
Inference from Quantized Data via Normal Variance-Mean Mixtures

Chenyu Gao¹ Zhexian Yang¹ Ziping Zhao¹

Abstract

Inference from quantized data has received significant attention in recent years due to its broad applications in machine learning and signal processing. Existing likelihood-based approaches are often restricted to Gaussian assumptions or low-bit quantization settings, limiting modeling flexibility and robustness under complex data distributions. In this work, we study inference from quantized observations under the general normal variance-mean mixture (NVMM) framework, which encompasses distributions including Gaussian, t , generalized hyperbolic skew- t , and generalized hyperbolic distributions. Optimization under the NVMM framework is challenging because the underlying likelihood function involves multidimensional integrals that are difficult to evaluate due to multidimensional quantization and latent mixture variables. To address this difficulty, we propose an expectation conditional maximization (ECM) algorithm with latent-variable augmentations for both quantization and mixture modeling. By leveraging the conditional Gaussian structure of the NVMM family, the proposed method admits closed-form updates for all model parameters at each iteration, leading to an efficient and tractable optimization procedure. We further establish global linear convergence guarantees for the proposed ECM algorithm. Beyond basic parameter estimation, the proposed framework naturally extends to several structured learning and recovery tasks under the NVMM framework, including quantized regression, matrix completion, compressive sensing, and covariance estimation. Numerical experiments demonstrate the effectiveness and robustness of the proposed framework across a variety of quantized inference problems.

¹School of Information Science and Technology, ShanghaiTech University, Shanghai, China. Correspondence to: Ziping Zhao <zipingzhao@shanghaitech.edu.cn>.

1. Introduction

Quantization, which represents data using a finite number of bits, provides an efficient mechanism for data acquisition, storage, transmission, and computation in resource-limited settings (Roberts, 1962; Gray & Neuhoﬀ, 2002; Widrow & Kollár, 2008; Kipnis & Reeves, 2021). By reducing the precision of continuous-valued observations, quantization enables scalable processing of large-scale and high-dimensional data while alleviating communication, memory, and hardware burdens. As a result, quantized observations arise naturally in a broad range of machine learning and signal processing applications. Representative examples include matrix recovery in recommendation systems (Davenport et al., 2014; Cai & Zhou, 2013; Bhaskar, 2016; Bottegal & Suykens, 2017; Gao et al., 2018), sparse signal recovery in compressed sensing (Boufounos & Baraniuk, 2008; Zymnis et al., 2009; Plan & Vershynin, 2013; Ai et al., 2014), and statistical learning tasks such as quantized regression and subspace estimation (Mayne, 1967; Györfi & Wegkamp, 2008; Chen et al., 2023; Chi & Fu, 2017; Dirksen et al., 2025). Quantization also plays a central role in wireless communications and sensing applications, including channel estimation, array detection, radar signal processing, spectrum sensing, and networked sensing (Choi et al., 2016; Stöckle et al., 2016; Plabst et al., 2018; Ren & Li, 2017; Ameri et al., 2019; Jin et al., 2020; Bar-Shalom & Weiss, 2002; Lu et al., 2024; Yang et al., 2025; Chi & Fu, 2017). In many of these applications, the objective is not merely to recover the quantized observations themselves, but rather to infer latent model structures and statistical quantities underlying the data generation process.

These applications can often be formulated through probabilistic modeling of latent quantities and structures from quantized observations. Depending on the application, the quantities of interest may correspond to low-order statistical quantities such as the mean (Papadopoulos et al., 2001; Dabeer & Masry, 2008) for distributed detection (Ribeiro & Giannakis, 2006a;b; Fang & Li, 2008), and the covariance matrix (Van Vleck & Middleton, 1966; Gray & Stockham, 1993; Liu & Lin, 2021; Dirksen et al., 2022; Xiao et al., 2023) for direction-of-arrival estimation (Eamazi et al., 2022; 2023; Bar-Shalom & Weiss, 2002), wireless power estimation (Mo et al., 2017), and spectrum sensing

(Yang et al., 2025). More generally, latent structures may also arise through low-rank, sparse, or low-dimensional representations, leading to structured inference problems such as matrix completion, compressive sensing, and subspace recovery.

However, quantization inevitably discards part of the information contained in the original observations, making inference from quantized data both challenging and practically important. In this paper, we study probabilistic inference of latent quantities and structures from quantized observations. Let $\mathbf{x} \in \mathbb{R}^d$ denote a random signal or latent variable of interest, and let $\mathbf{y}$ denote its quantized observation generated through

$$\mathbf{y} = \mathcal{Q}(\mathbf{x}), \quad (1)$$

where $\mathcal{Q} : \mathbb{R}^d \rightarrow \mathcal{K}^d$ denotes a quantization function, or quantizer, and $\mathcal{K} = \{k_1, \dots, k_e\}$ is a finite quantization alphabet containing e levels. Specifically, the i -th entry of $\mathcal{Q}(\mathbf{x})$ is defined as

$$[\mathcal{Q}(\mathbf{x})]_i = k_l, \quad \text{if } x_i \in [\tau_{l-1}, \tau_l), \quad (2)$$

where the quantization intervals satisfy

$$\bigcup_{l=1}^e [\tau_{l-1}, \tau_l) = \mathbb{R},$$

with $\tau_0 = -\infty$ and $\tau_e = \infty$. The number of bits required to represent each quantized value is given by $\log_2 e$. When $\log_2 e$ is small, the quantizer is referred to as coarse quantization. A particularly important special case is one-bit quantization, corresponding to $e = 2$. Under the additional conditions $k_1 = -1$, $k_2 = 1$, and $\tau_1 = 0$, the quantizer $\mathcal{Q}$ reduces to the element-wise signum function.

In this paper, we study maximum likelihood inference from quantized observations $\mathbf{y}$ under the normal variance-mean mixture (NVMM) framework. We assume that the latent signal $\mathbf{x}$ follows a general NVMM distribution (Barndorff-Nielsen et al., 1982; Barndorff-Nielsen, 1997; Polson & Scott, 2013), a flexible probabilistic family capable of modeling heavy tails, skewness, and heterogeneous covariance structures through its location, skewness, and scatter parameters. The NVMM framework encompasses several widely used distributions, including Gaussian, t , generalized hyperbolic skew- t (GHST), and generalized hyperbolic (GH) distributions. Unlike existing approaches that are primarily restricted to one-bit quantization or Gaussian assumptions, the proposed framework accommodates arbitrary multi-bit quantization functions and general distributions within the NVMM family.

To solve the resulting maximum likelihood problem, we develop an expectation conditional maximization (ECM)

algorithm that jointly handles multidimensional quantization and latent mixture variables through suitable latent-variable augmentations. The proposed method alternates between an expectation step (E-step), where a surrogate objective function is constructed via Jensen's inequality, and a conditional maximization step (CM-step), where the surrogate is optimized with respect to different parameter blocks while keeping the remaining parameters fixed. By leveraging the conditional Gaussian structure of the NVMM family, the proposed ECM algorithm admits closed-form updates for the location, skewness, and scatter parameters. We further establish global linear convergence guarantees for the proposed algorithm. Beyond parameter estimation under the NVMM framework, the proposed probabilistic inference framework naturally extends to several structured learning and recovery problems, including quantized regression, matrix completion, compressive sensing, and covariance estimation. Extensive numerical experiments demonstrate the effectiveness and robustness of the proposed framework across a variety of quantized inference tasks.

2. Problem Formulation

A random vector $\mathbf{x} \in \mathbb{R}^d$ following an NVMM model can be represented as

$$\mathbf{x} = \boldsymbol{\mu} + z\boldsymbol{\xi} + (z\boldsymbol{\Sigma})^{\frac{1}{2}}\boldsymbol{\epsilon}, \quad (3)$$

where $\boldsymbol{\mu}$ is the location parameter, $\boldsymbol{\xi}$ is the skewness parameter, $\boldsymbol{\Sigma}$ is the scatter parameter, z is a nonnegative random variable with density function $p(z)$, and $\boldsymbol{\epsilon}$ is a standard normal random vector with zero mean and identity covariance matrix, independent of z . Equivalently, conditional on z , the random vector $\mathbf{x}$ follows a Gaussian distribution with mean $\boldsymbol{\mu} + z\boldsymbol{\xi}$ and covariance matrix $z\boldsymbol{\Sigma}$. The conditional density function of $\mathbf{x}$ given z is therefore

$$p(\mathbf{x} | z; \boldsymbol{\theta}) = \frac{1}{(2\pi)^{\frac{d}{2}} \det(z\boldsymbol{\Sigma})^{\frac{1}{2}}} \times \exp\left(-\frac{1}{2z} \|\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi}\|_{\boldsymbol{\Sigma}^{-1}}^2\right), \quad (4)$$

where

$$\boldsymbol{\theta} = (\boldsymbol{\mu}, \boldsymbol{\xi}, \boldsymbol{\Sigma})$$

collects the model parameters¹, and $\|\mathbf{x}\|_{\mathbf{A}}^2 = \mathbf{x}^\top \mathbf{A} \mathbf{x}$. Following the quantization model in (1), the density function

¹In general, the distribution of z in an NVMM model may also depend on additional parameters, which are assumed to be known and are not included in the estimation procedure throughout this paper for simplicity.

of the quantized observation $\mathbf{y}$ is given by

$$\begin{aligned} p(\mathbf{y}; \boldsymbol{\theta}) &= \int_{\mathcal{Q}^{-1}(\mathbf{y})} p(\mathbf{x}; \boldsymbol{\theta}) d\mathbf{x} \\ &= \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^\infty p(\mathbf{x} | z; \boldsymbol{\theta}) p(z) dz d\mathbf{x}, \end{aligned} \quad (5)$$

where $\mathcal{Q}^{-1}(\mathbf{y})$ denotes the preimage of the quantized observation $\mathbf{y}$ under the quantizer $\mathcal{Q}$. In particular, $\mathcal{Q}^{-1}(\mathbf{y})$ forms a hyper-rectangle in $\mathbb{R}^d$, whose projection onto the i -th coordinate corresponds to the quantization interval associated with $[\mathbf{y}]_i$.

Given n independent and identically distributed observations $\mathbf{y}_1, \dots, \mathbf{y}_n$, the maximum likelihood (ML) estimation problem is formulated as

$$\max_{\boldsymbol{\theta}} L(\boldsymbol{\theta}) = \sum_{t=1}^n \log p(\mathbf{y}_t; \boldsymbol{\theta}). \quad (6)$$

Even evaluating the likelihood function in (5) is generally challenging, since it involves multidimensional integration induced jointly by the quantization region $\mathcal{Q}^{-1}(\mathbf{y})$ and the latent mixing variable z . In the special case where $d = 1$, z is deterministic, and one-bit quantization is applied, the resulting model reduces to a univariate Gaussian setting for which closed-form estimators are available (Ribeiro & Giannakis, 2006b). In contrast, for multivariate quantization and general NVMM models, the likelihood function typically does not admit a closed-form expression, making the associated optimization problem substantially more difficult.

3. Related Work

Building upon the ML optimization framework in (6), a variety of downstream tasks can be addressed. Among various research directions, two fundamental lines of inquiry concern the estimation of the mean and covariance of $\mathbf{x}$, which we adopt as our starting point. Furthermore, ML-based parameter estimation can be extended to the estimation of structured probabilistic models, which constitute standard practice in machine learning. In the following, we demonstrate applications of the proposed model to problems such as quantized regression, quantized matrix completion, and quantized compressive sensing.

3.1. Quantized Mean Estimation

In (Papadopoulos et al., 2001), in order to estimate signals from wireless sensor networks, an optimization problem is formulated for estimating the mean under a one-dimensional Gaussian distribution assumption, based on multi-bit quantized data. To address this optimization problem, the authors propose an ECM-based algorithm. Subsequently, (Ribeiro & Giannakis, 2006b) focuses on the

specific scenario of one-bit quantization under a Gaussian distribution and derives a closed-form expression for mean estimation in this case. Further, (Fang & Li, 2008) extends this line of research by introducing an adaptive threshold into the one-bit quantization function, thereby improving mean estimation accuracy. In addition, (Dabeer & Masry, 2008) generalizes the problem to the multidimensional setting. (Ribeiro & Giannakis, 2006a) extends the model distribution to the generalized Gaussian distribution under the one-dimensional assumption.

3.2. Quantized Linear Regression

Quantized linear regression characterizes the mapping of a feature vector $\mathbf{x}$ to a discrete output $\mathbf{y}$, modeled as

$$\mathbf{y} = \mathcal{Q}(\boldsymbol{\beta}^\top \mathbf{x} + \epsilon).$$

In this formulation, the linear component $\boldsymbol{\beta}\mathbf{x}$ is perturbed by an error term ϵ , the primary objective is to estimate the parameter $\boldsymbol{\beta}$, and the error term ϵ is treated as a zero-mean random variable. When ϵ is assumed as a Gaussian variable, an EM-based algorithm has been proposed (Finesso et al., 1999). However, as the Gaussian distribution assumption exhibits poor performance when data contain outliers, subsequent studies (Chen et al., 2023; 2024) center on robust estimation methods for heavy-tailed data. Furthermore, robust estimation in classical linear regression accounts for both asymmetric and heavy-tailed noise (Fan et al., 2021). For instance, the noise term ϵ can be modeled using a GH distribution (Kim & Browne, 2024), which captures both of these characteristics simultaneously. In this paper, our proposed algorithm facilitates parameter estimation using quantized measurements under the assumption of a normal variance-mean mixture. In the normal variance-mean mixture, the t distribution is heavy-tailed, and both the GHST and GH distributions exhibit heavy tails and asymmetry, this modeling framework enables robust estimation. In this case, the model becomes

$$\mathbf{y} = \mathcal{Q}\left(\boldsymbol{\beta}^\top \mathbf{x} + (z - \mathbb{E}[z])\xi + (z\sigma)^{\frac{1}{2}}\epsilon\right). \quad (7)$$

3.3. Quantized Probabilistic Matrix Completion

The goal of the low-rank matrix completion problem (Candes & Recht, 2012) is to recover an unknown low-rank matrix $\mathbf{M} \in \mathbb{R}^{d_1 \times d_2}$ from an observed, yet incomplete, matrix. Let $\mathbf{X}$ denote a matrix whose entries are drawn from a normal variance-mean mixture distribution. We use μ_{ij} and x_{ij} to represent the (i, j) -th entries of $\mathbf{M}$ and $\mathbf{X}$, respectively. Based on the normal variance-mean model, the relationship between μ_{ij} and x_{ij} is given by

$$x_{ij} = \mu_{ij} + z^{\frac{1}{2}}\sigma\epsilon, \quad (8)$$

where μ_{ij} serves as the mean of x_{ij} . In the quantization scenario (Davenport et al., 2014), however, the matrix $\mathbf{X}$

is not directly accessible; instead, we observe its quantized counterpart $\mathbf{Y} = \mathcal{Q}(\mathbf{X})$. Moreover, the entire matrix $\mathbf{Y}$ is not available for observation. Let $\mathcal{O}$ denote the index set of the observed entries; specifically, if $(i, j) \in \mathcal{O}$, then Y_{ij} is observed, otherwise Y_{ij} is missing. Our goal is to recover $\mathbf{M}$ from incomplete $\mathbf{Y}$. The study of matrix completion was popularized following the Netflix Million Dollar Challenge, which posed the task: accurately predicting the values of those entries with a user–item matrix in which entries represent item ratings.

To recover $\mathbf{M}$ from quantized and corrupted observations, a seminal study introduced an ML framework for one-bit low-rank matrix completion (Davenport et al., 2014). Building on this work, random dithering was incorporated into the quantization function to improve recovery performance (Eamaz et al., 2024). Both approaches employed the nuclear norm to relax the low-rank constraint and used projected gradient descent for optimization. An alternative strategy reformulated the low-rank constraint via matrix factorization, followed by projected gradient descent (Bhaskar & Javanmard, 2015). Similarly, factorization combined with a majorization–minimization method was used to derive a surrogate objective, which was solved via the Gauss–Newton method (Liu et al., 2025). The framework was further extended from one-bit to multi-bit quantization using low-rank factorization together with projected gradient descent (Bhaskar, 2016). The former works (Davenport et al., 2014; Bhaskar & Javanmard, 2015; Bhaskar, 2016) are based on the Gaussian distribution assumption. However, observed data often exhibit heavy-tailed characteristics, which are induced by the presence of outliers (Chen et al., 2024). To analyze such data, robust modeling techniques are frequently employed (Shen et al., 2019; Chen et al., 2024). The normal variance-mean mixture class encompasses numerous classical heavy-tailed distributions that are extensively utilized in robust modeling. Consequently, this study employs the t distribution and the symmetric GH distribution to perform robust estimation for the quantized matrix completion problem.

3.4. Quantized Compressive Sensing

One-bit compressive sensing ((Boufounos & Baraniuk, 2008)) aims to estimate a sparse signal lying in a known measurement subspace based on observed quantized data (Jacques et al., 2013; Chen et al., 2024). Existing methods for this problem generally fall into three categories. The first two primarily focus on 1-bit quantization without incorporating additive noise into the original measurements (Li et al., 2018). The first category, termed regularizer-class algorithms (Laska et al., 2011), introduces additional regularization terms to the classical compressive sensing recovery problem to enforce consistency between the sparse signal and the measurements. The second category, known

as penalty-class algorithms (Yan et al., 2012), models sign flips between quantized and recovered measurements as penalty terms. The third category assumes the presence of additive noise in the original measurements (Zymnis et al., 2009; Knudson et al., 2016). Our proposed model is developed as a further extension of the third approach. Similar to the quantized matrix completion problem, robust estimation represents a significant modeling framework in quantized compressed sensing (Plan & Vershynin, 2013; Dirksen & Mendelson, 2021; Jung et al., 2021). While previous research primarily focuses on characterizing the statistical errors of algorithms based on the properties of heavy-tailed data, this study utilizes the t and symmetric GH distributions to derive efficient algorithms for robust estimation. We employ a normal variance-mean mixture model for noise modeling. Let $\Phi \in \mathbb{R}^{d_1 \times d_2}$ be the known measurement matrix, and $\vartheta \in \mathbb{R}^{d_2}$ be the sparse signal to be estimated. Based on the normal variance-mean model, the quantized compressed sensing model is given by

$$\mathbf{y} = \mathcal{Q}(\mathbf{x}), \quad \mathbf{x} = \Phi\vartheta + z^{\frac{1}{2}}\sigma\epsilon, \quad (9)$$

where $\Phi\vartheta$ serves as the location parameter of $\mathbf{x}$.

