

On the Convergence of DEM's Linear Parameter Estimator

Meera, Ajith Anil; Wisse, Martijn

DOI

[10.1007/978-3-030-93736-2_49](https://doi.org/10.1007/978-3-030-93736-2_49)

Publication date

2022

Document Version

Final published version

Published in

Machine Learning and Principles and Practice of Knowledge Discovery in Databases

Citation (APA)

Meera, A. A., & Wisse, M. (2022). On the Convergence of DEM's Linear Parameter Estimator. In M. Kamp, M. Kamp, I. Koprinska, & E. A. (Eds.), *Machine Learning and Principles and Practice of Knowledge Discovery in Databases: Proceedings of the International Workshops of ECML PKDD 2021* (pp. 692-700). (Communications in Computer and Information Science; Vol. 1524 CCIS). Springer.
https://doi.org/10.1007/978-3-030-93736-2_49

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' - Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.



On the Convergence of DEM's Linear Parameter Estimator

Ajith Anil Meera^(✉) and Martijn Wisse

Cognitive Robotics, Delft Institute of Technology, Delft, The Netherlands
a.anilmeera@tudelft.nl

Abstract. The free energy principle from neuroscience provides an efficient data-driven framework called the Dynamic Expectation Maximization (DEM), to learn the generative model in the environment. DEM's growing potential to be the brain-inspired learning algorithm for robots demands a mathematically rigorous analysis using the standard control system tools. Therefore, this paper derives the mathematical proof of convergence for its parameter estimator for linear state space systems, subjected to colored noise. We show that the free energy based parameter learning converges to a stable solution for linear systems. The paper concludes by providing a proof of concept through simulation for a wide range of spring damper systems.

Keywords: Free energy principle · Dynamic expectation maximization · Parameter estimation · Linear state space systems

1 Introduction

The free energy principle (FEP) models the brain's perception and action as a gradient ascend over its free energy objective [7]. The action side of FEP, known as active inference [8], has already been applied to real robots including ground robots for SLAM [5], humanoid robots for body perception [12] and manipulator robots for pick and place operation [13]. Similarities with standard control technique like PID was also analyzed [3]. One of the variants of FEP, the Dynamic Expectation Maximization (DEM) [9], provides a model inversion framework for perception and system identification. DEM's distinctive feature lies in its capability to gracefully handle colored noise through the use of generalized coordinates [6], thereby rendering it with the potential to be the learning algorithm for robots. DEM was reformulated as a linear state and input observer under colored noise [11] and was validated for quadrotor flights [4]. A DEM based linear parameter estimator for system identification was developed by [2] and was applied for the perception of quadrotor in wind [1]. Since an estimator with convergence guarantees is preferred for safe robotics applications, we aim at paving way to DEM's practical use by mathematically analyzing it for its convergence properties. Moreover, it is of interest to the active inference community to develop active learning and control strategies with stability guarantees. The

presence of generalized coordinates, mean field terms and brain priors complicates the convergence proof and makes it different from other estimators like Expectation Maximization [10]. The goal of this paper is: 1) to show that DEM has convergence guarantees for linear systems with colored noise, and 2) to show that it can be applied to control system problems like the estimation of a mass-spring-damper system.

2 Preliminaries

Consider the linear plant dynamics (generative process) given in Eq. 1, where $\mathbf{A}$, $\mathbf{B}$ and $\mathbf{C}$ are constant system matrices, $\mathbf{x} \in \mathbb{R}^n$ is the hidden state, $\mathbf{v} \in \mathbb{R}^r$ is the input and $\mathbf{y} \in \mathbb{R}^m$ is the output.

$$\dot{\mathbf{x}} = \mathbf{A}\mathbf{x} + \mathbf{B}\mathbf{v} + \mathbf{w}, \quad \mathbf{y} = \mathbf{C}\mathbf{x} + \mathbf{z}. \quad (1)$$

Here $\mathbf{w} \in \mathbb{R}^n$ and $\mathbf{z} \in \mathbb{R}^m$ represent the process and measurement noise respectively. The notations of the plant are denoted in boldface, whereas its estimates are denoted in nonboldface letters. Since the brain has no access to the plant dynamics except through the sensory measurements $\mathbf{y}$, it maintains the copy of an approximate model of the generative process called the generative model. The noise color assumption (convolution of white noise with a Gaussian kernel) facilitates the differentiated form of the generative model as [9]:

$$\begin{aligned} x' &= Ax + Bv + w & y &= Cx + z \\ x'' &= Ax' + Bv' + w' & \dot{y} &= Cx' + z' \\ \dots & & \dots & \end{aligned} \quad (2)$$

