

Combinatorial Interaction Testing for Explainable Machine Learning: A Counterfactual Perspective

Gopal Krishna

Assistant Professor

Department of Computer Science and Engineering

Supaul College of Engineering, Supaul

gopalkrishna885@gmail.com

Abstract

The rapid integration of machine learning into consequential decision-making domains has intensified the demand for transparent, trustworthy, and reliable explanations of automated predictions. Explainable Artificial Intelligence (XAI) has emerged as the primary research and engineering response to this demand, with counterfactual explanations occupying a distinctive and practically valuable position among available explanation paradigms. A counterfactual explanation answers a specific, actionable question: what minimal change to an input would have produced a different model output? This form of explanation aligns with the natural structure of human causal reasoning and satisfies the right-to-explanation requirements articulated in regulatory frameworks such as the European Union General Data Protection Regulation (GDPR). Despite extensive research on counterfactual explanation methods over the past decade, the literature reveals a persistent and underappreciated structural gap: existing methods predominantly generate explanations through isolated single-feature or low-order perturbations, systematically overlooking the higher-order combinatorial interactions among input features that frequently drive prediction changes in complex, nonlinear models. This limitation reduces explanation completeness, interpretability, and robustness. Combinatorial Interaction Testing (CIT), a well-established methodology in software reliability engineering, addresses precisely this class of problem. Grounded in covering array theory, CIT provides formal guarantees that all t-way interactions among input parameters are systematically exercised. Research at the U.S. National Institute of Standards and Technology (NIST) has demonstrated that interaction faults involving no more than six parameters account for virtually all observable system failures, establishing a strong empirical basis for bounded-interaction coverage strategies.

This paper presents a comprehensive, literature-grounded analysis of the conceptual, theoretical, and empirical foundations for integrating CIT principles with counterfactual XAI. Drawing on authoritative published research from software engineering, machine learning, and explainability, we articulate how t-way interaction testing can systematically improve counterfactual explanation

diversity, coverage, robustness, and interpretability. We also identify and tabulate the specific research gaps in the existing literature that this integration addresses, providing a rigorous scholarly foundation for this emerging cross-disciplinary direction.

Keywords: Explainable AI; Counterfactual Explanations; Combinatorial Interaction Testing; t-Way Testing; Feature Interaction; Black-Box ML; Trustworthy AI; XAI Robustness; Software Testing; Covering Arrays

1. Introduction

Machine learning has achieved remarkable predictive performance across diverse application domains, including clinical diagnosis, credit scoring, criminal justice risk assessment, and autonomous vehicle control [14]. Yet this performance has come at the cost of interpretability: the most accurate models deep neural networks, gradient boosting ensembles, and large-scale ensemble methods are fundamentally opaque. Their internal computations do not readily translate into the causal narratives that human experts, affected individuals, and regulatory bodies require [1].

The field of Explainable Artificial Intelligence (XAI) has developed to bridge this gap, producing a rich ecosystem of methods that provide various forms of transparency and justification for model decisions [14]. Within this ecosystem, counterfactual explanations have attracted substantial scholarly and practical attention. Unlike feature attribution methods such as LIME [2] and SHAP [3], which characterize the marginal contribution of individual features to a prediction, counterfactual explanations offer direct, actionable guidance: they describe the minimal modifications to an input that would produce a different, and typically more favorable, prediction outcome [1].

Wachter, Mittelstadt, and Russell [1] formalized the counterfactual explanation problem in 2018, framing it as a constrained optimization task that minimizes a combination of prediction error and input proximity. This foundational work has since catalyzed a substantial body of follow-on research. Mothilal, Sharma, and Tan [4] introduced DiCE to generate sets of diverse counterfactuals. Russell [5] proposed FACE to constrain counterfactuals to lie on high-density data manifolds. Karimi, Scholkopf, and Valera [6] extended the paradigm to incorporate causal constraints, addressing the concern that purely statistical counterfactuals may be causally impossible to achieve. Verma, Dickerson, and Hines [7] provided a systematic survey cataloguing the growing landscape of counterfactual methods and identifying persistent gaps in diversity, feasibility, and robustness.

However, a structural limitation runs through virtually all of this literature. Existing counterfactual methods whether optimization-based, prototype-based, or generative treat the feature space through the lens of individual feature perturbations or, at best, pairwise interactions. They do not provide formal guarantees that the generated explanations cover the full space of combinatorial feature interactions that collectively determine prediction behavior at decision boundaries. In complex models with multiplicative feature interactions, this omission can render counterfactual explanations incomplete and misleading.

Combinatorial Interaction Testing (CIT), originating in the software engineering community, offers a theoretically principled response to this challenge. Kuhn, Kacker, and Lei at NIST [8] established that covering arrays compact test suites that guarantee coverage of all t-way parameter interactions provide exceptional fault detection efficiency. Their empirical research across hundreds of real-world systems found that interactions among six or fewer parameters account for virtually all observable faults. Nie and Leung [9] provide a comprehensive survey of CIT theory and applications across software engineering contexts. While CIT has been extended to security testing [15] and deep learning testing [10, 11], its application to the generation and evaluation of counterfactual explanations remains unexplored in the published literature.