3.5. Quantized Covariance/Correlation Estimation

In (Van Vleck & Middleton, 1966), the authors investigate the estimation of a correlation matrix from one-bit quantized zero-mean Gaussian measurements, based on the arcsine law (Lévy, 1940). However, in such one-bit zero-mean settings, the individual variance of each dimension cannot be determined, which prevents recovery of the full covariance matrix. To estimate the variance, the “dithering technique” is introduced, which is to set a threshold in the signum function (Liu & Lin, 2021). Based on the dithering technique, the variance can be obtained in closed form (Fang & Li, 2008), but the correlation matrix should be solved analytically. Consequently, subsequent studies (Dirksen et al., 2022; Eamaz et al., 2023; Xiao et al., 2023; Liu & Chou, 2025) have proposed various strategies for correlation estimation. The estimation approaches for correlation can be classified into two categories. The first category relies on the correlation coefficient function of the one-bit Gaussian distribution, with analyses usually restricted to pairwise interactions between dimensions in the multivariate setting. Since this function is generally computationally intractable, different approximate functions have been proposed to compute the correlation coefficient (Eamaz et al., 2023; Xiao et al., 2023; Liu & Chou, 2025). In (Eamaz et al., 2022; 2023), the authors introduce a known, time-varying threshold for the signum function and propose a modified arcsine law to express correlation coefficients as integrals, which are then approximated using Gauss–Legendre quadrature. The method in (Xiao et al., 2023) applies ML estimation to compute cor-

relation coefficients, which is equivalent to iteratively evaluating their approximate functions. The work (Liu & Chou, 2025) represents the correlation coefficient function as an infinite series via the one-bit Hermite law (Liu & Lin, 2021) and then approximates it using harmonic approximation. The second category (Dirksen et al., 2022; Chen et al., 2024) directly employs the sample covariance matrix computed using multiple thresholds. These methods are not restricted to one-bit quantization and can be generalized to multi-bit cases, though typically at the expense of reduced estimation accuracy. The studied model can also be applied to quantized structured covariance modeling, which focuses on recovering high-dimensional covariance matrices from quantized measurements by exploiting inherent geometric structures—such as Toeplitz (Liu & Lin, 2021; Eamam et al., 2022), sparse, or low-rank (Saha et al., 2023) properties (Maly et al., 2022).

3.6. Quantized Graphical Modeling

Quantized graphical modeling studies the recovery and inference of conditional dependence graphs when the observed variables are discrete or arise from quantization of latent continuous processes. Model construction typically proceeds either by directly specifying discrete Markov random fields (Fang & Li, 2010), or by introducing latent Gaussian variables combined with quantization (Lai, 2016; Tavassolipour et al., 2018). Besides, several studies focus on the robust design of graphs (Nobre & Frossard, 2019; Saad et al., 2021).

4. The ECM Algorithm

The optimization problem in (6) is challenging due to the multiple integrals in the density function (5). Nevertheless, the ECM framework can be applied by leveraging the models in (1) and (3), where $\mathbf{x}$ and $\mathbf{z}$ can be naturally treated as latent variables. In the following, we introduce an ECM procedure for (6).

4.1. E-step

In the E-step, we derive a surrogate function for the log-likelihood function $L(\theta)$. Based on (1), the observed signal $\mathbf{y}$ is conditioned on the hidden variable $\mathbf{x}$. Hence, given $[\mathbf{y}_1 \dots, \mathbf{y}_n]^\top$ and the corresponding hidden variables $[\mathbf{x}_1 \dots, \mathbf{x}_n]^\top$, we have²

$$\begin{aligned} L(\theta) &= \sum_{t=1}^n \log \int_{\mathbb{R}} p(\mathbf{y}_t | \mathbf{x}_t; \theta) p(\mathbf{x}_t; \theta) d\mathbf{x}_t \\ &\geq \sum_{t=1}^n \mathbb{E}_{\mathbf{x}_t | \mathbf{y}_t; \underline{\theta}} \log p(\mathbf{y}_t, \mathbf{x}_t; \theta) + \text{const.}, \end{aligned} \quad (10)$$

²Throughout this paper, underlined variables denote those whose values are given as constants.

where the inequality is based on Jensen’s inequality and const. is some constant term independent of θ . Under the models specified in (1), the joint density function is given by

$$p(\mathbf{y}_t, \mathbf{x}_t; \theta) = p(\mathbf{x}_t; \theta) \delta(\mathcal{Q}(\mathbf{x}_t) - \mathbf{y}_t), \quad (11)$$

where $\delta(\cdot)$ denotes the multivariate Dirac delta function, defined as $\delta(\mathbf{x}_t) = \begin{cases} 0, & \text{if } \mathbf{x}_t \neq \mathbf{0} \\ +\infty, & \text{if } \mathbf{x}_t = \mathbf{0} \end{cases}$. Since the term $\delta(\mathcal{Q}(\mathbf{x}_t) - \mathbf{y}_t)$ is independent of θ , we further obtain

$$\begin{aligned} L(\theta) &\geq \sum_{t=1}^n \mathbb{E}_{\mathbf{x}_t | \mathbf{y}_t; \underline{\theta}} \log p(\mathbf{x}_t; \theta) + \text{const.} \\ &= \sum_{t=1}^n \mathbb{E}_{\mathbf{x}_t | \mathbf{y}_t; \underline{\theta}} \log \int_0^\infty p(\mathbf{x}_t, z_t; \theta) dz_t + \text{const.} \end{aligned}$$

Regarding z as a hidden variable, applying Jensen’s inequality leads to

$$L(\theta) \geq \sum_{t=1}^n \mathbb{E}_{\mathbf{x}_t | \mathbf{y}_t; \underline{\theta}} [\mathbb{E}_{z_t | \mathbf{x}_t; \underline{\theta}} \log p(\mathbf{x}_t, z_t; \theta)] + \text{const.} \quad (12)$$

Since $z | \mathbf{x}$ is independent of $\mathbf{x} | \mathbf{y}$, we have $p(\mathbf{x} | \mathbf{y}; \theta) p(z | \mathbf{x}; \theta) = p(\mathbf{x}, z | \mathbf{y}; \theta)$. Given the $p(\mathbf{x} | z; \theta)$ in (4), the surrogate function $S(\theta; \underline{\theta})$ in (12) can be further expressed as follows:

$$\begin{aligned} S(\theta; \underline{\theta}) &= \sum_{t=1}^n \left[\mathbb{E}_{z_t | \mathbf{y}_t; \underline{\theta}} [\log p(z_t)] - \frac{1}{2} \log \det \Sigma \right. \\ &\quad \left. - \frac{1}{2} \text{Tr} \left(\mathbf{U}_t - 2\mathbf{v}_t \boldsymbol{\mu}^\top - 2\mathbf{w}_t \boldsymbol{\xi}^\top + 2\boldsymbol{\mu} \boldsymbol{\xi}^\top \right. \right. \\ &\quad \left. \left. + \boldsymbol{\iota}_t \boldsymbol{\mu} \boldsymbol{\mu}^\top + \zeta_t \boldsymbol{\xi} \boldsymbol{\xi}^\top \right) \Sigma^{-1} \right] + \text{const.} \end{aligned} \quad (13)$$

where

$$\begin{aligned} \mathbf{U}_t &= \mathbb{E}_{\mathbf{x}_t, z_t | \mathbf{y}_t; \underline{\theta}} [z_t^{-1} \mathbf{x} \mathbf{x}^\top], \quad \mathbf{v}_t = \mathbb{E}_{\mathbf{x}_t, z_t | \mathbf{y}_t; \underline{\theta}} [z_t^{-1} \mathbf{x}_t], \\ \mathbf{w}_t &= \mathbb{E}_{\mathbf{x}_t | \mathbf{y}_t; \underline{\theta}} [\mathbf{x}_t], \quad \boldsymbol{\iota}_t = \mathbb{E}_{z_t | \mathbf{y}_t; \underline{\theta}} [z_t^{-1}], \quad \zeta_t = \mathbb{E}_{z_t | \mathbf{y}_t; \underline{\theta}} [z_t]. \end{aligned}$$

The details for computing these expectations are given in Appendix A. In the term $\mathbb{E}_{z_t | \mathbf{y}_t; \underline{\theta}} [\log p(z_t)]$, some scalar parameters (such as the shape parameter ν in t distribution) are contained. In the practical implementation (Galarza et al., 2021), they are typically treated as given or estimated through the one-dimensional search.

4.2. CM-step

Based on the surrogate function (13), in the CM-step, we solve the following optimization problem:

$$\max_{\theta} S(\theta; \underline{\theta}), \quad (14)$$

where parameters μ , ξ , and Σ can be solved with closed-form solutions:

$$\begin{aligned}\mu &= \frac{\sum_{t=1}^n (v_t - \xi)}{\sum_{t=1}^n \iota_t}, & \xi &= \frac{\sum_{t=1}^n (w_t - \mu)}{\sum_{t=1}^n \zeta_t}, \\ \Sigma &= \frac{1}{n} \sum_{t=1}^n \left(U_t - 2v_t \mu^\top - 2w_t \xi^\top \right. \\ &\quad \left. + 2\mu \xi^\top + \iota_t \mu \mu^\top + \zeta_t \xi \xi^\top \right). \end{aligned} \quad (15)$$

Remark 1. In the context of quantization model estimation, a prototypical scenario is the one-bit Gaussian case (i.e., $e = 2$ and z is a constant). However, this model suffers from an inherent identifiability issue: for each dimension $i = 1, \dots, d$, only the ratio $\frac{\mu_i}{\sigma_i^2}$ can be determined, rather than the individual parameters. Hence, the estimated values of μ_i and σ_i^2 differ from their desirable values by an arbitrary scaling factor. To resolve this, existing estimation methods in the one-bit Gaussian setting typically fix one parameter (either the location vector or the scatter matrix) to identify the other. Moreover, the same ambiguity persists under one-bit quantization whenever x follows any elliptical distribution.

5. Convergence Analysis

In this section, we analyze the convergence properties of the proposed ECM algorithm. Based on (15), all parameters admit unique closed-form updates at each step. Consequently, starting from any feasible point, our ECM algorithm is guaranteed to converge to a stationary point in the feasible set (McLachlan & Krishnan, 2008). We first establish the convexity of each parameter individually.

Theorem 2. The optimization problem (6) is strictly concave with respect to each of the parameters μ , ξ , and Σ .

Based on Theorem 2, the individual convergence rates of the three parameters are given in the following.

Proposition 3. Denote the parameters in the k -th iteration as $\mu^{(k)}$, $\xi^{(k)}$, and $\Sigma^{(k)}$, and the stationary points as μ^* , ξ^* , and Σ^* . One parameter is selected and iteratively updated via (15) until convergence, while the remaining two parameters are kept constant. For the convergent sequences of these three parameters, the following relationship holds for any k :

$$\begin{aligned}\|\mu^{(k+1)} - \mu^*\|_2 &\leq c_\mu \|\mu^{(k)} - \mu^*\|_2, \\ \|\xi^{(k+1)} - \xi^*\|_2 &\leq c_\xi \|\xi^{(k)} - \xi^*\|_2, \\ \|\Sigma^{(k+1)} - \Sigma^*\|_F &\leq c_\Sigma \|\Sigma^{(k)} - \Sigma^*\|_F,\end{aligned}$$

where

$$\begin{aligned}c_\mu &= \max_{\theta} \left\| \frac{\sum_{t=1}^n \Sigma^{-1} \text{Cov}_{x_t, z_t | y_t; \theta} [z_t^{-1} x_t]}{\sum_{t=1}^n \mathbb{E}_{z_t | y_t; \theta} [z_t^{-1}]} \right\|_2 \in (0, 1), \\ c_\xi &= \max_{\theta} \left\| \frac{\sum_{t=1}^n \Sigma^{-1} \text{Cov}_{x_t, z_t | y_t; \theta} [x_t - z_t \xi]}{\sum_{t=1}^n \mathbb{E}_{z_t | y_t; \theta} [z_t]} \right\|_2 \in (0, 1), \\ c_\Sigma &= \max_{\theta} \left\| \frac{1}{2n} \sum_{t=1}^n (\Sigma^{-1} \otimes \Sigma^{-1}) \right. \\ &\quad \left. \text{Cov}_{x_t, z_t | y_t; \theta} [z_t^{-1} \text{vec}(x_t x_t^\top)] \right\|_2 \in (0, 1).\end{aligned}$$

Based on the result in Proposition 3, it follows that quotient linear convergence rates are achieved for each parameter when updated independently. For the global convergence rate, we introduce a result.

Denote $\theta^* = [\mu^{*\top}, \xi^{*\top}, \text{vec}(\Sigma^*)^\top]^\top$ as a stationary point of the optimization problem in (6). According to (Meng & Rubin, 1994), the global convergence rate of an ECM algorithm is the maximum of the component-wise rates of convergence. Then, we establish that the proposed ECM algorithm converges globally at a linear rate, as detailed in the following result.

Theorem 4. Denote the sequence generated by the ECM algorithm as $\{\theta^{(k)}\}$. Given an initial point $\theta^{(0)}$, there exists a constant $c_0 > 0$ such that, for any $k \geq 0$, we have

$$\|\theta^{(k)} - \theta^*\|_2 \leq c^k \|\theta^{(0)} - \theta^*\|_2, \quad (16)$$

where $c = \max\{c_\mu, c_\xi, c_\Sigma\} \in (0, 1 - c_0)$.

6. Experiments

In the previous sections, an algorithm for the ML estimation problem was proposed to model quantized data using a normal variance-mean mixture model. In this section, we present the practical applications of the proposed algorithm to quantized models and provide a comparative analysis between the proposed method and existing algorithms in these scenarios.

6.1. Convergence Analysis

Figure 1 presents a comparative analysis of the convergence rates of the algorithm across two statistical distributions: Gaussian and GH. All proposed methods achieve linear convergence rates.

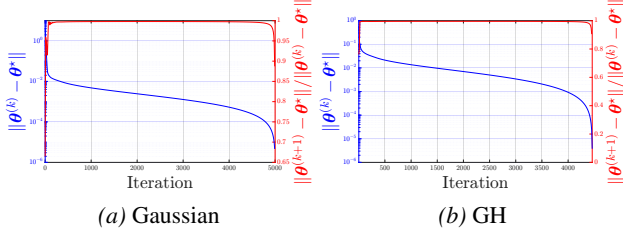


Figure 1. Convergence curve versus iterations

6.2. Quantized Linear Regression

Based on model (7), the corresponding optimization problem is formulated as follows:

$$\max_{\beta} \sum_{t=1}^n \log p(y_t | \beta). \quad (17)$$

Since the linear term $\beta^\top x_t$ in optimization problem (17) plays a role equivalent to the location parameter in the previously introduced ECM algorithm, the closed-form solution for the location parameter can be directly adapted to the linear regression framework. Specifically, the closed-form update for β at each ECM iteration is derived as:

$$\beta = \left(\sum_{t=1}^n \iota_t x_t x_t^\top \right)^{-1} \left(\sum_{t=1}^n (v_t - (1 + \iota_t E[z]) \xi) x_t \right).$$

The proposed method is evaluated against existing algorithms using synthetic data. The dimension of the parameters to be estimated is set to 10, and the number of samples is 1000. Three types of noise are considered: Gaussian noise, t noise (a heavy-tailed distribution with $\nu = 3$), and log-normal noise (a heavy-tailed and asymmetric distribution). The baseline algorithms for comparison include ordinary least squares (OLS) and the Gaussian EM algorithm; for OLS, each quantized observation is replaced by the representative value of its corresponding quantization interval. As illustrated in Figure 2, OLS and the Gaussian EM algorithm perform poorly under heavy-tailed noise, whereas the GHST and GH ECM algorithms exhibit significantly better performance in the presence of asymmetric noise.

6.3. Quantized Probabilistic Matrix Completion

Given (8) and $Y = Q(X)$, the ML estimation problem for M is given by

$$\begin{aligned} \max_M \quad & \sum_{(i,j) \in \mathcal{O}} \log p(y_{ij} | \mu_{ij}) \\ \text{s.t.} \quad & \text{rank}(M) \leq r. \end{aligned} \quad (18)$$

By applying a low-rank factorization to the matrix M , we express it as $M = AB^\top$, where $A \in \mathbb{R}^{d_1 \times r}$ and

$B \in \mathbb{R}^{d_2 \times r}$. Through the E-step in our proposed ECM algorithm, we have the surrogate function of (18) as

$$\sum_{(i,j) \in \mathcal{O}} \left(a_i b_j^\top - e_{ij} \right)^2, \quad (19)$$

where a_i and b_i are i -th row of the matrix A and B , respectively, and $e_{ij} = \frac{v_{ij} - \xi}{\iota_{ij}}$. The details of obtaining the surrogate (19) are given in the Appendix E.1. Then we can use the ECM algorithm to alternately update A and B .

For the experimental design, the MovieLens 1M dataset (Harper & Konstan, 2015) is employed, which contains 1,000,000 movie ratings provided by 6040 users for 3952 movies, with each rating ranging from 1 to 5. All ratings are randomly partitioned into training and test sets, accounting for 40% and 60% of the data, respectively. The training set is used to predict the completed rating matrix, and the entries overlapping between the completed rating matrix and the test set are then compared to evaluate prediction accuracy and root mean square error (RMSE), which are defined in Appendix F.

Existing methods encompass classic machine learning approaches for matrix completion, including the singular value decomposition (SVD (Sarwar et al., 2000)) model, the ℓ_2 -regularized matrix factorization (ℓ_2 -regularized (Patek, 2007)) model, and the nuclear norm minimization (nuclear norm (Cai et al., 2010)) model. Furthermore, several approaches are established upon random variable model assumptions. These include: directly modeling the Gaussian without quantization (Gaussian (Candes & Plan, 2010)); 1-bit quantization with a Gaussian distribution (1-bit Gaussian (Davenport et al., 2014)); and multi-bit quantization under a Gaussian assumption (multi-bit Gaussian (Bhaskar, 2016)). Our proposed method is regarded as an extension of the multi-bit Gaussian approach by extending the Gaussian assumption to a normal variance-mean mixture model, comprising t , GHST, and symmetric GH distributions.

The results are demonstrated in Table 1. As shown in the table, the proposed method consistently achieves superior performance in terms of both accuracy and RMSE. Among the evaluated models, the approach based on the symmetric GH distribution attains the best results, attributable to its flexibility and robustness in modeling skewness and heavy tails. Details regarding the experimental parameter settings are provided in the Appendix F.

