

中文图书分类号: O212.8

UDC: 519.2

学 校 代 码: 10005



硕 士 学 位 论 文

MASTERAL DISSERTATION

论 文 题 目: 基于一类 l_0 收缩指数先验的
模型选择方法研究

论 文 作 者: 邹庆春

学 科: 统 计 学

指 导 教 师: 韩 敏

论文提交日期: 2022 年 5 月

UDC : 519.2
中文图书分类号 : O212.8

学校代码 : 10005
学 号 : S201906108

北京工业大学理学硕士学位论文

题 目 : 基于一类 l_0 收缩指数先验的模型选择方法研究
英文题目 : THE RESEARCH OF MODEL SELECTION METHOD WITH A
CLASS OF l_0 SHRINKAGE EXPONENTIAL PRIOR

论 文 作 者 : 邹庆春
学 科 专 业 : 统 计 学
研 究 方 向 : 非参数统计与数据分析
申 请 学 位 : 理学硕士
指 导 教 师 : 韩 敏
所 在 单 位 : 理 学 部
答 辩 日 期 : 2022 年 5 月
授予学位单位 : 北京工业大学

摘要

随着信息技术的飞速发展，数据维度和样本数大幅增加，大量无关和冗余的数据特征蕴藏其中，为数据挖掘和机器学习算法的应用带来巨大挑战。因此，从众多特征中选择出重要特征是数据处理中的一项重要任务。

在过去的线性模型研究中，指数权重法 (Exponential Weights, EW) 在变量选择、估计和预测误差方面有了较为突出的结果。然而，EW 方法在偏低信噪比 (Signal to Noise Ratio, SNR) 数据情况下，往往不能够正确识别出真实模型。本文通过对模型假设空间大小的讨论，为模型空间引入合适的先验，进而在子同一模型子空间中，构造关于系数向量的支撑集的先验分布，我们将该先验分布求系数估计的方法称之为推广指数权重法 (Extended Exponential Weights, EEW)。本文从理论分析的角度，建立了在该稀疏先验分布下所选模型的预测误差的 Oracle 不等式。其次，证明了在满足可识别性条件下，只要系数向量非零元素足够大，使得后验概率最大化的支撑集依概率收敛到真实系数向量的支撑集。然后，建立了系数向量的贝叶斯估计在三种不同误差度量下的误差界，表明在本文方法下得到的参数估计是相合的，即当样本容量趋于无穷时，所求得的贝叶斯估计依概率收敛到真实系数向量。最后，本文通过数值模拟证实了，在偏低 SNR 情况下，本文方法能够得到更低估计误差的系数估计和更高的非零系数支撑恢复率，且在高 SNR 情况下，与其他变量选择方法相比，本文方法同样表现优异。

关键词：线性模型；指数权重；贝叶斯估计； l_0 范数；SNR

Abstract

With the rapid development of information technology, the data dimension and the number of samples have greatly increased, and a large number of irrelevant and redundant data features are hidden in it, which brings great challenges to the application of data mining and machine learning algorithms. Therefore, Selecting important features from many features is an important task in data processing.

In the past linear model research, when the noise of the data is low, the exponential weight method(EW) has more prominent results in variable selection, estimation and prediction errors. However, the EW often cannot correctly identify the real model in the case of low Signal-to-Noise Ratio(SNR) data. By discussing the size of the model hypothesis space, this thesis constructs a suitable prior to the model space, and then construct the prior distribution of the support set about the coefficient vector in the same subspace. We call the method for solving parameter estimation under the constructed prior distribution as the Extended Exponential Weights method(EEW). From the perspective of theoretical analysis, this thesis establishes the Oracle inequality of the chosen model under this sparse prior distribution. Secondly, it is proved that under the condition of identifiability, as long as the non-zero elements of the coefficient vector are large enough, the support that maximizes the posterior probability converges in probability to the support of the true coefficient vector. Furthermore, the error boundaries of the Bayesian estimation of the coefficient vector under three different measures are established, which shows that the parameter estimates obtained under the method in this thesis are consistent, i.e., when the sample size tends to infinity, the obtained Bayesian estimation converges in probability to the true coefficient vector. Finally, this thesis confirms through numerical simulation that under the condition of low SNR, the method in this thesis can obtain coefficient estimation with lower estimation error and higher non-zero coefficient support recovery rate. And under the condition of high SNR, compared with other variable selection methods, our method also performs well.

Keywords: linear model; exponential weight; Bayesian estimation; l_0 -norm; SNR

符号表

R^p	p 维欧氏空间
X_i	矩阵 X 的第 i 列
X_J	矩阵 X 关于集合 J 的子矩阵
I_n	n 阶单位矩阵
$\ \cdot\ _1$	l_1 范数
$\ \cdot\ _2$	欧几里得范数或 l_2 范数
$\ \cdot\ _\infty$	无穷范数或 l_∞ 范数
$ J $	集合 J 中的元素个数
s	集合 J 的元素个数
S	所有集合 J 组成的模型空间
L_J	由子矩阵 X_J 的各列张成的线性空间
ψ_J	线性空间 L_J 上的正交投影矩阵
$I_{\{0 < J < l\}}$	与 $ J $ 相关的示性函数
Ω_p	由整数 1 到 p 构成的集合
$K(\Omega_p)$	集合 Ω_p 的幂集
$r(X)$	矩阵 X 的秩
$tr(X)$	矩阵 X 的迹
$a \wedge b$	实数 a 和 b 中的较大值
$a \vee b$	实数 a 和 b 中的较小值

目 录

第 1 章 绪论.....	1
1.1 研究背景及意义.....	1
1.2 国内外研究现状.....	2
1.3 主要研究内容和结构安排	4
第 2 章 预备知识	5
2.1 模型介绍.....	5
2.2 贝叶斯估计.....	6
2.3 常见的线性模型稀疏学习方法	6
2.3.1 Lasso	6
2.3.2 SCAD	7
2.3.3 MCP	8
2.3.4 OMP	8
2.4 马尔科夫链蒙特卡罗法.....	9
2.5 指数权重法.....	10
2.6 SBR 算法	10
2.7 本章小结.....	11
第 3 章 基于 l_0 收缩指数先验的贝叶斯估计及其渐近性质	13
3.1 l_0 收缩指数先验分布与后验分布	13
3.2 贝叶斯估计.....	14
3.3 渐近性质.....	15
3.3.1 Oracle 不等式.....	16
3.3.2 支撑恢复.....	16
3.3.3 参数估计.....	17
3.4 理论证明.....	18
3.4.1 引理及其证明.....	18
3.4.2 定理证明.....	23
3.5 本章小结.....	34
第 4 章 稀疏学习算法及数值模拟	35
4.1 稀疏学习算法.....	35
4.2 数值模拟.....	36
4.3 本章小节.....	44
总结与展望.....	45
参考文献.....	47

第1章 绪论

本章内容主要介绍模型选择的研究背景,进行稀疏学习的意义以及在贝叶斯框架下进行线性模型的模型选择的当前国内外研究现状。概述了模型选择问题,对本文的结构框架进行了简要说明。

1.1 研究背景及意义

近年来现代科学技术飞速发展,对社会和科学研究产生了深远影响,数据的维度呈现高维的特性,且人们能够通过多种方式便捷的收集到大量的高维数据。随着数据维度的增加,必定会有大量的无用特征包含其中,增加了数据模型学习过程中的难度,为数据挖掘和机器学习算法的应用带来困难,因此数据降维成为日常生活中处理数据必不可少的环节。在对高维复杂数据进行建模分析时,稀疏学习技术在模型选择方面表现出较为突出的优势,如高光谱图像学习^[1]、雷达架构设计^[2]、机器学习^[3,4]等方面,其能够在众多冗余信息中进行筛选,学习到重要数据特征,实现稀疏学习的目的。而稀疏性,指的就是数据或者模型趋于简单。在研究具体问题过程中,数据的维度通常都非常高,甚至超过了所能收集到的样本数,且对问题本身的研究有用的特征往往只占一小部分,因此在对高维复杂数据进行建模分析的过程中,进行相应的模型选择是十分必要的。在进行模型选择问题研究时,通过增加额外的对模型的约束条件,能够使得所学模型既能拥有良好的数据保真度又有良好的稀疏性,甚至趋于真实的未知模型。

对于线性模型的模型选择问题研究,已有大量学者通过构造模型系数向量的稀疏先验分布或通过引入限制模型空间大小的约束条件,从而实现模型选择的目的,也取得了较为丰富的成果^[5-7]。然而,相较于其他一些连续的收缩先验形式估计方法,从 l_0 范数作为研究的切入点,在理论结果上要比大部分连续的估计方法要突出,但 l_0 问题的研究在方法的实现上较为困难,因为 l_0 范数为非连续且不可导的,但同样也有许多较为成熟的算法去解决这一问题,如混合整数法^[8,9],惩罚分解法^[10,11]等。本文基于 l_0 范数的特性,从贝叶斯的角度出发,构造关于模型 l_0 范数的收缩指数先验分布,进行线性模型的模型选择问题研究。

1.2 国内外研究现状

考虑如下标准的线性回归模型, 响应向量 $y \in R^n$, 数据矩阵 $X \in R^{n \times p}$, 参数向量 $\beta \in R^p$, 高斯白噪声 w , i.e. $w \sim N(0, \sigma^2 I_n)$ 。

$$y = X\beta + w, \quad (1-1)$$

在机器学习和统计文献中, 总是希望获得一个具有良好数据保真度的稀疏模型。对于线性模型的稀疏学习问题, 已经许多文章从不同方面给广大学者提供了十分丰富的思路。从 l_0 范数的稀疏约束优化的角度, Mallat^[12] 等人提出匹配追踪算法 (Matching Pursuit, MP), 其目的就是解决 l_0 范数约束下的线性回归模型的模型选择问题, 这种方法的核心思想在于将每次迭代过程能够使得残差减少最多的数据特征加入到响应向量 y 的表示支撑中, 直到残差满足经验误差阈值后, 将最终所选择的数据特征作为该算法学习到的系数向量支撑。然而, 对于同一特征, MP 算法可能会重复选择。基于此缺点, Tropp^[13] 等人提出正交匹配追踪算法 (Orthogonal Matching Pursuit, OMP), 能够保证在后续所选择的数据特征不会重复出现, 该方法已经被广泛的运用在目标检测^[14] 和信号处理^[15] 等领域。然而, OMP 算法在每次迭代更新的过程中只考虑将单个特征加入到表示响应向量的支撑集合中, 当数据维度非常高时, 算法的运行效率将会变低, 于是许多改进的方法也相继出现^[16, 17]。除此之外, 还有许多基于匹配思想而不断被提出的新方法^[18-20], 且被运用到车辆健康检测^[21] 和冲击波测试^[22] 等方面。与匹配思想相近, 最优子集估计^[23] 方法同样被广大学者青睐, 该方法通过寻求获得最多具有 k 个非零回归系数的最佳最小二乘拟合, 从而达到模型选择的目的, 且已经有大量的文章研究了该估计器的理论性质^[24-26]。

由于 l_0 范数是非连续且不可导的, 无法通过类似梯度优化的方式求解近似解, 因此运用 l_0 范数求解问题相对复杂。于是, 连续形式的惩罚方法吸引了许多学者进行研究, Tibshirani^[27] 提出最小绝对收缩与选择算子 (Least Absolute Shrinkage and Selection Operator, Lasso) 方法, 该方法通过引入参数向量的 l_1 范数惩罚, 求解惩罚最小二乘的解, 该惩罚能够将参数估计较小的变量压缩为 0, 获得关于模型的稀疏解, 实现模型选择的目的, 且能够运用到高维数据的变量选择问题研究中, 因此被许多学者所熟知。然而, 该方法缺乏 Oracle 性质, Fan 和 Li^[28] 在其基础上, 提出了平滑剪切绝对偏差 (Smoothly Clipped Absolute Deviation, SCAD) 方法, 该方法所得参数估计具有稀疏性, 且能够逼近真实模型系数。同样的, Zou^[29] 阐述了 Lasso 方法所得参数估计缺乏 Oracle 性质, 于是在其基础上提出了 Adaptive Lasso 方法, 对待估系数向量不同位置的参数配置不同的惩罚系数, 使得所求参数估计具有 Oracle 性质。随着研究的不断深入, 基于 Lasso 思想已经拓展出许多不同的版

本,逐渐完善了 Lasso 方法的不足,如 Relaxed Lasso^[30], Fused Lasso^[31], Group Lasso^[32]等。除了使用单一范数约束进行稀疏学习外,使用混合范数形式进行稀疏学习也有出现,Zou^[33]等人提出弹性网络算法 (Elastic Net, EN),通过引入待估系数向量的 l_1 范数和 l_2 范数惩罚,使得可以学习得到稀疏模型的同时,还具有岭回归的优点,使得该方法适用范围变得更广,如在对具有相关性的数据进行特征选择中表现优异。Mazumder^[34]等人研究了稀疏线性建模中 SNR 低时的过拟合方面,通过在最优子集选择^[23]方法中添加 l_1 或 l_2 范数惩罚形式来防止过拟合,并对该方法的预测特性进行了深入研究。Haider^[36]等人通过 l_0 和 l_1 范数的混合进行了欠采样稀疏信号的恢复研究。

在贝叶斯框架下,通过引入待估参数的先验信息,来求解线性模型的稀疏学习问题也是多年来的研究热点。基于 l_0 惩罚回归的指数权重方法^[37-42]已经在预测性能方面有了许多比较好的成果。Arias 和 Lounici^[43]同样在贝叶斯框架下提出了基于指数权重的稀疏先验,建立了预测误差和估计误差性能方面的 Oracle 不等式,证明了在满足可识别性条件且只要非零系数足够大的情况下,所学得的参数估计的支撑依概率收敛到真实支撑。此外,在贝叶斯稀疏学习领域, Spike-and-Slab^[44]形式的先验一直是比较优秀的标准,在过去的十几年在不断的改善,George 和 McCulloch^[45]将零均值高斯分布作为 Spike 部分的先验,并利用随机搜索思想,构造了相应的求解算法。Ishwaran 和 Rao^[46]研究了连续双峰先验对超方差参数建模的有用性,以及研究了惩罚项对后验均值的影响,并证明了使用重新缩放的 Spike-and-Slab 模型的后验来实现选择性收缩在风险错误分类方面进行有效变量选择的重要性。Rockova^[47]提出了 SS-Lasso(Spike-and-Slab Lasso)方法,通过引入新的自适应惩罚函数,进一步提高惩罚似然估计的性能。Polson^[48]等人考虑将两种不同的分布混合,阐述了贝叶斯正则化和近端更新之间的联系,推导证明了寻找的后验形式和具有不同正则化先验的后验均值之间的等价性,并基于 Soussen^[49]等人提出的单一最优替换 (Single Best Replacement, SBR) 算法进行求解混合分布下的等价优化问题。其他形式的先验,也常被用来求解线性模型的稀疏问题,如 Horseshoe 先验^[50]^[51], Dirichlet-Laplaces 先验^[52], Horseshoe+ 先验^[53]等。Point-mass 形式的先验同样是近年来研究的热点,Castillo 等人^[54]基于该先验研究了高维线性回归下的全贝叶斯方法,证明了在设计矩阵在满足可识别性条件下,能够以较快的速度恢复未知的稀疏向量。Martin^[55]等人考虑高维线性模型,从经验贝叶斯的方法的角度,证明了该方法进行模型选择的有效性。以 Point-mass 形式先验为基础,Belister 和 Ghosal^[56]研究了高维线性模型下的变量选择性质。Bai^[57]等人在高维线性模型推断中,提出了一种变分贝叶斯方法,对该方法所得参数估计的理论性质进行了较为详细的分析。

1.3 主要研究内容和结构安排

在线性模型下, 本文通过对模型空间的简要讨论, 构造了基于待估系数向量的 l_0 收缩指数先验分布, 通过构造具有先增后减趋势的关于模型子空间的先验分布, 从而引出本文的稀疏先验分布。该稀疏先验与其他常见的先验形式不同, 本文是先对支撑集所在模型子空间构造先验, 进而在同一子空间下, 再对系数向量的支撑集定义相应的先验分布, 在此先验信息下, 建立了在参数估计、预测误差和支撑恢复三个方面的理论结果。基于所建立的理论结果, 构造了相应的稀疏学习算法, 并应用该算法进行了充分的数值模拟。在 SNR 偏低时, 通过数值模拟验证了本文方法的相比于传统指数权重方法的优越性, 与此同时, 在高 SNR 情况下, 本文方法在大多数情况下不论是参数估计还是支撑恢复方面, 均优于其他常见的变量选择方法。本文的各个章节涉及内容如下:

第一章, 叙述了本文的研究背景, 进行稀疏学习的意义以及在贝叶斯框架下进行线性模型的模型选择的当前国内外研究现状, 叙述了本文的基本思路, 同时简要概括了 l_0 收缩指数先验下进行模型选择的问题, 并对本文的结构框架进行概述。

第二章, 阐明了本文研究的模型, 简述了贝叶斯估计的思想, 呈现了与本文相似的文献, 并对后续数值模拟过程中的对比方法进行了简要介绍。

第三章, 通过对模型空间的简要讨论, 构造了本文所研究的 l_0 收缩指数先验分布以及构造该形式先验的原因, 推导了该 l_0 收缩指数先验分布所对应的后验分布, 并展示了本文所构造的先验分布对应的贝叶斯估计的求解思路。最后, 建立了本文方法所得到的贝叶斯估计的理论性质并进行了详细的证明推导。

第四章, 在第三章的理论基础上, 构造了基于推广指数权重的单一最优替换 (Extended Exponential Weights with Single Best Replacement, EEW-SBR) 稀疏学习算法, 用于解决本文的模型选择问题, 完成相关数值模拟。在偏低 SNR 情况下, 与 Arias^[43] 等人的指数权重法进行对比, 突出本文的适用性。在高 SNR 情况下, 还与其他几种常见的线性模型变量选择方法进行比较, 验证本文方法的优越性。

第2章 预备知识

本章节主要介绍所研究的基本模型以及后续研究过程中，设计矩阵所满足的基本条件，简述了贝叶斯估计的思想，呈现了与本文相似的文献。此外，本章节还展示了后续章节出现的对比方法的基本形式以及对应的基础解。

2.1 模型介绍

本文研究如下形式的标准线性模型，

$$y = X\beta_* + w, \quad w \sim N(0, \sigma^2 I_n) \quad (2-1)$$

相应的介绍已在 1.2 节展示，在理论研究过程中，噪声方差 σ^2 被当作为已知量，可以通过预先设置的 SNR 求解得到。在对实际数据进行分析建模时， σ^2 可通过样本的方差进行估计。在数值模拟的实现中，不同的 σ^2 的值是通过设置不同的 SNR 来求得。