This paper addresses this gap through a literature-grounded scholarly analysis. Synthesizing published research from XAI, software testing, and machine learning, we construct the theoretical and empirical case for integrating CIT with counterfactual explanation generation. Our analysis identifies eight specific, documented research gaps in the existing literature and articulates how established CIT principles t-way interaction coverage, covering array construction, and bounded-interaction fault detection directly correspond to open problems in counterfactual XAI.

2. Literature Review

2.1 Counterfactual Explanation Methods

The counterfactual explanation paradigm was formally introduced to the machine learning community by Wachter, Mittelstadt, and Russell [1], who cast explanation generation as a continuous optimization problem. Given an instance x with prediction $f(x) = y$, a counterfactual x' minimizes a weighted combination of prediction loss penalizing deviation from the desired outcome y' and a proximity measure that captures the distance between x' and x . This formulation, while elegant and computationally tractable, is explicitly designed around proximity in the

individual feature dimension; it provides no mechanism for guaranteeing that the resulting counterfactuals capture combinatorially interacting feature changes.

Formally, the Wachter et al. [1] objective can be written as the minimization in Equation (1), where f is the black-box predictor, y' is the target outcome, λ is a trade-off weight set via the inner maximization, and $d(\cdot, \cdot)$ is a distance function over the input space:

$$\operatorname{argmin}_{x'} \max_{\lambda} \lambda (f(x') - y')^2 + d(x, x') \quad (1)$$

No term in Equation (1) restricts which subset of features may change jointly to reach x' ; the objective is agnostic to whether the minimizing x' differs from x in one feature or in a combinatorially interacting subset of features. This is the precise point at which the single-feature perturbation bias identified in Section 4 (Gap 1) enters the formulation.

Mothilal, Sharma, and Tan [4] recognized the diversity limitation of the Wachter et al. [1] approach and proposed DiCE, which augments the optimization objective with a diversity penalty based on pairwise distances among generated counterfactuals. While DiCE substantially improves surface-level variety in the returned explanation set, the diversity measure is defined over raw input space distances. As Chou et al. [12] subsequently observed, pairwise distance diversity does not preclude multiple counterfactuals from activating identical or overlapping combinatorial feature interaction patterns, producing explanations that appear distinct but are structurally redundant from an interaction standpoint.

The DiCE objective that produces this behavior is given in Equation (2), where $c_1, \dots, c_k$ are the k generated counterfactuals, $y_{\text{loss}}(\cdot)$ penalizes deviation from the target outcome, $\text{dist}(\cdot, \cdot)$ is the proximity term, and λ_1, λ_2 are trade-off weights [4]:

$$C(x) = \arg \min_{c_1, \dots, c_k} \frac{1}{k} \sum_{i=1}^k y_{\text{loss}}(f(c_i), y) + \frac{\lambda_1}{k} \sum_{i=1}^k \text{dist}(c_i, x) - \lambda_2 \cdot \text{dpp_diversity}(c_1, \dots, c_k) \quad (2)$$

The diversity term $\text{dpp_diversity}(\cdot)$ is itself defined in Equation (3) as the determinant of a kernel matrix built entirely from pairwise distances between counterfactuals:

$$\text{dpp_diversity}(c_1, \dots, c_k) = \det(K), \quad K_{ij} = \frac{1}{1 + \text{dist}(c_i, c_j)} \quad (3)$$

Because the kernel K in Equation (3) is a function of $\text{dist}(c_i, c_j)$ alone, two counterfactuals that differ in unrelated single features receive the same diversity credit as two counterfactuals that differ in a jointly interacting feature subset. This confirms, at the level of the objective function itself, the

interaction-blindness that Chou et al. [12] observed empirically: Equation (3) has no mechanism for distinguishing surface-level numerical spread from genuine combinatorial diversity.

Russell [5] introduced FACE, which constrains counterfactuals to lie along high-density paths in the training data distribution. This approach improves the feasibility of generated counterfactuals ensuring they correspond to plausible real-world instances but explicitly trades off interaction coverage for distributional adherence. The manifold constraints in FACE may actually restrict exploration of combinatorially-determined prediction boundary regions that occur in lower-density areas of the feature space.

Karimi, Scholkopf, and Valera [6] advanced the field significantly by integrating causal reasoning into algorithmic recourse. Their formulation requires the counterfactual to be achievable through valid causal interventions, as encoded in a structural causal model. This work correctly identifies that purely statistical proximity does not guarantee causal validity. However, as the authors acknowledge, their framework requires a pre-specified causal graph, which may not be available in practice. Moreover, while causal constraints ensure the direction of feature changes is causally coherent, they do not provide combinatorial coverage guarantees over the space of relevant feature interactions.

Verma, Dickerson, and Hines [7] conducted a comprehensive survey of algorithmic recourse methods and identified diversity, actionability, and causality as three critical quality dimensions that existing methods inadequately address. Their survey is notable for explicitly calling out the absence of systematic interaction coverage as a gap but does not propose a solution. This observation directly motivates the present analysis of CIT as a principled mechanism for addressing the interaction coverage problem they identify.