6.4. Quantized Compressive Sensing

Given (9) and the ℓ_1 regularization model from (Zymnis et al., 2009), we have the optimization problem

$$\max_{\vartheta} \log p(y | \vartheta) + \eta \|\vartheta\|_1. \quad (20)$$

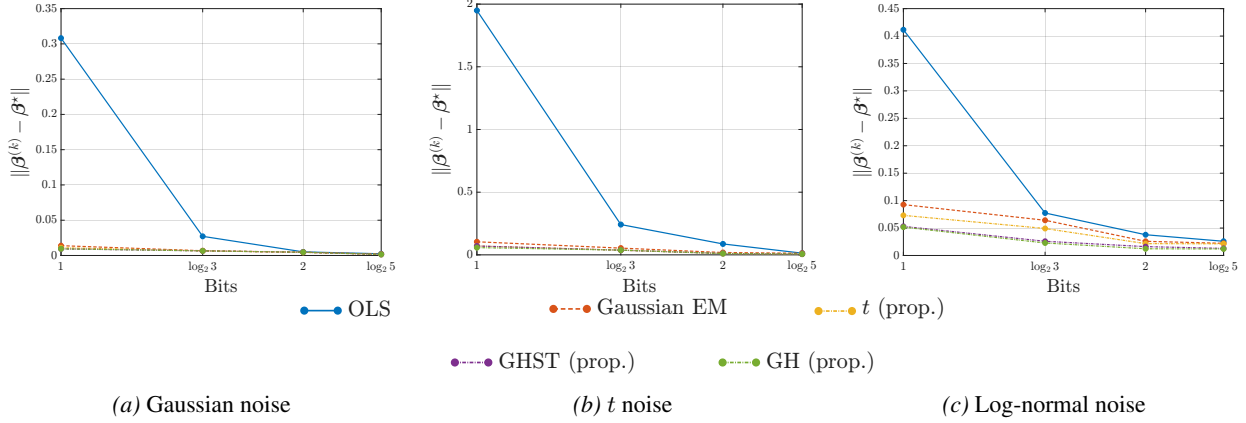


Figure 2. MSE comparisons for linear regression averaged over 10 Monte Carlo simulations.

Table 1. Accuracy and RMSE comparisons of matrix completion on MovieLens 1M

Method	Accuracy	RMSE	Time (s)
SVD	0.4261	0.9148	364.36
L2 regularization	0.4388	0.9355	165.08
Nuclear norm	0.3802	1.1662	867.39
Gaussian	0.4363	0.9486	25.26
1-bit Gaussian	0.4217	0.9659	110.44
Multi-bit Gaussian	0.4453	0.9151	90.84
Multi-bit t (prop.)	0.4464	0.9091	232.16
Multi-bit GH (prop.)	0.4503	0.8908	272.18

With the E-step in our proposed ECM algorithm, we have the surrogate function of (20) as

$$\frac{1}{2} \|\Phi \boldsymbol{\vartheta} - \mathbf{e}\|_2^2 + \eta \|\boldsymbol{\vartheta}\|_1,$$

where the details of deriving the surrogate function are given in the Appendix E.2. We can use the FISTA algorithm (Beck & Teboulle, 2009) to solve the surrogate. Performance comparisons are conducted using synthetic data, where the measurement dimension is denoted as $d_1 = 3000$ and the 30-sparse signal dimension as $d_2 = 1000$. To evaluate algorithmic robustness, an additive noise term exhibiting both skewed and heavy-tailed characteristics is introduced to the original measurements. The existing methods included in our comparison consist of restricted step shrinkage (RSS) (Laska et al., 2011), adaptive outlier pursuit (AOP) (Yan et al., 2012), and 1-bit Gaussian MLE (Zymnis et al., 2009). For each signal-to-noise ratio (SNR) level, experiments are repeated 10 times, and the average cosine similarity (Cos Sim) and computational time are reported. The results are summarized in Table 2, which demonstrates that the 1-bit GH ECM algorithm achieves the largest average cosine similarity. The definitions of SNR and cosine similarity are given in Appendix F.

6.5. Quantized Correlation Estimation

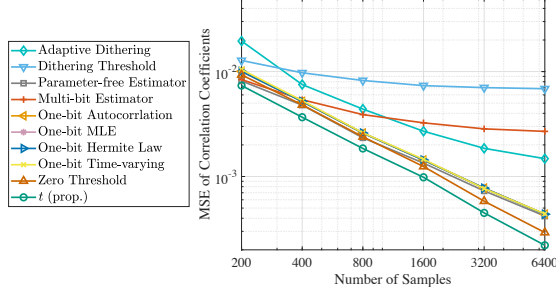
In this section, we address the problem of quantized covariance estimation. Existing approaches can be broadly categorized as follows:

- non-dithered method
 - arcsine law method (Van Vleck & Middleton, 1966) (Zero Threshold),
- non-zero threshold methods
 - one-bit autocorrelation estimation (Liu & Lin, 2021) (One-bit Autocorrelation);
 - one-bit Hermite law (Liu & Chou, 2025) (One-bit Hermite Law);
 - One-bit MLE (Xiao et al., 2023) (One-bit MLE);
- dithered methods
 - uniform dithering signal (Dirksen et al., 2022) (Dithering Threshold);
 - Gaussian dithering signal (Eamaz et al., 2022) (One-bit Time-varying);
 - adaptive dithering threshold (Dirksen & Maly, 2024) (Adaptive Dithering);
- dithered methods with multi-bit problem
 - multi-bit covariance estimator (Chen et al., 2024) (Multi-bit Estimator);
 - multi-bit parameter-free estimator (Chen & Ng, 2025) (Parameter-free Estimator).

We begin by comparing the estimation accuracy of the correlation matrix among our method and existing approaches, since some previous work cannot estimate the diagonal entries of one-bit data with zero mean (Van Vleck & Middleton, 1966). We generate samples from a t distribution with shape parameter $\nu = 3$ and estimate the correlation matrix using the t ECM and the aforementioned methods. Figure 3 illustrates the Frobenius norm error between the estimated and true correlation matrices across varying sample sizes, under the configuration $e = 4$, $\tau = 0.3$, and $d = 3$. The proposed ECM algorithm achieves a lower MSE compared

Table 2. Cosine similarity comparisons of 1-bit compressive sensing under different SNR levels.

Method	SNR = 0 dB		SNR = -5 dB		SNR = -10 dB	
	Cos Sim	Time (s)	Cos Sim	Time (s)	Cos Sim	Time (s)
RSS	0.7872	0.2639	0.6598	0.2651	0.3027	0.2546
AOP	0.7052	0.2689	0.4041	0.2951	0.2209	0.2738
1-bit Gaussian	0.8220	2.7704	0.5487	3.0875	0.2896	2.7202
1-bit t (prop.)	0.8396	2.8573	0.5678	3.1643	0.3031	3.3174
1-bit GH (prop.)	0.8906	2.7252	0.7705	3.2215	0.5386	3.5325


 Figure 3. Performance comparison under a t distribution.

to the existing methods. Notably, when comparing the use of different sample sizes for estimating the autocorrelation, the difference between One-bit Autocorrelation, One-bit MLE, One-bit Hermite Law, and One-bit Time-varying is quite small in the MSE.

7. Conclusion and Discussion

In this paper, we have proposed an ECM-based algorithm for parameter estimation in quantized models. The proposed method has exhibited broad applicability, handling problems from one-bit to multi-bit quantization and accommodating distributions ranging from Gaussian to the broader class of normal variance-mean mixtures. Experimental results have shown that our approach yields more accurate estimates than existing methods in certain cases (e.g., the one-bit Gaussian setting), while maintaining high accuracy across a wide range of extended scenarios. Furthermore, the proposed method has demonstrated strong potential in typical machine learning tasks, including quantized linear regression, quantized matrix completion, and quantized compressive sensing. An interesting future direction is to explore the statistical estimation properties of the proposed ECM algorithm for quantized maximum likelihood estimation problems.

Impact Statement

This paper presents work whose goal is to advance the field of machine learning. There are many potential societal consequences of our work, none of

which we feel must be specifically highlighted here.

References

- Ai, A., Lapanowski, A., Plan, Y., and Vershynin, R. One-bit compressed sensing with non-gaussian measurements. *Linear Algebra and its Applications*, 441:222–239, 2014. 1
- Ameri, A., Bose, A., Li, J., and Soltanalian, M. One-bit radar processing with time-varying sampling thresholds. *IEEE Transactions on Signal Processing*, 67(20):5297–5308, 2019. 1
- Bar-Shalom, O. and Weiss, A. J. DOA estimation using one-bit quantized measurements. *IEEE Transactions on Aerospace and Electronic Systems*, 38(3):868–884, 2002. 1
- Barndorff-Nielsen, O., Kent, J., and Sørensen, M. Normal variance-mean mixtures and z distributions. *International Statistical Review*, 50(2):145, 1982. 2
- Barndorff-Nielsen, O. E. Normal inverse gaussian distributions and stochastic volatility modelling. *Scandinavian Journal of Statistics*, 24(1):1–13, 1997. 2
- Beck, A. and Teboulle, M. A fast iterative shrinkage-thresholding algorithm for linear inverse problems. *SIAM journal on imaging sciences*, 2(1):183–202, 2009. 8
- Bhaskar, S. A. Probabilistic low-rank matrix completion from quantized measurements. *Journal of Machine Learning Research*, 17(60):1–34, 2016. 1, 4, 7
- Bhaskar, S. A. and Javanmard, A. 1-bit matrix completion under exact low-rank constraint. In *2015 49th Annual Conference on Information Sciences and Systems (CISS)*, pp. 1–6. IEEE, 2015. 4
- Bottegal, G. and Suykens, J. A. K. Probabilistic matrix factorization from quantized measurements. In *2017 International Joint Conference on Neural Networks (IJCNN)*, pp. 270–277, 2017. 1

- Boufounos, P. T. and Baraniuk, R. G. 1-bit compressive sensing. In *2008 42nd Annual Conference on Information Sciences and Systems*, pp. 16–21, 2008. 1, 4
- Brascamp, H. J. and Lieb, E. H. On extensions of the brunn-minkowski and prékopa-leindler theorems, including inequalities for log concave functions, and with an application to the diffusion equation. *Journal of functional analysis*, 22(4):366–389, 1976. 16
- Cai, J.-F., Candès, E. J., and Shen, Z. A singular value thresholding algorithm for matrix completion. *SIAM Journal on Optimization*, 20(4):1956–1982, 2010. 7
- Cai, T. and Zhou, W.-X. A max-norm constrained minimization approach to 1-bit matrix completion. *Journal of Machine Learning Research*, 14(1):3619–3647, 2013. 1
- Candès, E. and Plan, Y. Matrix completion with noise. *Proceedings of the IEEE*, 98(6):925–936, 2010. 7
- Candès, E. and Recht, B. Exact matrix completion via convex optimization. *Communications of the ACM*, 55(6):111–119, 2012. 3
- Casella, G. and Berger, R. L. *Statistical Inference*. Chapman and Hall/CRC, second edition, 2024. 19
- Chen, J. and Ng, M. K. A parameter-free two-bit covariance estimator with improved operator norm error rate. *Applied and Computational Harmonic Analysis*, 78:101774, 2025. 8
- Chen, J., Wang, Y., and Ng, M. K. Quantized low-rank multivariate regression with random dithering. *IEEE Transactions on Signal Processing*, 71:3913–3928, 2023. 1, 3
- Chen, J., Ng, M. K., and Wang, D. Quantizing heavy-tailed data in statistical estimation: (near) minimax rates, covariate quantization, and uniform recovery. *IEEE Transactions on Information Theory*, 70(3):2003–2038, 2024. 3, 4, 5, 8
- Chi, Y. and Fu, H. Subspace learning from bits. *IEEE Transactions on Signal Processing*, 65(17):4429–4442, 2017. 1
- Choi, J., Mo, J., and Heath, R. W. Near maximum-likelihood detector and channel estimator for uplink multiuser massive MIMO systems with one-bit ADCs. *IEEE Transactions on Communications*, 64(5):2005–2018, 2016. 1
- Dabeer, O. and Masry, E. Multivariate signal parameter estimation under dependent noise from 1-bit dithered quantized data. *IEEE Transactions on Information Theory*, 54(4):1637–1654, 2008. 1, 3
- Davenport, M. A., Plan, Y., Van Den Berg, E., and Wootters, M. 1-bit matrix completion. *Information and Inference*, 3(3):189–223, 2014. 1, 3, 4, 7
- Dirksen, S. and Maly, J. Tuning-Free One-Bit Covariance Estimation Using Data-Driven Dithering. *IEEE Transactions on Information Theory*, 70(7):5228–5247, July 2024. 8
- Dirksen, S. and Mendelson, S. Non-gaussian hyperplane tessellations and robust one-bit compressed sensing. *Journal of the European Mathematical Society*, 23(9):2913–2947, 2021. 4
- Dirksen, S., Maly, J., and Rauhut, H. Covariance estimation under one-bit quantization. *The Annals of Statistics*, 50(6):3538–3562, 2022. 1, 4, 5, 8
- Dirksen, S., Li, W., and Maly, J. Subspace estimation under coarse quantization. In *2025 International Conference on Sampling Theory and Applications (SampTA)*, pp. 1–5. IEEE, 2025. 1
- Eamaz, A., Yeganegi, F., and Soltanalian, M. Covariance recovery for one-bit sampled non-stationary signals with time-varying sampling thresholds. *IEEE Transactions on Signal Processing*, 70:5222–5236, 2022. 1, 4, 5, 8
- Eamaz, A., Yeganegi, F., and Soltanalian, M. Covariance recovery for one-bit sampled stationary signals with time-varying sampling thresholds. *Signal Processing*, 206:108899, 2023. 1, 4
- Eamaz, A., Yeganegi, F., and Soltanalian, M. Matrix completion from one-bit dither samples. *IEEE Transactions on Signal Processing*, pp. 1–14, 2024. 4
- Fan, J., Wang, W., and Zhu, Z. A shrinkage principle for heavy-tailed data: High-dimensional robust low-rank matrix recovery. *Annals of statistics*, 49(3):1239, 2021. 3
- Fang, J. and Li, H. Distributed adaptive quantization for wireless sensor networks: From delta modulation to maximum likelihood. *IEEE Transactions on Signal Processing*, 56(10):5246–5257, 2008. 1, 3, 4
- Fang, J. and Li, H. Distributed estimation of Gauss-Markov random fields with one-bit quantized data. *IEEE Signal Processing Letters*, 17(5):449–452, 2010. 5
- Finesso, L., Gerencsér, L., and Kmeics, I. A randomized EM-algorithm for estimating quantized linear Gaussian regression. In *Proceedings of the 38th IEEE Conference on Decision and Control*, volume 5, pp. 5100–5101. IEEE, 1999. 3

- Galarza, C. E., Lin, T.-I., Wang, W.-L., and Lachos, V. H. On moments of folded and truncated multivariate student-t distributions based on recurrence relations. *Metrika*, 84(6):825–850, 2021. 5
- Gao, P., Wang, R., Wang, M., and Chow, J. H. Low-rank matrix recovery from noisy, quantized, and erroneous measurements. *IEEE Transactions on Signal Processing*, 66(11):2918–2932, 2018. 1
- Gray, R. M. and Neuhoff, D. L. Quantization. *IEEE transactions on information theory*, 44(6):2325–2383, 2002. 1
- Gray, R. M. and Stockham, T. G. Dithered quantizers. *IEEE Transactions on Information Theory*, 39(3):805–812, 1993. 1
- Györfi, L. and Wegkamp, M. Quantization for nonparametric regression. *IEEE Transactions on Information Theory*, 54(2):867–874, 2008. 1
- Harper, F. M. and Konstan, J. A. The movielens datasets: History and context. *ACM transactions on interactive intelligent systems*, 5(4):1–19, 2015. 7
- Ho, H. J., Lin, T.-I., Chen, H.-Y., and Wang, W.-L. Some results on the truncated multivariate t distribution. *Journal of Statistical Planning and Inference*, 142(1):25–40, 2012. 13
- Jacques, L., Laska, J. N., Boufounos, P. T., and Baraniuk, R. G. Robust 1-bit compressive sensing via binary stable embeddings of sparse vectors. *IEEE Transactions on Information Theory*, 59(4):2082–2102, 2013. 4
- Jin, B., Zhu, J., Wu, Q., Zhang, Y., and Xu, Z. One-bit lfmw radar: Spectrum analysis and target detection. *IEEE Transactions on Aerospace and Electronic Systems*, 56(4):2732–2750, 2020. 1
- Jung, H. C., Maly, J., Palzer, L., and Stollenwerk, A. Quantized compressed sensing by rectified linear units. *IEEE transactions on information theory*, 67(6):4125–4149, 2021. 4
- Kim, N.-H. and Browne, R. P. Flexible mixture regression with the generalized hyperbolic distribution. *Advances in Data Analysis and Classification*, 18(1):33–60, 2024. 3
- Kipnis, A. and Reeves, G. Gaussian approximation of quantization error for estimation from compressed data. *IEEE Transactions on Information Theory*, 67(8):5562–5579, 2021. 1
- Knudson, K., Saab, R., and Ward, R. One-bit compressive sensing with norm estimation. *IEEE Transactions on Information Theory*, 62(5):2748–2758, 2016. 4
- Lai, W. L. A probabilistic graphical model based data compression architecture for Gaussian sources. Master’s thesis, Massachusetts Institute of Technology, 2016. URL <http://hdl.handle.net/1721.1/117322>. 5
- Laska, J. N., Wen, Z., Yin, W., and Baraniuk, R. G. Trust, but verify: Fast and accurate signal recovery from 1-bit compressive measurements. *IEEE Transactions on Signal Processing*, 59(11):5289–5301, 2011. 4, 8
- Lévy, P. Sur certains processus stochastiques homogènes. *Compositio mathematica*, 7:283–339, 1940. 4
- Li, Z., Xu, W., Zhang, X., and Lin, J. A survey on one-bit compressed sensing: Theory and applications. *Frontiers of Computer Science*, 12(2):217–230, 2018. 4
- Liu, C.-L. and Chou, Y.-H. Approximation and analysis of the one-bit Hermite law. In *2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pp. 1–5. IEEE, 2025. 4, 5, 8
- Liu, C.-L. and Lin, Z.-M. One-bit autocorrelation estimation with non-zero thresholds. In *2021 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pp. 4520–4524. IEEE, 2021. 1, 4, 5, 8
- Liu, X., Han, X., Chi, E. C., and Nadler, B. A majorization-minimization gauss-newton method for 1-bit matrix completion. *Journal of Computational and Graphical Statistics*, 34(3):1017–1029, 2025. 4
- Lu, X., Liu, W., and Alomainy, A. A 1.5-bit quantization scheme and its application to sparse array direction estimation. In *2024 19th International Symposium on Wireless Communication Systems (ISWCS)*, pp. 1–5, 2024. 1
- Maly, J., Yang, T., Dirksen, S., Rauhut, H., and Caire, G. New challenges in covariance estimation: multiple structures and coarse quantization. In *Compressed Sensing in Information Processing*, pp. 77–104. Springer, 2022. 5
- Mayne, D. Q. The use of quantization in regression analysis. *International Journal of Control*, 6(6):573–577, 1967. 1
- McLachlan, G. J. and Krishnan, T. *The EM algorithm and extensions*. John Wiley & Sons, 2008. 6, 18
- Meng, X.-L. and Rubin, D. B. On the global and componentwise rates of convergence of the em algorithm. *Linear Algebra and its Applications*, 199:413–425, 1994. 6
- Mo, J., Schniter, P., and Heath, R. W. Channel estimation in broadband millimeter wave MIMO systems with few-bit ADCs. *IEEE Transactions on Signal Processing*, 66(5):1141–1154, 2017. 1