在一般情况下，满足 (2-1) 的系数向量 β_* 是不唯一的，即模型是不可识别的，在理论结果的构建过程中，需要增加额外的条件来确保模型的可识别性。真实模型中系数向量的支撑集由 J_* 表示，相应的，其元素个数由 s_* 表示。本文所研究的情形是样本量 n 小于维度 p 且真实模型稀疏度 s_* 远小于 p 的情形，当然对于大 n 小 p 问题同样适用，本文在后续的模拟实验部分充分考虑了多种情形。对任意 $x = (x_1, \dots, x_p)^T \in R^p$ ，其中 $p \geq 1$ 和 $q \geq 1$ ，定义

$$\|x\|_q = \left(\sum_{j=1}^p |x_j|^q \right)^{1/q}, \quad \|x\|_\infty = \max_{1 \leq j \leq p} |x_j|. \quad (2-2)$$

对于设计矩阵 X ，其标准化后形式如下

$$\frac{1}{\sqrt{n}} \|X_j\|_2 = 1, \quad 1 \leq j \leq p. \quad (2-3)$$

本文从贝叶斯估计的角度出发，通过构造关于系数向量支撑集的 l_0 收缩指数先验分布，进而推导得到相应的稀疏贝叶斯估计，从而实现模型选择的目的。

2.2 贝叶斯估计

假设 $X = \{X_1, X_2, \dots, X_n\}$ 是已知的 n 个随机观测样本, θ 为连续型的未知参数或参数向量, 从贝叶斯角度, 参数 θ 是随机的, 且通常假定其服从某个分布 $\pi(\theta)$, 反映了人们对待估参数 θ 的经验认识^[7]。将样本在待估参数 θ 下的联合概率密度函数看作是关于 θ 的函数, 由贝叶斯公式,

$$\Pi(\theta) = P(\theta|X) = \frac{P(X|\theta)\pi(\theta)}{P(X)}, \quad (2-4)$$

参数 θ 的后验分布 $\Pi(\theta)$ 在 θ 先验信息的基础上, 将样本信息考虑在内, 是进行参数的贝叶斯估计的基础, 本文基于最大后验估计和后验期望估计^[66], 来对参数 θ 进行估计, 以下是对两种贝叶斯估计的概述。

首先介绍最大后验估计, 最大后验估计的基本思想在于将使得后验分布概率达到最大的 $\hat{\theta}_{map}$ 作为参数 θ 的估计^[66], 由如下等式得到,

$$\hat{\theta}_{map} = \arg \max_{\theta \in \Theta} \Pi(\theta) = \arg \max_{\theta \in \Theta} P(X|\theta)\pi(\theta) \quad (2-5)$$

事实上, $P(X)$ 与参数 θ 是无关的, 故在关于参数 θ 最大化过程中, 可看作是无影响的常数, 因此 (2-5) 式是成立的。对于后验期望估计 $\hat{\theta}_{mean}$, 是将后验期望作为参数的估计^[66], 因此有

$$\hat{\theta}_{mean} = \int_{\Theta} \theta \Pi(\theta) d\theta. \quad (2-6)$$

后验期望估计从后验分布角度出发, 通过求解积分形式得到其估计, 当 θ 为离散型随机向量时同理。

2.3 常见的线性模型稀疏学习方法

2.3.1 Lasso

Lasso^[27] 模型的基本形式为

$$\hat{\beta} = \arg \min_{\beta \in R^p} \frac{1}{2} \|y - X\beta\|_2^2 + \lambda \|\beta\|_1, \lambda \in [0, \infty) \quad (2-7)$$

其中等式右边的第一项衡量的模型拟合的优良性, 第二项为惩罚项, 一定程度上能够平衡拟合项与模型的复杂性。 λ 表示惩罚参数, 用于衡量拟合项与惩罚项的相对重要程度。 λ 越小, 所选择的变量就越多, 随着 λ 逐渐增大, 参数的收缩量变大,

最终全部收缩为 0。Lasso 方法以牺牲无偏性达到低方差的目的，其可以通过梯度优化的形式不断求解问题，例如利用坐标下降法不断迭代得到，当设计矩阵 X 列正交时，Lasso 方法所求得参数估计具有以下形式：

$$\hat{\beta}_i^{Lasso} = \text{sign}(\hat{\beta}_i)(|\hat{\beta}_i| - \lambda)_+ = \begin{cases} \hat{\beta}_i - \lambda, & \hat{\beta}_i > \lambda, \\ 0, & -\lambda \leq \hat{\beta}_i \leq \lambda, \\ \hat{\beta}_i + \lambda, & \hat{\beta}_i < -\lambda, \end{cases} \quad (2-8)$$

其中， $i = 1, \dots, p$ ， $\hat{\beta}$ 为参数最小二乘估计。由上述解的形式可知，当 $\hat{\beta}_i$ 处于 $[-\lambda, \lambda]$ 范围内，就会被压缩至 0，上述解的形象的解释了该方法能获得稀疏的原因，因此可以通过调整惩罚参数 λ 来获取稀疏解，达到变量选择的目的。

2.3.2 SCAD

Fan 和 Li^[28] 提出变量选择方法——SCAD，证明了该方法下参数估计具有 Oracle 性质，其模型基本形式为

$$\hat{\beta} = \arg \min_{\beta \in R^p} \frac{1}{2} \|y - X\beta\|_2^2 + P_\lambda(\beta), \quad (2-9)$$

惩罚函数 $P_\lambda(\beta)$ 定义为如下形式：

$$P_\lambda(\beta) = \begin{cases} \lambda|\beta|, & |\beta| \leq \lambda, \\ \frac{-(|\beta|^2 - 2a\lambda|\beta| + \lambda^2)}{2a-2}, & \lambda < |\beta| \leq a\lambda, \\ \frac{(a+1)\lambda^2}{2}, & a\lambda \leq |\beta|. \end{cases} \quad (2-10)$$

其中， $\lambda \geq 0$ 且 $a > 2$ ，当设计矩阵 X 列正交时，SCAD 方法的得到的参数估计值具有以下形式：

$$\hat{\beta}_i^{SCAD} = \begin{cases} \text{sign}(\hat{\beta}_i)(|\hat{\beta}_i| - \lambda)_+, & |\hat{\beta}_i| \leq 2\lambda, \\ \frac{(a-1)\hat{\beta}_i - \text{sign}(\hat{\beta}_i)a\lambda}{a-2}, & 2\lambda < |\hat{\beta}_i| \leq a\lambda, \\ \hat{\beta}_i, & a\lambda < |\hat{\beta}_i|, \end{cases} \quad (2-11)$$

其中， $\hat{\beta}$ 为参数的最小二乘估计。从 (2-11) 可知，当 $\hat{\beta}_i$ 取值很小时，与 Lasso 方法的解类似，该方法直接将 $\hat{\beta}_i$ 压缩为 0。当 $\hat{\beta}_i$ 的取值较大时，不会对 $\hat{\beta}_i$ 进行任何形式的处理。

2.3.3 MCP

极小极大凹惩罚方法 (Minimax Concave Penalty, MCP) 由 Zhang^[60] 提出, 其基本形式为

$$\hat{\beta} = \arg \min_{\beta \in R^p} \frac{1}{2} \|y - X\beta\|_2^2 + \sum_{i=1}^p \Gamma(|\beta_i|, \lambda, \gamma), \quad (2-12)$$

其中, λ, γ 均大于 0, $\Gamma(|\beta_i|, \lambda, \gamma)$ 是一个非凸函数, 在 0 处的紧邻域外为 0, 并且在 0 处具有一个非零的右导数, 如下所示:

$$\Gamma(|\beta_i|, \lambda, \gamma) = \begin{cases} \lambda |\beta_i| - \frac{|\beta_i|^2}{2\gamma}, & |\beta_i| \leq \gamma\lambda, \\ \frac{1}{2}\gamma\lambda^2, & |\beta_i| > \gamma\lambda. \end{cases} \quad (2-13)$$

其中, 参数 λ 的调整依赖于真实模型的稀疏度。

2.3.4 OMP

OMP^[13] 算法从 l_0 范数约束优化的角度出发, 该方法需要求解以下问题

$$\begin{aligned} & \min_{\beta} \|y - X\beta\|_2, \\ & s.t. \quad \|\beta\|_0 \leq s. \end{aligned} \quad (2-14)$$

在保证模型的稀疏度不大于 s 的约束条件下, 求解能够使得目标函数值达到最小的系数向量 β 。OMP 算法是基于贪婪搜索的思想, 通过将每次迭代过程能够使得误差达到最小的数据特征加入到响应向量 y 的表示支撑中, 直到残差满足一定的约束条件后, 即小于事先设定的阈值, 将最终所选择的数据特征作为该算法学习到的系数向量支撑^[66]。第 i 个分量对应的误差为

$$\begin{aligned} \varepsilon(i) &= \min_{\beta_i} \|x_i \beta_i - y\|_2^2 \\ &= \left\| \frac{x_i^T y}{\|x_i\|_2^2} x_i - y \right\|_2^2 \\ &= \|y\|_2^2 - \frac{(x_i^T y)^2}{\|x_i\|_2^2}. \end{aligned} \quad (2-15)$$

由上式可知, 通过计算设计矩阵的第 i 个分量与响应向量 y 的内积, 可找到使得误差达到最小的特征。记 $e^k = y - X\beta^k$ 表示第 k 步模型的残差向量。在迭代过程中逐步最小化误差, 利用残差向量对设计矩阵的第 i 个分量的上一步误差进行表示, 可

以得到

$$\begin{aligned}
 \varepsilon(i)_{k-1} &= \min_{\beta_i} \|x_i \beta_i - e^{k-1}\|_2^2 \\
 &= \left\| \frac{x_i^T e^{k-1}}{\|x_i\|_2^2} x_i - e^{k-1} \right\|_2^2 \\
 &= \|e^{k-1}\|_2^2 - \frac{(x_i^T e^{k-1})^2}{\|x_i\|_2^2}.
 \end{aligned} \tag{2-16}$$

因此, 在迭代求解过程中, 加入到表示响应向量 y 的设计矩阵分量是与残差做内积后最大的 x_{i_k} , 其伪代码总结在算法 2.1 中。

算法 2.1 OMP 算法^[66]

输入: 设计矩阵 X , 响应向量 y , 误差阈值 ε_0 。

输出: 迭代至满足停机准则的稀疏向量 β^k 。

- 1: 初始化: 迭代次数 $k=0$, 残差向量 $e^0=y$, 稀疏系数向量 $\beta^0=0$ 稀疏解的支撑集合 $J^0=\emptyset$ 。
- 2: 误差计算: $\forall i \in \{1, 2, \dots, p\}$, 计算 $\varepsilon(i)_{k-1} = \|e^{k-1}\|_2^2 - \frac{(x_i^T e^{k-1})^2}{\|x_i\|_2^2}$ 。
- 3: 支撑选取: 选取使得误差 $\varepsilon(i)_{k-1}$ 达到最小的支撑 i_k , $i_k = \arg \min_i \varepsilon(i)_{k-1}$ 。
- 4: 支撑集更新: 若 $\forall i \notin J^{k-1}$, 且 $\varepsilon(i_k) \leq \varepsilon_i$, 则 $J^k = J^{k-1} \cup i_k$ 。
- 5: 稀疏解求解: $\beta^k = \arg \min_{\beta_{J^k}} \|y - X\beta_{J^k}\|_2^2$ 。
- 6: 残差向量更新: $e^k = y - X\beta^k$ 。
- 7: 停止准则: 若 $\|e^k\|_2 < \varepsilon_0$, 终止循环, 输出 β^k 。否则, $k=k+1$, 返回步骤 2。

2.4 马尔科夫链蒙特卡罗法

马尔科夫链蒙特卡罗法 (Markov Chain Monte Carlo method, MCMC)^[66] 是在贝叶斯框架下, 引入马尔可夫链过程, 通过平稳分布上不断随机抽取样本, 最后通过平稳分布的样本进行近似计算与目标分布相关的高维积分。在实际应用中, 计算高维积分通常是有很大难度的, 基于此, MCMC 方法通过随机采样再取均值的方式, 当采样数足够多时, 能够对高维积分进行十分精确的近似, 以一种十分高效的方式解决了复杂高维积分的计算问题。对于参数的后验估计, MCMC 算法采样所得参数满足以下式子

$$\lim_{t \rightarrow \infty} \frac{1}{t-t_0} \sum_{n=t_0+1}^t \hat{\beta}_n = \hat{\beta}_{mean}, \tag{2-17}$$

其中, t 表示采样长度, t_0 表示在趋于平稳分布前的预烧采样次数, 可根据实际情况进行选择。随着采样次数 t 的不断增加趋于无穷, 对预烧后随机采样得到的参数序列再进行算术平均, 依据大数定律, 对随机采样结果取平均后会收敛到参数的后验期望估计, 因此该方法成为许多贝叶斯框架下求解后验期望估计的基础。更为详细的介绍请参考 Robert^[62] 的著作。

2.5 指数权重法

Arias 和 Lounici^[43] 的指数权重法 (Exponential Weights, EW) 使用一种快速压缩的方式, 构造关于待估系数支撑集的先验,

$$\pi(J) \propto \left(\frac{p}{|J|}\right)^{-1} e^{-m|J|} I_{\{|J| \leq l\}}, \quad (2-18)$$

其中 J 为待估系数的支撑集, m 为超参数, l 表示支撑大小的上界, 在没有其他更多信息的情况下通常设置为 p 。在求解最大后验估计时, 将求解最大后验估计问题转化求解以下等价的目标函数最小化问题,

$$f(J) = \frac{1}{2\sigma^2} \|y - X_J \beta_J\|^2 + m|J| + \log \left(\frac{p}{|J|}\right), \quad (2-19)$$

Arias 和 Lounici^[43] 证明了, 低噪声下, EW 法所求的参数估计在预测误差、参数估计和支撑恢复方面有较为理想的理论结果, 并利用 MCMC 方法进行了该先验分布下的后验期望估计的数值对比。然而, 在噪声偏大时, 即 SNR 偏低时, 相应的惩罚项的影响较高, 削弱了拟合部分的重要性, 从而导致参数估计以及支撑恢复的效果较为一般, 基于此, 本文引入了一个新的平衡项, 能够在低 SNR 情况下, 对过度的惩罚进行缓和。Arias 和 Lounici 并未对最大后验估计的效果进行模拟, 因此本文侧重于最大后验估计的模拟对比, 用到的求解算法将会在第四章介绍。

2.6 SBR 算法

SBR 算法^[49] 以最小化 l_0 范数惩罚回归函数为目标, 即以下式为优化函数,

$$f(J) = \|y - X_J \beta_J\|_2^2 + \lambda |J|, \quad (2-20)$$

其中 J 表示支撑集, 通过逐步选择最优支撑位置的方式, 贪婪搜索最优的支撑集, 在目标函数值不再减少时终止搜索。为了方便, 定义单次变换

$$J \cdot i = \begin{cases} J \cup i, & i \notin J, \\ J \setminus i, & i \in J. \end{cases} \quad (2-21)$$

SBR 算法步骤如下: 首先随机给定一个初始支撑集合 J , 通常为全集, 计算所有可能的单次变换下的目标函数值 $f(J \cdot i)$, $i \in \{1, 2, \dots, p\}$, 选择使得 p 个 $f(J \cdot i)$ 值中达到最小的支撑索引 i , 若此时 p 个值中的最小值不小于当前的目标函数值, 则

停止搜索，否则更新支撑集，即进行单次变换操作 $J \cdot i$ ，其伪代码总结在算法 2.2 中。

算法 2.2 SBR 算法 ^[49]

输入：设计矩阵 X ，响应向量 y ，超参数 λ 。

输出：迭代至满足停机准则的稀疏向量的支撑估计 $\hat{f}$ 以及系数估计 $\hat{\beta}$ 。

- 1: 随机初始化支撑集合: J_0 ，通常令 $J_0 = \emptyset$ 。
 - 2: 目标计算: $e_k = \min_{i \in \Omega_p} f(J_{k-1} \cdot i)$ ，以及 $j = \arg \min_{i \in \Omega_p} f(\beta_{J_{k-1} \cdot i})$ ， $J_k = J_{k-1} \cdot j$ 。
 - 3: 停机准则: 若 $e_k \geq e_{k-1}$ ，输出 $\hat{f} = J_{k-1}$ 并进行第 4 步；否则， $k = k + 1$ ，继续进行第 2 步。
 - 4: 系数计算: 若 $r(X_J) = |J|$ ，则 $\hat{\beta} = (X_J^T X_J)^{-1} X_J^T y$ ；否则 $\hat{\beta} = (X_J^T X_J)^{-} X_J^T y$ 。输出 $\hat{\beta}$ 。
-

2.7 本章小结

本章主要介绍了本文所研究的基本模型，贝叶斯估计的基本思想，马尔科夫链蒙特卡罗法的基本思想，以及与本文相关的研究方法。除此之外，简要介绍了后续数值模拟中所要对比方法的基本形式，并对 SBR 稀疏学习算法进行了介绍。

第3章 基于 l_0 收缩指数先验的贝叶斯估计及其渐近性质

通过对模型空间大小的讨论,本章构造关于系数向量的 l_0 收缩指数先验分布,并推导 l_0 收缩指数先验所对应的贝叶斯估计形式。将该稀疏先验分布下求解参数估计的方法称为推广指数权重法 (Extended Exponential Weights, EEW)。在理论性质方面,本章将会推导推广指数权重法下所得系数向量的贝叶斯估计的预测误差界限,待估系数支撑集的最大后验估计的一致性结论以及系数向量贝叶斯估计的误差边界。

