



Recursive Variational Inference for Total Least-Squares

FRIML, D.; VÁCLAVEK, P.

IEEE Control Systems Letters

Vol. 7, pp. 2839-2844

ISSN: 2475-1456

DOI: <https://doi.org/10.1109/LCSYS.2023.3289608>

Accepted manuscript

©2023 IEEE. Personal use of this material is permitted. Permission from IEEE must be obtained for all other uses, in any current or future media, including reprinting/republishing this material for advertising or promotional purposes, creating new collective works, for resale or redistribution to servers or lists, or reuse of any copyrighted component of this work in other works. FRIML, D.; VÁCLAVEK, P. „Recursive Variational Inference for Total Least-Squares“, IEEE Control Systems Letters, DOI: 10.1109/LCSYS.2023.3289608. Final version is available at <https://ieeexplore.ieee.org/document/10163935>

Recursive Variational Inference for Total Least-Squares

Dominik Friml, and Pavel Václavek *Senior Member, IEEE*

Abstract—This article analyzes methods for deriving credible intervals to facilitate errors-in-variables identification by expanding on Bayesian total least squares. The credible intervals are approximated employing Laplace and variational approximations of the intractable posterior density function. Three recursive identification algorithms providing an approximation of the credible intervals for inference with the Bingham and the Gaussian priors are proposed. The introduced algorithms are evaluated on numerical experiments, and a practical example of application on battery cell total capacity estimation compared to the state-of-the-art algorithms is presented.

Index Terms—Bayes methods, parameter estimation, identification, variational methods

I. INTRODUCTION

The identification of unknown coefficients θ in errors-in-variables problems has become ubiquitous in applied decision and control fields. The ordinary linear regression problem is as follows:

$$b \approx A\theta, \quad (1)$$

where the $\theta \in \mathbb{R}^{n-1}$ is the vector of unknown parameters, and the $A \in \mathbb{R}^{K \times n-1}$ is the design matrix consisting of K noiseless samples of the independent variable $a_k \in \mathbb{R}^{n-1}$ that are stacked underneath each other. Similarly, the $b \in \mathbb{R}^K$ is a vector of K corresponding noisy dependent variables. Compared to the ordinary linear regression, the errors-in-variables arise when the design matrix samples a_k , are noisy. The problem can be reformulated into

$$0 \approx X[\theta^T, -1]^T, \quad (2)$$

where the $X = [A, b] \in \mathbb{R}^{K \times n}$ is the vector of samples $x_k \in \mathbb{R}^n$, burdened with noise. A common assumption is that the noise embodies an independent, identically distributed Gaussian with zero mean $\mathcal{G}(\mathbf{0}, \sigma_n \mathbf{I})$ and known variance σ_n .

It is proved [1] that in the case of an independent, identically distributed Gaussian noise, the maximum likelihood solution is

*This work has been performed in the project RICAIP: Research and Innovation Centre on Advanced Industrial Production that has received funding from the European Union's Horizon 2020 research and innovation programme under grant agreement No 857306 and from Ministry of Education, Youth and Sports under OP RDE grant agreement No CZ.02.1.01/0.0/0.0/17_043/001/0085.

The completion of this paper was made possible by the grant No. FEKT-S-23-8451 - "Research on advanced methods and technologies in cybernetics, robotics, artificial intelligence, automation and measurement" financially supported by the Internal science fund of Brno University of Technology.

Dominik Friml is with the Faculty of Electrical Engineering and Communication, Brno University of Technology, and with the Central European Institute of Technology, Technická 12 Brno, Czech Republic dominik.friml@vutbr.cz, +420 54114 6454

Pavel Václavek is with the Faculty of Electrical Engineering and Communication, Brno University of Technology, and with the Central European Institute of Technology, Technická 12 Brno, Czech Republic pavel.vaclavek@vutbr.cz, +420 54114 6413

obtained with total least-squares [2]. Partially modified total least-squares [3] are the state-of-the-art in maximum likelihood identification, with readily available recursive algorithms [4], [5], [6].

Conversely, the statistical analysis is complicated when it comes to providing credible intervals for the errors-in-variables. The state-of-the-art methods [7] approximate the total least-squares credible intervals by covariance obtained as the inverse of the Fisher information matrix [8]. The likelihood is not Gaussian, as shown in [9]; thus, the approximation is flawed.

Further complications arise because the dimension of the inference grows linearly with the number of samples, K . Due to this fact, no recursive algorithm is currently available [10], [11], [12].

This article proposes to remedy both of the above-mentioned weaknesses by introducing real-time suitable recursive identification algorithms providing a numerically cheap and optimal approximation of the credible intervals.

The text is organized as follows: Section II presents a Bayesian analysis of the total least-squares and justifies the need of a credible interval approximation; sections III and IV analyze numerically efficient and optimal ways of posterior approximation, respectively; sections V and VI derive and evaluate the recursive identification algorithms, respectively; section VII discusses the practical use of the algorithm and compares them with the existing methods; and section VIII sets out the actual conclusion.

Throughout the article, the probability-density functions are denoted by $p(\cdot)$. For better readability, the Bingham distribution notation is abbreviated from the two standard parameters $\mathcal{B}(\cdot|M, Z)$ to $\mathcal{B}(\cdot|A)$, where $A = MZM^T$.