One of the key technique behind DEM to model the colored noise is to express the time varying components in generalized coordinates, denoted by a tilde operator. The colored noises can be expressed in generalized coordinates using their higher derivatives as $\tilde{z} = [z, z', z'', \dots]^T$ and $\tilde{w} = [w, w', w'', \dots]^T$. The generative model in Eq. 2 can be compactly written as [9]:

$$\dot{\tilde{x}} = D^x \tilde{x} = \tilde{A} \tilde{x} + \tilde{B} \tilde{v} + \tilde{w} \quad \tilde{y} = \tilde{C} \tilde{x} + \tilde{z} \quad (3)$$

$$\text{where } D^x = \begin{bmatrix} 0 & 1 & & \\ & 0 & 1 & \\ & & \ddots & \ddots \\ & & & 0 & 1 \\ & & & & 0 \end{bmatrix}_{(p+1) \times (p+1)} \otimes I_{n \times n}, \quad \tilde{A} = I_{p+1} \otimes A, \quad \tilde{B} = I_{p+1} \otimes B$$

and $\tilde{C} = I_{p+1} \otimes C$. Here $\otimes$ is the Kronecker tensor product. To facilitate the convergence proof later in the paper, we introduce a redefinition for Eq. 3 with all parameters grouped to the right side as θ :

$$\dot{\tilde{x}} = M\theta + \tilde{w}, \quad \tilde{y} = N\theta + \tilde{z}, \quad \theta = \begin{bmatrix} \text{vec}(A^T) \\ \text{vec}(B^T) \\ \text{vec}(C^T) \end{bmatrix}, \quad (4)$$

where

$$M = \begin{bmatrix} I_n \otimes x^T & I_n \otimes v^T & I_n \otimes O_{1 \times m} \\ I_n \otimes x'^T & I_n \otimes v'^T & I_n \otimes O_{1 \times m} \\ \dots & \dots & \dots \end{bmatrix}, N = \begin{bmatrix} I_n \otimes O_{1 \times n} & I_n \otimes O_{1 \times r} & I_m \otimes x^T \\ I_n \otimes O_{1 \times n} & I_n \otimes O_{1 \times r} & I_m \otimes x'^T \\ \dots & \dots & \dots \end{bmatrix}. \quad (5)$$

The goal of this paper is to mathematically prove that the DEM's estimate for θ converges while maximizing the free energy objective.

3 Parameter Learning as Free Energy Optimization

DEM postulates the parameter learning algorithm as the gradient ascend over the free energy action, which is the time integral of free energy $\bar{F} = \int F dt$. The parameter update equation can be expressed as the gradient [2, 9]:

$$\frac{\partial \theta}{\partial a} = k^\theta \frac{\partial \bar{F}}{\partial \theta} = -P^\theta (\theta - \eta^\theta) + \sum_t (-E_\theta + W_\theta^X), \quad (6)$$

where k^θ is the learning rate, $E_\theta = \frac{\partial E}{\partial \theta}$ is the gradient of precision weighed prediction error, $W_{\theta^i}^X = \frac{\partial W^X}{\partial \theta}$ is the gradient of state mean field term, η^θ is the prior parameters and P^θ is the prior parameter precision. Subscripts will be used for the derivative operator. E_θ for an LTI system can be simplified as:

$$E_\theta = \tilde{\epsilon}_\theta^T \tilde{\Pi} \tilde{\epsilon}, \text{ where } \tilde{\epsilon} = \begin{bmatrix} \tilde{\mathbf{y}} - N\theta \\ \tilde{v} - \tilde{\eta}^v \\ D^x \tilde{x} - M\theta \end{bmatrix} \text{ and } \tilde{\epsilon}_\theta = \begin{bmatrix} -N \\ O \\ -M \end{bmatrix} \quad (7)$$

are the prediction error and its gradient. Here $\tilde{\eta}^v$ is the prior on inputs with prior precision $\tilde{P}^v$, $\tilde{\Pi} = \text{diag}(\tilde{\Pi}^z, \tilde{P}^v, \tilde{\Pi}^w)$ is the generalized noise precision with Π^z and Π^w being the precisions (inverse covariance) for measurement and process noise. Here $\text{diag}()$ represents the block diagonal operation. Similarly, W_θ^X for an LTI system can be written as [2, 9]:

$$W_{\theta^i}^X = -\frac{1}{2} \text{tr}(\Sigma^X \tilde{\epsilon}_{X\theta^i}^T \tilde{\Pi} \tilde{\epsilon}_X), \quad \tilde{\epsilon} = \begin{bmatrix} \tilde{\mathbf{y}} - \tilde{C}\tilde{x} \\ \tilde{v} - \tilde{\eta}^v \\ D^x \tilde{x} - \tilde{A}\tilde{x} - \tilde{B}\tilde{v} \end{bmatrix}, \quad \tilde{\epsilon}_X = \begin{bmatrix} -\tilde{C} & O \\ O & I \\ D^x - \tilde{A} - \tilde{B} \end{bmatrix}. \quad (8)$$