3.1 l_0 收缩指数先验分布与后验分布

在 2.1 节已对所研究的模型进行过简要介绍,其模型形式为

$$y = X\beta_* + w \quad w \sim N(0, \sigma^2 I_n).$$

我们将模型空间 S 分为 p 个子空间 $S = \cup_{j=1}^p S_j$, 使得在任一子模型空间的集合具有相同的大小 j , 即具有相同个数的元素。此时子模型空间 S_j 中具有 $\binom{p}{j}$ 个模型, 对应于 p 个位置上排列组合个数。通常的做法是为同一个子模型空间中的子模型分配相同的概率, 即对任意 $J \in S_j$, $\pi(J|S_j) = 1/\binom{p}{j}$ 。而对于不同的子模型空间, 许多学者通常会根据越稀疏的模型或者说越稀疏的空间具有更高的概率的方式来构造合适的先验, 与之不同, 本文为子模型空间构造了一个具有先增后减趋势的先验, 在子空间的大小为 0 时, 模型的输出全部为噪声, 故在 0 处分配的先验概率为 0, 因此对于 S_j , 本文构造如下形式分布:

$$p(S_j) \propto |J|^{bm} e^{-m|J|} I_{\{0 < |J| \leq l\}}, \quad (3-1)$$

其中 $|J| = j$, 在后续的研究中为了方便, 我们记 $|J| = s$ 。通过上述对假设空间的先验分布的合理构建, 待估系数向量的支撑集 J 的先验分布形式将在后续给出。事实上, 系数向量的 l_0 范数为其支撑集的元素个数, 因此, 我们定义稀疏先验分布 $\pi(J)$ 为

$$\pi(J) \propto \frac{1}{\binom{p}{|J|}} |J|^{bm} e^{-m|J|} I_{\{0 < |J| \leq l\}}, \quad (3-2)$$

其中 l 为系数向量非 0 元个数的上界。该先验分布可通过设置 l , 来达到强制稀疏的目的, 在 $|J|$ 很大甚至与维度 p 接近时, 其概率趋于 0, 甚至在超过事先设定的

稀疏度上界 l 时, 强制为 0。在 (3-2) 中, m 与 b 是决定先验分布趋势的两个超参数, 满足 $m > 0$, $b \geq 0$ 。在实际分析过程中, 若无更多信息, 我们假定 $l = p$ 。在后续的模拟过程中, 高 SNR 对应于低噪声, 噪声 σ^2 大小可通过事先给定的 SNR 计算得到。记 X 为样本矩阵, 由贝叶斯公式, 我们可通过简单的计算转换得到后验分布 $\Pi(J)$ 为

$$\begin{aligned}\Pi(J) &= P(J|X) = \frac{P(X|J)\pi(J)}{P(X)} \\ &\propto P(X|J)\pi(J),\end{aligned}$$

对于似然函数 $P(X|J)$, 有

$$\begin{aligned}P(X|J) &= (2\pi)^{-\frac{n}{2}} \|\sigma^2 I_n\|^{-\frac{1}{2}} \exp\left(-\frac{1}{2}(y - X\beta_J)^T (\sigma^2 I_n)^{-1} (y - X\beta_J)\right) \\ &= (2\pi)^{-\frac{n}{2}} \sigma^{-n} \exp\left(-\frac{1}{2\sigma^2} (y - X\beta_J)^T (y - X\beta_J)\right) \\ &\propto \exp\left(-\frac{1}{2\sigma^2} \|y - X\beta_J\|_2^2\right).\end{aligned}\tag{3-3}$$

在本文中, 是先通过特定的准则寻找能够最大化后验概率的支撑集 J , 再其他方式对参数向量 β_* 进行估计。要最大化 $\Pi(J)$, 仅仅考虑支撑集 J 是无法实现的, 还需同时考虑支撑集对应位置上的取值。然而, 本文所构造的 l_0 收缩指数先验分布 (3-2) 中只包含系数向量支撑集的信息。因此, 当 $\Pi(J)$ 形式如 (3-2) 所示时, 无法仅通过选择支撑集而使得后验概率最大。我们注意到, 使得似然部分最大化的估计是该部分的最小二乘估计, 具有一个封闭的表达形式。因此, 本文考虑当支撑 J 固定时, 使得似然部分达到最大化的参数估计 $\hat{\beta}_J = (X_J^T X_J)^{-1} X_J^T y$, 将 $\hat{\beta}_J$ 代入似然部分, 可以得到 $\min_J \|y - X\beta_J\|_2^2 = \|\psi_J^+ y\|_2^2$ 。基于此, 通过计算后验概率最大化可以通过选择合适的支撑来解决, 因此, 处理之后的后验概率为

$$\Pi(J) \propto \pi(J) \exp\left(-\frac{1}{2\sigma^2} \|\psi_J^+ y\|_2^2\right).\tag{3-4}$$

在本文的理论研究中, 后验概率的形式均以 (3-4) 中的 $\Pi(J)$ 为准。

3.2 贝叶斯估计

在本文中, l_0 收缩指数先验分布是基于待估系数向量的支撑集的, 通过最大化 (3-4) 式得到的并不是系数向量的最大后验估计, 而是支撑集 J 的最大后验估计 $\hat{J}_{map}$ 。因此, 在得到 $\hat{J}_{map}$ 后, 再通过最小二乘估计, 即可以得到系数向量的参数估计。

因此, J 的最大后验估计由以下等式得到

$$\hat{J}_{map} = \arg \max_{J \in K(\Omega_p)} \Pi(J). \quad (3-5)$$

当设计矩阵的子矩阵 $X_{\hat{J}_{map}}$ 是列满秩时候, 即各列相互线性无关, 此时在 $\hat{J}_{map}$ 的基础上, 有

$$\hat{\beta}_{map}^* = (X_{\hat{J}_{map}}^T X_{\hat{J}_{map}})^{-1} X_{\hat{J}_{map}}^T y. \quad (3-6)$$

否则, 无法通过取逆矩阵的形式表示, 此时只能以广义逆形式表示, 此时有: $\hat{\beta}_{map}^* = (X_{\hat{J}_{map}}^T X_{\hat{J}_{map}})^- X_{\hat{J}_{map}}^T y$ 。由于支撑集 J 是一系列非零元所在位置组成的集合, 因此有 $\hat{\beta}_{map}^* \in R^{|\hat{J}_{map}|}$, 因此参数估计 $\hat{\beta}$ 需通过以零元素填充的形式得到, 在得到 $\hat{\beta}_{map}^*$ 后, 将 $\hat{\beta}_{map}^*$ 转换为 p 维向量。在新引入的 $p - |\hat{J}_{map}|$ 维度上, 系数的取值均设置为 0。

$$\hat{\beta}_{map}^{(i)} = \begin{cases} \hat{\beta}_{map}^{*(l_{(i)})} & i \in \hat{J}_{map}, \\ 0 & i \notin \hat{J}_{map}. \end{cases} \quad (3-7)$$

$l_{(i)}$ 表示 $\hat{\beta}$ 第 i 个分量 $\hat{J}_{map}$ 中的位置。因此, $\hat{\beta}_{map} \in R^p$ 。与此同时, 参数向量的后验期望估计具有以下形式

$$\hat{\beta}_{mean} = \sum_{J \in K(\Omega_p)} \Pi(J) \hat{\beta}_J, \quad (3-8)$$

3.3 渐近性质

本文所研究对象为标准的稀疏线性模型, 系数向量 β_* 是非常稀疏的, 因此通常假定有 $s_* \ll p$ 。本节从以下的三个方面, 分别建立了相应的理论性质: 首先, 在模型的预测误差方面, 建立了 $X\hat{\beta}_{map}$ 与 $X\hat{\beta}_{mean}$ 的误差边界, 即 Oracle 不等式; 其次, 在系数向量的支撑集恢复方面, 证明了在满足可识别性条件下, 只要系数向量非零元素足够大, 所求得的支撑集的最大后验估计 $\hat{J}_{map}$ 是相合的; 最后, 在估计误差方面, 推导了 $\hat{\beta}_{map}$ 在 l_1 , l_2 , l_∞ 范数误差下的估计误差边界, 同时在更为严格的条件下, 为 $\hat{\beta}_{mean}$ 在 l_∞ 范数误差下建立了相应的误差边界, 表明本文方法所得参数估计是相合的。我们注意到, 当 p 小于 n 时, 理论结果表现良好, 并且适用于小 n 大 p 问题, 更详细的结果在随后的定理中给出。

3.3.1 Oracle 不等式

首先, 基于本文构造的 l_0 收缩指数稀疏先验分布所选择的模型, 在其预测误差方面, 建立了 $X\hat{\beta}_{map}$ 与 $X\hat{\beta}_{mean}$ 的误差边界, 即预测误差的 Oracle 不等式。

定理 3.1 对任意常数 $c > 0$, 实数 $b \geq 0$, 对于给定的样本设计矩阵 $X \in R^{n \times p}$ 以及 $p \geq n$ 且满足 (2-3) 式, 若 $m = (31 + 6c) \log p$, 则对于 $\hat{\beta}_{map}$ 和 $\hat{\beta}_{mean}$ 至少以概率 $1 - p^{-c}$ 有,

$$\|X\hat{\beta}_{map} - X\beta_*\|_2^2 \leq 4\sigma^2(b+3)ms_*, \quad (3-9)$$

$$\|X\hat{\beta}_{mean} - X\beta_*\|_2^2 \leq \sigma^2(\sqrt{6(b+3)ms_*} + 1)^2. \quad (3-10)$$

该定理的详细推导证明见 (3-4) 节。从定理 3.1 可以看出, 本文的推广指数权重方法所得的贝叶斯估计的预测误差上界由 σ^2 , 系数向量的稀疏度 s_* 以及超参数 b 、 m 共同决定。

3.3.2 支撑恢复

在一般情况下, 若无其他额外的限制条件, 标准线性模型中的系数向量并不唯一, 即模型是不可识别的, 此时进行模型选择具有不确定性。而支撑恢复问题需要在模型可识别下进行, 要想达到模型选择的目的, 建立待估系数向量参数估计的性质, 得到突出的理论结果, 还需要附加额外的条件以保证 β_* 的唯一性。在研究支撑集估计时, 我们引入关于设计矩阵 X 的可识别性条件 $I(s)$ 。

$I(s)$: 对任意支撑集 $J \in A(J) := \{J \in K(\Omega_p) : |J| \leq s\}$, 设计矩阵 X 的子矩阵 X_J 是列满秩的^[66]。

为了后续的参数估计和待估系数的支撑恢复研究, 需要设计矩阵 X 满足条件 $I((2+\varepsilon)s_*)$, 其中 $\varepsilon > 0$ 。注意到条件 $I(s)$ 等价于 X 的任意子矩阵 X_J 的奇异值均为正, 为了后续研究方便, 定义:

$$v_s = \min_{J \subset \Omega_p: |J| \leq s} \min_{e \in R^{|J|}: \|e\|_2=1} \frac{1}{\sqrt{n}} \|X_J e\|_2. \quad (3-11)$$

对于矩阵 $\frac{1}{\sqrt{n}}X$, 当 $|J| \leq s$ 时, v_s 表示全体子矩阵 $\frac{1}{\sqrt{n}}X_J$ 中最小的奇异值, 基于该等价定义, 我们建立了关于支撑集最大后验估计的理论性质。

定理 3.2 任意常数 $c \geq 0$, 对于实数 b 和 ε , 若 $0 \leq b < \varepsilon \leq s_* - 1$, $m = \frac{1+\varepsilon}{\varepsilon-b}(21 + \frac{7}{2}c) \log p$, 设计矩阵 $X \in R^{n \times p}$ 以及 $p \geq n$ 且满足 (2-3) 式, 条件 $I((2+\varepsilon)s_*)$ 成立, 如果有下式成立

$$\min_{j \in J_*} |\beta_{*,j}| \geq \rho := \sqrt{\frac{4\sigma^2(2\varepsilon - b + 1)m}{n(1+\varepsilon)v_{(2+\varepsilon)s_*}^2}}, \quad (3-12)$$

则对任意的 J , $\Pi(J_*) \geq \Pi(J)$ 成立的概率不小于 $1 - 2p^{-c}$, 即至少以概率 $1 - 2p^{-c}$ 有 $J_* = \hat{J}_{map}$ 。

定理的详细证明过程在 3.4 节中。由定理 3.2 可知, 当设计矩阵 X 满足条件 $I((2+\varepsilon)s_*)$ 且 β_* 的非零元素取较大值时, 若有 $\sqrt{\log p/n}$ 随 n 增大而趋于 0, 此时表明使得后验概率达到最大的支撑集 $\hat{J}_{map}$ 是 J_* 的相合估计。

3.3.3 参数估计

在本节中, 主要研究关于推广指数权重方法所求得参数估计的误差上界。如预备知识所述, 我们考虑最大后验估计与后验均值估计的误差边界, 建立了本文方法下所求得的参数估计误差界的相关定理。我们注意到结果适用于小 n 大 p 问题。

定理 3.3 任意常数 $c > 0$, 对于实数 b 和 ε , 若 $0 \leq b < \varepsilon \leq 2b+1$, 设计矩阵 $X \in R^{n \times p}$ 以及 $p \geq n$ 且满足 (2-3) 式, $m = \frac{\varepsilon+1}{\varepsilon-b}(21 + \frac{7}{2}c) \log p$, 以不小于 $1 - 3p^{-c}$ 的概率, 有以下不等式成立,

$$\|\hat{\beta}_{map} - \beta_*\|_2 \leq \sigma \sqrt{\frac{4(b+3)ms_*}{nv_{(2+\varepsilon)s_*}^2}}. \quad (3-13)$$

若 (3-12) 式满足, 且 $\varepsilon \leq s_* - 1$, 则进一步有

$$\|\hat{\beta}_{map} - \beta_*\|_2 \leq \sigma \sqrt{\frac{4(b+3)ms_*}{nv_{s_*}^2}}. \quad (3-14)$$

定理 3.3 呈现了 β_* 与 $\hat{\beta}_{map}$ 在 l_2 范数下估计误差上界, 表明基于本文方法所学习得到的参数估计是相合的。

定理 3.4 任意常数 $c > 0$, 对于任意实数 b 和 ε , 若 $0 \leq b < \varepsilon \leq s_* - 1$, 设计矩阵 $X \in R^{n \times p}$ 以及 $p \geq n$ 且满足 (2-3) 式, 若有条件 $I((2+\varepsilon)s_*)$ 满足, $m = \frac{\varepsilon+1}{\varepsilon-b}(21 + \frac{7}{2}c) \log p$, 若 (3-12) 式成立, 以不小于 $1 - 3p^{-c}$ 的概率, 可以得到

$$\begin{aligned} \|\hat{\beta}_{map} - \beta_*\|_\infty &\leq \sigma \sqrt{\frac{2(c+1) \log p}{nv_{s_*}^2}}, \\ \|\hat{\beta}_{map} - \beta_*\|_1 &\leq \sigma s_* \sqrt{\frac{2(c+1) \log p}{nv_{s_*}^2}}. \end{aligned} \quad (3-15)$$

从定理 3.4 与定理 3.3 可知, 其条件是相同的, 因此在保证支撑估计 $\hat{J}_{map}$ 能够依概率收敛到真实支撑 J_* 时, 进而再建立 $\hat{\beta}_{map}$ 的误差界。定理 3.4 所表示的是系数向量的各分量的估计误差的上界。从该定理的理论结果可知, $\hat{\beta}_{map}$ 的 l_∞ 和 l_1 范数误差收敛阶数为 $\sqrt{\log p/n}$, 当 $\log p/n$ 随着 n 增大而趋于 0 时, 表明本文方法所得的参数估计是相合的。

定理 3.5 任意常数 $c > 0$, 对于实数 b 和 ε , 若 $0 \leq b < \varepsilon \leq s_* - 1$, 且设计矩阵 $X \in R^{n \times p}$ 以及 $p \geq n$ 满足 (2-3) 式与条件 $I(l + s_*)$, 若 $m = \frac{\varepsilon+1}{\varepsilon-b}(21 + \frac{7}{2}c) \log p$ 且 (3-12) 式子成立, 则至少以概率 $1 - 4p^{-c}$ 有

$$\|\hat{\beta}_{mean} - \beta_*\|_\infty \leq \sigma \sqrt{\frac{2(c+1) \log p}{nv_{s_*}^2}} + \frac{\sigma \sqrt{6(3+b)ms_*}}{\sqrt{nv_{l+s_*}}} + \frac{\sigma p^{-s_*}}{\sqrt{nv_{l+s_*}}}. \quad (3-16)$$

可以看到, 若条件 $I(l + s_*)$ 满足, 我们同样可以得到关于后验均值估计的一个较好的误差界。对于高斯设计矩阵 X , 条件 $I(l + s_*)$ 成立的概率趋近于 1^[43]。

3.4 理论证明

在对上一节的定理进行证明之前, 我们先给出几个理论证明里经常用到的引理, 并对引理进行了相应的详细证明。

3.4.1 引理及其证明

引理 3.1 ^[66] 对于任意常数 $c > 0$, 任意的支撑集 $J \in J_{s,k}$, 且满足条件 $k \leq s \wedge s_*$, 则至少以概率 $1 - p^{-c}$ 可以得到

$$w^T(\psi_J - \psi_{J_*})w \leq (15 + 3c)\sigma^2(s \vee s_* - k) \log p, \quad (3-17)$$

其中, $J_{s,k} = \{J \in K(\Omega_p) : |J| = s, |J \cap J_*| = k, J \neq J_*\}$, $w \sim N(0, \sigma^2 I_n)$ 。

证明

由 ψ_J 的定义, 可以得到 $\psi_{J \cap J_*} \psi_J = \psi_J \psi_{J \cap J_*} = \psi_{J \cap J_*}$, 且 ψ_J 为对称幂等阵, 即: $\psi_J^2 = \psi_J$, $\psi_J^T = \psi_J$, 因此有下式成立:

$$\begin{aligned} (\psi_J - \psi_{J \cap J_*})(\psi_J - \psi_{J \cap J_*}) &= \psi_J^2 - \psi_{J \cap J_*} \psi_J - \psi_J \psi_{J \cap J_*} + \psi_{J \cap J_*}^2 \\ &= \psi_J - \psi_{J \cap J_*} - \psi_{J \cap J_*} + \psi_{J \cap J_*} \\ &= \psi_J - \psi_{J \cap J_*}, \end{aligned}$$

即 $\psi_J - \psi_{J \cap J_*}$ 为幂等阵, 又有 $(\psi_J - \psi_{J \cap J_*})^T = \psi_J^T - \psi_{J \cap J_*}^T = \psi_J - \psi_{J \cap J_*}$, 因此 $\psi_J - \psi_{J \cap J_*}$ 为对称幂等阵, 因此 $\psi_J - \psi_{J \cap J_*}$ 为正交投影阵. 因此我们有下列等式成立

$$\begin{aligned} r(\psi_J - \psi_{J \cap J_*}) &= tr(\psi_J - \psi_{J \cap J_*}) \\ &= tr(\psi_J) - tr(\psi_{J \cap J_*}) \\ &= r(\psi_J) - r(\psi_{J \cap J_*}) \\ &= s - k. \end{aligned}$$

