

Practical and Matching Gradient Variance Bounds for Black-Box Variational Bayesian Inference

Kyurae Kim¹ Kaiwen Wu¹ Jisu Oh² Jacob R. Gardner¹

Abstract

Understanding the gradient variance of black-box variational inference (BBVI) is a crucial step for establishing its convergence and developing algorithmic improvements. However, existing studies have yet to show that the gradient variance of BBVI satisfies the conditions used to study the convergence of stochastic gradient descent (SGD), the workhorse of BBVI. In this work, we show that BBVI satisfies a *matching* bound corresponding to the *ABC* condition used in the SGD literature when applied to smooth and quadratically-growing log-likelihoods. Our results generalize to nonlinear covariance parameterizations widely used in the practice of BBVI. Furthermore, we show that the variance of the mean-field parameterization has provably superior dimensional dependence.

1. Introduction

Variational inference (VI; Jordan et al. 1999; Blei et al. 2017; Zhang et al. 2019) algorithms are fast and scalable Bayesian inference methods widely applied in fields of statistics and machine learning. In particular, black-box VI (BBVI; Ranganath et al. 2014; Titsias & Lázaro-Gredilla 2014) leverages stochastic gradient descent (SGD; Robbins & Monro 1951; Bottou 1999) for inference of non-conjugate probabilistic models. With the development of bijectors (Kucukelbir et al., 2017; Dillon et al., 2017; Fjelde et al., 2020), most of the methodological advances in BBVI have now been abstracted out through various probabilistic programming frame-

works (Carpenter et al., 2017; Ge et al., 2018; Dillon et al., 2017; Bingham et al., 2019; Salvatier et al., 2016).

Despite the advances of BBVI, little is known about its theoretical properties. Even when restricted to the location-scale family (Definition 2), it is unknown whether BBVI is guaranteed to converge without having to modify the algorithms used in practice, for example, by enforcing bounded domains, bounded support, bounded gradients, and such. This theoretical insight is necessary since BBVI methods are known to be less robust (Yao et al., 2018; Dhaka et al., 2020; Welandawe et al., 2022; Dhaka et al., 2021; Domke, 2020) compared to other inference methods such as Markov chain Monte Carlo. Although progress has been made to formalize the theory of BBVI with some generality, the gap between our understanding of BBVI and the convergence guarantees of SGD remains open. For example, Domke (2019; 2020) provided smoothness and gradient variance guarantees. Still, these results do not yet yield a full convergence guarantee and do not extend to *nonlinear* covariance parameterizations used in practice.

In this work, we investigate whether recent progress in relaxing the gradient variance assumptions used in SGD (Tseng, 1998; Vaswani et al., 2019; Schmidt & Roux, 2013; Bottou et al., 2018; Gower et al., 2019; 2021b; Nguyen et al., 2018) apply to BBVI. These extensions have led to new insights that the structure of the gradient bounds can have non-trivial interactions with gradient-adaptive SGD algorithms (Zhang et al., 2022). For example, when the “interpolation assumption” (the gradient noise converges to 0; Schmidt & Roux 2013; Ma et al. 2018; Vaswani et al. 2019) does not hold, ADAM (Kingma & Ba, 2015) provably diverges with certain stepsize combinations (Zhang et al., 2022). Until BBVI can be shown to conform to the assumptions used by these recent works, it is unclear how these results relate to BBVI.

While the variance of BBVI gradient estimators has been studied before (Xu et al., 2019; Domke, 2019; Mohamed et al., 2020a; Fujisawa & Sato, 2021), the connection with the conditions used in SGD has yet to be established. As such, we answer the following question:

Does the gradient variance of BBVI conform to the conditions assumed in convergence guaran-

¹Department of Computer and Information Sciences, University of Pennsylvania, Philadelphia, Pennsylvania, United States ²Department of Statistics, North Carolina State University, Raleigh, North Carolina, United States. Correspondence to: Kyurae Kim <kyrkim@seas.upenn.edu>, Jacob R. Gardner <jrgardner@seas.upenn.edu>.

Proceedings of the 40th International Conference on Machine Learning, Honolulu, Hawaii, USA. PMLR 202, 2023. Copyright 2023 by the author(s).

tees of SGD without modifying the implementations used in practice?

The answer is yes! Assuming the target log joint distribution is smooth and quadratically growing, we show that the gradient variance of BBVI satisfies the *ABC condition* (Assumption 2) used by Polyak & Tsykin (1973); Khaled & Richtárik (2023); Gower et al. (2021b). Our analysis extends the previous result of Domke (2019) to covariance parameterizations involving nonlinear functions for conditioning the diagonal (see Section 2.5), as commonly done in practice. Furthermore, we prove that the gradient variance of the mean-field parameterization (Peterson & Anderson, 1987; Peterson & Hartman, 1989; Hinton & van Camp, 1993) results in better dimensional dependence compared to full-rank ones.

Overall, our results should act as a key ingredient to obtaining a full convergence guarantees of BBVI, as recently done by Kim et al. (2023).

Our contributions are summarized as follows:

- ① We provide upper bounds on the gradient variance of BBVI that matches the *ABC condition* (Assumption 2) used for analyzing SGD.
 - Theorems 1 and 2 do not require any modification of the algorithms used in practice.
 - Theorem 3 achieves better constants under the stronger *bounded entropy* assumption.
- ② Our analysis applies to BBVI parameterizations (Section 2.5) widely used in practice (Table 1).
 - Lemma 1 enables the bounds to cover nonlinear covariance parameterizations.
 - Lemma 3 and Remark 4 shows that the gradient variance of the mean-field parameterization has superior dimensional scaling.
- ③ We provide a matching lower bound (Theorem 4) on the gradient variance, showing that, under the stated assumptions, the *ABC condition* is the weakest assumption applicable to BBVI.

2. Preliminaries

Notation Random variables are denoted in serif, while their realization is in regular font. (i.e., x is a realization of $\mathbf{x}$, $\mathbf{x}$ is a realization of the vector-valued $\mathbf{x}$.) $\|\mathbf{x}\|_2 = \sqrt{\langle \mathbf{x}, \mathbf{x} \rangle} = \sqrt{\mathbf{x}^\top \mathbf{x}}$ denotes the Euclidean norm, while $\|\mathbf{A}\|_F = \sqrt{\text{tr}(\mathbf{A}^\top \mathbf{A})}$ is the Frobenius norm, where $\text{tr}(\mathbf{A}) = \sum_{i=1}^d A_{ii}$ is the matrix trace.

2.1. Variational Inference

Variational inference (Peterson & Anderson, 1987; Hinton & van Camp, 1993) is a family of inference

algorithms devised to solve the problem

$$\underset{\lambda \in \mathbb{R}^p}{\text{minimize}} \quad D_{\text{KL}}(q_{\psi, \lambda}, \pi), \quad (1)$$

where $q_{\psi, \lambda}$ is called the “variational approximation”, while π is a distribution of interest, and D_{KL} is the (exclusive) Kullback-Leibler (KL) divergence.

For Bayesian inference, π is the posterior distribution

$$\pi(\mathbf{z}) \propto \ell(\mathbf{x} | \mathbf{z}) p(\mathbf{z}) = \ell(\mathbf{x}, \mathbf{z}),$$

where $\ell(\mathbf{x} | \mathbf{z})$ is the likelihood, and $p(\mathbf{z})$ is the prior. In practice, one only has access to the likelihood and the prior. Thus, Equation (1) cannot be directly solved. Instead, we can minimize the negative *evidence lower bound* (ELBO; Jordan et al. 1999) function $F(\lambda)$.

Evidence Lower Bound More formally, we solve

$$\underset{\lambda \in \mathbb{R}^p}{\text{minimize}} \quad F(\lambda),$$

where F is defined as

$$F(\lambda) \triangleq -\mathbb{E}_{\mathbf{z} \sim q_{\psi, \lambda}} [\log \ell(\mathbf{x}, \mathbf{z})] - H(q_{\psi, \lambda}), \quad (2)$$

$$= -\mathbb{E}_{\mathbf{z} \sim q_{\psi, \lambda}} [\log \ell(\mathbf{x} | \mathbf{z})] + D_{\text{KL}}(q_{\psi, \lambda}, p), \quad (3)$$

$\mathbf{z}$ is the latent (random) variable,
 $q_{\psi, \lambda}$ is the variational distribution,
 ψ is a bijector (support transformation), and
 H is the differential entropy.

The bijector ψ (Dillon et al., 2017; Fjelde et al., 2020; Leger, 2023) is a differentiable bijective map that is used to de-constrain the support of constrained random variables. For example, when $\mathbf{z}$ is expected to follow a gamma distribution, using $\eta = \psi(\mathbf{z})$ with $\psi(\mathbf{z}) = \log \mathbf{z}$ lets us work with η , which can be any real number, unlike $\mathbf{z}$. The use of ψ^{-1} corresponds to the automatic differentiation VI formulation (ADVI; Kucukelbir et al. 2017), which is now widespread.

2.2. Variational Family

In this work, we specifically consider the location-scale variational family with a standardized base distribution.

Definition 1 (Reparameterization Function). An affine mapping $\mathbf{t}_\lambda : \mathbb{R}^d \rightarrow \mathbb{R}^d$ defined as

$$\mathbf{t}_\lambda(\mathbf{u}) \triangleq \mathbf{C}\mathbf{u} + \mathbf{m}$$

with λ containing the parameters for forming the location $\mathbf{m} \in \mathbb{R}^d$ and scale $\mathbf{C} = \mathbf{C}(\lambda) \in \mathbb{R}^{d \times d}$ is called the (location-scale) *reparameterization function*.

Definition 2 (Location-Scale Family). Let φ be some d -dimensional distribution. Then, q_λ such that

$$\boldsymbol{\zeta} \sim q_\lambda \Leftrightarrow \boldsymbol{\zeta} \stackrel{d}{=} \mathbf{t}_\lambda(\mathbf{u}); \quad \mathbf{u} \sim \varphi$$

is said to be a member of the location-scale family indexed by the base distribution φ and parameter λ .

This family includes commonly used variational families, such as the mean-field Gaussian, full-rank Gaussian, Student-T, and other elliptical distributions.

Remark 1 (Entropy of Location-Scale Distributions). The differential entropy of a location-scale family distribution (Definition 2) is

$$H(q_\lambda) = H(\varphi) + \log |\mathbf{C}|.$$

Definition 3 (ADVI Family; Kucukelbir et al. 2017). Let q_λ be some d -dimensional distribution. Then, $q_{\psi,\lambda}$ such that

$$\mathbf{z} \sim q_{\psi,\lambda} \Leftrightarrow \mathbf{z} \stackrel{d}{=} \psi^{-1}(\boldsymbol{\zeta}); \quad \boldsymbol{\zeta} \sim q_\lambda$$

is said to be a member of the ADVI family with the base distribution q_λ parameterized with λ .

We impose assumptions on the base distribution φ .

Assumption 1 (Base Distribution). φ is a d -dimensional distribution such that $\mathbf{u} \sim \varphi$ and $\mathbf{u} = (u_1, \dots, u_d)$ with independently and identically distributed components. Furthermore, φ is (i) symmetric and standardized such that $\mathbb{E}u_i = 0$, $\mathbb{E}u_i^2 = 1$, $\mathbb{E}u_i^3 = 0$, and (ii) has finite kurtosis $\mathbb{E}u_i^4 = \kappa < \infty$.

These assumptions are already satisfied in practice by, for example, generating u_i from a univariate normal or Student-T with $\nu > 4$ degrees of freedom.

2.3. Reparameterization Trick

When restricted to location scale families (Definitions 2 and 3), we can invoke Change-of-Variable, or more commonly known as the “reparameterization trick,” such that

$$\begin{aligned} \mathbb{E}_{\mathbf{z} \sim q_{\psi,\lambda}} \log \ell(\mathbf{x}, \mathbf{z}) &= \mathbb{E}_{\boldsymbol{\zeta} \sim q_\lambda} \log \ell(\mathbf{x}, \psi^{-1}(\boldsymbol{\zeta})) \\ &= \mathbb{E}_{\mathbf{u} \sim \varphi} \log \ell(\mathbf{x}, \psi^{-1}(\mathbf{t}_\lambda(\mathbf{u}))) \end{aligned}$$

through the Law of the Unconscious Statistician. Differentiating this results in the *reparameterization* or *path* gradient, which often achieves lower variance than alternatives (Xu et al., 2019; Mohamed et al., 2020b).

Objective Function For generality, we represent our objective as a composite infinite sum problem:

Definition 4 (Composite Infinite Sum).

$$F(\lambda) = \mathbb{E}_{\mathbf{u} \sim \varphi} f(\mathbf{t}_\lambda(\mathbf{u})) + h(\lambda),$$

where $(\lambda, \mathbf{u}) \mapsto f \circ \mathbf{t}_\lambda : \mathbb{R}^p \times \mathbb{R}^d \rightarrow \mathbb{R}$ is some bivariate stochastic function of λ and the “noise source” $\mathbf{u}$, while h is a deterministic regularization term.

By appropriately defining f and h , we retrieve the two most common formulations of the ELBO in Equation (2) and Equation (3) respectively:

Definition 5 (ELBO Entropy-Regularized Form).

$$\begin{aligned} f_H(\boldsymbol{\zeta}) &= -\log \underbrace{\ell(\mathbf{x}, \psi^{-1}(\boldsymbol{\zeta}))}_{\text{Joint Likelihood}} - \log |\mathbf{J}_{\psi^{-1}}(\boldsymbol{\zeta})| \\ h_H(\lambda) &= -H(q_\lambda). \end{aligned}$$

Definition 6 (ELBO KL-Regularized Form).

$$\begin{aligned} f_{\text{KL}}(\boldsymbol{\zeta}) &= -\log \underbrace{\ell(\mathbf{x} | \psi^{-1}(\boldsymbol{\zeta}))}_{\text{Likelihood}} - \log |\mathbf{J}_{\psi^{-1}}(\boldsymbol{\zeta})| \\ h_{\text{KL}}(\lambda) &= D_{\text{KL}}(q_\lambda, p). \end{aligned}$$

Here, $\mathbf{J}_{\psi^{-1}}$ is the Jacobian of the bijector. Since $D_{\text{KL}}(q_\lambda, p)$ is seldomly available in tractable form, the entropy-regularized form is the most widely used, while the KL regularized is common for Gaussian processes and variational autoencoders.

Gradient Estimator We denote the M -sample estimator of the gradient of F as

$$\mathbf{g}_M(\lambda) \triangleq \frac{1}{M} \sum_{m=1}^M \mathbf{g}_m(\lambda), \text{ where} \quad (4)$$

$$\mathbf{g}_m(\lambda) \triangleq \nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}_m)) + \nabla h(\lambda); \quad \mathbf{u}_m \sim \varphi. \quad (5)$$

We will occasionally drop λ for clarity.

2.4. Gradient Variance Assumptions in Stochastic Gradient Descent

Gradient Variance Assumptions in SGD For a while, most convergence proofs in SGD have relied on the “bounded variance” assumption. That is, for a gradient estimator $\mathbf{g}$, $\mathbb{E}\|\mathbf{g}\|_2^2 \leq G$ for some finite constant G . This assumption is problematic because ① these types of global constants result in loose bounds, ② and it directly contradicts the strong-convexity assumption (Nguyen et al., 2018). Thus, retrieving previously known SGD convergence rates under weaker assumptions has been an important research direction (Tseng, 1998; Vaswani et al., 2019; Schmidt & Roux, 2013; Bottou et al., 2018; Gower et al., 2019; 2021b; Nguyen et al., 2018).

ABC Condition In this work, we focus on the recently rediscovered *expected smoothness*, or *ABC*, condition (Polyak & Tsykin, 1973; Gower et al., 2021b).

