

Defining information-based functional objectives for neurostimulation and control

Elham Ghazizadeh¹, Peng Yi¹ and ShiNung Ching¹

Abstract—Neurostimulation – the practice of applying exogenous excitation, e.g., via electrical current, to the brain – has been used for decades in clinical applications such as the treatment of motor disorders and neuropsychiatric illnesses. Over the past several years, more emphasis has been placed on understanding and designing neurostimulation from a systems-theoretic perspective, so as to better optimize its use. Particular questions of interest have included designing stimulation waveforms that best induce certain patterns of brain activity while minimizing expenditure of stimulus power. The pursuit of these designs faces a fundamental conundrum, insofar as they presume that the desired pattern (e.g., desynchronization of a neural population) is known *a priori*. In this paper, we present an alternative paradigm wherein the goal of the stimulation is not to induce a prescribed pattern, but rather to simply improve the functionality of the stimulated circuit/system. Here, the notion of functionality is defined in terms of an information-theoretic objective. Specifically, we seek closed loop control designs that maximize the ability of a controlled circuit to encode an afferent ‘hidden input,’ without prescription of dynamics or output. In this way, the control attempts only to make the system ‘effective’ without knowing beforehand the dynamics that are needed to be induced. We devote most of our effort to defining this framework mathematically, providing algorithmic procedures that demonstrate its solution and interpreting the results of this procedure for simple, prototypical dynamical systems. Simulation results are provided for more complex models, including an example involving control of a canonical neural mass model.

I. INTRODUCTION

Neurostimulation – the practice of exciting neural tissue using exogenous inputs (e.g., electrical fields) – is pervasive in both basic and clinical neuroscience [1], [2]. Usually, the objective of neurostimulation is to activate particular population of neurons within the brain and/or to induce a certain pattern of brain activity, such as oscillations. The complex dynamics present in neuronal circuits means that the design of such stimulation is quite nontrivial, especially as it relates to shaping open-loop stimulus waveforms [3], [4] and the design of closed-loop paradigms [5].

This paper deviates from prior efforts in neurostimulation optimization. Rather than attempt to design controls that induce a prescribed pattern, we face the conundrum that in many instances, there simply is no ‘target pattern’ to design

with. That is, we do not know if there is a ‘correct’ pattern we should be inducing. Instead, we might suppose that the actual goal of neurostimulation, especially in clinical contexts, is to improve the *functionality* of the stimulated system in a specifically defined yet general way, not tied to any one pattern of activity. How can such a goal, which is more abstract than a classical ‘tracking’ problem, be approached? Formalizing the answer to our question necessitates defining a notion of functionality that is applicable to neuronal circuits/systems.

To define such an objective, we turn to concepts from theoretical neuroscience and information theory. Indeed, information theory has been used extensively in the study of neural circuit function, under the premise that such circuits are highly efficient at encoding and transforming information toward higher cognitive functions [6]. In this vein, the concept of *empowerment* was first proposed in [6], [7] as a hypothetical, information-based utility function that might account for the emergent architecture and dynamics in sensory and motor neural circuits. Succinctly, empowerment can be thought of as the capacity of a stochastic dynamical system to encode an afferent input [8]. The greater the empowerment, the better a system can represent its input. Consistent with this notion, recently empowerment has been proposed as a general framework by which systems might operate without knowledge of predefined tasks [9], [10] and strategies have been proposed to calculate it for general dynamical systems [11].

In the present work, we approach *neurocontrol* in terms of empowerment, using it explicitly as a measure of a neural circuit’s functionality. Specifically, to improve the functionality of the stimulated system, we investigate closed-loop control designs that maximize the empowerment without additional prescription of the desirable dynamics. We devote most of our paper to present the mathematical framework of using variational bounds for estimation of empowerment as well as the algorithmic procedure. We demonstrate the methodology and the efficiency of the proposed algorithm via an example of a well-known neural mass model, i.e, the Wilson-Cowan model [12].

The remainder of the paper is organized as follows. In section II, we formalize the optimization problem and present the algorithmic procedure for performing the optimization. Section III provides simulation results illustrating the efficacy of empowerment maximization. Section IV concludes the paper.

*This work has been partially supported by grants 1537015 and 1653589 from the National Science Foundation

¹Elham Ghazizadeh, Peng Yi and ShiNung Ching are with the Department of Electrical and Systems Engineering, Washington University in St. Louis, 1 Brookings Drive, St. Louis, MO 63130, USA elham@wustl.edu, yipeng@wustl.edu and shinung@wustl.edu

II. PROBLEM FORMULATION AND LEARNING METHODS

In the following, we provide the dynamical system model as well as mathematical definition of mutual information and specifically empowerment as our objective function. Next, we define a parameterization to modify the dynamical system via optimizing the system empowerment. The optimization method and the algorithmic procedure are also discussed in detail.