II. BAYESIAN TOTAL LEAST-SQUARES

The authors' previous research [9], under the assumption of an orthogonal noise, formulates the total least-squares likelihood function

$$p(X|\theta) \propto \exp(-[\theta^T - 1]\Phi[\theta^T - 1]^T(\theta^T\theta + 1)^{-1}), \quad (3)$$

where $\Phi = XX^T/(2\sigma_n^2)$.

Using the normalization $v(\theta) = [\theta^T - 1]^T(\theta^T\theta + 1)^{-1/2}$, the likelihood becomes the Bingham distribution

$$p(X|v(\theta)) = \exp(-v(\theta)\Phi v(\theta)^T) = \mathcal{B}(v(\theta)|-\Phi). \quad (4)$$

The Bingham prior $p_B(v(\theta)) = \mathcal{B}(v(\theta)|-\beta)$ is conjugate and ensures the analytical posterior

$$p_B(v(\theta)|X) \propto \mathcal{B}(v(\theta)|-(\Phi + \beta)), \quad (5)$$

where the lower index denotes the Bingham prior.

The $v(\theta)$ can be denormalized at one's convenience via the $[\theta^T - 1]^T = -v(\theta)/v(\theta)_n$, where the $v(\theta)_n$ denotes the last element of the $v(\theta)$. After the normalization, the posterior is written as

$$p_B(\theta|X). \quad (6)$$

While this posterior has certain beneficial properties, as presented in [9], quantifying the θ uncertainty is computationally intensive due to the intractable normalization constant of the Bingham distribution. Although the literature offers algorithms for estimating the normalization constant [13] or enabling the sampling by working around the constant altogether [14], [15], the implementation in real-time is computationally intensive.

Expressing prior beliefs using the $p_B(v(\theta))$ can also be problematic, as one needs to state the belief about the $v(\theta)$ instead of the θ directly. The most frequent prior is the Gaussian $p_G(\theta) = \mathcal{G}(\theta|\mu_p, \Sigma_p)$. However, the resulting posterior is the unknown probability-density function with an unknown normalization constant,

$$p_G(\theta|X) \propto p(X|\theta)p_G(\theta), \quad (7)$$

where the lower index indicates the Gaussian prior.

In both cases, the uncertainty of the θ can be obtained by finding a suitable, preferably optimal, approximation. As the term $(\theta^T \theta + 1)^{-1}$ in the likelihood function (3) is the only difference from the normal distribution, the chapters below examine ways of finding the Gaussian approximations.

III. LAPLACE POSTERIOR APPROXIMATION

The simplest widely used probability-density function approximation providing a Gaussian result is the Laplace approximation [16], [17].

The surrogate posterior is obtained by finding the maximum a posteriori (MAP) estimate, μ^* , of the original posterior distribution, $p(\theta|X)$, computing the Hessian of the log posterior in the mode and then using these values to construct a Gaussian approximation to the posterior

$$q_L(\theta|\mu^*) = \mathcal{G}(\theta|\mu^*, H(\mu^*)), \quad (8)$$

where the $H(\mu^*)$ is the Hessian matrix of the log posterior, $\ln p(\theta|X)$, evaluated at μ^* .

The problematic posterior distribution is approximated by the Laplace surrogate posterior

$$q_L(\theta|\mu^*) \approx p(\theta|X). \quad (9)$$

In the case of the Bingham prior, the MAP μ^* is obtained by solving the singular value decomposition or the Rayleigh quotient iteration [5]

$$\mu^* := \arg \min_{\mu} \frac{[\mu^T - 1](\Phi + \beta)[\mu^T - 1]^T}{\mu^T \mu + 1}. \quad (10)$$

In the case of the Gaussian prior, a nonlinear optimization method is utilized in order to find the modus of the $p_G(\theta|X)$.

The derived Hessians for both the Bingham (20) and the Gaussian (21) priors are provided in the Appendix.

IV. FIXED-FORM VB POSTERIOR APPROXIMATION

While the Laplace approximation is easy to compute, it does not provide an optimal surrogate posterior. The optimal posterior approximation is obtained with variational Bayesian (VB) methods [18], [19]. Although the factorizing mean field variational approximation [16], also referred to as the free-form VB [20], embodies the most commonly used method, it is not applicable in this case, as the parameters are not separable due to the normalization term $(\theta^T \theta + 1)^{-1}$ in the likelihood $p(X|\theta)$. The utilization of the fixed-form VB is required.

As outlined in Section II, we decided to fix the functional form to the Gaussian $\mathcal{G}(\theta|\mu, \Sigma)$. The optimal approximation $q_{VB}(\theta|z)$ of the posterior

$$q_{VB}(\theta|\mu, \Sigma) \approx p(\theta|X) \quad (11)$$

is obtained by minimizing the Kullback-Liebler (KL) divergence, $D_{KL}(q_{VB}(\theta|z)||p(X|\theta)p(\theta))$, which is equivalently solved [20] by maximizing the negative variational free energy providing a lower bound on the marginal log-likelihood, frequently also recognized as evidence lower bound (ELBO),

$$\mathbf{L}(q_{VB}(\theta|z)) = \mathbb{E}_{\theta \sim q_{VB}(\theta|X)} \left(\ln \frac{p(X|\theta)p(\theta)}{q_{VB}(\theta|X)} \right), \quad (12)$$

where the functional form of the $q_{VB}(\theta|X)$ is fixed to a Gaussian with the mean μ and covariance Σ . This results in $z = [\mu^T, \text{vec}(\Sigma)^T]^T$. The prior can only be a Gaussian, $p_G(\theta) = \mathcal{G}(\theta|\mu_p, \Sigma_p)$, as for the Bingham prior the $\mathbf{L}(q_{VB}(\theta|z))$ diverges.