令 $M_J = \psi_J - \psi_{J \cap J_*}$, 此时满足 $\frac{\|M_J w\|_2^2}{\sigma^2} \sim \chi_{s-k}^2$. 对于自由度为 d 的卡方分布, 由切尔诺夫不等式^[64]有

$$P(\chi_d^2 > a) \leq \frac{E(e^{tx})}{e^{ta}},$$

其中 $t > 0$. 自由度为 d 的卡方分布的密度函数为

$$f(x) = \frac{1}{2^{\frac{d}{2}} \Gamma(\frac{d}{2})} e^{-\frac{x}{2}} x^{\frac{d}{2}-1} I_{\{x \geq 0\}},$$

其中

$$\begin{aligned} E(e^{tx}) &= \int_0^\infty \frac{1}{2^{\frac{d}{2}} \Gamma(\frac{d}{2})} e^{tx - \frac{x}{2}} x^{\frac{d}{2}-1} dx \\ &= \int_0^\infty \frac{1}{2^{\frac{d}{2}} \Gamma(\frac{d}{2})} e^{-\frac{y}{2}} y^{\frac{d}{2}-1} \left(\frac{1}{1-2t}\right)^{\frac{d}{2}} dy \\ &= (1-2t)^{-\frac{d}{2}}, \end{aligned}$$

令 $(t - \frac{1}{2})x = -\frac{1}{2}y$ 即得到上式中第二个等式, 其中 $0 \leq t \leq \frac{1}{2}$, 故有下式成立

$$P(\chi_d^2 > a) \leq (1-2t)^{-\frac{d}{2}} e^{-ta}.$$

令 $t = \frac{a-d}{2a}$, $d < a$, 由上式有

$$\log(P(\chi_d^2 > a)) \leq -\frac{d}{2} \left(\frac{a}{d} - 1 - \log\left(\frac{a}{d}\right) \right).$$

记关于 d 的函数 $h(d) = -\frac{d}{2} \left(\frac{a}{d} - 1 - \log\left(\frac{a}{d}\right) \right)$, 对 $h(d)$ 关于 d 求导并令其大于 0 得到函数 $h(d)$ 关于 d 的单调递增区间为 $(0, a)$. 当 $d < \frac{a}{10}$ 时, $h(d) < h(\frac{a}{10}) < -\frac{a}{3}$, 因此有下式成立,

$$\log(P(\chi_d^2 > a)) \leq -\frac{d}{2} \left(\frac{a}{d} - 1 - \log\left(\frac{a}{d}\right) \right) \leq -\frac{a}{3}. \quad (3-18)$$

根据不等式 (3-18) 可以得到 $P(\chi_d^2 > a) \leq \exp\{-\frac{a}{3}\}$, 因此,

$$\begin{aligned} P(\max_{J \in J_{s,k}} \frac{\|M_J w\|_2^2}{\sigma^2} > a) &\leq P(\bigcup_{J \in J_{s,k}} \frac{\|M_J w\|_2^2}{\sigma^2} > a) \\ &\leq \binom{s}{k} \binom{p-s}{s_*-k} P(\frac{\|M_J w\|_2^2}{\sigma^2} > a) \\ &\leq \binom{s}{k} \binom{p-s}{s_*-k} \exp(-\frac{a}{3}). \end{aligned} \quad (3-19)$$

由二项式系数的标准的上界: $\binom{n}{k} \leq (\frac{ne}{k})^k$, 且 $\binom{n}{k} = \binom{n}{n-k}$, 有

$$\binom{s}{k} \leq (\frac{se}{s-k})^{s-k} \leq (es)^{s-k}, \quad \binom{p-s}{s_*-k} \leq (\frac{e(p-s)}{s_*-k})^{s_*-k} \leq (ep)^{s_*-k},$$

由上述不等式进一步可推知

$$\begin{aligned} \log(\binom{s}{k} \binom{p-s}{s_*-k}) &\leq (s-k) \log(es) + (s_*-k) \log(ep) \\ &\leq (s \vee s_* - k) \log(spe^2) \\ &\leq 3(s \vee s_* - k) \log p. \end{aligned} \quad (3-20)$$

其中第二个不等式只需 $se^2 \leq p^2$ 成立, 由于 $s \leq p$, 故只需 $p > e^2$ 即可, 在高维线性模型中这通常都是成立的, 因此有 $\binom{s}{k} \binom{p-s}{s_*-k} \leq \exp(3(s \vee s_* - k) \log p)$ 。

令 $a = (15 + 3c)(s \vee s_* - k) \log p$, 有

$$\begin{aligned} P(\max_{J \in J_{s,k}} \|M_J w\|_2^2 > a\sigma^2) &\leq \exp(3(s \vee s_* - k) \log p - \frac{a}{3}) \\ &\leq \exp(-(2+c)(s \vee s_* - k) \log p). \end{aligned} \quad (3-21)$$

下证: $s \vee s_* - k \geq 1$ 。不妨假设 $s \vee s_* - k = 0$ 。由 k 的定义, 有 $k = s \vee s_* \leq s \wedge s_*$, 因此我们得到 $s = s_*$, 故 $J = J_*$, 而这与 $J_{s,k}$ 的定义相矛盾, 假设不成立, 故 $s \vee s_* - k \geq 1$, 因此有 $P(\max_{J \in J_{s,k}} \|M_J w\|_2^2 > a\sigma^2) \leq p^{-(2+c)}$ 。

由 l 的定义可知, 其表示模型系数向量的稀疏度上界, 因此有 $l \leq p$ 成立, 故我们可以得到

$$P(\max_{s,k} \max_{J \in J_{s,k}} \|M_J w\|_2^2 > a\sigma^2) \leq l(s \wedge s_* + 1)p^{-(2+c)} \leq p^{-c}, \quad (3-22)$$

且有

$$\begin{aligned}
 w^T(\psi_J - \psi_{J_*})w &= w^T(\psi_J - \psi_{J \cap J_*})w - w^T(\psi_{J_*} - \psi_{J \cap J_*})w \\
 &\leq w^T(\psi_J - \psi_{J \cap J_*})w \\
 &= \|M_J w\|_2^2.
 \end{aligned} \tag{3-23}$$

由 (3-22) 与 (3-23) 式可以得到

$$P(w^T(\psi_J - \psi_{J_*})w \leq (15 + 3c)\sigma^2(s \vee s_* - k) \log p) \geq 1 - p^{-c}.$$

综上, 引理 3.1 即证。

引理 3.2 ^[66] 对任意常数 $c > 0$, 任意支撑集 $J \in K(\Omega_p)$, 至少以概率 $1 - p^{-c}$ 可以得到

$$\|\psi_J w\|_2^2 \leq (15 + 3c)\sigma^2|J| \log p, \tag{3-24}$$

其中, $w \sim N(0, \sigma^2 I_n)$ 。

证明

在引理 3.1 中, 令 $s_* = 0$, 此时 $s \vee s_* = s$, 且 $k = 0$, 又因为 ψ_J 是正交投影阵, 有 $\psi_J = \psi_J^T = \psi_J T \psi_J$, 故至少以概率 $1 - p^{-c}$ 有

$$\|\psi_J w\|_2^2 = w^T(\psi_J - 0)w \leq (15 + 3c)\sigma^2 s \log p, \tag{3-25}$$

综上, 即证引理 3.2。

引理 3.3 ^[66] 对任意常数 $c > 0$, 任意支撑集 $J \in J_{s,k}$ 且满足条件 $k \leq s \wedge s_*$, 则至少以概率 $1 - p^{-c}$ 可以得到下式

$$\frac{\langle \psi_J^\perp X \beta_*, w \rangle^2}{\|\psi_J^\perp X \beta_*\|_2^2} \leq (10 + 2c)\sigma^2(s \vee s_* - k) \log p, \tag{3-26}$$

其中, $J_{s,k} = \{J \subset \Omega_p : |J| = s, |J \cap J_*| = k, J \neq J_*\}$, $w \sim N(0, \sigma^2 I_n)$ 。

证明

令 $\alpha_J = \|\psi_J^\perp X \beta_*\|_2$, $u_J = \langle \psi_J^\perp X \beta_*, w \rangle = (\psi_J^\perp X \beta_*)^T w$, 则 $u_J \sim N(0, \sigma^2 \alpha_J^2)$ 。再令

$\eta_J = \frac{u_J}{\sigma \alpha_J}$, 则 $\eta_J \sim N(0, 1)$ 。因此

$$\begin{aligned}
 P(\eta_J^2 > a) &= P(\eta_J > \sqrt{a}) + P(\eta_J < -\sqrt{a}) \\
 &= 2P(\eta_J > \sqrt{a}) \\
 &= 2 \int_{\sqrt{a}}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}} dx \\
 &< \frac{2}{\sqrt{2\pi}} \int_{\sqrt{a}}^{\infty} x e^{-\frac{x^2}{2}} dx \\
 &= \frac{2}{\sqrt{2\pi}} e^{-\frac{a}{2}} \\
 &< e^{-\frac{a}{2}},
 \end{aligned}$$

由引理 3.1 证明中的 (3-19) 与 (3-20) 式, 有

$$\begin{aligned}
 P(\max_{J \in J_{s,k}} \eta_J^2 > a) &\leq P(\bigcup_{J \in J_{s,k}} \eta_J^2 > a) \\
 &\leq \binom{s}{k} \binom{p-s}{s_*-k} P(\eta_J^2 > a) \\
 &\leq p^{3(s_* \vee s - k)} \exp\left\{-\frac{a}{2}\right\}.
 \end{aligned}$$

令 $a = (10 + 2c)(s_* \vee s - k) \log p$, 且因为 $s_* \vee s - k \geq 1$, 可得

$$\begin{aligned}
 P(\max_{J \in J_{s,k}} \eta_J^2 > a) &\leq p^{-(2+c)(s_* \vee s - k)} \\
 &\leq p^{-(2+c)}.
 \end{aligned} \tag{3-27}$$

因此

$$\begin{aligned}
 P(\max_{J \in J_{s,k}} \eta_J^2 > a) &\leq p^{-(2+c)(s_* \vee s - k)} \\
 &\leq l(s \wedge s_* + 1) p^{-(2+c)} \\
 &\leq p^{-c},
 \end{aligned} \tag{3-28}$$

故有

$$P\left(\frac{\langle \psi_J^\perp X \beta_*, w \rangle^2}{\|\psi_J^\perp X \beta_*\|_2^2} \leq (10 + 2c) \sigma^2 (s \vee s_* - k) \log p\right) \geq 1 - p^{-c}.$$

综上, 引理 3.3 得证。

3.4.2 定理证明

定理 3.1 的证明

证明

对任意支撑集 $J \in K(\Omega_p)$, 记 $\tau_J = \psi_J y - X\beta_*$. $\psi_J^\perp$ 表示子矩阵 X_{J_*} 的各列向量所张成线性空间正交补空间上的正交投影阵, 因此有 $\|\psi_J^\perp X\beta_*\|_2 = 0$, 可得 $\|\psi_J^\perp y\|_2^2 = \|\psi_J^\perp (X\beta_* + w)\|_2^2 = \|\psi_J^\perp w\|_2^2$. 且 $\psi_J^\perp w = (I - \psi_J)w = w - \psi_J(y - X\beta_*) = w - \tau_J$, $\psi_J^\perp y = (I - \psi_J)y = y - (\tau_J + X\beta_*) = w - \tau_J$, 因此,

$$\begin{aligned} \|\psi_J^\perp y\|_2^2 - \|\psi_J^\perp y\|_2^2 &= \|w - \tau_J\|_2^2 - \|w - \tau_J\|_2^2 \\ &= 2w^T(\tau_J - \tau_{J_*}) + \|\tau_{J_*}\|_2^2 - \|\tau_J\|_2^2, \end{aligned} \quad (3-29)$$

根据 τ_J 的定义, 其在子矩阵 $X_{J \cup J_*}$ 各列张成的线性空间内, 即 $\tau_J \in \text{span}(X_{J \cup J_*})$, 同理有 $\tau_{J_*} \in \text{span}(X_{J_*})$, 因此 $\tau_J - \tau_{J_*} \in \text{span}(X_{J \cup J_*})$, 进一步可得

$$\begin{aligned} \|w^T(\tau_J - \tau_{J_*})\|_2 &= \|w^T(\psi_{J \cup J_*} + \psi_{(J \cup J_*)^c})(\tau_J - \tau_{J_*})\|_2 \\ &= \|w^T \psi_{J \cup J_*}(\tau_J - \tau_{J_*})\|_2 \\ &= \|(\psi_{J \cup J_*} w)^T(\tau_J - \tau_{J_*})\|_2 \\ &\leq \|\psi_{J \cup J_*} w\|_2 \|\tau_J - \tau_{J_*}\|_2, \end{aligned} \quad (3-30)$$

由 (3-29) 和 (3-30) 式有

$$\|\psi_J^\perp y\|_2^2 - \|\psi_J^\perp y\|_2^2 \leq 2\|\psi_{J \cup J_*} w\|_2 \|\tau_J - \tau_{J_*}\|_2 + \|\tau_{J_*}\|_2^2 - \|\tau_J\|_2^2. \quad (3-31)$$

由 τ_J 的定义, 有 $\|\tau_{J_*}\|_2^2 = \|\psi_{J_*} y - X\beta_*\|_2^2 = \|\psi_{J_*}(X\beta_* + w) - X\beta_*\|_2^2 = \|\psi_{J_*} w\|_2^2$, 结合引理 3.2 条件与理论性质, 如下不等式成立的概率不小于 $1 - p^{-c}$,

$$\begin{aligned} \|\psi_J^\perp y\|_2^2 - \|\psi_J^\perp y\|_2^2 &\leq 2\|\psi_{J \cup J_*} w\|_2 \|\tau_J - \tau_{J_*}\|_2 + \|\tau_{J_*}\|_2^2 - \|\tau_J\|_2^2 \\ &\leq 2\|\psi_{J \cup J_*} w\|_2 (\|\tau_{J_*}\|_2 + \|\tau_J\|_2) + \|\tau_{J_*}\|_2^2 - \|\tau_J\|_2^2 \\ &\leq (2\|\psi_{J \cup J_*} w\|_2)^2 + \left(\frac{\|\tau_{J_*}\|_2 + \|\tau_J\|_2}{2}\right)^2 + \|\tau_{J_*}\|_2^2 - \|\tau_J\|_2^2 \\ &= 4\|\psi_{J \cup J_*} w\|_2^2 + \frac{5}{4}\|\tau_{J_*}\|_2^2 - \frac{3}{4}\|\tau_J\|_2^2 + \frac{\|\tau_{J_*}\|_2 \|\tau_J\|_2}{2} \\ &= 4\|\psi_{J \cup J_*} w\|_2^2 + \frac{3}{2}\|\tau_{J_*}\|_2^2 - \frac{1}{2}\|\tau_J\|_2^2 - \frac{1}{4}(\|\tau_{J_*}\|_2 - \|\tau_J\|_2)^2 \\ &\leq 4\|\psi_{J \cup J_*} w\|_2^2 + \frac{3}{2}\|\tau_{J_*}\|_2^2 - \frac{1}{2}\|\tau_J\|_2^2 \end{aligned}$$

$$\begin{aligned} &\leq 4(15+3c)\sigma^2(s+s_*-k)\log p + \frac{3}{2}(15+3c)\sigma^2 s_* \log p - \frac{1}{2}\|\tau_J\|_2^2 \\ &\leq \frac{11}{2}\sigma^2(15+3c)s_* \log p + (60+12c)\sigma^2 s \log p - \frac{1}{2}\|\tau_J\|_2^2. \end{aligned} \quad (3-32)$$

其中第三个不等式运用 $(2a_1)^2 + (\frac{a_2}{2})^2 \geq 2a_1a_2$ 得到, 只需取 $a_1 = \|\psi_{J \cup J_*} w\|_2$, $a_2 = \|\tau_{J_*}\|_2 + \|\tau_J\|_2$ 。当 $m \geq (31+6c)\log p$ 时, 根据引理 3.2, 对任意 $J \in \Theta := \{J \in K(\Omega_p) : \Pi(J) \geq \Pi(J_*)\}$, 如下不等式成立的概率不小于 $1 - p^{-c}$,

$$\begin{aligned} \frac{\Pi(J)}{\Pi(J_*)} &= \frac{\binom{p}{s_*}}{\binom{p}{s}} \exp\{m(s_* - s) + mb \log \frac{s}{s_*} + \frac{1}{2\sigma^2}(\|\psi_{J_*}^\perp y\|_2^2 - \|\psi_J^\perp y\|_2^2)\} \\ &\leq \frac{\binom{p}{s_*}}{\binom{p}{s}} \exp\{m(s_* - s) + mb \log \frac{s}{s_*} + \frac{11}{4}(15+3c)s_* \log p + (30+6c)s \log p - \frac{1}{4\sigma^2}\|\tau_J\|_2^2\} \\ &\leq \exp\{(m+bm + \frac{11}{4}(15+3c)\log p + \log p)s_* + s_* \log p + ((30+6c)\log p - m)s - \frac{1}{4\sigma^2}\|\tau_J\|_2^2\} \\ &\leq \exp\{(b+3)ms_* - \frac{1}{4\sigma^2}\|\tau_J\|_2^2\}. \end{aligned} \quad (3-33)$$

其中第二个不等式有 $C_p^{s_*} < p^{s_*}$, 且 $s \leq s_* e^{s_*}$, 这个条件通常都是成立的, 且也可以通过在稀疏先验中将稀疏度的上界 l 设置为 $s_* e^{s_*}$ 。因为 $J \in \Theta$, 则有

$$(b+3)ms_* - \frac{1}{4\sigma^2}\|\tau_J\|_2^2 \geq 0. \quad (3-34)$$

即,

$$\|\psi_J y - X\beta_*\|_2^2 = \|\tau_J\|_2^2 \leq 4\sigma^2(b+3)ms_*. \quad (3-35)$$

显然, $\hat{J}_{map} \in \Theta$, 因此定理的第一部分得证。

记 $\Theta_0 = \{J \in K(\Omega_p) : \|\tau_J\|_2 > \sigma\sqrt{6(3+b)ms_*}\}$, 则有

$$\begin{aligned} \|X\hat{\beta}_{mean} - X\beta_*\|_2 &= \|X(\sum_J \Pi(J)\hat{\beta} - \sum_J \Pi(J)\hat{\beta}_*)\|_2 \\ &= \|\sum_J \Pi(J)(X\hat{\beta} - X\beta_*)\|_2 \\ &\leq \sum_J \|\tau_J\|_2 \Pi(J) \\ &\leq \sigma\sqrt{6(b+3)ms_*} \sum_{J \notin \Theta_0} \Pi(J) + \sum_{J \in \Theta_0} \|\tau_J\|_2 \frac{\Pi(J)}{\Pi(J_*)}. \end{aligned} \quad (3-36)$$