Assumption 2 (Expected Smoothness; ABC). $\mathbf{g}$ is said to satisfy the expected smoothness condition if

$$\mathbb{E}\|\mathbf{g}_M(\lambda)\|_2^2 \leq 2A(F(\lambda) - F^*) + B\|\nabla F(\lambda)\|_2^2 + C.$$

for some finite $A, B, C \geq 0$, where $F^* = \inf_{\lambda \in \mathbb{R}^p} F(\lambda)$.

As shown by Khaled & Richtárik (2023), this condition is not only strictly weaker than many of the previously used assumptions but also *generalizes* them by retrieving known convergence rates when tweaking the constants.

With the *ABC* condition, for non-convex L -smooth functions, under a “appropriately chosen” stepsize (otherwise the bound may blow-up as explained by [Khaled & Richtárik](#)) of $\gamma \leq 1/LB$, SGD converges to a $\mathcal{O}(LC\gamma)$ neighborhood in a $\mathcal{O}\left((1+L\gamma^2A)^T/(\gamma T)\right)$ rate. Minor variants of the *ABC* condition have also been used to prove convergence of SGD for quasr-convex functions [Gower et al. \(2021a\)](#), stochastic heavy-ball/momentum methods [Liu & Yuan \(2022\)](#), and stochastic proximal methods [Li & Milzarek, \(2022\)](#). Given the influx of results based on the *ABC* condition, connecting with it would significantly broaden our theoretical understanding of BBVI.

2.5. Covariance Parameterizations

When using the location-scale family ([Definition 2](#)), the scale matrix $\mathbf{C}$ can be parameterized in different ways. Any parameterization that results in a positive definite covariance $\mathbf{C}\mathbf{C}^\top \in \mathbb{S}_{++}^d$ is valid. We consider multiple parameterizations as the choice can result in different theoretical properties. A brief survey on the use of different parameterizations is shown in [Table 1](#).

Linear Parameterization The previous results by [Domke \(2019\)](#) considered the matrix square root parameterization, which is linear with respect to the variational parameters.

Definition 7 (Matrix Square Root).

$$\mathbf{C}(\boldsymbol{\lambda}) = \mathbf{C},$$

where $\mathbf{C} \in \mathbb{R}^{d \times d}$ is a matrix, $\boldsymbol{\lambda}_\mathbf{C} = \text{vec}(\mathbf{C}) \in \mathbb{R}^{d^2}$ such that $\boldsymbol{\lambda} = (\mathbf{m}, \boldsymbol{\lambda}_\mathbf{C})$.

Note that $\mathbf{C}$ is not constrained to be symmetric so this is not a matrix square root in a narrow sense. Also, this parameterization does not guarantee $\mathbf{C}\mathbf{C}^\top$ to be positive definite (only positive semidefinite), which occasionally results in the entropy term $h_\mathbf{H}$ blowing up ([Domke, 2020](#)). [Domke](#) proposed to fix this by using proximal operators.

Nonlinear Parameterizations In practice, optimization is preferably done in unconstrained $\mathbb{R}^P$, which then positive definiteness can be ensured by explicitly mapping the diagonal elements to positive numbers. We denote this by the *diagonal conditioner* ϕ . (See [Table 1](#) for a brief survey on their use). The following two parameterizations are commonly used, where $\mathbf{D} = \text{diag}(\phi(\mathbf{s})) \in \mathbb{R}^{d \times d}$ denotes a diagonal matrix such that $D_{ii} = \phi(s_i) > 0$.

Table 1: Survey of Parameterizations Used in Black-Box Variational Inference

Framework	Version	Parameterizations	Conditioner	Code
TURING (Ge et al., 2018)	v0.23.2	Nonlinear Mean-field	softplus	link
STAN (Carpenter et al., 2017)	v2.31.0	Nonlinear Mean-field	exp	link
		Linear Cholesky		link
PYRO (Bingham et al., 2019)	v0.10.1	Nonlinear Mean-field	softplus	link
		Linear Cholesky ¹		link
PYMC3 (Salvatier et al., 2016)	v5.0.1	Nonlinear Mean-field	softplus	link
		Nonlinear Cholesky	softplus	link
GPYTORCH (Gardner et al., 2018)	v1.9.0	Linear Cholesky		link
		Linear Mean-field		link

¹ Numpyro also provides a low-rank Cholesky parameterization, which is non-linearly conditioned. But the full-rank Cholesky is linear.

* Tensorflow probability ([Dillon et al., 2017](#)) wasn’t included as it doesn’t provide a fully pre-configured variational family (although `tfp.experimental.vi.build_*posterior` exists, the parameterization is user-supplied).

Definition 8 (Mean-Field).

$$\mathbf{C}(\boldsymbol{\lambda}, \phi) = \text{diag}(\phi(\mathbf{s})),$$

where $\mathbf{s} \in \mathbb{R}^d$ and $\boldsymbol{\lambda} = (\mathbf{m}, \mathbf{s})$.

Definition 9 (Cholesky).

$$\mathbf{C}(\boldsymbol{\lambda}, \phi) = \text{diag}(\phi(\mathbf{s})) + \mathbf{L},$$

where $\mathbf{s} \in \mathbb{R}^d$, $\mathbf{L} \in \mathbb{R}^{d \times d}$ is a strictly lower triangular matrix, $\boldsymbol{\lambda}_\mathbf{L} = \text{vec}(\mathbf{L}) \in \mathbb{R}^{(d+1)d/2}$ such that $\boldsymbol{\lambda} = (\mathbf{m}, \mathbf{s}, \boldsymbol{\lambda}_\mathbf{L})$. The special case of $\phi(x) = x$ is called the “linear Cholesky” parameterization.

Diagonal conditioner For the diagonal conditioner, the softplus function $\phi(x) = \text{softplus}(x) \triangleq \log(1 + e^x)$ ([Dugas et al., 2000](#)) or the exponential function $\phi(x) = e^x$ is commonly used. While using these nonlinear functions significantly complicates the analysis, assuming ϕ to be 1-Lipschitz retrieves practical guarantees.

Assumption 3 (Lipschitz Diagonal Conditioner). The diagonal conditioner ϕ is 1-Lipschitz continuous.

Remark 2. The softplus function is 1-Lipschitz.

3. Main Results

3.1. Key Lemmas

The main challenge in studying BBVI is that the gradient of the composed function $\nabla_{\boldsymbol{\lambda}} f(\mathbf{t}_{\boldsymbol{\lambda}}(\mathbf{u}))$ is different from ∇f . For the matrix square root parameterization, [Domke \(2019\)](#) established the connection through Lemma 1 (restated as [Lemma 6](#) in [Appendix C.1](#)). We generalize this result to nonlinear parameterizations:

Lemma 1. Let $\mathbf{t}_\lambda : \mathbb{R}^d \rightarrow \mathbb{R}^d$ be a location-scale reparameterization function (Definition 1) with some differentiable function $f : \mathbb{R}^d \rightarrow \mathbb{R}$. Then, for $\mathbf{g}_f \triangleq \nabla f(\mathbf{t}_\lambda(\mathbf{u}))$,

(i) Mean-Field

$$\|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 = \|\mathbf{g}_f\|_2^2 + \mathbf{g}_f^\top \mathbf{U} \Phi \mathbf{g}_f,$$

(ii) Cholesky

$$\|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 = \|\mathbf{g}_f\|_2^2 + \mathbf{g}_f^\top \Sigma \mathbf{g}_f + \mathbf{g}_f^\top \mathbf{U} (\Phi - \mathbf{I}) \mathbf{g}_f,$$

where $\mathbf{U}, \Phi, \Sigma$ are diagonal matrices, which the diagonals are defined as

$$U_{ii} = u_i^2, \quad \Phi_{ii} = \phi'(s_i)^2, \quad \Sigma_{ii} = \sum_{j=1}^i u_j^2,$$

and ϕ is a diagonal conditioner for the scale matrix.

Proof. See the full proof in Appendix C.2.1.

Note that the relationships in this lemma are all equalities, which can be bounded with known quantities, as done in the next lemma. We note here that if any of our analyses were to be improved, this shall be done by obtaining tighter bounds on the equalities in Lemma 1.

Lemma 2. Let $\mathbf{t}_\lambda : \mathbb{R}^d \rightarrow \mathbb{R}^d$ be a location-scale reparameterization function (Definition 1), $f : \mathbb{R}^d \rightarrow \mathbb{R}$ be a differentiable function, and let ϕ satisfy Assumption 3.

(i) Mean-Field

$$\|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 \leq (1 + \|\mathbf{U}\|_F) \|\nabla f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2,$$

where $\mathbf{U}$ is a diagonal matrix such that $U_{ii} = u_i^2$.

(ii) Cholesky

$$\|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 \leq (1 + \|\mathbf{u}\|_2^2) \|\nabla f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2,$$

where the equality holds for the matrix square root parameterization.

Proof. See the full proof in Appendix C.2.2.

Lemma 1 act as the interface between the properties of the parameterization and the likelihood f .

Remark 3 (Variance Reduction Through ϕ). A nonlinear Cholesky parameterization with a 1-Lipschitz ϕ achieves lower or equal variance compared to the matrix square root and linear Cholesky, where the equality is achieved with the matrix square root parameterization.

Dimension Dependence of Mean-Field The superior dimensional dependence of the mean-field parameterization is given by the following lemma:

Lemma 3. Let the assumptions of Lemma 2 hold and $\mathbf{u} \sim \varphi$ satisfy Assumption 1. Then, for the mean-field parameterization,

$$\begin{aligned} & \mathbb{E} \|\mathbf{t}_\lambda(\mathbf{u}) - \mathbf{z}\|_2^2 (1 + \|\mathbf{U}\|_F) \\ & \leq (\sqrt{d\kappa} + \kappa\sqrt{d} + 1) \|\mathbf{m} - \mathbf{z}\|_2^2 + (2\kappa\sqrt{d} + 1) \|\mathbf{C}\|_F^2. \end{aligned}$$

Proof. See the full proof in Appendix C.2.3.

Remark 4 (Superior Variance of Mean-Field). The mean-field parameterization has $\mathcal{O}(\sqrt{d})$ dimensional dependence compared to the $\mathcal{O}(d)$ dimensional dependence of the full-rank parameterizations in Lemma 7.

Lastly, the following lemma is the basic building block for all of our upper bounds:

Lemma 4. Let $\mathbf{g}_M$ be the M -sample gradient estimator of F (Definition 4) for some function f, h and let $\mathbf{u}$ be some random variable. Then,

$$\mathbb{E} \|\mathbf{g}_M\|_2^2 \leq \frac{1}{M} \mathbb{E} \|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 + \|\nabla F(\lambda)\|_2^2.$$

Proof. See the full proof in Appendix C.2.4.

3.2. Upper Bounds

We restrict our analysis to the class of log-likelihoods that satisfy the following conditions:

Definition 10 (L -smoothness). A function $f : \mathbb{R}^d \rightarrow \mathbb{R}$ is L -smooth if it satisfies the following for all $\zeta, \zeta' \in \mathbb{R}^d$:

$$\|\nabla f(\zeta) - \nabla f(\zeta')\|_2 \leq L \|\zeta - \zeta'\|_2.$$

Definition 11 (Quadratic Functional Growth). A function $f : \mathbb{R}^d \rightarrow \mathbb{R}$ is μ -quadratically growing if

$$\frac{\mu}{2} \|\zeta - \bar{\zeta}\|_2^2 \leq f(\zeta) - f^*$$

for all $\zeta \in \mathbb{R}^d$, where $\bar{\zeta} = \Pi_f(\zeta)$ is a projection of ζ onto the set of minimizers of f and $f^* = \inf_{\zeta \in \mathbb{R}^d} f(\zeta)$.

The quadratic growth condition has first been used by (Anitescu, 2000) and is strictly weaker than the Polyak-Łojasiewicz inequality (see Karimi et al. 2016, Appendix A for the proof). Furthermore, for μ -strongly (quasar) convex functions (Hinder et al., 2020; Jin, 2020) automatically satisfy quadratic growth, but our analysis does not require (quasar) convexity.

Both assumptions are commonly used in SGD. For studying the gradient variance of BBVI, assuming both smoothness and quadratic growth is weaker than the assumptions of Xu et al. (2019) but stronger than those of Domke (2019), who assumed only smoothness. The additional assumption on growth is necessary to extend his results to establish the ABC condition.

For the variational family, we assume the followings:

Assumption 4. $q_{\psi, \lambda}$ is a member of the ADVI family (Definition 3), where the underlying q_λ is a member of the location-scale family (Definition 2) with its base distribution φ satisfying Assumption 1.

Entropy-Regularized Form First, we provide the upper bound for the ELBO in entropy-regularized form. This result does not require any modifications to vanilla SGD.

Theorem 1. Let $\mathbf{g}_M$ be an M -sample estimate of the gradient of the ELBO in entropy regularized form (Definition 5). Also, assume that Assumption 3 and 4 hold,

- f_H is L_H -smooth, and
- f_{KL} is μ_{KL} -quadratically growing.

Then,

$$\begin{aligned} \mathbb{E}\|\mathbf{g}_M\|_2^2 &\leq \frac{4L_H^2}{\mu_{KL}M} C(d, \kappa) (F(\lambda) - F^*) + \|\nabla F(\lambda)\|_2^2 \\ &\quad + \frac{2L_H^2}{M} C(d, \kappa) \|\bar{\xi}_{KL} - \bar{\xi}_H\|_2^2 \\ &\quad + \frac{4L_H^2}{\mu_{KL}M} C(d, \kappa) (F^* - f_{KL}^*), \end{aligned}$$

where

$C(d, \kappa) = 2\kappa\sqrt{d} + 1$ for mean-field,

$C(d, \kappa) = d + \kappa$ for the Cholesky and matrix square root,

$\bar{\xi}_{KL}, \bar{\xi}_H$ are the stationary points of f_{KL}, f_H , respectively, $F^* = \inf_{\lambda \in \mathbb{R}^p} F(\lambda)$, and $f_{KL}^* = \inf_{\xi \in \mathbb{R}^d} f(\xi)$.

Proof Sketch. From Lemma 4, we can see that the key quantity of upper bounding the gradient variance is to analyze $\mathbb{E}\|\nabla_{\lambda} f_H(\mathbf{t}_{\lambda}(\mathbf{u}))\|$. The bird's eye view of the proof is as follows:

- ① The relationship between $\|\nabla_{\lambda} f_H(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2$ and $\|\nabla f_H(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2$ is established through Lemma 2.
- ② Then, the L_H -smoothness of f_H relates $\|\nabla f_H(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2$ with $\|\mathbf{t}_{\lambda}(\mathbf{u}) - \bar{\xi}_H\|_2^2$, the average squared distance from f_H 's stationary point.
- ③ The average squared distance enables the simplification of stochastic terms through Lemmas 3 and 7. This step also introduces dimension dependence.

From here, we are now left with the $\mathbb{E}\|\mathbf{t}_{\lambda}(\mathbf{u}) - \bar{\xi}_H\|_2^2$ term. One might be tempted to assume the quadratic growth assumption on f_H and proceed as

$$\mathbb{E}\|\mathbf{t}_{\lambda}(\mathbf{u}) - \bar{\xi}_H\|_2^2 \leq \frac{2}{\mu} (f_H(\mathbf{t}_{\lambda}(\mathbf{u})) - f_H^*).$$

However, for the entropy-regularized form, this soon runs into a dead end since in

$$\begin{aligned} \mathbb{E}f_H(\mathbf{t}_{\lambda}(\mathbf{u})) - f_H^* &= F(\lambda) - h(\lambda) - f^* \\ &= (F(\lambda) - F^*) + (F^* - f^*) - h_H(\lambda), \end{aligned}$$

the negative entropy term h_H is not bounded unless we rely on assumptions that need modifications to the BBVI algorithms. (e.g., bounded support, bounded domain). Fortunately, the following inequality cleverly side-steps this problem:

$$\mathbb{E}\|\mathbf{t}_{\lambda}(\mathbf{u}) - \bar{\xi}_H\|_2^2 \leq 2\mathbb{E}\|\mathbf{t}_{\lambda}(\mathbf{u}) - \bar{\xi}_{KL}\|_2^2 + 2\|\bar{\xi}_{KL} - \bar{\xi}_H\|_2^2, \quad (6)$$

albeit at the cost of some looseness. By converting the entropy-regularized form into the KL-regularized form, the regularizer term becomes $h_{KL} = D_{KL}(q_{\lambda}, p) \geq 0$, which is bounded below by definition, unlike the entropic-regularizer h_H . The proof completes by

- ④ applying the quadratic growth assumption to relate the parameter distance with the function suboptimality gap, and
- ⑤ upper bounding the KL regularizer term.

