

Optimal Experiment Design in Markov Chains

ETH Zürich, Broad Institute
Mojmír Mutný

Andreas Krause



Tadeusz Janik



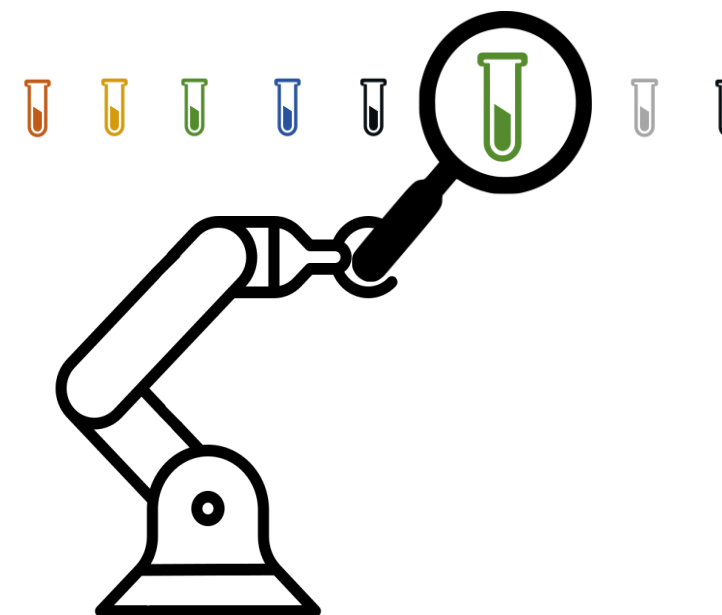
Manish Prajapat



Jose Folch

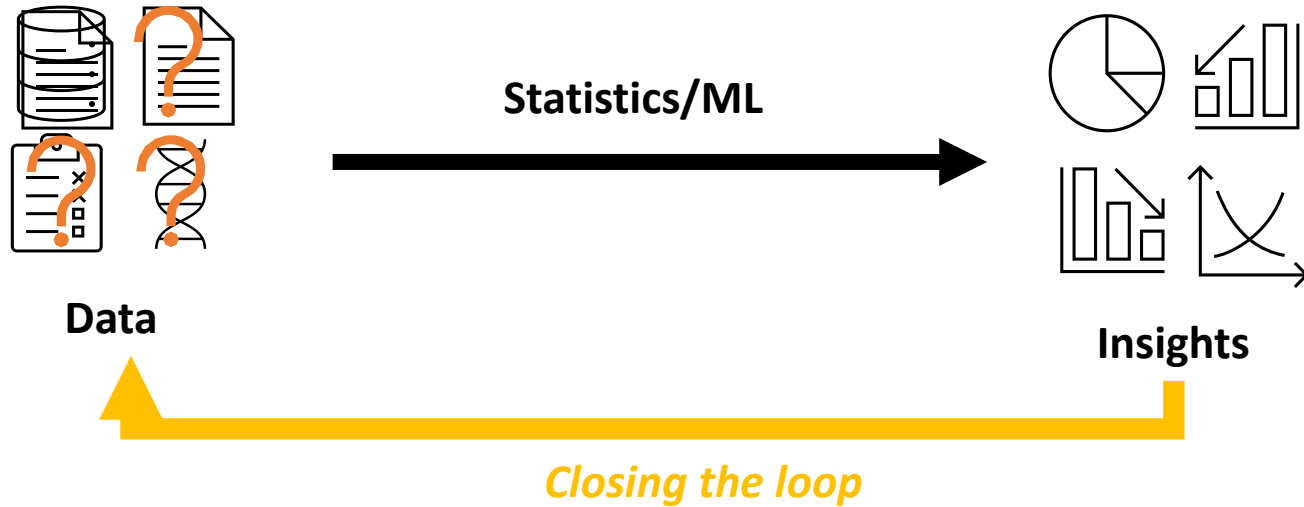


Ruth Misener



15 June 2026

Statistical Experiment Design?



What if?

- What if we can decide what data we acquire? → **Experiment Design**
- What if the downstream decisions influence what data we gather? → **Adaptive Experiment Design**

What data? Depends...

- How do we model the problem?
- What do we want to do with it?
- What data we can get?
- What information we already have?



Mathematical Model



Task



Search Space



What is uncertain?

In this talk: Experiments are temporally structured decision rules – policies.

Experiment Design and Estimators

General Setup

- **Unknown:** Unknown parameter θ_* (RKHS element)
- **Measurement:** x is a query to $y \sim p(\theta_*^\top \Phi(x))$
- **Evaluation functional** $f(x) = \Phi(x)^\top \theta_*$
- **Estimator:** Estimator $\hat{\theta}(\mathbf{X}, \mathbf{y})$
 - $\mathbf{X}$ stacked measurements, $|\mathbf{X}| < T$
 - and their outcomes $\mathbf{y}$
- **Utility:** $U(\hat{\theta}, \theta_*) : \mathcal{H}_k \times \mathcal{H}_k \rightarrow \mathbb{S}^{p \times p}$
Random quantity due to $\mathbf{y}$
- **Expected utility:**
 $\mathbf{E}_{t, \theta_*}(\mathbf{X}) = \mathbb{E}_{p_{\theta_*}(y_1|\mathbf{X}_1) \dots p_{\theta_*}(y_t|\mathbf{X}_{1:t}, y_{1:t-1})} [U(\hat{\theta}, \theta_*) | \mathcal{F}_{t-1}]$
- **Scalarization** $s(\mathbf{E}_{\theta_*})$, where $s : \mathcal{H}_k \rightarrow \mathbb{R}$