The univariate analytical solution and a sketch of its derivation are provided in the Appendix, (23), as the multivariate analytical solution is not derived at the current stage of the research.

The parameters of the optimal posterior Gaussian approximation are obtained as the solution to the optimization problem

$$z^* := \arg \min_z -\mathbf{L}(q_{VB}(\theta|z)) \quad \text{s.t.} \quad \Sigma \succ 0, \quad (13)$$

constrained by the positive definiteness of the surrogate covariance matrix, Σ , denoted by the $\succ 0$. The ELBO $\mathbf{L}(q_{VB}(\theta|z))$ is differentiable; therefore, this optimization problem can be evaluated by any nonlinear optimization method.

V. RECURSIVE IDENTIFICATION ALGORITHMS

This section presents recursive algorithms, as many applications benefit from sequentially incorporating newly acquired data. The algorithms are designed for real-time applications that demand an insight into the uncertainty of the θ . The MAP for the Bingham prior is recursively obtained via the inverse iteration-based recursive total least-squares algorithm specified in [5]. Extending this algorithm results in the algorithm below, which provides a Gaussian surrogate posterior using the Laplace approximation:

Algorithm 1 Inverse Iteration Recursive Laplace TLS

```

1:  $\mu \leftarrow \mu_0$ 
2:  $H \leftarrow H_B(\mu_0)$ 
3:  $P \leftarrow \beta^{-1}$  ▷ Initialize
4: for  $k \leftarrow 1$  to  $K$  do
5:    $P \leftarrow f(x_k)$  ▷ Incorporate sample
6:    $V \leftarrow P[\mu^T - 1]^T$ 
7:    $\mu \leftarrow V_{1:n-1}/v_n$  ▷ Obtain mean
8:    $H \leftarrow H_B(\mu)$  ▷ Obtain covariance
9: end for

```

For the Gaussian prior, we decided to exploit gradual convergence of the joint probability $p(\theta, X) = p(X|\theta)p(\theta)$ and gradual convergence of the optimization algorithm, similarly to Algorithm 1, resulting in two algorithms that provide the Laplace and the VB surrogate posteriors, respectively:

Algorithm 2 Recursive Laplace TLS

```

1:  $\mu \leftarrow \mu_0$ 
2:  $H \leftarrow H_G(\mu_0)$ 
3:  $\lambda \leftarrow \lambda_0$ 
4:  $\Phi \leftarrow \mathbf{0}$  ▷ Initialize
5: for  $k \leftarrow 1$  to  $K$  do
6:    $\Phi \leftarrow \Phi + (2\sigma_n^2)^{-1}x_kx_k^T$  ▷ Incorporate sample
7:   for  $i \leftarrow 1$  to  $i_{\max}$  do
8:      $J \leftarrow -p(X|\theta)\mathcal{G}(\mu|\mu_p, \Sigma_p)$  ▷ Evaluate joint at  $\mu$ 
9:      $\hat{\mu} \leftarrow \mu - (H_\mu + \lambda\mathbf{I})^{-1}g_\mu$  ▷ Update estimate
10:    if  $-p(X|\theta)\mathcal{G}(\hat{\mu}|\mu_p, \Sigma_p) > J$  then ▷ If  $\hat{\mu}$  invalid
11:       $\lambda \leftarrow \min(\lambda\iota, \lambda_{\max})$  ▷ Enlarge  $\lambda$ 
12:    else ▷ If  $\hat{\mu}$  valid
13:       $\lambda \leftarrow \max(\lambda/\iota, \lambda_{\min})$  ▷ Shrink  $\lambda$ 
14:       $\mu \leftarrow \hat{\mu}$  ▷ Accept mean estimate
15:       $H \leftarrow H_G(\mu)$  ▷ Update covariance
16:    end if
17:  end for
18: end for

```

Algorithm 3 Recursive VB TLS

```

1:  $z \leftarrow [\mu_0^T, \text{vec}(\Sigma_0)^T]^T$ 
2:  $\lambda \leftarrow \lambda_0$ 
3:  $\Phi \leftarrow \mathbf{0}$  ▷ Initialize
4: for  $k \leftarrow 1$  to  $K$  do
5:    $\Phi \leftarrow \Phi + (2\sigma_n^2)^{-1}x_kx_k^T$  ▷ Incorporate sample
6:   for  $i \leftarrow 1$  to  $i_{\max}$  do
7:      $L \leftarrow -\mathbf{L}(q_{\text{VB}}(\theta|z))$  ▷ Evaluate ELBO (23)
8:      $\hat{z} \leftarrow z - (H_L + \lambda\mathbf{I})^{-1}g_L$  ▷ Update estimate
9:     if  $-\mathbf{L}(q_{\text{VB}}(\theta|\hat{z})) > L$  or  $\Sigma(\hat{z}) \preceq \mathbf{0}$  then ▷  $\hat{z}$  invalid
10:       $\lambda \leftarrow \min(\lambda\iota, \lambda_{\max})$  ▷ Enlarge  $\lambda$ 
11:    else ▷  $\hat{z}$  valid
12:       $\lambda \leftarrow \max(\lambda/\iota, \lambda_{\min})$  ▷ Shrink  $\lambda$ 
13:       $z \leftarrow \hat{z}$  ▷ Accept estimate
14:    end if
15:  end for
16: end for

