

COUNTERFACTUAL PREDICTION WITH CONTINUOUS TREATMENTS

Anonymous authors

Paper under double-blind review

ABSTRACT

Continuous treatments replace two potential-outcome arms by an indexed family $\{Y(a) : a \in \mathcal{A}\}$, making support, smoothness, and the distinction between marginal dose response and a joint counterfactual process unavoidable. We introduce HCGM-Dose, a model that abduces a distribution over a latent unit representation from pretreatment covariates and evaluates a differentiable dose-indexed structural mechanism. With factorized Cauchy abduction and a mechanism affine in the abducted unit, the construction yields closed-form conditional predictive distributions, quantile surfaces, finite-dose location contrasts, and local dose derivatives. Under consistency, weak conditional exchangeability, dose-density positivity, and correct specification, factual population log-risk recovers the marginal response kernels on supported doses. It does not identify unsupported extrapolations or the joint law of the full counterfactual dose process. Pairwise individual contrast laws require an explicit cross-dose noise coupling. In a five-seed synthetic study we compare HCGM-Dose with its Gaussian twin, direct smooth-dose learners, and a discretized-dose baseline under randomization, observed confounding, outcome contamination, and a hidden-confounding negative control. Pinned official protocols produce a split result: VCNet/DRNet are stronger on Simu1, whereas HCGM-Dose-Gaussian has the lowest mean error over five frozen IHDP and News semi-synthetic splits. Rankings vary by split and data-generating process. Under increasing dose-selection bias, strong-support error remains stable while full-grid error rises as low-support regions expand. The experiments evaluate response surfaces, derivatives, finite-dose contrasts, and dose-policy regret using oracle fields withheld from training.

1 INTRODUCTION

A binary treatment asks a model to compare two intervention-indexed worlds. A continuous treatment asks for an entire family of worlds, such as drug dose, exposure intensity, or policy strength. This changes more than the number of labels. Every queried dose needs local treatment support; a useful estimator must vary coherently across dose; and a smooth curve outside observed support is still an extrapolation rather than an identified causal object.

Classical generalized propensity-score methods formalize identification with continuous exposures (Hirano & Imbens, 2004; Imai & van Dyk, 2004). Recent neural estimators improve the flexibility and continuity of learned dose-response curves (Schwab et al., 2020; Bica et al., 2020; Nie et al., 2021). Our question is complementary: can the HCGM grammar of latent abduction, treatment-indexed action, and analytic reduction produce a transparent distributional dose-response model whose estimands and counterfactual limits remain explicit?

We propose HCGM-Dose. An encoder uses only pretreatment covariates to define a factorized Cauchy distribution over a latent unit U . A differentiable mechanism maps dose a to an intercept, unit loading, and event scale, then generates

$$Y(a) = c(a) + w(a)^\top U + \sigma(a)E_a. \quad (1)$$

More precisely, the data-generating world contains one realization $U = u^*$, and its dose-indexed outcome is $Y_{u^*}(a) = f_\theta(a, E; u^*)$. The encoder treats $X = x$ as factual evidence and returns epistemic uncertainty about that same unobserved token. It does not draw a new identity for each

dose. Every factual or alternative dose query therefore reuses the abduction result obtained from x and changes only a in the shared mechanism. Affine Cauchy closure makes the marginal response at every queried dose analytic. The model therefore provides a coherent conditional location surface, finite-dose contrasts, and a differentiable local dose effect without discretizing the intervention.

Our contributions are:

1. We formulate HCGM-style abduction–intervention–reduction for a continuous treatment and derive closed-form conditional response and quantile surfaces.
2. We separate identified dose-indexed marginals from model-implied cross-dose contrast laws, and separate supported interpolation from unsupported extrapolation.
3. We establish marginal identification, population-risk recovery, analytic dose derivatives, a discretization error bound, and bounded Cauchy residual scores under explicit conditions.
4. We provide a clean-room implementation, a factual-only synthetic protocol measuring surface, derivative, finite-contrast, and policy recovery, pinned official Simu1/IHDP/News comparisons, and an exact-density support stress test.

HCGM-Dose does not claim that latent coordinates are physically identified, that abduction removes hidden confounding, or that the factual data determine a joint stochastic process over all counterfactual doses.

2 RELATED WORK

Continuous-treatment identification. The generalized propensity score extends balancing-score ideas to a continuous treatment through the conditional treatment density $r(a, x) = f_{A|X}(a | x)$ (Hirano & Imbens, 2004). Weak unconfoundedness identifies each dose-specific potential-outcome marginal without requiring joint independence of the entire potential-outcome process. Imai and van Dyk (Imai & van Dyk, 2004) provide a broader propensity-function formulation for nonbinary treatments. HCGM-Dose uses the same causal boundary: its outcome model is an estimator under these assumptions, not a substitute for them.