Experiment Design problem

$$\mathbf{X}^* = \arg \max_{\mathbf{X}, |\mathbf{X}| \leq T} s(\mathbf{E}_{\theta_*}(\mathbf{X})).$$

An Example (Regression with Poisson responses)

- **Unknown:** Unknown parameter θ_*
- **Measurement:** x is a query to $y = \theta_*^\top \Phi(x) + \epsilon_H$
- **Estimator:** heteroscedastic least-squares

$$\hat{\theta} \in \arg \min_{\theta \in \mathcal{H}_k} \text{LOSS}(\mathbf{X}, y)$$

- **Utility:** covariance of residuals

$$U(\hat{\theta}, \theta_*) = (\hat{\theta} - \theta_*)(\hat{\theta} - \theta_*)^\top$$

- **Expected utility:**
 $\mathbf{E}_{\theta_*} = \sum_{x \in \mathbf{X}} \frac{\Phi(x) \Phi(x)^\top}{\theta_*^\top \Phi(x)}$

- **Scalarization:** A-design
 $s(\mathbf{A}) = \frac{1}{\text{Tr}(\mathbf{A}^{-1})}$

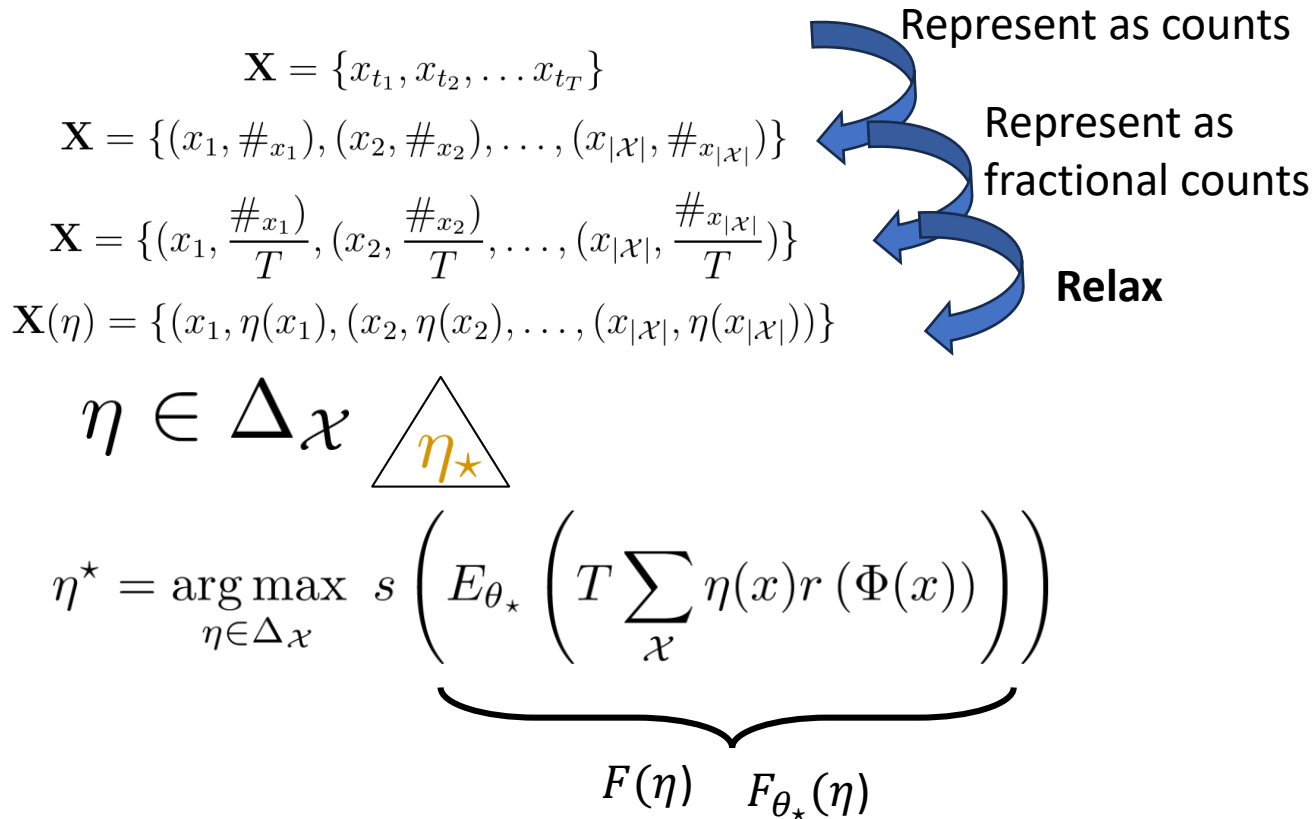
1. NP-hard Combinatorial problem

2. Depends on the unknown θ_*

Experiment Design Objectives: An Assortment

1. NP-hard Combinatorial problem

We resort to relaxation techniques!



