

**МЕТОДЫ ОЦЕНКИ ПРИЧИННО-СЛЕДСТВЕННЫХ ЭФФЕКТОВ
НА ОСНОВЕ МАШИННОГО ОБУЧЕНИЯ
METHODS FOR ASSESSING CAUSAL EFFECTS BASED
ON MACHINE LEARNING**

Михайлов А.М., студент

Щербакова И. В., к.т.н., доцент

**ФГБОУ ВО «Воронежский государственный лесотехнический университет
имени Г.Ф. Морозова»**

г. Воронеж, Россия

tg3381234@gmail.com

Mikhailov A. M., student

Shcherbakova I.V., CSc (Engineering), Associate Professor

**FSBEI HE «Voronezh State University of Forestry and Technologies
named after G.F. Morozov»**

Voronezh, Russian Federation

Аннотация: В статье проведен сравнительный анализ современных методов оценки причинно-следственных эффектов на базе машинного обучения: T-Learner, метода Гельмана и метод трансформации класса, которые помогают оценивать гетерогенные эффекты воздействия в наблюдательных выборках. В основе полученных результатов лежит проведенное экспериментальное исследование на реальных данных.

Summary: The article provides a comparative analysis of modern methods for assessing cause-and-effect effects based on machine learning: T-Learner, the Gelman method and the class transformation method, which help to evaluate heterogeneous effects of exposure in observational samples. The results obtained are based on an experimental study based on real data.

Ключевые слова: машинное обучение, Uplift, метод Гельмана, T-Learner, метод трансформации класса, treatment, Uplift - моделирование.

Keywords: machine learning, Uplift, Gelman's method, T-Learner, class transformation method, treatment, Uplift - modeling.

Одним из современных подходов в машинном обучении и причинно-следственном анализе является uplift-моделирование (моделирование прироста), который применяется для оценки индивидуального эффекта воздействия (treatment effect) на целевую метрику.

Рассмотрим задачу: компания готова предоставить скидки потенциальным клиентам, но не всем, а лишь тем, у которых высокая вероятность оттока. Возможность выявлять таких клиентов позволила бы увеличить прибыль и снизить отток пользователей. Построим модель, которая бы прогнозировала вероятность ухода покупателя к конкурентам (рисунок 1).

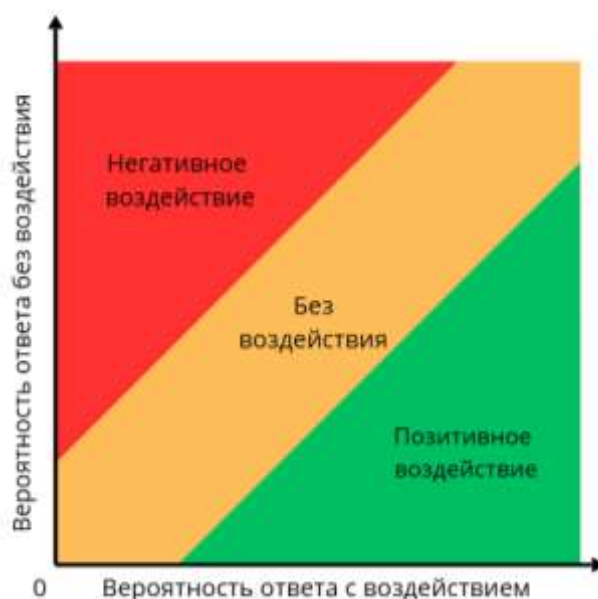


Рисунок 1 – Классификация реакций пользователя в зависимости от воздействия на него со стороны компании

Для решения этой задачи с использованием машинного обучения подходят несколько методов. Рассмотрим их применение на практике.

