

Using Causal ML to select articles for promotion on digital channels

Joel Persson, Research Associate @ ETH Zürich (Switzerland)
Dr. Cristina Kadar, Data Science & Machine Learning PO @ NZZ (Switzerland)

WAN-IFRA Data Science Expert Group
Webinar, November 2023

Why do we need causal modelling?

Digitalization has enabled data-driven decision making in the media industry

In the past decades we have witnessed vast improvements in data collection, storage, and processing.

This enabled common decision making use cases to be data-driven:

- **Which** content should be promoted on the frontpage, social media channels, etc.?
- **When** should a paywall be shown?
- **How much** should a certain subscription be priced?

Machine learning is increasingly used, yet it is not designed for decision making

Recent access to more data and computational power has enabled algorithms based on machine learning (ML)

Problem:

- ML learns a model from many past examples (fits a function) to predict outcomes from correlations. In standard application, it makes **predictions given the status quo** (e.g., forecasting).
- Decision-making requires predicting outcomes **caused by different actions** (“treatment” vs. “control”). Thus, we are interested in **predicting changes in outcomes** (i.e., the causal **treatment effect** or **uplift**).

Ex: In targeting users to reduce churn, we **should not** target those with highest predicted likelihood of churn, but those with the **largest reduction** in likelihood to churn if they receive an offer vs. if not.

=> For data-driven decision-making, ML needs causal inference (the study of designs and estimation methods needed for drawing causal conclusions from data).

Designs for getting data: RCTs and Observational Studies

Randomized controlled trials (A/B tests) are the “gold standard” for estimating treatment effects.

Pros:

- Randomization ensures causality
- No modelling required
- Easy to implement

Cons:

- Provides an average treatment effect
- “Costly” (wrt profits, logistics, ethics, etc.)
- Delayed learning (must wait until the end)

Observational studies (historical data with non-random treatment) sometimes the feasible approach.

Pros:

- No “cost” from randomization
- Often much larger sample size
- No delay; can use existing data

Cons:

- Req’s controlling for all confounders
- Req’s specifying a regression model
- => Prone to bias from misspecification, omitted variables, and estimation errors

Statistical methods for estimating treatment effects

Historically, methods for estimating treatment effects relied on parametric statistical models.

Ex: The **average treatment effect** is β in a linear regression model

- **RCT:** Outcome = intercept + β × treatment + controls + error
- **Obs study:** Outcome = intercept + β × treatment + confounders + controls + error

Ex: The **heterogeneous treatment effect** function is $\beta \times \text{moderators}$ in a linear regression model

- **RCT:** Outcome = intercept + $\beta \times \text{moderators}$ × treatment + controls + error
- **Obs study:** Outcome = intercept + $\beta \times \text{moderators}$ × treatment + confounders + controls + error

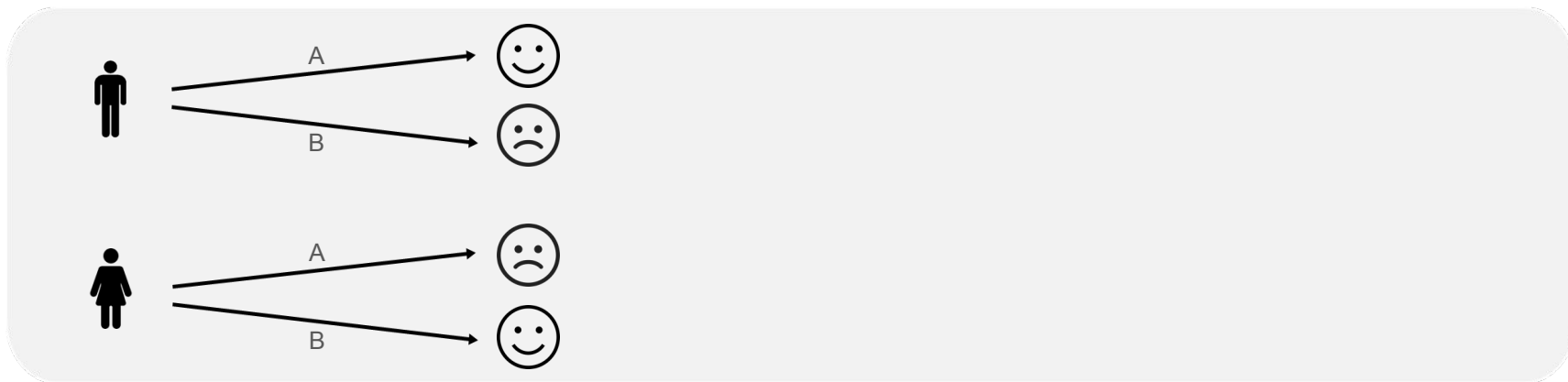
Challenge:

- Unbiased estimates (i.e., get true estimate on average) req's the **specified model to be correct**
- Valid inference (e.g., p-values, conf intervals) req's specifying the model **before seeing the data**

Latest estimation methods: ML

Lately, ML methods have been adapted to causal inference to address these challenges.

Causal ML improves estimation of heterogeneous treatment effects by framing it as a prediction problem.



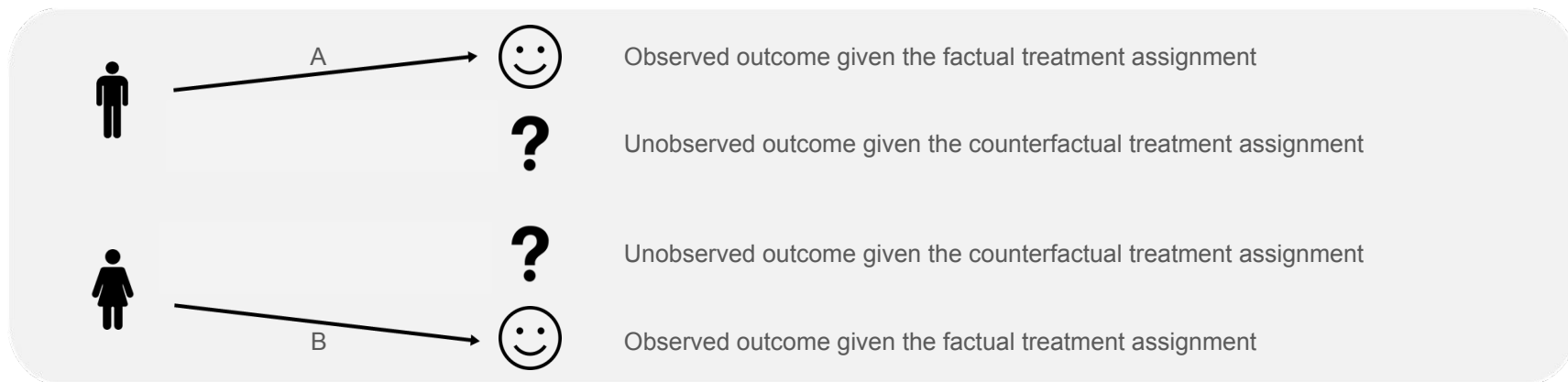
Benefits of Causal ML

- **Data-driven:** Learns the model that best fit the data => less prone to bias from incorrect model
- **Regularization:** Automatically selects which variables are moderators, controls, and confounders
- **Prediction performance:** Often better at out-of-sample prediction under treatment or control

Latest estimation methods: ML

Lately, ML methods have been adapted to causal inference to address these challenges.

Causal ML improves estimation of heterogeneous treatment effects by framing it as a prediction problem.



Benefits of Causal ML

- **Data-driven:** Learns the model that best fit the data => less prone to bias from incorrect model
- **Regularization:** Automatically selects which variables are moderators, controls, and confounders
- **Prediction performance:** Often better at out-of-sample prediction under treatment or control

Latest estimation methods: ML

Lately, ML methods have been adapted to causal inference to address these challenges.

Causal ML improves estimation of heterogeneous treatment effects by framing it as a prediction problem.



Benefits of Causal ML

- **Data-driven:** Learns the model that best fit the data => less prone to bias from incorrect model
- **Regularization:** Automatically selects which variables are moderators, controls, and confounders
- **Prediction performance:** Often better at out-of-sample prediction under treatment or control

Latest estimation methods: ML

Lately, ML methods have been adapted to causal inference to address these challenges.

Causal ML improves estimation of heterogeneous treatment effects by framing it as a prediction problem.