Continuous domain, solve and round to fractional

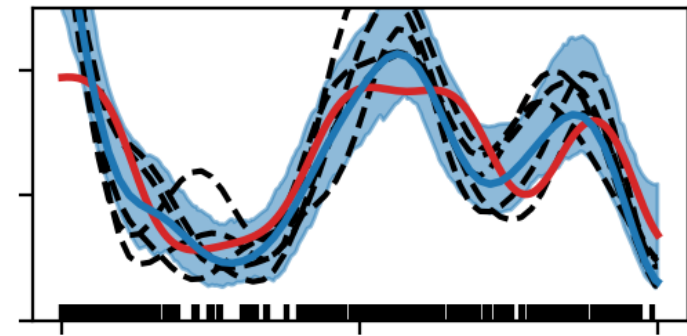
2. Depends on the unknown $\theta_{\star}$

- Build uncertainty set
- Be robust or average

$$\mathbf{X}^{\star} = \arg \max_{\mathbf{X}, |\mathbf{X}| \leq T} \min_{\theta \in \Theta} s(\mathbf{E}_{\theta_{\star}}(\mathbf{X}))$$

$$\mathbf{X}^{\star} = \arg \max_{\mathbf{X}, |\mathbf{X}| \leq T} \int_{\theta \in \Theta} s(\mathbf{E}_{\theta_{\star}}(\mathbf{X}))$$

- Θ contains all plausible models



Markov chains and Experiment Design

Markov chain process

- **Introducing:** state x and action a
- **Known** transition kernel $P(x'|x, a)$
- Observe Unknown function

$$y = \Phi(x, a)^\top \theta_\star + \epsilon$$

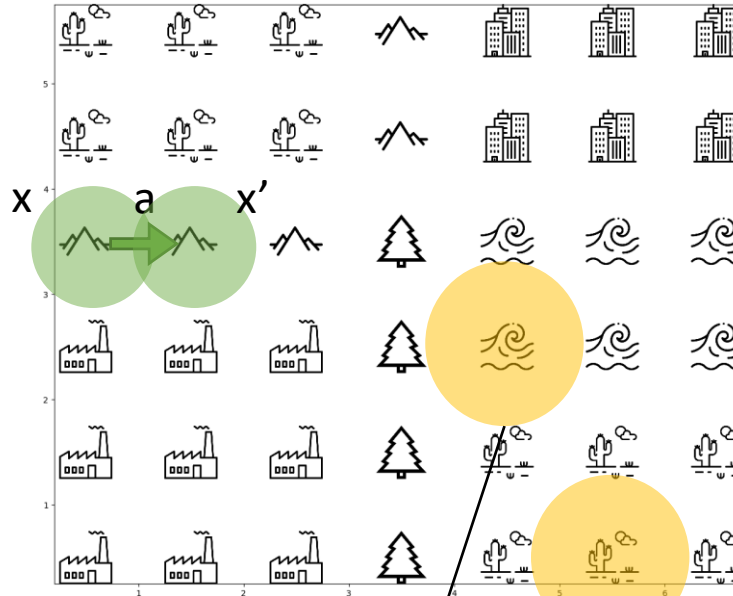
Interaction

- *Can only execute policy to interact with environment that respects the transitions*
- Classically we could pick any (x, a) !

Naïve solution:

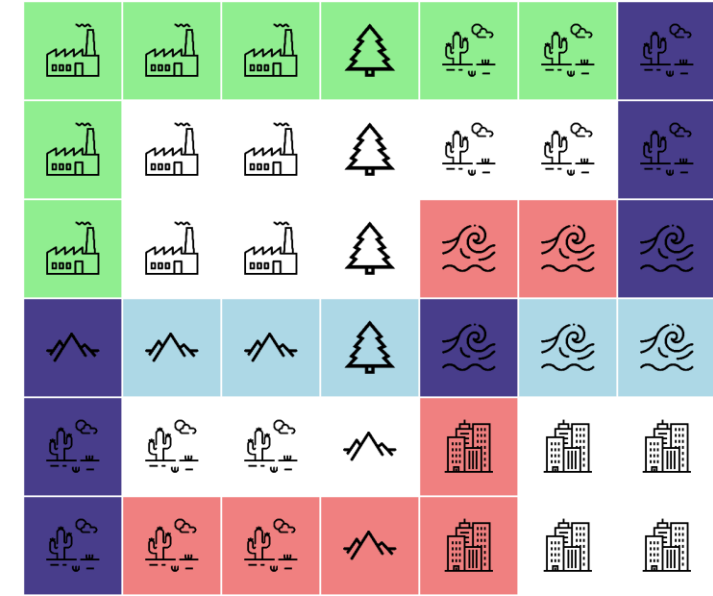
- Optimize over trajectories $\eta \in \Delta_{|\mathcal{P}|}$
 - Exponential blowup in variables ❌
 - Stochastic dynamics does not allow for execution of the trajectory ❌

Goal: Design a tractable sequence of policies such that the visited states and executed actions converge to the **optimal visitation** of the states and actions.



$$\begin{aligned}\Phi(x_1) &= (7, 20, \text{"water"}) \\ \Phi(x_1) &= (\text{pH}, \text{salinity}, \text{area type})\end{aligned}$$