4 Proof of Convergence for Parameter Estimator

If E_θ and W_θ^X can be expressed as linear in θ , in the form $E_\theta = E_1\theta + E_2$ and $W_\theta^X = W_1\theta + W_2$, Eq. 6 can be rewritten as:

$$\frac{\partial \theta}{\partial a} = -\left[P^\theta + \sum_t (E_1 - W_1)\right]\theta + \left[P^\theta \eta^\theta + \sum_t (-E_2 + W_2)\right]. \quad (9)$$

The differential equation given by Eq. 9 is of the form of a linear state space equation ($\dot{\theta} = A^\theta \theta + B^\theta \cdot 1$). From the basics of control theory, the solutions of this equation converges exponentially (stabilise) if $A^\theta = -[P^\theta + \sum_t (E_1 - W_1)]$ is negative definite (negative eigen values). This section aims to prove this result.

Lemma 1. *If $A, B \succ O$, then $A + B \succ O$.*

As per Lemma 1, the positive definiteness of $P^\theta - \sum_t W_1 + \sum_t E_1$ can be proved by proving the positive definiteness of the individual terms P^θ , $-W_1$ and E_1 . We know by definition that the prior parameter precision P^θ is positive definite. We now proceed to prove that $E_1 \succeq O$. Upon simplification of Eq. 7, E_θ can be written as $E_\theta = E_1 \theta + E_2$, where:

$$E_1 = N^T \tilde{\Pi}^z N + M^T \tilde{\Pi}^w M \text{ and } E_2 = -[N^T \tilde{\Pi}^z M^T \tilde{\Pi}^w D] \begin{bmatrix} \tilde{\mathbf{y}} \\ \tilde{\mathbf{x}} \end{bmatrix}. \quad (10)$$

Lemma 2. *If $A \succeq O$, then $B^T A B \succeq O$.*

Proof. By definition, if $A \succeq O$, there exists a square root $A^{\frac{1}{2}} \succeq O$. Therefore, $x^T (B^T A B) x = x^T (B^T A^{\frac{1}{2}} A^{\frac{1}{2}} B) x = (A^{\frac{1}{2}} B x)^T (A^{\frac{1}{2}} B x) \geq 0, \implies B^T A B \succeq O$.

Since $\tilde{\Pi}^z \succ O$ and $\tilde{\Pi}^w \succ O$ by definition, from Lemma 1 and 2, $E_1 = N^T \tilde{\Pi}^z N + M^T \tilde{\Pi}^w M \succeq O$. Therefore, E_1 is proved to be positive semi-definite.

The final term under consideration is W_1 . The rest of this section aims to prove that $W_1 \prec O$, which will conclude the entire convergence proof of parameter estimation. We rewrite the mean field term for parameter θ^i from Eq. 8 as:

$$\begin{aligned} W_{\theta^i}^X &= -\frac{1}{2} \text{tr}(\Sigma^X \tilde{\epsilon}_{X\theta^i}^T \tilde{\Pi} \tilde{\epsilon}_X), \\ &= -\frac{1}{2} \text{tr} \left[\begin{bmatrix} \Sigma^{\tilde{x}\tilde{x}} & \Sigma^{\tilde{x}\tilde{v}} \\ \Sigma^{\tilde{v}\tilde{x}} & \Sigma^{\tilde{v}\tilde{v}} \end{bmatrix} \begin{bmatrix} \tilde{C}_{\theta^i}^T \tilde{\Pi}^z \tilde{C} - \tilde{A}_{\theta^i}^T \tilde{\Pi}^w (D - \tilde{A}) & \tilde{A}_{\theta^i}^T \tilde{\Pi}^w \tilde{B} \\ -\tilde{B}_{\theta^i}^T \tilde{\Pi}^w (D - \tilde{A}) & \tilde{B}_{\theta^i}^T \tilde{\Pi}^w \tilde{B} \end{bmatrix} \right] \\ &= -\frac{1}{2} \text{tr} \left[\begin{bmatrix} \Sigma^{\tilde{x}\tilde{x}} & \Sigma^{\tilde{x}\tilde{v}} \\ \Sigma^{\tilde{v}\tilde{x}} & \Sigma^{\tilde{v}\tilde{v}} \end{bmatrix} \begin{bmatrix} \tilde{C}_{\theta^i}^T \tilde{\Pi}^z \tilde{C} + \tilde{A}_{\theta^i}^T \tilde{\Pi}^w \tilde{A} & \tilde{A}_{\theta^i}^T \tilde{\Pi}^w \tilde{B} \\ \tilde{B}_{\theta^i}^T \tilde{\Pi}^w \tilde{A} & \tilde{B}_{\theta^i}^T \tilde{\Pi}^w \tilde{B} \end{bmatrix} \right] \\ &\quad - \frac{1}{2} \text{tr} \left[\begin{bmatrix} \Sigma^{\tilde{x}\tilde{x}} & \Sigma^{\tilde{x}\tilde{v}} \\ \Sigma^{\tilde{v}\tilde{x}} & \Sigma^{\tilde{v}\tilde{v}} \end{bmatrix} \begin{bmatrix} -\tilde{A}_{\theta^i}^T \tilde{\Pi}^w D & O \\ -\tilde{B}_{\theta^i}^T \tilde{\Pi}^w D & O \end{bmatrix} \right]. \end{aligned} \quad (11)$$

