

Ranking Causal Effects – Causal Models @Uber

Ben Draves



About Me

My name is Benjamin Draves and I'm a Staff Applied Scientist at Uber in NYC. I work on building [ML recommendation and targeting systems](#) to grow Uber's consumer and earner businesses.

Data Science Topics that Excite Me

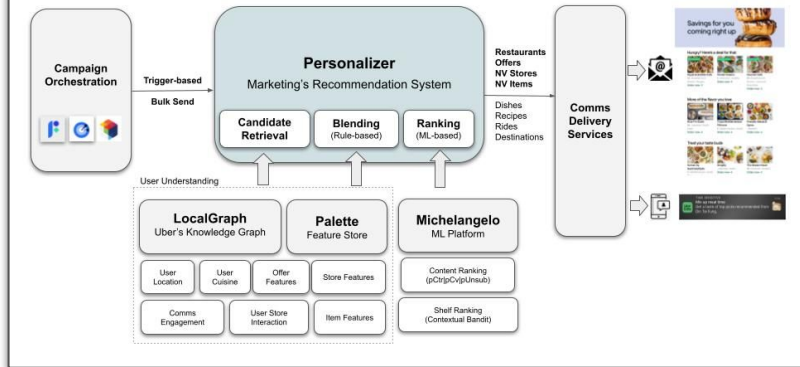
- Recommendation Systems
- Causal Machine Learning
- Graph Embeddings and Inference

Experiences & Education

- [2022-Present] Applied Scientist @ Uber
- [2017-2022] Ph.D. | M.A. in Stat @ Boston University
- [2014-2017] B.S. in Math @ Lafayette College

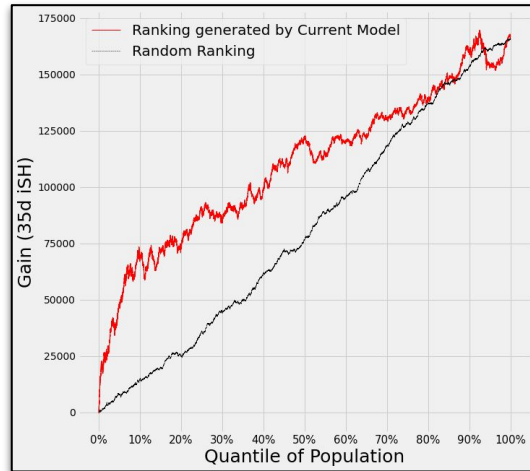
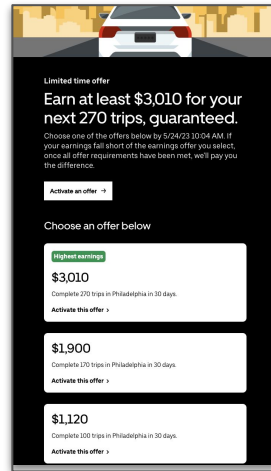
Growing Uber's Core Businesses

How | An extensible recommendation platform that aggregates user's lifetime & realtime signals from across Uber and optimizes for various marketing objectives.



Recommendation & Personalization

I work on personalizing customer's experiences in out-of-app channels to drive product awareness, adoption, and growth. This oftentimes manifests in **recommending restaurants, merchants, and offers** to customers to encourage product exploration.



Targeting & Orchestration

I develop communication targeting and orchestration systems to provide relevant and timely notifications to customers which service core business goals. This oftentimes manifests in marrying causal machine learning models and techniques from ads pacing systems.

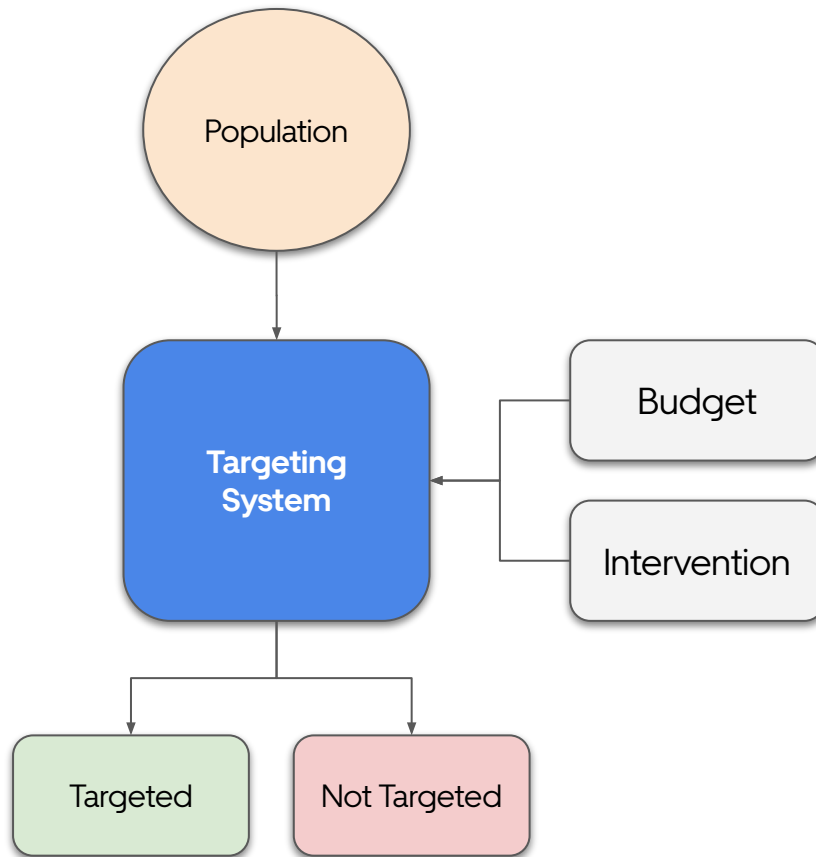
Motivating Problem | Marketing Intervention with Budget Constraint

Problem Statement

Given an a population P , intervention T , and budget B , how we do we optimally allocation treatment across the population?

Examples

- Free Uber One membership for 1M
- Bonus to drivers for completing 5 trips
- Weekly newsletter with restaurant promos



Agenda

01 Metalearners

02 DL & Robustness Extensions

03 Ranking Causal Effects

04 Conclusion & Discussion

Metalearners

Notation | Neyman-Rubin Potential Outcomes Framework

Data Generation Process

Assume we have the DGP $W = (X, T, Y) \sim \mathbb{P}$ where

- $X \in \mathcal{X} \subseteq \mathbb{R}^d$ represent covariates
- $T \in \{0, 1\}$ is a binary treatment
- $Y \in \mathbb{R}$ is a continuous outcome

We collect an i.i.d. sample $\{(x_i, t_i, y_i)\}_{i=1}^n \stackrel{i.i.d.}{\sim} \mathbb{P}$

CATE Definition

Each unit has a two potential outcomes $(Y(0), Y(1))$ where $Y(t)$ represents the outcome given $T = t$. The CATE is given by

$$\tau(x) = \mathbb{E}[Y(1) - Y(0) | X = x]$$

Identification Assumptions

We take the standard assumptions (i) Consistency (ii) Positivity (iii) Unconfoundedness so we can write CATE as

$$\tau(x) = \mu_1(x) - \mu_0(x)$$

where $\mu_t(x) = \mathbb{E}[Y(t) | X = x]$ are the conditional outcome regressions.

