

Demystifying TMLE: A Practical Approach for Clinical Data Scientists

Joshua Enxing, Phastar, Boston, United States

ABSTRACT

Targeted Maximum Likelihood Estimation (TMLE) is transforming causal inference in clinical data science by providing reliable effect estimates from complex observational and trial data. Traditional methods can produce biased or imprecise results, particularly when study designs are imperfect or confounding exists. TMLE addresses these challenges through a doubly-robust framework: if either the outcome model or treatment model is correctly specified, the estimator remains consistent, and if both are correct, it achieves semiparametric efficiency. Despite theoretical complexity, TMLE can be implemented practically using accessible tools. In this talk, we break down the TMLE framework for a clinical data audience, showcase studies where it improved precision and reliability over classical methods, and provide guidance on best practices for applying TMLE to real-world datasets. Attendees will gain actionable insights for using TMLE in clinical research, including strategies to reduce bias, maximize efficiency, and enhance reproducibility.

INTRODUCTION

Clinical trials are designed to answer intervention-based questions: *What would happen to outcomes if we implemented treatment strategy A rather than strategy B in the target population?* Randomization balances baseline confounders on average, but real trial analyses still face practical complexities: missing outcomes, intercurrent events, post-randomization biomarkers, non-adherence, and time-varying treatment decisions.

In many clinical settings, the default workflow is to (i) fit a regression model (logistic or Cox), (ii) include baseline and sometimes post-baseline covariates, and (iii) interpret the treatment coefficient as the treatment effect. This approach can fail in clinically relevant ways:

- **Estimand mismatch due to non-collapsibility:** logistic/Cox coefficients are conditional parameters and generally do not equal marginal causal effects even in randomized trials.
- **Post-treatment conditioning:** adjusting for biomarkers affected by treatment can block causal pathways and induce bias.
- **Informative missingness/censoring:** complete-case or naive modeling is biased when observation depends on covariates and treatment.
- **Complex strategies:** adaptive treatment rules are not representable by a single regression coefficient.

TMLE was developed to estimate marginal intervention-based estimands with valid uncertainty quantification, while allowing flexible nuisance modeling and providing robustness to model misspecification.

HOW TO READ THIS PAPER

This paper is structured to build intuition first and then add technical detail:

1. **Motivation and estimands:** define what we want to estimate in clinical trials.
2. **Why Classical Adjusted Analyses Can Fail:** collapsibility, post-treatment conditioning, and missingness.
3. **TMLE theory:** EIF, targeting, double robustness, cross-fitting, and inference.
4. **Worked Example 1: Baseline Adjustment in an RCT:** what changes under logistic regression vs TMLE.
5. **Worked Example 2: Longitudinal Strategy with Censoring:** where classical estimation fails.
6. **Implementation in R:** R packages, code details.
7. **Diagnostics for TMLE: Theory and Implementation:** address positivity, stability, transparency, and sensitivity.

MOTIVATION AND ESTIMANDS POINT-TREATMENT SETTING

Let $O = (W, A, Y) \sim P_0$, where W are baseline covariates, $A \in \{0, 1\}$ is treatment, and Y is outcome. Let $Y^{(a)}$ denote the potential outcome under treatment a .

We use standard identification assumptions:

1. **Consistency:** if $A = a$ then $Y = Y^{(a)}$.
2. **Exchangeability:** $Y^{(a)} \perp A \mid W$.
3. **Positivity:** $0 < P(A = 1 \mid W = w) < 1$ for relevant w .

Under these assumptions, the ATE (average treatment effect) is identified using counterfactuals:

$$\Psi(P_0) = \mathbb{E}[\mathbb{E}(Y \mid A = 1, W) - \mathbb{E}(Y \mid A = 0, W)].$$

LONGITUDINAL SETTING AND DYNAMIC STRATEGIES

Many clinical trial estimands are strategy-based. Consider $O = (W, A_0, L_1, A_1, C, Y)$ where L_1 is a post-baseline biomarker, A_1 is a treatment adaptation, C indicates observation/censoring, and Y is the final outcome.

