

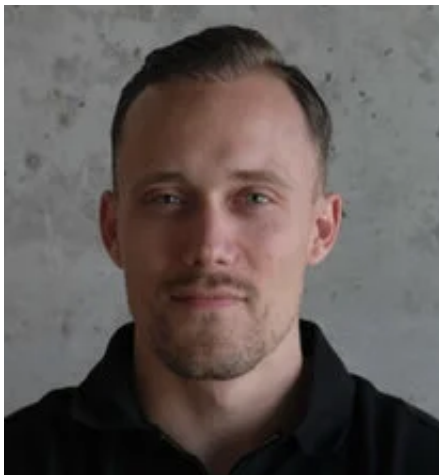
Kernel Methods for Causal Effect Estimation

Arthur Gretton

Gatsby Computational Neuroscience Unit,
Google DeepMind

Modern Kernel Methods, EMS 2026

Talk on behalf of Mattes Mollenhauer



Outline

Recap, **kernel** ridge regression

Causal effect estimation, **observed** covariates:

- Heterogeneous response curve (**HR**)

Causal effect estimation, **hidden** covariates:

- ... **instrumental** variables (and **proxy** variables)

What's new? What is it good for?

- Treatment A , covariates X , etc can be **multivariate, complicated...**
- ...by using **kernel** or **adaptive neural net** feature representations

Model fitting: kernel ridge regression

Learn $\gamma_0(x) := \mathbb{E}[Y|X = x]$ from features $\varphi(x_i)$ with outcomes y_i :

$$\hat{\gamma} = \arg \min_{\gamma \in \mathcal{H}} \left(\sum_{i=1}^n \left(y_i - \langle \gamma, \varphi(x_i) \rangle_{\mathcal{H}_X} \right)^2 + \lambda \|\gamma\|_{\mathcal{H}_X}^2 \right).$$

Model fitting: kernel ridge regression

Learn $\gamma_0(x) := \mathbb{E}[Y|X = x]$ from features $\varphi(x_i)$ with outcomes y_i :

$$\hat{\gamma} = \arg \min_{\gamma \in \mathcal{H}} \left(\sum_{i=1}^n \left(y_i - \langle \gamma, \varphi(x_i) \rangle_{\mathcal{H}_X} \right)^2 + \lambda \|\gamma\|_{\mathcal{H}_X}^2 \right).$$

Infinite dimensional solution at x :

$$\hat{\gamma}(x) = \widehat{C}_{YX} (\widehat{C}_X + \lambda)^{-1} \varphi(x)$$

$$\widehat{C}_{YX} = \frac{1}{n} \sum_{i=1}^n [y_i \varphi(x_i)]$$

$$\widehat{C}_X = \frac{1}{n} \sum_{i=1}^n [\varphi(x_i) \otimes \varphi(x_i)]$$

Notation : $(f \otimes g)h = f \langle g, h \rangle_{\mathcal{H}_X}$

Model fitting: kernel ridge regression

Learn $\gamma_0(x) := \mathbb{E}[Y|X = x]$ from **features** $\varphi(x_i)$ with outcomes y_i :

$$\hat{\gamma} = \arg \min_{\gamma \in \mathcal{H}} \left(\sum_{i=1}^n \left(y_i - \langle \gamma, \varphi(x_i) \rangle_{\mathcal{H}_X} \right)^2 + \lambda \|\gamma\|_{\mathcal{H}_X}^2 \right).$$

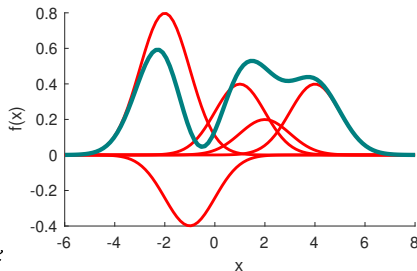
Kernel solution at x
(as weighted sum of y)

$$\hat{\gamma}(x) = \sum_{i=1}^n y_i \beta_i(x)$$

$$\beta(x) = (K_{XX} + \lambda I)^{-1} k_{Xx}$$

$$(K_{XX})_{ij} = k(x_i, x_j) = \langle \varphi(x_i), \varphi(x_j) \rangle_{\mathcal{H}_X}$$

$$(k_{Xx})_i = k(x_i, x)$$



KRR: consistency in RKHS norm

Define **population feature covariance**:

$$C_X := \mathbb{E} [\varphi(X) \otimes \varphi(X)] \quad \langle f, C_X g \rangle_{\mathcal{H}_X} = \mathbb{E} [f(X)g(X)]$$

$$C_X = \sum_{i=1}^{\infty} \eta_i (\sqrt{\eta_i} e_i) \otimes (\sqrt{\eta_i} e_i)$$

$$\eta_i e_i(x) = \int k(x, x') e_i(x') dP_X(x'), \quad \text{assume for simplicity: } \text{supp}(P_X) = \mathcal{X}$$

Assume problem **well specified**:

KRR: consistency in RKHS norm

Define population feature covariance:

$$C_X := \mathbb{E} [\varphi(X) \otimes \varphi(X)] \quad \langle f, C_X g \rangle_{\mathcal{H}_X} = \mathbb{E} [f(X)g(X)]$$

$$C_X = \sum_{i=1}^{\infty} \eta_i (\sqrt{\eta_i} e_i) \otimes (\sqrt{\eta_i} e_i)$$

$$\eta_i e_i(x) = \int k(x, x') e_i(x') dP_X(x'), \quad \text{assume for simplicity: } \text{supp}(P_X) = \mathcal{X}$$

Assume problem well specified:

$$\blacksquare \gamma_* \in \mathcal{H}_X^c \text{ where } \mathcal{H}_X^c \subset \mathcal{H}_X, \quad c \in (1, 2]$$

KRR: consistency in RKHS norm

Define **population feature covariance**:

$$C_X := \mathbb{E} [\varphi(X) \otimes \varphi(X)] \quad \langle f, C_X g \rangle_{\mathcal{H}_X} = \mathbb{E} [f(X)g(X)]$$

$$C_X = \sum_{i=1}^{\infty} \eta_i (\sqrt{\eta_i} e_i) \otimes (\sqrt{\eta_i} e_i)$$

$$\eta_i e_i(x) = \int k(x, x') e_i(x') dP_X(x'), \quad \text{assume for simplicity: } \text{supp}(P_X) = \mathcal{X}$$

Assume problem **well specified**:

■ $\gamma_* \in \mathcal{H}_X^c$ where $\mathcal{H}_X^c \subset \mathcal{H}_X$, $c \in (1, 2]$

$$\gamma_* = \sum_i \gamma_{*,i} e_i \quad \text{where} \quad \|\gamma_*\|_{\mathcal{H}_X^c}^2 = \sum_{i=1}^{\infty} \frac{\gamma_{*,i}^2}{\eta_i^c} < \infty$$

- Larger $c \implies$ smoother γ_* $\implies$ easier problem.

KRR: consistency in RKHS norm

Define **population feature covariance**:

$$C_X := \mathbb{E} [\varphi(X) \otimes \varphi(X)] \quad \langle f, C_X g \rangle_{\mathcal{H}_X} = \mathbb{E} [f(X)g(X)]$$

$$C_X = \sum_{i=1}^{\infty} \eta_i (\sqrt{\eta_i} e_i) \otimes (\sqrt{\eta_i} e_i)$$

$$\eta_i e_i(x) = \int k(x, x') e_i(x') dP_X(x'), \quad \text{assume for simplicity: } \text{supp}(P_X) = \mathcal{X}$$

Assume problem **well specified**:

■ $\gamma_* \in \mathcal{H}_X^c$ where $\mathcal{H}_X^c \subset \mathcal{H}_X$, $c \in (1, 2]$

$$\gamma_* = \sum_i \gamma_{*,i} e_i \quad \text{where} \quad \|\gamma_*\|_{\mathcal{H}_X^c}^2 = \sum_{i=1}^{\infty} \frac{\gamma_{*,i}^2}{\eta_i^c} < \infty$$

• Larger $c \implies$ smoother γ_* $\implies$ easier problem.

■ Eigenspectrum decay of input feature covariance, $\eta_j \sim j^{-b}$, $b \geq 1$

• Larger $b \implies$ easier problem

KRR: consistency in RKHS norm

Consistency for models in class $\mathcal{M}(\textcolor{red}{c}, \textcolor{brown}{b})$ [A, Theorem 1.ii]

$$\|\hat{\gamma} - \gamma_*\|_{L_2(\mathcal{X})}^2 = O_P\left(n^{-\frac{\textcolor{red}{c}}{\textcolor{red}{c}+1/\textcolor{brown}{b}}}\right) \quad \|\hat{\gamma} - \gamma_*\|_{\mathcal{H}_{\mathcal{X}}}^2 = O_P\left(n^{-\frac{\textcolor{red}{c}-1}{\textcolor{red}{c}+1/\textcolor{brown}{b}}}\right)$$

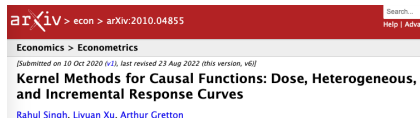
Example:

$\mathcal{X} = \mathbb{R}^d$, $\mathcal{H}_{\mathcal{X}} := W^{\nu,2}(\mathcal{X})$ Sobolev space with $\nu > \frac{d}{2}$ (Matérn kernel) gives $\textcolor{brown}{b} = \frac{2\nu}{d}$ [B, Corollary 2].

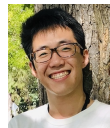
Then $\gamma_* \in \mathcal{H}_{\mathcal{X}}^{\textcolor{red}{c}}$ with $\textcolor{red}{c} \in (1, 2]$ gives consistency in RKHS and L_2 .