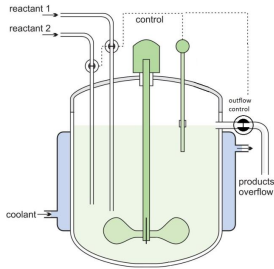
$$\Phi(x_2) = (8, 14, \text{"arid"})$$



Generating trajectories

Modern Experiment Design: Applications

Example 1: Reactor Optimization



model

Parametric Differential
Equations

task

$$\arg \max_x f(x)$$

randomness

measurement errors

search space

Space of non-Markov policies,
parametric controllers

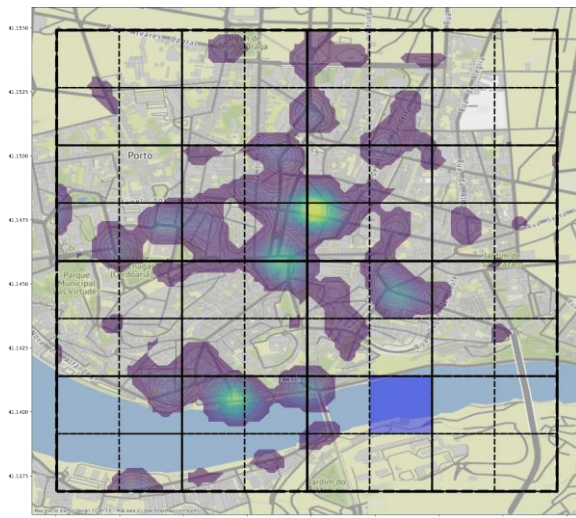
**Operating constraint, where next
measurement depends on prior state**

- Equilibrium constraints
- Smooth ramps
- Safety constraints

Mutny, Janik, Krause (AISTATS 2023)

Folch, ..., Mutny (NeurIPS 2024)

Example 2: Spatial Sensing



task

Levels sets identification

randomness

$$n(A) \sim \text{Poisson} \left(\int_A \lambda(x) dx \right)$$

model

Non-parametric
Poisson point process

action space

Hierarchical Borel sets

$$\int_A \lambda(x) dx$$

Sensing done with:

- Vehicles – need to respect roads
- Drones – spatial constraints
- Satellites – need to work with positioning

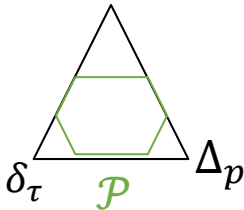
Mutny & Krause (ICML 2021)

Mutny & Krause (AISTATS 2022)

MDP structure of the experiments

How to deal with combinatorial space of all trajectories $\mathcal{P}$?

- η is over probability simplex of trajectories (exponentially many vertices -- very large)
- **Observation 1:** Our objectives depends only on pairs (x, a)



$$\eta^* = \arg \max_{\eta \in \Delta_{\mathcal{P}}} s \left(E_{\theta_*} \right)$$

Space of trajectory probabilities

As opposed to $\eta \in \Delta_{|\mathcal{P}|}$

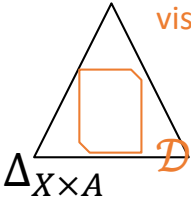
$$d^* = \arg \max_{d \in \Delta_{\mathcal{X} \times \mathcal{A}}} s \left(E_{\theta_*} \left(T \sum_{\mathcal{X} \times \mathcal{A}} d(x, a) r \left(\Phi(x, a) \right) \right) \right)$$

Proportion of states visiting x and executing a

- **Observation 2:** Not all states and actions proportions are possible:

Space of state-action visitation distributions

State-action visitation polytope [Putterman (1994)]



$$\mathcal{D}_h := \{d_h \mid d(x, a) \geq 0, \sum_{a, x} d_h(a, x) = 1, \sum_a d_h(x', a) = \sum_{x, a} d_{h-1}(x, a) p(x' | x, a)\}.$$

normalization

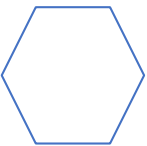
flow constraint

convex polytope

very well studied in RL

efficient solvers for linear problems

$\pi(a|x)$
Space of policies



Π

Forward propagation

$$d_h(x, a) = \left(\prod_{i=1}^h P_{\pi_h} d_0(x) \right) \pi_h(a|x).$$

Marginalization:

$$\pi_h(a|x) = \frac{d_h(a, x)}{\sum_a d_h(a, x)}$$

Lemma: Any d , can be realized by a Markov policy.

MDPs and Relaxations

State-action visitation distribution (Putterman (1994))

$$d^* = \arg \max_{d \in \mathcal{D} = \cup_{h=1}^H \mathcal{D}_h} s \left(E_{\theta_*} \left(T \sum_{x,a \in \mathcal{X} \times \mathcal{A}} d(x,a) r(\Phi(x,a)) \right) \right)$$

Optimal visitation of states

Same type of relaxation := F(d)

Good news: Its convex for many examples

Challenge: Constraint description might be complicated

Solution: Convex RL (Hazan et. al, 2019; Zahavy et al., 2021)

- **Franke Wolfe:** Decompose into a **sequence** of classical **RL problems**
- Solving each RL problem is *easy*: Policy iteration, policy gradient etc..