- Nobre, I. C. M. and Frossard, P. Optimized quantization in distributed graph signal filtering. *arXiv preprint arXiv:1909.12725*, 2019. [5](#)
- Papadopoulos, H. C., Wornell, G. W., and Oppenheim, A. V. Sequential signal encoding from noisy measurements using quantizers with dynamic bias control. *IEEE Transactions on Information Theory*, 47(3):978–1002, 2001. [1](#), [3](#)
- Paterek, A. Improving regularized singular value decomposition for collaborative filtering. In *Proceedings of the KDD Cup and Workshop*, pp. 394–401. ACM, 2007. [7](#)
- Plabst, D., Munir, J., Mezghani, A., and Nossek, J. A. Efficient non-linear equalization for 1-bit quantized cyclic prefix-free massive MIMO systems. In *2018 15th International Symposium on Wireless Communication Systems (ISWCS)*, pp. 1–6, 2018. [1](#)
- Plan, Y. and Vershynin, R. Robust 1-bit compressed sensing and sparse logistic regression: A convex programming approach. *IEEE Transactions on Information Theory*, 59(1):482–494, 2013. [1](#), [4](#)
- Polson, N. G. and Scott, J. G. Data augmentation for non-Gaussian regression models using variance-mean mixtures. *Biometrika*, 100(2):459–471, 2013. [2](#)
- Ren, J. and Li, J. One-bit digital radar. In *2017 51st Asilomar Conference on Signals, Systems, and Computers*, pp. 1142–1146, 2017. [1](#)
- Ribeiro, A. and Giannakis, G. B. Bandwidth-constrained distributed estimation for wireless sensor networks-Part II: Unknown probability density function. *IEEE Transactions on Signal Processing*, 54(7):2784–2796, 2006a. [1](#), [3](#)
- Ribeiro, A. and Giannakis, G. B. Bandwidth-constrained distributed estimation for wireless sensor networks-Part I: Gaussian case. *IEEE Transactions on Signal Processing*, 54(3):1131–1143, 2006b. [1](#), [3](#)
- Roberts, L. Picture coding using pseudo-random noise. *IRE Transactions on Information Theory*, 8(2):145–154, 1962. [1](#)
- Saad, L. B., Beferull-Lozano, B., and Isufi, E. Quantization analysis and robust design for distributed graph filters. *IEEE Transactions on Signal Processing*, 70:643–658, 2021. [5](#)
- Saha, R., Srivastava, V., and Pilanci, M. Matrix compression via randomized low rank and low precision factorization. *Advances in Neural Information Processing Systems*, 36:18828–18872, 2023. [5](#)
- Sarwar, B. M., Karypis, G., Konstan, J. A., and Riedl, J. T. Application of dimensionality reduction in recommender systems: A case study. In *WebKDD Workshop at the ACM SIGKDD*. ACM, 2000. [7](#)
- Shen, J., Awasthi, P., and Li, P. Robust matrix completion from quantized observations. In *The 22nd International Conference on Artificial Intelligence and Statistics*, pp. 397–407. PMLR, 2019. [4](#)
- Stöckle, C., Munir, J., Mezghani, A., and Nossek, J. A. Channel estimation in massive MIMO systems using 1-bit quantization. In *2016 IEEE 17th International Workshop on Signal Processing Advances in Wireless Communications (SPAWC)*, pp. 1–6, 2016. [1](#)
- Tavassolipour, M., Motahari, S. A., and Shalmani, M.-T. M. Learning of tree-structured Gaussian graphical models on distributed data under communication constraints. *IEEE Transactions on Signal Processing*, 67(1):17–28, 2018. [5](#)
- Van Vleck, J. and Middleton, D. The spectrum of clipped noise. *Proceedings of the IEEE*, 54(1):2–19, 1966. [1](#), [4](#), [8](#)
- Widrow, B. and Kollár, I. *Quantization noise: round-off error in digital computation, signal processing, control, and communications*. Cambridge University Press, 2008. [1](#)
- Xiao, Y.-H., Huang, L., Ramírez, D., Qian, C., and So, H. C. Covariance matrix recovery from one-bit data with non-zero quantization thresholds: Algorithm and performance analysis. *IEEE Transactions on Signal Processing*, 71:4060–4076, 2023. [1](#), [4](#), [8](#)
- Yan, M., Yang, Y., and Osher, S. Robust 1-bit compressive sensing using adaptive outlier pursuit. *IEEE Transactions on Signal Processing*, 60(7):3868–3875, 2012. [4](#), [8](#)
- Yang, J., Song, Z., Zhang, H., and Gao, Y. Compressive spectrum sensing with 1-bit ADCs. *IEEE Transactions on Vehicular Technology*, 2025. [1](#), [2](#)
- Zymnis, A., Boyd, S., and Candes, E. Compressed sensing with quantized measurements. *IEEE Signal Processing Letters*, 17(2):149–152, 2009. [1](#), [4](#), [7](#), [8](#), [26](#)

A. Expectation Evaluations on Specific Distribution Assumptions

From Section 4, we have established that when the random variable $\mathbf{x}$ in model (3) belongs to the normal variance-mean mixture family, the surrogate function (13) can be derived, along with its closed-form solutions with respect to $\boldsymbol{\mu}$, $\boldsymbol{\xi}$, and $\boldsymbol{\Sigma}$. Nevertheless, these closed-form solutions (15) depend on the expected values $\mathbf{U}_t$, $\mathbf{v}_t$, $\mathbf{w}_t$, ℓ_t , and ζ_t , which are determined by the distribution of the random variable $\mathbf{x}$. Therefore, in the following, we analyze the integrals corresponding to these expected values under various distributions of $\mathbf{x}$, and examine the convergence properties and statistical characteristics of the proposed algorithm under specific distributional assumptions. We first consider the fundamental case of the Gaussian distribution (i.e., $z = 1$ and $\boldsymbol{\xi} = \mathbf{0}$). Let $\mathbf{x} \sim \mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$. In this case, the first- and second-order moments are given by (Ho et al., 2012) as

$$\mathbb{E}_{\mathbf{x}|\mathbf{y};\boldsymbol{\theta}}[\mathbf{x}] = \boldsymbol{\mu} + p^{-1}(\mathbf{y}; \boldsymbol{\theta}) \boldsymbol{\Sigma} \mathbf{q}, \quad (21)$$

$$\mathbb{E}_{\mathbf{x}|\mathbf{y};\boldsymbol{\theta}}[\mathbf{x}\mathbf{x}^\top] = \boldsymbol{\mu}\boldsymbol{\mu}^\top + \boldsymbol{\Sigma} + p^{-1}(\mathbf{y}; \boldsymbol{\theta}) (2(\boldsymbol{\Sigma}\mathbf{q})\boldsymbol{\mu}^\top + \boldsymbol{\Sigma}(\mathbf{H} + \mathbf{D})\boldsymbol{\Sigma}). \quad (22)$$

Let the interval $\mathcal{Q}^{-1}(y_i)$ be $[l_i, u_i]$. The i -th elements of vector $\mathbf{q}$ is given by

$$q_i = p(x_i = l_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) - p(x_i = u_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}), \quad (23)$$

The matrix $\mathbf{H}$ is a matrix with all diagonal entries being zero and the (i, j) -th off-diagonal element being

$$\mathbf{H}_{ij} = h(l_i, l_j) - h(l_i, u_j) - h(u_i, l_j) + h(u_i, u_j), \quad (24)$$

with

$$h(c_i, c_j) = p(x_i = c_i, x_j = c_j, \mathbf{y}_{\setminus i, j}; \boldsymbol{\theta}).$$

The matrix $\mathbf{D}$ is a diagonal matrix with the diagonal entries:

$$D_{ii} = \frac{l_i - \mu_i}{\sigma_i^2} p(x_i = l_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) - \frac{u_i - \mu_i}{\sigma_i^2} p(x_i = u_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) - \frac{[\boldsymbol{\Sigma}\mathbf{H}]_{ii}}{\sigma_i^2}. \quad (25)$$

Now we consider the general normal variance-mean mixture with $\mathbf{x} = \boldsymbol{\mu} + z\boldsymbol{\xi} + (z\boldsymbol{\Sigma})^{\frac{1}{2}}\boldsymbol{\epsilon}$. In this case, the first and second moments satisfy

$$\mathbb{E}_{\mathbf{x}, z|\mathbf{y};\boldsymbol{\theta}}[\mathbf{x}] = \mathbb{E}_{z|\mathbf{y};\boldsymbol{\theta}}[\mathbb{E}_{\mathbf{x}|z, \mathbf{y};\boldsymbol{\theta}}[\mathbf{x}]],$$

and

$$\mathbb{E}_{\mathbf{x}, z|\mathbf{y};\boldsymbol{\theta}}[\mathbf{x}\mathbf{x}^\top] = \mathbb{E}_{z|\mathbf{y};\boldsymbol{\theta}}[\mathbb{E}_{\mathbf{x}|z, \mathbf{y};\boldsymbol{\theta}}[\mathbf{x}\mathbf{x}^\top]].$$

Based on $\mathbf{x} | z \sim \mathcal{N}(\boldsymbol{\mu} + z\boldsymbol{\xi}, z\boldsymbol{\Sigma})$ and the moments in (21) and (22), we have

$$\mathbb{E}_{\mathbf{x}|z, \mathbf{y};\boldsymbol{\theta}}[\mathbf{x}] = \boldsymbol{\mu} + z\boldsymbol{\xi} + p^{-1}(\mathbf{y} | z; \boldsymbol{\theta}) z\boldsymbol{\Sigma}\mathbf{q}_z, \quad (26)$$

$$\begin{aligned} \mathbb{E}_{\mathbf{x}|z, \mathbf{y};\boldsymbol{\theta}}[\mathbf{x}\mathbf{x}^\top] &= \boldsymbol{\mu}\boldsymbol{\mu}^\top + z^2\boldsymbol{\xi}\boldsymbol{\xi}^\top + 2z\boldsymbol{\mu}\boldsymbol{\xi}^\top + z\boldsymbol{\Sigma} + p^{-1}(\mathbf{y} | z; \boldsymbol{\theta}) (2z(\boldsymbol{\Sigma}\mathbf{q}_z)\boldsymbol{\mu}^\top \\ &\quad + 2z^2(\boldsymbol{\Sigma}\mathbf{q}_z)\boldsymbol{\xi}^\top + z^2\boldsymbol{\Sigma}(\mathbf{H}_z + \mathbf{D}_z)\boldsymbol{\Sigma}), \end{aligned} \quad (27)$$

where the i -th elements of vector $\mathbf{q}_z$ is denoted as

$$q_{z,i} = p(x_i = l_i, \mathbf{y}_{\setminus i} | z; \boldsymbol{\theta}) - p(x_i = u_i, \mathbf{y}_{\setminus i} | z; \boldsymbol{\theta}), \quad (28)$$

the matrix $\mathbf{H}_z$ is a matrix with all diagonal entries being zero and the (i, j) -th off-diagonal element being

$$\mathbf{H}_{z,ij} = h_z(l_i, l_j) - h_z(l_i, u_j) - h_z(u_i, l_j) + h_z(u_i, u_j), \quad (29)$$

with

$$h_z(c_i, c_j) = p(x_i = c_i, x_j = c_j, \mathbf{y}_{\setminus i, j} | z; \boldsymbol{\theta}),$$

and the matrix $\mathbf{D}_z$ is a diagonal matrix with the diagonal entries:

$$D_{z,ii} = \frac{l_i - \mu_i - z\xi_i}{\sigma_i^2} p(x_i = l_i, \mathbf{y}_{\setminus i} | z; \boldsymbol{\theta}) - \frac{u_i - \mu_i - z\xi_i}{\sigma_i^2} p(x_i = u_i, \mathbf{y}_{\setminus i} | z; \boldsymbol{\theta}) - \frac{[\boldsymbol{\Sigma}\mathbf{H}_z]_{ii}}{\sigma_i^2}. \quad (30)$$

Based on (26) and (27), we can obtain

$$\mathbf{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x}] = \boldsymbol{\mu} + \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z] \boldsymbol{\xi} + \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z p^{-1}(\mathbf{y} | z; \boldsymbol{\theta}) \boldsymbol{\Sigma} \mathbf{q}_z], \quad (31)$$

$$\begin{aligned} \mathbf{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x} \mathbf{x}^\top] &= \boldsymbol{\mu} \boldsymbol{\mu}^\top + \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z^2] \boldsymbol{\xi} \boldsymbol{\xi}^\top + 2 \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z] \boldsymbol{\mu} \boldsymbol{\xi}^\top + \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z] \boldsymbol{\Sigma} + \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [p^{-1}(\mathbf{y} | z; \boldsymbol{\theta}) \\ &\quad (2z (\boldsymbol{\Sigma} \mathbf{q}_z) \boldsymbol{\mu}^\top + 2z^2 (\boldsymbol{\Sigma} \mathbf{q}_z) \boldsymbol{\xi}^\top + z^2 \boldsymbol{\Sigma} (\mathbf{H}_z + \mathbf{D}_z) \boldsymbol{\Sigma})] \end{aligned} \quad (32)$$

We first compute $\mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z^k]$ as follows:

$$\mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z^k] = \int_0^{+\infty} p(z | \mathbf{y}; \boldsymbol{\theta}) z^k dz = \int_0^{+\infty} \frac{p(\mathbf{y} | z; \boldsymbol{\theta}) p(z)}{p(\mathbf{y}; \boldsymbol{\theta})} z^k dz.$$

Here, we introduce the size-biased distribution of order k of the positive random variable z , which has the density function $p_k(z) = \frac{z^k p(z)}{\mathbf{E}_z[z^k]}$. Hence, we have

$$\mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z^k] = \mathbf{E}_z[z^k] \frac{\int_0^{+\infty} p(\mathbf{y} | z; \boldsymbol{\theta}) p_k(z) dz}{p(\mathbf{y}; \boldsymbol{\theta})} = \mathbf{E}_z[z^k] \frac{p_k(\mathbf{y}; \boldsymbol{\theta})}{p(\mathbf{y}; \boldsymbol{\theta})}, \quad (33)$$

where we denote $p_k(\mathbf{y}; \boldsymbol{\theta}) = \int_0^{+\infty} p(\mathbf{y} | z; \boldsymbol{\theta}) p_k(z) dz$.

Similarly, we can obtain

$$\mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z^k p^{-1}(\mathbf{y} | z; \boldsymbol{\theta}) \mathbf{q}_z] = \mathbf{E}_z[z^k] p^{-1}(\mathbf{y}; \boldsymbol{\theta}) \mathbf{q}_k, \quad (34)$$

where $\mathbf{q}_k = \int_0^{+\infty} \mathbf{q}_z p_k(z) dz$ with the i -th element

$$q_{k,i} = p_k(x_i = l_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) - p_k(x_i = u_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}). \quad (35)$$

For the expectation of $z^k \mathbf{H}_z$ we have

$$\mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z^k p^{-1}(\mathbf{y} | z; \boldsymbol{\theta}) (\mathbf{H}_z + \mathbf{D}_z)] = p^{-1}(\mathbf{y}; \boldsymbol{\theta}) (\mathbf{H}_k + \mathbf{D}_k), \quad (36)$$

with the (i, j) entry of

$$\begin{aligned} \mathbf{H}_{k,ij} &= \mathbf{E}_z[z^k] (h_k(l_i, l_j) - h_k(l_i, u_j) - h_k(u_i, l_j) + h_k(u_i, u_j)), \\ h_k(c_i, c_j) &= p_k(x_i = c_i, x_j = c_j, \mathbf{y}_{\setminus i,j}; \boldsymbol{\theta}), \end{aligned}$$

and the matrix $\mathbf{D}_k$ is a diagonal matrix with the diagonal entries:

$$\begin{aligned} \mathbf{D}_{k,ii} &= \mathbf{E}_z[z^{k-1}] \frac{l_i - \mu_i}{\sigma_i^2} p_{k-1}(x_i = l_i, \mathbf{y}_{\setminus i} | z; \boldsymbol{\theta}) - \mathbf{E}_z[z^{k-1}] \frac{u_i - \mu_i}{\sigma_i^2} p_{k-1}(x_i = u_i, \mathbf{y}_{\setminus i} | z; \boldsymbol{\theta}) \\ &\quad - \mathbf{E}_z[z^k] \frac{\xi_i}{\sigma_i^2} q_{k,i} - \frac{[\boldsymbol{\Sigma} \mathbf{H}_k]_{ii}}{\sigma_i^2}. \end{aligned} \quad (37)$$

Therefore, the expectations in (31) and (32) become

$$\mathbf{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x}] = \boldsymbol{\mu} + \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z] \boldsymbol{\xi} + \mathbf{E}_z [z] p^{-1}(\mathbf{y}; \boldsymbol{\theta}) \boldsymbol{\Sigma} \mathbf{q}_1, \quad (38)$$

$$\begin{aligned} \mathbf{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x} \mathbf{x}^\top] &= \boldsymbol{\mu} \boldsymbol{\mu}^\top + \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z^2] \boldsymbol{\xi} \boldsymbol{\xi}^\top + 2 \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z] \boldsymbol{\mu} \boldsymbol{\xi}^\top + \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z] \boldsymbol{\Sigma} + p^{-1}(\mathbf{y}; \boldsymbol{\theta}) \\ &\quad \times \left(2 \mathbf{E}_z [z] (\boldsymbol{\Sigma} \mathbf{q}_1) \boldsymbol{\mu}^\top + 2 \mathbf{E}_z [z^2] (\boldsymbol{\Sigma} \mathbf{q}_2) \boldsymbol{\xi}^\top + \boldsymbol{\Sigma} (\mathbf{H}_2 + \mathbf{D}_2) \boldsymbol{\Sigma} \right). \end{aligned} \quad (39)$$

Meanwhile, we can also obtain

$$\mathbf{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [z^{-1} \mathbf{x}] = \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [\mathbf{E}_{\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}} [z^{-1} \mathbf{x}]] = \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z^{-1}] \boldsymbol{\mu} + \boldsymbol{\xi} + p^{-1}(\mathbf{y}; \boldsymbol{\theta}) \boldsymbol{\Sigma} \mathbf{q}, \quad (40)$$

$$(41)$$

$$\begin{aligned} \mathbf{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [z^{-1} \mathbf{x} \mathbf{x}^\top] &= \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [\mathbf{E}_{\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}} [z^{-1} \mathbf{x} \mathbf{x}^\top]] = \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z^{-1}] \boldsymbol{\mu} \boldsymbol{\mu}^\top + \mathbf{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z] \boldsymbol{\xi} \boldsymbol{\xi}^\top + 2 \boldsymbol{\mu} \boldsymbol{\xi}^\top + \boldsymbol{\Sigma} + p^{-1}(\mathbf{y}; \boldsymbol{\theta}) \times \\ &\quad \left(2 (\boldsymbol{\Sigma} \mathbf{q}) \boldsymbol{\mu}^\top + 2 \mathbf{E}_z [z] (\boldsymbol{\Sigma} \mathbf{q}_1) \boldsymbol{\xi}^\top + \boldsymbol{\Sigma} (\mathbf{H}_1 + \mathbf{D}_1) \boldsymbol{\Sigma} \right). \end{aligned} \quad (42)$$

B. The Proof of Theorem 2

Given the log-likelihood function as

$$L(\boldsymbol{\theta}) = \sum_{t=1}^n \log p(\mathbf{y}_t; \boldsymbol{\theta}).$$

in the following, we prove that $L(\boldsymbol{\theta})$ is component-wise concave with respect to $\boldsymbol{\mu}$, $\boldsymbol{\xi}$, and $\boldsymbol{\Sigma}$.