Proof. See the *full proof* in Appendix C.3.1.

Remark 5. If the bijector ψ is an identity function, $\bar{\xi}_{KL}$ and $\bar{\xi}_H$ are the maximum likelihood (ML) and maximum a-posteriori (MAP) estimates, respectively. Thus, with enough datapoints, the term $\|\bar{\xi}_{KL} - \bar{\xi}_H\|_2^2$ will be negligible since the ML and MAP estimates will be close.

Remark 6. It is also possible to tighten the constants by a factor of two. Instead of applying Equation (6), we can use the inequality

$$(a + b)^2 \leq (1 + \delta^2)a^2 + (1 + \delta^{-2})b^2,$$

for some $\delta > 0$. By setting $\delta^2 = b = \|\bar{\xi}_{KL} - \bar{\xi}_H\|_2$,

$$\mathbb{E}\|\mathbf{t}_{\lambda}(\mathbf{u}) - \bar{\xi}_H\|_2^2 \leq (1 + \delta^2) \mathbb{E}\|\mathbf{t}_{\lambda}(\mathbf{u}) - \bar{\xi}_{KL}\|_2^2 + \delta^4 + \delta^2.$$

Since $\delta \approx 0$ as explained in Remark 5, the constant in front of the first term is tightened almost by a factor of 2. However, the stated form is more convenient for theory since the first term does not depend on $\|\bar{\xi}_{KL} - \bar{\xi}_H\|_2$.

Remark 7. Let $\kappa_{\text{cond.}} = L_H/\mu_{KL}$ be the *condition number* of the problem. For the full-rank parameterizations, the variance is bounded as $\mathcal{O}(L_H\kappa_{\text{cond.}}(d + \kappa)/M)$. The variance depends linearly on

- ① the scaling of the problem L_H ,
- ② the conditioning of the problem $\kappa_{\text{cond.}}$,
- ③ the dimensionality of the problem d , and
- ④ the tail properties of the variational family κ ,

where the number of Monte Carlo samples M linearly reduces the variance.

KL-Regularized Form We now prove an equivalent result for the KL-regularized form. Here, we do not have to rely on Equation (6) since we already start from f_{KL} , which results in better constants.

Theorem 2. Let $\mathbf{g}_M$ be an M -sample estimator of the gradient of the ELBO in KL-regularized form (Definition 6). Also, assume that

- f_{KL} is L_{KL} -smooth,
- f_{KL} is μ_{KL} -quadratically growing,

and Assumption 3 and 4 hold. Then, the gradient vari-

ance is bounded above as

$$\begin{aligned} \mathbb{E}\|\mathbf{g}_M\|_2^2 &\leq \frac{2L_{\text{KL}}^2}{\mu_{\text{KL}}M}C(d, \kappa)(F(\lambda) - F^*) + \|\nabla F(\lambda)\|_2^2 \\ &\quad + \frac{2L_{\text{KL}}^2}{\mu_{\text{KL}}M}C(d, \kappa)(F^* - f_{\text{KL}}^*), \end{aligned}$$

where

$$\begin{aligned} C(d, \kappa) &= 2\kappa\sqrt{d} + 1 \text{ for mean-field,} \\ C(d, \kappa) &= d + \kappa \quad \text{for the Cholesky and matrix square root,} \\ F^* &= \inf_{\lambda \in \mathbb{R}^p} F(\lambda), \text{ and } f_{\text{KL}}^* = \inf_{\zeta \in \mathbb{R}^d} f(\zeta). \end{aligned}$$

Proof. See the [full proof](#) in [Appendix C.3.2](#).

3.3. Upper Bound Under Bounded Entropy

The bound in [Theorem 1](#) is slightly loose due to the use of [Equation \(6\)](#) and [Equation \(29\)](#). An alternative bound with slightly tighter constants, although the gains are marginal compared to [Remark 6](#), can be obtained by assuming the following:

Assumption 5 (Bounded Entropy). The regularization term is bounded below as $h_H(\lambda) \geq h_H^*$.

For the entropy-regularized form, this corresponds to the entropy being bounded above by some constant since $h(\lambda) = -H(q_\lambda)$. When using the nonlinear parameterizations ([Definitions 8](#) and [9](#)), this assumption can be practically enforced by bounding the output of ϕ by some large S .

Proposition 1. Let the diagonal conditioner ϕ be bounded as $\phi(x) \leq S$. Then, for any d -dimensional distribution q_λ in the location-scale family with the mean-field ([Definition 8](#)) or Cholesky ([Definition 9](#)) parameterizations,

$$h_H(\lambda) = -H(q_\lambda) \geq -H(\varphi) - \frac{d}{2} \log S.$$

Proof. From [Remark 1](#), $H(q_\lambda) = H(\varphi) + \log |\mathbf{C}|$. Since $\mathbf{C}$ under [Definitions 8](#) and [9](#) is a diagonal or triangular matrix, the log absolute determinant is the log sum of the diagonals. The conclusion follows from the fact that the diagonals $C_{ii} = \phi(s_i)$ are bounded by S . $\square$

This is essentially a weaker version of the bounded domain assumption, though only the diagonal elements of $\mathbf{C}$, $s_1, \dots, s_d$, are bounded. While this assumption results in an admittedly less realistic algorithm, it enables a tighter bound for the entropy-regularized form ELBO.

Theorem 3. Let $\mathbf{g}_M$ be an M -sample estimator of the gradient of the ELBO in entropy-regularized form ([Definition 5](#)). Also, assume that

- f_H is L_H -smooth,
- f_H is μ_H -quadratically growing,

- h_H is bounded as $h_H(\lambda) > h_H^*$ ([Assumption 5](#)),

and [Assumption 3](#) and [4](#) hold. Then, the gradient variance of $\mathbf{g}_M$ is bounded above as

$$\begin{aligned} \mathbb{E}\|\mathbf{g}_M\|_2^2 &\leq \frac{2L_H^2}{\mu_H M}C(d, \kappa)(F(\lambda) - F^*) + \|\nabla F(\lambda)\|_2^2 \\ &\quad + \frac{2L_H^2}{\mu_H M}C(d, \kappa)(F^* - f_H^* - h_H^*), \end{aligned}$$

where

$$\begin{aligned} C(d, \kappa) &= 2\kappa\sqrt{d} + 1 \text{ for mean-field,} \\ C(d, \kappa) &= d + \kappa \quad \text{for the Cholesky parameterization,} \\ F^* &= \inf_{\lambda \in \mathbb{R}^p} F(\lambda), \text{ and } f_H^* = \inf_{\zeta \in \mathbb{R}^d} f(\zeta). \end{aligned}$$

Proof Sketch. Instead of using [Equation \(6\)](#), we apply the quadratic assumption directly to f_H . The remaining entropic-regularizer term can now be bounded through the bounded entropy assumption.

Proof. See the [full proof](#) in [Appendix C.3.3](#).

3.4. Matching Lower Bound

Finally, we present a matching lower bound on the gradient variance of BBVI. Our lower bound holds broadly for smooth and strongly convex problem instances that are well-conditioned and high-dimensional.

Theorem 4. Let $\mathbf{g}_M$ be an M -sample estimator of the gradient of the ELBO in either the entropy- or KL-regularized form. Also, let [Assumption 4](#) hold where the matrix square root parameterization is used. Then, for all L -smooth and μ -strongly convex functions f such that $L/\mu < \sqrt{d+1}$, the variance of $\mathbf{g}_M$ is bounded below by some strictly positive constant as

$$\begin{aligned} \mathbb{E}\|\mathbf{g}_M\|_2^2 &\geq \frac{2\mu^2(d+1) - 2L^2}{ML}(F(\lambda) - F^*) + \|\nabla F(\lambda)\|_2^2 \\ &\quad + \frac{2\mu^2(d+1) - 2L^2}{ML}(\mathbb{E}f(\mathbf{t}_{\lambda^*}(\mathbf{u})) - f^*), \end{aligned}$$

as long as λ is in a local neighborhood around the unique global optimum $\lambda^* = \arg \min_{\lambda \in \mathbb{R}^p} F(\lambda)$, where $F^* = F(\lambda^*)$ and $f^* = \arg \min_{\zeta \in \mathbb{R}^d} f(\zeta)$.

Proof Sketch. We use the fact that, with the matrix square root parameterization, if f is L -smooth, $\mathbb{E}f(\mathbf{t}_\lambda(\mathbf{u}))$ is also L -smooth ([Domke, 2020](#)). From this, the parameter suboptimality can be related to the function suboptimality as

$$\|\lambda - \tilde{\lambda}\|_2^2 \geq (2/L)(\mathbb{E}f(\mathbf{t}_\lambda(\mathbf{u})) - f^*),$$

where $\tilde{\lambda} = (\tilde{\zeta}, \mathbf{0})$. For the entropy term, we circumvent the need to directly bound its value by restricting our interest to the neighborhood of the minimizer λ^* , where the contribution of $h(\lambda^*) - h(\lambda)$ will be marginal enough for the lower bound to hold.

Proof. See the [full proof](#) in [Appendix C.3.4](#).

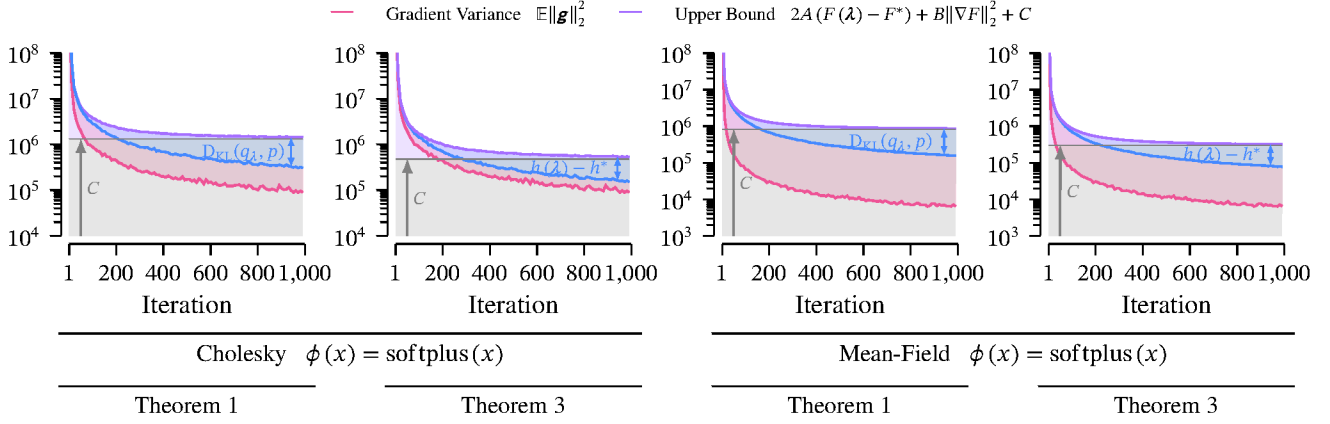


Figure 1: **Evaluation of the bounds for a perfectly conditioned quadratic target function.** The blue regions are the loosenesses resulting from either using (Theorem 1) or not using (Theorem 3) the bounded entropy assumption (Assumption 5), while the red regions are the remaining “technical loosenesses.” The gradient variance was estimated from 10^3 samples.

Remark 8 (Matching Dimensional Dependence). For well-conditioned problems such that $L/\mu < \sqrt{d} + 1$, a lower bound of the same dimensional dependence with our upper bounds holds near the optimum.

Remark 9 (Unimprovability of the ABC Condition). The lower bound suggests that the ABC gradient variance condition is unimprovable within the class of smooth, quadratically growing functions.

4. Simulations

We now evaluate our bounds and the insights gathered during the analysis through simulations. We implemented a bare-bones implementation of BBVI in Julia (Bezanson et al., 2017) with plain SGD. The stepsize were manually tuned so that all problems converge at similar speeds. For all problems, we use a unit Gaussian base distribution such that $\varphi(u) = \mathcal{N}(u; 0, 1)$ resulting in a kurtosis of $\kappa = 3$ and use $M = 10$ Monte Carlo samples.

4.1. Synthetic Problem

To test the *ideal* tightness of the bounds, we consider quadratics achieving the tightest bound for the constants $L_H, L_{KL}, \mu_H, \mu_{KL}$ given as

$$\log \ell(\mathbf{x} | \mathbf{z}) = -\frac{N}{\sigma^2} \|\mathbf{z} - \mathbf{z}^*\|_2^2, \quad \log p(\mathbf{z}) = -\frac{1}{\lambda} \|\mathbf{z}\|_2^2,$$

where N simulates the effect of the number of datapoints. We set the constants as $\sigma = 0.3$, $\lambda = 8.0$, and $N = 100$, the mode $\mathbf{z}^*$ is randomly sampled from a Gaussian, and the dimension of the problem is $d = 20$. For the bounded entropy case, we set $S = 2.0$ (the true standard deviation is in the order of $1e-3$).

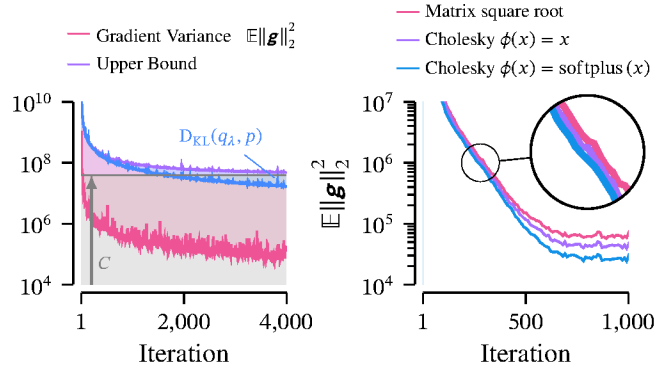


Figure 2: **Linear regression on the AIRFOIL dataset.** (left) Evaluation of the upper bound (Theorem 1). (right) Comparison of the variance of different parameterizations resulting in the same m, C .

Quality of Upper Bound The results for the Cholesky and mean-field parameterizations with a softplus bijector are shown in Figure 1. For the Cholesky parameterization, the bulk of the looseness comes from the treatment of the regularization term (blue region). The remaining “technical looseness” (red region) is relatively tight and can be shown to be tighter when using linear parameterizations ($\phi(x) = x$) and the square root parameterization, which is the tightest. However, for the mean-field parameterization, despite the superior constants (Remark 4), there is still room for improvement. Additional results for other parameterizations can be found in Appendix B.1.

4.2. Real Dataset

Model We now evaluate the theoretical results with real datasets. Given a regression dataset $(\mathbf{X}, \mathbf{y})$, we use the linear Gaussian model

$$\mathbf{y} \sim \mathcal{N}(\mathbf{X}\mathbf{w}, \sigma^2); \quad \mathbf{w} \sim \mathcal{N}(\mathbf{0}, \lambda\mathbf{I}),$$

where λ and σ are hyperparameters. The smoothness and quadratic growth constants for this model are given as the max- and minimum eigenvalues of $\sigma^{-2}\mathbf{X}^\top\mathbf{X} + \lambda^{-1}\mathbf{I}$ (for f_H) and $\sigma^{-2}\mathbf{X}^\top\mathbf{X}$ (for f_{KL}). f_{KL}^* and f_H^* are given as the mode of the likelihood and the posterior, while F^* is the negative marginal log-likelihood.

