estimator to random design and weakly sparse time series with application in now-casting. [84] considered non-parametric sparse additive model for linear autoregression. [23] studied Lasso estimators of high dimensional autoregressive generalized linear models, which was further extended in [45]. See also [22,26].

The phenomenon of heavy-tails is widely observed in time series data. It is one of the stylized facts in financial econometrics that financial returns and macroeconomic variables have high excess kurtosis. Large scale imaging data in biology, such as neural spike recordings (see, for example, [7,15,54]), are often corrupted by non-Gaussian noises. The conventional regression estimator may fare poorly or even be inconsistent when the observations are heavy tailed and/or contaminated by outliers in the predictors and/or the response variable. Therefore, it is important to study effective principles for dealing with heavy-tailed or noisy time series data.

The origin of robust statistics dates back to the fundamental works of John Tukey [70,71], Peter Huber [35,36] and Frank Hampel [24,25]. In general, robustness can be defined in two ways; model misspecification and outliers. For example, Tukey's work [70] is about robustness to a misspecification of the Gaussian model, while Hodges' work [30] is robustness to contamination of the dataset by extreme outliers or robustness to heavy-tailed distributions in the model that lead to the appearance of some aberrant data. It is well known that if the covariates and/or the errors deviate more wildly from the sub-Gaussian distribution, the linear regression estimator based on the least squares loss no longer converges at the optimal rates. Intuitively, an outlier in the covariates may cause the corresponding M -estimator to behave arbitrarily badly. This motivates the use of generalized M -estimators that downweight high-leverage observations. In the classical theory of robust regression in low dimensions, many weighting functions are introduced, such as Mallows estimator [44], Hill-Ryan estimator [29], and Schweppe estimator [50]. In this paper, we focus on heavy-tailed covariates and heavy-tailed errors for generalized linear model. We also extend the robust M -estimator to high dimensional time series.

Driven by a wide range of contemporary scientific applications, robust regression of high dimensional data is of substantial research interest. Indeed, several papers have shed new light on high dimensional robust M -estimator when the population distribution is heavy tailed or noisy. [11] considered estimation of the mean of heavy-tailed distributions via a robust empirical loss, which is insensitive to extreme values. Cantoni's mean estimator is further extended in [8] to empirical risk minimization. [31,51] applied "median of mean" estimator to high dimensional sparse regression. [19] introduces a simple principle for robust high dimensional low rank matrix recovery via an appropriate shrinkage on the data. [17] developed estimation bounds for penalized robust regression with the Huber loss function. [41] gave a general framework for robust regularized M -estimators under both convex and non-convex loss functions. However, all prior works focused on the setting where samples are i.i.d. To the best of our knowledge, the existing procedures cannot readily be applied to high dimensional time series data.

The goal of this paper is threefold: (i) To lay a theoretical foundation for high dimensional general-

ized linear models (GLMs) in situations in which the errors or the covariates can be serially dependent; (ii) To develop novel robust estimation of GLMs for serially dependent data in the case that the dimension can be much larger than the sample size, and to provide a solid theoretical guarantee; (iii) To derive sharp inequalities for tail probabilities for dependent and/or non-sub-Gaussian processes under some mild and easily verifiable conditions. It is worth emphasizing that our model is different from the autoregressive generalized linear model considered in [23] and [45]. It is expected that our framework, inequalities and tools will be useful in other high dimensional problems that involve temporal dependent data.

In this paper, we propose to appropriately shrink the feature variables before calculating the M -estimator to achieve the robustness for high dimensional time series regression. Let X_i be a p -dimensional vector of covariates. If X_i is heavy-tailed, the basic idea is to shrink each feature X_{ij} ($1 \leq j \leq p$) to a predetermined threshold level τ . We show that the regularized robust regression functions continue to enjoy good behavior. Our first contribution is to provide the asymptotic behavior of the estimated GLM coefficients of the Lasso penalized method for both the original time series data and shrinkage heavy-tailed data. It is shown that an appropriate truncation does not induce significant bias. Under only bounded moment conditions for either noise or covariates, our robust estimator can nearly achieve the error bound for i.i.d. sub-Gaussian data, modulo a price for temporal dependence. The performance of the proposed shrinkage Lasso estimator is thus shown to be much better than that of the vanilla estimator. The allowed dimension p can be as large as $\exp(n^c)$, where n is the sample size and $0 < c < 1$. This means that shrinkage not only overcomes heavy-tailed corruption, but also mitigates the curse of dimensionality. We also establish support recovery guarantees and ℓ_∞ bounds for both methods. Furthermore, unlike the usual robust quasi-likelihood estimators in low dimension, which is non-convex, our method still maintains convexity, thus enjoys certain computational advantages.

In addition, our robust estimator can also be applied to the usual linear regression setting for high dimensional time series. For weakly temporal dependence and heavy-tailed data, our robust method achieves nearly the minimax optimal rate of ℓ_2 -norm established by [59] for i.i.d sub-Gaussian data. The difference lies in the scaling requirements on p , n , the sparsity condition, and the additional logarithmic factor of n in the rate, which is induced by the temporal dependence. It is also worth noting that we provide new concentration inequalities, which extend Talagrand's inequality and Bousquet's inequality [6] to high dimensional time series. The extension is of independent interest.

Besides the theoretical properties, we also study the numerical performance of the proposed robust procedure using both simulated and real data. Section 5 considers the simulation studies and shows that our robust procedure performs well numerically in the presence of both symmetric and asymmetric heavy-tailed covariates and/or errors. In particular, the robust procedure significantly outperforms the standard Lasso method, especially in ℓ_1 and ℓ_2 losses of the GLM coefficients. We also illustrate our procedure with an application to high-frequency stock trading for predicting price changes in consecutive transactions via a multinomial logistic regression. Our method leads to marked improvements in prediction compared with the existing methods in financial econometrics.

The rest of the article is organized as follows. Section 2 introduces the standard Lasso procedure and the robust procedure for GLM when the covariates or/and the errors have heavy tails. The framework of high dimensional time series is presented in Section 3.1. Theoretical properties of both robust and non-robust GLM estimators are also investigated in Section 3. After the basic assumptions of Section 3.2, Sections 3.3 and 3.4 study the convergence rates of the standard Lasso procedure and the robust procedure, respectively. Section 4 discusses the conventional linear regression with robust estimator and time series data. Section 5 investigates the numerical performance of the proposed robust procedure and compares it with that of the standard Lasso procedure. Concentration inequalities for high dimensional time series, real data analysis, and all the proofs are given in the supplementary material [28].

1.1. Notation

Throughout the paper, for a vector $x = (x_1, \dots, x_p)'$, define $|x|_2 = (x_1^2 + \dots + x_p^2)^{1/2}$, $|x|_\infty = \max\{|x_1|, \dots, |x_p|\}$. For a matrix $A = (a_{ij}) \in \mathbb{R}^{p \times m}$, denote entrywise max norm $|A|_\infty = \max_{i,j} |a_{ij}|$ and induced matrix-operator norm $\|A\|_1 = \max_{1 \leq j \leq m} \sum_{i=1}^p |a_{ij}|$, $\|A\|_\infty = \max_{1 \leq i \leq p} \sum_{j=1}^m |a_{ij}|$. If $A \in \mathbb{R}^{p \times p}$ is a square matrix, let $\lambda_1(A) \geq \lambda_2(A) \geq \dots \geq \lambda_p(A)$ be its eigenvalues in descending order. Also denote $\lambda_{\max}(A) = \lambda_1(A)$ and $\lambda_{\min}(A) = \lambda_p(A)$. For a random variable ξ , let $\|\xi\|_m = (\mathbb{E}|\xi|^m)^{1/m}$ and write $\xi \in \mathcal{L}^m$, $m \geq 1$, if $\|\xi\|_m < \infty$. For simplicity, denote $\|\xi\| = \|\xi\|_2$. For a set S , write $|S|$ as its cardinality. For two sequences of real numbers $\{a_n\}$ and $\{b_n\}$, write $a_n = O(b_n)$ if there exists a constant C such that $|a_n| \leq C|b_n|$ holds for all sufficiently large n , write $a_n = o(b_n)$ if $\lim_{n \rightarrow \infty} a_n/b_n = 0$, and write $a_n \asymp b_n$ if there are positive constants c and C such that $c \leq a_n/b_n \leq C$ for all sufficiently large n . Denote $a \wedge b = \min\{a, b\}$ and $a \vee b = \max\{a, b\}$. We use $C, C_1, C_2, \dots$ to denote positive constants whose values may differ from place to place. A constant with a symbolic subscript is used to emphasize the dependence of the value on the subscript. We assume $p = p_n \rightarrow \infty$ as $n \rightarrow \infty$.

2. The model

2.1. Generalized linear models and the loss function

Consider n observations $\{(X_i, Y_i)\}_{i=1}^n$, where $X_i \in \mathcal{X} \subset \mathbb{R}^p$ is a p -dimensional vector of covariate variables, and $Y_i \in \mathcal{Y} \subset \mathbb{R}$ is the response variable. We model the dependence of the mean of Y_i on X_i via the linear function $f_{\beta^*}(X_i) = X_i^\top \beta^*$, where β^* is a vector of unknown coefficients and $X_i^\top$ is the transpose of X_i . The goal is to estimate β^* . In a high dimensional model, the number of covariates p can be much larger than the number of observations n . Let $R: \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ be a loss function.

We consider the following estimator of empirical risk minimization with Lasso penalty

$$\hat{\beta} := \arg \min_{\beta} \left\{ \frac{1}{n} \sum_{i=1}^n R(f_{\beta}(X_i), Y_i) + \lambda |\beta|_1 \right\}, \quad (1)$$

where $f_{\beta}(X_i) = X_i^\top \beta$. Assume the response variable Y_i is from an exponential family with the probability density function taking the canonical form

$$h_Y(y; \mu) = \exp[y\mu - r(\mu) + b(y)]$$

for some known functions $r(\cdot)$ and $b(\cdot)$, and unknown function μ . The function μ is usually called the canonical or natural parameter. The mean response is $r'(\mu)$, the first derivative of $r(\mu)$ with respect to μ . The generalized linear model assumes the form:

$$\mathbb{E}(Y|X) = r'(\mu(X)) = r'(X^\top \beta^*).$$

The canonical link function is thus defined as $g := (r')^{-1}$. Let $z = \mu(X)$. The maximum (marginal) log-likelihood loss function is then

$$R(z, y) = -yz + r(z), \quad y \in \mathcal{Y}, \quad z \in \mathbb{R}. \quad (2)$$

Define the objective empirical risk function as

$$\mathcal{R}_n(\beta) = \frac{1}{n} \sum_{i=1}^n R(f_{\beta}(X_i), Y_i).$$

The gradient and Hessian of $\mathcal{R}_n(\beta)$ are respectively

$$\nabla \mathcal{R}_n(\beta) = \frac{1}{n} \sum_{i=1}^n (r'(f_\beta(X_i)) - Y_i) X_i, \quad (3)$$

$$\mathbf{H}_n(\beta) = \nabla^2 \mathcal{R}_n(\beta) = \frac{1}{n} \sum_{i=1}^n r''(f_\beta(X_i)) X_i X_i^\top. \quad (4)$$

We also define the symmetric Bregman divergence as

$$D_{\mathcal{R}}(\beta, \beta^*) = (\beta - \beta^*)^\top (\nabla \mathcal{R}_n(\beta) - \nabla \mathcal{R}_n(\beta^*)), \quad (5)$$

which can be viewed as symmetric, partial Kullback-Leibler distance between the log-likelihood at β and β^* .

In this paper, we focus on $\hat{\beta}$ of the empirical risk minimization of (1) and (2). Extension to quasi-likelihood loss will be discussed in the supplementary material [28].

2.2. Robust Lasso estimator

Inspired by the theory on robust estimation for linear regression (see e.g. [19]), we study regularized versions of high dimensional robust GLM estimators and establish their statistical guarantees. In order to deal with heavy-tailed data, we propose a robust estimator to be used in (1) by the simple and classical principle of truncation [19], or more generally shrinkage. Our approach is simple: we truncate or shrink appropriately the heavy-tailed covariates or/and the response variable. Intuitively, shrinkage reduces sensitivity of the estimator to data corruption caused by the heavy-tailed distributions. However, shrinkage leads to bias. We shall find an appropriate shrinkage level to balance the induced bias and the statistical error rate. The resulting estimator is then defined as follows:

$$\hat{\beta} := \arg \min_{\beta} \left\{ \frac{1}{n} \sum_{i=1}^n R_\tau(f_\beta(X_i), Y_i) + \lambda |\beta|_1 \right\}, \quad (6)$$

where τ is a predetermined threshold level,

$$R_\tau(f_\beta(X_i), Y_i) := R(f_\beta(\tilde{X}_i), Y_i),$$

and $\tilde{X}_i$ is a truncated version of the covariates X_i if they are heavy-tailed and equals the original covariates (truncation threshold goes to infinity) if they are light-tailed. When the covariates X_i are heavy-tailed, we choose

$$\tilde{X}_{ij} = \text{sgn}(X_{ij})(|X_{ij}| \wedge \tau), \quad 1 \leq j \leq p,$$

where $a \wedge b = \min(a, b)$. In practice, we may need to first normalize X_i for each covariate X_{ij} , for $1 \leq j \leq p$, before applying truncation. As discussed in [27], such normalization will not lose signal information when the time series is stationary.