Since the second trace term in Eq. 11 is independent of θ^i , it is lumped into the $W_2^{\theta^i}$ term. Equation 11 is further simplified as:

$$\begin{aligned} W_{\theta^i}^X &= -\frac{1}{2} \left[\text{tr}(\Sigma^{\tilde{x}\tilde{x}} \tilde{C}_{\theta^i}^T \tilde{\Pi}^z \tilde{C}) + \text{tr}(\Sigma^{\tilde{x}\tilde{x}} \tilde{A}_{\theta^i}^T \tilde{\Pi}^w \tilde{A}) + \text{tr}(\Sigma^{\tilde{x}\tilde{v}} \tilde{B}_{\theta^i}^T \tilde{\Pi}^w \tilde{A}) \right. \\ &\quad \left. + \text{tr}(\Sigma^{\tilde{v}\tilde{x}} \tilde{A}_{\theta^i}^T \tilde{\Pi}^w \tilde{B}) + \text{tr}(\Sigma^{\tilde{v}\tilde{v}} \tilde{B}_{\theta^i}^T \tilde{\Pi}^w \tilde{B}) \right] + W_2^{\theta^i} \end{aligned} \quad (12)$$

We aim to separate θ out so that the mean field term can be expressed in the form $W_\theta^X = W_1 \theta + W_2$. We proceed by first introducing the transpose of the generalized parameter matrices $\tilde{A}$, $\tilde{B}$ and $\tilde{C}$ to Eq. 12 and then moving them out of the trace terms.

Lemma 3. *If A , B , C and D are matrices, then $\text{tr}(ABCD) = \text{tr}(C^T B^T A^T D^T)$*

Proof. $\text{tr}(ABCD) = \text{tr}((ABCD)^T) = \text{tr}(D^T C^T B^T A^T) = \text{tr}(C^T B^T A^T D^T)$.

Lemma 4. *If A , B and C are matrices, then $\text{tr}(ABC) = \text{vec}(A^T)^T(I \otimes B)\text{vec}(C)$.*

Applying Lemma 3 throughout Eq. 12 results in:

$$W_{\theta^i}^X = -\frac{1}{2} \left[\text{tr}(\tilde{\Pi}^z \tilde{C}_{\theta^i} \Sigma^{\tilde{x}\tilde{x}T} \tilde{C}^T) + \text{tr}(\tilde{\Pi}^w \tilde{A}_{\theta^i} \Sigma^{\tilde{x}\tilde{x}T} \tilde{A}^T) + \text{tr}(\tilde{\Pi}^w \tilde{B}_{\theta^i} \Sigma^{\tilde{x}\tilde{v}T} \tilde{A}^T) \right. \\ \left. + \text{tr}(\tilde{\Pi}^w \tilde{A}_{\theta^i} \Sigma^{\tilde{v}\tilde{x}T} \tilde{B}^T) + \text{tr}(\tilde{\Pi}^w \tilde{B}_{\theta^i} \Sigma^{\tilde{v}\tilde{v}T} \tilde{B}^T) \right] + W_2^{\theta^i}, \quad (13)$$

which upon further expansion using Lemma 4 and grouping yields:

$$W_{\theta^i}^X = -\frac{1}{2} \left[\left(\text{vec}(\tilde{A}_{\theta^i}^T \tilde{\Pi}^w)^T (I \otimes \Sigma^{\tilde{x}\tilde{x}T}) + \text{vec}(\tilde{B}_{\theta^i}^T \tilde{\Pi}^w)^T (I \otimes \Sigma^{\tilde{x}\tilde{v}T}) \right) \text{vec}(\tilde{A}^T) \right. \\ \left. + \left(\text{vec}(\tilde{A}_{\theta^i}^T \tilde{\Pi}^w)^T (I \otimes \Sigma^{\tilde{v}\tilde{x}T}) + \text{vec}(\tilde{B}_{\theta^i}^T \tilde{\Pi}^w)^T (I \otimes \Sigma^{\tilde{v}\tilde{v}T}) \right) \text{vec}(\tilde{B}^T) \right. \\ \left. + \left(\text{vec}(\tilde{C}_{\theta^i}^T \tilde{\Pi}^z)^T (I \otimes \Sigma^{\tilde{x}\tilde{x}T}) \right) \text{vec}(\tilde{C}^T) \right] + W_2^{\theta^i}. \quad (14)$$