A. Dynamical system model

We consider a dynamical system driven by a random input as follows:

$$\dot{\mathbf{x}}(t) = f_{\mathbf{K}}(\mathbf{x}(t)) + \mathbf{u}(t) \quad (1)$$

where $f_{\mathbf{K}}(\cdot) : \mathbb{R}^d \rightarrow \mathbb{R}^d$ denotes the vector field (parameterized by $\mathbf{K}$), while $\mathbf{x}(t) \in \mathbb{R}^d$ and $\mathbf{u}(t) \in \mathbb{R}^d$ are the state and the random input, respectively. We discretize the system in (1) with a fixed interval Δt as follows:

$$\mathbf{x}(t + \Delta t) = \mathbf{x}(t) + (f_{\mathbf{K}}(\mathbf{x}(t)) + \mathbf{u}(t)) \Delta t \quad (2)$$

For ease of notation, we rewrite (2):

$$\mathbf{x}_{t+1} = \mathbf{x}_t + (f_{\mathbf{K}}(\mathbf{x}_t) + \mathbf{u}_t) \Delta t, \quad (3)$$

where we refer to $\mathbf{x}_t$ and $\mathbf{x}_{t+1}$ as the current state and the next state at time t , respectively.

The discrete model in (3) is a one-step autoregressive equation that describes the effect of the input on the state of the system at the subsequent time. Following the application of an impulsive input $\mathbf{u}_t$, we allow the system to evolve (unforced) for some number of future steps. That is, over n time steps we obtain the terminal state as:

$$\mathbf{x}_{t+n} = \mathbf{x}_t + f_{\mathbf{K}}(\dots f_{\mathbf{K}}(f_{\mathbf{K}}(\mathbf{x}_t) + \mathbf{u}_t)) \Delta t \quad (4)$$

where the number of recursions in (4) is n . We further introduce stochasticity to the above dynamics by adding noise to the read-out of the system states, i.e.,

$$\hat{\mathbf{x}}_t = \mathbf{x}_t + \mathbf{w}_t \quad (5)$$

where $\mathbf{w}_t$ is an uncorrelated noise process.

B. Mutual information and empowerment

The Shannon mutual information is a fundamental information-theoretic quantity that measures the statistical dependency between two random variables. Since we assume the input of the system, $\mathbf{u}_t$, is a random vector, the dynamical system (2) can be treated as a stochastic channel between input and system state. Given the current state, the mutual information between the input and the final state is:

$$I(\mathbf{x}_{t+n}; \mathbf{u}_t | \mathbf{x}_t) = \quad (6)$$

$$\iint p(\mathbf{x}_{t+n} | \mathbf{u}_t, \mathbf{x}_t) \omega(\mathbf{u}_t | \mathbf{x}_t) \log \frac{p(\mathbf{x}_{t+n}, \mathbf{u}_t | \mathbf{x}_t)}{p(\mathbf{x}_{t+n} | \mathbf{x}_t) \omega(\mathbf{u}_t | \mathbf{x}_t)} d\mathbf{x}_{t+n} d\mathbf{u}_t$$

where we denote $\omega(\mathbf{u}_t | \mathbf{x}_t)$ as the *input distribution* and $p(\mathbf{x}_{t+n} | \mathbf{u}_t, \mathbf{x}_t)$ as the *transition distribution*. Note that $\omega(\mathbf{u}_t | \mathbf{x}_t)$ can also be denoted as *prior distribution* over the input in the sense that it does not have the knowledge of the final state compared to posterior distribution i.e., $p(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t)$.

Empowerment is defined as the channel capacity [7], [13], or, the maximum information that the input emits to the dynamical system by manipulating the final state. Particularly, for a given current state $\mathbf{x}_t$, the empowerment is:

$$\mathcal{E}(\mathbf{x}_t) = \max_{\omega} I(\mathbf{x}_{t+n}; \mathbf{u}_t | \mathbf{x}_t). \quad (7)$$

Here, $\mathbf{u}_t$ can be seen as a ‘hidden’ exogenous input that is uncorrelated with the dynamics of system. Indeed, maximization of mutual information with respect to all possible distributions of $\mathbf{u}_t$ leads to the optimal encoding of the optimal input within the states of the system. Hence, empowerment quantifies the maximum potential information flow emanating from this input through the system, or equivalently, the channel capacity of the system at hand.

C. Dynamical System Parameterization

Thus far, we have framed empowerment of the system as the maximal mutual information by optimizing the input distribution, $\mathbf{u}_t$. Our problem formulation proceeds to consider the maximization of empowerment by means of parametrization design.