事实上, 从 (3-32) 的倒数第二个不等式, 我们可以计算得到 $(m+bm + \frac{11}{4}(15+3c)\log p + \log p)s_* \leq (b+\frac{5}{2})ms_*$ 。由 (3-32) 和 (3-33) 的推导过程以及引理 3.2, 当

$J \in \Theta_0$ 时, 至少以概率 $1 - p^{-c}$, 下式成立

$$\begin{aligned}
 \|\tau_J\|_2 \frac{\Pi(J)}{\Pi(J_*)} &= \|\tau_J\|_2 \frac{\binom{p}{s_*}}{\binom{p}{s}} \exp\{m(s_* - s) + mb \log \frac{s}{s_*} + \frac{1}{2\sigma^2} (\|\psi_{J_*}^\perp y\|_2^2 - \|\psi_J^\perp y\|_2^2)\} \\
 &\leq \frac{\|\tau_J\|_2}{\binom{p}{s}} \exp\{(b + \frac{5}{2})ms_* + ((30 + 6c) \log p - m)s - \frac{1}{4\sigma^2} \|\tau_J\|_2^2\} \\
 &= \frac{1}{\binom{p}{s}} \exp\{(b + \frac{5}{2})ms_* + ((30 + 6c) \log p - m)s - \frac{1}{6\sigma^2} \|\tau_J\|_2^2\} \cdot \|\tau_J\|_2 \exp\{-\frac{1}{12\sigma^2} \|\tau_J\|_2^2\} \\
 &\leq \frac{1}{\binom{p}{s}} \sigma \exp\{(b + \frac{5}{2})ms_* + ((30 + 6c) \log p - m)s - (b + 3)ms_*\} \\
 &= \frac{1}{\binom{p}{s}} \sigma \exp\{((30 + 6c) \log p - m)s - \frac{1}{2}ms_*\}.
 \end{aligned} \tag{3-37}$$

记关于 x 的函数 $h(x) = x \exp(-x^2)$, 对其求导并令导数大于 0, 可得到 $h(x)$ 关于 x 的单调增区间为 $(0, \frac{\sqrt{2}}{2})$, 因此 $h(x) \leq h(\frac{\sqrt{2}}{2}) = \frac{\sqrt{2}}{2} \exp\{-\frac{1}{2}\} \leq \frac{\sqrt{12}}{12}$, 因此 $\|\tau_J\|_2^2 \exp\{-\frac{1}{12} \|\tau_J\|_2^2\} \leq 1$, 即得到上述第二个不等式。进一步, 下式至少以概率 $1 - p^{-c}$ 成立

$$\begin{aligned}
 \sum_{J \in \Theta_0} \|\tau_J\|_2 \frac{\Pi(J)}{\Pi(J_*)} &\leq \sum_{s=0}^l \sum_{|J|=s} \frac{1}{\binom{p}{s}} \sigma \exp\{((30 + 6c) \log p - m)s - \frac{1}{2}ms_*\} \\
 &\leq \sum_{s=0}^l \sigma \exp\{((30 + 6c) \log p - m)s - \frac{1}{2}ms_*\} \\
 &\leq \sum_{s=0}^l \sigma p^{-s} \cdot \exp\{-\frac{1}{2}ms_*\} \\
 &\leq \frac{\sigma}{1 - p^{-1}} \cdot \frac{1}{p^{(15+3c)s_*}} \\
 &\leq \sigma p^{-s_*} \leq \sigma.
 \end{aligned} \tag{3-38}$$

由 (3-36) 和 (3-38) 式有

$$\|X\hat{\beta}_{mean} - X\beta_*\|_2 \leq \sigma \sqrt{6(3+b)ms_*} + \sigma = \sigma(\sqrt{6(3+b)ms_*} + 1), \tag{3-39}$$

即 $\|X\hat{\beta}_{mean} - X\beta_*\|_2^2 \leq \sigma^2(\sqrt{6(b+3)ms_*} + 2)^2$.

综上, 定理 3.1 得证。

定理 3.2 的证明

证明

将 J 所属范围分为三个部分, 分别证明定理的结论。记 $A_1 = \{J \in K(\Omega_p) : |J| \geq$

$(1+\varepsilon)s_*$, $A_2 = \{J \in K(\Omega_p) : s_* < |J| < (1+\varepsilon)s_*\}$, $A_3 = \{J \in K(\Omega_p) : |J| \leq s_*\}$ 。根据定理 3.1 的证明推导有

$$\begin{aligned}
 \|\psi_{J_*}^\perp w\|_2^2 - \|\psi_J^\perp y\|_2^2 &= \|\psi_{J_*}^\perp y\|_2^2 - \|\psi_J^\perp y\|_2^2 \\
 &= y^T(I - \psi_{J_*})y - y^T(I - \psi_J)y \\
 &= y^T(\psi_J - \psi_{J_*})y \\
 &= (X\beta_* + w)^T(\psi_J - \psi_{J_*})(X\beta_* + w) \\
 &= -\|\psi_J^\perp X\beta_*\|_2^2 - 2\langle \psi_J^\perp X\beta_*, w \rangle + w^T(\psi_J - \psi_{J_*})w.
 \end{aligned} \tag{3-40}$$

回顾引理 3.1 与引理 3.3 的理论结果, 我们可以得到, 以下不等式成立的概率不小于 $1 - 2p^{-c}$, 即

$$y^T(\psi_J - \psi_{J_*})y \leq -\alpha_J^2 + 2\alpha_J\sigma\sqrt{(10+2c)(s \vee s_* - k)\log p} + (15+3c)\sigma^2(s \vee s_* - k)\log p. \tag{3-41}$$

根据不等式 $2a_1a_2 \leq a_1^2 + a_2^2$, 取 $a_1 = \frac{\alpha_J}{\sqrt{2}}$, $a_2 = \sigma\sqrt{(20+4c)(s \vee s_* - k)\log p}$, 由 (3-40) 进一步可推知

$$\begin{aligned}
 y^T(\psi_J - \psi_{J_*})y &\leq -\alpha_J^2 + \frac{1}{2}\alpha_J^2 + (35+7c)\sigma^2(s \vee s_* - k)\log p \\
 &= (35+7c)\sigma^2(s \vee s_* - k)\log p - \frac{1}{2}\alpha_J^2 \\
 &\leq (35+7c)\sigma^2(s \vee s_* - k)\log p,
 \end{aligned} \tag{3-42}$$

因此, 当 $J \in A_1$ 时, 至少以概率 $1 - 2p^{-c}$ 有

$$\frac{\Pi(J)}{\Pi(J_*)} \leq \sum_{J \in J_{s,k}} \frac{\Pi(J)}{\Pi(J_*)} \leq \sum_{J \in J_{s,k}} \frac{\binom{p}{s_*}}{\binom{p}{s}} \exp\{m(s_* - s) + mb \log \frac{s}{s_*} + \frac{35+7c}{2}(s - k)\log p\}. \tag{3-43}$$

而 $|J_{s,k}| = \binom{s_*}{k} \binom{p-s_*}{s-k}$, 且下式成立

$$\frac{\binom{s_*}{k} \binom{p-s_*}{s-k} \binom{p}{s_*}}{\binom{p}{s_*}} = \binom{s}{k} \binom{p-s}{s_*-k}. \tag{3-44}$$

由 (3-20) 式可知

$$\binom{s}{k} \binom{p-s}{s_*-k} \leq \exp\{3(s_* \vee s - k)\log p\},$$

因此

$$\begin{aligned}
 \frac{\Pi(J)}{\Pi(J_*)} &\leq \sum_{J \in J_{s,k}} \frac{\binom{p}{s_*}}{\binom{p}{s}} \exp\{m(s_* - s) + mb \log \frac{s}{s_*} + \frac{35+7c}{2}(s-k) \log p\} \\
 &\leq \exp\{m(s_* - s) + mb \log \frac{s}{s_*} + (21 + \frac{7c}{2})(s-k) \log p\} \\
 &\leq \exp\{-m\epsilon s_* + mbs_* + (21 + \frac{7c}{2})(1+\epsilon)s_* \log p\} \\
 &\leq \exp\{-s_*((21 + \frac{7c}{2})(1+\epsilon) \log p + (\epsilon - b)m)\} \\
 &= 1.
 \end{aligned} \tag{3-45}$$

将 $m = \frac{1+\epsilon}{\epsilon-b}(21 + \frac{7c}{2}) \log p$ 带入第四个不等式, 即可得到最后一个等式。综上所述, 当 $J \in A_1$ 时, 有 $\Pi(J) \leq \Pi(J_*)$, 使得后验概率达到最大的支撑估计就是真实支撑。事实上, 此时并未对待估系数的值有其他额外的条件即可得到最大后验能够保证恢复真实的支撑。

当 $J \in A_2$ 时, 记 $X = [X_1, X_2]$, 而矩阵 X 的最小奇异值用 δ 表示, 记 $\beta = (\beta_1, \beta_2)$, 同时用 ψ_2 表示 X_2 的列向量张成线性空间的正交投影阵。于是可以得到如下关于任意向量 β_1 的关系式,

$$\|(I - \psi_2)X_1\beta_1\|_2^2 = \min_{\beta_2} \|X_1\beta_1 + X_2\beta_2\|_2^2 = \min_{\beta_2} \beta^T X^T X \beta.$$

对于实对称阵 $X^T X$, 存在正交矩阵 Q , 使得其可以对角化, 即有下列式子成立^[66],

$$\begin{aligned}
 \beta^T X^T X \beta &= \beta^T Q^T \Lambda Q \beta \\
 &\geq \beta^T Q^T \text{diag}(\delta^2, \dots, \delta^2) Q \beta \\
 &\geq \delta^2 \|Q\beta\|_2^2 = \delta^2 \|\beta\|_2^2,
 \end{aligned}$$

因此,

$$\|(I - \psi_2)X_1\beta_1\|_2^2 \geq \min_{\beta_2} \delta^2 \|\beta\|_2^2 = \delta^2 \|\beta_1\|_2^2.$$

由于 $|J| \in A_2$, 我们可以得到 $|J_*|/|J| \leq s_*$, 进一步可得 δ 满足 $\delta \geq \sqrt{n}v$, 其中 $v = v_{(2+\epsilon)s_*}$, 故有

$$\alpha_J = \|(I - \psi_J)X_{J_*}\beta_*\|_2 = \|(I - \psi_J)X_{J_*/J}\beta_{J_*/J}\|_2 \geq \sqrt{n}v \|\beta_{J_*/J}\|_2.$$

因此,

$$\alpha_J \geq \rho v \sqrt{n(s_* - k)}. \tag{3-46}$$

由 (3-42) 和 (3-46), 进一步有

$$\begin{aligned} y^T(\psi_J - \psi_{J_*})y &\leq (35 + 7c)\sigma^2(s \vee s_* - k)\log p - \frac{1}{2}\alpha_J^2 \\ &\leq (35 + 7c)\sigma^2(s \vee s_* - k)\log p - \frac{1}{2}\rho^2 v^2 n(s_* - k)I_{(s \leq k + (1+\varepsilon)s_*)}, \end{aligned} \quad (3-47)$$

由 $J \in A_2$, 有 $s_* < s < (1 + \varepsilon)s_*$. 令 $s = (1 + t)s_*$, 则 $t \in (0, \varepsilon)$, 至少以概率 $1 - p^{-c}$ 下式成立

$$\begin{aligned} \frac{\Pi(J)}{\Pi(J_*)} &\leq \exp\{m(s_* - s) + bm\log\left(\frac{s}{s_*}\right) + (21 + \frac{7c}{2})(s - k)\log p - \frac{\rho^2 v^2 n}{4\sigma^2}(s_* - k)\} \\ &= \exp\{m(s_* - s) + bm\log\left(\frac{s}{s_*}\right) + (21 + \frac{7c}{2})s\log p - \frac{\rho^2 v^2 n}{4\sigma^2}s_* + (\frac{\rho^2 v^2 n}{4\sigma^2} - (21 + \frac{7c}{2})\log p)k\} \\ &\leq \exp\{m(s_* - s) + bm\log\left(\frac{s}{s_*}\right) + (21 + \frac{7c}{2})s\log p - \frac{\rho^2 v^2 n}{4\sigma^2}s_* + (\frac{\rho^2 v^2 n}{4\sigma^2} - (21 + \frac{7c}{2})\log p)s_*\} \\ &\leq \exp\{-ms_*t + bm\log(1 + t) + (21 + \frac{7c}{2})s_*t\log p\}. \end{aligned} \quad (3-48)$$

由于 $\frac{\rho^2 v^2 n}{4\sigma^2} - (21 + \frac{7c}{2})\log p = m \geq 0$, 因此有 $m < ms_*$, 即可得到第二个不等式. 记关于 t 的函数 $g(t) = -ms_*t + bm\log(1 + t) + (21 + \frac{7c}{2})s_*t\log p$, $t \in (0, \varepsilon)$. 对 $g(t)$ 关于 t 求导得到下式

$$\frac{\partial g(t)}{\partial t} = -(m - (21 + \frac{7c}{2})\log p)s_* + \frac{bm}{1 + t} < bm - (m - (21 + \frac{7c}{2})\log p),$$

令 $bm - (m - (21 + \frac{7c}{2})\log p) < 0$, 有

$$b < \frac{m - (21 + \frac{7c}{2})\log p}{m}s_* = \frac{1 + b}{1 + \varepsilon}s_*,$$

根据条件 $s_* - 1 > \varepsilon$, 则 $\frac{1+b}{1+\varepsilon}s_* > 1 + b > b$, 因此 $bm - (m - (21 + \frac{7c}{2})\log p) < 0$ 恒成立, 即函数 $g(t)$ 在 $(0, \varepsilon)$ 上单调递减, 因此

$$g(t) < g(0) = 0, \quad (3-49)$$

由 (3-49)、(3-50) 式可得, 当 $J \in A_2$ 时, 至少以概率 $1 - 2p^{-c}$ 有

$$\frac{\Pi(J)}{\Pi(J_*)} \leq \exp\{-ms_*t + bm\log(1 + t) + (21 + \frac{7c}{2})s_*t\log p\} \leq 1, \quad (3-50)$$

即使得后验概率达到最大的支撑集依概率收敛到真实支撑集。

当 $J \in A_3$ 时, 有 $\log \frac{s}{s_*} < 0$, 故 $bm \log \frac{s}{s_*} < 0$, 有

$$\begin{aligned} \frac{\Pi(J)}{\Pi(J_*)} &\leq \exp\{m(s_* - s) + bm \log(\frac{s}{s_*}) + (21 + \frac{7c}{2})(s - k) \log p - \frac{\rho^2 v^2 n}{4\sigma^2}(s_* - k)\} \\ &\leq \exp\{m(s_* - s) + (21 + \frac{7c}{2})(s - k) \log p - \frac{\rho^2 v^2 n}{4\sigma^2}(s_* - k)\}. \end{aligned} \quad (3-51)$$

由于 $\frac{\rho^2 v^2 n}{4\sigma^2} - (21 + \frac{7c}{2}) \log p = m$, 则

$$m(s_* - s) + (21 + \frac{7c}{2})(s - k) \log p - \frac{\rho^2 v^2 n}{4\sigma^2}(s_* - k) = -m(s - k) < 0. \quad (3-52)$$

故有 $\frac{\Pi(J)}{\Pi(J_*)} \leq 1$, 即 $\Pi(J_*) \geq \Pi(J)$ 。

综上所述, 定理 3.2 得证

定理 3.3 的证明

证明

由定理 3.3 条件, $b \leq \varepsilon \leq 1 + 2b$, 则有 $\frac{1 + \varepsilon}{\varepsilon - b} \geq 2$, 此时 $m = \frac{1 + \varepsilon}{\varepsilon - b}(21 + \frac{7c}{2}) > (31 + 6c) \log p$ 。因此与定理 3.1 相比, 定理 3.3 的条件更为严格, 于是定理 3.1 的理论结果, 在定理 3.3 中同样具有。同时, 根据定理 3.2 的 (3-45) 证明推导, 当 $J \in A_1 = \{J : |J| \geq (1 + \varepsilon)s_*\}$ 时, 至少以概率 $1 - 2p^{-c}$ 有, $\Pi(J) \leq \Pi(J_*)$ 。即对任意 $|J| > (1 + \varepsilon)s_*$, $|\hat{f}_{map}| \leq (1 + \varepsilon)s_*$ 成立的概率不小于 $1 - 2p^{-c}$ 。因此, 综合上述讨论, 以下证明过程成立概率的将不会小于 $1 - 3p^{-c}$ 。

根据 v_l 的定义, $v_l^2 \leq \lambda_{\min}(\frac{1}{n}X_J^T X_J)$ 且 $\lambda_{\min}(\frac{1}{n}X_J^T X_J)$ 表示矩阵 $\frac{1}{n}X_J^T X_J$ 最小特征根, 其中, $|J| \leq l$ 。由瑞丽熵^[43]的性质, 我们可以得到

$$v_l^2 \leq \lambda_{\min}(\frac{1}{n}X_J^T X_J) \leq \frac{\beta^T (\frac{1}{n}X_J^T X_J) \beta}{\beta^T \beta} = \frac{\|X\beta\|_2^2}{n\|\beta\|_2^2}, \quad (3-53)$$

其中, β 是 p 维非零向量且有: $\|\beta\|_0 \leq l$ 。

由于 $\|\hat{\beta}_{map} - \beta_*\|_0 \leq (1 + \varepsilon)s_* + s_* = (2 + \varepsilon)s_*$, 即

$$v_{(2+\varepsilon)s_*}^2 \leq \frac{\|X(\hat{\beta}_{map} - \beta_*)\|_2^2}{n\|\hat{\beta}_{map} - \beta_*\|_2^2}, \quad (3-54)$$

进一步可有

$$\|\hat{\beta}_{map} - \beta_*\|_2^2 \leq \frac{\|X(\hat{\beta}_{map} - \beta_*)\|_2^2}{nv_{(2+\varepsilon)s_*}^2} \leq \frac{4\sigma^2(b+3)ms_*}{nv_{(2+\varepsilon)s_*}^2}. \quad (3-55)$$

若 (3-12) 式满足, 且 $\varepsilon \leq s_* - 1$, 则定理 3.2 的条件, 定理 3.3 也满足, 即至少以概率 $1 - 2p^{-c}$ 有 $\hat{J}_{map} = J_*$, 则进一步有 $\|\hat{\beta}_{map} - \beta_*\|_0 \leq s_*$, 与 (3-54)、(3-55) 类似, 我们可以得到

$$\|\hat{\beta}_{map} - \beta_*\|_2^2 \leq \frac{4\sigma^2(b+3)ms_*}{nv_{s_*}^2}, \quad (3-56)$$