Causal effects, observed covariates

Kernel features (Biometrika 2023):



NN features (ICLR 2023):



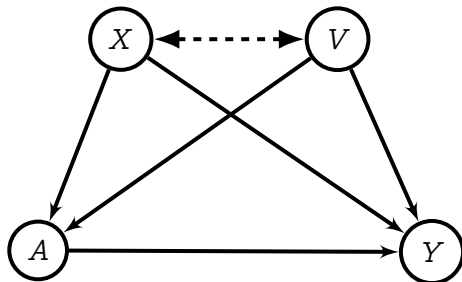
Code for NN and kernel causal estimation with observed covariates:

<https://github.com/liyuan9988/DeepFrontBackDoor/>

Heterogeneous response curve

US job corps:

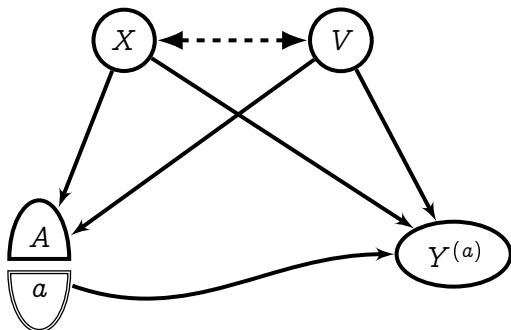
- X : confounder/context (education, marital status, ...)
- A : treatment (training hours)
- Y : outcome (percent employed)
- V : age



Heterogeneous response curve

US job corps:

- X : confounder/context (education, marital status, ...)
- A : treatment (training hours)
- Y : outcome (percent employed)
- V : age



Intervention: $\mathbb{E} \left[Y^{(a)} \mid V = v \right]$

Assume: (1) Stable Unit Treatment Value Assumption (aka “no interference”), (2)

Conditional exchangeability $Y^{(a)} \perp\!\!\!\perp A \mid X$. (3) Overlap.

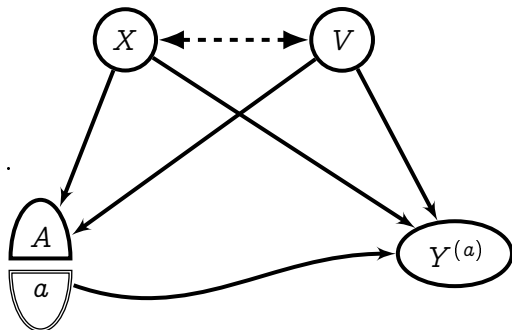
Heterogeneous response curve

Well-specified setting:

$$\begin{aligned}\mathbb{E}[Y|a, x, v] &=: \gamma_*(a, x, v) \\ &= \langle \gamma_*, \varphi(a) \otimes \varphi(x) \otimes \varphi(v) \rangle.\end{aligned}$$

Heterogeneous response:

$$\begin{aligned}\text{HR}(a, v) \\ = \mathbb{E} \left[Y^{(a)} \mid V = v \right]\end{aligned}$$



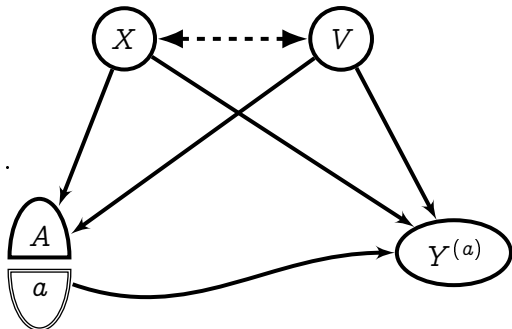
Heterogeneous response curve

Well-specified setting:

$$\begin{aligned}\mathbb{E}[Y|a, x, v] &=: \gamma_*(a, x, v) \\ &= \langle \gamma_*, \varphi(a) \otimes \varphi(x) \otimes \varphi(v) \rangle.\end{aligned}$$

Heterogeneous response:

$$\begin{aligned}\text{HR}(a, v) &= \mathbb{E} \left[Y^{(a)} \mid V = v \right] \\ &= \mathbb{E} \left[\langle \gamma_*, \varphi(a) \otimes \varphi(X) \otimes \varphi(V) \rangle \mid V = v \right]\end{aligned}$$



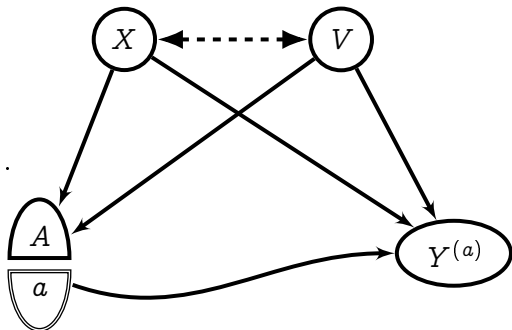
Heterogeneous response curve

Well-specified setting:

$$\begin{aligned}\mathbb{E}[Y|a, x, v] &=: \gamma_*(a, x, v) \\ &= \langle \gamma_*, \varphi(a) \otimes \varphi(x) \otimes \varphi(v) \rangle.\end{aligned}$$

Heterogeneous response:

$$\begin{aligned}\text{HR}(a, v) &= \mathbb{E} \left[Y^{(a)} \mid V = v \right] \\ &= \mathbb{E} \left[\langle \gamma_*, \varphi(a) \otimes \varphi(X) \otimes \varphi(V) \rangle \mid V = v \right] \\ &= \dots?\end{aligned}$$



How to take conditional expectation?

