
Causal Inference in the Closed-Loop: Marginal Structural Models for Sequential Excursion Effects

Alexander W. Levis*
Carnegie Mellon University
alevis@cmu.edu

Gabriel Loewinger*
National Institutes of Health
gloewinger@gmail.com

Francisco Pereira
National Institutes of Health
francisco.pereira@nih.gov

Abstract

Optogenetics is widely used to study the effects of neural circuit manipulation on behavior. However, the paucity of causal inference methodological work on this topic has resulted in analysis conventions that discard information, and constrain the scientific questions that can be posed. To fill this gap, we introduce a nonparametric causal inference framework for analyzing “closed-loop” designs, which use dynamic policies that assign treatment based on covariates. In this setting, standard methods can introduce bias and occlude causal effects. Building on the sequentially randomized experiments literature in causal inference, our approach extends history-restricted marginal structural models for dynamic regimes. In practice, our framework can identify a wide range of causal effects of optogenetics on trial-by-trial behavior, such as, fast/slow-acting, dose-response, additive/antagonistic, and floor/ceiling. Importantly, it does so without requiring negative controls, and can estimate how causal effect magnitudes evolve across time points. From another view, our work extends “excursion effect” methods—popular in the mobile health literature—to enable estimation of causal contrasts for treatment sequences greater than length one, in the presence of positivity violations. We derive rigorous statistical guarantees, enabling hypothesis testing of these causal effects. We demonstrate our approach on data from a recent study of dopaminergic activity on learning, and show how our method reveals relevant effects obscured in standard analyses.

1 Introduction

Optogenetics is a neuroscience technique to “turn on/off” neurons *in vivo* in real-time, with millisecond time resolution. It works by shining lasers on neurons that have been genetically modified through viral infection to express a light-sensitive protein. It is one of the most popular assays with roughly 700 references to it in 2023 alone.² Optogenetics is often applied to study the causal effect of manipulating specific brain circuits while animals (e.g., mice) perform behavioral tasks to study, for example, learning and decision-making. These tasks are typically composed of a sequence of trials, $t \in \{1, 2, \dots, T\}$, each of which involves presentation of stimuli and an opportunity for a behavioral response. For example, a trial might begin with a cue (e.g., a light), which indicates that a lever press will trigger delivery of a food reward. Investigators might want to know, for instance, whether applying optogenetic stimulation on a random subset of trials alters the rate at which mice press the lever. On trial t , an animal’s behavioral outcome, Y_t , time-varying covariates, X_t , and optogenetic (treatment) indicator, A_t , are observed. Experiments often include both treatment ($G = 1$) and negative-control ($G = 0$) groups, with animals assigned randomly to each. While the laser (i.e., the sequential treatment, A_t) is often applied on a random³ subset of trials in

* Authors contributed equally, names alphabetized

²699 papers on Web of Science mention “optogenetics” in 2023 and between 526-796 in years 2015-2023.

³Some studies deterministically set laser/no-laser trials but we focus on “stochastic” experimental designs.

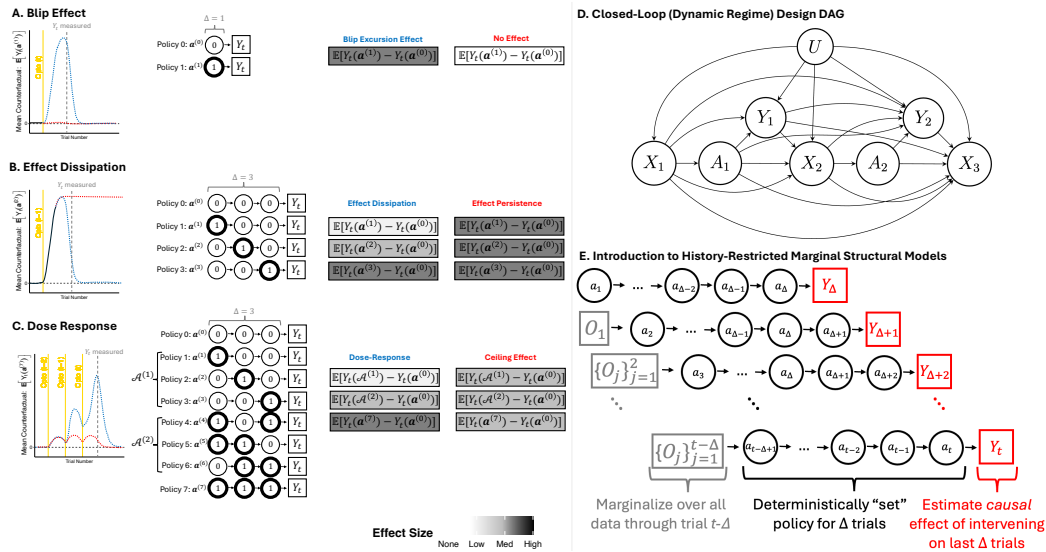


Figure 1: *Sequential Excursion Effects*. [A]-[C] The left panels show one setting where a sequence of laser simulations **do** or **do not** have the indicated effect on the outcome. The middle panel shows deterministic static policies that could be used to construct a causal contrast to probe the effect. The right panel shows what the anticipated effect size (darker is larger) of the contrast might be if the effect **was** or **was not** present. [A] **Blip Effect**: the effect of a single stimulation vs. no treatment on a recent trial. [B] **Effect Dissipation**: Whether the effect of a single stimulation causes an effect that rises and dissipates after a few trials, or persists. [C] **Dose Response**: Do successive stimulations increase the response in a dose-dependent fashion? [D] Closed-loop design DAG for two trials. U is an unmeasured variable. [E] HR-MSM illustration inspired by figure in [6].

both groups, only treatment group animals ($G = 1$) express the protein that enables the laser to trigger the target neural response. The control group thus controls for “off-target” effects such as the laser heating the brain, and the optogenetic insertion surgery. To answer the question above, investigators often estimate the effect of optogenetic manipulation through comparisons such as $\psi_t = \mathbb{E}[Y_t | G = 1] - \mathbb{E}[Y_t | G = 0]$. It is common to test whether $\psi_t = 0$ at specific timepoints like the end of the study ($t = T$), or to conduct inference on summaries (e.g., $\bar{\psi} = \frac{1}{T} \sum_t \psi_t$). These *between-group* comparisons assess the intervention impact based on simple long-term, or “macro”/“global” longitudinal effects. When studies randomly deliver treatment at each trial (i.e., with stochastic policies), *within-group* comparisons between laser and no-laser trials are also common (e.g., $\tilde{\psi} = \sum_t \{\mathbb{E}[Y_t | A_t = 1, G = 1] - \mathbb{E}[Y_t | A_t = 0, G = 1]\}$).

Importantly, such comparisons do not lend themselves to testing *within-group* “micro”/“local” longitudinal effects related to specific treatment sequence patterns. For example, one might ask whether there is a dose-dependent relationship between the outcome and the number of stimulations in the last five trials, or whether stimulation on two consecutive trials has a synergistic effect that is greater than if stimulation instead occurred on two non-consecutive trials. Figures 1A-C shows some representative micro longitudinal effects that are identifiable in many optogenetics studies, yet typically are not explored. Critically, such effects may be present even in studies in which one fails to detect the macro effects commonly tested. However, no formal causal inference framework has been applied to these studies, resulting in analysis conventions that limit the scope of questions researchers can ask.

Furthermore, certain experimental designs can complicate the use and interpretation of even standard analysis approaches. In “closed-loop” (referred to as “dynamic regimes” in the causal inference literature) designs, stimulation is applied depending on the behavior of the animal. For example, say a study tests if lever pressing for food, Y_t , decreases if optogenetic stimulation ($A_t = 1$) is applied, with positive probability, only when animals approach the lever ($X_t = 1$). Since A_t is randomized conditional on X_t , one must incorporate X_t into their analysis, but standard strategies like including X_t as a covariate in a regression can obscure effects and induce bias. This is because 1) X_t influences the probability of *both* the outcome and treatment, and thus can be cast as a time-varying confounder, and 2) X_t *also* mediates the effect of prior treatments [7]; see the illustrative DAG in Figure 1D, though note that we generalize this setting later on to allow treatment to depend on the

complete history of previously measured variables at each time point. In this case, since treatment also influences both Y_t and X_t on subsequent trials, closed-loop designs induce “treatment–confounder feedback” [7], which can lead to bias with standard analyses. We include an example in Appendix B to show how, when the treatment has opposing effects on Y_t and X_t , treatment and control groups can exhibit *identical* average (observed) outcome levels even if the laser causes a large immediate effect. Furthermore, standard regression approaches can actually induce collider-bias, and block mediators of the treatment effect [7]. Finally, if treatment policies deterministically rule out treatment (e.g., when $X_t = 0$), certain effects are not identifiable: the positivity violation inherent to these designs precludes estimation of certain counterfactual distributions. Closed-loop designs therefore require specialized algorithms for valid causal inference.

More broadly, there have been a number of high profile calls for more rigorous definitions of causality and causal inference in neuroscience [1, 33, 4, 16]. However, to the best of our knowledge, existing methodological work [25, 10, 14] focuses on instrumental variable-based approaches to estimate causal effects of optogenetics on neural activity. Unlike our setting, these methods are restricted to datasets that include both measurements of the activity of the neurons stimulated by optogenetics, and the neurons those cells interact with. One can then conceptualize the neural activity of the stimulated neurons as treatment variables, and the optogenetics sequence as instruments. In addition to focusing on behavioral outcomes, we explicitly deal with sequentially randomized (and closed-loop) designs, whereas prior work treats each trial as an exchangeable draw, ignoring the sequential nature of trials.

Our contributions are (1) proposing the first formal counterfactual-based causal framing of these behavioral optogenetics designs, (2) developing an analysis framework based on history-restricted marginal structural models that enables the estimation of “sequential excursion effects” that capture the local causal contrasts described above, (3) expanding excursion effect methodology to account for positivity violations, and to accommodate treatment sequences greater than length one, (4) providing estimators with efficient computational implementations and strong theoretical guarantees under minimal nonparametric conditions (verified in simulations), and (5) applying our methods to data from a high profile *Nature* paper, and showing how they reveal effects obscured by standard methods.

2 Notation and Related Work

In this section, we (i) provide the necessary notation and a brief review of relevant work, and (ii) describe the key methodological gap in the current literature: existing methods cannot estimate causal effects of proximal treatment sequences longer than one timepoint in closed-loop designs.

Notation Let $\mathcal{O}_t = \{X_t, A_t, Y_t\}$ be the vector of *observed* variables for an animal on trial t . We denote T as the number of trials and $[T]$ as the set $\{1, 2, \dots, T\}$. A sample of subjects $i = 1, 2, \dots, n$ is collected but, as subjects are exchangeable, we often suppress indices to reduce notational burden. We express *counterfactual* variables, or *potential outcomes*, with parentheses. For example, $Y_t(\mathbf{a}_t)$, represents the potential outcome that would be observed at trial t if a subject received the treatment sequence, $\mathbf{a}_t = (a_1, \dots, a_t)$. Overbars represent all history up to and including a given trial. For example, $\overline{B}_j = (B_1, \dots, B_j)$, for any sequence of variables $\{B_t\}_{t=1}^T$, and any $j \in [T]$. Finally, we define $H_t = (\overline{X}_t, \overline{A}_{t-1}, \overline{Y}_{t-1})$, so H_t includes all information prior to the treatment “decision” at t .

Relevant Literature *Marginal structural models* (MSMs) are often used to model the mean counterfactuals $\mathbb{E}[Y_t(\mathbf{a}_t)]$ [28, 30, 29] in sequentially randomized experiments, though these typically do not perform well [21] when there are a large number of time points (e.g., as in many optogenetics studies): the variance of the model coefficients can grow prohibitively large. *History-restricted* MSMs [21] (HR-MSMs) model $\mathbb{E}[Y_t(\mathbf{a}_{\Delta,t})]$ for some $\mathbf{a}_{\Delta,t} = (a_{t-\Delta+1}, \dots, a_t)$, typically with $\Delta \ll t$. That is, HR-MSMs model the mean counterfactual outcome at time t , under an intervention defined on a proximal (often short) treatment sequence. As $\mathbb{E}[Y_t(\mathbf{a}_{\Delta,t})] = \mathbb{E}[Y_t(\overline{A}_{t-\Delta}, \mathbf{a}_{\Delta,t})]$, by a consistency assumption, these estimands implicitly marginalize over the observed treatment sequence, $\overline{A}_{t-\Delta}$, prior to the first point of intervention. However, any Markov-like assumptions made by the causal framework follow directly from the experimental design: HR-MSMs (and, by extension, our proposed methods) allow for X_t, Y_t to be causally affected by *all* prior trials (i.e., $\mathcal{O}_j$ for $j \in [t-1]$). By placing structure on $\mathbb{E}[Y_t(\mathbf{a}_{\Delta,t})]$, the HR-MSM can borrow strength across treatment sequences $\mathbf{a}_{\Delta,t}$, which can increase power when there are many trials. Figure 1E provides a graphical illustration of HR-MSMs. HR-MSMs are typically fit by using generalized estimating equations (GEE) with

inverse probability of treatment weighting (IPW). IPW resolves the dilemma with standard regression techniques in sequentially randomized experiments, outlined above, where failure to condition on time-varying confounders, X_t , biases estimates (as treatment is randomized conditional on X_t in closed-loop designs), but conditioning on X_t induces confounding (X_t are colliders on the path between past treatments and subsequent outcomes, through unmeasured confounders, U , as shown in the DAG in Figure 1D)) [7]. HR-MSMs can also incorporate time-varying effect modifiers (e.g., see [23] and references), to test, for example, whether causal effects vary across trials, or animal-specific covariate levels.