对式子 (3-54 与 (3-55) 两边开根号即得到想要的结果, 定理 3.3 得证。

定理 3.4 的证明

证明

对于任意给定的实数 $a > 0$, 可以得到

$$\begin{aligned} P(\|\hat{\beta}_{map} - \beta_*\|_\infty > a) &= P(\|\hat{\beta}_{J_*} - \beta_*\|_\infty > a, \hat{J}_{map} = J_*) + P(\|\hat{\beta}_{map} - \beta_*\|_\infty > a, \hat{J}_{map} \neq J_*) \\ &\leq P(\|\hat{\beta}_{J_*} - \beta_*\|_\infty > a, \hat{J}_{map} = J_*) + P(\hat{J}_{map} \neq J_*). \end{aligned} \quad (3-57)$$

对于 (3-27) 式右边的第一项, 当只考虑支撑集上的系数向量的最小二乘估计时, 其为 $\hat{\beta}_{J_*} = (X_{J_*}^T X_{J_*})^{-1} X_{J_*}^T y$, 而真正的最小二乘估计需要通过以零元素填充的形式得到, 满足 $\hat{\beta}_{J_*} \sim N(\beta_{J_*}, \frac{\sigma^2}{n} \Phi_*^{-1})$, 其中矩阵 $\Phi_* = \frac{1}{n} (X_{J_*}^T X_{J_*})$, 为实对称矩阵。因而对任意的 $j \in J_*$, $\hat{\beta}_{J_*}$ 的任一分量均服从如下的正态分布

$$\hat{\beta}_{J_*,j} - \beta_{J_*,j} \sim N(0, \frac{\sigma^2 \xi_j^2}{n}), \quad (3-58)$$

容易看出, ξ_j^2 为矩阵 Φ_*^{-1} 的第 j 个对角元素。

令非零向量 u 为以下的特殊向量

$$u_{(l)} = \begin{cases} 1, & l = j, \\ 0, & \text{else.} \end{cases}$$

则 Φ_*^{-1} 的第 j 个对角元素 $\xi_j^2 = \frac{u^T \Phi_*^{-1} u}{u^T u}$ 。由瑞利熵^[43]的性质, 可知 $\xi_j^2 \leq \lambda_{\max}(\Phi_*^{-1})$ 。因为 Φ_* 是实对称阵, 可知 $\lambda_{\max}(\Phi_*^{-1}) = (\lambda_{\min}(\Phi_*))^{-1}$ 。因此可以得到如下的不等式关系

$$\xi_j^2 \leq (\lambda_{\min}(\Phi_*))^{-1} \leq (v_{s_*}^2)^{-1},$$

集合 (3-58), 可进一步得到 $\text{var}(\hat{\beta}_{J_*,j}) \leq \frac{\sigma^2}{nv_{s_*}^2}$ 。

对于 (3-57) 不等号右侧的第一项, 通过简单的放缩变换可以得到下式,

$$\begin{aligned}
 P(|\hat{\beta}_{J_*} - \beta_{J_*}|_\infty > a, \hat{J}_{map} = J_*) &= P(\max_j |\hat{\beta}_{J_*,j} - \beta_{J_*,j}| > a, \hat{J}_{map} = J_*) \\
 &\leq P(\cup_j |\hat{\beta}_{J_*,j} - \beta_{J_*,j}| > a, \hat{J}_{map} = J_*) \\
 &\leq \sum_j P(|\hat{\beta}_{J_*,j} - \beta_{J_*,j}| > a),
 \end{aligned} \tag{3-59}$$

其中,

$$\begin{aligned}
 P(|\hat{\beta}_{J_*,j} - \beta_{J_*,j}| > a) &= P\left(\frac{|\hat{\beta}_{J_*,j} - \beta_{J_*,j}|}{\sqrt{\sigma^2 \xi_j^2/n}} > \frac{a}{\sqrt{\sigma^2 \xi_j^2/n}}\right) \\
 &\leq P\left(\frac{|\hat{\beta}_{J_*,j} - \beta_{J_*,j}|}{\sqrt{\sigma^2 \xi_j^2/n}} > \frac{a}{\sqrt{\sigma^2/nv_{s_*}^2}}\right).
 \end{aligned} \tag{3-60}$$

我们记 $\zeta_j = \frac{\hat{\beta}_{J_*,j} - \beta_{J_*,j}}{\sqrt{\sigma^2 \xi_j^2/n}}$, 结合 (3-58), 可以得到 $\zeta_j \sim N(0, 1)$ 。因此 (3-60) 式可以简化为如下形式,

$$P(|\hat{\beta}_{J_*,j} - \beta_{J_*,j}| > a) \leq P(|\zeta_j| > \frac{\sqrt{nav_{s_*}}}{\sigma}). \tag{3-61}$$

由于 ζ_j 服从标准正态分布, 故有 $P(|\zeta_j| > t) = 2P(\zeta_j > t)$ 。因此,

$$\begin{aligned}
 P(\zeta_j > t) &= \int_t^{+\infty} \frac{1}{\sqrt{2\pi}} \exp(-\frac{\zeta_j^2}{2}) d\zeta_j \\
 &\leq \frac{1}{\sqrt{2\pi}} \int_t^{+\infty} \frac{\zeta_j}{t} \exp(-\frac{\zeta_j^2}{2}) d\zeta_j \\
 &= \frac{1}{\sqrt{2\pi}t} \exp(-\frac{t^2}{2}).
 \end{aligned} \tag{3-62}$$

由于 (3-62) 中 ζ_j 的积分区域为 t 到 $+\infty$, 在积分区域中 $\frac{\zeta_j}{t} \geq 1$, 因此 (3-62) 中的不等号成立。

令 $a = \sigma \sqrt{\frac{2(c+1)\log p}{nv_{s_*}^2}}$, 由 (3-60)、(3-62) 可推导得到

$$\begin{aligned}
 P(|\zeta_j| > \frac{\sqrt{nav_{s_*}}}{\sigma}) &= 2P(\zeta_j > \frac{\sqrt{nav_{s_*}}}{\sigma}) \\
 &\leq 2 \frac{1}{\sqrt{2\pi}} \frac{1}{\frac{\sqrt{nav_{s_*}}}{\sigma}} \exp(-\frac{1}{2} \frac{nv_{s_*}^2 a^2}{\sigma^2}) \\
 &= \frac{p^{-(c+1)}}{\sqrt{\pi(c+1)\log p}},
 \end{aligned} \tag{3-63}$$

由 a 的取值可以得到, $\frac{\sqrt{nav_{s_*}}}{\sigma} = \sqrt{2(c+1)\log(p)} > 1$, 结合 (3-63) 式, 可以进一步得到

$$\begin{aligned}
 P(\|\hat{\beta}_{J_*} - \beta_*\|_\infty > a, \hat{J}_{map} = J_*) &\leq \sum_j P(|\hat{\beta}_{J_*,j} - \beta_{*,j}| > a) \\
 &\leq \sum_j P(|\zeta_j| > \frac{\sqrt{nav_{s_*}}}{\sigma}) \\
 &= \frac{s_*}{\sqrt{\pi(c+1)\log p}} p^{-(c+1)} \\
 &\leq p^{-c}.
 \end{aligned} \tag{3-64}$$

回顾定理 3.2 的结论, 可知 $\hat{J}_{map} = J_*$ 成立的概率不小于 $1 - 2p^{-c}$, 因此可以得到 $P(\hat{J}_{map} \neq J_*) \leq 2p^{-c}$ 。

结合 (3-64) 与上述推导过程, 可得

$$\begin{aligned}
 P(\|\hat{\beta}_{map} - \beta_*\|_\infty > a) &\leq P(\|\hat{\beta}_{J_*} - \beta_*\|_\infty > a, \hat{J}_{map} = J_*) + P(\hat{J}_{map} \neq J_*) \\
 &\leq p^{-c} + 2p^{-c} \\
 &= 3p^{-c}.
 \end{aligned} \tag{3-65}$$

因此, 下列不等式成立的概率不小于 $1 - 3p^{-c}$, 即

$$\|\hat{\beta}_{map} - \beta_*\|_\infty \leq \sigma \sqrt{\frac{2(c+1)\log p}{nv_{s_*}^2}}. \tag{3-66}$$

定理的第一部分得证。再次回顾定理 3.2 的结论, $\hat{J}_{map} = J_*$ 成立的概率不小于 $1 - 2p^{-c}$, 等价于至少以概率 $1 - 2p^{-c}$ 有 $\|\hat{\beta}_{map} - \beta_*\|_0 = s_*$ 成立, 因此

$$\begin{aligned}
 \|\hat{\beta}_{map} - \beta_*\|_1 &= \sum_{j=1}^p |\hat{\beta}_{map,j} - \beta_{*,j}| \\
 &\leq s_* \|\hat{\beta}_{map} - \beta_*\|_\infty \\
 &\leq \sigma s_* \sqrt{\frac{2(c+1)\log p}{nv_{s_*}^2}}.
 \end{aligned} \tag{3-67}$$

综上所述, 定理 3.4 得证。

定理 3.5 的证明

证明

由 (3-8) 式, 有下式成立

$$\begin{aligned}
 \|\hat{\beta}_{mean} - \beta_*\|_\infty &= \left\| \sum_J (\hat{\beta}_J - \beta_*) \Pi(J) \right\|_\infty \\
 &\leq \sum_J \|\hat{\beta}_J - \beta_*\|_\infty \Pi(J) \\
 &\leq \|\hat{\beta}_{J_*} - \beta_*\|_\infty \Pi(J_*) + \sum_{J \neq J_*} \|\hat{\beta}_J - \beta_*\|_\infty \Pi(J) \\
 &\leq \|\hat{\beta}_{J_*} - \beta_*\|_\infty + \sum_{J \neq J_*} \|\hat{\beta}_J - \beta_*\|_\infty \Pi(J),
 \end{aligned} \tag{3-68}$$

其中第一个不等式由柯西施瓦茨不等式得到。当 $v_{l+s_*} > 0$ 成立时, 对任意 $J \in K(\Omega_p)$ 且 $|J| < l$, 结合定理 3.3, 下列式子成立的概率不小于 $1 - 3p^{-c}$, 即

$$\|\hat{\beta}_J - \beta_*\|_\infty \leq \|\hat{\beta}_J - \beta_*\|_2 \leq \frac{\|X\hat{\beta}_J - X\beta_*\|_2}{\sqrt{nv_{l+s_*}}}. \tag{3-69}$$

由于 $\Theta_0 = \{J \in K(\Omega_p) : \|\tau_J\|_2 > \sigma \sqrt{6(3+b)ms_*}\}$, 且 $\tau_J = \psi_J y - X\beta_* = X\hat{\beta}_J - X\beta_*$. 结合 (3-69) 和 (3-70) 式, 我们可以得到

$$\begin{aligned}
 \|\hat{\beta}_{mean} - \beta_*\|_\infty &\leq \|\hat{\beta}_{J_*} - \beta_*\|_\infty + \sum_{J \neq J_*} \|\hat{\beta}_J - \beta_*\|_\infty \Pi(J) \\
 &\leq \|\hat{\beta}_{J_*} - \beta_*\|_\infty + \frac{1}{\sqrt{nv_{l+s_*}}} \sum_{J \notin \Theta_0, J \neq J_*} \|\tau_J\|_2 \Pi(J) + \frac{1}{\sqrt{nv_{l+s_*}}} \sum_{J \in \Theta_0, J \neq J_*} \|\tau_J\|_2 \Pi(J) \\
 &\leq \|\hat{\beta}_{J_*} - \beta_*\|_\infty + \frac{\sigma \sqrt{6(3+b)ms_*}}{\sqrt{nv_{l+s_*}}} + \frac{1}{\sqrt{nv_{l+s_*}}} \sum_{J \in \Theta_0} \|\tau_J\|_2 \Pi(J).
 \end{aligned} \tag{3-70}$$

回顾 (3-38) 式, 至少以概率 $1 - p^{-c}$, 我们有

$$\sum_{J \in \Theta_0} \|\tau_J\|_2 \Pi(J) < \sigma p^{-s_*}. \tag{3-71}$$

事实上, 当 $l \gg s_*$ 时, 条件 $I((l+s_*))$ 是比 $I((2+\varepsilon)s_*)$ 要严格的。因此定理 3.5 的条件同样满足定理 3.4, 结合 (3-66)、(3-70)、(3-71), 至少以概率 $1 - 4p^{-c}$, 我们可以得到

$$\|\hat{\beta}_{mean} - \beta_*\|_\infty \leq \sigma \sqrt{\frac{2(c+1) \log p}{nv_{s_*}^2}} + \frac{\sigma \sqrt{6(3+b)ms_*}}{\sqrt{nv_{l+s_*}}} + \frac{\sigma p^{-s_*}}{\sqrt{nv_{l+s_*}}}. \tag{3-72}$$

综上即证定理 3.5。

3.5 本章小结

本章通过对模型空间大小的讨论,构造了关于系数向量的 l_0 收缩指数先验分布,该先验分布是对系数向量支撑集的先验信息的定义,在理论性质方面,本章推导了推广指数权重法下所得系数向量的贝叶斯估计的预测误差限,证明了待估系数向量支撑集的最大后验估计依概率收敛到真实的系数向量支撑集,同时表明,推广指数权重方法所得参数估计是相合的,只要满足当 n 趋于无穷时, $\sqrt{\log p/n}$ 趋于 0,最后对理论结果进行了相应的证明。

第4章 稀疏学习算法及数值模拟

4.1 稀疏学习算法

本节主要以 l_0 收缩指数先验所对应的最大后验估计为目标, 将后验概率最大化问题转化为等价的结构风险最小化问题, 通过求解使得结构风险最小化的解作为后验概率最大化的解。本节求最优支撑和参数估计的优化目标函数为下式

$$f_{EEW}(J) = \frac{1}{2\sigma^2} \|\psi_J^\top y\|^2 + m|J| + \log \binom{p}{|J|} - bm \log |J|. \quad (4-1)$$

事实上, 上述目标函数是在 $|J| < l$ 条件下的转换, 上述结构风险最小化表述为常见的线性模型加上基于 l_0 范数变体的惩罚项。相比于指数权重方法, 新增加的一项 $-bm \log |J|$ 在一定条件下能够防止过度压缩, 在数值模拟部分将会呈现。受 SBR 算法的启发, 本文基于后验分布 (3-4) 的等价目标函数 $f_{EEW}(J)$, 构造了基于推广指数权重的单一最优替换 (Extended Exponential Weights with Single Best Replacement, EEW-SBR) 稀疏学习算法, 通过该算法得到 $f_{EEW}(J)$ 的局部最优解。在 EEW-SBR 算法中, 通过先初始化原始支撑集为非空稀疏集, 通过逐步增加或删除某一支撑索引, 得到所有可能的支撑集合, 再计算得到使得当前所有可能支撑集下的目标函数值 $f_{EEW}(J \cdot i)$ 最小的支撑索引 i , 在当前支撑集合的基础上将索引 i 加入或删除, 并进入下一步迭代。当没有单个替换能够使得目标函数值再一次降低时, 迭代停止。在特殊情况下, 如 $m = 0$, 此时 EEW-SBR 算法退化为广为人知的正交最小二乘算法^[61], 此时单一变换操作只会增加而不会删除支撑元素, 因为删除操作会增加拟合误差。其伪代码总结在算法 4.1 中。

算法 4.1 EEW-SBR 稀疏学习算法

输入: 设计矩阵 X , 响应向量 y , 支撑上界 l , 超参数 m, b , 随机误差的方差 σ^2 。

输出: 待估系数向量的支撑估计 $\hat{f}_{map}$ 以及系数估计 $\hat{\beta}_{map}$ 。

- 1: 随机初始化非空支撑集合: J_0 。
 - 2: 定义单次替换操作: $J \cdot i$, 若 $i \in J$, 则 $J \cdot i = J \setminus i$; 否则 $J \cdot i = J \cup i, \forall i \in \Omega_p := \{1, 2, \dots, p\}$ 。
 - 3: 计算: $J_k = J_{k-1} \cdot j, e_k = \min_{i \in \Omega_p} f_{EEW}(J_{k-1} \cdot i)$, 以及 $j = \arg \min_{i \in \Omega_p} f_{EEW}(J_{k-1} \cdot i)$ 。
 - 4: 停机准则: 若 $e_k \geq e_{k-1}$, 输出 $\hat{f}_{map} = J_{k-1}$ 并进行第 5 步; 否则, $k = k + 1$, 继续进行第 3 步。
 - 5: 系数计算: 若 $r(X_{\hat{f}_{map}}) = |\hat{f}_{map}|$, $\hat{\beta}_{map} = (X_{\hat{f}_{map}}^T X_{\hat{f}_{map}})^{-1} X_{\hat{f}_{map}}^T y$; 否则 $\hat{\beta}_{map} = (X_{\hat{f}_{map}}^T X_{\hat{f}_{map}})^{-1} X_{\hat{f}_{map}}^T y$ 。输出 $\hat{\beta}_{map}$ 。
-

其中 $r(X_{f_{map}})$ 表示矩阵 $X_{f_{map}}$ 的秩。由 EEW-SBR 算法的过程可知, 同一支撑集 J , 不会重复被学习到, 而对于所有的 $J \in K(\Omega_p)$, 有 2^p 个支撑集, 即支撑集的个数是有限的, 因此 EEW-SBR 算法会在有限次迭代后终止, 得到一个局部最优解。

4.2 数值模拟

本小节采用模拟数据, 方差大小 σ^2 可以通过预先给定的 $\text{SNR} = \|X\beta_\star\|_2^2 / \|w\|_2^2$ 计算得到。在实际问题的研究过程中, 真实方差可以通过样本方差替代。在比较 EW 方法与 EEW 方法时, 超参数 m 与 b 的选择十分重要。对于 EW 方法, 单个超参数 m 直接决定了其正则化项的重要程度。通过引入的这一平衡项, 一定程度上加大了先验反应真实模型的信息。在其选取技巧上, 本文先通过对 EW 方法采取 K 折交叉验证的随机搜索方法选取超参数 m , 经过多次模拟, 确定最优 m , 并在此基础上设置与真实稀疏度相关的超参数 b 。估计的学习好坏是由支撑集决定的, 因此在本章节的模拟过程中, 所定义的真系数向量 $\beta_\star$ 的各个分量具有以下形式

$$\beta_{\star,j} = \begin{cases} 1, & j \leq s, \\ 0, & j > s. \end{cases} \quad (4-2)$$