Density estimation for $p(X \mid V = v)$? Sample from $p(X \mid V = v)$?

Heterogeneous response curve

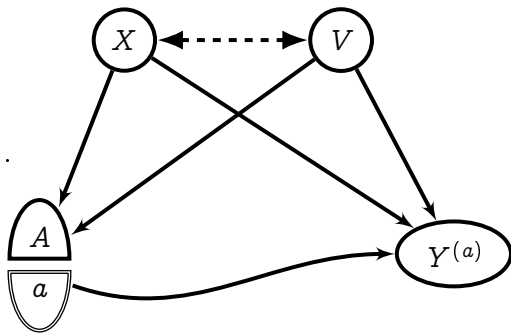
Well-specified setting:

$$\begin{aligned}\mathbb{E}[Y|a, x, v] &=: \gamma_*(a, x, v) \\ &= \langle \gamma_*, \varphi(a) \otimes \varphi(x) \otimes \varphi(v) \rangle.\end{aligned}$$

Heterogeneous response:

$$\begin{aligned}\text{HR}(a, v) &= \mathbb{E} \left[Y^{(a)} \mid V = v \right] \\ &= \mathbb{E} \left[\langle \gamma_*, \varphi(a) \otimes \varphi(X) \otimes \varphi(V) \rangle \mid V = v \right] \\ &= \langle \gamma_*, \varphi(a) \otimes \underbrace{\mathbb{E}[\varphi(X) \mid V = v]}_{\mu_{X|V=v}} \otimes \varphi(v) \rangle\end{aligned}$$

Learn **conditional mean embedding**: $\mu_{X|V=v} := \mathbb{E}_X [\varphi(X) \mid V = v]$



Regressing from feature space to feature space

Our goal: an operator $F_* : \mathcal{H}_{\mathcal{Y}} \rightarrow \mathcal{H}_{\mathcal{X}}$ such that

$$F_* \varphi(v) = \mathbb{E}_X [\varphi(X) | V = v] =: \mu_{X|V=v}$$

Song, Huang, Smola, Fukumizu (2009). Hilbert space embeddings of conditional distributions with applications to dynamical systems.

Grunewalder, Lever, Baldassarre, Patterson, G, Pontil (2012). Conditional mean embeddings as regressors.

Grunewalder, G, Shawe-Taylor (2013) Smooth operators.

Li, Meunier, Mollenhauer, G (2022), Optimal Rates for Regularized Conditional Mean Embedding Learning

Regressing from feature space to feature space

Our goal: an operator $F_* : \mathcal{H}_{\mathcal{V}} \rightarrow \mathcal{H}_{\mathcal{X}}$ such that

$$F_* \varphi(v) = \mathbb{E}_{\mathcal{X}} [\varphi(X) | V = v] =: \mu_{X|V=v}$$

Assume F_* Hilbert-Schmidt,

$$F_* \in S_2(\mathcal{H}_{\mathcal{V}}, \mathcal{H}_{\mathcal{X}})$$

Implied smoothness assumption:

$$\mathbb{E}[h(X) | V = v] \in \mathcal{H}_{\mathcal{V}} \quad \forall h \in \mathcal{H}_{\mathcal{X}}$$

Song, Huang, Smola, Fukumizu (2009). Hilbert space embeddings of conditional distributions with applications to dynamical systems.

Grunewalder, Lever, Baldassarre, Patterson, G, Pontil (2012). Conditional mean embeddings as regressors.

Grunewalder, G, Shawe-Taylor (2013) Smooth operators.

Li, Meunier, Mollenhauer, G (2022), Optimal Rates for Regularized Conditional Mean Embedding Learning

Regressing from feature space to feature space

Our goal: an operator $F_* : \mathcal{H}_\mathcal{V} \rightarrow \mathcal{H}_\mathcal{X}$ such that

$$F_* \varphi(\boldsymbol{v}) = \mathbb{E}_X [\varphi(\boldsymbol{X}) | \boldsymbol{V} = \boldsymbol{v}] =: \mu_{\boldsymbol{X} | \boldsymbol{V} = \boldsymbol{v}}$$

Assume F_* Hilbert-Schmidt,

$$F_* \in S_2(\mathcal{H}_\mathcal{V}, \mathcal{H}_\mathcal{X})$$

Implied smoothness assumption:

$$\mathbb{E}[h(\boldsymbol{X}) | \boldsymbol{V} = \boldsymbol{v}] \in \mathcal{H}_\mathcal{V} \quad \forall h \in \mathcal{H}_\mathcal{X}$$

A Smooth Operator

Song, Huang, Smola, Fukumizu (2009). Hilbert space embeddings of conditional distributions with applications to dynamical systems.