In particular, suppose that $\mathbf{K}$ in (1) represents a parametrization enacted through an external control input. Then, tuning $\mathbf{K}$ could be performed so as to alter the transition distribution in such a way that empowerment is maximized. For a given set of current states, $\mathbf{x}_t^{(m)} \in \{\mathbf{x}_t^{(1)}, \dots, \mathbf{x}_t^{(M)}\}$, we posit this optimization problem as follows:

$$\max_{\mathbf{K}} \frac{1}{M} \sum_m \mathcal{E}(\mathbf{x}_t^{(m)}) \triangleq \max_{\mathbf{K}} \max_{\omega} \frac{1}{M} \sum_m I(\mathbf{x}_{t+n}; \mathbf{u}_t | \mathbf{x}_t^{(m)}). \quad (8)$$

Alternatively, $\mathbf{K}$ can be interpreted as a degree of freedom for the controlled dynamical system that alters the dynamics of system based on the empowerment maximization objective.

D. Learning for empowerment lower bound maximization

The fundamental algorithm to obtain the channel capacity is the Blahut-Arimoto (BA) algorithm [14]. However, since the BA algorithm is based on enumeration, it cannot be applied for evaluating mutual information over the continuous domain of variables. The issue that now arises is that calculating the quantities within (8) is highly challenging. These quantities involve integrating over the continuous domain of all states and inputs. Also, the marginal transition distribution i.e. $p(\mathbf{x}_{t+n} | \mathbf{x}_t)$ in (6) is not tractable. To circumvent the mentioned difficulties, instead of optimizing the exact empowerment, we can maximize a variational lower bound on the value of empowerment [10], [11]. To obtain

this lower bound, we first approximate mutual information by simplifying (6) as:

$$\iint p(\mathbf{x}_{t+n}, \mathbf{u}_t | \mathbf{x}_t) \log \frac{p(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t)}{\omega(\mathbf{u}_t | \mathbf{x}_t)} d\mathbf{x}_{t+n} d\mathbf{u}_t \quad (9)$$

where $p(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t)$ presents the *posterior distribution* over the input once the final state is observed. This posterior distribution, too, is intractable. However, using the variational approximation introduced in [15], mutual information can be approximated as:

$$\begin{aligned} \hat{I}(\mathbf{x}_{t+n}; \mathbf{u}_t | \mathbf{x}_t) = \\ \iint p(\mathbf{x}_{t+n}, \mathbf{u}_t | \mathbf{x}_t) \log \frac{q(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t)}{\omega(\mathbf{u}_t | \mathbf{x}_t)} d\mathbf{x}_{t+n} d\mathbf{u}_t \end{aligned} \quad (10)$$

where $q(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t)$ presents a *variational distribution* to approximate the true posterior distribution $p(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t)$. Using the variational approximation, obtaining the posterior distribution can be considered as an optimization problem, where $q_\xi(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t)$ is a variational family of distributions parameterized by ξ [16]. If the variational distribution expressively represents the true posterior distribution, a tight variational lower bound can be obtained. Specifically, following the same mathematical argument provided in [10]:

$$\begin{aligned} I(\mathbf{x}_{t+n}; \mathbf{u}_t | \mathbf{x}_t) - \hat{I}(\mathbf{x}_{t+n}; \mathbf{u}_t | \mathbf{x}_t) = \\ \iint p(\mathbf{x}_{t+n}, \mathbf{u}_t | \mathbf{x}_t) \log \frac{p(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t)}{\omega(\mathbf{u}_t | \mathbf{x}_t)} d\mathbf{x}_{t+n} d\mathbf{u}_t \\ - \iint p(\mathbf{x}_{t+n}, \mathbf{u}_t | \mathbf{x}_t) \log \frac{q(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t)}{\omega(\mathbf{u}_t | \mathbf{x}_t)} d\mathbf{x}_{t+n} d\mathbf{u}_t = \\ \int p(\mathbf{x}_{t+n} | \mathbf{x}_t) \int p(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t) \log \frac{p(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t)}{q(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t)} d\mathbf{x}_{t+n} d\mathbf{u}_t = \\ \mathbb{E}_{p(\mathbf{x}_{t+n} | \mathbf{x}_t)} [\text{KL}(p(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t) || q(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t))], \end{aligned} \quad (11)$$

where $KL(\cdot)$ is the KL divergence. Therefore, a variational lower bound on empowerment can be achieved as follows:

$$\hat{\mathcal{E}}(\mathbf{x}_t) = \max_{\omega, q} \hat{I}(\mathbf{x}_{t+n}; \mathbf{u}_t | \mathbf{x}_t). \quad (12)$$

It is worthwhile to mention that we can approach the optimization problem in (12) using different methods of approximate inference, such as Markov Chain Monte Carlo (MCMC) method and mean-field approximation. However, the mentioned methods are computationally slow for evaluation of the probability terms in (10) for each given current state $\mathbf{x}_t$.