```

In Algorithm 1, the v_n denotes the last element of the vector $V = [V_{1:n-1}^T v_n]^T$; the covariance is calculated using the Hessian

$H_B(\mu)$ provided in the Appendix, (20), and the sample is incorporated using

$$f(x_k) = P - \frac{\frac{1}{2\sigma_n^2}P(x_kx_k^T)P}{1 + \frac{1}{2\sigma_n^2}(x_k^TPx_k)}. \quad (14)$$

We determined experimentally that performing one step, $i_{\max} = 1$, of the optimization algorithm in each sample is sufficient for the convergence. We propose algorithms inspired by Levenberg-Marquardt [21], [22], [23] for both the Laplace and the VB approximation. A variable learning rate using the damping factor, λ , in connection with rejecting undesirable estimates ensures a steady decrease of the optimization goal and a positive definite surrogate covariance matrix, $\Sigma \succ 0$.

In Algorithm 2, the surrogate covariance, $H_G(\mu)$, is provided in the Appendix, (21); the learning gradient and learning Hessian are

$$g_\mu := -\frac{\partial}{\partial \theta} \left(p(X|\theta)p_{\mathcal{G}}(\theta) \right) \quad (15)$$

$$H_\mu := -\frac{\partial^2}{\partial \theta^2} \left(p(X|\theta)p_{\mathcal{G}}(\theta) \right), \quad (16)$$

respectively. Although the derivation is straightforward, the resulting form is too space-intensive to be included in this article.

The parameters μ_0 and Σ_0 are the initial estimates of the surrogate mean and covariance, and $\lambda > 0$ is the Levenberg-Marquardt damping factor, adjusted in each optimization step. The scaling parameter $\iota > 1$ controls the magnitude of the scaling parameter adjustment. The damping factor is bounded by the damping factor limits, $\lambda_{\min}$ and $\lambda_{\max}$, which can be tailored to the data type used in the implementation.

In Algorithm 3, the learning gradient and learning Hessian are

$$g_L := -\frac{\partial}{\partial z} \left(\mathbf{L}(q_{\text{VB}}(\theta|z)) \right), \quad (17)$$

$$H_L := -\frac{\partial^2}{\partial z^2} \left(\mathbf{L}(q_{\text{VB}}(\theta|z)) \right), \quad (18)$$

where the $\mathbf{L}(q_{\text{VB}}(\theta|z))$ is provided in the Appendix, (23). Similarly to 15 and 16, the derivation is straightforward, but the result is too long to be incorporated herein.

VI. NUMERICAL RESULTS

The derived recursive approximation methods are evaluated on numerical simulations using MATLAB. Although the conclusions drawn in this chapter apply to a wide range of simulation settings, the parameters used are provided in the Appendix. This evaluation aims to show that the proposed algorithms converge to results calculated by numerically intensive but precise optimization algorithms. The experiment is repeated 1,000 times. The mean and variance of the convergence are presented, along with the outcome of a single run.

The synthetic errors-in-variables dataset X of $K = 500$ samples is generated by sampling from

$$x_k = [a_k \bar{\theta} a_k]^T + \mathcal{G}(\eta|\mathbf{0}, \sigma_n \mathbf{I}), \quad (19)$$

where the a_k is sampled from the $\mathcal{G}(a_k|\mathbf{0}, 1)$ and $\bar{\theta} = 1.8$.

The algorithms for both the Bingham and the Gaussian priors are evaluated. Each prior is analyzed separately.

In the case of the Bingham prior, only the Laplace approximation method can be utilized, and an optimal surrogate can not be derived.

The KL divergence, $D_{\text{KL}}(q_L(\theta|\mu^*)||q_L(\theta|\mu))$, is employed to quantify the distance from the Laplace estimate $q_L(\theta|z^*)$, with the μ^* obtained via singular decomposition to the surrogate $q_L(\theta|\mu)$ calculated using Algorithm 1.

As is shown in Fig. 1, Algorithm 1 converges rapidly, and the estimate $q_L(\theta|\mu)$ is numerically close to the $q_L(\theta|\mu^*)$.

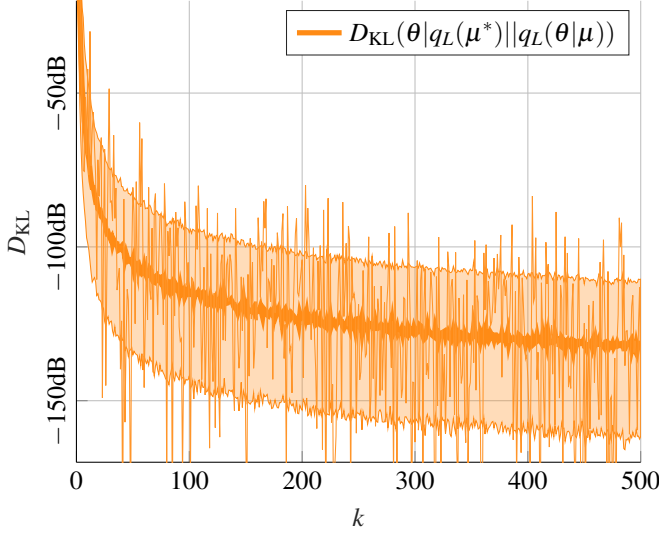


Fig. 1. The Kullback-Liebler divergence from the true MAP-centered Laplace estimate to the estimate calculated by Algorithm 1. The mean and variance over 1,000 runs are represented with the bold line and transparent area, respectively, and the D_{KL} of a single run is denoted by the thin line.