A dynamic regime is $\bar{a} = (a_0, a_1(L_1))$. The policy estimand is the marginal mean under this intervention:

$$\Psi_{\bar{a}}(P_0) = \mathbb{E}[Y^{\bar{a}}].$$

WHY CLASSICAL ADJUSTED ANALYSES CAN FAIL COLLAPSIBILITY AND NON-COLLAPSIBILITY

A central concept underlying covariate adjustment is *collapsibility*. Informally, an effect measure is collapsible if adjusting for additional covariates that are *not confounders* does not change the value of the effect estimate. Formally, let Z be a covariate that is not a confounder of the treatment–outcome relationship. An effect measure θ is *collapsible over Z* if

$$\theta\{P(Y \mid A)\} = \sum_z \theta\{P(Y \mid A, Z = z)\} P(Z = z).$$

Mean differences and risk differences are collapsible under standard conditions: stratifying on a purely prognostic Z does not change the estimand (though it may improve precision). In contrast, odds ratios and hazard ratios are *non-collapsible*. For these measures, even when Z is not a confounder (e.g., in a randomized trial), the conditional effect (e.g., $OR(Y, A | Z)$) typically differs from the marginal effect $OR(Y, A)$.

The key point is that non-collapsibility is a *mathematical* property of the effect measure (driven by the nonlinearity of the odds or hazard transformation), not a consequence of confounding. Therefore, adding prognostic covariates to a logistic or Cox model can change what is being estimated *even under randomization and correct model specification*. This is an *estimand change*, not a bias correction.

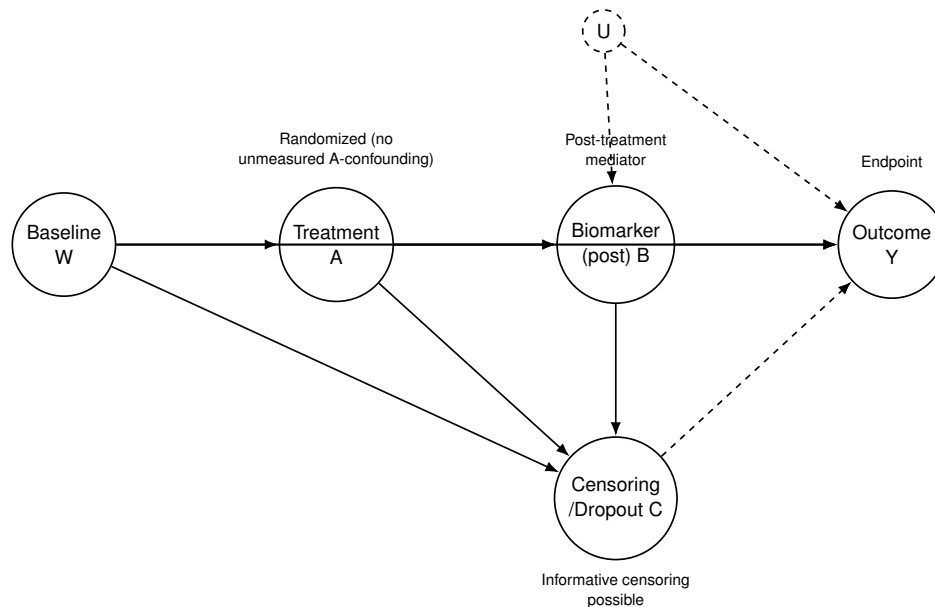
IMPLICATIONS FOR CLINICAL TRIALS

Clinical trial estimands are typically marginal and population-level, corresponding to questions such as: *What would the outcome be if all trial participants received treatment versus if none did?* Common examples include marginal risk differences, marginal mean differences, and marginal survival contrasts. By contrast, coefficients from adjusted logistic or Cox models correspond to *conditional* parameters (e.g., conditional odds ratios or hazard ratios given covariates). Because odds ratios and hazard ratios are non-collapsible, covariate adjustment changes the estimand, so two analysts using different covariate sets can legitimately obtain different treatment coefficients—each answering a different question.

In regulatory contexts (e.g., ICH E9(R1)), where the estimand must be clearly specified, interpreting an adjusted logistic or Cox coefficient as the “trial effect” can be misleading if the intended estimand is marginal. Adjustment may improve prediction and precision, but on non-collapsible scales it also changes the target.

POST-TREATMENT COVARIATES AND INFORMATIVE CENSORING

The challenges become more severe in longitudinal settings where covariates are affected by prior treatment. Let B denote a post-baseline biomarker influenced by A . Conditioning on B in an outcome model can: (i) block part of the causal effect of A operating through B (estimating a controlled direct effect rather than the total effect), and (ii) induce collider/selection bias when outcome observation or censoring depends on B and treatment. Consequently, complete-case regression or naive covariate adjustment can be biased even in randomized trials, while inverse-probability weighting can be unstable in finite samples.