In general, under the GLM setting, the response (e.g., binary data) is considered not to be corrupted by heavy tailed noise. Thus, in Section 3, we focus on the case when the data generating distribution of the response is indeed the model distribution so that we only need to trim the covariates. If the response is also corrupted by random noises, we shall also truncate the response properly. For the linear regression model with least squares loss in Section 4, it is common in the literature to consider

the case that the response might have heavy tails. Then we need to further truncate the response to achieve robustness.

With the aforementioned data robustifications, the proposed methodology yields an estimator that, under a bounded moment condition on the covariates or/and the response, has the similar statistical rate as that of the estimator available in the literature for sub-Gaussian distributions. Our study gives a formal theoretical consideration of both the original estimator in (1) and the robust one in (6). As shown in Sections 3 and 4, compared with [19], our rates are nearly optimal up to an additional multiplicative $\log n$ term.

The first and most important advantage of our robust method is to maintain the convexity of (negative) log-likelihood loss. There are several alternatives to robust estimation in the context of low dimensional generalized linear model. For example, [10] proposed a class of robust quasi-likelihood loss function. [3] constructed robust estimator for the logistic regression by bounded deviance, which was further extended to other generalized linear models. [81] introduced a class of robust estimators for generalized linear models motivated by the Bregman divergence. However, most of these robust estimators in the low dimensional case are non-convex.

Our robust method is also much easier to implement than many existing ones, as it only needs to truncate or shrink the data before applying the standard method to the transformed data. The tuning parameter τ plays a key role by adapting to covariates and/or errors with different shapes and tails. In practice, the optimal values of tuning parameters τ and λ can be chosen by a two-dimensional grid search using an information-based criterion or time series cross-validation, e.g., the Akaike information criterion or Bayesian information criterion. Specifically, we may partition a rectangle in the scale of $(\log(\tau), \log(\lambda))$ to form our search grid. Then the optimal values are achieved by the combination of the two parameters that minimizes the cross-validated measurement, the Akaike information criterion or Bayesian information criterion. In addition, the robust cross-validation proposed in [19] can also be used to tune (λ, τ) by replacing the K -fold cross-validation with time series rolling forecast.

3. Asymptotic properties

We now consider the properties of the standard Lasso method (1) and the robust Lasso method (6). We first show the statistical error rates and the support recovery guarantees of the estimated GLM coefficients for the original time series data and then demonstrate that the convergence rates of robust estimator for shrinkage heavy-tailed data are significantly improved.

3.1. High dimensional time series

Let $\varepsilon_i, i \in \mathbb{Z}$, be i.i.d. random vectors and $\mathcal{F}_i = \sigma(\dots, \varepsilon_{i-1}, \varepsilon_i)$. We assume that the covariate process $(X_i, i = 1, \dots, n)$ is high dimensional and stationary in the form

$$X_i = (g_1(\mathcal{F}_i), \dots, g_p(\mathcal{F}_i))^\top, \quad (7)$$

and the response Y_i assumes the form

$$Y_i = g_y(\mathcal{F}_i), \quad (8)$$

where $g_1(\cdot), \dots, g_p(\cdot)$ and $g_y(\cdot)$ are measurable functions in $\mathbb{R}$ such that X_i and Y_i are well-defined. In the scalar case with $p = 1$, (7) and (8) include a very general class of stationary processes; c.f. [55,60,68,69,75,77].

Following [77], at lag $i \geq 0$, we define the functional dependence measure

$$\begin{aligned}\delta_{i,q,j} &= \|X_{ij} - X_{ij,\{0\}}\|_q = \|g_j(\mathcal{F}_i) - g_j(\mathcal{F}_{i,\{0\}})\|_q, \\ \delta_{i,q,y} &= \|Y_i - Y_{i,\{0\}}\|_q = \|g_y(\mathcal{F}_i) - g_y(\mathcal{F}_{i,\{0\}})\|_q,\end{aligned}$$

where the coupled process $X_{ij,\{0\}} = g_j(\mathcal{F}_{i,\{0\}})$ and $Y_{i,\{0\}} = g_y(\mathcal{F}_{i,\{0\}})$ with $\mathcal{F}_{i,\{0\}} = \sigma(\dots, \varepsilon_{-1}, \varepsilon'_0, \varepsilon_1, \dots, \varepsilon_{i-1}, \varepsilon_i)$ and $\varepsilon'_0, \varepsilon_l, l \in \mathbb{Z}$, being i.i.d. random elements. The dependence measure $\delta_{i,q,j}$ quantifies the q -th moment of the difference between the original process X_{ij} and the coupled process $X_{ij,\{0\}}$ with ε_0 replaced by ε'_0 and all the other innovations kept the same. We assume short-range dependence so that

$$\Delta_{m,q,j} := \sum_{i=m}^{\infty} \delta_{i,q,j} < \infty,$$

$$\Delta_{m,q,y} := \sum_{i=m}^{\infty} \delta_{i,q,y} < \infty.$$

Then for fixed m , $\Delta_{m,q,j}$ and $\Delta_{m,q,y}$ measure the cumulative effect of ε_0 on $(X_{ij})_{i \geq m}$ and $(Y_i)_{i \geq m}$.

We introduce the following dependence adjusted norms

$$\|X_{\cdot j}\|_{q,\alpha} = \sup_{m \geq 0} (m+1)^\alpha \Delta_{m,q,j}, \quad \alpha \geq 0, \quad (9)$$

$$\|X_{\cdot j}\|_{q,\text{GMC}} = \sup_{m \geq 0} \rho_j^{-m} \Delta_{m,q,j}, \quad \text{for some } 0 < \rho_j < 1, \quad (10)$$

$$\|X_{\cdot j}\|_{\psi_\nu} = \sup_{q \geq 2} q^{-\nu} \Delta_{0,q,j}. \quad (11)$$

Similarly, we can define $\|Y_{\cdot}\|_{q,\alpha}$, $\|Y_{\cdot}\|_{q,\text{GMC}}$ and $\|Y_{\cdot}\|_{\psi_\nu}$. Both $\|X_{\cdot j}\|_{q,\alpha}$ and $\|X_{\cdot j}\|_{q,\text{GMC}}$ are called the q -th dependence adjusted norms. By (9), if $\|X_{\cdot j}\|_{q,\alpha} < \infty$, then $\Delta_{m,q,j} = \sum_{i=m}^{\infty} \delta_{i,q,j} = O(m^{-\alpha})$ so that a larger α indicates weaker temporal dependence. In other words, $\|X_{\cdot j}\|_{q,\alpha}$ assumes polynomial decay of the dependence measure with a larger value of α leading to weaker temporal dependence. In contrast, $\|X_{\cdot j}\|_{q,\text{GMC}}$ allows exponential decay of the functional dependence measure $\delta_{i,q,j}$. Property (10) is also called geometric moment contraction (GMC(q)); c.f. [63,80,82]. In addition, $\|X_{\cdot j}\|_{\psi_\nu}$ represents the $\psi_{1/\nu}$ Orlicz norm in the dependence case. In the special case that X_{ij} are i.i.d. with mean 0, we have $\delta_{i,q,j} = 0$ for all $i \geq 1$, and $\delta_{0,q,j} = \|X_{0j} - X_{0j,\{0\}}\|_q$. In this case, the q -th dependence adjusted norms $\|X_{\cdot j}\|_{q,\alpha}$ and $\|X_{\cdot j}\|_{q,\text{GMC}}$ are equivalent to the $\mathcal{L}^q$ norm $\|X_{ij}\|_q$ in the sense that $\|X_{ij}\|_q \leq \delta_{0,q,j} \leq \|X_{ij}\|_q + \|X_{ij,\{0\}}\|_q = 2\|X_{ij}\|_q$. Moreover, if $\nu = 1/2$ (resp. $\nu = 1$), $\|X_{\cdot j}\|_{\psi_\nu}$ is a sub-Gaussian (resp. sub-exponential) norm of the random variable X_{ij} . Hence, the dependence adjusted norms $\|\cdot\|_{q,\alpha}$, $\|\cdot\|_{q,\text{GMC}}$ and $\|\cdot\|_{\psi_\nu}$ can be naturally interpreted as the polynomial moment and the exponential moment accounting for dependence, respectively.

Remark 3.1. If $\|X_{ij}\|_{q^*} < \infty$ for some $q^* > 0$ and $\|X_{\cdot j}\|_{q_0,\text{GMC}} < \infty$ holds for the process (X_{ij}) for some $0 < q_0 \leq q^*$ and $\rho_j \in (0, 1)$, then $\|X_{\cdot j}\|_{q,\text{GMC}} < \infty$ also satisfies with the same ρ_j for all $q \in (0, q^*]$. The above property of geometric moment contraction follows by Lemma 2 in [79].

We provide an example of high dimensional time series below, for which we can calculate the bounds of the dependence adjust norms defined in (9), (10) and (11).

Example 3.1. Let ε_{ij} , $i, j \in \mathbb{Z}$, be i.i.d. random variables with mean 0 and $\|\varepsilon_{ij}\|_q < \infty$ for some $q \geq 2$. Define the p -dimensional linear process

$$X_i = \sum_{k=0}^{\infty} A_k \varepsilon_{i-k}, \quad (12)$$

where $\varepsilon_i = (\varepsilon_{i1}, \dots, \varepsilon_{ip})^\top$, A_k are $p \times p$ real coefficient matrices such that $\sum_{k=0}^{\infty} \text{tr}(A_k A_k^\top) < \infty$. Then by Kolmogorov's three series theorem, the linear process (12) is well defined. Let $A_{k,j\cdot}$ be the j th row of A_k . By Rosenthal's inequality [61], we can obtain

$$\delta_{i,q,j} = \|A_{i,j\cdot} \varepsilon_0\|_q \leq (q-1)^{1/2} \|A_{i,j\cdot}\|_2 \|\varepsilon_{00}\|_q.$$

(i). If there exist $\theta > 1$ and $K > 0$ such that $\max_{1 \leq j \leq p} \|A_{i,j\cdot}\|_2 \leq K(i+1)^{-\theta}$ for all $i \geq 0$, then with $\alpha = \theta - 1$, we have

$$\|X_j\|_{q,\alpha} = \sup_{m \geq 0} (m+1)^\alpha \sum_{i=m}^{\infty} \delta_{i,q,j} \leq C_{\theta,q} K \|\varepsilon_{00}\|_q,$$

where the constant $C_{\theta,q}$ only depends on θ and q . If $\alpha' > \theta - 1$, then $\|X_j\|_{q,\alpha'}$ may be ∞ .

(ii). If there exist $\rho_j \in (0, 1)$ and $K > 0$ such that $\max_{1 \leq j \leq p} \|A_{i,j\cdot}\|_2 \leq K \rho_j^i$ holds for all $i \geq 0$, then since $\sum_{i=m}^{\infty} \rho_j^i = \rho_j^m / (1 - \rho_j)$, we have

$$\|X_j\|_{q,\text{GMC}} = \sup_{m \geq 0} \rho_j^{-m} \sum_{i=m}^{\infty} \delta_{i,q,j} \leq \frac{K(q-1)^{1/2} \|\varepsilon_{00}\|_q}{1 - \rho_j}.$$

(iii). Suppose $\max_{|\theta|_2=1} q^{-\nu} \|\varepsilon_i^\top \theta\|_q \leq C_\nu < \infty$ for all $q \geq 2$. If $\max_{1 \leq j \leq p} \sum_{i=0}^{\infty} \|A_{i,j\cdot}\|_2 \leq K$, for some $K > 0$, then

$$\|X_j\|_{\psi_\nu} = \sup_{q \geq 2} q^{-\nu} \sum_{i=0}^{\infty} \delta_{i,q,j} \leq C_\nu K.$$

When $\nu = 1/2$, it is similar to the sub-Gaussian Orlicz norm in the i.i.d. case; see for example [16,17].

To account for the cross-sectional dependence of the p -dimensional stationary process (X_i) , we define the $\mathcal{L}^\infty$ functional dependence measure and its corresponding dependence adjusted norm (cf. [13,14,83])

$$\omega_{i,q} = \left\| \max_{1 \leq j \leq p} |X_{ij} - X_{ij,\{0\}}| \right\|_q = \| |X_i - X_{i,\{0\}}| \|_q,$$

$$\| |X_\cdot|_\infty \|_{q,\alpha} = \sup_{m \geq 0} (m+1)^\alpha \Omega_{m,q}, \quad \alpha \geq 0, \quad \text{and} \quad \Omega_{m,q} = \sum_{i=m}^{\infty} \omega_{i,q}.$$

Additionally, we define

$$\| |X_\cdot|_\infty \|_{\psi_\nu} = \sup_{q \geq 2} q^{-\nu} \Omega_{0,q}.$$

In this paper, we use the above dependence adjusted norms to study the limiting properties of Lasso estimators in the presence of serial dependence. The framework of functional dependence measure allows many linear and nonlinear time series models; see [63,77,80] for examples.