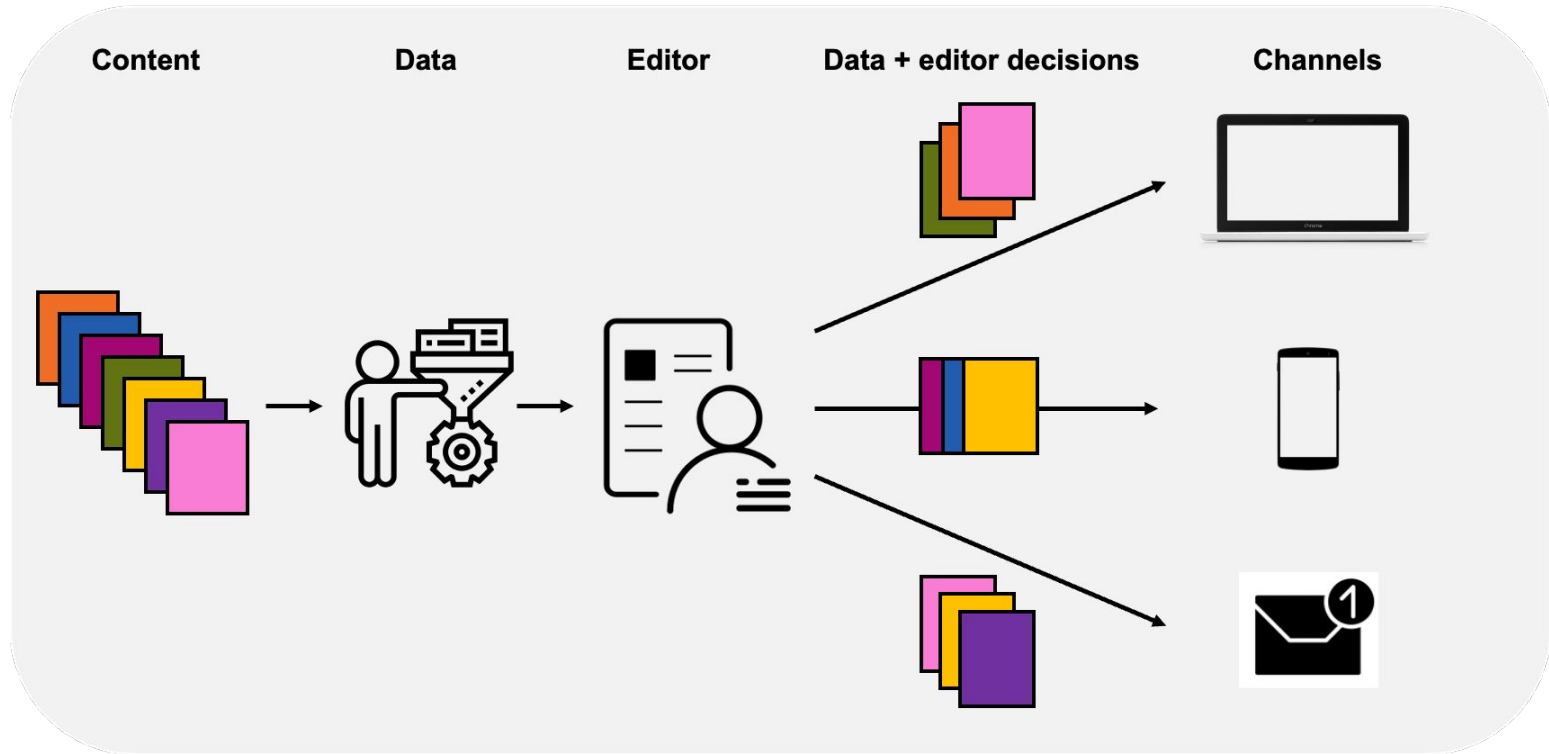
Benefits of Causal ML

- **Data-driven:** Learns the model that best fit the data => less prone to bias from incorrect model
- **Regularization:** Automatically selects which variables are moderators, controls, and confounders
- **Prediction performance:** Often better at out-of-sample prediction under treatment or control

NZZ's use case:

Editorial content promotion on the digital channels,
such as the CH and DE frontpages