NUMERIC TOY EXAMPLE: NON-COLLAPSIBILITY OF THE ODDS RATIO (EVEN UNDER RANDOMIZATION)

We construct an example where $A \perp Z$ (as in a randomized trial), Z is *purely prognostic* for Y , and the conditional odds ratio is constant across strata of Z , yet the marginal odds ratio differs from that common conditional odds ratio.

Let $P(Z = 1) = P(Z = 0) = 0.5$, and $P(A = 1 | Z) = P(A = 1) = 0.5$. Define the stratum-specific risks:

$$\begin{aligned} P(Y = 1 | A = 0, Z = 0) &= 0.10, & P(Y = 1 | A = 1, Z = 0) &= 0.1818, \\ P(Y = 1 | A = 0, Z = 1) &= 0.30, & P(Y = 1 | A = 1, Z = 1) &= 0.4615. \end{aligned}$$

Then within each stratum, the (conditional) odds ratio is

$$\text{OR}(Y, A | Z = z) = \frac{\frac{P(Y=1|A=1,Z=z)}{1-P(Y=1|A=1,Z=z)}}{\frac{P(Y=1|A=0,Z=z)}{1-P(Y=1|A=0,Z=z)}} = 2 \quad \text{for } z \in \{0, 1\}.$$

(Indeed, for $Z = 0$: odds under $A = 0$ is $0.10/0.90 = 0.1111$ and under $A = 1$ is $0.1818/0.8182 = 0.2222$, ratio = 2; for $Z = 1$: odds under $A = 0$ is $0.30/0.70 = 0.4286$ and under $A = 1$ is $0.4615/0.5385 = 0.8571$, ratio = 2.)

Now compute the *marginal* risks by averaging over Z :

$$P(Y = 1 | A = 0) = 0.5(0.10) + 0.5(0.30) = 0.20,$$

$$P(Y = 1 | A = 1) = 0.5(0.1818) + 0.5(0.4615) = 0.32165.$$

Thus the marginal odds ratio is

$$\text{OR}(Y, A) = \frac{\frac{0.32165}{1-0.32165}}{\frac{0.20}{1-0.20}} = \frac{0.32165/0.67835}{0.20/0.80} = \frac{0.4742}{0.25} \approx 1.90 \neq 2.$$

Therefore, even though treatment is randomized ($A \perp Z$) and the conditional odds ratio is constant across strata, the marginal odds ratio differs. This is exactly the sense in which the odds ratio is non-collapsible: adjustment for a purely prognostic Z changes the odds-ratio estimand.

Clinical interpretation. If a trial estimand is marginal (e.g., risk difference or marginal mean), reporting $\exp(\hat{\beta}_1)$ from an adjusted logistic regression targets a conditional odds ratio that is not equal to the marginal effect of treatment in the trial population. Methods that explicitly target marginal estimands—for example, TMLE targeting $E[Y^{(1)}] - E[Y^{(0)}]$ —can use W for efficiency gains while remaining aligned with the estimand.

TMLE THEORY

EFFICIENT INFLUENCE FUNCTION FOR THE AVERAGE TREATMENT EFFECT

Let $O = (W, A, Y)$ denote a single observed data unit, where W represents baseline covariates, $A \in \{0, 1\}$ is a binary treatment indicator, and Y is the outcome of interest. Let P_0 denote the true (unknown) data-generating distribution of O , and let $\mathbb{E}_{P_0}$ denote expectation under P_0 .

The estimand of interest is the *average treatment effect* (ATE), defined as

$$\Psi(P_0) = \mathbb{E}_{P_0} [Y^{(1)} - Y^{(0)}],$$

where $Y^{(a)}$ denotes the potential outcome under treatment level a . Under standard identification assumptions (consistency, conditional exchangeability given W , and positivity), the ATE can be expressed as a functional of the observed-data distribution:

$$\Psi(P_0) = \mathbb{E}_{P_0}[\mathbb{E}_{P_0}(Y \mid A = 1, W) - \mathbb{E}_{P_0}(Y \mid A = 0, W)].$$

Define the following *nuisance functions*:

- The *outcome regression*

$$Q_0(a, w) = \mathbb{E}_{P_0}(Y \mid A = a, W = w),$$

- The *treatment mechanism* (or propensity score)

$$g_0(a \mid w) = P_0(A = a \mid W = w).$$

In the nonparametric model, the ATE is a pathwise differentiable parameter, and its *efficient influence function* (EIF) is given by

$$D^*(O; P_0) = \frac{\mathbb{I}(A = 1)}{g_0(1 \mid W)} \{Y - Q_0(1, W)\} - \frac{\mathbb{I}(A = 0)}{g_0(0 \mid W)} \{Y - Q_0(0, W)\} + Q_0(1, W) - Q_0(0, W) - \Psi(P_0),$$

where $\mathbb{I}(\cdot)$ denotes the indicator function.