3.2. Assumptions

To prove theoretical properties of our method, we make the following assumptions.

Assumption 1 (Convex Loss). Throughout this paper, the map

$$z \mapsto R(z, y)$$

is convex for all $y \in \mathcal{Y}$.

This assumption is important from a computational perspective; see e.g. [41, 72]. It also plays a crucial role in our theory, as it allows us to prove that the estimator $\hat{\beta}$ is in a neighborhood of β^* .

Assumption 2. Consider the function r in the maximum log-likelihood loss function (2). It holds that $|r'(x)| \leq M_1 < \infty$, $|r''(x)| \leq M_2 < \infty$ for any $x \in \mathbb{R}$. Moreover, $|r''(x_1) - r''(x_2)| \leq M_3|x_1 - x_2|$ for any $x_1, x_2 \in \mathbb{R}$ with $M_3 < \infty$.

Assumption 2 describes some smoothness conditions. Similar assumptions were imposed in [1, 16, 42]. Assumption 2 requires that the second order derivative is Lipschitz. It is used to control the symmetric Bregman divergence and Hessian of the log-likelihood in a neighborhood of β^* . In fact, except for Poisson regression, many examples of generalized linear models and loss functions satisfy this condition, such as logistic regression, multinomial logistic regression, huber loss, etc.

Assumption 3. Let $S = \{j : \beta_j^* \neq 0\}$. Assume $|S| \leq s$.

Assumption 4. Assume $\mathbb{E}X_i = 0$ and there exists a constant $\kappa_H > 0$ such that $\lambda_{\min}(\mathbb{E}\mathbf{H}_n(\beta^*)) = \lambda_{\min}(\mathbb{E}r''(X_i^\top \beta^*)X_i X_i^\top) \geq \kappa_H$.

Assumption 4 is similar to the compatibility condition in [9, 73]. It is well known that the restricted strong convexity (RSC) of the loss function underpins the statistical guarantee of the M -estimator [52]. In high dimensional sparse linear regression, RSC is implied by the restricted eigenvalue condition [4]. However, in generalized linear models, the Hessian matrix $\mathbf{H}_n(\beta)$ depends on β , which creates some technical difficulty of verifying RSC for the loss function. To address the issue, we need to establish local RSC (LRSC) of $\mathcal{R}_n(\beta)$ within a neighborhood of β^* , which has been shown to be sufficient for statistical guarantee of regularized M -estimators in the high dimensional regime [18]. The LRSC is also closely related to the quadratic margin condition in [9, 73]. It can be verified by our Assumptions 2 and 4 in every case of this paper.

Formally, we say a loss function $\mathcal{R}_n(\beta)$ satisfies LRSC($C(S), \mathcal{N}, \kappa_{\mathcal{R}}, \varphi_{\mathcal{R}}$) if any $\Delta \in C(S)$ and $\beta \in \mathcal{N}$,

$$\mathcal{R}_n(\beta + \Delta) - \mathcal{R}_n(\beta) - (\nabla \mathcal{R}_n(\beta))^\top \Delta \geq \kappa_{\mathcal{R}}|\Delta|_2^2 - \varphi_{\mathcal{R}}, \quad (13)$$

where $\mathcal{N}$ is a neighborhood of β^* , $\kappa_{\mathcal{R}}$ is a positive constant, $\varphi_{\mathcal{R}}$ is a small tolerance term, $C(S)$ is the restricted cone for $S \subset \{1, 2, \dots, p\}$ and $|S| = s$,

$$C(S) := \{\Delta \in \mathbb{R}^p : |\Delta_{S^c}|_1 \leq 3|\Delta_S|_1\}. \quad (14)$$

Note that, according to [52], when $\lambda \geq 2|\nabla \mathcal{R}_n(\beta^*)|_\infty$, $\hat{\beta} - \beta^*$ falls in the restricted cone $C(S)$.

3.3. Rate of convergence for the standard Lasso procedure

In this section, we present ℓ_1 , ℓ_2 and ℓ_∞ bounds of the standard Lasso estimator $\hat{\beta}$ and the model selection consistency. We first establish the LRSC of Lasso method for the original time series data under the finite polynomial moment case.

Lemma 3.1. *Suppose Assumptions 3 and 4 hold. Assume $|\beta^*|_1 \leq L < \infty$, $|r''(x)| \leq M_2 < \infty$ and $|r''(x_1) - r''(x_2)| \leq M_3|x_1 - x_2|$ with $M_3 < \infty$. Furthermore, assume that $\max_{|v|_2=1} \mathbb{E}|X_i^\top v|^\gamma \leq c_0 < \infty$ and $\|X\|_\infty \|X\|_{\gamma, \alpha_X} < \infty$, where $\gamma > 32/7$, $\alpha_X > 21/2 - 8/\gamma$. Let $\alpha_2 = \alpha_X/7 - 1$. Suppose*

$$\lambda \asymp a_{p,4} \sqrt{\frac{\log p}{n}} + a_{p,5} n^{16/(7\gamma)-1} (\log p)^{3/2},$$

where

$$a_{p,4} = \|X\|_\infty \|X\|_{\gamma, \alpha_X}^{1/7} \max_j \|X_{\cdot j}\|_{\gamma, \alpha_X}^2 + \max_j \|X_{\cdot j}\|_{\gamma, \alpha_2}^2, \quad (15)$$

$$a_{p,5} = \|X\|_\infty \|X\|_{\gamma, \alpha_X}^{1/7} \|X\|_\infty \|X\|_{\gamma, \alpha_X}^2 + \|X\|_\infty \|X\|_{\gamma, \alpha_2}^2. \quad (16)$$

Let $\mathcal{N} = \{\beta \in \mathbb{R}^p : |\beta - \beta^*|_2^2 \leq C_1 s \lambda^2, \beta - \beta^* \in C(S)\}$, where $C(S)$ is defined in (14). Then, as long as $n^{2/\gamma} \sqrt{s} \lambda + s \lambda \leq C_L$, $\mathcal{R}_n(\beta)$ satisfies $\text{LRSC}(C(S), \mathcal{N}, \kappa_H/2, 0)$ in an event with probability at least $1 - C_2(\log p)^{-7/(16\gamma)} - C_3 n^{-1} - p^{-C_4}$, where $C_1, \dots, C_4 > 0$ are constants and $C_L > 0$ depends on L .

Based on the above lemma, we can derive the statistical rate of the Lasso estimators. For any Lasso solution $\hat{\beta}$ of Equation (1), the following theorem provides the rates of convergence of $|\hat{\beta} - \beta^*|_1$ and $|\hat{\beta} - \beta^*|_2$ by the dependence adjusted norms.

Theorem 3.1. *Suppose Assumptions 1, 2, 3 and 4 hold. Assume $|\beta^*|_1 \leq L < \infty$. Also assume $\max_{|v|_2=1} \mathbb{E}|X_i^\top v|^\gamma \leq c_0 < \infty$, $\|X\|_\infty \|X\|_{\gamma, \alpha_X} < \infty$ and $\|Y\|_{q, \alpha_Y} < \infty$, where $\gamma > 32/7$, $q > 2$, $\alpha_X > 21/2 - 8/\gamma$, $\alpha_Y > 1/2 - 1/\gamma - 1/q > 0$. Let $\alpha_1 = \min\{\alpha_X, \alpha_Y\}$ and $\alpha_2 = \alpha_X/7 - 1$. Define*

$$a_{p,1} = \max_j \|X_{\cdot j}\|_{\gamma, \alpha_1} \|Y\|_{q, \alpha_1} + \|X\|_\infty \|X\|_{\gamma, \alpha_X}^{1/7} \max_j \|X_{\cdot j}\|_{\gamma, \alpha_X} + \max_j \|X_{\cdot j}\|_{\gamma, \alpha_2},$$

$$a_{p,2} = \|X\|_\infty \|X\|_{\gamma, \alpha_1} \|Y\|_{q, \alpha_1},$$

$$a_{p,3} = \|X\|_\infty \|X\|_{\gamma, \alpha_X}^{1/7} \|X\|_\infty \|X\|_{\gamma, \alpha_X} + \|X\|_\infty \|X\|_{\gamma, \alpha_2}.$$

Suppose that

$$\lambda \asymp (a_{p,1} + a_{p,4}) \sqrt{\frac{\log p}{n}} + \frac{a_{p,2}(\log p)^{3/2}}{n^{1-1/q-1/\gamma}} + \frac{a_{p,3}(\log p)^{3/2}}{n^{1-8/(7\gamma)}} + \frac{a_{p,5}(\log p)^{3/2}}{n^{1-16/(7\gamma)}}, \quad (17)$$

where $a_{p,4}$ and $a_{p,5}$ are defined in (15) and (16). Then, as long as $n^{2/\gamma} \sqrt{s} \lambda + s \lambda \leq C_L$, in an event with probability at least $1 - C_1(\log p)^{-q\gamma/(q+\gamma)} - C_2(\log p)^{-7\gamma/16} - n^{-1} - p^{-C_3}$,

$$|\hat{\beta} - \beta^*|_2^2 \leq C_4 s \lambda^2, \quad (18)$$

$$|\hat{\beta} - \beta^*|_1 \leq C_5 s \lambda, \quad (19)$$

where $C_1, \dots, C_5 > 0$ are constants, and $C_L > 0$ only depends on L .

Note that quantities $a_{p,1}, \dots, a_{p,5}$ characterize the dependence adjusted norms and may depend on the dimension p . The temporal dependence contributes some additional factors in the statistical error rates and in the sample size requirement. The first term in the order of λ in (17) resembles the well known sub-Gaussian rate $\sqrt{(\log p)/n}$ in the i.i.d. case. The quantities $\alpha_1, \alpha_2, \alpha_X$ measure the temporal dependence strength, and the moment condition γ quantifies the heaviness of the tail. Thus the remaining terms in (17), which are introduced by the heavy-tailedness, will induce a polynomial term of p if $\|X_{\cdot}\|_{\infty} \|Y_{\cdot}\|_{\gamma, \alpha_X} \asymp p^{1/\gamma}$.

Here we assume that the short-range dependence condition holds, that is,

$$\|X_{\cdot}\|_{\infty} \|Y_{\cdot}\|_{\gamma, \alpha_X} < \infty \quad \text{and} \quad \|Y_{\cdot}\|_{q, \alpha_Y} < \infty.$$

If it fails, the processes (X_i) and (Y_i) may exhibit some long-range dependence, and the asymptotic behavior can be quite different.

In Theorem 3.1, both α_X and α_Y control the decaying speed of the functional dependence measures. To gain more insight into the effects of temporal dependence, we provide the statistical error rates in the following proposition under the i.i.d. setting. The quantity λ in (17) can be substantially simplified with $(\log p)^{3/2}$ there being replaced by $\log p$, and $a_{p,1}, \dots, a_{p,5}$ replaced by the moments that do not involve temporal dependencies.

Proposition 3.1. *Suppose Assumptions 1, 2, 3 and 4 hold. Assume that the observations (X_i, Y_i) are i.i.d. and $\max_{|v|_2=1} \mathbb{E}|X_i^\top v|^\gamma \leq c_0 < \infty$, $\|Y_i\|_q < \infty$, where $\gamma \geq 4$, $q > 2$. Suppose that*

$$\begin{aligned} \lambda \asymp & (\max_j \|X_{ij}\|_\gamma \|Y_i\|_q + \max_j \|X_{ij}\|_\gamma^2) \sqrt{\frac{\log p}{n}} + \frac{(\log p) \|\max_j |X_{ij}|\|_\gamma \|Y_i\|_q}{n^{1-1/q-1/\gamma}} \\ & + \frac{(\log p) \|\max_j |X_{ij}|\|_\gamma^2}{n^{1-2/\gamma}}. \end{aligned}$$

Then, as long as $n^{2/\gamma} \sqrt{s} \lambda + s \lambda \leq C_0$, in an event with probability at least $1 - C_1(\log p)^{-\gamma/2} - C_2(\log p)^{-q\gamma/(q+\gamma)} - n^{-1} - p^{-C_3}$, the statistical error rates in (18) and (19) continue to hold, where $C_0, C_1, C_2, C_3 > 0$ are constants.

Remark 3.2 (Effects of heavy-tailedness). In both Theorem 3.1 and Proposition 3.1, the tuning parameter λ includes a polynomial term of dimension p and sample size n , indicating how the dimension breaks down if the moment condition weakens or the dependence becomes stronger. In addition, error bounds (18) and (19) with the new rates for λ show that the convergence rate of the estimated coefficient $\hat{\beta}$ can be much slower than that in the i.i.d. sub-Gaussian setting.

When the process X_i has an exponential type tail bound, we verify the LRSC in the following Lemma 3.2. Differently from Theorem 3.1, sharper statistical rates of $\hat{\beta}$ are established in Theorem 3.2.