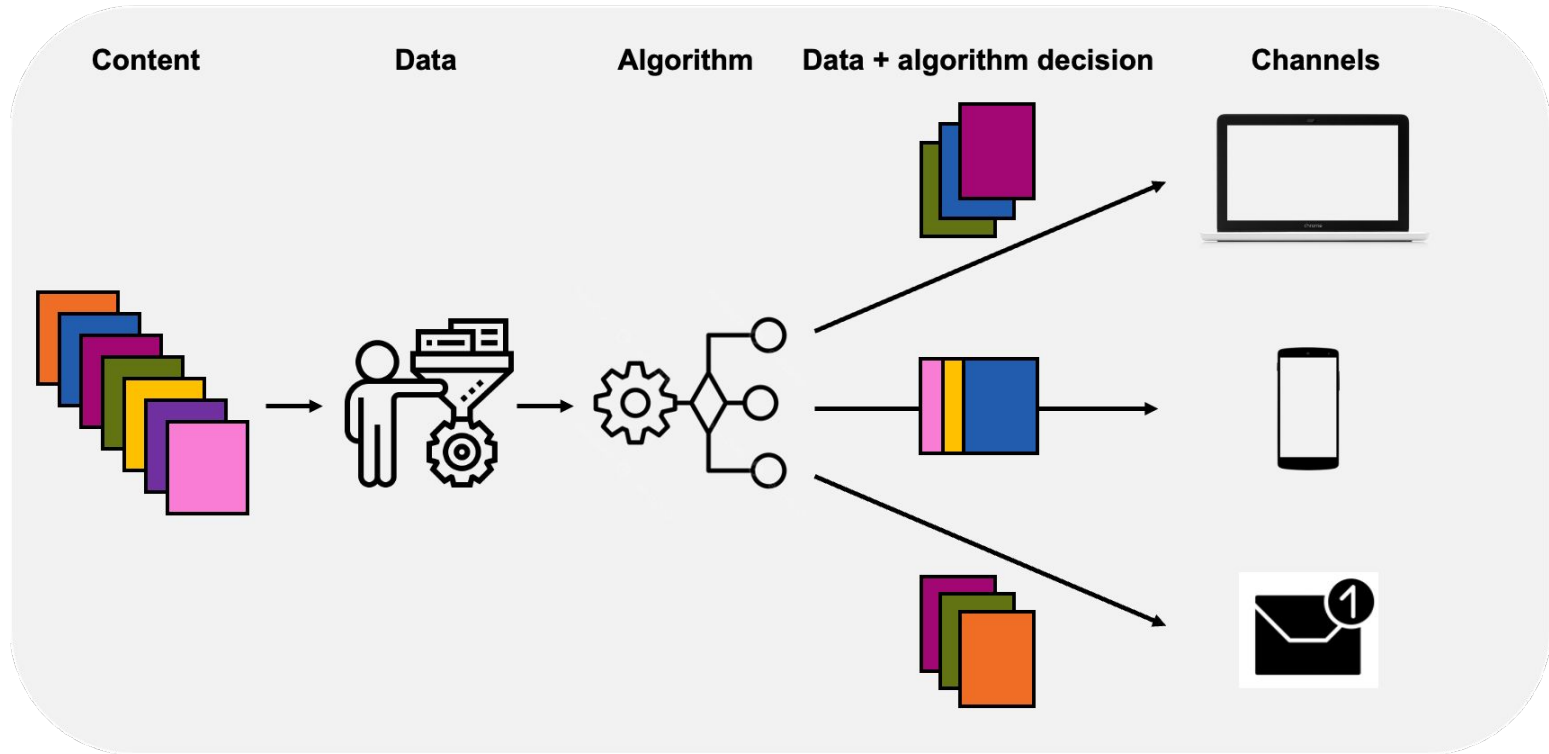
Current setup



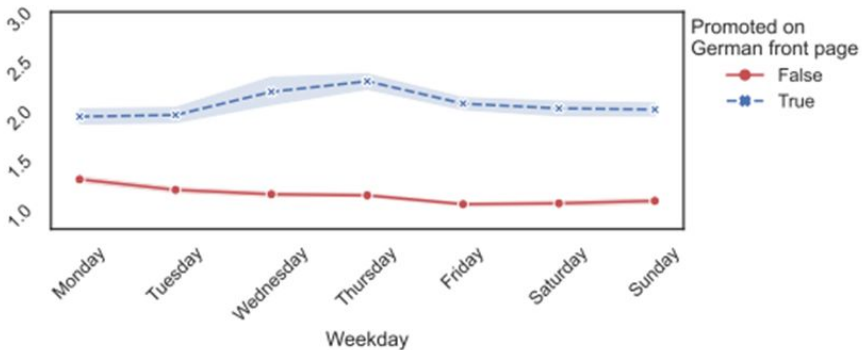
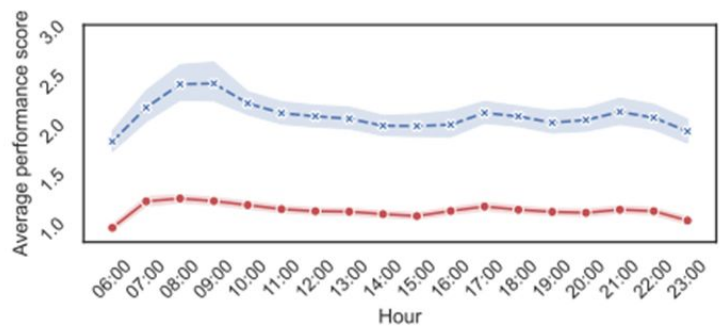