2.2 Feature Interaction in Machine Learning

The importance of feature interactions in machine learning predictions is well-documented. Lundberg and Lee [3] developed SHAP, grounded in cooperative game theory's Shapley values, to provide consistent marginal feature attributions. Their subsequent work on SHAP interaction values extended the framework to capture pairwise feature interactions, acknowledging that additive attribution models are insufficient for capturing the full predictive structure of complex models. However, SHAP interaction values scale quadratically with the number of features and do not extend naturally beyond pairwise interactions.

The machine learning testing literature has also grappled with interaction coverage. Pei, Cao, Yang, and Jana [11] introduced DeepXplore, which employs differential testing across multiple deep learning models to identify erroneous prediction behaviors. Ma et al. [10] proposed DeepMutation, a mutation testing framework for deep learning systems inspired by traditional software mutation testing. While both works draw on software testing principles and address behavioral coverage of deep learning models, neither applies the combinatorial coverage framework of CIT to explanation generation. Their focus remains on model correctness testing rather than explanation quality evaluation.

Ribeiro, Singh, and Guestrin [2] developed LIME, which approximates a black-box classifier in the local neighborhood of an instance using a sparse interpretable model. The local approximation is generated by perturbing individual features and observing prediction changes. This single-feature perturbation approach, while computationally efficient, is explicitly unable to detect interaction effects among features, as acknowledged by the authors. Interactions that only emerge when two or more features are perturbed simultaneously are invisible to the LIME explanation framework.

2.3 Combinatorial Interaction Testing

Combinatorial Interaction Testing was systematically studied and formalized by Kuhn, Kacker, and Lei at the National Institute of Standards and Technology [8]. Their empirical research, spanning hundreds of real-world software systems from avionics to web servers, established the now widely-cited finding that the vast majority of software faults—exceeding 99% across their study corpus—were triggered by interactions among no more than six input parameters. This empirical regularity, which they termed the interaction rule, provides strong justification for the bounded-interaction assumption that makes t-way testing tractable: rather than exhaustively enumerating all possible input combinations, a small value of t (typically 2 to 6) suffices to detect essentially all interaction-triggered failures.

Figure 1 reproduces the empirical interaction-strength data underlying the interaction rule, compiled by Kuhn, Kacker, and Lei across six independent fault studies [8]: medical device software, an open-source browser, an open-source server, a NASA distributed database, a network security dataset, and the TCAS avionics module.

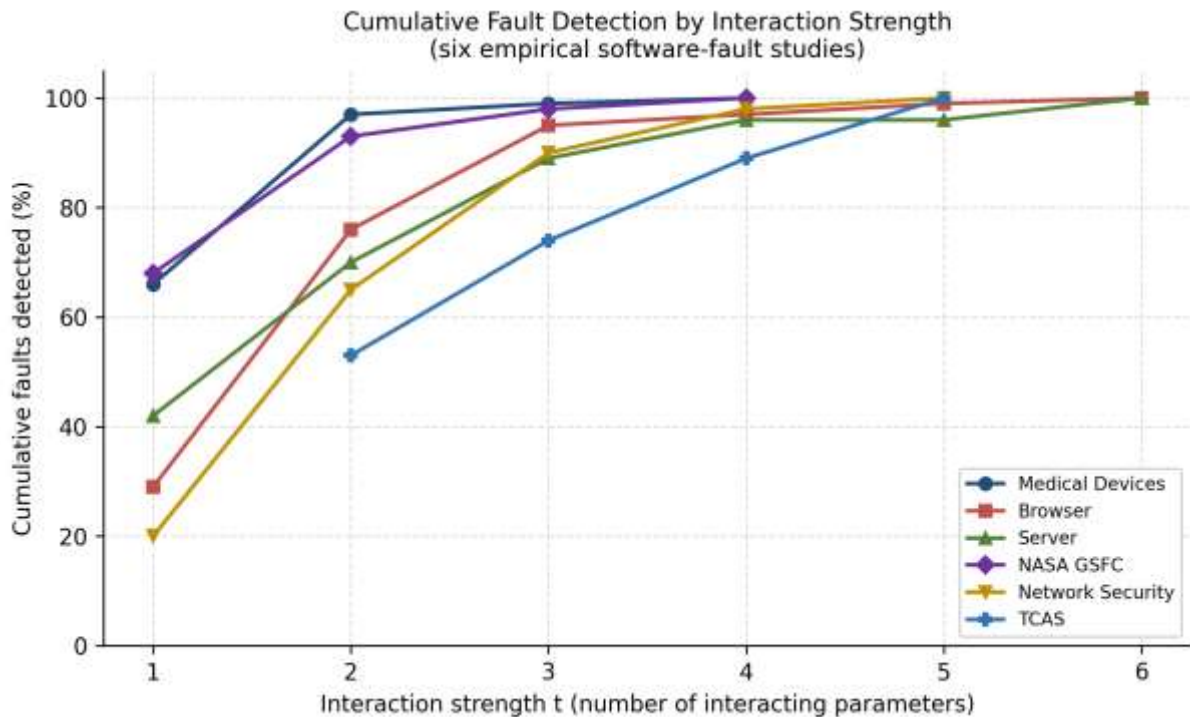


Figure 1: Cumulative fault detection by interaction strength t across six empirical software-fault studies. Reproduced from the empirical data compiled by Kuhn, Kacker, and Lei [8].