For the sake of online estimation of empowerment and efficient sampling from the probability distributions, we perform the optimization for ω and q , (12), using multilayer feedforward neural networks. We can obtain the prior and posterior distributions using deep neural networks parameterized by ϕ and ξ , respectively. Using (artificial) multilayer neural networks (NNs) for learning a mapping from

the states to the distribution parameters allows us to only deal with finite number of neural network parameters (i.e. weights and biases) instead of learning a separate prior distribution and variational distribution for each system's state, which shares the similar idea as the amortized variational inference [17]. Here, we choose prior and posterior distributions from the Gaussian family i.e.,

$$\omega(\mathbf{u}_t | \mathbf{x}_t^{(m)}) = \mathcal{N}(\mu_\phi(\mathbf{x}_t^{(m)}), \sigma_\phi^2(\mathbf{x}_t^{(m)}) I)$$

$$q(\mathbf{u}_t | \mathbf{x}_t^{(m)}, \mathbf{x}_{t+n}) = \mathcal{N}(\mu_\xi(\mathbf{x}_t^{(m)}, \mathbf{x}_{t+n}), \sigma_\xi^2(\mathbf{x}_t^{(m)}, \mathbf{x}_{t+n}) I) \quad (13)$$

where mean μ and variance σ are also parameterized by NNs. We denote $\theta = \{\phi, \xi\}$ as joint parameter set which are learned via (12). By joint optimization of variational bound w.r.t. parameters of ω and q , we can obtain the prior distribution that maximizes the mutual information and the variational distribution that is responsible for the tightness of the variational lower bound [10].

Using the empowerment variational lower bound, for a given set of current states, $\mathbf{x}_t^{(m)} \in \{\mathbf{x}_t^{(1)}, \dots, \mathbf{x}_t^{(M)}\}$, we consider our optimization problem as:

$$\max_{\mathbf{K}} \max_{\theta} J(\mathbf{K}, \theta), \quad (14)$$

where

$$J(\mathbf{K}, \theta) = \frac{1}{M} \sum_m \hat{I}(\mathbf{x}_{t+n}; \mathbf{u}_t | \mathbf{x}_t^{(m)}) \quad (15)$$

and

$$\hat{I}(\mathbf{x}_{t+n}; \mathbf{u}_t | \mathbf{x}_t^{(m)}) = \quad (16)$$

$$\mathbb{E}_{p(\mathbf{x}_{t+n}, \mathbf{u}_t | \mathbf{x}_t^{(m)})} [\log q(\mathbf{u}_t | \mathbf{x}_{t+n}, \mathbf{x}_t^{(m)}) - \log \omega(\mathbf{u}_t | \mathbf{x}_t^{(m)})]$$

Since the expectation in (16) is taken over the joint distribution of $\mathbf{u}_t$ and $\mathbf{x}_{t+n}$, i.e. $p(\mathbf{x}_{t+n}, \mathbf{u}_t | \mathbf{x}_t^{(m)})$, we can use joint Monte Carlo sampling method and the reparameterisation-trick proposed in [18], [19] to evaluate the stochastic gradients of the objective function with respect to the decision variables $\mathbf{K}$ and θ as follows:

$$\frac{\partial}{\partial \theta} \hat{I}(\mathbf{x}_{t+n}; \mathbf{u}_t | \mathbf{x}_t^{(m)}) \approx$$

$$\frac{1}{L} \sum_l \frac{\partial}{\partial \theta} [\log q(\mathbf{u}_t^{(l)} | \mathbf{x}_{t+n}^{(l)}, \mathbf{x}_t^{(m)}) - \log \omega(\mathbf{u}_t^{(l)} | \mathbf{x}_t^{(m)})] \quad (17)$$

and

$$\frac{\partial}{\partial \mathbf{K}} \hat{I}(\mathbf{x}_{t+n}; \mathbf{u}_t | \mathbf{x}_t^{(m)}) \approx \quad (18)$$

$$\frac{1}{L} \sum_l \frac{\partial}{\partial \mathbf{K}} [\log q(\mathbf{u}_t^{(l)} | \mathbf{x}_{t+n}^{(l)}, \mathbf{x}_t^{(m)}) - \log \omega(\mathbf{u}_t^{(l)} | \mathbf{x}_t^{(m)})]$$

where L is the number of Monte Carlo samples for estimating $\hat{I}$. Therefore, we formulate the optimization problem posed in (14) as *empowerment variational lower bound maximization* for θ and $\mathbf{K}$.

Algorithm 1 summarizes the optimization procedure. Given that the system's dynamics are known and differentiable, the decision variables can be updated using stochastic gradient ascent.