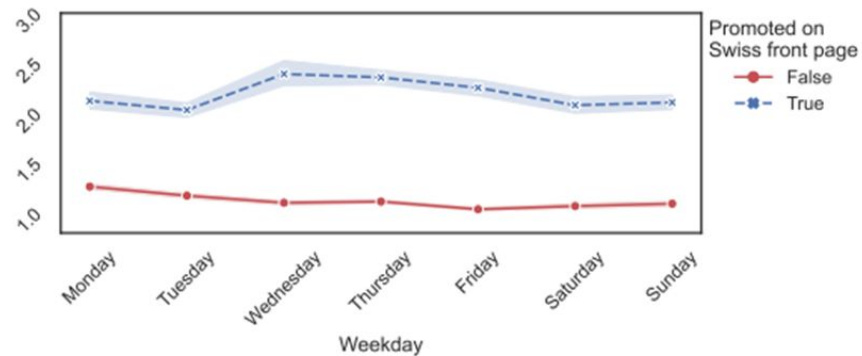
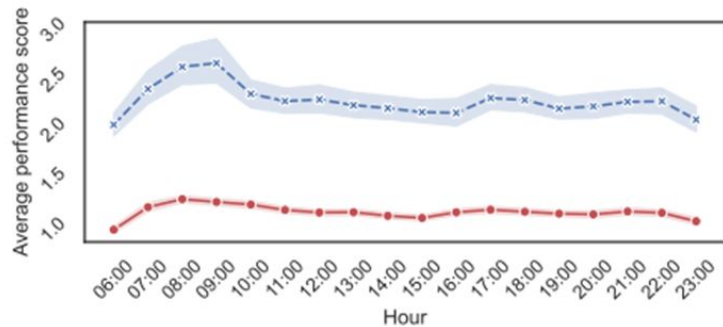
Current system used by the editors to decide what to promote when and where