Метод двух моделей (англ. Two-Model Approach или T-Learner) предполагает независимую разработку двух моделей: первая модель – для контрольной группы (группа, которая не получила воздействие), вторая модель – для промо-группы (тестовая группа, для которой было применено воздействие). Значение Uplift для каждого объекта вычисляется, как разница между двумя моделями.

$$Uplift(x) = \hat{Y}_1(x) - \hat{Y}_0(x),$$

где $\hat{Y}_1(x)$ – значения, полученные для тестовой группы, $\hat{Y}_0(x)$ – значения, полученные для контрольной группы.

Преимущества данного метода:

1. Простота построения моделей, т.к. возможно использование любых ML-алгоритмов.

2. Каждая модель имеет свои настройки, специфичные для группы.

К недостаткам можно отнести:

1. Возникновение двойной ошибки моделирования, потребность в большем объеме данных.

2. Данный метод не учитывает зависимости между treatment'ом.

Рассмотри метод трансформации класса (ClassTransformation Method).

Введем целевую переменную:

$$Z_i = Y_i \cdot T_i + (1 - Y_i) \cdot (1 - T_i),$$

где Y_i – исходная целевая переменная (например, конверсия: 1 или 0), T_i – индикатор treatment'a (1 – тестовая группа, 0 – контрольная группа).

Ситуация $Z_i = 1$ возможна, если:

– клиент, который получил воздействие, совершил целевое действие ($Y_i = 1, T_i = 1$);

– клиент, который не получил воздействие, не совершил покупку ($Y_i = 0, T_i = 0$).

Ситуация $Z_i = 0$ возможна, если:

– клиент, который получил воздействие, не совершил покупку ($Y_i = 0, T_i = 1$);

– клиент, который не получил воздействие, совершил покупку. ($Y_i = 1, T_i = 0$).

После обучения uplift будем считать по формуле:

$$Uplift(x) = 2 \cdot P(Z = 1 | X = x) - 1.$$

Достоинствами данного метода является:

1. Простота модели.

2. Возможность учета взаимодействия признаков.

Недостатки метода:

1. Работает только для бинарных задач.

2. Не очень большая точность предсказаний на сложных данных.

Метод Гельмана (байесовский подход Гельмана) предполагает моделирование с применением алгоритма построения uplift-дерева, в котором разбиение узлов (splits) выбирается так, чтобы максимизировать разницу в uplift между дочерними узлами:

$$\Delta_{gain} = D_{after-split}(P^T, P^C) - D_{before-split}(P^T, P^C),$$

где $D_{after-split}$ и $D_{before-split}$ – меры различия между распределениями в тестовой (P^T) и контрольной (P^C) группах до и после разбиения.

Δ_{gain} показывает, насколько разбиение улучшает разделение эффектов treatment'a. На рисунке 2 приведен пример разбиения uplift-дерева в зависимости от эффекта treatment'a.