Lemma 3.2. *Suppose Assumptions 3 and 4 hold. Assume $|\beta^*|_1 \leq L < \infty$, $|r''(x)| \leq M_2 < \infty$ and $|r''(x_1) - r''(x_2)| \leq M_3|x_1 - x_2|$ with $M_3 < \infty$. Furthermore, assume $\max_{1 \leq j \leq p} \|X_{\cdot j}\|_{\psi_\iota} < \infty$, and $\sup_{\gamma \geq 1} \max_{|\theta|_2=1} \gamma^{-\iota} \|X_i^\top \theta\|_\gamma \leq c_0 < \infty$, where $\iota > 0$. Let $\mathcal{N} = \{\beta \in \mathbb{R}^p : |\beta - \beta^*|_2^2 \leq C_1 s \lambda^2, \beta - \beta^* \in C(S)\}$, where $C(S)$ is defined in (14).*

(i). *Suppose*

$$\lambda \asymp \max_j \|X_{\cdot j}\|_{\psi_\iota}^3 \frac{(\log p)^{1/2+3\iota}}{\sqrt{n}}.$$

Then, as long as $(\log n)^t \sqrt{s} \lambda + s \lambda \leq C_L$, $\mathcal{R}_n(\beta)$ satisfies LRSC($C(S), \mathcal{N}, \kappa_H/2, 0$) in an event with probability at least $1 - n^{-C_2} - p^{-C_3}$, where $C_1, C_2, C_3 > 0$ are constants and $C_L > 0$ depends on L .

(ii). Assume the process X_i also satisfies the geometric moment contraction, so that $\|X_{\cdot j}\|_{6, \text{GMC}} < \infty$ for some constant $0 < \rho_j < 1$, and $\rho = \min_j \{\rho_j\} \in (0, 1)$. Suppose

$$\lambda \asymp \max_j \|X_{\cdot j}\|_{6, \text{GMC}}^3 \sqrt{\frac{\log p}{n}} + \frac{(\log p + \log n)^{1+2\iota}}{n}.$$

Then, so long as $(\log n)^t \sqrt{s} \lambda + s \lambda \leq C_L$, $\mathcal{R}_n(\beta)$ satisfies LRSC($C(S), \mathcal{N}, \kappa_H/2, 0$) in an event with probability at least $1 - n^{-C_5} - p^{-C_6}$, where $C_1, C_5, C_6 > 0$ are constants and $C_L > 0$ depends on L .

Theorem 3.2. Suppose Assumptions 1, 2, 3 and 4 hold. Assume $|\beta^*|_1 \leq L < \infty$. Also assume $\sup_{\gamma \geq 1} \max_{|\theta|_2=1} \gamma^{-t} \|X_i^\top \theta\|_\gamma \leq c_0 < \infty$, $\max_{1 \leq j \leq p} \|X_{\cdot j}\|_{\psi_\iota} < \infty$, $\|Y_{\cdot}\|_{\psi_\nu} < \infty$, where $\iota, \nu > 0$.

(i). Suppose

$$\lambda \asymp \max_j \|X_{\cdot j}\|_{\psi_\iota} \|Y_{\cdot}\|_{\psi_\nu} \frac{(\log p)^{1/2+\nu+\iota}}{\sqrt{n}} + \max_j \|X_{\cdot j}\|_{\psi_\iota}^3 \frac{(\log p)^{1/2+3\iota}}{\sqrt{n}}. \quad (20)$$

Then, so long as $(\log n)^t \sqrt{s} \lambda + s \lambda \leq C_L$, in an event with probability at least $1 - n^{-C_1} - p^{-C_2}$, we have

$$|\hat{\beta} - \beta^*|_2^2 \leq C_3 s \lambda^2, \quad (21)$$

$$|\hat{\beta} - \beta^*|_1 \leq C_4 s \lambda, \quad (22)$$

where $C_1, C_2, C_3, C_4 > 0$ are constants, and $C_L > 0$ only depends on L .

(ii). Assume the processes X_i and Y_i also satisfy geometric moment contraction, so that $\|X_{\cdot j}\|_{6, \text{GMC}} < \infty$ for some constant $0 < \rho_j < 1$, $\|Y_{\cdot}\|_{4, \text{GMC}} < \infty$ for some constant $0 < \rho_Y < 1$, and $\rho = \min\{\rho_j, \rho_Y\} \in (0, 1)$. Suppose

$$\begin{aligned} \lambda \asymp & (\max_j \|X_{\cdot j}\|_{4, \text{GMC}} \|Y_{\cdot}\|_{4, \text{GMC}} + \max_j \|X_{\cdot j}\|_{6, \text{GMC}}^3) \sqrt{\frac{\log p}{n}} \\ & + \frac{(\log p + \log n)^{1+\iota+\nu}}{n} + \frac{(\log p + \log n)^{1+2\iota}}{n}. \end{aligned} \quad (23)$$

Then, as long as $(\log n)^t \sqrt{s} \lambda + s \lambda \leq C_L$, in an event with probability at least $1 - n^{-C_1} - p^{-C_2}$,

$$|\hat{\beta} - \beta^*|_2^2 \leq C_3 s \lambda^2, \quad (24)$$

$$|\hat{\beta} - \beta^*|_1 \leq C_4 s \lambda, \quad (25)$$

where $C_1, C_2, C_3, C_4 > 0$ are constants, and $C_L > 0$ only depends on L .

Theorem 3.2 describes how the rate of convergence depends on the sample size n , the dimension p , and the dependence adjusted norms which are characterized by ι and ν . It suggests that, under the short-range dependence with exponential moment conditions, we can take $\lambda \asymp (\log p)^{c_1}/n^{c_2}$ for some positive constants c_1, c_2 and $c_1 < c_2$. Based on the scaling condition $(\log n)^t \sqrt{s} \lambda + s \lambda \leq C_L$, p is allowed to be of ultra high dimension in that $\exp(n^c)$ for some $0 < c < 1$.

Remark 3.3. In the setting of Theorem 3.2(ii), besides exponential moment, we also assume X_i and Y_i are weakly dependent and satisfy the geometric moment contraction, which is defined in (10). In

other words, if X_i and Y_i have exponentially decay speed of the functional dependence measures, then, a sharper convergence rate of $\hat{\beta}$ can be achieved. In the order of the tuning parameter λ in (23), $\max_j \|X_{\cdot j}\|_{4,\text{GMC}} \|Y_{\cdot}\|_{4,\text{GMC}} + \max_j \|X_{\cdot j}\|_{6,\text{GMC}}^3$ contributes some additional dependence adjusted norm terms. Meanwhile, the term 1 in the power of the terms $(\log p + \log n)^{1+\iota+\nu} + (\log p + \log n)^{1+2\iota}$ is introduced by the geometric moment contraction. The following Proposition 3.2 shows the results in the case that the observations (X_i, Y_i) are i.i.d. and have exponential tail bounds.

Proposition 3.2. *Suppose Assumptions 1, 2, 3 and 4 hold. Assume that the observations (X_i, Y_i) are i.i.d. and $\sup_{\gamma \geq 1} \max_{|\theta|_2=1} \gamma^{-\iota} \|X_i^\top \theta\|_\gamma \leq c_0 < \infty$, $\sup_{q \geq 1} q^{-\nu} \|Y_i\|_q \leq c_0 < \infty$, where $\iota, \nu > 0$. Suppose that*

$$\lambda \asymp \sqrt{\frac{\log p}{n}} + \frac{(\log p + \log n)^{\iota+\nu}}{n} + \frac{(\log p + \log n)^{2\iota}}{n}.$$

Then, as long as $(\log n)^\iota \sqrt{s} \lambda + s \lambda \leq C_0$, in an event with probability at least $1 - n^{-C_1} - p^{-C_2}$, (24) and (25) continue to hold, where $C_0, C_1, C_2 > 0$ are constants.

Again, both Theorem 3.2 and Proposition 3.2 show that the range of dimension p is narrower than the range $\log(p) = o(n)$ in the i.i.d. sub-Gaussian setting, and the convergence rates of the estimated coefficient $\hat{\beta}$ are also slower.

Theorems 3.1 and 3.2 also lead to the following result on the ℓ_∞ bound of $\hat{\beta}$ and support recovery.

Corollary 3.1. *Let $H(\beta^*) = \mathbb{E}H_n(\beta^*)$. Suppose the following incoherence condition holds*

$$\left\| H(\beta^*)_{S^c S} (H(\beta^*)_{SS})^{-1} \right\|_\infty \leq \eta < 1,$$

where $\|\cdot\|_\infty$ denotes the ℓ_∞ matrix operator norm. Further, suppose the conditions of Theorem 3.1 (resp. Theorem 3.2(i) or 3.2(ii)) are satisfied. Then, as long as $n^{2/\gamma} s \lambda + s^{3/2} \lambda \leq C_L$ (resp. $(\log n)^\iota s \lambda + s^{3/2} \lambda \leq C_L$), in an event with probability at least $1 - C_1(\log p)^{-\chi} - C_2(\log p)^{-7\gamma/16} - n^{-1} - p^{-C_3}$ (resp. $1 - n^{-C_1} - p^{-C_2}$), we have $\hat{\beta}_{S^c} = 0$ and

$$\|\hat{\beta} - \beta^*\|_\infty \leq C_0 \| (H(\beta^*)_{SS})^{-1} \|_\infty \lambda,$$

where $C_0, \dots, C_3 > 0$, $C_L > 0$ depends on L , λ is defined in (17) (resp. (20) or (23)).

Similarly to Corollary 3 in [42], the required scaling condition for support recovery would be stronger than that for Theorem 3.1 or Theorem 3.2. It is also worth noting that the incoherence condition in Corollary 3.1 can be removed by using various non-convex regularizers, as shown in [42].

Remark 3.4. For i.i.d. data, the asymptotic behavior of Lasso estimator for GLM was studied in [73]. Many other papers investigated high dimensional robust M-estimator, such as [17] and [41]. A key technique used in these articles is Massart's inequality [46], Bousquet's inequality [6] or other similar inequalities for empirical process. These inequalities cannot be directly used for serially dependent data. It is noteworthy that the proofs of Theorem 3.1 and 3.2 require new concentration inequalities for high dimensional time series. We establish Bousquet-type inequality for time dependent data, which is shown in the supplementary material [28].

Remark 3.5. The setting in our Theorems 3.1 and 3.2 is very general as it allows dependent and/or non sub-Gaussian processes and it also allows heteroscedasticity in that the error process and the covariate process can be dependent. Comparing with the conditions in the i.i.d setting (Propositions 3.1 and 3.2), our Theorems 3.1 and 3.2 require an additional condition $|\beta^*| \leq L < \infty$ to derive functional dependence measures for temporal dependent processes.

3.4. Rate of convergence for the robust procedure

To tackle the problem of heavy-tailed data, we propose to use the robust regularization method in Section 2.2, and analyze the robust estimator in this section. Under similar assumptions, we establish LRSC of the proposed robust procedure.

Lemma 3.3. Assume $|\beta^*|_1 \leq L < \infty$, $|r''(x)| \leq M_2 < \infty$ and $|r''(x_1) - r''(x_2)| \leq M_3|x_1 - x_2|$ with $M_3 < \infty$. Also assume $\mathbb{E}X_{ij}^6 \leq C < \infty$, for any $1 \leq j \leq p$, $\max_{|v|_2=1} \mathbb{E}|X_i^\top v|^2 \leq c_0 < \infty$, $\|X_{\cdot j}\|_{6,\text{GMC}} < \infty$ for some constant $0 < \rho_j < 1$, and $\rho = \min_j \rho_j \in (0, 1)$. Suppose $\lambda \asymp (\log n)\sqrt{(\log p)/n}$ and $\tau \asymp n^{1/4}(\log p)^{-1/4}(\log n)^{-1/2}$. Let $\mathcal{N} = \{\beta \in \mathbb{R}^p : |\beta - \beta^*|_2^2 \leq C_1 s \lambda^2, \beta - \beta^* \in C(S)\}$, where $C(S)$ is defined in (14). Then, as long as $s^2(\log n)(\log p/n)^{1/2} \leq C_2$, $\mathcal{R}_n(\beta)$ satisfies LRSC($C(S), \mathcal{N}, \kappa_H/2, 0$) in an event with probability at least $1 - p^{-C_3}$, where C_1 and C_2 are positive constants and $C_3 > 0$ depends on L, ρ , and $\max_j \|X_{\cdot j}\|_{6,\text{GMC}}$.

The next theorem shows the rate of convergence of $|\hat{\beta} - \beta^*|_1$ and $|\hat{\beta} - \beta^*|_2$ by the dependence adjusted norms. In Theorem 3.3, the response is generated from a particular distribution in the exponential family. Therefore, there does not exist any model misspecification issue, and the response has sub-exponential tails.