Grunewalder, Lever, Baldassarre, Patterson, G, Pontil (2012). Conditional mean embeddings as regressors.

Grunewalder, G, Shawe-Taylor (2013) Smooth operators.

Li, Meunier, Mollenhauer, G (2022), Optimal Rates for Regularized Conditional Mean Embedding Learning

Regressing from feature space to feature space

Our goal: an operator $F_* : \mathcal{H}_\mathcal{V} \rightarrow \mathcal{H}_\mathcal{X}$ such that

$$F_* \varphi(\mathbf{v}) = \mathbb{E}_\mathbf{X} [\varphi(\mathbf{X}) | \mathbf{V} = \mathbf{v}] =: \mu_{\mathbf{X} | \mathbf{V}=\mathbf{v}}$$

Assume F_* Hilbert-Schmidt,

$$F_* \in S_2(\mathcal{H}_\mathcal{V}, \mathcal{H}_\mathcal{X})$$

Implied smoothness assumption:

$$\mathbb{E}[h(\mathbf{X}) | \mathbf{V} = \mathbf{v}] \in \mathcal{H}_\mathcal{V} \quad \forall h \in \mathcal{H}_\mathcal{X}$$

Kernel ridge regression from $\varphi(\mathbf{v})$ to infinite features $\varphi(\mathbf{x})$:

$$\hat{F} = \operatorname{argmin}_{F \in S_2(\mathcal{H}_\mathcal{V}, \mathcal{H}_\mathcal{X})} \sum_{\ell=1}^n \|\varphi(\mathbf{x}_\ell) - F \varphi(\mathbf{v}_\ell)\|_{\mathcal{H}_\mathcal{X}}^2 + \lambda_2 \|F\|_{HS}^2$$

Song, Huang, Smola, Fukumizu (2009). Hilbert space embeddings of conditional distributions with applications to dynamical systems.

Grunewalder, Lever, Baldassarre, Patterson, G, Pontil (2012). Conditional mean embeddings as regressors.

Grunewalder, G, Shawe-Taylor (2013) Smooth operators.

Li, Meunier, Mollenhauer, G (2022), Optimal Rates for Regularized Conditional Mean Embedding Learning

Regressing from feature space to feature space

Our goal: an operator $F_* : \mathcal{H}_V \rightarrow \mathcal{H}_X$ such that

$$F_* \varphi(v) = \mathbb{E}_X [\varphi(X) | V = v] =: \mu_{X|V=v}$$

Assume F_* Hilbert-Schmidt,

$$F_* \in S_2(\mathcal{H}_V, \mathcal{H}_X)$$

Implied smoothness assumption:

$$\mathbb{E}[h(X) | V = v] \in \mathcal{H}_V \quad \forall h \in \mathcal{H}_X$$

Kernel ridge regression from $\varphi(v)$ to infinite features $\varphi(x)$:

$$\hat{F} = \operatorname{argmin}_{F \in S_2(\mathcal{H}_V, \mathcal{H}_X)} \sum_{\ell=1}^n \|\varphi(x_\ell) - F \varphi(v_\ell)\|_{\mathcal{H}_X}^2 + \lambda_2 \|F\|_{HS}^2$$

Ridge regression solution:

$$\mu_{X|V=v} := \mathbb{E}[\varphi(X) | V = v] \approx \hat{F} \varphi(v) = \sum_{\ell=1}^n \varphi(x_\ell) \beta_\ell(v)$$
$$\beta(v) = [K_{VV} + \lambda_2 I]^{-1} k_{Vv}$$

Consistency of Vector Valued Least Squares



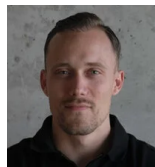
NeurIPS (2022)



NeurIPS (2024)



JMLR (2024)



Consistency of Vector Valued Least Squares

Define C_V as covariance of features $\varphi(\mathbf{v})$,

$$C_V = \mathbb{E}(\varphi(\mathbf{V}) \otimes \varphi(\mathbf{V})), \quad \mathbb{E}_V(f(\mathbf{V})g(\mathbf{V})) = \langle f, C_V g \rangle_{\mathcal{H}_V}$$

- Assume problem well specified [C, Assumption 4]

$$F_* \in S_2(\mathcal{H}_V^{c_1}, \mathcal{H}_X), \quad c_1 \in (1, 2],$$

larger $c_1 \implies$ smoother $F_* \implies$ easier problem.

- Eigenspectrum of C_V decays as $\eta_j(C_V) \sim j^{-b_1}$, $b_1 \geq 1$.

[A] Li, Meunier, Mollenhauer, G (2022), Optimal Rates for Regularized CME Learning

[B] Li, Meunier, Mollenhauer, G (2024), Towards Optimal Sobolev Norm Rates for VV RLS Algorithm

[C] Singh, Xu, G (2023)

Grunewalder, Lever, Baldassarre, Patterson, G, Pontil (2012).

Caponnetto, De Vito (2007).