B.1. Concavity of Location Parameter

The second order derivative of $\log p(\mathbf{y}; \boldsymbol{\theta})$ with respect to $\boldsymbol{\mu}$ is given as

$$\nabla_{\boldsymbol{\mu}}^2 \log p(\mathbf{y}; \boldsymbol{\theta}) = \frac{\nabla_{\boldsymbol{\mu}}^2 p(\mathbf{y}; \boldsymbol{\theta}) \cdot p(\mathbf{y}; \boldsymbol{\theta}) - \nabla_{\boldsymbol{\mu}} p(\mathbf{y}; \boldsymbol{\theta}) \nabla_{\boldsymbol{\mu}}^\top p(\mathbf{y}; \boldsymbol{\theta})}{p^2(\mathbf{y}; \boldsymbol{\theta})}. \quad (43)$$

To analyze the above expression, we should compute the first and second order derivatives of $p(\mathbf{y}; \boldsymbol{\theta})$ with respect to $\boldsymbol{\mu}$ at first. The first order one is given as follows:

$$\nabla_{\boldsymbol{\mu}} p(\mathbf{y}; \boldsymbol{\theta}) = \nabla_{\boldsymbol{\mu}} \int_{\mathcal{Q}^{-1}(\mathbf{y})} p(\mathbf{x}; \boldsymbol{\theta}) d\mathbf{x} = \int_{\mathcal{Q}^{-1}(\mathbf{y})} \nabla_{\boldsymbol{\mu}} p(\mathbf{x}; \boldsymbol{\theta}) d\mathbf{x}. \quad (44)$$

Since $p(\mathbf{x}; \boldsymbol{\theta}) = \int_0^{+\infty} p(z) p(\mathbf{x} | z; \boldsymbol{\theta}) dz$, (44) becomes

$$\nabla_{\boldsymbol{\mu}} p(\mathbf{y}; \boldsymbol{\theta}) = \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} p(z) \nabla_{\boldsymbol{\mu}} p(\mathbf{x} | z; \boldsymbol{\theta}) dz d\mathbf{x} = \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} \frac{p(z)}{z} p(\mathbf{x} | z; \boldsymbol{\theta}) \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi}) dz d\mathbf{x}.$$

Then the second order one can be further computed as

$$\begin{aligned} \nabla_{\boldsymbol{\mu}}^2 p(\mathbf{y}; \boldsymbol{\theta}) &= \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} \frac{p(z)}{z} \nabla_{\boldsymbol{\mu}} (p(\mathbf{x} | z; \boldsymbol{\theta}) \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})) dz d\mathbf{x} \\ &= \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} \boldsymbol{\Sigma}^{-1} p(\mathbf{x} | z; \boldsymbol{\theta}) p(z) (-z^{-1} \mathbf{I} + z^{-2} (\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})^\top \boldsymbol{\Sigma}^{-1}) dz d\mathbf{x}. \end{aligned} \quad (45)$$

Lemma 5. For a function $f(\mathbf{x}, z)$, if $\int_{\mathcal{Q}^{-1}(\mathbf{y})} p(\mathbf{x} | z; \boldsymbol{\theta}) f(\mathbf{x}, z) d\mathbf{x}$ is integrable, we have

$$\int_{\mathbb{R}^d} \int_0^{+\infty} p(z) p(\mathbf{x}, \mathbf{y} | z; \boldsymbol{\theta}) f(\mathbf{x}, z) dz d\mathbf{x} = p(\mathbf{y}; \boldsymbol{\theta}) \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [f(\mathbf{x}, z)].$$

Proof. Given that

$$\int_{\mathcal{Q}^{-1}(\mathbf{y})} p(\mathbf{x} | z; \boldsymbol{\theta}) f(\mathbf{x}, z) d\mathbf{x} = \int_{\mathbb{R}^d} p(\mathbf{x}, \mathbf{y} | z; \boldsymbol{\theta}) f(\mathbf{x}, z) d\mathbf{x},$$

we can obtain that

$$\int_{\mathbb{R}^d} \int_0^{+\infty} p(z) p(\mathbf{x}, \mathbf{y} | z; \boldsymbol{\theta}) f(\mathbf{x}, z) dz d\mathbf{x} = p(\mathbf{y}; \boldsymbol{\theta}) \int_{\mathbb{R}^d} \int_0^{+\infty} p(\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}) f(\mathbf{x}, z) dz d\mathbf{x} = p(\mathbf{y}; \boldsymbol{\theta}) \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [f(\mathbf{x}, z)].$$

□

Based on Lemma 5, the first and second order derivatives of $p(\mathbf{y}; \boldsymbol{\theta})$ becomes

$$\nabla_{\boldsymbol{\mu}} p(\mathbf{y}; \boldsymbol{\theta}) = p(\mathbf{y}; \boldsymbol{\theta}) \boldsymbol{\Sigma}^{-1} \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [z^{-1} (\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})], \quad (46)$$

and

$$\nabla_{\boldsymbol{\mu}}^2 p(\mathbf{y}; \boldsymbol{\theta}) = p(\mathbf{y}; \boldsymbol{\theta}) \boldsymbol{\Sigma}^{-1} (-\mathbb{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z^{-1}] \mathbf{I} + \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [z^{-2} (\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})^\top] \boldsymbol{\Sigma}^{-1}). \quad (47)$$

Substituting (46) and (47) into the second order derivative of $\log p(\mathbf{y}; \boldsymbol{\theta})$ in (43) becomes

$$\begin{aligned} \nabla_{\boldsymbol{\mu}}^2 \log p(\mathbf{y}; \boldsymbol{\theta}) &= \boldsymbol{\Sigma}^{-1} \left(-\mathbb{E}_{z|\mathbf{y};\boldsymbol{\theta}} [z^{-1}] \mathbf{I} + \mathbb{E}_{x,z|\mathbf{y};\boldsymbol{\theta}} [z^{-2}(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})^\top] \boldsymbol{\Sigma}^{-1} \right. \\ &\quad \left. - \boldsymbol{\Sigma}^{-1} \mathbb{E}_{x,z|\mathbf{y};\boldsymbol{\theta}} [z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})] \mathbb{E}_{x,z|\mathbf{y};\boldsymbol{\theta}}^\top [z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})] \boldsymbol{\Sigma}^{-1} \right) \\ &= \boldsymbol{\Sigma}^{-1} \left(-\mathbb{E}_{z|\mathbf{y};\boldsymbol{\theta}} [z^{-1}] \boldsymbol{\Sigma} + \text{Cov}_{x,z|\mathbf{y};\boldsymbol{\theta}} [z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})] \right) \boldsymbol{\Sigma}^{-1}, \end{aligned} \quad (48)$$

Since $\boldsymbol{\Sigma}^{-1}$ on the both sides of the parentheses is positive definite, the definiteness of $\nabla_{\boldsymbol{\mu}}^2 \log p(\mathbf{y}; \boldsymbol{\theta})$ is determined by the term in the parentheses. For the covariance term, based on law of total variance, we have

$$\text{Cov}_{x,z|\mathbf{y};\boldsymbol{\theta}} [z^{-1}(\mathbf{x} - \boldsymbol{\mu})] = \mathbb{E}_{x,z|\mathbf{y};\boldsymbol{\theta}} [\text{Cov}_{x|z,\mathbf{y};\boldsymbol{\theta}} [z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})]] + \text{Cov}_{x,z|\mathbf{y};\boldsymbol{\theta}} [\mathbb{E}_{x|z,\mathbf{y};\boldsymbol{\theta}} [z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})]]. \quad (49)$$

Due to $\mathbb{E}_{x|z,\mathbf{y};\boldsymbol{\theta}} [\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi}] = \mathbf{0}$, we only consider the first term in (49). To describe the relations between the expectation term and the covariance term, we introduce the Brascamp–Lieb inequality.

Lemma 6 ((Brascamp & Lieb, 1976)). *Consider a probability density function $p(\mathbf{X} | \mathbf{Y})$ which is log-concave to $\mathbf{X}$. The Brascamp–Lieb inequality is given by*

$$\text{Cov}_{\mathbf{X}|\mathbf{Y}}[f(\mathbf{X})] \preceq \mathbb{E}_{\mathbf{X}|\mathbf{Y}} [\nabla_{\mathbf{X}}^\top f(\mathbf{X}) [\nabla_{\mathbf{X}}^2 (-\log p(\mathbf{X} | \mathbf{Y}))]^{-1} \nabla_{\mathbf{X}} f(\mathbf{X})],$$

where the equality is obtained when $\log p(\mathbf{X} | \mathbf{Y})$ is linear with respect to $\mathbf{X}$.

Let $\mathbf{X} = \mathbf{x}$, $\mathbf{X} | \mathbf{Y} = \mathbf{x} | z, \mathbf{y}$, and $f(\mathbf{X}) = z^{-1}\mathbf{x}$. Given that

$$\nabla_{\mathbf{x}}^2 \log p(\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}) = \nabla_{\mathbf{x}}^2 \log \left(\frac{p(\mathbf{x}, z; \boldsymbol{\theta}) p(\mathbf{y} | \mathbf{x}; \boldsymbol{\theta})}{p(\mathbf{y}, z; \boldsymbol{\theta})} \right) = -\frac{\boldsymbol{\Sigma}^{-1}}{z} \prec \mathbf{0},$$

and $\nabla_{\mathbf{x}} z^{-1}(\mathbf{x} - \boldsymbol{\mu}) = z^{-1}\mathbf{1}$, based on the Brascamp–Lieb inequality in Lemma 6, we have

$$\text{Cov}_{x|z,\mathbf{y};\boldsymbol{\theta}} [z^{-1}(\mathbf{x} - \boldsymbol{\mu})] \preceq \mathbb{E}_{x|z,\mathbf{y};\boldsymbol{\theta}} [z^{-1}z\boldsymbol{\Sigma}z^{-1}] = z^{-1}\boldsymbol{\Sigma},$$

Since $p(\mathbf{x} | z, \mathbf{y})$ is not linear to $\mathbf{x}$, the above inequality is strict. Hence we have

$$\text{Cov}_{x,z|\mathbf{y};\boldsymbol{\theta}} [z^{-1}(\mathbf{x} - \boldsymbol{\mu})] \prec \mathbb{E}_{x,z|\mathbf{y};\boldsymbol{\theta}} [z^{-1}] \boldsymbol{\Sigma},$$

i.e., $\nabla_{\boldsymbol{\mu}}^2 \log p(\mathbf{y}; \boldsymbol{\theta})$ is negative definite.

B.2. Concavity of Skewness Parameter

The second order derivative of $\log p(\mathbf{y}; \boldsymbol{\theta})$ with respect to $\boldsymbol{\xi}$ is given as

$$\nabla_{\boldsymbol{\xi}}^2 \log p(\mathbf{y}; \boldsymbol{\theta}) = \frac{\nabla_{\boldsymbol{\xi}}^2 p(\mathbf{y}; \boldsymbol{\theta}) \cdot p(\mathbf{y}; \boldsymbol{\theta}) - \nabla_{\boldsymbol{\xi}} p(\mathbf{y}; \boldsymbol{\theta}) \nabla_{\boldsymbol{\xi}}^\top p(\mathbf{y}; \boldsymbol{\theta})}{p^2(\mathbf{y}; \boldsymbol{\theta})}. \quad (50)$$

To analyze the above expression, we should compute the first and second order derivatives of $p(\mathbf{y}; \boldsymbol{\theta})$ with respect to $\boldsymbol{\xi}$ at first. The first order one is given as follows:

$$\nabla_{\boldsymbol{\xi}} p(\mathbf{y}; \boldsymbol{\theta}) = \nabla_{\boldsymbol{\xi}} \int_{\mathcal{Q}^{-1}(\mathbf{y})} p(\mathbf{x}; \boldsymbol{\theta}) d\mathbf{x} = \int_{\mathcal{Q}^{-1}(\mathbf{y})} \nabla_{\boldsymbol{\xi}} p(\mathbf{x}; \boldsymbol{\theta}) d\mathbf{x}. \quad (51)$$

Since $p(\mathbf{x}; \boldsymbol{\theta}) = \int_0^{+\infty} p(z) p(\mathbf{x} | z; \boldsymbol{\theta}) dz$, (51) becomes

$$\nabla_{\boldsymbol{\xi}} p(\mathbf{y}; \boldsymbol{\theta}) = \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} p(z) \nabla_{\boldsymbol{\xi}} p(\mathbf{x} | z; \boldsymbol{\theta}) dz d\mathbf{x} = \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} (p(z) p(\mathbf{x} | z; \boldsymbol{\theta}) \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})) dz d\mathbf{x}.$$

Then the second order derivative can be further computed as

$$\begin{aligned} \nabla_{\boldsymbol{\xi}}^2 p(\mathbf{y}; \boldsymbol{\theta}) &= \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} p(z) \nabla_{\boldsymbol{\xi}} (p(\mathbf{x} | z; \boldsymbol{\theta}) \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})) dz d\mathbf{x} \\ &= \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} \boldsymbol{\Sigma}^{-1} p(\mathbf{x} | z; \boldsymbol{\theta}) p(z) (-z\mathbf{I} + (\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})^\top \boldsymbol{\Sigma}^{-1}) dz d\mathbf{x}. \end{aligned} \quad (52)$$

Based on Lemma 5, the derivatives of $p(\mathbf{y}; \boldsymbol{\theta})$ becomes

$$\nabla_{\boldsymbol{\xi}} p(\mathbf{y}; \boldsymbol{\theta}) = p(\mathbf{y}; \boldsymbol{\theta}) \boldsymbol{\Sigma}^{-1} \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi}], \quad (53)$$

and

$$\nabla_{\boldsymbol{\xi}}^2 p(\mathbf{y}; \boldsymbol{\theta}) = p(\mathbf{y}; \boldsymbol{\theta}) \boldsymbol{\Sigma}^{-1} (-\mathbb{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z] \mathbf{I} + \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})(\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi})^\top] \boldsymbol{\Sigma}^{-1}). \quad (54)$$

Substituting (53) and (54) into $\log p(\mathbf{y}; \boldsymbol{\theta})$ in (50), we have that

$$\nabla_{\boldsymbol{\xi}}^2 \log p(\mathbf{y}; \boldsymbol{\theta}) = \boldsymbol{\Sigma}^{-1} (-\mathbb{E}_{z | \mathbf{y}; \boldsymbol{\theta}} [z] \boldsymbol{\Sigma} + \text{Cov}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi}]) \boldsymbol{\Sigma}^{-1}. \quad (55)$$

Based on the law of total variance, the covariance term becomes

$$\text{Cov}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi}] = \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [\text{Cov}_{\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi}]] + \text{Cov}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [\mathbb{E}_{\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi}]]. \quad (56)$$

where the second term satisfies $\text{Cov}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [\mathbb{E}_{\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi}]] = \mathbf{0}$.

Setting $\mathbf{X} = \mathbf{x}$, $\mathbf{X} = \mathbf{x} | z, \mathbf{y}$, and $f(\mathbf{X}) = \mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi}$, since $p(\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta})$ is a log-concave function with respect to $\mathbf{x}$. Based on the Brascamp–Lieb inequality in Lemma 6, we have

$$\text{Cov}_{\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x} - \boldsymbol{\mu} - z\boldsymbol{\xi}] \prec \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [z] \boldsymbol{\Sigma},$$

and hence $\nabla_{\boldsymbol{\xi}}^2 p(\mathbf{y}; \boldsymbol{\theta})$ is negative definite.

B.3. Concavity of Scatter Matrix

Since $\log p(\mathbf{x} | z; \boldsymbol{\theta})$ is linear with respect to $\boldsymbol{\Sigma}^{-1}$, the second order derivative of $\log p(\mathbf{y}; \boldsymbol{\theta})$ with respect to $\boldsymbol{\Sigma}^{-1}$ is easier to be obtained. We first compute $\frac{\partial \text{vec}(\nabla_{\boldsymbol{\Sigma}^{-1}} \log p(\mathbf{y}; \boldsymbol{\theta}))}{\partial \text{vec}(\boldsymbol{\Sigma}^{-1})}$ and transform it later.