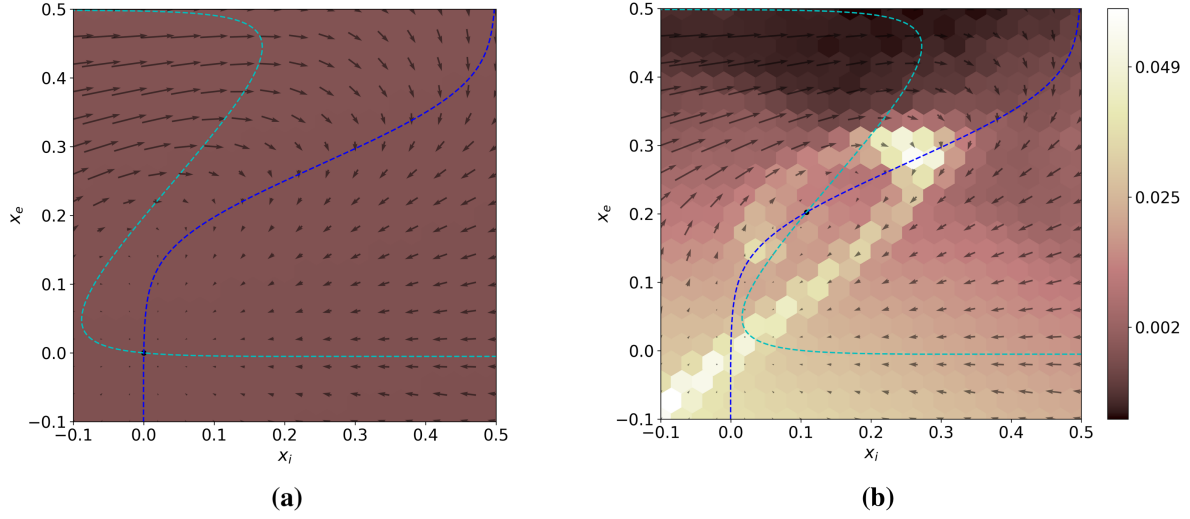


Fig. 1. The vector field, nullclines and empowerment landscape (the empowerment values are in *nats*). The Wilson-Cowan parameters are: $k_j = 1$, $c_1 = 16$, $c_2 = 12$, $c_3 = 15$, $c_4 = 3$, $a_e = 1.3$, $a_i = 2$, $\theta_e = 4$ and $\theta_i = 3.7$. In plot (a), $P = 0$ and there exist a stable fixed point at the origin. In plot (b) $P = 1.25$ and the Wilson-Cowan model manifests a limit cycle.

Algorithm 1 maximization of empowerment variational lower bound, w.r.t. $\theta = \{\phi, \xi\}$ and K

```

initialize uniform samples from state space,
 $\{\mathbf{x}_t^{(m)}\}_{m=1,\dots,M}$ 
while not converged do
  for each  $\mathbf{x}_t^{(m)}$  do
    draw one sample with  $\omega(\mathbf{u}_t|\mathbf{x}_t^{(m)})$ :
     $\mathbf{u}_t \sim \omega(\mathbf{u}_t|\mathbf{x}_t^{(m)})$ 
    transition to the next state  $\mathbf{x}_{t+n}$ 
  end for
   $J = \frac{1}{M} \sum_m \log q(\mathbf{u}_t|\mathbf{x}_t^{(m)}, \mathbf{x}_{t+n}) - \log \omega(\mathbf{u}_t|\mathbf{x}_t^{(m)})$ 
   $\theta \leftarrow \theta + \eta_\theta \nabla_\theta J$ , for  $r$  epochs
   $K \leftarrow K + \eta_K \nabla_K J$ 
end while

```

III. RESULTS

In this section, we apply the problem setup to the problem of neurocontrol. We first present the Wilson-Cowan model and the setup for parametrization of the dynamical system at hand. Next, we proceed by providing results for two scenarios of variational empowerment maximization and provide interpretation of the ensuing results.

A. The Wilson-Cowan Model

We use the well-known Wilson-Cowan model that provides a coarse-grained description of the overall activity of large populations of neurons [12], [20]. In particular, the Wilson-Cowan model represents excitatory and inhibitory activity in synaptically coupled neuronal networks. The two-dimensional system describing the nonlinear dynamics of

these populations is:

$$\begin{aligned}\dot{x}_e &= -x_e + (k_e - r_e x_e) \mathcal{S}_e(c_1 x_e - c_2 x_i + I_e + P) \\ \dot{x}_i &= -x_i + (k_i - r_i x_i) \mathcal{S}_i(c_3 x_e - c_4 x_i + I_i + Q)\end{aligned}\quad (19)$$

where x_e and x_i represent the overall activity in the excitatory and inhibitory populations. Here, k_j , $j \in \{i, e\}$, presents a dimensional constant and r_j is a constant describing the refractory period. c_1 , c_2 , c_3 and c_4 present the strength of excitatory and inhibitory interactions. P and Q are the excitation level in the system. Also, $\mathcal{S}_j$ is a sigmoid function:

$$\mathcal{S}_j(x) = \frac{1}{1 + \exp[-a_j(x - \theta_j)]} - \frac{1}{1 + \exp(a_j \theta_j)} \quad (20)$$

where a and θ are free parameters. Further, I_e and I_i are external control inputs currents that impinge on the respective populations.