Supplements

Let the propensity score be given by

$$e(x) = \mathbb{P}(T = 1 | X = x)$$

Let the nuisance functions be given by

$$\eta = (\mu_1, \mu_0, e)$$

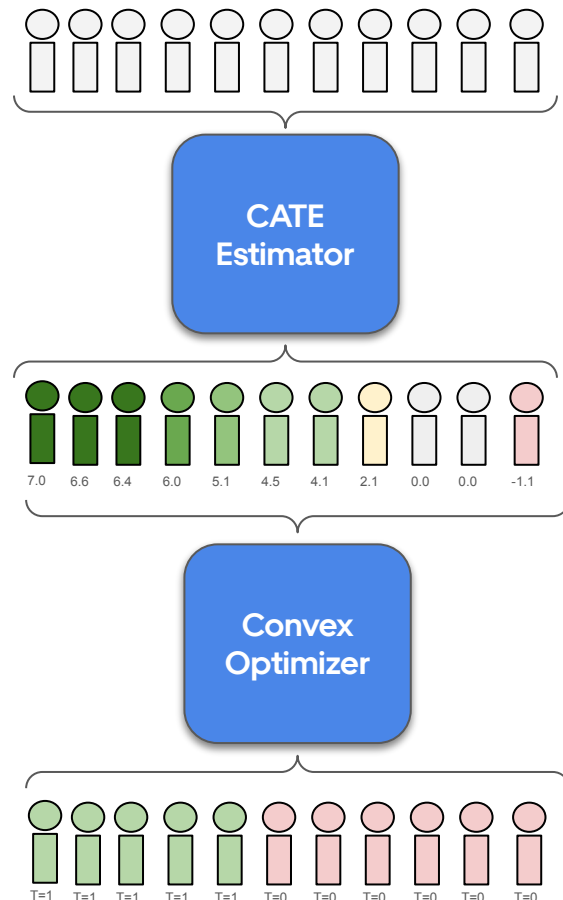
Motivating Problem | Marketing Intervention with Budget Constraint

Problem Statement

Given an a population **P**, intervention **T**, and budget **B**, how we do we optimally allocation treatment across the population?

Solution Formulation

1. Identify growth metric of interest Y
2. Estimate inc. increase in growth metric $\hat{\tau}(x_i)$
3. Use convex optimization to assign treatment



Motivating Problem | Marketing Intervention with Budget Constraint

Problem Statement

Given an a population P , intervention T , and budget B , how we do we optimally allocation treatment across the population?

Solution Formulation

1. Identify growth metric of interest Y
2. Estimate inc. increase in growth metric $\hat{\tau}(x_i)$
3. Use convex optimization to assign treatment

$$\begin{aligned} & \underset{z_i \in \{0,1\}}{\text{maximize}} && \sum_{i=1}^N z_i \hat{\tau}(x_i) \\ & \text{subject to} && \sum_{i=1}^N z_i \hat{c}(x_i) \leq B \end{aligned}$$

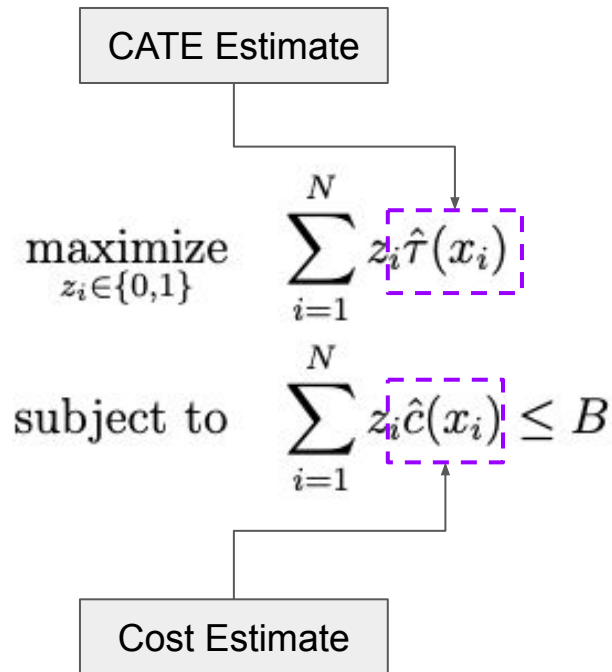
Motivating Problem | Marketing Intervention with Budget Constraint

Problem Statement

Given an a population P , intervention T , and budget B , how we do we optimally allocation treatment across the population?

Solution Formulation

1. Identify growth metric of interest Y
2. Estimate inc. increase in growth metric $\hat{\tau}(x_i)$
3. Use convex optimization to assign treatment



Metalearners | S/T/X Learners

S Learner

- Fit **one** ML model with the treatment indicator as a covariates

$$\mu(x, t) = \mathbb{E}[Y^{\text{obs}} | X = x, T = t]$$

- Estimate CATE via two prediction calls

$$\hat{\tau}_S(x) = \hat{\mu}(x, 1) - \hat{\mu}(x, 0)$$

```
metalearners > s_learner.py
1  """S-Learner: estimates CATE using a single model with treatment as a feature."""
2
3  import numpy as np
4  from BaseLearner import BaseLearner
5
6  class SLearner:
7      def __init__(self, base_learner: BaseLearner = None):
8          self.model = base_learner
9
10     def fit(self, X, T, Y):
11         #augment X with treatment
12         XT = np.column_stack([X, T])
13         self.model.fit(XT, Y)
14         return self
15
16     def predict_cate(self, X):
17         n = len(X)
18         X1 = np.column_stack([X, np.ones(n)])
19         X0 = np.column_stack([X, np.zeros(n)])
20         return self.model.predict(X1) - self.model.predict(X0)
21
```

Metalearners | S/T/X Learners

T Learner

- Fit **two** ML models with the treatment / control data respectively

$$\mu_1(x) = \mathbb{E}[Y(1)|X = x] \quad \mu_0(x) = \mathbb{E}[Y(0)|X = x]$$

- Estimate CATE via differencing

$$\hat{\tau}_T(x) = \hat{\mu}_1(x) - \hat{\mu}_0(x)$$

```
metalearners > t_learner.py
1  """T-Learner: estimates CATE using separate models for T/C groups."""
2
3  import numpy as np
4  from BaseLearner import BaseLearner
5
6
7  class TLearner:
8      def __init__(self, base_learner: BaseLearner = None):
9          factory = lambda: base_learner
10         self.model_t = factory()
11         self.model_c = factory()
12
13     def fit(self, X, T, Y):
14         self.model_t.fit(X[T == 1], Y[T == 1])
15         self.model_c.fit(X[T == 0], Y[T == 0])
16         return self
17
18     def predict_cate(self, X):
19         return self.model_t.predict(X) - self.model_c.predict(X)
20
```

Metalearners | S/T/X Learners

X Learner

- Fit **two** ML models with the treatment / control data respectively

$$\mu_1(x) = \mathbb{E}[Y(1)|X = x] \quad \mu_0(x) = \mathbb{E}[Y(0)|X = x]$$

- Impute ITE for each unit

$$\tilde{D}(1) = Y(1) - \hat{\mu}_0(X) \quad \tilde{D}(0) = \hat{\mu}_1(X) - Y(0)$$

- Fit **two** models to estimate impute ITEs

$$\hat{\tau}_1(x) = \mathbb{E}[\tilde{D}^{(1)}|X = x] \quad \hat{\tau}_0(x) = \mathbb{E}[\tilde{D}^{(0)}|X = x]$$

- Estimates CATE via

$$\hat{\tau}_X(x) = \hat{e}(x)\hat{\tau}_0(x) + (1 - \hat{e}(x))\hat{\tau}_1(x)$$

```
metalearners > x_learner.py
1  """X-Learner: uses pseudo-outcomes and propensity weighting to handle imbalanced treatment."""
2
3  import numpy as np
4  from BaseLearner import BaseLearner
5
6  class XLearner:
7      def __init__(self,
8                  outcome_learner: BaseLearner = None,
9                  effect_learner: BaseLearner = None,
10                 propensity_learner: BaseLearner = None
11                 ):
12         self.mu0 = outcome_learner
13         self.mu1 = outcome_learner
14         self.tau0 = effect_learner
15         self.tau1 = effect_learner
16         self.propensity = propensity_learner
17
18     def fit(self, X, T, Y):
19         X0, Y0 = X[T == 0], Y[T == 0]
20         X1, Y1 = X[T == 1], Y[T == 1]
21
22         # Stage 1: fit outcome models on each group
23         self.mu0.fit(X0, Y0)
24         self.mu1.fit(X1, Y1)
25
26         # Stage 2: pseudo-outcomes cross-impute
27         D1 = Y1 - self.mu0.predict(X1) # treated: actual - counterfactual(control)
28         D0 = self.mu1.predict(X0) - Y0 # control: counterfactual(treated) - actual
29
30         self.tau1.fit(X1, D1)
31         self.tau0.fit(X0, D0)
32
33         # Stage 3: propensity model for weighting
34         self.propensity.fit(X, T)
35         return self
36
37     def predict_cate(self, X):
38         e = self.propensity.predict_proba(X)[:, 1] # P(T=1 | X)
39         return e * self.tau0.predict(X) + (1 - e) * self.tau1.predict(X)
40
```