The derivative term $\frac{\partial \text{vec}(\nabla_{\boldsymbol{\Sigma}^{-1}} \log p(\mathbf{y}; \boldsymbol{\theta}))}{\partial \text{vec}(\boldsymbol{\Sigma}^{-1})}$ can be computed as

$$\frac{\partial \text{vec}(\nabla_{\boldsymbol{\Sigma}^{-1}} \log p(\mathbf{y}; \boldsymbol{\theta}))}{\partial \text{vec}(\boldsymbol{\Sigma}^{-1})} = \frac{\frac{\partial \text{vec}(\nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{y}; \boldsymbol{\theta}))}{\partial \text{vec}(\boldsymbol{\Sigma}^{-1})} \cdot p(\mathbf{y}; \boldsymbol{\theta}) - \nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{y}; \boldsymbol{\theta}) \otimes \nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{y}; \boldsymbol{\theta})}{p^2(\mathbf{y}; \boldsymbol{\theta})}. \quad (57)$$

Since the first and second order derivatives of $p(\mathbf{y}; \boldsymbol{\theta})$ with respect to $\boldsymbol{\Sigma}^{-1}$ satisfy

$$\nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{y}; \boldsymbol{\theta}) = \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} \nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}) dz d\mathbf{x},$$

and

$$\frac{\partial \text{vec}(\nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{y}; \boldsymbol{\theta}))}{\partial \text{vec}(\boldsymbol{\Sigma}^{-1})} = \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} \frac{\partial \text{vec}(\nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}))}{\partial \text{vec}(\boldsymbol{\Sigma}^{-1})} dz d\mathbf{x},$$

based on Lemma 5, the expression (57) becomes

$$\begin{aligned} & \frac{\frac{\partial \text{vec}(\nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{y}; \boldsymbol{\theta}))}{\partial \text{vec}(\boldsymbol{\Sigma}^{-1})} \cdot p(\mathbf{y}; \boldsymbol{\theta}) - \nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{y}; \boldsymbol{\theta}) \otimes \nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{y}; \boldsymbol{\theta})}{p^2(\mathbf{y}; \boldsymbol{\theta})} \\ &= \frac{\int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} \frac{\partial \text{vec}(\nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}))}{\partial \text{vec}(\boldsymbol{\Sigma}^{-1})} dz d\mathbf{x}}{p(\mathbf{y}; \boldsymbol{\theta})} - \frac{\int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} \nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}) dz d\mathbf{x}}{p(\mathbf{y}; \boldsymbol{\theta})} \otimes \frac{\int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} \nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}) dz d\mathbf{x}}{p(\mathbf{y}; \boldsymbol{\theta})} \\ &= \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\frac{\frac{\partial \text{vec}(\nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}))}{\partial \text{vec}(\boldsymbol{\Sigma}^{-1})}}{p(\mathbf{x}, z; \boldsymbol{\theta})} \right] - \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\frac{\nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta})}{p(\mathbf{x}, z; \boldsymbol{\theta})} \right] \otimes \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\frac{\nabla_{\boldsymbol{\Sigma}^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta})}{p(\mathbf{x}, z; \boldsymbol{\theta})} \right]. \end{aligned} \quad (58)$$

Since we have

$$\nabla_{\Sigma^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}) = \frac{p(\mathbf{x}, z; \boldsymbol{\theta})}{2} (\Sigma - z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top),$$

and

$$\begin{aligned} \frac{\partial \text{vec}(\nabla_{\Sigma^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}))}{\partial \text{vec}(\Sigma^{-1})} &= \frac{p(\mathbf{x}, z; \boldsymbol{\theta})}{4} (\Sigma - z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top) \\ &\quad \otimes (\Sigma - z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top) - \frac{p(\mathbf{x}, z; \boldsymbol{\theta})}{2} \Sigma \otimes \Sigma, \end{aligned}$$

the expression (58) can further be derived as

$$\begin{aligned} &\mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\frac{\partial \text{vec}(\nabla_{\Sigma^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}))}{\partial \text{vec}(\Sigma^{-1})} \frac{1}{p(\mathbf{x}, z; \boldsymbol{\theta})} \right] - \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\frac{\nabla_{\Sigma^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta})}{p(\mathbf{x}, z; \boldsymbol{\theta})} \right] \otimes \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\frac{\nabla_{\Sigma^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta})}{p(\mathbf{x}, z; \boldsymbol{\theta})} \right] \\ &= \frac{1}{4} \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\text{vec}(\Sigma - z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top) \text{vec}(\Sigma - z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top)^\top \right] \\ &\quad - \frac{1}{2} \Sigma \otimes \Sigma - \frac{1}{4} \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\text{vec}(\Sigma - z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top) \right] \\ &\quad \times \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\text{vec}(\Sigma - z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top)^\top \right] \\ &= \frac{1}{4} \text{Cov}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\text{vec}(\Sigma - z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top) \right] - \frac{1}{2} \Sigma \otimes \Sigma. \end{aligned} \tag{59}$$

For the covariance term, by the law of total variance, we have

$$\begin{aligned} &\text{Cov}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\text{vec}(\Sigma - z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top) \right] \\ &= \text{Cov}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\text{vec}(z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top) \right] \\ &= \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\text{Cov}_{\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}} \left[\text{vec}(z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top) \right] \right] \\ &\quad + \text{Cov}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} \left[\mathbb{E}_{\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}} \left[\text{vec}(z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top) \right] \right]. \end{aligned}$$

The second term is equal to a zero matrix based on $\mathbb{E}_{\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}} [\mathbf{x} - \boldsymbol{\mu} - z\xi] = \mathbf{0}$. Setting $\mathbf{X} = \mathbf{x} - \boldsymbol{\mu} - z\xi$, since $\nabla_{\mathbf{x} - \boldsymbol{\mu} - z\xi}^2 \log p(\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}) = -\frac{\Sigma^{-1}}{z}$, $p(\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta})$ is a log-concave function with respect to $\mathbf{x} - \boldsymbol{\mu} - z\xi$. Based on the Brascamp–Lieb inequality in Lemma 6, we have

$$\begin{aligned} \text{Cov}_{\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}} (z^{-1} \text{vec}((\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top)) &\prec \mathbb{E}_{\mathbf{x} | z, \mathbf{y}; \boldsymbol{\theta}} [z^{-1} \text{vec}((\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top)] \otimes \Sigma \\ &= (\mathbf{I}_{d \times d} + \mathbf{K}_{d \times d}) \Sigma \otimes \Sigma, \end{aligned}$$

where $\mathbf{K}_{d \times d}$ is the commutation matrix. Hence, the expression (59) can further becomes

$$\frac{\partial \text{vec}(\nabla_{\Sigma^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}))}{\partial \text{vec}(\Sigma^{-1})} \prec (\mathbf{K}_{d \times d} - \mathbf{I}_{d \times d}) \Sigma \otimes \Sigma,$$

which is negative definite.

C. The Proof of Proposition 3

In this section, we prove the global linear convergence of our proposed ECM algorithm with respect to the three parameters: $\boldsymbol{\mu}$, Σ , and ξ . Since the ECM algorithm can always converge to the stationary points (McLachlan & Krishnan, 2008), we define $\boldsymbol{\mu}^*$, Σ^* , and ξ^* as the stationary points of the optimization problem (6). In the following, we give the convergence analysis of these three parameters with the convergence rate, respectively.

C.1. Convergence Rate of Location Parameter

Given the update rule of $\underline{\mu}$ ³ in (15), we have

$$\underline{\mu} = \frac{\sum_{t=1}^n (\mathbf{v}_t - \underline{\xi})}{\sum_{t=1}^n \iota_t} = \frac{\sum_{t=1}^n (\mathbb{E}_{\mathbf{x}_t, z_t | \mathbf{y}_t; \underline{\theta}} [z^{-1} \mathbf{x}] - \underline{\xi})}{\sum_{t=1}^n \mathbb{E}_{z_t | \mathbf{y}_t; \underline{\theta}} [z^{-1}]}. \quad (60)$$

Based on (33) and (40), the update rule of $\underline{\mu}$ in (60) becomes

$$\underline{\mu} = \underline{\mu} + \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \underline{\theta})}{p(\mathbf{y}_i; \underline{\theta})} \right)^{-1} \sum_{t=1}^n p^{-1}(\mathbf{y}_t; \underline{\theta}) \underline{\Sigma} \mathbf{q}_t \quad (61)$$

Hence, the distance between $\underline{\mu}$ at the current iteration and $\underline{\mu}^*$ is given by

$$\|\underline{\mu} - \underline{\mu}^*\|_2 = \left\| \underline{\mu} - \underline{\mu}^* + \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \underline{\theta})}{p(\mathbf{y}_i; \underline{\theta})} \right)^{-1} \sum_{t=1}^n p^{-1}(\mathbf{y}_t; \underline{\theta}) \underline{\Sigma} \mathbf{q}_t \right\|_2. \quad (62)$$

To further analyze the term on the right side in (62), we first give a result of $\mathbf{q}$.

Lemma 7. *When $\mathbf{x}$ follows a normal variance-mean mixture and $\mathbf{q}$ is defined in (23), we have*

$$\mathbf{q} = \nabla_{\underline{\mu}} p(\mathbf{y}; \underline{\theta}).$$

Proof. Consider the partial derivative of $p(\mathbf{y}; \underline{\theta})$ with respect to the i -th element of $\underline{\mu}$. We have that

$$\nabla_{\mu_i} p(\mathbf{y}; \underline{\theta}) = \nabla_{\mu_i} \int_{\mathcal{Q}^{-1}(\mathbf{y})} p(\mathbf{x}; \underline{\theta}) d\mathbf{x} = \int_{\mathcal{Q}^{-1}(\mathbf{y})} \nabla_{\mu_i} p(\mathbf{x}; \underline{\theta}) d\mathbf{x},$$

where the interchange of the integration and differentiation operations in the second equal sign is valid as the Leibniz integral rule (Casella & Berger, 2024). Since the partial derivatives of $p(\mathbf{x}; \underline{\theta})$

$$\nabla_{\mu_i} p(\mathbf{x} | z; \underline{\theta}) = -\nabla_{x_i} p(\mathbf{x} | z; \underline{\theta}),$$

we can obtain that

$$\begin{aligned} \nabla_{\mu_i} p(\mathbf{y}; \underline{\theta}) &= - \int_0^{+\infty} \int_{\mathcal{Q}^{-1}(\mathbf{y}_{\setminus i})} \int_{l_i}^{u_i} \nabla_{x_i} p(\mathbf{x} | z; \underline{\theta}) p(z) dx_i d\mathbf{x}_{\setminus i} dz \\ &= \int_0^{+\infty} \int_{\mathcal{Q}^{-1}(\mathbf{y}_{\setminus i})} p(z) (p(\mathbf{x}_{\setminus i}, x_i = l_i | z; \underline{\theta}) - p(\mathbf{x}_{\setminus i}, x_i = u_i | z; \underline{\theta})) d\mathbf{x}_{\setminus i} dz \\ &= p(x_i = l_i, \mathbf{y}_{\setminus i}, \underline{\theta}) - p(x_i = u_i, \mathbf{y}_{\setminus i}, \underline{\theta}), \end{aligned}$$

which is equivalent to the definition of the i -th element of $\mathbf{q}$ in (23).

Based on Lemma 7, we have $p^{-1}(\mathbf{y}; \underline{\theta}) \mathbf{q} = \nabla_{\underline{\mu}} \log p(\mathbf{y}; \underline{\theta})$ and $\mathbf{q}^* = \mathbf{0}$. Then the distance (62) becomes

$$\begin{aligned} \|\underline{\mu} - \underline{\mu}^*\|_2 &= \left\| \underline{\mu} - \underline{\mu}^* + \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \underline{\theta})}{p(\mathbf{y}_i; \underline{\theta})} \right)^{-1} \sum_{t=1}^n \underline{\Sigma} \nabla_{\underline{\mu}} \log p(\mathbf{y}_t; \underline{\theta}) \right. \\ &\quad \left. - \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \underline{\theta}^*)}{p(\mathbf{y}_i; \underline{\theta}^*)} \right)^{-1} \sum_{t=1}^n \underline{\Sigma} \nabla_{\underline{\mu}^*} \log p(\mathbf{y}_t; \underline{\theta}^*) \right\|_2 \\ &= \left\| \underline{\mu} - \underline{\mu}^* + \int_0^1 \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \tilde{\underline{\theta}}_{\mu})}{p(\mathbf{y}_i; \tilde{\underline{\theta}}_{\mu})} \right)^{-1} \sum_{t=1}^n \underline{\Sigma} \nabla_{\tilde{\underline{\mu}}}^2 \log p(\mathbf{y}_t; \tilde{\underline{\theta}}_{\mu}) (\tilde{\underline{\mu}} - \underline{\mu}^*) d\beta \right\|_2 \\ &\leq \sup_{\beta \in (0,1]} \left\| \mathbf{I} + \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \tilde{\underline{\theta}}_{\mu})}{p(\mathbf{y}_i; \tilde{\underline{\theta}}_{\mu})} \right)^{-1} \sum_{t=1}^n \underline{\Sigma} \nabla_{\tilde{\underline{\mu}}}^2 \log p(\mathbf{y}_t; \tilde{\underline{\theta}}_{\mu}) \right\|_2 \|\underline{\mu} - \underline{\mu}^*\|_2, \end{aligned}$$

³For simplicity, we use $\underline{\mu}$, $\underline{\xi}$, and $\underline{\Sigma}$ to denote the parameters $\underline{\mu}^{(k)}$, $\underline{\xi}^{(k)}$, and $\underline{\Sigma}^{(k)}$ before the k -th update, and use $\underline{\mu}$, $\underline{\xi}$, and $\underline{\Sigma}$ to denote the updated parameters $\underline{\mu}^{(k+1)}$, $\underline{\xi}^{(k+1)}$, and $\underline{\Sigma}^{(k+1)}$.

where the second equation is given by the mean value theorem of integrals and $\tilde{\theta}_\mu = \{\tilde{\mu}, \underline{\Sigma}, \underline{\xi}\}$ with $\tilde{\mu} = \beta \underline{\mu} + (1 - \beta) \mu^*$ and $\beta \in (0, 1]$.

To bound the term $\sup_{\beta \in (0, 1]} \left\| \mathbf{I} + \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \tilde{\theta}_\mu)}{p(\mathbf{y}_i; \tilde{\theta}_\mu)} \right)^{-1} \sum_{t=1}^n \underline{\Sigma} \nabla_{\tilde{\mu}}^2 \log p(\mathbf{y}_t; \tilde{\theta}_\mu) \right\|_2$, we have a result.

Lemma 8. *The second order derivative of $\log p(\mathbf{y}; \theta)$ with respect to μ , i.e., $\nabla_{\mu}^2 \log p(\mathbf{y}; \theta)$, satisfy that $\mathbf{I} + \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \theta)}{p(\mathbf{y}_i; \theta)} \right)^{-1} \sum_{t=1}^n \underline{\Sigma} \nabla_{\mu}^2 \log p(\mathbf{y}_t; \theta)$ is a positive definite matrix, for all θ in the feasible set.*

Proof. Substituting $\nabla_{\mu}^2 \log p(\mathbf{y}; \theta)$ in (48) into $\mathbf{I} + \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \theta)}{p(\mathbf{y}_i; \theta)} \right)^{-1} \sum_{t=1}^n \underline{\Sigma} \nabla_{\mu}^2 \log p(\mathbf{y}_t; \theta)$, we have

$$\mathbf{I} + \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \theta)}{p(\mathbf{y}_i; \theta)} \right)^{-1} \sum_{t=1}^n \underline{\Sigma} \nabla_{\mu}^2 \log p(\mathbf{y}_t; \theta) = \frac{\sum_{t=1}^n \underline{\Sigma}^{-1} \text{Cov}_{\mathbf{x}, z | \mathbf{y}_t; \theta} [z^{-1} \mathbf{x}]}{\sum_{t=1}^n \mathbb{E}_{z | \mathbf{y}_t; \theta} [z^{-1}]}.$$

Since both $\underline{\Sigma}$ and $\text{Cov}_{\mathbf{x}, z | \mathbf{y}; \theta} [z^{-1} \mathbf{x}]$ are positive definite matrices, $\mathbf{I} + \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \theta)}{p(\mathbf{y}_i; \theta)} \right)^{-1} \sum_{t=1}^n \underline{\Sigma} \nabla_{\mu}^2 \log p(\mathbf{y}_t; \theta)$ is positive definite. $\square$

Based on Theorem 2 and Lemma 8, we have that all the eigenvalues of matrix $\mathbf{I} + \underline{\Sigma} \nabla_{\mu}^2 \log p(\mathbf{y}; \theta)$ for any θ are belongs to $(0, 1)$, hence we have

$$\begin{aligned} \sup_{\beta \in (0, 1]} \left\| \mathbf{I} + \left(\sum_{i=1}^n \frac{p_{-1}(\mathbf{y}_i; \tilde{\theta}_\mu)}{p(\mathbf{y}_i; \tilde{\theta}_\mu)} \right)^{-1} \sum_{t=1}^n \underline{\Sigma} \nabla_{\tilde{\mu}}^2 \log p(\mathbf{y}_t; \tilde{\theta}_\mu) \right\|_2 &= \sup_{\beta \in (0, 1]} \left\| \frac{\sum_{t=1}^n \underline{\Sigma}^{-1} \text{Cov}_{\mathbf{x}, z | \mathbf{y}_t; \tilde{\theta}_\mu} [z^{-1} \mathbf{x}]}{\sum_{t=1}^n \mathbb{E}_{z | \mathbf{y}_t; \tilde{\theta}_\mu} [z^{-1}]} \right\|_2 \\ &\leq \max_{\theta} \left\| \frac{\sum_{t=1}^n \underline{\Sigma}^{-1} \text{Cov}_{\mathbf{x}, z | \mathbf{y}_t; \theta} [z^{-1} \mathbf{x}]}{\sum_{t=1}^n \mathbb{E}_{z | \mathbf{y}_t; \theta} [z^{-1}]} \right\|_2 \\ &\triangleq c_\mu \in (0, 1). \end{aligned}$$

$\square$

C.2. Convergence Rate of skewness parameter

Given the update rule of ξ in (15), we have

$$\xi = \frac{\sum_{t=1}^n (\mathbf{w}_t - \mu)}{\sum_{t=1}^n \zeta_t} = \frac{\sum_{t=1}^n (\mathbb{E}_{\mathbf{x} | \mathbf{y}_t; \theta} [\mathbf{x}] - \mu)}{\sum_{t=1}^n \mathbb{E}_{z_t | \mathbf{y}_t; \theta} [z]}. \quad (63)$$

Based on (38), we have

$$\mathbb{E}_{\mathbf{x} | \mathbf{y}; \theta} [\mathbf{x}] - \mu = \mathbb{E}_{z | \mathbf{y}; \theta} [z] \xi + p^{-1}(\mathbf{y}; \theta) \underline{\Sigma} \mathbf{q}_1.$$

Hence, the update rule of ξ in (63) becomes

$$\xi = \underline{\xi} + \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \theta)}{p(\mathbf{y}_i; \theta)} \right)^{-1} \sum_{t=1}^n p^{-1}(\mathbf{y}_t; \theta) \underline{\Sigma} \mathbf{q}_{1,t} \quad (64)$$

Hence, the distance between ξ at the current iteration and ξ^* is given by

$$\|\xi - \xi^*\|_2 = \left\| \underline{\xi} - \xi^* + \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \theta)}{p(\mathbf{y}_i; \theta)} \right)^{-1} \sum_{t=1}^n p^{-1}(\mathbf{y}_t; \theta) \underline{\Sigma} \mathbf{q}_{1,t} \right\|_2. \quad (65)$$

To further analyze the term on the right side in (65), we first give a result of $\mathbf{q}_1$.

Lemma 9. *When $\mathbf{x}$ follows a normal variance-mean mixture and $\mathbf{q}_1$ is defined in (23), we have*

$$\mathbf{q}_1 = \nabla_{\xi} p(\mathbf{y}; \theta).$$

Proof. Consider the partial derivative of $p(\mathbf{y}; \boldsymbol{\theta})$ with respect to the i -th element of $\boldsymbol{\xi}$. We have that

$$\nabla_{\xi_i} p(\mathbf{y}; \boldsymbol{\theta}) = \nabla_{\xi_i} \int_{\mathcal{Q}^{-1}(\mathbf{y})} p(\mathbf{x}; \boldsymbol{\theta}) d\mathbf{x} = \int_{\mathcal{Q}^{-1}(\mathbf{y})} \nabla_{\xi_i} p(\mathbf{x}; \boldsymbol{\theta}) d\mathbf{x}.$$

Since the partial derivatives of $p(\mathbf{x}; \boldsymbol{\theta})$

$$\nabla_{\xi_i} p(\mathbf{x} \mid z; \boldsymbol{\theta}) = -z \nabla_{x_i} p(\mathbf{x} \mid z; \boldsymbol{\theta}),$$

we can obtain that

$$\begin{aligned} \nabla_{\xi_i} p(\mathbf{y}; \boldsymbol{\theta}) &= - \int_0^{+\infty} \int_{\mathcal{Q}^{-1}(\mathbf{y}_{\setminus i})} \int_{l_i}^{u_i} z \nabla_{x_i} p(\mathbf{x} \mid z; \boldsymbol{\theta}) p(z) dx_i dz \\ &= \int_0^{+\infty} \int_{\mathcal{Q}^{-1}(\mathbf{y}_{\setminus i})} \frac{p(z)}{z} z (p(\mathbf{x}_{\setminus i}, x_i = l_i \mid z; \boldsymbol{\theta}) - p(\mathbf{x}_{\setminus i}, x_i = u_i \mid z; \boldsymbol{\theta})) d\mathbf{x}_{\setminus i} dz \\ &= p(x_i = l_i, \mathbf{y}_{\setminus i} \mid z, \boldsymbol{\theta}) - p(x_i = u_i, \mathbf{y}_{\setminus i} \mid z, \boldsymbol{\theta}). \end{aligned}$$

which is equivalent to the definition of the i -th element of $\mathbf{q}_1$ from (35).