在本小节中, 采用如下三个评价指标,

$$\begin{aligned} e_\infty &= \frac{1}{N} \sum_{i=1}^N \|\hat{\beta}^i - \beta_\star\|_\infty, \\ e_1 &= \frac{1}{N} \sum_{i=1}^N \|\hat{\beta}^i - \beta_\star\|_1, \\ e_2 &= \frac{1}{N} \sum_{i=1}^N \|\hat{\beta}^i - \beta_\star\|_2. \end{aligned} \quad (4-3)$$

(4-3) 中所展示的三个关于参数估计的评价指标均可用于衡量估计的效果。同时, 在支撑恢复方面, 本节同样定义了相应的评价指标。首先, 我们定义支撑的正向选择率为 $C = |J \cap J_\star| / s_\star$, 以及零元素的正向选择率为 $C_0 = |(\Omega_p \setminus J) \cap (\Omega_p \setminus J_\star)| / (p - s_\star)$ 。当 C 值越接近于 1, 表示对于重要特征, 本文方法能正确识别出来, 表明支撑恢复效果突出。当 C_0 值越接近于 1, 表示能够正确的剔除模型中的无关变量。通过上述关于 C 与 C_0 评价指标的定义, 可以对支撑恢复的好坏起到进行一定的衡量作用。

以下的数值模拟分为三个部分, 每个部分的模拟设置均不相同, 我们均进行了详细的介绍。

首先,为了比较本文的 EEW 方法与 EW 方法的优劣,展示低 SNR 情况下 EW 方法的估计的不足,我们通过展示模拟所得参数估计与支撑恢复随 SNR 变化的趋势图来说明,此部分只考虑了 e_2 误差以及支撑的正向选择率 C 这两个指标。在经过大量数值模拟结果对比以及分析,我们将 EW 方法中的超参数 m 的随机采样区间限定为 $[0,3]$,随机采样 20 次,再通过 5 折交叉验证法选出能够学习到较低估计误差的超参数 m 。而对于 EEW 方法,则是通过随机数对来实现。对于 EEW 方法的超参数 m ,与 EW 方法的采样区间相同,但仅进行 5 次的 m 随机均匀采样。而超参数 b ,根据第三章的理论性质分析, b 的取值不应超过真实模型稀疏度,故我们限制 b 的随机均匀采样区间为 $[0,s_*]$,对于每个 m 进行 10 次的 b 采样,得到 50 个超参数对,再从中最优超参数对,最后进行相应的估计计算。在此部分,设置初始 SNR 为 1,并逐步以 0.25 的大小数值逐步相加,得到一个具有 20 个 SNR 值的序列,对于每一个 SNR 值,均随机模拟 100 次,即 $N = 100$ 。我们取 $n = 50$ 和 $p = 60$, $s_* = 6$ 进行模拟,其相应的模拟趋势图如下:

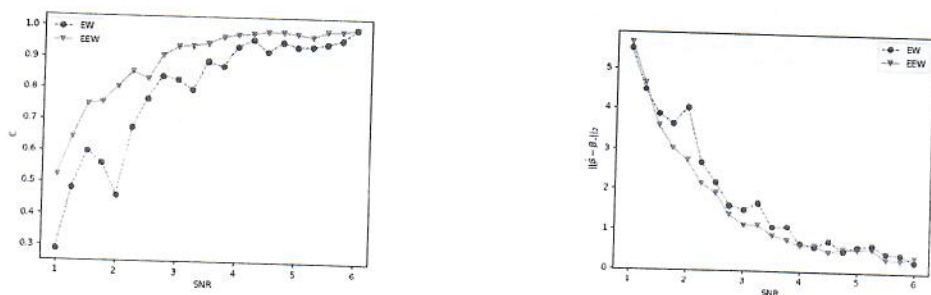


图 4-1 $n = 50$, $p = 60$, $s_* = 6$, 变量选择方法 EW, EEW, 所得参数估计和支撑正向选择率随 SNR 变化趋势图。左: C , 右: e_2 。

从上图可以看出,EEW 方法在偏低 SNR 时,能够得到更高的支撑正向选择率,以及偏低的 e_2 误差,随着 SNR 的增大,差距逐渐缩小,因此 EEW 方法在偏低 SNR 情况下相较于 EW 具有一定的改进作用。

其次,为了更为细致的展示 EEW 在 EW 方法上的改进,本部分进行了多维度的展开,即进行了多维度的模拟比较。EW 方法的超参数 m 的随机采样区间限定为 $[0,3]$,随机采样 30 次,同样的,再基于 5 折交叉验证选出能够学习到较低估计误差的超参数 m 。而对于 b 的选择,与之前的不同, $b = \frac{s_*}{10} \cdot i$, 其中 $i \in \{0,1,2,3,4\}$ 。注意到当 $b = 0$ 时,EEW 方法与 EW 方法是相同的,故在图表中, $b = 0$ 表示 EW 方法,其他值为本文的 EEW 方法。在此处,我们对两种方法均进行 100 次数值模拟,即 $N = 100$ 。同时,我们将 SNR 的值预先设置为 2,3,4。SNR 固定,在 (n, p, s_*) 分别为 $(100, 120, 10)$, $(50, 200, 5)$ 时进行实验。以下表为数值模拟对比结果。

表 4-1 $SNR = 2, n = 100, p = 120, s_* = 10$ 下, 两种方法指标 $e_1, e_2, e_\infty, C, C_0$ 的值.Table 4-1 The values of $e_1, e_2, e_\infty, C, C_0$ for two methods ($SNR = 2, n = 100, p = 120, s_* = 10$).

i	e_1	e_2	e_∞	C	C_0
$i = 0$	3.89686	2.62417	0.79405 _{0.08395}	0.860	0.98891
$i = 1$	3.90319	2.62445	0.78813 _{0.08021}	0.854	0.98964
$i = 2$	3.86806	2.51713	0.78196 _{0.07275}	0.874	0.98745
$i = 3$	3.98960	2.62609	0.79245 _{0.07787}	0.866	0.98727
$i = 4$	4.00854	2.56293	0.79046 _{0.06767}	0.880	0.98509

表 4-2 $SNR = 3, n = 100, p = 120, s_* = 10$ 下, 两种方法指标 $e_1, e_2, e_\infty, C, C_0$ 的值.Table 4-2 The values of $e_1, e_2, e_\infty, C, C_0$ for two methods ($SNR = 3, n = 100, p = 120, s_* = 10$).

i	e_1	e_2	e_∞	C	C_0
$i = 0$	2.54237	1.17178	0.61885 _{0.07253}	0.962	0.98964
$i = 1$	2.47317	1.10324	0.61098 _{0.06951}	0.968	0.98964
$i = 2$	2.61870	1.17346	0.62057 _{0.06447}	0.968	0.98782
$i = 3$	2.76488	1.28887	0.61631 _{0.06230}	0.968	0.98582
$i = 4$	2.73863	1.25451	0.62001 _{0.05928}	0.970	0.98582

表 4-3 $SNR = 4, n = 100, p = 120, s_* = 10$ 下, 两种方法指标 $e_1, e_2, e_\infty, C, C_0$ 的值.Table 4-3 The values of $e_1, e_2, e_\infty, C, C_0$ for two methods ($SNR = 4, n = 100, p = 120, s_* = 10$).

i	e_1	e_2	e_∞	C	C_0
$i = 0$	2.02038	0.76494	0.47187 _{0.06002}	0.976	0.99309
$i = 1$	2.00994	0.76270	0.47372 _{0.06018}	0.976	0.99309
$i = 2$	2.03044	0.76846	0.47103 _{0.06000}	0.976	0.99291
$i = 3$	1.99991	0.71611	0.48253 _{0.05762}	0.982	0.99236
$i = 4$	2.04745	0.73101	0.48490 _{0.05029}	0.984	0.99127

表 4-4 $SNR = 2, n = 50, p = 200, s_* = 5$ 下, 两种方法指标 $e_1, e_2, e_\infty, C, C_0$ 的值.Table 4-4 The values of $e_1, e_2, e_\infty, C, C_0$ for two methods ($SNR = 2, n = 50, p = 200, s_* = 5$).

i	e_1	e_2	e_∞	C	C_0
i = 0	3.86840	3.45099	0.94781 _{0.11804}	0.620	0.99344
i = 1	4.02379	3.63638	0.94993 _{0.13987}	0.612	0.99272
i = 2	4.18775	3.75554	0.97274 _{0.10171}	0.600	0.99179
i = 3	3.91300	3.42587	0.95954 _{0.10257}	0.644	0.99251
i = 4	3.94986	3.39367	0.92232 _{0.10925}	0.660	0.99128

表 4-5 $SNR = 3, n = 50, p = 200, s_* = 5$ 下, 两种方法指标 $e_1, e_2, e_\infty, C, C_0$ 的值.Table 4-5 The values of $e_1, e_2, e_\infty, C, C_0$ for two methods ($SNR = 3, n = 50, p = 200, s_* = 5$).

i	e_1	e_2	e_∞	C	C_0
i = 0	1.83263	1.27842	0.55987 _{0.11802}	0.896	0.99621
i = 1	1.91550	1.25826	0.57110 _{0.11994}	0.896	0.99579
i = 2	1.85578	1.21425	0.54867 _{0.12238}	0.904	0.99600
i = 3	2.18688	1.50381	0.59964 _{0.12827}	0.876	0.99487
i = 4	2.32539	1.59644	0.60585 _{0.11788}	0.872	0.99364

表 4-6 $SNR = 4, n = 50, p = 200, s_* = 5$ 下, 两种方法指标 $e_1, e_2, e_\infty, C, C_0$ 的值.Table 4-6 The values of $e_1, e_2, e_\infty, C, C_0$ for two methods ($SNR = 4, n = 50, p = 200, s_* = 5$).

i	e_1	e_2	e_∞	C	C_0
i = 0	1.83112	1.18110	0.50865 _{0.12579}	0.912	0.99467
i = 1	1.90152	1.25914	0.52387 _{0.13414}	0.904	0.99436
i = 2	1.65474	0.94503	0.49641 _{0.09604}	0.932	0.99487
i = 3	2.13346	1.48619	0.54345 _{0.14212}	0.888	0.99333
i = 4	2.01821	1.34041	0.53475 _{0.13243}	0.904	0.99354

从表 4-1 至表 4-3 可知,在不同的 SNR 设置下,本文的 EEW 方法在参数估计和非零位置支撑正向选择率方面均优于 EW 方法,在零位置正向选择率方面,EEW 方法也能表现的和 EW 方法一样好。从表 4-4 至表 4-6 可知,当 $\text{SNR}=2$ 和 $\text{SNR}=3$ 时,在 e_1 误差和零元素正向选择率上,表现较好,在其他方面 EEW 方法表现更突出。当 $\text{SNR}=4$ 时,EEW 方法全面优于 EW。上述模拟结果表明,在偏低 SNR 下,本文的 EEW 方法能够改进 EW 方法,得到更低的参数估计误差以及更为优秀的支撑恢复效果,且稳定性也偏高。

最后,除了在 SNR 偏低时,对 EEW 方法与 EW 方法进行了对比,本文同样在高 SNR 情况下,与其他变量选择方法进行了模拟比较。此部分的超参数 m 的采样范围限定在 $[3,6]$,采样 10 次,对于每次采样的 m ,再进行 10 次的超参数 b 的采样,共生成 100 组超参数对。由定理中的条件可知,超参数 b 要小于真实支撑大小,因此对于不同的支撑大小,超参数 b 的随机采样的区间相应也在变化,即限定在 $[1, s_*/2]$ 之间。通过比较超参数组合中使得待估系数支撑回复较好的组合作为最终的模拟超参数。在此部分,SCAD 与 MCP 方法的模拟实现由 R 软件完成,Lasso 方法的模拟由 Python 软件的 sklearn 包完成,EEW 和 OMP 方法均在 Python 软件中自行编码实现。在此部分,同样对所有方法均模拟 100 次,我们将 SNR 的值分别选取为 8 和 10。固定 SNR,在 (n, p, s_*) 分别为 $(50, 200, 5)$ 、 $(100, 120, 10)$ 、 $(500, 100, 10)$ 下进行实验,分别对应于 $n \ll p$, $n < p$ 和 $n > p$ 三种情况。所有模拟数据结果均呈现在下列图和表中。

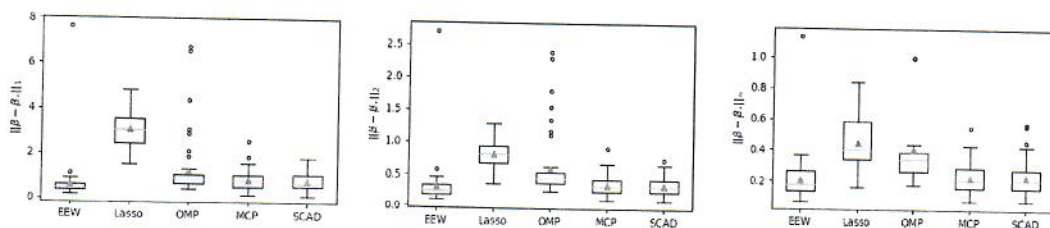


图 4-2 $\text{SNR}=8$, $n=50$, $p=200$, $s_*=5$, 在 100 次模拟数据下, 变量选择方法 EEW, Lasso, OMP, MCP, SCAD 所得参数估计误差箱型图。左: e_1 , 中: e_2 , 右: e_∞ 。

表 4-7 $SNR = 8, n = 50, p = 200, s_* = 5$ 下, 各个方法下指标 $e_1, e_2, e_\infty, C, C_0$ 的值.

Table 4-8 The values of $e_1, e_2, e_\infty, C, C_0$ for different methods ($SNR = 8, n = 50, p = 200, s_* = 5$).

	e_1	e_2	e_∞	C	C_0
EEW	0.65217	0.31538	0.20845 _{0.15319}	0.984	0.99959
LASSO	3.08184	0.81352	0.44560 _{0.15301}	1	0.87702
OMP	1.29016	0.59757	0.40567 _{0.24805}	0.944	0.99333
MCP	0.87285	0.35497	0.22864 _{0.09852}	1	0.98123
SCAD	0.85516	0.35167	0.23257 _{0.11555}	1	0.97476

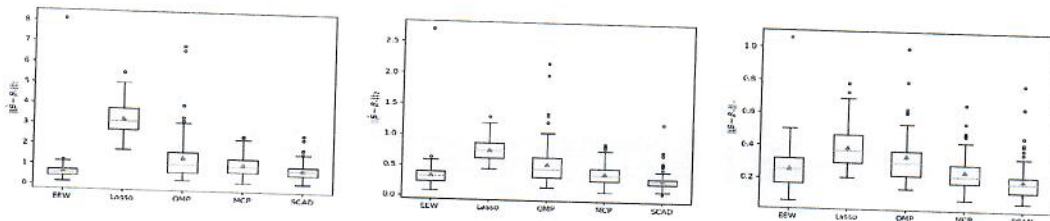


图 4-3 $SNR = 10, n = 50, p = 200, s_* = 5$, 在 100 次模拟数据下, 变量选择方法 EEW, Lasso, OMP, MCP, SCAD 所得参数估计误差箱型图. 左: e_1 , 中: e_2 , 右: e_∞ .

表 4-8 $SNR = 10, n = 50, p = 200, s_* = 5$, 各个方法下指标 $e_1, e_2, e_\infty, C, C_0$ 的值.

Table 4-8 The values of $e_1, e_2, e_\infty, C, C_0$ for different methods ($SNR = 10, n = 50, p = 200, s_* = 5$).

	e_1	e_2	e_∞	C	C_0
EEW	0.66782	0.34367	0.25925 _{0.12800}	1	0.99697
LASSO	3.26168	0.76642	0.39210 _{0.11839}	1	0.84764
OMP	1.38854	0.55521	0.33963 _{0.17653}	1	0.98487
MCP	1.07926	0.40467	0.25119 _{0.10781}	1	0.96882
SCAD	0.87334	0.31287	0.20066 _{0.10848}	1	0.95425

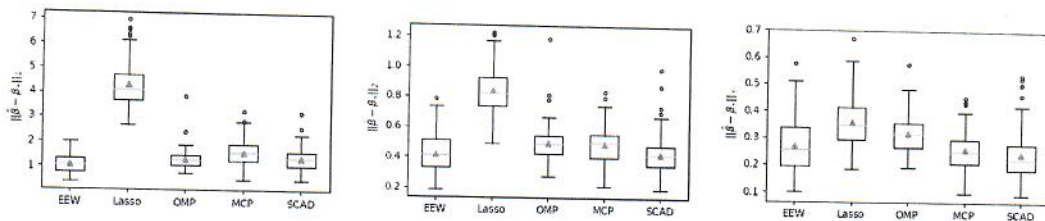
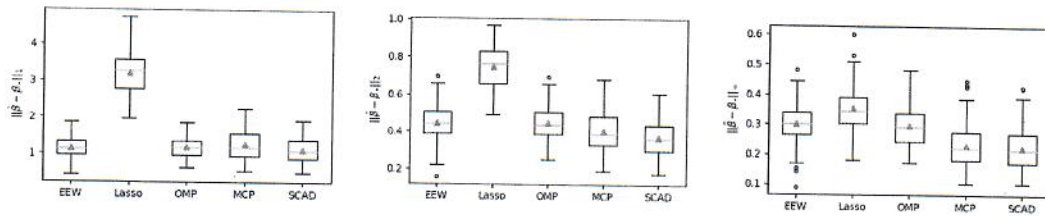


图 4-4 $SNR = 8, n = 100, p = 120, s_* = 10$, 在 100 次模拟数据下, 变量选择方法 EEW, Lasso, OMP, MCP, SCAD 所得参数估计误差箱型图. 左: e_1 , 中: e_2 , 右: e_∞ .

