

Instrumental Variable Methods for Sequential Experiment Design

Elisabeth Ailer

Vollständiger Abdruck der von der TUM School of Computation, Information and Technology der
Technischen Universität München zur Erlangung einer
Doktorin der Naturwissenschaften (Dr. rer. nat.)
genehmigten Dissertation.

Vorsitz: Prof. Dr. Mathias Drton

Prüfende der Dissertation:

1. Prof. Dr. Niki Kilbertus
2. Prof. Ricardo Silva, Ph.D.

Die Dissertation wurde am 17.04.2025 bei der Technischen Universität München eingereicht und
durch die TUM School of Computation, Information and Technology am 26.10.2025 angenommen.

Acknowledgments

During the final months of writing the thesis, one of my guilty pleasures (if one can even call it that) was reading the “Acknowledgments” of other PhD thesis. These acknowledgments allowed me to peek into these lives that all circled around writing their thesis. They wrote about challenges they had to overcome, support, love and guidance they received. Some eloquently thanked their partners, collaborators, friends, families. Every sentence felt like it was poured directly from their soul and touched my heart. Some only dedicated the thesis to their parents, their grandma or their brother, carrying a deep feeling only via those three simple words. *To my mom.* While reading, my subconscious already formulated its own acknowledgments, paragraph by paragraph, chapter by chapter. Thanking each and every individual who helped me in those last four years, thanking my supervisor, my parents, my partner, my sister. With each further paragraph, a warm feeling of love rolled over me and just thinking about writing this chapter left me utterly grateful. Writing my thesis made me feel lonely at times, but it was exactly in those moments that I realized that this could not be further from the truth. I am so very grateful for all these wonderful, amazing, loving, and brilliant people in my life who would blow this page completely out of proportion and into a whole book of its own.

But I will be brief.

Thank you Jason Hartford, for lightening up days with inspirational words and making me never lose fun learning new things during this whole experience.

Thank you Niki Kilbertus, for being an utterly inspiring leader, for consistently providing invaluable guidance, support and knowledge.

Thank you Mom and Dad, for becoming my greatest cheerleaders.

Thank you Maria for always knowing what to say. Always.

Thank you Daniel for loving me without holding back.

And thank you to the many unnamed people here. Thank you for making this page look lousy compared to the book I feel.

Abstract

This thesis focuses on quantifying causal relationships in complex, high-dimensional systems subject to hidden confounders by developing instrumental variable techniques for sequential experimental design.

We are motivated by challenges that arise in the context of cause-effect estimation in biological datasets. These challenges include—among others—the high-dimensionality and the hidden confounding of the datasets, the nonlinear effects, and the limited experimental resources to explore causal relationships. Instrumental variable methods can help infer causal effects in such settings. Therefore, we first identify limitations of current instrumental variable methods and then extend them to also accommodate sequential estimation settings.

We address this in two steps.

First, we note that the high-dimensionality of the treatment space and the lack of equally high-dimensional instruments limit current IV methods that require at least fully specified settings (i.e. the instrument space is at least as large as the treatment space). Therefore, we initially develop a method that can compute an estimate for the underspecified linear setting, i.e. where the treatment space is larger than the instrument space.

Although we show that we can only learn a specific part of the causal effect in those underspecified settings, we also identify an explicit representation of this subspace, enabling us to combine estimates of different experimental runs from different data and instruments. Based on the current combined estimate, we select the next experiment from a predefined set which provides us with targeted information about the (yet) unknown parts of the causal effect. Moreover, we derive a stopping criterion, informing us when the full causal effect is identified, even if only a part of the treatment space was covered by the instruments.

In the second part, we address two key constraints: we extend the setting from linear to nonlinear, and from predefined experiment selection to structured experiment design. Although the shift to nonlinear settings means that complete effect identification is not

feasible, it underscores an important scientific insight: queries do not always necessitate full effect determination. In many cases, possessing specific knowledge about an aspect of the causal effect is sufficient to answer causal research questions. Thus, we concentrate on estimating bounds for a specific aspect of the causal effect. A key component of this approach is the adaptive instrument sampling strategy which leverages the bounds to inform the gradient updates to get a subsequent experiment that would then produce even tighter bounds. Thus, experiment-by-experiment, we devise the most informative policy, effectively reducing the gap between the bounds.

In summary, this work contributes to the theory and practice of causal inference by advancing the IV framework, emphasizing its role in experimental design, and enhancing its applicability to real-world, high-dimensional data settings.

Contents

1	Introduction	1
1.1	Motivation	1
1.2	Outline and Contribution	4
1.3	Publications	4
2	Background	5
2.1	Instrumental Variables	5
2.1.1	Two-Stage Methods	10
2.1.2	Minimax Estimation	16
2.1.3	Targeted Identification of Specific Effect Properties	19
2.1.4	Assumptions	20
2.1.5	Instrumental Variables as Experiments	22
2.2	Sequential Experiment Design	26
2.3	Causality in Real Data	29
2.3.1	Microbiome Data	30
2.3.2	Genomic Data	32
2.3.3	Compositional Data	33
3	Sequential Underspecified Instrument Selection for Cause-Effect Estimation	35
3.1	Problem Setting	35
3.2	Method	36
3.3	Contribution	36
3.4	Future Work	37
3.5	Individual Contribution	37
4	Targeted Sequential Indirect Experiment Design	38
4.1	Problem Setting	38
4.2	Method	39
4.3	Contribution	39

4.4	Future Work	40
4.5	Individual Contribution	40
5	Conclusion	41
A	Appendix	43
A.1	Sequential Underspecified Instrument Selection for Cause-Effect Estimation .	43
A.2	Targeted Sequential Indirect Experiment Design	58
A.3	Instrumental Variable Estimation for Compositional Treatments	95
	List of Figures	121
	List of Tables	122
	Bibliography	123

1

Introduction

1.1 Motivation

One of the drivers of scientific research is understanding cause-effect relationships, particularly in high-dimensional systems. Most questions in the natural sciences are inherently causal: How does our gut microbiome influence health and disease? How does the transcriptome of a cell influence its function? What genes are responsible for an increased risk of breast cancer? In order to answer those questions, we need to develop methods that can handle high-dimensional data and nonlinear relationships. Although general machine learning methods skillfully address these challenges, i.e. neural networks are both universal function approximators and are able to process such data sets, this might still not suffice to answer above questions. Such models are primarily designed to identify patterns and correlations in large datasets, which allow them to perform well on prediction tasks. However, this ability to learn correlation is also their weakness: While exploiting correlational relationships to achieve their prediction goal, they overlook the foundational causal mechanisms that drive the prediction target.

Before readily dismissing these models with the oft-cited remark “Correlation is not causation,” we want to provide a more detailed motivation for causal models.

In many cases, the high-dimensional input and the outcome are confounded via unknown mechanisms: Both their gut microbiome and people’s health are influenced by, for example, their genetics, their lifestyle, their family, and their pets. The transcriptome and the cell function are affected by epigenetic modifications, cell cycle stages, and environmental stressors. Genes and breast cancer might both be affected by environmental influences, age, and hormonal changes. Although certain confounding variables can be measured or eliminated through well-structured experiments, in each of the previous examples, there are confounders that are hard to fully quantify or that are even entirely unmeasurable.

If we trained a general machine learning model on such data, without considering the confounders, the model might and probably will perform exceptionally well; as long as the confounders do not change. As soon as the model is used in a different context (for instance, in a different country with different lifestyles and environmental factors or in a different age group resulting in different hormonal levels, social backgrounds), the unobserved confounders might affect the microbiome and the health of patients in opposite ways causing the model to fail. This statement should not be interpreted as a general critique of the model's adequacy; rather, it highlights the model's intentional design for an alternative function: It is constructed to identify and learn on correlation, not to establish causal relationships. Therefore, correlational models excel in tasks like language translation and recommendation engines, among many others, they only are inadequate to be a universal solution to all problems.

We illustrate the difference by an example (see Fig. 1.1). A study collected data on gut microbiome composition and risk of depression from thousands of individuals. Training a machine learning model on this dataset, we might find an association showing that certain types of gut bacteria are linked to a higher risk of depression. The more abundant these bacteria are in an individual's gut, the higher their likelihood of experiencing depressive symptoms. Although this model might be useful for identifying individuals who would benefit from mental health monitoring, it cannot prove that changes in the gut microbiome directly reduce the risk of depression. There could be several unobserved confounding factors that explain this correlation. For example, people with a healthier gut microbiome may have different dietary habits, exercise routines, and lifestyle factors that influence both their microbial composition and their mental well-being. Consequently, the presence of certain bacteria might ultimately have no causal impact on depression risk at all. The model may still reliably indicate which individuals have a higher risk of depression, but it is unfit to suggest that altering the gut microbiome would reduce depression risk.

The model is not constructed to learn the causal mechanisms in the dataset, and thus, we will most likely not be able to solve the question of causality by building larger models. Therefore, how can we address the question of causality?

The answer is two-fold. First, we can leverage experimental access to the system in question. Interventional data is able to cut off any dependence on the confounder. In theory, an interventional dataset will look as follows: We would "set" the microbiome to one specific composition which in turn makes it independent from confounders like dietary habits and lifestyle choices. In practice, interventions such as these are rarely possible; the gut microbiome cannot be accessed directly. However, while we do not have direct control of the microbiome composition, we may have control over external perturbations such as probiotics or antibiotics. A second measure to get causal insights is to integrate prior

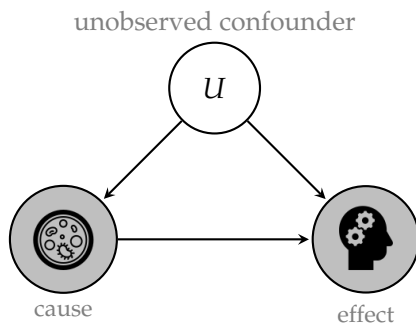


Figure 1.1 Confounded estimation of the effect of gut bacteria composition and depression.

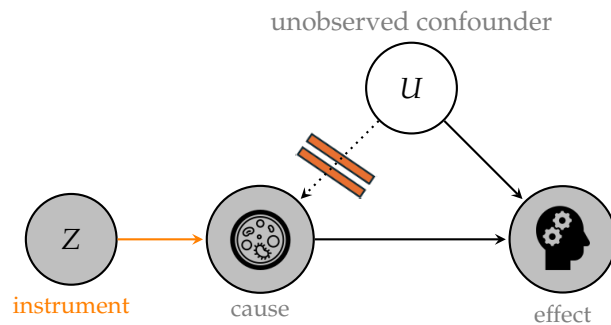


Figure 1.2 Additional instrument to enable unconfounded cause-effect estimation of gut bacteria composition on depression.

knowledge and to make assumptions about the setup by using structural equation models. In the example above, the assumption of “no unobserved confounding” trivially turns the correlational model into a causal one. In the following, we will concentrate on an approach that is able to take from both approaches: We introduce a nontrivial setting of assumptions which additionally can integrate data from indirect experiments to gain causal insights: the Instrumental Variable Setting (see Fig. 1.2).

For instrumental variables, the experimental setting has to fulfill a few assumptions while then providing the practitioner with an estimator of the causal effect, which—depending on the underlying modeling assumptions—is consistent, asymptotically normal, or at least convergent. These assumptions include that the instrument has a measurable impact on the cause, the instrument affects the effect but only via the cause, and the instrument is not affected by the unobserved confounding. If those are fulfilled, the instrument induces some variation in the cause variable that is independent of the unobserved confounding, and therefore allows causal estimation.

Although experiments help to infer causal relationships in a system that is subject to hidden confounding, sequential experimentation additionally helps to overcome the complexity of high-dimensional biological systems: A single experiment might not suffice to identify the causal effect in question. For example, we would have to use a lot of antibiotics (as instruments) to induce enough variation in all dimensions of the cause variable, which may kill or permanently damage the studied microbiome. Moreover, in biological settings, experiments are often costly, time-consuming, and sometimes ethically constrained. Therefore, we rely on sequential experiments that allow us to iteratively adapt hypothesis and structure experiments to optimally use resources and identify the true causal effect.

1.2 Outline and Contribution

This thesis explores various challenges within the instrumental variable framework that are targeted towards experimental design in biological datasets, i.e., how the instrumental variable (IV) estimate from one experiment can be used to suggest the optimal subsequent experiment in order to improve the running IV estimate within high-dimensional datasets. The thesis is based on two subsequent papers.

Section 3: Especially for high-dimensional data, the instrumental variable is lower dimensional than the cause variable. We investigate whether and what can still be learned about the causal effect in this underspecified setting. We examine how this can be set up in a sequential manner, such that subsequent experiments provide insights on the yet-unknown parts of the causal effect.

Section 4: For the nonlinear IV setting, we additionally are confronted with a general unidentifiability of the full causal effect (even in the case of matching dimensionality of instrument and cause). We propose a sequential algorithm that bounds a specific aspect of the causal effect and, based on those bounds, updates the sampling policy of the instruments for the next experiment to narrow the gap between those bounds.

1.3 Publications

This thesis is based on the following publications:

SEQUENTIAL UNDERSPECIFIED INSTRUMENT SELECTION FOR CAUSE-EFFECT ESTIMATION

Elisabeth Ailer, Jason Hartford, Niki Kilbertus.

International Conference on Machine Learning (ICML), 2023

TARGETED SEQUENTIAL INDIRECT EXPERIMENT DESIGN

Elisabeth Ailer, Jason Hartford, Niki Kilbertus.

Neural Information Processing Systems (NeurIPS), 2024

One publication which is referenced in the thesis but not part of the examination:

INSTRUMENTAL VARIABLE ESTIMATION FOR COMPOSITIONAL TREATMENTS

Elisabeth Ailer, Christian L. Müller, Niki Kilbertus.

Scientific Reports (Nature), 2025

One publication which I have contributed to but is not included in this thesis:

SUPERVISED LEARNING AND MODEL ANALYSIS WITH COMPOSITIONAL DATA

Shimeng Huang, Elisabeth Ailer, Niklas Pfister, Niki Kilbertus.

PLOS Computational Biology, vol. 19, 2022

2

Background

2.1 Instrumental Variables

History. Instrumental variables allow us to quantify the causal relationship between two variables that are subject to unobserved confounding. The term *Instrumental variables* refers to one specific set of assumptions and its collection of methods that were first mentioned in the 1920s by [Wright \(1928\)](#). Later, [Haavelmo \(1943\)](#) and [Reiersøl \(1950\)](#) also applied a similar approach in the context of errors-in-variable models and contributed to the development of IVs unaware of the contributions of the Wrights. Wrights' confounded relationship of interest was the supply and demand for butter and olive oil. As supply and demand are measured simultaneously, it is impossible to determine the causal effect of supply on demand. Simple regression analysis merely leads to biased estimates. Only by using a third variable, which is called the *instrument* and which is represented by rainfall in this example, the confounded relationship becomes quantifiable. By assuming that the instrument (rainfall) is not connected to the unobserved confounding but influences the supply of butter and olive oil, the instrument produces external variation via an as-if random relationship that resembles the one in randomized control trials, the gold standard for scientific discovery.

After the first mention in [Wright \(1928\)](#), an important breakthrough for the rise of IV methods was the development of the two-stage estimator by [Theil \(1953a;b\)](#) and [Basmann \(1957\)](#). The two-stage estimator laid the foundation for IV methods that, in consequence, have become increasingly influential since the 1950s. Especially in econometrics, the search and justification of instruments in observed settings became both an art and an incredible resource for causal insights; often followed by intense discussions. For instance, [Angrist & Krueger \(1991\)](#) use the Vietnam draft lottery to estimate the effect of school education on later wages. They argue that since the lottery-assigned priority numbers are based on birthdays, it makes them a natural experiment. Moreover, as a consequence of the draft lottery, a common strategy to defer military service was to enroll in college, which explained

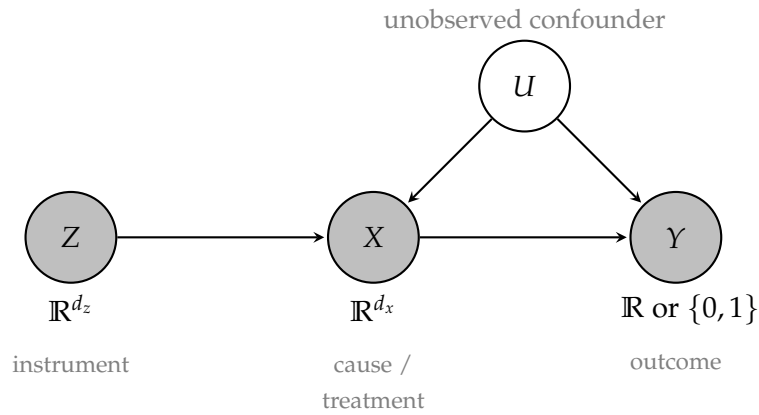


Figure 2.1 Cause-effect estimation of $X \rightarrow Y$ via an instrumental variable Z for a general treatment variable X . Note that the terms *cause* and *treatment* for X are used interchangeably.

the impact of the lottery on school education. Another controversial example of IV methods is the use of distance from college to quantify the relationship between years of schooling and later earnings (Card, 1993). Card (1993) argue that the closer potential students live to college, the more likely they are to attend (e.g. due to lower costs) while at the same time the distance to college is assumed to be unaffected by other confounding factors. Furthermore, Aizer & Doyle (2015) look at the random assignment of judges to quantify the effect of juvenile incarceration on high school completion and adult recidivism (with the assumption of the individual judge having a bias towards each sentence) to name only a few of the debated IV examples.

In econometrics, IV methods were utilized primarily in purely observational settings. However, IV methods are also helpful in actual randomized control trials (RCT) which initially they were designed to mimic. For example, studies are often confronted with the problem of potential non-compliers, i.e. individuals not following their treatment assignment, which introduces stochasticity to the treatment administration itself. In this context, instrumental variables prove to be beneficial. By considering the treatment *assignment* as the instrumental variable, researchers can accommodate scenarios in which participants deviate from their assigned treatment and still achieve valid and unbiased estimates for the treatment effect. Note that not only in the context of RCTs, the terms *treatment* and *cause* are used interchangeably, so that the treatment variable does not even need to be a literal treatment (Fig. 2.1). In this thesis, we will adopt this implicit understanding and equivalently use both terms.

Overall, although instrumental variables were commonly viewed as “natural experiments”, IV methods are now also widely used in other scientific experimental design contexts. Consequently, in recent years theoretical research on instrumental variables has regained attention: for example, in the area of relaxation of assumptions (Kilbertus et al., 2020; Padh

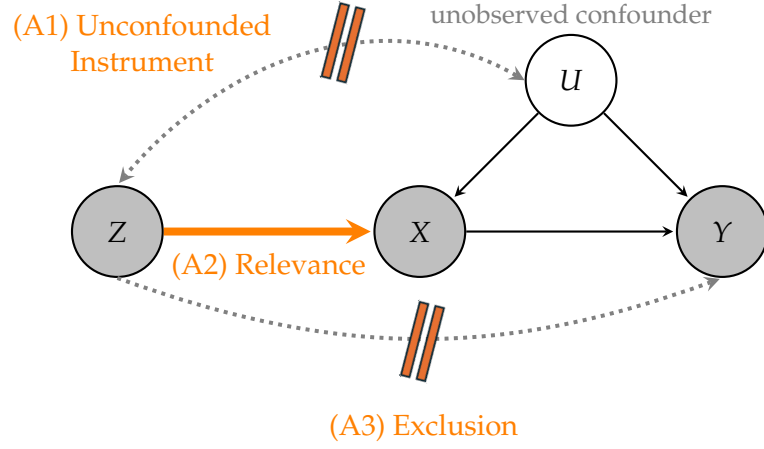


Figure 2.2 Assumptions of IV methods: (A1)-(A3) explicitly illustrate the IV assumptions which are necessary for a valid IV estimation.

et al., 2022; Bennett et al., 2023) or by taking advantage of the developments in kernel and general machine learning (Singh et al., 2019; Bennett et al., 2019; Hartford et al., 2017).

Formalization. Instrumental variables are used to quantify the causal relationship between the *treatment* variable resp. the *cause* X and the *outcome* variable Y when both are subject to unobserved confounding U . The goal is to estimate a direct causal effect of X on Y , written as $E[Y|do(x)]$ in the do-calculus notation (Pearl, 2009) or as $E[Y(x)]$ in the potential outcome framework (Imbens & Rubin, 2015), where $Y(x)$ denotes the potential outcome for treatment value x . In the remainder of the thesis, we will employ the do-calculus notation.

In order to recover the causal effect, the IV setting requires access to an additional discrete or continuous variable Z . The so-called instrument Z has to fulfill the following assumptions to be considered a valid instrument (cf. Fig. 2.2 for a visual illustration):

- (A1) *Unconfounded Instrument*, $Z \perp\!\!\!\perp U$: the confounder is independent of the instrument.
- (A2) *Relevance*, $Z \not\perp\!\!\!\perp X$: the instrument induces relevant variation in the cause variable.
- (A3) *Exclusion Restriction*, $Z \perp\!\!\!\perp Y \mid \{X, U\}$: the instrument influences the outcome only via the cause.

Moreover, the IV setting implies the following functional dependencies:

$$X = h(Z, U), \quad Y = f(X, U), \quad (2.1)$$

with $Z \in \mathbb{R}^{d_z}$, $X \in \mathbb{R}^{d_x}$, $Y \in \mathbb{R}$ as random variables and $U \in \mathbb{R}^{d_u}$ as the unobserved random confounder. In addition, $f \in \mathcal{F}$ and $h \in \mathcal{H}$ represent the functional dependencies, where $\mathcal{F}$ and $\mathcal{H}$ denote the respective function spaces.

To recover the causal effect f , we have to solve the implied Fredholm integral of the first kind (Kress, 1989).

$$\mathbb{E}[Y | Z] = \int f(x) dP(X | Z) \quad (2.2)$$

This equation is generally ill-posed. Without further restrictions on f and h , the causal effect is not identified (Pearl, 1995; Bonet, 2001; Gunsilius, 2021). A frequently used assumption to achieve identification is the *additive noise* model, namely $Y = f(X) + U$ with $\mathbb{E}[U] = 0$. Note that $X \perp\!\!\!\perp U$ still holds.

In general, the linear case (i.e. f and h are linear functions) is well understood, together with its consistent estimators (Angrist & Pischke, 2008). But even for the nonlinear case (and certain additional regularity conditions), the IV problem can be solved consistently, as demonstrated in (Newey & Powell, 2003; Blundell et al., 2007) and more recently (Singh et al., 2019; Muandet et al., 2019; Zhang et al., 2020c; Bennett et al., 2023). In the nonlinear case, nonetheless, the problem is typically underspecified in the sense that multiple f are compatible with the observed data and regularization techniques are used to obtain a unique solution—typically the smallest compatible f according to some norm. Besides, we generally consider further conditions to ensure identification and well-definedness when working in nonlinear regimes.

Completeness Condition. The completeness condition is key for effect identification in nonparametric settings. It ensures that the instrument Z provides “enough variation” in X . Formally, X is complete for Z , if for all measurable functions f with $\mathbb{E}[f(X)] = 0$ (D’Haultfoeulle, 2006):

$$\mathbb{E}[f(X) | Z] = 0 \text{ a.s.} \implies f(X) = 0 \text{ a.s.}$$

Note that completeness is equivalent to the injectivity of the conditional expectation operator. Moreover, Severini & Tripathi (2006) can show the following:

Lemma 1 (Severini & Tripathi (2006)). *The conditional distribution $P(X | Z)$ is complete if and only if for each function $f(x)$ such that $\mathbb{E}[f(x)] = 0$ and $\text{Var}(f(x)) > 0$, there exists a function $h(z)$ such that $f(x)$ and $h(z)$ are correlated.*

Therefore, completeness can also be interpreted as a measure of correlation between the model space $\mathcal{F}$ and the instrument space $\mathcal{H}$. For the linear multivariate case this reduces to a full rank condition on $\text{Cov}(Z, X)$. For nonparametric cases, Newey & Powell (2003) point out that the exponential family can provide a sufficient condition for identification:

Lemma 2 (Newey & Powell (2003)). *Assume that the structural equations are given as in Eq. (2.1). Assume further that these include an additive noise model for Y with $\mathbb{E}[U] = 0$. If x is absolutely continuous with the density $p(x | z) = s(x, z)t(z) \exp\{v(z)^\top \tau(x, z)\}$ as a part of the exponential*

family (with $s(x, z) > 0$, $\tau(x, z)$ is one-to-one in x and the support of $v(z)$ is an open set) then f is identified.

Source Condition. The source condition is essential for ill-posed inverse problems encountered in nonparametric IV estimation. At its core, the source condition ensures that the true causal function $f(X)$ is sufficiently smooth, allowing the convergence of estimation methods using regularization techniques such as Tikhonov regularization, which are necessary to obtain a unique solution in underspecified, nonlinear settings.

It will be useful to express Eq. (2.2) using the linear operator $\mathcal{T}$. Note that we use the space of square integrable functions for model and instrument space, i.e. $\mathcal{L}_2(X) \subset \mathcal{F}$ and $\mathcal{L}_2(Z) \subset \mathcal{H}$, a common requirement for nonparametric IV methods, e.g. (Bennett et al., 2023; Liao et al., 2020; Singh et al., 2019).¹ Let $\mathcal{T}$ be a bounded linear operator with $\mathcal{T} : \mathcal{L}_2(X) \rightarrow \mathcal{L}_2(Z)$ given as $\mathcal{T}f = \mathbb{E}[f(x) | Z]$. We denote $\mathcal{T}^* : \mathcal{L}_2(Z) \rightarrow \mathcal{L}_2(X)$ as the adjoint.

A general formulation of the source condition is provided in Bennett et al. (2022):

Definition 1. (β -Source Condition) Assume $\mathcal{T}$ to be the linear operator and $\mathcal{T}^*$ its adjoint, and we denote $\mathcal{R}(\cdot)$ as the range space. We say that the β -source condition is satisfied, if f satisfies the following for a positive exponent $\beta > 0$:

$$f \in \mathcal{R}(\mathcal{T}^*\mathcal{T})^{\frac{\beta}{2}}.$$

Unlike the completeness condition, which necessitates that the instrument exhibits “sufficient variation”, the source condition stipulates that f does not require additional variation. To illustrate the notion of smoothness for f , let us assume that the linear operator $\mathcal{T}$ admits a singular value decomposition (Bennett et al., 2023). Now, if β is large, the true function f shows greater smoothness; that is, the eigenvectors of $\mathcal{T}$ have smaller eigenvalues which decay faster.

Closedness Condition. In contrast to two-stage IV estimators (cf. Sec. 2.1.1), current minimax estimators (cf. Sec. 2.1.2) need a variant of the closedness condition to hold in order to guarantee solving a well-defined problem. Minimax estimators are constructed as a bi-level optimization problem consisting of (1) a minimization that incorporates the IV assumptions and (2) a maximization that ensures that the estimator still performs well against the worst-case scenario (for more details, we refer to Section 2.1.2). Therefore, we first need to introduce a function g and its function class $\mathcal{G}$ as the discriminator, resp. witness function, which guarantees robustness to worst-case scenarios. The closedness condition requires $\mathcal{R}(\mathcal{T})$ to be closed in the function space under consideration, which in turn ensures that the problem is well-defined.

¹ $\mathcal{L}_2(X)$ is the $\mathcal{L}_2$ -space of scalar-valued functions of X with respect to the distribution of X .

There are different variants of the closedness condition, the most general as follows:

Definition 2. (Closedness Condition) Assume $\mathcal{T}$ to be a linear operator, f as a function with corresponding function class $\mathcal{F}$, $\mathcal{G}$ as the function class of the witness functions. The closedness condition can then be written as:

$$\mathcal{T}f \subseteq \mathcal{G}, \forall f \in \mathcal{F}.$$

Bennett et al. (2022) use the adjoint for their closedness condition $\mathcal{T}^*\mathcal{G} \subseteq \mathcal{F}$, while Dikkala et al. (2020) and Liao et al. (2020) require $\mathcal{T}\mathcal{F} \subseteq \mathcal{G} + r_0$ with $r_0(Z) = \mathbb{E}[Y | Z]$.

2.1.1 Two-Stage Methods

The origins of two-stage methods date back to the 1950s with the development of the linear 2SLS estimator by Theil (1953a;b) and Basmann (1957). The expression *two-stage method* refers to the underlying methodological process which is based essentially on two sequential regression analyses. First, X is estimated using the instrument Z (first stage), followed by regressing Y on the estimated X values from the first step (second stage). Conceptually, the first stage removes the effect of U on X , isolating the exogenous variation of X , i.e. $\hat{X}$, which then allows us to estimate the causal (unconfounded) effect on Y . Note that this intuition still lacks some technicalities necessary to ensure a valid estimate, which we will provide in the following.

First, we present the linear two-stage estimator, followed by three nonlinear methodologies that consistently adhere to the two-stage framework. Nonetheless, this overview is not comprehensive, as a variety of methods have been developed for nonlinear IV frameworks, particularly in recent years, e.g. in Newey & Powell (2003); Darolles et al. (2011); Hartford et al. (2017); Lewis & Syrgkanis (2018); Singh et al. (2019); Zhang et al. (2020c); Xu et al. (2020); Saengkyongam et al. (2022).

We will overload notation and flexibly use α and β as parameters for the first, resp. second stage.

Linear 2SLS. The original two-stage methods were developed for the purely linear setting. The functional dependencies f, h follow the linear additive noise model:

$$\begin{aligned} X &= h(Z, U) = Z\alpha + U \\ Y &= f(X, U) = X\beta + U \end{aligned} \tag{2.3}$$

with $\alpha \in \mathbb{R}^{d_z \times d_x}$, $\beta \in \mathbb{R}^{d_x \times 1}$ representing the structural equations. Equivalently, we can write Eq. (2.3) by replacing the confounder U with the correlated noise variables ϵ_X, ϵ_Y :

$$\begin{aligned} X &= h(Z, \epsilon_X) = Z\alpha + \epsilon_X \\ Y &= f(X, \epsilon_Y) = X\beta + \epsilon_Y \end{aligned} \quad (2.4)$$

Consequently, we assume $\mathbb{E}[\epsilon_X] = \mathbb{E}[\epsilon_Y] = 0$.

In the linear setting, the assumptions of the standard instrumental variable (A1)-(A3) reduce to $\alpha \neq 0$, $Z \perp\!\!\!\perp \{\epsilon_X, \epsilon_Y\}$, and $Z \perp\!\!\!\perp Y \mid \{X, \epsilon_Y\}$ and estimating the causal effect corresponds to estimating the parameter β . Note that for the linear case, the completeness condition is equivalent to the rank condition and corresponds to α having full rank. Consequently, if $d_z \geq d_x$ and the rank condition on the covariance of Z and X are satisfied, β is point-identified and can be estimated consistently by the 2SLS estimator (Angrist & Pischke, 2008):

$$\widehat{\beta}_{2SLS} = (X^T P_Z X)^{-1} X^T P_Z y \quad (2.5)$$

where $P_Z = Z(Z^T Z)^{-1} Z^T$ is the projection matrix and Z, X, y the sample matrices with $\mathbb{R}^{n \times d_z}, \mathbb{R}^{n \times d_x}, \mathbb{R}^{n \times 1}$ and n the number of samples.

The 2SLS estimator is consistent and asymptotically normal with $\widehat{\beta}_{2SLS} \xrightarrow{p} \beta$ and $\sqrt{n}(\widehat{\beta}_{2SLS} - \beta) \xrightarrow{d} \mathcal{N}(\beta, \sigma^2 (X^T P_Z X)^{-1})$ with n the number of samples, and σ^2 the variance of ϵ_Y (Hahn & Hausman, 2005). Despite these properties, the original 2SLS estimator is biased (Nagar, 1959; Buse, 1992). However, current research now also includes different variants to obtain an unbiased linear IV estimator (Kymn, 1974; Chao & Swanson, 2007; Andrews & Armstrong, 2017).

Note that the standard linear IV setting only considers the just- or overidentified case, i.e. $d_z = d_x$ or $d_z > d_x$. In Chapter 3, we examine the underspecified setting for which we also allow $d_z < d_x$.

Sieve IV - Basis Function Expansion. Newey & Powell (2003) significantly advance the two-stage framework towards nonlinear models. They introduce the idea of using basis function approximations to enable nonlinear IV estimation based on the structural equations:

$$\begin{aligned} X &= h(Z, \epsilon_X) \\ Y &= f(X, \epsilon_Y) = f(X) + \epsilon_Y \end{aligned}$$

This estimator relying on basis function approximations seems to be the natural progression of the linear 2SLS estimator: The basis functions capture the nonlinearity of the function spaces while the estimation of their preceding coefficients still preserves some linearity.

The first regression stage involves estimation of the conditional means via the basis functions which are then used in turn for the second regression to approximate f . Note that the validity of the estimator is based on the underlying assumption that f can reasonably be written as a finite combination of basis functions. Consequently, we can justify writing $f = \sum_{k=1}^K \beta_k \phi_k(x)$ with $\phi_k(x)$ as the k -th basis function and $\beta_k, k \in [K]$ the coefficients to be estimated. Table 2.1 shows some examples for $\phi_k(x), k \in [K]$.

Due to the “linearity in coefficients” for the coefficients $\beta_k, k \in [K]$ the estimator resembles the 2SLS estimator in Eq. (2.5). To make this relationship more apparent, we introduce the operator $\Phi : X \rightarrow \phi_{[K]}(X)$ which represents the mapping of X onto the space spanned by the finite basis function expansion of $\phi_k, k \in [K]$. This allows us to express the nonlinear estimator in the following form:

$$\widehat{\beta_{\text{NP-2SLS}}} = \left(\Phi(X) P_{\Phi(Z)} \Phi(X) \right)^{-1} \Phi(X) P_{\Phi(Z)} y$$

As a result, the estimated causal effect can be written in terms of the basis functions and the estimated coefficients: $\hat{f}(x) = \widehat{\beta_{\text{NP-2SLS}}} \Phi(x)$.

	Basis Function	Hilbert Space	Weight Function
<i>Polynomial</i>	$\phi_k(x) = x^k$	—	—
<i>Hermite</i>	$\phi_k(x) = (-1)^k e^{x^2} \frac{d^k}{dx^k} e^{-x^2}$	$\mathcal{L}_2((-\infty, \infty))$	$v(x) = e^{-x^2}$
<i>Chebyshev</i>	$\phi_k(x) = T_k(x),$ $T_0(x) = 1, T_1(x) = 2x,$ $T_{k+1}(x) = 2xT_k(x) - T_{k-1}(x)$	$\mathcal{L}_2([-1, 1])$	$v(x) = \frac{1}{\sqrt{1-x^2}}$
<i>Legendre</i>	$\phi_k(x) = \frac{1}{2^k k!} \frac{d^k}{dx^k} (x^2 - 1)^k$	$\mathcal{L}_2([-1, 1])$	$v(x) = 1$

Table 2.1 Basis function approximation terms for $k \in [K]$: Note that certain basis functions form an orthogonal basis of the Hilbert space with inner product given by $\langle f, g \rangle = \int f(x) \overline{g(x)} v(x) dx$.

Two challenges remain in the form of the source and the completeness condition: In Newey & Powell (2003) the completeness condition is fulfilled via the completeness property of the exponential family for $P(x|z)$. This assumption guarantees the completeness of the operator $\mathcal{T}$. Moreover, they assume that the true causal function f belongs to a compact set of sufficiently smooth functions and restrict $\hat{f}$ to this set, thus also ensuring that the source condition is fulfilled.

Newey & Powell (2003) also prove the consistency of the estimator: By assuming the uniqueness of the solution and by putting some further constraints on the parameter spaces (cf. Assumptions 1-5 in Newey & Powell (2003)), they ensure $\|\beta_{\text{NP-2SLS}} - \beta_0\| \xrightarrow{p} 0$, with β_0 as the exact solution, and $\|\cdot\|$ as the norm for which the assumptions are fulfilled.

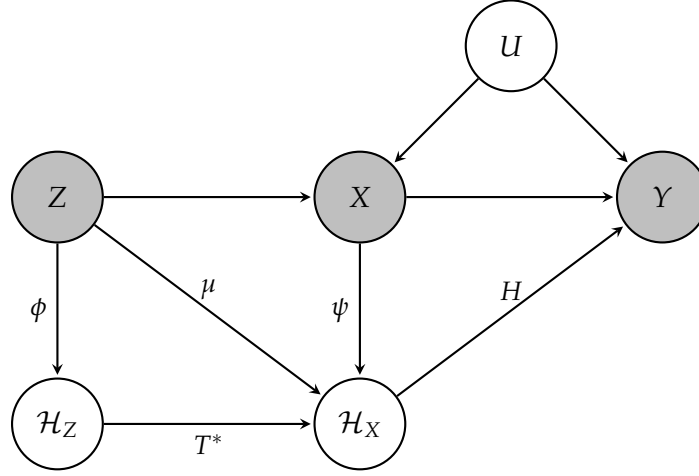


Figure 2.3 Cause-Effect Estimation with IVs in a nonlinear setting: Z and X are mapped to nonlinear function spaces $\mathcal{H}_Z$ and $\mathcal{H}_X$. The operators T^* and H between those function spaces are linear.

However, in practice, it remains difficult to pick the appropriate basis for Φ and the finite number K of approximation terms, as theoretically justified methods for choosing appropriate basis functions are not yet available (Horowitz, 2011).

KIV. Further progress in nonlinear methods is made by Singh et al. (2019). As opposed to the finite basis functions for SieveIV, they use more general *infinite* basis functions derived from kernels. Fig. 2.3 illustrates the general intuition for KIV (which is similar to Sieve IV). The functions ϕ and ψ map Z and X to the nonlinear function spaces $\mathcal{H}_Z$ and $\mathcal{H}_X$, while the operators between these function spaces remain linear.

We define the scalar-valued reproducing kernel Hilbert spaces (RKHS) as $\mathcal{H}_X$ and $\mathcal{H}_Z$ with the corresponding measurable and positive definite kernel functions $k_X : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ and $k_Z : \mathcal{Z} \times \mathcal{Z} \rightarrow \mathbb{R}$. The feature maps are defined as follows: $\psi : \mathcal{Z} \rightarrow \mathcal{H}_Z$, mapping z to $k_Z(z, \cdot)$, and $\phi : \mathcal{X} \rightarrow \mathcal{H}_X$, mapping x to $k_X(x, \cdot)$. Furthermore, we introduce the linear conditional expectation operator $T : \mathcal{H}_X \rightarrow \mathcal{H}_Z$ so that for any f , $Tf(z)$ is given by $\mathbb{E}[f(X) | z]$. This operator contains the same information as the conditional mean embedding specified by $\mu : \mathcal{Z} \rightarrow \mathcal{H}_X$ with $\mu(z) = \mathbb{E}[\phi(X) | z]$ (Singh et al., 2019).

Observe that in linear contexts, the first stage regression yields the prediction $\hat{X} = \hat{\alpha}^\top Z$. Within the kernel framework, linearity is maintained between Hilbert spaces, and the conditional mean embedding can be expressed similarly as $\mu(z) = T^*\phi(z)$. Therefore, to approximate the first stage, we choose the optimal T that minimizes the expected discrepancy in the Hilbert space with $\mathcal{L}(\mathcal{H}_X, \mathcal{H}_Z)$ as the space of Hilbert-Schmidt operators from $\mathcal{H}_X$ to

$\mathcal{H}_Z$ and T^* as the adjoint of the linear operator T :

$$\min_{T \in \mathcal{L}(\mathcal{H}_X, \mathcal{H}_Z)} \|\mathbb{E}[\psi(X) - T^*\phi(Z)]\|_{\mathcal{H}_X}^2 \quad (2.6)$$

Note that the equation would imply that the source condition is satisfied. Therefore, [Singh et al. \(2019\)](#) refer to Eq. (2.6) only as the surrogate risk. The surrogate risk upper bounds the natural risk for the conditional expectation, where the bound becomes tight when $Tf(x) \in \mathcal{H}_Z, \forall f \in \mathcal{H}_X$. Note that tight bounds would then correspond to a fulfilled source condition. The completeness condition, on the other hand, is satisfied by the requirement of k_X being a characteristic kernel function: For a sufficiently rich RKHS, the mean embedding is injective and the conditional mean embedding characterizes the full distribution of $P(X | Z)$, not just the conditional mean.

Based on the estimated conditional mean embedding of the first stage, [Singh et al. \(2019\)](#) estimate the causal effect $f \in \mathcal{H}_X$.

$$\min_{H \in \mathcal{H}_\Omega} \|\mathbb{E}[Y - H\mu(Z)]\|_{\mathcal{Y}}^2 \quad (2.7)$$

with the linear operator $H : \mathcal{H}_X \rightarrow \mathcal{Y}$, $H \in \mathcal{H}_\Omega$ and $\mathcal{H}_\Omega$ a scalar-valued RKHS.

However, solving the problem via Eq. (2.6) and Eq. (2.7) for an infinite dimensional domain is ill-posed. Therefore, [Singh et al. \(2019\)](#) additionally require smoothness and Tikhonov regularization (with ζ and λ as regularization parameters) to ensure that the problem is well defined and that the solution is unique:

$$\begin{aligned} \hat{T}^* &= \min_{T \in \mathcal{L}(\mathcal{H}_X, \mathcal{H}_Z)} \|\mathbb{E}[\psi(X) - T^*\phi(Z)]\|_{\mathcal{H}_X}^2 + \lambda \|T\|_{\mathcal{L}(\mathcal{H}_X, \mathcal{H}_Z)} \\ \hat{H} &= \min_{H \in \mathcal{H}_\Omega} \|\mathbb{E}[Y - H\mu(Z)]\|_{\mathcal{Y}}^2 + \zeta \|H\|_{\mathcal{H}_\Omega} \end{aligned} \quad (2.8)$$

[Singh et al. \(2019\)](#) also provide the empirical analogue for Eq. (2.8). Based on the empirical gram matrices, they derive a closed-form expression for the kernel estimator:

$$\begin{aligned} W &= K_{XX}(K_{ZZ} + n\lambda I)^{-1}K_{ZZ} \\ \hat{H} &= (WW^\top + m\zeta K_{XX})^{-1}y \end{aligned}$$

with K_{XX} and K_{ZZ} as the empirical gram matrices. The resulting estimator is consistent under the requirement that the RKHS regression is correctly specified and under additional mild assumptions on the function spaces $\mathcal{H}_Z, \mathcal{H}_X$ and $\mathcal{Y}$ (cf. Hypothesis 4-9 in [Singh et al. \(2019\)](#)). We can now write the estimated causal function as $\hat{f}(x) = \hat{H}K_{XX}$ analogous to 2SLS and SieveIV.

Deep IV. [Hartford et al. \(2017\)](#) incorporate neural networks into the instrumental variable framework. They propose a methodology that uses neural networks for each step to estimate the causal effect f . Despite adopting a two-stage procedure, this method stands out from previous ones, since it does not rely on a form of “linearity in coefficients” between function spaces.

The general minimization is as follows:

$$\min_{f \in \mathcal{F}} \left\| \left(Y - \int f(X) dP(X | z) \right) \right\|^2 \quad (2.9)$$

For the first stage, [Hartford et al. \(2017\)](#) need to estimate the unknown conditional distribution of the treatment given the instruments $P(X | Z)$ which then serves as input for Eq. (2.9). Note that instead of the conditional mean ([Newey & Powell, 2003](#); [Angrist & Pischke, 2008](#)) or the conditional mean embedding ([Singh et al., 2019](#)), [Hartford et al. \(2017\)](#) are required to learn the full conditional distribution $P_\alpha(X | Z)$ with α as the network parameters that parameterize the conditional distribution. Depending on whether the conditional distribution is categorical or continuous, P_α is modeled either via a deep neural network, for which the softmax output corresponds to the probability of each category, or via a mixture density network (cf. §5.6. in [Bishop \(2006\)](#)).

In the second stage regression, they recover the causal effect f by integrating with the estimated conditional distribution $\hat{P}_\alpha(X | Z)$. They parameterize f by a neural network with parameters β and learn the causal effect by the following objective function:

$$\min_{\beta} \left\| \left(Y - \int f_\beta(X) d\hat{P}_\alpha(X | z) \right) \right\|^2$$

[Hartford et al. \(2017\)](#) give the empirical analogue for the first and second stage which they validate on a held-out data sample of the empirical data set. The first stage reduces to a standard probability density estimation problem:

$$\min_{\alpha} \sum_{i=1}^n -\log P_\alpha(x_i | z_i)$$

The second stage builds on the estimated conditional distribution and the integral can either be solved exact or via Monte Carlo sampling from the estimated $\hat{P}_\alpha$:

$$\min_{\beta} \frac{1}{n} \sum_{i=1}^n \left(y_i - \int f_\beta(x_i) d\hat{P}_\alpha(x_i | z_i) \right)^2.$$

2.1.2 Minimax Estimation

In addition to two-stage estimators, we can formulate a solution to the instrumental variable condition $\mathbb{E}[Y - f(X) | Z] = 0$ by using minimax optimization (Muandet et al., 2019; Bennett et al., 2019; Liao et al., 2020; Dikkala et al., 2020; Bennett et al., 2022; 2023; Zhang et al., 2023).

Wald (1940) formalized the general minimax principle that only later found its way into econometrics and instrumental variable methods. The underlying goal is to make optimal decisions under uncertainty: The minimax principle aims to construct an estimator which performs well against the worst-case scenario within a pre-specified model class.

For instrumental variables, the minimax optimization approach builds on two insights. First, we observe that we only want to consider functions f , such that the residuals $U_f := Y - f(X)$ are independent of the instrument Z to adhere to the IV assumptions and, secondly, that if two random variables are independent, in our case $U_f \perp\!\!\!\perp Z$, we have $E[g(Z)U_f] = E[g(Z)]E[U_f]$ for all test functions g with $g \in \mathcal{G}$. Consequently, we search for an f that minimizes the residual U_f while maintaining the independence constraint which we find by solving a minimax optimization that minimizes f over a maximum worst case g . As a result, this approach ensures that the solution is more robust to violations of assumptions.

In this context, it is important to highlight that two-stage estimators and the minimax formulation do not represent *completely separate* classes of estimators. Specifically, also the two-stage least squares (2SLS) approach can be interpreted as a minimax estimation: It minimizes the maximum possible estimation error, given the moment conditions, which is nothing more than ensuring the orthogonality between the instruments and the error term $(Y - f(X))$ for linear g and f . At least for the linear case, both solutions are equivalent. However, a key distinction between these approaches is their underlying objective: (1) two-stage estimators focus on the act of estimation, while (2) minimax estimation emphasizes achieving robustness.

DeepGMM. Although the generalized method of moments (GMM) is not a minimax formulation per se, Bennett et al. (2019) show that GMMs can be interpreted as a minimax problem and thus pave the way for a more general minimax formulation within instrumental variable settings.

First, we introduce the standalone concept of GMMs. The method is essential when working with less rich function classes, as it guarantees more robustness and more flexibility towards the error distribution. Consequently, even though we introduce GMMs in the context of instrumental variables, the method extends well beyond this scope.

GMMs leverage the moment conditions implied by the IV setting. As we assume to have access to a valid instrument, the moment conditions are derived from $Z \perp\!\!\!\perp U$: If $\mathbb{E}[U] = 0$,

we know that $\mathbb{E}[g_k(Z)U] = \mathbb{E}[g_k(Z)]\mathbb{E}[U] = 0$ for all test functions $g_k \in \mathcal{G}$. Thus, for given functions $g_k, k \in [K]$ and β as a parameterization for the causal effect f , the moment conditions are the following:

$$\zeta(g; \beta) := \mathbb{E}[g_k(Z)(Y - f_\beta(X))] = 0, \forall k \in [K] \quad (2.10)$$

For the estimation of β , we consider the empirical counterpart of Eq. (2.10), i.e. $\zeta(g; \beta) = \frac{1}{n} \sum_{i=1}^n g(Z_i)(Y_i - f_\beta(X_i))$ and aim to minimize all moment conditions simultaneously:

$$\widehat{\beta}_{\text{GMM}} \in \arg \min_{\beta} \left\| (\zeta(g_1; \beta), \zeta(g_2; \beta), \dots, \zeta(g_K; \beta)) \right\|^2$$

with $\|v\|^2 = v^\top v$ the Euclidean norm. Other vector norms are also viable. [Lewis & Syrgkanis \(2018\)](#), for example, propose the maximum norm while [Hansen \(1982\)](#) show that the inverse covariance weighting is the optimal weighting in terms of minimal variance.

There is one caveat: A large function class for g requires possibly infinitely many moment conditions, and standard GMM estimators have trouble handling this for either norm. Attributing equal importance to each of the moment conditions makes optimization infeasible, while using an optimal weighting will cancel out most moment conditions due to a vanishing pseudoinverse. In order to allow flexible function classes despite these limitations, [Bennett et al. \(2019\)](#) reformulate the above GMM equation to mimic a smooth zero-sum game between f and the witness function g . Note that we include the parameter γ to parameterize the witness function class $g \in \mathcal{G}$.

$$\begin{aligned} \widehat{\beta}_{\text{Deep-GMM}} &\in \arg \min_{\beta} \sup_{\gamma} \mathbb{E}[(Y - f_\beta(X))g_\gamma(Z)] - \frac{1}{4} \mathbb{E}[(Y - f_\beta(X))^2 g_\gamma^2(Z)] \\ \widehat{\beta}_{\text{Deep-GMM}} &\in \arg \min_{\beta} \sup_{\gamma} \frac{1}{n} \sum_{i=1}^n (Y_i - f_\beta(X_i))g_\gamma(Z_i) - \frac{1}{4n} \sum_{i=1}^n (Y_i - f_\beta(X_i))^2 g_\gamma^2(Z_i) \end{aligned}$$

Moreover, [Bennett et al. \(2019\)](#) can prove consistency for $\widehat{\beta}_{\text{Deep-GMM}}$ under particular assumptions, such as the uniqueness of the solution β , as well as a source and closedness condition which they achieve by restricting the function classes of $\mathcal{G}$ and $\mathcal{F}$.

General Minimax Formulation. Inspired by [Bennett et al. \(2019\)](#), we write the nonparametric IV problem as a general minimax optimization problem of the following form ([Lewis & Syrgkanis, 2018](#); [Muandet et al., 2019](#); [Dikkala et al., 2020](#); [Liao et al., 2020](#); [Bennett et al., 2023](#)):

$$\arg \inf_{f \in \mathcal{F}} \sup_{g \in \mathcal{G}} L(f, g) = \arg \inf_{f \in \mathcal{F}} \sup_{g \in \mathcal{G}} \mathbb{E}[(Y - f(X))g(Z)]$$

where $\mathcal{L}(f, g)$ is an objective function mapping from $\mathcal{F} \times \mathcal{G}$ to $\mathbb{R}$.

In the subsequent discussion, we introduce the more general minimax estimator following the work of [Bennett et al. \(2023\)](#): We assume f to be from a set of hypothesis functions $\mathcal{F} \subset \mathcal{L}_2(\mathcal{X})$ and g to be part of the witness function class $\mathcal{G} \subset \mathcal{L}_2(\mathcal{Z})$. Again, we employ the notation of $r_0 := E[Y | Z]$ and $\mathcal{T}f = \mathbb{E}[f(X) | Z]$, with $\mathcal{T}$ as a bounded linear functional to facilitate computations.