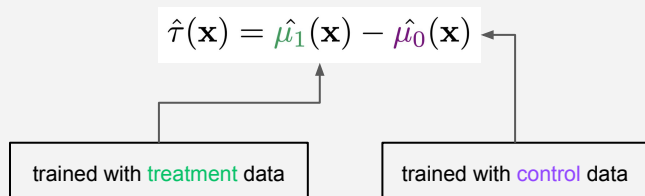
Metalearners | Statistical Guarantees & Empirical Evidence

Meta Learner	Statistical Guarantees	Strengths	Weaknesses
S	Converges if base learners do: $\ \hat{\mu}(x, t) - \mu(x, t)\ = O(r_n)$ implies $\ \hat{\tau}_S(x) - \tau(x)\ = O(r_n)$	• Simple to implement • Lower variance due to treatment group pooling	• Model can regularize away treatment covariate
T	Converges if base learners do: $\ \hat{\mu}_0(x) - \mu_0(x)\ = O(r_{0,n})$ $\ \hat{\mu}_1(x) - \mu_1(x)\ = O(r_{1,n})$ implies $\ \hat{\tau}_T(x) - \tau(x)\ = O(r_{1,n} + r_{0,n})$	• Simple to implement • Lower bias due to no treatment group pooling	• Model adapts poorly to class imbalances
X	Converges if base learners do: $\ \hat{\tau}_X(x) - \tau(x)\ = O(r_{\mu,n} + r_{\tau,n} + n^{-1/2})$	• Handles class imbalance • If the CATE curve is simple, then the rate improves to the larger group	• Complex model management

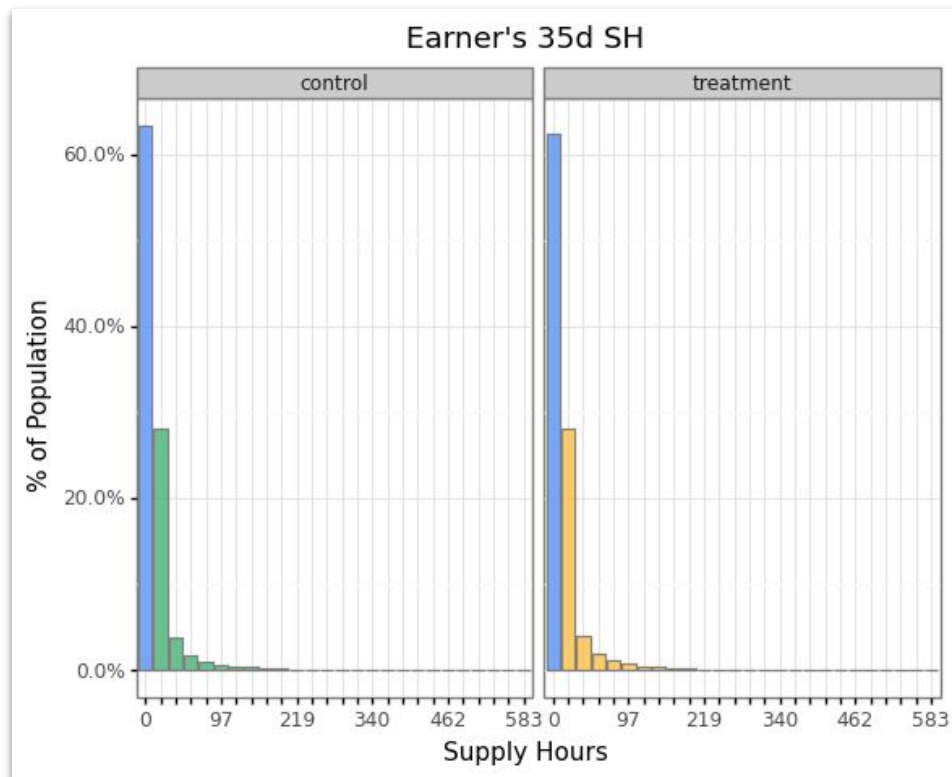
Metalearners | Uber Example - Earner Incentives

T-Learner

- Fit two SH regression models
- Incremental SH prediction is delta between the two model predictions



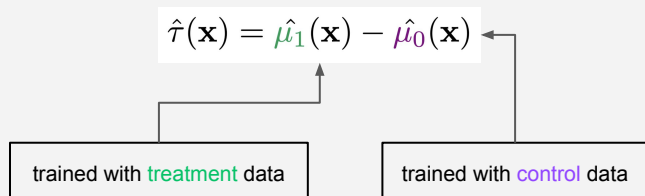
- Base learners are XGBoost Regressor fit with a [tweedie loss](#)



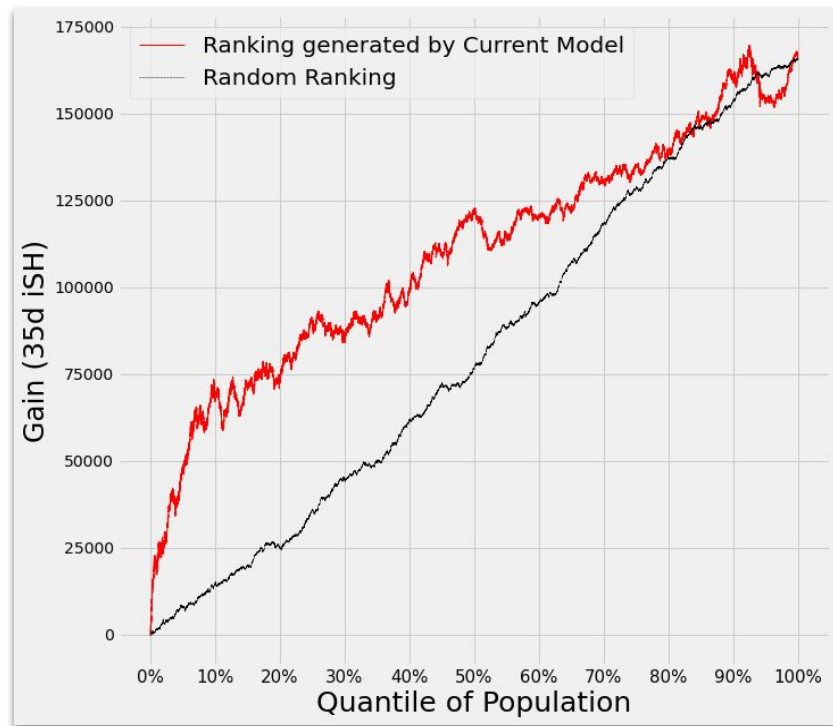
Metalearners | Uber Example - Earner Incentives

T-Learner

- Fit two SH regression models
- Incremental SH prediction is delta between the two model predictions



- Base learners are XGBoost Regressor fit with a [tweedie loss](#)



DL & Robustness Extensions

Modern Improvements | Areas for Improvement

Model Complexity

Maintaining multiple models in practice can be complex. Can we simplify the engineering design?