Consistency of Vector Valued Least Squares

Define C_V as covariance of features $\varphi(v)$,

$$C_V = \mathbb{E}(\varphi(V) \otimes \varphi(V)), \quad \mathbb{E}_V(f(V)g(V)) = \langle f, C_V g \rangle_{\mathcal{H}_V}$$

- Assume problem well specified [C, Assumption 4]

$$F_* \in S_2(\mathcal{H}_V^{c_1}, \mathcal{H}_X), \quad c_1 \in (1, 2],$$

larger $c_1 \implies$ smoother $F_* \implies$ easier problem.

- Eigenspectrum of C_V decays as $\eta_j(C_V) \sim j^{-b_1}$, $b_1 \geq 1$.

Consistency [A, Theorem 2, Theorem 3] (minimax)

$$\left\| \widehat{F} - F_* \right\|_{\text{HS}}^2 = O_P \left(n^{-\frac{c_1-1}{c_1+1/b_1}} \right),$$

best rate is $O_P(n^{-1/4})$.

[A] Li, Meunier, Mollenhauer, G (2022), Optimal Rates for Regularized CME Learning

[B] Li, Meunier, Mollenhauer, G (2024), Towards Optimal Sobolev Norm Rates for VV RLS Algorithm

[C] Singh, Xu, G (2023)

Grunewalder, Lever, Baldassarre, Patterson, G, Pontil (2012).

Caponnetto, De Vito (2007).

Heterogeneous Response from Samples

Empirical HR:

$$\begin{aligned} & \widehat{\text{HR}}(a, \mathbf{v}) \\ &= Y^\top (K_{AA} \odot K_{XX} \odot K_{VV} + n\lambda I)^{-1} (K_{Aa} \odot \underbrace{K_{XX}(K_{VV} + n\lambda_1 I)^{-1} K_{V\mathbf{v}}}_{\text{from } \hat{\mu}_{X|V=\mathbf{v}}} \odot K_{Vv}) \end{aligned}$$

Heterogeneous Response from Samples

Empirical HR:

$$\begin{aligned} \widehat{\text{HR}}(a, v) \\ = Y^\top (K_{AA} \odot K_{XX} \odot K_{VV} + n\lambda I)^{-1} (K_{Aa} \odot \underbrace{K_{XX}(K_{VV} + n\lambda_1 I)^{-1} K_{Vv}}_{\text{from } \hat{\mu}_{X|V=v}} \odot K_{Vv}) \end{aligned}$$

Consistency: [A, Theorem 2]

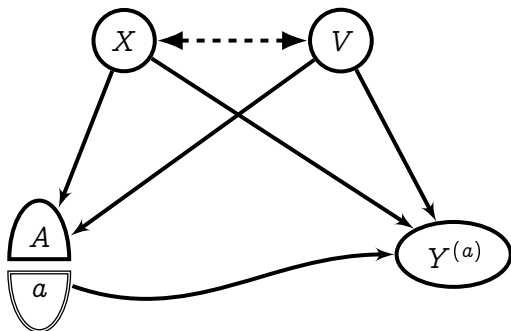
$$\|\widehat{HR} - \text{HR}\|_\infty = O_P \left(n^{-\frac{1}{2} \frac{c-1}{c+1/b}} + n^{-\frac{1}{2} \frac{c_1-1}{c_1+1/b_1}} \right).$$

Follows from consistency of $\hat{F}$ and $\hat{\gamma}$

Heterogeneous response: example

US job corps:

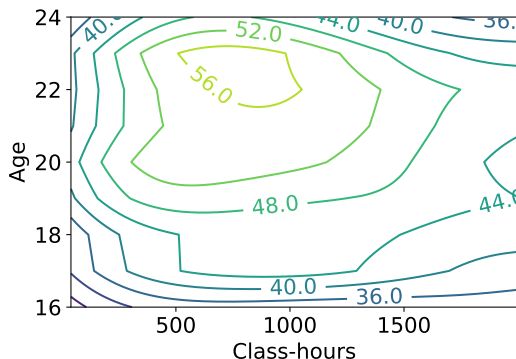
- X : confounder/context (education, marital status, ...)
- A : treatment (training hours)
- Y : outcome (percent employed)
- V : age



Empirical HR:

$$\widehat{\text{HR}}(a, \mathbf{v}) = \langle \hat{\gamma}_0, \varphi(a) \otimes \underbrace{\hat{F} \varphi(\mathbf{v})}_{\hat{\mathbb{E}}[\varphi(X) | V=\mathbf{v}]} \otimes \varphi(\mathbf{v}) \rangle$$

Heterogeneous response: results



Average percentage employment $Y^{(a)}$ for class hours a , **conditioned on age v** . Given around 12-14 weeks of classes:

- 16 y/o: employment increases from 28% to at most 36%.
- 22 y/o: percent employment increases from 40% to 56%.