Note, that if we assume the existence of a solution for the IV condition $\mathcal{T}f = r_0$, there might be infinitely many solutions for f . [Bennett et al. \(2023\)](#) propose to target the solution that achieves the least norm.

$$f = \arg \min_{f \in \mathcal{L}_2(\mathcal{X})} \frac{1}{2} \langle f, f \rangle_{\mathcal{L}_2(\mathcal{X})} \quad \text{such that } \mathcal{T}f = r_0.$$

which is then reformulated based on Lagrange multipliers:

$$f = \arg \min_{f \in \mathcal{L}_2(\mathcal{X})} \sup_{g \in \mathcal{L}_2(\mathcal{Z})} \langle f, f \rangle_{\mathcal{L}_2(\mathcal{X})} + \langle r_0 - \mathcal{T}f, g \rangle_{\mathcal{L}_2(\mathcal{Z})}. \quad (2.11)$$

To construct an estimator based on Eq. (2.11), they rewrite the inner products into expectations with respect to the distribution of X and Z and impose further restrictions on the function classes:

$$\hat{f} = \arg \min_{f \in \mathcal{F}} \sup_{g \in \mathcal{G}} \frac{1}{2} \mathbb{E}_n[f^2(X)] + \mathbb{E}_n[(Y - f(X))g(Z)]$$

with $\mathbb{E}_n[\cdot]$ as the empirical expectation based on the data set $\{X_i, Z_i, Y_i\}_{i \in [n]}$.

This result of [Bennett et al. \(2023\)](#) holds under a source condition on $\mathcal{T}$ (i.e. $r_0 \in \mathcal{R}(\mathcal{T}\mathcal{T}^*)$) and mild realizability assumptions regarding the capacity of the (statistically restricted) function classes $\mathcal{F}$ and $\mathcal{G}$ ([Bennett et al., 2023](#), Assumptions 2, 3, 4).

Various other works leverage the minimax formulation differing mostly with respect to their assumptions, regularization techniques and convergence rates: For instance, [Dikkala et al. \(2020\)](#) and [Liao et al. \(2020\)](#) both need a source condition of $f \in \mathcal{R}(\mathcal{T}^*\mathcal{T})$, a realizability condition (the functional classes $\mathcal{G}$ and $\mathcal{F}$ are well-specified and do contain some solution saddle points) and a closedness condition with $\mathcal{T}\mathcal{F} \subset \mathcal{G} + r_0$. However, [Liao et al. \(2020\)](#) can prove $\mathcal{L}_2$ -convergence with a rate of $n^{-\frac{1}{6}}$ compared to [Dikkala et al. \(2020\)](#) who only achieve convergence with respect to the projected MSE. Nevertheless, [Liao et al. \(2020\)](#) obtain this stronger convergence only by assuming the uniqueness of the solution to $\mathcal{T}f = r_0$ and thus requiring that $\mathcal{T}$ is injective. Since [Dikkala et al. \(2020\)](#) provide a less strong convergence guarantee, they, however, do not rely on this arguably strong uniqueness assumption. Subsequently, [Bennett et al. \(2022\)](#) build on these results and relax the closedness

condition. Finally, in the presented approach, [Bennett et al. \(2023\)](#) can completely abandon any assumptions about the closedness condition by leveraging Lagrange multipliers.

Using the minimax formulation on real data is quite challenging due to the bi-level optimization. Consequently, in addition to theoretical advances, [Dikkala et al. \(2020\)](#) provide a result that facilitates estimation. They show that if $\mathcal{G}$ is generated based on kernel functions, the inner optimization can be solved in closed form:

$$\sup_{g \in \mathcal{G}} \langle T f_{\beta}, g \rangle = \left(\frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n (y_i - f_{\beta}(x_i)) k(z_i, z_j) (y_j - f_{\beta}(x_j)) \right)^{\frac{1}{2}}.$$

2.1.3 Targeted Identification of Specific Effect Properties

For instrumental variable methods, full effect identification is not always feasible. When considering real world examples, however, full effect identification might not even be necessary as scientific questions are typically more narrow: ‘Can I expect changes in body mass index when I increase the relative abundance of Lactobacillus bacteria in the gut? How does the susceptibility of breast cancer depend on BRCA1 expression levels?’ Such targeted scientific queries do not require complete knowledge of the underlying mechanism f , but focus only on a specific aspect of f . We can formally express this as a functional Q of f .

Functionals define a mapping from the function space to real numbers: Consider $\mathcal{F}$ as a vector space of functions; then a functional Q is defined by assigning a real number to each function, i.e. $Q : \mathcal{F} \rightarrow \mathbb{R}$. Generic functionals can measure both local as well as global properties of any $f \in \mathcal{F}$ such as simple average treatment effects via $Q[f] := \mathbb{E}[Y | do(X = x^*)] = f(x^*)$ or projections of f onto a fixed basis function f^* via $Q[f] := \langle f^*, f \rangle_{\mathcal{F}} / \|f\|_{\mathcal{F}}^2$ assuming $\mathcal{F}$ is a Hilbert space or a local causal effect of an individual component X_k on Y at x^* via $Q[f] := (\partial_k f)(x^*)$.² In this case, x^* could be the mean of the observed treatments, i.e., representing a “base gene expression level” and our interest is how a local change away from that base level affects the outcome. Namely, examples for functionals (also) include:

1. *Evaluation Functional*: $Q[f] = f(x^*)$.
2. *Gradient Functional*: $Q[f] = \partial_k f(x^*)$ for $k \in [d_x]$.
3. *Fourier Coefficient Functional*: $Q[f] = \int f(x) \exp\{i\omega x\} dx$ at frequency ω .
4. *Integral Transform*: $Q[f] = \int_{u_1}^{u_2} f(u) k(u, x) du$ with k as the kernel of the integral transform.

[Severini & Tripathi \(2006\)](#) state an identification result connecting the function space $\mathcal{F}$ and the space of all linear bounded continuous functionals on $\mathcal{F}$:

²We write $(\partial_k f)(x^*)$ for the partial derivative of f with respect to the k -th argument evaluated at x^* .

Theorem 1 (Severini & Tripathi (2006)). *f is identified if and only if all bounded linear functionals of f are identified.*

However, the main question for the identification of a causal aspect is: When can we identify the functional but not the function itself? Severini & Tripathi (2006) provide an illustrative example, considering the scenario in which the conditional distribution $p(x | z)$ is not complete and therefore the full causal effect $f(x)$ is not identified. They state that if we consider the expectation functional of the form $\mathbb{E}[\omega(x)f(x)]$ for a known weight function $\omega(x)$, we can still identify it under the following assumptions:

Lemma 3. *The expectation functional $Q[f] = \mathbb{E}[\omega(x)f(x)]$ is identified if and only if $\omega(x) \in \mathcal{N}(\mathcal{T})^\perp$ with $\omega(x)$ a known weight function and $\mathbb{E}[\omega(x)^2] < \infty$ and $\mathcal{N}(\mathcal{T}) = \{f \in \mathcal{H}_{\mathcal{X}} : f \perp \mathcal{H}_{\mathcal{Z}}\}$.*

This example is generalized by the following lemma:

Lemma 4. *We have two Hilbert spaces $\mathcal{H}_{\mathcal{Z}}$ and $\mathcal{H}_{\mathcal{X}}$ with $\mathcal{H}_{\mathcal{Z}}$ the instrument space and $\mathcal{H}_{\mathcal{X}}$ as the treatment space. The functional $Q[f]$ is identified if and only if $f \in \mathcal{N}(\mathcal{T})$ then also $Q[f] = 0$.*

Therefore, we can identify the functional $Q[f]$ if it lies in $\mathcal{N}(\mathcal{T})^\perp$ and is zero otherwise. In other words, we can identify the functional $Q[f]$ if it can be completely determined by the subspace of $\mathcal{X}$ that is spanned by the exogeneous variation induced by Z , which we will later also call the instrumented subspace (Chapter 3).

2.1.4 Assumptions

The IV methods are based on strong assumptions. Although in certain cases they are testable (Pearl, 1995; Swanson et al., 2018; Gunsilius, 2021; Kédagni & Mourifié, 2020), in general, their verification depends on domain knowledge and qualitative arguments (Cinelli & Hazlett, 2022). However, verifying IV assumptions and conducting sensitivity analysis are crucial because, if violated, the resulting estimates may exhibit greater bias than the confounded OLS estimate (Bound et al., 1995). Despite numerous methods available (DiPrete & Gangl, 2004; Altonji et al., 2005; Small, 2007; Small & Rosenbaum, 2008; Conley et al., 2012; Wang et al., 2018; Kang et al., 2021; Cinelli et al., 2019), their practical application remains uncertain (Davies et al., 2013).

In the following, we focus on two concepts that stem from violations of the IV assumptions: *weak instruments* and *forbidden regression*. We refer to Ailer et al. (2025) for a more detailed discussion on the impact of these two violations on the IV estimation with a compositional treatment variable.

Weak Instruments. “Strong instruments” are a prerequisite for successful two-stage estimation in the instrumental variable setting and one of the key discussion points in real world

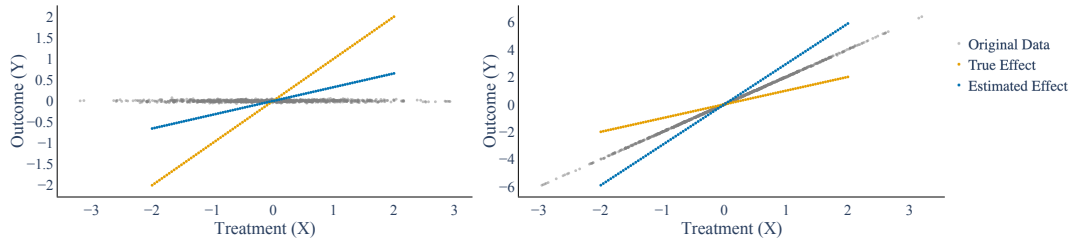


Figure 2.4 Two scenarios for weak instrument bias in a linear IV setting with $d_z = d_x = 1$: The left hand side shows a scenario with an F-statistic of 3.4 for the first stage regression and an estimated effect that is biased towards the single confounded regression. The right hand side shows a scenario where the estimation bias with instrumental variables is even worse than the single regression (F-statistic 0.4).

applications of IV. Consequently, the term “weak instrument” denotes the contrary. It refers to the lack or insufficient association of the treatment variable with the instrument.

Those “weak instruments” can lead to heavily biased estimates (Staiger & Stock, 1994; Harding et al., 2016; Hahn & Hausman, 2005) and poor finite sample properties (Hall et al., 1996). The question of how to measure instrument validity is not unambiguously clarified; there is no clear definition on which instruments are weak and which instruments are strong. Both categories rely on heuristics and empirically derived best practices. For example, in the linear setting, the instrument strength for $d_x = 1$ can be approximated using a first stage F-statistic (cf. Fig. 2.4). A value greater than 10 is generally considered sufficient to avoid weak instrument bias in 2SLS (Staiger & Stock, 1994; Andrews et al., 2019). For $d_x > 1$, measuring the strength of the instrument is more challenging even in the linear case; although the F-statistic still provides a good rule-of-thumb (Sanderson & Windmeijer, 2016; Sanderson et al., 2019).

For the nonlinear case, the detection and definition of “weak instruments” become even more tiresome. Assumptions like the completeness property or the source condition try to address this aspect (e.g. in the linear case completeness guarantees that the $Cov(X, Z)$ has full rank, but not whether the coefficients are “large enough”), but in general the verification is not quantifiable and relies on other sources: In experimental settings, for example, one can rely on already certified and independently known associations for the first stage to ensure strong association, i.e. the impact of specific antibiotics for which the effect on the microbiome is known, specific genome modifications for which the effect on the targeted gene expression has been previously seen, or, trivially, the treatment assignment on the treatment itself.

Nevertheless, an F-statistic < 10 or other methods that indicate a weak association between Z and X do not inherently prevent IV estimation; various recent methods show the possibility of IV regression even in the presence of weak instruments (Moreira, 2003; Kang et al., 2016;

Huang et al., 2021; Guo et al., 2022; Bibaut et al., 2024; Oprescu & Kallus, 2024; Wang et al., 2024).

Forbidden Regression. The term *forbidden regression* stems from the incorrect use of nonlinear models for first and second stages without considering whether the residuals of the first stage are truly uncorrelated with the instruments Z .

The notion of “forbidden regression” was first introduced by Hausman (2001). Since then, this fairly expressive term has caused some fear and confusion about instrumental variables and their use in general. It therefore counteracts the also slightly misleading notion induced by the name “two-stage” regression, which seems to invite a (careless) mix and match of different regression methods for either one of the two stages without taking into account the necessity of the instrument being uncorrelated with the first stage residuals.

To clarify both terms, we first note that for the purely linear IV setting, the estimates are consistent even for model misspecifications. For nonlinear settings, this is generally not the case; hence “forbidden regression”. One can only blindly combine different models, i.e. a logistic first stage with a linear second one, and expect consistent estimates if both stages are correctly specified, and therefore if the first stage is perfectly replicated by i.e. a logistic model. In real world settings, perfect specification is rarely ever the case. Despite this lack of consistency for general nonlinear models, there exists now a rich literature on the conditions under which “manual 2SLS” with nonlinear first (and/or second) stage can yield consistent causal estimators, originating with the work in Kelejian (1971) and continuing with cf. Sec. 2.1.1.

2.1.5 Instrumental Variables as Experiments

Two variables have a causal relationship if the change in one variable causes a change in a second variable. With the words of Pearl (2009): “There is no causation without manipulation”. For a system to reveal its underlying mechanisms, one has to *perturb* it, resp. change parts of its distribution. However, various concepts in the causal literature are related to the term “experiment”: Although randomized control trials (RCT) are still mentioned as the *gold standard* of scientific discovery, other terms such as *interventions*, *environments*, and more recently *perturbations* support cause-effect identification as well. Their common denominator is the goal to affect a system and to change the (conditional) distribution of X to learn about some causal effect of X on Y , see Fig. 2.5 for a visual illustration.

In the following, let us clarify how those terms relate to the IV setting. Note first that the terms are not only used within a strictly causal context and consequently assemble a broad range of problem settings under their name, and the IV setting might only address a subcategory within this range. In the subsequent sections, we will rely on Pearl’s *do-calculus* to build a bridge between the different terms and connect them to the IV setting.

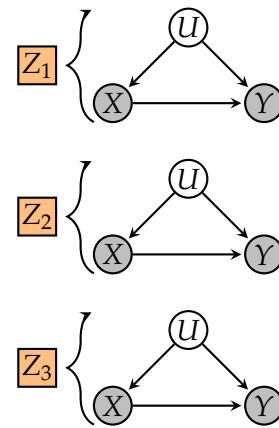
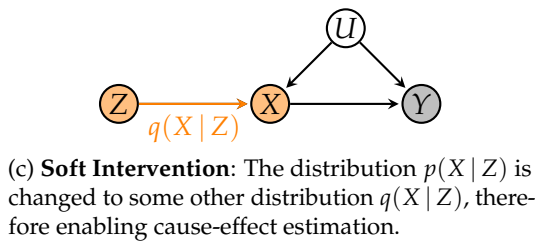
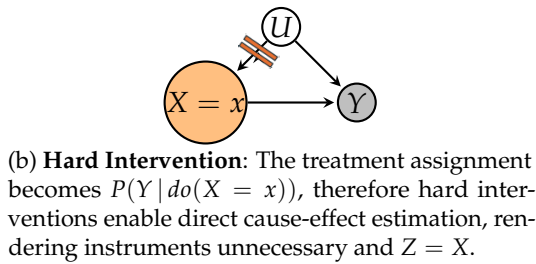
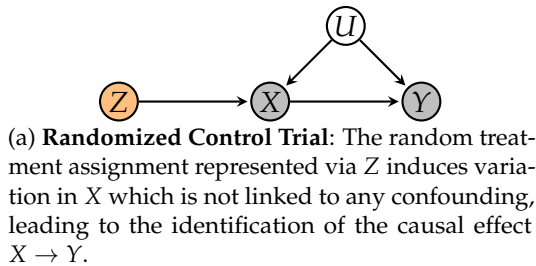


Figure 2.5 Different notions of *experiments* that enable cause-effect estimation. The shaded orange indicate the variables that are exogenously changed resp. determined by the user.

Randomized Control Trials. Randomized control trials (RCT) are a specific experimental study design that randomly assigns participants into two or more groups. Due to the random assignment, each participant has equal chances of ending up in either group and thus the effect of the treatment can be measured without being affected by confounding. Although RCTs appear to be a standalone solution to cause-effect quantification, they rely on IV methods due to issues such as treatment contamination or defier bias (Heckman, 1995; Lousdal, 2018; Sussman & Hayward, 2010) (Fig. 2.5a). In those cases, instrumental variable methods help to identify the causal effect: The random treatment *assignment* acts as the instrumental variable. Note that the random distribution of participants into treatment and control groups guarantees the independence of the assignment from any confounding factors. Therefore, assumption (A1), *Unconfounded Instrument*, is fulfilled by design. Furthermore, the assumption of *Relevance* (A2) is justified since the random assignment is a strong indicator that the participant actually receives treatment. For (A3) *Exclusion Restriction*, we reasonably assume that the random assignment has no direct impact on the outcome, since participants (and other actors in the study) are blind to the individual treatment assignment. Note that for RCTs, some of the assumptions seem to be trivially fulfilled by design (though still need to be carefully checked), emphasizing their power to perform cause-effect estimation.

Intervention. In causal inference, the term intervention has many flavors: hard interventions, soft interventions, fat-hand interventions, activity intervention, to name just a few (Meinshausen et al., 2016; Eaton & Murphy, 2007; Mooij & Heskes, 2013). The first two are the most common.

Hard interventions refer to changes in a system for which a variable's value is forcibly set to a specific constant, independent of its natural causes or parent variables in the causal graph. For instance, in the context of instrumental variable methods and referring to the earlier example for RCTs, hard interventions can mean that the "treatment assignment-treatment relationship" is considered deterministic and therefore free from treatment contamination and defier bias. The distribution becomes $P(Y | do(X = x))$ with x being the assigned treatment value. Knowing this distribution, we subsequently know the effect of $X = x$ on Y (Fig. 2.5b). Consequently, hard interventions directly enable cause-effect estimation, rendering instruments unnecessary, i.e. $Z = X$.

A less strict version than *hard interventions* are *soft interventions*. Soft interventions only modify the mechanism that generates the variable, but still allow to retain some of its dependence on other variables. Instead of forcing X to a specific value, it is set to be a random variable (Meinshausen et al., 2016); therefore, changing the distribution $p(X | Z)$ to some other distribution $q(X | Z)$ (Markowitz et al., 2005) (Fig. 2.5c). Soft interventions occur more often than hard interventions, as hard interventions are nearly impossible to achieve in practice, e.g. dictating one specific genetic code or replacing the microbiome to obtain

one specific composition is infeasible. Even CRISPR techniques fail to work flawlessly (Fu et al., 2013; Zhang et al., 2015). Instrumental variable settings can be seen as a special form of soft intervention. Soft interventions modify the distribution of $p(X | Z)$ which resembles the effect the choice of instrument Z has on the treatment variable X . If we ensure that the choice of distribution, resp. Z , is not influenced by unobserved confounding and $Z \perp\!\!\!\perp Y | \{X, U\}$, we can estimate the causal effect using standard IV methods.

Environments. The term *environment* is strongly linked to the concept of invariant mechanisms. It is based on the idea that a causal relationship remains consistent over different contexts, such as geographical regions or demographic groups, i.e. environments. In contrast to *interventions*, this concept has the advantage that the perturbed distribution does not need to be perfectly specified. It only requires (1) that the intervention within the environment does not affect the conditional distribution of the target Y and (2) that the different environments show variation in their interventions of the system (Peters et al., 2016; Meinshausen et al., 2016; Heinze-Deml et al., 2018; Creager et al., 2021) (Fig. 2.5d). The goal is to find an estimator that performs well across all environments, resp. distributions. This robustness to different distributional shifts (under the assumption of no hidden confounding, with a few exceptions) is supposed to recover the underlying causal function. Although we can generally allow for unobserved confounding between X and Y in the IV setting, we need to ensure another assumption: the environment only affects the treatment variable X . Given this, we may reinterpret different environments as different levels of the instrument Z , that is, the distribution $p_{Z_j}(x)$ changes over the different context (with Z_j as the indicator variable and $j \in [J]$ the number of environments). Based on the change in distribution, we can estimate the causal effect using standard IV methods. Peters et al. (2016), for instance, utilize this connection the other way around. They apply their algorithm on data previously analyzed in IV methods, by leveraging the IV levels to split the data into different experimental groups, i.e. environments.

Perturbations. Perturbations are often used as a collective term for distributional changes in a system (Meinshausen et al., 2016; Bühlmann, 2020). Their formulation is quite general, as it can refer to both a deliberate and a natural change in the system. The goal is to infer causal relationships by observing how outcomes respond to perturbations in specific variables. Especially in biology, the hope is to interpret gene perturbations as “soft interventions” and thus use methods from causal inference to estimate a causal effect (Dong et al., 2023; Tejada-Lapuerta et al., 2023). According to this reinterpretation, perturbations are distributional changes of $p(X | Z)$. Consequently, IV methods would be applicable again to estimate the causal effect with the previous argument for soft interventions. However, we must bear in mind that these so-interpreted “soft interventions” originated from real world data and perturbations in biological systems often not only affect the conditional distribution $p(X | Z)$,

but also have effects on the confounder or/and the outcome itself, thus being quite likely to violate IV assumptions such as the exclusion restriction.

In conclusion, *environments* and *perturbations* require more careful checks to comply with IV assumptions, while the connection to *interventions* shows that IV methods are valuable beyond traditional “cause-effect estimation via experiments” for RCTs and quasi-experiments in observational data.

2.2 Sequential Experiment Design

Sequential experiment design is a methodological framework for the adaptive planning and execution of experiments. The algorithm uses information collected from previous experiments (and their corresponding estimates) to choose or design the next experiment. Especially in biological research with limited data and limited resources sequential experiments seem like the perfect approach: The algorithm indicates to the user which data or experiment will most effectively improve the current estimate thereby making the learning faster and more precise.

The literature on causal adaptive learning comes closest to achieving our objective of deducing causal effects through sequential experiments. Moreover, some work already links adaptive experiment design and the IV framework: For example, [Hahn et al. \(2011\)](#) use the initial results of a previous instrumental variable experiment and alter conditional treatment assignment probabilities in the subsequent experiment to estimate the causal treatment effect more precisely. [Hüyük et al. \(2024\)](#) iteratively update the random treatment assignment in an exploratory clinical trial design with the goal of identifying subpopulations that exhibit positive treatment effects. This is particularly aimed at scenarios for which an overall positive treatment effect might not be applicable. [Chandak et al. \(2024\)](#) develop a strategy for targeted sequential experimentation supported by influence functions that selects instruments to efficiently estimate the full causal treatment effect f with a focus on sample efficiency and variance reduction. Moreover, adaptive learning has also been used in another context for cause-effect estimation by “deciding what to observe” assuming the availability of different data sources ([Gupta et al., 2021](#)).

However, the notion of sequential experiment design also connects to other methods of well-known settings such as *active learning*, *reinforcement learning*, and *bayesian optimization*. We will briefly clarify the differences and similarities and thus differentiate the key concepts necessary for targeted experimentation for cause-effect estimation with unobserved confounders.

Causal Active Learning. Active learning aims to improve a model’s current estimate by choosing the most informative data points for labeling. There are various methods that

assess the unlabeled data points with respect to their potential of, for example, reducing model uncertainty or stabilizing the model. For causal active learning, the choice of which data points to label becomes the choice of which interventions to perform: Causal active learning approaches propose the intervention that will best help to identify the underlying causal structure. Note that most of the literature on active experimentation for causality aims to learn the causal structure (instead of estimating properties of f) from interventions, resp. direct experiments (He & Geng, 2008; Sverchkov & Craven, 2017; von Kügelgen et al., 2019; Agrawal et al., 2019; Gamella & Heinze-Deml, 2020; Zemplenyi & Miller, 2021; Elahi et al., 2024). The current goal of causal active learning therefore differs from the sequential experiment design we have in mind since our main goal is the quantification (instead of the discovery) of causal structure.

However, the work of Singh (2023) serves as a starting point to connect cause-effect estimation and active learning (in the IV setting). They examine scenarios in which a learner cannot directly manipulate the treatment X but can only influence it indirectly through a control variable, resp. an instrument Z . Their objective is to predict the causal effect of $f(x_0)$ at a specific treatment value x_0 by using sequential experiments with a fixed budget of T experimental iterations. Each subsequent choice of Z values to label, i.e. $(Z_{i,t}, X_{i,t}, Y_{i,t})_{i \in [n]}$ for each experiment $t \in [T]$, depends on their estimated likelihood to generate values that are close to x_0 and consequently are most effective for the estimation of $f(x_0)$.

Causal Bayesian Optimization. Bayesian optimization is a sequential design strategy for the optimization of black-box functions where evaluation is costly and gradients are unavailable. The two main components are a Gaussian process that maintains a belief over the outcome and an acquisition function that determines the most promising setting for the next experiment based on the current distribution in order to find the global maximum of some function f . Building on this, Causal Bayesian optimization (CBO) additionally integrates the causal structure into the model and uses interventions to maximize its target (Aglietti et al., 2020; 2021; Gultchin et al., 2023; Sussex et al., 2022). Consequently, Causal Bayesian optimization is related to sequential experimentation as their higher-level approach is similar: CBO acquires new data samples from interventions in order to minimize uncertainty over some outcome; its goal is to design the subsequent experiment that is most informative of the outcome of interest. However, they differ when taken closer: First, CBO is focused on the maximization of some target Y instead of gaining information about a causal mechanism (the function that determines Y). Moreover, CBO considers direct interventions on all graph variables, which is often not feasible in indirect experiment designs. Last but not least, the objective remains unchanged over all iterations, while sequential experimentation for causal effects would require a path-dependent strategy.

Causal Reinforcement Learning. Reinforcement Learning (RL) is another subfield of machine learning in which the algorithm itself influences the course of the learning process. In RL, the goal is to find a sequence of actions that maximizes some reward. Unlike active learning, which chooses the next batch of data for labeling, reinforcement learning learns a sequence of actions, i.e. a policy π_θ over time T parameterized by θ , that maximizes some expected cumulative reward $J(\pi_\theta)$. Reinforcement learning gains knowledge through interacting with its environment using its feedback in the form of cumulative rewards R (Sutton, 2018). Some argue that RL is inherently causal (Gershman, 2017; Szepesvari, 2018; Kaddour et al., 2022; Weichwald et al., 2022; Deng et al., 2023). For instance, Markov Decision Processes are part of the formalization of the RL setting, i.e. the RL problem is (among others) defined by a state space $\mathcal{S}$, an action space $\mathcal{A}$ and a cumulative reward function R for said actions. Despite this similarity, the goals differ: The RL community has its focus on reward maximization, while the causal community aims to identify the causal structure and quantify causal effects. However, there also exists a notion of Causal RL even though a precise definition is missing: Generally, Causal RL uses traditional RL methods and incorporates additional assumptions, prior knowledge (in form of causal structure) or interventional data to further learn the causal structure (Zhang et al., 2020b; Zhang, 2020; Ruan et al., 2023).

Among other approaches (Pu & Zhang, 2021; Fu et al., 2022), Liao et al. (2021) are the first to connect reinforcement learning and instrumental variables. They use instrumental variables to find the optimal policy in an offline setting where they only have access to confounded data. However, for sequential experiments, we focus on online settings as we aim to iteratively update the experiments.

In the following, we introduce one of the main principles in reinforcement learning, which is the *REINFORCE* algorithm. Introduced by Williams (1992), *REINFORCE* was one of the first successful policy gradient methods, optimizing the policy π directly via the gradient of the expected reward $J(\pi)$, and therefore differs from early reinforcement learning methods, which focus on value-based methods such as Q-learning and the Bellman equation.

The objective function in policy gradient methods is the expected return of some policy π_θ parametrized by θ :

$$J(\pi_\theta) = \mathbb{E}_{\tau \sim \pi_\theta} [R(\tau)]$$

with τ being the trajectory of the states, $R(\cdot)$ the reward function, and actions being sampled by the policy π_θ . The current policy is then updated via the gradient step (and the learning rate α_t):

$$\pi_{\theta_{t+1}} = \pi_{\theta_t} + \alpha_t \nabla_{\pi_\theta} J|_{\pi_\theta = \pi_{\theta_t}}$$

Therefore, in order to obtain an optimal policy π , we need to compute the gradient $\nabla_{\pi_\theta} J(\pi_\theta)$. *REINFORCE* leverages the log-derivative trick to derive the policy gradient. The log-derivative trick is based on the differentiation rule for the logarithm, i.e. $\nabla_\theta \pi_\theta(\tau) = \pi_\theta(\tau) \nabla_\theta \log \pi_\theta(\tau)$, and allows pulling the gradient operator inside the expectation.

$$\nabla_{\pi_\theta} J(\pi_\theta) = \nabla_{\pi_\theta} \mathbb{E}_{\tau \sim \pi_\theta} [R(\tau)] = \mathbb{E}_{\tau \sim \pi_\theta} [R(\tau) \nabla_{\pi_\theta} \log \pi_\theta(\tau)]$$

Consequently, we can approximate $\nabla_{\pi_\theta} J(\pi_\theta)$ by Monte Carlo sampling, i.e. sampling the trajectories of the policy and using their cumulative rewards $R(\tau)$ to weight the gradients.

Connection to targeted sequential estimation. Specifically in this thesis, we will rely on different concepts from these settings: the Explore-Exploit strategy from active learning settings (cf. [Singh \(2023\)](#)), the *REINFORCE* algorithm, and a heuristic integrating prior knowledge which can be extended to Bayesian settings (see [Chapter 3](#) and [Chapter 4](#)).

2.3 Causality in Real Data

Ultimately, the goal of causal inference methods is to derive causal effects from real data. Compared to large deep learning models and their versatile applications ([Liu et al., 2023](#)), causal methods seem limited when it comes to deploying them to real world settings ([Guo et al., 2020](#)). Even though we could argue that the goals of both are different—causal models are not used for prediction tasks, but to answer questions “why” and “what if” which are supposedly harder—we need to be aware of the challenges that arise for causality in real world data. Causal methods need careful consideration (of assumptions) and adaptation to the application domain (to meet assumptions) in order to be valuable.

First, real world data is heavily confounded and causal inference algorithms are difficult to evaluate ([Parikh et al., 2022](#)). Usually data sets with known ground truth are not available. While even unsupervised learning methods can be evaluated by metrics such as the homogeneity of the clusters, for causal models, without ground truth, there is no straightforward way to even rank them. We rely on the (argued) causal assumptions to guarantee that the method is learning something meaningful from the data. However, these assumptions are difficult to verify in practice. We briefly reiterate on the IV assumptions as an example: While (A2) *Relevance* is—at least in the linear case—approximated by the F-statistic ([Sanderson & Windmeijer, 2016](#); [Andrews et al., 2019](#)), (A3) *Exclusion Restriction* and (A1) *Unconfoundedness of the Instrument* are nearly impossible to verify. Nevertheless, for biological data, there is the resort of randomized control trials and other carefully designed experiments which allow for a stronger qualitative and quantitative backing of IV assumptions.

This thesis is strongly motivated by the problems that arise when doing cause-effect estimation within biological data such as the microbiome or further genomic datasets. In the following, we explain their structure and the challenges that arise when trying to establish cause-effect relationships within these data sets. Furthermore, we describe the current experimental configurations associated with these frameworks and provide a few examples of already established causal methods within them.

2.3.1 Microbiome Data

Microbiome Data Structure. Bacterial microbiome measurements typically come in the form of counts of operational taxonomic units (OTUs) or amplicon sequencing variants (ASVs) derived from high-throughput sequencing of 16S ribosomal RNA (rRNA) (Johnson et al., 2019b) and are summarized as taxonomic compositions, for example, on the species, genus, or family level.

One way of dealing with this available relative abundance information is to normalize these read counts by their respective totals, resulting in *compositional data*. Compositional data comprises the proportions of some whole, implying that data points live on the unit simplex $\mathbb{S}^{d_x-1} := \{x \in \mathbb{R}_{\geq 0}^{d_x} \mid \sum_{k=1}^{d_x} x_k = 1\}$. To avoid dealing with the peculiarities of compositional data, often the compositional vector x is condensed into a one-dimensional summary statistic. However, we note that summary statistics do not provide causal meaning; cf. Sec. 2.3.3.

Challenges for Cause-Effect Estimation. Much of the microbiome research to date has been observational in nature (Pasolli et al., 2016), with studies comparing the microbial composition of healthy individuals to those with various diseases. These studies have uncovered numerous associations between specific microbial taxa or functions and diseases such as inflammatory bowel disease (IBD) (Morgan et al., 2012; 2015), obesity (Turnbaugh et al., 2006; 2009), diabetes (Qin et al., 2012), and neurological disorders like depression and autism (Liu et al., 2019). However, establishing whether these associations are causal remains a significant challenge (Arnold et al., 2020; Breskin & Murray, 2020).

One major issue are confounding factors that influence both the treatment (gut microbiome) and outcome (such as disease status), leading to biased estimates of the causal effect (Van-Every et al., 2023). For example, diet is a well-known confounder in microbiome studies (Asnicar et al., 2021; Johnson et al., 2019a), as it can directly influence the composition of the gut microbiome while also being associated with disease risk. Other potential confounders include age (Yatsunenkov et al., 2012), medication use, and host genetics (Kolde et al., 2018).

Another significant challenge is reverse causality. It is often unclear whether changes in the microbiome cause disease or whether the disease causes changes in the microbiome. For

example, for inflammatory diseases (IBD), it is difficult to determine whether dysbiosis (an imbalance in the microbial community) contributes to the onset of the disease or if it is a consequence of the disease itself.

Overall, the main challenges for the microbiome example are the high-dimensional and strongly confounded data which make experimental studies essential to establish causality (Fritz et al., 2013).

Experiments in Microbiome Data. Due to the inherent randomness found in nature, we can observe experiments “for free”: Mendelian randomization (MR) is a method that uses genetic variants to identify the causal relationship between a modifiable exposure (in our case the composition of the microbiome) and an outcome. By leveraging this natural randomization, MR can provide stronger evidence of causality than traditional observational studies by using these genetic variants as instrumental variables. For example, in various studies, the causal connection of the microbiome and obesity could be established via MR (Liu et al., 2024).

Another effective experimental strategy in microbiome research is utilizing germ-free animal models (Trexler, 1999). Germ-free mice, which are born and raised in a sterile environment without any microbiota, provide a controlled setting for studying the causal effects of specific microbial taxa or communities. Researchers can colonize these mice with microbial communities from humans or other animals and observe the resulting physiological changes (Schulfer et al., 2019).

In addition to germ-free models, human microbiome transplantation studies, such as fecal microbiota transplantation (FMT), provide a way to experimentally manipulate the gut microbiome in humans. FMT has been shown to be highly effective in treating recurrent *Clostridium difficile* infections (Eiseman et al., 1958) and is being explored as a potential treatment for other conditions, such as inflammatory bowel disease and metabolic disorders (Ley et al., 2005).

Causal Models in Microbiome Data. Apart from the instrumental variable setting (Ailer et al., 2025), there are also other models that have been used to derive causal effects from microbiome data: In fact, several recent works have studied the causal mediation effect of the microbiome on health-related outcomes, assuming that all relevant covariates are observed and can be controlled for (Sohn & Li, 2019; Carter et al., 2020; Wang et al., 2020; Xia, 2021; Sohn et al., 2022; Wang et al., 2023; Zhang et al., 2020a), or vice versa, the effect of environmental factors on the microbiome (Sommer et al., 2022). However, in practice, there is little hope of measuring *all* latent factors in these complex interactions.

2.3.2 Genomic Data

Genomic Data Structure. The availability of large-scale genomic data sets, such as those generated by genome-wide association studies (GWAS), whole-exome sequencing studies (WES) and whole-genome sequencing studies (WGS), has significantly improved our understanding of how genetic variation influences health and disease (Timpson et al., 2018). GWAS test the association of millions of genetic variants, typically single-nucleotide polymorphisms (SNPs), with a given trait (Brookes & Robinson, 2015; Timpson et al., 2018). A variety of statistical downstream analyses can be implemented to identify the genetic variants that cause the disease (Yang et al., 2011; 2012). It is because of those association studies that we have a better understanding of which genetic traits are, for example, correlated with autism and ADHD (Peyre et al., 2021), cardiovascular risk (Sun et al., 2021) or diabetes 1 (Barrett et al., 2009).

Challenges for Cause-Effect Estimation. However, one of the central challenges in genomic research is distinguishing between correlation and causation (Grosz et al., 2020; Hernán, 2018): Many genetic associations discovered through observational studies provide valuable information, but do not necessarily reveal whether genetic variants cause specific diseases or traits (Callaway, 2017; Hu et al., 2018). In observational genomic studies, confounding factors such as population stratification (differences in allele frequencies between populations due to ancestry) and environmental influences can create spurious associations, leading to biased or inaccurate conclusions (Timpson et al., 2018).

Similar to the specific microbiome example, high-dimensionality and confounded data make causal inference methods essential (Pingault et al., 2022).

Experiments in Genomic Data. Also for general genomic data, Mendelian randomization (MR) creates randomness in observed data and allows the use of causal methods (Pingault et al., 2022; Timpson et al., 2018). Genetic variants are randomly inherited at conception and therefore serve as proxies for instrumental variables: For example, a specific variant is responsible for an individual not being able to effectively metabolize alcohol (Richardson et al., 2021). As a consequence, those individuals drink less often or abstain from alcohol completely, as they are less able to clear alcohol from their system. In the Asian population (where it is polymorphic), this has been pivotal in showing the effect of alcohol on health risk factors such as increased risk of stroke (Cho et al., 2015; Millwood et al., 2019) and high blood pressure (Chen et al., 2008).

In recent years, advances in genome editing technologies, particularly CRISPR-Cas9, have made it possible to directly test the effects of specific genetic variants on gene function and disease phenotypes. Targeted gene knockouts and CRISPR-Cas9 gene editing (Fu et al., 2013; Zhang et al., 2015) allow researchers to precisely manipulate the genome at specific loci. By observing the phenotypic effects of this modification, researchers can confirm whether the

variant is truly causal and gain insight into its functional role. However, we note that this intervention is still subject to some error and thus is not a *hard intervention* in the sense of Pearl (2009), cf. (Hsu et al., 2013; Adikusuma et al., 2018; Tsuchida et al., 2023; Lazar et al., 2024).

Causal Models in Genomic Data. Despite and because of these challenges, the field of establishing causality in genomic data is large. Most causal approaches to date have focused on Gene Regulatory Network inference, therefore, learning connections that go from a regulator to a regulated gene (Tam et al., 2013; Belyaeva et al., 2021; Ke et al., 2023; Wen et al., 2023). Recently, perturbation data (generated by MR but more importantly by CRIPR-Cas9) have sparked the hope to use such interventional data with more thorough causal models (Dong et al., 2023; Tejada-Lapuerta et al., 2023).

2.3.3 Compositional Data

We briefly present a method to address the compositionality of the data in microbiome and other genomic studies.

In concrete terms, compositional data comprise the proportions of some whole, implying that the data points live on the unit simplex $\mathbb{S}^{d_x-1} := \{x \in \mathbb{R}_{\geq 0}^{d_x} \mid \sum_{k=1}^{d_x} x_k = 1\}$. Consequently, for instrumental variable settings, the treatment variable X lives on the simplex $\mathbb{S}^{d_x-1}$ instead of $\mathbb{R}^{d_x}$ and the d_x entries of a composition remain dependent via the unit sum constraint.

Although popular books and research articles alike seem to suggest that (bio)diversity, such as α -diversity, e.g. $\alpha_{\text{Simpson}} = -\sum_{k=1}^{d_x} (x_k)^2$ or $\alpha_{\text{Shannon}} = -\sum_{k=1}^{d_x} x_k \log x_k$, is indeed an important *causal driver* of ecosystem functioning and human health (Chapin et al., 2000; Blaser, 2014), two aspects contradict this approach: (1) When considering the proposed causal effect estimand $\mathbb{E}[Y \mid \text{do}(\alpha = \alpha^*)]$ directly, one implicitly assumes the existence of clearly defined interventions on α . However, changing the diversity by a fixed amount $\Delta\alpha$ is highly ambiguous (as the change is achieved at the microbiota level) and different ways of achieving this change must be expected to have different implications on the outcome Y . (2) The definition of α -diversity is not unique, which could lead to a potential search for positive results by using a different metric (Kers & Saccenti, 2022) or to contradictory causal claims. Consequently, it is necessary to consider the compositional vector instead of the summary statistic such as α -diversity to derive causal effects from microbiome (and other genomic) data sets.

To adequately address the compositionality, different invertible log-based transformations have been proposed, for example the additive log ratio (alr), centered log ratio (clr) (Aitchison, 1982), and isometric log ratio (ilr) (Egozcue et al., 2003) transformations. The key advantage of such coordinate transformations is that they allow us to use regular multivariate data analysis methods (typically tailored to Euclidean space) for compositional data. For

example, we can directly fit a linear model $y = \beta_0 + \beta^T \text{ilr}(x) + \epsilon$ on the ilr coordinates via ordinary least squares (OLS) regression.

The log transformations are as follows:

$$\begin{aligned} \text{alr}(x) &:= \left(\log \frac{x_1}{x_{d_x}}, \dots, \log \frac{x_{d_x-1}}{x_{d_x}} \right) = \log(x) \cdot \begin{bmatrix} 1 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 1 \\ -1 & -1 & \cdots & -1 \end{bmatrix} \\ \text{clr}(x) &:= \left(\log \frac{x_1}{g(x)}, \dots, \log \frac{x_{d_x}}{g(x)} \right) = \frac{\log(x)}{D} \cdot \begin{bmatrix} D-1 & -1 & \cdots & -1 \\ -1 & D-1 & \cdots & -1 \\ \vdots & \vdots & \ddots & \vdots \\ -1 & -1 & \cdots & D-1 \end{bmatrix} \\ \text{ilr}(x) &:= \log(x) V_i \end{aligned}$$

with $g(x) := \sqrt[d_x]{x_1 \cdots x_{d_x}}$. For the isometric log-transform a typical choice for V_i^T is the so-called Helmert matrix with the first row removed. Therefore $V_i \in \mathbb{R}^{d_x \times d_x-1}$ with $V_i^T V_i = \mathbb{I}_{d_x-1}$ provides an orthonormal basis of $\mathbb{R}^{d_x-1}$.

While alr is a vector space isomorphism that preserves a one-to-one correspondence between all components except for one, which is chosen as a fixed reference point to reduce the dimensionality (we choose x_{d_x} , but any other component works), it is not an isometry, i.e., it does not preserve distances or scalar products. In contrast, both clr and ilr are isometries, but clr only maps onto a subspace of $\mathbb{R}^{d_x}$, often rendering measure theoretic objects such as distributions degenerate. As an isometry between $\mathbb{S}^{d_x-1}$ and $\mathbb{R}^{d_x-1}$, ilr allows for an orthonormal coordinate representation of compositions. However, it is difficult to assign meaning to the individual components of $\text{ilr}(x)$, which all entangle a different subset of relative abundances in x , which poses challenges for interpretability (Greenacre & Grunsky, 2019). Therefore, alr remains a useful tool in statistical analyses where interpretability is required despite the lack of the isometric property.

Ailer et al. (2025) examine an instrumental variable setting with a compositional treatment and discuss potential methods involving log ratio transformations.

Sequential Underspecified Instrument Selection for Cause-Effect Estimation

3.1 Problem Setting

The goal is to estimate the causal effect of treatments $X \in \mathbb{R}^{d_x}$ on a scalar outcome $Y \in \mathbb{R}$. We assume access to valid instruments $Z \in \mathbb{R}^{d_z}$ and focus on the linear setting:

$$X = Z\alpha + \epsilon_X, \quad Y = X\beta + \epsilon_Y,$$

where $\alpha \in \mathbb{R}^{d_z \times d_x}$ and $\beta \in \mathbb{R}^{d_x}$ represent the linear structural functions, cf. Eq. (2.4). The noise variables ϵ_X, ϵ_Y are typically not independent and we assume $\mathbb{E}[\epsilon_X] = \mathbb{E}[\epsilon_Y] = 0$. The standard instrumental variable assumptions (A1)-(A3) become $\alpha \neq 0$, $Z \perp\!\!\!\perp \{\epsilon_X, \epsilon_Y\}$, and $Z \perp\!\!\!\perp Y \mid \{X, \epsilon_Y\}$. Estimating the causal effect corresponds to estimating β . Naturally, the OLS estimator for β will be strongly biased, but if $d_z \geq d_x$ and a rank condition on the covariance of Z and X are satisfied, β is point-identified and can be estimated consistently by linear two-stage methods cf. Section 2.1.1. Note that traditional linear IV methods assume the access to as many instrument as treatment variables.

We address three main challenges that arise within linear cause-effect estimation in high-dimensional biological datasets.

(C1) Generally, the full identification of β requires just- or overidentification, i.e. $d_z \geq d_x$. This is not feasible in scenarios, e.g. with (tens of) thousands of organisms/genes and much fewer antibiotics/drugs to experiment with. **(C2)** Experimentation is typically slow and expensive, putting natural limits on the number of experiments that can be run. Therefore, even if we can perform experiments sequentially, we might still end up with a pooled dimension of Z that is still smaller than the treatment space. **(C3)** The cost of an individual experiment depends on the number of randomized instruments. There may be limits to how

many instruments can be sensibly combined in one single experiment. For example, too many drugs, gene knockouts, or antibiotics at once may kill or permanently damage the studied organism.

3.2 Method

Based on the just- and overidentified linear IV estimator, we derive an estimator for the underspecified setting via a singular value decomposition (SVD) and the resulting pseudo-inverse. The SVD returns the explicit description of the subspace on which we can identify the causal effect exactly, i.e. the instrumented subspace. Additionally, this explicit knowledge of the instrumented subspace allows us to combine estimators from different and not necessary orthogonal experiments and can be used to describe the overall instrumented subspace over multiple sequential experiments. Similarly, by additionally taking advantage of the linearity, we identify the subspace for which we know nothing about the causal effect. The description of this orthogonal complement of the instrumented subspace builds the groundwork for a selection heuristic that proposes the next (most informative) experiment. Furthermore, an estimator based on [Janzing & Schölkopf \(2018\)](#) and [Janzing \(2019\)](#) enables the algorithm to decide when no further experiments are necessary for full identification, even if $d_z < d_x$, thus serving as a natural stopping criterion.

3.3 Contribution

- We show that in the underspecified linear IV setting, we can estimate the orthogonal projection of the causal effect onto the “instrumented subspace” and construct a $\sqrt{n}$ -consistent, asymptotically normal estimator.
- We combine such estimates from experiments with different instruments to consistently recover the estimate had we randomized all IVs simultaneously (without actually having to perform this experiment).
- We develop an algorithm that sequentially proposes subsets of the available instruments to maximally identify the treatment effect from the combined estimate across all experimental rounds. The algorithm trades off the information gained from multiple instruments with the increasing cost of including them in a single experiment based on a pre-specified similarity between instruments.
- We develop techniques to keep track of which components of β have been identified reliably at each round, upper bound the absolute error in unidentified components, and propose a stopping criterion based purely on observational data ($p(x, y)$) that—when reached—guarantees full identification of β under mild assumptions.

3.4 Future Work

The linear IV setting is still widely used in econometrics and health, since it reliably captures dominant effects even in noisy settings and still attracts attention with novel results recently (Pfister & Peters, 2022; Rothenhäusler et al., 2018), however, the linearity assumption may appear restrictive. Therefore, the theoretical understanding in this work serves as a necessary foundation to extend the method and the notion of the instrumented subspace to nonlinear settings.

Another direction of future work is geared towards its use in real world data: Sequential experiments require the access to special and expensive facilities, making it naturally difficult to evaluate the method. However, an evaluation in such settings will be crucial for further development of the method. Moreover, explicit domain knowledge provides new avenues to refine the current instrument selection process.

3.5 Individual Contribution

NK and JH formed and developed the idea. EA performed the proofs, developed the methodology and conducted the simulation. NK, JH and EA wrote the manuscript.

Targeted Sequential Indirect Experiment Design

4.1 Problem Setting

Previously, we aimed for cause-effect identification in the underspecified linear setting. We introduced a sequential estimation routine that selects the instruments for the subsequent experiment which will maximally cover the instrumented subspace from a predefined set of experiments. Building on these insights, we now extend the sequential estimation to the nonlinear setting. Due to the general non-identifiability of the full causal effect in the nonlinear setting, we solely seek to bound a specific property of the causal effect. Note, however, that for the subsequent experiment, we now aim to *design* the best experiment (by constructing the best instrument sampling policy) rather than relying on heuristic and predefined experiments.

We assume the following data generating process of the nonlinear additive noise setting:

$$X = h(Z, \epsilon_X), \quad Y = f(X) + \epsilon_Y,$$

where we have experimental access to $Z \in \mathcal{Z} \subset \mathbb{R}^{d_z}$ and $X \in \mathcal{X} \subset \mathbb{R}^{d_x}$, $Y \in \mathcal{Y} \subset \mathbb{R}$ and $\mathbb{E}[\epsilon_X] = \mathbb{E}[\epsilon_Y] = 0$. As in typical indirect experiment settings, we assume that Z consists of valid instrumental variables: (A1) $Z \perp\!\!\!\perp U$, i.e., our choice of experiments is not influenced by confounding variables, (A2) $Z \not\perp\!\!\!\perp X$, i.e., we do perturb the distribution of X by experimenting on Z , and (A3) $Z \perp\!\!\!\perp Y \mid X, U$, i.e., Z only affects Y via X and does not have a direct effect on Y . When X is multivariate and f, h is allowed to be nonlinear, f is commonly not identified from observational data. However, we argue that in such settings, one should aim for a more targeted query, that is, one that does not aim to identify f fully but concentrates on the or one relevant aspect $Q[f]$. Our goal is to learn as much as possible

about $Q[f]$ by estimating upper and lower bounds. In each subsequent experiment we update the sampling policy of Z .

4.2 Method

The method builds on a bi-level optimization problem which, (1) requires to estimate the bounds $Q_t^\pm$ based on one experiment and then (2) updates the current experiment sampling policy over the instruments $\pi_t(Z)$ to tighten the bounds over $t \in [T]$, for T the total number of experiments.

Estimation of Bounds. The algorithm is initialized by the first experiment with the sampling policy $Z \sim \pi_1$. For the estimation of the bounds we use standard IV minimax methods (Dikkala et al., 2020; Bennett et al., 2022). Generally, minimax estimation is designed to satisfy the moment conditions inherent in the IV framework, concurrently incorporating robustness through an adversarial test function $g \in \mathcal{G}$. Note that due to the nonlinearity, the solution is not unique. We integrate the functional Q into the objective of the minimax method to favor solutions that have the lowest resp. highest values for the functional $Q[f]$. To facilitate estimation, we show that if we assume that f lies within a reproducing kernel Hilbert space (RKHS), there are queries $Q[f]$ for which we can solve the minimax estimation in closed form.

Policy Update. For the policy update of π_t , we use the REINFORCE algorithm with horizon 1: An intuitive and effective method for updating policies is to employ gradient descent on the difference between the upper and lower bound $\delta(\pi_t) = Q_t^+ - Q_t^-$. This approach leverages the gradient to methodically refine policy decisions. By using the log-derivative trick and splitting the samples in random batches (to generate variance within the experiment), we compute the update step.