Theorem 3.3. Suppose Assumptions 1, 2, 3 and 4 hold. Assume $|\beta^*|_1 \leq L < \infty$. Also assume $\mathbb{E}X_{ij}^6 \leq C < \infty$, for any $1 \leq j \leq p$, $\max_{|v|_2=1} \mathbb{E}|X_i^\top v|^2 \leq c_0 < \infty$, and $\mathbb{E}Y_i^4 \leq C < \infty$. Let $\|X_{\cdot j}\|_{6,\text{GMC}} < \infty$ for some constant $0 < \rho_j < 1$, $\|Y_{\cdot}\|_{4,\text{GMC}} < \infty$ for some constant $0 < \rho_y < 1$, $\|Y_{\cdot}\|_{\psi_v} < \infty$ for $v > 0$, and $\rho = \min\{\rho_j, \rho_y\} \in (0, 1)$. Choose $\tau \asymp n^{1/4}(\log p)^{-1/4}(\log n)^{-1/2}$, and $\lambda = C_1(\log n)\sqrt{(\log p)/n}$. Then, as long as $(\log n)^{2\nu+1}(\log p/n)^{1/2} + s^2(\log n)(\log p/n)^{1/2} \leq C_2$, in an event with probability at least $1 - n^{-C_3} - p^{-C_4}$,

$$|\hat{\beta} - \beta^*|_2^2 \leq C_5 s \left(\frac{(\log n)^2 \log p}{n} \right), \quad (26)$$

$$|\hat{\beta} - \beta^*|_1 \leq C_6 s \left(\frac{(\log n)^2 \log p}{n} \right)^{1/2}, \quad (27)$$

where $C_1, \dots, C_6 > 0$ are constants depending on $L, \rho, \max_{1 \leq j \leq p} \|X_{\cdot j}\|_{6,\text{GMC}}$ and $\|Y_{\cdot}\|_{4,\text{GMC}}$.

Corollary 3.2. If $\max_i |Y_i| < C < \infty$, such as categorical variables, then, as long as $s^2(\log n)(\log p/n)^{1/2} \leq C_2$, in an event with probability at least $1 - p^{-C_4}$, the upper bounds (26) and (27) hold.

Corollary 3.3. Let $H(\beta^*) = \mathbb{E}H_n(\beta^*)$. Suppose the following incoherence condition holds

$$\left\| H(\beta^*)_{S^c S} (H(\beta^*)_{SS})^{-1} \right\|_{\infty} \leq \eta < 1.$$

Also, suppose the conditions of Theorem 3.3 are satisfied. Choose $\lambda = C_1(\log n)\sqrt{(\log p)/n}$ and $\tau \asymp n^{1/4}(\log p)^{-1/4}(\log n)^{-1/2}$. Then, as long as $(\log n)^{2\nu+1}(\log p/n)^{1/2} + s^3(\log n)(\log p/n)^{1/2} \leq C_2$, in an event with probability at least $1 - n^{-C_3} - p^{-C_4}$, we have $\hat{\beta}_{Sc} = 0$ and

$$|\hat{\beta} - \beta^*|_\infty \leq C_0 \| (H(\beta^*)_{SS})^{-1} \|_\infty (\log n) \sqrt{(\log p)/n},$$

where $C_0, \dots, C_4 > 0$ and depend on $L, \rho, \max_{1 \leq j \leq p} \|X_{\cdot j}\|_{6, \text{GMC}}$ and $\|Y_{\cdot}\|_{4, \text{GMC}}$.

Theorem 3.3 indicates that the robust estimator admits nearly the same rate as that of the i.i.d. sub-Gaussian setting. It is much sharper than the deviation bounds in Theorem 3.1 under finite dependence adjusted norms, which will include a polynomial term of p if $\|X_{\cdot}\|_{\gamma, \alpha_X} \asymp p^{1/\gamma}$. In comparison with the robust procedure in the i.i.d. case, we have an additional $\log n$ factor when the result is generalized to the time series setting. Under the scaling condition $(\log n)^{2\nu+1}(\log p/n)^{1/2} + s^2(\log n)(\log p/n)^{1/2} \leq C_2$, our robust regularization method in (6) can handle the ultra high dimension case with $\log p = o(s^{-4}(\log n)^{-2}n + (\log n)^{-4\nu-2}n)$, which is much wider than the range of Theorem 3.1, which only allows a polynomial increase with n . In the robust procedure, the order of τ is not related to the finite moment condition. Only the constant term of τ should be adapted to the degree of heavy-tailedness. Comparing Corollary 3.3 with Corollary 3.1, the support recovery guarantee and the ℓ_∞ bound are also improved by the robust procedure.

Remark 3.6. The requirement of geometric moment contraction (exponentially decaying dependence measure) is needed. In the high dimensional setting, concentration inequalities are the cornerstone for theoretical analysis. [12] showed that exponential decay bounds for large deviation type tail probabilities do not hold generally with polynomial decaying functional dependence measures as defined in (9), i.e. $\sum_{i=m}^\infty \delta_{i,q,j} = O(m^{-\alpha})$, even if the process is uniformly bounded. In other words, if the temporal dependence is not very weak, any robust regularized M -estimators for observations with finite moments are not able to achieve statistical error rates close to the optimal rates under the i.i.d. sub-Gaussian case, and can only handle the dimensionality $p = o(n^c)$ for some $c > 0$. This is significantly different from the low dimensional M -estimates, where polynomial decaying dependence measures and even long-range dependent process can be applied; see for example [78].

4. Linear regression

Robust estimation of linear regression can be regarded as a generalized linear model with quadratic loss. In this special case, although the first derivative of the loss function is not bounded, we still have an explicit concentration result for our robust estimator. Consider the usual linear regression setup for the response variable Y_i and the covariate vector X_i ,

$$Y_i = X_i^\top \beta^* + e_i, \quad (28)$$

where $\beta^* \in \mathbb{R}^p$ is the unknown parameter vector to be estimated and e_i is the error term. Let the loss function $R(f_\beta(X_i), Y_i) = (Y_i - f_\beta(X_i))^2/2$. Differently from Theorem 3.3 in Section 3, if the error has heavy tails, the response Y_i also has heavy tails. This motivates us to truncate both the heavy tailed covariates X_i and the response Y_i under the ℓ_2 loss. Then, similarly to (6), we propose to use the following M -estimator of β^* with the generalized ℓ_2 loss to robustify the estimation:

$$\hat{\beta} := \arg \min_{\beta} \left\{ \frac{1}{2n} \sum_{i=1}^n (\tilde{Y}_i - f_\beta(\tilde{X}_i))^2 + \lambda |\beta|_1 \right\}, \quad (29)$$

where

$$f_{\beta}(\tilde{X}_i) = \tilde{X}_i^{\top} \beta. \quad (30)$$

Here, we choose $\tilde{Y}_i = \tilde{Y}_i(\tau_1) = \text{sgn}(Y_i)(|Y_i| \wedge \tau_1)$ and $\tilde{X}_{ij} = \text{sgn}(X_{ij})(|X_{ij}| \wedge \tau_2)$ for all $1 \leq j \leq p$, where both τ_1 and τ_2 are predetermined thresholds.

To be solvable in the high dimensional regression setting, β^* is usually assumed to be sparse or weakly sparse, *i.e.*, many elements of β^* are 0 or small. In particular, we impose the following weak sparsity condition and a condition on the covariance matrix of the covariates.

Assumption 5 (Weak Sparsity Condition). There exists some $0 \leq \vartheta < 1$, with a uniform radius K_{ϑ} , such that

$$\sum_{j=1}^p |\beta_j^*|^{\vartheta} \leq K_{\vartheta}. \quad (31)$$

Note that K_{ϑ} might depend on n and p . In the special case $\vartheta = 0$, this quantity corresponds to an exact sparsity constraint—that is, β^* has at most K_0 nonzero entries.

Assumption 6. Suppose $\mathbb{E}X_i = 0$ and there exists a constant $\kappa_0 > 0$ such that $\lambda_{\min}(\mathbb{E}X_i X_i^{\top}) \geq \kappa_0$.

To derive the statistical error rate of $\hat{\beta}$, we establish, in the following lemma, the restricted strong convexity of the robust procedure.

Lemma 4.1. Assume $\mathbb{E}X_{ij}^4 \leq C < \infty$, for any $1 \leq j \leq p$. Let $\|X_{\cdot j}\|_{4,\text{GMC}} < \infty$ for some constant $0 < \rho_j < 1$. Choose $\tau_2 \asymp n^{1/4}(\log p)^{-1/4}(\log n)^{-1/2}$. Then, for some constants $C_1, C_2 > 0$, which only depend on $\min_j \rho_j$ and $\max_j \|X_{\cdot j}\|_{4,\text{GMC}}$, we have

$$\mathbb{P}\left(\frac{1}{n} \sum_{i=1}^n \beta^{\top} \tilde{X}_i \tilde{X}_i^{\top} \beta \geq \beta^{\top} (\mathbb{E}X_i X_i^{\top}) \beta - C_1(\log n) \sqrt{\frac{\log p}{n}} |\beta|_1^2, \quad \forall \beta \in \mathbb{R}^p\right) \leq p^{-C_2}. \quad (32)$$

Finally, in the following Theorem 4.1, we present the statistical error rate of $\hat{\beta}$ as defined in (29). We show that $|\hat{\beta} - \beta^*|_2$ and $|\hat{\beta} - \beta^*|_1$ are upper bounded by nearly optimal rates under light tails and i.i.d. data so long as the tuning parameters τ_1 , τ_2 and λ are properly chosen. Corollary 4.1 concerns support recovery and ℓ_{∞} bounds.

Theorem 4.1. Suppose Assumptions 5 and 6 hold. Assume $|\beta^*|_1 \leq L < \infty$. Also assume $\max_j \mathbb{E}X_{ij}^4 \leq C < \infty$ and $\mathbb{E}Y_i^4 \leq C < \infty$. Let $\|X_{\cdot j}\|_{4,\text{GMC}} < \infty$ for some constant $0 < \rho_j < 1$, $\|Y\|_{4,\text{GMC}} < \infty$ for some constant $0 < \rho_Y < 1$, and $\rho = \min\{\rho_j, \rho_Y\} \in (0, 1)$. Choose $\tau_1 \asymp \tau_2 \asymp n^{1/4}(\log p)^{-1/4}(\log n)^{-1/2}$ and $\lambda = C_1(\log n) \sqrt{(\log p)/n}$. Then, as long as $K_{\vartheta}[(\log n)^2(\log p)/n]^{(1-\vartheta)/2} \leq C_2$, we have, in an event with probability at least $1 - p^{-C_3}$,

$$|\hat{\beta} - \beta^*|_2^2 \leq C_4 K_{\vartheta} \left(\frac{(\log n)^2 \log p}{n} \right)^{1-\vartheta/2}, \quad (33)$$

$$|\hat{\beta} - \beta^*|_1 \leq C_5 K_{\vartheta} \left(\frac{(\log n)^2 \log p}{n} \right)^{(1-\vartheta)/2}, \quad (34)$$

where $C_1, \dots, C_5 > 0$ depend on L , ρ , $\max_{1 \leq j \leq p} \|X_{\cdot j}\|_{4,\text{GMC}}$ and $\|Y\|_{4,\text{GMC}}$.

Corollary 4.1. Consider the exact sparsity case with $K_0 = s$. Let $\Sigma = \mathbb{E}(X_t X_t^\top)$. Suppose the following incoherence condition holds

$$\left\| \Sigma_{S^c S} (\Sigma_{SS})^{-1} \right\|_\infty \leq \eta < 1.$$

Suppose also the conditions of Theorem 4.1 are satisfied. Choose $\lambda = C_1(\log n)\sqrt{(\log p)/n}$ and $\tau_1 \asymp \tau_2 \asymp n^{1/4}(\log p)^{-1/4}(\log n)^{-1/2}$. Then, as long as $s[(\log n)^2(\log p)/n]^{1/2} \leq C_2$, in an event with probability at least $1 - p^{-C_3}$, we have $\hat{\beta}_{S^c} = 0$ and

$$|\hat{\beta} - \beta^*|_\infty \leq C_0 \|(\Sigma_{SS})^{-1}\|_\infty (\log n) \sqrt{(\log p)/n},$$

where $C_0, \dots, C_3 > 0$ depend on $L, \rho, \max_{1 \leq j \leq p} \|X_{\cdot j}\|_{4, \text{GMC}}$ and $\|Y\|_{4, \text{GMC}}$.

Theorem 4.1 reveals that $|\hat{\beta} - \beta^*|_2$ has a convergence rate $K_\vartheta^{1/2}[(\log n)^2(\log p)/n]^{1/2-\vartheta/4}$. In comparison with the optimal rate for weak sparsity models under light tails and i.i.d data setting [59], our result concerns high dimensional time series and relaxes the Gaussian/sub-Gaussian assumption to the existence of finite moments at the minor cost of a logarithmic factor of n in the convergence rate. In a special case that $\vartheta = 0$, $|\hat{\beta} - \beta^*|_2$ converges at the rate $K_0^{1/2}(\log n)((\log p)/n)^{1/2}$, where K_0 is the number of non-zero elements of β^* . It suggests that our robust method does not lose much information for heavy-tailed data with exponentially decaying functional dependence measures. In addition, to achieve the desired statistical rate, we need the scaling condition $K_\vartheta(\log n)(\log p/n)^{(1-\vartheta)/2} \leq c_2$, which also includes an additional $\log n$ factor.