Neural dose-response estimation. DRNet partitions the dose range and uses treatment/dose-specific prediction heads (Schwab et al., 2020). SCIGAN learns individualized response curves through a generative-adversarial construction (Bica et al., 2020). VCNet uses varying coefficients to preserve continuity and adds functional targeted regularization for average dose-response estimation (Nie et al., 2021). HCGM-Dose shares the goal of continuous response surfaces but differs in its explicit latent abduction, analytic distributional reduction, and emphasis on conditional location/quantile estimands under heavy-tailed response models. Our direct and binned baselines test whether any empirical gain comes from this factorization rather than merely from a robust loss or smooth dose features.

Robustness and hidden confounding. Flexible response models remain vulnerable to unmeasured confounding. Sensitivity and partial-identification approaches for continuous treatments make this limitation explicit (Marmarelis et al., 2023). We include hidden confounding only as a negative control. The bounded residual scores of the Cauchy likelihood address gross outcome contamination; they do not solve assignment bias.

3 SETUP AND ESTIMANDS

We observe i.i.d. factual records $\mathcal{D}_n = \{(X_i, A_i, Y_i)\}_{i=1}^n$, where X contains only pretreatment covariates, $A \in \mathcal{A} = [a_{\min}, a_{\max}] \subset \mathbb{R}$ is a continuous dose, and $Y = Y(A)$. The subscript i is only bookkeeping for records; it does not define a different unit variable or response function for every row. At model level U is the unit-selection variable and $U = u^*$ is the actual token/unit embedding associated with a record.

Assumption 1 (Dose-indexed marginal identification). *For every a in the target dose region, we assume consistency, no interference, weak conditional exchangeability $Y(a) \perp\!\!\!\perp A | X$, and a regular conditional density satisfying $f_{A|X}(a | x) > 0$ for the target covariate–dose pairs.*

Weak exchangeability is pointwise in a ; it does not assert that the complete process $\{Y(a) : a \in \mathcal{A}\}$ is jointly independent of assignment. Stable estimation needs more than pointwise positivity. We call a region $\mathcal{S}_\epsilon = \{(x, a) : f_{A|X}(a | x) \geq \epsilon\}$ a strong-support region. All causal claims are restricted to supported doses.

Let

$$F_a(y | x) = \mathbb{P}\{Y(a) \leq y | X = x\}, \quad (2)$$

$$Q_a(p | x) = \inf\{y : F_a(y | x) \geq p\}, \quad p \in (0, 1). \quad (3)$$

Our primary surface and finite-dose contrast are

$$m(a, x) = Q_a(1/2 | x), \quad (4)$$

$$\tau(a, a'; x) = m(a, x) - m(a', x). \quad (5)$$

When m is differentiable in dose, the local dose effect is

$$g(a, x) = \partial_a m(a, x). \quad (6)$$

The derivative is a local causal estimand only where neighboring doses are supported. Smoothness alone cannot identify it across a support gap.

For a bounded utility $u(y, a)$ and a declared feasible set $\mathcal{A}_\epsilon(x) = \{a : (x, a) \in \mathcal{S}_\epsilon\}$, a support-aware policy may select

$$\pi(x) \in \arg \max_{a \in \mathcal{A}_\epsilon(x)} \mathbb{E}\{u(Y(a), a) | X = x\}. \quad (7)$$

Our synthetic location-policy diagnostic takes u to be the response location. It is not a clinical recommendation.

4 HCGM-DOSE: LATENT ABDUCTION ALONG AN INTERVENTION PATH

HCGM-Dose first abduces the actual token from pretreatment evidence. In the general grammar this result may be a point estimate, a Gaussian distribution expressing local finite-variance uncertainty, or a Cauchy distribution expressing heavy-tailed uncertainty over candidate embeddings for the same token. This paper uses Cauchy as its primary model and Gaussian as a matched ablation; the point case is a deterministic boundary rather than a reported estimator:

$$q_\phi(u | X = x) = \prod_{j=1}^d \text{Cauchy}(u_j; \mu_{\phi,j}(x), \gamma_{\phi,j}(x)), \quad \gamma_{\phi,j}(x) > 0. \quad (8)$$

This is learner-side epistemic uncertainty about the fixed actual realization u^* selected by U , not a population distribution that resamples physical identity. We retain μ only as the legacy implementation field name; mathematically these parameters are Cauchy locations/medians rather than means, and γ parameters are scales rather than standard deviations; non-degenerate Cauchy coordinates have no finite means or variances. Heavy tails keep remote candidate embeddings plausible relative to a Gaussian family, but do not make every outcome attainable; that remains controlled by the mechanism and event-noise support.