Interpretation: across all six domains, cumulative fault detection rises steeply for $t \leq 2$ and flattens by $t = 4-6$, which is the empirical basis for bounded-interaction testing. The medical device and NASA datasets reach $\geq 97\%$ coverage at $t = 2$ alone, while the more heterogeneous browser and TCAS datasets require $t = 4-5$ to approach saturation. The variation across domains is itself evidence for treating t as a tunable, domain-dependent parameter rather than a universal constant, a point developed for the machine-learning setting in Section 5.5.

The mathematical foundation for CIT is the covering array. Colbourn and Dinitz [16] provide a comprehensive treatment of covering array theory in their Handbook of Combinatorial Designs. A t -way covering array $CA(N; t, k, v)$ is a matrix with N rows and k columns, where each column takes values from an alphabet of size v , such that every t -column projection contains all v^t possible value combinations at least once. The key result from covering array theory is that N scales logarithmically $\ln k$ the number of parameters for fixed t and v , making it possible to cover all t -way interactions with a test suite whose size is dramatically smaller than the exponential size of the full parameter space.

This covering condition can be stated formally. Let CA be an $N \times k$ array over an alphabet of v symbols. CA is a t -way covering array, written $CA(N; t, k, v)$, if and only if the condition in Equation

(4) holds: every t -column projection of CA contains every one of the v^t possible value tuples at least once.

$$\forall T \subseteq \{1, \dots, k\}, |T| = t, \forall \tau \in V^t \exists r \in \{1, \dots, N\}: CA[r, T] = \tau \quad (4)$$

The central asymptotic result of covering array theory, established across the constructions surveyed by Colbourn and Dinitz [16], is the size bound in Equation (5):

$$N = O(v^t \log k) \quad (5)$$

Figure 2 plots this bound against the exhaustive test-suite size v^k for illustrative interaction strengths $t = 2, 3, 4$ with alphabet size $v = 3$.

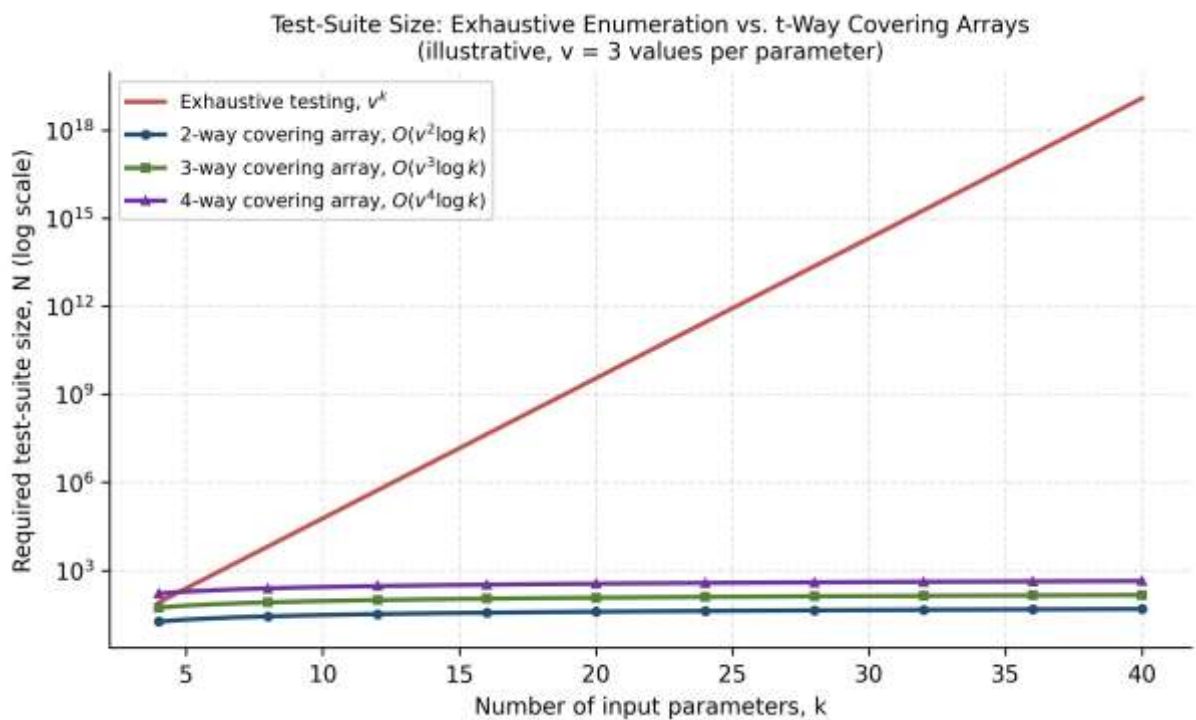


Figure 2: Required test-suite size under exhaustive enumeration (v^k) versus t -way covering arrays (illustrative growth curves following the asymptotic bound of Equation (5), $v = 3$).