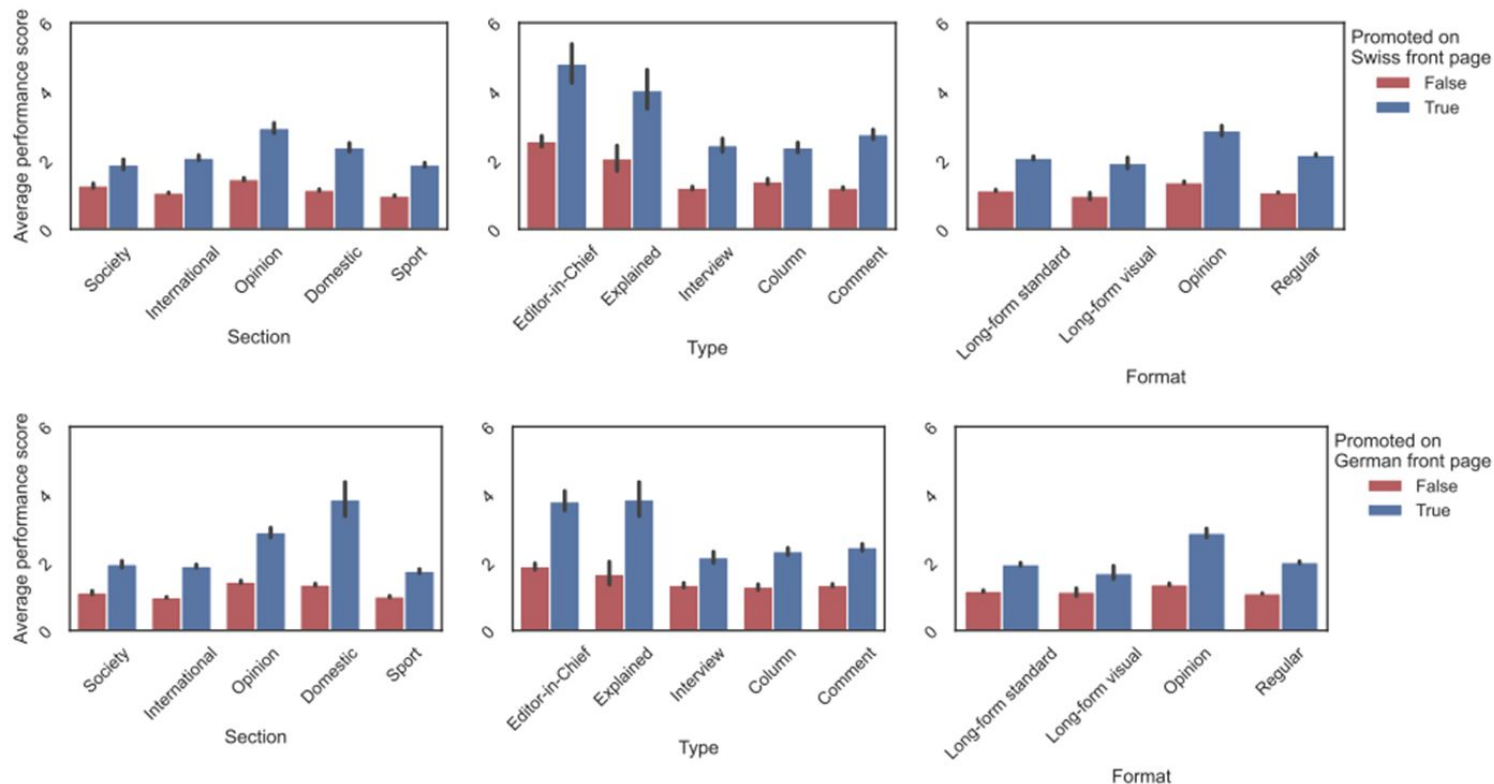
Envisioned setup (e.g. with human in the loop)



Descriptive statistics: promoted articles have higher performance score than non-promoted articles on the CH and DE frontpages



Descriptive statistics: promoted articles have higher performance score than non-promoted articles across sections, types, and formats



Causal inference approach

To learn what content is best to promote where and when, we must learn the heterogeneous uplift of promoting content. This implicitly requires thinking about counterfactuals.

*What are the potential **outcomes** (performance scores) for the different **choices** (promoted Y/N)?*

*Given the potential outcomes from the different choices, take the action with the expected greatest **uplift** and promote those articles that would receive the greatest benefit from being featured on the frontpage.*

Step 1: Estimation of the nuisance models on training data

Outcome model:

Gradient Boosting Regressor

$$\text{performance_hat}_{a,t} = f(\text{promotion}_{a,t}, \text{metrics}_{a,t-1}, \text{article_covariates}_a, \text{time_covariates}_t)$$

Propensity score model:

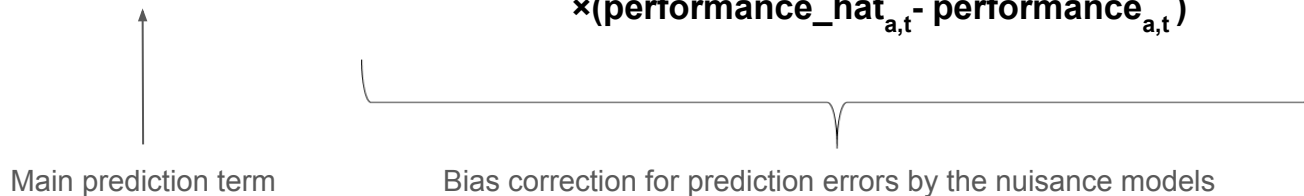
Gradient Boosting Classifier

$$\text{promotion_hat}_{a,t} = g(\text{metrics}_{a,t-1}, \text{article_covariates}_a, \text{time_covariates}_t)$$