The EIF plays a central role in semiparametric estimation: it characterizes the lowest achievable asymptotic variance among all regular, asymptotically linear estimators of $\Psi(P_0)$ and forms the basis for both estimation and inference in targeted maximum likelihood estimation.

EIF EXPANSION AND DOUBLE ROBUSTNESS

A defining feature of targeted maximum likelihood estimation (TMLE) is that the estimator admits a first-order expansion in terms of the EIF. Let $\hat{Q}$ and $\hat{g}$ denote estimators of Q_0 and g_0 , respectively, and let $\mathbb{P}_n$ denote the empirical distribution of the observed data. Then the TMLE estimator $\hat{\Psi}$ satisfies

$$\hat{\Psi} - \Psi(P_0) = \mathbb{P}_n D^*(O; \hat{Q}, \hat{g}) + R_2(\hat{Q}, \hat{g}),$$

where $R_2(\hat{Q}, \hat{g})$ is a second-order remainder term.

For the ATE, this remainder satisfies

$$R_2(\hat{Q}, \hat{g}) = O_p\left(\|\hat{Q} - Q_0\| \|\hat{g} - g_0\|\right),$$

where $\|\cdot\|$ denotes an appropriate $L_2(P_0)$ norm and $O_p(\cdot)$ denotes stochastic order in probability.

This product-of-errors structure implies the celebrated *double robustness* property:

- If either the outcome regression $\hat{Q}$ or the treatment mechanism $\hat{g}$ is consistently estimated, the remainder term $R_2(\hat{Q}, \hat{g})$ converges to zero.
- If both nuisance estimators converge at rates faster than $n^{-1/4}$, then $\sqrt{n}R_2(\hat{Q}, \hat{g}) \rightarrow 0$, and the TMLE estimator is asymptotically linear with influence function $D^*(O; P_0)$.

As a consequence, TMLE achieves $\sqrt{n}$ -consistent, asymptotically normal inference even when flexible, data-adaptive methods (e.g., machine learning) are used to estimate Q_0 and g_0 , provided appropriate rate conditions

are met.

TARGETING AND PLUG-IN ESTIMATION

TMLE proceeds in two conceptually distinct stages. First, an *initial estimator* $\hat{Q}_0$ of the outcome regression is obtained by minimizing a loss function (e.g., Bernoulli log-likelihood for binary outcomes), possibly using flexible machine learning methods to optimize predictive performance. Second, this initial estimator is *targeted* toward the parameter of interest by updating it along a *least favorable submodel*.

For binary outcomes, a common fluctuation submodel is

$$\text{logit}\{\hat{Q}_\varepsilon(A, W)\} = \text{logit}\{\hat{Q}_0(A, W)\} + \varepsilon H(A, W),$$

where

$$H(A, W) = \frac{\mathbb{I}(A = 1)}{\hat{g}(1 | W)} - \frac{\mathbb{I}(A = 0)}{\hat{g}(0 | W)}$$

is the so-called *clever covariate*. The fluctuation parameter ε is estimated by fitting a univariate regression with offset $\text{logit}\{\hat{Q}_0(A, W)\}$.

This targeting step enforces the empirical EIF estimating equation

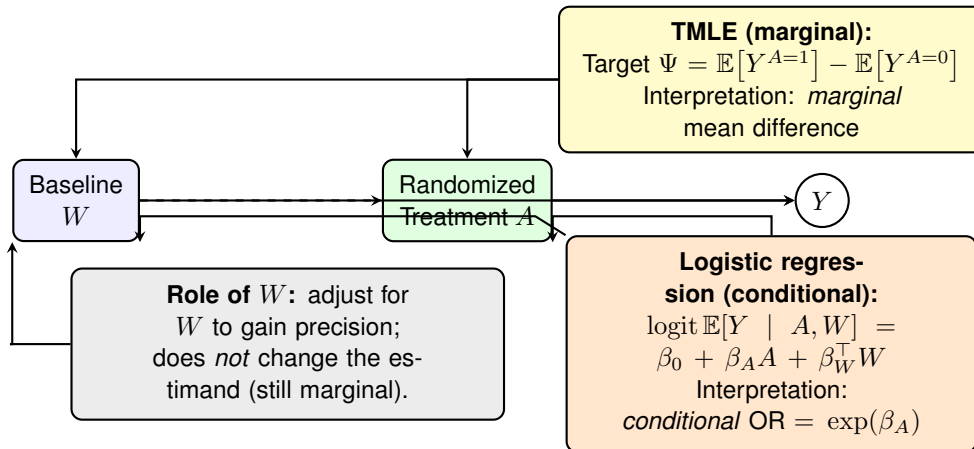
$$\mathbb{P}_n D^*(O; \hat{Q}^*, \hat{g}) \approx 0,$$