The gaps: Sequential effects and positivity violations In designs that assign treatment randomly at each trial, HR-MSMs can be used to estimate the causal effect of specific deterministic treatment sequences $\mathbf{a}_{\Delta,t}$ that may differ from the observed sequence $\bar{A}_t$ close to trial t , and are compatible with the experimental treatment rule (“policy”). Importantly, this enables estimation of interpretable causal parameters, such as the effect of treatment on the most recent trial, $\mathbb{E}[Y_t(a_t = 1) - Y_t(a_t = 0)]$. These causal contrasts have grown popular recently in the analysis of mobile health studies [5], where they are referred to as “excursion effects.” However, current methods are restricted to estimating excursion effects for the $\Delta = 1$ case in experimental designs like ours, and thus preclude estimation of effects defined only for $\Delta > 1$ (e.g., the micro longitudinal effects in Figures 1 and 4). Mobile health studies often include treatment rules with positivity violations: due to ethical or practical constraints, treatment must be withheld in certain cases (e.g., no phone notifications while driving). [5] use the notation that treatment is withheld when the time-varying “availability” indicator, I_t , equals zero. Similarly, in “closed-loop” optogenetics experiments, $I_t = 1$ when the conditions are met such that neural manipulation may occur (e.g., when the animal approaches the lever in the example in Section 1). There have been proposals for methods intended to account for such implied positivity violations [19, 5, 26], such as the availability-conditional estimand [5]: $\mathbb{E}[Y_t(a_t = 1) - Y_t(a_t = 0) \mid I_t = 1]$. However, estimands proposed for these settings are defined only for $\Delta = 1$. Thus, in the presence of these positivity violations, there is currently no methodology to conduct causal inference for longer proximal treatment sequences. We note that machine learning based causal methods including causal transformers [18], counterfactual recurrent networks [3], and recurrent marginal structural networks [15] are comparable to HR-MSMs that condition on all measured variables prior to the first intervention timepoint ($t - \Delta + 1$). These methods target effects of static treatment sequences and require a positivity assumption, and thus cannot be applied in closed-loop designs. They also do not provide tools for statistical inference.

3 Methods

To fill the gaps identified above, we propose HR-MSMs for proximal sequences of dynamic treatment regimes, designed to be compatible with treatment availability restrictions in this scientific context. These estimands are defined for any $\Delta \geq 1$, can incorporate time-varying effect modifiers, and can dissect more intricate patterns of treatment over time, compared to standard excursion effects.

3.1 HR-MSMs for Dynamic Treatment Regimes

Adopting the notation from [5], we define $I_t := \mathbb{1}(\mathbb{P}[A_t = 1 \mid H_t] > 0)$ as an “availability indicator”, i.e., $I_t = 0$ if and only if active treatment (e.g., laser stimulation) is prohibited by design. Define $\mathcal{D}_t = \{d_t : \mathcal{H}_t \rightarrow \{0, 1\} \mid d_t(H_t) = 0 \text{ if } I_t = 0\}$, for any t , to be the class of treatment rules at time t compatible with I_t . In particular, we will consider the deterministic rules $\mathcal{D}_t^* = \{d_t^{(0)}, d_t^{(1)}\} \subset \mathcal{D}_t$, where $d_t^{(0)} \equiv 0$, $d_t^{(1)} \equiv I_t$. In words, $d_t^{(0)}$ fixes $A_t = 0$, and $d_t^{(1)}$ sets A_t equal to I_t . The treatment rules $d_t^{(0)}, d_t^{(1)} \in \mathcal{D}_t$ represent the two most extreme policies whose effects remain identifiable. We can combine these time-specific rules to construct multiple time-point analogs of excursion effects compatible with availability restrictions: for $\Delta \in \mathbb{N}$, we let $\bar{\mathcal{D}}_{\Delta,t}$ be a subset of $\mathcal{D}_{t-\Delta+1}^* \times \dots \times \mathcal{D}_t^*$, taking $\mathbf{d}_{\Delta,t} = (d_{t-\Delta+1}, \dots, d_t) \in \bar{\mathcal{D}}_{\Delta,t}$ to be a sequence of Δ treatment rules (compatible with availability restrictions) for trials $j \in \{t - \Delta + 1, \dots, t\}$. The counterfactual outcome under this policy sequence is defined to be

$$Y_t(\mathbf{d}_{\Delta,t}) = Y_t(A_1, \dots, A_{t-\Delta}, d_{t-\Delta+1}(H_{t-\Delta+1}), \dots, d_t(H_t(\mathbf{d}_{\Delta-1,t-1}))). \quad (1)$$

That is, $Y_t(\mathbf{d}_{\Delta,t})$ is the counterfactual outcome under an intervention that leaves the natural value of treatment for the first $t - \Delta$ trials, then sequentially determines treatment by applying $d_{t-\Delta+j}$ to $H_{t-\Delta+j}(\mathbf{d}_{j-1,t-\Delta+j-1})$, for $j \in [\Delta]$, where $\mathbf{d}_{j-1,t-\Delta+j-1} = (d_{t-\Delta+1}, \dots, d_{t-\Delta+j-1})$.

Letting $V_t \subseteq H_t$ be a set of effect modifiers at trial t , we seek to estimate $\mathbb{E}[Y_t(\mathbf{d}_{\Delta,t}) \mid V_{t-\Delta+1}]$, the counterfactual mean outcome, conditional on effect modifiers that are observed *before* the treatment decision of trial $t - \Delta + 1$. By construction, these estimands are identifiable under standard causal assumptions (see Section 3.2). We discuss their interpretation, and compare with existing proposals in Appendix C.1. When $\Delta > 1$ and studies have many trials, there may be many potential treatment rule sequence combinations. We thus propose to estimate effects of these interventions with an MSM on the (conditional) means of the counterfactuals (1): $m(t, \mathbf{d}_{\Delta,t}, V_{t-\Delta+1}; \beta) \approx \mathbb{E}[Y_t(\mathbf{d}_{\Delta,t}) \mid V_{t-\Delta+1}]$, where m is a fixed known function. We aim to conduct inference on the MSM parameters, β , but we do not assume that the model is necessarily well-specified, and thus treat the MSM parameters as projections onto the working model m [20, 32]:

$$\beta_0 = \arg \min_{\beta \in \mathbb{R}^q} \sum_{t=\Delta}^T \sum_{\mathbf{d}_{\Delta,t} \in \overline{\mathcal{D}}_{\Delta,t}} \mathbb{E} \left(h(t, \mathbf{d}_{\Delta,t}, V_{t-\Delta+1}) \{Y_t(\mathbf{d}_{\Delta,t}) - m(t, \mathbf{d}_{\Delta,t}, V_{t-\Delta+1}; \beta)\}^2 \right), \quad (2)$$

for some fixed non-negative weight function h . This projection approaches lies between a fully parametric strategy, that assumes m is correctly specified, and a fully nonparametric approach, that places no structure across the target causal quantities. The target β_0 is defined as the parameter of the best fitting working model m (i.e., closest in $L_2(\mathbb{P})$). In practice, the choice between considering m as a working model or as a correctly specified model amounts to a trade-off between bias and variance—see the discussions in [13, 12] where analogous projection parameters are proposed.

3.2 Identification and Estimation

In this section, we first describe the causal assumptions under which the effects of interest are identified. We then develop an inverse probability-weighted estimator of the MSM parameters, and derive their asymptotic properties. While we focus on dynamic regime HR-MSMs below, our results also apply to static regime HR-MSMs in the case that there are no availability issues (i.e., $I_t \equiv 1$). There, the treatment rule $\mathbf{d}_{\Delta,t}$ reduces to a corresponding static sequence $\mathbf{a}_{\Delta,t}$.

For each t , define the treatment probability function $\pi_t(a; H_t) := \mathbb{P}[A_t = a \mid H_t]$. We make the following standard assumptions, which are expected to hold in many optogenetics designs:

Assumption 3.1. *Consistency:* $Y_t(\mathbf{d}_{\Delta,t}) = Y_t$, whenever $A_j = d_j(H_j)$, for all $j \in \{t-\Delta+1, \dots, t\}$

Assumption 3.2. *Positivity:* For all $t \in \{\Delta, \dots, T\}$, and $d_t \in \mathcal{D}_t^*$, $\pi_t(d_t(H_t); H_t) \geq \epsilon$, w.p. 1

Assumption 3.3. *Sequential randomization:* $A_s \perp\!\!\!\perp Y_t(\mathbf{d}_{\Delta,t}) \mid H_s$, for all $t \in \{\Delta, \dots, T\}$, $s \in \{t - \Delta + 1, \dots, t\}$

We provide a detailed discussion of these assumptions in practice in Appendix C.2. The following result says that these three assumptions are sufficient for identification of the counterfactual means $\mathbb{E}[Y_t(\mathbf{d}_{\Delta,t}) \mid V_{t-\Delta+1}]$, and of the MSM parameters β_0 .

Proposition 3.4. *Under Assumptions 3.1–3.3, we have*

$$\mathbb{E}(Y_t(\mathbf{d}_{\Delta,t}) \mid V_{t-\Delta+1}) = \mathbb{E}_{\mathbb{P}} \left(\prod_{j=t-\Delta+1}^t \frac{\mathbb{1}(A_j = d_j(H_j))}{\pi_j(A_j; H_j)} Y_t \mid V_{t-\Delta+1} \right).$$

Recall that $Z_i = \{\mathcal{O}_{t,i}\}_{t=1}^T$ is the totality of data observed on subject i ; suppressing subject-specific index for clarity, define $\phi(Z, \cdot) : \mathbb{R}^q \rightarrow \mathbb{R}^q$ via

$$\begin{aligned} \phi(Z, \beta) = \sum_{t=\Delta}^T \sum_{\mathbf{d}_{\Delta,t} \in \overline{\mathcal{D}}_{\Delta,t}} h(t, \mathbf{d}_{\Delta,t}, V_{t-\Delta+1}) M(t, \mathbf{d}_{\Delta,t}, V_{t-\Delta+1}; \beta) \\ \times \left[\prod_{j=t-\Delta+1}^t \frac{\mathbb{1}(A_j = d_j(H_j))}{\pi_j(A_j; H_j)} \right] \{Y_t - m(t, \mathbf{d}_{\Delta,t}, V_{t-\Delta+1}; \beta)\}, \end{aligned}$$

where $M(t, \mathbf{d}_{\Delta,t}, V_{t-\Delta+1}; \beta) = \nabla_{\beta} m(t, \mathbf{d}_{\Delta,t}, V_{t-\Delta+1}; \beta)$. Then, assuming the solution to (2) is unique, and the working model m is differentiable in β , the MSM parameters β_0 are identified through the estimating equation $\mathbf{0} = \mathbb{E}_{\mathbb{P}}(\phi(Z, \beta_0))$.

The result of Proposition 3.4 is a population inverse probability-weighted estimating equation for the target parameters β_0 . This estimating equation motivates a corresponding IPW estimator, $\hat{\beta}$, solving the empirical IPW estimating equation $\mathbb{P}_n[\phi(Z, \hat{\beta})] = \mathbf{0}$. In the optogenetics applications of interest, the propensity scores π_t are known by design, and can be plugged in when estimating $\hat{\beta}$.

We now prove asymptotic normality of the our estimator, $\hat{\beta}$, under mild conditions. We require the following notation: define $\mathbf{A}(\beta) = \mathbb{E}[\phi(Z, \beta)\phi(Z, \beta)^T]$ and $\mathbf{B}(\beta) = \mathbb{E}[\nabla_{\beta} \phi(Z, \beta)]$.

Theorem 3.5. *Suppose Assumptions 3.1–3.3 and the following conditions hold:*

- (i) *The minimizer β_0 in (2) is unique;*
- (ii) *$m(t, \mathbf{d}_{\Delta, t}, V_{t-\Delta+1}; \beta)$ is Donsker in β , continuously differentiable at β_0 , uniformly in $V_{t-\Delta+1}$;*
- (iii) *In a neighborhood around β_0 , $\mathbf{A}(\beta)$ and $\mathbf{B}(\beta)$ are finite-valued, and $\mathbf{B}(\beta)$ is non-singular;*
- (iv) *$\hat{\beta} \xrightarrow{P} \beta_0$.*

Then $\sqrt{n}(\hat{\beta} - \beta_0) \xrightarrow{d} \mathcal{N}(\mathbf{0}, \mathbf{V}(\beta_0))$, where $\mathbf{V}(\beta) = \mathbf{B}(\beta)^{-1} \mathbf{A}(\beta) \mathbf{B}(\beta)^{-1}$.