Instrumental variable regression

Instrumental variable regression with kernels

- X : unobserved confounder.
- A : treatment
- Y : outcome
- Z : instrument

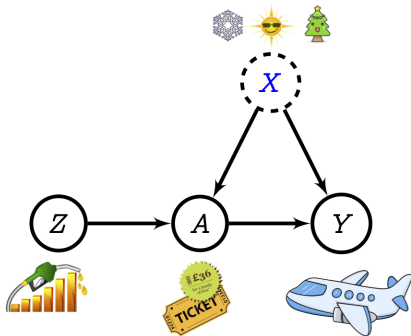
Assumptions

$$\mathbb{E}[X|Z] = 0$$

$$Z \not\perp\!\!\!\perp A$$

$$(Y \perp\!\!\!\perp Z|A)_{G_{\bar{A}}}$$

$$Y = \langle \gamma, \phi(A) \rangle + X$$



Instrumental variable regression with kernels

- X : unobserved confounder.
- A : treatment
- Y : outcome
- Z : instrument

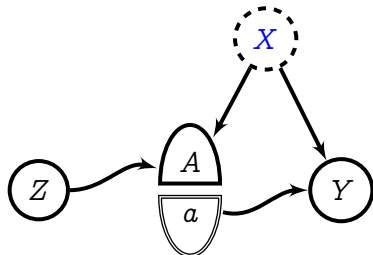
Assumptions

$$\mathbb{E}[X|Z] = 0$$

$$Z \not\perp\!\!\!\perp A$$

$$(Y \perp\!\!\!\perp Z|A)_{G_{\bar{A}}}$$

$$Y = \langle \gamma, \phi(A) \rangle + X$$

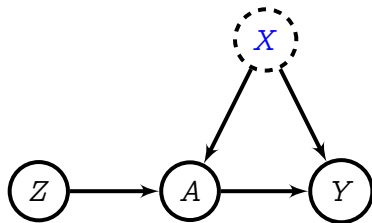


Dose-response curve:

$$\text{DR}(a) = \int \mathbb{E}(Y|X, a) dp(X) = \langle \gamma, \phi(a) \rangle$$

Instrumental variable regression with kernels

- X : unobserved confounder.
- A : treatment
- Y : outcome
- Z : instrument



Assumptions

$$\mathbb{E}[X|Z] = 0$$

$$Z \not\perp\!\!\!\perp A$$

$$(Y \perp\!\!\!\perp Z|A)_{G_{\bar{A}}}$$

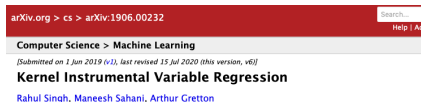
$$Y = \langle \gamma, \phi(A) \rangle + X$$

IV regression: Condition both sides on Z ,

$$\mathbb{E}[Y|Z] = \mathbb{E}[\langle \gamma, \phi(A) \rangle | Z] + \underbrace{\mathbb{E}[X|Z]}_{=0}$$

Two-stage least squares for IV regression

Kernel features (NeurIPS 2019):



Consistency (submitted):



Code for NN and kernel IV methods:

<https://github.com/liyuan9988/DeepFeatureIV/>

IV using kernel features

Stage 2 regression (IV): Solve for γ with RR loss:

$$\hat{\gamma} = \underset{\gamma \in \mathcal{H}_{\mathcal{A}}}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{\ell=1}^n \left[\left(y_{\ell} - \left\langle \gamma, \hat{\mathbb{E}}[\varphi(A)|z_{\ell}] \right\rangle_{\mathcal{H}_{\mathcal{A}}} \right)^2 \right] + \lambda_2 \|\gamma\|_{\mathcal{H}_{\mathcal{A}}}^2$$

Stage 1 regression: define $\hat{F} \in S_2(\mathcal{H}_{\mathcal{Z}}, \mathcal{H}_{\mathcal{A}})$:

$$\hat{\mathbb{E}}[\varphi(A)|z] = \hat{F}\varphi(z)$$

Two stage least squares

Xu, Chen, Srinivasan, De Freitas, Doucet, G. (2021) Learning Deep Features in Instrumental Variable Regression

Consistency: goal

Strong error control:

$$\|\hat{\gamma} - \gamma_*\|_{L_2(A)}$$

Consistency: goal

Strong error control:

$$\|\hat{\gamma} - \gamma_*\|_{L_2(A)}$$

Stage 2 loss naturally expressed in weak metric:

$$\|\mathcal{T}(\hat{\gamma} - \gamma_*)\|_{L_2(Z)}$$

where $\mathcal{T}\gamma = \mathbb{E}(\gamma(A)|Z)$.

Consistency: goal

Strong error control:

$$\|\hat{\gamma} - \gamma_*\|_{L_2(A)}$$

Stage 2 loss naturally expressed in weak metric:

$$\|\mathcal{T}(\hat{\gamma} - \gamma_*)\|_{L_2(Z)}$$

where $\mathcal{T}\gamma = \mathbb{E}(\gamma(A)|Z)$.

By Jensen,

$$\|\mathcal{T}(\hat{\gamma} - \gamma_*)\|_{L_2(Z)} \leq \|\hat{\gamma} - \gamma_*\|_{L_2(A)}$$

Weak control does not imply strong control: need additional link assumption.