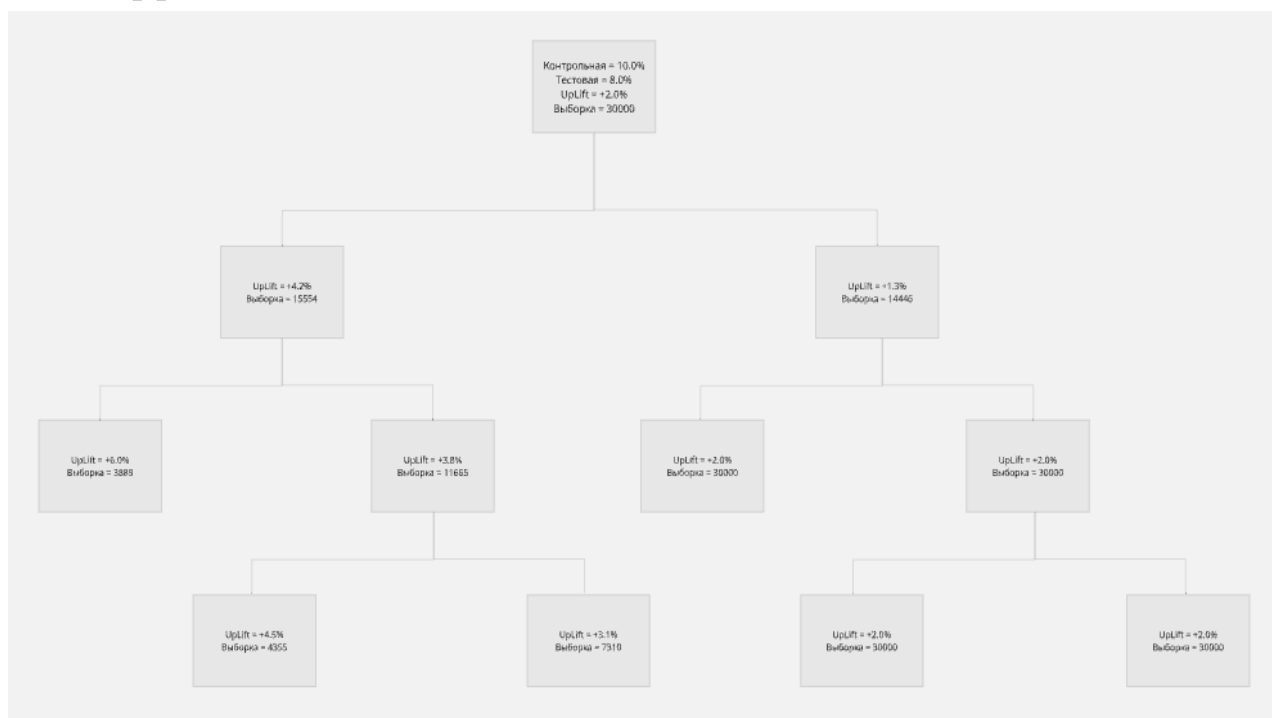


Рисунок 2 – Разбиение uplift-дерева в зависимости от эффекта treatment'a

Для разделения можно использовать KL-дивергенцию или другие методы. KL-дивергенция (расстояние Кульбака–Лейблера) используется в uplift-деревьях как один из возможных критериев разделения (splitting criteria). Она помогает выбрать такое разбиение узла, которое максимально увеличивает неоднородность эффекта treatment'a между подгруппами (левой и правой ветвями дерева).

$$KL(P:Q) = \sum_{k=Left,Right} p_k \log \frac{p_k}{q_k},$$

где P – распределение конверсий в тестовой группе, Q – распределение в контрольной группе.

По данному критерию можно сделать следующий вывод: чем больше $KL(P:Q)$, тем сильнее различается эффект treatment'a в подгруппах.

Алгоритм выбирает разбиение, которое максимизирует KL-дивергенцию между P и Q (т.е. делает эффект treatment'a максимально неоднородным).

Для экспериментального исследования выберем метод Гельмана и метод T-Learner, так целевые значения используемого дата-сета не бинарные. В дата-сете содержится информация о 64000 клиентах, которые совершили покупку перед email-рассылкой в последний раз в течение 12 месяцев (рисунок 3).