where $\hat{Q}^* = \hat{Q}_\varepsilon$ denotes the updated outcome regression. The final TMLE is then obtained as a *plug-in estimator*:

$$\hat{\Psi} = \frac{1}{n} \sum_{i=1}^n \{\hat{Q}^*(1, W_i) - \hat{Q}^*(0, W_i)\}.$$

Because TMLE is plug-in, it respects natural parameter bounds (e.g., predicted probabilities remain in $[0, 1]$), and because it solves the EIF equation, it achieves semiparametric efficiency under correct nuisance estimation. This combination of targeting, double robustness, and valid inference under data-adaptive estimation distinguishes TMLE from classical regression adjustment and from one-step or inverse-probability-weighted estimators.

WORKED EXAMPLE 1: BASELINE ADJUSTMENT IN AN RCT SETUP



The reported effect is $\exp(\beta_1)$.

$$\text{logit } P(Y = 1 \mid A, W) = \beta_0 + \beta_1 A + \beta^\top W.$$

The reported effect is $\exp(\beta_1)$.

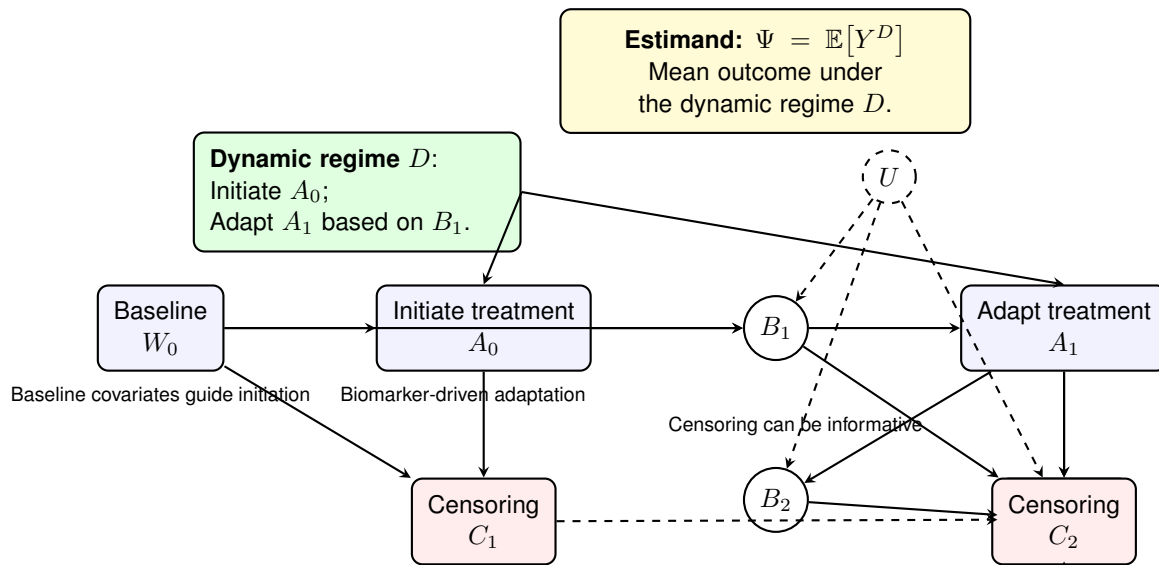
WHAT QUESTION DOES THIS ANSWER?

$\exp(\beta_1)$ is a conditional odds ratio. If the clinical trial estimand is marginal risk difference or marginal mean, this is an estimand mismatch. TMLE can target the marginal estimand directly while using W for precision gains.

WHERE TMLE FITS IN PRACTICE

In an RCT, g is known by design; however, estimating g can still stabilize finite-sample performance and integrates naturally with missingness modeling (e.g., inverse probability of censoring). TMLE provides a coherent framework that (i) targets the marginal estimand and (ii) yields EIF-based uncertainty.

WORKED EXAMPLE 2 (COMPLEX): LONGITUDINAL STRATEGY WITH CENSORING STRATEGY ESTIMAND



Consider initiating treatment ($A_0 = 1$) and intensifying if biomarker worsens ($A_1 = \mathbb{I}\{B_1 > c\}$). The estimand is $\Psi_{\bar{a}} = \mathbb{E}[Y^{(1, \mathbb{I}\{B_1 > c\})}]$.