We used our variational technique to compute the empowerment for the WC model in two parameter regimes, the first corresponding to a single stable equilibrium and the second corresponding to a stable limit cycle. As shown in Figure 1, the landscapes are markedly different. In the former case, empowerment is almost zero since all trajectories approach the same point asymptotically. In the latter, we see regions of higher empowerment that likely correspond to a high density of isochrons so that nearby points tend asymptotically to spatially disparate locations on the limit cycle.

B. Controlling Empowerment

Our goal now is to apply exogenous input to the system in order to improve its overall empowerment according to the aforementioned algorithm. We consider both feedforward and feedback design scenarios. For feedback, we assume

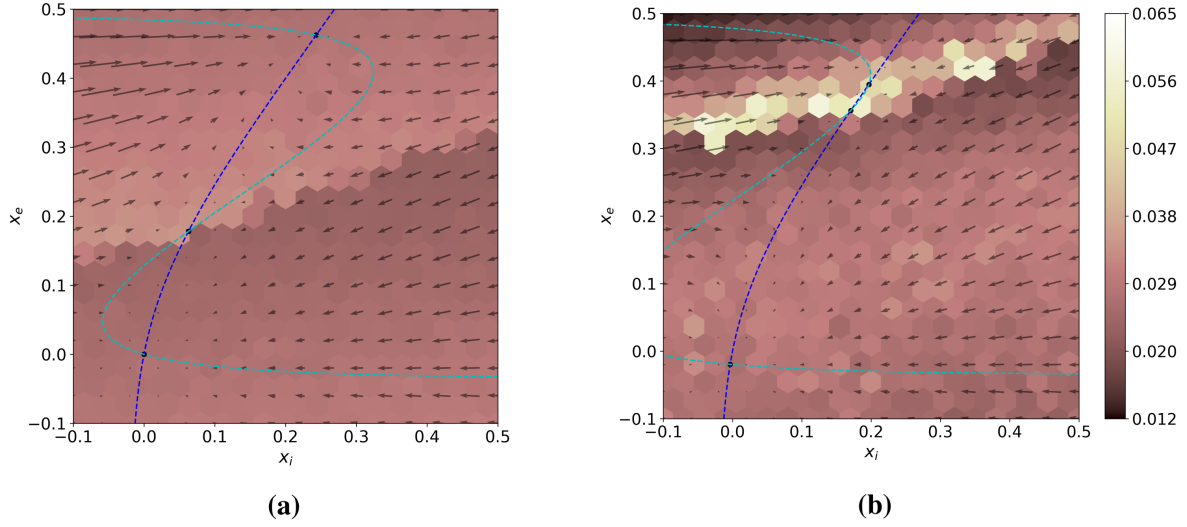


Fig. 2. The vector field, nullclines and empowerment landscape (the empowerment values are in *nats*). The Wilson-Cowan model parameters are: $k_j = 1$, $c_1 = 12$, $c_2 = 4$, $c_3 = 13$, $c_4 = 11$, $a_e = 1.2$, $a_i = 1$, $\theta_e = 2.8$ and $\theta_i = 4$. (a) $P = 0$, (b) $P = -0.5$ (after optimization)

a linear state-feedback via these inputs (i.e., $I_e = k_1 x_e + k_2 x_i$)¹, resulting in the closed loop dynamics:

$$\begin{aligned}\dot{x}_e &= -x_e + (k_e - r_e x_e) \mathcal{S}_e(c_1 x_e - c_2 x_i + k_1 x_e + k_2 x_i + P) \\ \dot{x}_i &= -x_i + (k_i - r_i x_i) \mathcal{S}_i(c_3 x_e - c_4 x_i + k_3 x_e + k_4 x_i + Q)\end{aligned}\quad (21)$$

On the other hand, for the feedforward strategy we simply design I_e as a constant offset to P (We already know from Figure 1 that such a strategy can work).

C. Efficacy of the feedback control via empowerment variational lower bound

We proceed to demonstrate the ability of feedback control to improve the system's empowerment.

First, we study the case wherein the ‘open-loop’ Wilson-Cowan model exhibits an asymptotically stable limit cycle. It is important to note that the dynamical system equation is discretized as mentioned in (4). In terms of aligning the Wilson-Cowan model with (1), note that the stochastic input (i.e., the source, $\mathbf{u}$) is modeled as an additional additive term within $\mathcal{S}_e$ and $\mathcal{S}_i$. Further, $\mathbf{x} = [x_e \ x_i]^T$, respectively.

a) Implementation details: We provide optimization parameters in the following: Optimization is done over 6000 epochs and the learning rates were $\eta_\theta = 0.001$ and $\eta_K = 0.01$. The NNs for representation of the prior and posterior distributions have the same structure; 2 hidden layers, 16 hidden nodes and 2 output nodes (i.e. distribution mean and logarithm of variance). Also, given that we choose the number of samples, i.e., M large enough, L can be set to 1 [18], as presented in Algorithm 1. In our case, we have used