Remark 4.1. Theorem 2 in [19] studied robust linear regression in the i.i.d. case and suggested $\lambda \asymp \sqrt{\log p/n}$ to achieve the minimax optimal rate for $\hat{\beta}$. In comparison, there is an additional multiplicative $\log n$ term in the order of λ and also the convergence rates for the dependent case, which is induced by the additional order $(\log n)^2$ in the Bernstein-type inequality under geometric moment contraction (see the supplementary materials [28]), which serves as the main technical tool. To the best of our knowledge, the sharpest available Bernstein-type inequalities for weakly dependent random variables without any structural assumptions (such as linear process) involve additional $\log n$ terms [49, 82].

5. Simulation study

In this section, we expound upon some concrete instances of our theoretical results and provide some simulation results. We assess the finite sample performance of the robust procedure and compare it with the standard Lasso procedure in logistic regression and linear regression. The implementation of our robust procedure is simple: truncate or shrink the data appropriately, then apply the standard procedure to the transformed data. The simulation results are based on 5000 independent Monte Carlo replications. We select the optimal values of tuning parameters λ and τ by a two dimensional grid search using Bayesian information criterion. It is worth noting that the robust cross-validation proposed in [19] can also be used by replacing the K -fold cross-validation with time series rolling forecast.

We first specify the parameters of the logistic regression. We generate data from independent AR(1) processes, say

$$Z_{ij} = \phi_j Z_{i-1,j} + \xi_{ij}, \quad 1 \leq j \leq p, \quad (35)$$

where $\phi_j \sim U[0.2, 0.6]$ or $\phi_j \sim U[-0.6, -0.2]$, and the innovations a_{ij} is given below. To generate cross-dependence, let

$$\Sigma = (\sigma_{jk}) = (\rho^{|j-k|}), \quad 0 < \rho < 1,$$

and $\Sigma^{1/2}$ be the square-root matrix of Σ . We use $X_i = \Sigma^{1/2}Z_i$, $1 \leq i \leq n$, where $Z_i = (Z_{i1}, \dots, Z_{ip})'$. We choose the true regression coefficient vector as

$$\beta^* = (3, \dots, 3, 0, \dots, 0)',$$

where the first 20 elements are all 3 and the rest are all 0. Let $\theta_i = 1/(1 + \exp(-X_i'\beta^*))$ be the probability of success of the Bernoulli distribution of Y_i . Thus, Y_i is a random draw from $\text{Bernoulli}(\theta_i)$. We run the simulations for sample sizes $n = 200, 300, 400, 500, 800, 1000, 2000, 3000$, and choose the number of parameters p to be 400, dependence parameter $\rho = 0.5$. For each case, additional 500 observations are generated and used for out-of-sample predictions. To entertain various shapes of covariate distributions, we consider the following two scenarios for ξ_{ij} of (35):

1. $\xi_{ij} = 0.1t_5$, i.e. the Student- t distribution with 5 degrees of freedom divided by 10;
2. $\xi_{ij} = 0.2 \log \text{Normal}(0, 0.5^2)$, i.e. a log-normal distribution with parameters 0 and 0.25^2 divided by 5.

They represent heavy-tailed symmetric and asymmetric distributions, respectively. To meet the model assumptions, the covariates are standardized to have *mean* 0. The constants used are chosen to ensure appropriate signal-to-noise ratio and θ_i not trivially equals to either 0 or 1 for better presentation. The numerical performance of the robust procedure and standard Lasso procedure under the two scenarios is evaluated by the following five measurements.

1. ℓ_2 error, which is defined as $|\hat{\beta} - \beta^*|_2$;
2. ℓ_1 error, which is defined as $|\hat{\beta} - \beta^*|_1$;
3. the number of false positive results, FP, which is the number of noise covariates that are selected;
4. the number of false negative results, FN, which is the number of signal covariates that are not selected;
5. one-step-ahead forecast errors of a total of 500 out-of-sample observations, FE, which is the misclassification rate.

For the robust Lasso, we choose the optimal tuning parameters λ and τ on the basis of 100 independent validation data sets. For each case, we run a two-dimensional grid search to find the best (λ, τ) pair that minimizes the misclassification rate of the 100 validation data sets. Then the optimal pair is used in the simulation. Similar methods are applied in choosing the tuning optimal parameters in other models. The means, over 5000 repetitions, of the five performance measures are summarized in Table 1.

The results of Table 1 show that our robust Lasso method has certain advantages over the standard Lasso method when the covariates are heavy-tailed. The results are in agreement with the theorems. As the sample size increases, the performance measures improve. In both symmetric and asymmetric covariates cases, our robust method has smaller ℓ_1 and ℓ_2 errors. The advantage of the proposed robust method is more pronounced when the sample size is large. In addition, FP increases slightly with the sample size n , but FN approaches zero as n increases.

We also investigate the empirical properties of the proposed method in linear regression. We again generate data from independent AR(1) model (35). Analogously to the logistic regression, we set $X_i = \Sigma^{1/2}Z_i$. For response, we generate time series process Y_i from the model,

$$Y_i = X_i'\beta^* + e_i, \tag{36}$$

where e_i is given below. The following two scenarios for ξ_{ij} are considered:

1. the Student- t distribution with 5 degrees of freedom, t_5 ;
2. a log-normal distribution with parameters 0 and 0.25^2 , $\log \text{Normal}(0, 0.25^2)$.

Table 1. Simulation results of Lasso and robust Lasso (RLasso) for logistic regression ($p = 400, \rho = 0.5$), where n is the sample size. The results are averages over 5000 replications

n	Scenario	Student t_5		LogNormal(0, 0.25)	
		Lasso	RLasso	Lasso	RLasso
200	ℓ_2 loss	11.53	11.32	11.59	11.39
	ℓ_1 loss	50.14	48.38	50.30	48.60
	FP	0.96	0.88	0.88	0.79
	FN	11.30	10.85	11.84	11.41
	FE	28.11%	27.24%	29.08%	28.50%
300	ℓ_2 loss	10.22	9.85	10.28	9.94
	ℓ_1 loss	43.69	40.89	43.80	41.15
	FP	1.69	1.51	1.61	1.44
	FN	5.98	5.37	6.59	5.92
	FE	22.02%	20.82%	23.09%	22.00%
400	ℓ_2 loss	9.46	8.95	9.48	9.05
	ℓ_1 loss	40.21	36.62	40.10	36.85
	FP	2.28	2.04	2.24	1.97
	FN	3.53	3.08	4.08	3.44
	FE	20.33%	19.24%	21.34%	20.23%
500	ℓ_2 loss	8.87	8.28	8.88	8.34
	ℓ_1 loss	37.56	33.51	37.38	33.61
	FP	2.64	2.32	2.54	2.23
	FN	2.23	1.84	2.66	2.14
	FE	19.44%	18.37%	20.40%	19.33%
800	ℓ_2 loss	7.66	6.82	7.67	6.93
	ℓ_1 loss	32.39	27.29	32.15	27.52
	FP	3.28	2.94	3.12	2.71
	FN	0.56	0.41	0.79	0.53
	FE	18.18%	17.17%	19.22%	18.13%
1000	ℓ_2 loss	7.15	6.20	7.11	6.26
	ℓ_1 loss	30.27	24.71	29.85	24.82
	FP	3.49	3.13	3.38	2.96
	FN	0.24	0.17	0.35	0.22
	FE	17.83%	16.82%	18.83%	17.77%
2000	ℓ_2 loss	5.67	4.42	5.61	4.43
	ℓ_1 loss	24.15	17.48	23.69	17.51
	FP	3.82	3.50	3.69	3.34
	FN	0.0018	0.0008	0.0060	0.0014
	FE	17.19%	16.22%	18.12%	17.05%
3000	ℓ_2 loss	4.92	3.53	4.84	3.52
	ℓ_1 loss	21.07	13.87	20.50	13.90
	FP	3.90	3.56	3.79	3.41
	FN	0	0	0	0
	FE	16.94%	15.98%	17.97%	16.93%

For the distribution of error e_i , we choose:

1. $e_i = 20t_3$, i.e. a Student- t distribution with 3 degrees of freedom multiplied by 20, the standard deviation of which is about 34.64;
2. $e_i = 20\log\text{Normal}(0, 0.5^2)$, i.e. a log-normal distribution with parameters 0 and 0.5^2 times 20, the standard deviation of which is about 12.08.

Again, the covariates and the errors are standardized to have *mean* 0, and the constants used are chosen to ensure appropriate signal-to-noise ratio for better presentation. We set $n = 50, 100, 200, 300, 400, 800, 1000, 2000$, and use the root mean squared forecast error (RMSE) to measure one-step-ahead forecasts of a total of 200 out-of-sample predictions. The results are reported in Table 2.

The results of Table 2 are also in agreement with the theorem. In particular, as expected, the RMSE approaches the standard deviation of e_i as the sample size increases. In general, similarly to the logistic regression, the robust estimator outperforms the non-robust one. This is particularly so for the case of heavy-tailed noises $e_i \sim 20t_3$. But as the sample size increases, the difference between the robust procedure and the standard procedure gradually decreases. The out-of-sample predictions seem to work well when the sample size $n \geq p$. For log-normal errors, FN seems to be higher than FP. Both are sizable when the sample size is small.

In conclusion, our robust method is more flexible than the standard Lasso. The above two simulation studies show clearly that the robust procedure outperforms the standard procedure under the setting with heavy-tailed covariates and errors. The truncation parameter enables the robust method to render consistently satisfactory results under all scenarios considered in our simulation.

Next, we assess the sensitivity of the robust procedures in linear regression with respect to the thresholds τ_1 and τ_2 . We use the same model setting as before for (36), and consider a heavy-tail case with t_5 features and t_3 noises, and a light-tail case with normally distributed covariates and noises. We set $n = 100, 800$, $p = 400$, and choose different quantiles of the feature values and responses as the thresholds τ_1, τ_2 . Table 3 shows the ℓ_2 and ℓ_1 loss of the estimated coefficients $\hat{\beta}$ using different values of threshold τ_1, τ_2 over 1000 replications. When τ_1 and τ_2 are set to be the 100% quantile, the procedure becomes the original Lasso method without any shrinkage. For the heavy-tail setting, our robust procedure has smaller ℓ_2 and ℓ_1 estimation error than the vanilla Lasso method in all cases. For the light-tail setting, our shrinkage method loses some efficiency on ℓ_2 and ℓ_1 estimation error. Moreover, if τ_1, τ_2 are set to be above the 90% quantiles of the original feature values and responses, the difference between robust and non-robust methods in terms of estimation error is less than 2%.

Finally, we compare our proposed truncation method with the standard method under different degrees of the cross-sectional dependence of the covariates. The setup is the same as that used before in (36). We choose $n = 200$, $p = 400$, but vary $\rho = 0.1, 0.2, 0.3, 0.4, 0.5, 0.6, 0.7, 0.8, 0.9, 0.98$. The numerical result is based on 1000 independent Monte Carlo experiments. It can be seen from Table 4 that the proposed truncation method shows significant improvement in the statistical performance over the standard method in all cases of ρ . In addition, the differences in ℓ_2 and ℓ_1 loss for different values of ρ are minor. Note that the condition number of the population covariance increases as ρ grows, but is not too large.

6. Discussion

In this paper, we first established the statistical error bounds and the support recovery guarantees of high dimensional generalized linear models in time series setting. Both finite polynomial moment case and exponential moment case are studied. In the finite moment case with geometric moment contraction, we also proposed a new robust M-estimator by shrinking the feature variables (and the response in

Table 2. Simulation results of Lasso and robust Lasso (RLasso) for linear regression ($p = 400, \rho = 0.5$), where n is the sample size and the results are averages over 5000 replications

n	Scenario	Student t		LogNormal	
		Lasso	RLasso	Lasso	RLasso
50	ℓ_2 loss	24.01	22.31	35.11	31.90
	ℓ_1 loss	148.01	140.39	218.33	201.09
	FP	44.26	44.12	47.52	47.31
	FN	12.23	12.16	15.12	14.94
	RMSE	49.86	48.10	16.46	15.78
100	ℓ_2 loss	25.29	23.44	40.86	37.04
	ℓ_1 loss	200.83	187.76	331.69	302.13
	FP	51.31	50.79	57.69	57.12
	FN	8.05	7.73	11.52	11.18
	RMSE	49.72	47.66	17.19	16.39
200	ℓ_2 loss	12.46	11.04	24.52	22.78
	ℓ_1 loss	66.12	52.68	205.78	187.80
	FP	11.89	7.83	33.56	31.73
	FN	7.19	5.85	11.06	10.64
	RMSE	40.84	39.13	15.10	14.73
300	ℓ_2 loss	9.67	8.03	12.64	10.36
	ℓ_1 loss	39.98	36.83	55.58	52.13
	FP	3.49	3.34	7.81	7.91
	FN	5.78	4.47	10.30	10.25
	RMSE	38.25	37.37	13.34	13.26
400	ℓ_2 loss	8.77	8.04	11.72	9.24
	ℓ_1 loss	37.53	32.60	50.27	47.80
	FP	2.99	2.82	1.50	1.71
	FN	3.67	2.67	10.27	9.51
	RMSE	37.04	36.18	13.13	13.01
800	ℓ_2 loss	6.42	5.78	9.44	8.95
	ℓ_1 loss	25.13	22.34	38.17	36.04
	FP	2.56	2.52	1.37	1.61
	FN	0.86	0.41	5.94	4.67
	RMSE	35.11	34.65	12.58	12.51
1000	ℓ_2 loss	5.79	5.15	8.75	8.28
	ℓ_1 loss	22.53	19.84	34.74	32.82
	FP	2.58	2.49	0.71	0.90
	FN	0.45	0.16	4.29	3.29
	RMSE	34.82	34.42	12.46	12.40
2000	ℓ_2 loss	4.11	3.56	6.69	6.22
	ℓ_1 loss	15.85	12.60	25.65	23.94
	FP	2.22	2.13	0.39	0.49
	FN	0.0376	0.0014	1.10	0.67
	RMSE	34.39	34.14	12.26	12.23