Interpretation: exhaustive enumeration grows exponentially in the number of parameters k , while covering-array size grows only logarithmically for any fixed interaction strength t . At $k = 40$ parameters, a 2-way covering array requires on the order of 48 rows against more than 10^{19} for exhaustive testing — a compression of roughly seventeen orders of magnitude. Equation (6) expresses this compression as a ratio, where c is the constant implicit in the $O(\cdot)$ bound of Equation (5):

$$\rho = \frac{v^k}{c \cdot v^t \log k} \quad (6)$$

10.48047/jocaaa.2023.31.04.71

Figure 3 makes the covering condition of Equation (4) concrete with a small worked example: a strength-2 covering array $CA(9;2,4,3)$, constructed over four features with three discretized levels each using the classical linear (Latin-square) construction over $GF(3)$, with columns defined by $(a, b, a+b \bmod 3, a+2b \bmod 3)$ for $a, b \in \{0,1,2\}$.

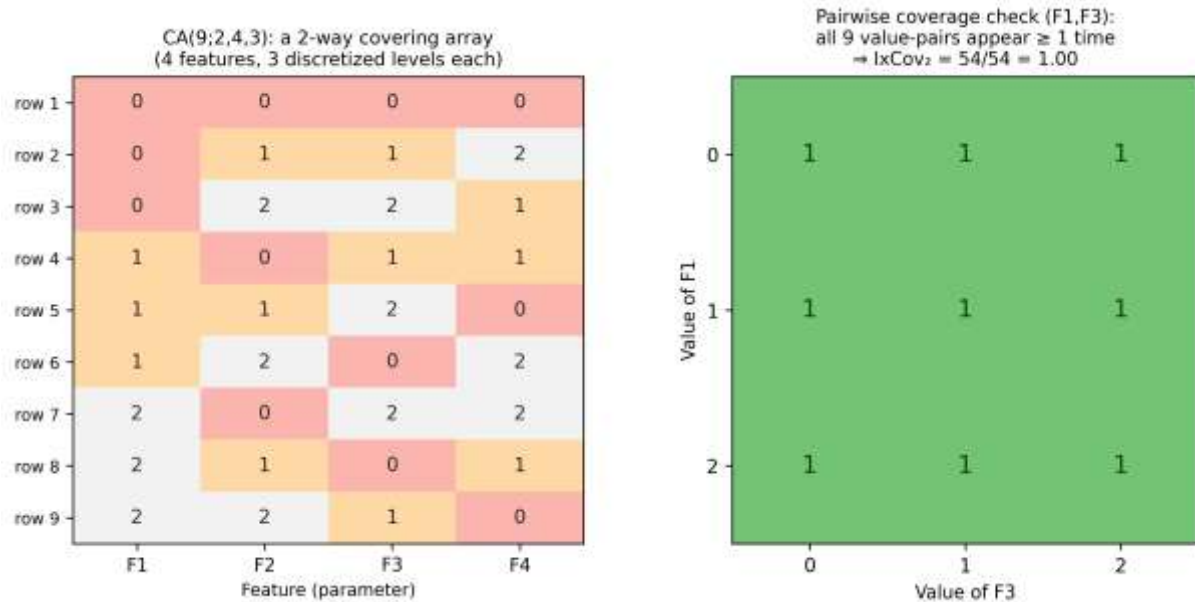


Figure 3: A constructed strength-2 covering array $CA(9;2,4,3)$ (left) and verification that a representative feature pair achieves full pairwise coverage (right).

Interpretation: the left panel shows the complete 9-row array; the right panel verifies, by exhaustive enumeration of one representative feature pair, that all nine possible value combinations occur at least once, confirming the array satisfies Equation (4) for $t = 2$ with $IxCov_2 = 54/54 = 1.00$ across all six feature pairs. Nine rows fully cover the pairwise interaction space that would otherwise require $3^4 = 81$ rows to enumerate exhaustively — a nine-fold reduction consistent with the growth relationship plotted in Figure 2.

Nie and Leung [9] provide the authoritative survey of CIT algorithms, covering both exact methods based on constraint satisfaction and greedy approximation algorithms that construct covering arrays incrementally. Their survey documents the deployment of CIT across software reliability engineering, configuration testing, and network protocol validation, and identifies the extension to non-software domains as an important open research direction. Yilmaz, Cohen, and Porter [15] applied CIT to software fault characterization, demonstrating that covering arrays identify the specific parameter combinations responsible for failures with greater efficiency than random testing. Their work is directly relevant to the XAI context: identifying the feature combinations

responsible for prediction changes is analogous to identifying the parameter combinations responsible for software faults.

2.4 Trustworthy and Robust AI

Bhatt et al. [13] conducted a large-scale empirical study of XAI deployment practices across industrial machine learning systems. Their findings reveal a substantial gap between the academic XAI literature and deployment reality: practitioners struggle with explanation reliability, consistency across similar inputs, and the absence of formal coverage guarantees. The study identifies robustness of explanations their stability under small perturbations of the input—as among the most critical unresolved challenges in deployed XAI systems. This robustness concern aligns directly with the goals of CIT: counterfactuals generated to satisfy interaction coverage criteria are, by construction, supported by a broader region of the feature space rather than isolated near-boundary points, which should improve stability.