¹This is a strong assumption in the context of current neurostimulation technologies, since it is not easy to independently actuate such populations. Further, obtaining the signals x_e and x_i from recording modalities is far from easy. Our goal here is to illustrate the theoretical framework in this paper, so this assumption should be regarded as mostly pedagogical.

exponential linear units (ELU) as the activation function for training the NNs. For our simulations, we have used the TensorFlow [21] framework to benefit from automatic differentiation and facilitate performing stochastic gradient ascent and backpropagation in the high-dimensional space of decision variables.

b) Open-loop design: We initialized the model to $P = 0$ and ran Algorithm 1 for 10 iterations (as distinct from epochs, see Algorithm 1). Figure 2 illustrates the outcome of one instance of this optimization, where an inhibitory solution is found ($P = -0.5$). This is near a bifurcation point (note proximity of e and i nullclines), leading to a locus of high empowerment relative to the uncontrolled system.

c) Closed-loop design: Figure 3 illustrates the results of Algorithm 2 for the case of the closed loop design (21). Note that here the state space has been downsampled relative to Figure 1. We observe that the design has altered the shape of the nullclines so that now a stable fixed point occurs with very slow damped asymptotic oscillations. The high empowerment regions may reflect initial conditions whose terminal states remain separable over our specified time horizon (i.e. $n = 300$).

IV. CONCLUSIONS BLUE AND DISCUSSION

The paper has introduced a new concept for control of neural systems: empowerment maximization. The idea is that rather than prescribing the output or dynamics of a system *a priori*, we instead define an objective only in terms of its functionality. In the present context, empowerment characterizes the capacity of a system of encode information from afferent inputs, which may be a general surrogate for the efficacy of neuronal networks to transform information.

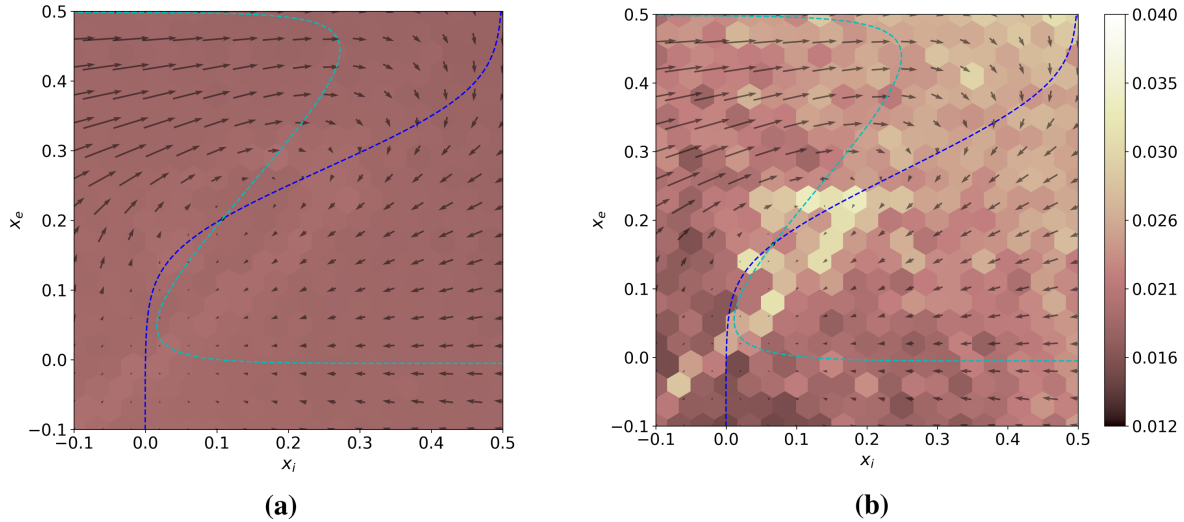


Fig. 3. The vector field, nullclines and empowerment landscape (the empowerment values are in *nats*). The Wilson-Cowan model parameters are: $k_j = 1$, $c_1 = 16$, $c_2 = 12$, $c_3 = 15$, $c_4 = 3$, $a_e = 1.3$, $a_i = 2$, $\theta_e = 4$ and $\theta_i = 3.7$. (a) $k_1 = 0, k_2 = 0, k_3 = 0$ and $k_4 = 0$ (b) $k_1 = -1.41, k_2 = -1.41, k_3 = 1.40$ and $k_4 = 1.42$ (after optimization).

The specific contributions of this paper are theoretical in nature. We formulated the general problem of empowerment maximization and provided the mathematical framework and an algorithmic procedure for optimization by using an empowerment variational lower bound. We demonstrated the results of this procedure on a simple neural mass model and highlighted how empowerment maximization affects the vector field of this model.