The dose is not part of abduction. A shared differentiable mechanism produces $c_\theta(a) \in \mathbb{R}$, $w_\theta(a) \in \mathbb{R}^d$, and $\sigma_\theta(a) > 0$, and acts by

$$Y(a) = c_\theta(a) + w_\theta(a)^\top U + \sigma_\theta(a) E_a, \quad E_a \sim \text{Cauchy}(0, 1). \quad (9)$$

Fixing $U = u$ gives the token-level generator $Y_u(a) = f_\theta(a, E; u)$. For learner-side prediction, a candidate $\tilde{U} \sim q_\phi(\cdot | x)$ is propagated through this generator. All dose queries use this same factual abduction distribution; an alternative dose is never fed back to the encoder to infer a different unit. The implementation makes this literal: `abduce(X)` constructs one batched belief object, and every `response_parameters_from_unit` call receives that same object while only the dose changes. For every queried dose, $E_a \perp\!\!\!\perp U | X$. A joint relationship among $\{E_a\}$ is unnecessary for marginal response estimation and is deliberately left unspecified.

Affine Cauchy closure gives the following predictive distribution:

$$m_{\vartheta}(a, x) = c_{\theta}(a) + w_{\theta}(a)^{\top} \mu_{\phi}(x), \quad (10)$$

$$s_{\vartheta}(a, x) = \sum_{j=1}^d |w_{\theta,j}(a)| \gamma_{\phi,j}(x) + \sigma_{\theta}(a), \quad (11)$$

$$Y(a) \mid X = x \sim \text{Cauchy}\{m_{\vartheta}(a, x), s_{\vartheta}(a, x)\}. \quad (12)$$

Consequently,

$$Q_{\vartheta,a}(p \mid x) = m_{\vartheta}(a, x) + s_{\vartheta}(a, x) \tan\{\pi(p - 1/2)\}. \quad (13)$$

If c_{θ} and w_{θ} are differentiable, the location derivative is

$$g_{\vartheta}(a, x) = c'_{\theta}(a) + w'_{\theta}(a)^{\top} \mu_{\phi}(x). \quad (14)$$

The implementation parameterizes the dose mechanism with a shared multilayer network over polynomial and Fourier dose features. This preserves continuity without assigning a separate head to each discretized interval.

For a factual observation (x_i, a_i, y_i) , analytic reduction yields

$$\ell_i(\vartheta) = \log\{\pi s_{\vartheta}(a_i, x_i)\} + \log\left[1 + \left\{\frac{y_i - m_{\vartheta}(a_i, x_i)}{s_{\vartheta}(a_i, x_i)}\right\}^2\right]. \quad (15)$$

We minimize the unweighted empirical factual risk. Oracle response surfaces are used only after training for synthetic evaluation.

Pairwise model-implied contrasts. For two doses a, a' , using the same U and independent event draws gives

$$Y(a) - Y(a') \mid x \sim \text{Cauchy}\{\tau_{\vartheta}(a, a'; x), \kappa_{\text{ind}}(a, a'; x)\}, \quad (16)$$

$$\kappa_{\text{ind}} = \sum_j |w_j(a) - w_j(a')| \gamma_j(x) + \sigma(a) + \sigma(a'). \quad (17)$$

Reusing one event realization instead replaces the last two terms by $|\sigma(a) - \sigma(a')|$. Both laws are coupling choices. Neither is identified from dose-indexed factual marginals alone.

5 IDENTIFICATION AND PROPERTIES

Full proofs are in Appendix A.

Proposition 1 (Marginal dose-response identification). *Under Assumption 1, for every supported dose a ,*

$$F_a(y \mid x) = \mathbb{P}(Y \leq y \mid A = a, X = x). \quad (18)$$

Thus $Q_a(p \mid x)$, $m(a, x)$, and finite supported contrasts are identified P_X -almost surely. Derivatives additionally require local support and regularity permitting differentiation of the identified surface.

Proposition 2 (Analytic response and derivative). *Under Equation (8), conditional coordinate independence, and $E_a \perp\!\!\!\perp U \mid X$, Equation (9) yields Equations (10)–(13). If c, w are differentiable, its location derivative is Equation (14). Where every nonzero $w_j(a)$ remains away from zero and σ is differentiable, the quantile derivative is*

$$\partial_a Q_a(p \mid x) = g(a, x) + \left\{ \sum_j \text{sign}(w_j(a)) w'_j(a) \gamma_j(x) + \sigma'(a) \right\} \tan\{\pi(p - 1/2)\}. \quad (19)$$

At a zero crossing of a loading, the scale surface can have a kink.

Theorem 1 (Fisher consistency on supported doses). *Suppose a parameter ϑ^* realizes the true factual conditional Cauchy kernel and all candidate cross-entropies are finite. For the unpenalized population log-risk,*

$$R(\vartheta) - R(\vartheta^*) = \mathbb{E}_{X,A}[D_{\text{KL}}\{p^*(\cdot \mid X, A) \parallel p_{\vartheta}(\cdot \mid X, A)\}] \geq 0. \quad (20)$$

Table 1: Claim boundary for continuous-dose quantities.