4.3 Contribution

- We formalize the problem of indirect adaptive experiments in nonlinear, multivariate, confounded settings as a sequential underspecified instrumental variable estimation, where we seek to adaptively tighten estimated bounds on the target query.
- We derive closed form expressions for the bound estimates when assuming f to lie within a reproducing kernel Hilbert space (RKHS) for linear Q .
- We develop adaptive strategies for the outer optimization to iteratively tighten these bounds and demonstrate empirically that our method robustly selects informative experiments leading to identification of $Q[f]$ (collapsing bounds) when it can be identified via allowed experimentation.

4.4 Future Work

We propose two avenues for further work: More theoretical contributions concentrate on the method's applicability by addressing conditions which will provide valid bounds at all steps and by relaxing IV assumptions and providing sensitivity analysis for their violation. Moreover, we currently assume learning a static function, which sets the foundation to extend the framework to dynamically evolving systems. In addition, we can further optimize the method by integrating more advanced techniques for variance reduction and optimal balancing of exploration and exploitation.

Secondly, and more importantly, a special emphasis should be put on real world evaluations; although challenging due to access and resource needs, they offer a valuable opportunity to refine and validate our approach. In addition, by explicitly integrating knowledge about the specific data setting, we can propose even more targeted experiments.

4.5 Individual Contribution

NK, JH and EA formed and developed the idea. EA performed the proofs, developed the methodology and conducted the simulation. NK, JH and EA wrote the manuscript.

5

Conclusion

Summary. In this dissertation, we put instrumental variables in a larger context and extend their scope beyond social and economic pseudo-experiments and randomized control trials.

We identified various challenges at the intersection of instrumental variables methods in the context of experiments and biological data sets, which are—among others—the high-dimensionality of the datasets, the nonlinear effects, and the limited experimental resources to explore causal relationships. Motivated by these challenges, we identified two subsequent settings with the common goal of contributing to answer causal questions in sequential experiment settings with the support of IV methods.

First, we note that to increase the impact of IVs in general experimental design, a sequential estimator is essential. Therefore, we extended the linear setting towards a sequential estimation. We developed a method that can compute an estimate for the underspecified linear setting, i.e. $d_z < d_x$. Although we showed that we can only learn that part of the causal effect that is projected onto the instrumented subspace, we also identified an explicit representation of this instrumented subspace, enabling us to combine estimates of different experimental runs from different data and instruments. Based on the current combined estimate, we selected the next experiment from a predefined set and, given some heuristics, this subsequent experiment provided us with targeted information about the (yet) unknown parts of the causal effect. Moreover, we derived a stopping criterion, informing us when the full causal effect is identified, even if only a part of the treatment space was covered by the instruments.

In the second part, we addressed two key constraints: we wanted to extend the setting from linear to nonlinear, and from predefined experiment selection to structured experiment design. Although the shift to nonlinear settings meant that complete effect identification was not feasible, it underscored an important scientific insight: queries do not always necessitate full effect determination. In many cases, possessing specific knowledge about an

aspect of the causal effect is sufficient to answer research questions. Thus, we concentrated on estimating bounds for a specific aspect of the causal effect. A key component of this approach was the adaptive strategy which leveraged the bounds to inform the gradient updates to get a subsequent experiment that would then produce even tighter bounds. Thus, experiment-by-experiment, we devised the most informative instrument sampling policy, effectively reducing the gap between the bounds.

Limitations. Although the promise of causality through sequential IV experiments makes a compelling case, these methods come with two main constraints.

First of all, their evaluation on real world data is essential; especially when the goal is to integrate them into applications involving decision-making in medicine. Unfortunately, ground truth datasets for cause-effect estimation are hard to come by, notably when they should involve high-dimensional data. The scarcity of real data is particularly critical when it comes to the evaluation of sequential experimentation: the generation of real world datasets would additionally require access (and full control) over an actual experimental setup as well as the time and resources required to conduct these experiments.

Secondly, IV methods require the verification of strong assumptions. Although randomization of the instrument can mitigate the issue, most IV assumptions cannot be tested in practice. Nevertheless, fulfilling these conditions is crucial, as their violation can lead to significant errors in estimates. The possibility of such deviations highlights the need for careful consideration and scrutiny of the assumptions in any experimental implication.

Outlook. In light of these limitations, there are various possibilities to further the current methods. One way to address the evaluation of IV assumption is to use stress test methods for performing sensitivity analysis. In addition, it will be useful to develop methods that allow for some violation of the assumptions or that concentrate on bounding the causal effect. Secondly, the adaption of those methods to other peculiarities of real world data (scarce effect, missing data) will help to make them more valuable for real world datasets.

Overall, the use of instrumental variable methods in sequential experimental design represents a promising avenue for theoretical research and practical evaluation. Especially in biological settings, IV methods can offer a novel perspective on answering causal questions. Moreover, although real world evaluations are expensive, it is these very methods that will decrease the costs of experimentation as they provide more targeted and more efficient estimators.

A

Appendix

A.1 Sequential Underspecified Instrument Selection for Cause-Effect Estimation

Sequential Underspecified Instrument Selection for Cause-Effect Estimation

Elisabeth Ailer^{1,2} Jason Hartford^{3,4} Niki Kilbertus^{1,2,5}

Abstract

Instrumental variable (IV) methods are used to estimate causal effects in settings with unobserved confounding, where we cannot directly experiment on the treatment variable. Instruments are variables which only affect the outcome indirectly via the treatment variable(s). Most IV applications focus on low-dimensional treatments and crucially require at least as many instruments as treatments. This assumption is restrictive: in the natural sciences we often seek to infer causal effects of high-dimensional treatments (e.g., the effect of gene expressions or microbiota on health and disease), but can only run few experiments with a limited number of instruments (e.g., drugs or antibiotics). In such underspecified problems, the full treatment effect is not identifiable in a single experiment even in the linear case. We show that one can still reliably recover the projection of the treatment effect onto the instrumented subspace and develop techniques to consistently combine such partial estimates from different sets of instruments. We then leverage our combined estimators in an algorithm that iteratively proposes the most informative instruments at each round of experimentation to maximize the overall information about the full causal effect.

1. Introduction

Motivation. Understanding cause-effect relationships in high-dimensional systems is a common challenge in various scientific areas. For example, how does our gut microbiome

(treatment X) causally influence health and disease (outcome Y)? How does the transcriptome of a cell (treatment X) causally influence its function (outcome Y)? Typically, the high-dimensional treatment and the outcome are heavily confounded via unknown mechanisms, rendering the strong assumptions required for cause-effect estimation, i.e., computing $p(y | do(x))$, from observational data sampled from $p(x, y)$ untenable. Hence, experimentation is indispensable for the ultimate goal of identifying and estimating these causal relationships. Whenever we can intervene directly on the treatment, we have direct access to $p(y | do(x))$ and cause-effect relationships can be captured by mere association in experiments. We are motivated by two crucial realizations: (a) oftentimes practically feasible experiments do not intervene directly on the treatment X but some other variable Z ; (b) still, the scientific goal is to estimate the effect of the high-dimensional X on the outcome Y (instead of the effect of the actual intervention on Z).

Regarding (a), administering certain antibiotics has a strong and highly predictable effect on the gut microbiome, but it does not break causal links from potentially unobserved confounders to X ; similarly, applying various drug (dosages) to cell cultures influences the transcriptome, but again does not directly intervene on it. Even targeted gene knockouts or CRISPR/Cas gene editing (Zhang et al., 2015; Fu et al., 2013) do not constitute interventions as defined in causal inference (Pearl, 2009). They do not strip the microbiome or transcriptome free of any other causal influences, i.e., they may still be confounded with the outcome of interest. As for (b), ultimately we give a certain antibiotic to learn about *how the microbiome causally influences disease* ($p(y | do(x))$) and not merely about what overall effect the antibiotic has on disease ($p(y | do(z))$).

In such settings, the variable experimented on (antibiotics, drugs, gene knockout/edits) can at most serve as an *instrument* Z for the treatment: **(A1)** Z (strongly) affects the treatment ($Z \not\perp\!\!\!\perp X$). **(A2)** Z is independent of unobserved confounders U ($Z \perp\!\!\!\perp U$). **(A3)** Z “only affects the outcome via the treatment” ($Y \perp\!\!\!\perp Z | \{X, U\}$). Whether we can ascertain these conditions depends on the setting. For example, orally administered sub therapeutic dosages of antibiotics (such that no antibiotics can be detected in the bloodstream) may satisfy this condition (Ailer et al., 2021). Similarly,

¹HelmholtzAI, Helmholtz Munich, Munich, Germany ²School of Computation, Information and Technology, Technical University Munich, Munich, Germany ³MILA - Quebec Artificial Intelligence Institute, Montreal, Quebec, Canada ⁴Recursion, Montreal Quebec Canada ⁵Munich Center for Machine Learning, Munich, Germany. Correspondence to: Elisabeth Ailer <elisabeth.ailer@helmholtz-munich.de>.

genetic interventions (or drugs) may indeed only influence cell function via the transcriptome.

As a result, even with access to experimentation, we must resort to instrumental variable (IV) techniques to estimate the treatment effect. A number of challenges arise in this setting even in the fully linear case: **(B1)** At least as many instruments as treatment variables are required for the causal effect to be identified (just- or overspecified setting). This is not feasible in our scenarios with (tens of) thousands of organisms/genes and much fewer antibiotics/drugs to experiment with. **(B2)** Experimentation is typically slow and expensive, putting natural bounds on how many experiments can be run. **(B3)** The cost of an individual experiment depends on the number of randomized instruments and there may be limits to how many instruments can be sensibly combined in one single experiment. For example, too many drugs, gene knockouts, or antibiotics at once may kill or permanently damage the studied organism.

Main goal and contributions. Constrained by (B1)-(B3), we formulate our main goal: *Can we select a bounded number of instruments in a bounded number of sequential experiments and combine the results for an informative estimate of a high-dimensional causal effect?*¹ In answering this question, we make the following contributions:

- We show that in the underspecified linear IV setting, we can estimate the orthogonal projection of the causal effect onto the “instrumented subspace” and construct a $\sqrt{n}$ -consistent, asymptotically normal estimator.
- We combine such estimates from experiments with different instruments to consistently recover the estimate had we randomized all IVs simultaneously (without actually having to apply all perturbations at once).
- We develop an algorithm that sequentially proposes subsets of the available instruments to maximally identify the treatment effect from the combined estimate across all experimental rounds. The algorithm trades off the information gained from multiple instruments with the increasing cost of including them in a single experiment based on a pre-specified similarity between instruments.
- We develop techniques to keep track of which components of β have been identified reliably at each round, upper bound the absolute error in unidentified components, and propose a stopping criterion based purely on observational data $(p(x, y))$ that—when reached—guarantees full identification of β under mild assumptions.

Related work. We build on the theory of (linear) instrumental variable estimators, with a specific focus on two-stage methods (Angrist & Pischke, 2008). Instrumental variables

have been used since 1928 by Philip G. Wright (Wright, 1928; Stock & Trebbi, 2003) and are an essential part of the econometrics toolkit (Angrist & Pischke, 2008). They have also received renewed interest from the machine learning community recently (Hartford et al., 2017; Zhang et al., 2020; Kilbertus et al., 2020; Singh et al., 2019; Bennett et al., 2019; Padh et al., 2022; Saengkyongam et al., 2022; Muandet et al., 2019). Despite the difficulty of finding valid instruments for a given target effect in practice (Hernán & Robins, 2006), instrumental variable estimation has been applied successfully in genetics via Mendelian randomization (Sanderson et al., 2022; Didelez & Sheehan, 2007) and recently on microbiome data (Sohn & Li, 2019; Wang et al., 2020; Ailer et al., 2021).

Our work is also related to ideas in experiment design for causal structure learning (Hyttinen et al., 2013; Gamella & Heinze-Deml, 2020; Sussex et al., 2021; Tigas et al., 2022). Two key differences are that those works focus on sequential selection of *interventions* (not just instruments) and that they seek to identify causal structure (i.e., the causal graph) instead of a specific causal effect. To the best of our knowledge, there is no literature on adaptively selecting instruments in sequential experiments to identify a causal effect from high-dimensional treatments.

Finally, variants of underspecified instrumental variable settings have been studied by Pfister & Peters (2022). They assume the causal effect from X on Y to be sparse, which allows them to relax standard identifiability assumptions in the linear IV setting. Rothenhäusler et al. (2018) make use of exogenous variables to provide an estimator that interpolates between the ordinary least-squares (OLS) and two stage least-squares (2SLS) estimates, even when IV assumptions are not fully satisfied and point-identification is not guaranteed. Their proposal can be interpreted as “choosing the best performing (in terms of mean squared error) estimate among the compatible ones”. Thus, both works pursue goals orthogonal to ours.

2. Background and Problem Setting

We aim to estimate the causal effect of treatments $X \in \mathbb{R}^{d_x}$ on a scalar outcome $Y \in \mathbb{R}$. There may be unobserved confounding between X and Y , but we assume access to valid instruments $Z \in \mathbb{R}^{d_z}$. We focus on the linear setting

$$X = Z\alpha + \epsilon_X, \quad Y = X\beta + \epsilon_Y, \quad (1)$$

where $\alpha \in \mathbb{R}^{d_z \times d_x}$ and $\beta \in \mathbb{R}^{d_x}$ represent the linear structural functions. We interpret β flexibly as a row or column vector as needed. Because of unobserved confounding, the noise variables ϵ_X, ϵ_Y are typically not independent, but we assume $\mathbb{E}[\epsilon_X] = \mathbb{E}[\epsilon_Y] = 0$. The standard instrumental variable assumptions (A1)-(A3) become $\alpha \neq 0$, $Z \perp\!\!\!\perp \{\epsilon_X, \epsilon_Y\}$, and $Z \perp\!\!\!\perp Y \mid \{X, \epsilon_Y\}$.

¹We highlight that while we call the high-dimensional X the “treatment”, we experiment (or intervene) on lower-dimensional instruments Z .

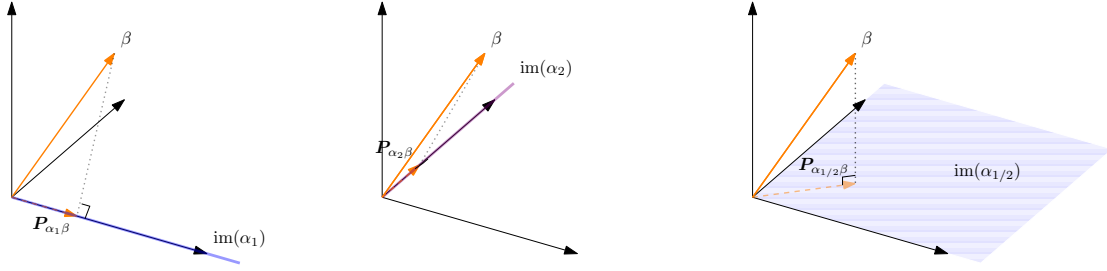


Figure 1. The illustration shows the possible experimental settings, including their estimates in a setting with $d_z = 2, d_x = 3$: (1) The experiments with one of the two instruments respectively, denoted by α_1 and α_2 and (2) - if feasible - with both instruments at once, i.e. $\alpha_{1/2}$. In each figure, the true causal effect β (orange) is projected onto the corresponding instrumented spaces, i.e. $\text{im}(\alpha_1)$, $\text{im}(\alpha_2)$ and $\text{im}(\alpha_{1/2})$ (shaded blue) by $P_{\alpha_1}\beta$, $P_{\alpha_2}\beta$ and $P_{\alpha_{1/2}}\beta$ (dashed orange).

Therefore, estimating the causal effect simply corresponds to estimating β . The OLS estimator for β will be biased, but if $d_x \leq d_z$ and a rank condition on the covariance of Z and X is satisfied, β is point-identified and can be estimated consistently, for example via the standard 2SLS estimator (Angrist & Pischke, 2008). When collecting n i.i.d. observations in matrices $\mathbf{X} \in \mathbb{R}^{n \times d_x}$, $\mathbf{Z} \in \mathbb{R}^{n \times d_z}$, and $\mathbf{y} \in \mathbb{R}^n$, the 2SLS estimator for β is given by

$$\hat{\beta}_{2\text{SLS}} = (\mathbf{X}^T \mathbf{P}_Z \mathbf{X})^{-1} \mathbf{X}^T \mathbf{P}_Z \mathbf{y}, \quad (2)$$

where $\mathbf{P}_Z = \mathbf{Z}(\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T$ is a projection matrix. The 2SLS estimator can be viewed as (a) regressing X on Z , and (b) regressing Y on the predicted X values from the first-stage regression. Intuitively, the influence of ϵ_X on X is “regressed out” in the first stage, leaving only the direct causal effect of X on Y in the second stage.

Following traditional application settings, most proposed IV methods assume the *just-identified* setting $d_x = d_z$ (as well as the rank condition on $\text{cov}(X, Z)$, which translates into a full rank condition on α). Typically, these methods can also accommodate the *overidentified* (or *overspecified*) setting $d_z > d_x$. Having many instruments available in the overspecified setting can also be exploited for consistent estimators under relaxed assumptions such as correlated or weak instruments (Hahn & Hausman, 2005; Kang et al., 2016). In contrast, the *underidentified* (or *underspecified*) case $d_z < d_x$, where there are fewer instruments than treatments, has received little attention in the literature. In this case, since $\text{rk}(\mathbf{P}_Z) \leq d_z$, we have that $\text{rk}(\mathbf{X}^T \mathbf{P}_Z \mathbf{X}) \leq d_z < d_x$ and consequently cannot take the inverse in Equation (2). More generally, in this situation the causal effect β is not fully identified. However, the available instruments still constrain the possible values of β . We will later estimate and exploit those confines.

In our setting, we assume access to a fixed number of $N_{\text{IV}} \in \mathbb{N}$ instruments (e.g., available antibiotics, or drugs) and—in high-dimensional treatment settings—typically have $N_{\text{IV}} <$

d_x . Because we will also consider subsets of instruments, we generically denote by $d_z \in [N_{\text{IV}}] := \{1, \dots, N_{\text{IV}}\}$ the number of IVs in a given estimation (or experiment). Further, we assume experimental access to the instruments in that we can run experiments in which a chosen set of instruments is randomized (e.g., in mouse studies or on cell cultures). Since the specific distribution of Z does not affect identifiability, we assume without loss of generality that the components of Z all independently follow a Rademacher distribution. This can be interpreted as whether a drug or antibiotic is applied.² With this choice, α fully characterizes the effect of Z on X and we consequently also call the rows of $\alpha \in \mathbb{R}^{N_{\text{IV}} \times d_x}$ “the instruments”.³ Since the components of Z are jointly independent, the full rank condition on $\text{cov}(X, Z)$ translates to α having full rank.

To maximally constrain the effect estimate, one would randomize all N_{IV} available instruments simultaneously. Due to (B3), this is typically infeasible in practice. We model this via a cost function, which can also incorporate hard constraints such as limiting the number of instruments per experiment to at most $N_{\text{IV}/\text{exp}} < N_{\text{IV}}$. For (B2) we limit the total number of possible experiments to $T \in \mathbb{N}$.

We proceed as follows. In Section 3, we first construct a consistent, asymptotically normal estimator for the orthogonal projection of β onto the image of α viewed as a linear map $\alpha : \mathbb{R}^{d_z} \rightarrow \mathbb{R}^{d_x}$. We then establish a method to combine multiple such estimates obtained from (different) subsets of IVs to obtain a consistent estimate for the orthogonal projection of β on the linear subspace spanned by all instruments combined. Furthermore, we introduce a method to determine which components of β have been successfully identified and a condition that guarantees full identification

²Other illustrative choices are $\frac{1}{2}$ -Bernoulli or the uniform distribution on a finite interval representing (standardized) dosage levels. All our results immediately transfer.

³Generically, $\alpha \in \mathbb{R}^{d_z \times d_x}$ with $d_z \leq N_{\text{IV}}$ represents some choice of d_z instruments in a given estimation/experiment.

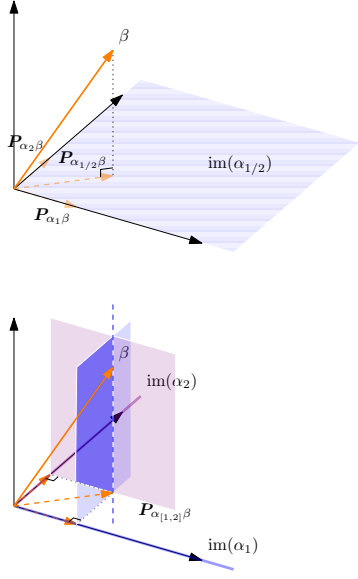


Figure 2. The illustration shows a setting with $d_z = 2, d_x = 3$. It contrasts the single estimation steps with the combination step. Upper panel, single estimation: the true causal effect β (orange) is projected onto the individual instrumented spaces by $P_{\alpha_1}\beta$, $P_{\alpha_2}\beta$ and $P_{\alpha_{1/2}}\beta$. Lower panel, combination step: The intersection (dashed blue line) of the orthogonal complements of $\{\text{im}(\alpha_1), \text{im}(\alpha_2)\}$ represents all vectors that are compatible with both individual estimates. Equation (4) then selects the minimum norm point on that line. In the illustration it comes to show that the combination of individual estimators $P_{\alpha_{[1,2]}}\beta$ recovers the effect of the single estimation with both instruments $P_{\alpha_{1/2}}\beta$.

of β under mild assumptions. In Section 4, we then leverage all these findings to develop a procedure that proposes subsets of instruments for sequential experimentation to maximally identify β under the given constraints.

Figure 1 and Figure 2 illustrate the estimation resp. the combination step and the role of linearity in this context visually. We suggest to consult these illustrations for intuition when reading the formal statements in the following section.

3. Underidentified IV Estimates

3.1. Estimator for a Single Experiment

Consider a set of instruments $\alpha \in \mathbb{R}^{d_z \times d_x}$ with $d_z < d_x$ and $\text{rk}(\alpha) = d_z$. The predicted treatment values via OLS $\hat{X} = Z(Z^T Z)^{-1} Z^T X = P_Z X \in \mathbb{R}^{n \times d_x}$ are confined to a proper linear subspace of $\mathbb{R}^{d_x}$, rendering $X^T P_Z X$ singular and $\hat{\beta}_{2\text{SLS}}$ in Equation (2) ill-defined in the underspecified setting. With $\hat{\alpha}$ being the first-stage OLS estimate, we have $\hat{X}_i = P_{\hat{\alpha}} X_i \in \text{im}(\hat{\alpha}) := \{\hat{\alpha}z \mid z \in \mathbb{R}^{d_z}\} \subset \mathbb{R}^{d_x}$ for all $i \in [n]$, where $P_{\hat{\alpha}} = \hat{\alpha}^T (\hat{\alpha} \hat{\alpha}^T)^{-1} \hat{\alpha} \in \mathbb{R}^{d_x \times d_x}$ denotes the orthogonal projection onto $\text{im}(\hat{\alpha})$. That is, all first-stage predictions lie in the image of $\hat{\alpha}$ and we can consider

the second stage as a mapping $\mathbb{R}^{d_x} \supset \text{im}(\hat{\alpha}) \rightarrow \mathbb{R}$. We call $\text{im}(\hat{\alpha})$ the *instrumented subspace* (in this experiment). Intuitively, while the low-dimensional Z cannot “shake” or “instrument” the entire treatment space to identify β , it still “instruments” a non-trivial subspace, inducing non-trivial constraints on β . In particular, one may expect to recover “the part of β within the instrumented subspace”. The following statement formalizes this intuition.

Proposition 1. Let $\alpha \in \mathbb{R}^{d_z \times d_x}$ have full rank and assume we have i.i.d. data Z, X, y from an experiment in which all d_z instruments have been randomized. Then

$$\widehat{P_{\alpha}}\beta := (X^T P_Z X)^+ X^T P_Z y \xrightarrow{d} \mathcal{N}(P_{\alpha}\beta, \Sigma) \quad (3)$$

with $\Sigma := \frac{1}{n} \alpha^+ \Sigma_Z^{-1} (\alpha^T)^+ \text{Var}[\epsilon_Y] = \frac{1}{n} (\alpha^T \alpha)^+ \text{Var}[\epsilon_Y]$,

where $(\cdot)^+$ denotes the Moore-Penrose pseudoinverse, Σ_Z is the covariance matrix of Z , and the simple form of Σ (in gray) applies when Z are i.i.d. Rademacher variables.

Proof. For $d_z \geq d_x$, the pseudoinverse in Equation (3) can be replaced with a regular matrix inverse and the result follows from the asymptotic normality of $\hat{\beta}_{2\text{SLS}}$ in Equation (2) (Angrist & Pischke, 2008, Sec. 4.2.1). For the underspecified case $d_z < d_x$, we start with the singular value decomposition (SVD) $\hat{X} = U D V^T$, with which $\hat{X}^T \hat{X} = V D^T D V^T$. Assuming that the singular values in the rectangular diagonal matrix $D \in \mathbb{R}^{n \times d_x}$ are sorted in non-ascending order, only the first d_z entries are nonzero, because $\text{rk}(\hat{X}) = d_z$. Accordingly, the first d_z rows of V form an orthonormal basis of $\text{im}(\hat{\alpha})$. We write $V_{\hat{\alpha}} \in \mathbb{R}^{d_x \times d_z}$ for the first d_z columns of V , $D_{\hat{\alpha}} \in \mathbb{R}^{d_z \times d_z}$ for the upper left $d_z \times d_z$ block of D , and $U_{\hat{\alpha}} \in \mathbb{R}^{n \times d_z}$ for the first d_z columns of U . With $\hat{X} = P_Z X$ we compute

$$\begin{aligned} \widehat{P_{\alpha}}\beta &= (X^T P_Z^T P_Z X)^+ X^T P_Z^T P_Z (X\beta + \epsilon_Y) \\ &= (\hat{X}^T \hat{X})^+ (\hat{X}^T \hat{X}\beta + \hat{X}^T \epsilon_Y) \\ &= V (D^T D)^+ D^T D V^T \beta + V (D^T)^+ U^T \epsilon_Y \\ &= V_{\hat{\alpha}} V_{\hat{\alpha}}^T \beta + V_{\hat{\alpha}} D_{\hat{\alpha}}^{-1} U_{\hat{\alpha}}^T \epsilon_Y, \end{aligned}$$

where we have used

$$(D^T D)^+ D^T D = \begin{pmatrix} I_{d_z \times d_z} & \mathbf{0}_{d_z \times (d_x - d_z)} \\ \mathbf{0}_{(d_x - d_z) \times d_z} & \mathbf{0}_{(d_x - d_z) \times (d_x - d_z)} \end{pmatrix}.$$

From the asymptotic normality of the first-stage OLS estimate $\hat{\alpha}$, it follows that $V_{\hat{\alpha}}$ and $D_{\hat{\alpha}}$ are $\sqrt{n}$ -consistent estimators for V_{α} and D_{α} respectively (Bura & Pfeiffer, 2008). Therefore, $P_{\hat{\alpha}} = V_{\hat{\alpha}} V_{\hat{\alpha}}^T$ converges in probability to P_{α} in the large sample limit. With $\mathbb{E}[\epsilon_Y] = 0$ we thus have

$$\mathbb{E}[\widehat{P_{\alpha}}\beta] = \mathbb{E}[V_{\hat{\alpha}} V_{\hat{\alpha}}^T] \beta \xrightarrow{p} P_{\alpha} \beta.$$

Similarly, with $\hat{X} = Z \hat{\alpha}$ we compute the covariance

$$\begin{aligned} \text{Cov}[\widehat{P_{\alpha}}\beta] &= \text{Cov}[V_{\hat{\alpha}} D_{\hat{\alpha}}^{-1} U_{\hat{\alpha}}^T \epsilon_Y] = \hat{X}^+ (\hat{X}^T)^+ \text{Var}[\epsilon_Y] \\ &= \hat{\alpha}^+ (Z^T Z)^{-1} (\hat{\alpha}^T)^+ \text{Var}[\epsilon_Y]. \end{aligned}$$

Since the matrix inverse is continuous (on invertible matrices), the continuous mapping theorem yields that $\widehat{\Sigma}_Z^{-1}$ consistently estimates Σ_Z^{-1} . Again using the asymptotic normality of $\hat{\alpha}$, we thus have (for centered) $\text{Cov}[\widehat{\mathbf{P}}_{\alpha}\beta] \xrightarrow{p} \frac{1}{n}\alpha^+\Sigma_Z^{-1}(\alpha^T)^+ \text{Var}[\epsilon_Y]$. For i.i.d. Rademacher instruments with covariance matrix $\Sigma_Z = \mathbf{I}_{d_z \times d_z}$ we have $\text{Cov}[\widehat{\mathbf{P}}_{\alpha}\beta] \xrightarrow{p} \frac{1}{n}(\alpha^T\alpha)^+$. $\square$

Proposition 1 substantiates the first intuition that in the underidentified setting we can still consistently estimate the orthogonal projection of β onto the instrumented subspace.

We refer to Figure 1 for an illustration of the estimation in a linear setting with $d_z = 2$, $d_x = 3$, $T = 2$. For visualization purposes, we choose axis-aligned $\alpha_1 = (1, 0, 0)$ and $\alpha_2 = (0, 1, 0)$. For both instruments, we estimate the projection of β onto the respective instrumented spaces ($\mathbf{P}_{\alpha_1}\beta$ and $\mathbf{P}_{\alpha_2}\beta$). The illustration on the right hand side also includes the estimation which uses both instruments ($\alpha_{1/2}$) at once. In the following, this is the effect we want to recover solely based on the individual estimates.

3.2. Combined Estimator

In sequential experiments, we select pairwise disjoint instruments $\alpha_1, \alpha_2, \dots, \alpha_T$, where each α_i is a subset of the N_{IV} available ones. Each corresponding individual estimator $\widehat{\mathbf{P}}_{\alpha_1}\beta, \widehat{\mathbf{P}}_{\alpha_2}\beta, \dots, \widehat{\mathbf{P}}_{\alpha_T}\beta$ estimates an orthogonal projection of β onto the linear subspace $\text{im}(\hat{\alpha}_i)$. Our goal is to reconstruct the estimate we would have obtained, had we used all instruments in the sequential experiments at once in a single one. Denoting the union (or concatenation) of all available instruments by a single matrix $\alpha_{[T]}$, we aim to reconstruct $\mathbf{P}_{\alpha_{[T]}}\beta$ from the individual $\widehat{\mathbf{P}}_{\alpha_i}\beta$.⁴ In other words, we are looking for the least-square vector compatible with all projections

$$\min_{\gamma \in \mathbb{R}^{d_x}} \|\gamma\|_2^2 \quad \text{s.t.} \quad \widehat{\mathbf{P}}_{\alpha_i}\beta = \mathbf{V}_{\hat{\alpha}_i} \mathbf{V}_{\hat{\alpha}_i}^T \gamma \quad \text{for all } i \in [T]. \quad (4)$$

Proposition 2. *Let $\alpha \in \mathbb{R}^{N_{IV} \times d_x}$ have full rank and let $(\mathbf{Z}_1, \mathbf{X}_1, \mathbf{y}_1), \dots, (\mathbf{Z}_T, \mathbf{X}_T, \mathbf{y}_T)$ be i.i.d. datasets from T experiments with disjoint subsets $\alpha_1, \dots, \alpha_T$ of randomized instruments. Then the solution of Equation (4) is a consistent estimator for $\mathbf{P}_{\alpha_{[T]}}\beta$.*

Proof. Let $\mathbf{A} := \mathbf{V}_{\hat{\alpha}_{[T]}} \mathbf{V}_{\hat{\alpha}_{[T]}}^T = (\mathbf{V}_{\hat{\alpha}_1} \mathbf{V}_{\hat{\alpha}_1}^T | \dots | \mathbf{V}_{\hat{\alpha}_T} \mathbf{V}_{\hat{\alpha}_T}^T) \in \mathbb{R}^{(Td_x) \times d_x}$ and $\mathbf{b} := (\widehat{\mathbf{P}}_{\alpha_1}\beta, \dots, \widehat{\mathbf{P}}_{\alpha_T}\beta) \in \mathbb{R}^{Td_x}$, where “ $|$ ” denotes concatenation. Then the solution to Equation (4) is simply $\mathbf{A}^+\mathbf{b}$, i.e., the least-squares solution to $\mathbf{A}\gamma = \mathbf{b}$. Hence the optimal γ is the orthogonal projection of $\mathbf{b}$ onto $\text{im}(\mathbf{A})$, which by construction is just $\text{im}(\hat{\alpha}_{[T]})$. By the

⁴We use similar notation $\hat{\alpha}_{[t]}$ for the concatenation of the instruments (as matrices) used up until and including round t .

consistency of $\hat{\alpha}_i$ (and thus the consistency of $\mathbf{A}$ and $\mathbf{b}$ for their respective expressions without hats as in the proof of Proposition 1), asymptotically $\mathbf{A}^+\mathbf{b} \xrightarrow{p} \mathbf{P}_{\alpha_{[T]}}\beta$. $\square$

Note that the instrument sets need not be disjoint for the proof of Proposition 2. However, since we do not gain further information by including the same instrument in two distinct experiments, which would instrument the same subspace twice, it is reasonable to keep instrument sets distinct across experiments for efficiency. The least-squares problem with linear equality constraints in Equation (4) can be solved efficiently (in high dimensions and for many constraints) by simply solving a linear system (Nocedal & Wright, 2006, Sec. 16.1). In practice, we keep a running estimate $\widehat{\mathbf{P}}_{\alpha_{[t]}}\beta$, which is the combined estimate for all experiments up to and including round $t \in [T]$. Being able to combine estimates from individual sets of instruments “as if we had run a single experiment using the union of instruments at once” forms the basis for leveraging sequential experiments with different sets of randomized instruments to optimally constrain the overall causal effect in Section 4.

3.3. Full Identification and Identified Components

Full identification. As an interesting special case of Proposition 1, we note that regardless of d_x , a single instrument $d_z = 1$ in principle suffices to identify $\beta \in \mathbb{R}^{d_x}$, namely when $\alpha \in \mathbb{R}^{d_x \times 1}$ is parallel to β (viewed as vectors in $\mathbb{R}^{d_x}$), which implies $\mathbf{P}_{\alpha}\beta = \beta$. Of course, this is unlikely to “happen accidentally” in practice. However, it highlights that if we had full control over the instruments, we could fully identify β by crafting $\alpha \in \mathbb{R}^{d_x \times 1}$ such as to maximize the 2-norm of the resulting estimated $\widehat{\mathbf{P}}_{\alpha}\beta$. This can be seen from the fact that $\|\mathbf{P}_{\alpha}\beta\|_2 \leq \|\beta\|_2$ for any α with equality implying $\beta \in \text{im}(\alpha)$. Therefore, when running sequential experiments, the 2-norm of the combined estimate is a proxy for whether we are still gaining relevant information. We summarize these results in the following Corollary.

Corollary 3. *In the setting of Proposition 1, the following are equivalent: (i) $\beta \in \text{im}(\alpha)$; (ii) $\|\widehat{\mathbf{P}}_{\alpha}\beta\|_2 \rightarrow \|\beta\|_2$ as $n \rightarrow \infty$; (iii) $\widehat{\mathbf{P}}_{\alpha}\beta$ consistently estimates β (not just $\mathbf{P}_{\alpha}\beta$).*

In our underspecified setting, it is generally impossible to determine whether any (and therefore all) of the statements in Corollary 3 hold true. However, recently Janzing & Schölkopf (2018) have shown that under mild assumptions on the confounding model (i.e., on the joint distribution of ϵ_X, ϵ_Y), one can estimate what they call the *confounding strength*. Their estimator has been further refined by Rendsburg et al. (2022). The confounding strength can then be leveraged to estimate $\|\beta\|_2$ itself from purely observational data (Janzing, 2019). This means that we can consistently estimate $\|\beta\|_2$ from data $(\mathbf{X}, \mathbf{y})$, where no instruments have been randomized in our setting (Janzing, 2019, step 5 in

Alg. ConCorr).⁵ Let us denote this estimator by $\widehat{\|\beta\|_2}$. We outline the assumptions on the confounding model, which are satisfied in our empirical evaluation, in Section 5.

Following Corollary 3 we can use the difference of the 2-norm of our running estimate $\|\widehat{P_{\alpha[t]}}\beta\|_2$ and $\widehat{\|\beta\|_2}$ as a heuristic for whether we have already fully identified β , i.e., whether the overall instrumented subspace contains β .⁶ In future work, with the asymptotic normality of $\widehat{P_{\alpha}}\beta$ in Proposition 1 and thus an explicit expression for the asymptotic distribution of $\|\widehat{P_{\alpha}}\beta\|_2^2$ (Moschopoulos, 1985) as well as the asymptotic behavior of $\widehat{\|\beta\|_2}$ in Rendsburg et al. (2022), one could derive (asymptotic) confidence intervals for whether $\|\widehat{P_{\alpha[t]}}\beta\|_2 = \widehat{\|\beta\|_2}$. In practice, we propose

$$\left| \widehat{\|\beta\|_2} - \|\widehat{P_{\alpha[t]}}\beta\|_2 \right| < \epsilon \quad (5)$$

for a fixed tolerance $\epsilon > 0$ as a stopping criterion of our algorithm. When this condition is reached, we can conclude (for sufficiently large n) that β is near fully identified, even though we may have used fewer than d_x instruments.

Identified components. When the above stopping criterion is not reached within our T rounds of experimentation, learning an orthogonal projection of β may still be informative in itself. However, in many situations one is interested in precise values of individual components β_i for $i \in [d_x]$. Notably, when $\beta \notin \text{im}(\alpha)$, the individual components of $P_{\alpha}\beta$ are not easily interpretable. For example, neither holds $(P_{\alpha}\beta)_i = 0 \Rightarrow \beta_i = 0$ nor the other way round. Therefore, we now devise a method to determine when we can trust any given component of our running estimate $\widehat{P_{\alpha[t]}}\beta$.

Corollary 4. *Let $(e_i)_{i \in [d_x]}$ be the standard basis of $\mathbb{R}^{d_x}$. In the setting of Proposition 2, when $V_{\alpha[t]} V_{\alpha[t]}^T e_i = e_i$, then $(\widehat{P_{\alpha[t]}}\beta)_i$ consistently estimates β_i .*

Proof. This follows from Proposition 2 and the linearity of projections when considering $\beta = \sum_{i=1}^{d_x} \beta_i e_i$. $\square$

We note that we already compute $V_{\hat{\alpha}_t}$ at each round as part of $\widehat{P_{\alpha_t}}\beta$. Since $V_{\hat{\alpha}_t} V_{\hat{\alpha}_t}^T \xrightarrow{P} P_{\alpha[t]}$, the check for identified components in Corollary 4 can therefore be performed efficiently in practice by checking for (approximate) equality of $V_{\hat{\alpha}_t} V_{\hat{\alpha}_t}^T e_i \approx e_i$. In our empirical evaluation, we use the absolute value of the cosine similarity, denoted by cdist , between the two vectors as a continuous measure

⁵Since we assume experimental access, we can also collect purely observational (X, Y) data about the system of interest.

⁶Recall that the 2-norm of an orthogonal projection of β is always smaller or equal to the 2-norm of β . This provides additional motivation for why we choose to combine estimates according to Equation (4) using a minimum 2-norm objective.

for whether β_i has been identified. As an aggregate metric, we report the percentage of identified components

$$\frac{1}{d_x} \left| \{i \in [d_x] \mid \text{cdist}(V_{\alpha[t]} V_{\alpha[t]}^T e_i, e_i) < \delta\} \right| \quad (6)$$

for some fixed tolerance $\delta > 0$.

Finally, we can estimate an upper bound on the absolute error in each non-identified component. Let $\beta = P_{\alpha}\beta + \nu$ be the orthogonal decomposition of β into the instrumented subspace $\text{im}(\alpha)$ and its orthogonal complement. Since $\|\beta\|_2^2 = \|P_{\alpha}\beta\|_2^2 + \|\nu\|_2^2$ and $|\nu_i| \leq \|\nu\|_2$ for all components $i \in [d_x]$, we have

$$|\nu_i| \leq \sqrt{\|\beta\|_2^2 - \|P_{\alpha}\beta\|_2^2} \text{ for all } i \in [d_x]. \quad (7)$$

With our consistent estimates $\widehat{\|\beta\|_2}$ and $\widehat{P_{\alpha}}\beta$, we can thus upper bound all remaining unidentified components.

We return now to Figure 2 which illustrates the estimation (upper panel) and combination steps (lower panel) in our linear setting for $d_z = 2, d_x = 3, T = 2$. For both instruments, we estimate the projection of β onto the respective instrumented spaces ($P_{\alpha_1}\beta$ and $P_{\alpha_2}\beta$), including the effect $P_{\alpha_{1/2}}\beta$ which is the one that should be recovered. In the lower panel, the two planes are the orthogonal complements of the instrumented spaces and their intersection corresponds to all vectors that are compatible, i.e., would be projected onto $\text{im}(\alpha_1)$ and $\text{im}(\alpha_2)$ respectively. This corresponds to the constraints in Equation (4). Among those, we then select the vector with the smallest norm as our combined estimate. The figure both illustrates the necessity of linearity for the combination of estimates and the increasing norm of the combined estimate converging to $\|\beta\|_2$.

4. Sequential Selection of Instruments

At each step $t \in [T]$, we seek to select the most informative subset of instruments α_t out of the pool of N_{IV} available choices. After each round we combine the newly obtained estimate $\widehat{P_{\alpha_t}}\beta$ with all previous ones to obtain $\widehat{P_{\alpha[t]}}\beta$. Regarding consideration (B3), we assume a cost function $c : [N_{IV}] \rightarrow \mathbb{R}_{\geq 0}$, where $c(d_z)$ is the cost of running a single experiment with d_z randomized instruments. For instance, this may be the actual monetary and logistic cost of randomly administering the selected drugs to a collection of n cell cultures (Z), sequencing the cells (X), and measuring the outcome of interest (y) for each culture. However, the cost may also incorporate a hard limit $N_{IV/\text{exp}}$ on the number of instruments that can sensibly be combined in a single experiment (e.g., without killing the organism), by setting $c(d) = \infty$ for $d > N_{IV/\text{exp}}$.

In light of Equation (5), the ultimate goal is to select instruments that maximize $\|\widehat{P_{\alpha[T]}}\beta\|_2$. However, in round t

we cannot anticipate by how much $\|\widehat{P_{\alpha_{[t-1]}}\beta}\|_2$ is going to increase for a candidate set of instruments α_t without actually performing the experiment. Therefore, we must rely on another signal to sequentially select subsets of instruments. Without any information about the available instruments, we cannot do better than selecting instruments (uniformly) at random at each step. While in practice one may not have precise information about the actual effects of individual instruments, information about the similarity of different antibiotics or drugs is typically still available. For example, certain antibiotics may have related active agents (high similarity) or certain drugs may target similar pathways (high similarity). We assume that pairwise normalized similarities $\text{sim}_{i,j} \in [0, 1]$ are provided for all available instruments $i, j \in [N_{IV}]$. Here, $\text{sim}_{i,j} = \text{sim}_{j,i}$ and $\text{sim}_{i,i} = 1$. Such similarities could also be derived from a set of features known about the instruments. To optimally explore the treatment space, it is then natural to attempt to sequentially select highly dissimilar instruments.

Hence, to evaluate the expected gain of adding a new set of instruments $\mathcal{I} \subset [N_{IV}]$ to the already used instruments $\mathcal{J} \subset [N_{IV}]$, we define the following gain function

$$\begin{aligned} \text{gain} : 2^{[N_{IV}]} \times 2^{[N_{IV}]} &\rightarrow \mathbb{R}_{\geq 0}, \\ (\mathcal{I}, \mathcal{J}) &\mapsto \frac{1}{|\mathcal{I}| + |\mathcal{J}| - 1} \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{I} \cup \mathcal{J}} (1 - \text{sim}_{i,j}). \end{aligned} \quad (8)$$

This takes into account the similarities of the newly proposed instrument set within itself as well as with respect to the previously used ones. For example, when all instruments are maximally dissimilar ($\text{sim}_{i,j} = \delta_{i,j}$), $\text{gain}(\mathcal{I}, \mathcal{J}) = |\mathcal{I}|$ regardless of $\mathcal{J}$. When all instruments are equal ($\text{sim}_{i,j} = 1$), $\text{gain}(\mathcal{I}, \mathcal{J}) = 0$ for all inputs. Finally, we define the score of the set $\mathcal{I}$ when the set $\mathcal{J}$ has already been used by

$$\text{score}(\mathcal{I}, \mathcal{J}) := \text{gain}(\mathcal{I}, \mathcal{J}) - c(|\mathcal{I}|). \quad (9)$$

These considerations lead to the following setting. At each round t , we select a subset $\mathcal{I}_t \subset [N_{IV}] \setminus \mathcal{I}_{[t-1]}$ of still unused instruments that maximize the score function given the already used instruments. We run a randomized experiment with those instruments to collect data, estimate (a projection of) β from this data (Proposition 1), and combine the estimate with the previous ones (Proposition 2). By convention, we have $\mathcal{I}_{[0]} = \mathcal{I}_\emptyset = \emptyset$ and we overload terminology to call both $\alpha_t \in \mathbb{R}^{|\mathcal{I}_t| \times d_x}$ and $\mathcal{I}_t \subset [N_{IV}]$ “the instruments selected at round t ”. We similarly use $\alpha_{[t]}$ and $\mathcal{I}_{[t]}$ for the instruments selected up to (and including) round t . When the stopping criterion (Equation (5)) is satisfied, we return our current estimate as an estimate of the full β . Otherwise, we return the estimate after T experiments together with the identified components (Corollary 4). We outline our sequential instrument selection (SIS) procedure in Algorithm 1.

We remark that since $\text{gain}(\mathcal{I}, \mathcal{J}) \leq |\mathcal{I}|$ it makes sense to

use a sublinear cost function (until a potential hard limit $N_{IV/\text{exp}}$). Intuitively, while it becomes more expensive to include multiple randomized variables in a single experiment, it is still cheaper than running an individual randomized experiment for each instrument separately. The precise choice of the cost function is informed by the actual experimental setting and determines the trade-off between randomizing many instruments at once (to increase the dimensionality of the instrumented subspace) and the cost of doing so.

Algorithm 1 Sequential selection of instrument sets

Require: maximum rounds T , pairwise similarities $\text{sim} \in [0, 1]^{N_{IV} \times N_{IV}}$, cost function c , tolerance $\epsilon > 0$

- 1: collect observational data $\mathbf{X}, \mathbf{y}$
- 2: compute $\|\beta\|_2$ from $\mathbf{X}, \mathbf{y}$ $\triangleright$ Janzing (2019, ConCorr)
- 3: $\mathcal{C} \leftarrow \emptyset$ $\triangleright$ set of identified components
- 4: **for** $t \in [T]$ **do** $\triangleright$ experimental rounds
- 5: $\mathcal{I}_t \leftarrow \arg \max_{\mathcal{I} \subset [N_{IV}] \setminus \mathcal{I}_{[t-1]}} \text{score}(\mathcal{I}, \mathcal{I}_{[t-1]})$
- 6: collect $\mathbf{Z}_t, \mathbf{X}_t, \mathbf{y}_t$ $\triangleright$ run experiment with $\mathcal{I}_t$
- 7: $\widehat{P_{\alpha_t}}\beta \leftarrow (\mathbf{X}_t^T \mathbf{P}_{\mathbf{Z}_t} \mathbf{X}_t)^+ \mathbf{X}_t^T \mathbf{P}_{\mathbf{Z}_t} \mathbf{y}_t$ $\triangleright$ Proposition 1
- 8: $\widehat{P_{\alpha_{[t]}}}\beta \leftarrow \arg \min_{\gamma \in \mathbb{R}^{d_x}} \|\gamma\|_2$ $\triangleright$ Proposition 2
 s.t. $\widehat{P_{\alpha_\tau}}\beta = \mathbf{V}_{\alpha_\tau} \mathbf{V}_{\alpha_\tau}^T \gamma$ for all $\tau \in [t]$
- 9: **if** $\|\widehat{P_{\alpha_{[t]}}}\beta - \|\beta\|_2\|_2 < \epsilon$ **then** $\triangleright$ fully identified, (5)
- 10: $\mathcal{C} \leftarrow [d_x]$
- 11: **return** $\widehat{P_{\alpha_{[t]}}}\beta, \mathcal{C}$
- 12: $\mathcal{C} \leftarrow \{i \in [d_x] \mid \mathbf{V}_{\alpha_{[t]}} \mathbf{V}_{\alpha_{[t]}}^T e_i \approx e_i\}$ $\triangleright$ Corollary 4
- 13: **return** $\widehat{P_{\alpha_{[T]}}}\beta, \mathcal{C}$ $\triangleright$ estimate, identified components

5. Empirical Evaluation

Setup. Since a real-world evaluation of our approach would require access to sequential randomized experimentation in a complex setting, we are restricted to simulation studies. We first illustrate the properties of our proposed (combined) causal effect estimators in the underspecified IV setting and then evaluate our full sequential instrument selection method. We generate the parameters α, β randomly (Appendix A) in order to avoid parameter selection bias. Next, we fix a mixing matrix $\mathbf{M} \in \mathbb{R}^{d_x \times d_x}$ with all entries sampled from independent standard Gaussians as well as a direction $v \in \mathbb{R}^{d_x}$ as a uniform sample from the sphere $\mathbb{S}^{d_x-1}$. In Equation (1), we then sample instrument components independently from a Rademacher distribution, and set $\epsilon_X = \mathbf{M}e$, $\epsilon_Y = v^T e$ with all components of $e \in \mathbb{R}^{d_x}$ sampled independently from standard Gaussians for each sample. This confounding model satisfies the assumptions required to estimate $\|\beta\|_2$ (Janzing & Schölkopf, 2018; Janzing, 2019; Rendsburg et al., 2022). Loosely speaking, the

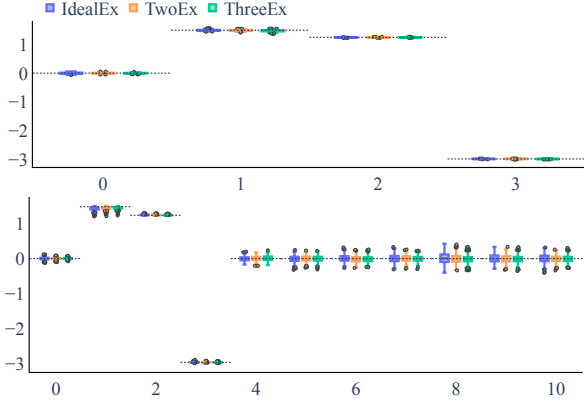


Figure 3. Estimates of β over 500 runs with $d_z = d_x = 3$ (top) and $d_z = 3$ and $d_x = 10$ (bottom). The dotted lines are the true β and boxplots show the median (horizontal line), first and third quartile (box height) and 10/90 percentiles (whiskers).

confounder affects both the treatment and the outcome as (independent) random mixtures of the same independent noise sources. We ensure in our data generation that $\beta \in \text{im}(\alpha)$. Hence, β can be recovered fully in principle. Moreover, we also guarantee that there is a subset of less than $N_{\text{IV}/\text{exp}} \cdot T$ instruments that suffices to fully identify β (see Appendix A for details). Therefore, even if we cannot use all instruments throughout the experiments ($N_{\text{IV}/\text{exp}} \cdot T < N_{\text{IV}}$) a good selection algorithm can in principle fully identify β .