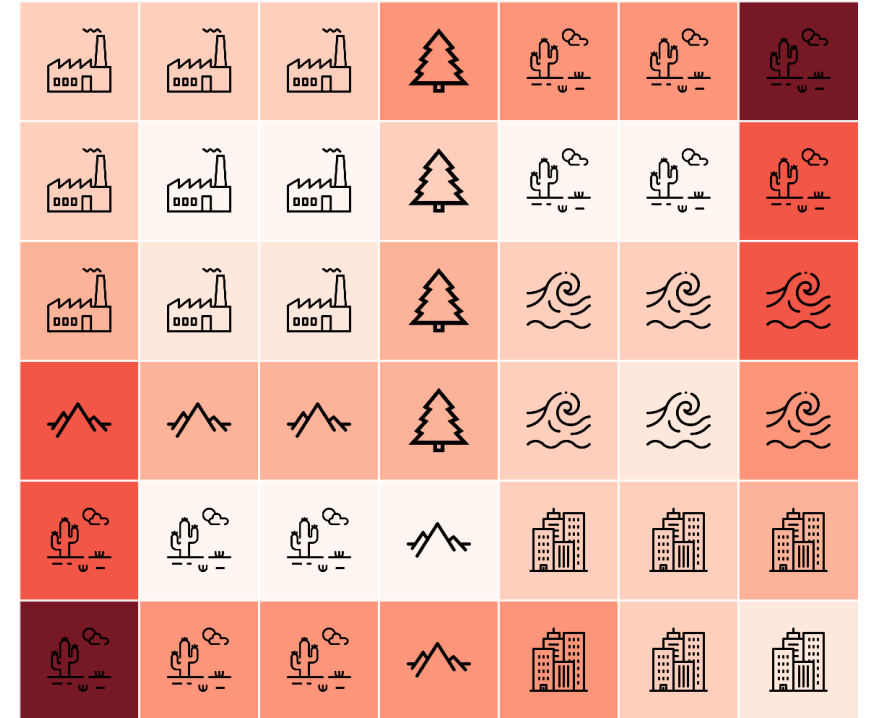
Mixture policy: $\{\alpha_i, \pi_i\}_{i=1}^n \quad \sum_{i=1}^n \alpha_i = 1, \alpha \geq 0$

$$d_{\pi_{\text{mix}}, n} = \sum_{i=1}^n \alpha_i d_{\pi_i}$$

Linearization at current mixture

Linear Maximization Oracle

$$d_{\pi_{n+1}} = \arg \max_{d \in \mathcal{D}} \sum_{x,a} \nabla F(d_{\pi_{\text{mix}}})(x,a) d(x,a) = \arg \max_{d_{\pi} \in \mathcal{D}} \sum_{x,a} r(x,a) d(x,a)$$



D-design

Reward from RL

Adaptivity into the picture

So far: No interaction with the system, **Now:** Interaction and unknowns!

In each step we optimize only correction:

- In each round we calculate and execute a **new Markov policy** π_t

$$G_t(d) = \min_{\theta \in \Theta_t} F_\theta \left(\frac{1}{t+1}d + \frac{t}{t+1}\hat{d}_t \right) \quad d_{\pi_t} = \arg \max_{d \in \mathcal{D}} G_t(d)$$

Robust design

Empirical visitation of past interactions (counts)

Historical sidenote: This idea appears in Pukelsheim's textbook on Bayesian designs

Note:

- Each policy is optimized to be **correction** to the empirical visitation
- The set of policies $\{\pi_1, \pi_2, \dots, \pi_T\}$ is **history dependent**
- The history dependence is via **the objective**, namely empirical state visitation
- The **uncertainty** is time-dependent
- The overall policy over episodes/rounds is **non-Markovian!** $\rightarrow$ *but not optimal non-Markovian*

Theory and Convergence

- **Theorem:** “Best possible explanation of a teacher to a student”
- **Theory (Assumption):**
 - Concave $F_\theta(\eta)$ in η
 - Smooth $F_\theta(\eta)$ in θ
 - “Compatibility of the objective with uncertainty reduction”

best empirical

$$F_{\theta_\star}(d_\star) - F_{\theta_\star}(d_\tau)$$

Comparing to the unknown $\theta_\star$

The diagram shows two orange circles. The left circle is labeled 'best' and contains the expression $F_{\theta_\star}(d_\star)$. The right circle is labeled 'empirical' and contains the expression $F_{\theta_\star}(d_\tau)$. An orange arrow points from the left circle to the right circle, with the text 'Comparing to the unknown $\theta_\star$ ' written below it.

Then as the number of experiments τ ,

$$F_{\theta_\star}(d_\star) - F_{\theta_\star}(d_\tau) \leq \frac{1}{\tau} \left(\underbrace{F_{\theta_\star}(\eta^\star) - F_{\theta_\star}(\eta_0)}_{\text{Initial error (regularization)}} + \underbrace{\sum_{t=1}^{\tau} \epsilon_t^S}_{\text{Rounding Error } O(\sqrt{\tau})} + \underbrace{\sum_{t=1}^{\tau} \epsilon_t^A}_{\text{Approximation error}} + \cancel{\sum_{t=1}^{\tau} \frac{L}{1+\tau}} \right)$$

$\epsilon_A = |l_{v_t}^\top (\theta_\star - \bar{\theta}_t)|$
Confidence set width at time t

Conclusion: As $\Theta_t \rightarrow 0$, this error decreases at $O(1/\sqrt{\tau})$ and **optimal constant** in many cases.