Step 2: Estimation of the uplift model on the training data

We combine the predictions from the outcome and propensity score models to a “doubly robust (DR) score” of the predicted performance over the coming hour:

$$\text{DR_score}_{a,t} = \text{performance_hat}_{a,t} + (1\{\text{promotion decision in the data} = a\}_{.t} / \text{promotion_hat}_{a,t}) \times (\text{performance_hat}_{a,t} - \text{performance}_{a,t})$$



We get the predicted uplift by taking the difference in DR_score given promotion or not:

$$\text{Predicted uplift}_t = \text{DR_score}_{a = \text{promoted},t} - \text{DR_score}_{a = \text{not promoted},t}$$

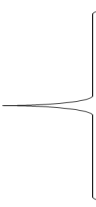
Finally, we fit an uplift model so that we can predict uplifts for content in the test data:

$$\text{Predicted uplift}_t = h(\text{metrics}_{a,t-1}, \text{article_covariates}_a, \text{time_covariates}_t)$$

ForestDRLearner

Step 3: Estimate optimal policy via simulation on the test data

Recommend top-N articles with the highest predicted uplift as per the model in step 2 for promotion on the frontpage in the next hour.

Optimal promotion policy_{a,t} =  **IF predicted uplift rank_{a,t} > N-1, Promote**
ELSE, Do not promote

Compare this optimal policy to the behavior policy (current practice of the editors).

In the simulation, our method makes promotion decisions that outperform the editors' historical decisions at all points in time