For the similarities between instruments, in our experiments we take the absolute value of the cosine distance $\text{sim}_{i,j} = |\alpha_i^T \alpha_j| / (\|\alpha_i\|_2 \|\alpha_j\|_2)$. Crucially, we thereby do not assume access to α —only the similarities enter the algorithm. While in a given application, the similarities would be provided by practitioners and domain experts, in our simulated study we have to choose *some* similarity measure that is informed by α . We account for uncertainties in the provided similarities by computing them on a noisy version of α , where we add independent standard Gaussian noise to all entries. For the cost function we choose $c(d) = \log(d)$ for $d \leq N_{\text{IV}/\text{exp}}$ and $c(d) = \infty$ otherwise, effectively limiting the maximum number of instruments per round to $N_{\text{IV}/\text{exp}}$. While pairwise similarities can help the sequential selection to converge quickly, we note that our findings from Section 3 are extremely useful for the (often strong) baseline of selecting instruments randomly.

Finite sample properties of our estimators. We first empirically analyze the finite sample behavior of our estimators from Propositions 1 and 2. For illustration, Figure 3 shows a setting with $d_x = d_z = 3$ on the upper panel and $d_x = 10$, $d_z = 3$ on the lower panel. The x-axis shows component indices of $\beta \in \mathbb{R}^{d_x}$ (component 0 is the offset); dotted lines are the ground truth values of β , and boxplots show distribution of estimated values ($\widehat{P}_\alpha \beta$) over 500 random seeds. We compare three different estimators. *IdealEx* randomizes all

three available instruments simultaneously and shows the single estimate (from Proposition 1). *TwoEx* combines (via Proposition 2) two individual estimates using only the first two and the last instrument, respectively. *ThreeEx* combines three individuals estimates, obtained by randomizing each of them separately. Figure 3 shows that all methods correctly identify all three components of β (plus the offset) on average with low variance in the low-dimensional $d_x = 3$ setting. The $d_x = 10$ setting (lower panel) in which 3 instruments suffice for full identification, shows similar results for all components. Even though the variance increases compared to the lower-dimensional setting, estimates are still consistent. All details are provided in Appendix A.

Sequential instrument selection. For our instrument selection algorithm we compare the following methods. *IdealEx*: the hypothetical ideal where all instruments are used at once in a single experiment. *Random*: a baseline that selects one of the allowed (smaller or equal $N_{\text{IV}/\text{exp}}$) subsets uniformly at random from the remaining instruments at each round. *Sequential Instrument Selector (SIS)*: our Algorithm 1. We remark that the random baseline, to its advantage, does not respect the cost per experiment, but is allowed to select the maximum number of possible instruments in each round.

We set $d_z = 30$, $d_{\text{id}} = 15$ and use two different treatment dimensions $d_x = 50$ and $d_x = 150$. Further, we allow $N_{\text{IV}/\text{exp}} = 3$ instruments per round with a total budget of $T = 6$ experimental rounds. Note that $N_{\text{IV}/\text{exp}} \cdot T \leq d_{\text{id}}$, i.e., the algorithm has the budget to fully identify β .

In Figure 4 (left) we show boxplots of the mean squared error (MSE) of our estimates for the nonzero components after each optimization round. *SIS* outperforms the random baseline in terms of MSE. Additionally, each boxplot in the lower panel shows the percentage of uncertain components from Equation (6) for $\delta = 0.3$. In Figure 4 (right), the norm of our estimator $\|\widehat{P}_\alpha \beta\|$ converges towards the norm of the actual β (and its estimate $\|\widehat{\beta}\|_2$), i.e., the stopping criterion is reached. Figure 5 compares these estimates for the nonzero components of β after the last round of experiments. Our greedy optimization indeed identifies each of the components just as well as the hypothetical ideal of a single experiment that uses all instruments at once. Empirically, the finite sample properties remain unaffected by our combination procedure. We refer to Appendix B for results on a $d_x = 150$ setting.

6. Conclusion

In this work we made multiple contributions aiming at inferring causal effects of high-dimensional treatments under unobserved confounding by sequential experimentation, where we cannot intervene on the treatments directly, but can only randomize instruments. We proposed consistent, asymptoti-

Sequential Underspecified Instrument Selection for Cause-Effect Estimation

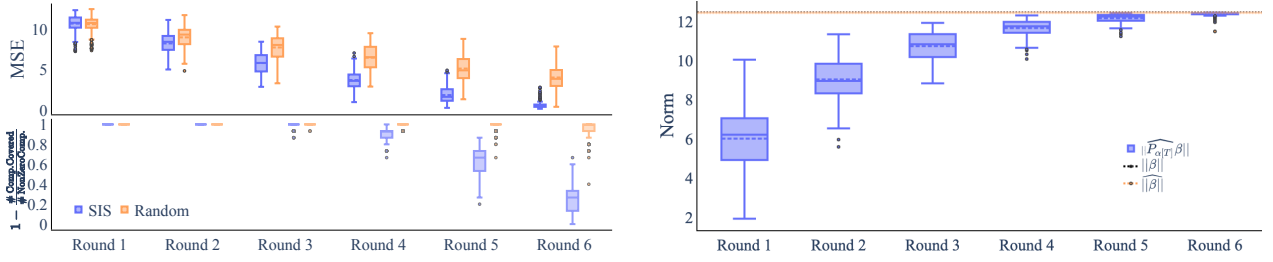


Figure 4. Results for the sequential selection of instruments with $d_z = 30$, $d_{id} = 15$, $N_{IV/exp} = 3$, $T = 6$ and $d_x = 50$: each boxplot shows the median and mean (solid resp. dashed line), first and third quartile (box height) and 10/90 percentiles (whiskers) over $n_{runs} = 250$. **Left:** The upper panel shows the of squared error $\widehat{P}_{\alpha[t]}\beta - \beta$ over rounds $t \in \{1, \dots, 6\}$. The lower panel shows the corresponding percentages of unidentified components. **Right:** The $\|\widehat{P}_{\alpha[t]}\beta\|$ increases for $t \in \{1, \dots, 6\}$ approaching $\|\beta\|$ (dotted orange line), which perfectly estimates the true $\|\beta\|$ in this case (dotted black line).

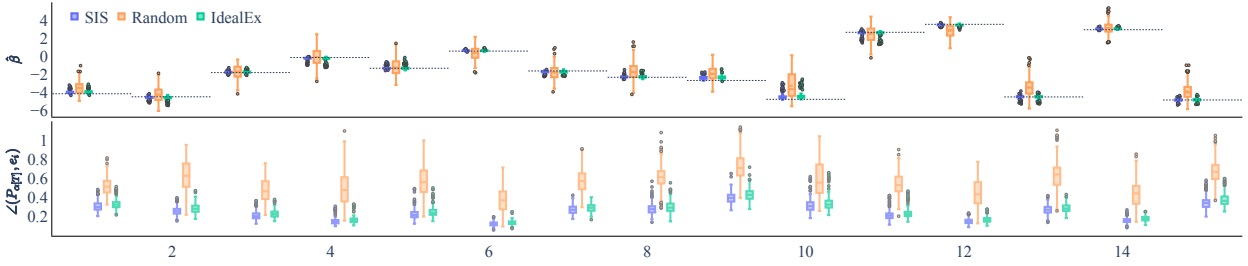


Figure 5. Results for the sequential selection of instruments with $d_z = 30$, $d_{id} = 15$, $N_{IV/exp} = 3$, $T = 6$ and $d_x = 50$: each boxplot shows the median and mean (solid resp. dotted line), first and third quartile (box height) and 10/90 percentiles (whiskers) over $n_{runs} = 250$ after the last round $T = 6$. **Top:** estimates of the nonzero components of β with the black dotted line being ground truth. **Bottom:** distribution of cosine similarities in Equation (6) for each component.

cally normal estimators for the orthogonal projection of a treatment effect onto the instrumented subspace in the linear setting and introduced a method to consistently combine such estimates from separate experiments. Surprisingly, neither the (perhaps intuitive) estimator in Proposition 1, nor the geometric intuition around instrumented subspaces in the underspecified setting can be found in existing literature. These estimators may be of independent interest as a contribution to the largely ignored underspecified IV setting. We then developed an algorithm to sequentially propose subsets of instruments from a given pool that flexibly trades off the expected information gain (informed by provided similarities) with the cost of each experiment. Moreover, we integrated a stopping criterion for when the sequential selection has fully identified the causal effect, a method to keep track of all components that are consistently estimated, and an upper bound on the absolute error of unidentified components: these additions inform the practitioner about whether (and which parts) of the estimate can be trusted.

The linearity assumption may appear restrictive. However, the linear IV setting is still heavily used in econometrics and health, as it reliably captures dominant effects even in noisy settings, and still attracts attention with novel results recently (Pfister & Peters, 2022; Rothenhäusler et al., 2018). The thorough theoretical understanding developed in this

work is a challenging and necessary foundation for experiment design via instrument selection. Extensions of our method and of our notion of the instrumented subspace to (certain) non-linear settings is an important direction for future work. A second limitation of our work is inherent to the problem setting: missing real-world experiments due to a lack of access to the required expensive, specialized facilities. Finally, we highlight that independent testing and verification is paramount when using algorithmically obtained insights to inform consequential decisions such as actual clinical treatment decisions.

Code. The implementation as well as experimental details are publicly available on Github: <https://github.com/EAILer/underspecified-iv>.

Acknowledgments. We thank Sören Becker for useful suggestions and Dominik Janzing for fruitful discussions on estimating $\|\beta\|_2$. EA is supported by the Helmholtz Association under the joint research school “Munich School for Data Science - MUDS”. JH is supported by the Natural Sciences and Engineering Research Council of Canada (NSERC) and Recursion Pharmaceuticals. This work has been supported by the Helmholtz Association’s Initiative and Networking Fund through CausalCellDynamics (grant # Interlabs-0029).

References

- Ailer, E., Müller, C. L., and Kilbertus, N. A causal view on compositional data. *arXiv preprint arXiv:2106.11234*, 2021. 1, 2
- Angrist, J. D. and Pischke, J.-S. *Mostly harmless econometrics: An empiricist's companion*. Princeton university press, 2008. 2, 3, 4
- Bennett, A., Kallus, N., and Schnabel, T. Deep generalized method of moments for instrumental variable analysis. In Wallach, H., Larochelle, H., Beygelzimer, A., d'Alché-Buc, F., Fox, E., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. URL <https://proceedings.neurips.cc/paper/2019/file/15d185eaa7c954e77f5343d941e25fbd-Paper.pdf>. 2
- Bura, E. and Pfeiffer, R. On the distribution of the left singular vectors of a random matrix and its applications. *Statistics & Probability Letters*, 78(15):2275–2280, 2008. 4
- Didelez, V. and Sheehan, N. Mendelian randomization as an instrumental variable approach to causal inference. *Statistical Methods in Medical Research*, 16(4):309–330, 2007. 2
- Fu, Y., Foden, J., Khayter, C., Maeder, M., Reyon, D., Joung, J., and Sander, J. High-frequency off-target mutagenesis induced by crispr-cas nucleases in human cells. *Nature biotechnology*, 31, 06 2013. doi: 10.1038/nbt.2623. 1
- Gamella, J. L. and Heinze-Deml, C. Active invariant causal prediction: Experiment selection through stability. *Advances in Neural Information Processing Systems*, 33: 15464–15475, 2020. 2
- Hahn, J. and Hausman, J. Estimation with valid and invalid instruments. *Annales d'Économie et de Statistique*, 79/80: 25–57, 2005. ISSN 0769489X, 22726497. URL <http://www.jstor.org/stable/20777569>. 3
- Hartford, J., Lewis, G., Leyton-Brown, K., and Taddy, M. Deep iv: A flexible approach for counterfactual prediction. In *International Conference on Machine Learning*, pp. 1414–1423, 2017. 2
- Hernán, M. A. and Robins, J. M. Instruments for causal inference: an epidemiologist's dream? *Epidemiology*, pp. 360–372, 2006. 2
- Hyttinen, A., Eberhardt, F., and Hoyer, P. O. Experiment selection for causal discovery. *Journal of Machine Learning Research*, 14:3041–3071, 2013. 2
- Janzing, D. Causal regularization. *Advances in Neural Information Processing Systems*, 32, 2019. 5, 7
- Janzing, D. and Schölkopf, B. Detecting non-causal artifacts in multivariate linear regression models. In *International Conference on Machine Learning*, pp. 2245–2253. PMLR, 2018. 5, 7, 12
- Kang, H., Zhang, A., Cai, T. T., and Small, D. S. Instrumental variables estimation with some invalid instruments and its application to mendelian randomization. *Journal of the American statistical Association*, 111(513):132–144, 2016. 3
- Kilbertus, N., Kusner, M. J., and Silva, R. A class of algorithms for general instrumental variable models. In *Advances in Neural Information Processing Systems*, volume 33, 2020. 2
- Moschopoulos, P. G. The distribution of the sum of independent gamma random variables. *Annals of the Institute of Statistical Mathematics*, 37(1):541–544, 1985. 6
- Muandet, K., Mehrjou, A., Lee, S. K., and Raj, A. Dual instrumental variable regression. *arXiv preprint arXiv:1910.12358*, 2019. 2
- Nocedal, J. and Wright, S. J. *Numerical optimization*. Springer, 2006. 5
- Padh, K., Zeitler, J., Watson, D., Kusner, M., Silva, R., and Kilbertus, N. Stochastic causal programming for bounding treatment effects, 2022. URL <https://arxiv.org/abs/2202.10806>. 2
- Pearl, J. *Causality*. Cambridge university press, 2009. 1
- Pfister, N. and Peters, J. Identifiability of sparse causal effects using instrumental variables. *arXiv preprint arXiv:2203.09380*, 2022. 2, 9
- Rendsburg, L., Vankadara, L. C., Ghoshdastidar, D., and von Luxburg, U. A consistent estimator for confounding strength. *arXiv preprint arXiv:2211.01903*, 2022. 5, 6, 7
- Rothenhäusler, D., Bühlmann, P., Meinshausen, N., and Peters, J. Anchor regression: Heterogeneous data meet causality. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 83, 01 2018. doi: 10.1111/rssb.12398. 2, 9
- Saengkyongam, S., Henckel, L., Pfister, N., and Peters, J. Exploiting independent instruments: Identification and distribution generalization. *arXiv preprint arXiv:2202.01864*, 2022. 2
- Sanderson, E., Glymour, M. M., Holmes, M. V., Kang, H., Morrison, J., Munafò, M. R., Palmer, T., Schooling, C. M., Wallace, C., Zhao, Q., and Davey Smith,

- G. Mendelian randomization. *Nature Reviews Methods Primers*, 2(1):6, 2022. [2](#)
- Singh, R., Sahani, M., and Gretton, A. Kernel instrumental variable regression. In *Advances in Neural Information Processing Systems*, pp. 4593–4605, 2019. [2](#)
- Sohn, M. B. and Li, H. Compositional mediation analysis for microbiome studies. *Annals of Applied Statistics*, 13(1):661–681, 2019. ISSN 19417330. doi: 10.1214/18-AOAS1210. [2](#)
- Stock, J. H. and Trebbi, F. Retrospectives: Who invented instrumental variable regression? *Journal of Economic Perspectives*, 17(3):177–194, 2003. [2](#)
- Sussex, S., Uhler, C., and Krause, A. Near-optimal multi-perturbation experimental design for causal structure learning. *Advances in Neural Information Processing Systems*, 34:777–788, 2021. [2](#)
- Tigas, P., Annadani, Y., Jesson, A., Schölkopf, B., Gal, Y., and Bauer, S. Interventions, where and how? experimental design for causal models at scale. *arXiv preprint arXiv:2203.02016*, 2022. [2](#)
- Wang, C., Hu, J., Blaser, M. J., Li, H., and Birol, I. Estimating and testing the microbial causal mediation effect with high-dimensional and compositional microbiome data. *Bioinformatics*, 2020. ISSN 14602059. doi: 10.1093/bioinformatics/btz565. [2](#)
- Wright, P. G. *Tariff on animal and vegetable oils*. Macmillan Company, New York, 1928. [2](#)
- Zhang, R., Imaizumi, M., Schölkopf, B., and Muandet, K. Maximum moment restriction for instrumental variable regression. *arXiv preprint arXiv:2010.07684*, 2020. [2](#)
- Zhang, X.-H., Tee, L. Y., Wang, X.-G., Huang, Q.-S., and Yang, S.-H. Off-target effects in crispr/cas9-mediated genome engineering. *Molecular Therapy - Nucleic Acids*, 4:e264, 2015. [1](#)

A. Details of Experiments

Data Generation The data generation process follows the model described in (Janzing & Schölkopf, 2018). We adopt their data setup in order to estimate $\|\beta\|$ consistently:

$$X = Z\alpha + \epsilon_X, \quad Y = X\beta + \epsilon_Y, \quad (10)$$

ϵ_X and ϵ_Y are confounded via the common variable e .

$$\epsilon_X = eM, \quad \epsilon_Y = ev, \quad (11)$$

with $Z \sim \text{Rademacher}(0.5)$ and $e \sim \mathcal{N}(0, Id_l)$.

Scenario Generation. For each figure, i.e. simulation, we generate random scenarios in order to avoid introducing involuntary bias in the parameter setting.

Each generated scenario is based on a seed and the constants n, d_x, d_z and d_{id} , i.e. the number of instruments it will take to identify the causal effect in full. Further, we assume $d_x = l$, i.e. $M \in \mathbb{R}^{d_x \times d_x}$. Note that this choice is based on the setting in the simulation studies of (Janzing & Schölkopf, 2018).

We sample d_{id} nonzero components of $\alpha_j, j \in \{1, \dots, d\}$ and the d_{id} nonzero components of β from a uniformly distributed random variable $\mathcal{U}(-5.0, 5.0)$. The sparsity is introduced by setting the remaining components to 0.0. Our motivation for this choice is to reliably generate a setup for which we are guaranteed to identify the causal effect by a maximum of d_{id} instruments. In addition to d_{id} identifying instruments, we generate $d - d_{\text{id}}$ further instruments by picking two of the necessary instruments on top of which we add a p -dimensional Gaussian noise. Overall, we end up with $d_{\text{id}} - 2$ necessary instruments and two clusters from which we can pick any instrument in order to identify the remaining components of β . For the confounder we sample all entries of $M \in \mathbb{R}^{d_x \times d_x}$ from independent standard Gaussians and the direction $v \in \mathbb{R}^{d_x}$ as a uniform sample from the sphere $\mathbb{S}^{d_x-1}$.

Scenario Generation for Figure 3. For showcasing the finite sample properties of our estimator, we sampled from the scenario generation above with seed 253 and $d_z = 3$ and $d_x = 3$ resp. $d_x = 10$. This choice of parameters left us with two settings (1) being just-identified $d_z = d_x = 3$ and (2) being underspecified $d_z = 3, d_x = 10$. Moreover, we set β to only have 3 non-zero components in order to be able to identify the causal effect in full. Note that this choice is for illustration purpose and does not affect the method’s applicability in a setting where we can only identify parts of the causal effect. However, in those setting where some β_i -components are not part of the instrumented subspace $\text{im}(\alpha)$, full causal recovery is in general impossible.

B. Additional Experiments

In the optimization we compare the baseline which uses all instruments at once (*IdealEx*) to the developed sequential optimization routine (*SIS*) and the random baseline. We include results for the same setting as Figure 4 and Figure 5, except with an increased treatment dimension, i.e. $d_x = 150$ in Figure 6:

$$d_x = 150, d_z = 30, d_{\text{id}} = 15, N_{\text{IV}/\text{exp}} = 3, T = 6$$

Moreover, it might not always be the case that the stopping criterion in (Janzing & Schölkopf, 2018) works as nicely as in the previous scenarios, see Fig. Figure 7 with parameter setting:

$$d_x = 50, d_z = 30, d_{\text{id}} = 20, N_{\text{IV}/\text{exp}} = 4, T = 6$$

Sequential Underspecified Instrument Selection for Cause-Effect Estimation

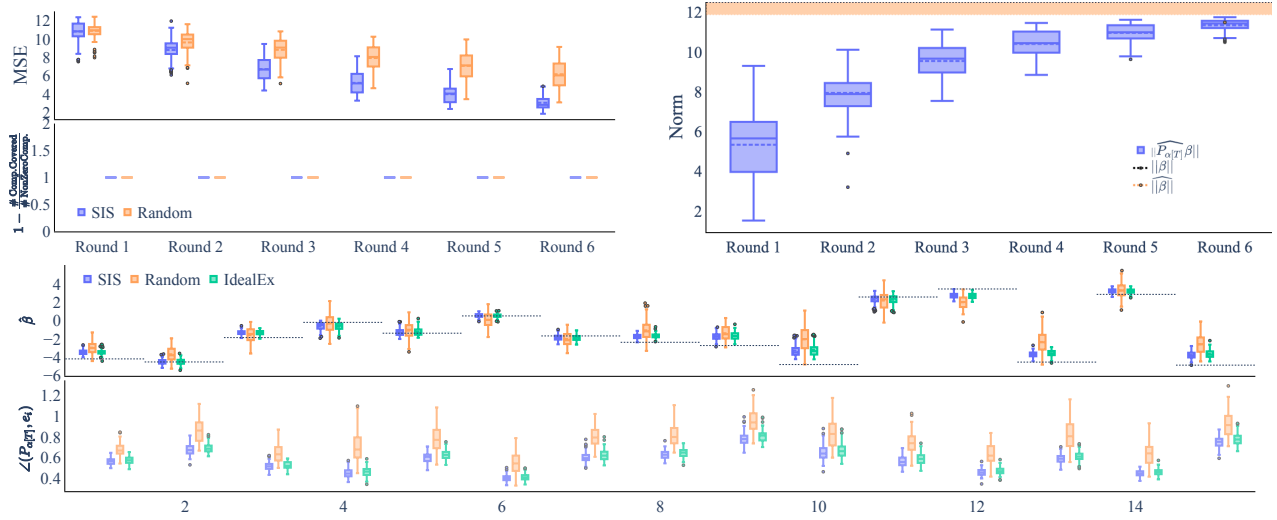


Figure 6. Results for the sequential selection of instruments with $d_z = 30, d_{id} = 15, N_{IV/exp} = 3, T = 6$ and $d_x = 150$: each boxplot shows the median and mean (solid resp. dashed line), first and third quartile (box height) and 10/90 percentiles (whiskers) over $n_{runs} = 250$. **Top Left:** The upper panel shows the of squared error $\widehat{P}_{\alpha[t]}\beta - \beta$ over rounds $t \in \{1, \dots, 6\}$. The lower panel shows the corresponding percentages of unidentified components. As we are in the scenario with $p = 150$, we would need to adjust the threshold $\delta = 0.3$ to a higher value. **Top Right:** The $\|\widehat{P}_{\alpha[t]}\beta\|$ increases for $t \in \{1, \dots, 6\}$ approaching $\|\widehat{\beta}\|$ (shaded orange block), which does not perfectly estimate the true $\|\beta\|$ in this case (dotted black line), but performs still reasonably well. **Bottom:** *Upper Panel:* estimates of the nonzero components of β with the black dotted line being ground truth. *Lower Panel:* distribution of cosine similarities in Equation (6) for each component.

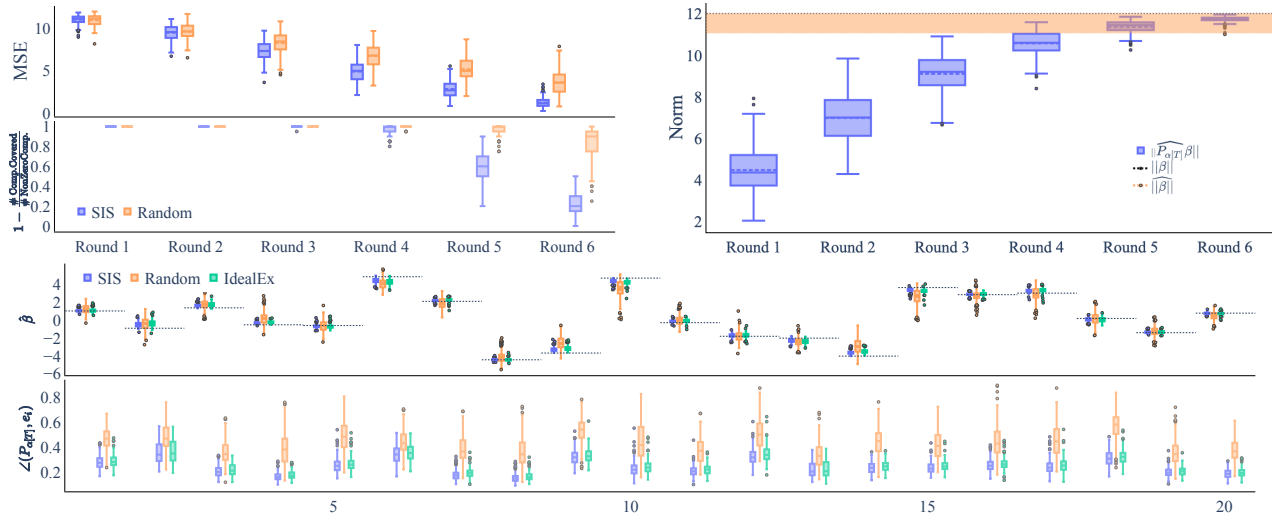


Figure 7. Results for the sequential selection of instruments with $d_z = 30, d_{id} = 20, N_{IV/exp} = 4, T = 6$ and $d_x = 50$: each boxplot shows the median and mean (solid resp. dashed line), first and third quartile (box height) and 10/90 percentiles (whiskers) over $n_{runs} = 250$. **Top Left:** The upper panel shows the of squared error $\widehat{P}_{\alpha[t]}\beta - \beta$ over rounds $t \in \{1, \dots, 6\}$. The lower panel shows the corresponding percentages of unidentified components. **Top Right:** The $\|\widehat{P}_{\alpha[t]}\beta\|$ increases for $t \in \{1, \dots, 6\}$ approaching $\|\widehat{\beta}\|$ (shaded orange block), which does not perfectly estimate the true $\|\beta\|$ in this case (dotted black line), but performs still reasonably well. **Bottom:** *Upper Panel:* estimates of the nonzero components of β with the black dotted line being ground truth. *Lower Panel:* distribution of cosine similarities in Equation (6) for each component.

Who holds the Copyright on a ICML paper

According to U.S. Copyright Office's page [What is a Copyright](#), When you create an original work you are the author and the owner and hold the copyright, unless you have an agreement to transfer the copyright to a third party such as the company or school you work for.

Authors do not transfer the copyright of their paper to ICML, instead they grant ICML a non-exclusive, perpetual, royalty-free, fully-paid, fully-assignable license to copy, distribute and publicly display all or part of the paper

Figure A.1 ICML License Copyright Info Statement

A.2 Targeted Sequential Indirect Experiment Design

Targeted Sequential Indirect Experiment Design

Elisabeth Ailer

Technical University of Munich
Helmholtz Munich

Munich Center for Machine Learning (MCML)

Niclas Dern

Technical University of Munich

Jason Hartford

Valence Labs

Niki Kilbertus

Technical University of Munich
Helmholtz Munich

Munich Center for Machine Learning (MCML)

Abstract

Scientific hypotheses typically concern specific aspects of complex, imperfectly understood or entirely unknown mechanisms, such as the effect of gene expression levels on phenotypes or how microbial communities influence environmental health. Such queries are inherently causal (rather than purely associational), but in many settings, experiments can not be conducted directly on the target variables of interest, but are indirect. Therefore, they perturb the target variable, but do not remove potential confounding factors. If, additionally, the resulting experimental measurements are multi-dimensional and the studied mechanisms nonlinear, the query of interest is generally not identified. We develop an adaptive strategy to design indirect experiments that optimally inform a targeted query about the ground truth mechanism in terms of sequentially narrowing the gap between an upper and lower bound on the query. While the general formulation consists of a bi-level optimization procedure, we derive an efficiently estimable analytical kernel-based estimator of the bounds for the causal effect, a query of key interest, and demonstrate the efficacy of our approach in confounded, multivariate, nonlinear synthetic settings.

1 Introduction

Experimentation is the ultimate arbiter of scientific discovery. While advances in machine learning (ML) and ever increasing amounts of observational data promise to accelerate scientific discovery in virtually all scientific disciplines, all hypotheses ultimately have to be supported or falsified by experiments. But the space of possible experiments is combinatorial, and as a result, experimentation in the physical world may become a major bottleneck of the scientific process and critically inhibit, for example, fast turnaround in drug discovery loops even in fully automated labs. Efficient adaptive data-driven strategies to propose the most useful experiments are thus of vital importance.

Targeted experimentation requires a well-formed hypothesis about the underlying system. We focus on hypotheses that are typically expressed in terms of causal effects or properties of functional relationships. For example, we may posit that there is some ground truth mechanism/function f that describes how a given phenotype depends on gene expression levels in a cell, or how the composition of microbes affects certain environmental health indicators. In the natural sciences, such mechanisms are typically (a) nonlinear, (b) dependent on multi-variate inputs, and (c) confounded by additional, unobserved quantities. For example, the way in which the gut microbiome composition affects energy metabolism in humans is likely confounded by various environmental and lifestyle choices. Due to

these factors, f is generally unidentifiable, i.e., we cannot hope to infer f from (even infinite amounts of) observational data alone.

Scientific queries are typically more narrow: ‘Can I expect changes in body mass index when I increase the relative abundance of *Lactobacillus* bacteria in the gut? How does the susceptibility of breast cancer depend on BRCA1 expression levels?’ Such targeted scientific queries do not require full knowledge of the underlying mechanism f , but only concern a specific aspect of f , which we can express mathematically as a functional Q of f . Due to potential unobserved confounding, even targeted queries $Q[f]$ may not be identified from observational data alone (Bennett et al., 2022)—experimentation is required. In these examples, the inputs to f of interest – which we call *treatments* – are multi-variate measurements such as gene expression levels or microbiome abundances. Randomizing these treatment variables is typically thought of as the gold standard for removing confounding effects, but perfect randomization is typically impossible: we can seldom perform hard interventions that set a distribution over these variables’ outcomes independently of confounding factors. Instead, experimental access in the natural sciences often amounts to limited possible perturbations of the system that induce distribution shifts of the treatments. For example, while we can administer antibiotics with very predictable effects on certain microbes, such an intervention does not remove confounding effects from lifestyle choices such as nutrition and exercise or environmental factors. Therefore, we think of experimentation as a source of indirect interventions, which strongly perturb the distribution over the treatment variables of interest.

Our primary goal is to design experiments that maximally inform the query of interest, $Q[f]$, within a fixed budget of experimentation. For multi-variate treatments and nonlinear f the query $Q[f]$ may not be identifiable at all, or only be identified after an infeasible number of experiments. We address this by maintaining an upper and lower bound on the query, and sequentially selecting experiments that minimize the gap between these bounds. We show that by treating experiments as instrumental variables (IV) we can estimate the bounds by building on existing techniques in (underspecified) nonlinear IV estimation. Our procedure involves a bi-level optimization, where an inner optimization estimates the bounds, and an outer optimization seeks to minimize the gap between the bounds. We show that if we are prepared to assume that f lies within a reproducing kernel Hilbert space (RKHS), then there are queries $Q[f]$ for which we can solve the inner estimation in closed form. In summary, we make the following contributions:

- We formalize the problem of indirect adaptive experiments in nonlinear, multi-variate, confounded settings as a sequential underspecified instrumental variable estimation, where we seek to adaptively tighten estimated bounds on the target query.
- We derive closed form expressions for the bound estimates when assuming f to lie within a reproducing kernel Hilbert space (RKHS) for linear Q .
- We develop adaptive strategies for the outer optimization to iteratively tighten these bounds and demonstrate empirically that our method robustly selects informative experiments leading to identification of $Q[f]$ (collapsing bounds) when it can be identified via allowed experimentation.

2 Problem Setting and Related Work

2.1 Problem Setting

Data generating process. We assume the following data generating process

$$X = h(Z, U), \quad Y = f_0(X) + U, \quad (1)$$

where we have experimental access to $Z \in \mathcal{Z} \subset \mathbb{R}^{d_z}$, $X \in \mathcal{X} \subset \mathbb{R}^{d_x}$ are the actual inputs to the mechanism of interest, i.e., the ‘treatments’ (in the cause-effect estimation sense) or ‘targets’ (in the indirect experimentation sense), $Y \in \mathcal{Y} \subset \mathbb{R}$ is the scalar outcome, e.g., the phenotype of interest, and U is an unobserved confounding variable with $\mathbb{E}[U] = 0$. The function h determines how experiments (choices of Z) affect treatments X and can be arbitrary. The function f_0 is the mechanism of interest for our scientific query and can also be arbitrary. As in typical indirect experiment settings, we assume that Z consists of valid instrumental variables: (1) $Z \not\perp\!\!\!\perp X$, i.e., we do perturb the distribution of X by experimenting on Z , (2) $Z \perp\!\!\!\perp U$, i.e., our choice of experiments is not influenced by the confounding variables, and (3) $Z \perp\!\!\!\perp Y \mid X, U$, i.e., Z only affects Y via X and does not have a direct effect on Y . When X is multi-variate and f_0, h are allowed to be non-linear, f_0 is commonly not

identified from observational data. However, we argue that in such settings one should be aiming for a more targeted query, i.e., does not aim to identify f_0 fully, but only a certain aspect of the mechanism.

Scientific queries. We represent the scientific query of interest by a functional $Q : \mathcal{F} \rightarrow \mathbb{R}$ where $\mathcal{F} \subset L_2(X)$ is the space of considered mechanisms f .¹ Generic functionals can measure both local as well as global properties of any $f \in \mathcal{F}$: simple average treatment effects via $Q[f] := \mathbb{E}[Y | do(X = x^*)] = f(x^*)$; when $\mathcal{F}$ is a Hilbert space, projections of f onto a fixed basis function f^* via $Q[f] := \langle f^*, f \rangle_{\mathcal{F}} / \|f^*\|_{\mathcal{F}}^2$; the local causal effect of an individual component X_i on Y at x^* via $Q[f] := (\partial_i f)(x^*)$.² Here, x^* could be the mean of the observed treatments, i.e., represent a ‘base gene expression level’ and we are interested in how a local change away from that base level affects the outcome. While our methodology applies to any functional, we focus our empirical experiments on causal effects of individual components as key scientific queries, such as how does up- or down-regulating a given gene affect the phenotype?

Learning experiments. Our goal is to sequentially learn a policy $\pi \in \mathcal{P}(Z)$ ³ such that data sampled from the joint distribution $P_\pi(X, Y, Z) = \pi(Z)P(X | Z)P(Y | X)$ induced by the model in eq. (1) optimally informs $Q[f_0]$. We highlight that depending on f_0, h and the distribution of U , there may exist a policy π such that $Q[f_0]$ is identified from $P_\pi(X, Y, Z)$, but in general it may remain unidentified for all policies even in the infinite data limit. For a non-optimal policy, i.e., non-informative experimentation, we should expect $Q[f_0]$ to be partially identified from $P(X, Y, Z)$ at best. Therefore, we aim to estimate upper and lower bounds $Q^+(\pi), Q^-(\pi)$ of $Q[f_0]$ and sequentially learn the policy π that minimizes $\Delta(\pi) := Q^+(\pi) - Q^-(\pi)$. In each round $t \in [T] := \{1, \dots, T\}$, we observe n i.i.d. samples from $P_{\pi_t}(X, Y, Z)$ using policy π_t , which we use (potentially together with data collected under previous policies $\pi_{<t}$) to estimate $Q^+(\pi_t), Q^-(\pi_t)$ and then aim to propose a new policy π_{t+1} that yields a smaller gap in bounds: $\Delta(\pi_{t+1}) < \Delta(\pi_t)$.

Connections to other settings. Besides the ‘indirect experimentation’ lens, the setting in eq. (1) can be viewed from different perspectives. While they all share the same mathematical formulation, they start from different conceptions of Z . The literature on *invariant causal learning* (Peters et al., 2016; Heinze-Deml et al., 2018) may interpret Z in eq. (1) as an environment indicator and f_0 is invariant across these environments. Each policy corresponds to an environment (Z cannot be designed, but we can collect data for different Z) with its own distribution of X . The heterogeneity across environments can be leveraged to learn about the invariant mechanism f . In *instrumental variables* (Pearl, 2009; Angrist & Pischke, 2008), one typically assumes to encounter a situation as in eq. (1) with the assumptions (1)-(3) where X, Y, Z are observed ‘in the wild’. Valid instruments Z then sometimes allow for (partial) identification and estimation of f_0 despite the unobserved confounder U . Again, there is typically no active experimentation in IV estimation, but Z is taken ‘as is’. Finally, experimentation on Z in our setting can also be interpreted as *soft interventions* (Jaber et al., 2020; Pearl, 2009) on X , where we intervene on X setting it to follow a given distribution instead of a fixed value. Instead of allowing arbitrary soft interventions, eq. (1) restricts us to distributions $P(X | Z)\pi(Z)$ for feasible experimental policies $\pi(Z)$. Importantly, unlike proper soft interventions, we do not get rid of unobserved confounding. Ultimately, we face a sequence of IV estimation tasks. Experimental access to Z justifies the instrumental variable assumption (2) $Z \perp\!\!\!\perp U$, and indirect experiments are setup by design such that (1) $Z \not\perp\!\!\!\perp X$ holds. Assumption (3) $Z \perp\!\!\!\perp Y | X, U$ instead may be hard to justify in practice and limits the types of experimentation allowed in our framework.

2.2 Related Work

IV estimation. Instrumental variables are a cornerstone of cause-effect estimation in econometrics (Angrist & Pischke, 2008) at least since their use by Philip G. Wright in 1928 (Wright, 1928; Stock & Trebbi, 2003). Even though valid instruments are hard to come by in practice (Hernán & Robins, 2006), generalizations to settings with relaxed assumptions have sparked renewed interest in IV estimation from the machine learning community not least due to successful applications in Mendelian randomization (Sanderson et al., 2022; Didelez & Sheehan, 2007; Legault et al., 2024) or on microbiome data (Sohn & Li, 2019; Wang et al., 2020; Ailer et al., 2021). In particular, we build on recent advances in nonlinear/nonparametric IV estimation (Newey & Powell, 2003; Lewis & Syrgkanis, 2018; Singh et al., 2019; Hartford et al., 2017; Saengkyongam et al., 2022; Zhang

¹ $L_2(X)$ is the L_2 -space of scalar-valued functions of X with respect to the distribution of X .

² We write $(\partial_i f)(x^*)$ for the partial derivative of f with respect to the i -th argument evaluated at x^* .

³ Here, $\mathcal{P}(\mathbb{R}^{d_z})$ denotes the space of probability distributions over $\mathbb{R}^{d_z}$.

et al., 2020; Xu et al., 2020), with a specific focus on the minimax formulation using the generalized method of moments (Liao et al., 2020; Dikkala et al., 2020; Bennett et al., 2023; Zhang et al., 2023; Liao et al., 2020; Bennett et al., 2019; Muandet et al., 2019; Bennett et al., 2022). In addition, we do not assume identifiability in line with recent approaches to partial identification and bounding causal effects in IV settings relaxing various discreteness and additivity assumptions (Kilbertus et al., 2020; Padh et al., 2022; Frauen et al., 2024; Melnychuk et al., 2024; Wang & Tchetgen Tchetgen, 2018; Gunsilius, 2019; Hu et al., 2021; Zhang & Bareinboim, 2021).

Sequential experiment design. The work closest to ours is perhaps by Ailer et al. (2023), where they consider a similar setting of sequential indirect experiment design via IV estimation. The key differences are that they only consider a fully linear setting (h, f_0 linear), aim for full identification of f_0 (even though the setting is underspecified), and only provide non-adaptive strategies, i.e., the sequence of experiments does not depend on data collected throughout the experiments. Other work on indirect experimentation either assumes no unobserved confounding (Singh, 2023) or focuses primarily on sample efficiency and variance reduction when aiming for full identification of f (Chandak et al., 2024). Adaptive learning has also been used in another context for cause effect estimation by ‘deciding what to observe’ Gupta et al. (2021). While clearly related, we highlight that most of the literature on active experimentation for causality aims at learning the causal structure (instead of estimating properties of f) from access to interventions, i.e., direct experiments (instead of more realistic indirect experiments), e.g., (von Kügelgen et al., 2019; Sverchkov & Craven, 2017; Agrawal et al., 2019; He & Geng, 2008; Gamella & Heinze-Deml, 2020; Elahi et al., 2024; Zemplenyi & Miller, 2021). Additionally, we require a strategy that is path-dependent, while in active learning the objective function remains the same over the different iterations.

3 Methodology

3.1 Minimax Instrumental Variable Estimation

For the general instrumental variable setting in eq. (1) that we face in each round, we aim to solve $\mathbb{E}[Y - f(X) | Z] = 0$. To solve this, we follow the conditional moment formulation (Bennett et al., 2019; Dikkala et al., 2020; Bennett et al., 2022, 2023), which builds on the observation that we only want to consider functions f , such that the residuals, $U_f := Y - f(X)$ are independent of the instrument Z . If two random variables are independent, $U_f \perp\!\!\!\perp Z$, then $E[g(Z)U_f] = E[g(Z)]E[U_f]$ for all test functions g . We search for f that minimize the residual while maintaining the independence constraint by solving a minimax optimization that minimizes f over a worst case g . The integral equation $\mathbb{E}[Y - f(X) | Z] = 0$ for f can be written as $Tf = r_0$, where T is a linear, bounded operator mapping $f : \mathcal{X} \rightarrow \mathbb{R}$ to $Tf := \mathbb{E}[f(X) | Z] : \mathcal{Z} \rightarrow \mathbb{R}$ and $r_0 = \mathbb{E}[Y | Z]$. This is typically an ill-posed inverse problem where both T and r_0 are not known but have to be estimated from i.i.d. data $\mathcal{D} = \{(x_i, y_i, z_i)\}_{i=1}^n$. Assuming f to come from a set of hypothesis functions $\mathcal{F} \subset L_2(X)$ and $\mathcal{G} \subset L_2(Z)$ a class of ‘witness functions’ or ‘test functions’, we can write the non-parametric IV problem as a minimax optimization problem of the form (Dikkala et al., 2020; Bennett et al., 2023)

$$\arg \inf_{f \in \mathcal{F}} \sup_{g \in \mathcal{G}} L(f, g) \quad (2)$$

for some objective function L mapping tuples of hypotheses f and test functions g to $\mathbb{R}$. While most existing approaches start by assuming that there exists a unique solution of $Tf = r_0$ (Lewis & Syrgkanis, 2018; Liao et al., 2020; Muandet et al., 2019; Dikkala et al., 2020; Bennett et al., 2019), based on Lagrange multipliers Bennett et al. (2023) recently have shown that under a source condition on T and mild realizability assumptions regarding the capacity of the (statistically restricted) function classes $\mathcal{F}$ and $\mathcal{G}$ (Bennett et al., 2023, Assumptions 2, 3, 4), the solution of $Tf = r_0$ is given by

$$f_0 = \arg \min_{f \in \mathcal{F}} \max_{g \in \mathcal{G}} \frac{1}{2} \|f\|_{L_2(X)}^2 + \langle r_0 - Tf, g \rangle_{L_2(Z)} = \arg \min_{f \in \mathcal{F}} \max_{g \in \mathcal{G}} \frac{1}{2} \mathbb{E}[f^2(X)] + \mathbb{E}[(Y - f(X))g(Z)]. \quad (3)$$

This formulation of the objective is immediately amenable to empirical estimation from $\mathcal{D}$ by using empirical averages for the expectations. Eq. (3) can be viewed as a penalized optimization, where $\frac{1}{2} \mathbb{E}[f^2(X)]$ penalizes the solution towards the minimum L_2 norm.

The inner maximization over test functions ensures that we only consider hypotheses $f \in \mathcal{F}$ compatible with the conditional moment restrictions that enforce the required independence of the residuals. While there may still be multiple hypotheses $f \in \mathcal{F}$ compatible with observational data in

this way, [Bennett et al.](#) pick the f with the least L_2 norm, yielding a unique solution. In the following, we build on this intuition to estimate bounds on a chosen scientific query Q among all f compatible with the observed data and assumed structure (via the moment restrictions).

3.2 Sequential Policy Learning for Bounding Functionals

The source condition required for eq. (3) is that r_0 lies in the range of TT^* , where T^* is the adjoint of T . As [Bennett et al.](#) point out, this is a nontrivial restriction and paramount in ensuring uniqueness in eq. (3). Without this source condition the IV problem is still ill-posed in that there may be multiple solutions $f \in \mathcal{F}$ compatible with the observed data $\mathcal{D}$ (even in the infinite data limit). One of our key realizations is that we can still meaningfully make use of the penalized optimization formulation in eq. (3) to compute bounds on targeted queries about f .

Let $Q : \mathcal{F} \rightarrow \mathbb{R}$ be a functional on hypotheses $\mathcal{F}$ that captures the scientific query. For any given joint distribution $P_\pi(X, Y, Z)$ induced by experimental policy π , to bound $Q[f_0]$, we can compute

$$Q^\pm(\pi) = \inf_{f \in \mathcal{F}} \sup_{g \in \mathcal{G}} \pm Q[f] + \mathbb{E}[(Y - f(X))g(Z)], \quad (4)$$

where again the inner maximization aims at only considering hypotheses compatible with the conditional moment restrictions $\mathcal{F}_c \subset \mathcal{F}$ and the outer minimization aims at finding the $f \in \mathcal{F}_c$ with the largest/smallest value of $Q[f]$. The realizability assumption that $f_0 \in \mathcal{F}$ then clearly implies that $f_0 \in \mathcal{F}_c$ and the bounds $Q^\pm$ will contain $Q[f_0]$. We note that specially when X is of higher dimension than Z , in practice neither f nor the much more narrow query $Q[f]$ are guaranteed to be identified. However, depending on π , we may shape P_π such that $\mathcal{F}_c$ becomes restricted enough to obtain informative bounds $Q^\pm$ on $Q[f_0]$. [Bennett et al. \(2022\)](#) work out in detail when a linear functional of f_0 , but not f_0 itself, is identified. If P_π is such that T satisfies the source condition, certain queries Q (such as the $Q[f] = \frac{1}{2} \|f\|_{L_2}^2$ as a canonical example) may actually be identified and the bounds $Q^\pm$ will coincide.

We can now formulate our main policy learning goal as

$$\inf_{\pi \in \Pi(\mathcal{Z})} \Delta(\pi), \quad \text{where } \Delta(\pi) = Q^+(\pi) - Q^-(\pi), \quad (5)$$

with $\Pi(\mathcal{Z}) \subset \mathcal{P}(\mathcal{Z})$ being a subset of implementable experimentation policies. In practice, we aim at approximating this optimization via sequential updates of π_t over T rounds, where we observe n i.i.d. samples of P_{π_t} at each round. This means that given π_t , we seek to choose π_{t+1} such that $\Delta(\pi_{t+1}) < \Delta(\pi_t)$. The final output of our method is the interval $[Q^-(\pi_T), Q^+(\pi_T)] \ni Q[f_0]$. A key advantage of this formulation is that we need not assume identifiability of f nor of $Q[f]$, will obtain valid (albeit potentially loose) bounds when the experimentation budget T is restricted, and that potential identifiability will be ‘captured automatically’ by the bounds coinciding. If after T rounds the bounds are still not informative, the experimenter can choose more expressive experiments (different $\Pi(\mathcal{Z})$ and even different $\mathcal{Z}$), a more specific query Q , or put stronger assumptions on the hypothesis class $\mathcal{F}$.

3.3 Solving the Minimax Optimization Problem

Let us unpack the sequential policy learning problem. At each round $t \in [T]$, we need to solve two minimax problems for $Q^\pm(\pi_t)$ where the objective is estimated from finite data and then take a minimization step over the difference of the two solutions. To further complicate this ‘tri-level’ (min + minimax) optimization problem, we are optimizing over two function spaces $(\mathcal{F}, \mathcal{G})$ and a family of distributions $(\Pi(\mathcal{Z}))$. In the remainder of this section we develop techniques to render these optimizations feasible in practice for suitable choices of $\mathcal{F}, \mathcal{G}, \Pi(\mathcal{Z})$.

Existing solutions for the minimax problems typically employ adversarial optimization (iterative gradient-based optimization over parametric function spaces) or kernel-based techniques. Since the former are usually hard and expensive to optimize reliably even without the outer minimization, we focus on the latter approach using reproducing kernel Hilbert spaces (RKHS) for test functions $\mathcal{G} = \mathcal{H}_Z$ induced by a characteristic kernel $k_Z : \mathcal{Z} \times \mathcal{Z} \rightarrow \mathbb{R}$. Then there exists a closed form expression for a consistent estimate of the inner supremum.

Theorem 1 (Zhang et al. (2020)). If $\mathbb{E}[(Y - f(X))^2 k_Z(Z, Z)] < \infty$,

$$\left(\frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n (y_i - f(x_i)) k(z_i, z_j) (y_j - f(x_j)) \right)^{\frac{1}{2}} \quad (6)$$

consistently estimates $\sup_{g \in \mathcal{H}_Z, \|g\| \leq 1} \langle r_0 - Tf, g \rangle_{\mathcal{H}_Z}$ from data $\mathcal{D}$.

Hence, for n samples $\mathcal{D} \stackrel{\text{iid}}{\sim} P_\pi$ we reduce the minimax optimizations in eq. (4) to

$$Q^\pm(\pi) = \inf_{f \in \mathcal{F}} \mp Q[f] + \left(\frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n (y_i - f(x_i)) k(z_i, z_j) (y_j - f(x_j)) \right)^{\frac{1}{2}}. \quad (7)$$

One option to solve eq. (7) efficiently is via gradient-based methods leveraging highly efficient auto-differentiation packages for parametric function families, such as choosing neural networks for $\mathcal{F}$, i.e., we optimize over a finite-dimensional real-vector θ —the weights of the neural network. Our framework may also benefit from other existing approaches to the minimax optimization from the literature, e.g., in Muandet et al. (2019); Dikkala et al. (2020). Parametric functions may have the advantage that they render the computation of $Q[f_\theta]$ feasible. For example, the causal effect of X_i on Y at x^* given by $Q[f_\theta] = (\partial_i f_\theta)(x^*)$ can be obtained directly from automatic differentiation as well. Global functionals such as $Q[f_\theta] = \int_{\mathcal{X}} \psi(x) f_\theta(x) dx$ for some function ψ may pose greater difficulties as the integral over $\mathcal{X}$ is typically infeasible. Still, this approach ultimately requires a bi-level optimization with a substantial amount of tuning.

In order to also obtain a closed-form estimate of the minimization over $\mathcal{F}$, we have to specify the functional Q . To this end, we focus on bounded linear functionals. One example is the causal-effect of an individual input $(\partial_i f)(x^*)$, because it is arguably one of the most common and relevant scientific queries: how does a small change of a specific variable affect the outcome? When following a fully non-parametric approach by also assuming hypotheses to come from an RKHS $\mathcal{F} = \mathcal{H}_X$ induced by a characteristic kernel $k_X : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$, we can also estimate the minimization over $\mathcal{F}$ in closed form for the causal effect—a local, linear functional. The key realization is that functionals of functions represented as linear combinations of RKHS functions can be written as linear functions of the coefficients.