Quantity	Status
$F_a, Q_a, m(a, x)$ on support	Identified under Assumption 1
$g(a, x)$	Identified with local support and differentiability
Fitted $m_{\vartheta}, s_{\vartheta}$	Population-risk recovery under correct unpenalized specification
Unsupported dose surface	Model extrapolation, not identified
Pairwise or full cross-dose joint law	Coupling-dependent and not identified by marginals
Latent coordinates	Model-implied and gauge non-identified

Every finite-risk global minimizer recovers the conditional response kernel $P_{X,A}$ -almost surely. If the true and fitted location/scale surfaces are continuous in a and the conditional treatment density is positive throughout an interval, equality extends from almost every dose to every dose in that interval. This does not identify a unique latent factorization.

Proposition 3 (Unsupported doses are not recovered by factual risk). *If an open dose set B has zero conditional assignment probability for a positive-probability covariate set, two response models can agree on factual support and differ on B while attaining identical factual risk. Smooth parameterization selects an extrapolation but does not identify it.*

Proposition 4 (Pairwise contrast laws and non-identification). *The independent-resampling and shared-realization couplings give the pairwise Cauchy laws described after Equation (17). They have the same location contrast but can have different scales and probabilities of benefit. Dose-indexed factual marginals do not identify either coupling, and arbitrary pairwise couplings need not define a coherent process over all doses.*

Proposition 5 (Discretization approximation bound). *If $m(\cdot, x)$ is L_x -Lipschitz and a bin of width at most h is represented by its center $b(a)$, then*

$$|m(a, x) - m\{b(a), x\}| \leq L_x h / 2. \quad (21)$$

For two binned doses, the induced contrast error is at most $L_x h$ when both bins have width at most h . A piecewise-constant representation has zero within-bin derivative and cannot recover a nonzero local dose effect there.

Proposition 6 (Bounded Cauchy residual scores). *For the loss in Equation (15) and $s > 0$,*

$$\left| \frac{\partial \ell}{\partial m} \right| \leq \frac{1}{s}, \quad \left| \frac{\partial \ell}{\partial \log s} \right| \leq 1. \quad (22)$$

This bounds residual scores, not causal bias from confounding and not parameter gradients without additional Jacobian bounds.

6 EXPERIMENTS

The experiment asks four separate questions: can the response surface be learned under random assignment; does observed dose selection degrade recovery; does a Cauchy likelihood protect against gross outcome contamination; and does hidden confounding break all factual outcome learners as the theory predicts?

Nonlinear synthetic SCM. We draw six observed covariates independently from a standard Gaussian and use the conditional location surface

$$m(a, x) = 0.45 \sin x_1 + 0.18x_2^2 - 0.22x_4x_5 + 0.15x_6 \quad (23)$$

$$+ \{0.85 + 0.30 \tanh x_1\} \sin(\pi a) + 0.30x_2a - 0.55\{a - \text{sigmoid}(0.55x_3)\}^2. \quad (24)$$

The oracle derivative is computed analytically. Factual outcomes add Gaussian noise with standard deviation 0.22, so this DGP is deliberately outside the primary Cauchy outcome family.

In the randomized scenario, $A \sim \text{Unif}(0.02, 0.98)$. In the observational scenario,

$$A = \text{sigmoid}(0.85X_1 - 0.60X_2 + 0.30X_3X_4 + 1.05\varepsilon_A), \quad \varepsilon_A \sim \mathcal{N}(0, 1). \quad (25)$$

Table 2: Integrated location-surface MAE (mean $\pm$ SE over five seeds; lower is better).

Model	Randomized	Observed conf.	Observed + outliers	Hidden conf. (neg.)
HCGM-Dose-Cauchy	0.066 $\pm$ 0.002	0.083 $\pm$ 0.006	0.079 $\pm$ 0.005	0.715 $\pm$ 0.016
HCGM-Dose-Gaussian	0.057 $\pm$ 0.001	0.072 $\pm$ 0.006	0.386 $\pm$ 0.009	0.723 $\pm$ 0.014
Direct-Cauchy	0.059 $\pm$ 0.002	0.071 $\pm$ 0.005	0.076 $\pm$ 0.003	0.723 $\pm$ 0.017
Direct-Gaussian	0.052 $\pm$ 0.004	0.062 $\pm$ 0.005	0.420 $\pm$ 0.025	0.725 $\pm$ 0.014
Binned-Cauchy	0.131 $\pm$ 0.004	0.152 $\pm$ 0.005	0.155 $\pm$ 0.005	0.687 $\pm$ 0.017

The contaminated scenario independently replaces 5% of training and 5% of validation labels by adding $+12$ or -12 with equal probability. The oracle test surface remains clean. In the hidden-confounding negative control, an unobserved $H \sim \mathcal{N}(0, 1)$ enters both dose assignment with coefficient 0.95 and the factual response with coefficient $0.75 + 0.55A$.

Models. We compare Cauchy HCGM-Dose, an architecture-matched Gaussian HCGM-Dose, direct smooth dose-response networks with Cauchy or Gaussian likelihood, and a five-bin Cauchy model. The direct learners receive (X, A) and use the same polynomial and Fourier dose features but no latent abduction/mechanism factorization. The binned model has separate piecewise-constant heads.