Quality of Upper Bound Section 4.1 shows the result on the AIRFOIL dataset (Dua & Graff, 2017). The constants are $L_H = 3.520 \times 10^4$, $\mu_{KL} = 2.909 \times 10^3$. Due to poor conditioning, the bound is much looser compared to the quadratic case. We note that generalizing our bounds to utilize matrix smoothness and matrix-quadratic growth as done by (Domke, 2019) would tighten the bounds. But the theoretical gains would be marginal. Detailed information about the datasets and additional results for other parameterizations can be found in Appendix B.2.

Comparison of Parameterizations Section 4.1 compares the gradient variance resulting from the different parameterizations. For a fair comparison, the gradient is estimated on the λ that results in the same $\mathbf{m}, \mathbf{C}$ for all three parameterizations. This shows the gradual increase in variance by (i) not using a nonlinear conditioner (linear Cholesky) (ii) and increasing the number of variational parameters (matrix square root).

5. Related Works

Controlling Gradient Variance The main algorithmic challenge in BBVI is to control the gradient noise (Ranganath et al., 2014). This has led to various methods for reducing the variance of VI gradient estimators using control variates (Ranganath et al., 2014; Miller et al., 2017; Geffner & Domke, 2018), ensembling of estimators (Geffner & Domke, 2020), modifying the differentiation procedure (Roeder et al., 2017), quasi-Monte Carlo (Buchholz et al., 2018; Liu & Owen, 2021), and multilevel Monte Carlo (Fujisawa & Sato, 2021). Cultivating a deeper understanding of the properties of gradient variance could further extend this list.

Convergence Guarantees Obtaining full convergence guarantees has been an important task for understanding BBVI algorithms. However, most guarantees so far have relied on strong assumptions such as that the log-likelihood is Lipschitz (Chérif-Abdellatif et al., 2019; Alquier, 2021), that the gradient variance is bounded by constant (Liu & Owen, 2021; Buchholz et al., 2018; Domke, 2020; Hoffman & Ma, 2020), and that the support of q_λ is bounded (Fujisawa & Sato, 2021). Our result shows that similar results can be obtained under relaxed assumptions. Meanwhile, Bhatia et al. (2022) have recently proven a full complexity guarantee for a variant of BBVI. But similarly to Hoffman & Ma (2020), they only optimize

the scale matrix $\mathbf{C}$, and the specifics of the algorithm diverge from the usual BBVI implementations as it uses the stochastic power iterations instead of SGD.

Gradient Variance Guarantees Studying the actual gradient variance properties of BBVI has only started to make progress recently. Fan et al. (2015) first provided bounds by assuming the log-likelihood to be Lipschitz. Under more general conditions, Domke (2019) provided tight bounds for smooth log-likelihoods, which our work builds upon. Domke’s result can also be seen as a direct generalization of the results of Xu et al. (2019), which are restricted to quadratic log-likelihoods and the mean-field family. Lastly, Mohamed et al. (2020a) provides a conceptual evaluation of gradient estimators used in BBVI.

6. Discussions

Conclusions In this work, we have proven upper bounds on the gradient variance of BBVI with the location-scale family for smooth, quadratically-growing log-likelihoods. Specifically, we have provided bounds for both the ELBO in entropy-regularized and KL-regularized forms. Our guarantees work without a single modification to the algorithms used in practice, although stronger assumptions establish a tighter bound for the entropy-regularized form ELBO. Also, our bounds corresponds to the ABC condition (Section 2.4) and the *expected residual* (ER) condition, where the latter is a special case of the former with $B = 1$. The ER condition has been used by Gower et al. (2021a) for proving convergence of SGD on quasir convex functions, which generalize convex functions. The results of this paper are used by Kim et al. (2023) to establish convergence of BBVI through the results of Khaled & Richtárik (2023).

Limitations Our results have the following limitations: ❶ Our results only apply to smooth and quadratically-growing log likelihoods and ❷ the location-scale ADVI family. Also, ❸ our bounds cannot distinguish the variance of the Cholesky and matrix square root parameterizations, ❹ and empirically, the bounds for the mean-field parameterization appear loose. Furthermore, ❺ our results only work with 1-Lipschitz diagonal conditioners such as the softplus function. Unfortunately, assuming both smoothness and quadratic growth is quite restrictive, as it leaves a very small number of known distributions. Also, in practice, non-Lipschitz conditioners such as the exponential functions are widely used. While obtaining similar bounds with such conditioners would be challenging, constructing a theoretical framework that extends to such would be an important future research direction.

Acknowledgements

This work was supported by NSF award IIS-2145644.

References

- Alquier, P. Non-Exponentially Weighted Aggregation: Regret Bounds for Unbounded Loss Functions. In *Proceedings of the International Conference on Machine Learning*, volume 193 of *PMLR*, pp. 207–218. ML Research Press, July 2021. (page 9)
- Anitescu, M. Degenerate nonlinear programming with a quadratic growth condition. *SIAM Journal on Optimization*, 10(4):1116–1135, January 2000. (page 5)
- Bezanson, J., Edelman, A., Karpinski, S., and Shah, V. B. Julia: A fresh approach to numerical computing. *SIAM review*, 59(1):65–98, 2017. (page 8)
- Bhatia, K., Kuang, N. L., Ma, Y.-A., and Wang, Y. Statistical and computational trade-offs in variational inference: A case study in inferential model selection. (arXiv:2207.11208), July 2022. (page 9)
- Bingham, E., Chen, J. P., Jankowiak, M., Obermeyer, F., Pradhan, N., Karaletsos, T., Singh, R., Szerlip, P., Horsfall, P., and Goodman, N. D. Pyro: Deep universal probabilistic programming. *Journal of Machine Learning Research*, 20(28):1–6, 2019. (pages 1, 4)
- Blei, D. M., Kucukelbir, A., and McAuliffe, J. D. Variational inference: A review for statisticians. *Journal of the American Statistical Association*, 112(518):859–877, April 2017. (page 1)
- Bottou, L. On-line learning and stochastic approximations. In *On-Line Learning in Neural Networks*, pp. 9–42. Cambridge University Press, first edition, January 1999. (page 1)
- Bottou, L., Curtis, F. E., and Nocedal, J. Optimization methods for large-scale machine learning. *SIAM Review*, 60(2):223–311, January 2018. (pages 1, 3)
- Buchholz, A., Wenzel, F., and Mandt, S. Quasi-Monte Carlo variational inference. In *Proceedings of the International Conference on Machine Learning*, volume 80 of *PMLR*, pp. 668–677. ML Research Press, July 2018. (page 9)
- Carpenter, B., Gelman, A., Hoffman, M. D., Lee, D., Goodrich, B., Betancourt, M., Brubaker, M., Guo, J., Li, P., and Riddell, A. Stan: A probabilistic programming language. *Journal of Statistical Software*, 76(1), 2017. (pages 1, 4)
- Chérif-Abdellatif, B.-E., Alquier, P., and Khan, M. E. A generalization bound for online variational inference. In *Proceedings of the Asian Conference on Machine Learning*, volume 101 of *PMLR*, pp. 662–677. ML Research Press, October 2019. (page 9)
- Dhaka, A. K., Catalina, A., Andersen, M. R., ns Magnusson, M., Huggins, J., and Vehtari, A. Robust, accurate stochastic optimization for variational inference. In *Advances in Neural Information Processing Systems*, volume 33, pp. 10961–10973. Curran Associates, Inc., 2020. (page 1)
- Dhaka, A. K., Catalina, A., Welandawe, M., Andersen, M. R., Huggins, J., and Vehtari, A. Challenges and opportunities in high dimensional variational inference. In *Advances in Neural Information Processing Systems*, volume 34, pp. 7787–7798. Curran Associates, Inc., 2021. (page 1)
- Dillon, J. V., Langmore, I., Tran, D., Brevdo, E., Vasudevan, S., Moore, D., Patton, B., Alemi, A., Hoffman, M., and Saurous, R. A. TensorFlow distributions. (arXiv:1711.10604), November 2017. (pages 1, 2, 4)
- Domke, J. Provable gradient variance guarantees for black-box variational inference. In *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. (pages 1, 2, 4, 5, 9, 13, 16, 17, 18, 21)
- Domke, J. Provable smoothness guarantees for black-box variational inference. In *Proceedings of the International Conference on Machine Learning*, volume 119 of *PMLR*, pp. 2587–2596. ML Research Press, July 2020. (pages 1, 4, 7, 9, 17, 23)
- Dua, D. and Graff, C. UCI machine learning repository. 2017. (page 9)
- Dugas, C., Bengio, Y., Bélisle, F., Nadeau, C., and Garcia, R. Incorporating second-order functional knowledge for better option pricing. In *Advances in Neural Information Processing Systems*, volume 13. MIT Press, 2000. (page 4)
- Fan, K., Wang, Z., Beck, J., Kwok, J., and Heller, K. A. Fast second order stochastic backpropagation for variational inference. In *Advances in Neural Information Processing Systems*, volume 28. Curran Associates, Inc., 2015. (page 9)
- Fjelde, T. E., Xu, K., Tarek, M., Yalburgi, S., and Ge, H. Bijectors.jl: Flexible transformations for probability distributions. In *Proceedings of The Symposium on Advances in Approximate Bayesian Inference*, volume 118 of *PMLR*, pp. 1–17. ML Research Press, February 2020. (pages 1, 2)
- Fujisawa, M. and Sato, I. Multilevel Monte Carlo variational inference. *Journal of Machine Learning Research*, 22(278):1–44, 2021. (pages 1, 9)

- Gardner, J., Pleiss, G., Weinberger, K. Q., Bindel, D., and Wilson, A. G. GPyTorch: Blackbox matrix-matrix Gaussian process inference with GPU acceleration. In *Advances in Neural Information Processing Systems*, volume 31. Curran Associates, Inc., 2018. (page 4)
- Ge, H., Xu, K., and Ghahramani, Z. Turing: A language for flexible probabilistic inference. In *Proceedings of the International Conference on Machine Learning*, volume 84 of *PMLR*, pp. 1682–1690. ML Research Press, 2018. (pages 1, 4)
- Geffner, T. and Domke, J. Using large ensembles of control variates for variational inference. In *Advances in Neural Information Processing Systems*, volume 31. Curran Associates, Inc., 2018. (page 9)
- Geffner, T. and Domke, J. A rule for gradient estimator selection, with an application to variational inference. In *Proceedings of the International Conference on Artificial Intelligence and Statistics*, volume 108 of *PMLR*, pp. 1803–1812. ML Research Press, August 2020. (page 9)
- Gower, R., Sebbouh, O., and Loizou, N. SGD for structured nonconvex functions: Learning rates, minibatching and interpolation. In *Proceedings of The International Conference on Artificial Intelligence and Statistics*, volume 130 of *PMLR*, pp. 1315–1323. ML Research Press, March 2021a. (pages 4, 9)
- Gower, R. M., Loizou, N., Qian, X., Sailanbayev, A., Shulgin, E., and Richtárik, P. SGD: General analysis and improved rates. In *Proceedings of the International Conference on Machine Learning*, volume 97 of *PMLR*, pp. 5200–5209. ML Research Press, June 2019. (pages 1, 3)
- Gower, R. M., Richtárik, P., and Bach, F. Stochastic quasi-gradient methods: Variance reduction via Jacobian sketching. *Mathematical Programming*, 188(1): 135–192, July 2021b. (pages 1, 2, 3)
- Hinder, O., Sidford, A., and Sohoni, N. Near-optimal methods for minimizing star-convex functions and beyond. In *Proceedings of Conference on Learning Theory*, volume 125 of *PMLR*, pp. 1894–1938. ML Research Press, July 2020. (page 5)
- Hinton, G. E. and van Camp, D. Keeping the neural networks simple by minimizing the description length of the weights. In *Proceedings of the Annual Conference on Computational Learning Theory*, pp. 5–13, Santa Cruz, California, United States, 1993. ACM Press. (page 2)
- Hoffman, M. and Ma, Y. Black-box variational inference as a parametric approximation to Langevin dynamics. In *Proceedings of the International Conference on Machine Learning*, PMLR, pp. 4324–4341. ML Research Press, November 2020. (page 9)
- Jin, J. On the convergence of first order methods for quasar-convex optimization. (arXiv:2010.04937), October 2020. (page 5)
- Jordan, M. I., Ghahramani, Z., Jaakkola, T. S., and Saul, L. K. An introduction to variational methods for graphical models. *Machine Learning*, 37(2):183–233, 1999. (pages 1, 2)
- Karimi, H., Nutini, J., and Schmidt, M. Linear convergence of gradient and proximal-gradient methods under the Polyak-Łojasiewicz condition. In *Machine Learning and Knowledge Discovery in Databases*, Lecture Notes in Computer Science, pp. 795–811, Cham, 2016. Springer International Publishing. (page 5)
- Khaled, A. and Richtárik, P. Better theory for SGD in the nonconvex world. *Transactions of Machine Learning Research*, 2023. (pages 2, 3, 4, 9)
- Kim, K., Wu, K., Oh, J., Ma, Y., and Gardner, J. R. Black-box variational inference converges. (arXiv:2305.15349), May 2023. (pages 2, 9)
- Kingma, D. P. and Ba, J. Adam: A Method for Stochastic Optimization. In *Proceedings of the International Conference on Learning Representations*, San Diego, California, USA, 2015. (page 1)
- Kucukelbir, A., Tran, D., Ranganath, R., Gelman, A., and Blei, D. M. Automatic differentiation variational inference. *Journal of Machine Learning Research*, 18(14): 1–45, 2017. (pages 1, 2, 3)
- Leger, J.-B. Parametrization cookbook: A set of bijective parametrizations for using machine learning methods in statistical inference. (arXiv:2301.08297), January 2023. (page 2)
- Li, X. and Milzarek, A. A unified convergence theorem for stochastic optimization methods. In *Advances in Neural Information Processing Systems*, October 2022. (page 4)
- Liu, J. and Yuan, Y. On almost sure convergence rates of stochastic gradient methods. In *Proceedings of the Conference on Learning Theory*, volume 178 of *PMLR*, pp. 2963–2983. ML Research Press, June 2022. (page 4)
- Liu, S. and Owen, A. B. Quasi-Monte Carlo quasi-Newton in Variational Bayes. *Journal of Machine Learning Research*, 22(243):1–23, 2021. (page 9)
- Ma, S., Bassily, R., and Belkin, M. The power of interpolation: Understanding the effectiveness of SGD in modern