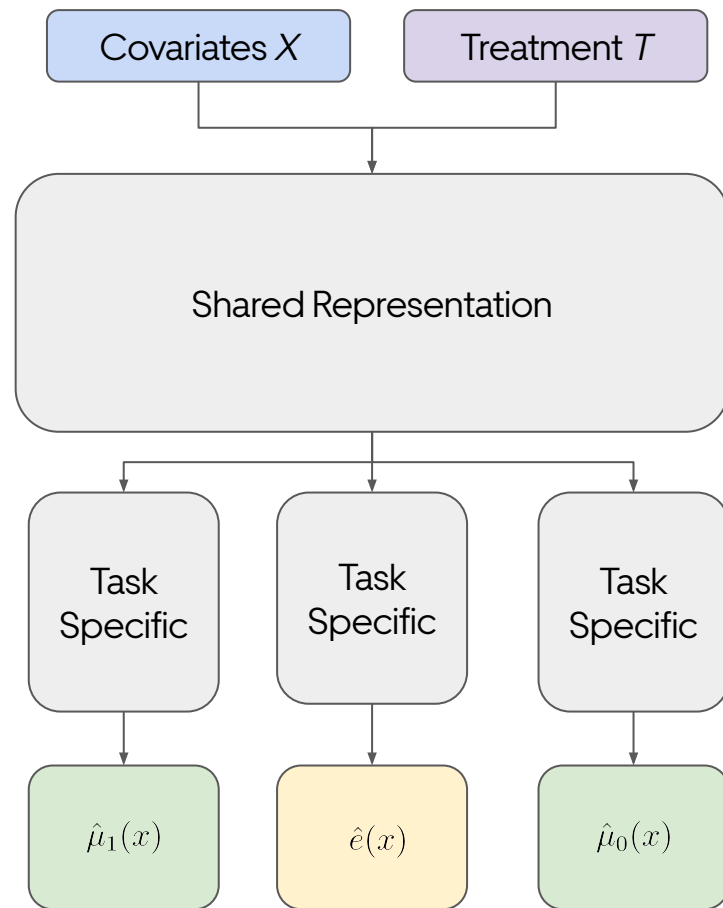
Statistical Robustness

S/T/X learners suffer from “plugin” bias. Can make our CATE estimators more robust to upstream estimation errors?

Modern Improvements | Tarnet, CFR, DragoNet, ...

DL Architectures

- For each S/T/X learner the first stage needs to estimate the nuisance functions $\eta = (\mu_1, \mu_0, e)$ from the covariates
- Can we use an multi-output DL model to make this prediction simpler to manage?
- Many DL methods available ([survey](#)) produce these estimates. We discuss DragonNet
 - DragonNet: $\hat{\eta}_{DN}(x) = (\hat{\mu}_1(x), \hat{\mu}_0(x), \hat{e}(x))$



Modern Improvements | Tarnet, CFR, DragonNet, ...

DL Architectures

- For each S/T/X learner the first stage needs to estimate the nuisance functions $\eta = (\mu_1, \mu_0, e)$ from the covariates
- Can we use an multi-output DL model to make this prediction simpler to manage?
- Many DL methods available ([survey](#)) produce these estimates. We discuss DragonNet
 - DragonNet: $\hat{\eta}_{DN}(x) = (\hat{\mu}_1(x), \hat{\mu}_0(x), \hat{e}(x))$

metalearners > dragonnet.py

```
1  """
2  DragonNet: shared-representation network with joint outcome + propensity heads.
3
4  Reference: Shi et al. 2019 - "Adapting Neural Networks for the Estimation
5  of Treatment Effects" (https://arxiv.org/abs/1906.02120)
6  """
7
8  import numpy as np
9  import torch
10 import torch.nn as nn
11
12 class _DragonNet(nn.Module):
13     def __init__(self, input_dim, hidden_dim, dropout):
14         super().__init__()
15         self.trunk = nn.Sequential(
16             nn.Linear(input_dim, hidden_dim), nn.ELU(),
17             nn.Linear(hidden_dim, hidden_dim), nn.ELU(),
18             nn.Linear(hidden_dim, hidden_dim), nn.ELU(),
19             nn.Dropout(dropout),
20         )
21         self.y0 = nn.Sequential(nn.Linear(hidden_dim, hidden_dim // 2), nn.ELU(),
22                                 nn.Linear(hidden_dim // 2, 1))
23         self.y1 = nn.Sequential(nn.Linear(hidden_dim, hidden_dim // 2), nn.ELU(),
24                                 nn.Linear(hidden_dim // 2, 1))
25         self.propensity = nn.Sequential(nn.Linear(hidden_dim, 1), nn.Sigmoid())
26
27     def forward(self, x):
28         z = self.trunk(x)
29         return self.y0(z).squeeze(-1), self.y1(z).squeeze(-1), self.propensity(z).squeeze(-1)
30
31
```

Modern Improvements | Training

Training Details

- Multiobjective loss
 - Propensity estimation
 - Outcome loss

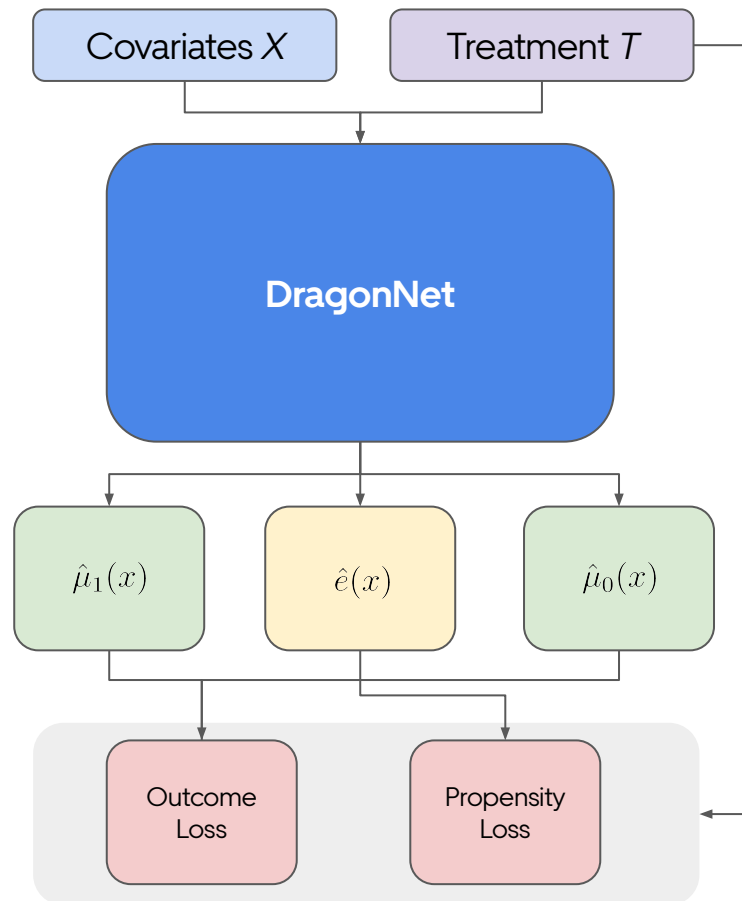
$$\mathcal{L} = \mathcal{L}_{out} + \mathcal{L}_{prop}$$

- Only compute loss based on labeled data for outcomes

$$\mathcal{L}_{out} = -\frac{1}{N} \sum_{i=1}^N [I(t_i = 1)(y_i - \hat{\mu}_1(x_i))^2 + I(t_i = 0)(y_i - \hat{\mu}_0(x_i))^2]$$

- Propensity loss computed for all labels

$$\mathcal{L}_{prop} = -\frac{1}{N} \sum_{i=1}^N y_i \log(\hat{e}(x_i)) + (1 - y_i) \log(1 - \hat{e}(x_i))$$



Modern Improvements | Training

Training Details

- Multiobjective loss
 - Propensity estimation
 - Outcome loss

$$\mathcal{L} = \mathcal{L}_{out} + \mathcal{L}_{prop}$$

- Only compute loss based on labeled data for outcomes

$$\mathcal{L}_{out} = -\frac{1}{N} \sum_{i=1}^N [I(t_i = 1)(y_i - \hat{\mu}_1(x_i))^2 + I(t_i = 0)(y_i - \hat{\mu}_0(x_i))^2]$$