We have now separated all the generalized parameters out of the trace terms in their vector forms. These vectors can be grouped such that the mean field term

is linear with respect to the generalized parameter vector $\tilde{\theta} = \begin{bmatrix} \text{vec}(\tilde{A}^T) \\ \text{vec}(\tilde{B}^T) \\ \text{vec}(\tilde{C}^T) \end{bmatrix}$ as:

$$W_{\theta^i}^X = -\frac{1}{2} \left[\text{vec}(\tilde{A}_{\theta^i}^T \tilde{\Pi}^w)^T (I \otimes \Sigma^{\tilde{x}\tilde{x}T}) + \text{vec}(\tilde{B}_{\theta^i}^T \tilde{\Pi}^w)^T (I \otimes \Sigma^{\tilde{x}\tilde{v}T}), \right. \\ \text{vec}(\tilde{A}_{\theta^i}^T \tilde{\Pi}^w)^T (I \otimes \Sigma^{\tilde{v}\tilde{x}T}) + \text{vec}(\tilde{B}_{\theta^i}^T \tilde{\Pi}^w)^T (I \otimes \Sigma^{\tilde{v}\tilde{v}T}), \quad (15) \\ \left. \text{vec}(\tilde{C}_{\theta^i}^T \tilde{\Pi}^z)^T (I \otimes \Sigma^{\tilde{x}\tilde{x}T}) \right] \tilde{\theta} + W_2^{\theta^i}.$$

Lemma 5. *If A and B are matrices, then $\text{vec}(AB)^T = \text{vec}(A)^T(B \otimes I)$.*

We use Lemma 5 to further simplify Eq. 15 as:

$$W_{\theta^i}^X = -\frac{1}{2} \left[\text{vec}(\tilde{A}_{\theta^i}^T)^T (\tilde{\Pi}^w \otimes I) (I \otimes \Sigma^{\tilde{x}\tilde{x}T}) + \text{vec}(\tilde{B}_{\theta^i}^T)^T (\tilde{\Pi}^w \otimes I) (I \otimes \Sigma^{\tilde{x}\tilde{v}T}), \right. \\ \text{vec}(\tilde{A}_{\theta^i}^T)^T (\tilde{\Pi}^w \otimes I) (I \otimes \Sigma^{\tilde{v}\tilde{x}T}) + \text{vec}(\tilde{B}_{\theta^i}^T)^T (\tilde{\Pi}^w \otimes I) (I \otimes \Sigma^{\tilde{v}\tilde{v}T}), \\ \left. \text{vec}(\tilde{C}_{\theta^i}^T)^T (\tilde{\Pi}^z \otimes I) (I \otimes \Sigma^{\tilde{x}\tilde{x}T}) \right] \tilde{\theta} + W_2^{\theta^i}. \quad (16)$$

Since the parameters A, B and C are independent of each other, their derivatives with respect to each other are zeros, resulting in $\text{vec}(\tilde{A}_{\theta^i}^T) = O, \forall \theta^i \in \{B, C\}$, $\text{vec}(\tilde{B}_{\theta^i}^T) = O, \forall \theta^i \in \{A, C\}$ and $\text{vec}(\tilde{C}_{\theta^i}^T) = O, \forall \theta^i \in \{A, B\}$. This simplifies the expression for $W_{\theta^i}^X$ in Eq. 16. The total mean field term W_{θ}^X can be computed by vertically stacking the individual mean field contributions $W_{\theta^i}^X$ from each parameter θ^i as:

$$W_{\theta}^X = -\frac{1}{2}W_3\tilde{\theta} + W_2, \quad (17)$$

where $W_3 = \begin{bmatrix} W_4 & O \\ O & W_5 \end{bmatrix}$ with $W_5 = \text{vec}(\tilde{C}^T)_{\text{vec}C^T}^T (\tilde{\Pi}^z \otimes I)(I \otimes \Sigma^{\tilde{x}\tilde{x}T})$ and

$$W_4 = \begin{bmatrix} \text{vec}(\tilde{A}^T)_{\text{vec}A^T}^T (\tilde{\Pi}^w \otimes I)(I \otimes \Sigma^{\tilde{x}\tilde{x}T}) & \text{vec}(\tilde{A}^T)_{\text{vec}A^T}^T (\tilde{\Pi}^w \otimes I)(I \otimes \Sigma^{\tilde{v}\tilde{x}T}) \\ \text{vec}(\tilde{B}^T)_{\text{vec}B^T}^T (\tilde{\Pi}^w \otimes I)(I \otimes \Sigma^{\tilde{x}\tilde{v}T}) & \text{vec}(\tilde{B}^T)_{\text{vec}B^T}^T (\tilde{\Pi}^w \otimes I)(I \otimes \Sigma^{\tilde{v}\tilde{v}T}) \end{bmatrix}.$$