WHY CLASSICAL REGRESSION FAILS

A complete-case regression of Y on (A_0, A_1, W, B_1) conditions on a post-treatment variable and on selection ($C = 1$). Even if predictive, it does not identify the marginal strategy estimand.

WHY MSM MAY BE UNSTABLE

IPW-based MSMs can be approximately unbiased but are sensitive to weight instability and model misspecification (practical positivity violations). This instability is common in adaptive strategies where treatment probabilities become extreme for certain biomarker strata.

WHY LONGITUDINAL TMLE WORKS

Longitudinal TMLE combines sequential regression (g-formula structure) and targeting steps that incorporate treatment and censoring mechanisms over time, yielding EIF-based inference for the strategy estimand.

DEMONSTRATION: WHEN CLASSICAL ESTIMATION IS WRONG SIMULATION

We simulate data with post-treatment L_1 and censoring depending on L_1 and A_1 :

$$\begin{aligned} W &\sim N(0, 1), \\ A_0 &\sim \text{Bernoulli}(0.5), \\ L_1 \mid A_0, W &\sim N(0.8A_0 + 0.9W, 1), \\ A_1 \mid L_1, W, A_0 &\sim \text{Bernoulli}(\text{expit}(-0.5 + 1.2\mathbb{I}\{L_1 > 0\} + 0.2W)), \\ C \mid L_1, A_1, W &\sim \text{Bernoulli}(\text{expit}(0.8 - 1.3L_1 + 0.4A_1 + 0.2W)), \\ Y \mid A_0, A_1, L_1, W &\sim \text{Bernoulli}(\text{expit}(-2 + 0.7A_0 + 0.9A_1 + 1.1L_1 + 0.8W)). \end{aligned}$$

TYPICAL RESULTS

Table 1: Typical performance (1000 replications; $n = 5000$).

Method	Bias	RMSE	Coverage (95%)
Covariate-adjusted GLM (complete case)	-0.11	0.18	0.62
MSM (IPW; stabilized)	-0.02	0.27	0.88
Parametric g-computation	-0.01	0.12	0.93
TMLE (SL + cross-fitting)	0.00	0.09	0.95

IMPLEMENTATION IN R PACKAGE ECOSYSTEM

- `tmle`: point-treatment TMLE.
- `tmle3`: modular TMLE specification (likelihood + parameter objects).
- `sl3`: Super Learner learners/pipelines.
- `origami`: cross-fitting and CV folds.
- `ltmle`: longitudinal TMLE for regimes and censoring.

CODE EXAMPLES Point-treatment TMLE

```
library(tmle)
fit <- tmle(Y=Y, A=A, W=W, family="binomial")
fit$estimates$ATE$psi
fit$estimates$ATE$CI
```

Longitudinal TMLE (dynamic regime)

```
library(ltmle)
fit <- ltmle(data=df,
  Anodes=c("A0", "A1"), Lnodes=c("L1"),
  Cnodes=c("C"), Ynodes=c("Y"),
```

```

abare=list(c(1, function(L1) as.numeric(L1>0)),
           c(0, function(L1) as.numeric(L1>0))),
SL.library=c("SL.glm", "SL.glmnet", "SL.ranger"))
fit$effect.measures$ATE
fit$effect.measures$ATE.CI

```

DIAGNOSTICS FOR TMLE: THEORY AND IMPLEMENTATION

TMLE inherits its asymptotic guarantees from (i) identifiability conditions (e.g., positivity), (ii) stability of the estimating equations driven by the efficient influence function (EIF), and (iii) control of the second-order remainder term

$$\hat{\Psi} - \Psi(P_0) = \mathbb{P}_n D^*(O; \hat{Q}, \hat{g}) + R_2(\hat{Q}, \hat{g}), \quad R_2(\hat{Q}, \hat{g}) = O_p(\|\hat{Q} - Q_0\| \|\hat{g} - g_0\|).$$

The diagnostic checks below map directly to these theoretical requirements and provide practical safeguards when implementing TMLE with flexible learners and cross-fitting.