These results and interpretations are early steps and there are plenty of open questions about scalability and practicality of the approach. We are especially curious as to whether, there is a general principle that links empowerment to particular forms of dynamics.

REFERENCES

- [1] R. Nardone, Y. Höller, S. Leis, P. Höller, N. Thon, A. Thomschewski, S. Golaszewski, F. Brigo, and E. Trinka, “Invasive and non-invasive brain stimulation for treatment of neuropathic pain in patients with spinal cord injury: a review,” *The journal of spinal cord medicine*, vol. 37, no. 1, pp. 19–31, 2014.
- [2] M. Asllani, P. Expert, and T. Carletti, “A minimally invasive neurostimulation method for controlling abnormal synchronisation in the neuronal activity,” *PLoS computational biology*, vol. 14, no. 7, p. e1006296, 2018.
- [3] J. T. Ritt and S. Ching, “Neurocontrol: Methods, models and technologies for manipulating dynamics in the brain,” in *American Control Conference (ACC), 2015*. IEEE, 2015, pp. 3765–3780.
- [4] W. M. Grill, “Model-based analysis and design of waveforms for efficient neural stimulation,” in *Progress in brain research*. Elsevier, 2015, vol. 222, pp. 147–162.
- [5] L. Grose, J. H. Marshel, and K. Deisseroth, “Closed-loop and activity-guided optogenetic control,” *Neuron*, vol. 86, no. 1, pp. 106–139, 2015.
- [6] A. S. Klyubin, D. Polani, and C. L. Nehaniv, “Keep your options open: an information-based driving principle for sensorimotor systems,” *PloS one*, vol. 3, no. 12, p. e4018, 2008.
- [7] —, “Empowerment: A universal agent-centric measure of control,” in *Evolutionary Computation, 2005. The 2005 IEEE Congress on*, vol. 1. IEEE, 2005, pp. 128–135.
- [8] S. Tiomkin, D. Polani, and N. Tishby, “Control capacity of partially observable dynamic systems in continuous time,” *arXiv preprint arXiv:1701.04984*, 2017.
- [9] C. Salge, C. Glackin, and D. Polani, “Changing the environment based on empowerment as intrinsic motivation,” *Entropy*, vol. 16, no. 5, pp. 2789–2819, 2014.
- [10] M. Karl, M. Soelch, P. Becker-Ehmck, D. Benbouazid, P. van der Smagt, and J. Bayer, “Unsupervised real-time control through variational empowerment,” *arXiv preprint arXiv:1710.05101*, 2017.
- [11] S. Mohamed and D. J. Rezende, “Variational information maximisation for intrinsically motivated reinforcement learning,” in *Advances in neural information processing systems*, 2015, pp. 2125–2133.
- [12] H. R. Wilson and J. D. Cowan, “Excitatory and inhibitory interactions in localized populations of model neurons,” *Biophysical journal*, vol. 12, no. 1, pp. 1–24, 1972.
- [13] C. Salge, C. Glackin, and D. Polani, “Empowerment—an introduction,” in *Guided Self-Organization: Inception*. Springer, 2014, pp. 67–114.
- [14] R. Blahut, “Computation of channel capacity and rate-distortion functions,” *IEEE transactions on Information Theory*, vol. 18, no. 4, pp. 460–473, 1972.
- [15] D. B. F. Agakov, “The im algorithm: a variational approach to information maximization,” *Advances in Neural Information Processing Systems*, vol. 16, p. 201, 2004.
- [16] R. Ranganath, S. Gerrish, and D. Blei, “Black box variational inference,” in *Artificial Intelligence and Statistics*, 2014, pp. 814–822.
- [17] D. Ritchie, P. Horsfall, and N. D. Goodman, “Deep amortized inference for probabilistic programs,” *arXiv preprint arXiv:1610.05735*, 2016.
- [18] D. P. Kingma and M. Welling, “Auto-encoding variational bayes,” *arXiv preprint arXiv:1312.6114*, 2013.
- [19] D. J. Rezende, S. Mohamed, and D. Wierstra, “Stochastic backpropagation and approximate inference in deep generative models,” *arXiv preprint arXiv:1401.4082*, 2014.
- [20] H. R. Wilson and J. D. Cowan, “A mathematical theory of the functional dynamics of cortical and thalamic nervous tissue,” *Kybernetik*, vol. 13, no. 2, pp. 55–80, 1973.
- [21] M. Abadi, P. Barham, J. Chen, Z. Chen, A. Davis, J. Dean, M. Devin, S. Ghemawat, G. Irving, M. Isard *et al.*, “Tensorflow: a system for large-scale machine learning,” in *OSDI*, vol. 16, 2016, pp. 265–283.
- [22] H. Touchette and S. Lloyd, “Information-theoretic approach to the study of control systems,” *Physica A: Statistical Mechanics and its Applications*, vol. 331, no. 1-2, pp. 140–172, 2004.