Protocol and metrics. For each of five frozen seeds, we generate 2,400 records and use a 65/15/20 train/validation/test split. Training and early stopping see factual records only. We evaluate on 31 doses from 0.05 to 0.95 after model selection. Primary surface error is integrated absolute location error. Secondary metrics are the mean absolute derivative error, the absolute error of the 0.8 versus 0.2 location contrast, and oracle location-policy regret on the same dose grid. The binned model’s derivative is zero inside bins and undefined at boundaries; we score its within-bin derivative as zero. Reported values are means and standard errors over seeds.

Pinned official benchmark protocols. We additionally run the public VCNet repository at pinned commit a2e9538. Its official Simul generator supplies 500 factual training records, 200 test covariates, and an oracle average dose-response grid. We use five independently generated datasets. The repository’s IHDP protocol uses 25 real covariates, a semi-synthetic continuous treatment and response, 471/276 train/test splits, and five frozen splits of one generated response surface. The News protocol uses 498 real document covariates, a semi-synthetic continuous treatment and response, 2,000/993 train/test splits, and five frozen splits of one generated response surface. For all three benchmarks we run the repository’s VCNet-TR and DRNet-TR classes with their published 800-epoch hyperparameters. HCGM-Dose receives the identical factual matrices and is evaluated on the identical oracle grids. The common metric is average dose-response MSE; this tests population-curve prediction, not the full conditional distributional target of HCGM-Dose. The official generators use symmetric zero-location Gaussian outcome noise, so their oracle conditional mean and median surfaces coincide; the Cauchy HCGM variant contributes its conditional location before averaging over test covariates. External source, commit, runtime, and license boundaries are recorded in `implementation/OFFICIAL_BENCHMARK_PROVENANCE.md`.

Support-sensitivity protocol. To separate supported interpolation from smooth extrapolation, we use the exact conditional density induced by the observational assignment $A = \text{sigmoid}\{s\eta(X) + 1.05\epsilon_A\}$, where $s \in \{0.5, 1, 2, 3\}$ controls selection strength. We define strong support by $f_{A|X}(a | x) \geq 0.10$, cross each held-out covariate context with 49 doses on $[0.02, 0.98]$, and separately score the full grid, strong-support cells, and low-support cells. Five seeds and 3,000 records per seed are used for HCGM-Dose-Cauchy and Direct-Cauchy. We also restrict both the learned and oracle location policies to that context’s strong-support dose set. This protocol was run on a second NVIDIA GB10 host; all 40 fits produced finite outputs despite a runtime warning that the installed PyTorch build declared support only through compute capability 12.0 while the device reports 12.1.

Table 3: Contaminated observational scenario (mean $\pm$ SE; lower is better).

Model	Derivative	Contrast	Policy regret
HCGM-Dose-Cauchy	0.347 ± 0.023	0.088 ± 0.010	0.005 ± 0.001
HCGM-Dose-Gaussian	1.569 ± 0.087	0.265 ± 0.044	0.124 ± 0.043
Direct-Cauchy	0.375 ± 0.020	0.082 ± 0.011	0.006 ± 0.001
Direct-Gaussian	1.355 ± 0.069	0.338 ± 0.053	0.083 ± 0.018
Binned-Cauchy	1.907 ± 0.006	0.432 ± 0.012	0.085 ± 0.011

Table 4: Official VCNet protocols: average dose-response MSE (mean $\pm$ descriptive SE over five generated datasets or frozen splits; lower is better).

Model	Simul	IHDP semi-synthetic	News semi-synthetic
HCGM-Dose-Cauchy	0.043 ± 0.014	0.168 ± 0.030	0.039 ± 0.012
HCGM-Dose-Gaussian	0.038 ± 0.016	0.107 ± 0.018	0.034 ± 0.010
VCNet-TR-official	0.016 ± 0.003	0.373 ± 0.143	0.047 ± 0.011
DRNet-TR-official	0.025 ± 0.004	0.743 ± 0.165	0.038 ± 0.003