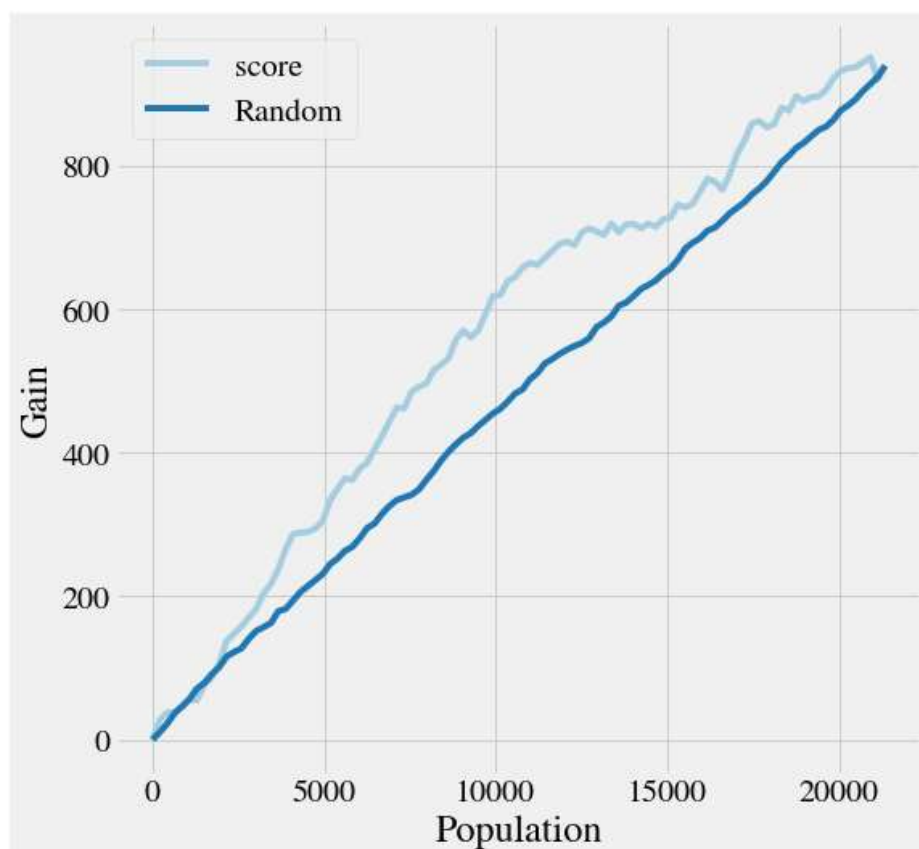


Рисунок 5 – Результаты после обучения по методу T-Learner

Линия score (Uplift Model) показывает, как растет полезный эффект (Gain) при таргетинге на клиентов, отобранных моделью. Чем выше кривая, тем лучше модель выделяет «чувствительных» к воздействию пользователей.

Линия (Random) – это ожидаемый эффект при случайном выборе клиентов для воздействия. Чем больше разница между uplift-моделью и Random, тем выше качество модели.

После реализации uplift-модели, построенной методом Гельмана, были получены следующие результаты (рисунок 6).

Модель показала лучшие результаты относительно метода T-Learner. Почти весь gain достигается на малой популяции (около 5000 клиентов). Это говорит о точном ранжировании: лучшие клиенты сосредоточены в начале.

На основании анализа графиков можно сделать следующие выводы: uplift-деревья демонстрируют быстрый рост gain на малой популяции, что указывает на способность точно выделять наиболее «отзывчивых» пользователей. Пиковый эффект (~800) достигается раньше, чем у T-Learner. Исследования показывают, что метод Гельмана оптимален для задач с ярко выраженной гетерогенностью эффекта, где небольшая группа клиентов дает основной прирост.