1. OVERLAP / POSITIVITY: PROPENSITY SCORE HISTOGRAMS AND TRUNCATION SENSITIVITY

Positivity requires the treatment mechanism to be bounded away from 0 and 1:

$$0 < g_0(1 | W = w) < 1 \quad \text{for all } w \text{ with support under } P_0.$$

In finite samples, violations manifest as *practical positivity problems*: regions of covariate space where $\hat{g}(1 | W)$ is extremely small or large. This directly inflates the EIF through inverse weights. For the ATE EIF,

$$D^*(O; P) \supset \frac{\mathbb{I}(A = 1)}{g(1 | W)}(Y - Q(1, W)) - \frac{\mathbb{I}(A = 0)}{g(0 | W)}(Y - Q(0, W)),$$

so small $g(1 | W)$ (or $g(0 | W)$) produces large contributions and unstable variance. When censoring is present, the issue compounds because the EIF typically contains additional inverse-probability-of-censoring terms.

From the remainder perspective, extreme $\hat{g}$ can also degrade R_2 by increasing the effective magnitude of the product-of-errors contribution in regions with poor support.

Implementation.

- Plot histograms (or density plots) of $\hat{g}(1 | W)$ by treatment arm; for longitudinal TMLE, also plot censoring probabilities $\hat{\pi}(\cdot)$ if modeled.
- Identify mass near 0 or 1 (e.g., below 0.01 or above 0.99). These indicate limited overlap.
- Apply *truncation* (a.k.a. bounding) to enforce $\hat{g}_\delta = \min(\max(\hat{g}, \delta), 1 - \delta)$ for $\delta \in \{0.01, 0.02, 0.05\}$ and re-run TMLE.

Reporting sensitivity. A common practice is to report $\hat{\Psi}$ across several truncation thresholds and note whether conclusions are stable. Large swings indicate that the estimand may not be well supported by the data, and that the estimator is relying heavily on extrapolation.

2. STABILITY: CLEVER COVARIATE DISTRIBUTION AND EIF OUTLIERS

The targeting step in TMLE uses the *clever covariate* (for the ATE)

$$H(A, W) = \frac{\mathbb{I}(A = 1)}{\hat{g}(1 | W)} - \frac{\mathbb{I}(A = 0)}{\hat{g}(0 | W)}.$$

When $\hat{g}$ approaches 0 or 1, $H(A, W)$ becomes large in magnitude. This impacts:

- the fluctuation regression (numerical instability, over-influence of a few observations),
- the EIF variance (large empirical variance),
- finite-sample behavior (high leverage points dominate inference).

Moreover, inference is based on the empirical variance of the EIF:

$$\widehat{\text{Var}}(\hat{\Psi}) = \frac{1}{n} \widehat{\text{Var}}\left(D^*(O; \hat{Q}^*, \hat{g})\right).$$

If a small number of EIF values are extreme, Wald-type confidence intervals can be misleading and standard errors can become unstable.

Implementation.

- Plot the distribution of the clever covariate $H(A, W)$ (histogram/boxplot).
- Compute summary diagnostics: max, 99th percentile, and fraction exceeding a threshold (e.g., $|H| > 20$).
- Plot the EIF values $D_i = D^*(O_i; \hat{Q}^*, \hat{g})$ and identify outliers (e.g., via boxplot or standardized residuals $D_i/\hat{\sigma}_D$).
- Investigate whether outliers correspond to near-positivity violations, rare covariate patterns, or model extrapolation.

Mitigations. Common mitigations include stronger truncation of $\hat{g}$, simplifying or regularizing the propensity model, reducing model complexity in regions with sparse support, and checking whether a restricted estimand (or alternative analysis set) is more appropriate.

3. TRANSPARENCY: LEARNER LIBRARY, TUNING, FOLDS, AND SEEDS

When machine learning is used for nuisance estimation, TMLE relies on the fact that the first-order term is governed by the EIF and that nuisance estimation error enters primarily through the second-order remainder R_2 . Achieving valid inference typically requires:

- nuisance estimators with adequate convergence rates (often heuristically supported by cross-validated risk minimization),
- *cross-fitting* (sample splitting) to avoid empirical process restrictions and reduce overfitting-induced bias in the EIF evaluation.

Because these properties depend on algorithmic choices (model classes, hyperparameters, folds), reproducibility and interpretability require complete transparency in what was fit.

Implementation.

At minimum, report:

- The candidate learner library for Q and g (e.g., generalized linear model, penalized regression, random forest, gradient boosting).
- Hyperparameter tuning approach (grid, internal CV, defaults) and the final selected parameters.
- Cross-fitting scheme: number of folds K , whether folds are stratified by treatment, and whether nuisance fits were trained on $K - 1$ folds and evaluated on held-out folds.
- Random seeds (global and, if feasible, learner-specific) to ensure exact reproducibility.