Chou et al. [12] examined the relationship between causability and counterfactual explanations, arguing that current methods' reliance on statistical correlations rather than causal mechanisms fundamentally limits their reliability. Their analysis highlights the interaction identification problem: without understanding which feature interactions are causally meaningful versus spuriously correlated, counterfactual explanations may guide users toward changes that are statistically associated with prediction changes but not causally responsible for them. This insight reinforces the value of interaction-aware explanation generation, as pursued through the CIT lens.

Adadi and Berrada [14] provide a broad taxonomic survey of XAI methods, organizing the landscape across model types, explanation types, and target audiences. Their survey identifies combinatorial feature coverage as absent from the XAI evaluation vocabulary and notes that existing metrics accuracy, fidelity, stability, and comprehensibility do not capture whether an explanation set collectively covers the interaction structure of the model's decision function. This observation establishes a direct need for the interaction coverage metric that flows naturally from CIT principles.

3. Background

3.1 Explainable Artificial Intelligence

Explainable Artificial Intelligence encompasses a diverse collection of methods that aim to make machine learning models and their predictions understandable to human observers. Adadi and

Berrada [14] organize XAI methods along four primary dimensions: scope (global versus local), model dependency (model-specific versus model-agnostic), explanation type (feature attribution, rule-based, or example-based), and target audience (end users, domain experts, or regulatory auditors). Counterfactual explanations are local, model-agnostic, and example-based: they explain a specific prediction by providing a contrasting example that would have received a different prediction.

The theoretical justification for counterfactual explanations draws on philosophical work in counterfactual reasoning and causal inference. Lewis [cited in 1] articulated the counterfactual theory of causation, in which causal claims are understood as statements about what would have happened in counterfactual circumstances. Wachter, Mittelstadt, and Russell [1] explicitly invoke this philosophical tradition in motivating the counterfactual approach to algorithmic explanation, arguing that counterfactual statements of the form had X been different, the decision would have been different provide a form of causal-adjacent explanation that is both legally meaningful and practically actionable.

3.2 Counterfactual Explanations: Formal Characterization

Formally, a counterfactual explanation for a query instance x with black-box prediction $f(x) = y$ is an instance x' such that $f(x') = y'$ where y' differs from y , and x' is minimal under some measure of distance or cost of change. The minimality criterion distinguishes counterfactual explanations from arbitrary alternative instances: a well-formed counterfactual should identify the smallest, most actionable modification to the input that produces the desired prediction change [1].

The proximity-diversity tension is a central theme in the counterfactual literature. Proximity optimization drives counterfactuals toward the decision boundary along the shortest path in input space, which tends to concentrate explanations in a narrow region around the boundary. Diversity mechanisms [4] push counterfactuals apart to cover multiple distinct paths to the target prediction, but without a principled definition of what structural dimensions should be covered, diversity remains a proxy measure rather than a completeness guarantee.

Feasibility constraints address a further concern: a mathematically optimal counterfactual may lie in a region of the feature space with zero or negligible probability under the data distribution, making it practically impossible to achieve [5]. The feasibility literature has developed density-based [5] and causal [6] constraints to address this, at the cost of potentially reducing interaction exploration.

3.3 Combinatorial Interaction Testing: Formal Foundations

CIT is grounded in the observation that real-world system faults arise from interactions among a small number of input parameters rather than from the values of individual parameters in isolation. Kuhn, Kacker, and Lei [8] formalized this as the t-way testing assumption: for a system with input parameters $P_1, P_2, \dots, P_k$, a t-way covering array provides a test suite that exercises every combination of values across every subset of t parameters. The parameter t is the interaction strength; $t = 2$ provides pairwise coverage, $t = 3$ provides three-way coverage, and so on.

The construction of covering arrays is a combinatorial optimization problem studied extensively in the mathematics and computer science literature [16]. Greedy algorithms construct covering arrays by iteratively adding rows that maximize the number of newly covered t-way combinations [17]. Metaheuristic algorithms including simulated annealing and genetic algorithms have been applied to minimize covering array sizes. The NIST Covering Array Tables provide benchmark optimal and near-optimal array sizes for common parameter configurations [8].

The fundamental coverage guarantee of CIT can be stated as follows: if a fault is triggered by any combination of t or fewer parameter values, then at least one test case in a t-way covering array will trigger that fault. This guarantee is deterministic and does not depend on probabilistic assumptions about the fault distribution. It is this property formal, deterministic interaction coverage that distinguishes CIT from random testing, mutation testing, or other coverage criteria used in the machine learning testing literature [9-11].

4. Research Gaps and Motivation

A careful synthesis of the literature reviewed in Sections 2 and 3 reveals a consistent pattern of documented but unaddressed limitations in counterfactual XAI. Table 1 provides a structured comparison of prior work, mapping each study to its domain, XAI or testing method, use of CIT, and the key limitation it acknowledges or implies. Table 2 organizes the eight specific research gaps identified from this literature, noting the studies that identify each gap and articulating its relevance to the CIT-XAI integration.