- over-parametrized learning. In *Proceedings of the International Conference on Machine Learning*, volume 80 of *PMLR*, pp. 3325–3334. ML Research Press, July 2018. (page 1)
- Miller, A., Foti, N., D’Amour, A., and Adams, R. P. Reducing reparameterization gradient variance. In *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017. (page 9)
- Mohamed, S., Rosca, M., Figurnov, M., and Mnih, A. Monte Carlo gradient estimation in machine learning. *Journal of Machine Learning Research*, 21(132):1–62, 2020a. (pages 1, 9)
- Mohamed, S., Rosca, M., Figurnov, M., and Mnih, A. Monte Carlo gradient estimation in machine learning. *Journal of Machine Learning Research*, 21(132):1–62, 2020b. (page 3)
- Nguyen, L., Nguyen, P. H., van Dijk, M., Richtarik, P., Scheinberg, K., and Takac, M. SGD and Hogwild! Convergence without the bounded gradients assumption. In *Proceedings of the International Conference on Machine Learning*, volume 80 of *PMLR*, pp. 3750–3758. ML Research Press, July 2018. (pages 1, 3)
- Peterson, C. and Anderson, J. R. A mean field theory learning algorithm for Neural Networks. *Complex Systems*, 1(5):995–1019, 1987. (page 2)
- Peterson, C. and Hartman, E. Explorations of the mean field theory learning algorithm. *Neural Networks*, 2(6): 475–494, January 1989. (page 2)
- Polyak, B. T. and Tsytkin, Y. Z. Pseudogradient adaptation and training algorithms. *Automatic Remote Control*, 34(3):45–68, 1973. (pages 2, 3)
- Ranganath, R., Gerrish, S., and Blei, D. Black box variational inference. In *Proceedings of the International Conference on Artificial Intelligence and Statistics*, volume 33 of *PMLR*, pp. 814–822. ML Research Press, April 2014. (pages 1, 9)
- Robbins, H. and Monroe, S. A stochastic approximation method. *The Annals of Mathematical Statistics*, 22(3): 400–407, September 1951. (page 1)
- Roeder, G., Wu, Y., and Duvenaud, D. K. Sticking the landing: Simple, lower-variance gradient estimators for variational inference. In *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017. (page 9)
- Salvatier, J., Wiecki, T. V., and Fonnesbeck, C. Probabilistic programming in Python using PyMC3. *PeerJ Computer Science*, 2:e55, April 2016. (pages 1, 4)
- Schmidt, M. and Roux, N. L. Fast convergence of stochastic gradient descent under a strong growth condition. (arXiv:1308.6370), August 2013. (pages 1, 3)
- Titsias, M. and Lázaro-Gredilla, M. Doubly stochastic variational Bayes for non-conjugate inference. In *Proceedings of the International Conference on Machine Learning*, volume 32 of *PMLR*, pp. 1971–1979. ML Research Press, June 2014. (page 1)
- Tseng, P. An incremental gradient(-projection) method with momentum term and adaptive stepsize rule. *SIAM Journal on Optimization*, 8(2):506–531, May 1998. (pages 1, 3)
- Vaswani, S., Bach, F., and Schmidt, M. Fast and faster convergence of SGD for over-parameterized models and an accelerated perceptron. In *Proceedings of the International Conference on Artificial Intelligence and Statistics*, volume 89 of *PMLR*, pp. 1195–1204. ML Research Press, April 2019. (pages 1, 3)
- Welandawe, M., Andersen, M. R., Vehtari, A., and Huggins, J. H. Robust, automated, and accurate black-box variational inference. (arXiv:2203.15945), March 2022. (page 1)
- Xu, M., Quiroz, M., Kohn, R., and Sisson, S. A. Variance reduction properties of the reparameterization trick. In *Proceedings of the International Conference on Artificial Intelligence and Statistics*, volume 89 of *PMLR*, pp. 2711–2720. ML Research Press, April 2019. (pages 1, 3, 5, 9)
- Yao, Y., Vehtari, A., Simpson, D., and Gelman, A. Yes, but did it work?: Evaluating variational inference. In *Proceedings of the International Conference on Machine Learning*, PMLR, pp. 5581–5590. ML Research Press, July 2018. (page 1)
- Zhang, C., Butepage, J., Kjellstrom, H., and Mandt, S. Advances in variational inference. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 41(8):2008–2026, August 2019. (page 1)
- Zhang, Y., Chen, C., Shi, N., Sun, R., and Luo, Z.-Q. Adam can converge without any modification on update rules. In *Advances in Neural Information Processing Systems*, 2022. (page 1)

TABLE OF CONTENTS

1 Introduction	1	6 Discussions	9
2 Preliminaries	2	A Detailed Comparison Against Domke 2019	13
2.1 Variational Inference	2	B Additional Simulation Results	14
2.2 Variational Family	2	B.1 Synthetic Problem	14
2.3 Reparameterization Trick	3	B.2 Real Datasets	14
2.4 Gradient Variance Assumptions in Stochastic Gradient Descent	3	C Proofs	16
2.5 Covariance Parameterizations	4	C.1 External Lemmas	16
3 Main Results	4	C.2 Proof of Key Lemmas	17
3.1 Key Lemmas	4	C.2.1 Proof of Lemma 1	17
3.2 Upper Bounds	5	C.2.2 Proof of Lemma 2	18
3.3 Upper Bound Under Bounded Entropy	7	C.2.3 Proof of Lemma 3	19
3.4 Matching Lower Bound	7	C.2.4 Proof of Lemma 4	20
4 Simulations	8	C.3 Proof of Theorems	21
4.1 Synthetic Problem	8	C.3.1 Proof of Theorem 1	21
4.2 Real Dataset	8	C.3.2 Proof of Theorem 2	22
5 Related Works	9	C.3.3 Proof of Theorem 3	23
		C.3.4 Proof of Theorem 4	23

A. Detailed Comparison Against Domke 2019

Under the assumption that f_H is L_H -smooth and the linear full-rank Cholesky parameterization, under our notation, (Domke, 2019) prove the following bound:

$$\mathbb{E}\|\mathbf{g}_{M=1}\|_2^2 \leq L_H^2 \left((d+1)\|\mathbf{m} - \bar{\xi}_H\|_2^2 + (d+\kappa)\|\mathbf{C}\|_F^2 \right).$$

We extend Domke’s analysis in three original directions.

❶ Generalization to Nonlinear Parameterizations First, we generalize the bounds to support nonlinear parameterizations. In particular, Lemma 1 and Lemma 2 generalize Lemma 1 of Domke (2019) to 1-Lipschitz nonlinear conditioners. From here, the analysis becomes identical to Domke’s setup, until we reach our original analysis we discuss in Item ❸.

❷ Tighter Bound for the Mean-Field Parameterization Second, for the mean-field parameterization, we prove a bound that is tighter in the large d regime,

$$\mathbb{E}\|\mathbf{g}_{M=1}\|_2^2 \leq L_H^2 \left((\sqrt{d\kappa} + \kappa\sqrt{d} + 1)\|\mathbf{m} - \bar{\xi}_H\|_2^2 + (2\kappa\sqrt{d} + 1)\|\mathbf{C}\|_F^2 \right),$$

as a direct consequence of Lemma 5.

❸ Connecting with the ABC Condition Furthermore, we extend the bounds above and establish the ABC condition (Assumption 2) for the ELBO, through the quadratic function growth condition (Definition 11). Specifically, in our proof of Theorem 1, the derivation past Equation (27) is original.

B. Additional Simulation Results

B.1. Synthetic Problem

We provide additional results for the simulations with quadratics in [Section 4.1](#).

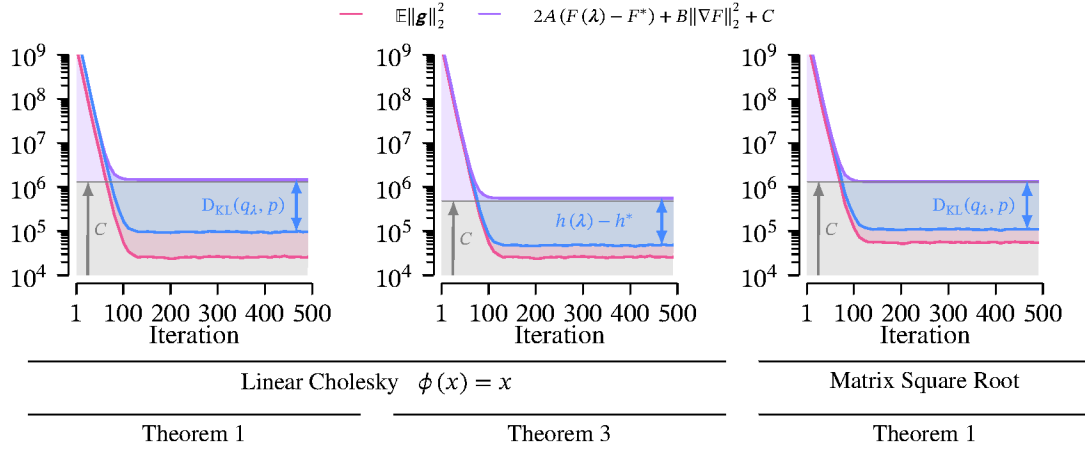


Figure 3: **Evaluation of the bounds for a perfectly conditioned quadratic target.** The blue regions are the loosenesses resulting from either using ([Theorem 1](#)) or not using ([Theorem 3](#)) the bounded entropy assumption ([Assumption 5](#)), while the red regions are the remaining “technical loosenesses.” The gradient variance was estimated from 10^3 samples.

B.2. Real Datasets

We provide detailed information and additional results for the linear regression problem in [Section 4.2](#). The constants for the linear regression datasets are shown in [Table 2](#), while additional results for the nonlinear Cholesky ([Figure 4](#)), linear Cholesky ([Figure 5](#)), nonlinear mean-field ([Figure 6](#)), and matrix square root ([Figure 7](#)) parameterizations are displayed.

Table 2: Properties of the Linear Regression Datasets

Dataset	d	N	L_H	μ_{KL}	$\kappa_{\text{cond.}}$	$\ \bar{\mathbf{S}}_{\text{KL}} - \bar{\mathbf{S}}_H\ _2^2$	Constants for Theorem 1	
							A	C
FERTILITY	9	100	1.840×10^3	5.017×10^2	4	5.167×10^{-9}	1.620×10^4	1.313×10^6
PENDULUM	9	630	1.525×10^4	1.897×10^3	8	1.243×10^{-10}	2.942×10^5	2.858×10^7
AIRFOIL	5	1,503	3.520×10^4	2.909×10^3	12	2.937×10^{-10}	6.815×10^5	3.936×10^7
WINE	11	1,599	5.526×10^4	1.786×10^3	31	6.628×10^{-9}	4.787×10^6	6.054×10^8

* N is the number of datapoints in the dataset, $\kappa_{\text{cond.}} = L_H/\mu_{\text{KL}}$ is the condition number.

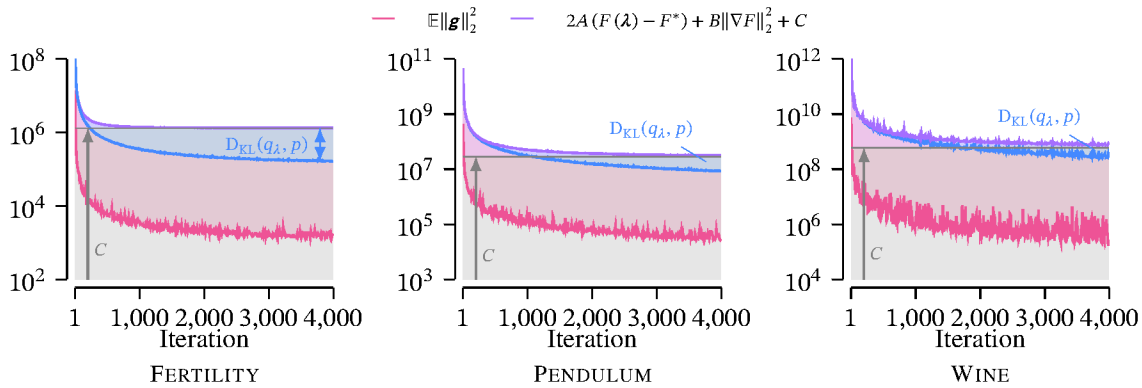


Figure 4: **Evaluation of [Theorem 1](#) with the nonlinear Cholesky ($\phi(x) = \text{softplus}(x)$) parameterization on linear regression datasets.** The gradient variance was estimated from 4×10^3 samples.

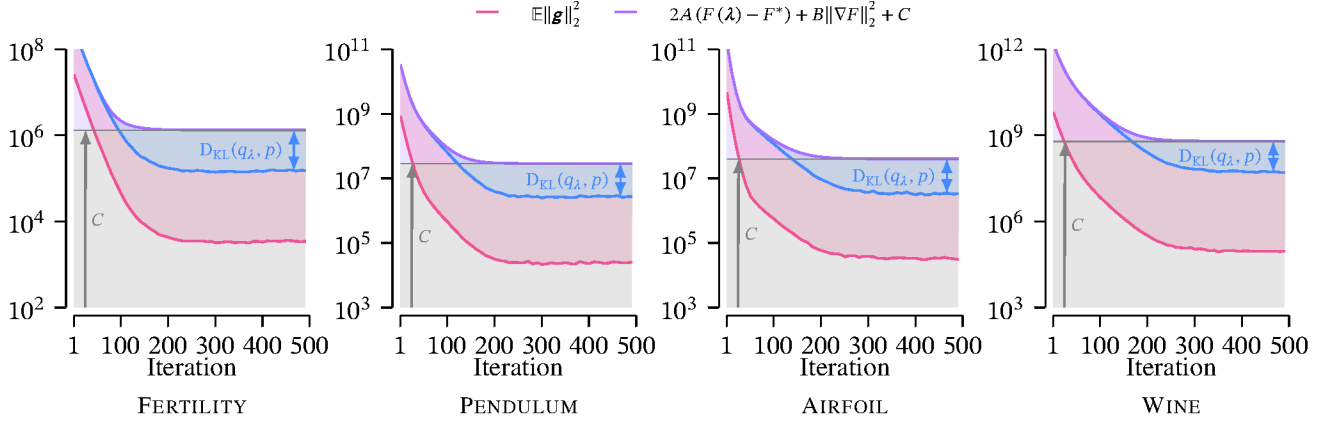


Figure 5: **Evaluation of Theorem 1 with the linear Cholesky ($\phi(x) = x$) parameterization on linear regression datasets.** The gradient variance was estimated from 4×10^3 samples.

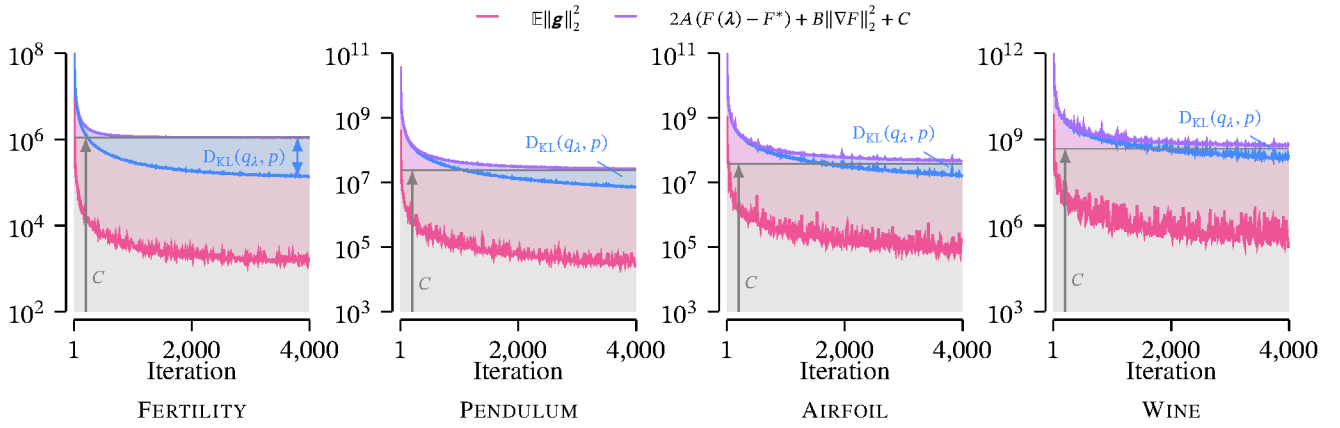


Figure 6: **Evaluation of Theorem 1 with the nonlinear mean-field ($\phi(x) = \text{softplus}(x)$) parameterization on linear regression datasets.** The gradient variance was estimated from 4×10^3 samples.