Theorem 2. Assume functions in $\mathcal{H}_Z, \mathcal{H}_X$ to have bounded variation, (w.l.o.g.) images contained in $[-1, 1]$, and for $h \in \mathcal{H}_Z$ also $-h \in \mathcal{H}_Z$. Denote by $K_{XX}, K_{ZZ} \in \mathbb{R}^{n \times n}$ the empirical kernel matrices of $\mathcal{D}$, let k_X be continuously differentiable, and let $\lambda_g, \lambda_f, \lambda_c \geq 0$. Further, fix $Q[f]$ to be a bounded linear functional, and write $f_\theta(\cdot) = \sum_{i=1}^n \theta_i k(x_i, \cdot)$. Then the solutions to the regularized minimax problems

$$f^\pm = \arg \min_{f \in \mathcal{H}_X} \mp Q[f] + \lambda_c \sup_{g \in \mathcal{H}_Z} \langle r_0 - Tf, g \rangle_{\mathcal{H}_Z} + \lambda_f \|f\|_{\mathcal{H}_X} - \lambda_g \|g\|_{\mathcal{H}_Z} \quad (8)$$

are consistently estimated by $f_{\hat{\theta}^\pm}$ with

$$\hat{\theta}^\pm = (K_{XX} K_{ZZ} K_{XX} + 4 \frac{\lambda_g \lambda_f}{\lambda_c} K_{XX})^+ \left(K_{XX} K_{ZZ} y \mp \frac{1}{\lambda_c} Q[K_X \cdot](x^*) \right), \quad (9)$$

where $(Q[K_X \cdot])(x^*) \in \mathbb{R}^n$ is the vector $((Q[k_X(x_1, \cdot)])(x^*), \dots, (Q[k_X(x_n, \cdot)])(x^*))$. Note that these general bounded linear functionals can be computed analytically for many common kernels such as linear, polynomial, or radial basis function (RBF) kernels.

One common example is the partial derivative at x^* :

$$\hat{\theta}^\pm = (K_{XX} K_{ZZ} K_{XX} + 4 \frac{\lambda_g \lambda_f}{\lambda_c} K_{XX})^+ \left(K_{XX} K_{ZZ} y \mp \frac{1}{\lambda_c} (\partial_i K_X \cdot)(x^*) \right), \quad (10)$$

where $(\partial_i K_X \cdot)(x^*) \in \mathbb{R}^n$ is the vector $((\partial_i k_X(x_1, \cdot))(x^*), \dots, (\partial_i k_X(x_n, \cdot))(x^*))$.

We defer all proofs of our theoretical statements to Appendix A.

We can now substitute these closed-form estimates for the minimax problems into our policy learning objective $\Delta(\pi) = Q[f_{\hat{\theta}^+}] - Q[f_{\hat{\theta}^-}]$ in eq. (5)

$$\Delta(\pi) = -\frac{2}{\lambda_c} \left((K_{XX} K_{ZZ} K_{XX} + 4 \frac{\lambda_g \lambda_f}{\lambda_c} K_{XX})^+ (Q[K_X \cdot])(x^*) \right)^\top (Q[K_X \cdot])(x^*). \quad (11)$$

(In)validity of bounds. Our general policy learning formulation in eqs. (4) and (5) involves nonlinear, nonconvex optimizations over function spaces and families of distributions. Since we cannot provably obtain global optima to those, we cannot guarantee valid bounds. Even for the RKHS setting in Theorem 2 for linear Q , while the estimates $\hat{\theta}^\pm$ are consistent, for finite sample estimation we have to regularize both the hypotheses and witness functions ($\lambda_g, \lambda_f > 0$) to obtain numerically stable and reliable estimates. Clearly, those pose restrictions on the search spaces for f, g and may thus lead to invalid bounds. We may expect real-world mechanisms f_0 to be relatively smooth such that these regularization terms do not render our estimates invalid in practice. Moreover, since we made no explicit assumptions about the functional Q , using $Q[f]$ as a penalty (as compared to, e.g., $\frac{1}{2}\|f\|_{L_2}$) may not yield a unique solution to eq. (8) even in the infinite data limit and for $\lambda_g = \lambda_f = 0$. While uniqueness can be recovered for strictly convex, coercive, and lower semicontinuous Q , this implies non-trivial restrictions on the allowed scientific queries, which we may not be willing to tolerate. Therefore, we introduce a ‘penalization parameter’ λ_c (in front of the supremum term rather than $Q[f]$ for convenience) that allows us to empirically trade off the scale/importance of $Q[f]$ and the conditional moment restriction in practice. Enforcing conditional moment restrictions via a minimax formulation from finite samples typically exhibits high variance. Therefore, a lack of provably valid bounds due to practical regularization does not impair the usefulness of our method in a setting where we aim to learn about a potentially unidentified scientific query from limited experimental access. In particular, we argue that overly conservative bounds are still more useful than likely invalid point estimates based on one-size-fits-all assumptions required for identifiability.

Continuing with the closed-form expression in eq. (11) for the gap between maximally and minimally possible values of the scientific query – here the causal effect of changing one treatment component around a base level X^* – over all hypotheses that are compatible with the observed data, we can now seek to represent the experimentation strategy π in a way that allows us to sequentially improve it.

3.4 Sequential Experiment Selection

We now describe different approaches to sequentially update the experimentation policy. The experimentation budget T and the number of samples per round n are fixed. We denote by π_t the policy in round t and by $\mathcal{D}^{(t)}$ ($\mathcal{D}^{(\leq t)}$) the data collected in (up to) round t . We denote by $Q_t^\pm$ and Δ_t the bounds and gap between them, respectively, estimated in round t and explicitly mention which data was used for the estimates. Some strategies rely on a ‘meta-distribution’ over policies $\Pi(\mathcal{Z})$, which we denote by Γ . A sample $\pi \sim \Gamma$ is thus a policy in $\Pi(\mathcal{Z}) \subset \mathcal{P}(\mathcal{Z})$. While the ‘final result’ of all strategies are just the final bounds $Q_T^\pm$, we report $Q_t^\pm$ also for suitable $t < T$ for different strategies to compare efficiency.

Simple baseline. In modern methods for experiment design, random exploration is often surprisingly competitive (Ailer et al., 2023). We consider a simple entirely non-adaptive baseline called **random**, which independently samples $\pi_t \sim \Gamma$ for $t \in [T]$ and estimates $Q_t^\pm$ on $\mathcal{D}^{(\leq t)}$.

Locally guided strategies. Next, we look to simple adaptive strategies that leverage the locality of Q . If Q is determined in a small neighborhood around some $x^* \in \mathcal{X}$, a useful guiding principle for π is to aim at concentrating the mass of the marginal $P_\pi(X)$ around x^* . The causal effect $(\partial_i f)(x^*)$ or the average treatment effect $\mathbb{E}[Y | do(X = x^*)]$ are examples of local queries. Estimating such local relationships in indirect experimentation without unobserved confounding has recently been studied from a theoretical perspective by Singh (2023). They propose a simple explore-then-exploit strategy and prove favorable minimax convergence rates depending on the complexities of h and f_0 as well as the noise levels in the unconfounded setting, i.e., X and Y have independent noises and h does not depend on U . We highlight that these strategies do not use the intermediate estimated bounds as feedback to select the next policy π . Pseudocode for the different strategies is in Appendix G.

1. **Explore then exploit (EE)** (inspired by Singh (2023)): Split the budget into $T = T_1 + T_2$. Sample $\pi_t \sim \Gamma$ for $t \in [T_1]$. Then, fit π as a (parametric) distribution on the Z values of the $K \leq T_1 \cdot n$ samples in $\mathcal{D}^{(\leq T_1)}$ closest to x^* (i.e., their X values have the smallest distance to x^* for some suitable distance measure on $\mathcal{X}$). Set $\pi_t := \pi$ for $t \in \{T_1 + 1, \dots, T\}$ and estimate $Q_t^\pm$ from $\bigcup_{t=T_1+1}^T \mathcal{D}^{(t)}$. The split between T_1 and T_2 may be informed by n to maximize exploration while retaining sufficiently many samples during exploitation for a low variance estimate of $Q_T^\pm$.
2. **Alternating explore exploit (AEE)**: Throughout all rounds, keep a list of all observed samples sorted by the distance of their X values to x^* . For odd $t \in [T] \setminus \{T\}$ independently sample

$\pi_t \sim \Gamma$. For even $t \in [T]$ and $t = T$ fit π_t as a (parametric) distribution to the $\min\{K, n \cdot t\}$ nearest samples to x^* for some $K \in \mathbb{N}$ and estimate $Q_t^\pm$ on $\bigcup_{\text{even } t} \mathcal{D}^{(t)}$.

There is some freedom in which data is used to estimate $Q_t^\pm$ and using $\mathcal{D}^{(\leq t)}$ is always a valid choice. For local queries it can still be beneficial in practice to discard data far away from x^* . More broadly, the above strategies could potentially be further improved by weighing samples in $\mathcal{D}^{(t)}$ inversely proportional to the distance of (the mean of) π_t from x^* when estimating $Q_t^\pm$ from $\mathcal{D}^{(\leq t)}$. We only report the vanilla versions described above in our experiments and leave further variance reduction techniques for future work. In all experiments, we use $K = n$ and the Euclidean distance on $\mathcal{X}$.

While these locally guided strategies appear rather limited in which information from previous rounds is used in updating the policy, [Singh et al. \(2019, Sec. 7\)](#) make a convincing case for why the simple types of active learning in EE and AEE greatly improve over the naive random strategy and may be competitive with fully adaptive experiments. For the fixed policy in π and the meta-distribution Γ one should arguably err on the side of ‘uninformative’ priors, i.e., covering all of $\mathcal{Z}$ mostly uniformly.

Targeted Adaptive Strategy. We now develop an adaptive strategy denoted by **adaptive** that actually uses the intermediate estimates of $Q_t^\pm$ to update π_t . A natural choice for policy updates is via gradient descent on our objective Δ

$$\pi_{t+1} = \pi_t - \alpha_t \nabla_\pi \Delta(\pi)|_{\pi=\pi_t}, \quad (12)$$

for step sizes $\alpha_t > 0$ at round t . In practice, we assume parametric policies $\pi_t := \pi_{\phi_t}$ with parameters $\phi_t \in \Phi \subset \mathbb{R}^d$. With the log-derivative trick ([Williams, 1992](#)) we then compute

$$\nabla_\phi \Delta(\pi_\phi) = \mathbb{E}_{X,Y,Z \sim P_{\pi_\phi}(X,Y,Z)} [\Delta(\pi_\phi) \nabla_\phi \log \pi_\phi]. \quad (13)$$

and can update π_{ϕ_t} akin to the REINFORCE algorithm with horizon one. Since the gradient in eq. (12) is evaluated at π_{ϕ_t} from which we have samples available, we can directly use the empirical counterpart of the expectation in eq. (13). One caveat is that our objective Δ cannot be computed on the ‘instance level’ for individual samples (x_i, y_i, z_i) , but is already an estimate based on multiple samples. In practice, we therefore split the n observations per round into random batches, and average over the $\Delta(\pi_\phi)$ estimate for each batch times the mean score function across the batch in eq. (13).

An extension to use data $\mathcal{D}^{(\tau)}$ from a previous round $\tau < t$ is to reweigh the term within the expectation in eq. (13) by $\pi_{\phi_t}(Z)/\pi_{\phi_\tau}(Z)$ (and then average across rounds). While this allows us to use more data for the gradient estimates, this approach does not yield an unbiased estimator and typically also increases variance due to arbitrarily small reweighing terms. In practice, we may want to only consider data from a small number of previous rounds, where we expect the policy not to have changed too much. Again, we believe our technique could benefit from further variance reduction techniques ([Chandak et al., 2024](#)), but leave these for future work.

While our formulation in eqs. (12) and (13) can be efficiently implemented for any parametric choice of π_ϕ , in our experiments we choose a multivariate Gaussian mixture model (GMM) $\pi_\phi(z) = \sum_{m=1}^M \gamma_m \mathcal{N}(z; \mu_m, \Sigma_m)$ with weights $\gamma \in [0, 1]^M$ with $\sum_{m=1}^M \gamma_m = 1$, means $\mu_m \in \mathbb{R}^{d_z}$, and covariances $\Sigma \in \mathbb{R}^{d_z \times d_z}$ for $i \in [M]$. Hence, the parameters are $\phi = (\gamma, \mu_1, \dots, \mu_M, \Sigma_1, \dots, \Sigma_M)$. With the score function $\nabla_\phi \log \pi_\phi$ known analytically for GMMs, eq. (13) can be estimated efficiently (see Appendix B for details).

4 Experiments

Setup. We consider a low-dimensional setting (for visualization purpose) with $d_x = d_z = 2$ and

$$h_j(Z, U) = \alpha \cdot (\sin(Z_j)(1 + U)), \quad f(X) = \beta \sum_{j=1}^{d_x} \exp(X_j) \sin(X_j) + U, \quad U \sim \mathcal{N}(0, 1), \quad (14)$$

where we use $\alpha = \beta = 20$. The local point of interest is $x^* = 0 \in \mathbb{R}^{d_x}$ and we easily check that $\partial_i f(x^*) = \beta$. We use $n = 250$ samples in each round over a total of $T = 16$ experiments.

Method parameters. We use radial basis function (RBF) kernels $k(x_1, x_2) = \exp(\rho \|x_1 - x_2\|^2)$ with a fixed $\rho = 1$. Note, that the three hyperparameters are relative weights. Thus we set $\lambda_s := \frac{\lambda_g \lambda_f}{\lambda_c} = 0.01$ and $\lambda_c = 0.04$, we refer to the Appendix C for a further comparison of different

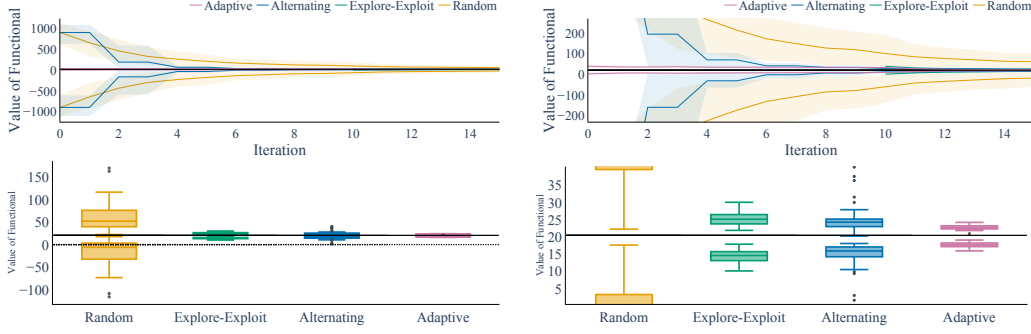


Figure 1: We compare the different strategies in our synthetic setting. Left and right only differ in the range of the y-axis. The black constant line represents the true value of $Q[f_0]$. **Top:** Estimated upper and lower bounds $Q_t^\pm$ over $t \in [T]$ for $n_{\text{seeds}} = 50$ and two different ‘zoom levels’ on the y-axis. Lines are means and shaded regions are (10, 90)-percentiles. **Bottom:** The final estimated bounds $Q_T^\pm$ at $T = 16$. The dotted line is $y = 0$. Both locally guided heuristic (explore-then-exploit, alternating explore exploit) confidently bound the target query away from zero with a relatively narrow gap between them. Our targeted adaptive strategy is even better and essentially identifies the target query $Q[f_0] = 20$ after $T = 16$ rounds.

hyperparameter choices. For all strategies the variance of our policy is set at $\sigma_e = 0.001$. For explore then exploit, we chose $T_1 = 10, T_2 = 6$ with $\pi_t = \mathcal{N}(\mu_t, \sigma_e \text{Id}_{d_z})$ and independent $\mu_t \sim \mathcal{N}(\mathbf{0}_{d_z}, \text{Id}_{d_z})$ for $t \in [T_1]$. The Gaussian mixture of the adaptive strategy is initialized with $M = 3, \gamma = (1/3, 1/3, 1/3), \Sigma_m = \text{Id}_{d_z}$ and independent $\mu_m \sim \mathcal{N}(\mathbf{0}_{d_z}, \text{Id}_{d_z})$. We use a constant learning rate $\alpha_t = 0.01$ for all t and restrict ourselves to learning only weights γ_m , means μ_m and the diagonal entries of the covariances Σ_m .

Results. We perform $n_{\text{seeds}} = 50$ runs for different random seeds and report means with the 10- and 90-percentiles in Figure 1. The simple non-adaptive random baseline starts out with poor performance as expected, and improves mildly over multiple rounds. We note that the ultimate performance of the random baseline depends primarily on how much mass the random exploration (defined by the meta-distribution Γ) puts near x^* , i.e., whether we collect sufficiently many samples near x^* over the T rounds. Explore then exploit performs quite well as soon as we start exploitation. Again, whether sufficient informative examples have been collected during the exploration stage depends on the choice of Γ . However, explore then exploit does outperform random (with the same exploration distribution) indicating that the heuristic of focusing on the informative samples (the ones near x^*) does provide substantial improvements (as also found by (Singh, 2023)). The alternating strategy AEE provides bounds across all rounds and clearly shows step-wise improvement from the first round on—ultimately performing similar to EE. Finally, the adaptive strategy quickly narrows the bounds and achieves a fairly narrow gap Δ already after few rounds. At $T = 16$ it has essentially identified the target query $Q[f_0] = 20$ with low variance across seeds. For completeness, we provide a runtime comparison of the different methods in Appendix D. The code to reproduce results is available at <https://github.com/EAILer/targeted-iv-experiments>.

5 Discussion and Conclusion

Summary. We formalized designing optimal experiments to learn about a scientific query as sequential instrument design with the goal of minimizing the gap between estimated upper and lower bounds of a target functional Q of the ground truth mechanism f_0 . We only assume indirect experimental access, allow for unobserved confounding, and consider nonlinear f_0 with multi-variate inputs such that both f_0 and $Q[f_0]$ may be unidentified. For a broad set of queries, we derive closed form estimators for the bounds within each round of experimentation when f lies within an RKHS. Based on these estimates, we then develop adaptive strategies to sequentially narrow the bounds on the scientific query of interest and demonstrate the efficacy of our method in synthetic experiments. Given the increased amount of data collected in fully or partially automated labs, for example for drug

discovery, we believe that efficient, adaptive experiment design strategies will be a vital component of data-driven scientific discovery.

Limitations and future work. We have not validated our method in a real-world setting, as it would require access (and full control) over an actual experimental setup as well as the time and resources required to conduct these experiments. A key technical limitation of our work is that neither our general policy learning formulation in eqs. (4) and (5) nor the concrete closed form estimates in Theorem 2 obtain provably valid bounds for all queries Q from finite samples (see discussion right before Section 3.4). While our approach is reliable in synthetic experiments, we believe thorough theoretical analysis of necessary and sufficient conditions for (asymptotically) valid bounds (including ideally asymptotic normality with known asymptotic covariance or even finite sample guarantees) is a useful direction for future work beyond the scope of our current manuscript. In practice, the applicability of our approach may also be limited by the assumptions that the underlying system is well described via a fixed, static function $f_0 : \mathcal{X} \rightarrow \mathcal{Y}$ as opposed to, say, a temporally evolving and interacting systems governed by a differential equation. Similarly, while IV assumptions (1) and (2) can arguably be justified in our setting, the third assumption $Z \perp\!\!\!\perp Y \mid X, U$ limits the types of experimentation we can consider (see discussion at the end of Section 1).

Methodologically, we believe that our approach could benefit from advanced variance reduction techniques and be sped up by more efficient estimators (Chandak et al., 2024). Along these lines, it is worthwhile future work to analyze the optimal trade-off between exploration and retaining a large number of samples for exploitation and thus lower variance estimates. Finally, our current empirical validation is limited to linear local functionals, rather simple parametric choices for policies (such as (mixtures of) Gaussians), and we have not optimized the kernel choices and hyperparameters. Hence, a validation where all these choices have been tailored to a real-world application scenario is an important direction for future work.

Acknowledgments and Disclosure of Funding

EA is supported by the Helmholtz Association under the joint research school “Munich School for Data Science – MUDS”. This work has been supported by the Helmholtz Association’s Initiative and Networking Fund through the Helmholtz international lab “CausalCellDynamics” (grant # Interlabs-0029).

References

- Agrawal, R., Squires, C., and Yang, K. Abcd-strategy: Budgeted experimental design for targeted causal structure discovery. In *Proceedings of the 22nd International Conference on Artificial Intelligence and Statistics*, pp. 2874–2882. PMLR, 2019. 4
- Ailer, E., Müller, C. L., and Kilbertus, N. A causal view on compositional data. *arXiv preprint arXiv:2106.11234*, 2021. 3
- Ailer, E., Hartford, J., and Kilbertus, N. Sequential underspecified instrument selection for cause-effect estimation. In *Proceedings of the 40th International Conference on Machine Learning, ICML’23*. JMLR.org, 2023. 4, 7, 17
- Andy, Y., Morris, J., and Mark, G. *Slurm: Simple linux utility for resource management*. Workshop on job scheduling strategies for parallel processing, 2003. 24
- Angrist, J. D. and Pischke, J.-S. *Mostly harmless econometrics: An empiricist’s companion*. Princeton university press, 2008. 3
- Bennett, A., Kallus, N., and Schnabel, T. Deep generalized method of moments for instrumental variable analysis. In Wallach, H., Larochelle, H., Beygelzimer, A., d’Alché-Buc, F., Fox, E., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. URL <https://proceedings.neurips.cc/paper/2019/file/15d185eaa7c954e77f5343d941e25fbd-Paper.pdf>. 4
- Bennett, A., Kallus, N., Mao, X., Newey, W., Syrgkanis, V., and Uehara, M. Inference on strongly identified functionals of weakly identified functions. *arXiv preprint arXiv:2208.08291*, 2022. 2, 4, 5

- Bennett, A., Kallus, N., Mao, X., Newey, W., Syrgkanis, V., and Uehara, M. Minimax instrumental variable regression and l_2 convergence guarantees without identification or closedness. In *The Thirty Sixth Annual Conference on Learning Theory*, pp. 2291–2318. PMLR, 2023. 4, 5
- Bradbury, J., Frostig, R., Hawkins, P., Johnson, M. J., Leary, C., Maclaurin, D., Necula, G., Paszke, A., VanderPlas, J., Wanderman-Milne, S., and Zhang, Q. JAX: Composable transformations of Python+NumPy programs, 2018. 24
- Chandak, Y., Shankar, S., Syrgkanis, V., and Brunskill, E. Adaptive instrument design for indirect experiments. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=4Zz5UELkIt>. 4, 8, 10
- Didelez, V. and Sheehan, N. Mendelian randomization as an instrumental variable approach to causal inference. *Statistical Methods in Medical Research*, 16(4):309–330, 2007. 3
- Dikkala, N., Lewis, G., Mackey, L., and Syrgkanis, V. Minimax estimation of conditional moment models. *Advances in Neural Information Processing Systems*, 33:12248–12262, 2020. 4, 6, 14
- Elahi, M. Q., Wei, L., Kocaoglu, M., and Ghasemi, M. Adaptive online experimental design for causal discovery, 2024. 4
- Frauen, D., Melnychuk, V., and Feuerriegel, S. Sharp bounds for generalized causal sensitivity analysis. *Advances in Neural Information Processing Systems*, 36, 2024. 4
- Gamella, J. and Heinze-Deml, C. Active invariant causal prediction: Experiment selection through stability. In *Advances in Neural Information Processing Systems*, 2020. 4
- Gunsilius, F. A path-sampling method to partially identify causal effects in instrumental variable models. *arXiv preprint arXiv:1910.09502*, 2019. 4
- Gupta, S., Lipton, Z. C., and Childers, D. Efficient Online Estimation of Causal Effects by Deciding What to Observe. Papers 2108.09265, arXiv.org, August 2021. URL <https://ideas.repec.org/p/arx/papers/2108.09265.html>. 4
- Harris, C. R., Millman, K. J., van der Walt, S. J., Gommers, R., Virtanen, P., Cournapeau, D., Wieser, E., Taylor, J., Berg, S., Smith, N. J., Kern, R., Picus, M., Hoyer, S., van Kerkwijk, M. H., Brett, M., Haldane, A., del Río, J. F., Wiebe, M., Peterson, P., Gérard-Marchant, P., Sheppard, K., Reddy, T., Weckesser, W., Abbasi, H., Gohlke, C., and Oliphant, T. E. Array programming with NumPy. *Nature*, 2020. 24
- Hartford, J., Lewis, G., Leyton-Brown, K., and Taddy, M. Deep iv: A flexible approach for counterfactual prediction. In *International Conference on Machine Learning*, pp. 1414–1423, 2017. 3
- He, Y. and Geng, Z. Active learning of causal networks with intervention experiments and optimal designs. *Journal of Machine Learning Research*, 9:2523–2547, 2008. 4
- Heinze-Deml, C., Peters, J., and Meinshausen, N. Invariant causal prediction for nonlinear models. *Journal of Causal Inference*, 6(2):20170016, 2018. 3
- Hernán, M. A. and Robins, J. M. Instruments for causal inference: an epidemiologist’s dream? *Epidemiology*, pp. 360–372, 2006. 3
- Hu, Y., Wu, Y., Zhang, L., and Wu, X. A generative adversarial framework for bounding confounded causal effects. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pp. 12104–12112, 2021. 4
- Hunter, J. D. Matplotlib: A 2D graphics environment. *Computing in Science & Engineering*, 2007. 24
- Jaber, A., Kocaoglu, M., Shanmugam, K., and Bareinboim, E. Causal discovery from soft interventions with unknown targets: Characterization and learning. *Advances in neural information processing systems*, 33:9551–9561, 2020. 3

- Kilbertus, N., Kusner, M. J., and Silva, R. A class of algorithms for general instrumental variable models. In *Advances in Neural Information Processing Systems*, volume 33, 2020. 4
- Kluyver, T., Ragan-Kelley, B., Pérez, F., Granger, B., Bussonnier, M., Frederic, J., Kelley, K., Hamrick, J., Grout, J., Corlay, S., Ivanov, P., Avila, D., Abdalla, S., Willing, C., and Jupyter Development Team. Jupyter Notebooks—a publishing format for reproducible computational workflows. In *IOS Press*, pp. 87–90. IOS Press, 2016. doi: 10.3233/978-1-61499-649-1-87. 24
- Legault, M.-A., Hartford, J., Arsenault, B. J., Yang, A. Y., and Pineau, J. A novel and efficient machine learning mendelian randomization estimator applied to predict the safety and efficacy of sclerostin inhibition. *medRxiv*, 2024. doi: 10.1101/2024.01.30.24302021. URL <https://www.medrxiv.org/content/early/2024/01/31/2024.01.30.24302021>. 3
- Lewis, G. and Syrgkanis, V. Adversarial generalized method of moments. *arXiv preprint arXiv:1803.07164*, 2018. 3, 4
- Liao, L., Chen, Y.-L., Yang, Z., Dai, B., Kolar, M., and Wang, Z. Provably efficient neural estimation of structural equation models: An adversarial approach. *Advances in Neural Information Processing Systems*, 33:8947–8958, 2020. 4
- Melnychuk, V., Frauen, D., and Feuerriegel, S. Partial counterfactual identification of continuous outcomes with a curvature sensitivity model. *Advances in Neural Information Processing Systems*, 36, 2024. 4
- Muandet, K., Mehrjou, A., Lee, S. K., and Raj, A. Dual instrumental variable regression. *arXiv preprint arXiv:1910.12358*, 2019. 4, 6
- Newey, W. K. and Powell, J. L. Instrumental variable estimation of nonparametric models. *Econometrica*, 71(5):1565–1578, 2003. 3
- Padh, K., Zeitler, J., Watson, D., Kusner, M., Silva, R., and Kilbertus, N. Stochastic causal programming for bounding treatment effects, 2022. URL <https://arxiv.org/abs/2202.10806>. 4
- pandas development team. Pandas, 2020. URL <https://doi.org/65910.5281/zenodo.3509134>. 24
- Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., Killeen, T., Lin, Z., Gimelshein, N., Antiga, L., Desmaison, A., Köpf, A., Yang, E., DeVito, Z., Raison, M., Tejani, A., Chilamkurthy, S., Steiner, B., Fang, L., Bai, J., and Chintala, S. Pytorch: An imperative style, high-performance deep learning library, 2019. 24
- Pearl, J. *Causality*. Cambridge university press, 2009. 3
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M., and Duchesnay, É. Scikit-learn: Machine learning in Python. *JMLR*, 2011. 24
- Peters, J., Bühlmann, P., and Meinshausen, N. Causal inference by using invariant prediction: identification and confidence intervals. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 78(5):947–1012, 2016. 3
- Petersen, K. B. and Pedersen, M. S. The matrix cookbook. *Technical University of Denmark*, 7(15): 510, 2008. 15
- Saengkyongam, S., Henckel, L., Pfister, N., and Peters, J. Exploiting independent instruments: Identification and distribution generalization. *arXiv preprint arXiv:2202.01864*, 2022. 3
- Sanderson, E., Glymour, M. M., Holmes, M. V., Kang, H., Morrison, J., Munafò, M. R., Palmer, T., Schooling, C. M., Wallace, C., Zhao, Q., and Davey Smith, G. Mendelian randomization. *Nature Reviews Methods Primers*, 2(1):6, 2022. 3
- Severini, T. A. and Tripathi, G. Some identification issues in nonparametric linear models with endogenous regressors. *Econometric Theory*, 22(2):258–278, 2006. doi: 10.1017/S0266466606060117. 18

- Severini, T. A. and Tripathi, G. Efficiency bounds for estimating linear functionals of nonparametric regression models with endogenous regressors. *Journal of Econometrics*, 170(2):491–498, 2012. ISSN 0304-4076. doi: <https://doi.org/10.1016/j.jeconom.2012.05.018>. URL <https://www.sciencedirect.com/science/article/pii/S0304407612001303>. Thirtieth Anniversary of Generalized Method of Moments. 18
- Singh, R., Sahani, M., and Gretton, A. Kernel instrumental variable regression. In *Advances in Neural Information Processing Systems*, pp. 4593–4605, 2019. 3, 8
- Singh, S. Nonparametric indirect active learning. In Ruiz, F., Dy, J., and van de Meent, J.-W. (eds.), *Proceedings of The 26th International Conference on Artificial Intelligence and Statistics*, volume 206 of *Proceedings of Machine Learning Research*, pp. 2515–2541. PMLR, 25–27 Apr 2023. URL <https://proceedings.mlr.press/v206/singh23a.html>. 4, 7, 9
- Sohn, M. B. and Li, H. Compositional mediation analysis for microbiome studies. *Annals of Applied Statistics*, 13(1):661–681, 2019. ISSN 19417330. doi: 10.1214/18-AOAS1210. 3
- Stock, J. H. and Trebbi, F. Retrospectives: Who invented instrumental variable regression? *Journal of Economic Perspectives*, 17(3):177–194, 2003. 3
- Sverchkov, Y. and Craven, M. A review of active learning approaches to experimental design for uncovering biological networks. *PLOS Computational Biology*, 13(6):1–26, 06 2017. doi: 10.1371/journal.pcbi.1005466. URL <https://doi.org/10.1371/journal.pcbi.1005466>. 4
- Van Rossum, G. and Drake, F. *Python 3 Reference Manual: (Python Documentation Manual Part 2)*. Documentation for Python. CreateSpace Independent Publishing Platform, 2009. ISBN 9781441412690. URL <https://books.google.de/books?id=KIybQQAACAAJ>. 24
- Virtanen, P., Gommers, R., Oliphant, T. E., Haberland, M., Reddy, T., Cournapeau, D., Burovski, E., Peterson, P., Weckesser, W., Bright, J., et al. Scipy 1.0: fundamental algorithms for scientific computing in python. *Nature methods*, 17(3):261–272, 2020. 24
- von Kügelgen, J., Rubenstein, P. K., Schölkopf, B., and Weller, A. Optimal experimental design via bayesian optimization: active causal structure learning for gaussian process networks. *arXiv preprint arXiv:1910.03962*, 2019. 4
- Wang, C., Hu, J., Blaser, M. J., Li, H., and Birol, I. Estimating and testing the microbial causal mediation effect with high-dimensional and compositional microbiome data. *Bioinformatics*, 2020. ISSN 14602059. doi: 10.1093/bioinformatics/btz565. 3
- Wang, L. and Tchetgen Tchetgen, E. Bounded, efficient and multiply robust estimation of average treatment effects using instrumental variables. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 80(3):531–550, 2018. 4
- Williams, R. J. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Machine learning*, 8:229–256, 1992. 8
- Wright, P. G. *Tariff on animal and vegetable oils*. Macmillan Company, New York, 1928. 3
- Xu, L., Chen, Y., Srinivasan, S., de Freitas, N., Doucet, A., and Gretton, A. Learning deep features in instrumental variable regression. *arXiv preprint arXiv:2010.07154*, 2020. 4
- Zemplenyi, M. and Miller, J. W. Bayesian optimal experimental design for inferring causal structure. *Bayesian Analysis*, 2021. doi: 10.1214/22-ba1335. 4
- Zhang, J. and Bareinboim, E. Bounding causal effects on continuous outcome. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pp. 12207–12215, 2021. 4
- Zhang, R., Imaizumi, M., Schölkopf, B., and Muandet, K. Maximum moment restriction for instrumental variable regression. *arXiv preprint arXiv:2010.07684*, 2020. 3, 6
- Zhang, R., Imaizumi, M., Schölkopf, B., and Muandet, K. Instrumental variable regression via kernel maximum moment loss. *Journal of Causal Inference*, 11(1):20220073, 2023. doi: doi: 10.1515/jci-2022-0073. URL <https://doi.org/10.1515/jci-2022-0073>. 4

A Proofs

In this section we provide proofs of the statements in the main paper. Before proving Theorem 2 we restate a result from the literature that provides closed form consistent estimates of the minimax problem with the L_2 penalty term (instead of penalizing the functional).

Theorem 3 (Proposition 10 of (Dikkala et al., 2020)). *Assume without loss of generality that functions in $\mathcal{H}_Z, \mathcal{H}_X$ have bounded variation and their images are contained in $[-1, 1]$. Additionally, assume that for each $h \in \mathcal{H}_Z$ also $-h \in \mathcal{H}_Z$. Denote by K_{XX} and K_{ZZ} the empirical kernel matrices from $\mathcal{D}$. Then we can consistently estimate*

$$f = \arg \inf_{f \in \mathcal{F}} \sup_{g \in \mathcal{G}} \mathbb{E}[(Y - f(X))g(Z)] - \lambda_g \|g\|_{\mathcal{G}} + \lambda_f \|f\|_{\mathcal{F}} \quad (15)$$

via

$$\hat{f}(\cdot) = \sum_{i=1}^n \hat{\theta}_i k(X_i, \cdot) \quad \hat{\theta} = (K_{XX} M K_{XX} + 4\lambda_g \lambda_f K_{XX})^+ K_{XX} M y \quad (16)$$

with $M = K_{ZZ}$. If the function classes $\mathcal{F}, \mathcal{G}$ are already norm constrained, the estimator needs no penalization.

Proof. Note that we have the optimum of the inner supremum

$$\sup_{g \in \mathcal{G}} \mathbb{E}[(Y - f(X))g(Z)] - \lambda_g \|g\|_{\mathcal{G}} \quad (17)$$

at

$$\frac{1}{4\lambda_g} (Y - f(X))^\top M (Y - f(X)) \quad (18)$$

Therefore, we are left to solve the outer minimization of

$$\hat{f} = \min_{f \in \mathcal{F}} (Y - f(X))^\top M (Y - f(X)) + 4\lambda_g \lambda_f \|f\|_{\mathcal{F}}^2 \quad (19)$$

$$\begin{aligned} \hat{f} &= \min_{f \in \mathcal{F}} (Y - f(X))^\top M (Y - f(X)) + 4\lambda_g \lambda_f \|f\|_{\mathcal{F}}^2 \\ &= \min_{\theta \in \mathbb{R}^n} (y - K_{XX} \theta)^\top M (y - K_{XX} \theta) + 4\lambda_g \lambda_f \|K_X \theta\|_2^2 && \text{Replacing the function with the (empirical) kernels} \\ &= \min_{\theta \in \mathbb{R}^n} \theta^\top K_{XX} M K_{XX} \theta - 2y^\top M K_{XX} \theta + 4\lambda_g \lambda_f \theta^\top K_X \theta && \text{Writing out product} \\ &= \min_{\theta \in \mathbb{R}^n} \theta^\top (K_{XX} M K_{XX} + 4\lambda_g \lambda_f K_{XX}) \theta - 2y^\top M K_{XX} \theta && \text{Aggregation of terms wrt } \theta \\ \hat{f} &= (K_{XX} M K_{XX} + 4\lambda_g \lambda_f K_{XX})^+ K_{XX} M y && \text{Solution for convex minimization problems} \end{aligned}$$

□

Theorem 2. *Assume functions in $\mathcal{H}_Z, \mathcal{H}_X$ to have bounded variation, (w.l.o.g.) images contained in $[-1, 1]$, and for $h \in \mathcal{H}_Z$ also $-h \in \mathcal{H}_Z$. Denote by $K_{XX}, K_{ZZ} \in \mathbb{R}^{n \times n}$ the empirical kernel matrices of $\mathcal{D}$, let k_X be continuously differentiable, and let $\lambda_g, \lambda_f, \lambda_c \geq 0$. Further, fix $Q[f]$ to be a bounded linear functional, and write $f_\theta(\cdot) = \sum_{i=1}^n \theta_i k(x_i, \cdot)$. Then the solutions to the regularized minimax problems*

$$f^\pm = \arg \min_{f \in \mathcal{H}_X} \mp Q[f] + \lambda_c \sup_{g \in \mathcal{H}_Z} \langle r_0 - T f, g \rangle_{\mathcal{H}_Z} + \lambda_f \|f\|_{\mathcal{H}_X} - \lambda_g \|g\|_{\mathcal{H}_Z} \quad (8)$$

are consistently estimated by $f_{\hat{\theta}^\pm}$ with

$$\hat{\theta}^\pm = (K_{XX} K_{ZZ} K_{XX} + 4 \frac{\lambda_g \lambda_f}{\lambda_c} K_{XX})^+ \left(K_{XX} K_{ZZ} y \mp \frac{1}{\lambda_c} Q[K_X \cdot](x^*) \right), \quad (9)$$

where $(Q[K_X \cdot])(x^*) \in \mathbb{R}^n$ is the vector $((Q[k_X(x_1, \cdot)])(x^*), \dots, (Q[k_X(x_n, \cdot)])(x^*))$. Note that these general bounded linear functionals can be computed analytically for many common kernels such as linear, polynomial, or radial basis function (RBF) kernels.

Proof. The solution of the inner supremum remains the same as in Theorem 3. This leaves us to solve the outer minimization

$$\min_{f \in \mathcal{H}_X} Q[f] + 4\lambda_g \lambda_f \|f\|_{\mathcal{H}_X} + \lambda_c (y - K_{XX}\theta)^\top M (y - K_{XX}\theta), \quad (20)$$

where $M = K_{ZZ}$. We write out f as a linear combination of kernel basis functions (at the observed data points) and use the linearity of the gradient functional, which applied to the kernel basis functions is just $(\partial_i K_{X\cdot})(x^*)$. This yields for the overall functional $Q[f_\theta] = (\partial_i K_{X\cdot})(x^*)\theta$. We are thus seeking

$$\min_{\theta \in \mathbb{R}^n} (\partial_i K_{X\cdot})(x^*)\theta + \lambda_c \theta^\top K_{XX} M K_{XX} \theta - 2\lambda_c y^\top M K_{XX} \theta + 4\lambda_g \lambda_f \theta^\top K_{XX} \theta, \quad (21)$$

which we rewrite as

$$\min_{\theta \in \mathbb{R}^n} \theta^\top \left(K_{XX} M K_{XX} + 4 \frac{\lambda_g \lambda_f}{\lambda_c} K_{XX} \right) \theta + \left(\frac{1}{\lambda_c} (\partial_i K_{X\cdot})(x^*) - 2y^\top M K_{XX} \right) \theta. \quad (22)$$

Building on the arguments in the proof of Theorem 3 where we replace the linear part in θ by the partial derivative functional, we obtain the minimizer

$$\left(K_{XX} M K_{XX} + 4 \frac{\lambda_g \lambda_f}{\lambda_c} K_{XX} \right)^+ \left(K_{XX} M y - \frac{1}{\lambda_c} (\partial_i K_{X\cdot})(x^*) \right) \quad (23)$$

as a consistent estimator. $\square$

Proof. We re-iterate the same arguments to prove the theorem for general bounded linear functionals. We assume $Q[f]$ to be a linear bounded functional acting on f .

Theorem 4. Riesz representation theorem *Let $Q[f] : \mathcal{H}_X \rightarrow \mathbb{R}$ be a continuous linear functional on a separable Hilbert space $\mathcal{H}_X$. Then there exists a $q \in \mathcal{H}_X$ such that $Q[f] = \langle f, q \rangle_{\mathcal{H}_X}$, $\forall f \in \mathcal{H}_X$.*

Assuming $f \in \mathcal{H}_X$, via the Riesz-representer theorem for linear functionals, there exists a unique element $q \in \mathcal{H}_X$ such that for any $f \in \mathcal{H}_X$. The bounded linear functionals on $\mathcal{H}_X$ form themselves a hilbert space called the dual space.

$$Q[f] = \langle f, q \rangle_{\mathcal{H}_X} \quad (24)$$

We assume $\mathcal{H}_X$ to be an RKHS. Therefore, we can write out $f(x)$ as a linear combination of the kernel basis functions (at the observed data points):

We replace $Q[f] = \langle f, q \rangle_{\mathcal{H}_X} = \langle \sum_{i=1}^n \theta_i k(x_i, \cdot), q \rangle = \sum_{i=1}^n \theta_i q(x_i)$ as $q \in \mathcal{H}_X$. Due to the RKHS, we can write the empirical version as $q(\cdot) = \sum_{i=1}^n Q[k(x_i, \cdot)]$.

$$\min_{f \in \mathcal{H}_X} \theta^\top Q[k(x_i, \cdot)] + 4\lambda_g \lambda_f \theta^\top K_{XX} \theta + \lambda_c (y - K_{XX}\theta)^\top M (y - K_{XX}\theta) \quad (25)$$

Thus we obtain the minimizer

$$\left(K_{XX} M K_{XX} + 4 \frac{\lambda_g \lambda_f}{\lambda_c} K_{XX} \right)^+ \left(K_{XX} M y - \frac{1}{\lambda_c} Q[K_{X\cdot}(x^*)] \right) \quad (26)$$

$\square$

B Details on Adaptive Policies

For completeness, we here recall the score function of Gaussian mixture models, i.e., the gradients with respect to the mixture weights, means, and covariances of the log-likelihood function. We write π for the GMM for brevity. First, we define the ‘responsibilities’ ζ_{im} , the probability that a sample z_i comes from component m via (see, e.g., [Petersen & Pedersen \(2008\)](#))

$$\zeta_{im} = \frac{\gamma_m \mathcal{N}(z_i | \mu_m, \Sigma_m)}{\sum_{j=1}^M \gamma_j \mathcal{N}(z_i | \mu_j, \Sigma_j)}. \quad (27)$$

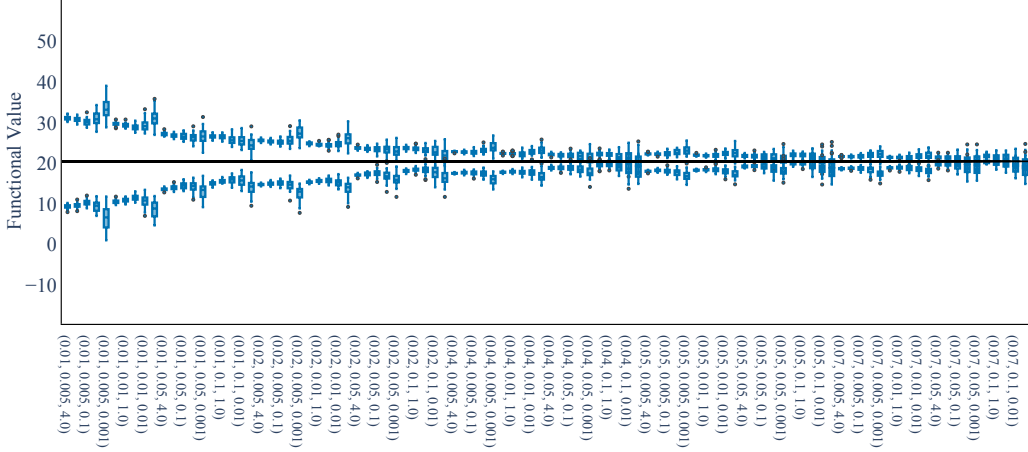


Figure 2: Adaptive Method for different hyperparameter settings. The x-axis shows them in the following order: $(\lambda_c, \lambda_s, \alpha_t)$

Then, we have for a sample z_i and mixture component $m \in [M]$ that

$$\nabla_{\gamma_m} \log \pi(z_i) = \frac{\zeta_{im}}{\gamma_m}, \quad (28)$$

$$\nabla_{\mu_m} \log \pi(z_i) = \zeta_{im} \Sigma_m^{-1} (z_i - \mu_m), \quad (29)$$

$$\nabla_{\Sigma_m} \log \pi(z_i) = -\frac{1}{2} \zeta_{im} (\Sigma_m^{-1} - \Sigma_m^{-1} (z_i - \mu_m)(z_i - \mu_m)^T \Sigma_m^{-1}). \quad (30)$$

C Hyperparameter Tuning

In Eq. (11) we end up with three hyperparameters. Note however, that those three mentioned hyperparameters are relative weights. Thus, we can set $\lambda_s := \frac{\lambda_f \lambda_g}{\lambda_c}$ and only tune λ_s and λ_c . Intuitively, λ_s regularizes the smoothness of the function spaces and λ_c weighs the functional Q relative to the moment conditions. In Fig. 2, we show the dependence of our bounds on these hyperparameters: very low values of λ_s lead to conservative bounds whereas large values of λ_s yield narrow bounds throughout, as λ_s effectively controls the size of the search space via regularity restrictions. For the learning rate, a lot of learning rates below a certain value actually work for the method. More gradient update steps (and therefore more experiments) are required for very small learning rates, very high learning rates result in drastic updates of the policy; something that is not very favorable in real-world experiments.

D Runtimes

For completeness, we show wallclock runtimes of the different methods in our empirical evaluation in Figure 3. All methods were run on a MacBook Pro with Intel CPU. The runtimes are per iteration and clearly increase for methods that accumulate data over previous rounds. Even the most expensive passive baseline is relatively fast to compute without specialized hardware and we generally imagine computational cost to be negligible compared to the actual cost of physical experimentation required in each round in most applications.

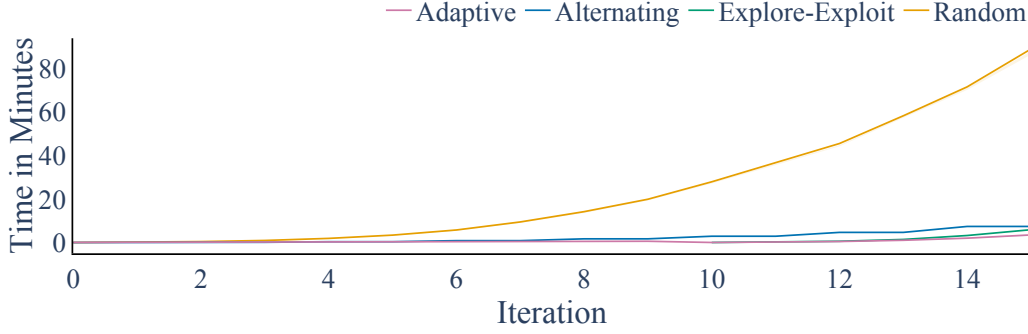


Figure 3: Wallclock runtimes of the different methods.

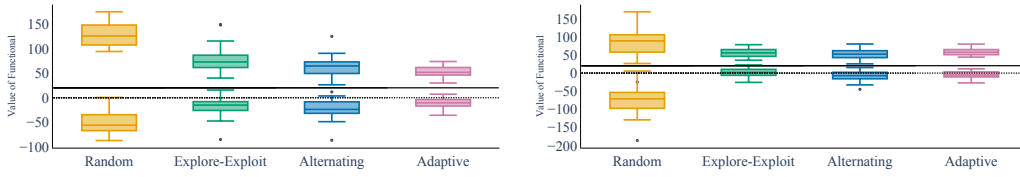


Figure 4: We compare the different strategies in our synthetic setting. The black constant line represents the true value of $Q[f_0]$. Both plots show the estimated upper and lower bounds $Q_t^\pm$ at $T = 16$ for $n_{\text{seeds}} = 50$. Lines are means and shaded regions are (10, 90)-percentiles. **Left:** $d_z = 5, d_x = 20$, **Right:** $d_z = d_x = 20$.

E Additional Experiments

We assume the same setting as the low-dimensional one, only increasing the dimensions.

$$h_j(Z, U) = \begin{cases} \alpha \cdot \sin(Z_j) \cdot (1 + U), & \text{if } j \leq d_z, \\ 1 + U, & \text{if } j > d_z. \end{cases} \quad (31)$$

$$f(X) = \beta \sum_{j=1}^{d_x} \exp(X_j) \sin(X_j) + U, \quad U \sim \mathcal{N}(0, 1), \quad (32)$$

with $d_z = 5, d_x = 20$ and $d_z = 20, d_x = 20$. The setting guarantees, that though we have a mismatch in dimensions, i.e. $d_z < d_x$, we would—in principle—still be able to identify the full causal effect, as the completeness property is fulfilled. We refer to Fig. 4 for the results of the different methods. Moreover, we note that the increase in time is still manageable, c.f. Fig 5. The main driver for the computation time are the samples within each experiment (T_1, T_2) due to the kernel approach and, therefore, the inversion of gram matrices based on the sample size.

For $d_z = 5, d_x = 20$, we used $\lambda_s := \frac{\lambda_g \lambda_f}{\lambda_c} = 0.04$ and $\lambda_c = 0.1$. For $d_z = d_x = 20$, we used $\lambda_s := \frac{\lambda_g \lambda_f}{\lambda_c} = 0.05$ and $\lambda_c = 0.1$. Intuitively, the optimization requires stronger hyperparameters in higher dimensional settings as with dimensions, also the function space increases.

F Underspecification

Additionally, we include some examples along the lines of underspecification. We refer to [Ailer et al. \(2023\)](#) for a thorough discussion on underspecification, but give some motivation for the examples in the following:

Since f is the solution to a linear inverse problem, the problem is often ill-posed and thus f is not identified. However, underspecification means something more explicit. It can be seen as a violation of the completeness property.

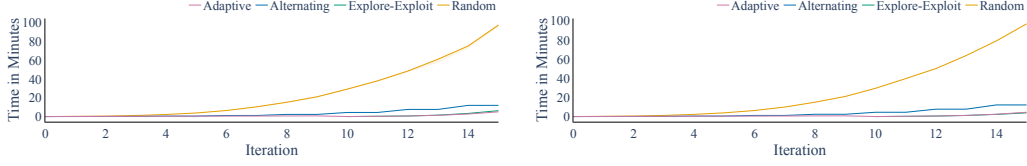


Figure 5: We compare the wallclock time in a higher dimensional setting. The time increase is still mild. The main driver is the number of samples in each experiment, instead of the dimensionality itself. **Left:** $d_z = 5, d_x = 20$, **Right:** $d_z = d_x = 20$.

Lemma 5 (Severini & Tripathi (2006)). *The conditional distribution of $p(X | Z)$ is complete if and only if for each function $f(x)$ such that $\mathbb{E}[f(x)] = 0$ and $\text{Var}(f(x)) > 0$, there exists a function $g(z)$ such that $f(x)$ and $g(z)$ are correlated.*

In most papers, the completeness property is trivially fulfilled by assuming an exponential family for $p(x | z)$ with non-zero variance for all components. This means that both T and T^* are injective which again guarantees the identifiability of f .