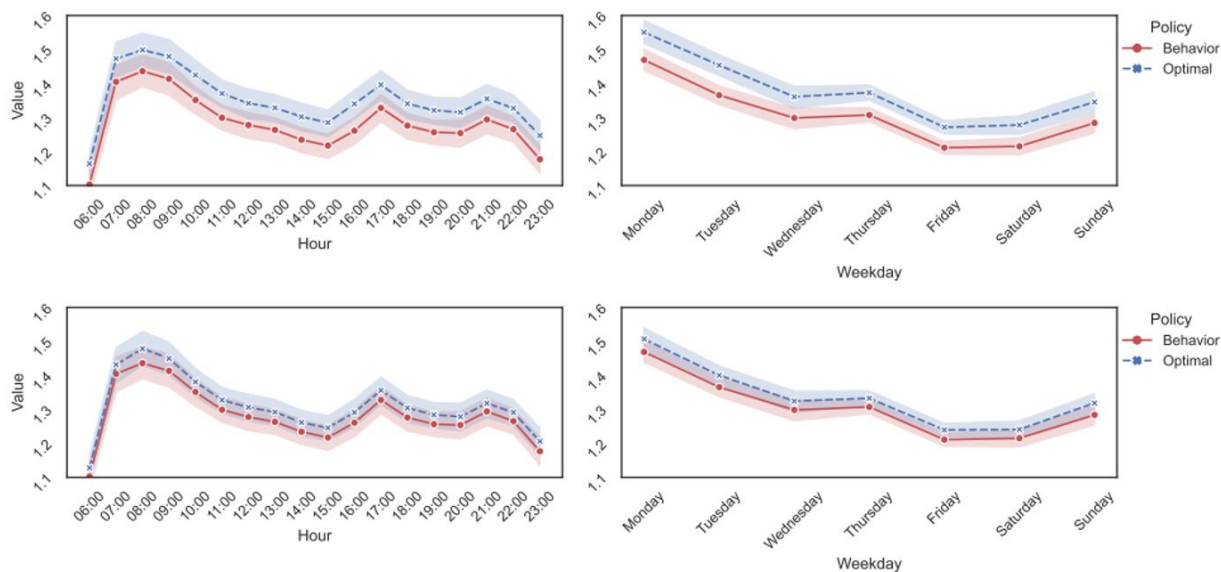
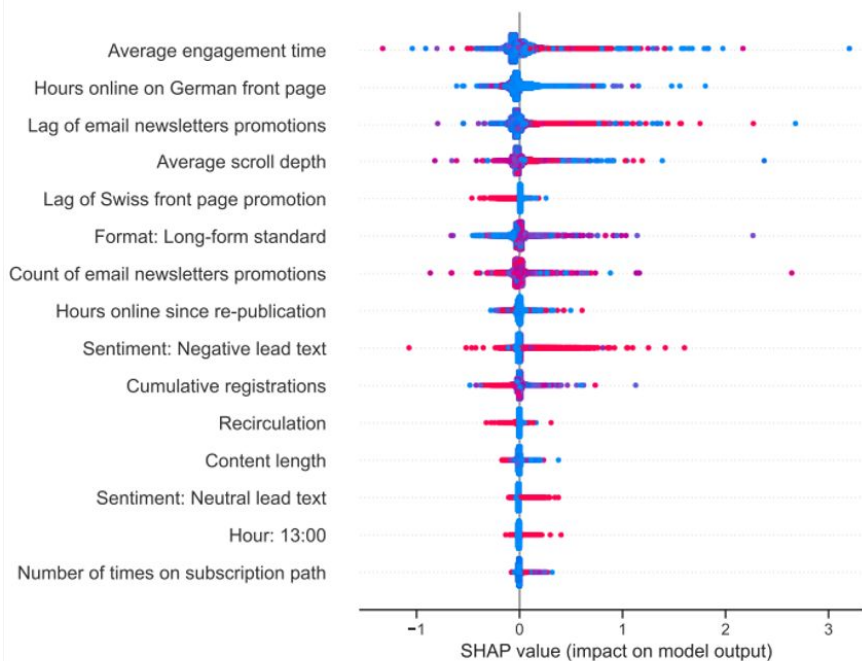
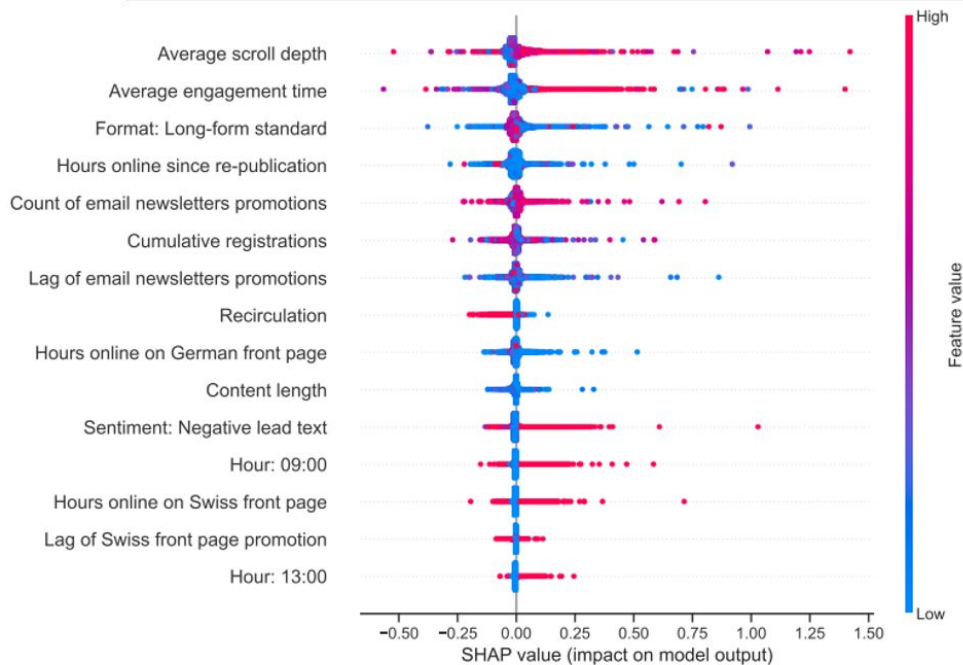


Figure 3.4: Value of the behavior policy (current practice) and our optimal policy by hour and day of the week. Shaded regions are 95% confidence intervals across 1000 bootstrap runs. Top: Swiss front page. Bottom: German front page.

Which content features explain the heterogeneous treatment effect of the promotion and thus the optimal promotion policy?



(a) Swiss front page



(b) German front page

Thank you for your attention!