7 RESULTS

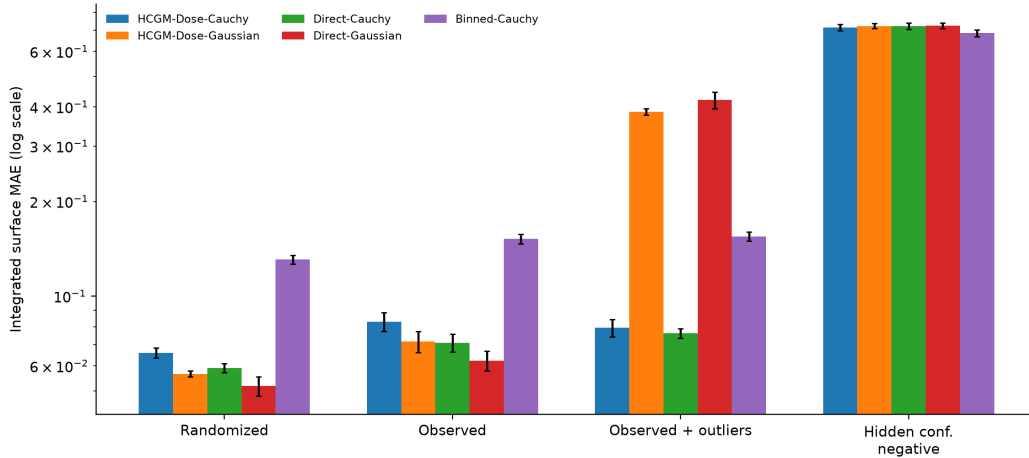


Figure 1: Integrated location-surface MAE on a logarithmic axis (mean $\pm$ standard error over five seeds; lower is better). Hidden confounding is a negative control and is shown on the same metric, not as a claimed identifiable regime.

On clean Gaussian outcomes, Gaussian likelihoods are appropriately competitive and the direct Gaussian learner has the smallest mean surface error. HCGM-Dose is not universally superior to a direct smooth-dose learner. Both continuous model families substantially improve over five-bin discretization on surface and derivative recovery.

Gross-label contamination changes the comparison. Cauchy HCGM-Dose remains close to its clean observational error, while its matched Gaussian model degrades sharply. Direct Cauchy is also robust and is slightly more accurate than HCGM-Dose in this DGP, showing that much of the contamination advantage comes from the likelihood rather than uniquely from latent factorization. The factorization’s value here is analytic distributional reduction and explicit intervention structure, not a claim of universal predictive dominance.

All methods fail in the hidden-confounding negative control: surface, derivative, contrast, and policy errors increase together. This confirms the intended boundary that abduction and factual likelihood do not restore exchangeability.

Table 5: HCGM-Dose support stress under exact treatment density (mean $\pm$ SE over five seeds; lower error is better).

Strength	Supported fraction	Supported MAE	Low-support MAE	Full MAE
0.5	0.916 $\pm$ 0.000	0.065 $\pm$ 0.004	0.124 $\pm$ 0.010	0.070 $\pm$ 0.004
1.0	0.849 $\pm$ 0.001	0.063 $\pm$ 0.003	0.148 $\pm$ 0.017	0.076 $\pm$ 0.005
2.0	0.681 $\pm$ 0.002	0.065 $\pm$ 0.002	0.171 $\pm$ 0.013	0.099 $\pm$ 0.005
3.0	0.545 $\pm$ 0.002	0.062 $\pm$ 0.003	0.157 $\pm$ 0.009	0.105 $\pm$ 0.005

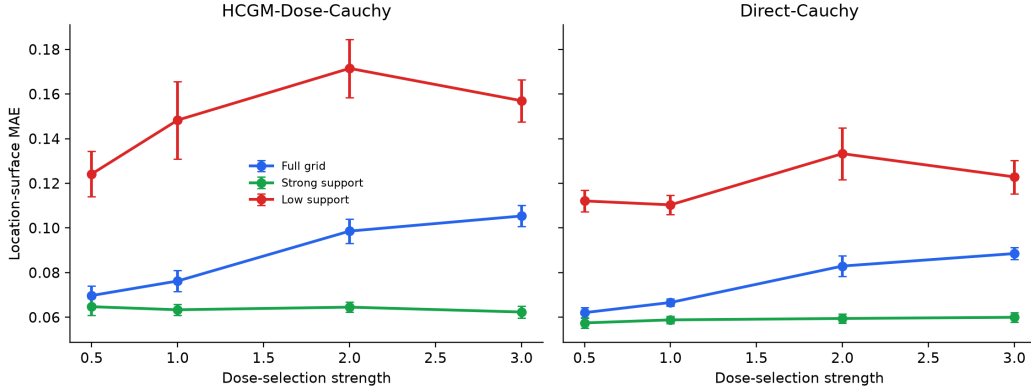


Figure 2: Full-grid, strong-support, and low-support location-surface error as dose-selection bias increases. Strong support uses the exact threshold $f_{A|X}(a | x) \geq 0.10$.

The official protocols reverse the ranking across data-generating processes and sometimes across splits. On Simu1, VCNet-TR and DRNet-TR have lower average-curve error than both HCGM-Dose variants. On IHDP semi-synthetic covariates, HCGM-Dose-Gaussian has the lowest mean error and HCGM-Dose-Cauchy is second, while the two official target-regularized models have larger errors and variability in these five splits. On the 498-covariate News benchmark, HCGM-Dose-Gaussian also has the lowest mean error, followed closely by DRNet-TR and HCGM-Dose-Cauchy. However, DRNet-TR or VCNet-TR wins two of the five individual News splits, and the HCGM standard errors are larger than DRNet’s. These are small-sample descriptive benchmark results, not a universal superiority claim or a significance test. They show that the factorization’s inductive bias can help or hurt depending on the response surface and covariate geometry. They also separate the robust Cauchy question from the Gaussian-efficiency question: without contamination, the matched Gaussian HCGM variant has the lowest mean error on two of the three official protocols.