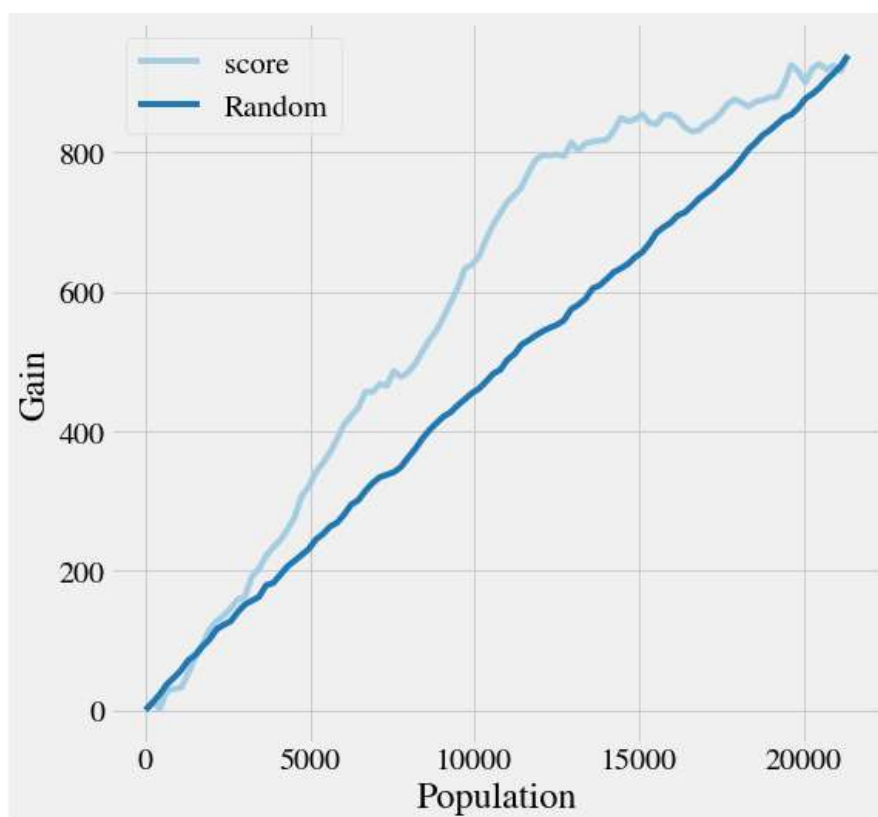


Рисунок 6 – Результаты после обучения по методу Гельмана

У T-learner модели gain растет более плавно, требуя охвата большей популяции (10,000 клиентов для 80% эффекта). Однако данный метод универсален и прост в реализации.

В заключение можно сказать, что выбор метода зависит от бизнес-целей. Если ресурсы ограничены (например, бюджет на скидки), uplift-деревья предпочтительнее, они быстрее выявляют «золотую» аудиторию. Если важна интерпретируемость и простота, можно применять метод T-Learner. Uplift-деревья лучше подходят для точечного таргетинга, а T-Learner – для базового анализа.

Список литературы

1. Pearl, J. Causal Inference in Statistics: A Primer / J. Pearl, M. Glymour, N.P. Jewell. – Wiley, 2016. – 160 p.
2. Imbens, G.W. Causal Inference for Statistics, Social, and Biomedical Sciences / G.W. Imbens, D.B. Rubin. – Cambridge University Press, 2015. – 640 p.
3. Bayesian Data Analysis / A. Gelman, J.B. Carlin, H.S. Stern [et al.]. – 3rd ed. – Chapman & Hall/CRC, 2013. – 675 p.
4. Gutierrez, P. Uplift Modeling: From Causal Inference to Personalization. – URL: <https://arxiv.org/abs/1705.08492> (дата обращения: 01.04.2025).
5. CausalML Documentation. – URL: <https://causalml.readthedocs.io/> (дата обращения: 01.04.2025).

6. Scikit-uplift Library. – URL: <https://scikit-uplift.readthedocs.io/> (дата обращения: 01.04.2025).

References

1. Pearl, J. Causal Inference in Statistics: A Primer / J. Pearl, M. Glamour, N.P. Jewell. – Wiley, 2016. – 160 p.

2. Imbens, G.W. Causal Inference for Statistics, Social, and Biomedical Sciences / G.W. Imbens, D.B. Rubin. – Cambridge University Press, 2015. – 640 p.

3. Bayesian Data Analysis / A. Gelman, J.B. Karlin, H.S. Stern [et al.]. – 3rd ed. – Chapman & Hall/CRC, 2013. – 675 p.

4. Gutierrez, P. Uplift Modeling: From Causal Inference to Personalization. – URL: <https://arxiv.org/abs/1705.08492> (accessed: 01.04.2025).

5. CausalML Documentation. – URL: <https://causalml.readthedocs.io/> (accessed: 01.04.2025).

6. Scikit-uplift Library. – URL: <https://scikit-uplift.readthedocs.io/> (accessed: 01.04.2025).