Theorem 3.5 gives the asymptotic distribution of the estimator $\hat{\beta}$. The conditions (i)–(iv) are relatively mild; see Appendix C.3 for a discussion. Theorem 3.5 provides a strategy to construct asymptotically valid Wald-based confidence intervals for the MSM parameters β_0 : for any β we can take

$$\hat{\mathbf{A}}(\beta) = \mathbb{P}_n[\phi(Z, \beta)\phi(Z, \beta)^T], \quad \hat{\mathbf{B}}(\beta) = \mathbb{P}_n[\nabla_{\beta} \phi(Z, \beta)],$$

and define $\hat{\mathbf{V}} = \hat{\mathbf{B}}(\hat{\beta})^{-1} \hat{\mathbf{A}}(\hat{\beta}) \hat{\mathbf{B}}(\hat{\beta})^{-1}$, which is consistent for $\mathbf{V}(\beta_0)$. Then, for $j \in [q]$, an $(1 - \alpha)$ confidence interval for $\beta_{j,0}$ is given by $\hat{\beta}_j \pm z_{1-\alpha/2} \sqrt{\hat{V}_{jj}/n}$, where $z_{1-\alpha/2}$ is the $(1 - \alpha/2)$ -quantile of the standard normal distribution, and $\hat{V}_{jj}$ is the j -th diagonal element of $\hat{\mathbf{V}}$. Confidence intervals for any linear combination of the β parameters can be constructed in a similar fashion.

We provide an implementation that builds the necessary dataset (with each observation copied once for every regime in $\overline{\mathcal{D}}_{\Delta, t}$), calculates the corresponding IPW weights, and estimates the HR-MSM parameters β by solving the estimating equation in expression 2 using the `rootSolve` R package [11]. The process takes about 10 seconds on a standard laptop, for $> 100,000$ total (pre-copy) trials.

4 Experiments

4.1 Simulation Studies

Experimental Setup We sought to assess performance of the proposed estimator, and identify variance estimators that yield nominal coverage in the small n settings common in optogenetics studies. To evaluate the accuracy of our framework in estimating mean counterfactuals, we designed the simulations such that the target estimands—contrasts of mean counterfactuals—corresponded to regression coefficients from the true HR-MSM. The data were simulated to mimic closed-loop optogenetics designs with positivity violations: we drew i) $X_0 \sim \text{Bernoulli}(1/2)$; ii) $A_t \mid X_t \sim \text{Bernoulli}(\frac{1}{2}X_t)$, for $t \in \{0, \dots, T\}$; iii) $X_t \mid A_{t-1} \sim \text{Bernoulli}(0.4 + 0.4A_{t-1})$, for $t \in [T]$; and iv) $Y_t \mid X_{t-1}, A_{t-1}, X_t, A_t \sim \mathcal{N}(\alpha_1 X_{t-1} + \alpha_2 A_{t-1} + \alpha_3 X_t + \alpha_4 A_t, \sigma_t^2)$, for $t \in [T]$, where $(\alpha_1, \alpha_2, \alpha_3, \alpha_4) = (0.25, 2, 1.75, 0.5)$, and $\sigma_t^2 = 1$ for all t . These set availability indicator $I_t \equiv X_t$, for all t , and result in marginal probabilities $\mathbb{P}[X_t = 1] = \frac{1}{2}$, $\mathbb{P}[A_t = 1] = \frac{1}{4}$, for all t . We obtain a closed form for the parameters of the saturated two time-point dynamic treatment regime HR-MSM: letting $\mathbf{d}_{2,t} = (d_{t-1}, d_t) \in \overline{\mathcal{D}}_{2,t}$ be arbitrary, and defining $J_{t-1} := \mathbb{1}(d_{t-1} \equiv d_{t-1}^{(1)})$, $J_t := \mathbb{1}(d_t \equiv d_t^{(1)})$, we derive in Appendix D.1 that $\mathbb{E}(Y_t(\mathbf{d}_{2,t})) = \beta_0 + \beta_1 J_{t-1} + \beta_2 J_t + \beta_3 J_{t-1} J_t$, where $\beta_0 = 0.5\alpha_1 + 0.4\alpha_3$, $\beta_1 = 0.5\alpha_2 + 0.2\alpha_3$, $\beta_2 = 0.4\alpha_4$, and $\beta_3 = 0.2\alpha_4$. Aggregating $\beta = (\beta_0, \beta_1, \beta_2, \beta_3)$, the HR-MSM given by

$$m(t, \mathbf{d}_{2,t}; \beta) = \beta_0 + \beta_1 J_{t-1} + \beta_2 J_t + \beta_3 J_{t-1} J_t \quad (3)$$

is correctly specified under this data generating process, and we can evaluate the performance of the proposed estimator relative to these true values.

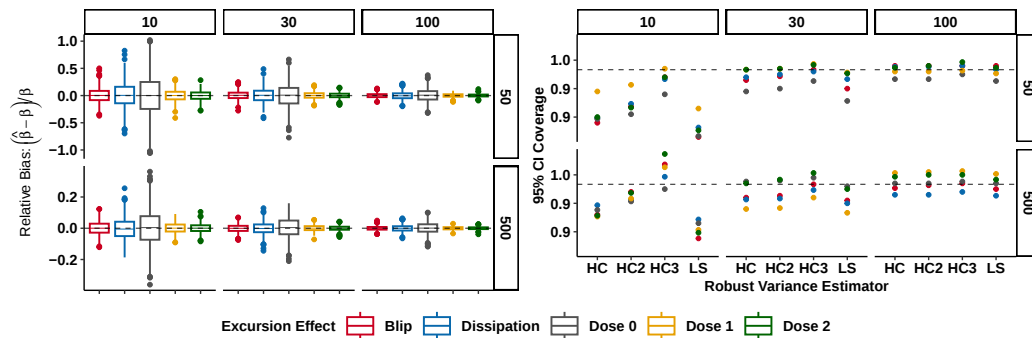


Figure 2: **Simulation Results** Panel columns indicate sample sizes, n (10, 30, or 100), and panel rows indicate number of trials, T (cluster sizes, 50 or 500). [Left] Relative bias associated with each sequential excursion effect. These results show that our estimator is consistent for the target parameters. [Right] 95% Confidence interval (CI) coverage for the sequential excursion effects. The coverage of 95% CIs constructed using one of three established robust variance estimators and our robust large sample (shown as *LS*) variance estimator. The nominal coverage is reached for either large n or large t for all estimators.

To show we can conduct valid inference on sequential excursion effects, we estimated the three estimands illustrated in Figure 1: (1) the “blip” effect of an additional exposure opportunity at the more proximal trial t , while keeping treatment at $t - 1$ fixed at the control condition ($\beta_2 = \mathbb{E}[Y_t(d_{t-1}^{(0)}, d_t^{(1)}) - Y_t(d_{t-1}^{(0)}, d_t^{(0)})]$); (2) “effect dissipation,” comparing the effect of an exposure opportunity at one versus two trials prior to the measurement of the outcome ($\beta_2 - \beta_1 = \mathbb{E}[Y_t(d_{t-1}^{(0)}, d_t^{(1)}) - Y_t(d_{t-1}^{(1)}, d_t^{(0)})]$); and (3) the “dose response” curve of exposure opportunities, where the two possible sequences for a single opportunity are averaged (the sequence $(\beta_0, \beta_0 + \frac{1}{2}\{\beta_1 + \beta_2\}, \sum_{j=0}^3 \beta_j)$). This setup also illustrates how HR-MSMs are easily specified such that sequential excursion effects can be calculated as linear combinations of the β parameters.

Although the HR-MSM (3) is correctly specified, we still estimated the β coefficients as projection parameters. We proceeded as if we started by defining β as the minimizers in (2), with $V_{t-\Delta+1} = \emptyset$ (i.e., no effect modifiers), and $h(t, \mathbf{d}_{\Delta,t}) \equiv 1$ (i.e., constant weight function). We applied the IPW point estimator for the HR-MSM parameters, $\hat{\beta}$, described in Section 3.2. We assessed the coverage of 95% CIs constructed from our large sample variance estimator, and the small sample size-adjusted HC, HC2, and HC3 variance estimators (using the *sandwich* package in R [36]). We tested performance with sample size $n \in \{6, 10, 30, 100\}$, and number of trials $T \in \{10, 50, 500\}$.

Results Simulation results in Figures 2 and Appendix Figure 6 show that our estimator is unbiased for the target sequential excursion effects. We present 95% CI coverage in Table 1, and HR-MSM coefficient estimate bias and MSE in Appendix Tables 2 and 3. Sandwich estimators can yield small-sample bias [36], but in small n and T settings, the sample size-adjusted HC3-based CIs achieve 95% coverage. All CIs achieve 95% coverage when n is large, showing that we can conduct valid inference in all settings.

4.2 Application: Optogenetic Study

Data [17] tested whether optogenetically stimulating dopamine (DA) release in the dorsolateral striatum while an animal engaged in a specific “pose” (e.g., exploring, rearing, grooming) could “teach” mice to exhibit that movement more frequently. This study was foundational in identifying the role this region plays in learning. To that end, the researchers implanted mice with optogenetics machinery, and filmed them freely-moving in a behavioral chamber. They used a pre-trained hidden Markov model to estimate an animal’s pose online in real-time. They first measured the animals’ target pose frequency on a baseline session without optogenetics. Then, on a subsequent treatment session, they applied the laser on a random subset of the target pose occurrences. They repeated this experiment for six target poses, in both the optogenetics and control (laser has no effect) groups.

Effect	CI	$T = 50$				$T = 500$			
		$n = 6$	$n = 10$	$n = 30$	$n = 100$	$n = 6$	$n = 10$	$n = 30$	$n = 100$
Blip	HC	0.94 ± 0.01	0.94 ± 0.01	0.95 ± 0.01	0.95 ± 0.01	0.99 ± 0.00	0.97 ± 0.01	0.95 ± 0.01	0.95 ± 0.01
	LS	0.86 ± 0.01	0.88 ± 0.01	0.93 ± 0.01	0.95 ± 0.01	0.86 ± 0.01	0.89 ± 0.01	0.93 ± 0.01	0.94 ± 0.01
Dissip	HC	0.93 ± 0.01	0.94 ± 0.01	0.95 ± 0.01	0.95 ± 0.01	0.97 ± 0.01	0.96 ± 0.01	0.94 ± 0.01	0.94 ± 0.01
	LS	0.85 ± 0.01	0.89 ± 0.01	0.94 ± 0.01	0.95 ± 0.01	0.88 ± 0.01	0.91 ± 0.01	0.93 ± 0.01	0.94 ± 0.01
Dose 0	HC	0.92 ± 0.01	0.92 ± 0.01	0.94 ± 0.01	0.94 ± 0.01	0.94 ± 0.01	0.94 ± 0.01	0.96 ± 0.01	0.95 ± 0.01
	LS	0.86 ± 0.01	0.88 ± 0.01	0.92 ± 0.01	0.94 ± 0.01	0.86 ± 0.01	0.91 ± 0.01	0.95 ± 0.01	0.95 ± 0.01
Dose 1	HC	0.95 ± 0.01	0.95 ± 0.01	0.96 ± 0.01	0.95 ± 0.01	1.00 ± 0.00	0.97 ± 0.01	0.94 ± 0.01	0.96 ± 0.01
	LS	0.88 ± 0.01	0.91 ± 0.01	0.95 ± 0.01	0.95 ± 0.01	0.86 ± 0.01	0.90 ± 0.01	0.92 ± 0.01	0.96 ± 0.01
Dose 2	HC	0.95 ± 0.01	0.94 ± 0.01	0.96 ± 0.01	0.96 ± 0.01	1.00 ± 0.00	0.98 ± 0.00	0.96 ± 0.01	0.96 ± 0.01
	LS	0.86 ± 0.01	0.89 ± 0.01	0.95 ± 0.01	0.95 ± 0.01	0.87 ± 0.01	0.90 ± 0.01	0.94 ± 0.01	0.96 ± 0.01

Table 1: **Simulation Results: CI Coverage** We achieve 95% confidence interval (CI) coverage using either small sample size-adjusted *HC3* (shown as *HC*), or our large sample (shown as *LS*) sandwich variance estimators. Mean of $R = 1000$ replicates is shown ($\pm$ standard error). We recommend *HC3* when n is low. When n is high, *LS* achieves nominal coverage, confirming our asymptotic theory.