Now, we denote $\mathcal{P}_{\mathcal{F}}$ as the orthogonal projection onto the function space $\mathcal{F}$ and $\mathcal{N}(T)$ as the null space of the linear operator $T : \mathcal{H}_X \rightarrow \mathcal{H}_Z$ with $Tf = T[f(X) | Z]$. Following this notation, Severini & Tripathi (2012) argue that $\mathcal{P}_{\mathcal{N}(T)^\perp} f$ is the identifiable part of f . This leads to the following lemma:

Lemma 6 (Severini & Tripathi (2012)). *$\mathbb{E}[Q[f]]$ is identified if and only if $Q \in \mathcal{N}(T)^\perp$.*

Examples. Let us introduce an example for an unidentified $Q[f]$: Consider $Z, Y \in \mathbb{R}, X \in \mathbb{R}^2$ and $h(z) = (z, 0)$, $f(x_1, x_2) = 0.5x_1 + 2x_2$ and $Q[f] = \partial_2 f(x^*)$. In this fully linear setting, due to the structure of h the instrument can only ever perturb the first component of X regardless of what policy we choose. Hence, it is impossible to identify e.g. $Q[f]$ w.r.t. the second argument with $\partial_2 f(x^*) = 2\forall x^*$. Even a policy with full support will not (even partially) identify $Q[f_0]$. In contrast, let us discuss a fully identifiable scenario, which is identifiable for all $Q[f]$, but still depends on the policy: Consider a similar setting with $f(x_1, x_2) = x_1^2 + x_2$ and $Q[f] = \partial_1 f(x^*)$. Assume we have $x^* = (6, 1)$, then $Q[f_0] = 2 \cdot 6 = 12$. Any policy π that has positive density in a neighborhood of $z = 6$ will be able to fully identify $Q[f_0]$. However, any policy that puts no mass (or in practice, little mass) near $z = 6$, will not identify $Q[f_0]$. This would be an uninformative policy.

G Algorithmic Boxes

We present pseudocode for the ‘explore then exploit’ (EE) strategy in Algorithm 1, for ‘alternating explore exploit’ (AEE) in Algorithm 2, and for the adaptive strategy in Algorithm 3.

Algorithm 1 Explore then exploit (EE)

```
1: Input: Budget  $T$ ,  $T_1$ ,  $T_2$ ,  $n$ , distance measure  $d$  on  $\mathcal{X}$ , target  $x^*$ 
2: Initialize:  $\mathcal{D}^{(\leq T_1)} \leftarrow \emptyset$ 
3: Split budget:  $T = T_1 + T_2$ 
4: for  $t = 1$  to  $T_1$  do
5:   Sample  $\pi_t \sim \Gamma$ 
6:   Observe sample  $(X_t, Z_t)$ 
7:    $\mathcal{D}^{(\leq T_1)} \leftarrow \mathcal{D}^{(\leq T_1)} \cup \{(X_t, Z_t)\}$ 
8: end for
9: Find  $K$  nearest samples to  $x^*$  in  $\mathcal{D}^{(\leq T_1)}$  based on distance  $d$ 
10: Fit  $\pi$  as a parametric distribution on the  $Z$  values of these  $K$  samples
11: for  $t = T_1 + 1$  to  $T$  do
12:   Set  $\pi_t := \pi$ 
13:   Observe sample  $(X_t, Z_t)$ 
14:    $\mathcal{D}^{(t)} \leftarrow \{(X_t, Z_t)\}$ 
15: end for
16: Estimate  $Q_t^\pm$  from  $\bigcup_{t=T_1+1}^T \mathcal{D}^{(t)}$ 
```

Algorithm 2 Alternating explore exploit (AEE)

```
1: Input: Budget  $T$ ,  $K$ ,  $n$ , distance measure  $d$  on  $\mathcal{X}$ , target  $x^*$ 
2: Initialize: Sorted list of observed samples  $\mathcal{L} \leftarrow \emptyset$ 
3: for  $t \in [T]$  do
4:   if  $t$  is odd and  $t \neq T$  then
5:     Sample  $\pi_t \sim \Gamma$ 
6:     Observe sample  $(X_t, Z_t)$ 
7:      $\mathcal{L} \leftarrow \mathcal{L} \cup \{(X_t, Z_t)\}$ 
8:     Sort  $\mathcal{L}$  by distance  $d(X, x^*)$ 
9:   else if  $t$  is even or  $t = T$  then
10:    Find the  $\min\{K, n \cdot t\}$  nearest samples to  $x^*$  in  $\mathcal{L}$ 
11:    Fit  $\pi_t$  as a parametric distribution on the  $Z$  values of these samples
12:    Observe sample  $(X_t, Z_t)$ 
13:     $\mathcal{L} \leftarrow \mathcal{L} \cup \{(X_t, Z_t)\}$ 
14:    Sort  $\mathcal{L}$  by distance  $d(X, x^*)$ 
15:   end if
16: end for
17: Estimate  $Q_t^\pm$  on  $\bigcup_{\text{even } t} \mathcal{D}^{(t)}$ 
```

Algorithm 3 Adaptive

```
1: Input: Budget  $T$ ,  $T_1$ ,  $T_2$ ,  $n$ 
2: Initialize:  $\Phi_t$ 
3: Split budget:  $T = T_1 + T_2$ 
4: for  $t = 1$  to  $T_1$  do
5:   Sample  $Z_t \sim \pi_{\Phi_t}$ 
6:   Observe sample  $(X_t, Z_t)$ 
7:   Compute gradient update  $\nabla(\Phi_t)\Delta(\pi_{\Phi_t})$  by splitting observed samples
8: end for
9: for  $t = T_1 + 1$  to  $T$  do
10:   Set  $\pi_t := \pi_{\Phi_{T_1}}$ 
11:   Observe sample  $(X_t, Z_t)$ 
12:    $\mathcal{D}^{(t)} \leftarrow \{(X_t, Z_t)\}$ 
13: end for
14: Estimate  $Q_t^\pm$  from  $\bigcup_{t=T_1+1}^T \mathcal{D}^{(t)}$ 
```

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 abstract clearly states our contributions, the additional *Limitation* block marks additionally all goals that are motivational and states which of those are further projects.

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 conclusion includes an additional paragraph stating the limitations of our work. Moreover, we acknowledge the open question and explain a potential way how to address those.

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: (1) Assumptions: we state them directly in the main paper, moreover we provide references to the papers on which our results are based on, including the section of their proofs. (2) Proofs: we sketch the proofs in the main paper and reference to the full proof in the supplementary.

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: The code will be provided in a public repository. The experiments are performed via a command line run file. The results are saved to a npy-formated data file. One jupyter notebook does perform the visualization. Additionally, we listed the seeds that are relevant for the data generation to ensure full reproducibility.

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: The data is based on a simulation. The code, as well as the exact command line arguments are on github. Moreover, all code to reproduce the results is openly available and commented.

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: The code, as well as instructions are openly available. The code and results can be reproduced via a python file that is run via the command line. Hyperparameters are listed and the choice justified in the Experimental section of the main paper and the appendix.

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: The experiments are run multiple times over $n_{\text{seeds}} = 25$ different seeds. We state barplots in the results.

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: [Yes]

Justification: The computing was done at an internal academic cluster. We provide estimated time in the appendix.

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: Our data is evaluated on simulated data, moreover, our algorithm tries to move away from correlational assumptions and tries to identify a part of the causal quantity. Even though our claim is to recover part of the causal effect, we state its limitations. We provide sufficient documentation to reproduce the results, but also to understand and comprehend our claims. In conclusion we are convinced that we comply to the NeurIPS Code of Conduct.

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: [NA]

Justification: Our paper focuses on foundational research.

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: Neither the simulated dataset, nor the algorithm have a high risk of misuse. The algorithm only adds value as soon as it is used on real world experimental set-ups. Currently, it is only based on our simulation studies.

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: [NA]

Justification: The theoretical contributions on which our work is based are correctly cited in the main paper, the code is written from scratch and solely showcases our proposed algorithm. As packages we are using:

Table 1: Overview of resources used in our work.

Name	Reference	License
Python	Van Rossum & Drake (2009)	PSF License
PyTorch	Paszke et al. (2019)	BSD-style license
Numpy	Harris et al. (2020)	BSD-style license
Pandas	pandas development team (2020)	BSD-style license
Jupyter	Kluyver et al. (2016)	BSD-style license
Matplotlib	Hunter (2007)	modified PSF (BSD compatible)
Scikit-learn	Pedregosa et al. (2011)	BSD 3-Clause
SciPy	Virtanen et al. (2020)	BSD 3-Clause
JAX	Bradbury et al. (2018)	
SLURM	Andy et al. (2003)	Apache 2.0

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: [\[Yes\]](#)

Justification: While our main contribution is not the code that we provide alongside the paper, but only provided to ensure reproducibility, the code is still well documented and ready to use. The consent is given, as the authors of the papers are owners of the code as well.

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: The paper does not involve any crowdsourcing or research with human objects.

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: The paper does not involve any crowdsourcing or research with human objects.

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.



Attribution 4.0 International

Deed

Canonical URL : <https://creativecommons.org/licenses/by/4.0/>

[See the legal code](#)

You are free to:

Share — copy and redistribute the material in any medium or format for any purpose, even commercially.

Adapt — remix, transform, and build upon the material for any purpose, even commercially.

The licensor cannot revoke these freedoms as long as you follow the license terms.

Under the following terms:

Attribution — You must give **appropriate credit** , provide a link to the license, and **indicate if changes were made** . You may do so in any reasonable manner, but not in any way that suggests the licensor endorses you or your use.

No additional restrictions — You may not apply legal terms or **technological measures** that legally restrict others from doing anything the license permits.

Notices:

You do not have to comply with the license for elements of the material in the public domain or where your use is permitted by an applicable **exception or limitation** .

No warranties are given. The license may not give you all of the permissions necessary for your intended use. For example, other rights such as **publicity, privacy, or moral rights** may limit how you use the material.

Notice

This deed highlights only some of the key features and terms of the actual license. It is not a license and has no legal value. You should

carefully review all of the terms and conditions of the actual license before using the licensed material.

Creative Commons is not a law firm and does not provide legal services. Distributing, displaying, or linking to this deed or the license that it summarizes does not create a lawyer-client or any other relationship.

Creative Commons is the nonprofit behind the open licenses and other legal tools that allow creators to share their work. Our legal tools are free to use.

- [Learn more about our work](#)
- **[Learn more about CC Licensing](#)**
- [Support our work](#)
- [Use the license for your own material.](#)
- [Licenses List](#)
- [Public Domain List](#)

Footnotes

appropriate credit — If supplied, you must provide the name of the creator and attribution parties, a copyright notice, a license notice, a disclaimer notice, and a link to the material. CC licenses prior to Version 4.0 also require you to provide the title of the material if supplied, and may have other slight differences.

- [More info](#)

indicate if changes were made — In 4.0, you must indicate if you modified the material and retain an indication of previous modifications. In 3.0 and earlier license versions, the indication of changes is only required if you create a derivative.

- [Marking guide](#)
- [More info](#)

technological measures — The license prohibits application of effective technological measures, defined with reference to Article 11 of the WIPO Copyright Treaty.

- [More info](#)

exception or limitation — The rights of users under exceptions and limitations, such as fair use and fair dealing, are not affected by the CC licenses.

- [More info](#)

publicity, privacy, or moral rights — You may need to get additional permissions before using the material as you intend.

- [More info](#)

Attribution 4.0 International

=====
Creative Commons Corporation ("Creative Commons") is not a law firm and does not provide legal services or legal advice. Distribution of Creative Commons public licenses does not create a lawyer-client or other relationship. Creative Commons makes its licenses and related information available on an "as-is" basis. Creative Commons gives no warranties regarding its licenses, any material licensed under their terms and conditions, or any related information. Creative Commons disclaims all liability for damages resulting from their use to the fullest extent possible.

Using Creative Commons Public Licenses

Creative Commons public licenses provide a standard set of terms and conditions that creators and other rights holders may use to share original works of authorship and other material subject to copyright and certain other rights specified in the public license below. The following considerations are for informational purposes only, are not exhaustive, and do not form part of our licenses.

Considerations for licensors: Our public licenses are intended for use by those authorized to give the public permission to use material in ways otherwise restricted by copyright and certain other rights. Our licenses are irrevocable. Licensors should read and understand the terms and conditions of the license they choose before applying it. Licensors should also secure all rights necessary before applying our licenses so that the public can reuse the material as expected. Licensors should clearly mark any material not subject to the license. This includes other CC-licensed material, or material used under an exception or limitation to copyright. More considerations for licensors: wiki.creativecommons.org/Considerations_for_licensors

Considerations for the public: By using one of our public licenses, a licensor grants the public permission to use the licensed material under specified terms and conditions. If the licensor's permission is not necessary for any reason—for example, because of any applicable exception or limitation to copyright—then that use is not regulated by the license. Our licenses grant only permissions under copyright and certain other rights that a licensor has authority to grant. Use of the licensed material may still be restricted for other reasons, including because others have copyright or other rights in the material. A licensor may make special requests, such as asking that all changes be marked or described. Although not required by our licenses, you are encouraged to respect those requests where reasonable. More considerations for the public: wiki.creativecommons.org/Considerations_for_licensees

Creative Commons Attribution 4.0 International Public License

By exercising the Licensed Rights (defined below), You accept and agree

to be bound by the terms and conditions of this Creative Commons Attribution 4.0 International Public License ("Public License"). To the extent this Public License may be interpreted as a contract, You are granted the Licensed Rights in consideration of Your acceptance of these terms and conditions, and the Licensor grants You such rights in consideration of benefits the Licensor receives from making the Licensed Material available under these terms and conditions.

Section 1 -- Definitions.

- a. Adapted Material means material subject to Copyright and Similar Rights that is derived from or based upon the Licensed Material and in which the Licensed Material is translated, altered, arranged, transformed, or otherwise modified in a manner requiring permission under the Copyright and Similar Rights held by the Licensor. For purposes of this Public License, where the Licensed Material is a musical work, performance, or sound recording, Adapted Material is always produced where the Licensed Material is synched in timed relation with a moving image.
- b. Adapter's License means the license You apply to Your Copyright and Similar Rights in Your contributions to Adapted Material in accordance with the terms and conditions of this Public License.
- c. Copyright and Similar Rights means copyright and/or similar rights closely related to copyright including, without limitation, performance, broadcast, sound recording, and Sui Generis Database Rights, without regard to how the rights are labeled or categorized. For purposes of this Public License, the rights specified in Section 2(b)(1)-(2) are not Copyright and Similar Rights.
- d. Effective Technological Measures means those measures that, in the absence of proper authority, may not be circumvented under laws fulfilling obligations under Article 11 of the WIPO Copyright Treaty adopted on December 20, 1996, and/or similar international agreements.
- e. Exceptions and Limitations means fair use, fair dealing, and/or any other exception or limitation to Copyright and Similar Rights that applies to Your use of the Licensed Material.
- f. Licensed Material means the artistic or literary work, database, or other material to which the Licensor applied this Public License.
- g. Licensed Rights means the rights granted to You subject to the terms and conditions of this Public License, which are limited to all Copyright and Similar Rights that apply to Your use of the Licensed Material and that the Licensor has authority to license.
- h. Licensor means the individual(s) or entity(ies) granting rights under this Public License.
- i. Share means to provide material to the public by any means or process that requires permission under the Licensed Rights, such as reproduction, public display, public performance, distribution, dissemination, communication, or importation, and to make material available to the public including in ways that members of the

public may access the material from a place and at a time individually chosen by them.

- j. Sui Generis Database Rights means rights other than copyright resulting from Directive 96/9/EC of the European Parliament and of the Council of 11 March 1996 on the legal protection of databases, as amended and/or succeeded, as well as other essentially equivalent rights anywhere in the world.
- k. You means the individual or entity exercising the Licensed Rights under this Public License. Your has a corresponding meaning.

Section 2 -- Scope.

a. License grant.

1. Subject to the terms and conditions of this Public License, the Licensor hereby grants You a worldwide, royalty-free, non-sublicensable, non-exclusive, irrevocable license to exercise the Licensed Rights in the Licensed Material to:
 - a. reproduce and Share the Licensed Material, in whole or in part; and
 - b. produce, reproduce, and Share Adapted Material.
2. Exceptions and Limitations. For the avoidance of doubt, where Exceptions and Limitations apply to Your use, this Public License does not apply, and You do not need to comply with its terms and conditions.
3. Term. The term of this Public License is specified in Section 6(a).
4. Media and formats; technical modifications allowed. The Licensor authorizes You to exercise the Licensed Rights in all media and formats whether now known or hereafter created, and to make technical modifications necessary to do so. The Licensor waives and/or agrees not to assert any right or authority to forbid You from making technical modifications necessary to exercise the Licensed Rights, including technical modifications necessary to circumvent Effective Technological Measures. For purposes of this Public License, simply making modifications authorized by this Section 2(a) (4) never produces Adapted Material.
5. Downstream recipients.
 - a. Offer from the Licensor -- Licensed Material. Every recipient of the Licensed Material automatically receives an offer from the Licensor to exercise the Licensed Rights under the terms and conditions of this Public License.
 - b. No downstream restrictions. You may not offer or impose any additional or different terms or conditions on, or apply any Effective Technological Measures to, the Licensed Material if doing so restricts exercise of the Licensed Rights by any recipient of the Licensed

Material.

6. No endorsement. Nothing in this Public License constitutes or may be construed as permission to assert or imply that You are, or that Your use of the Licensed Material is, connected with, or sponsored, endorsed, or granted official status by, the Licensor or others designated to receive attribution as provided in Section 3(a)(1)(A)(i).

b. Other rights.

1. Moral rights, such as the right of integrity, are not licensed under this Public License, nor are publicity, privacy, and/or other similar personality rights; however, to the extent possible, the Licensor waives and/or agrees not to assert any such rights held by the Licensor to the limited extent necessary to allow You to exercise the Licensed Rights, but not otherwise.
2. Patent and trademark rights are not licensed under this Public License.
3. To the extent possible, the Licensor waives any right to collect royalties from You for the exercise of the Licensed Rights, whether directly or through a collecting society under any voluntary or waivable statutory or compulsory licensing scheme. In all other cases the Licensor expressly reserves any right to collect such royalties.

Section 3 -- License Conditions.

Your exercise of the Licensed Rights is expressly made subject to the following conditions.

a. Attribution.

1. If You Share the Licensed Material (including in modified form), You must:
 - a. retain the following if it is supplied by the Licensor with the Licensed Material:
 - i. identification of the creator(s) of the Licensed Material and any others designated to receive attribution, in any reasonable manner requested by the Licensor (including by pseudonym if designated);
 - ii. a copyright notice;
 - iii. a notice that refers to this Public License;
 - iv. a notice that refers to the disclaimer of warranties;
 - v. a URI or hyperlink to the Licensed Material to the extent reasonably practicable;
 - b. indicate if You modified the Licensed Material and

retain an indication of any previous modifications; and

- c. indicate the Licensed Material is licensed under this Public License, and include the text of, or the URI or hyperlink to, this Public License.
2. You may satisfy the conditions in Section 3(a)(1) in any reasonable manner based on the medium, means, and context in which You Share the Licensed Material. For example, it may be reasonable to satisfy the conditions by providing a URI or hyperlink to a resource that includes the required information.
3. If requested by the Licensor, You must remove any of the information required by Section 3(a)(1)(A) to the extent reasonably practicable.
4. If You Share Adapted Material You produce, the Adapter's License You apply must not prevent recipients of the Adapted Material from complying with this Public License.

Section 4 -- Sui Generis Database Rights.

Where the Licensed Rights include Sui Generis Database Rights that apply to Your use of the Licensed Material:

- a. for the avoidance of doubt, Section 2(a)(1) grants You the right to extract, reuse, reproduce, and Share all or a substantial portion of the contents of the database;
- b. if You include all or a substantial portion of the database contents in a database in which You have Sui Generis Database Rights, then the database in which You have Sui Generis Database Rights (but not its individual contents) is Adapted Material; and
- c. You must comply with the conditions in Section 3(a) if You Share all or a substantial portion of the contents of the database.

For the avoidance of doubt, this Section 4 supplements and does not replace Your obligations under this Public License where the Licensed Rights include other Copyright and Similar Rights.

Section 5 -- Disclaimer of Warranties and Limitation of Liability.

- a. UNLESS OTHERWISE SEPARATELY UNDERTAKEN BY THE LICENSOR, TO THE EXTENT POSSIBLE, THE LICENSOR OFFERS THE LICENSED MATERIAL AS-IS AND AS-AVAILABLE, AND MAKES NO REPRESENTATIONS OR WARRANTIES OF ANY KIND CONCERNING THE LICENSED MATERIAL, WHETHER EXPRESS, IMPLIED, STATUTORY, OR OTHER. THIS INCLUDES, WITHOUT LIMITATION, WARRANTIES OF TITLE, MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE, NON-INFRINGEMENT, ABSENCE OF LATENT OR OTHER DEFECTS, ACCURACY, OR THE PRESENCE OR ABSENCE OF ERRORS, WHETHER OR NOT KNOWN OR DISCOVERABLE. WHERE DISCLAIMERS OF WARRANTIES ARE NOT ALLOWED IN FULL OR IN PART, THIS DISCLAIMER MAY NOT APPLY TO YOU.
- b. TO THE EXTENT POSSIBLE, IN NO EVENT WILL THE LICENSOR BE LIABLE TO YOU ON ANY LEGAL THEORY (INCLUDING, WITHOUT LIMITATION, NEGLIGENCE) OR OTHERWISE FOR ANY DIRECT, SPECIAL, INDIRECT,

INCIDENTAL, CONSEQUENTIAL, PUNITIVE, EXEMPLARY, OR OTHER LOSSES, COSTS, EXPENSES, OR DAMAGES ARISING OUT OF THIS PUBLIC LICENSE OR USE OF THE LICENSED MATERIAL, EVEN IF THE LICENSOR HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH LOSSES, COSTS, EXPENSES, OR DAMAGES. WHERE A LIMITATION OF LIABILITY IS NOT ALLOWED IN FULL OR IN PART, THIS LIMITATION MAY NOT APPLY TO YOU.

- c. The disclaimer of warranties and limitation of liability provided above shall be interpreted in a manner that, to the extent possible, most closely approximates an absolute disclaimer and waiver of all liability.

Section 6 -- Term and Termination.

- a. This Public License applies for the term of the Copyright and Similar Rights licensed here. However, if You fail to comply with this Public License, then Your rights under this Public License terminate automatically.
- b. Where Your right to use the Licensed Material has terminated under Section 6(a), it reinstates:
 - 1. automatically as of the date the violation is cured, provided it is cured within 30 days of Your discovery of the violation; or
 - 2. upon express reinstatement by the Licensor.

For the avoidance of doubt, this Section 6(b) does not affect any right the Licensor may have to seek remedies for Your violations of this Public License.

- c. For the avoidance of doubt, the Licensor may also offer the Licensed Material under separate terms or conditions or stop distributing the Licensed Material at any time; however, doing so will not terminate this Public License.
- d. Sections 1, 5, 6, 7, and 8 survive termination of this Public License.

Section 7 -- Other Terms and Conditions.

- a. The Licensor shall not be bound by any additional or different terms or conditions communicated by You unless expressly agreed.
- b. Any arrangements, understandings, or agreements regarding the Licensed Material not stated herein are separate from and independent of the terms and conditions of this Public License.

Section 8 -- Interpretation.

- a. For the avoidance of doubt, this Public License does not, and shall not be interpreted to, reduce, limit, restrict, or impose conditions on any use of the Licensed Material that could lawfully be made without permission under this Public License.
- b. To the extent possible, if any provision of this Public License is

deemed unenforceable, it shall be automatically reformed to the minimum extent necessary to make it enforceable. If the provision cannot be reformed, it shall be severed from this Public License without affecting the enforceability of the remaining terms and conditions.

- c. No term or condition of this Public License will be waived and no failure to comply consented to unless expressly agreed to by the Licensor.
- d. Nothing in this Public License constitutes or may be interpreted as a limitation upon, or waiver of, any privileges and immunities that apply to the Licensor or You, including from the legal processes of any jurisdiction or authority.

=====
Creative Commons is not a party to its public licenses. Notwithstanding, Creative Commons may elect to apply one of its public licenses to material it publishes and in those instances will be considered the "Licensor." The text of the Creative Commons public licenses is dedicated to the public domain under the CC0 Public Domain Dedication. Except for the limited purpose of indicating that material is shared under a Creative Commons public license or as otherwise permitted by the Creative Commons policies published at creativecommons.org/policies, Creative Commons does not authorize the use of the trademark "Creative Commons" or any other trademark or logo of Creative Commons without its prior written consent including, without limitation, in connection with any unauthorized modifications to any of its public licenses or any other arrangements, understandings, or agreements concerning use of licensed material. For the avoidance of doubt, this paragraph does not form part of the public licenses.

Creative Commons may be contacted at creativecommons.org.

A.3 Instrumental Variable Estimation for Compositional Treatments



OPEN Instrumental variable estimation for compositional treatments

Elisabeth Ailer^{1,2,3}✉, Christian L. Müller^{1,3,4,5} & Niki Kilbertus^{1,2,3}

Many scientific datasets are compositional in nature. Important biological examples include species abundances in ecology, cell-type compositions derived from single-cell sequencing data, and amplicon abundance data in microbiome research. Here, we provide a causal view on compositional data in an instrumental variable setting where the composition acts as the cause. First, we crisply articulate potential pitfalls for practitioners regarding the interpretation of compositional causes from the viewpoint of interventions and warn against attributing causal meaning to common summary statistics such as diversity indices in microbiome data analysis. We then advocate for and develop multivariate methods using statistical data transformations and regression techniques that take the special structure of the compositional sample space into account while still yielding scientifically interpretable results. In a comparative analysis on synthetic and real microbiome data we show the advantages and limitations of our proposal. We posit that our analysis provides a useful framework and guidance for valid and informative cause-effect estimation in the context of compositional data.

Keywords Causality, Cause-effect estimation, Compositional data, Instrumental variable, Microbial diversity

The statistical modeling of compositional (or relative abundance) data plays a pivotal role in many areas of science, ranging from the analysis of mineral samples or rock compositions in earth sciences¹ to correlated topic modeling in large text corpora^{2,3}. Recent advances in biological high-throughput sequencing techniques, including single-cell RNA-Seq and microbial amplicon sequencing^{4,5}, have triggered renewed interest in compositional data analysis. Since only a limited total number of transcripts can be captured in a sample by current sequencing technologies, the resulting count data provides relative abundance information about mRNA transcripts or microbial amplicon sequences, respectively^{6,7}.

For example, in microbiome sequencing, this stems from the fact that one cannot easily control for the total number of microbes entering the measurement process. Bacterial microbiome measurements typically come in the form of counts of operational taxonomic units (OTUs) or amplicon sequencing variants (ASVs) derived from high-throughput sequencing of 16S ribosomal RNA (rRNA)⁸ and are summarized as taxonomic compositions, e.g., on the species, genus, or family level.

One way of dealing with the available relative abundance information is to normalize read counts by their respective totals, resulting in *compositional data*. Compositional data comprises the proportions of some whole, implying that data points live on the unit simplex $\mathbb{S}^{p-1} := \{x \in \mathbb{R}_{\geq 0}^p \mid \sum_{j=1}^p x_j = 1\}$.

In the microbiome example, assume there are p different microbial taxa that have been identified in a human gut microbiome experiment. A specific gut microbiome measurement is then represented by a vector x , where x_j denotes the relative abundance of taxon j (under an arbitrary ordering of taxa). An increase in x_1 within this composition could correspond to an actual increase in the absolute abundance of the first taxon, while the rest remained constant. However, it could equally result from a decrease of the absolute abundance of the first species with the remaining ones having decreased even more.

Statisticians have recognized the significance of compositional data early on (dating back to Karl Pearson) and tailored models to naturally account for compositionality via simplex arithmetic¹. Despite these efforts, adjusting predictive statistical and machine learning methods to compositional data remains an active field of research^{9–18}.

This work focuses on estimating the causal effect of a composition on a categorical or continuous outcome. Only recently have the fundamental challenges in interpreting causal effects of compositions been acknowledged explicitly^{19,20} with little work on how to estimate such effects from observational data. Our work provides

¹Helmholtz Munich, Ingolstädter Landstraße 1, 85764 Neuherberg, Germany. ²TUM School of Computation, Information and Technology, Technical University of Munich, Boltzmannstraße 3, 85748 Garching, Germany. ³Munich Center for Machine Learning (MCML), Munich, Germany. ⁴Department of Statistics, Ludwig-Maximilian University, Geschwister-Scholl-Platz 1, 80539 Munich, Germany. ⁵Center for Computational Mathematics, Flatiron Institute, 162 5th Ave, 10010 New York, NY, United States. ✉email: elisabeth.ailer@helmholtz-munich.de; elisabeth.ailer@gmx.de

scalable methods that *enable practitioners to answer the simple question: “What is the causal effect of a composition on some outcome of interest?”* In the following, we develop methods which can solve this causal questions while also creating an understanding of common pitfalls in causal settings.

Pitfalls with summary statistics

First, let us motivate the *compositional* aspect of the question. In microbiome research specifically, species diversity became the center of attention to an extent that asking “what is the causal effect of *the diversity* of a composition X on the outcome Y ?” appears more intuitive than asking for the causal effect of individual abundances. In fact, popular books and research articles alike seem to suggest that (bio-)diversity is indeed an important *causal driver* of ecosystem functioning and human health, even though these claims are largely grounded in observational, non-experimental data^{21,22}. Similar summary statistics or low-dimensional representations have been proposed in other domains such as in single-cell RNA data²³. We now explain why, even in situations where summary statistics appear to be useful proxies, no causal conclusions can be drawn from them.

Let us consider α -diversity as an example of a one-dimensional summary statistic of a microbiome measurement, e.g. $\alpha_{\text{Simpson}} = -\sum_{j=1}^p (x_j)^2$ or $\alpha_{\text{Shannon}} = -\sum_{j=1}^p x_j \log x_j$. The “causal effect” of the diversity α on some outcome of interest Y (e.g., health or disease indicator) is usually considered to be the expected value of Y under an *intervention on the diversity*, i.e., externally setting the diversity to a chosen value α^* , with all host and environmental factors unchanged. This causal effect is commonly denoted by $\mathbb{E}[Y \mid \text{do}(\alpha = \alpha^*)]$. When $\mathbb{E}[Y \mid \text{do}(\alpha = \alpha_1)] < \mathbb{E}[Y \mid \text{do}(\alpha = \alpha_2)]$ for two diversity values α_1, α_2 with $\alpha_1 < \alpha_2$, one would then be tempted to conclude that “increasing diversity α causes an increase in the outcome Y ”, which is often loosely translated to “diversity is a causal driver for health”. We now highlight critical issues with this approach.

(a) When considering the proposed causal effect estimand $\mathbb{E}[Y \mid \text{do}(\alpha = \alpha^*)]$ directly, one presupposes the existence of clearly defined interventions on α . However, there are infinitely many ways of changing the diversity of a composition by a fixed amount. This ‘many-to-one’ nature prevents a consistent conceptualization of external interventions. In particular, for a given value of α , there is a $(p-2)$ -dimensional subspace of $\mathbb{S}^{p-1}$ with that value of α . Hence, an intervention to “increase the diversity of a given composition” by some $\Delta\alpha$ is highly ambiguous. The different ways of achieving this change must be expected to have different implications for the outcome Y . Similarly, most common diversity measures are invariant under permutations of components and the above approach would require us to conclude that all $p!$ permutations of a composition are functionally completely equivalent with regard to the outcome Y —an abstruse claim. Hence, assigning causal powers to diversity by estimating $\mathbb{E}[Y \mid \text{do}(\alpha)]$ is highly ambiguous and does not carry the intended meaning. This concern is further exacerbated by the difficulty and ambiguity in measuring α -diversity in the first place^{7,24,25}.

(b) The definition of α -diversity is not unique, which could lead to a potential search for positive results by using a different metric²⁶ or contradictory causal claims. Consider two different one-dimensional summary statistics α_1, α_2 on $\mathbb{S}^{p-1}$. These can be defined in terms of their contours, i.e., the collection of $(p-2)$ -dimensional subspaces of $\mathbb{S}^{p-1}$ of constant values of α_1 and α_2 respectively. Since they are different, there will be a contour line of α_1 along which α_2 either increases or decreases. Along this path through compositions, we would have to conclude that the causal effect of one summary statistic is zero, while it is non-zero for the other. See Fig. 1 for a visualization. In typical scenarios, there is no “one correct” summary statistic, such that reliable claims even about the sign of the causal effect of a summary statistic of a composition become void.

Cause-effect estimation with instrumental variables

While researchers continue to develop predictive methods for compositional data¹⁶, in most scientific contexts causal effects are of greater interest. For example, the human microbiome co-evolves with its host and the external environment through diet, activity, climate, or geography, etc. leading to plentiful microbiome-host-environment interactions²⁷. Carefully designed studies may allow us to control for certain environmental factors

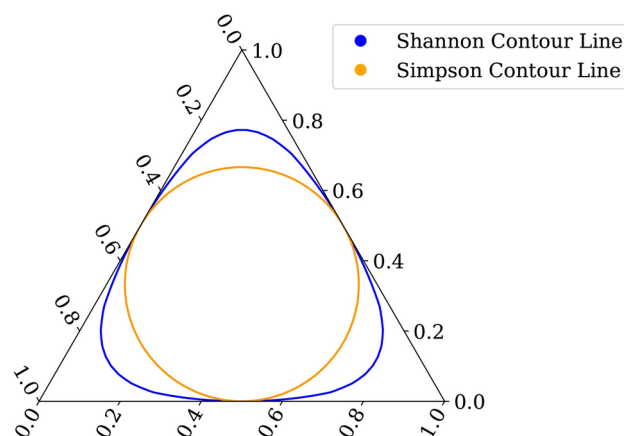


Fig. 1. The ternary plot shows an exemplary scenario with $p = 3$. The orange contour contains compositions for which the Simpson diversity is constant, while the blue contour shows compositions for which the Shannon diversity is constant. Shannon diversity changes along contours of Simpson diversity and vice versa.

and specifics of the host. In fact, several recent works studied the causal mediation effect of the microbiome on health-related outcomes, assuming all relevant covariates are observed and can be controlled for^{28–34}, or vice versa, the effect of environmental factors on the microbiome³⁵. However, in practice there is little hope of measuring *all* latent factors in these complex interactions. In such a situation, a purely predictive model will suffer from bias due to the unobserved confounders. Such unobserved confounders are a major hurdle in cause-effect estimation broadly and also specifically for compositional causes.

Concretely, without further assumptions, the direct causal effect $X \rightarrow Y$ is not identified from observational data in the presence of unobserved confounding $X \leftarrow U \rightarrow Y$ ³⁶. One common way to still identify the causal effect from purely observational data is through so-called instrumental variables (IV)³⁷. An *instrumental variable* Z is a variable that has an effect on the cause X ($Z \rightarrow X$), but is independent of the confounder ($Z \perp U$), and conditionally independent of the outcome given the cause and the confounder ($Z \perp Y \mid \{X, U\}$). In practice, it can be hard to find valid instruments for a target effect³⁸, but when they do exist, instrumental variables often render efficient cause-effect estimation possible.

In this work, we develop interpretable methods to estimate the direct causal effect of a *compositional cause* X on a continuous or categorical outcome Y within the IV setting. The question of whether and how cause-effect estimation for compositional treatments under unobserved confounding is possible remains unanswered in the literature, motivating our in-depth analysis of two-stage methods for interpretable cause-effect estimation of individual relative abundances on the outcome. In the analysis, we focus on a careful selection and combination of existing approaches and a thorough examination of potential pitfalls and mis-usage. Our extensive empirical evaluations carefully assess assumptions (additive noise, strong instruments) and model misspecification as a potential obstacle to interpretable and reliable effect estimates. We evaluate the efficacy and robustness of our proposed methods on both synthetic and real data from a mouse experiment, examining how the gut microbiome (X) affects body weight (Y) instrumented by sub-therapeutic antibiotic treatment (STAT) (Z).

The rest of the manuscript proceeds as follows. First, we introduce the concepts of compositional data and instrumental variables in detail. Following this introduction of our methods, we provide some simulation to study the advantage and potential pitfalls in using high-dimensional compositional data in instrumental variable settings. Last but not least, we then apply the methods to a real world dataset.

Methods

Instrumental variables

We briefly recap the assumptions of the instrumental variable setting as depicted in Fig. 2. For an *outcome* (or *effect*) Y , a *treatment* (or *cause*) X , and potential *unobserved confounders* U , we assume access to a discrete or continuous *instrument* $Z \in \mathbb{R}^q$ satisfying (i) $Z \perp U$ (the confounder is independent of the instrument), (ii) $Z \not\perp X$ (“the instrument influences the cause”), and (iii) $Z \perp Y \mid \{X, U\}$ (“the instrument influences the outcome only through the cause”). Our goal is to estimate the direct causal effect of X on Y , written as $\mathbb{E}[Y|do(x)]$ in the do-calculus notation³⁶ or as $\mathbb{E}[Y(x)]$ in the potential outcome framework³⁹, where $Y(x)$ denotes the potential outcome for treatment value x . The functional dependencies are $X = g(Z, U)$, $Y = f(X, U)$. While Z, X, Y denote random variables, we also consider a dataset of n i.i.d. samples $\mathcal{D} = \{(z_i, x_i, y_i)\}_{i=1}^n$ from their joint distribution. We arrange observations in matrices or vectors denoted by $\mathbf{X} \in \mathbb{R}^{n \times p}$ or $\mathbf{X} \in (\mathbb{S}^{p-1})^n$, $\mathbf{Z} \in \mathbb{R}^{n \times q}$, $\mathbf{y} \in \mathbb{R}^n$.

Without further restrictions on f and g , the causal effect is not identified^{40–42}. The most common assumption leading to identification is that of *additive noise*, namely $Y = f(X) + U$ with $\mathbb{E}[U] = 0$ but not necessarily $X \perp U$. Here, we overload the symbols f and g for simplicity. The implied Fredholm integral equation of first kind $\mathbb{E}[Y \mid Z] = \int f(x) dP(X \mid Z)$ is generally ill-posed. While the linear case is well understood³⁷, under certain regularity conditions the IV problem can be solved consistently even for non-linear f , see e.g.,^{43,44} and more recently^{45–48}. However, in this case the problem is typically under-specified in that multiple f are compatible with the observed data and regularization techniques are typically used to obtain a unique solution—typically the smallest compatible f according to some norm. It is thus difficult to interpret estimates of non-linear causal-effects in a way that aids understanding of the underlying processes.

In the simplest case, where $X \in \mathbb{R}^p$ and f, g are linear, a standard *instrumental variable estimator* is

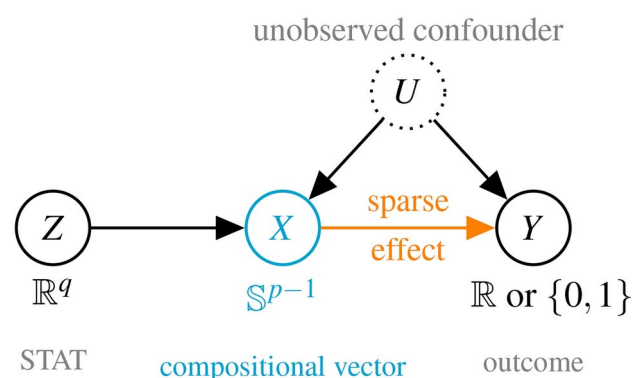


Fig. 2. Cause-effect estimation of $X \rightarrow Y$ via an instrumental variable Z for compositional X .

$$\hat{\beta}_{iv} = (X^T P_Z X)^{-1} X^T P_Z y \quad (1)$$

with $P_Z = Z(Z^T Z)^{-1} Z^T$ ³⁷. For the *just-identified* case $q = p$ as well as the *over-identified* case $q > p$, this estimator is consistent and asymptotically unbiased, albeit not unbiased. In the *under-identified* case $q < p$, where there are fewer instruments than treatments, the orthogonality of Z and U does not imply a unique solution. Again, regularization or other objectives such as sparsity assumptions have been proposed to obtain unique a unique solution within the space of compatible β ^{49–51}. The estimator $\hat{\beta}_{iv}$ can also be interpreted as the outcome of a *two-stage least squares* (2SLS) procedure consisting of (1) regressing X on Z via OLS $\hat{\delta} = (Z^T Z)^{-1} Z^T X$, and (2) regressing y on the predicted values $\hat{X} = Z\hat{\delta}$ via OLS, again resulting in $\hat{\beta}_{iv}$. Practitioners are typically discouraged from using the manual two-stage approach, because the OLS standard errors of the second stage are wrong—a correction is needed³⁷. However, we note that the point estimator obtained by the manual two-stage procedure is equivalent to Eq. (1).

Moreover, the two-stage description suggests that the two-stages are independent problems and thereby seems to invite us to mix and match different regression methods as we see fit.

The authors in³⁷ highlight that the asymptotic properties of $\hat{\beta}_{iv}$ rely on the fact that for OLS the residuals of the first stage are uncorrelated with the instruments Z . Hence, for OLS we achieve consistency of $\hat{\beta}_{iv}$ *even when the first stage is misspecified*. For a non-linear first stage regression we may only hope to achieve uncorrelated residuals asymptotically when the model is correctly specified. Replacing the OLS first stage with a non-linear model is known as the “forbidden regression”, a term commonly attributed to Prof. Jerry Hausmann.

Angrist and Pischke acknowledge that the practical relevance of the forbidden regression is not well understood. However, the rather strong term tries to warn against a careless use of two consecutive regression models for the two stages. First, when also the second stage is assumed to be non-linear, one would require independence of the first stage residuals from Z . Secondly, if this requirement could not be guaranteed, both models would have to be specified correctly. Something that—in practice—can seldom be guaranteed. Nevertheless, nonlinear models for the IV setting are definitely feasible.

Starting with⁵² there is now a rich literature on the circumstances under which “manual 2SLS” with non-linear first (and/or second) stage can yield consistent causal estimators. Primarily interested in *compositional* treatments X , we cannot directly use OLS for either stage. Since there is no theoretical guidance for this case, we assess our options empirically, paying great attention to potential issues due to the “forbidden regression” and misspecification in our proposed methods.

Compositional data

Simplex geometry

The authors in¹ introduced the *perturbation* and *power transformation* as the simplex $\mathbb{S}^{p-1}$ counterparts to addition and scalar multiplication of Euclidean vectors in $\mathbb{R}^p$:

$$\begin{array}{c|c} \text{Perturbation} & \text{Power transformation} \\ \oplus : \mathbb{S}^{p-1} \times \mathbb{S}^{p-1} \rightarrow \mathbb{S}^{p-1} & \odot : \mathbb{R} \times \mathbb{S}^{p-1} \rightarrow \mathbb{S}^{p-1} \\ x \oplus w = C(x_1 w_1, \dots, x_p w_p) & a \odot x := C(x_1^a, x_2^a, \dots, x_p^a) \end{array}$$

Here, the *closure operator* $C : \mathbb{R}_{\geq 0}^p \rightarrow \mathbb{S}^{p-1}$ normalizes a p -dimensional, non-negative vector to the simplex $C(x) := x / \sum_{i=1}^p x_i$. Together with the dot-product

$$\langle x, w \rangle := \frac{1}{2p} \sum_{i,j=1}^p \log\left(\frac{x_i}{x_j}\right) \log\left(\frac{w_i}{w_j}\right) \quad (2)$$

the tuple $(\mathbb{S}^{p-1}, \oplus, \odot, \langle \cdot, \cdot \rangle)$ forms a finite-dimensional real Hilbert space⁵³ allowing to transfer usual geometric notions such as lines and circles from Euclidean space to the simplex.

Coordinate representations

The p entries of a composition remain dependent via the unit sum constraint, leading to $\mathbb{S}^{p-1}$ having dimension $p - 1$. To deal with this fact, different invertible log-based transformations have been proposed, for example the additive log ratio, centered log ratio¹, and isometric log ratio⁵⁴ transformations

$$\text{alr}(x) = V_a \log(x) \in \mathbb{R}^{p-1}, \quad \text{clr}(x) = V_c \log(x) \in \mathbb{R}^p, \quad \text{ilr}(x) = V_i \log(x) \in \mathbb{R}^{p-1}, \quad (3)$$

where the logarithm is applied element-wise and the matrices $V_a, V_i \in \mathbb{R}^{(p-1) \times p}$ and $V_c \in \mathbb{R}^{p \times p}$ are defined in Supplementary Material S2. While alr is a vector space isomorphism that preserves a one-to-one correspondence between all components except for one, which is chosen as a fixed reference point to reduce the dimensionality (we choose x_p , but any other component works), it is not an isometry, i.e., it does not preserve distances or scalar products. Both clr and ilr are also isometries, but clr only maps onto a subspace of $\mathbb{R}^p$, which often renders measure theoretic objects such as distributions degenerate. As an isometry between $\mathbb{S}^{p-1}$ and $\mathbb{R}^{p-1}$, ilr allows for an orthonormal coordinate representation of compositions. However, it is hard to assign meaning to the

individual components of $\text{ilr}(x)$, which all entangle a different subset of relative abundances in x leading to challenges for interpretability⁵⁵. Therefore, alr remains a useful tool in statistical analyses where interpretability is required despite the lack of the isometric property.

Log-contrast estimation

The key advantage of such coordinate transformations is that they allow us to use regular multivariate data analysis methods (typically tailored to Euclidean space) for compositional data. For example, we can directly fit a linear model $y = \beta_0 + \beta^T \text{ilr}(x) + \epsilon$ on the ilr coordinates via ordinary least squares (OLS) regression. However, in real-world datasets, p is often a large number capturing “all possible components in a measurement”, leading to $p \gg n$ with each of the n measurements being sparse, i.e., a substantial fraction of x being zero. Moreover, in many (especially high-dimensional) situations only few components exert direct causal influence on the outcome. Both overparameterization $p \gg n$ as well as assuming sparse effects call for regularization. The problem with enforcing sparsity in a “linear-in- ilr ” model is that a zero entry in β does not correspond directly to a zero effect of the relative abundance of any single component. This motivates *log-contrast* estimation⁵⁶ with a sparsity penalty^{57–59}

$$\min_{\beta} \sum_{i=1}^n \mathcal{L}(x_i, y_i, \beta) + \lambda \|\beta\|_1 \quad \text{s.t.} \quad \sum_{i=1}^p \beta_i = 0. \quad (4)$$

In our examples, we focus mostly on continuous $y \in \mathbb{R}$ and the squared loss $\mathcal{L}(x, y, \beta) = (y - \beta^T \log(x))^2$. However, our framework also supports the Huber loss for robust log-contrast regression as well as an optional joint concomitant scale estimation for both losses^{59,60}. Moreover, for classification tasks with $y \in \{0, 1\}$, we can directly use the squared Hinge loss (or a “Huberized” version thereof) for $\mathcal{L}$, see Supplementary Material S7 for details. These flexible estimation formulations respect the compositional nature of x while retaining the association between the entry β_i and the relative abundance of the individual component x_i . Even though, due to the additional sum constraint, individual components of β are still not—and can never be—entirely disentangled.

Logs and zeros

In the previous paragraphs, we introduced multiple log-based coordinate representations for compositions and at the same time claimed that measurements are often sparse in relevant settings. Since the logarithm is undefined for zero entries, a simple strategy is to add a small constant to all absolute counts, so called *pseudo-counts*^{61,62}. These pseudo-counts are particularly popular in the microbiome and single-cell RNA literature where there are many more possible taxa/genes (up to tens of thousands) that occur in any given sample. Despite the simplicity of adding a constant pseudo-count, for example 0.5, recent work gives theoretical and empirical evidence for this approach⁶³, which we also use here.

Summary statistics

Traditionally, interpretability issues around compositions have been circumvented by focusing on summary statistics instead of individual relative abundances. One of the key measures to describe ecological populations is *diversity*. Diversity captures the variation within a composition and is in this context often called α -diversity. There is no unique definition of α -diversity. Among the most common ones in the literature are *richness*, i.e. the number of non-zero entries denoted as $\|x\|_0$, *Shannon diversity* $-\sum_{j=1}^p x_j \log(x_j)$ and *Simpson diversity* $-\sum_{j=1}^p x_j^2$. Especially in the microbial context, there exist entire families of diversity measures taking into account species, functional, or phylogenetic similarities between taxa and tracing out continuous parametric profiles for varying sensitivity to highly-abundant taxa. See for example^{64–66} for an overview of the possibilities and choices of estimating α -diversity in a specific application. While the popularity of α -diversity for assessing the impact and health of microbial compositions⁶⁷ seemingly renders it a natural choice for causal queries, we argue that such claims are misleading and void of a solid foundation.

Methods for higher dimensional causes

In this section we develop methods to reason about the effects of hypothetical interventions on the relative abundance of individual components from observational data.

- **2SLS:** As the first baseline, we run 2SLS from Eq. (1) directly on $X \in \mathbb{S}^{p-1}$ ignoring its compositional nature.
- **Only LC** For completeness, as the second baseline, we run log-contrast (LC) estimation for the second stage only, thereby entirely ignoring confounding.
- **2SLS_{ILR}:** 2SLS with $\text{ilr}(X) \in \mathbb{R}^{p-1}$ as the treatment; since OLS minima do not depend on the chosen basis, parameter estimates for different log-transformations of X are related via fixed linear transformations. Hence, as long as no sparsity penalty is added, ilr and alr regression yield equivalent results. The isometric ilr coordinates are useful due to the consistency guarantees of 2SLS given that $Z^T \text{ilr}(X)$ has full rank. For interpretability, alr coordinates can be beneficial as individual coordinates correspond to individual components (given a reference). The respective coordinate transformations are given in Supplementary Material S2.

- KIV_{ILR} : Following⁴⁵ we replace OLS in $2SLS_{ILR}$ with kernel ridge regression in both stages to allow for non-linearities. Like $2SLS_{ILR}$, KIV_{ILR} cannot enforce sparsity in an interpretable fashion.
- $ILR+LC$: To account for sparsity, we use sparse log-contrast estimation (see Eq. (4)) for the second stage, while retaining OLS to ilr coordinates for the first stage. Log-contrast estimation conserves interpretability in that the estimated parameters correspond directly to the effects of individual relative abundances.
- $DIR+LC$: Finally, we circumvent log-transformations entirely and deploy regression methods that naturally work on compositional data in both stages. For the first stage, we use a Dirichlet distribution—a common choice for modeling compositional data—where $X | Z \sim \text{Dirichlet}(\alpha_1(Z), \dots, \alpha_p(Z))$ with density $B(\alpha_1, \dots, \alpha_p)^{-1} \prod_{j=1}^p x_j^{\alpha_j-1}$ where we drop the dependence of $\alpha = (\alpha_1, \dots, \alpha_p) \in \mathbb{R}^p$ on Z for simplicity. With the mean of the Dirichlet distribution given by $\alpha / \sum_{j=1}^p \alpha_j$, we account for the Z -dependence via $\log(\alpha_j(Z_i)) = \omega_{0,j} + \omega_j^T Z_j$. We then estimate the newly introduced parameters $\omega_{0,j} \in \mathbb{R}$ and $\omega_j \in \mathbb{R}^q$ via maximum likelihood estimation with ℓ_1 regularization. For the second stage we again resort to sparse log-contrast estimation. If the non-linear first stage is misspecified, the “forbidden regression” bias may distort effect estimates of this approach. This is contrasted by Dirichlet regression potentially resulting in a better fit of the data than linearly modeling log-transformations. We highlight that only $ILR+LC$ and $DIR+LC$ accommodate all relevant requirements: (i) unobserved confounding, (ii) compositional treatments, (iii) sparse effects, and (iv) interpretable estimates. Additionally, we wish to emphasize the concept of “forbidden regression” as discussed in³⁷. This issue arises specifically when a nonlinear first stage is improperly combined with a subsequent regression stage. With regard to the proposed methods, this issue will only affect $DIR+LC$, which we will further assess in the simulations. Both $2SLS$ and $Only LC$ are inherently consistent. Meanwhile, $ILR+LC$, $2SLS_{ILR}$, and KIV_{ILR} derive their consistency from linear techniques resp.⁴⁵, as the log-transformation is itself linear. Therefore, the linear transformation ensures that there is no “forbidden” non-linearity introduced in the first stage regression step.