- Propensity loss computed for all labels

$$\mathcal{L}_{prop} = -\frac{1}{N} \sum_{i=1}^N y_i \log(\hat{e}(x_i)) + (1 - y_i) \log(1 - \hat{e}(x_i))$$

metalearners > dragonnet.py

```
32 class DragonNet:
33
34     def fit(self, X, T, Y):
35         X_t = torch.tensor(X, dtype=torch.float32)
36         T_t = torch.tensor(T, dtype=torch.float32)
37         Y_t = torch.tensor(Y, dtype=torch.float32)
38
39         self.net = _DragonNet(X.shape[1], self.hidden_dim, self.dropout)
40         opt = torch.optim.Adam(self.net.parameters(), lr=self.lr)
41         mse = nn.MSELoss()
42         bce = nn.BCELoss()
43
44         n = len(X_t)
45         self.net.train()
46         for _ in range(self.epochs):
47             idx = torch.randperm(n)
48             for start in range(0, n, self.batch_size):
49                 b = idx[start : start + self.batch_size]
50                 xb, tb, yb = X_t[b], T_t[b], Y_t[b]
51
52                 y0, y1, prop = self.net(xb)
53
54                 # Outcome loss: use y1 for treated, y0 for controls
55                 y_hat = torch.where(tb == 1, y1, y0)
56                 outcome_loss = mse(y_hat, yb)
57                 prop_loss = bce(prop, tb)
58
59                 loss = outcome_loss + self.alpha * prop_loss
60                 opt.zero_grad()
61                 loss.backward()
62                 opt.step()
63
64         self.net.eval()
65         return self
66
67     def predict_cate(self, X):
68         self.net.eval()
69         with torch.no_grad():
70             X_t = torch.tensor(X, dtype=torch.float32)
71             y0, y1, _ = self.net(X_t)
72             return (y1 - y0).numpy()
73
74
75
76
77
78
79
80
81
82
83
84
85
86
87
88
89
```

Modern Improvements | Identifying Robustness

What is “Robustness”

- We are interested in estimate the CATE
$$\tau(x) = \mathbb{E}[Y(1)|X = x] - \mathbb{E}[Y(0)|X = x]$$
but oftentimes suffer “plug-in” bias
- Can we devise “second stage” estimation procedures which estimate $\tau(x)$ that are robust to certain errors
- What type of errors do we want to be robust to?

Neyman Orthogonality

If the nuisance estimates are biased, overfit, or converging slowly, we want $\hat{\tau}(x_i)$ to still be consistent and converge quickly.

Double Robustness

If only one of our nuisance estimators (propensity scores or surface function) are correctly specified, then our estimator $\hat{\tau}(x_i)$ is still unbiased and consistent.

Modern Improvements | Takeaways

Most of the best available methods today follow a two stage approach:

- **Stage 1:** Estimate nuisance parameters & necessary parameters for inference with **best methods** from ML.
- **Stage 2:** Use nuisances estimates as an input to the second optimization problem which should be **robust to first stage noise**.

Ranking Causal Effects

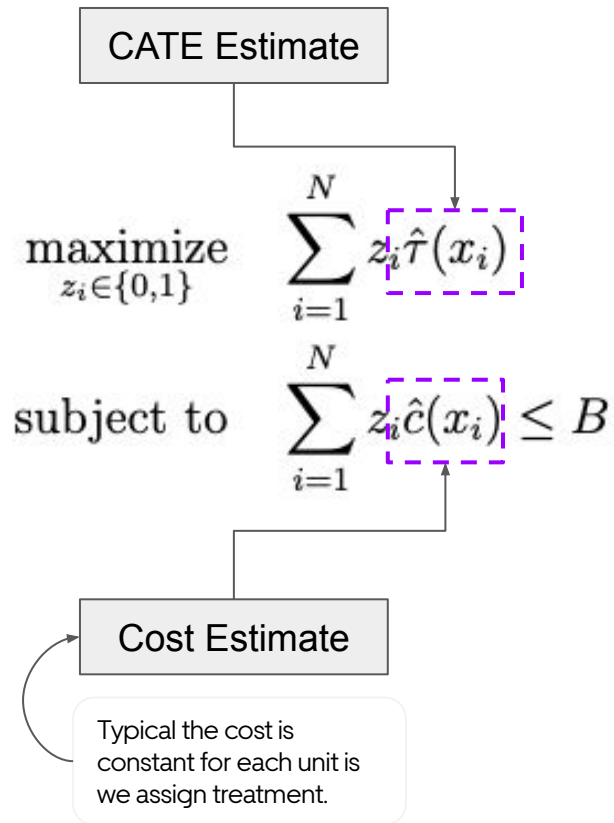
Ranking Causal Effects | Revisiting Our Motivating Example

Problem Statement

Given an a population P , intervention T , and budget B , how we do we optimally allocation treatment across the population?

Solution Formulation

1. Identify growth metric of interest Y
2. Estimate inc. increase in growth metric $\hat{\tau}(x_i)$
3. Use convex optimization to assign treatment



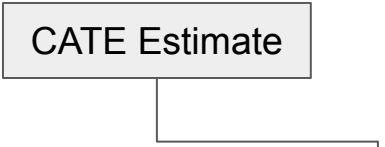
Ranking Causal Effects | Revisiting Our Motivating Example

Problem Statement

Given an a population P , intervention T , and budget B , how we do we optimally allocation treatment across the population?

Solution Formulation

1. Identify growth metric of interest Y
2. Estimate inc. increase in growth metric $\hat{\tau}(x_i)$
3. Use convex optimization to assign treatment


$$\begin{aligned} & \text{maximize}_{z_i \in \{0,1\}} \sum_{i=1}^N z_i \hat{\tau}(x_i) \\ & \text{subject to} \sum_{i=1}^N z_i c \leq B \end{aligned}$$

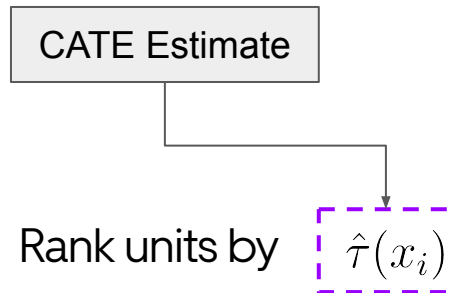
Ranking Causal Effects | Revisiting Our Motivating Example

Problem Statement

Given an a population P , intervention T , and budget B , how we do we optimally allocation treatment across the population?

Solution Formulation

1. Identify growth metric of interest Y
2. Estimate inc. increase in growth metric $\hat{\tau}(x_i)$
3. Select top $k = \lfloor B/c \rfloor$ ranked units



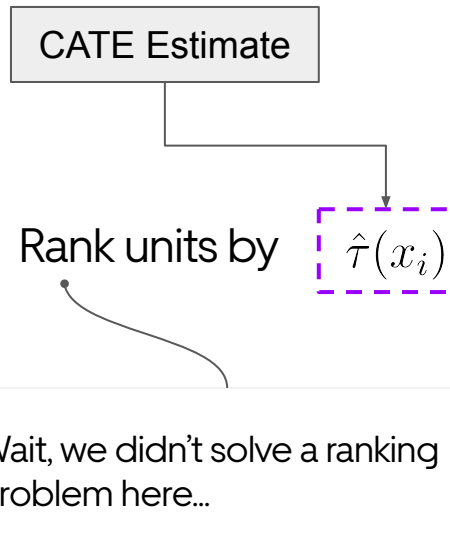
Ranking Causal Effects | Revisiting Our Motivating Example

Problem Statement

Given an a population P , intervention T , and budget B , how we do we optimally allocation treatment across the population?

Solution Formulation

1. Identify growth metric of interest Y
2. Estimate inc. increase in growth metric $\hat{\tau}(x_i)$
3. Select top $k = \lfloor B/c \rfloor$ ranked units



Ranking Causal Effects | Revisiting Our Motivating Example

Problem Statement

Given an a population $\mathcal{P}$, intervention T , and budget B , how we do we optimally allocation treatment across the population?

Solution Formulation

1. Identify growth metric of interest Y
2. Estimate inc. increase in growth metric $\hat{\tau}(x_i)$
3. Select top $k = \lfloor B/c \rfloor$ ranked units

Stage 1: Estimate nuisance parameters & necessary parameters for inference with **best methods** from ML.

$$\eta = (\mu_1, \mu_0, e)$$

Stage 2: Use nuisances estimates as a second optimization problem which should be **robust to first stage noise**.

Ranking units by $\hat{\tau}(x_i)$

Is this true?