For the case of the Gaussian prior, Algorithms 2 and 3 are proposed; both of them are compared to the optimal estimate $q_{\text{VB}}(\theta|z^*)$, with the z^* acquired by numerically solving (13).

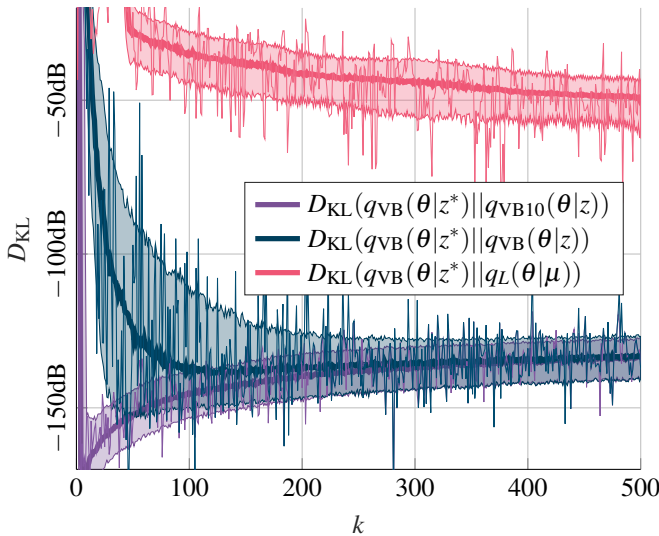


Fig. 2. The Kullback-Liebler divergence from the optimal posterior estimate $q_{\text{VB}}(\theta|z^*)$ to the various estimation methods. The mean and variance over 1,000 runs are represented with the bold line and transparent area, respectively; the D_{KL} of a single run is denoted by the thin line.

The numerical closeness of all the estimation methods is displayed in Fig. 2. The Laplace estimation is computationally cheaper, but the $q_L(\theta|\mu)$ converges to the sub-optimal Laplace estimate $q_L(\theta|\mu^*) \neq q_{\text{VB}}(\theta|z^*)$. The VB approximation exhibits convergence to the optimal approximation $q_{\text{VB}}(\theta|z^*)$ for multiple iterations, $i_{\text{max}} = 10$ per sample $q_{\text{VB}10}(\theta|z)$, and for a single iteration, $i_{\text{max}} = 1$ per sample $q_{\text{VB}}(\theta|z)$.

The visual comparison of the surrogate posteriors depicted in Fig. 3 shows the closeness of the Laplace and VB-based surrogate posteriors. It is clear that while the Laplace estimate entails the MAP and the curvature of the peak, the VB surrogate optimally entails the whole posterior by shifting the mean to explain the skewness.

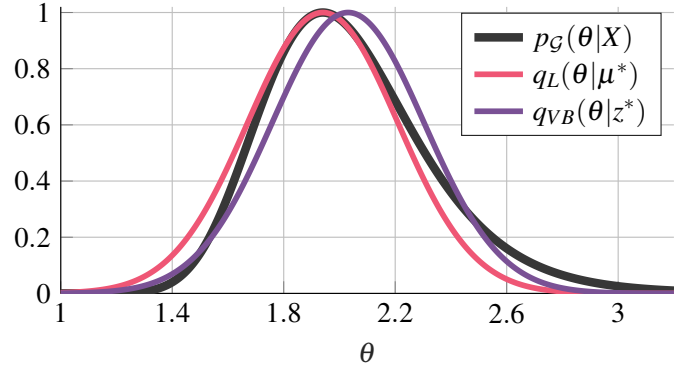


Fig. 3. The visual comparison of the posterior $p_G(\theta|X)$, the Laplace surrogate posterior $q_L(\theta|\mu^*)$, and VB surrogate posterior $q_{\text{VB}}(\theta|z^*)$ for $K = 50$ samples, all normalized such that their maxima are equal to one.

VII. APPLICATION EXAMPLE

The numerical example demonstrates such a practical application of the proposed algorithms for estimating battery cell total capacity where information about the estimate's uncertainty can be crucial for safety. An alternative approach to allow the comparison and detailed subject analysis is provided in [7]. For a survey of other applications, see [2].

The battery cell total capacity, θ , measured in ampere-hours arises in the linear structure, $q = \Delta z \theta$. The linearly dependent variables are the accumulated ampere hour measurement, $q = \int_{t_1}^{t_2} \frac{\eta i(\tau)}{3600} d\tau$, and the state of charge difference, $\Delta z = \frac{z(t_2) - z(t_1)}{100}$. The $z(t)$ is the percentage of the battery cell state of charge at time t , η denotes the unitless efficiency factor, and $i(t)$ represents the battery cell current measured in amperes at time t . For a full explanation of the parameters, refer to [7].

Both the q and the Δz embody measured variables and are therefore burdened with noise. Under the common assumption, the noise components are an independent and identically distributed Gaussian noise with zero mean and the proportional variances σ_q and σ_z , respectively. Due to the proportionality, the measurement scaling makes the total least-squares the optimal estimator of the θ .