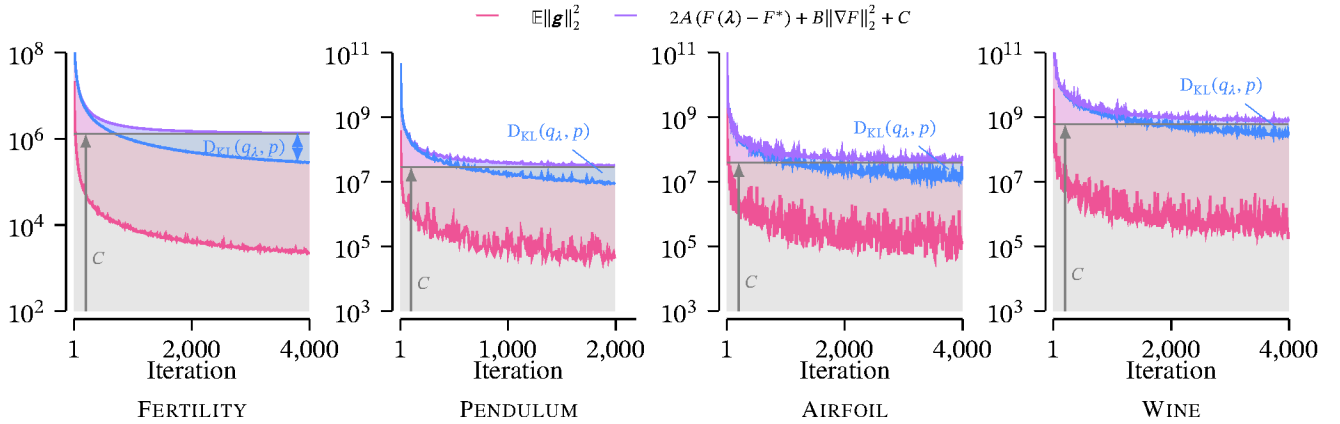


Figure 7: **Evaluation of Theorem 1 matrix square root parameterization on linear regression datasets.** The gradient variance was estimated from 4×10^3 samples.

C. Proofs

C.1. External Lemmas

Lemma 5 (Domke 2019, Lemma 9). *Let $\mathbf{u} = (u_1, u_2, \dots, u_d)$ be a d -dimensional vector-valued random variable with zero-mean independently and identically distributed components. Then,*

$$\begin{aligned}\mathbb{E}\mathbf{u}\mathbf{u}^\top &= (\mathbb{E}u_i^2) \mathbf{I} \\ \mathbb{E}\|\mathbf{u}\|_2^2 &= d \mathbb{E}u_i^2 \\ \mathbb{E}\mathbf{u} \left(1 + \|\mathbf{u}\|_2^2\right) &= (\mathbb{E}u_i^3) \mathbf{1} \\ \mathbb{E}\mathbf{u}\mathbf{u}^\top \mathbf{u}\mathbf{u}^\top &= \left((d-1) (\mathbb{E}u_i^2)^2 + \mathbb{E}u_i^4\right) \mathbf{I}.\end{aligned}$$

Lemma 6 (Domke 2019, Lemma 1). *Let $\mathbf{t}_\lambda : \mathbb{R}^d \rightarrow \mathbb{R}^d$ be a location-scale reparameterization function (Definition 1). Also, let $f : \mathbb{R}^d \mapsto \mathbb{R}$ be some differentiable function. Then,*

$$\begin{aligned}\|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 \\ = \|\nabla f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 (1 + \|\mathbf{u}\|_2^2).\end{aligned}$$

Lemma 7 (Domke 2019, Lemma 1). *Let $\mathbf{t}_\lambda : \mathbb{R}^d \rightarrow \mathbb{R}^d$ be a location-scale reparameterization function (Definition 1). Also, let $\mathbf{z} \in \mathbb{R}^d$ be some vector and $\mathbf{u} \sim \varphi$ satisfy Assumption 1. Then,*

$$\mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \mathbf{z}\|_2^2 (1 + \|\mathbf{u}\|_2^2) = (d+1) \|\mathbf{m} - \mathbf{z}\|_2^2 + (d+\kappa) \|\mathbf{C}\|_F^2.$$

Lemma 8. *Let $\mathbf{t}_\lambda : \mathbb{R}^d \rightarrow \mathbb{R}^d$ be a location-scale reparameterization function (Definition 1). Also, let $\mathbf{z} \in \mathbb{R}^d$ be some vector, and let $\mathbf{u} \sim \varphi$ satisfy Assumption 1. Then,*

$$\mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \mathbf{z}\|_2^2 = \|\mathbf{m} - \mathbf{z}\|_2^2 + \|\mathbf{C}\|_F^2.$$

Proof.

$$\begin{aligned}\mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \mathbf{z}\|_2^2 &= \mathbb{E}\|\mathbf{C}\mathbf{u} + \mathbf{m} - \mathbf{z}\|_2^2 \\ &= \mathbb{E}\mathbf{u}^\top \mathbf{C}^\top \mathbf{C} \mathbf{u} + 2 \mathbb{E}\mathbf{u}^\top \mathbf{C}^\top \mathbf{m} - 2 \mathbb{E}\mathbf{u}^\top \mathbf{C}^\top \mathbf{z} \\ &\quad + \mathbf{m}^\top \mathbf{m} - 2 \mathbf{m}^\top \mathbf{z} + \mathbf{z}^\top \mathbf{z}.\end{aligned}\tag{7}$$

The first three terms follow as

$$\begin{aligned}\mathbb{E}\mathbf{u}^\top \mathbf{C}^\top \mathbf{C} \mathbf{u} + 2 \mathbb{E}\mathbf{u}^\top \mathbf{C}^\top \mathbf{m} - 2 \mathbb{E}\mathbf{u}^\top \mathbf{C}^\top \mathbf{z} \\ = \mathbb{E}\text{tr}(\mathbf{u}^\top \mathbf{C}^\top \mathbf{C} \mathbf{u}) + 2 \mathbb{E}\mathbf{u}^\top \mathbf{C}^\top \mathbf{m} - 2 \mathbb{E}\mathbf{u}^\top \mathbf{C}^\top \mathbf{z} \\ = \text{tr}(\mathbf{C}^\top \mathbf{C} \mathbb{E}\mathbf{u}\mathbf{u}^\top) + 2 \mathbb{E}\mathbf{u}^\top \mathbf{C}^\top \mathbf{m} - 2 \mathbb{E}\mathbf{u}^\top \mathbf{C}^\top \mathbf{z},\end{aligned}$$

applying Lemma 5,

$$\begin{aligned}&= \text{tr}(\mathbf{C}^\top \mathbf{C}) \\ &= \|\mathbf{C}\|_F^2.\end{aligned}$$

Applying this to Equation (7),

$$\begin{aligned}\mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \mathbf{z}\|_2^2 &= \mathbf{m}^\top \mathbf{m} - 2 \mathbf{m}^\top \mathbf{z} + \mathbf{z}^\top \mathbf{z} + \|\mathbf{C}\|_F^2 \\ &= \|\mathbf{m} - \mathbf{z}\|_2^2 + \|\mathbf{C}\|_F^2.\end{aligned}$$

□

C.2. Proof of Key Lemmas

C.2.1. PROOF OF LEMMA 1

Lemma 1. Let $\mathbf{t}_\lambda : \mathbb{R}^d \rightarrow \mathbb{R}^d$ be a location-scale reparameterization function (Definition 1) with some differentiable function $f : \mathbb{R}^d \rightarrow \mathbb{R}$. Then, for $\mathbf{g}_f \triangleq \nabla f(\mathbf{t}_\lambda(\mathbf{u}))$,

(i) Mean-Field

$$\|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 = \|\mathbf{g}_f\|_2^2 + \mathbf{g}_f^\top \mathbf{U} \Phi \mathbf{g}_f,$$

(ii) Cholesky

$$\|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 = \|\mathbf{g}_f\|_2^2 + \mathbf{g}_f^\top \Sigma \mathbf{g}_f + \mathbf{g}_f^\top \mathbf{U} (\Phi - \mathbf{I}) \mathbf{g}_f,$$

where $\mathbf{U}, \Phi, \Sigma$ are diagonal matrices, which the diagonals are defined as

$$U_{ii} = u_i^2, \quad \Phi_{ii} = \phi'(s_i)^2, \quad \Sigma_{ii} = \sum_{j=1}^i u_j^2,$$

and ϕ is a diagonal conditioner for the scale matrix.

Proof. The proof starts by applying the Chain Rule and then computing the quadratic norm of the gradient as

$$\begin{aligned} \|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 &= \left(\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} \nabla f(\mathbf{t}_\lambda(\mathbf{u})) \right)^\top \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} \nabla f(\mathbf{t}_\lambda(\mathbf{u})) \\ &= \nabla f^\top(\mathbf{t}_\lambda(\mathbf{u})) \left(\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} \right)^\top \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} \nabla f(\mathbf{t}_\lambda(\mathbf{u})) \\ &= \mathbf{g}_f^\top \left(\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} \right)^\top \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} \mathbf{g}_f. \end{aligned} \quad (8)$$

Naturally, the derivative of the reparameterization function will depend on the specific parameterization used.

Proof for Cholesky Let p denote the number of scalar variational parameters such that $\lambda = (\lambda_1, \dots, \lambda_p)$. Then,

$$\begin{aligned} \left(\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} \right)^\top \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} &= \sum_{i=1}^d \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial m_i} \left(\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial m_i} \right)^\top + \sum_{i=1}^d \sum_{j \leq i} \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda_{C_{ij}}} \left(\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda_{C_{ij}}} \right)^\top, \end{aligned}$$

where $\lambda_{C_{ij}}$ denote the parameter responsible for the ij -th entry of $\mathbf{C}$, C_{ij} . Notice that, unlike for the matrix square root parameterization (Domke, 2019), the sum for C_{ij} is only over the lower triangular section.

For the derivatives with respect to m_i and C_{ij} , Domke (2020; 2019) show that

$$\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial m_i} = \mathbf{e}_i, \quad \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda_{C_{ij}}} = \mathbf{e}_i u_j, \quad (9)$$

where $\mathbf{e}_i$ is the unit basis of the i th component.

Therefore,

$$\begin{aligned} \left(\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} \right)^\top \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} &= \sum_{i=1}^d \mathbf{e}_i \mathbf{e}_i^\top + \sum_{i=1}^d \sum_{j \leq i} \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda_{C_{ij}}} \left(\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda_{C_{ij}}} \right)^\top \\ &= \mathbf{I} + \underbrace{\sum_{i=1}^d \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda_{C_{ii}}} \left(\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda_{C_{ii}}} \right)^\top}_{\text{diagonal of } \mathbf{C}} \\ &\quad + \underbrace{\sum_{i=1}^d \sum_{j < i} \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda_{C_{ij}}} \left(\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda_{C_{ij}}} \right)^\top}_{\text{off-diagonal of } \mathbf{C}}, \end{aligned} \quad (10)$$

leaving us with the derivatives of the scale term.

The gradient with respect to $\lambda_{C_{ij}}$, however, depends on the parameterization. That is,

$$\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda_{C_{ij}}} = \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial C_{ij}} \frac{\partial C_{ij}}{\partial \lambda_{C_{ij}}} = \mathbf{e}_i u_j \frac{\partial C_{ij}}{\partial \lambda_{C_{ij}}}. \quad (11)$$

For the diagonal elements, $\lambda_{C_{ii}} = s_i$. Thus,

$$\frac{\partial C_{ii}}{\partial s_i} = \frac{\partial \phi(s_i)}{\partial s_i} = \phi'(s_i). \quad (12)$$

And for the off-diagonal elements, $\lambda_{L_{ij}} = L_{ij}$, and

$$\frac{\partial C_{ij}}{\partial L_{ij}} = 1. \quad (13)$$

Plugging Equations (11) to (13) into Equation (10),

$$\begin{aligned} \left(\frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} \right)^\top \frac{\partial \mathbf{t}_\lambda(\mathbf{u})}{\partial \lambda} &= \mathbf{I} + \underbrace{\sum_{i=1}^d (u_i \phi'(s_i))^2 \mathbf{e}_i \mathbf{e}_i^\top}_{\text{diagonal of } \mathbf{C}} + \underbrace{\sum_{i=1}^d \sum_{j=1, j < i} u_j^2 \mathbf{e}_i \mathbf{e}_i^\top}_{\text{off-diagonal of } \mathbf{C}} \\ &= \mathbf{I} + \underbrace{\sum_{i=1}^d u_i^2 (\phi'(s_i))^2 \mathbf{e}_i \mathbf{e}_i^\top}_{\text{diagonal of } \mathbf{C}} + \underbrace{\sum_{i=1}^d \sum_{j \leq i} u_j^2 \mathbf{e}_i \mathbf{e}_i^\top - \sum_{i=1}^d u_i^2 \mathbf{e}_i \mathbf{e}_i^\top}_{\text{off-diagonal of } \mathbf{C}} \\ &= \mathbf{I} + \underbrace{\mathbf{U} \Phi}_{\text{diagonal of } \mathbf{C}} + \underbrace{\Sigma - \mathbf{U}}_{\text{off-diagonal of } \mathbf{C}} \\ &= (\mathbf{I} + \Sigma) + \mathbf{U} (\Phi - \mathbf{I}), \end{aligned} \quad (14)$$

where $\mathbf{U}, \Phi, \Sigma$ are diagonal matrices defined as

$$\begin{aligned}\Phi &= \text{diag}\left(\left[\phi'(s_1)^2, \dots, \phi'(s_d)^2\right]\right) \\ \mathbf{U} &= \text{diag}\left(\left[u_1^2, \dots, u_d^2\right]\right) \\ \Sigma &= \text{diag}\left(\left[u_1^2, u_1^2 + u_2^2, \dots, \sum_{i=1}^d u_i^2\right]\right).\end{aligned}$$

The major difference with the proof of Domke (2019, Lemma 8) for the matrix square root case is that we only sum the $u_j^2 \mathbf{e}_j \mathbf{e}_i^\top$ terms over the *lower diagonal elements*. This is the variance reduction effect we get from using the Cholesky parameterization.

Coming back to Equation (8),

$$\begin{aligned}\|\nabla_{\lambda} f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 &= \mathbf{g}_f^\top \left(\frac{\partial \mathbf{t}_{\lambda}(\mathbf{u})}{\partial \lambda} \right)^\top \frac{\partial \mathbf{t}_{\lambda}(\mathbf{u})}{\partial \lambda} \mathbf{g}_f \\ &= \mathbf{g}_f^\top \left((\mathbf{I} + \Sigma) + \mathbf{U}(\Phi - \mathbf{I}) \right) \mathbf{g}_f \\ &= \|\mathbf{g}_f\|_2^2 + \mathbf{g}_f^\top \Sigma \mathbf{g}_f + \mathbf{g}_f^\top \mathbf{U}(\Phi - \mathbf{I}) \mathbf{g}_f.\end{aligned}\quad (15)$$

Proof for Mean-field For the mean-field variational family, the covariance has only diagonal elements. Therefore, Equation (14) becomes

$$\left(\frac{\partial \mathbf{t}_{\lambda}(\mathbf{u})}{\partial \lambda} \right)^\top \frac{\partial \mathbf{t}_{\lambda}(\mathbf{u})}{\partial \lambda} = \mathbf{I} + \mathbf{U}\Phi,$$

and Equation (15) becomes

$$\|\nabla_{\lambda} f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 = \mathbf{g}_f^\top (\mathbf{I} + \mathbf{U}\Phi) \mathbf{g}_f = \|\mathbf{g}_f\|_2^2 + \mathbf{g}_f^\top \mathbf{U}\Phi \mathbf{g}_f.$$

□

C.2.2. PROOF OF LEMMA 2

Lemma 2. Let $\mathbf{t}_{\lambda} : \mathbb{R}^d \rightarrow \mathbb{R}^d$ be a location-scale reparameterization function (Definition 1), $f : \mathbb{R}^d \rightarrow \mathbb{R}$ be a differentiable function, and let ϕ satisfy Assumption 3.

(i) *Mean-Field*

$$\|\nabla_{\lambda} f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 \leq (1 + \|\mathbf{U}\|_F) \|\nabla f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2,$$

where $\mathbf{U}$ is a diagonal matrix such that $U_{ii} = u_i^2$.

(ii) *Cholesky*

$$\|\nabla_{\lambda} f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 \leq (1 + \|\mathbf{u}\|_2^2) \|\nabla f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2,$$

where the equality holds for the matrix square root parameterization.

Proof. The proof continues from the result of Lemma 1.

Proof for Cholesky Lemma 1 shows that

$$\|\nabla_{\lambda} f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 = \|\mathbf{g}_f\|_2^2 + \mathbf{g}_f^\top \Sigma \mathbf{g}_f + \mathbf{g}_f^\top \mathbf{U}(\Phi - \mathbf{I}) \mathbf{g}_f,$$

where $\mathbf{g}_f = \nabla f(\mathbf{t}_{\lambda}(\mathbf{u}))$.