Ranking Causal Effects | [Rank-learner: Arno et al. 2026.](#)

Orthogonal Ranking of Treatment Effects

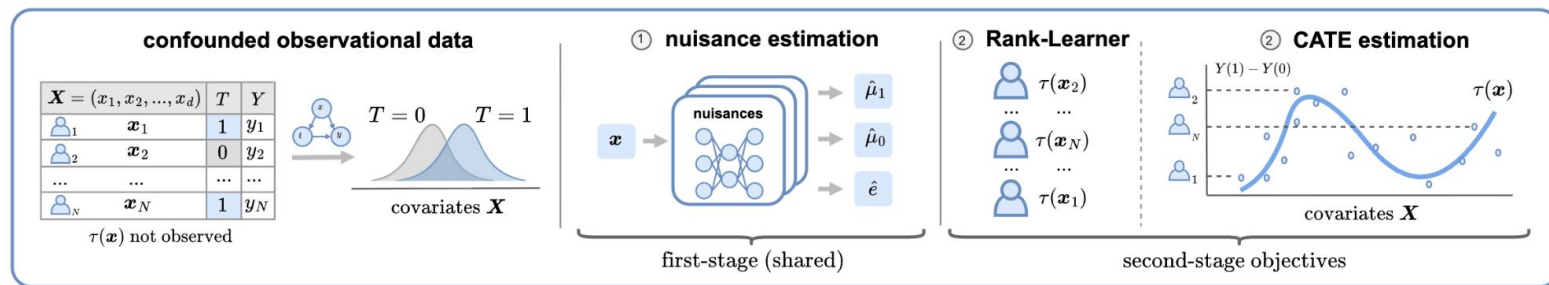


Figure 1. Two-stage learners for ranking treatment effects. *Left:* Confounded observational data with unobserved treatment effects. *Center:* First-stage estimation of nuisance functions (i.e., response surfaces and propensity score). *Right:* Second-stage objectives: our *Rank-Learner* vs. standard CATE estimation. Here, our proposed *Rank-Learner* directly optimizes a pairwise, Neyman-orthogonal ranking objective that targets the relative ordering of treatment effects. In contrast, standard CATE estimation optimizes a pointwise regression objective to recover treatment effect magnitudes, which is harder than necessary for ranking.

Ranking Causal Effects | [Rank-learner: Arno et al. 2026.](#)

Rank Learner Overview

- CATE estimation is generally a much more difficult problem than ranking.
- Authors propose **optimizing a pairwise ranking loss** in the second stage that is **Neyman Orthogonal**
- This allows for
 - **Stage 1** - use the best from ML
 - **Stage 2** - optimizing ranking directly

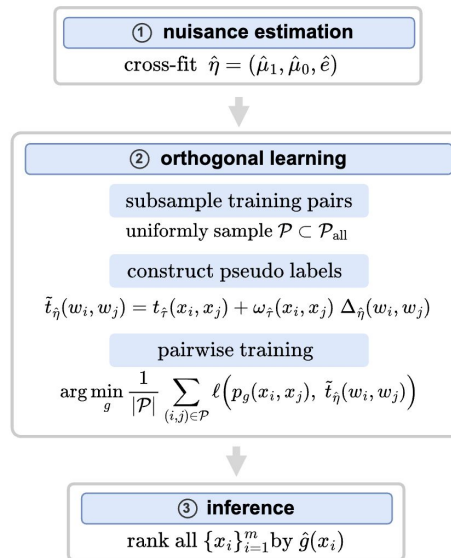


Figure 2. Overview of Rank-Learner. In the first stage, we estimate nuisance functions (response surfaces and propensity score) via cross-fitting. In the second stage, we subsample training pairs, construct the pseudo labels, and learn a scoring function by minimizing the orthogonal ranking objective. Both stages can be instantiated with arbitrary machine learning models. At inference time, we rank individuals directly by their learned scores $\hat{g}(x)$.

Ranking Causal Effects | [Rank-learner: Arno et al. 2026.](#)

Rank Learner Overview

- CATE estimation is generally a much more difficult problem than ranking.
- Authors propose **optimizing a pairwise ranking loss** in the second stage that is **Neyman Orthogonal**
- This allows for
 - **Stage 1** - use the best from ML
 - **Stage 2** - **optimizing ranking directly**

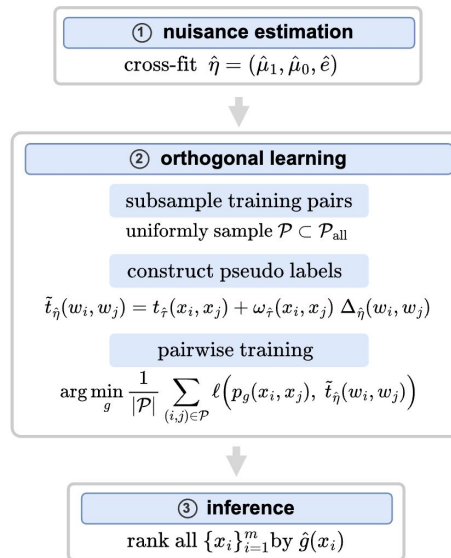


Figure 2. Overview of Rank-Learner. In the first stage, we estimate nuisance functions (response surfaces and propensity score) via cross-fitting. In the second stage, we subsample training pairs, construct the pseudo labels, and learn a scoring function by minimizing the orthogonal ranking objective. Both stages can be instantiated with arbitrary machine learning models. At inference time, we rank individuals directly by their learned scores $\hat{g}(x)$.

Ranking Causal Effects | Learning the Ranker

Ranking Loss Construcion

- **Goal:** Learn a ranking function $g(\cdot) : \mathcal{X} \rightarrow \mathbb{R}$ such that relative ordering is maintained
$$\tau(x) > \tau(x') \implies g(x) > g(x')$$
- **Loss:** Authors take a pairwise ranking approach
 - Use estimated $\hat{\tau}(x_i)$ to derive labels
 - Learn $g(\cdot)$ by estimating those labels

$$\mathcal{L}^{(bin)}(g, \eta) = \mathbb{E}_{X, X'} [\ell(p_g(X, X'), b_\tau(X, X'))]$$

Ranking Causal Effects | Learning the Ranker

Ranking Loss Construcion

- **Goal:** Learn a ranking function $g(\cdot) : \mathcal{X} \rightarrow \mathbb{R}$ such that relative ordering is maintained
 $\tau(x) > \tau(x') \implies g(x) > g(x')$
- **Loss:** Authors take a pairwise ranking approach
 - Use estimated $\hat{\tau}(x_i)$ to derive labels
 - Learn $g(\cdot)$ by estimating those labels

$$\mathcal{L}^{(bin)}(g, \eta) = \mathbb{E}_{X, X'}[\ell(p_g(X, X'), b_\tau(X, X'))]$$

Binary cross entropy loss

$$l(p, t) = -t \log(p) + (1 - t) \log(1 - p)$$

Ranking Causal Effects | Learning the Ranker

Ranking Loss Construcion

- **Goal:** Learn a ranking function $g(\cdot) : \mathcal{X} \rightarrow \mathbb{R}$ such that relative ordering is maintained
 $\tau(x) > \tau(x') \implies g(x) > g(x')$
- **Loss:** Authors take a pairwise ranking approach
 - Use estimated $\hat{\tau}(x_i)$ to derive labels
 - Learn $g(\cdot)$ by estimating those labels

Diagram illustrating the relationship between the probability of one causal effect being greater than another and a 0/1 label, which are used in a binary loss function.

prob of $\tau(x) > \tau(x')$ 0/1 label

$\mathcal{L}^{(bin)}(g, \eta) = \mathbb{E}_{X, X'} [\ell(p_g(X, X'), b_\tau(X, X'))]$

Ranking Causal Effects | Learning the Ranker

Ranking Loss Construcion

- **Goal:** Learn a ranking function $g(\cdot) : \mathcal{X} \rightarrow \mathbb{R}$ such that relative ordering is maintained
 $\tau(x) > \tau(x') \implies g(x) > g(x')$
- **Loss:** Authors take a pairwise ranking approach
 - Use estimated $\hat{\tau}(x_i)$ to derive labels
 - Learn $g(\cdot)$ by estimating those labels

$$\mathcal{L}^{(bin)}(g, \eta) = \mathbb{E}_{X, X'} [\ell(p_g(X, X'), b_\tau(X, X'))]$$