Table 3. Sensitivity of different values of thresholds τ_1 and τ_2 in linear regression case using the upper quantiles of the feature values and responses. The results are based on 1000 replications

	70%	75%	80%	85%	90%	95%	98%	99%	99.5%	100%
Student t and $n = 100, p = 400$										
ℓ_2 loss	23.51	23.39	23.30	23.25	23.29	23.45	23.82	24.15	24.53	25.70
ℓ_1 loss	188.56	187.21	186.08	185.24	185.35	186.49	189.38	192.26	195.41	204.39
Student t and $n = 800, p = 400$										
ℓ_2 loss	6.08	5.98	5.88	5.82	5.76	5.80	5.91	6.00	6.11	6.62
ℓ_1 loss	23.59	23.15	22.77	22.48	22.35	22.44	22.91	23.29	23.72	26.05
Normal and $n = 100, p = 400$										
ℓ_2 loss	20.72	20.53	20.36	20.20	20.09	19.97	19.90	19.91	19.92	19.91
ℓ_1 loss	166.34	164.38	162.44	160.72	159.35	158.14	157.52	157.41	157.36	157.26
Normal and $n = 800, p = 400$										
ℓ_2 loss	5.77	5.62	5.50	5.36	5.24	5.10	5.03	5.02	5.00	5.00
ℓ_1 loss	22.50	21.90	21.34	20.79	20.32	19.80	19.52	19.45	19.39	19.38

linear regression) before solving the empirical risk minimization. The shrinkage method works as if we have sub-Gaussian data, and does not require any algorithmic adaptation. Our robust procedure marks a significant improvement over the existing literature in the time series setting by relaxing the sub-Gaussian condition to the existence of finite moments but retaining a nearly i.i.d. sub-Gaussian deviation bound.

It is worth mentioning that high dimensional time series creates overwhelming technical difficulties in theoretical analysis. Commonly used concentration inequalities in the i.i.d. case do not hold. For example, the sharpest available Bernstein-type inequalities for weakly dependent random variables without any structural assumptions (such as linear process) involve additional logarithmic factors. Whether the large deviation bound under the i.i.d. setting is achievable or not remains a challenging open prob-

Table 4. Sensitivity of different values of the cross-sectional dependence of the covariates (ρ) in linear regression. The results are based on 1000 replications

ρ	0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9	0.98
Standard procedure										
ℓ_2 loss	12.36	12.29	12.43	12.54	12.23	12.63	12.30	12.33	12.54	12.78
ℓ_1 loss	65.11	64.10	65.48	67.43	63.23	67.68	64.87	65.24	67.85	70.75
Robust procedure										
ℓ_2 loss	10.93	10.88	10.77	10.78	10.81	11.11	10.88	11.03	11.04	10.89
ℓ_1 loss	51.47	51.15	48.89	49.40	49.64	53.84	51.06	53.44	53.14	51.28

lem. We conjecture that the scaling condition in the main theorems may be improved by developing sharper concentration inequalities.

Besides the concentration inequalities, there are many future research directions to pursue. The first direction is to study the single index model or the additive regression in time series setting using the framework of the functional dependence measure. Another interesting problem is the robust regularized estimation in some special high dimensional time series models, such as vector autoregressive (VAR) model, vector autoregressions with heteroskedasticity, etc. By exploiting the model structure more explicitly, better statistical error bounds may be established. In addition, though many inferential theories have been developed recently for high dimensional VAR models, the statistical inference problem in high dimensional time series regression remains an important and challenging task to investigate.

Acknowledgements

We would like to thank the Editor, the Associate Editor and the anonymous referees for their detailed and constructive reviews, which helped to improve the paper substantially. We would also like to thank Professor Cun-Hui Zhang for helpful comments and discussions.

Supplementary Material

Supplementary Material to “High dimensional generalized linear models for temporal dependent data” (DOI: [10.3150/21-BEJ1451SUPP](https://doi.org/10.3150/21-BEJ1451SUPP); .pdf). In the supplementary material, we shall provide a real data application, extension to quasi log-likelihood loss, the proofs of main results in the paper and some lemmas that are useful in proofs of the paper. The readers are referred to Appendix A for concentration inequalities for high dimensional time series. Appendix B discusses quasi log-likelihood loss. Appendix C includes a real data application. The proofs of main theorems are presented in Appendix D.

References

- [1] Avella-Medina, M. and Ronchetti, E. (2018). Robust and consistent variable selection in high-dimensional generalized linear models. *Biometrika* **105** 31–44. [MR3768863 https://doi.org/10.1093/biomet/asx070](https://doi.org/10.1093/biomet/asx070)
- [2] Basu, S. and Michailidis, G. (2015). Regularized estimation in sparse high-dimensional time series models. *Ann. Statist.* **43** 1535–1567. [MR3357870 https://doi.org/10.1214/15-AOS1315](https://doi.org/10.1214/15-AOS1315)
- [3] Bianco, A.M. and Yohai, V.J. (1996). Robust estimation in the logistic regression model. In *Robust Statistics, Data Analysis, and Computer Intensive Methods (Schloss Thurnau, 1994)*. *Lect. Notes Stat.* **109** 17–34. New York: Springer. [MR1491394 https://doi.org/10.1007/978-1-4612-2380-1_2](https://doi.org/10.1007/978-1-4612-2380-1_2)
- [4] Bickel, P.J., Ritov, Y. and Tsybakov, A.B. (2009). Simultaneous analysis of lasso and Dantzig selector. *Ann. Statist.* **37** 1705–1732. [MR2533469 https://doi.org/10.1214/08-AOS620](https://doi.org/10.1214/08-AOS620)
- [5] Blundell, C., Beck, J. and Heller, K.A. (2012). Modelling reciprocating relationships with Hawkes processes. In *Advances in Neural Information Processing Systems* 2600–2608.
- [6] Bousquet, O. (2002). A Bennett concentration inequality and its application to suprema of empirical processes. *C. R. Math. Acad. Sci. Paris* **334** 495–500. [MR1890640 https://doi.org/10.1016/S1631-073X\(02\)02292-6](https://doi.org/10.1016/S1631-073X(02)02292-6)
- [7] Brown, E.N., Kass, R.E. and Mitra, P.P. (2004). Multiple neural spike train data analysis: State-of-the-art and future challenges. *Nat. Neurosci.* **7** 456.
- [8] Brownlees, C., Joly, E. and Lugosi, G. (2015). Empirical risk minimization for heavy-tailed losses. *Ann. Statist.* **43** 2507–2536. [MR3405602 https://doi.org/10.1214/15-AOS1350](https://doi.org/10.1214/15-AOS1350)
- [9] Bühlmann, P. and van de Geer, S. (2011). *Statistics for High-Dimensional Data: Methods, Theory and Applications*. *Springer Series in Statistics*. Heidelberg: Springer. [MR2807761 https://doi.org/10.1007/978-3-642-20192-9](https://doi.org/10.1007/978-3-642-20192-9)

- [10] Cantoni, E. and Ronchetti, E. (2001). Robust inference for generalized linear models. *J. Amer. Statist. Assoc.* **96** 1022–1030. [MR1947250](#) <https://doi.org/10.1198/016214501753209004>
- [11] Catoni, O. (2012). Challenging the empirical mean and empirical variance: A deviation study. *Ann. Inst. Henri Poincaré Probab. Stat.* **48** 1148–1185. [MR3052407](#) <https://doi.org/10.1214/11-AIHP454>
- [12] Chen, L. and Wu, W.B. (2017). Concentration inequalities for empirical processes of linear time series. *J. Mach. Learn. Res.* **18** Paper No. 231, 46 pp. [MR3845530](#)
- [13] Chen, X., Xu, M. and Wu, W.B. (2013). Covariance and precision matrix estimation for high-dimensional time series. *Ann. Statist.* **41** 2994–3021. [MR3161455](#) <https://doi.org/10.1214/13-AOS1182>
- [14] Chen, X., Xu, M. and Wu, W.B. (2016). Regularized estimation of linear functionals of precision matrices for high-dimensional time series. *IEEE Trans. Signal Process.* **64** 6459–6470. [MR3566612](#) <https://doi.org/10.1109/TSP.2016.2605079>
- [15] Ding, M., Mo, J., Schroeder, C.E. and Wen, X. (2011). Analyzing coherent brain networks with granger causality. In *2011 Annual International Conference of the IEEE Engineering in Medicine and Biology Society* 5916–5918. IEEE.
- [16] Fan, J., Gong, W. and Zhu, Z. (2019). Generalized high-dimensional trace regression via nuclear norm regularization. *J. Econometrics* **212** 177–202. [MR3994013](#) <https://doi.org/10.1016/j.jeconom.2019.04.026>
- [17] Fan, J., Li, Q. and Wang, Y. (2017). Estimation of high dimensional mean regression in the absence of symmetry and light tail assumptions. *J. R. Stat. Soc. Ser. B. Stat. Methodol.* **79** 247–265. [MR3597972](#) <https://doi.org/10.1111/rssb.12166>
- [18] Fan, J., Liu, H., Sun, Q. and Zhang, T. (2018). I-LAMM for sparse learning: Simultaneous control of algorithmic complexity and statistical error. *Ann. Statist.* **46** 814–841. [MR3782385](#) <https://doi.org/10.1214/17-AOS1568>
- [19] Fan, J., Wang, W. and Zhu, Z. (2021). A shrinkage principle for heavy-tailed data: High-dimensional robust low-rank matrix recovery. *Ann. Statist.* **49** 1239–1266. [MR4298863](#) <https://doi.org/10.1214/20-aos1980>
- [20] Fan, Y. and Lv, J. (2013). Asymptotic equivalence of regularization methods in thresholded parameter space. *J. Amer. Statist. Assoc.* **108** 1044–1061. [MR3174683](#) <https://doi.org/10.1080/01621459.2013.803972>
- [21] Genkin, A., Lewis, D.D. and Madigan, D. (2007). Large-scale Bayesian logistic regression for text categorization. *Technometrics* **49** 291–304. [MR2408634](#) <https://doi.org/10.1198/004017007000000245>
- [22] Guo, S., Wang, Y. and Yao, Q. (2016). High-dimensional and banded vector autoregressions. *Biometrika* **103** 889–903. [MR3620446](#) <https://doi.org/10.1093/biomet/asw046>
- [23] Hall, E.C., Raskutti, G. and Willett, R.M. (2019). Learning high-dimensional generalized linear autoregressive models. *IEEE Trans. Inf. Theory* **65** 2401–2422. [MR3928208](#) <https://doi.org/10.1109/TIT.2018.2884673>
- [24] Hampel, F.R. (1971). A general qualitative definition of robustness. *Ann. Math. Stat.* **42** 1887–1896. [MR0301858](#) <https://doi.org/10.1214/aoms/1177693054>
- [25] Hampel, F.R. (1974). The influence curve and its role in robust estimation. *J. Amer. Statist. Assoc.* **69** 383–393. [MR0362657](#)
- [26] Han, F., Lu, H. and Liu, H. (2015). A direct estimation of high dimensional stationary vector autoregressions. *J. Mach. Learn. Res.* **16** 3115–3150. [MR3450535](#)
- [27] Han, Y. and Tsay, R.S. (2020). High-dimensional linear regression for dependent data with applications to nowcasting. *Statist. Sinica* **30** 1797–1827. [MR4260745](#) <https://doi.org/10.5705/ss.202018.0044>
- [28] Han, Y., Tsay, R.S. and Wu, W.B. (2023). Supplement to “High dimensional generalized linear models for temporal dependent data.” <https://doi.org/10.3150/21-BEJ1451SUPP>
- [29] Hill, R.W. (1977). *Robust Regression when There Are Outliers in the Carriers*. Ann Arbor, MI: ProQuest LLC. Thesis (Ph.D.)—Harvard University. [MR2940734](#)
- [30] Hodges, J.L. Jr. (1967). Efficiency in normal samples and tolerance of extreme values for some estimates of location. In *Proc. Fifth Berkeley Sympos. Math. Statist. and Probability (Berkeley, Calif., 1965/66)*, Vol. I 163–186. Berkeley, CA: Univ. California Press. [MR0214251](#)
- [31] Hsu, D. and Sabato, S. (2016). Loss minimization and parameter estimation with heavy tails. *J. Mach. Learn. Res.* **17** Paper No. 18, 40 pp. [MR3491112](#)
- [32] Huang, J., Sun, T., Ying, Z., Yu, Y. and Zhang, C.-H. (2013). Oracle inequalities for the LASSO in the Cox model. *Ann. Statist.* **41** 1142–1165. [MR3113806](#) <https://doi.org/10.1214/13-AOS1098>
- [33] Huang, J. and Zhang, C.-H. (2012). Estimation and selection via absolute penalized convex minimization and its multistage adaptive applications. *J. Mach. Learn. Res.* **13** 1839–1864. [MR2956344](#)