W_3 can be simplified as:

$$W_3 = \frac{\partial \tilde{\theta}^T}{\partial \theta} \begin{bmatrix} \tilde{\Pi}^w \otimes I & O & O \\ O & \tilde{\Pi}^w \otimes I & O \\ O & O & \tilde{\Pi}^z \otimes I \end{bmatrix} \begin{bmatrix} I \otimes \Sigma^{\tilde{x}\tilde{x}T} & I \otimes \Sigma^{\tilde{v}\tilde{x}T} & O \\ I \otimes \Sigma^{\tilde{x}\tilde{v}T} & I \otimes \Sigma^{\tilde{v}\tilde{v}T} & O \\ O & O & I \otimes \Sigma^{\tilde{x}\tilde{x}T} \end{bmatrix}, \quad (18)$$

where $\frac{\partial \tilde{\theta}}{\partial \theta} = \text{diag}(\text{vec}\tilde{A}_{\text{vec}A^T}^T, \text{vec}\tilde{B}_{\text{vec}B^T}^T, \text{vec}\tilde{C}_{\text{vec}C^T}^T)$. Since the generalized parameter vector $\tilde{\theta}$ is linear in parameter vector θ , we can write:

$$\tilde{\theta} = \frac{\partial \tilde{\theta}}{\partial \theta} \theta = \begin{bmatrix} \text{vec}\tilde{A}_{\text{vec}A^T}^T & O & O \\ O & \text{vec}\tilde{B}_{\text{vec}B^T}^T & O \\ O & O & \text{vec}\tilde{C}_{\text{vec}C^T}^T \end{bmatrix} \theta. \quad (19)$$

Substituting Eq. 18 and 19 in Eq. 17 yields:

$$W_{\theta}^X = W_1\theta + W_2,$$

$$W_1 = -\frac{1}{2} \frac{\partial \tilde{\theta}^T}{\partial \theta} \begin{bmatrix} \tilde{\Pi}^w \otimes I & O & O \\ O & \tilde{\Pi}^w \otimes I & O \\ O & O & \tilde{\Pi}^z \otimes I \end{bmatrix} \begin{bmatrix} I \otimes \Sigma^{\tilde{x}\tilde{x}T} & I \otimes \Sigma^{\tilde{v}\tilde{x}T} & O \\ I \otimes \Sigma^{\tilde{x}\tilde{v}T} & I \otimes \Sigma^{\tilde{v}\tilde{v}T} & O \\ O & O & I \otimes \Sigma^{\tilde{x}\tilde{x}T} \end{bmatrix} \frac{\partial \tilde{\theta}}{\partial \theta}. \quad (20)$$

Therefore, the mean field term W_{θ}^X is linear in θ . For the parameter estimator to provide a converging solution, we need to prove that $W_1 \prec O$. Lemma 2 could be applied to the expression for W_1 to prove that $W_1 \prec O$ if:

$$W_6 = \begin{bmatrix} \tilde{\Pi}^w \otimes I & O & O \\ O & \tilde{\Pi}^w \otimes I & O \\ O & O & \tilde{\Pi}^z \otimes I \end{bmatrix} \begin{bmatrix} I \otimes \Sigma^{\tilde{x}\tilde{x}T} & I \otimes \Sigma^{\tilde{v}\tilde{x}T} & O \\ I \otimes \Sigma^{\tilde{x}\tilde{v}T} & I \otimes \Sigma^{\tilde{v}\tilde{v}T} & O \\ O & O & I \otimes \Sigma^{\tilde{x}\tilde{x}T} \end{bmatrix} \succ O \quad (21)$$

Lemma 6. *If $A, B \succeq O$ and A is invertible, then $AB \succeq O$.*

Proof. $AB = A^{\frac{1}{2}}(A^{\frac{1}{2}}BA^{\frac{1}{2}})A^{-\frac{1}{2}}$, implies AB and $A^{\frac{1}{2}}BA^{\frac{1}{2}}$ are similar matrices, sharing all eigen values. Using Lemma 2, since $B \succeq O$, $A^{\frac{1}{2}}BA^{\frac{1}{2}} \succeq O \implies AB \succeq O$.

Using Lemma 6 it is straightforward to see that $W_6 \succeq O$ because: $\tilde{\Pi}^z \succ O, \tilde{\Pi}^w \succ O, \implies \tilde{\Pi}^z \otimes I \succ O$ and $\tilde{\Pi}^w \otimes I \succ O, I \otimes \Sigma^X \succ O$. Therefore, $W_1 \preceq O$. This completes the proof that the parameter estimation of DEM converges for an LTI system with colored noise.