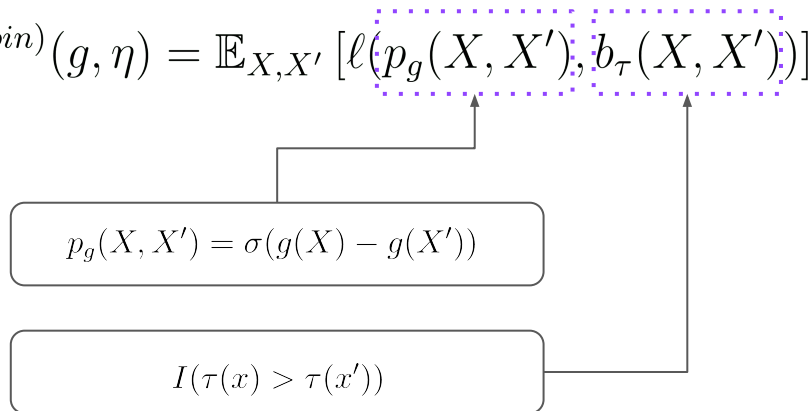
The diagram illustrates the components of the binary ranking loss. It shows two boxes at the bottom: the top box contains the predicted probability $p_g(X, X') = \sigma(g(X) - g(X'))$ and the bottom box contains the true indicator $I(\tau(x) > \tau(x'))$. Arrows from these boxes point to the $p_g(X, X')$ and $b_\tau(X, X')$ terms in the loss function equation above, which is enclosed in a dashed purple box.

Ranking Causal Effects | Learning the Ranker

Ranking Loss Construcion

- **Goal:** Learn a ranking function $g(\cdot) : \mathcal{X} \rightarrow \mathbb{R}$ such that relative ordering is maintained
 $\tau(x) > \tau(x') \implies g(x) > g(x')$
- **Loss:** Authors take a pairwise ranking approach
 - Use estimated $\hat{\tau}(x_i)$ to derive labels
 - Learn $g(\cdot)$ by estimating those labels
- This loss is minimized for any $g(\cdot)$ that preserves ordering of CATE - Much easier than CATE estimation!

$$\mathcal{L}^{(bin)}(g, \eta) = \mathbb{E}_{X, X'} [\ell(p_g(X, X'), b_\tau(X, X'))]$$



Ranking Causal Effects | Learning the Orthogonal Ranker

Orthogonalizing the Loss

- **Goal:** Can we update the loss so the function we learn $g(\cdot)$ is robust to Stage 1 noise?

$$\mathcal{L}^{(bin)}(g, \eta) = \mathbb{E}_{X, X'} [\ell(p_g(X, X'), b_\tau(X, X'))]$$

Ranking Causal Effects | Learning the Orthogonal Ranker

Orthogonalizing the Loss

- **Goal:** Can we update the loss so the function we learn $g(\cdot)$ is robust to Stage 1 noise?
- Authors change the loss in two ways
 1. Smoothing - ensure the loss is differentiable with respect to nuisance

$$\mathcal{L}^{(bin)} \rightarrow \mathcal{L}^{(soft)}$$

$$\mathcal{L}^{(bin)}(g, \eta) = \mathbb{E}_{X, X'} [\ell(p_g(X, X'), b_\tau(X, X'))]$$

$$\mathcal{L}^{(soft)}(g, \eta) = \mathbb{E}_{X, X'} [\ell(p_g(X, X'), t_\tau(X, X'))]$$

“Smooth” Labels

$$t_\tau(X, X') = \sigma \left(\frac{\tau(X) - \tau(X')}{\kappa} \right)$$

Ranking Causal Effects | Learning the Orthogonal Ranker

Orthogonalizing the Loss

- **Goal:** Can we update the loss so the function we learn $g(\cdot)$ is robust to [Stage 1](#) noise?
- Authors change the loss in two ways

1. Smoothing - ensure the loss is differentiable with respect to nuisance

$$\mathcal{L}^{(bin)} \rightarrow \mathcal{L}^{(soft)}$$

2. Orthogonalization - augment the labels with an orthogonal correction

$$\mathcal{L}^{(soft)} \rightarrow \mathcal{L}^{(orth)}$$

$$\mathcal{L}^{(bin)}(g, \eta) = \mathbb{E}_{X, X'} [\ell(p_g(X, X'), b_\tau(X, X'))]$$

$$\mathcal{L}^{(soft)}(g, \eta) = \mathbb{E}_{X, X'} [\ell(p_g(X, X'), t_\tau(X, X'))]$$

$$\mathcal{L}^{(orth)}(g, \eta) = \mathbb{E}_{X, X'} [\ell(p_g(X, X'), \tilde{t}_\eta(X, X'))]$$

“Orthogonalized” Labels

$$\tilde{t}_\eta(X, X') = t_\tau(X, X') + \omega_\tau(X, X') \Delta_\eta(X, X')$$

Ranking Causal Effects | Training the Orthogonal Ranker

Training

Phase 1

1. Estimate nuisance functions with cross fitting

$$\hat{\eta} = (\hat{\mu}_1, \hat{\mu}_0, \hat{e})$$

Phase 2

1. Sample training pairs

$$\mathcal{P} = \{(i, j) : i \in [n], j \in [n], i \neq j\} \subseteq \mathcal{P}_{all}$$

2. Construct pseudo labels

$$\tilde{t}_{\eta}(x_i, x_j) = t_{\tau}(x_i, x_j) + \omega_{\tau}(x_i, x'_j) \Delta_{\eta}(x_i, x'_j)$$

3. Supervised training loop over loss

$$\mathcal{L}^{(orth)}(g, \hat{\eta}) = \frac{1}{|\mathcal{P}|} \sum_{(i,j) \in \mathcal{P}} \ell(p_g(x_i, x_j), \tilde{t}_{\hat{\eta}}(x_i, x_j))$$

Ranking Causal Effects | Training the Orthogonal Ranker

Training

Phase 1

1. Estimate nuisance functions with cross fitting

$$\hat{\eta} = (\hat{\mu}_1, \hat{\mu}_0, \hat{e})$$

Phase 2

1. Sample training pairs

$$\mathcal{P} = \{(i, j) : i \in [n], j \in [n], i \neq j\} \subseteq \mathcal{P}_{all}$$

2. Construct pseudo labels

$$\tilde{t}_{\eta}(x_i, x_j) = t_{\tau}(x_i, x_j) + \omega_{\tau}(x_i, x'_j) \Delta_{\eta}(x_i, x'_j)$$

3. Supervised training loop over loss

$$\mathcal{L}^{(orth)}(g, \hat{\eta}) = \frac{1}{|\mathcal{P}|} \sum_{(i,j) \in \mathcal{P}} \ell(p_g(x_i, x_j), \tilde{t}_{\hat{\eta}}(x_i, x_j))$$

metalearners >  mlp_ranker.py

```
1  """
2  MLPRanker: default ranker for RankLearner.
3  Reference: https://arxiv.org/abs/2602.03517
4  """
5
6  import torch
7  import torch.nn as nn
8  import torch.nn.functional as F
9
10
11 class _MLP(nn.Module):
12     def __init__(self, input_dim, hidden_dim, dropout):
13         super().__init__()
14         self.net = nn.Sequential(
15             nn.Linear(input_dim, hidden_dim), nn.ELU(),
16             nn.Linear(hidden_dim, hidden_dim), nn.ELU(),
17             nn.Dropout(dropout),
18             nn.Linear(hidden_dim, 1),
19         )
20
21     def forward(self, x):
22         return self.net(x).squeeze(-1)
23
```

Ranking Causal Effects | Training the Orthogonal Ranker

Training

Phase 1

1. Estimate nuisance functions with cross fitting

$$\hat{\eta} = (\hat{\mu}_1, \hat{\mu}_0, \hat{e})$$

Phase 2

1. Sample training pairs

$$\mathcal{P} = \{(i, j) : i \in [n], j \in [n], i \neq j\} \subseteq \mathcal{P}_{all}$$

2. Construct pseudo labels

$$\tilde{t}_{\eta}(x_i, x_j) = t_{\tau}(x_i, x_j) + \omega_{\tau}(x_i, x'_j) \Delta_{\eta}(x_i, x'_j)$$