Simulation studies

Data generation

For the evaluation of our methods we require the ground truth causal effect to be known. Since unobserved confounders (and thus counterfactuals) are never observed in practice (by definition), this can only be achieved via synthetic data. We simulate data (in two different settings) to maintain control over ground truth effects, confounding strength, potential misspecification, and the strength of instruments (see Fig. 3).

- **Setting A:** The first setting is

$$\begin{aligned} Z_j &\sim \text{Unif}(0, 1), \quad U \sim \mathcal{N}(\mu_c, 1), \\ \text{ilr}(X) &= \gamma_0 + \gamma^T Z + U c_X, \quad Y = \beta_0 + \beta^T \text{ilr}(X) + U c_Y, \end{aligned} \quad (5)$$

where we model $\text{ilr}(X) \in \mathbb{R}^{p-1}$ directly and $\mu_c, c_Y \in \mathbb{R}$, $\gamma_0, c_X \in \mathbb{R}^{p-1}$, $\gamma \in \mathbb{R}^{q \times (p-1)}$ are fixed up front. Our goal is to estimate the causal parameters $\beta \in \mathbb{R}^{p-1}$ and the intercept $\beta_0 \in \mathbb{R}$. This setting satisfies the standard 2SLS assumptions (linear, additive noise) and all our linear methods are thus *wellspecified*. To explore effects of *misspecification*, we also consider the same setting only replacing (using $\mathbf{1} = (1, \dots, 1)$)

$$Y = \beta_0 + \frac{1}{100} \mathbf{1}^T (\text{ilr}(X) + 1)^2 + 10 \cdot \mathbf{1}^T \sin(\text{ilr}(X)) + c_Y U. \quad (6)$$

- **Setting B:** We consider a sparse effect model for $X \in \mathbb{S}^{p-1}$ which is more realistic for higher-dimensional compositions. Note that some parameter dimensions are different, i.e., the same symbols have different meanings in the settings A and B. With $\mu = \gamma_0 + \gamma^T Z$ for fixed $\gamma_0 \in \mathbb{R}^p$, $\gamma \in \mathbb{R}^{q \times p}$ we use

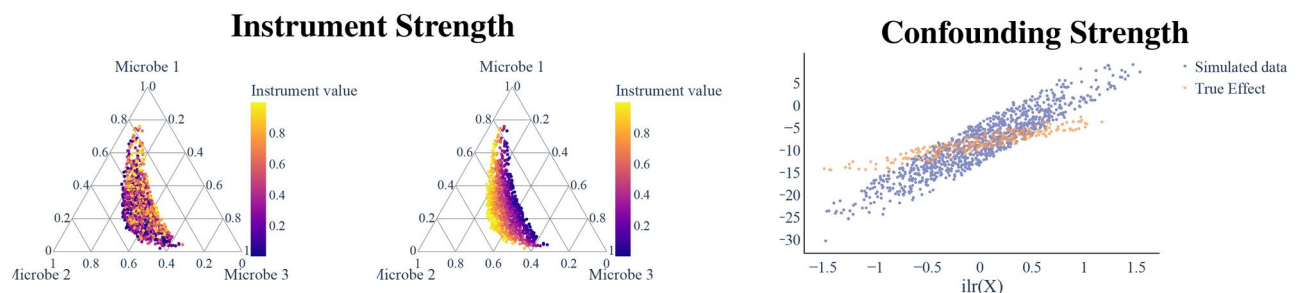


Fig. 3. Visualization of a Setting A ($p = 3, q = 2$): The left panel shows both a weak (left) and a strong (right) instrument, i.e. with the *instrument value* Z either barely influencing the the composition of the microbiome (left) or strongly impacting the composition of the microbiome (right). The right panel shows a discrepancy between the true causal effect and the observed effect which stems from a confounding factor.

$$\begin{aligned}
 Z_j &\sim \text{Unif}(Z_{\min}, Z_{\max}), & U &\sim \text{Unif}(U_{\min}, U_{\max}), \\
 X &\sim C(\text{ZINB}(\mu, \Sigma, \theta, \eta)) \oplus (U \odot \Omega_C), \\
 Y &= \beta_0 + \beta^T \log(X) + c_Y^T \log(U \odot \Omega_C).
 \end{aligned}
 \tag{7}$$

The treatment X is assumed to follow a zero-inflated negative binomial (ZINB) distribution⁶⁸, commonly used for modelling count data with excess zeros⁶⁹. Here, $\eta \in (0, 1)$ is the probability of zero entries, $\Sigma \in \mathbb{R}^{p \times p}$ is the covariance matrix, and $\theta \in \mathbb{R}$ the shape parameter. The confounder $U \in [U_{\min}, U_{\max}]$ perturbs this base composition in the direction of another fixed composition $\Omega_C \in \mathbb{S}^{p-1}$ scaled by U . In simplex geometry $x_0 \oplus (U \odot x_1)$ corresponds to a line starting at x_0 and moving along x_1 by a fraction U . A linear combination of the log-transformed perturbation enters Y additively with weights $c_Y \in \mathbb{R}^p$ controlling confounding strength. All other parameter choices are given in Supplementary Material S6. This setting is linear in how Z enters μ and how U enters X and Y in the simplex geometry. All our two-stage models are (intentionally) misspecified in the first stage for setting B.

The precise choices of all parameters for the different empirical evaluations are described in the appendix (Supplementary Material S6). All relevant code is available at <https://github.com/EAILer/causal-compositions>.

Metrics and evaluation

Appropriate evaluation metrics are key for cause-effect estimation tasks. We aim at capturing the average causal effect (under interventions) and the causal parameters when warranted by modeling assumptions. When the true effect is linear in $\log(X)$, we can compare the estimated causal parameters $\hat{\beta}$ from 2SLS_{ILR}, ILR+LC, and DIR+LC with the ground truth β directly. In these linear settings, we report causal effects of individual relative abundances X_j on the outcome Y via the mean squared difference (β -MSE) between the true and estimated parameters β and $\hat{\beta}$. Moreover, we also report the number of falsely predicted non-zero entries (FNZ) and falsely predicted zero entries (FZ), which are most informative in sparse settings and metrics of key interest to biostatisticians.

In the general case, where a measure for identification of the interventional distribution $P(Y \mid do(X))$ is not straightforward to evaluate, we focus on the *out of sample error* (OOS MSE): For the true causal effect we first draw an i.i.d. sample $\{x_i\}_{i=1}^m$ from the data generating distribution (that are not in the training set, i.e., out of sample) and compute $\mathbb{E}_U[f(x_i, U)]$ for the known $f(X, U)$, the expected Y under intervention $do(x_i)$. We use $m = 250$ for all experiments. OOS MSE is then the mean square difference to our second-stage predictions $\hat{f}(x_i)$ on these out of sample x_i . Because in real observational data we do not have access to $P(Y \mid do(X))$ (but only the conditional distribution $P(Y \mid X)$), we can not evaluate OOS MSE in real-world observational data.

We run each method for 50 random seeds in setting A (Eq. (5)), and 20 random seeds in setting B (Eq. (7)). In result tables, we report mean and standard error over these runs. The sample size is $n = 1000$ in the low-dimensional case ($p = 3$) and $n = 10,000$ in the higher-dimensional cases ($p = 30, p = 250$). Additionally, we report results for an overparameterized setting with $n = 100$ and $p = 250$. Supplementary Material S6 and S8 contain further explanations and more detailed results. Note, that since alr coordinates for X yield equivalent optimization minima as ILR+LC, we only report results from ILR+LC. All numbers match precisely for ALR+LC in our empirical evaluation.

Results for low-dimensional compositions

We first consider settings A and B with $p = 3$ and $q = 2$. The top section of Tables 1 and 2 shows our metrics for all methods. First, effect estimates are far off when ignoring the compositional nature (2SLS) or the confounding (Only LC) as expected. Also, recent non-linear IV methods such as^{47,70,71} could not overcome the issues of 2SLS in this setting. Without sparsity in the second stage, 2SLS_{ILR} and ILR+LC yield equivalent estimates in this low-dimensional linear setting—we only report ILR+LC. ILR+LC (and equivalent methods) succeed in cause-effect estimation under unobserved confounding: they recover the true causal parameters with high precision on average (low β -MSE) and thus achieve low OOS MSE. While DIR+LC performs reasonably well in setting A, setting B surfaces that despite being a seemingly plausible approach with powerful regression techniques, DIR+LC suffers substantially under a misspecified first-stage.

Results for high-dimensional compositions

We now consider the challenging cases $p = 30$ and $p = 250$ with $q = 10$ and sparse ground truth β for settings A and B (8 non-zeros: 3 times -5 and 5 and once -10 and 10) in the bottom sections of Tables 1 and 2. ILR+LC deals well with sparsity: unlike Only LC, it identifies non-zero parameters perfectly (FZ = 0) and rarely predicts false non-zeros. It also identifies the true β and accordingly predicts interventional effects (OOS MSE) well. DIR+LC and 2SLS_{ILR} fail entirely in these settings because the optimization does not converge. While we could get KIV_{ILR} to return a solution, tuning the kernel hyperparameters for high-dimensional ilr coordinates becomes increasingly challenging, which is reflected in poor OOS MSE. In Fig. 4 we show detailed results for the most challenging setting (setting B with $p = 250$ and $q = 10$) including the OOS MSE (left), recovery of individual non-zero coefficients (middle), and recovery of zero coefficients (right). Analogous plots for all other settings can be found in Supplementary Material S8.

Robustness checks

Due to the inherent entanglement via the unit sum constraint, analyses involving compositional data are generically hard to interpret. Causal analyses involving compositions in the instrumental variable setting are

		Setting A, Equation (5)			
Dim.	Method	OOS MSE	β -MSE	FZ	FNZ
$p = 3q = 2$	DIR+LC	0.58±0.08	1.6±0.17	0.0	0.0
	ILR+LC [†]	0.37±0.07	1.1±0.15	0.0	0.0
	KIV _{ILR}	0.37±0.07	–	–	–
	Only LC	15.03±0.20	32.6±0.14	0.0	0.0
	2SLS	> 200	> 5k	0.0	0.0
$p = 30q = 10$	ILR+LC	0.42±0.08	0.22±0.01	0.0	12.0
	KIV _{ILR}	240.6±35.7	–	–	–
	Only LC	24.4±0.37	1.9±0.00	0.0	12.3
$p = 250q = 10$	ILR+LC	0.67±0.14	0.22±0.02	0.0	0.0
	KIV _{ILR}	5060.5±1196.2	–	–	–
	Only LC	30.8±0.48	143.3±0.27	3.0	1.0

Table 1. Results for setting A (fully linear in $\text{ilr}(X)$). Bold values indicate the lowest OOS MSE resp. β -MSE in the corresponding dimension scenario. [†] Identical to 2SLS_{ILR} in low-dimensional setting without sparsity

		Setting B, Equation (7)			
Dim.	Method	OOS MSE	β -MSE	FZ	FNZ
$p = 3q = 2$	DIR+LC	> 10k	> 2k	0.0	0.0
	ILR+LC [†]	20.3±4.85	9.9±3.3	0.0	0.0
	KIV _{ILR}	19.2±4.42	–	–	–
	Only LC	269.0±6.85	129.8±2.13	0.0	0.0
	2SLS	> 15k	> 300k	0.0	0.0
$p = 30q = 10$	ILR+LC	120.4±25.1	37.0±15.1	0.0	13.3
	KIV _{ILR}	287.2±19.1	–	–	–
	Only LC	3863.8±166.3	458.8±12.2	3.4	15.8
$p = 250q = 10$	ILR+LC	99.1±7.9	24.4±4.30	0.13	0.39
	KIV _{ILR}	622.8±30.1	–	–	–
	Only LC	3366.0±166.3	498.3±17.2	6.9	1.9

Table 2. Results for setting B (first stage ZINB with sparse effects in higher dimensions), where all our two-stage methods are (intentionally) misspecified in the first stage. Bold values indicate the lowest OOS MSE resp. β -MSE in the corresponding dimension scenario. [†] Identical to 2SLS_{ILR} in low-dimensional setting without sparsity

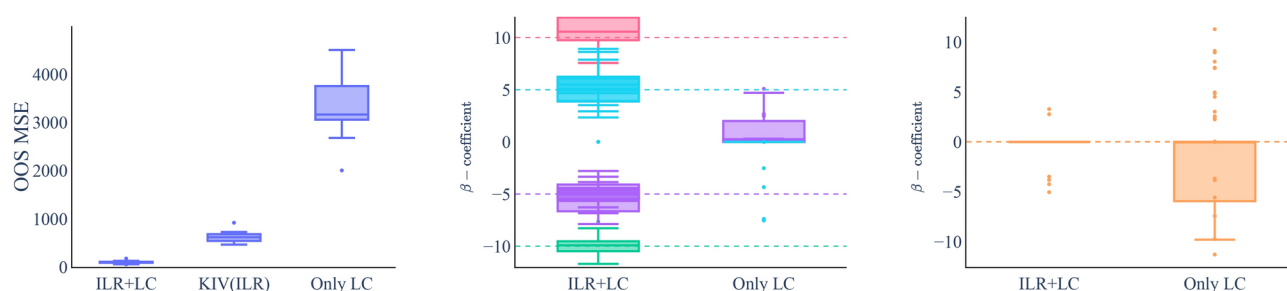


Fig. 4. Boxplots of the results for setting B in Table 2 with $p = 250, q = 10$. We show OOS MSE (left), recovery of non-zero β coefficients (middle), and recovery of zero β coefficients.

further challenged by potential violations of assumptions such as weak instruments or misspecification. We assess the sensitivity of our proposed methodology to such potential pitfalls in the following scenarios.

Weak instruments

“Strong instruments” are a prerequisite for successful two-stage estimation in the instrumental variable setting and one of the key discussion points in real-world applications of IV. Nevertheless, how to measure instrument validity is not unambiguously clarified, there is no clear definition on which instruments are weak and which

instruments are strong. The categories rely on heuristics and empirically derived best practices. In the linear setting, instrument strength for $p = 1$ can be approximated via a first-stage F-statistic with a value greater than 10 generally being considered sufficient to avoid weak instrument bias in 2SLS⁷². For $p > 1$, measuring instrument strength is more challenging even in the linear case⁷³. Therefore, we report first-stage F-statistics for each dimension of X as a proxy for instrument strength.

When instruments are weak, the estimation bias can theoretically become arbitrarily large (even in the limit of infinite data). To assess the sensitivity of our methods to weak instrument bias, we re-analyze setting A ($p = 3$ and $q = 2$) only changing the dependence of X and Z to be weak with first-stage F-statistic values of 6.9 and 4.7 for the two components of $\text{ilr}(X)$. In the linear setting, we can directly control instrument strength via α (see Eq. (5)).

The first row in Table 3 summarizes our results for weak instruments: the two-stage methods have a substantially higher variation in their estimates, both for OOS MSE and β compared to the strong instrument setting in Tables 1 and 2. As the second stage has not changed, Only LC performs equally bad. Notably, while the wellspecified two-stage methods ILR+LC and 2SLS_{ILR} seem to do worse than Only LC, the large OOS MSE and β -MSE are mostly due to outliers. Fig. 5 shows that the range of β estimates still cover the true values for ILR+LC and 2SLS_{ILR}, while Only LC is systematically off with low variance (confidently wrong). DIR+LC now not only suffers from the misspecified first stage but also the weak instrument resulting in virtually useless estimates. The surprisingly good performance of KIV_{ILR} in this specific setting is unexpected and cannot be consistently reproduced over different weak instrument scenarios: the performance is highly volatile and often worse than ILR+LC. Therefore, despite the good performance for these specific parameters, we find that more flexible methods are also affected heavily by weak instruments. In general, while two-stage estimates generally cannot be broadly trusted when instruments are weak, reverting to Only LC is potentially even more detrimental because the estimated coefficients are systematically off.

Non-linear second stage

Well-specification is typically impossible to ascertain in practice and most real-world examples are likely not perfectly linear even when the linearity assumption can be defended. Therefore, we introduce a non-linear f for setting A with $p = 3$ and $q = 2$ (Equation (6)), resulting in a misspecified second stage for all our methods except KIV_{ILR}, which can in principle capture non-linearities. Note that β cannot be interpreted directly as causal parameters when the true causal effect depends non-linearly on $\text{ilr}(X)$. The results in the second row of Table 3 show that DIR+LC (doubly misspecified) and 2SLS (ignoring compositionality) again fail. Moreover, in this non-linear scenario KIV_{ILR} beats ILR+LC (both still outperforming Only LC) and we expect the difference to grow as the non-linearity of f increases.

Scarce data

Finally, we return to the original setting A ($p = 250$, $q = 10$, linear in both stages), but mimic a scarce data scenario with $n = 100$. The third row in Table 3 clearly highlights again how the lack of regularization becomes problematic for 2SLS_{ILR} and KIV_{ILR}. Compared to the larger dataset, also our regularized two-stage methods naturally exhibit higher variation in their estimates. Notably, Only LC appears to compare favorably to ILR+LC in OOS MSE, but β -MSE surfaces its failure to accurately recover causal parameters. We thus conclude that despite increased variability, the ILR+LC is still better equipped to recover β in the small data regime (see Fig. 5).

Case study on murine sub-therapeutic antibiotic treatment

We consider the mouse dataset described by⁷⁴ and analyzed in³⁰ using causal mediation. A total of 57 newborn mice were assigned randomly to a sub-therapeutic antibiotic treatment (STAT) during their early stages of

Scenario	Method	OOS MSE	β -MSE
Weak Instruments	DIR+LC	7.1±2.6	43.2±10.3
	ILR+LC [†]	3.6±1.3	26.0±5.6
	KIV _{ILR}	2.7±0.9	–
	Only LC	15.9±0.20	52.2±0.18
	2SLS	> 100	> 5k
Non-Linearity	DIR+LC	135.6±6.34	–
	ILR+LC [†]	92.0±1.2	–
	KIV _{ILR}	73.4±2.28	–
	Only LC	104.1±1.43	–
	2SLS	> 300	–
Scarce Data	ILR+LC	45.1±7.90	72.8±8.3
	KIV _{ILR}	290.8±62.5	–
	Only LC	40.5±1.00	196.4±8.7
	2SLS _{ILR}	> 10k	> 2 · 10 ²⁴

Table 3. Results for various robustness checks. Bold values indicate the lowest OOS MSE resp. β -MSE in the relevant dimension scenario. [†] Identical to 2SLS_{ILR} in low-dimensional setting without sparsity

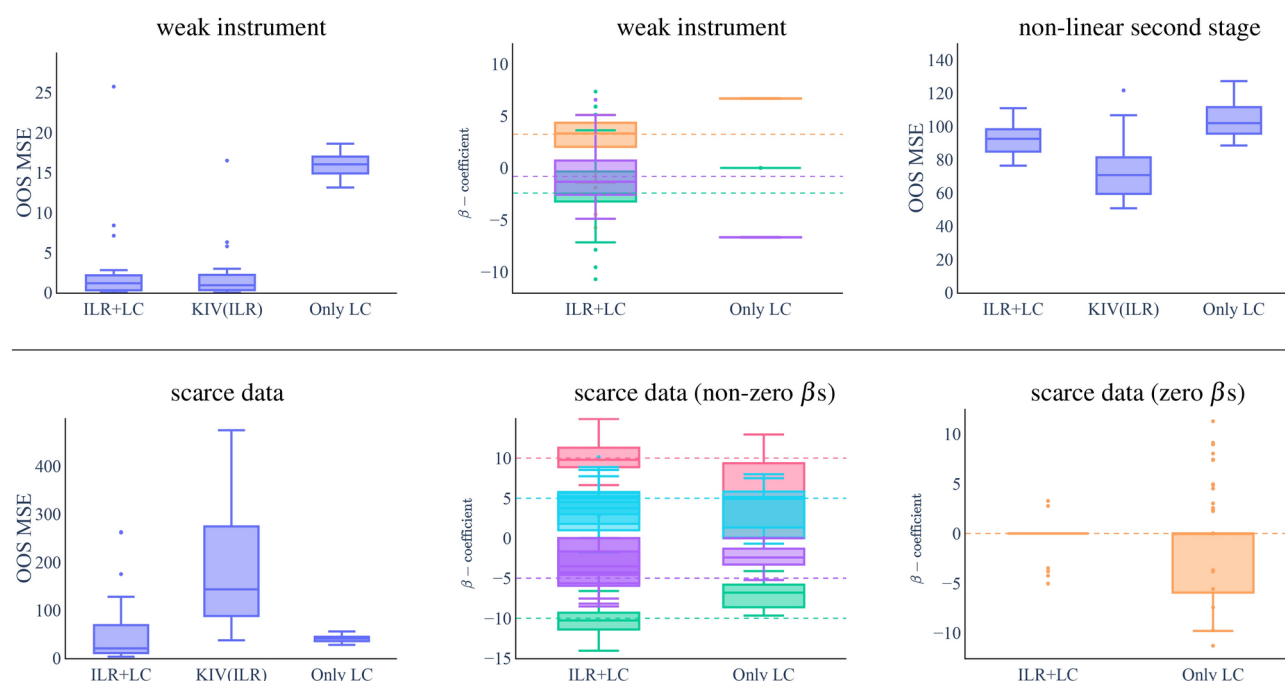


Fig. 5. We display OOS MSE and β -MSE (for non-zero coefficients and where applicable) for our robustness checks. All results and further visualizations are in Supplementary Material S8.

development. Sub-therapeutic antibiotic treatment means that the administered doses of antibiotics are too small to be detectable in the mice' bloodstream. There were 35 mice in the treatment group and 22 mice in the control group. After 21 days, the gut microbiome composition of each mouse was recorded. We are interested in the causal effect of the gut microbiome composition on body weight $Y \in \mathbb{R}$ of the mice (at sacrifice).

We assume a valid instrument due to the following characteristics in the data generation: The random assignment of the antibiotic treatment ensures independence of potential confounders such as genetic factors ($Z \perp U$). The sub-therapeutic dose implies that antibiotics can not be detected in the mice' blood, providing reason to assume no effect of the antibiotics on the weight other than through its effect on the gut microbiome ($Z \perp Y \mid \{U, X\}$).

Finally, we observe empirically, that there are statistically significant differences of microbiome compositions between the treatment and control groups ($Z \not\perp X$) based on the first stage F-statistic. Thus, the sub-therapeutic antibiotic treatment is a good candidate for an instrument $Z \in \{0, 1\}$ in estimating the effect $X \rightarrow Y$. Note, however, that this work is focused on methods rather than novel biological insights as more scrutiny of the IV assumptions would be required for substantive biological claims.

Figure 6 highlights the two most influential microbes on the genus level for our two-stage ILR+LC estimator and to the non-causal baseline Only LC, respectively. In the causal setting, we estimate the log-ratio of *Blautia* to *Anaerostipes* to be most influential for weight gain whereas standard log-contrast regression deems the ratio of an unclassified *Enterobacteria* genus to *Lactobacillus* to be the most predictive genus pair. This discrepancy suggests that the second stage might be subject to confounding. However, the mediation analysis on the same dataset in³⁰ posits a negative mediation effect of *Lactobacillus* on weight gain, consistent with the non-causal baseline model. This highlights the fact that different causal models provide alternative interpretation of the data that can only resolved by follow-up biological experiments.

Finally, we also assessed the influence of taxonomic aggregation levels and different loss functions as well as different values of zero counts on the results (see Supplementary Material S5). We observed that our causal model is robust to the choice of the loss function in terms of selected taxa whereas the baseline model found loss-function dependent sets of predictive taxa (see Supplementary Material Figure S2).

Discussion

In this work, we developed and analyzed methods for cause-effect estimation with compositional causes under unobserved confounding in instrumental variable settings. First, we succinctly expose that the common portrayal of summary statistics as a decisive (rather than merely descriptive) description of compositions is misguided. Instead, we advocate for causal effects to be estimated from the entire composition vector directly to establish meaningful and interpretable causal links. As a result, analysts cannot tap into a collection of well established cause-effect estimation tools for scalar data, but are instead faced with a large number of possible components (calling for sparsity-enforcing methods) and typically have to deal with unobserved confounding. Given the potentially profound impact of microbiome or single cell RNA data on advancing human health or of species abundances on global health, it is of vital importance that we face these challenges and develop interpretable methods to obtain causal insights from compositional data.

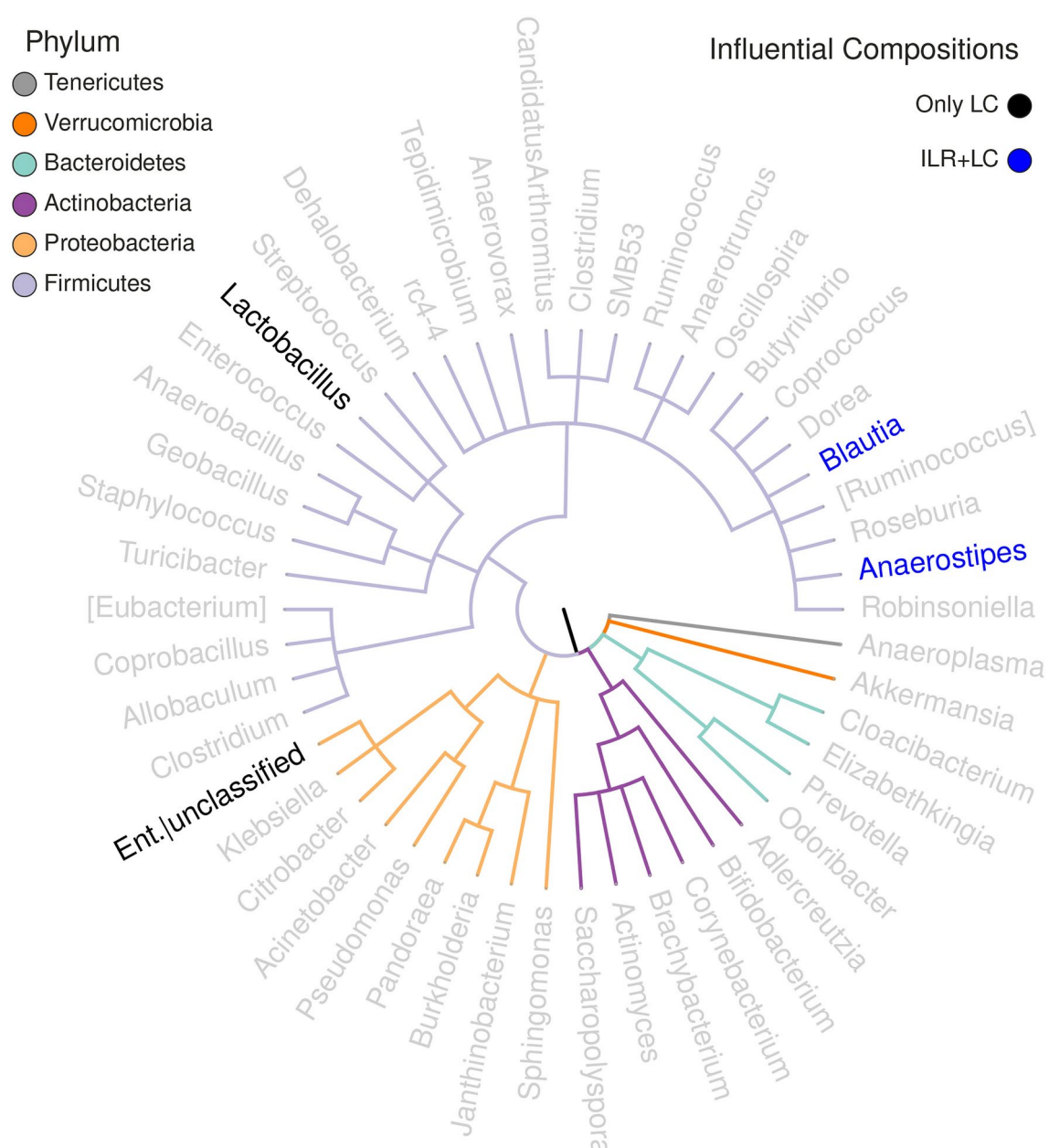


Fig. 6. Taxonomic tree of the microbiome data at genus level. The influential log-ratios for both Only Second LC and ILR+LC are highlighted in black and blue, respectively.

To this end, we carefully developed and assessed the effectiveness of various methods to not only reliably recover causal effects (OOS MSE), but also yield *interpretable and sparse* effect estimates for individual abundances (β -MSE, FZ, FNZ) whenever applicable. We also put special emphasis on how IV assumptions (misspecification, weak instrument bias) interact with compositionality. Our extensive empirical results for different two-stage methods highlight that accounting for the compositional nature as well as confounding is not optional. The overall failure of DIR+LC shows that not any seemingly suitable compositional technique can be trusted to yield reliable estimates in a manual two-stage procedure—careful analysis is needed. We have identified ILR+LC, to work reliably in wellspecified sparse and non-sparse settings as well as being relatively robust to first- and (small) second-stage misspecifications (i.e., non-linearities) and scarce data. It also yields interpretable estimates for individual components. When interpretability is not required or second-stage non-linearities are strong, KIV_{ILR} can still perform well under these relaxed assumptions albeit being challenging to tune for large p and unable to incorporate sparsity. As expected, valid instruments are required for all our two-stage methods. Taken together, our results on the efficacy and robustness of our methods in simulation and on real microbiome data provide first recommendations for practitioners to fully integrate compositional data into cause-effect estimation.

Data availability

All relevant code and data to reproduce the experiments, results and visualizations is available at <https://github.com/EAiler/causal-compositions>. The folder “/input/data” includes the data files (.Rdata) as well as the preprocessing steps (.Rmd). The folder “/notebooks” includes all relevant steps to understand the analysis steps and the visualizations, while the folder “/src” includes relevant background code.

Received: 10 June 2024; Accepted: 4 February 2025

Published online: 12 February 2025

References

- Aitchison, J. The statistical analysis of compositional data. *J. Roy. Stat. Soc.: Ser. B (Methodol.)* **44**, 139–160 (1982).
- Blei, D. M. & Lafferty, J. D. Correlated topic models. *Advances in Neural Information Processing Systems* 147–154 (2005).
- Blei, D. M. & Lafferty, J. D. A correlated topic model of Science. *Ann. Appl. Stat.* **1**, 17–35. <https://doi.org/10.1214/07-aos114> (2007) [arXiv:0708.3601v2](https://arxiv.org/abs/0708.3601v2).
- Rozenblatt-Rosen, O., Stubbington, M. J., Regev, A. & Teichmann, S. A. The human cell atlas: From vision to reality. *Nat. News* **550**, 451 (2017).
- Turnbaugh, P. J. et al. The human microbiome project. *Nature* **449**, 804–810 (2007).
- Quinn, T. P., Erb, I., Richardson, M. F. & Crowley, T. M. Understanding sequencing data as compositions: an outlook and review. *Bioinformatics* **34**, 2870–2878. <https://doi.org/10.1093/bioinformatics/bty175> (2018).
- Gloor, G. B., Macklaim, J. M., Pawlowsky-Glahn, V. & Egozcue, J. J. Microbiome datasets are compositional: and this is not optional. *Front. Microbiol.* **8**, 2224 (2017).
- Johnson, J. S. et al. Evaluation of 16s rrna gene sequencing for species and strain-level microbiome analysis. *Nat. Commun.* **10**, 1–11 (2019).
- Rivera-Pinto, J. et al. Balances: a new perspective for microbiome analysis. *MSystems* **3**, e00053-18 (2018).
- Bates, S. & Tibshirani, R. Log-ratio lasso: scalable, sparse estimation for log-ratio models. *Biometrics* **75**, 613–624 (2019).
- Cammarota, G. et al. Gut microbiome, big data and machine learning to promote precision medicine for cancer. *Nat. Rev. Gastroenterol. Hepatol.* **17**, 635–648. <https://doi.org/10.1038/s41575-020-0327-3> (2020).
- Quinn, T. P., Nguyen, D., Rana, S., Gupta, S. & Venkatesh, S. DeepCoDA: personalized interpretability for compositional health data. *arXiv* (2020). [arXiv:2006.01392](https://arxiv.org/abs/2006.01392).
- Oh, M. & Zhang, L. Deepmicro: deep representation learning for disease prediction based on microbiome data. *Scient. Rep.* <https://doi.org/10.1038/s41598-020-63159-5> (2020).
- Buettner, M., Ostner, J., Mueller, C. L., Theis, F. J. & Schubert, B. scCODA is a bayesian model for compositional single-cell data analysis. *Nat. Commun.* **12**, 6876 (2021).
- Park, J., Yoon, C., Park, C. & Ahn, J. Kernel methods for radial transformed compositional data with many zeros. In *International Conference on Machine Learning*, 17458–17472 (PMLR, 2022).
- Huang, S., Ailer, E., Kilbertus, N. & Pfister, N. Supervised learning and model analysis with compositional data. *PLoS Comput. Biol.* **19**, e1011240 (2023).
- Taba, N., Fischer, K., research team, E. B., Org, E. & Aasmets, O. A novel framework for assessing causal effect of microbiome on health: long-term antibiotic usage as an instrument. *medRxiv* <https://doi.org/10.1101/2023.09.20.23295831> (2023). <https://www.medrxiv.org/content/early/2023/12/11/2023.09.20.23295831.full.pdf>.
- K, X. et al. Causal Effects of Gut Microbiome on Systemic Lupus Erythematosus: A Two-Sample Mendelian Randomization Study. *Frontiers in Immunology* **12**, <https://doi.org/10.3389/fimmu.2021.667097> Format: (2021).
- Arnold, K. F., Berrie, L., Tennant, P. W. & Gilthorpe, M. S. A causal inference perspective on the analysis of compositional data. *Int. J. Epidemiol.* **49**, 1307–1313 (2020).
- Breskin, A. & Murray, E. J. Commentary: Compositional data call for complex interventions. *Int. J. Epidemiol.* **49**, 1314–1315 (2020).
- Chapin, F. S. et al. Consequences of changing biodiversity. *Nature* **405**, 234–242. <https://doi.org/10.1038/35012241> (2000).
- Blaser, M. J. *Missing Microbes: How the overuse of antibiotics is fueling our modern plagues* 1st edn. (Henry Holt and Company, 2014).
- Heumos, L. et al. Best practices for single-cell analysis across modalities. *Nature Reviews Genetics* 1–23 (2023).
- Shade, A. Diversity is the question, not the answer. *ISME J.* **11**, 1–6 (2017).
- Willis, A. Rarefaction, alpha diversity, and statistics. *Front. Microbiol.* **10**, 2407. <https://doi.org/10.3389/fmicb.2019.02407> (2019).
- Kers, J. G. & Saccenti, E. The power of microbiome studies: Some considerations on which alpha and beta metrics to use and how to report results. *Front. Microbiol.* <https://doi.org/10.3389/fmicb.2021.796025> (2022).
- Vujkovic-Cvijin, I. et al. Host variables confound gut microbiota studies of human disease. *Nature* **2020**, 1–7. <https://doi.org/10.1038/s41586-020-2881-9> (2020).
- Sohn, M. B. & Li, H. Compositional mediation analysis for microbiome studies. *Ann. Appl. Statist.* **13**, 661–681. <https://doi.org/10.1214/18-AOAS1210> (2019).
- Carter, K. M., Lu, M., Jiang, H. & An, L. An information-based approach for mediation analysis on high-dimensional metagenomic data. *Front. Genet.* **11**, 148 (2020).
- Wang, C., Hu, J., Blaser, M. J., Li, H. & Birol, I. Estimating and testing the microbial causal mediation effect with high-dimensional and compositional microbiome data. *Bioinformatics* <https://doi.org/10.1093/bioinformatics/btz565> (2020).
- Xia, Y. Mediation analysis of microbiome data and detection of causality in microbiome studies. *Inflammation, Infection, and Microbiome in Cancers: Evidence, Mechanisms, and Implications* 457–509 (2021).
- Sohn, M. B., Lu, J. & Li, H. A compositional mediation model for a binary outcome: Application to microbiome studies. *Bioinformatics* **38**, 16–21 (2022).
- Wang, C. et al. A microbial causal mediation analytic tool for health disparity and applications in body mass index. *Microbiome* **11**, 164. <https://doi.org/10.1186/s40168-023-01608-9> (2023).
- Zhang, H. et al. Mediation effect selection in high-dimensional and compositional microbiome data. *Statist. Med.* <https://doi.org/10.1002/sim.8808> (2020).
- Sommer, A. J. et al. A randomization-based causal inference framework for uncovering environmental exposure effects on human gut microbiota. *PLoS Comput. Biol.* **18**, e1010044 (2022).
- Pearl, J. *Causality* (Cambridge University Press, 2009).
- Angrist, J. D. & Pischke, J.-S. *Mostly harmless econometrics: An empiricist's companion* (Princeton University Press, 2008).
- Hernán, M. A. & Robins, J. M. Instruments for causal inference: an epidemiologist's dream? *Epidemiology* 360–372 (2006).
- Imbens, G. W. & Rubin, D. B. *Causal inference in statistics, social, and biomedical sciences* (Cambridge University Press, 2015).
- Pearl, J. On the testability of causal models with latent and instrumental variables. In *Proceedings of the Eleventh conference on Uncertainty in artificial intelligence*, 435–443 (Morgan Kaufmann Publishers Inc., 1995).
- Bonnet, B. Instrumentality tests revisited. In *Proceedings of the 17th Conference on Uncertainty in Artificial Intelligence*, 48–55 (2001).

42. Gunsilius, F. Testability of instrument validity under continuous endogenous variables. *arXiv preprint* [arXiv:1806.09517](https://arxiv.org/abs/1806.09517) (2018).
43. Newey, W. K. & Powell, J. L. Instrumental variable estimation of nonparametric models. *Econometrica* **71**, 1565–1578 (2003).
44. Blundell, R., Chen, X. & Kristensen, D. Semi-nonparametric iv estimation of shape-invariant engel curves. *Econometrica* **75**, 1613–1669 (2007).
45. Singh, R., Sahani, M. & Gretton, A. Kernel instrumental variable regression. In *Advances in Neural Information Processing Systems*, 4593–4605 (2019).
46. Muandet, K., Mehrjou, A., Lee, S. K. & Raj, A. Dual instrumental variable regression. *arXiv preprint* [arXiv:1910.12358](https://arxiv.org/abs/1910.12358) (2019).
47. Zhang, R., Imaizumi, M., Schölkopf, B. & Muandet, K. Maximum moment restriction for instrumental variable regression. *arXiv preprint* [arXiv:2010.07684](https://arxiv.org/abs/2010.07684) (2020).
48. Bennett, A. et al. Minimax instrumental variable regression and l2 convergence guarantees without identification or closedness. *arXiv preprint* [arXiv:2302.05404](https://arxiv.org/abs/2302.05404) (2023).
49. Rothenhäusler, D., Meinshausen, N., Bühlmann, P. & Peters, J. Anchor regression: Heterogeneous data meet causality. *J. R. Stat. Soc. Ser. B Stat Methodol.* **83**, 215–246 (2021).
50. Pfister, N. & Peters, J. Identifiability of sparse causal effects using instrumental variables. In *Uncertainty in Artificial Intelligence*, 1613–1622 (PMLR, 2022).
51. Ailer, E., Hartford, J. & Kilbertus, N. Sequential underspecified instrument selection for cause-effect estimation. In *Proceedings of the 40th International Conference on Machine Learning*, vol. 202 of *Proceedings of Machine Learning Research*, 408–420 (PMLR, 2023).
52. Kelejian, H. H. Two-stage least squares and econometric systems linear in parameters but nonlinear in the endogenous variables. *J. Am. Stat. Assoc.* **66**, 373–374 (1971).
53. Pawłowsky-Glahn, V. & Egozcue, J. J. Geometric approach to statistical analysis on the simplex. *Stoch. Env. Res. Risk Assess.* **15**, 384–398 (2001).
54. Egozcue, J. J., Pawłowsky-Glahn, V., Mateu-Figueras, G. & Barcelo-Vidal, C. Isometric logratio transformations for compositional data analysis. *Math. Geol.* **35**, 279–300 (2003).
55. Greenacre, M. & Grunsky, E. The isometric logratio transformation in compositional data analysis: a practical evaluation. *preprint* (2019).
56. Aitchison, J. & Bacon-Shone, J. Log contrast models for experiments with mixtures. *Biometrika* **71**, 323–330. <https://doi.org/10.1093/biomet/71.2.323> (1984).
57. Lin, W., Shi, P., Feng, R. & Li, H. Variable selection in regression with compositional covariates. *Biometrika* **101**, 785–797. <https://doi.org/10.1093/biomet/asu031> (2014).
58. Shi, P., Zhang, A. & Li, H. Regression analysis for microbiome compositional data. *Ann. Appl. Statist.* **10**, 1019–1040. <https://doi.org/10.1214/16-AOAS928> (2016).
59. Combettes, P. & Müller, C. Regression models for compositional data: General log-contrast formulations, proximal optimization, and microbiome data applications. *Stat. Biosci.* <https://doi.org/10.1007/s12561-020-09283-2> (2021).
60. Combettes, P. L. & Müller, C. L. Perspective maximum likelihood-type estimation via proximal decomposition. *Electron. J. Statist.* **14**, 207–238. <https://doi.org/10.1214/19-EJS1662> (2020).
61. Kaul, A., Mandal, S., Davidov, O. & Peddada, S. D. Analysis of microbiome data in the presence of excess zeros. *Front. Microbiol.* **8**, 2114 (2017).
62. Lin, H. & Peddada, S. D. Analysis of microbial compositions: a review of normalization and differential abundance analysis. *NPJ Biofilms Microbiomes* **6**, 1–13 (2020).
63. Shi, P., Zhou, Y. & Zhang, A. R. High-dimensional log-error-in-variable regression with applications to microbial compositional data analysis. *Biometrika* **109**, 405–420 (2022).
64. Leinster, T. & Cobbold, C. Measuring diversity: The importance of species similarity. *Ecology* **93**, 477–89. <https://doi.org/10.2307/23143936> (2012).
65. Chao, A., Chiu, C.-H. & Jost, L. Unifying species diversity, phylogenetic diversity, functional diversity, and related similarity and differentiation measures through hill numbers. *Annu. Rev. Ecol. Evol. Syst.* **45**, 297–324 (2014).
66. Daly, A. J., Baetens, J. M. & De Baets, B. Ecological diversity: Measuring the unmeasurable. *Mathematics* **6**, 119 (2018).
67. Bello, M. G. D., Knight, R., Gilbert, J. A. & Blaser, M. J. Preserving microbial diversity. *Science* <https://doi.org/10.1126/science.aau8816> (2018).
68. Greene, W. H. Accounting for excess zeros and sample selection in poisson and negative binomial regression models. *NYU working paper no. EC-94-10* (1994).
69. Xu, L., Paterson, A. D., Turpin, W. & Xu, W. Assessment and selection of competing models for zero-inflated microbiome data. *PLoS One* **10**, e0129606 (2015).
70. Hartford, J., Lewis, G., Leyton-Brown, K. & Taddy, M. Deep iv: A flexible approach for counterfactual prediction. In *International Conference on Machine Learning*, 1414–1423 (2017).
71. Bennett, A., Kallus, N. & Schnabel, T. Deep generalized method of moments for instrumental variable analysis. In Wallach, H. et al. (eds.) *Advances in Neural Information Processing Systems*, vol. 32 (Curran Associates, Inc., 2019).
72. Andrews, I., Stock, J. H. & Sun, L. Weak instruments in instrumental variables regression: Theory and practice. *Ann. Rev. Econ.* **11**, 727–753 (2019).
73. Sanderson, E. & Windmeijer, F. A weak instrument f-test in linear iv models with multiple endogenous variables. *Journal of Econometrics* **190**, 212–221. <https://doi.org/10.1016/j.jeconom.2015.06.004> (2016). Endogeneity Problems in Econometrics.
74. Schulfer, A. et al. The impact of early-life sub-therapeutic antibiotic treatment (stat) on excessive weight is robust despite transfer of intestinal microbes. *ISME J.* **13**, 1. <https://doi.org/10.1038/s41396-019-0349-4> (2019).

Acknowledgements

We thank Dr. Chan Wang and Dr. Huilin Li, NYU Langone Medical Center, for kindly providing the pre-processed murine amplicon and associated phenotype data used in this study. We thank Léo Simpson, TU München, and Alice Sommer, LMU München, for kindly and patiently providing their technical and scientific support.

Author contributions

NK, CLM and EA wrote the manuscript. NK and CLM reviewed the manuscript. EA conducted the analysis.

Funding

Open Access funding enabled and organized by Projekt DEAL.

EA is supported by the Helmholtz Association under the joint research school “Munich School for Data Science - MUDS”.

Declarations

Competing interests

No competing interest is declared.

Additional information

Supplementary Information The online version contains supplementary material available at <https://doi.org/10.1038/s41598-025-89204-9>.

Correspondence and requests for materials should be addressed to E.A.

Reprints and permissions information is available at www.nature.com/reprints.

Publisher's note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

© The Author(s) 2025



Attribution 4.0 International

Deed

Canonical URL : <https://creativecommons.org/licenses/by/4.0/>

[See the legal code](#)

You are free to:

Share — copy and redistribute the material in any medium or format for any purpose, even commercially.

Adapt — remix, transform, and build upon the material for any purpose, even commercially.

The licensor cannot revoke these freedoms as long as you follow the license terms.

Under the following terms:

Attribution — You must give **appropriate credit** , provide a link to the license, and **indicate if changes were made** . You may do so in any reasonable manner, but not in any way that suggests the licensor endorses you or your use.

No additional restrictions — You may not apply legal terms or **technological measures** that legally restrict others from doing anything the license permits.

Notices:

You do not have to comply with the license for elements of the material in the public domain or where your use is permitted by an applicable **exception or limitation** .

No warranties are given. The license may not give you all of the permissions necessary for your intended use. For example, other rights such as **publicity, privacy, or moral rights** may limit how you use the material.

Notice

This deed highlights only some of the key features and terms of the actual license. It is not a license and has no legal value. You should

carefully review all of the terms and conditions of the actual license before using the licensed material.

Creative Commons is not a law firm and does not provide legal services. Distributing, displaying, or linking to this deed or the license that it summarizes does not create a lawyer-client or any other relationship.

Creative Commons is the nonprofit behind the open licenses and other legal tools that allow creators to share their work. Our legal tools are free to use.

- [Learn more about our work](#)
- **[Learn more about CC Licensing](#)**
- [Support our work](#)
- [Use the license for your own material.](#)
- [Licenses List](#)
- [Public Domain List](#)

Footnotes

appropriate credit — If supplied, you must provide the name of the creator and attribution parties, a copyright notice, a license notice, a disclaimer notice, and a link to the material. CC licenses prior to Version 4.0 also require you to provide the title of the material if supplied, and may have other slight differences.

- [More info](#)

indicate if changes were made — In 4.0, you must indicate if you modified the material and retain an indication of previous modifications. In 3.0 and earlier license versions, the indication of changes is only required if you create a derivative.

- [Marking guide](#)
- [More info](#)

technological measures — The license prohibits application of effective technological measures, defined with reference to Article 11 of the WIPO Copyright Treaty.

- [More info](#)

exception or limitation — The rights of users under exceptions and limitations, such as fair use and fair dealing, are not affected by the CC licenses.

- [More info](#)

publicity, privacy, or moral rights — You may need to get additional permissions before using the material as you intend.

- [More info](#)

Attribution 4.0 International

=====
Creative Commons Corporation ("Creative Commons") is not a law firm and does not provide legal services or legal advice. Distribution of Creative Commons public licenses does not create a lawyer-client or other relationship. Creative Commons makes its licenses and related information available on an "as-is" basis. Creative Commons gives no warranties regarding its licenses, any material licensed under their terms and conditions, or any related information. Creative Commons disclaims all liability for damages resulting from their use to the fullest extent possible.

Using Creative Commons Public Licenses

Creative Commons public licenses provide a standard set of terms and conditions that creators and other rights holders may use to share original works of authorship and other material subject to copyright and certain other rights specified in the public license below. The following considerations are for informational purposes only, are not exhaustive, and do not form part of our licenses.

Considerations for licensors: Our public licenses are intended for use by those authorized to give the public permission to use material in ways otherwise restricted by copyright and certain other rights. Our licenses are irrevocable. Licensors should read and understand the terms and conditions of the license they choose before applying it. Licensors should also secure all rights necessary before applying our licenses so that the public can reuse the material as expected. Licensors should clearly mark any material not subject to the license. This includes other CC-licensed material, or material used under an exception or limitation to copyright. More considerations for licensors: wiki.creativecommons.org/Considerations_for_licensors

Considerations for the public: By using one of our public licenses, a licensor grants the public permission to use the licensed material under specified terms and conditions. If the licensor's permission is not necessary for any reason—for example, because of any applicable exception or limitation to copyright—then that use is not regulated by the license. Our licenses grant only permissions under copyright and certain other rights that a licensor has authority to grant. Use of the licensed material may still be restricted for other reasons, including because others have copyright or other rights in the material. A licensor may make special requests, such as asking that all changes be marked or described. Although not required by our licenses, you are encouraged to respect those requests where reasonable. More considerations for the public: wiki.creativecommons.org/Considerations_for_licensees

Creative Commons Attribution 4.0 International Public License

By exercising the Licensed Rights (defined below), You accept and agree

to be bound by the terms and conditions of this Creative Commons Attribution 4.0 International Public License ("Public License"). To the extent this Public License may be interpreted as a contract, You are granted the Licensed Rights in consideration of Your acceptance of these terms and conditions, and the Licensor grants You such rights in consideration of benefits the Licensor receives from making the Licensed Material available under these terms and conditions.

Section 1 -- Definitions.

- a. Adapted Material means material subject to Copyright and Similar Rights that is derived from or based upon the Licensed Material and in which the Licensed Material is translated, altered, arranged, transformed, or otherwise modified in a manner requiring permission under the Copyright and Similar Rights held by the Licensor. For purposes of this Public License, where the Licensed Material is a musical work, performance, or sound recording, Adapted Material is always produced where the Licensed Material is synched in timed relation with a moving image.
- b. Adapter's License means the license You apply to Your Copyright and Similar Rights in Your contributions to Adapted Material in accordance with the terms and conditions of this Public License.
- c. Copyright and Similar Rights means copyright and/or similar rights closely related to copyright including, without limitation, performance, broadcast, sound recording, and Sui Generis Database Rights, without regard to how the rights are labeled or categorized. For purposes of this Public License, the rights specified in Section 2(b)(1)-(2) are not Copyright and Similar Rights.
- d. Effective Technological Measures means those measures that, in the absence of proper authority, may not be circumvented under laws fulfilling obligations under Article 11 of the WIPO Copyright Treaty adopted on December 20, 1996, and/or similar international agreements.
- e. Exceptions and Limitations means fair use, fair dealing, and/or any other exception or limitation to Copyright and Similar Rights that applies to Your use of the Licensed Material.
- f. Licensed Material means the artistic or literary work, database, or other material to which the Licensor applied this Public License.
- g. Licensed Rights means the rights granted to You subject to the terms and conditions of this Public License, which are limited to all Copyright and Similar Rights that apply to Your use of the Licensed Material and that the Licensor has authority to license.
- h. Licensor means the individual(s) or entity(ies) granting rights under this Public License.
- i. Share means to provide material to the public by any means or process that requires permission under the Licensed Rights, such as reproduction, public display, public performance, distribution, dissemination, communication, or importation, and to make material available to the public including in ways that members of the

public may access the material from a place and at a time individually chosen by them.