Strong/weak geometry, link condition

Strong geometry:

$$\begin{aligned}\|\gamma\|_{L_2(A)}^2 &= \mathbb{E} \left[\langle \gamma, \varphi(A) \rangle_{\mathcal{H}_A}^2 \right] \\ &= \langle \gamma, C_A \gamma \rangle_{\mathcal{H}_A}\end{aligned}$$

where

$$C_A = \mathbb{E}_A [\varphi(A) \otimes \varphi(A)]$$

Strong/weak geometry, link condition

Strong geometry:

$$\begin{aligned}\|\gamma\|_{L_2(A)}^2 &= \mathbb{E} \left[\langle \gamma, \varphi(A) \rangle_{\mathcal{H}_A}^2 \right] \\ &= \langle \gamma, C_A \gamma \rangle_{\mathcal{H}_A}\end{aligned}$$

where

$$C_A = \mathbb{E}_A [\varphi(A) \otimes \varphi(A)]$$

Weak geometry:

$$\begin{aligned}\|\mathcal{T}\gamma\|_{L_2(Z)}^2 &= \mathbb{E} \left[\mathbb{E} \left[\langle \gamma, \varphi(A) \rangle_{\mathcal{H}_A} \mid Z \right]^2 \right] \\ &= \langle \gamma, C_M \gamma \rangle_{\mathcal{H}_A}\end{aligned}$$

where

$$\begin{aligned}C_M &= \mathbb{E} \left[\mu_{A|Z} \otimes \mu_{A|Z} \right], \\ \mu_{A|Z} &:= \mathbb{E} [\varphi(A) \mid Z].\end{aligned}$$

Strong/weak geometry, link condition

Strong geometry:

$$\begin{aligned}\|\gamma\|_{L_2(A)}^2 &= \mathbb{E} \left[\langle \gamma, \varphi(A) \rangle_{\mathcal{H}_A}^2 \right] \\ &= \langle \gamma, C_A \gamma \rangle_{\mathcal{H}_A}\end{aligned}$$

where

$$C_A = \mathbb{E}_A [\varphi(A) \otimes \varphi(A)]$$

Weak geometry:

$$\begin{aligned}\|\mathcal{T}\gamma\|_{L_2(Z)}^2 &= \mathbb{E} \left[\mathbb{E} \left[\langle \gamma, \varphi(A) \rangle_{\mathcal{H}_A} \mid Z \right]^2 \right] \\ &= \langle \gamma, C_M \gamma \rangle_{\mathcal{H}_A}\end{aligned}$$

where

$$\begin{aligned}C_M &= \mathbb{E} \left[\mu_{A|Z} \otimes \mu_{A|Z} \right], \\ \mu_{A|Z} &:= \mathbb{E} [\varphi(A) \mid Z].\end{aligned}$$

Link condition: $\exists \zeta > 0, \theta \geq 1$ such that

$$C_A \preceq \zeta C_M^{1/\theta}$$

Earlier link conditions two-sided $C_A^\gamma \asymp C_M$: Nair, Pereverzev, Tautenhahn (2005); Chen, Reiss (2011).

Convergence result

Model class $\mathcal{M}(c, b, \theta)$: c smoothness of $\gamma_* \in \mathcal{H}_{\mathcal{A}}^c$, b effective dimension, θ ill-posedness.

Convergence result

Model class $\mathcal{M}(c, b, \theta)$: c smoothness of $\gamma_* \in \mathcal{H}_{\mathcal{A}}^c$, b effective dimension, θ ill-posedness.

Assume: Stage 1 samples $m = n^\alpha$ sufficiently large relative to Stage 2 samples n . Then

$$\|\hat{\gamma} - \gamma_*\|_{L_2(A)}^2 = O_P(n^{-\frac{c}{c+1/b+\theta-1}}).$$

Convergence result

Model class $\mathcal{M}(c, b, \theta)$: c smoothness of $\gamma_* \in \mathcal{H}_{\mathcal{A}}^c$, b effective dimension, θ ill-posedness.

Assume: Stage 1 samples $m = n^\alpha$ sufficiently large relative to Stage 2 samples n . Then

$$\|\hat{\gamma} - \gamma_*\|_{L_2(A)}^2 = O_P(n^{-\frac{c}{c+1/b+\theta-1}}).$$

Lower bound (minimax):

$$\inf_{\hat{\gamma} \text{ NPIV models in } \mathcal{M}(c, b, \theta)} \sup \|\hat{\gamma} - \gamma_*\|_{L_2(A)}^2 \gtrsim n^{-\frac{c}{c+1/b+\theta-1}}$$

For $\theta = 1$ recover nonparametric rate for regression.

Conclusions

Kernel (and neural net) solutions:

- ...for dose-response, heterogeneous response, dynamic treatment effects
- ...with treatment A , covariates X , V , multivariate, “complicated”
- Convergence guarantees

Unobserved confounding:

- IV (and proxy) methods

Code available for all methods

Research support

Work supported by:

The Gatsby Charitable Foundation



Google DeepMind



Questions?