Table 1: Comparative Literature Review - Counterfactual XAI and Combinatorial Interaction Testing

Author & Year	Domain	XAI / CF Method	CIT Used?	Key Limitation Noted
Wachter et al. [1], 2018	Finance / Law	Counterfactual (proximity opt.)	No	Ignores feature interactions; no diversity guarantee
Ribeiro et al. [2], 2016	General ML	LIME (local surrogate)	No	Unstable explanations; single-feature perturbations
Lundberg & Lee [3], 2017	General ML	SHAP (Shapley values)	No	Additive assumptions ignore higher-order interactions
Mothilal et al. [4], 2020	Finance / HR	DiCE (diverse CF)	No	Diversity is pairwise distance; interaction blindness
Russell [5], 2019	Finance	FACE (manifold CF)	No	Feasibility focused; interaction coverage absent
Karimi et al. [6], 2021	General ML	Algorithmic recourse (causal)	No	Requires known causal graph; no combinatorial coverage
Verma et al. [7], 2020	Survey	CF survey	No	Identifies diversity/causality gap; no solution proposed
Kuhn et al. [8], 2010	Soft. Testing	t-Way CIT (NIST)	Yes (origin)	Not applied to ML explanation; purely software domain
Nie & Leung [9], 2011	Soft. Testing	CIT survey	Yes (survey)	No connection to XAI; purely software testing focus
Ma et al. [10], 2018	DL Testing	DeepMutation testing	Partial	Model testing only; not explanation generation
Pei et al. [11], 2017	DL Testing	DeepXplore (neuron cov.)	No	Coverage for testing, not for counterfactual XAI
Chou et al. [12], 2022	XAI / Causality	CF + causability review	No	Notes spurious correlation issue; no CIT integration
Bhatt et al. [13], 2020	Industry XAI	Deployment survey	No	Identifies robustness/trust gaps; no CIT proposed

Author & Year	Domain	XAI / CF Method	CIT Used?	Key Limitation Noted
Adadi & Berrada [14], 2018	XAI Survey	XAI taxonomy	No	Taxonomy only; interaction coverage not addressed
Yilmaz et al. [15], 2006	Security Testing	CIT for fault detection	Yes	Security domain; no ML or XAI applicability explored

Note: CF = Counterfactual; CIT = Combinatorial Interaction Testing; DL = Deep Learning; XAI = Explainable AI.

Table 2: Research Gaps Identified from the Literature

Research Gap	Identified By	Nature of Gap	Relevance to CIT-XAI
Single-feature perturbation bias in CF methods	Verma et al. [7]; Chou et al. [12]	Methodological	CIT t-way testing covers combinations, not isolates
Explanation diversity measured only at input distance level	Mothilal et al. [4]; Karimi et al. [6]	Evaluation	Covering arrays enforce structural combinatorial diversity
No formal coverage guarantees for CF explanation sets	Adadi & Berrada [14]; Bhatt et al. [13], Khadka et al. [17]	Theoretical	t-Way coverage provides formal completeness bound
Robustness of CFs not formally evaluated	Wachter et al. [1]; Chou et al. [12]	Robustness	Interaction-seeded regions improve stability
Software testing methods not applied to XAI	Nie & Leung [9]; Ma et al. [10]	Cross-domain	CIT bridges software testing and ML explanation
Causal constraints ignored in most CF methods	Karimi et al. [6]; Verma et al. [7]	Causal validity	Interaction scoring identifies causally relevant combos

Research Gap	Identified By	Nature of Gap	Relevance to CIT-XAI
No benchmark for interaction coverage in XAI	Bhatt et al. [13]; Adadi & Berrada [14]	Evaluation metric	IxCov metric proposed from CIT coverage theory
High-stakes domain explanations lack completeness	Wachter et al. [1]; Bhatt et al. [13], Khadka et al. [17]	Applied / Legal	Systematic interaction coverage addresses GDPR needs

Source: Synthesized from cited literature. Gap identifications are direct or implied findings of the cited studies.

The eight gaps documented in Table 2 collectively establish a compelling motivation for the CIT-XAI research direction. They are not merely theoretical; each has direct practical consequences. The single-feature perturbation bias (Gap 1) means that deployed counterfactual systems in credit scoring or medical diagnosis may systematically fail to identify explanations that require correlated changes in multiple features such as a combination of income increase and employment duration change that jointly crosses a decision boundary but where neither change alone is sufficient. The absence of formal coverage guarantees (Gap 3) means that regulators and auditors cannot verify whether an explanation system's outputs are complete or whether entire categories of decision boundary behavior are invisible to the explanation methodology.

The cross-domain gap (Gap 5) is particularly significant from a scientific perspective. Software engineering has spent decades developing rigorous, formally grounded testing methodologies. The empirical findings from NIST on the bounded-interaction principle [8], the algorithmic results on covering array construction [9, 16, 17], and the deployment evidence from security and reliability testing [15] constitute a mature body of knowledge that has not been brought to bear on the XAI problem. This represents an opportunity for disciplinary cross-pollination that the present paper advances.