The simulated errors-in-variables data are generated with the true battery cell total capacity $\bar{\theta} = 10$ and $\sigma_q = \sigma_z = 0.5$. The proposed algorithm settings are identical to those in the previous section and can be studied in the Appendix.

For comparison, the recursive least-squares (LS) and recursive total least-squares (TLS) algorithms from [7] are displayed. The LS algorithm assumes no error in the Δz . Both of the comparison algorithms (LS and TLS) approximate the posterior credible intervals by the inverse of the Fisher information matrix. This approximation is too simplistic because the likelihood and posterior densities are not Gaussian. The algorithms are initialized with a synthetic true value measurement, $\Delta z = 1$, $q = \bar{\theta}$.

The results of the recursive estimation can be observed in Fig. 4. It is apparent that the mean of the surrogates $q_{VB}(\theta|z)$ and $q_L(\theta|\mu)$ provided by Algorithms 2 and 3, respectively, coincides with the mean of the TLS algorithm, which is the maximum-likelihood solution. While the mean value of the algorithms converges to the true value, $\bar{\theta}$, the mean of the LS algorithm provides biased results; this is due to Δz noise negligence of the LS algorithm.

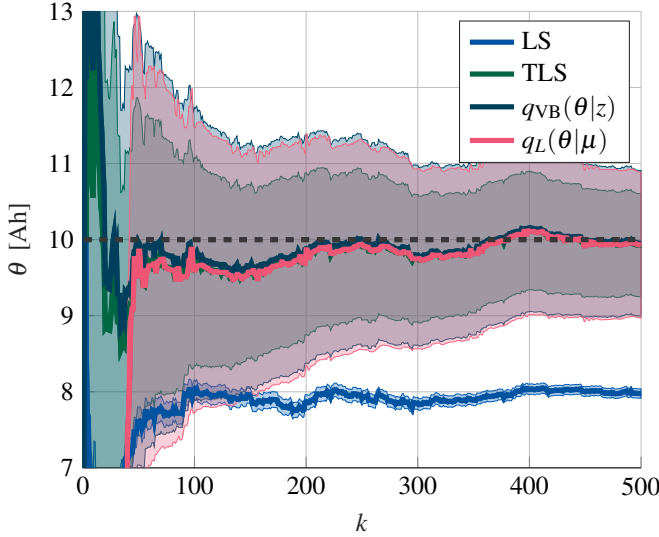


Fig. 4. The battery cell total capacity estimation results for the different algorithms. The mean and 3σ interval approximations are represented by the bold line and transparent area, respectively.

The closest credible interval approximation is expected for the $q_{VB}(\theta|z)$, as the surrogate posterior is an optimal approximation of the true posterior. While the $q_L(\theta|\mu)$ exhibits results similar to those of the $q_{VB}(\theta|z)$, the TLS and LS algorithm intervals are underestimated; this is due to the oversimplified approach using the Fisher information matrix.

VIII. CONCLUSION

Applying analytical Bayesian inference in solving problems with errors-in-variables has numerous disadvantages that complicate further derivations. The disadvantages include the intractable normalization constant of the posterior distribution, absence of computationally efficient algorithms allowing for implementation in real-time applications, and unknown credible intervals for the posterior distribution.

Several procedures to solve the drawbacks are proposed in this article. The intractable posterior complicating the derivation of the credible intervals is solved by a Gaussian approximation of the posterior density function for the Bingham and

the Gaussian priors. The proposed surrogate posteriors are recovered as a Laplace approximation and a variational Bayes approximation with the functional form fixed to a Gaussian, as the factorizing free-form VB would not entail covariance and the unknown parameters are not factorizable.

Inverse iteration and Levenberg-Marquardt inspired algorithms are proposed; they are suitable for the implementation and practical real-time use. All the three algorithms allow batch and recursive identification and facilitate credible interval approximation.

The most numerically efficient option is Algorithm 1; it provides a suboptimal Laplace approximation to the surrogate posterior for the case of the Bingham prior. Expressing prior beliefs via the Bingham distribution can be more challenging than with the Gaussian prior.

A suboptimal Laplace posterior approximation for the Gaussian prior is enabled by Algorithm 2; however, the computational cost is higher than in Algorithm 1.

The numerically least efficient algorithm is Algorithm 3, which delivers an optimal Gaussian posterior approximation for the Gaussian prior. The Faddeeva function appearing in the Hessian evaluated at each step of Algorithm 3, together with the complicated optimization problem, embodies the leading cause of the increased computational cost.

The proposed approximation methods and the corresponding algorithms are numerically evaluated for both of the analyzed prior alternatives (Bingham and Gaussian). The evaluation shows convergence and numerical stability.

The practical use of the proposed algorithms and the comparison with the existing methods are presented on a simulated estimation of battery cell total capacity. The results expose that the maximum a posteriori estimates of the proposed algorithms are compatible with the state-of-the-art maximum-likelihood outcomes of the recursive total least-squares algorithm. Additionally, the proposed algorithms provide an approximation of the credible intervals and, compared to the inverse of the Fisher information matrix, do not underestimate the variance.