Implementation note. For regulated environments, it is often appropriate to: (i) pin package versions, (ii) log session info, and (iii) archive model objects or at least the fitted predictions $\hat{Q}$ and $\hat{g}$ by fold to support audit trails.

4. SENSITIVITY: COMPARE AGAINST G-COMPUTATION AND MSM UNDER ALTERNATE TRUNCATION

TMLE is doubly robust and efficient under conditions, but it is not immune to finite-sample instability or poor support. Comparing TMLE to alternative estimators helps disentangle where sensitivity arises:

- **g-computation** (outcome regression plug-in) relies primarily on correct specification (or adequate learning) of Q and can fail if Q extrapolates poorly.
- **Marginal structural models (MSM) via inverse probability weighting (IPW)** rely primarily on correct specification of treatment/censoring models and are especially sensitive to positivity and weight instability.
- **TMLE** combines both and targets the estimand, so discrepancies often indicate either (a) practical positivity problems, (b) nuisance model instability, or (c) an estimand mismatch in the comparator.

Truncation sensitivity provides a structured way to evaluate reliance on sparse regions: if conclusions are stable across truncation thresholds, inference is less likely to be driven by a small set of high-leverage observations.

Implementation. An example workflow:

1. Fit TMLE at multiple truncation levels $\delta \in 0.01, 0.02, 0.05$ for propensity (and censoring, if relevant).
2. Fit outcome-regression g-computation using the same Q learner(s) used inside TMLE, and compute

$$\hat{\Psi}_{\text{gcomp}} = \frac{1}{n} \sum_{i=1}^n (\hat{Q}(1, W_i) - \hat{Q}(0, W_i)).$$

3. Fit MSM/IPW with stabilized weights (and the same truncation set), and compute the marginal contrast.
4. Compare point estimates and uncertainty intervals. Large divergence between TMLE and g-computation suggests g -driven instability or overlap issues; large divergence between TMLE and MSM suggests weight instability or misspecification in the treatment/censoring model.

Reporting. Provide a small table summarizing estimates across methods and truncation thresholds; this is often more persuasive than technical prose because it directly communicates robustness (or lack thereof) of conclusions.

CONCLUSION

TMLE provides a principled approach for estimating marginal, intervention-based estimands in complex clinical settings. For clinical trials, TMLE is useful beyond observational confounding adjustment: it supports covariate adjustment for precision, explicit handling of missing outcomes, and estimation of strategy-based estimands aligned with ICH E9(R1). By targeting the estimand directly and using EIF-based inference (with cross-fitting), TMLE avoids common pitfalls of classical adjusted analyses—including non-collapsibility and post-treatment conditioning—while enabling flexible nuisance modeling.

REFERENCES

- [1] van der Laan, M.J. and Rose, S. *Targeted Learning: Causal Inference for Observational and Experimental Data*. Springer Series in Statistics, New York, NY: Springer; 2011.
- [2] Talbot, D., Diop, A. et al. Guidelines and best practices for the use of targeted maximum likelihood and machine learning when estimating causal effects of exposures on time-to-event outcomes. *Statistics in Medicine*, **44**: e70034. <https://doi.org/10.1002/sim.70034>
- [3] Havlir, D.V., Balzer, L.B., Charlebois, E.D., et al. HIV testing and treatment with the use of a community health approach in rural Africa. *The New England Journal of Medicine*. 2019; **381**(3):219–229.
- [4] Balzer, L.B., Ayieko, J., Kwarisiima, D., et al. Far from MCAR: obtaining population-level estimates of HIV viral suppression. *Epidemiology*. 2020;31(5):620–627. doi:10.1097/EDE.0000000000001215.
- [5] Hejazi, N.S., van der Laan, M.J., Janes, H.E., Gilbert, P.B., and Benkeser, D.C. Efficient nonparametric inference on the effects of stochastic interventions under two-phase sampling, with applications to vaccine efficacy trials. *Biometrics*. 2021;77(4):1241–1253. doi:10.1111/biom.13375.

ACKNOWLEDGMENTS

The author thanks Phastar for financing their involvement in this conference.

CONTACT INFORMATION

Your comments and questions are valued and encouraged! Contact the author at:

Author Name: Joshua Enxing

Company: Phastar

Address: 2D Bollo Ln, Chiswick, London, W4 5LE

Work Phone: +44 (0)20 7183 7062

Email: joshua.enxing@phastar.com

Website: <https://www.linkedin.com/in/joshua-enxing-01b3b162/>

Brand and product names are trademarks of their respective companies.