- Examples:**
- **Rounding error with** sampling from optimized measure $\rightarrow$ martingale $O(\sqrt{\tau})$
 - **Approximation error with** D-design and Gaussian likelihood is zero

Experiment Design with MDPs: Difficulty & Relaxation



The whole process summarized

- In each round we calculate and execute a **Markov policy** π_t
- The set of policies $\{\pi_1, \pi_2, \dots, \pi_T\}$ is **history dependent**
- The history dependence is via **the objective**, namely empirical state visitation

$$G_t(d) = \min_{\theta \in \Theta_t} F_{\theta} \left(\frac{1}{t+1}d + \frac{t}{t+1} \hat{d}_t \right)$$

Empirical visitation of past interactions (induces history dependence)

- The overall policy over episodes/rounds is **non-Markovian**.

Competitive with respect to the unrelaxed objective?

$$\underbrace{\mathbb{E}_{\eta_i \sim \pi_i^*, i \in [\tau]}}_{\text{Optimal policy of the unrelaxed objective}} \left[F_{\theta_*} \left(\frac{1}{\tau} \sum_{i=1}^{\tau} \hat{\eta}_i \right) \right] - \underbrace{F_{\theta_*}(\hat{\eta}_{\tau})}_{\text{Empirical state-visitation (i.e. integral lattice point)}} \stackrel{\text{Jensen}}{\leq} F_{\theta_*}(\eta^*) - \underbrace{F_{\theta_*}(\hat{\eta}_{\tau})}_{\text{Empirical state-visitation (i.e. integral lattice point)}} \leq \mathcal{O}(\tau^{-1/2})$$

Optimal policy of the unrelaxed objective

Difficulty of the unrelaxed problem (in terms of combinatorial optimization)

- Results Known θ_* , Additional structure: **Submodularity (D-design)**
- Optimal policy is non-Markovian and deterministic (interesting)
- Optimal problem is **NP-hard to mult. approximate** with constant factor

Def: Experiment Design problem (on MDP)

$$\arg \max_{\pi_1, \dots, \pi_T} \mathbb{E}_{\pi_i} [s(\mathbf{E}_{\theta_*})(\mathbf{X})]$$

$\mathbf{X}$ contains concatenated trajectories

Continuous Domains and Model Predictive Control (MPC)

So far, the decision space was discrete

- Continuous decision space $\rightarrow$ no change in terms of methodology
- Problem is still convex in the probability distribution space
 - infinite dimensional convex optimization problem

$$d^* = \arg \max_{d \in \mathcal{D} \subset \mathcal{P}(\mathcal{X}, \mathcal{A})} s \left(E_{\theta_*} \left(T \int_{\mathcal{X}} d(x, a) r(\Phi(x, a)) dx da \right) \right)$$

- Need to resort to some form of representation.
 - Mixture of Gaussians, Normalizing flows, Neural Networks
 - **Model Predictive Control: Parametrize in terms of actions**
- **Suppose Linear Dynamics** $x_{h+1} = Ax_h + Ba_h$
- The policy is parametrized by the set of actions a_h

Past visitation of states

$$\arg \max_{a_0, \dots, a_H} \sum_{h=0}^H \nabla F_{t,0}(d_{\pi_{\text{mix},t}})(x_h, a_h)$$