By the 1-Lipschitz assumption, the entries of the diagonal matrix Φ satisfy

$$\Phi_{ii} = \phi'(d_i)^2 \leq 1,$$

which means

$$\Phi \leq \mathbf{I} \Rightarrow \mathbf{U}(\Phi - \mathbf{I}) \leq 0 \Rightarrow \mathbf{g}_f^\top \mathbf{U}(\Phi - \mathbf{I}) \mathbf{g}_f \leq 0.$$

Therefore, for the full-rank Cholesky parameterization and a 1-Lipschitz conditioner ϕ ,

$$\begin{aligned}\|\nabla_{\lambda} f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 &= \|\nabla f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 + \mathbf{g}_f^\top \Sigma \mathbf{g}_f + \mathbf{g}_f^\top \mathbf{U}(\Phi - \mathbf{I}) \mathbf{g}_f \\ &\leq \|\nabla f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 + \mathbf{g}_f^\top \Sigma \mathbf{g}_f \\ &\leq \|\nabla f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 + \|\Sigma\|_{2,2} \|\nabla f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 \\ &= \|\nabla f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 + \left(\sum_{i=1}^d u_i^2 \right) \|\nabla f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 \\ &= (1 + \|\mathbf{u}\|_2^2) \|\nabla f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2,\end{aligned}$$

where $\|\mathbf{U}\|_{2,2}$ is the L_2 operator norm of $\mathbf{U}$. This upper bound coincides with that of the matrix square root parameterization. Thus, unfortunately, this bound fails to acknowledge the lower variance of the Cholesky parameterization, coinciding with that of the matrix square root parameterization.

Proof for Mean-field (Definition 8) For the mean-field parameterization, Lemma 1 shows that

$$\|\nabla_{\lambda} f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 = \|\mathbf{g}_f\|_2^2 + \mathbf{g}_f^{\top} \mathbf{U} \Phi \mathbf{g}_f.$$

For the second term,

$$\mathbf{g}_f^{\top} \mathbf{U} \Phi \mathbf{g}_f \leq \|\mathbf{U}\|_{2,2} \|\Phi\|_{2,2} \|\mathbf{g}_f\|_2^2.$$

By the 1-Lipschitzness of ϕ ,

$$\|\Phi\|_{2,2} = \sigma_{\max}(\Phi) = \max_{i=1,\dots,d} \phi'(s_i)^2 \leq 1.$$

Then,

$$\mathbf{g}_f^{\top} (\mathbf{U} \Phi) \mathbf{g}_f \leq \|\mathbf{U}\|_{2,2} \|\mathbf{g}_f\|_2^2 \quad (16)$$

$$\leq \|\mathbf{U}\|_{\mathbb{F}} \|\mathbf{g}_f\|_2^2, \quad (17)$$

which gives the result. Here, unlike the bounds on Φ , the bounds in Equations (16) and (17) are quite loose, and become looser as the dimensionality increases.

□

C.2.3. PROOF OF LEMMA 3

Lemma 3. *Let the assumptions of Lemma 2 hold and $\mathbf{u} \sim \varphi$ satisfy Assumption 1. Then, for the mean-field parameterization,*

$$\begin{aligned} \mathbb{E} \|\mathbf{t}_{\lambda}(\mathbf{u}) - \mathbf{z}\|_2^2 (1 + \|\mathbf{U}\|_{\mathbb{F}}) \\ \leq (\sqrt{d\kappa} + \kappa\sqrt{d} + 1) \|\mathbf{m} - \mathbf{z}\|_2^2 + (2\kappa\sqrt{d} + 1) \|\mathbf{C}\|_{\mathbb{F}}^2. \end{aligned}$$

Proof. The key idea is to prove a similar result as Lemma 7, but with better constants to reflect that the mean-field parameterization has a lower variance.

First,

$$\begin{aligned} \mathbb{E} \|\mathbf{t}_{\lambda}(\mathbf{u}) - \mathbf{z}\|_2^2 (1 + \|\mathbf{U}\|_{\mathbb{F}}) \\ = \mathbb{E} \|\mathbf{t}_{\lambda}(\mathbf{u}) - \mathbf{z}\|_2^2 + \mathbb{E} \|\mathbf{U}\|_{\mathbb{F}} \|\mathbf{t}_{\lambda}(\mathbf{u}) - \mathbf{z}\|_2^2, \end{aligned}$$

applying Lemma 8,

$$= \|\mathbf{m} - \mathbf{z}\|_2^2 + \|\mathbf{C}\|_{\mathbb{F}} + \mathbb{E} \|\mathbf{U}\|_{\mathbb{F}} \|\mathbf{t}_{\lambda}(\mathbf{u}) - \mathbf{z}\|_2^2. \quad (18)$$

The last term decomposes as

$$\begin{aligned} \mathbb{E} \|\mathbf{U}\|_{\mathbb{F}} \|\mathbf{t}_{\lambda}(\mathbf{u}) - \mathbf{z}\|_2^2 &= \underbrace{\mathbb{E} \|\mathbf{U}\|_{\mathbb{F}} \mathbf{u}^{\top} \mathbf{C}^{\top} \mathbf{C} \mathbf{u}}_{\text{Term 1}} \\ &\quad + 2 \underbrace{\mathbb{E} \|\mathbf{U}\|_{\mathbb{F}} \mathbf{u}^{\top} \mathbf{C}^{\top} (\mathbf{m} - \mathbf{z})}_{\text{Term 2}} \\ &\quad + \underbrace{\mathbb{E} \|\mathbf{U}\|_{\mathbb{F}}}_{\text{Term 3}} \|\mathbf{m} - \mathbf{z}\|_2^2. \end{aligned} \quad (19)$$

We will now focus on the stochastic terms 1-3 one by one.

First, for Term 1, notice that the mean-field parameterization implies that $\mathbf{C} = \text{diag}(c_1, \dots, c_d)$. Thus,

$$\begin{aligned} \mathbb{E} \|\mathbf{U}\|_{\mathbb{F}} \mathbf{u}^{\top} \mathbf{C}^{\top} \mathbf{C} \mathbf{u} &= \mathbb{E} \left(\sqrt{\sum_{i=1}^d u_i^4} \right) \left(\sum_{i=1}^d c_i^2 u_i^2 \right) \\ &= \sum_{i=1}^d c_i^2 \mathbb{E} \left(\sqrt{\sum_{j=1}^d u_j^4} \right) u_i^2, \end{aligned}$$

applying Cauchy-Schwarz inequality for expectations,

$$\leq \sum_{i=1}^d c_i^2 \sqrt{\left(\mathbb{E} \sum_{j=1}^d u_j^4 \right) (\mathbb{E} u_i^4)}$$

and given Assumption 1,

$$\begin{aligned} &= \sum_{i=1}^d c_i^2 \sqrt{d\kappa^2} \\ &= \kappa\sqrt{d} \|\mathbf{C}\|_{\mathbb{F}}^2. \end{aligned} \quad (20)$$

Term ② can be bounded as

$$\mathbb{E} \|\mathbf{U}\|_F \mathbf{u}^\top \mathbf{C}^\top (\mathbf{m} - \mathbf{z})$$

using the Cauchy-Schwarz inequality for vectors as

$$\leq \mathbb{E} \|\mathbf{U}\|_F \|\mathbf{C}\mathbf{u}\|_2 \|\mathbf{m} - \mathbf{z}\|_2,$$

again, applying the inequality for expectations,

$$\begin{aligned} &= \sqrt{\mathbb{E} \|\mathbf{U}\|_F^2 \mathbb{E} \|\mathbf{C}\mathbf{u}\|_2^2} \|\mathbf{m} - \mathbf{z}\|_2 \\ &= \sqrt{\mathbb{E} \left(\sum_{i=1}^d u_i^4 \right) \text{tr}(\mathbf{C}^\top \mathbf{C} \mathbb{E} \mathbf{u} \mathbf{u}^\top)} \|\mathbf{m} - \mathbf{z}\|_2, \end{aligned}$$

from Assumption 1,

$$\begin{aligned} &= \sqrt{d\kappa \text{tr}(\mathbf{C}^\top \mathbf{C})} \|\mathbf{m} - \mathbf{z}\|_2 \\ &= \sqrt{d\kappa} \|\mathbf{C}\|_F \|\mathbf{m} - \mathbf{z}\|_2 \\ &= \sqrt{d\kappa} \sqrt{\|\mathbf{C}\|_F^2 \|\mathbf{m} - \mathbf{z}\|_2^2}, \end{aligned}$$

and by the arithmetic mean-geometric mean inequality,

$$= \frac{\sqrt{d\kappa}}{2} \left(\|\mathbf{C}\|_F^2 + \|\mathbf{m} - \mathbf{z}\|_2^2 \right). \quad (21)$$

Finally, Term ③ follows as

$$\mathbb{E} \|\mathbf{U}\|_F = \mathbb{E} \sqrt{\sum_{i=1}^d u_i^4},$$

using Jensen's inequality,

$$\begin{aligned} &\leq \sqrt{\mathbb{E} \sum_{i=1}^d u_i^4} \\ &= \sqrt{d\kappa}. \end{aligned} \quad (22)$$

Combining all the results, Equation (18) becomes

$$\begin{aligned} &\mathbb{E} \|\mathbf{t}_\lambda(\mathbf{u}) - \mathbf{z}\|_2^2 (1 + \|\mathbf{U}\|_F) \\ &\leq \|\mathbf{m} - \mathbf{z}\|_2^2 + \|\mathbf{C}\|_F^2 \\ &\quad + \mathbb{E} \|\mathbf{U}\|_F \mathbf{u}^\top \mathbf{C}^\top \mathbf{C} \mathbf{u} \\ &\quad + 2 \mathbb{E} \|\mathbf{U}\|_F \mathbf{u}^\top \mathbf{C}^\top \|\mathbf{m} - \mathbf{z}\|_2^2 \\ &\quad + \mathbb{E} \|\mathbf{U}\|_F \|\mathbf{m} - \mathbf{z}\|_2^2 \end{aligned}$$

and applying Equations (20) to (22),

$$\begin{aligned} &\leq \|\mathbf{m} - \mathbf{z}\|_2^2 + \|\mathbf{C}\|_F^2 \\ &\quad + \kappa \sqrt{d} \|\mathbf{C}\|_F \\ &\quad + \kappa \sqrt{d} \left(\|\mathbf{C}\|_F^2 + \|\mathbf{m} - \mathbf{z}\|_2^2 \right) \\ &\quad + \sqrt{d\kappa} \|\mathbf{m} - \mathbf{z}\|_2^2 \\ &= \left(\sqrt{d\kappa} + \kappa \sqrt{d} + 1 \right) \|\mathbf{m} - \mathbf{z}\|_2^2 + \left(2\kappa \sqrt{d} + 1 \right) \|\mathbf{C}\|_F^2. \end{aligned}$$

□

C.2.4. PROOF OF LEMMA 4

Lemma 4. Let $\mathbf{g}_M$ be the M -sample gradient estimator of F (Definition 4) for some function f, h and let $\mathbf{u}$ be some random variable. Then,

$$\mathbb{E} \|\mathbf{g}_M\|_2^2 \leq \frac{1}{M} \mathbb{E} \|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 + \|\nabla F(\lambda)\|_2^2.$$

Proof. From the definition of variance,

$$\begin{aligned} &\mathbb{E} \|\mathbf{g}_M\|_2^2 \\ &= \text{tr} \mathbb{V}[\mathbf{g}_M] + \|\mathbb{E} \mathbf{g}_M\|_2^2, \end{aligned}$$

following the definition in Equation (4),

$$= \text{tr} \mathbb{V} \left[\frac{1}{M} \sum_{m=1}^M \mathbf{g}_m \right] + \|\nabla F(\lambda)\|_2^2,$$

and then the definition in Equation (5),

$$= \text{tr} \mathbb{V} \left[\frac{1}{M} \sum_{m=1}^M \nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}_m)) + \nabla h(\lambda) \right] + \|\nabla F(\lambda)\|_2^2,$$

by the linearity of variance,

$$\begin{aligned} &= \frac{1}{M} \text{tr} \mathbb{V}[\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))] + \|\nabla F(\lambda)\|_2^2 \\ &= \frac{1}{M} \left(\mathbb{E} \|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 - \|\mathbb{E} \nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 \right) \\ &\quad + \|\nabla F(\lambda)\|_2^2 \\ &\leq \frac{1}{M} \mathbb{E} \|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 + \|\nabla F(\lambda)\|_2^2. \end{aligned} \quad (23)$$

□

C.3. Proof of Theorems

C.3.1. PROOF OF THEOREM 1

Theorem 1. Let $\mathbf{g}_M$ be an M -sample estimate of the gradient of the ELBO in entropy regularized form (Definition 5). Also, assume that Assumption 3 and 4 hold,

- f_H is L_H -smooth, and
- f_{KL} is μ_{KL} -quadratically growing.

Then,

$$\begin{aligned} \mathbb{E}\|\mathbf{g}_M\|_2^2 &\leq \frac{4L_H^2}{\mu_{KL}M} C(d, \kappa) (F(\lambda) - F^*) + \|\nabla F(\lambda)\|_2^2 \\ &\quad + \frac{2L_H^2}{M} C(d, \kappa) \|\bar{\xi}_{KL} - \bar{\xi}_H\|_2^2 \\ &\quad + \frac{4L_H^2}{\mu_{KL}M} C(d, \kappa) (F^* - f_{KL}^*), \end{aligned}$$

where

$C(d, \kappa) = 2\kappa\sqrt{d} + 1$ for mean-field,

$C(d, \kappa) = d + \kappa$ for the Cholesky and matrix square root,

$\bar{\xi}_{KL}, \bar{\xi}_H$ are the stationary points of f_{KL}, f_H , respectively, $F^* = \inf_{\lambda \in \mathbb{R}^p} F(\lambda)$, and $f_{KL}^* = \inf_{\xi \in \mathbb{R}^d} f(\xi)$.

Proof. The proof uses the L_H -smoothness of f_H such that

$$\begin{aligned} \mathbb{E}\|\nabla f_H(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 &= \mathbb{E}\|\nabla f_H(\mathbf{t}_\lambda(\mathbf{u})) - \nabla f_H(\bar{\xi}_H)\|_2^2 \\ &\leq L_H^2 \mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \bar{\xi}_H\|_2^2, \end{aligned} \quad (24)$$

where $\bar{\xi}_H$ is a stationary point of f_H such that $\nabla f_H(\bar{\xi}_H) = \mathbf{0}$. These steps have been previously used by Domke (2019, Theorem 3) to prove the special case for the matrix square root parameterization.