- [34] Huang, S.-J. and Shih, K.-R. (2003). Short-term load forecasting via arma model identification including non-Gaussian process considerations. *IEEE Trans. Power Syst.* **18** 673–679.
- [35] Huber, P.J. (1964). Robust estimation of a location parameter. *Ann. Math. Stat.* **35** 73–101. [MR0161415 https://doi.org/10.1214/aoms/1177703732](https://doi.org/10.1214/aoms/1177703732)
- [36] Huber, P.J. (1967). The behavior of maximum likelihood estimates under nonstandard conditions. In *Proc. Fifth Berkeley Sympos. Math. Statist. and Probability (Berkeley, Calif., 1965/66), Vol. I: Statistics* 221–233. Berkeley, CA: Univ. California Press. [MR0216620](https://doi.org/10.1214/aoms/1177703732)
- [37] Ivanoff, S., Picard, F. and Rivoirard, V. (2016). Adaptive Lasso and group-Lasso for functional Poisson regression. *J. Mach. Learn. Res.* **17** Paper No. 55, 46 pp. [MR3504615](https://doi.org/10.1214/aoms/1177703732)
- [38] Jiang, X., Raskutti, G. and Willett, R. (2015). Minimax optimal rates for Poisson inverse problems with physical constraints. *IEEE Trans. Inf. Theory* **61** 4458–4474. [MR3372365 https://doi.org/10.1109/TIT.2015.2441072](https://doi.org/10.1109/TIT.2015.2441072)
- [39] Krishnapuram, B., Carin, L., Figueiredo, M.A.T. and Hartemink, A.J. (2005). Sparse multinomial logistic regression: Fast algorithms and generalization bounds. *IEEE Trans. Pattern Anal. Mach. Intell.* **27** 957–968.
- [40] Linderman, S. and Adams, R. (2014). Discovering latent network structure in point process data. In *International Conference on Machine Learning* 1413–1421.
- [41] Loh, P.-L. (2017). Statistical consistency and asymptotic normality for high-dimensional robust M -estimators. *Ann. Statist.* **45** 866–896. [MR3650403 https://doi.org/10.1214/16-AOS1471](https://doi.org/10.1214/16-AOS1471)
- [42] Loh, P.-L. and Wainwright, M.J. (2017). Support recovery without incoherence: A case for nonconvex regularization. *Ann. Statist.* **45** 2455–2482. [MR3737898 https://doi.org/10.1214/16-AOS1530](https://doi.org/10.1214/16-AOS1530)
- [43] Lokhorst, J. (1999). *The Lasso and Generalised Linear Models. Honors Project*. Australia: The Univ. Adelaide.
- [44] Mallows, C.L. (1975). On some topics in robustness. Unpublished memorandum, Bell Telephone Laboratories, Murray Hill, NJ.
- [45] Mark, B., Raskutti, G. and Willett, R. (2019). Network estimation from point process data. *IEEE Trans. Inf. Theory* **65** 2953–2975. [MR3951378 https://doi.org/10.1109/TIT.2018.2875766](https://doi.org/10.1109/TIT.2018.2875766)
- [46] Massart, P. (2000). About the constants in Talagrand’s concentration inequalities for empirical processes. *Ann. Probab.* **28** 863–884. [MR1782276 https://doi.org/10.1214/aop/1019160263](https://doi.org/10.1214/aop/1019160263)
- [47] McCullagh, P. and Nelder, J.A. (1989). *Generalized Linear Models. Monographs on Statistics and Applied Probability*. London: CRC Press. [MR3223057 https://doi.org/10.1007/978-1-4899-3242-6](https://doi.org/10.1007/978-1-4899-3242-6)
- [48] Meier, L., van de Geer, S. and Bühlmann, P. (2008). The group Lasso for logistic regression. *J. R. Stat. Soc. Ser. B. Stat. Methodol.* **70** 53–71. [MR2412631 https://doi.org/10.1111/j.1467-9868.2007.00627.x](https://doi.org/10.1111/j.1467-9868.2007.00627.x)
- [49] Merlevède, F., Peligrad, M. and Rio, E. (2011). A Bernstein type inequality and moderate deviations for weakly dependent sequences. *Probab. Theory Related Fields* **151** 435–474. [MR2851689 https://doi.org/10.1007/s00440-010-0304-9](https://doi.org/10.1007/s00440-010-0304-9)
- [50] Merrill, H.M. and Schweppe, F.C. (1971). Bad data suppression in power system static state estimation. *IEEE Trans. Power Appar. Syst.* **6** 2718–2725.
- [51] Minsker, S. (2015). Geometric median and robust estimation in Banach spaces. *Bernoulli* **21** 2308–2335. [MR3378468 https://doi.org/10.3150/14-BEJ645](https://doi.org/10.3150/14-BEJ645)
- [52] Negahban, S.N., Ravikumar, P., Wainwright, M.J. and Yu, B. (2012). A unified framework for high-dimensional analysis of M -estimators with decomposable regularizers. *Statist. Sci.* **27** 538–557. [MR3025133 https://doi.org/10.1214/12-STS400](https://doi.org/10.1214/12-STS400)
- [53] Ogata, Y. (1999). Seismicity analysis through point-process modeling: A review. In *Seismicity Patterns, Their Statistical Significance and Physical Meaning* 471–507. Springer.
- [54] Pillow, J.W., Shlens, J., Paninski, L., Sher, A., Litke, A.M., Chichilnisky, E. and Simoncelli, E.P. (2008). Spatio-temporal correlations and visual signalling in a complete neuronal population. *Nature* **454** 995.
- [55] Priestley, M.B. (1988). *Nonlinear and Nonstationary Time Series Analysis*. London: Academic Press [Harcourt Brace Jovanovich, Publishers]. [MR0991969](https://doi.org/10.1109/TSP.2011.2157913)
- [56] Raginsky, M., Jafarpour, S., Harmany, Z.T., Marcia, R.F., Willett, R.M. and Calderbank, R. (2011). Performance bounds for expander-based compressed sensing in Poisson noise. *IEEE Trans. Signal Process.* **59** 4139–4153. [MR2865974 https://doi.org/10.1109/TSP.2011.2157913](https://doi.org/10.1109/TSP.2011.2157913)

- [57] Raginsky, M., Willett, R.M., Harmany, Z.T. and Marcia, R.F. (2010). Compressed sensing performance bounds under Poisson noise. *IEEE Trans. Signal Process.* **58** 3990–4002. [MR2780163](#) <https://doi.org/10.1109/TSP.2010.2049997>
- [58] Raginsky, M., Willett, R.M., Horn, C., Silva, J. and Marcia, R.F. (2012). Sequential anomaly detection in the presence of noise and limited feedback. *IEEE Trans. Inf. Theory* **58** 5544–5562. [MR2957451](#) <https://doi.org/10.1109/TIT.2012.2201375>
- [59] Raskutti, G., Wainwright, M.J. and Yu, B. (2011). Minimax rates of estimation for high-dimensional linear regression over ℓ_q -balls. *IEEE Trans. Inf. Theory* **57** 6976–6994. [MR2882274](#) <https://doi.org/10.1109/TIT.2011.2165799>
- [60] Rosenblatt, M. (1971). *Markov Processes. Structure and Asymptotic Behavior. Die Grundlehren der Mathematischen Wissenschaften* **184**. New York: Springer. [MR0329037](#)
- [61] Rosenthal, H.P. (1970). On the subspaces of L^p ($p > 2$) spanned by sequences of independent random variables. *Israel J. Math.* **8** 273–303. [MR0271721](#) <https://doi.org/10.1007/BF02771562>
- [62] Roth, V. (2004). The generalized LASSO. *IEEE Trans. Neural Netw.* **15** 16–28. <https://doi.org/10.1109/TNN.2003.809398>
- [63] Shao, X. and Wu, W.B. (2007). Asymptotic spectral theory for nonlinear time series. *Ann. Statist.* **35** 1773–1801. [MR2351105](#) <https://doi.org/10.1214/009053606000001479>
- [64] Shevade, S.K. and Keerthi, S.S. (2003). A simple and efficient algorithm for gene selection using sparse logistic regression. *Bioinformatics* **19** 2246–2253.
- [65] Silva, J. and Willett, R. (2008). Hypergraph-based anomaly detection of high-dimensional co-occurrences. *IEEE Trans. Pattern Anal. Mach. Intell.* **31** 563–569.
- [66] Tibshirani, R. (1996). Regression shrinkage and selection via the lasso. *J. Roy. Statist. Soc. Ser. B* **58** 267–288. [MR1379242](#)
- [67] Tibshirani, R. (1997). The lasso method for variable selection in the Cox model. *Stat. Med.* **16** 385–395.
- [68] Tong, H. (1990). *Nonlinear Time Series: A Dynamical System Approach. Oxford Statistical Science Series* **6**. New York: Oxford Univ. Press, The Clarendon Press. [MR1079320](#)
- [69] Tsay, R.S. (2005). *Analysis of Financial Time Series*, 2nd ed. *Wiley Series in Probability and Statistics*. Hoboken, NJ: Wiley Interscience. [MR2162112](#) <https://doi.org/10.1002/0471746193>
- [70] Tukey, J.W. (1960). A survey of sampling from contaminated distributions. In *Contributions to Probability and Statistics* 448–485. Stanford, CA: Stanford Univ. Press. [MR0120720](#)
- [71] Tukey, J.W. (1962). The future of data analysis. *Ann. Math. Stat.* **33** 1–67. [MR0133937](#) <https://doi.org/10.1214/aoms/1177704711>
- [72] van de Geer, S. and Müller, P. (2012). Quasi-likelihood and/or robust estimation in high dimensions. *Statist. Sci.* **27** 469–480. [MR3025129](#) <https://doi.org/10.1214/12-STS397>
- [73] van de Geer, S.A. (2008). High-dimensional generalized linear models and the lasso. *Ann. Statist.* **36** 614–645. [MR2396809](#) <https://doi.org/10.1214/0090536070000000929>
- [74] Vere-Jones, D. and Ozaki, T. (1982). Some examples of statistical estimation applied to earthquake data. *Ann. Inst. Statist. Math.* **34** 189–207.
- [75] Wiener, N. (1958). *Nonlinear Problems in Random Theory. Technology Press Research Monographs*. New York: Wiley. [MR0100912](#)
- [76] Wu, W.-B. and Wu, Y.N. (2016). Performance bounds for parameter estimates of high-dimensional linear models with correlated errors. *Electron. J. Stat.* **10** 352–379. [MR3466186](#) <https://doi.org/10.1214/16-EJS1108>
- [77] Wu, W.B. (2005). Nonlinear system theory: Another look at dependence. *Proc. Natl. Acad. Sci. USA* **102** 14150–14154. [MR2172215](#) <https://doi.org/10.1073/pnas.0506715102>
- [78] Wu, W.B. (2007). M -Estimation of linear models with dependent errors. *Ann. Statist.* **35** 495–521. [MR2336857](#) <https://doi.org/10.1214/009053606000001406>
- [79] Wu, W.B. and Min, W. (2005). On linear processes with dependent innovations. *Stochastic Process. Appl.* **115** 939–958. [MR2138809](#) <https://doi.org/10.1016/j.spa.2005.01.001>
- [80] Wu, W.B. and Shao, X. (2004). Limit theorems for iterated random functions. *J. Appl. Probab.* **41** 425–436. [MR2052582](#) <https://doi.org/10.1239/jap/1082999076>
- [81] Zhang, C., Guo, X., Cheng, C. and Zhang, Z. (2014). Robust-BD estimation and inference for varying-dimensional general linear models. *Statist. Sinica* **24** 653–673. [MR3235393](#)

- [82] Zhang, D. (2021). Robust estimation of the mean and covariance matrix for high dimensional time series. *Statist. Sinica* **31** 797–820. [MR4286195](#) <https://doi.org/10.5705/ss.20>
- [83] Zhang, D. and Wu, W.B. (2017). Gaussian approximation for high dimensional time series. *Ann. Statist.* **45** 1895–1919. [MR3718156](#) <https://doi.org/10.1214/16-AOS1512>
- [84] Zhou, H.H. and Raskutti, G. (2019). Non-parametric sparse additive auto-regressive network models. *IEEE Trans. Inf. Theory* **65** 1473–1492. [MR3923181](#) <https://doi.org/10.1109/TIT.2018.2849988>
- [85] Zhou, K., Zha, H. and Song, L. (2013). Learning social infectivity in sparse low-rank networks using multi-dimensional Hawkes processes. In *Artificial Intelligence and Statistics* 641–649.

Received December 2020 and revised November 2021