While the proposed algorithms are suitable for practical applications, further optimizations and enhancements are possible, such as minimizing the computational complexity and memory requirements or implementing forgetting factors; the last of these modifications then allows for identifying the time-variant parameters, similarly to [24], [25], [26]. Providing a tractable posterior density function also facilitates deriving novel methods for estimation, smoothing, filtering, decision, and control problems with the errors-in-variables.

APPENDIX

The Hessian of the $p(X|\theta)p(\theta)$ in the case of the Bingham prior is derived as

$$H_B(\mu) = 4 \frac{t_1 T_\theta}{t_0^3} - 2 \frac{T_\theta \Xi_\theta - \mu \Xi_n^T}{t_0^2} + 2 \frac{\Xi_\theta T_\theta - \Xi_n \mu^T}{t_0^2} - \frac{\mathbf{I} t_1}{t_0^2} + \frac{\Xi_\theta}{t_0}, \quad (20)$$

where the $t_0 = \mu^T \mu + 1$, $t_1 = [\mu^T - 1] \Xi [\mu^T - 1]^T$, $T_\theta = \mu \mu^T$; $\Xi_\theta \in \mathbb{R}^{n \times n}$, $\Xi_n \in \mathbb{R}^{n \times 1}$ are the submatrices of $\Xi = \begin{bmatrix} \Xi_\theta & \Xi_n \\ \Xi_n^T & \xi_n^2 \end{bmatrix}$ and $\Xi = P^{-1}$.

In the case of the Gaussian prior,

$$H_G(\mu) = H_B(\mu) + \Sigma_p^{-1}, \quad (21)$$

where the $\Xi = \Phi$ and Σ_p is the Gaussian prior covariance matrix.

The univariate analytical variational free energy (ELBO) for the joint $p(X|\theta)p_G(\theta)$ in the case of the Gaussian prior can be evaluated as

$$\begin{aligned} \mathbf{L}(q_{VB}(\theta|z)) &= \int_{-\infty}^{\infty} q(\theta|z) \ln p(\Phi|\theta) d\theta + \\ &+ \int_{-\infty}^{\infty} q(\theta|z) \ln p(\theta) d\theta - \int_{-\infty}^{\infty} q(\theta|z) \ln q(\theta|z) d\theta. \end{aligned} \quad (22)$$

While the evaluation of the latter two terms is trivial, the first term is solved as a convolution problem similar to the Voigt profile [27], resulting in

$$\begin{aligned} \int_{-\infty}^{\infty} q(\theta|z) \ln p(\Phi|\theta) d\theta &= -\frac{1}{2} \left(\frac{\sigma^2 + (\mu - \mu_p)^2}{\sigma_p^2} - \ln \left(\frac{\sigma^2 2\pi e}{e^a} \right) \right) + \\ &+ \frac{\sqrt{\pi}}{\sqrt{2}\sigma^2} \left[(c-a) V \left(\frac{-\mu}{\sqrt{2}\sigma^2}, \frac{1}{\sqrt{2}\sigma^2} \right) - bL \left(\frac{-\mu}{\sqrt{2}\sigma^2}, \frac{1}{\sqrt{2}\sigma^2} \right) \right], \end{aligned} \quad (23)$$

where the μ and σ are the surrogate posterior mean and variance, respectively, and the μ_p and σ_p are the Gaussian prior mean and variance, respectively. The a , b , and c are elements of the $\Phi = \begin{bmatrix} a & b \\ b & c \end{bmatrix}$. The $V(\alpha, \beta)$ and $L(\alpha, \beta)$ are the real and the imaginary parts of the Faddeeva function $w(\gamma)$ [28] respectively, which can be estimated efficiently [29] or with a guaranteed accuracy [30]. The derivation of the ELBO gradient g_L , and the Hessian H_L is straightforward, as

$$\frac{\partial w(\gamma)}{\partial \gamma} = \frac{2i}{\sqrt{\pi}} - 2\gamma w(\gamma). \quad (24)$$

All the derivatives are implemented with precision instead of covariance to improve the numerical stability.

For the numerical tests, the following constants were used: $\mu_p = 5$, $\Sigma_p = \sigma_p^2 = 100$, $\mu_0 = \mu_p$, $z_0 = [\mu_p, \Sigma_p]^T$, $\beta = \text{diag}([10, 10])$, $\lambda_{\min} = 10^{-10}$, $\lambda_{\max} = 10^{10}$, $\lambda_0 = \lambda_{\min}$, $\iota = 2$.

REFERENCES

- [1] L. A. Stefanski, "Measurement error models," in *Statistics in the 21st Century*. Wiley, 2001, pp. 461–470, iSSN: 01621459.
- [2] S. Van Huffel and P. Lemmerling, *Total least squares and errors-in-variables modeling: analysis, algorithms and applications*. Springer Science & Business Media, 2013.
- [3] I. Markovsky and S. V. Huffel, "Overview of total least-squares methods," *Signal Processing*, vol. 87, pp. 2283–2302, 2007.
- [4] S. Van Huffel, *The Generalized Total Least Squares Problem : Formulation, Algorithm and Properties*. Springer Berlin Heidelberg, 1991.
- [5] S. Rhode, F. Bleimund, and F. Gauterin, "Recursive generalized total least squares with noise covariance estimation," in *IFAC Proceedings Volumes (IFAC-PapersOnline)*, vol. 19, Jan. 2014, pp. 4637–4643, issue: 3 ISSN: 14746670.