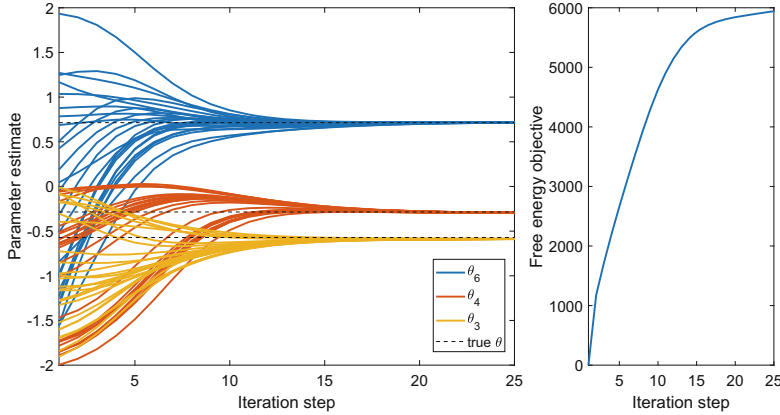


Fig. 1. The parameter estimates of DEM converges to the correct value of $\theta_3 = -\frac{k}{m} = -0.5714$, $\theta_4 = -\frac{b}{m} = -0.2857$ and $\theta_6 = \frac{1}{m} = 0.7143$, marked in black, for a set of 25 experiments, despite being initialized by randomly sampled priors such that $\eta^{\theta^i} \in [-2, 2]$ and that the prior A matrix is stable. The parameter estimation proceeds by maximizing the free energy objective as shown on the right (sample realization).

5 Proof of Concept: Mass-Spring-Damper System

This section aims at providing a proof of concept for the convergence of DEM's parameter estimator, through realistic simulations. A mass-spring-damper system with mass $m = 1.4$ kg, spring constant $k = 0.8$ N/m and damping coefficient $b = 0.4$ Ns/m, is considered in the state space form given by:

$$\begin{bmatrix} \dot{x} \\ \dot{\dot{x}} \end{bmatrix} = \begin{bmatrix} 0 & 1 \\ -\frac{k}{m} & -\frac{b}{m} \end{bmatrix} \begin{bmatrix} x \\ \dot{x} \end{bmatrix} + \begin{bmatrix} 0 \\ \frac{1}{m} \end{bmatrix} v, \quad y = \begin{bmatrix} 1 & 0 \end{bmatrix} \begin{bmatrix} x \\ \dot{x} \end{bmatrix}. \quad (22)$$

A Gaussian bump input $v = e^{-0.25(t-12)^2}$, centred around 12s and sampled at $dt = 0.1$ s for $T = 32$ s was used. To generate the colored noise, the white noise ($\Pi^w = e^6 I_2$ and $\Pi^z = e^6$) was convoluted using a Gaussian kernel with a width of $\sigma = 0.5$ s. A partially known system with unknown $\theta_3 = -\frac{k}{m}$, $\theta_4 = -\frac{b}{m}$ and $\theta_6 = \frac{1}{m}$ was considered. Using the output y generated from the spring damper system, parameter estimation was performed using DEM for 25 experiments with different η^θ . The parameter priors η^θ for unknown parameters were randomly sampled from $[-2, 2]$ such that the resulting prior A matrix is stable. A low prior precision ($P^{\theta_i} = e^{-4}$) was used for known parameters, and a high precision

($P^{\theta_i} = e^{32}$) was used for unknown parameters. The order of generalized motion of $p = 6$ and $d = 2$ were used for the states and inputs respectively. The result for DEM's parameter estimation is shown in Fig. 1. Despite being initialized by random wrong priors, DEM's parameter estimates exponentially converges to the correct values, by maximizing the free energy objective.

Next, we proceed to show that the estimate converges for a wide range of systems. The same experiment was repeated for 25 different randomly selected stable mass-spring-damper systems. Although the convergence applies to unstable systems, sampling was restricted to stable systems within the range $[-1, 1]$ ($\theta_3, \theta_4 \in [-1, 0]$ and $\theta_6 \in [0, 1]$) for better visualization. DEM was initialized with the same priors for all experiments ($\eta^{\theta_6} = 2$, $\eta^{\theta_4} = -1$ and $\eta^{\theta_3} = -2$). Figure 2 shows that DEM is capable of providing converging solutions for a wide range of stable spring-damper systems, that are influenced by colored noise. Note that the numerical analysis is restricted to the dynamics of spring damper systems for demonstrative purposes, and can be extended to other systems. In summary, DEM can provide converging parameter estimates for linear systems with colored noise, by maximizing the free energy objective.