To define “trials,” the authors spliced the time-series of estimated pose classifications into intervals of consecutive timepoints with the same pose classification. If mice exhibited the target pose on trial t , they were considered “available” for optogenetic stimulation, $I_t = 1$, and were “unavailable” otherwise, $I_t = 0$. The laser was applied ($A_t = 1$) with the dynamic policy, $\mathbb{P}(A_t = 1 | I_t) = 0.75I_t$. Denoting Y_t^0 and Y_t as a binary indicator that an animal engaged in the target pose on trial t of the *baseline* and *treatment* sessions, respectively, the authors estimated treatment effects of the form $\psi = (\mathbb{E}[\bar{Y}^1 | G = 1] - \mathbb{E}[\bar{Y}^0 | G = 1]) - (\mathbb{E}[\bar{Y}^1 | G = 0] - \mathbb{E}[\bar{Y}^0 | G = 0])$ where $\bar{Y}^0 = \sum_{t=1}^{T_0} Y_t^0$, $\bar{Y}^1 = \sum_{t=1}^T Y_t$, and $T, T_0 \in \mathbb{N}$ are the trial numbers in treatment and baseline sessions, respectively.⁴ There were $n_1 = 28$ and $n_0 = 12$ animals in the optogenetics and control groups, respectively. T ranged across animals/sessions from 1207-4876, with a mean of 3612 and IQR = [3341, 3940]. The authors reported a (pooled across target poses) positive optogenetics treatment effect estimate akin to $\hat{\psi}$, suggesting DA stimulation causes an increase in target pose frequency.

We argue this analysis procedure leaves many scientific questions untested. Conceptualizing optogenetics like a “study drug,” we question whether stimulation immediately “taught” the animal the target pose, or whether the treatment effect on learning had a lagged onset. Similarly, did the effect of a single stimulation persist or dissipate across trials? Did more treatments lead to more learning monotonically, or is there an antagonistic effect or non-monotonic dose-response curve in learning?

Application Methods We applied our framework to provide a nuanced trial-by-trial characterization of the causal effects of DA stimulation, and formally answer the questions above. Specifically, we tested the causal effect of specific sequences of deterministic dynamic policies, $d_{\Delta,t}$ (occurring on trials $t \in \{t - \Delta + 1, \dots, t\}$), on the mean counterfactual $\mathbb{E}[Y_t(d_{\Delta,t})]$. We defined the outcome Y_t as an indicator that the mouse exhibited the target pose on trial $t + 2$, the next trial on which mice could exhibit the target pose if they were available for stimulation on trial t . We fit a set of HR-MSMs illustrating that we can reliably estimate the types of excursion effects in Figure 1. We describe the models we fit below, and relegate code, data and pre-processing details to Appendix Section F.

Results Our first question was whether standard methods reveal significant treatment effects when assessed with the estimands commonly tested in optogenetics studies. We applied a GEE with mean model, $\log(\mathbb{E}[\bar{Y}^s | G = g, S = s]) = \gamma_0 + \gamma_1 g + \gamma_2 s + \gamma_3 g \times s$, where $S \in \{0, 1\}$ indicates baseline and optogenetics sessions, respectively. The estimate $\hat{\gamma}_3$, shown in Figure 3F, thus provides a treatment effect estimate for the *observed* stochastic dynamic policy in [17]. We adopted a Poisson working model, since [17] analyzed $\bar{Y}^1, \bar{Y}^0 \in \mathbb{N}$. We tested these (macro longitudinal) effects for each pose individually, rather than pooling over them, as in [17]. The model yielded no significant effects for any individual pose. We show boxplots in Appendix Figure 10 of the subject-level summary $\bar{Y}^1 - \bar{Y}^0$ that is compared across groups in this model. Outcome levels are similar across groups for most poses, further highlighting how standard outcome summaries can obscure effects.

To assess an analogous “local” treatment effect using our method, we tested the impact of a single stimulation opportunity. We further evaluated whether the effect had a lagged onset and/or dissipated

⁴The authors used a Mann Whitney U Test applied to a summary across poses but, in keeping with the mean counterfactual-based causal estimands, we describe it in terms of means (not medians) and individual poses.

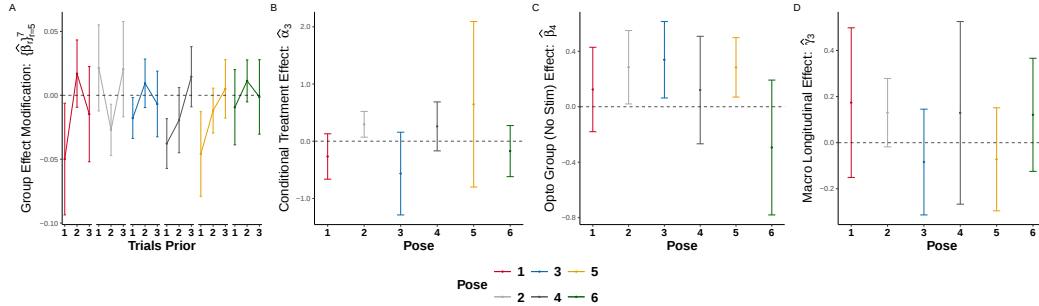


Figure 3: Optogenetics Analyses. Plots show coefficient estimates (error bars show 95% CIs). Columns/colors indicate the target pose. [A] Interaction term between G and sequential excursion effect of a single “dose” occurring $r = 1, 2$, or 3 trials prior to the proximal outcome that the “dose” occurred on. The excursion effects are significant for poses 1-5 (at, at least, one lag level). [B] Availability conditional estimate of interaction $A \times G$: laser $\times$ group interaction. [C] Main effect of G under a “no-recent-treatment opportunity” policy; this reflects the average causal effect of group among a population that has received no laser opportunities in the last $\Delta = 3$ trials. [D] Macro longitudinal analysis, similar to original paper, identifies no significant effects.

across trials. We included group, G , as an effect modifier, to test whether the causal effect of this treatment opportunity was larger in one of the groups. Setting $\Delta = 3$, and restricting the regimes of interest to those with at most one treatment opportunity “dose”, $\mathbf{d}_{3,t} \in \{(d_{t-2}, d_{t-1}, d_t) : \sum_{j=t-2}^t \sigma_j(d_j) \leq 1\}$, where $\sigma_j(d_j) = \mathbb{1}(d_j = d_j^{(1)})$, we fit the HR-MSM

$$\text{logit}(\mathbb{E}[Y_t(\mathbf{d}_{\Delta,t}) \mid G = g]) = \beta_0 + \sum_{r=0}^2 \beta_{r+1} \sigma_{t-r}(d_{t-r}) + \beta_4 g + \sum_{r=0}^2 \beta_{5+r} g \times \sigma_{t-r}(d_{t-r}). \quad (4)$$

Thus, $\hat{\beta}_r$ with $r \in [3]$ is an estimate of the log odds ratio comparing the mean counterfactual of Y_t under a treatment sequence with a single dose (on $r = 1, 2$, or 3 trials prior) vs. a treatment sequence with zero dose. This permits assessment of effect dissipation or persistence. Figure 1B illustrates the analogous effect under a static regime. The interaction terms, $\hat{\beta}_r$ with $r \in \{5, 6, 7\}$, quantify how these causal effects of a recent treatment opportunity differ between the two groups.

The results from our model (4) reveal that stimulation opportunities in the treatment group tend to *reduce* the odds of the outcome, compared to the control group. As shown in Figure 3C, these effects are significantly negative for at least one lag level in five out of six target poses. In personal communications, the authors of [17] stated that this result appeared consistent with their finding that animal exploration increased right after stimulation (quantified as higher pose entropy). Figure 3D shows the main effect of group under a treatment sequence of dose zero. In essence, this provides an estimate of the “long-term” effect of DA stimulation: $\hat{\beta}_4$ is the log odds ratio of treatment group under a regime of dose zero (i.e., a “no recent stimulation” policy). We fit comparable models for sequences as long as $\Delta = 7$ and found results were similar across Δ values.

Next, we fit the analogous model for the availability-conditional estimand [5] to determine whether current excursion effect methods (i.e., those confined to $\Delta = 1$ policies) identify the same treatment effects: $\text{logit}(\mathbb{E}[Y_t(a_t) \mid I_t = 1, G = g]) = \alpha_0 + \alpha_1 a_t + \alpha_2 g + \alpha_3 g \times a_t$. Figure 3C shows that the effect estimates, $\hat{\alpha}_3$, are significant in only one pose. These results highlight how our approach can uncover a greater number of significant effects.

Finally, in Appendix Section E, we include analyses showing that the laser exhibits a dose-response curve in both groups: more treatment opportunities (on some poses) in the last $\Delta = 5$ trials causes the animal to exhibit the target pose more often. Additionally, there is significant effect modification by baseline responding: the laser has a larger effect in animals who exhibited high baseline pose frequency. Together, these results show we can reliably estimate sequential excursion effects.

5 Discussion

We propose the first, to the best of our knowledge, formal causal inference framework for closed-loop optogenetics behavioral studies. We introduce a nonparametric excursion effect framework, an associated IPW estimator (with valid CIs), with a scalable implementation, and proved its consistency

and asymptotic normality under mild assumptions. Methodologically, our proposed sequential excursion effects represent an expansion of the conditional estimands proposed in [5] to longitudinal policies ($\Delta \geq 1$), in the presence of positivity violations. Our methods also directly apply to “open-loop” (static policies) designs, as they arise as a special case when $I_t = 1$ for all $t \in [T]$.

HR-MSMs are powerful and useful models, but have their limitations. As has been discussed in the causal inference literature, these estimands marginalize over all treatments for trials $t \in [t - \Delta]$, and thus depend on the protocol used in the design [31, 6]. Moreover, while contrasts of our estimands are null under the sharp null of no causal effect of treatment (e.g., optogenetic laser stimulation), effects should generally still be interpreted in terms of treatment *opportunities*. Finally, while our implementation is computationally efficient, we anticipate computational challenges for very large Δ .

The model m and the number of intervention timepoints Δ represent key choices for practitioners. Our inferential results (i.e., Theorem 3.5) are valid for a large class of working mean models m , and notably do not rely on any distributional assumptions. The “Donsker” requirement (condition (ii) in Theorem 3.5) is satisfied outright by generalized linear models such as those we use in our application [35], as well as some formulations of random forests [34] and kernel estimators [2]. In future work, we will study how inference can be obtained for more flexible models that do not satisfy the Donsker assumption. Likewise, the value of Δ plays a significant role as it determines the nature of the effects being estimated, and should be chosen on the basis of subject matter expertise. That said, we found that 2-3 intervention timepoints are often sufficient to capture a rich set of sequential excursion effects, and in our application that results were relatively stable across a range of Δ values.

The application highlights the drawbacks of standard optogenetics analysis methods. Our finding that macro longitudinal estimates of individual target poses show almost no effect between groups highlights how “treatment–confounder” feedback can obscure strong treatment effects in closed-loop designs, even when inspecting simple averages of observed outcomes. Our methods account for this by careful causal adjustment with IPW. In personal communications with the authors of [17], they agreed with our findings and remarked at how these methods reveal a collection of causal effects that are difficult to uncover without sophisticated causal inference methods.

Our analyses reveal immediate *negative* effects (detectable on the next trial) and *positive* slower effects of DA stimulation (i.e., in treatment relative to control animals). We also find the control group exhibits positive, off-target effects of the laser. Together the *opposing* signs of these “fast”/“slow” and on/off-target causal effects may further dilute the magnitude of macro longitudinal effects that summarize the outcome across many trials (e.g., total pose counts). Finally, by enabling estimation of *sequential* excursion effects (i.e., $\Delta > 1$), we can reveal effect profiles (e.g., dose-response curves) not possible with availability-conditional estimands whose definition is confined to $\Delta = 1$ regimes. As we observed, the optogenetics group sometimes exhibits an excursion effect not present in the control group. Thus, by combining different sequential excursion effects, analysts can, for example, disentangle laser on-target from off-target effects. When off-target effects are not a major concern, our framework enables estimation of causal effects *without* having to collect data in a control group, thereby potentially reducing the number of animals required in a study.

Although we focus on optogenetics here, our proposed methods are relevant for a wide range of mobile health, neuroscience and psychology experiments for which the “local/micro” longitudinal structure is of scientific interest. Indeed, “closed-loop” designs are common in many behavioral studies in human neuroimaging and cognitive sciences (e.g., when stimuli are conditionally randomized). We hope our methods constitute a useful methodological contribution to the causal inference literature, and will help applied researchers exploit the rich information contained in their experiments.

Acknowledgements

This research was supported by the Intramural Research Program of the National Institute of Mental Health (NIMH), project ZIC-MH002968. This study utilized the high-performance computational capabilities of the Biowulf Linux cluster at the National Institutes of Health, Bethesda, MD (<http://biowulf.nih.gov>). First, we would like to thank the authors of “Spontaneous behaviour is structured by reinforcement without explicit reward,” Drs. Jeffrey Markowitz and Sandeep Robert Datta, for sharing their data, helping us conduct analyses, and interpret the results. This work would not have been possible without their generosity, commitment to open science and scientific rigor. We would also like to thank the NIMH Machine Learning Team for helpful feedback on our project.