s.t. $\|a_h\| \leq a_{\max} x_h \in [-0.5, 0.5]$
 $x_{h+1} = Ax_h + Ba_h$

Non-linear MPC

- Non-convex objective
- Adaptivity means: **receding horizon replanning**

Experiment Design with MDPs: Examples I

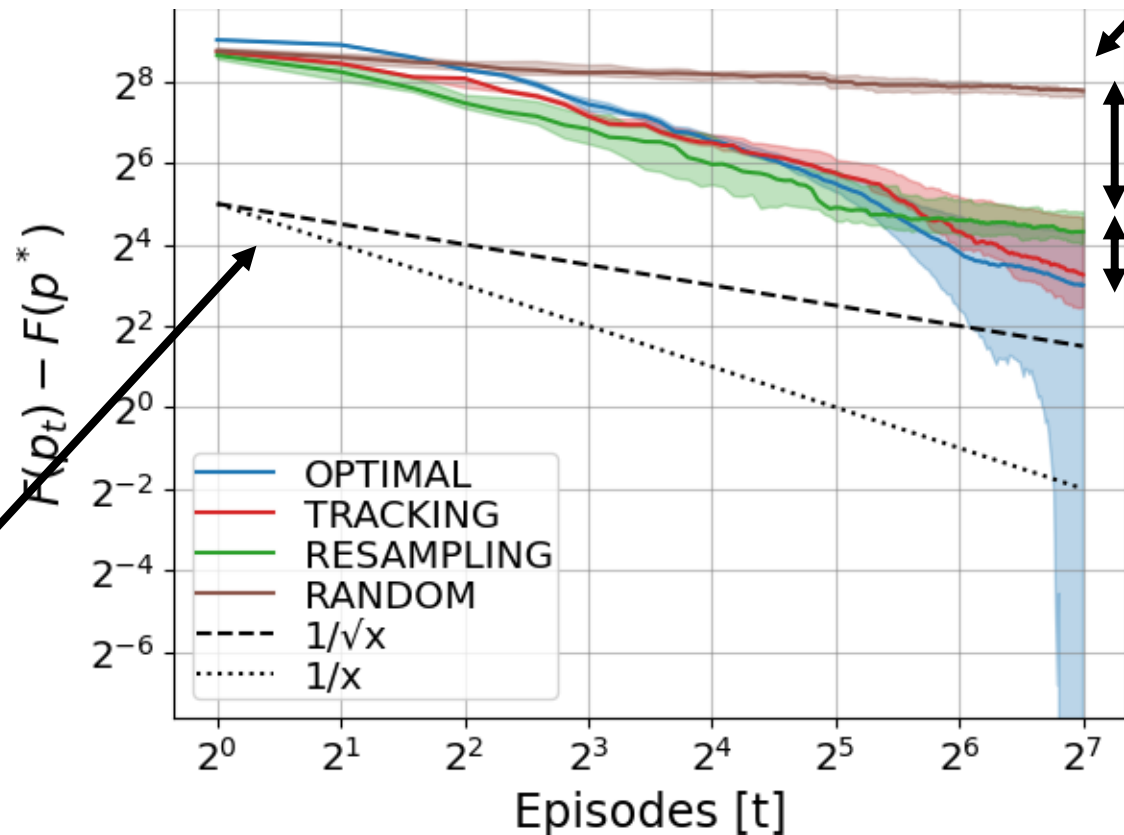
Spatial Sensing

- **Goal:** Estimate the spatial distribution of species
- Drone flying over and local constraints
- Features: elevation and slope

“The y-axis is value of the constant:”

$$\frac{F(p_t)}{t}$$

Theory well-represented



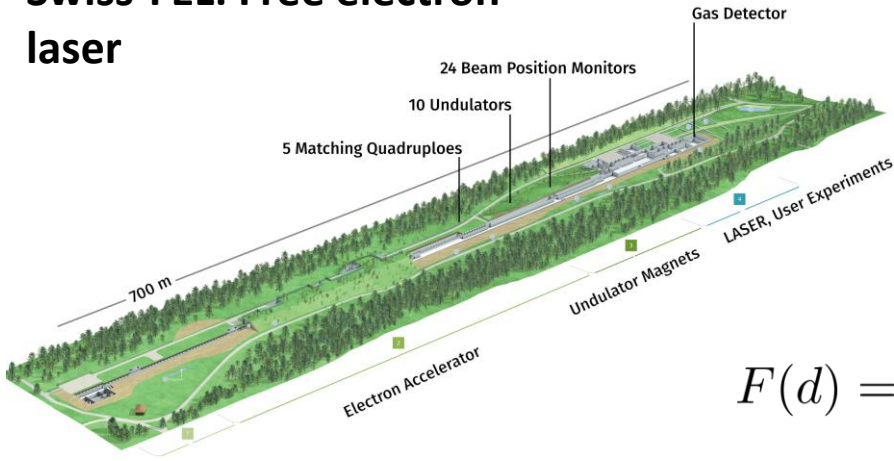
Random flatlines;
It can't compete with
optimal policy

Gap due to ED

Adaptivity Gap

Experiment Design with MDPs: Examples III

Swiss-FEL: Free electron laser

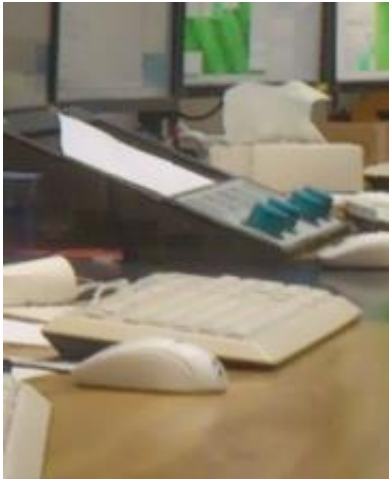


- **Goal:** Identify best-parameter
- By changing parameters, system is not equilibrium.
- The larger the parameter change the more noise

$$F(d) = \max_{z, z' \in \mathcal{Z}} \|\Phi(z) - \Phi(z')\|_{\mathbf{V}(d)^{-1}}^2$$

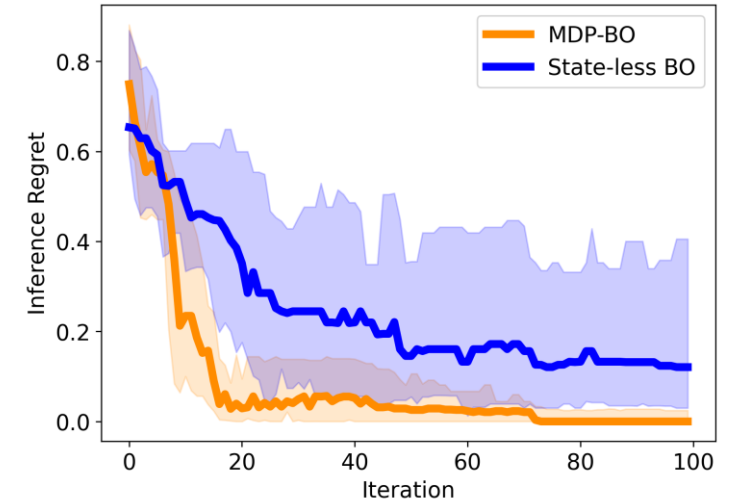
$$\mathbf{V}(\eta) = \left(\sum_{x, a \in \mathcal{X} \times \mathcal{A}} \frac{d(x, a) \Phi(x, a) \Phi(x, a)^\top}{\sigma^2(x, a)} + \lambda \mathbf{I} \right)$$