Based on Lemma 9, we have $p^{-1}(\mathbf{y}; \boldsymbol{\theta}) \mathbf{q}_1 = \nabla_{\boldsymbol{\xi}} \log p(\mathbf{y}; \boldsymbol{\theta})$ and $\mathbf{q}_1^* = \mathbf{0}$. Then the distance (62) becomes

$$\begin{aligned} \|\boldsymbol{\xi} - \boldsymbol{\xi}^*\|_2 &= \left\| \boldsymbol{\xi} - \boldsymbol{\xi}^* + \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \boldsymbol{\theta})}{p(\mathbf{y}_i; \boldsymbol{\theta})} \right)^{-1} \sum_{t=1}^n \boldsymbol{\Sigma} \nabla_{\boldsymbol{\xi}} \log p(\mathbf{y}; \boldsymbol{\theta}) \right. \\ &\quad \left. - \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \boldsymbol{\theta}^*)}{p(\mathbf{y}_i; \boldsymbol{\theta}^*)} \right)^{-1} \sum_{t=1}^n \boldsymbol{\Sigma}^* \nabla_{\boldsymbol{\xi}^*} \log p(\mathbf{y}; \boldsymbol{\theta}^*) \right\|_2 \\ &= \left\| \boldsymbol{\xi} - \boldsymbol{\xi}^* + \int_0^1 \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \tilde{\boldsymbol{\theta}}_{\xi})}{p(\mathbf{y}_i; \tilde{\boldsymbol{\theta}}_{\xi})} \right)^{-1} \sum_{t=1}^n \tilde{\boldsymbol{\Sigma}} \nabla_{\tilde{\boldsymbol{\xi}}}^2 \log p(\mathbf{y}_t; \tilde{\boldsymbol{\theta}}_{\xi}) (\tilde{\boldsymbol{\xi}} - \boldsymbol{\xi}^*) d\beta \right\|_2 \\ &\leq \sup_{\beta \in (0,1]} \left\| \mathbf{I} + \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \tilde{\boldsymbol{\theta}}_{\xi})}{p(\mathbf{y}_i; \tilde{\boldsymbol{\theta}}_{\xi})} \right)^{-1} \sum_{t=1}^n \tilde{\boldsymbol{\Sigma}} \nabla_{\tilde{\boldsymbol{\xi}}}^2 \log p(\mathbf{y}_t; \tilde{\boldsymbol{\theta}}_{\xi}) \right\|_2 \|\boldsymbol{\xi} - \boldsymbol{\xi}^*\|_2, \end{aligned}$$

where the second equation is given by the mean value theorem of integrals and $\tilde{\boldsymbol{\theta}}_{\xi} = \{\boldsymbol{\mu}, \boldsymbol{\Sigma}, \tilde{\boldsymbol{\xi}}\}$ with $\tilde{\boldsymbol{\xi}} = \beta \boldsymbol{\xi} + (1 - \beta) \boldsymbol{\xi}^*$, and $\beta \in (0, 1]$.

To bound the term $\sup_{\beta \in (0,1]} \left\| \mathbf{I} + \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \tilde{\boldsymbol{\theta}}_{\xi})}{p(\mathbf{y}_i; \tilde{\boldsymbol{\theta}}_{\xi})} \right)^{-1} \sum_{t=1}^n \tilde{\boldsymbol{\Sigma}} \nabla_{\tilde{\boldsymbol{\xi}}}^2 \log p(\mathbf{y}_t; \tilde{\boldsymbol{\theta}}_{\xi}) \right\|_2$, we need some results for the second order derivative of $\log p(\mathbf{y}_t; \boldsymbol{\theta})$ at first.

Lemma 10. *The second order derivative of $\log p(\mathbf{y}; \boldsymbol{\theta})$ with respect to $\boldsymbol{\xi}$, i.e., $\nabla_{\boldsymbol{\xi}}^2 \log p(\mathbf{y}; \boldsymbol{\theta})$, satisfy that $\mathbf{I} + \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \boldsymbol{\theta})}{p(\mathbf{y}_i; \boldsymbol{\theta})} \right)^{-1} \sum_{t=1}^n \boldsymbol{\Sigma} \nabla_{\boldsymbol{\xi}}^2 \log p(\mathbf{y}_t; \boldsymbol{\theta})$ is a positive definite matrix.*

Proof. Substituting $\nabla_{\boldsymbol{\xi}}^2 \log p(\mathbf{y}; \boldsymbol{\theta})$ in (55) into $\mathbf{I} + \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \boldsymbol{\theta})}{p(\mathbf{y}_i; \boldsymbol{\theta})} \right)^{-1} \sum_{t=1}^n \boldsymbol{\Sigma} \nabla_{\boldsymbol{\xi}}^2 \log p(\mathbf{y}_t; \boldsymbol{\theta})$, we have

$$\mathbf{I} + \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \boldsymbol{\theta})}{p(\mathbf{y}_i; \boldsymbol{\theta})} \right)^{-1} \sum_{t=1}^n \boldsymbol{\Sigma} \nabla_{\boldsymbol{\xi}}^2 \log p(\mathbf{y}_t; \boldsymbol{\theta}) = \frac{\sum_{t=1}^n \boldsymbol{\Sigma}^{-1} \text{Cov}_{\mathbf{x}, z | \mathbf{y}_t, \boldsymbol{\theta}} [\mathbf{x} - z \boldsymbol{\xi}]}{\sum_{t=1}^n \mathbb{E}_{z | \mathbf{y}_t, \boldsymbol{\theta}} [z]}.$$

Since both $\boldsymbol{\Sigma}$ and $\text{Cov}_{\mathbf{x}, z | \mathbf{y}_t, \boldsymbol{\theta}} [\mathbf{x} - z \boldsymbol{\xi}]$ are positive definite matrices, $\mathbf{I} + \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \boldsymbol{\theta})}{p(\mathbf{y}_i; \boldsymbol{\theta})} \right)^{-1} \sum_{t=1}^n \boldsymbol{\Sigma} \nabla_{\boldsymbol{\xi}}^2 \log p(\mathbf{y}_t; \boldsymbol{\theta})$ is positive definite. $\square$

Based on Theorem 2 and Lemma 10, we have that all the eigenvalues of matrix $\mathbf{I} + \Sigma \nabla_{\xi}^2 \log p(\mathbf{y}; \theta)$ for any θ are belong to $(0, 1)$, hence we have

$$\begin{aligned} \sup_{\beta \in (0,1)} \left\| \mathbf{I} + \left(\sum_{i=1}^n \frac{p_1(\mathbf{y}_i; \tilde{\theta}_{\xi})}{p(\mathbf{y}_i; \tilde{\theta}_{\xi})} \right)^{-1} \sum_{t=1}^n \Sigma \nabla_{\xi}^2 \log p(\mathbf{y}_t; \tilde{\theta}_{\xi}) \right\|_2 &= \sup_{\beta \in (0,1)} \left\| \frac{\sum_{t=1}^n \Sigma^{-1} \text{Cov}_{\mathbf{x}, z | \mathbf{y}_t; \tilde{\theta}_{\xi}} [\mathbf{x} - z \xi]}{\sum_{t=1}^n \mathbb{E}_{z | \mathbf{y}_t; \tilde{\theta}_{\xi}} [z]} \right\|_2 \\ &\leq \max_{\theta} \left\| \frac{\sum_{t=1}^n \Sigma^{-1} \text{Cov}_{\mathbf{x}, z | \mathbf{y}_t; \theta} [\mathbf{x} - z \xi]}{\sum_{t=1}^n \mathbb{E}_{z | \mathbf{y}_t; \theta} [z]} \right\|_2 \\ &\triangleq c_{\xi} \in (0, 1). \end{aligned}$$

□

C.3. Convergence Rate of Scatter Matrix

Given the update rule of Σ in (15), we have

$$\begin{aligned} \Sigma &= \frac{1}{n} \sum_{t=1}^n \left((\mathbf{U}_t - 2\mathbf{v}_t \underline{\mu}^{\top} + \iota_t \underline{\mu} \underline{\mu}^{\top}) - 2(\mathbf{w}_t - \underline{\mu}) \underline{\xi}^{\top} + \zeta_t \underline{\xi} \underline{\xi}^{\top} \right) \\ &= \frac{1}{n} \sum_{t=1}^n \left(\mathbb{E}_{\mathbf{x}_t, z_t | \mathbf{y}_t; \underline{\theta}} [z^{-1} \mathbf{x} \mathbf{x}^{\top}] - 2\mathbb{E}_{\mathbf{x}_t, z_t | \mathbf{y}_t; \underline{\theta}} [z^{-1} \mathbf{x}_t] \underline{\mu}^{\top} + \mathbb{E}_{z_t | \mathbf{y}_t; \underline{\theta}} [z^{-1}] \underline{\mu} \underline{\mu}^{\top} - 2(\mathbb{E}_{\mathbf{x} | \mathbf{y}_t; \underline{\theta}} [\mathbf{x}] - \underline{\mu}) \underline{\xi}^{\top} + \mathbb{E}_{z_t | \mathbf{y}_t; \underline{\theta}} [z] \underline{\xi} \underline{\xi}^{\top} \right). \end{aligned} \quad (66)$$

In the following, we derive all five terms in the summation in (66).

Term 1: Based on (42), we have

$$\begin{aligned} \mathbb{E}_{\mathbf{x}, z | \mathbf{y}; \underline{\theta}} [z^{-1} \mathbf{x} \mathbf{x}^{\top}] &= \mathbb{E}_{z | \mathbf{y}; \underline{\theta}} [z^{-1}] \underline{\mu} \underline{\mu}^{\top} + \mathbb{E}_{z | \mathbf{y}; \underline{\theta}} [z] \underline{\xi} \underline{\xi}^{\top} + 2\underline{\mu} \underline{\xi}^{\top} + \underline{\Sigma} \\ &\quad + p^{-1}(\mathbf{y}; \theta) \left(2(\underline{\Sigma} \mathbf{q}) \underline{\mu}^{\top} + 2\mathbb{E}_z [z] (\underline{\Sigma} \mathbf{q}_1) \underline{\xi}^{\top} + \underline{\Sigma} (\mathbf{H}_1 + \mathbf{D}_1) \underline{\Sigma} \right), \end{aligned}$$

where $\mathbf{q}$, $\mathbf{q}_1$, $\mathbf{H}_1$ and $\mathbf{D}_1$ are defined in (23), (34), (36) and (37), respectively.

Term 2: Based on (40), we can obtain

$$-2\mathbb{E}_{\mathbf{x}_t, z_t | \mathbf{y}_t; \underline{\theta}} [z_t^{-1} \mathbf{x}_t] \underline{\mu}^{\top} = -2 \left(\mathbb{E}_{z | \mathbf{y}; \underline{\theta}} [z^{-1}] \underline{\mu} \underline{\mu}^{\top} + \underline{\xi} + p^{-1}(\mathbf{y}; \theta) \underline{\Sigma} \mathbf{q} \right) \underline{\mu}^{\top}.$$

Term 4: Based on (38), we have

$$-2(\mathbb{E}_{\mathbf{x}_t | \mathbf{y}_t; \underline{\theta}} [\mathbf{x}_t] - \underline{\mu}) \underline{\xi}^{\top} = -2\mathbb{E}_{z | \mathbf{y}; \theta} [z] \underline{\xi} \underline{\xi}^{\top} - 2p^{-1}(\mathbf{y}; \theta) (\underline{\Sigma} \mathbf{q}_1) \underline{\xi}^{\top}$$

Upon cancellation of the opposing terms, the update rule in (66) reduces to:

$$\Sigma = \underline{\Sigma} + \frac{1}{n} \sum_{t=1}^n p^{-1}(\mathbf{y}_t; \theta) \underline{\Sigma} (\mathbf{H}_{1,t} + \mathbf{D}_{1,t}) \underline{\Sigma}, \quad (67)$$

Hence, the distance between Σ at the current iteration and Σ^* is given by

$$\|\Sigma - \Sigma^*\|_{\text{F}} = \left\| \underline{\Sigma} - \Sigma^* + \frac{1}{n} \sum_{t=1}^n p^{-1}(\mathbf{y}_t; \theta) \underline{\Sigma} (\mathbf{H}_{1,t} + \mathbf{D}_{1,t}) \underline{\Sigma} \right\|_{\text{F}}. \quad (68)$$

To further analyze the term on the right side, we first give a result of $\mathbf{H}_{1,t} + \mathbf{D}_{1,t}$.

Lemma 11. When $\mathbf{x}$ follows a normal mean variance mixture and $\mathbf{H}_1$ and $\mathbf{D}_1$ are defined in (36) and (37), respectively, we have

$$\frac{1}{2}(\mathbf{H}_1 + \mathbf{D}_1) = \nabla_{\Sigma} p(\mathbf{y}; \theta).$$

Proof. Considering the partial derivative of $p(\mathbf{y}; \boldsymbol{\theta})$ with respect to the (i, j) -th element of $\boldsymbol{\Sigma}$, we have that

$$\nabla_{\boldsymbol{\Sigma}_{ij}} p(\mathbf{y}; \boldsymbol{\theta}) = \nabla_{\boldsymbol{\Sigma}_{ij}} \int_{\mathcal{Q}^{-1}(\mathbf{y})} p(\mathbf{x}; \boldsymbol{\theta}) d\mathbf{x} = \int_{\mathcal{Q}^{-1}(\mathbf{y})} \nabla_{\boldsymbol{\Sigma}_{ij}} p(\mathbf{x}; \boldsymbol{\theta}) d\mathbf{x}.$$

The partial derivatives of $p(\mathbf{x}; \boldsymbol{\theta})$ satisfy the following equation:

$$\nabla_{\boldsymbol{\Sigma}_{ij}} p(\mathbf{x} \mid z; \boldsymbol{\theta}) = \frac{z}{2} \nabla_{x_i x_j}^2 p(\mathbf{x} \mid z; \boldsymbol{\theta}), \quad (69)$$

Hence, for $i \neq j$, we can obtain that

$$\begin{aligned} \nabla_{\boldsymbol{\Sigma}_{ij}} p(\mathbf{y}; \boldsymbol{\theta}) &= \int_{\mathcal{Q}^{-1}(\mathbf{y}_{\setminus i,j})} \int_{l_i}^{u_i} \int_{l_j}^{u_j} \int_0^{+\infty} \frac{z}{2} \nabla_{x_i x_j}^2 p(\mathbf{x} \mid z; \boldsymbol{\theta}) p(z) dz dx_i dx_j d\mathbf{x}_{\setminus i,j} \\ &= \int_{\mathcal{Q}^{-1}(\mathbf{y}_{\setminus i,j})} \int_0^{+\infty} \frac{z p(z)}{2} (p(\mathbf{x}_{\setminus i,j}, x_i = l_i, x_j = l_j \mid z; \boldsymbol{\theta}) - p(\mathbf{x}_{\setminus i,j}, x_i = l_i, x_j = u_j \mid z; \boldsymbol{\theta}) \\ &\quad - p(\mathbf{x}_{\setminus i,j}, x_i = u_i, x_j = l_j \mid z; \boldsymbol{\theta}) + p(\mathbf{x}_{\setminus i,j}, x_i = u_i, x_j = u_j \mid z; \boldsymbol{\theta})) dz d\mathbf{x}_{\setminus i,j} \\ &= \frac{1}{2} \mathbb{E}_z(z) \left(p_1(\mathbf{y}_{\setminus i,j}, x_i = l_i, x_j = l_j; \boldsymbol{\theta}) - p_1(\mathbf{y}_{\setminus i,j}, x_i = l_i, x_j = u_j; \boldsymbol{\theta}) \right. \\ &\quad \left. - p_1(\mathbf{y}_{\setminus i,j}, x_i = u_i, x_j = l_j; \boldsymbol{\theta}) + p_1(\mathbf{y}_{\setminus i,j}, x_i = u_i, x_j = u_j; \boldsymbol{\theta}) \right) \\ &= \frac{1}{2} \mathbf{H}_{1,ij}. \end{aligned}$$

For the case $i = j$, we first introduce a result.