References

- [1] D. L. Barack, E. K. Miller, C. I. Moore, A. M. Packer, L. Pessoa, L. N. Ross, and N. C. Rust. A call for more clarity around causality in neuroscience. *Trends in neurosciences*, 45(9):654–655, 2022.
- [2] E. Beutner and H. Zähle. Donsker results for the empirical process indexed by functions of locally bounded variation and applications to the smoothed empirical process. *Bernoulli*, 29(1):205–228, 2023.
- [3] I. Bica, A. M. Alaa, J. Jordon, and M. van der Schaar. Estimating counterfactual treatment outcomes over time through adversarially balanced representations. *arXiv preprint arXiv:2002.04083*, 2020.
- [4] R. Biswas and E. Shlizerman. Statistical perspective on functional and causal neural connectomics: A comparative study. *Frontiers in Systems Neuroscience*, 16:817962, 2022.
- [5] A. Boruvka, D. Almirall, K. Witkiewitz, and S. A. Murphy. Assessing time-varying causal effect moderation in mobile health. *Journal of the American Statistical Association*, 113(523):1112–1121, 2018.
- [6] F. R. Guo, T. S. Richardson, and J. M. Robins. Discussion of ‘Estimating time-varying causal excursion effects in mobile health with binary outcomes’. *Biometrika*, 108(3):541–550, 08 2021.
- [7] M. Hernan and J. Robins. *Causal Inference: What If*. Chapman & Hall/CRC Monographs on Statistics & Applied Probab. CRC Press, 2023.
- [8] P. J. Huber. Robust estimation of a location parameter. *The Annals of Mathematical Statistics*, pages 73–101, 1964.
- [9] P. J. Huber et al. The behavior of maximum likelihood estimates under nonstandard conditions. In *Proceedings of the fifth Berkeley symposium on mathematical statistics and probability*, volume 1, pages 221–233. Berkeley, CA: University of California Press, 1967.
- [10] Z. Jiang, S. Chen, and P. Ding. An instrumental variable method for point processes: generalized wald estimation based on deconvolution. *Biometrika*, page asad005, 2023.
- [11] Karline Soetaert. *rootSolve: Nonlinear root finding, equilibrium and steady-state analysis of ordinary differential equations*, 2009. R package 1.6.
- [12] E. Kennedy, S. Balakrishnan, and L. Wasserman. Semiparametric counterfactual density estimation. *Biometrika*, page asad017, 2023.
- [13] E. H. Kennedy, S. Lorch, D. S. Small, et al. Robust causal inference with continuous instruments using the local instrumental variable curve. *Journal of the Royal Statistical Society Series B*, 81(1):121–143, 2019.
- [14] M. E. Lepperød, T. Stöber, T. Hafting, M. Fyhn, and K. P. Kording. Inferring causal connectivity from pairwise recordings and optogenetics. *PLoS Computational Biology*, 19(11):e1011574, 2023.
- [15] B. Lim. Forecasting treatment responses over time using recurrent marginal structural networks. *Advances in Neural Information Processing Systems*, 31, 2018.
- [16] I. E. Marinescu, P. N. Lawlor, and K. P. Kording. Quasi-experimental causality in neuroscience and behavioural research. *Nature Human Behaviour*, 2(12):891–898, 2018.
- [17] J. E. Markowitz, W. F. Gillis, M. Jay, J. Wood, R. W. Harris, R. Cieszkowski, R. Scott, D. Brann, D. Koveal, T. Kula, C. Weinreb, M. A. M. Osman, S. R. Pinto, N. Uchida, S. W. Linderman, B. L. Sabatini, and S. R. Datta. Spontaneous behaviour is structured by reinforcement without explicit reward. *Nature*, 614(7946):108–117, 2023.

- [18] V. Melnychuk, D. Frauen, and S. Feuerriegel. Causal transformer for estimating counterfactual outcomes. In *ICML*, pages 15293–15329. PMLR, 2022.
- [19] K. Moore, R. Neugebauer, F. Lurmann, J. Hall, V. Brajer, S. Alcorn, and I. Tager. Ambient ozone concentrations cause increased hospitalizations for asthma in children: An 18-year study in southern california. *Environmental Health Perspectives*, 116(8):1063–1070, 2008.
- [20] R. Neugebauer and M. van der Laan. Nonparametric causal effects based on marginal structural models. *Journal of Statistical Planning and Inference*, 137(2):419–434, 2007.
- [21] R. Neugebauer, M. J. van der Laan, M. M. Joffe, and I. B. Tager. Causal inference in longitudinal studies with history-restricted marginal structural models. *Electronic journal of statistics*, 1:119, 2007.
- [22] W. K. Newey. Uniform convergence in probability and stochastic equicontinuity. *Econometrica: Journal of the Econometric Society*, pages 1161–1167, 1991.
- [23] M. L. Petersen, S. G. Deeks, J. N. Martin, and M. J. van der Laan. History-adjusted Marginal Structural Models for Estimating Time-varying Effect Modification. *American Journal of Epidemiology*, 166(9):985–993, 09 2007.
- [24] D. Pollard. *Convergence of stochastic processes*. Springer Science & Business Media, 2012.
- [25] D. Pospisil, M. Aragon, and J. Pillow. From connectome to effectome: learning the causal interaction map of the fly brain. *bioRxiv*, 2023.
- [26] T. Qian, H. Yoo, P. Klasnja, D. Almirall, and S. A. Murphy. Estimating time-varying causal excursion effects in mobile health with binary outcomes. *Biometrika*, 108(3):507–527, 09 2020.
- [27] J. Robins. A new approach to causal inference in mortality studies with a sustained exposure period—application to control of the healthy worker survivor effect. *Mathematical modelling*, 7(9-12):1393–1512, 1986.
- [28] J. M. Robins. Marginal structural models versus structural nested models as tools for causal inference. In *Statistical models in epidemiology, the environment, and clinical trials*, pages 95–133. Springer, 2000.
- [29] J. M. Robins, S. Greenland, and F.-C. Hu. Estimation of the causal effect of a time-varying exposure on the marginal mean of a repeated binary outcome. *Journal of the American Statistical Association*, 94(447):687–700, 1999.
- [30] J. M. Robins, M. A. Hernan, and B. Brumback. Marginal structural models and causal inference in epidemiology. *Epidemiology*, pages 550–560, 2000.
- [31] J. M. Robins, M. A. Hernán, and A. Rotnitzky. Invited Commentary: Effect Modification by Time-varying Covariates. *American Journal of Epidemiology*, 166(9):994–1002, 09 2007.
- [32] M. Rosenblum and M. J. van der Laan. Targeted maximum likelihood estimation of the parameter of a marginal structural model. *The international journal of biostatistics*, 6(2), 2010.
- [33] L. N. Ross and D. S. Bassett. Causation in neuroscience: keeping mechanism meaningful. *Nature Reviews Neuroscience*, pages 1–10, 2024.
- [34] E. Scornet. On the asymptotics of random forests. *Journal of Multivariate Analysis*, 146:72–83, 2016.
- [35] A. W. Van der Vaart. *Asymptotic statistics*, volume 3. Cambridge university press, 2000.
- [36] A. Zeileis, S. Köll, and N. Graham. Various versatile variances: An object-oriented implementation of clustered covariances in r. *Journal of Statistical Software*, 95(1):1–36, 2020.

A Micro Longitudinal Effects

In Appendix Figure 4, we illustrate additional micro longitudinal effects that can be probed with our sequential excursion effect framework. This figure has the same layout as Figure 1A-C.

B Illustration of Treatment-Confounder Feedback

We provide in this section a synthetic example in which two independent groups exhibit identical mean outcome patterns over time, but where the treatment (e.g., turning on laser in the brain) has a substantial effect in one group but not the other. As our construction will demonstrate, this phenomenon manifests due to treatment-confounder feedback leading to effects canceling out. In a similar fashion, one can similarly construct scenarios where effects are exaggerated.

Suppose $G \in \{0, 1\}$ represents an experimentally manipulable marker (e.g., animals expressing opsin in the brain), and counterfactual outcomes under $G = g$ are denoted Y_t^g . We will suppose that potential outcomes generated in the active setting ($G = 1$) are given by

$$Y_t^1 \sim \mathcal{N}(\gamma_{0t} + \gamma_1 X_{t-1} + \gamma_2 A_{t-1} + \gamma_3 X_t + \gamma_4 A_t, \sigma_t^2),$$

and potential outcomes in the control condition ($G = 0$) are given by

$$Y_t^0 \sim \mathcal{N}(\gamma_{0t} + \gamma_1 X_{t-1} + \gamma_3 X_t, \sigma_t^2),$$

i.e., the treatment (e.g., laser) has an effect when $G = 1$, but not when $G = 0$.

Suppose further that a behavior X_t is measured at all time points t , and determines whether or not treatment will be administered with positive probability. Like the outcomes, this behavior will be affected by the laser only when $G = 1$:

$$X_t^g \sim \text{Bernoulli}(0.7 - 0.5 A_{t-1} g), \text{ for } t \in \{1, \dots, T\},$$

and $X_0^g \sim \text{Bernoulli}(\frac{1}{2})$ at baseline.

Now we consider a study where animals are randomly assigned at baseline to either $G = 1$ or $G = 0$. At each time point t , the behavior X_t is measured, and treatment is then drawn according to $A_t \sim \text{Bernoulli}(0.8X_t)$. By induction, $\mathbb{E}(A_t | G = 1) = 0.4$ and $\mathbb{E}(X_t | G = g) = 0.7 - 0.2g$, for all t . It follows that

$$\begin{aligned} \mathbb{E}(Y_t | G = g) &= \gamma_{0t} + \gamma_1 \mathbb{E}(X_{t-1} | G = g) + \gamma_2 \mathbb{E}(A_{t-1} | G = 1)g + \gamma_3 \mathbb{E}(X_t | G = g) + \gamma_4 \mathbb{E}(A_t | G = 1)g \\ &= \{\gamma_{0t} + 0.7(\gamma_1 + \gamma_3)\} + \{-0.2(\gamma_1 + \gamma_3) + 0.4(\gamma_2 + \gamma_4)\}g. \end{aligned}$$

Thus, the “macro”/“global” between-group mean difference trajectory is given by

$$\mathbb{E}(Y_t | G = 1) - \mathbb{E}(Y_t | G = 0) = -0.2(\gamma_1 + \gamma_3) + 0.4(\gamma_2 + \gamma_4),$$

which will be null if $\gamma_2 + \gamma_4 = 0.5(\gamma_1 + \gamma_3)$. Notice that this cancellation is possible even if the immediate effect of treatment on the outcome is quite strong, say if γ_2 and γ_4 are large and positive. The cancellation is made possible through the opposing effects of treatment on the intermediate behavior and the outcome: when $G = 1$, A_{t-1} negatively impacts X_t but positively impacts Y_t . More generally, these X_t - A_t feedback loops can lead to dilution or exaggeration of the actual effect of treatments when only analyzing observed mean outcomes.

We note that in the data generating scenario described in this appendix, the proposed dynamic treatment regime HR-MSM methodology would pick out non-null effects of treatment within the active group ($G = 1$), and show differing effects between groups, even if the condition above held such that observed mean outcomes were identical. This example thus serves to illustrate both the challenges with closed-loop designs and, despite these challenges, the ability of the proposed methodology to elucidate effects.

Example Analysis on Synthetic Data To illustrate the above, we provide an example on a simulated dataset, taking $n = 100$, $T = 500$, $\gamma_1 = \gamma_3 = 1$, $\gamma_2 = \gamma_4 = 0.5$.