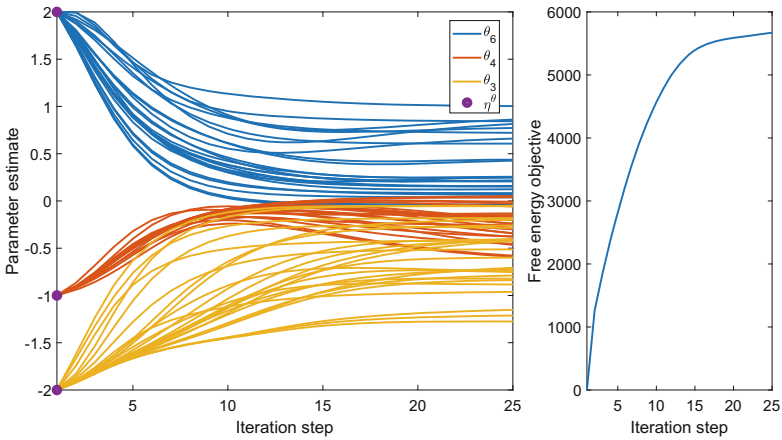


Fig. 2. DEM's parameter estimates for 25 different randomly sampled stable mass-spring-damper systems. The estimates for all the experiments started from the same prior of $\eta^{\theta_6} = 2$, $\eta^{\theta_4} = -1$ and $\eta^{\theta_3} = -2$, and converged, while maximizing the free energy objective. Therefore, the estimator converges for a wide range of systems.

6 Conclusion

DEM has the potential to be a bioinspired learning algorithm for future robots, due to its capability to robustly handle colored noise. Its superior performance in state estimation under colored noise was proven by [11] and was experimentally validated by [4]. In this paper, we derived a mathematical proof of convergence

for DEM's parameter estimator, applied to linear systems with colored noise. We proved that a perception scheme based on the gradient ascend of the free energy action, provides a converging solution. Since a convergence proof is mandatory for the safe and reliable application of DEM on real robots, this work widens its applicability in robotics. The applicability of DEM for real control system problem was demonstrated through rigorous simulations on the estimation problem for mass-spring-damper systems. The future research will focus on the conditions for unbiased estimation and on applying DEM to real robots.

References

1. Meera, A.A., Wisse, M.: A brain inspired learning algorithm for the perception of a quadrotor in wind. arXiv preprint [arXiv:2109.11971](https://arxiv.org/abs/2109.11971) (2021)
2. Anil Meera, A., Wisse, M.: Dynamic expectation maximization algorithm for estimation of linear systems with colored noise. *Entropy* **23**(10), 1306 (2021)
3. Baltieri, M., Buckley, C.L.: PID control as a process of active inference with linear generative models. *Entropy* **21**(3), 257 (2019)
4. Bos, F., Meera, A.A., Benders, D., Wisse, M.: Free energy principle for state and input estimation of a quadcopter flying in wind. arXiv preprint [arXiv:2109.12052](https://arxiv.org/abs/2109.12052) (2021)
5. Çatal, O., Verbelen, T., Van de Maele, T., Dhoedt, B., Safron, A.: Robot navigation as hierarchical active inference. *Neural Netw.* **142**, 192–204 (2021)
6. Friston, K.: Hierarchical models in the brain. *PLoS Comput. Biol.* **4**(11), e1000211 (2008)
7. Friston, K.: The free-energy principle: a unified brain theory? *Nat. Rev. Neurosci.* **11**(2), 127–138 (2010)
8. Friston, K., Mattout, J., Kilner, J.: Action understanding and active inference. *Biol. Cybern.* **104**(1), 137–160 (2011)
9. Friston, K.J., Trujillo-Barreto, N., Daunizeau, J.: DEM: a variational treatment of dynamic systems. *Neuroimage* **41**(3), 849–885 (2008)
10. Mader, W., Linke, Y., Mader, M., Sommerlade, L., Timmer, J., Schelter, B.: A numerically efficient implementation of the expectation maximization algorithm for state space models. *Appl. Math. Comput.* **241**, 222–232 (2014)
11. Meera, A.A., Wisse, M.: Free energy principle based state and input observer design for linear systems with colored noise. In: 2020 American Control Conference (ACC), pp. 5052–5058. IEEE (2020)
12. Oliver, G., Lanillos, P., Cheng, G.: Active inference body perception and action for humanoid robots. arXiv preprint [arXiv:1906.03022](https://arxiv.org/abs/1906.03022) (2019)
13. Pezzato, C., Ferrari, R., Corbato, C.H.: A novel adaptive controller for robot manipulators based on active inference. *IEEE Robot. Autom. Lett.* **5**(2), 2973–2980 (2020)