Support restriction changes the interpretation of apparent degradation. For HCGM-Dose-Cauchy, increasing selection strength from 0.5 to 3.0 reduces the strong-support fraction from 0.916 to 0.545 and increases full-grid error by a factor of 1.51. Yet strong-support error is essentially unchanged (0.065 versus 0.062), while low-support error at strength 3.0 is 2.52 times the strong-support error. Direct-Cauchy shows the same qualitative separation. This does not prove that density thresholding removes all bias; it demonstrates that a smooth full-grid curve mixes supported prediction with unsupported extrapolation.

8 DISCUSSION AND LIMITATIONS

HCGM-Dose turns a continuous treatment into a differentiable index of a shared structural mechanism rather than an ordinary predictor coordinate. In the Cauchy-affine family this supplies an unusually direct map from modeling assumptions to response distributions, quantile surfaces, finite-dose contrasts, and local derivatives.

Its unit ontology is fixed across that path. The world contains one actual $U = u^*$; $q_\phi(u | x)$ records uncertainty about this same token after seeing factual pretreatment evidence. A counterfactual dose

query keeps this abduction result fixed and changes only the dose supplied to $Y_u(a) = f_\theta(a, E; u)$. Re-abducting from an alternative dose would change the unit rather than interrogate its mechanism.

The same transparency exposes limitations. First, unweighted conditional outcome regression may be statistically weak where the treatment density is small even when positivity holds in principle. Density-aware diagnostics, balancing, or doubly robust continuous-treatment objectives are natural next steps, but are not evidence in this paper. Second, a flexible dose network can extrapolate smoothly outside support; such smoothness is a modeling choice, not causal identification. Third, a pairwise contrast distribution depends on how event noise is coupled across doses. Defining all pairwise couplings does not automatically define a coherent continuous stochastic process.

The support experiment makes the second limitation operational: full-grid error grows mainly because the low-density portion of the query grid expands, not because error grows inside the declared strong-support region. A useful continuous-treatment system should therefore return a response prediction together with a treatment-density or support status, and policy optimization should be constrained to supported doses unless extrapolation is explicitly accepted as a modeling decision.

The current study is synthetic and semi-synthetic, and deliberately narrow. It establishes a mathematical and computational wedge, not clinical utility. The pinned Simu1, IHDP, and high-dimensional News comparisons plus the exact-density stress test strengthen the empirical position, but direct SCIGAN execution, the TCGA benchmark, and a real application whose design supports a continuous-dose causal interpretation remain missing. Multiple simultaneous treatments, longitudinal dosing, time-varying confounding, instrumental variables, and hidden-confounding partial identification remain separate research problems.

REFERENCES

- Ioana Bica, James Jordon, and Mihaela van der Schaar. Estimating the effects of continuous-valued interventions using generative adversarial networks. In *Advances in Neural Information Processing Systems*, volume 33, pp. 16434–16445, 2020.
- Keisuke Hirano and Guido W. Imbens. The propensity score with continuous treatments. In Andrew Gelman and Xiao-Li Meng (eds.), *Applied Bayesian Modeling and Causal Inference from Incomplete-Data Perspectives*, pp. 73–84. Wiley, 2004.
- Kosuke Imai and David A. van Dyk. Causal inference with general treatment regimes: Generalizing the propensity score. *Journal of the American Statistical Association*, 99(467):854–866, 2004.
- Myrl G. Marmarelis, Elizabeth Haddad, Andrew Jesson, Neda Jahanshad, Aram Galstyan, and Greg Ver Steeg. Partial identification of dose responses with hidden confounders. In *Proceedings of the Thirty-Ninth Conference on Uncertainty in Artificial Intelligence*, pp. 1368–1379. PMLR, 2023.
- Lizhen Nie, Mao Ye, Qiang Liu, and Dan Nicolae. VCNet and functional targeted regularization for learning causal effects of continuous treatments. *arXiv preprint arXiv:2103.07861*, 2021.
- Patrick Schwab, Lorenz Linhardt, and Walter Karlen. Learning counterfactual representations for estimating individual dose-response curves. In *AAAI Conference on Artificial Intelligence*, volume 34, pp. 5612–5619, 2020.

A PROOFS AND ADDITIONAL THEORY

A.1 PROOF OF PROPOSITION 1

Consistency gives $Y = Y(a)$ on the regular conditional event $A = a$. Weak exchangeability gives $\mathcal{L}\{Y(a) \mid A = a, X = x\} = \mathcal{L}\{Y(a) \mid X = x\}$. Combining the two equalities proves the conditional CDF identity wherever the conditional density is positive. Quantiles and locations are functionals of that marginal. A finite contrast uses two identified marginal locations. Differentiation requires the surface to be identified in a neighborhood and the derivative/interchange regularity stated in the main text. $\square$