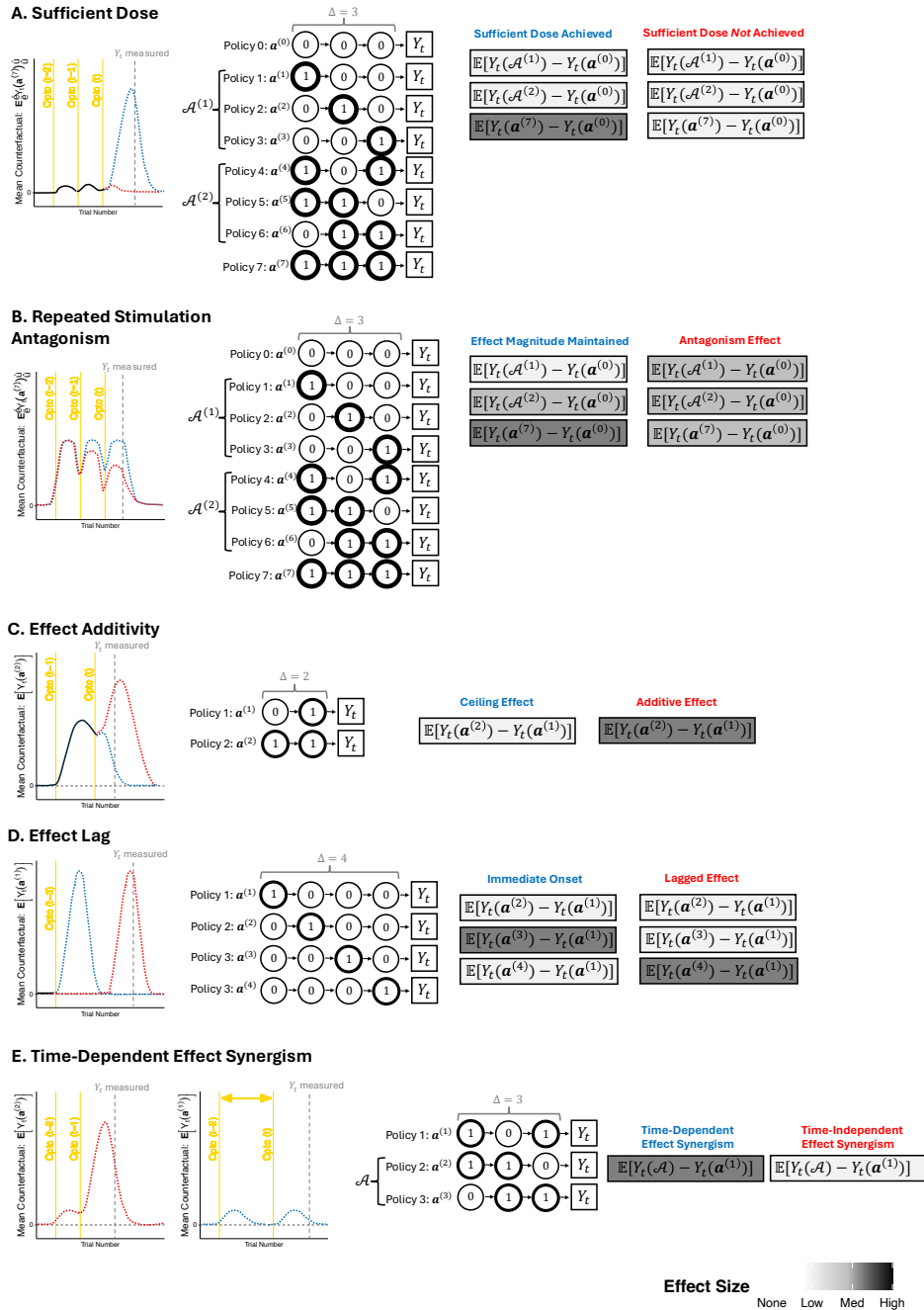


Figure 4: Example Sequential Excursion Effects. The left panels show one setting where a sequence of laser simulations **do** or **do not** have the indicated effect on the outcome. The middle panel shows deterministic static policies that could be used to construct a causal contrast to probe the effect. The right panel shows what the anticipated effect size (darker is larger) of those contrasts might be if there **is** or **is not** the indicated effect profile. [A] Sufficient dose. The red line shows how three successive stimulations is required to trigger a large effect, whereas the effect profile in blue shows that the sufficient dose has not been reached. [B] Repeated stimulation antagonism. The red line shows a negative dose-response, and the blue line shows a stable effect size. [C] Effect additivity. The red line shows a second stimulation triggers a larger response, whereas the blue shows that the second stimulation does not increase the response substantially beyond that of the first stimulation. [D] Effect Lag. The red line shows that the causal effect of stimulation is not visible until after a lag period. The blue line shows a setting where the effect is immediate. [E] Time-dependent effect synergism. The red line shows a setting where the effect is additive provided the stimulations occur close enough together (red line), but if stimulations occur far apart, this synergism does not occur (blue line).

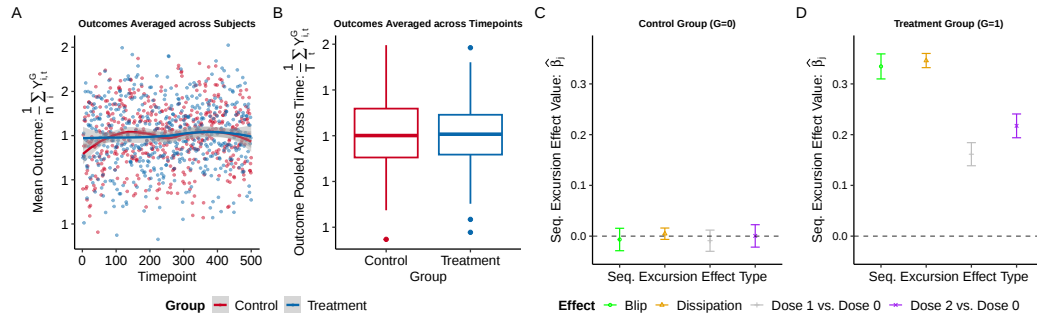


Figure 5: Treatment-Confounder Feedback Example Sequential excursion effects reveal causal effects obscured in “macro” summaries. Analysis results from a simulated dataset following the argument above (in Appendix B), taking $n = 100$, $T = 500$, $\gamma_1 = \gamma_3 = 1$, $\gamma_2 = \gamma_4 = 0.5$. [A] Each dot is an outcome value, $Y_{i,t}^G$, for subject i at timepoint t from “control” ($G = 0$), or from “treatment” ($G = 1$) groups. Lines are timepoint-specific means (averaged across subjects), estimated using a linear smoother (loess). (B) Same data as (A), but each point in boxplot is a subject’s mean outcome value (averaged across timepoints). In (A)-(B), “macro” summaries show no differences due to treatment-confounder feedback: mean outcome values (averaged across subjects or timepoints) are nearly identical in both groups. (C)-(D) Point estimates and 95% CIs (error bars) of sequential excursion effects reveal “local” causal effects (in Treatment group only), obscured in “macro” summaries (shown in (A)-(B)).

C Additional Details for Section 3

C.1 Interpretation of Sequential Excursion Effects

The interpretation of the mean counterfactual quantity $\mathbb{E}[Y_t(\mathbf{d}_{\Delta,t}) \mid V_{t-\Delta+1}]$ is somewhat subtle, and warrants further discussion. When $\Delta = 1$, we can express a contrast of these estimands in terms of the effect of exposure in a certain subgroup:

$$\mathbb{E}[Y_t(d_t^{(1)}) \mid V_t] - \mathbb{E}[Y_t(d_t^{(0)}) \mid V_t] = \mathbb{E}[Y_t(a_t = 1) - Y_t(a_t = 0) \mid V_t, I_t = 1] \mathbb{P}[I_t = 1 \mid V_t].$$

That is, the mean contrast in counterfactual outcomes for $d_t^{(1)}$ versus $d_t^{(0)}$ is the mean effect of $A_t = 1$ versus $A_t = 0$ among those with $I_t = 1$ —the availability-conditional estimand proposed by [5]—diluted by the probability of availability. Even in this case with $\Delta = 1$, it may not always be clear for whom the availability-conditional estimand generalizes to, i.e., the group $I_t = 1$ may be highly idiosyncratic and not of particular interest. On the other hand, the parameters we are proposing summarize the effects of plausible interventions on the whole population, acknowledging that for some individuals active treatment (i.e., $A_t = 1$) is not possible.

The comparison just described is somewhat akin to the duality in clinical trials of per-protocol (or complier-specific) effects, and intention-to-treat effects. Thus, in practice when $\Delta = 1$, we would recommend assessing both the availability-conditional estimand, as in [5], as well as our proposed population-level effect. When $\Delta > 1$, it is not clear whether an analogous availability-conditional estimand exists; our approach is viable for arbitrary Δ . In general, our estimands have the population-level (possibly conditional on effect modifiers) interpretation of summarizing how outcomes would be affected if the experimental protocol were changed to match $\mathbf{d}_{\Delta,t}$ for the Δ trials leading up to the outcome. Finally, we note that, as for all excursion effects or history-restricted marginal structural models, the estimands under study are dependent on the treatment protocol [6].

C.2 Discussion of Causal Assumptions

Consistency (Assumption 3.1) states that for any of the regimes $\mathbf{d}_{\Delta,t}$ under study, the counterfactual outcome $Y_t(\mathbf{d}_{\Delta,t})$ equals the observed outcome Y_t when observed treatment values correspond to assignment under $\mathbf{d}_{\Delta,t}$. Positivity (Assumption 3.2) states that treatment probabilities are bounded away from zero—this is required for the asymptotic analysis of the proposed estimator later on. Note that, by definition of the availability indicator I_t , and the regimes $\mathcal{D}_t^*$ in Section 3.1, we are allowing $\mathbb{P}[A_t = 1 \mid H_t] = 0$ in some cases (i.e., when $I_t = 0$), but Assumption 3.2 rules out $\mathbb{P}[A_t = 1 \mid H_t] = 1$. This positivity assumption holds in many open- and closed-loop optogenetic studies. In practice, in such experiments, one can ensure that Assumption 3.2 holds by design when choosing the treatment assignment probabilities. Finally, Assumption 3.3 says that treatments are randomly assigned at each time t , based on all previously measured data H_t . In the sequential optogenetic experiments that motivate this work, this assumption would hold by design. In observational studies, one will have to assess the plausibility of Assumption 3.3 (as well as Assumptions 3.1 and 3.2) on a case-by-case basis, ideally based on subject matter knowledge; it may be harder to justify Assumption 3.3 due to the possible presence of unmeasured confounders.

C.3 Conditions of Theorem 3.5

Conditions (i) through (iv) are standard conditions for asymptotic normality of M-estimators [8, 9]. For condition (ii), we expect the working model m to be differentiable in β for most common models. Moreover, for standard generalized linear models, m will be appropriately Donsker—see [35] for formal definitions. Condition (iii) is satisfied under mild conditions, e.g., if the weight functions h , the model m and its derivative M , and the outcomes Y_t are uniformly bounded, and no haphazard degeneracy in $\mathbf{B}$ exists that could cause singularity. Lastly, condition (iv) is also quite weak, only requiring convergence of $\hat{\beta}$ at an arbitrarily slow rate, and would hold under some stochastic equicontinuity conditions [22, 24].

C.4 Proofs of results in Section 3.2

Proof of Proposition 3.4. This result follows the usual g -formula identification argument [27]: defining $\mathbf{A}_{\Delta,t} = (A_{t-\Delta+1}, \dots, A_t)$, $\mathbf{d}_{\Delta,t}(H_t) = (d_{t+\Delta+1}(H_{t+\Delta+1}), \dots, d_t(H_t))$,

$$\begin{aligned} & \mathbb{E}(Y_t(\mathbf{d}_{\Delta,t}) \mid V_{t-\Delta+1}) \\ &= \mathbb{E}(\mathbb{E}(Y_t(\mathbf{d}_{\Delta,t}) \mid H_{t-\Delta+1}) \mid V_{t-\Delta+1}) \\ &= \mathbb{E}(\mathbb{E}(Y_t(\mathbf{d}_{\Delta,t}) \mid H_{t-\Delta+1}, A_{t-\Delta+1} = d_{t-\Delta+1}(H_{t-\Delta+1})) \mid V_{t-\Delta+1}) \\ & \dots \\ &= \mathbb{E}(\mathbb{E}(\dots \mathbb{E}(Y_t(\mathbf{d}_{\Delta,t}) \mid H_t, \mathbf{A}_{\Delta,t} = \mathbf{d}_{\Delta,t}(H_t)) \dots \mid H_{t-\Delta+1}, A_{t-\Delta+1} = d_{t-\Delta+1}(H_{t-\Delta+1})) \mid V_{t-\Delta+1}), \end{aligned}$$

where we repeatedly invoke iterated expectations and Assumption 3.3 (justified by Assumption 3.2), then use Assumption 3.1 in the last equality. We can then rewrite this formula in an equivalent IPW form:

$$\begin{aligned} & \mathbb{E}(\mathbb{E}(\dots \mathbb{E}(Y_t(\mathbf{d}_{\Delta,t}) \mid H_t, \mathbf{A}_{\Delta,t} = \mathbf{d}_{\Delta,t}(H_t)) \dots \mid H_{t-\Delta+1}, A_{t-\Delta+1} = d_{t-\Delta+1}(H_{t-\Delta+1})) \mid V_{t-\Delta+1}) \\ &= \mathbb{E} \left(\mathbb{E} \left(\frac{\mathbb{1}(A_{t-\Delta+1} = d_{t-\Delta+1}(H_{t-\Delta+1}))}{\pi_{t-\Delta+1}(A_{t-\Delta+1}; H_{t-\Delta+1})} \dots \mathbb{E} \left(\frac{\mathbb{1}(A_t = d_t(H_t))}{\pi_t(A_t; H_t)} Y_t(\mathbf{d}_{\Delta,t}) \mid H_t \right) \dots \mid H_{t-\Delta+1} \right) \mid V_{t-\Delta+1} \right) \\ &= \mathbb{E} \left(\prod_{j=t-\Delta+1}^t \frac{\mathbb{1}(A_j = d_j(H_j))}{\pi_j(A_j; H_j)} Y_t \mid V_{t-\Delta+1} \right), \end{aligned}$$

where the last equality is achieved again by iterated expectations. The second statement in Proposition 3.4 is obtained by differentiating (2) with respect to β , setting this to zero, then invoking the first statement of Proposition 3.4 (which we have just proved). $\square$

Proof of Theorem 3.5. This is an immediate application of Theorem 5.31 in [35]. $\square$