3. Supervised training loop over loss

$$\mathcal{L}^{(orth)}(g, \hat{\eta}) = \frac{1}{|\mathcal{P}|} \sum_{(i,j) \in \mathcal{P}} \ell(p_g(x_i, x_j), \tilde{t}_{\hat{\eta}}(x_i, x_j))$$

```
25 class MLPRanker:
36     def fit(self, X, phi, tau):
37         """
38         X : covariates
39         phi : doubly robust scores  $\phi_{\eta}(W)$ 
40         tau : CATE estimates  $\hat{t}(X) = \mu_1(X) - \mu_0(X)$ 
41         """
42         X_t = torch.tensor(X, dtype=torch.float32)
43         phi_t = torch.tensor(phi, dtype=torch.float32)
44         tau_t = torch.tensor(tau, dtype=torch.float32)
45         n = len(X_t)
46
47         self.mlp = _MLP(X.shape[1], self.hidden_dim, self.dropout)
48         opt = torch.optim.Adam(self.mlp.parameters(), lr=self.lr)
49
50         self.mlp.train()
51         for _ in range(self.epochs):
52             idx = torch.randint(n, (self.batch_pairs, 2))
53             i, j = idx[:, 0], idx[:, 1]
54
55             # pseudo-label  $\tilde{t}_{\eta} = t_{\tau} + \omega_{\tau} \cdot \Delta_{\eta}$  (clamped for BCE)
56             t_tau = torch.sigmoid((tau_t[i] - tau_t[j]) / self.kappa)
57             omega = t_tau * (1 - t_tau) / self.kappa
58             Delta = (phi_t[i] - tau_t[i]) - (phi_t[j] - tau_t[j])
59             t_tilde = torch.clamp(t_tau + omega * Delta, 0.0, 1.0)
60
61             # Predicted pairwise probability  $p_g(X, X') = \sigma(g(X) - g(X'))$ 
62             p_g = torch.sigmoid(self.mlp(X_t[i]) - self.mlp(X_t[j]))
63
64             # Back propagation loss
65             loss = F.binary_cross_entropy(p_g, t_tilde)
66             opt.zero_grad()
67             loss.backward()
68             opt.step()
69
70         self.mlp.eval()
71         return self
```

Ranking Causal Effects | Training the Orthogonal Ranker

Inference

1. Collect features $\{x_i : i = 1, 2, \dots, m\}$
2. Call model $\hat{r}_i = \hat{g}(x_i) \quad i = 1, 2, \dots, m$
3. Rank based on $\{\hat{r}_i\}_{i=1}^m$

Table 4. **Semi-synthetic benchmarks:** Test AUOC (mean $\pm$ std dev over five seeds) with training size $n = 1,000$. *Higher is better.* The oracle row reports AUOC obtained by ranking the test set using the true treatment effects. Best mean is shown in **bold**.

Method	MOVIELENS	MIMIC-III	CPS
oracle	1.39	1.22	1.01
T-learner	1.31 $\pm$ 0.03	1.12 $\pm$ 0.05	0.87 $\pm$ 0.08
DR-learner	1.34 $\pm$ 0.02	1.16 $\pm$ 0.02	0.92 $\pm$ 0.02
Plug-in ranker	1.30 $\pm$ 0.03	1.11 $\pm$ 0.05	0.87 $\pm$ 0.08
Rank-learner (<i>ours</i>)	1.35 $\pm$ 0.01	1.18 $\pm$ 0.02	0.95 $\pm$ 0.01

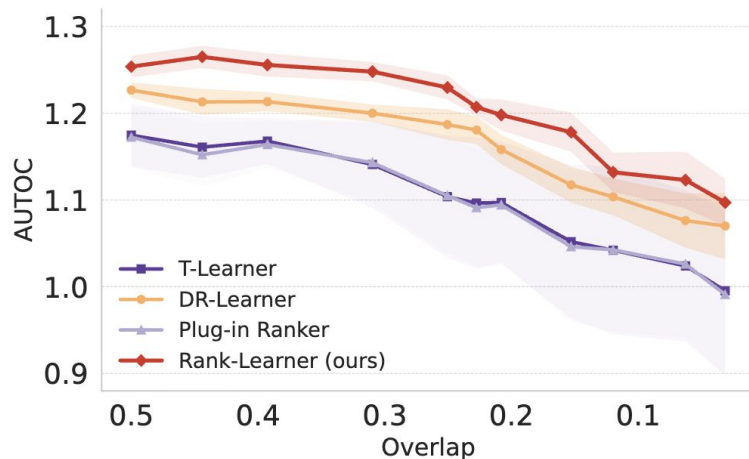


Figure 4. **Synthetic benchmark (overlap sensitivity).** Test AUOC (mean $\pm$ s.e. over five seeds) as a function of overlap (decreasing left to right) for $n = 500$. *Higher is better.* Overlap is varied by changing treatment assignment, remaining components of the synthetic data-generating process are kept fixed.

Conclusion

Conclusion | Discussion

Takeaways

- CausalML in industry is a mixture of engineering feasibility & causal inference best practices
- Metalearners allow us to leverage the best ideas from the ML community flexibly
- Second stage inference approaches that are (i) robust and (ii) problem specific perform the best

Questions?

Contact Information

- [Linkedin](#)
- Email
 - benjamin.draves@gmail.com
 - ben.draves@uber.com

Modern Improvements | Identifying Robustness

What is “Robustness”

- We are interested in estimate the CATE
$$\tau(x) = \mathbb{E}[Y(1)|X = x] - \mathbb{E}[Y(0)|X = x]$$
but oftentimes suffer “plug-in” bias
- Can we devise “second stage” estimation procedures which estimate $\tau(x)$ that are robust to certain errors
- What type of errors do we want to be robust to?

Neyman Orthogonality

If the nuisance estimates are biased, overfit, or converging slowly, we want $\hat{\tau}(x_i)$ to still be consistent and converge quickly.

Double Robustness

If only one of our nuisance estimators (propensity scores or surface function) are correctly specified, then our estimator $\hat{\tau}(x_i)$ is still unbiased and consistent.

Modern Improvements | R Learner

R Learner

- Fit **two** ML models to predict nuisance

$$m(x) = \mathbb{E}[Y|X = x] \quad e(x) = \mathbb{P}(Y = 1|X = x)$$

- Residualize outcome and treatment

$$\tilde{Y} = Y - \hat{m}(X) \quad \tilde{T} = T - \hat{e}(X)$$

- Estimate $\hat{\tau}(x_i)$ by minimizing

$$\hat{\tau}(\cdot) = \arg \min_{\tau} \sum_{i=1}^N \{(\tilde{y}_i - \tilde{t}_i \tau(x_i))^2\}$$

```
metalearners > r_learner.py
1  """
2  R-Learner: residualises both outcome and treatment before fitting CATE.
3  Reference: Nie & Wager 2017 - https://arxiv.org/abs/1712.04912
4  """
5
6  import numpy as np
7  from BaseLearner import BaseLearner
8
9  class RLearner:
10     def __init__(self,
11                 outcome_learner: BaseLearner = None,
12                 propensity_learner: BaseLearner = None,
13                 effect_learner: BaseLearner = None):
14         self.outcome_model = outcome_learner
15         self.propensity_model = propensity_learner
16         self.effect_model = effect_learner
17
18     def fit(self, X, T, Y):
19         # Stage 1: cross-fit m(X) and e(X) to avoid overfitting
20         m_hat = self._cross_fit_outcome(X, T, Y)
21         e_hat = self._cross_fit_propensity(X, T)
22
23         # Stage 2: compute residuals
24         Y_res = Y - m_hat          # outcome residual
25         T_res = T - e_hat          # treatment residual
26
27         # Stage 3: weighted regression tau(X) on pseudo-targets
28         # target = Y_res / T_res, weight = T_res^2
29         weights = T_res ** 2
30         targets = np.where(np.abs(T_res) > 1e-6, Y_res / T_res, 0.0)
31
32         self.effect_model.fit(X, targets, sample_weight=weights)
33         return self
34
35     def predict_cate(self, X):
36         return self.effect_model.predict(X)
37
```