A.2 PROOF OF PROPOSITION 2

For independent $Z_j \sim \text{Cauchy}(\mu_j, \gamma_j)$, stability of the Cauchy law implies

$$r + \sum_j v_j Z_j \sim \text{Cauchy} \left(r + \sum_j v_j \mu_j, \sum_j |v_j| \gamma_j \right).$$

Apply this identity to $U_1, \dots, U_d, E_a$ conditional on $X = x$ to obtain the location and scale in Equations (10) and (11). The Cauchy quantile formula gives Equation (13). Differentiating the location is direct. Away from zeros of w_j , $\partial_a |w_j| = \text{sign}(w_j) w'_j$, which gives the displayed quantile derivative. $\square$

A.3 PROOF OF THEOREM 1

Condition on (X, A) . Subtracting the true conditional cross-entropy from a candidate cross-entropy yields the conditional KL divergence. Integrating gives Equation (20). Nonnegativity and the equality condition of KL divergence establish recovery $P_{X,A}$ -almost surely. For almost every x , positive density on an interval makes the exceptional dose set Lebesgue-null. If two continuous location or scale functions differed at one dose, continuity would make them differ on a neighborhood of positive measure, contradicting almost-everywhere equality. Latent non-identification follows because multiple encoder/mechanism factorizations can induce the same response kernels. $\square$

A.4 PROOF OF PROPOSITION 3

Let B be an open set outside conditional factual support. Modify a fitted location surface by a nonzero smooth bump function whose support lies inside B , leaving its scale unchanged. The original and modified conditional densities agree for every factual (X, A) and hence have identical factual risk, while their response surfaces differ on B . $\square$

A.5 PROOF OF PROPOSITION 4

Subtract Equation (9) at a' from its value at a . Under independent event draws, the coefficients of the independent Cauchy variables are $w_j(a) - w_j(a')$, $\sigma(a)$, and $-\sigma(a')$; Cauchy stability gives Equation (17). Under a shared event realization the last coefficient becomes $\sigma(a) - \sigma(a')$. Both constructions preserve the same marginals but generally produce different contrast scales, proving non-identification. Pairwise choices can violate the finite-dimensional consistency conditions required for a stochastic process. $\square$

A.6 PROOF OF PROPOSITION 5

The bin center lies within $h/2$ of every dose in its bin, so Lipschitz continuity gives the first bound. Apply the triangle inequality to two doses to obtain the contrast bound. Piecewise-constant functions have zero derivative in the interior of every bin and are nondifferentiable at jumps. $\square$

A.7 PROOF OF PROPOSITION 6

Write $z = (y - m)/s$. Direct differentiation gives

$$\frac{\partial \ell}{\partial m} = -\frac{2z}{s(1+z^2)}, \quad \frac{\partial \ell}{\partial \log s} = \frac{1-z^2}{1+z^2}.$$

The inequalities follow from $2|z| \leq 1+z^2$ and $|1-z^2| \leq 1+z^2$. $\square$

B EXPERIMENTAL DETAILS

All networks use SiLU activations. The HCGM-Dose encoder has two 32-unit hidden layers and latent dimension six. Its dose mechanism has two 24-unit layers and uses the six-dimensional map

$$\psi(a) = \{a, a^2, \sin(\pi a), \cos(\pi a), \sin(2\pi a), \cos(2\pi a)\}.$$

Direct models use two 48-unit layers over the concatenated covariate and dose features. The five-bin baseline uses two 40-unit layers followed by five location and scale heads.

AdamW uses learning rate 2×10^{-3} , weight decay 10^{-5} , batches of 256, gradient clipping at norm 10, at most 160 epochs, and patience 25. Model initialization, minibatch order, DGP generation, split construction, and

contamination are assigned explicit deterministic seeds. Validation uses the same contamination rate as training in the outlier scenario; no clean model-selection oracle is supplied.

The evaluation tensor crosses every held-out covariate vector with 31 doses on $[0.05, 0.95]$. Neither this tensor nor any oracle surface or derivative enters training. Policy regret compares the true location at the model-selected grid dose with the true best location on the same finite grid; it does not evaluate off-grid optimization or a stochastic clinical utility.

Machine-readable per-seed rows, aggregate statistics, and the frozen configuration are stored with the confirmatory outputs. The artifact renderer checks the expected scenario–model–seed grid before producing manuscript tables and figures.

The official Simu1, IHDP, and News runs use the same pinned external model source but different published optimizer settings. Simu1 uses VCNet learning rate 10^{-4} and DRNet learning rate 0.05; IHDP uses 10^{-3} and 0.02, respectively. News uses learning rate 5×10^{-4} for both models, batches of 500, and the repository’s model-specific target-regularization weights. All use SGD with momentum 0.9, model weight decay 5×10^{-3} , functional targeted regularization, and 800 epochs. Simu1 was run on the local CPU runtime. IHDP and News used an NVIDIA GB10 for the HCGM-Dose fits; the official external models remained on CPU because their dynamic-layer code constructs CPU tensors and uses NumPy indexing. Machine-readable runtime fields are stored beside each benchmark summary.