D Additional Simulation Details and Results

D.1 HR-MSM for Simulation Data-Generating Mechanism

In this section, we derive the form of the HR-MSM in (3), and show that it is implied by the data-generating mechanism of the simulation study. First, observe that for $t \geq 2$,

$$Y_t(d_{t-1}, d_t) = \alpha_1 X_{t-1} + \alpha_2 d_{t-1}(X_{t-1}) + \alpha_3 X_t(d_{t-1}) + \alpha_4 d_t(X_t(d_{t-1})) + \epsilon_t,$$

for some exogenous $\epsilon_t \sim \mathcal{N}(0, \sigma_t^2)$, where $X_t(d_{t-1})$ is the potential X_t value under the intervention setting A_{t-1} to $d_{t-1}(X_{t-1})$. Note that $d_{t-1}(X_{t-1}) = J_{t-1}X_{t-1}$, and by our structural equations, $X_t(d_{t-1}) \sim \text{Bernoulli}(0.4 + 0.4d_{t-1}(X_{t-1}))$, so that $\mathbb{E}(X_t(d_{t-1})) = 0.4 + 0.2J_{t-1}$, recalling that $\mathbb{E}(X_{t-1}) = 0.5$. Finally, $d_t(X_t(d_{t-1})) = J_t \cdot \text{Bernoulli}(0.4 + 0.4d_{t-1}(X_{t-1}))$, which gives $\mathbb{E}(d_t(X_t(d_{t-1}))) = \{0.4 + 0.2J_{t-1}\}J_t$. Putting everything together, we obtain

$$\begin{aligned} \mathbb{E}(Y_t(d_{t-1}, d_t)) &= 0.5\alpha_1 + 0.5J_{t-1}\alpha_2 + \{0.4 + 0.2J_{t-1}\}\{\alpha_3 + J_t\alpha_4\} \\ &= \{0.5\alpha_1 + 0.4\alpha_3\} + \{0.5\alpha_2 + 0.2\alpha_3\}J_{t-1} + 0.4\alpha_4 J_t + 0.2\alpha_4 J_{t-1}J_t \\ &\equiv \beta_0 + \beta_1 J_{t-1} + \beta_2 J_t + \beta_3 J_{t-1}J_t, \end{aligned}$$

as claimed.

D.2 Further Simulation Results

In Appendix Figure 6 we present the same results as in Figure 2 but with more sample sizes, n and trials, T . We also present these same simulation results in terms of the β coefficients of HR-MSM 3. In the main text we presented results in terms of sequential excursion effect parameters, which are linear combinations of these HR-MSM regression coefficients.

E Additional Application Results

E.1 Dose-Response Excursion Effects

History-Restricted MSM We fit an HR-MSM within the treatment group ($G = 1$) to estimate the causal effect of “dose,” the number of treatment opportunities in the previous $\Delta = 5$ trials:

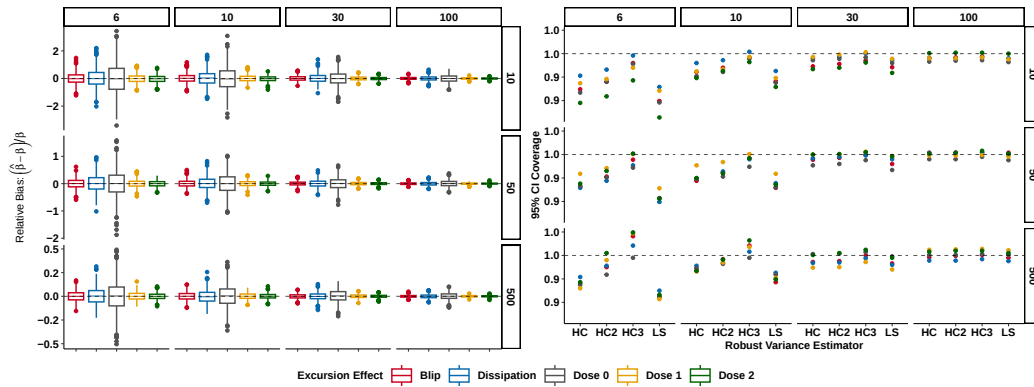


Figure 6: **Simulation Results** Panel columns indicate sample sizes, n , and panel rows indicate number of trials, T (cluster sizes). [Left] Relative bias associated with each sequential excursion effect. These results show that our estimator is consistent for the target parameters. [Right] 95% Confidence interval (CI) coverage for the sequential excursion effects. The coverage of 95% CIs constructed using one of three established robust variance estimators and our robust large sample (shown as *LS*) variance estimator. The nominal coverage is reached for either large n or large t for all estimators.

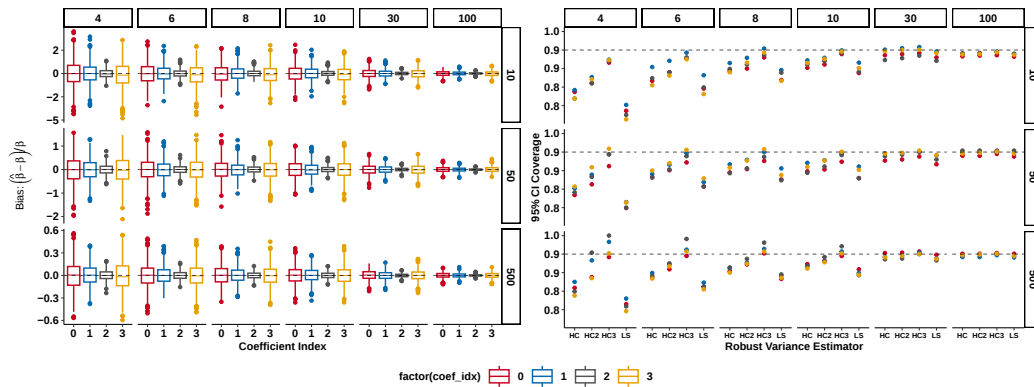


Figure 7: **Simulation Results in Terms of MSM Coefficients** Relative bias and 95% Confidence Interval (CI) coverage of regression coefficients of history-restricted marginal structural model (MSM) 3. Panel columns indicate sample sizes, n , and panel rows indicate number of trials, T (cluster sizes). [Left] Relative bias associated with each HR-MSM regression coefficient. These results show that our estimator is consistent for the target parameters. [Right] 95% CI coverage for the MSM coefficients. The coverage of 95% CIs constructed using one of three established robust variance estimators and our robust large sample (shown as *LS*) variance estimator. The nominal coverage is reached for either large n or large t for all estimators.

$$\text{logit}(\mathbb{E}[Y_t(\mathbf{d}_{\Delta,t}) \mid G = 1]) = \beta_0 + \sum_{r=1}^3 \beta_r \mathbb{1}\left(\sum_{j=t-\Delta+1}^t \sigma_j(d_j) = r\right), \quad (5)$$

where $\sigma_j(d_j) = \mathbb{1}(d_j = d_j^{(1)})$. The coefficient $\hat{\beta}_r$ is an estimate of the log odds ratio comparing the mean counterfactual of Y_t for a treatment sequence of dose $r \in [3]$ compared to a sequence of dose zero (see Figure 1C for an illustration of the static regime analogue). A dose of three is the maximum feasible dose for $\Delta = 5$ since the same pose cannot occur on two consecutive trials.

The dose-response effect estimates, $\{\hat{\beta}_r\}_{r=1}^3$, from HR-MSM (5) are shown in Figure 8A. This illustrates the capacity of our approach to identify a clear dose-response effect: within the past $\Delta = 5$ trials, each additional opportunity for a stimulation *causes* an increase in the odds of engaging in the

Effect	$T = 50$				$T = 500$			
	$n = 6$	$n = 10$	$n = 30$	$n = 100$	$n = 6$	$n = 10$	$n = 30$	$n = 100$
Blip	1.91 ± 0.09	1.10 ± 0.05	0.30 ± 0.02	0.03 ± 0.01	0.11 ± 0.01	0.03 ± 0.01	0.00 ± 0.00	0.00 ± 0.00
Dissip	2.11 ± 0.09	1.23 ± 0.06	0.33 ± 0.02	0.04 ± 0.01	0.11 ± 0.01	0.03 ± 0.01	0.00 ± 0.00	0.00 ± 0.00
Dose 0	0.73 ± 0.04	0.40 ± 0.02	0.08 ± 0.01	0.00 ± 0.00	0.02 ± 0.00	0.00 ± 0.00	0.00 ± 0.00	0.00 ± 0.00
Dose 1	0.90 ± 0.04	0.51 ± 0.03	0.10 ± 0.01	0.00 ± 0.00	0.02 ± 0.00	0.00 ± 0.00	0.00 ± 0.00	0.00 ± 0.00
Dose 2	3.52 ± 0.15	2.05 ± 0.09	0.58 ± 0.03	0.12 ± 0.01	0.25 ± 0.02	0.11 ± 0.01	0.01 ± 0.00	0.00 ± 0.00

Table 2: **Simulation Results: MSE** Our estimator’s MSE decreases to 0 as T or n grows. Denoting the estimated effect j (e.g., $j = \text{“Blip”}$) for replicate r as $\hat{\beta}_{j,r}$, we show $\text{MSE}_j := \frac{1}{R} \sum_{r=1}^R (\hat{\beta}_{j,r} - \beta_j)^2$ for $R = 1000$ simulation replicates ($\pm$ SE) for a sample size, n , and timepoints, T . Values are scaled by 100 for readability (e.g., 0.01 is shown in the table as 1.0). Thus 0 indicates a value $< 1e - 4$.

Effect	$T = 50$				$T = 500$			
	$n = 6$	$n = 10$	$n = 30$	$n = 100$	$n = 6$	$n = 10$	$n = 30$	$n = 100$
Blip	0.23 ± 0.55	0.42 ± 0.43	0.42 ± 0.24	0.10 ± 0.13	0.07 ± 0.17	0.05 ± 0.13	0.08 ± 0.08	0.03 ± 0.04
Dissip	0.82 ± 0.93	1.04 ± 0.71	0.45 ± 0.39	0.26 ± 0.22	0.17 ± 0.27	0.22 ± 0.21	0.10 ± 0.13	0.06 ± 0.07
Dose 0	0.28 ± 1.48	0.15 ± 1.17	0.12 ± 0.68	0.42 ± 0.36	0.28 ± 0.48	0.23 ± 0.35	0.34 ± 0.20	0.13 ± 0.11
Dose 1	0.12 ± 0.42	0.05 ± 0.32	0.28 ± 0.18	0.06 ± 0.10	0.10 ± 0.13	0.08 ± 0.10	0.04 ± 0.06	0.01 ± 0.03
Dose 2	0.41 ± 0.35	0.09 ± 0.27	0.12 ± 0.15	0.10 ± 0.08	0.03 ± 0.11	0.05 ± 0.08	0.03 ± 0.05	0.01 ± 0.03

Table 3: **Simulation Results: Bias** Our estimator is unbiased. Moreover, the absolute relative bias decreases to 0 as T and/or n grows. Denoting the estimated effect j (e.g., $j = \text{“Blip”}$) for replicate r as $\hat{\beta}_{j,r}$, we show $\text{Absolute Relative Bias}_j := \frac{1}{R} \sum_{r=1}^R (\hat{\beta}_{j,r} - \beta_j) / \beta_j$ for $R = 1000$ replicates ($\pm$ standard error). Values are scaled by 100 for readability (e.g., 0.01 is shown in the table as 1.0).

target pose on the next trial. The effects are significant for at least one dose value in all but two target poses. Interestingly, the effect is also significantly negative for one target pose.

Conditional Excursion Effect We next estimate an availability-conditional estimand [5], to determine if existing excursion effect methods have the capacity to reveal the effects identified with our method. We estimate this in the MSM

$$\text{logit}(\mathbb{E}[Y_t(a_t) \mid I_t = 1, G = 1]) = \alpha_0 + \alpha_1 a_t. \quad (6)$$

Figure 8B shows the availability-conditional treatment effect estimates, $\hat{\alpha}_1$ estimated in model (6). It identifies no significant effects for any target pose. The conditional estimand, often referred to as a “blip effect” (see Figure 1A for an illustration) is only defined for the effect of applying the laser on the most recent trial (i.e., a dose of 1), and thus cannot estimate dose-response profiles. In contrast, our approach can test *sequential* excursion effects (i.e., for policies with $\Delta > 1$), enabling the estimation of a dose-response profile that reveals treatment effects here. Importantly, the effect estimates $\hat{\beta}_1$ and $\hat{\alpha}_1$ have different interpretations because $\hat{\beta}_1$ reflects a causal effect of a single treatment opportunity for any of the last $\Delta = 5$ trials, and $\hat{\beta}_1$ is not interpreted as conditional on availability.