The noise is **proportional to the distance between $x \rightarrow a$**



Johannes Kirschner

Nicole Hiller



Experiment Design with MDPs: Examples IV

Optimizing preference models with MDPs

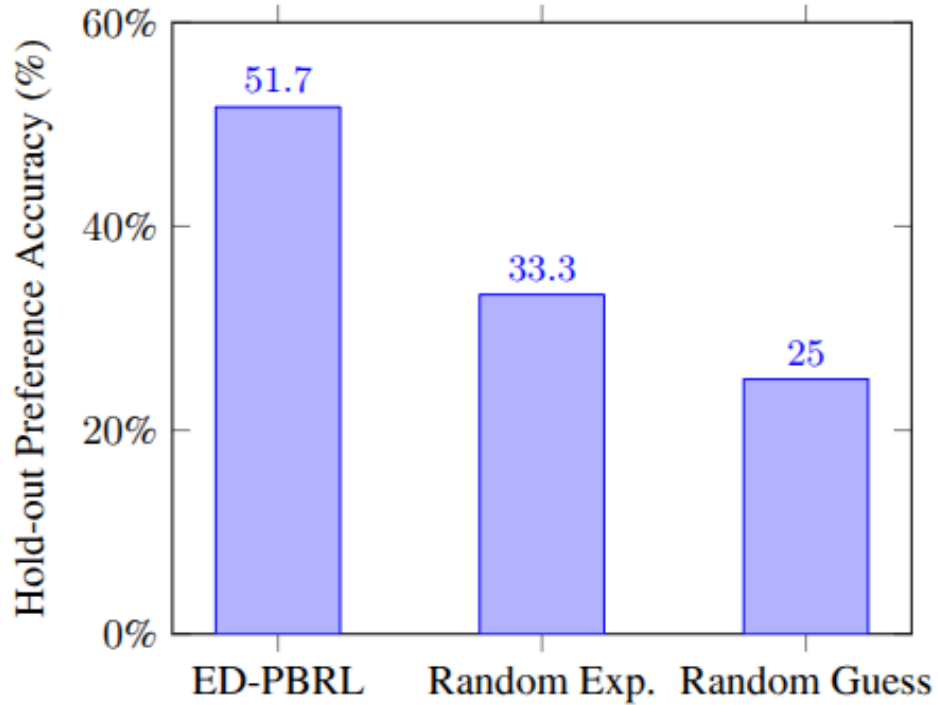


Figure: Human Experiment

- Our algorithm – ED-PBRL
- Budget: 100 queries
 - Accuracy of model trained on 100 queries
 - 20 participants



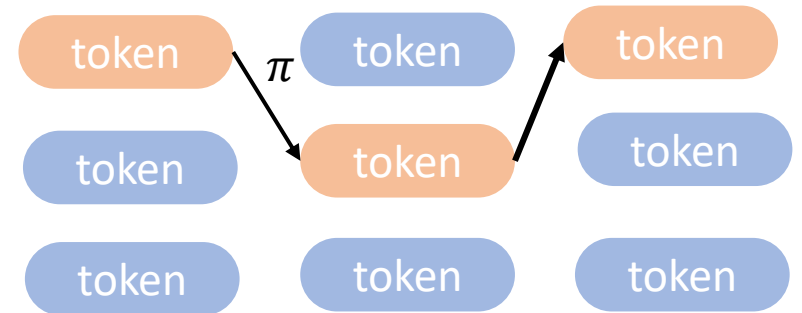
Q: Can we learn a model?

$F(\text{user}, \text{prompt}) \rightarrow \text{\#hashtag1} \text{\#hashtag2}$

Challenges:

- Per-user interaction
→ **data efficiency**
- Feedback is only relative
→ **preference feedback**
- Structured feedback
→ **LLMs sampling is a Markov chain!**

What prompts we show to the user?



$$\pi_1, \pi_2 = \operatorname{argmin} F(d_{\pi_1}, d_{\pi_2})$$

Thank you for your attention

Andreas Krause



Tadeusz Janik



Manish Prajapat



Antanas Murelis



Johannes Kirschner



Jose Folch



Ruth Misener

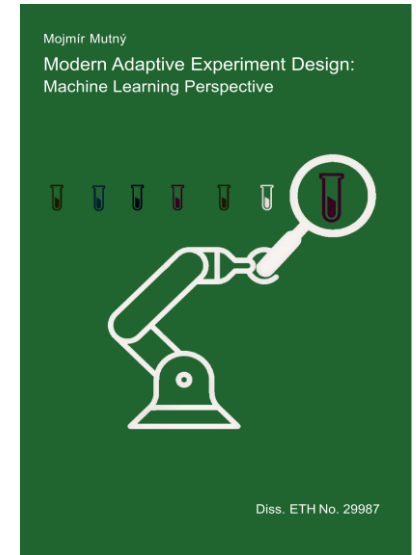


Nicole Hiller



Eidgenössische Technische Hochschule Zürich
Swiss Federal Institute of Technology Zurich

Imperial College
London



Basic framework in the thesis,
more recent works on my
website.

Experiment Design with MDPs: Examples II

Optimizing batch chemical reactors

- Complicated reaction mechanism (non-linear)
 - Not first-order reaction
- **Goal:** Identify best-reaction parameters (maximization identification)
- Opening a valve & changing residence time is continuous
 - → Transient information
 - → Noisy information
 - → Optimal path in parameter space

$$\text{ratio} = \frac{\text{flow reactant 1}}{\text{flow reactant 2}}$$