- [6] S. Rhode, K. Usevich, I. Markovsky, and F. Gauterin, "A recursive restricted total least-squares algorithm," *IEEE Transactions on Signal Processing*, vol. 62, no. 21, pp. 5652–5662, Nov 2014.
- [7] G. L. Plett, "Recursive approximate weighted total least squares estimation of battery cell total capacity," *Journal of Power Sources*, vol. 196, no. 4, pp. 2319–2331, Feb 2011.
- [8] J. L. Crassidis and Y. Cheng, "Maximum likelihood analysis of the total least squares problem with correlated errors," *Journal of Guidance Control and Dynamics*, vol. 42, no. 6, pp. 1204–1217, Jun 2019.
- [9] D. Friml and P. Vaclavek, "Bayesian inference of total least-squares with known precision," *IEEE, Dec 2022*, pp. 203–208.
- [10] O. Nestares, D. Fleet, and D. Heeger, "Likelihood functions and confidence bounds for total-least-squares problems," vol. 1. *IEEE Comput. Soc*, 2000, pp. 523–530.
- [11] N. A. Rozliman, A. I. N. Ibrahim, and R. M. Yunus, "Bayesian approach to errors-in-variables in regression models," in *AIP Conference Proceedings*, vol. 1842, May 2017, p. 030018, issue: 1 ISSN: 15517616.
- [12] X. Fang, B. Li, H. Alkhatib, W. Zeng, and Y. Yao, "Bayesian inference for the errors-in-variables model," *Studia Geophysica et Geodaetica*, vol. 61, no. 1, pp. 35–52, Jan 2017.
- [13] A. Kume and S. G. Walker, "On the Bingham distribution with large dimension," *Journal of Multivariate Analysis*, vol. 124, pp. 345–352, Feb. 2014, publisher: Academic Press.
- [14] C. J. Fallaize and T. Kyraios, "Exact bayesian inference for the bingham distribution," *Statistics and Computing*, vol. 26, no. 1-2, pp. 349–360, Jan 2016.
- [15] M. G. Tsionas, "Note on posterior inference for the bingham distribution," *Communications in Statistics - Theory and Methods*, vol. 47, no. 12, pp. 3022–3028, Jun 2018.
- [16] C. M. Bishop, *Pattern Recognition and Machine Learning (Information Science and Statistics)*. Berlin, Heidelberg: Springer-Verlag, 2006.
- [17] P. E. Rossi, G. M. Allenby, and R. McCulloch, "Bayesian Statistics and Marketing," *Bayesian Statistics and Marketing*, pp. 1–348, Oct. 2006, publisher: wiley ISBN: 9780470863695.
- [18] C. W. Fox and S. J. Roberts, "A tutorial on variational Bayesian inference," *Artificial Intelligence Review 2011 38:2*, vol. 38, no. 2, pp. 85–95, Jun. 2011, publisher: Springer.
- [19] I. Goodfellow, Y. Bengio, and A. Courville, "Approximate Inference," in *Deep Learning*. MIT Press, 2016, pp. 629–650, section: 19.
- [20] W. Penny, S. Kiebel, and K. Friston, *Variational bayes*. Elsevier, London, 2006.
- [21] D. Marquardt, "An algorithm for least-squares estimation of nonlinear parameters," *J. Soc. Ind. Appl. Math.*, vol. 11, pp. 431–441, 1963.
- [22] J. J. Moré, "The Levenberg-Marquardt algorithm: Implementation and theory," in *Numerical Analysis: Proceedings of the Biennial Conference Held at Dundee*. Springer, Berlin, Heidelberg, 1978, pp. 105–116.
- [23] S. Bellavia, S. Gratton, and E. Riccietti, "A levenberg-marquardt method for large nonlinear least-squares problems with dynamic accuracy in functions and gradients," *Numerische Mathematik*, vol. 140, no. 3, pp. 791–825, Nov 2018.
- [24] J. Dokoupil and P. Vaclavek, "Regularized Estimation with Variable Exponential Forgetting," in *Proceedings of the IEEE Conference on Decision and Control*, vol. 2020-Decem, Dec. 2020, pp. 312–318, iSSN: 25762370.
- [25] J. Dokoupil, A. Voda, and P. Vaclavek, "Regularized extended estimation with stabilized exponential forgetting," *IEEE Transactions on Automatic Control*, vol. 62, no. 12, pp. 6513–6520, Dec 2017.
- [26] J. Dokoupil and P. Vaclavek, "Recursive Identification of Time-Varying Hammerstein Systems With Matrix Forgetting," *IEEE Transactions on Automatic Control*, May 2022, publisher: Institute of Electrical and Electronics Engineers Inc.
- [27] F. W. J. Olver, D. W. Lozier, R. F. Boisvert, and C. W. Clark, *NIST Handbook of Mathematical Functions*. Cambridge university press, 2010.
- [28] G. Navas-Palencia, "Extending error function and related functions to complex arguments," 2016.
- [29] G. P. M. Poppe and C. M. J. Wijers, "More efficient computation of the complex error function," *ACM Trans. Math. Softw.*, vol. 16, no. 1, pp. 38–46, Mar 1990.
- [30] M. Al Azah and S. N. Chandler-Wilde, "Computation of the complex error function using modified trapezoidal rules," *SIAM Journal on Numerical Analysis*, vol. 59, no. 5, pp. 2346–2367, Jan 2021.