表 4-9 $SNR = 8, n = 100, p = 120, s_* = 10$ 下, 各个方法下指标 $e_1, e_2, e_\infty, C, C_0$ 的值.Table 4-9 The values of $e_1, e_2, e_\infty, C, C_0$ for different methods ($SNR = 8, n = 100, p = 120, s_* = 10$).

	e_1	e_2	e_∞	C	C_0
EEW	1.06642	0.41389	0.26856 _{0.09522}	1	0.99709
LASSO	4.33356	0.84268	0.35908 _{0.08639}	1	0.70036
OMP	1.31568	0.49692	0.31731 _{0.06969}	1	0.99072
MCP	1.61066	0.48833	0.25901 _{0.07607}	1	0.94281
SCAD	1.39361	0.42793	0.24244 _{0.08380}	1	0.93390

图 4-5 $SNR = 10, n = 100, p = 120, s_* = 10$, 在 100 次模拟数据下, 变量选择方法 EEW, Lasso, OMP, MCP, SCAD 所得参数估计误差箱型图。左: e_1 , 中: e_2 , 右: e_∞ 。表 4-10 $SNR = 10, n = 100, p = 120, s_* = 10$ 下, 各个方法下指标 $e_1, e_2, e_\infty, C, C_0$ 的值.Table 4-10 The values of $e_1, e_2, e_\infty, C, C_0$ for different methods ($SNR = 10, n = 100, p = 120, s_* = 10$).

	e_1	e_2	e_∞	C	C_0
EEW	1.15436	0.44960	0.30111 _{0.06891}	1	0.99336
LASSO	3.20410	0.75158	0.35414 _{0.07074}	1	0.80399
OMP	1.17730	0.45219	0.29536 _{0.06628}	1	0.99190
MCP	1.26238	0.41041	0.23165 _{0.07224}	1	0.96000
SCAD	1.13204	0.37586	0.22407 _{0.06861}	1	0.95500

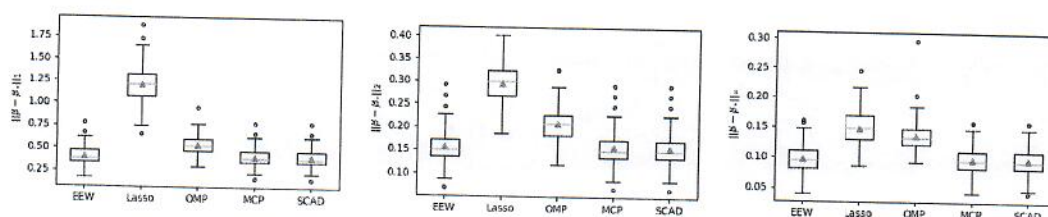


图 4-6 $SNR = 8, n = 500, p = 100, s_* = 10$, 在 100 次模拟数据下, 变量选择方法 EEW, Lasso, OMP, MCP, SCAD 所得参数估计误差箱型图。左: e_1 , 中: e_2 , 右: e_∞ 。

表 4-11 $SNR = 8, n = 500, p = 100, s_* = 10$ 下, 各个方法下指标 $e_1, e_2, e_\infty, C, C_0$ 的值。

Table 4-11 The values of $e_1, e_2, e_\infty, C, C_0$ for different methods ($SNR = 8, n = 500, p = 50, s_* = 10$).

	e_1	e_2	e_∞	C	C_0
EEW	0.40246	0.15610	0.09699 _{0.02229}	1	1
LASSO	1.20483	0.29561	0.14801 _{0.02999}	1	0.79788
OMP	0.53824	0.20836	0.13612 _{0.02632}	1	0.98955
MCP	0.40695	0.15737	0.09780 _{0.02235}	1	0.99900
SCAD	0.40579	0.15693	0.09735 _{0.02199}	1	0.99900

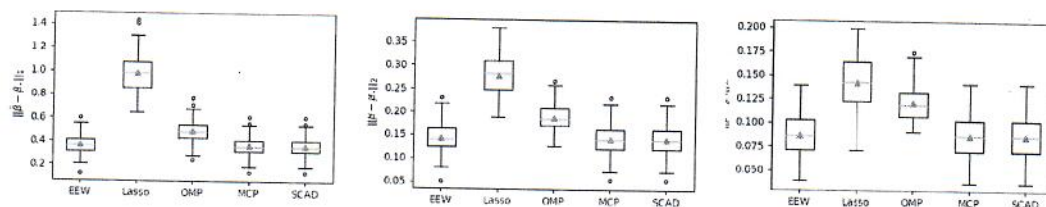


图 4-7 $SNR = 10, n = 500, p = 100, s_* = 10$, 在 100 次模拟数据下, 变量选择方法 EEW, Lasso, OMP, MCP, SCAD 所得参数估计误差箱型图。左: e_1 , 中: e_2 , 右: e_∞ 。

表 4-12 $SNR = 10, n = 500, p = 100, s_* = 10$ 下, 各个方法下指标 $e_1, e_2, e_\infty, C, C_0$ 的值。

Table 4-12 The values of $e_1, e_2, e_\infty, C, C_0$ for different methods ($SNR = 10, n = 500, p = 50, s_* = 10$).

	e_1	e_2	e_∞	C	C_0
EEW	0.37062	0.14340	0.08655 _{0.02147}	1	1
LASSO	0.98549	0.27846	0.14304 _{0.02748}	1	0.87044
OMP	0.49520	0.18951	0.12127 _{0.01949}	1	0.98977
MCP	0.37122	0.14356	0.08664 _{0.02180}	1	0.99933
SCAD	0.37114	0.14356	0.08664 _{0.02180}	1	0.99933

观察图 4-2, 在 $\text{SNR}=8$, $n=50$, $p=200$, $s_*=5$ 情况下, 此时 n 和 p 相差较大, 可以明显的看到, 本文的 EEW 方法在参数估计方面要明显优于 Lasso 方法, 且相较于 OMP、MCP 和 SCAD, 也具有略微的优势。观察表 4-7, 除了在估计方面表现优异以外, 在零元素的正向选择率方面, 均优于其他方法。其次, 在支撑的正向选择率方面, MCP、SCAD 和 Lasso 方法均达到了最优值, 且无穷范数误差方面, MCP 方法的更为稳定。观察图 4-3, Lasso 方法在参数估计方面全面劣于其他方法, 而 EEW 与 SCAD 方法表现相当, 更为细致的结果展示在表 4-8 中, 我们可以看到在 e_1 误差下, EEW 得到了更小的估计误差, 其次是 SCAD 方法。而在 e_2 和 e_∞ 下, SCAD 方法的参数估计表现更为优异。在支撑的正向选择率上, 五种方法均能正确的选择到正确的支撑位置。在零元素的正向选择率上, EEW 方法优势更为明显。

观察图 4-4, 在 $\text{SNR}=8$, $n=100$, $p=120$, $s_*=5$ 情况下, 此时 n 与 p 相差较小。与前面的结果类似, EEW 方法在参数估计方面要明显优于 Lasso 方法。观察表 4-9, 在 e_1 和 e_2 下, EEW 方法均优于其他方法, 其次是 SCAD 方法, 在 e_∞ 下, EEW 方法稍微劣于 SCAD 和 MCP 方法, 而 OMP 方法能够得到更为稳定的解。同样的, 五种方法均能正确的选择到正确的支撑位置。在零元素的正向选择率上, EEW 方法优势略微优于 OMP 方法, 明显优于其他方法。当 $\text{SNR}=10$ 时, 结果呈现在图 4-5 和表 4-10 中, 与 $\text{SNR}=8$ 时不同, SCAD 方法在参数估计方面均要优于其他方法, 同样的, OMP 方法在无穷范数估计的稳定性上更好。在支撑的正向选择率上, 五种方法均能正确的选择到正确的支撑位置。

而在 $n > p$ 时, $n=500$, $p=100$, $s_*=10$, 当 $\text{SNR}=8$ 和 $\text{SNR}=10$ 时, 不论是在参数估计还是支撑恢复方面, EEW 方法全面优于其他方法, 其次 MCP 和 SCAD 的表现是一致的, 详细结果呈现在图 4-6、图 4-7、表 4-11 以及表 4-12 中。

从模拟结果的整体表现来看, 不论是 $n > p$ 还是 $n < p$ 亦或是 $n \ll p$, 在高 SNR 情况下, EEW 方法的参数估计误差均小于 Lasso 方法, 而且在零元素的正向选择率方面, EEW 方法比其他方法均要好, 表明 EEW 方法在多数情况下均能获得更为稀疏的解。

4.3 本章小节

本章基于 l_0 收缩指数先验对应的最大后验的等价目标函数构建了 EEW-SBR 算法, 在低 SNR 和 $n < p$ 情况下, 通过数值模拟将 EW 和 EEW 方法进行了对比, 表明 EEW 方法不论是在参数估计还是支撑的正向选择率上均要优于 EW 方法。其次, 在高 SNR 情况下, 同时考虑了三种不同的情况, 即 $n > p$, $n < p$ 和 $n \ll p$, 将 EEW 求解参数估计的方法与其他四种方法进行对比, 突出了本文 EEW 方法进行模型选择的有效性。

总结与展望

本文从贝叶斯框架的角度,通过讨论模型空间大小的变化,构造了 l_0 收缩指数先验,利用 l_0 收缩指数先验对线性模型的模型选择问题进行研究。在 SNR 较大时, l_0 收缩指数先验能够尽可能地收缩系数向量的支撑大小,从而实现稀疏学习与模型选择的目标,且在 SNR 较小时,防止收缩过度,产生过稀疏。本文的推广指数权重法,在数据矩阵满足一定条件时,能够得到正确的稀疏的系数向量支撑估计,从而选择出正确的模型。在理论性质方面,本文建立了 l_0 收缩指数先验下的贝叶斯估计在模型的预测、支撑集学习、参数估计方面具有良好的理论性质。首先,建立了所构造稀疏先验分布对应贝叶斯估计的 Oracle 不等式。其次,证明了在满足可识别性条件下,只要系数向量非零元素足够大,使得后验概率最大化的支撑集依概率收敛到真实系数向量的支撑集。然后,在满足可识别性条件下,建立了系数向量的贝叶斯估计在三种不同误差度量下的误差界,表明在推广指数权重法下得到的参数估计是相合的,即当样本容量趋于无穷时,所求得的贝叶斯估计依概率收敛到真实系数向量。最后,本文通过数值模拟证实了,在偏低 SNR 情况下,本文方法能够得到更低估计误差的系数估计和更高的非零系数支撑恢复率,且在高 SNR 情况下,与其他变量选择方法相比,本文方法同样表现优异。同时,本文利用 EEW-SBR 算法来解决最大后验估计的对应的等价优化问题。通过数值模拟,在 SNR 偏低时,将本文所提方法与 EW 方法进行对比,突出了本文方法的适用性与优越性,同时在高 SNR 时,对 n 和 p 的三种不同情形下,与其他模型选择方法进行比较,表明了本文方法表现优异。

对本文内容可以进一步研究的课题有:

- (1) 本文专注点在于系数向量的支撑恢复与系数估计,在 SNR 较小时如何控制预测误差也是值得研究的点。
- (2) 本文研究的是线性模型,进一步考虑是否可以将本文的方法面向非线性模型上进行推广。
- (3) 本文研究的是二维表格数据,而张量形式的数据逐渐走进人们的生活。可考虑能否将本文方法推广到张量数据下的模型选择问题。

参考文献

- [1] FANG X, CAI Y, CAI Z, et al. Sparse feature learning of hyperspectral imagery via multiobjective-based extreme learning machine[J]. *Sensors*, 2020, 20(5):1262.
- [2] HAN S, PALLOTTA L, HUANG X, et al. A sparse learning approach to the design of radar tunable architectures with enhanced selectivity properties[J]. *IEEE Transactions on Aerospace and Electronic Systems*, 2020, 56(5):3840-3853.
- [3] LI N, ZHANG H H. Sparse learning with non-convex penalty in multi-classification[J]. *Journal of Data Science*, 2021(3):56-74.
- [4] WANG Z, LI F, SHI G, et al. Network pruning using sparse learning and genetic algorithm ScienceDirect[J]. *Neurocomputing*, 2020, 404:247-256.
- [5] BUHLMANN P. Statistics for high-dimensional data: methods, theory and applications[M]. Springer Science and Business Media, 2011.
- [6] HASTIE T, TIBSHIRANI R, WAINWRIGHT M. Statistical learning with sparsity: the lasso and generalizations[J]. *International statistical review*, 2015, 84(1): 156-157.
- [7] TURLACH B A, VENABLES W N, WRIGHTS S J. Simultaneous variable selection[J]. *Technometrics*, 2005, 47(3): 349-363.
- [8] BERTISIMAS D, WEISMANTEL R. Optimization over integers[M]. Belmont: Dynamic Ideas, 2005.
- [9] NEMHAUSER G L, SIGISMONDI G C. Constraint classification for mixed integer programming formulations[M]. Department of Mathematics and Computing Science University of Technology, 1991.
- [10] BAI Y, LIANG R, YANG Z. Splitting augmented lagrangian method for optimization problems with a cardinality constraint and semicontinuous variables. *Optimization Methods and Software*, 2016, 31(5): 1089-1109.
- [11] GANGI L D, LAPUCCI M, SCHOEN F, et al. An efficient optimization approach for best subset selection in linear regression, with application to model selection and fitting in autoregressive time-series[J]. *Computational Optimization and Applications*, 2019, 74(3): 919-948.
- [12] MALLAT S G, ZHANG Z, et al. Matching pursuits with time-frequency dictionaries[J]. *IEEE Transactions on Signal Processing*, 1993, 41(12): 3397-3415.
- [13] TROPP J A, GILBERT A C. Signal recovery from random measurements via orthogonal matching pursuit[J]. *IEEE Transactions on Information Theory*, 2007, 53(12): 4655-4666.
- [14] LESICKA A, KAWALEC A. An application of the orthogonal matching pursuit algorithm in space-time adaptive processing[J]. *Sensors (Basel, Switzerland)*, 2020, 20(12): 3468.
- [15] RUI Q, DIWEI Y, YUJIE Z, et al. On recovery of block sparse signals via block generalized orthogonal matching pursuit[J]. *Signal Processing*, 2018, 153: 34-46.

- [16] NEEDELL D, VERSHYNIN R. Signal recovery from incomplete and inaccurate measurements via regularized orthogonal matching pursuit[J]. IEEE Journal of Selected Topics in Signal Processing, 2007, 4(2): 310-316.
- [17] DONOHO D L, TSAIG Y, DRORI I, et al. Sparse solution of underdetermined systems of linear equations by stagewise orthogonal matching pursuit[J]. IEEE Transactions on Information Theory, 2012, 58(2): 1094-1121.
- [18] DAI W, MILENKOVIC O. Subspace pursuit for compressive sensing signal reconstruction[J]. IEEE Transactions on Information Theory, 2009, 55(5): 2230-2249.
- [19] NEEDELL D, TROOP J A. CoSaMP: Iterative signal recovery from incomplete and inaccurate samples[J]. Applied and Computational Harmonic Analysis, 2008, 26(3): 301-321.
- [20] COHEN A, DAHMEN W, DEVORE R. Orthogonal matching pursuit under the restricted isometry property[J]. Constructive Approximation, 2017, 45(1): 113-127.
- [21] FENG W. Filtering-based regularized sparsity variable step-size matching pursuit and its applications in vehicle health monitoring[J]. Applied Sciences, 2021, 11(11): 4816.
- [22] ZHANG J, WANG X, JU M, et al. An improved sparsity adaptive matching pursuit algorithm and its application in shock wave testing[J]. Mathematical Problems in Engineering, 2021, 2021(8): 1-10.
- [23] MILLER A. Subset selection in regression[M]. Chapman and hall/CRC, 2002.
- [24] ZHANG C H, ZHANG T. A general theory of concave regularization for high-dimensional sparse estimation problems[J]. Statistical Science, 2012, 27(4): 576-593.
- [25] RASKUTTI G, WAINWRIGHT M, Yu B. Minimax rates of estimation for high-dimensional linear regression over L -balls[J]. IEEE Transactions on Information Theory, 2011, 57(10): 6976-6994.
- [26] GREENSHTEIN E. Best subset selection, persistence in high-dimensional statistical learning and optimization under L_1 constraint[J]. The Annals of Statistics, 2006, 34(5): 2367-2386.
- [27] TIBSHIRANI R. Regression shrinkage and selection via the lasso[J]. Journal of the Royal Statistical Society(Series B), 1996, 58: 267-288.
- [28] FAN J, LI R. Variable selection via nonconcave penalized likelihood and its oracle properties[J]. Journal of the American Statistical, 2001, 96: 1348-1360.
- [29] ZOU H. The adaptive lasso and its oracle properties[J]. Journal of the Royal Statistical Society(Series B), 2007, 69: 143-161.
- [30] MEINSHAUSEN N. Lasso with relaxation[J]. Computational Statistics Data Analysis, 2007, 52: 374-393.
- [31] TIBSHIRANI R, SAUNDERS M, ROSSET S, et al. Sparsity and smoothness via the fused lasso[J]. Journal of the Royal Statistical Society(Series B), 2005, 67: 91-108.
- [32] YUAN M, LIN Y. Model selection and estimation in regression with grouped variables[J]. Journal of the Royal Statistical Society(Series B), 2006, 68: 49-67.

- [33] ZOU H, HASTIE T. Redularization and variable selection via the elastic net[J]. Journal of the Royal Statistical Society(Series B), 2005, 67: 301-320.
- [34] MAZUMDER R, RADCHENKO P, DEDIEU A. Subset selection with shrinkage: Sparse linear modeling when the SNR is low[J]. Arxiv preprint arXiv:1708.03288, 2017.
- [35] MILLER R, MILLER A, MILLER J, et al. Subset selection in regression[M]. Chapman and hall/CRC, 2002.
- [36] HAIDER H, SHAH J A, KADIR K, et al. Recovery of undersampled sparse signals using combined smoothed $l_0 - l_1$ norm[C]. International Conference on Signal and Image Processing Applications (ICSIPA). IEEE, 2019.
- [37] LOUNICI K. Generalized mirror averaging and d-convex aggregation[J]. Mathematical Methods of Statistics, 2007, 16(3): 246-259.
- [38] JUDITSKY A, RIGOLLET P, TSYBAKOV A B. Learning by mirror averaging[J]. The annals of statistics, 2008, 36(5): 2183-2206.
- [39] GIRAUD C. Mixing least-squares estimators when the variance is unknown[J]. Bernoulli: Official Journal of the Bernoulli Society for Mathematical Statistics and Probability, 2008, 14(4): 1089-1107.
- [40] SALMON J, DALALYAN A S. Optimal aggregation of affine estimators[J]. Journal of Machine Learning Research, 2011, 19: 635-660.
- [41] CATONI O. Statistical learning theory and stochastic optimization[M]. Lecture Notes in Mathematics Springer-verlag, 2004.
- [42] ALQUIER P, LOUNICI K. Pac-bayesian bounds for sparse regression estimation with exponential weights[J]. Working thesiss, 2010, 5(5): 127-145.
- [43] ARIAS C A, LOUNICI K. Estimation and variable selection with exponential weights[J]. Electronic Journal of Statistics, 2014, 8(1): 328-354.
- [44] MITCHELL T J, BEAUCHAMP J J. Bayesian variable selection in linear regression[J]. Journal of the American Statistical Association, 1988, 83(401), 1023-1032.
- [45] GEORGE E I, MCCULLOCH R E. Variable selection via Gibbs sampling[J]. Journal of the American Statistical Association, 1993, 88(423): 881-889.