Lemma 12. *When $\mathbf{x}$ follows a normal variance-mean mixture, we have the following equation:*

$$\sum_{k=1}^d \boldsymbol{\Sigma}_{ik} \nabla_{x_k x_i}^2 p_1(\mathbf{x} \mid z; \boldsymbol{\theta}) = \nabla_{x_i} \left(-\frac{p_1(\mathbf{x} \mid z; \boldsymbol{\theta})}{z} (x_i - \mu_i - z\xi_i) \right). \quad (70)$$

Proof. We begin with

$$\sum_{k=1}^d \boldsymbol{\Sigma}_{ik} \nabla_{x_k x_i}^2 p_1(\mathbf{x} \mid z; \boldsymbol{\theta}) = \nabla_{x_i} \left(\sum_{k=1}^d \boldsymbol{\Sigma}_{ik} \nabla_{x_k} p_1(\mathbf{x} \mid z; \boldsymbol{\theta}) \right)$$

For the term $\sum_{k=1}^d \boldsymbol{\Sigma}_{ik} \nabla_{x_k} p_1(\mathbf{x} \mid z; \boldsymbol{\theta})$, we have

$$\sum_{k=1}^d \boldsymbol{\Sigma}_{ik} \nabla_{x_k} p_1(\mathbf{x} \mid z; \boldsymbol{\theta}) = \sum_{k=1}^d \boldsymbol{\Sigma}_{ik} \left(-\frac{p_1(\mathbf{x} \mid z; \boldsymbol{\theta})}{z} \sum_{l=1}^d \boldsymbol{\Sigma}_{kl}^{-1} (x_l - \mu_l - z\xi_l) \right) = -\frac{p_1(\mathbf{x} \mid z; \boldsymbol{\theta})}{z} \sum_{l=1}^d \left(\sum_{k=1}^d \boldsymbol{\Sigma}_{ik} \boldsymbol{\Sigma}_{kl}^{-1} (x_l - \mu_l - z\xi_l) \right).$$

Since $\sum_{k=1}^d \boldsymbol{\Sigma}_{ik} \boldsymbol{\Sigma}_{kl}^{-1}$ is equivalent to the (i, l) -entry of the matrix $\boldsymbol{\Sigma} \boldsymbol{\Sigma}^{-1}$, we can obtain that

$$\sum_{k=1}^d \boldsymbol{\Sigma}_{ik} \boldsymbol{\Sigma}_{kl}^{-1} = \begin{cases} 0, & i \neq l, \\ 1, & i = l. \end{cases}$$

Hence, we have

$$-\frac{p_1(\mathbf{x} \mid z; \boldsymbol{\theta})}{z} \sum_{k=1}^d \left(\sum_{l=1}^d \boldsymbol{\Sigma}_{ik} \boldsymbol{\Sigma}_{kl}^{-1} (x_l - \mu_l - z\xi_l) \right) = -\frac{p_1(\mathbf{x} \mid z; \boldsymbol{\theta})}{z} (x_i - \mu_i - z\xi_i).$$

Therefore, we prove the equation (70). $\square$

Computing the integral of the left side of (70), we have

$$\begin{aligned}
 & \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} \sum_{k=1}^d \Sigma_{ik} \nabla_{x_k x_i}^2 p_1(\mathbf{x} \mid z; \boldsymbol{\theta}) p(z) dz d\mathbf{x} \\
 &= \int_0^{+\infty} \int_{\mathcal{Q}^{-1}(\mathbf{y})} \left(\sigma_i \nabla_{x_i}^2 p_1(\mathbf{x} \mid z; \boldsymbol{\theta}) + \sum_{k \neq i} \Sigma_{ik} \nabla_{x_k x_i}^2 p_1(\mathbf{x} \mid z; \boldsymbol{\theta}) \right) p(z) dz d\mathbf{x} \\
 &= 2\mathbb{E}_z[z^{-1}] \sigma_i \nabla_{\sigma_i} p(\mathbf{y}; \boldsymbol{\theta}) + [\Sigma \mathbf{H}_1]_{ii}.
 \end{aligned} \tag{71}$$

Then the integral of the right side of (70) is given by

$$\begin{aligned}
 & \int_{\mathcal{Q}^{-1}(\mathbf{y})} \int_0^{+\infty} \nabla_{x_i} \left(-\frac{p_1(\mathbf{x}; \boldsymbol{\theta})}{z} (x_i - \mu_i - z\xi_i) \right) dz d\mathbf{x} = \int_{\mathcal{Q}^{-1}(\mathbf{y}_{\setminus i})} \int_{l_i}^{u_i} \int_0^{+\infty} \nabla_{x_i} \left(-\frac{p_1(\mathbf{x}; \boldsymbol{\theta})}{z} (x_i - \mu_i - z\xi_i) \right) dz dx_i d\mathbf{x}_{\setminus i} \\
 &= \mathbb{E}_z(z^{-1}) \left((l_i - \mu_i) p(x_i = l_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) - \xi_i p_1(x_i = l_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) \right. \\
 &\quad \left. - (u_i - \mu_i) p(x_i = u_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) + \xi_i p_1(x_i = u_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) \right).
 \end{aligned} \tag{72}$$

Since (71) and (72) are equivalent, we have

$$\begin{aligned}
 \nabla_{\sigma_i} p(\mathbf{y}; \boldsymbol{\theta}) &= \frac{1}{2\sigma_i} \left((l_i - \mu_i) p(x_i = l_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) - \xi_i p_1(x_i = l_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) \right. \\
 &\quad \left. - (u_i - \mu_i) p(x_i = u_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) + \xi_i p_1(x_i = u_i, \mathbf{y}_{\setminus i}; \boldsymbol{\theta}) - \mathbb{E}_z[z] [\Sigma \mathbf{H}_1]_{ii} \right) = \frac{1}{2} \mathbf{D}_{1,ii}.
 \end{aligned}$$

Therefore, $\frac{1}{2}(\mathbf{H}_1 + \mathbf{D}_1) = \nabla_{\Sigma} p(\mathbf{y}; \boldsymbol{\theta})$ is valid. $\square$

Based on Lemma 11, we have that

$$p^{-1}(\mathbf{y}; \boldsymbol{\theta}) \Sigma (\mathbf{H}_1 + \mathbf{D}_1) \Sigma = -2 \nabla_{\Sigma^{-1}} \log p(\mathbf{y}; \boldsymbol{\theta}).$$

Hence, the distance (68) becomes

$$\begin{aligned}
 \|\Sigma - \Sigma^*\|_F &= \left\| \underline{\Sigma} - \Sigma^* - \frac{2}{n} \sum_{t=1}^n \nabla_{\underline{\Sigma}^{-1}} \log p(\mathbf{y}; \underline{\boldsymbol{\theta}}) + \frac{2}{n} \sum_{t=1}^n \nabla_{\Sigma^{*-1}} \log p(\mathbf{y}; \boldsymbol{\theta}^*) \right\|_F \\
 &= \left\| \underline{\Sigma} - \Sigma^* - \frac{2}{n} \sum_{t=1}^n \int_0^1 \nabla_{\tilde{\Sigma}^{-1}, \tilde{\Sigma}}^2 \log p(\mathbf{y}_t; \tilde{\boldsymbol{\theta}}_{\Sigma}) (\tilde{\Sigma} - \Sigma^*) d\beta \right\|_F \\
 &\leq \sup_{\beta \in (0,1]} \left\| \mathbf{I}_{d \times d} - \frac{2}{n} \sum_{t=1}^n \frac{\partial \text{vec}(\nabla_{\tilde{\Sigma}^{-1}} \log p(\mathbf{y}_t; \tilde{\boldsymbol{\theta}}_{\Sigma}))}{\partial \text{vec}(\tilde{\Sigma})} \right\|_2 \|\underline{\Sigma} - \Sigma^*\|_F.
 \end{aligned} \tag{73}$$

where the second equation is given by the mean value theorem of integrals and $\tilde{\boldsymbol{\theta}}_{\Sigma} = \{\underline{\boldsymbol{\mu}}, \tilde{\Sigma}, \underline{\boldsymbol{\xi}}\}$ with $\tilde{\Sigma} = \beta \underline{\Sigma} + (1 - \beta) \Sigma^*$ and $\beta \in (0, 1]$.

To bound the term $\sup_{\beta \in (0,1]} \left\| \mathbf{I}_{d \times d} - \frac{2}{n} \sum_{t=1}^n \frac{\partial \text{vec}(\nabla_{\tilde{\Sigma}^{-1}} \log p(\mathbf{y}_t; \tilde{\boldsymbol{\theta}}_{\Sigma}))}{\partial \text{vec}(\tilde{\Sigma})} \right\|_2$, we need some results for the term $\frac{\partial \text{vec}(\nabla_{\Sigma^{-1}} \log p(\mathbf{y}_t; \boldsymbol{\theta}))}{\partial \text{vec}(\Sigma)}$ at first.

Lemma 13. The term $\frac{\partial \text{vec}(\nabla_{\Sigma^{-1}} \log p(\mathbf{y}; \boldsymbol{\theta}))}{\partial \text{vec}(\Sigma)}$ satisfy that $\mathbf{I}_{d \times d} - \frac{2}{n} \sum_{t=1}^n \frac{\partial \text{vec}(\nabla_{\Sigma^{-1}} \log p(\mathbf{y}_t; \boldsymbol{\theta}))}{\partial \text{vec}(\Sigma)}$ is a positive definite matrix.

Proof. Based on

$$\frac{\partial \text{vec}(\nabla_{\Sigma^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}))}{\partial \text{vec}(\Sigma)} = -\frac{\partial \text{vec}(\nabla_{\Sigma^{-1}} p(\mathbf{x}, z; \boldsymbol{\theta}))}{\partial \text{vec}(\Sigma^{-1})} (\Sigma^{-1} \otimes \Sigma^{-1}),$$

we can obtain that $\frac{\partial \text{vec}(\nabla_{\Sigma^{-1}} p(\mathbf{x}; \boldsymbol{\theta}))}{\partial \text{vec}(\Sigma)}$ is positive definite and

$$\begin{aligned} & \mathbf{I}_{d \times d} - \frac{2}{n} \sum_{t=1}^n \frac{\partial \text{vec}(\nabla_{\Sigma^{-1}} p(\mathbf{y}_t; \boldsymbol{\theta}))}{\partial \text{vec}(\Sigma)} \\ &= \frac{1}{2n} \sum_{t=1}^n (\Sigma^{-1} \otimes \Sigma^{-1}) \text{Cov}_{\mathbf{x}, z | \mathbf{y}; \boldsymbol{\theta}} [\text{vec}(z^{-1}(\mathbf{x} - \boldsymbol{\mu} - z\xi)(\mathbf{x} - \boldsymbol{\mu} - z\xi)^\top)] \succ \mathbf{0}. \end{aligned}$$

□

Based on Theorem 2 and Lemma 13, we have

$$\begin{aligned} \sup_{\beta \in (0,1]} \left\| \mathbf{I}_{d \times d} - \sum_{t=1}^n \frac{\partial \text{vec}(\nabla_{\tilde{\Sigma}^{-1}} \log p(\mathbf{y}_t; \tilde{\boldsymbol{\theta}}_\Sigma))}{\partial \text{vec}(\tilde{\Sigma})} \right\|_2 &= \sup_{\beta \in (0,1]} \left\| \frac{1}{2n} \sum_{t=1}^n (\tilde{\Sigma}^{-1} \otimes \tilde{\Sigma}^{-1}) \text{Cov}_{\mathbf{x}, z | \mathbf{y}_t; \tilde{\boldsymbol{\theta}}_\Sigma} [\text{vec}(\mathbf{x}\mathbf{x}^\top)] \right\|_2 \\ &= \max_{\boldsymbol{\theta}} \left\| \frac{1}{2n} \sum_{t=1}^n (\Sigma^{-1} \otimes \Sigma^{-1}) \text{Cov}_{\mathbf{x}, z | \mathbf{y}_t; \boldsymbol{\theta}} (\text{vec}(\mathbf{x}\mathbf{x}^\top)) \right\|_2 \\ &\triangleq c_\Sigma \in (0, 1). \end{aligned}$$

D. The Proof of Theorem 4

The component-wise convergence rates c_μ , c_Σ , and c_ξ given in Proposition 3 are the upper bound for all iterates. Since all of these rates are in $(0, 1)$, there exists a constant $c_0 > 0$, we have

$$c_\mu \in (0, 1 - c_0], \quad c_\xi \in (0, 1 - c_0], \quad \text{and } c_\Sigma \in (0, 1 - c_0].$$

E. Derivations of Surrogate Functions in Quantized Matrix Completion and Compressive Sensing

E.1. Surrogate Function Derivation in Quantized Matrix Completion

The goal of the low-rank matrix completion problem is to recover an unknown low-rank matrix $\mathbf{M} \in \mathbb{R}^{d_1 \times d_2}$ from an observed, yet incomplete, matrix. Let $\mathbf{X}$ denote a matrix whose entries are drawn from a normal variance-mean mixture distribution. We use μ_{ij} and x_{ij} to represent the (i, j) -th entries of $\mathbf{M}$ and $\mathbf{X}$, respectively. Based on the normal variance-mean model, the relationship between μ_{ij} and x_{ij} is given by

$$x_{ij} = \mu_{ij} + z^{\frac{1}{2}} \sigma \epsilon,$$

where μ_{ij} serves as the location parameter of x_{ij} and both ξ and σ are given constants. In the quantization scenario, we have $\mathbf{Y} = \mathcal{Q}(\mathbf{X})$. The entire matrix $\mathbf{Y}$ is not available for observation. Let $\mathcal{O}$ denote the index set of the observed entries. Our goal is to recover $\mathbf{M}$ from incomplete $\mathbf{Y}$, which is equivalent to estimating all the μ_{ij} in the matrix $\mathbf{M}$.

Since x_{ij} follows a univariate normal variance-mean model and $\mathbf{Y} = \mathcal{Q}(\mathbf{X})$, the density function of y_{ij} is identical with (5) under the univariate case. Hence, the negative log-likelihood function for y_{ij} with $ij \in \mathcal{O}$ is

$$\sum_{(i,j) \in \mathcal{O}} \log p(y_{ij} | \mathbf{M}),$$

which is the objective function of (18). Applying Jensen's inequality as in (12), we have the surrogate function

$$\begin{aligned} S(\mathbf{M}; \underline{\mathbf{M}}) &= \sum_{(i,j) \in \mathcal{O}} \left[\mathbb{E}_{z_t | \mathbf{y}_t; \underline{\boldsymbol{\theta}}} [\log p(z_t)] - \frac{1}{2} \log \det \sigma^2 - \frac{1}{2\sigma^2} (u_{ij} - 2v_{ij}(\mu_{ij} - \mathbb{E}(z)\xi) \right. \\ &\quad \left. + \iota_{ij}(\mu_{ij} - \mathbb{E}(z)\xi)^2 - 2(w_{ij} - (\mu_{ij} - \mathbb{E}(z)\xi))\xi + \zeta_{ij}\xi^2) \right] + \text{const.}, \end{aligned}$$

where u_{ij} , v_{ij} , and w_{ij} are the univariate versions of U_{ij} , v_{ij} , and w_{ij} , respectively. Ignoring the terms which are independent with μ_{ij} , the surrogate function becomes

$$S(\mathbf{M}; \underline{\mathbf{M}}) = \frac{1}{2\sigma^2} \sum_{(i,j) \in \mathcal{O}} (2v_{ij}\mu_{ij} - \iota_{ij}\mu_{ij}^2 - 2\mu_{ij}\xi) + \text{const.}$$

Making a low-rank factorization to the matrix $\mathbf{M}$ as $\mathbf{M} = \mathbf{A}\mathbf{B}^\top$, we have $\mu_{ij} = \mathbf{a}_i \mathbf{b}_j^\top$. Setting $e_{ij} = \frac{v_{ij} - \xi}{\iota_{ij}}$, we have that maximizing $S(\mathbf{M}; \underline{\mathbf{M}})$ is equivalent to maximize

$$\sum_{(i,j) \in \mathcal{O}} \left(\mathbf{a}_i \mathbf{b}_j^\top - e_{ij} \right)^2,$$

which is identical with (19).

E.2. Surrogate Function Derivation in Quantized Compressive Sensing

In quantized compressive sensing, the base model with normal variance-mean mixture noise is given by

$$\mathbf{y} = \mathcal{Q}(\mathbf{x}), \quad \mathbf{x} = \Phi \boldsymbol{\vartheta} + z \boldsymbol{\xi} + z^{\frac{1}{2}} \sigma \boldsymbol{\epsilon}.$$

In the above model, the term $\Phi \boldsymbol{\vartheta}$ can be regarded as the location parameter of $\mathbf{x}$. We can use the ML estimation method to recover the sparse signal $\boldsymbol{\vartheta}$, which has the optimization problem

$$\max_{\boldsymbol{\vartheta}} \log p(\mathbf{y} | \boldsymbol{\vartheta}).$$

Following the suggestions from (Zymnis et al., 2009), we add a ℓ_1 -regularization term to force the solution of $\boldsymbol{\vartheta}$ to be sparse. Hence, we obtain the optimization problem

$$\max_{\boldsymbol{\vartheta}} \log p(\mathbf{y} | \boldsymbol{\vartheta}) + \eta \|\boldsymbol{\vartheta}\|_1.$$

To solve the above optimization problem through ECM algorithm, we can apply the E-step as in (12) to obtain the surrogate function

$$\begin{aligned} S(\boldsymbol{\vartheta}; \underline{\boldsymbol{\vartheta}}) &= \mathbb{E}_{z_t | \mathbf{y}_t, \underline{\boldsymbol{\vartheta}}} [\log p(z_t)] - \frac{1}{2} \log \det \sigma^2 \\ &\quad - \frac{1}{2\sigma^2} \left(\left(\mathbf{u} - 2\mathbf{v}^\top \mathbf{A} \boldsymbol{\vartheta} + \iota \boldsymbol{\vartheta}^\top \mathbf{A}^\top \mathbf{A} \boldsymbol{\vartheta} \right) - 2(\mathbf{w} - \mathbf{A} \boldsymbol{\vartheta})^\top \boldsymbol{\xi} + \zeta \boldsymbol{\xi}^\top \boldsymbol{\xi} \right) + \eta \|\boldsymbol{\vartheta}\|_1 + \text{const.} \end{aligned}$$

Ignoring the terms which are independent with $\boldsymbol{\vartheta}$, the surrogate function becomes

$$S(\boldsymbol{\vartheta}; \underline{\boldsymbol{\vartheta}}) = \frac{1}{2\sigma^2} \left(2\mathbf{v}^\top \mathbf{A} \boldsymbol{\vartheta} - \iota \boldsymbol{\vartheta}^\top \mathbf{A}^\top \mathbf{A} \boldsymbol{\vartheta} - 2\boldsymbol{\vartheta}^\top \mathbf{A}^\top \boldsymbol{\xi} \right) + \eta \|\boldsymbol{\vartheta}\|_1 + \text{const.}$$

Setting $\mathbf{e} = \frac{\mathbf{v} - \boldsymbol{\xi}}{\iota}$, we have that maximizing $S(\boldsymbol{\vartheta}; \underline{\boldsymbol{\vartheta}})$ is equivalent to maximize

$$\|\mathbf{A} \boldsymbol{\vartheta} - \mathbf{e}\|_2^2 + \eta \|\boldsymbol{\vartheta}\|_1,$$

which is identical with the surrogate function in quantized compressive sensing.

F. Experiment Details

F.1. Benchmark Settings

Quantized matrix completion: For all algorithms presented in Table 1, the complete matrix $\mathbf{M}$ is reconstructed from the training data. Denote by $\mathcal{O}_{\text{test}}$ the set of indices corresponding to entries in the test data. The definitions of the two benchmarks are given as follows.

1. Accuracy: $\frac{1}{|\mathcal{O}_{\text{test}}|} \sum_{(i,j) \in \mathcal{O}_{\text{test}}} I(Q(\mu_{ij}) = y_{ij})$;
2. RMSE: $\sqrt{\frac{1}{n_{\text{test}}} \sum_{(i,j) \in \mathcal{O}_{\text{test}}} (\mu_{ij} - y_{ij})^2}$,

where $I(\cdot)$ is the indicator function.

Quantized compressive sensing: The signal-to-noise ratio (SNR) is defined as the ratio of the expectation of the original measurements to the variance of the noise term, which is

$$\text{SNR} = \frac{\mathbb{E}[\mathbf{A}\boldsymbol{\vartheta}]}{\text{Var}[\mathbf{x} - \mathbf{A}\boldsymbol{\vartheta}]}.$$

Denote the ground truth sparse signal as $\boldsymbol{\vartheta}_{\text{gd}}$, the estimated sparse signal as $\boldsymbol{\vartheta}_{\text{est}}$. The cosine similarity is defined as

$$\text{Cos Sim} = \frac{\boldsymbol{\vartheta}_{\text{gd}}^\top \boldsymbol{\vartheta}_{\text{est}}}{\|\boldsymbol{\vartheta}_{\text{gd}}\|_2 \|\boldsymbol{\vartheta}_{\text{est}}\|_2}.$$