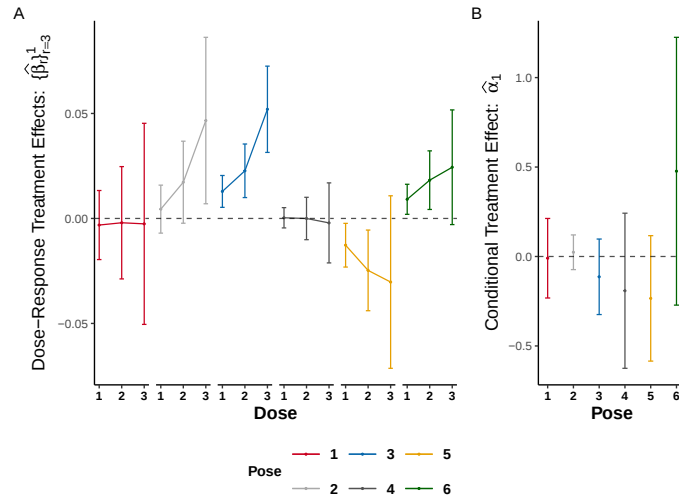


Figure 8: **Our method enables estimation of dose effects.** Plots show coefficient estimates (error bars show 95% CIs) as a function of dose. Columns and colors indicate the dose. [A] Main effects of stimulation opportunity from HR-MSM (5). [B] Availability-conditional effects of treatment estimated in MSM (6).

E.1.1 Effect Modification by Baseline Behavior

Marginal Effect Modification Parameter Next we show the capacity of our method to estimate effect modification in the form of interactions between covariates and functions of deterministic policies (i.e., treatment opportunity dose). We estimated effect modification of total target pose counts on baseline sessions, $\bar{Y}^0$ (defined in Section 4.2), because [17] estimated the treatment effect of the laser by comparing the mean change in outcome levels between treatment and baseline sessions.

Augmenting the HR-MSM in the previous section with effect modifier, $\bar{Y}^0$, we fit the HR-MSM

$$\text{logit}(\mathbb{E}[Y_t(d_{\Delta,t}) | \bar{Y}^0, G = 1,]) = \beta_0 + \sum_{r=1}^3 \beta_r \mathbb{1}(\sum_{j=t-\Delta+1}^t \sigma_j(d_j) = r) + \beta_4 \bar{Y}^0 + \sum_{r=1}^3 \beta_{4+r} \bar{Y}^0 \times \mathbb{1}(\sum_{j=t-\Delta+1}^t \sigma_j(d_j) = r). \quad (7)$$

Similarly, we fit an analogous model for the conditional estimand,

$$\text{logit}(\mathbb{E}[Y_t(a_t) | I_t = 1, G = 1, \bar{Y}^0]) = \alpha_0 + \alpha_1 a_t + \alpha_2 \bar{Y}^0 + \alpha_3 \bar{Y}^0 \times a_t. \quad (8)$$

We centered and scaled the effect modifier to make the coefficients easier to interpret.

Figure 9A-B shows how our method can be used to probe effect-modification. The figures illustrate that our approach estimates that for two target poses, there is a statistically significant effect modification by baseline responding at, at least, one of the doses. Specifically, animals who exhibited higher levels of responding on *baseline* sessions, exhibited a larger effect of dose. Figure 9C-D shows that the availability-conditional estimator does not identify any significant main effect of stimulation or interactions with baseline responding. These results show how our framework enables one to probe effect modification of sequential excursion effects.

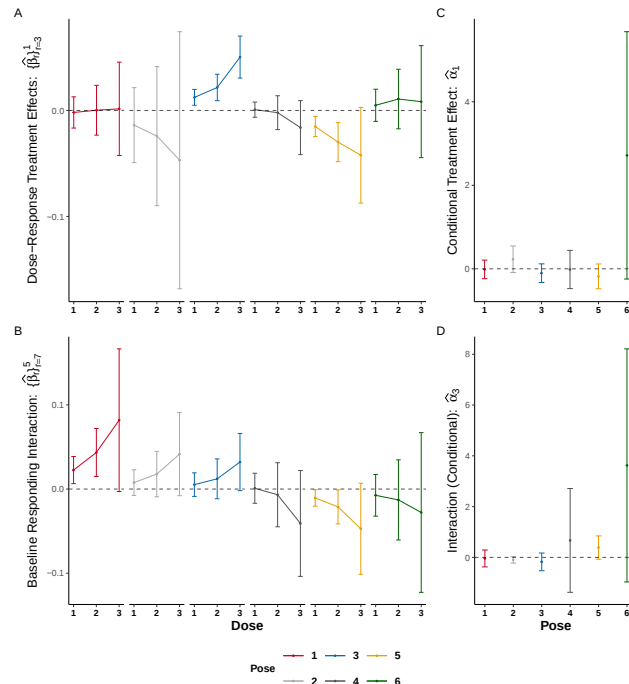


Figure 9: **Our method enables estimation of effect modification of baseline behavioral responding.** Plots show coefficient estimates (error bars show 95% CIs) as a function of dose. Columns and colors indicate pose. [A] Main effects of treatment opportunity for doses 1-3 at mean baseline responding levels with marginal HR-MSM (our approach), $\Delta = 5$ $k = 1$. [B] Interaction with baseline (pre-stimulation session) responding levels with marginal HR-MSM (our approach). [C] Main effects of stimulation on past trial at mean baseline responding levels with availability-conditional approach. [D] Interaction between stimulation and baseline responding with availability-conditional approach.

F Data, Pre-processing and Code Availability

We provide code to reproduce all analyses and figures on anonymous GitHub Repo: https://anonymous.4open.science/r/causal_opto-52CD/README.md. We downloaded the open-source dataset from [17] from <https://zenodo.org/records/7274803>. We used the open-source pre-processing code provided by the authors on Github repo <https://github.com/dattalab/dopamine-reinforces-spontaneous-behavior>. We analyzed data from all animals in the on-line dataset (both ChR2 and Chrimson animals). We constructed trials as described in [17]. That is, we defined trials as consecutive timepoints when the animal was classified to be in a given pose. For our HR-MSM analyses of the treatment (opto) sessions, we classified “target pose” trials only if they met the criteria of [17], which required that the hidden Markov model predictions had sufficiently high forward algorithm probabilities of the latent states. This indicator was provided in the opto session dataset provided by the authors. The baseline session data did not, however, include this indicator since no optogenetic stimulation was applied. Thus when recreating the “standard” between-group (macro longitudinal) analyses that compared baseline and opto session data, we did not classify target pose trials based on whether it met this criteria: we classified the pose based on the most likely latent state prediction but did not require the forward algorithm probabilities met the threshold set by the authors (for either baseline or opto sessions to be consistent). We corresponded closely with the authors to ensure we pre-processed the data correctly.

There was a small percentage of trials that the authors described eliminating because they were deemed too short. We did not eliminate these trials because this created inconsistencies in the pattern of trials: it allowed two consecutive trials to be of the same trial type which broke with the pattern in the remainder of the dataset. This was a very small percentage of the dataset. We compared results with and without this criteria and the decision appeared to have negligible effects on analysis results.

Finally, to the best of our understanding, the original authors' hypothesis tests were conducted on further processed version of the data that first calculated the number of target pose occurrences in each 30 second bin of the experiment (period). From our understanding, this was done to provide a smoothed time-series of outcome frequency across the course of the opto sessions. We conducted similar analyses to make sure our pre-processing yielded comparable results, but we did not use these pre-processing steps in our HR-MSM analyses or replication of the “standard” between-group (macro longitudinal) analyses as it appeared to “double-count” target pose occurrences that began before the end of one 30 second bin and ended after the start of the subsequent bin.

Finally, as described in [17], the experiment included two 30 minute replicates of both opto and baseline sessions. We constructed trials on each replicate separately (to account for the discontinuity in time between replicates) and then pooled the replicate datasets together to be consistent with the analysis procedures in [17]. We accounted for the longitudinal structure by using sandwich variance estimators in all of our hypothesis tests.

F.1 Replication of Original Author Analysis

Finally, Figure 3D shows that the “standard” macro longitudinal effects exhibit no significant changes, emphasizing how estimands that marginalize over the stochastic dynamic (closed-loop) policies can obscure effects. We visualize the within-subject differences in target pose counts between treatment (opto) and baseline sessions in Figure 10: $\sum_{t=1}^T Y_t - \sum_{t=1}^{T_0} Y_t^0$, where $Y_t, Y_t^0 \in \{0, 1\}$ are indicators that the animal exhibited the target pose on trial t of the treatment and baseline sessions, respectively; $T, T_0 \in \mathbb{Z}$ denote the total number of trials in treatment and baseline sessions, respectively. Because of the trial definition, the total number of trials usually differed within-subject between treatment and baseline sessions (i.e., $T \neq T_0$), but we found comparable analyses of $\frac{1}{T} \sum_{t=1}^T Y_t$ and $\frac{1}{T_0} \sum_{t=1}^{T_0} Y_t^0$ yielded similar results to analyses of total counts $\sum_{t=1}^T Y_t - \sum_{t=1}^{T_0} Y_t^0$. For that reason, we present results in terms of total outcome counts to be consistent with the analyses presented in [17]. We showed these results to the authors of [17] and they confirmed that these analyses aligned with theirs.

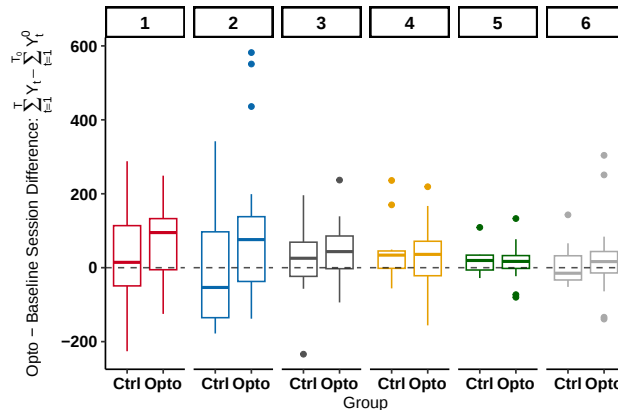


Figure 10: **Difference between target pose counts within-subject between baseline and treatment (opto) sessions.** Each point in the boxplot shows $\sum_{t=1}^T Y_t - \sum_{t=1}^{T_0} Y_t^0$, where $Y_t, Y_t^0 \in \{0, 1\}$ are indicators that the animal exhibited the target pose on trial t of the treatment and baseline sessions, respectively. $T, T_0 \in \mathbb{Z}$ were the total number of trials in treatment and baseline sessions, respectively. Columns and colors indicate target pose. *Ctrl* and *Opto* indicates control and treatment group subjects, respectively.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The methodological, theoretical, and experimental contributions described in the abstract are each addressed in their own main text sections. We only make claims in the abstract that we carefully justify and explain in the paper.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: The Discussion section includes limitations of our proposed method and the interpretation of the method's results. We state the mathematical assumptions clearly in the main text and discuss their justification in the Appendix. We also discuss method scalability in the Methods (Section 3). Our simulation results highlight the limitations of our approach in small sample settings: small sample bias, and the requirement of sample size-adjusted covariance estimators to achieve nominal coverage. We intentionally include very small sample sizes in our simulations to emphasize the trade-offs in these settings even though our data application has datasets with a much larger sample size. Our work focuses on animal data and thus does not include problems of privacy or fairness.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.

- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: Complete proofs of Proposition 3.4 and Theorem 3.5 are provided in Appendix C.4. All remaining technical details are discussed in the paper, with additional details provided in the appendices when relevant.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We provide all details needed to reproduce our results in the main text and appendix. We include all code used to produce results, pre-process data and make figures. We provide links to the open-source datasets that we analyze.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.

- (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
- (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [\[Yes\]](#)

Justification: We include a data pre-processing and data availability section (data from our application is open-source) in the Appendix section. We also include a link to an anonymous GitHub repository with an implementation of all results presented in the main text and Appendix.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [\[Yes\]](#)

Justification: We discuss the optimizer used. Our method has no tuning parameters. No data splits were used anywhere in the paper since our method is focused on statistical inference and hypothesis testing (not on prediction).

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: We include 95% confidence intervals (CIs) for our results, e.g., represented by error bars in figures. We verify in simulations the validity of these CIs (i.e., that they achieve the nominal coverage).

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [No]

Justification: For the vast majority of applications, our methods are not computationally intensive and can be run quickly with little memory. We briefly discuss but we do not provide details since we ran everything on a standard laptop with no GPUs or parallelization.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: The authors have reviewed the NeurIPS Code of Ethics, and have ensured that the paper conforms in all respects with these guidelines.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: We discuss that our method has the capacity to reduce the number of animal subjects used in experiments.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The paper is focused on methodology for animal studies, and we do not foresee our approaches being misused or used irresponsibly. The data used are publicly available, coming from a high profile *Nature* paper.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [\[Yes\]](#)

Justification: We properly cite existing packages and datasets.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New Assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[No\]](#)

Justification: We are not releasing any packages, benchmarks, or other assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and Research with Human Subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[NA\]](#)

Justification: Data used in the paper did not involve crowdsourcing nor human subjects. Data were publicly available—see Appendix F for details.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: Data used in the paper did not involve crowdsourcing nor human subjects. Data were publicly available—see Appendix F for details.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.