- j. Sui Generis Database Rights means rights other than copyright resulting from Directive 96/9/EC of the European Parliament and of the Council of 11 March 1996 on the legal protection of databases, as amended and/or succeeded, as well as other essentially equivalent rights anywhere in the world.
- k. You means the individual or entity exercising the Licensed Rights under this Public License. Your has a corresponding meaning.

Section 2 -- Scope.

a. License grant.

1. Subject to the terms and conditions of this Public License, the Licensor hereby grants You a worldwide, royalty-free, non-sublicensable, non-exclusive, irrevocable license to exercise the Licensed Rights in the Licensed Material to:
 - a. reproduce and Share the Licensed Material, in whole or in part; and
 - b. produce, reproduce, and Share Adapted Material.
2. Exceptions and Limitations. For the avoidance of doubt, where Exceptions and Limitations apply to Your use, this Public License does not apply, and You do not need to comply with its terms and conditions.
3. Term. The term of this Public License is specified in Section 6(a).
4. Media and formats; technical modifications allowed. The Licensor authorizes You to exercise the Licensed Rights in all media and formats whether now known or hereafter created, and to make technical modifications necessary to do so. The Licensor waives and/or agrees not to assert any right or authority to forbid You from making technical modifications necessary to exercise the Licensed Rights, including technical modifications necessary to circumvent Effective Technological Measures. For purposes of this Public License, simply making modifications authorized by this Section 2(a) (4) never produces Adapted Material.
5. Downstream recipients.
 - a. Offer from the Licensor -- Licensed Material. Every recipient of the Licensed Material automatically receives an offer from the Licensor to exercise the Licensed Rights under the terms and conditions of this Public License.
 - b. No downstream restrictions. You may not offer or impose any additional or different terms or conditions on, or apply any Effective Technological Measures to, the Licensed Material if doing so restricts exercise of the Licensed Rights by any recipient of the Licensed

Material.

6. No endorsement. Nothing in this Public License constitutes or may be construed as permission to assert or imply that You are, or that Your use of the Licensed Material is, connected with, or sponsored, endorsed, or granted official status by, the Licensor or others designated to receive attribution as provided in Section 3(a)(1)(A)(i).

b. Other rights.

1. Moral rights, such as the right of integrity, are not licensed under this Public License, nor are publicity, privacy, and/or other similar personality rights; however, to the extent possible, the Licensor waives and/or agrees not to assert any such rights held by the Licensor to the limited extent necessary to allow You to exercise the Licensed Rights, but not otherwise.
2. Patent and trademark rights are not licensed under this Public License.
3. To the extent possible, the Licensor waives any right to collect royalties from You for the exercise of the Licensed Rights, whether directly or through a collecting society under any voluntary or waivable statutory or compulsory licensing scheme. In all other cases the Licensor expressly reserves any right to collect such royalties.

Section 3 -- License Conditions.

Your exercise of the Licensed Rights is expressly made subject to the following conditions.

a. Attribution.

1. If You Share the Licensed Material (including in modified form), You must:
 - a. retain the following if it is supplied by the Licensor with the Licensed Material:
 - i. identification of the creator(s) of the Licensed Material and any others designated to receive attribution, in any reasonable manner requested by the Licensor (including by pseudonym if designated);
 - ii. a copyright notice;
 - iii. a notice that refers to this Public License;
 - iv. a notice that refers to the disclaimer of warranties;
 - v. a URI or hyperlink to the Licensed Material to the extent reasonably practicable;
 - b. indicate if You modified the Licensed Material and

retain an indication of any previous modifications; and

- c. indicate the Licensed Material is licensed under this Public License, and include the text of, or the URI or hyperlink to, this Public License.
2. You may satisfy the conditions in Section 3(a)(1) in any reasonable manner based on the medium, means, and context in which You Share the Licensed Material. For example, it may be reasonable to satisfy the conditions by providing a URI or hyperlink to a resource that includes the required information.
3. If requested by the Licensor, You must remove any of the information required by Section 3(a)(1)(A) to the extent reasonably practicable.
4. If You Share Adapted Material You produce, the Adapter's License You apply must not prevent recipients of the Adapted Material from complying with this Public License.

Section 4 -- Sui Generis Database Rights.

Where the Licensed Rights include Sui Generis Database Rights that apply to Your use of the Licensed Material:

- a. for the avoidance of doubt, Section 2(a)(1) grants You the right to extract, reuse, reproduce, and Share all or a substantial portion of the contents of the database;
- b. if You include all or a substantial portion of the database contents in a database in which You have Sui Generis Database Rights, then the database in which You have Sui Generis Database Rights (but not its individual contents) is Adapted Material; and
- c. You must comply with the conditions in Section 3(a) if You Share all or a substantial portion of the contents of the database.

For the avoidance of doubt, this Section 4 supplements and does not replace Your obligations under this Public License where the Licensed Rights include other Copyright and Similar Rights.

Section 5 -- Disclaimer of Warranties and Limitation of Liability.

- a. UNLESS OTHERWISE SEPARATELY UNDERTAKEN BY THE LICENSOR, TO THE EXTENT POSSIBLE, THE LICENSOR OFFERS THE LICENSED MATERIAL AS-IS AND AS-AVAILABLE, AND MAKES NO REPRESENTATIONS OR WARRANTIES OF ANY KIND CONCERNING THE LICENSED MATERIAL, WHETHER EXPRESS, IMPLIED, STATUTORY, OR OTHER. THIS INCLUDES, WITHOUT LIMITATION, WARRANTIES OF TITLE, MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE, NON-INFRINGEMENT, ABSENCE OF LATENT OR OTHER DEFECTS, ACCURACY, OR THE PRESENCE OR ABSENCE OF ERRORS, WHETHER OR NOT KNOWN OR DISCOVERABLE. WHERE DISCLAIMERS OF WARRANTIES ARE NOT ALLOWED IN FULL OR IN PART, THIS DISCLAIMER MAY NOT APPLY TO YOU.
- b. TO THE EXTENT POSSIBLE, IN NO EVENT WILL THE LICENSOR BE LIABLE TO YOU ON ANY LEGAL THEORY (INCLUDING, WITHOUT LIMITATION, NEGLIGENCE) OR OTHERWISE FOR ANY DIRECT, SPECIAL, INDIRECT,

INCIDENTAL, CONSEQUENTIAL, PUNITIVE, EXEMPLARY, OR OTHER LOSSES, COSTS, EXPENSES, OR DAMAGES ARISING OUT OF THIS PUBLIC LICENSE OR USE OF THE LICENSED MATERIAL, EVEN IF THE LICENSOR HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH LOSSES, COSTS, EXPENSES, OR DAMAGES. WHERE A LIMITATION OF LIABILITY IS NOT ALLOWED IN FULL OR IN PART, THIS LIMITATION MAY NOT APPLY TO YOU.

- c. The disclaimer of warranties and limitation of liability provided above shall be interpreted in a manner that, to the extent possible, most closely approximates an absolute disclaimer and waiver of all liability.

Section 6 -- Term and Termination.

- a. This Public License applies for the term of the Copyright and Similar Rights licensed here. However, if You fail to comply with this Public License, then Your rights under this Public License terminate automatically.
- b. Where Your right to use the Licensed Material has terminated under Section 6(a), it reinstates:
 - 1. automatically as of the date the violation is cured, provided it is cured within 30 days of Your discovery of the violation; or
 - 2. upon express reinstatement by the Licensor.

For the avoidance of doubt, this Section 6(b) does not affect any right the Licensor may have to seek remedies for Your violations of this Public License.

- c. For the avoidance of doubt, the Licensor may also offer the Licensed Material under separate terms or conditions or stop distributing the Licensed Material at any time; however, doing so will not terminate this Public License.
- d. Sections 1, 5, 6, 7, and 8 survive termination of this Public License.

Section 7 -- Other Terms and Conditions.

- a. The Licensor shall not be bound by any additional or different terms or conditions communicated by You unless expressly agreed.
- b. Any arrangements, understandings, or agreements regarding the Licensed Material not stated herein are separate from and independent of the terms and conditions of this Public License.

Section 8 -- Interpretation.

- a. For the avoidance of doubt, this Public License does not, and shall not be interpreted to, reduce, limit, restrict, or impose conditions on any use of the Licensed Material that could lawfully be made without permission under this Public License.
- b. To the extent possible, if any provision of this Public License is

deemed unenforceable, it shall be automatically reformed to the minimum extent necessary to make it enforceable. If the provision cannot be reformed, it shall be severed from this Public License without affecting the enforceability of the remaining terms and conditions.

- c. No term or condition of this Public License will be waived and no failure to comply consented to unless expressly agreed to by the Licensor.
- d. Nothing in this Public License constitutes or may be interpreted as a limitation upon, or waiver of, any privileges and immunities that apply to the Licensor or You, including from the legal processes of any jurisdiction or authority.

=====
Creative Commons is not a party to its public licenses. Notwithstanding, Creative Commons may elect to apply one of its public licenses to material it publishes and in those instances will be considered the "Licensor." The text of the Creative Commons public licenses is dedicated to the public domain under the CC0 Public Domain Dedication. Except for the limited purpose of indicating that material is shared under a Creative Commons public license or as otherwise permitted by the Creative Commons policies published at creativecommons.org/policies, Creative Commons does not authorize the use of the trademark "Creative Commons" or any other trademark or logo of Creative Commons without its prior written consent including, without limitation, in connection with any unauthorized modifications to any of its public licenses or any other arrangements, understandings, or agreements concerning use of licensed material. For the avoidance of doubt, this paragraph does not form part of the public licenses.

Creative Commons may be contacted at creativecommons.org.

List of Figures

1.1	Confounded estimation of the effect of gut bacteria composition and depression.	3
1.2	Additional instrument to enable unconfounded cause-effect estimation of gut bacteria composition on depression.	3
2.1	Cause-effect estimation of $X \rightarrow Y$ via an instrumental variable Z for a general treatment variable X . Note that the terms <i>cause</i> and <i>treatment</i> for X are used interchangeably.	6
2.2	Assumptions of IV methods: (A1)-(A3) explicitly illustrate the IV assumptions which are necessary for a valid IV estimation.	7
2.3	Cause-Effect Estimation with IVs in a nonlinear setting: Z and X are mapped to nonlinear function spaces $\mathcal{H}_Z$ and $\mathcal{H}_X$. The operators $\mathcal{T}^*$ and H between those function spaces are linear.	13
2.4	Two scenarios for weak instrument bias in a linear IV setting with $d_z = d_x = 1$: The left hand side shows a scenario with an F-statistic of 3.4 for the first stage regression and an estimated effect that is biased towards the single confounded regression. The right hand side shows a scenario where the estimation bias with instrumental variables is even worse than the single regression (F-statistic 0.4).	21
2.5	Different notions of <i>experiments</i> that enable cause-effect estimation. The shaded orange indicate the variables that are exogenously changed resp. determined by the user.	23
A.1	ICML Licencse Copyright Info Statement	57

List of Tables

2.1	Basis function approximation terms for $k \in [K]$: Note that certain basis functions form an orthogonal basis of the Hilbert space with inner product given by $\langle f, g \rangle = \int f(x) \overline{g(x)} v(x) dx$	12
-----	---	----

Bibliography

- Adikusuma, F., Piltz, S., Corbett, M. A., Turvey, M., McColl, S. R., Helbig, K. J., Beard, M. R., Hughes, J., Pomerantz, R. T., and Thomas, P. Q. Large deletions induced by cas9 cleavage. *Nature*, 560(7717):E8–E9, 2018.
- Aglietti, V., Lu, X., Paleyes, A., and González, J. Causal bayesian optimization. In *International Conference on Artificial Intelligence and Statistics*, pp. 3155–3164. PMLR, 2020.
- Aglietti, V., Dhir, N., González, J., and Damoulas, T. Dynamic causal bayesian optimization. *Advances in Neural Information Processing Systems*, 34:10549–10560, 2021.
- Agrawal, R., Squires, C., and Yang, K. Abcd-strategy: Budgeted experimental design for targeted causal structure discovery. In *Proceedings of the 22nd International Conference on Artificial Intelligence and Statistics*, pp. 2874–2882. PMLR, 2019.
- Ailer, E., Müller, C. L., and Kilbertus, N. Instrumental variable estimation for compositional treatments. *Scientific Reports*, 15(1):5158, 2025.
- Aitchison, J. The statistical analysis of compositional data. *Journal of the Royal Statistical Society: Series B (Methodological)*, 44(2):139–160, 1982.
- Aizer, A. and Doyle, Joseph J., J. Juvenile Incarceration, Human Capital, and Future Crime: Evidence from Randomly Assigned Judges *. *The Quarterly Journal of Economics*, 130(2): 759–803, 02 2015. ISSN 0033-5533. doi: 10.1093/qje/qjv003. URL <https://doi.org/10.1093/qje/qjv003>.
- Altonji, J. G., Elder, T. E., and Taber, C. R. An evaluation of instrumental variable strategies for estimating the effects of catholic schooling. *Journal of Human resources*, 40(4):791–821, 2005.
- Andrews, I. and Armstrong, T. B. Unbiased instrumental variables estimation under known first-stage sign. *Quantitative Economics*, 8(2):479–503, 2017.
- Andrews, I., Stock, J. H., and Sun, L. Weak instruments in instrumental variables regression: Theory and practice. *Annual Review of Economics*, 11:727–753, 2019.
- Angrist, J. and Krueger, A. Does compulsory school attendance affect schooling and earnings? *The Quarterly Journal of Economics*, 106:979–1014, 02 1991. doi: 10.2307/2937954.
- Angrist, J. D. and Pischke, J.-S. *Mostly harmless econometrics: An empiricist’s companion*. Princeton university press, 2008.
- Arnold, K. F., Berrie, L., Tennant, P. W., and Gilthorpe, M. S. A causal inference perspective on the analysis of compositional data. *International journal of epidemiology*, 49(4):1307–1313, 2020.

- Asnicar, F., Berry, S. E., Valdes, A. M., Nguyen, L. H., Piccinno, G., Drew, D. A., Leeming, E., Gibson, R., Le Roy, C., Khatib, H. A., et al. Microbiome connections with host metabolism and habitual diet from 1,098 deeply phenotyped individuals. *Nature medicine*, 27(2): 321–332, 2021.
- Barrett, J. C., Clayton, D. G., Concannon, P., Akolkar, B., Cooper, J. D., Erlich, H. A., Julier, C., Morahan, G., Nerup, J., Nierras, C., et al. Genome-wide association study and meta-analysis find that over 40 loci affect risk of type 1 diabetes. *Nature genetics*, 41(6):703–707, 2009.
- Basman, R. L. A generalized classical method of linear estimation of coefficients in a structural equation. *Econometrica*, 25(1):77–83, 1957.
- Belyaeva, A., Squires, C., and Uhler, C. Dci: learning causal differences between gene regulatory networks. *Bioinformatics*, 37(18):3067–3069, 2021.
- Bennett, A., Kallus, N., and Schnabel, T. Deep generalized method of moments for instrumental variable analysis. In Wallach, H., Larochelle, H., Beygelzimer, A., d'Alché-Buc, F., Fox, E., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. URL <https://proceedings.neurips.cc/paper/2019/file/15d185eaa7c954e77f5343d941e25fbd-Paper.pdf>.
- Bennett, A., Kallus, N., Mao, X., Newey, W., Syrgkanis, V., and Uehara, M. Inference on Strongly Identified Functionals of Weakly Identified Functions. Papers 2208.08291, arXiv.org, August 2022. URL <https://ideas.repec.org/p/arx/papers/2208.08291.html>.
- Bennett, A., Kallus, N., Mao, X., Newey, W., Syrgkanis, V., and Uehara, M. Minimax instrumental variable regression and l_2 convergence guarantees without identification or closedness. In *The Thirty Sixth Annual Conference on Learning Theory*, pp. 2291–2318. PMLR, 2023.
- Bibaut, A., Chou, W., Ejdemyr, S., and Kallus, N. Learning the covariance of treatment effects across many weak experiments. In *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, pp. 153–162, 2024.
- Bishop, C. M. *Pattern Recognition and Machine Learning (Information Science and Statistics)*. Springer-Verlag, Berlin, Heidelberg, 2006. ISBN 0387310738.
- Blaser, M. J. *Missing Microbes: How the Overuse of Antibiotics Is Fueling Our Modern Plagues*. Henry Holt and Company, New York, first edit edition, 2014.
- Blundell, R., Chen, X., and Kristensen, D. Semi-nonparametric iv estimation of shape-invariant engel curves. *Econometrica*, 75(6):1613–1669, 2007.
- Bonnet, B. Instrumentality tests revisited. In *Proceedings of the 17th Conference on Uncertainty in Artificial Intelligence*, pp. 48–55, 2001.
- Bound, J., Jaeger, D. A., and Baker, R. M. Problems with instrumental variables estimation when the correlation between the instruments and the endogenous explanatory variable is weak. *Journal of the American statistical association*, 90(430):443–450, 1995.
- Breskin, A. and Murray, E. J. Commentary: Compositional data call for complex interventions. *International Journal of Epidemiology*, 49(4):1314–1315, 2020.
- Brookes, A. J. and Robinson, P. N. Human genotype–phenotype databases: aims, challenges and opportunities. *Nature Reviews Genetics*, 16(12):702–715, 2015.

- Buse, A. The bias of instrumental variable estimators. *Econometrica: Journal of the Econometric Society*, pp. 173–180, 1992.
- Bühlmann, P. Toward causality and improving external validity. *Proceedings of the National Academy of Sciences*, 117:25963–25965, 10 2020. doi: 10.1073/pnas.2018002117.
- Callaway, E. Genome studies attract criticism. *Nature*, 546(7659):463–463, 2017.
- Card, D. Using geographic variation in college proximity to estimate the return to schooling, 1993.
- Carter, K. M., Lu, M., Jiang, H., and An, L. An information-based approach for mediation analysis on high-dimensional metagenomic data. *Frontiers in Genetics*, 11:148, 2020.
- Chandak, Y., Shankar, S., Syrgkanis, V., and Brunskill, E. Adaptive instrument design for indirect experiments. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=4Zz5UELkIt>.
- Chao, J. and Swanson, N. R. Alternative approximations of the bias and mse of the iv estimator under weak identification with an application to bias correction. *Journal of Econometrics*, 137(2):515–555, 2007.
- Chapin, F. S., Zavaleta, E. S., Eviner, V. T., Naylor, R. L., Vitousek, P. M., Reynolds, H. L., Hooper, D. U., Lavorel, S., Sala, O. E., Hobbie, S. E., Mack, M. C., and Díaz, S. Consequences of changing biodiversity. *Nature*, 405(6783):234–242, 2000. ISSN 00280836. doi: 10.1038/35012241.
- Chen, L., Davey Smith, G., Harbord, R. M., and Lewis, S. J. Alcohol intake and blood pressure: a systematic review implementing a mendelian randomization approach. *PLoS medicine*, 5(3):e52, 2008.
- Cho, Y., Shin, S.-Y., Won, S., Relton, C. L., Davey Smith, G., and Shin, M.-J. Alcohol intake and cardiovascular risk factors: a mendelian randomisation study. *Scientific Reports*, 5(1): 18422, 2015.
- Cinelli, C. and Hazlett, C. An omitted variable bias framework for sensitivity analysis of instrumental variables. *Available at SSRN 4217915*, 2022.
- Cinelli, C., Kumor, D., Chen, B., Pearl, J., and Bareinboim, E. Sensitivity analysis of linear structural causal models. In *International conference on machine learning*, pp. 1252–1261. PMLR, 2019.
- Conley, T. G., Hansen, C. B., and Rossi, P. E. Plausibly exogenous. *Review of Economics and Statistics*, 94(1):260–272, 2012.
- Creager, E., Jacobsen, J.-H., and Zemel, R. Environment inference for invariant learning. In Meila, M. and Zhang, T. (eds.), *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pp. 2189–2200. PMLR, 18–24 Jul 2021. URL <https://proceedings.mlr.press/v139/creager21a.html>.
- Darolles, S., Fan, Y., Florens, J.-P., and Renault, E. Nonparametric instrumental regression. *Econometrica*, 79(5):1541–1565, 2011.
- Davies, N. M., Smith, G. D., Windmeijer, F., and Martin, R. M. Issues in the reporting and conduct of instrumental variable studies: a systematic review. *Epidemiology*, 24(3):363–369, 2013.

- Deng, Z., Jiang, J., Long, G., and Zhang, C. Causal reinforcement learning: A survey. *arXiv preprint arXiv:2307.01452*, 2023.
- D'Haultfoeuille, X. On the completeness condition in nonparametric instrumental problems. *Econometric Theory*, 27:460–471, 01 2006. doi: 10.1017/S0266466610000368.
- Dikkala, N., Lewis, G., Mackey, L., and Syrgkanis, V. Minimax estimation of conditional moment models. *Advances in Neural Information Processing Systems*, 33:12248–12262, 2020.
- DiPrete, T. A. and Gangl, M. Assessing bias in the estimation of causal effects: Rosenbaum bounds on matching estimators and instrumental variables estimation with imperfect instruments. *Sociological methodology*, 34(1):271–310, 2004.
- Dong, M., Wang, B., Wei, J., de O. Fonseca, A. H., Perry, C. J., Frey, A., Ouerghi, F., Foxman, E. F., Ishizuka, J. J., Dhodapkar, R. M., and van Dijk, D. Causal identification of single-cell experimental perturbation effects with cinema-ot. *Nature Methods*, 20(11):1769–1779, 2023. doi: 10.1038/s41592-023-02040-5. URL <https://doi.org/10.1038/s41592-023-02040-5>.
- Eaton, D. and Murphy, K. Exact bayesian structure learning from uncertain interventions. In *Artificial intelligence and statistics*, pp. 107–114. PMLR, 2007.
- Egozcue, J. J., Pawlowsky-Glahn, V., Mateu-Figueras, G., and Barcelo-Vidal, C. Isometric logratio transformations for compositional data analysis. *Mathematical Geology*, 35(3): 279–300, 2003.
- Eiseman, B., Silen, W., GS, B., and AJ, K. Fecal enema as an adjunct in the treatment of pseudomembranous enterocolitis. *Surgery*, 44(5):854–859, 1958.
- Elahi, M. Q., Wei, L., Kocaoglu, M., and Ghasemi, M. Adaptive online experimental design for causal discovery, 2024.
- Fritz, J. V., Desai, M. S., Shah, P., Schneider, J. G., and Wilmes, P. From meta-omics to causality: experimental models for human microbiome research. *Microbiome*, 1:1–15, 2013.
- Fu, Y., Foden, J., Khayter, C., Maeder, M., Reyon, D., Joung, J., and Sander, J. High-frequency off-target mutagenesis induced by crispr-cas nucleases in human cells. *Nature biotechnology*, 31, 06 2013. doi: 10.1038/nbt.2623.
- Fu, Z., Qi, Z., Wang, Z., Yang, Z., Xu, Y., and Kosorok, M. R. Offline reinforcement learning with instrumental variables in confounded markov decision processes. *arXiv preprint arXiv:2209.08666*, 2022.
- Gamella, J. and Heinze-Deml, C. Active invariant causal prediction: Experiment selection through stability. In *Advances in Neural Information Processing Systems*, 2020.
- Gershman, S. J. Reinforcement learning and causal models. *The Oxford handbook of causal reasoning*, 1(295):4–2, 2017.
- Greenacre, M. and Grunsky, E. The isometric logratio transformation in compositional data analysis: a practical evaluation. *preprint*, 2019. URL <https://ideas.repec.org/p/upf/upfgen/1627.html>.
- Grosz, M. P., Rohrer, J. M., and Thoemmes, F. The taboo against explicit causal inference in nonexperimental psychology. *Perspectives on Psychological Science*, 15(5):1243–1255, 2020.
- Gultchin, L., Aglietti, V., Bellot, A., and Chiappa, S. Functional causal bayesian optimization. In *Uncertainty in Artificial Intelligence*, pp. 756–765. PMLR, 2023.

- Gunsilius, F. F. Nontestability of instrument validity under continuous treatments. *Biometrika*, 108(4):989–995, 2021.
- Guo, R., Cheng, L., Li, J., Hahn, P. R., and Liu, H. A survey of learning causality with data: Problems and methods. *ACM Computing Surveys (CSUR)*, 53(4):1–37, 2020.
- Guo, Z., Zheng, M., and Bühlmann, P. Robustness against weak or invalid instruments: Exploring nonlinear treatment models with machine learning. *arXiv e-prints*, pp. arXiv–2203, 2022.
- Gupta, S., Lipton, Z. C., and Childers, D. Efficient Online Estimation of Causal Effects by Deciding What to Observe. Papers 2108.09265, arXiv.org, August 2021. URL <https://ideas.repec.org/p/arx/papers/2108.09265.html>.
- Haavelmo, T. The statistical implications of a system of simultaneous equations. *Econometrica, Journal of the Econometric Society*, pp. 1–12, 1943.
- Hahn, J. and Hausman, J. Estimation with valid and invalid instruments. *Annales d’Economie et de Statistique*, pp. 25–57, 2005.
- Hahn, J., Hirano, K., and Karlan, D. Adaptive experimental design using the propensity score. *Journal of Business & Economic Statistics*, 29(1):96–108, 2011.
- Hall, A. R., Rudebusch, G. D., and Wilcox, D. W. Judging instrument relevance in instrumental variables estimation. *International Economic Review*, pp. 283–298, 1996.
- Hansen, L. P. Large sample properties of generalized method of moments estimators. *Econometrica*, 50(4):1029–1054, 1982. ISSN 00129682, 14680262. URL <http://www.jstor.org/stable/1912775>.
- Harding, M., Hausman, J., and Palmer, C. J. Finite sample bias corrected iv estimation for weak and many instruments. In *Essays in honor of Aman Ullah*, pp. 245–273. Emerald Group Publishing Limited, 2016.
- Hartford, J., Lewis, G., Leyton-Brown, K., and Taddy, M. Deep iv: A flexible approach for counterfactual prediction. In *International Conference on Machine Learning*, pp. 1414–1423, 2017.
- Hausman, J. Mismeasured variables in econometric analysis: problems from the right and problems from the left. *Journal of Economic perspectives*, 15(4):57–67, 2001.
- He, Y.-B. and Geng, Z. Active learning of causal networks with intervention experiments and optimal designs. *Journal of Machine Learning Research*, 9(Nov):2523–2547, 2008.
- Heckman, J. J. Randomization as an instrumental variable, 1995.
- Heinze-Deml, C., Peters, J., and Meinshausen, N. Invariant causal prediction for nonlinear models. *Journal of Causal Inference*, 6(2):20170016, 2018. doi: doi:10.1515/jci-2017-0016. URL <https://doi.org/10.1515/jci-2017-0016>.
- Hernán, M. A. The c-word: scientific euphemisms do not improve causal inference from observational data. *American journal of public health*, 108(5):616–619, 2018.
- Horowitz, J. L. Applied nonparametric instrumental variables estimation. *Econometrica*, 79(2):347–394, 2011.
- Hsu, P. D., Scott, D. A., Weinstein, J. A., Ran, F. A., Konermann, S., Agarwala, V., Li, Y., Fine, E. J., Wu, X., Shalem, O., et al. Dna targeting specificity of rna-guided cas9 nucleases. *Nature biotechnology*, 31(9):827–832, 2013.

- Hu, P., Jiao, R., Jin, L., and Xiong, M. Application of causal inference to genomic analysis: advances in methodology. *Frontiers in Genetics*, 9:238, 2018.
- Huang, A., Chandra, M., and Malkhasyan, L. Weak instrumental variables: Limitations of traditional 2sls and exploring alternative instrumental variable estimators. *arXiv preprint arXiv:2104.12370*, 2021.
- Hüyük, A., Qian, Z., and van der Schaar, M. Adaptive experiment design with synthetic controls. In *International Conference on Artificial Intelligence and Statistics*, pp. 1180–1188. PMLR, 2024.
- Imbens, G. W. and Rubin, D. B. *Causal inference in statistics, social, and biomedical sciences*. Cambridge University Press, 2015.
- Janzing, D. Causal regularization. *Advances in Neural Information Processing Systems*, 32, 2019.
- Janzing, D. and Schölkopf, B. Detecting non-causal artifacts in multivariate linear regression models. In *International Conference on Machine Learning*, pp. 2245–2253. PMLR, 2018.
- Johnson, A. J., Vangay, P., Al-Ghalith, G. A., Hillmann, B. M., Ward, T. L., Shields-Cutler, R. R., Kim, A. D., Shmagel, A. K., Syed, A. N., Walter, J., et al. Daily sampling reveals personalized diet-microbiome associations in humans. *Cell host & microbe*, 25(6):789–802, 2019a.
- Johnson, J. S., Spakowicz, D. J., Hong, B.-Y., Petersen, L. M., Demkowicz, P., Chen, L., Leopold, S. R., Hanson, B. M., Agresta, H. O., Gerstein, M., et al. Evaluation of 16s rna gene sequencing for species and strain-level microbiome analysis. *Nature communications*, 10(1):1–11, 2019b.
- Kaddour, J., Lynch, A., Liu, Q., Kusner, M. J., and Silva, R. Causal machine learning: A survey and open problems. *arXiv preprint arXiv:2206.15475*, 2022.
- Kang, H., Zhang, A., Cai, T. T., and Small, D. S. Instrumental variables estimation with some invalid instruments and its application to mendelian randomization. *Journal of the American statistical Association*, 111(513):132–144, 2016.
- Kang, H., Jiang, Y., Zhao, Q., and Small, D. S. Ivmodel: an r package for inference and sensitivity analysis of instrumental variables models with one endogenous variable. *Observational Studies*, 7(2):1–24, 2021.
- Ke, N. R., Dunn, S.-J., Bornschein, J., Chiappa, S., Rey, M., Lespiau, J.-B., Cassirer, A., Wang, J., Weber, T., Barrett, D., et al. Discogen: Learning to discover gene regulatory networks. *arXiv preprint arXiv:2304.05823*, 2023.
- Kédagni, D. and Mourifié, I. Generalized instrumental inequalities: testing the instrumental variable independence assumption. *Biometrika*, 107(3):661–675, 2020.
- Kelejian, H. H. Two-stage least squares and econometric systems linear in parameters but nonlinear in the endogenous variables. *Journal of the American Statistical Association*, 66(334):373–374, 1971.
- Kers, J. G. and Saccenti, E. The power of microbiome studies: Some considerations on which alpha and beta metrics to use and how to report results. *Frontiers in Microbiology*, 12, 2022. ISSN 1664-302X. doi: 10.3389/fmicb.2021.796025. URL <https://www.frontiersin.org/articles/10.3389/fmicb.2021.796025>.
- Kilbertus, N., Kusner, M. J., and Silva, R. A class of algorithms for general instrumental variable models. In *Advances in Neural Information Processing Systems*, volume 33, 2020.

- Kolde, R., Franzosa, E. A., Rahnavard, G., Hall, A. B., Vlamakis, H., Stevens, C., Daly, M. J., Xavier, R. J., and Huttenhower, C. Host genetic variation and its microbiome interactions within the human microbiome project. *Genome medicine*, 10:1–13, 2018.
- Kress, R. Linear integral equations, 1989.
- Kymn, K. O. An asymptotically unbiased two-stage least squares estimator of the structural disturbance variances and the bias. *Metroeconomica*, 26(1-3):245–255, 1974.
- Lazar, N. H., Celik, S., Chen, L., Fay, M. M., Irish, J. C., Jensen, J., Tillinghast, C. A., Urbanik, J., Bone, W. P., Gibson, C. C., et al. High-resolution genome-wide mapping of chromosome-arm-scale truncations induced by crispr-cas9 editing. *Nature Genetics*, pp. 1–12, 2024.
- Lewis, G. and Syrgkanis, V. Adversarial generalized method of moments, March 2018. URL <https://www.microsoft.com/en-us/research/publication/adversarial-generalized-method-moments/>. arXiv.
- Ley, R. E., Bäckhed, F., Turnbaugh, P., Lozupone, C. A., Knight, R. D., and Gordon, J. I. Obesity alters gut microbial ecology. *Proceedings of the national academy of sciences*, 102(31):11070–11075, 2005.
- Liao, L., Chen, Y.-L., Yang, Z., Dai, B., Kolar, M., and Wang, Z. Provably efficient neural estimation of structural equation models: An adversarial approach. *Advances in Neural Information Processing Systems*, 33:8947–8958, 2020.
- Liao, L., Fu, Z., Yang, Z., Wang, Y., Kolar, M., and Wang, Z. Instrumental variable value iteration for causal offline reinforcement learning. *arXiv preprint arXiv:2102.09907*, 2021.
- Liu, S., Li, E., Sun, Z., Fu, D., Duan, G., Jiang, M., Yu, Y., Mei, L., Yang, P., Tang, Y., et al. Altered gut microbiota and short chain fatty acids in chinese children with autism spectrum disorder. *Scientific reports*, 9(1):287, 2019.
- Liu, S., Li, F., Cai, Y., Ren, L., Sun, L., Gang, X., and Wang, G. Unraveling the mystery: a mendelian randomized exploration of gut microbiota and different types of obesity. *Frontiers in Cellular and Infection Microbiology*, 14:1352109, 2024.
- Liu, Y., Han, T., Ma, S., Zhang, J., Yang, Y., Tian, J., He, H., Li, A., He, M., Liu, Z., et al. Summary of chatgpt-related research and perspective towards the future of large language models. *Meta-Radiology*, pp. 100017, 2023.
- Lousdal, M. L. An introduction to instrumental variable assumptions, validation and estimation. *Emerging themes in epidemiology*, 15(1):1, 2018.
- Markowitz, F., Grossmann, S., and Spang, R. Probabilistic soft interventions in conditional gaussian networks. In *International Workshop on Artificial Intelligence and Statistics*, pp. 214–221. PMLR, 2005.
- Meinshausen, N., Hauser, A., Mooij, J., Peters, J., Versteeg, P., and Bühlmann, P. Methods for causal inference from gene perturbation experiments and validation. *Proceedings of the National Academy of Sciences*, 113:7361–7368, 07 2016. doi: 10.1073/pnas.1510493113.
- Millwood, I. Y., Walters, R. G., Mei, X. W., Guo, Y., Yang, L., Bian, Z., Bennett, D. A., Chen, Y., Dong, C., Hu, R., et al. Conventional and genetic evidence on alcohol and vascular disease aetiology: a prospective study of 500 000 men and women in china. *The Lancet*, 393(10183):1831–1842, 2019.
- Mooij, J. and Heskes, T. Cyclic causal discovery from continuous equilibrium data. *arXiv preprint arXiv:1309.6849*, 2013.

- Moreira, M. J. A conditional likelihood ratio test for structural models. *Econometrica*, 71(4): 1027–1048, 2003.
- Morgan, X. C., Tickle, T. L., Sokol, H., Gevers, D., Devaney, K. L., Ward, D. V., Reyes, J. A., Shah, S. A., LeLeiko, N., Snapper, S. B., et al. Dysfunction of the intestinal microbiome in inflammatory bowel disease and treatment. *Genome biology*, 13:1–18, 2012.
- Morgan, X. C., Kabakchiev, B., Waldron, L., Tyler, A. D., Tickle, T. L., Milgrom, R., Stempak, J. M., Gevers, D., Xavier, R. J., Silverberg, M. S., et al. Associations between host gene expression, the mucosal microbiome, and clinical outcome in the pelvic pouch of patients with inflammatory bowel disease. *Genome biology*, 16:1–15, 2015.
- Muandet, K., Mehrjou, A., Lee, S. K., and Raj, A. Dual instrumental variable regression. *arXiv preprint arXiv:1910.12358*, 2019.
- Nagar, A. L. The bias and moment matrix of the general k-class estimators of the parameters in simultaneous equations. *Econometrica: Journal of the Econometric Society*, pp. 575–595, 1959.
- Newey, W. K. and Powell, J. L. Instrumental variable estimation of nonparametric models. *Econometrica*, 71(5):1565–1578, 2003.
- Opreescu, M. and Kallus, N. Estimating heterogeneous treatment effects by combining weak instruments and observational data. *arXiv preprint arXiv:2406.06452*, 2024.
- Padh, K., Zeitler, J., Watson, D., Kusner, M., Silva, R., and Kilbertus, N. Stochastic causal programming for bounding treatment effects, 2022. URL <https://arxiv.org/abs/2202.10806>.
- Parikh, H., Varjao, C., Xu, L., and Tchetgen, E. T. Validating causal inference methods. In *International conference on machine learning*, pp. 17346–17358. PMLR, 2022.
- Pasolli, E., Truong, D. T., Malik, F., Waldron, L., and Segata, N. Machine learning meta-analysis of large metagenomic datasets: tools and biological insights. *PLoS computational biology*, 12(7):e1004977, 2016.
- Pearl, J. On the testability of causal models with latent and instrumental variables. In *Proceedings of the Eleventh conference on Uncertainty in artificial intelligence*, pp. 435–443. Morgan Kaufmann Publishers Inc., 1995.
- Pearl, J. *Causality*. Cambridge university press, 2009.
- Peters, J., Bühlmann, P., and Meinshausen, N. Causal Inference by using Invariant Prediction: Identification and Confidence Intervals. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 78(5):947–1012, 10 2016. ISSN 1369-7412. doi: 10.1111/rssb.12167. URL <https://doi.org/10.1111/rssb.12167>.
- Peyre, H., Schoeler, T., Liu, C., Williams, C. M., Hoertel, N., Havdahl, A., and Pingault, J.-B. Combining multivariate genomic approaches to elucidate the comorbidity between autism spectrum disorder and attention deficit hyperactivity disorder. *Journal of Child Psychology and Psychiatry*, 62(11):1285–1296, 2021.
- Pfister, N. and Peters, J. Identifiability of sparse causal effects using instrumental variables. *arXiv preprint arXiv:2203.09380*, 2022.
- Pingault, J.-B., Richmond, R., and Smith, G. D. Causal inference with genetic data: Past, present, and future. *Cold Spring Harbor perspectives in medicine*, 12(3):a041271, 2022.

- Pu, H. and Zhang, B. Estimating optimal treatment rules with an instrumental variable: A partial identification learning approach. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 83(2):318–345, 2021.
- Qin, J., Li, Y., Cai, Z., Li, S., Zhu, J., Zhang, F., Liang, S., Zhang, W., Guan, Y., Shen, D., et al. A metagenome-wide association study of gut microbiota in type 2 diabetes. *nature*, 490(7418):55–60, 2012.
- Reiersøl, O. Identifiability of a linear relation between variables which are subject to error. *Econometrica: Journal of the Econometric Society*, pp. 375–389, 1950.
- Richardson, T. G., Zheng, J., and Gaunt, T. R. Computational tools for causal inference in genetics. *Cold Spring Harbor Perspectives in Medicine*, 11(6):a039248, 2021.
- Rothenhäusler, D., Bühlmann, P., Meinshausen, N., and Peters, J. Anchor regression: Heterogeneous data meet causality. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 83, 01 2018. doi: 10.1111/rssb.12398.
- Ruan, K., Zhang, J., Di, X., and Bareinboim, E. Causal imitation learning via inverse reinforcement learning. In *The Eleventh International Conference on Learning Representations*, 2023.
- Saengkyongam, S., Henckel, L., Pfister, N., and Peters, J. Exploiting independent instruments: Identification and distribution generalization. *arXiv preprint arXiv:2202.01864*, 2022.
- Sanderson, E. and Windmeijer, F. A weak instrument f-test in linear iv models with multiple endogenous variables. *Journal of Econometrics*, 190(2):212–221, 2016. ISSN 0304-4076. doi: <https://doi.org/10.1016/j.jeconom.2015.06.004>. URL <https://www.sciencedirect.com/science/article/pii/S0304407615001736>. Endogeneity Problems in Econometrics.
- Sanderson, E., Davey Smith, G., Windmeijer, F., and Bowden, J. An examination of multi-variable mendelian randomization in the single-sample and two-sample summary data settings. *International journal of epidemiology*, 48(3):713–727, 2019.
- Schulfer, A., Schluter, J., Zhang, Y., Brown, Q., Pathmasiri, W., McRitchie, S., Sumner, S., Li, H., Xavier, J., and Blaser, M. The impact of early-life sub-therapeutic antibiotic treatment (stat) on excessive weight is robust despite transfer of intestinal microbes. *The ISME Journal*, 13:1, 01 2019. doi: 10.1038/s41396-019-0349-4.
- Severini, T. A. and Tripathi, G. Some identification issues in nonparametric linear models with endogenous regressors. *Econometric Theory*, 22(2):258–278, 2006. doi: 10.1017/S0266466606060117.
- Singh, R., Sahani, M., and Gretton, A. Kernel instrumental variable regression. In *Advances in Neural Information Processing Systems*, pp. 4593–4605, 2019.
- Singh, S. Nonparametric indirect active learning. In *International Conference on Artificial Intelligence and Statistics*, pp. 2515–2541. PMLR, 2023.
- Small, D. S. Sensitivity analysis for instrumental variables regression with overidentifying restrictions. *Journal of the American Statistical Association*, 102(479):1049–1058, 2007.
- Small, D. S. and Rosenbaum, P. R. War and wages: the strength of instrumental variables and their sensitivity to unobserved biases. *Journal of the American Statistical Association*, 103(483):924–933, 2008.
- Sohn, M. B. and Li, H. Compositional mediation analysis for microbiome studies. *Annals of Applied Statistics*, 13(1):661–681, 2019. ISSN 19417330. doi: 10.1214/18-AOAS1210.

- Sohn, M. B., Lu, J., and Li, H. A compositional mediation model for a binary outcome: Application to microbiome studies. *Bioinformatics*, 38(1):16–21, 2022.
- Sommer, A. J., Peters, A., Rommel, M., Cyrys, J., Grallert, H., Haller, D., Müller, C. L., and Bind, M.-A. C. A randomization-based causal inference framework for uncovering environmental exposure effects on human gut microbiota. *PLoS computational biology*, 18(5):e1010044, 2022.
- Staiger, D. O. and Stock, J. H. Instrumental variables regression with weak instruments, 1994.
- Sun, L., Pennells, L., Kaptoge, S., Nelson, C. P., Ritchie, S. C., Abraham, G., Arnold, M., Bell, S., Bolton, T., Burgess, S., et al. Polygenic risk scores in cardiovascular risk prediction: A cohort study and modelling analyses. *PLoS medicine*, 18(1):e1003498, 2021.
- Sussex, S., Makarova, A., and Krause, A. Model-based causal bayesian optimization. *arXiv preprint arXiv:2211.10257*, 2022.
- Sussman, J. B. and Hayward, R. A. An iv for the rct: using instrumental variables to adjust for treatment contamination in randomised controlled trials. *Bmj*, 340, 2010.
- Sutton, R. S. Reinforcement learning: An introduction. *A Bradford Book*, 2018.
- Sverchkov, Y. and Craven, M. A review of active learning approaches to experimental design for uncovering biological networks. *PLOS Computational Biology*, 13(6):1–26, 06 2017. doi: 10.1371/journal.pcbi.1005466. URL <https://doi.org/10.1371/journal.pcbi.1005466>.
- Swanson, S. A., Hernán, M. A., Miller, M., Robins, J. M., and Richardson, T. S. Partial identification of the average treatment effect using instrumental variables: review of methods for binary instruments, treatments, and outcomes. *Journal of the American Statistical Association*, 113(522):933–947, 2018.
- Szepesvari, C. Causality from the perspective of reinforcement learning. In *Machine Learning for Causal Inference, Counterfactual Prediction, and Autonomous Action (CausalML) Workshop, ICML*, 2018.
- Tam, G. H. F., Chang, C., and Hung, Y. S. Gene regulatory network discovery using pairwise granger causality. *IET systems biology*, 7(5):195–204, 2013.
- Tejada-Lapueta, A., Bertin, P., Bauer, S., Aliee, H., Bengio, Y., and Theis, F. J. Causal machine learning for single-cell genomics. *arXiv preprint arXiv:2310.14935*, 2023.
- Theil, H. Repeated least-squares applied to complete equation systems. *Mimeo. The Hague: Central Planning Bureau*, 1953a.
- Theil, H. Estimating and simultaneous correlation in complete equation systems. *Mimeo. The Hague: Central Planning Bureau*, 1953b.
- Timpson, N. J., Greenwood, C. M., Soranzo, N., Lawson, D. J., and Richards, J. B. Genetic architecture: the shape of the genetic contribution to human traits and disease. *Nature Reviews Genetics*, 19(2):110–124, 2018.
- Trexler, P. Development of gnotobiotics and contamination control in laboratory animal science. 1999.
- Tsuchida, C. A., Brandes, N., Bueno, R., Trinidad, M., Mazumder, T., Yu, B., Hwang, B., Chang, C., Liu, J., Sun, Y., et al. Mitigation of chromosome loss in clinical crispr-cas9-engineered t cells. *Cell*, 186(21):4567–4582, 2023.

- Turnbaugh, P. J., Ley, R. E., Mahowald, M. A., Magrini, V., Mardis, E. R., and Gordon, J. I. An obesity-associated gut microbiome with increased capacity for energy harvest. *nature*, 444(7122):1027–1031, 2006.
- Turnbaugh, P. J., Hamady, M., Yatsunencko, T., Cantarel, B. L., Duncan, A., Ley, R. E., Sogin, M. L., Jones, W. J., Roe, B. A., Affourtit, J. P., et al. A core gut microbiome in obese and lean twins. *nature*, 457(7228):480–484, 2009.
- VanEvery, H., Franzosa, E. A., Nguyen, L. H., and Huttenhower, C. Microbiome epidemiology and association studies in human health. *Nature Reviews Genetics*, 24(2):109–124, 2023.
- von Kügelgen, J., Rubenstein, P. K., Schölkopf, B., and Weller, A. Optimal experimental design via bayesian optimization: active causal structure learning for gaussian process networks. *arXiv preprint arXiv:1910.03962*, 2019.
- Wald, A. The fitting of straight lines if both variables are subject to error. *Annals of Mathematical Statistics*, 11:284–300, 1940. URL <https://api.semanticscholar.org/CorpusID:123294717>.
- Wang, C., Hu, J., Blaser, M. J., Li, H., and Birol, I. Estimating and testing the microbial causal mediation effect with high-dimensional and compositional microbiome data. *Bioinformatics*, 2020. ISSN 14602059. doi: 10.1093/bioinformatics/btz565.
- Wang, C., Ahn, J., Tarpey, T., Yi, S. S., Hayes, R. B., and Li, H. A microbial causal mediation analytic tool for health disparity and applications in body mass index. *Microbiome*, 11(1): 164, July 2023. ISSN 2049-2618. doi: 10.1186/s40168-023-01608-9.
- Wang, R., Zhang, H., and Lin, X. Debiased estimating equation method for versatile and efficient mendelian randomization using large numbers of correlated weak and invalid instruments. *arXiv preprint arXiv:2408.05386*, 2024.
- Wang, X., Jiang, Y., Zhang, N. R., and Small, D. S. Sensitivity analysis and power for instrumental variable studies. *Biometrics*, 74(4):1150–1160, 2018.
- Weichwald, S., Mogensen, S. W., Lee, T. E., Baumann, D., Kroemer, O., Guyon, I., Trimpe, S., Peters, J., and Pfister, N. Learning by doing: Controlling a dynamical system using causality, control, and reinforcement learning. In *NeurIPS 2021 Competitions and Demonstrations Track*, pp. 246–258. PMLR, 2022.
- Wen, Y., Huang, J., Guo, S., Elyahu, Y., Monsonogo, A., Zhang, H., Ding, Y., and Zhu, H. Applying causal discovery to single-cell analyses using causalcell. *Elife*, 12:e81464, 2023.
- Williams, R. J. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Machine learning*, 8:229–256, 1992.
- Wright, P. *The Tariff on Animal and Vegetable Oils, by Philip G. Wright, with the Aid of the Council and Staff of the Institute of Economics*. [Preface by Harold G. Moulton.]. Macmillan Company, 1928. URL <https://books.google.de/books?id=CS9vXwAACAAJ>.
- Xia, Y. Mediation analysis of microbiome data and detection of causality in microbiome studies. *Inflammation, Infection, and Microbiome in Cancers: Evidence, Mechanisms, and Implications*, pp. 457–509, 2021.
- Xu, L., Chen, Y., Srinivasan, S., de Freitas, N., Doucet, A., and Gretton, A. Learning deep features in instrumental variable regression. *arXiv preprint arXiv:2010.07154*, 2020.

- Yang, J., Lee, S. H., Goddard, M. E., and Visscher, P. M. Gcta: a tool for genome-wide complex trait analysis. *The American Journal of Human Genetics*, 88(1):76–82, 2011.
- Yang, J., Ferreira, T., Morris, A. P., Medland, S. E., of ANthropometric Traits (GIANT) Consortium, G. I., Replication, D. G., analysis (DIAGRAM) Consortium, M., Madden, P. A., Heath, A. C., Martin, N. G., Montgomery, G. W., et al. Conditional and joint multiple-snp analysis of gwas summary statistics identifies additional variants influencing complex traits. *Nature genetics*, 44(4):369–375, 2012.
- Yatsunenکو, T., Rey, F. E., Manary, M. J., Trehan, I., Dominguez-Bello, M. G., Contreras, M., Magris, M., Hidalgo, G., Baldassano, R. N., Anokhin, A. P., et al. Human gut microbiome viewed across age and geography. *nature*, 486(7402):222–227, 2012.
- Zemplenyi, M. and Miller, J. W. Bayesian optimal experimental design for inferring causal structure. *Bayesian Analysis*, 2021. doi: 10.1214/22-ba1335.
- Zhang, H., Chen, J., Feng, Y., Wang, C., Li, H., and Liu, L. Mediation effect selection in high-dimensional and compositional microbiome data. *Statistics in Medicine*, 40, 11 2020a. doi: 10.1002/sim.8808.
- Zhang, J. Designing optimal dynamic treatment regimes: A causal reinforcement learning approach. In *International conference on machine learning*, pp. 11012–11022. PMLR, 2020.
- Zhang, J., Kumor, D., and Bareinboim, E. Causal imitation learning with unobserved confounders. *Advances in neural information processing systems*, 33:12263–12274, 2020b.
- Zhang, R., Imaizumi, M., Schölkopf, B., and Muandet, K. Maximum moment restriction for instrumental variable regression. *arXiv preprint arXiv:2010.07684*, 2020c.
- Zhang, R., Imaizumi, M., Schölkopf, B., and Muandet, K. Instrumental variable regression via kernel maximum moment loss. *Journal of Causal Inference*, 11(1):20220073, 2023. doi: doi:10.1515/jci-2022-0073. URL <https://doi.org/10.1515/jci-2022-0073>.
- Zhang, X.-H., Tee, L. Y., Wang, X.-G., Huang, Q.-S., and Yang, S.-H. Off-target effects in crispr/cas9-mediated genome engineering. *Molecular Therapy - Nucleic Acids*, 4:e264, 2015.