For the mean-field parameterization, we start from Lemma 2 and apply Equation (24) as

$$\begin{aligned} \mathbb{E}\|\nabla_\lambda f_H(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 &\leq \mathbb{E}\|\nabla f_H(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 (1 + \|\mathbf{u}\|_F) \\ &\leq L_H^2 \mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) + \bar{\xi}_H\|_2^2 (1 + \|\mathbf{u}\|_F), \end{aligned}$$

applying Lemma 3,

$$\begin{aligned} &\leq L_H^2 (\kappa\sqrt{d} + \sqrt{\kappa d} + 1) \|\mathbf{m} - \bar{\xi}_H\|_2^2 \\ &\quad + L_H^2 (2\kappa\sqrt{d} + 1) \|\mathbf{C}\|_F^2, \end{aligned}$$

and since the kurtosis satisfies $\kappa \geq 1$ and thus $\kappa \geq \sqrt{\kappa}$,

$$\leq L_H^2 (2\kappa\sqrt{d} + 1) (\|\mathbf{m} - \bar{\xi}_H\|_2^2 + \|\mathbf{C}\|_F^2). \quad (25)$$

Similarly, for the full-rank parameterizations, we start from

Lemma 2 and apply Equation (24) as

$$\begin{aligned} \mathbb{E}\|\nabla_\lambda f_H(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 &\leq \mathbb{E}\|\nabla f_H(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 (1 + \|\mathbf{u}\|_2^2), \\ &\leq L_H^2 \mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \bar{\xi}_H\|_2^2 (1 + \|\mathbf{u}\|_2^2), \end{aligned} \quad (26)$$

applying Lemma 7,

$$= L_H^2 \left((d+1) \|\mathbf{m} - \bar{\xi}_H\|_2^2 + (d+\kappa) \|\mathbf{C}\|_F^2 \right),$$

and since the kurtosis satisfies $\kappa \geq 1$,

$$\leq L_H^2 (d+\kappa) (\|\mathbf{m} - \bar{\xi}_H\|_2^2 + \|\mathbf{C}\|_F^2). \quad (27)$$

Both Equations (25) and (27) can now be denoted as

$$\mathbb{E}\|\nabla_\lambda f_H(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 \leq L_H^2 C(d, \kappa) (\|\mathbf{m} - \bar{\xi}_H\|_2^2 + \|\mathbf{C}\|_F^2),$$

where by Lemma 8,

$$= L_H^2 C(d, \kappa) \mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \bar{\xi}_H\|_2^2, \quad (28)$$

and the constants are $C(d, \kappa) = \kappa\sqrt{d} + 1$ for mean-field and $C(d, \kappa) = d + \kappa$ for the full-rank parameterizations.

As mentioned in the sketch, it is necessary to convert the entropy-regularized form into the KL-regularized form through the following inequality:

$$\mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \bar{\xi}_H\|_2^2 \leq 2 \mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \bar{\xi}_{KL}\|_2^2 + 2 \|\bar{\xi}_{KL} - \bar{\xi}_H\|_2^2,$$

where $\bar{\xi}_{KL} = \Pi_{f_{KL}}(\bar{\xi}_H)$ is a projection of $\bar{\xi}_H$ to the set of minimizers of f_{KL} . Note that the KL-regularized form does not need to be tractable; only its existence suffices. We can now apply the quadratic growth assumption as

$$\begin{aligned} \mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \bar{\xi}_{KL}\|_2^2 &\leq \frac{2}{\mu_{KL}} (\mathbb{E}f_{KL}(\mathbf{t}_\lambda(\mathbf{u})) - f_{KL}^*) \\ &= \frac{2}{\mu_{KL}} (F(\lambda) - h_{KL}(\lambda) - f_{KL}^*), \end{aligned}$$

and since $-h_{KL}(\lambda) = -D_{KL}(q_\lambda, p) \leq 0$ by definition,

$$\leq \frac{2}{\mu_{KL}} (F(\lambda) - f_{KL}^*) \quad (29)$$

$$= \frac{2}{\mu_{KL}} ((F(\lambda) - F^*) + (F^* - f_{KL}^*)). \quad (30)$$

Combining Equation (28) with Equation (6),

$$\begin{aligned} \mathbb{E}\|\nabla_\lambda f_H(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 &\leq 2 L_H^2 C(d, \kappa) \mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \bar{\xi}_H\|_2^2 + 2 L_H^2 C(d, \kappa) \|\bar{\xi}_{KL} - \bar{\xi}_H\|_2^2, \\ &\text{and applying Equation (30),} \\ &\leq \frac{4 L_H^2}{\mu_{KL}} C(d, \kappa) ((F(\lambda) - F^*) + (F^* - f_{KL}^*)) \\ &\quad + 2 L_H^2 C(d, \kappa) \|\bar{\xi}_{KL} - \bar{\xi}_H\|_2^2 \end{aligned}$$

Plugging this into [Lemma 4](#) yields the result. $\square$

C.3.2. PROOF OF THEOREM 2

Theorem 2. Let $\mathbf{g}_M$ be an M -sample estimator of the gradient of the ELBO in KL-regularized form ([Definition 6](#)). Also, assume that

- f_{KL} is L_{KL} -smooth,
- f_{KL} is μ_{KL} -quadratically growing,

and [Assumption 3](#) and [4](#) hold. Then, the gradient variance is bounded above as

$$\begin{aligned} \mathbb{E} \|\mathbf{g}_M\|_2^2 &\leq \frac{2L_{\text{KL}}^2}{\mu_{\text{KL}}M} C(d, \kappa) (F(\boldsymbol{\lambda}) - F^*) + \|\nabla F(\boldsymbol{\lambda})\|_2^2 \\ &\quad + \frac{2L_{\text{KL}}^2}{\mu_{\text{KL}}M} C(d, \kappa) (F^* - f_{\text{KL}}^*), \end{aligned}$$

where

$$C(d, \kappa) = 2\kappa\sqrt{d} + 1 \text{ for mean-field,}$$

$$C(d, \kappa) = d + \kappa \quad \text{for the Cholesky and matrix square root,}$$

$$F^* = \inf_{\boldsymbol{\lambda} \in \mathbb{R}^p} F(\boldsymbol{\lambda}), \text{ and } f_{\text{KL}}^* = \inf_{\boldsymbol{\zeta} \in \mathbb{R}^d} f(\boldsymbol{\zeta}).$$

Proof. This proof uses the smoothness of f_{KL} instead of f_{H} . That is,

$$\begin{aligned} \mathbb{E} \|\nabla f_{\text{KL}}(\mathbf{t}_{\boldsymbol{\lambda}}(\mathbf{u}))\|_2^2 &= \mathbb{E} \|\nabla f_{\text{KL}}(\mathbf{t}_{\boldsymbol{\lambda}}(\mathbf{u})) - \nabla f_{\text{KL}}(\bar{\boldsymbol{\zeta}}_{\text{KL}})\|_2^2 \\ \text{applying Equation (24),} \\ &\leq L_{\text{KL}}^2 \mathbb{E} \|\mathbf{t}_{\boldsymbol{\lambda}}(\mathbf{u}) - \bar{\boldsymbol{\zeta}}_{\text{KL}}\|_2^2, \end{aligned} \quad (31)$$

where $\bar{\boldsymbol{\zeta}}_{\text{KL}}$ is a stationary point of f_{KL} .

Substituting [Equation \(31\)](#) in [Equation \(28\)](#),

$$\begin{aligned} \mathbb{E} \|\nabla_{\boldsymbol{\lambda}} f_{\text{KL}}(\mathbf{t}_{\boldsymbol{\lambda}}(\mathbf{u}))\|_2^2 &\leq L_{\text{KL}}^2 C(d, \kappa) \mathbb{E} \|\mathbf{t}_{\boldsymbol{\lambda}}(\mathbf{u}) - \bar{\boldsymbol{\zeta}}_{\text{KL}}\|_2^2, \\ \text{and by applying Equation (30) for } f_{\text{KL}}, \\ &= \frac{2L_{\text{KL}}^2}{\mu_{\text{KL}}} C(d, \kappa) (F(\boldsymbol{\lambda}) - F^*) + \frac{2L_{\text{KL}}^2}{\mu_{\text{KL}}} C(d, \kappa) (F^* - f_{\text{KL}}^*). \end{aligned}$$

Plugging this to [Lemma 4](#) proves the result. $\square$

C.3.3. PROOF OF THEOREM 3

Theorem 3. Let $\mathbf{g}_M$ be an M -sample estimator of the gradient of the ELBO in entropy-regularized form (Definition 5). Also, assume that

- f_H is L_H -smooth,
- f_H is μ_H -quadratically growing,
- h_H is bounded as $h_H(\lambda) > h_H^*$ (Assumption 5),

and Assumption 3 and 4 hold. Then, the gradient variance of $\mathbf{g}_M$ is bounded above as

$$\begin{aligned} \mathbb{E}\|\mathbf{g}_M\|_2^2 &\leq \frac{2L_H^2}{\mu_H M} C(d, \kappa) (F(\lambda) - F^*) + \|\nabla F(\lambda)\|_2^2 \\ &\quad + \frac{2L_H^2}{\mu_H M} C(d, \kappa) (F^* - f_H^* - h_H^*), \end{aligned}$$

where

$$C(d, \kappa) = 2\kappa\sqrt{d} + 1 \text{ for mean-field,}$$

$$C(d, \kappa) = d + \kappa \quad \text{for the Cholesky parameterization,}$$

$$F^* = \inf_{\lambda \in \mathbb{R}^p} F(\lambda), \text{ and } f_H^* = \inf_{\zeta \in \mathbb{R}^d} f(\zeta).$$

Proof. The proof is similar to that of Theorem 1. As mentioned in the *proof sketch*, we use the fact that the entropic regularizer is bounded such that

$$-h_H(\lambda) < -h_H^*.$$

By applying the quadratic growth assumption directly to f_H ,

$$\begin{aligned} \mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \bar{\zeta}_H\|_2^2 &\leq \frac{2}{\mu_H} (\mathbb{E}f_H(\mathbf{t}_\lambda(\mathbf{u})) - f_H^*) \\ &= \frac{2}{\mu_H} (F(\lambda) - h_H(\lambda) - f_H^*), \end{aligned}$$

and by Assumption 5,

$$\leq \frac{2}{\mu_H} (F(\lambda) - F^*) + \frac{2}{\mu_H} (F^* - f_H^* - h_H^*). \quad (32)$$

The proof resumes from Equation (28) as

$$\begin{aligned} \mathbb{E}\|\nabla_\lambda f_H(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 \\ \leq L_H^2 C(d, \kappa) \mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \bar{\zeta}_H\|_2^2, \end{aligned}$$

and by applying Equation (32),

$$= \frac{2L_H^2}{\mu_H} C(d, \kappa) (F(\lambda) - F^*) + \frac{2L_H^2}{\mu_H} C(d, \kappa) (F^* - f_H^* - h_H^*).$$

Plugging this to Lemma 4 proves the result. $\square$

C.3.4. PROOF OF THEOREM 4

Theorem 4. Let $\mathbf{g}_M$ be an M -sample estimator of the gradient of the ELBO in either the entropy- or KL-regularized form. Also, let Assumption 4 hold where the matrix square root parameterization is used. Then, for all L -smooth and μ -strongly convex functions f such that $L/\mu < \sqrt{d+1}$, the variance of $\mathbf{g}_M$ is bounded below by some strictly positive constant as

$$\begin{aligned} \mathbb{E}\|\mathbf{g}_M\|_2^2 &\geq \frac{2\mu^2(d+1) - 2L^2}{ML} (F(\lambda) - F^*) + \|\nabla F(\lambda)\|_2^2 \\ &\quad + \frac{2\mu^2(d+1) - 2L^2}{ML} (\mathbb{E}f(\mathbf{t}_{\lambda^*}(\mathbf{u})) - f^*), \end{aligned}$$

as long as λ is in a local neighborhood around the unique global optimum $\lambda^* = \arg \min_{\lambda \in \mathbb{R}^p} F(\lambda)$, where $F^* = F(\lambda^*)$ and $f^* = \arg \min_{\zeta \in \mathbb{R}^d} f(\zeta)$.

Proof. When using the matrix square root parameterization, Domke (2020) have shown that if f is L -smooth, $\mathbb{E}f(\mathbf{t}_\lambda(\mathbf{u}))$ is also L -smooth. Therefore, we have

$$\|\mathbb{E}\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 \leq 2L (\mathbb{E}f(\mathbf{t}_\lambda(\mathbf{u})) - f^*). \quad (33)$$

Furthermore, let $\bar{\zeta}$ be the minimizer of f , namely $f^* = f(\bar{\zeta})$. From Lemma 6, we have

$$\mathbb{E}\|\nabla_\lambda f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 = \mathbb{E}\|\nabla f(\mathbf{t}_\lambda(\mathbf{u}))\|_2^2 (1 + \|\mathbf{u}\|_2^2),$$

by the μ -strong convexity of f ,

$$\begin{aligned} &\geq 2\mu (\mathbb{E}f(\mathbf{t}_\lambda(\mathbf{u})) - f^*) (1 + \|\mathbf{u}\|_2^2) \\ &\geq \mu^2 \mathbb{E}\|\mathbf{t}_\lambda(\mathbf{u}) - \bar{\zeta}\|_2^2 (1 + \|\mathbf{u}\|_2^2), \end{aligned}$$

applying Lemma 7,

$$= \mu^2 \left((d+1) \|\mathbf{m} - \bar{\zeta}\|_2^2 + (d+\kappa) \|\mathbf{C}\|_F^2 \right),$$

and by the property of the kurtosis that $\kappa \geq 1$,

$$\geq \mu^2 (d+1) \|\lambda - \bar{\lambda}\|_2^2,$$

where $\bar{\lambda} = (\bar{\zeta}, \mathbf{O})$.

Observe that $\bar{\lambda}$ is the minimizer of $\mathbb{E}f(\mathbf{t}_\lambda(\mathbf{u}))$ such that

$$\mathbb{E}f(\mathbf{t}_{\bar{\lambda}}(\mathbf{u})) = f(\bar{\zeta}) = f^* \leq \mathbb{E}f(\mathbf{t}_\lambda(\mathbf{u}))$$

for any λ . Furthermore, from the L -smoothness of $\mathbb{E}f(\mathbf{t}_\lambda(\mathbf{u}))$, we have

$$\begin{aligned} &\mu^2 (d+1) \|\lambda - \bar{\lambda}\|_2^2 \\ &\geq \frac{2\mu^2(d+1)}{L} (\mathbb{E}f(\mathbf{t}_\lambda(\mathbf{u})) - \mathbb{E}f(\mathbf{t}_{\bar{\lambda}}(\mathbf{u}))). \end{aligned}$$

Thus, we have

$$\mathbb{E}\|\nabla_{\lambda} f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 \geq \frac{2\mu^2(d+1)}{L} (\mathbb{E}f(\mathbf{t}_{\lambda}(\mathbf{u})) - f^*). \quad (34)$$

Now, from Equation (23),

$$\mathbb{E}\|\mathbf{g}_M\|_2^2 = \frac{1}{M} \left(\mathbb{E}\|\nabla_{\lambda} f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 - \|\mathbb{E}\nabla_{\lambda} f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 \right) + \|\nabla F(\lambda)\|_2^2,$$

applying Equation (33),

$$\geq \frac{1}{M} \left(\mathbb{E}\|\nabla_{\lambda} f(\mathbf{t}_{\lambda}(\mathbf{u}))\|_2^2 - 2L^2 (\mathbb{E}f(\mathbf{t}_{\lambda}(\mathbf{u})) - f^*) \right) + \|\nabla F(\lambda)\|_2^2$$

applying Equation (34),

$$\begin{aligned} &\geq \frac{2\mu^2(d+1) - 2L^2}{ML} (\mathbb{E}f(\mathbf{t}_{\lambda}(\mathbf{u})) - f^*) + \|\nabla F(\lambda)\|_2^2 \\ &\geq \frac{2\mu^2(d+1) - 2L^2}{ML} (F(\lambda) - h(\lambda) - f^*) + \|\nabla F(\lambda)\|_2^2 \\ &= \frac{2\mu^2(d+1) - 2L^2}{ML} (F(\lambda) - F^*) + \|\nabla F(\lambda)\|_2^2 \\ &\quad + \frac{2\mu^2(d+1) - 2L^2}{ML} (F^* - f^* - h(\lambda)). \end{aligned}$$

The last term

$$\frac{2\mu^2(d+1) - 2L^2}{ML} (F^* - f^* - h(\lambda))$$

can be shown to be positive if λ is sufficiently close to the optimum. Let $\lambda^* = \arg \min_{\lambda} F(\lambda)$ be the minimizer of F . Then, we have

$$\begin{aligned} F^* - f^* - h(\lambda) &= \mathbb{E}f(\mathbf{t}_{\lambda^*}(\mathbf{u})) + h(\lambda^*) - f^* - h(\lambda) \\ &= (\mathbb{E}f(\mathbf{t}_{\lambda^*}(\mathbf{u})) - f^*) + (h(\lambda^*) - h(\lambda)), \end{aligned}$$

where the first term is strictly positive and the second term goes to zero as $\lambda \rightarrow \lambda^*$. $\square$