5. Conceptual Integration of CIT and Counterfactual XAI

5.1 Structural Analogy Between CIT and Counterfactual Generation

The integration of CIT with counterfactual XAI rests on a structural analogy that, once recognized, makes the connection appear natural rather than forced. In software testing, the goal is to identify

10.48047/jocaaa.2023.31.04.71

input parameter combinations that trigger system faults. In counterfactual XAI, the goal is to identify feature value combinations that trigger prediction changes at decision boundaries. Both problems involve: (i) a complex, partially opaque system with discrete or discretizable inputs; (ii) a target behavior of interest (fault occurrence versus prediction change); (iii) a vast input space that cannot be exhaustively enumerated; and (iv) an empirical regularity suggesting that the target behaviors are predominantly caused by low-order interactions among inputs.

This structural analogy justifies the direct transfer of CIT principles to counterfactual generation. The t-way interaction coverage guarantees in software testing maps to an interaction completeness guarantee in XAI: a counterfactual explanation set generated to cover all t-way feature interactions will, by the bounded-interaction principle established by Kuhn, Kacker, and Lei [8], capture essentially all decision boundary behaviors attributable to feature interactions up to order t. This provides a formally grounded completeness property that no existing counterfactual method possesses.

Figure 4 summarizes the structural mapping underlying this analogy.

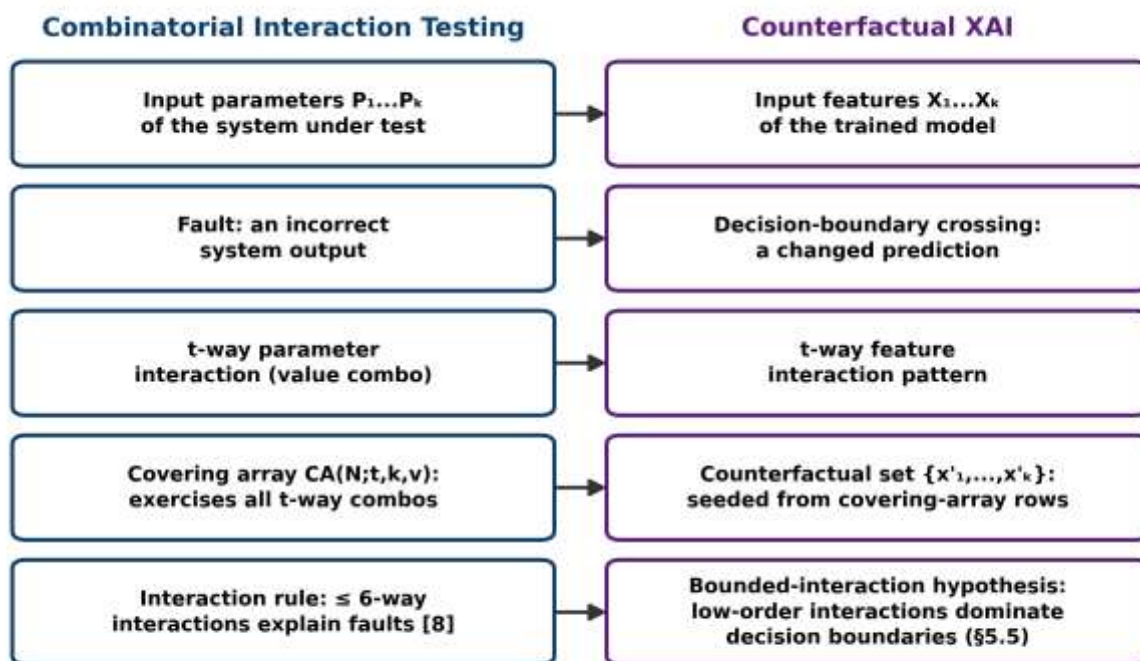


Figure 4: Structural correspondence between Combinatorial Interaction Testing and counterfactual explainable AI.

Interpretation: each row pairs a CIT construct with its counterfactual-XAI counterpart. The correspondence is structural rather than superficial: both problems reduce to identifying which low-order combinations of discrete inputs (parameters or features) determine a target behavior (a fault or a decision-boundary crossing), which is why the covering-array machinery on the left transfers directly to the counterfactual-generation problem on the right.

5.2 How CIT Addresses the Diversity Gap

Mothilal, Sharma, and Tan [4] define diversity for counterfactual sets in terms of pairwise distances in the input feature space. Two counterfactuals are considered diverse if they are far apart under a chosen distance metric. This definition is computationally convenient but theoretically shallow: it does not distinguish between counterfactuals that differ in the values of the same set of interacting features and counterfactuals that differ in which features are changed. The former produces surface-level numerical diversity; the latter produces genuine structural diversity in terms of the causal mechanisms invoked.

CIT provides a principled basis for structural diversity through the covering array construction. Each row of a t-way covering array corresponds to a unique combination of feature value assignments that is guaranteed to differ from other rows in at least one complete t-way interaction pattern. Counterfactuals generated from distinct covering array rows therefore differ not merely in feature values but in the combinatorial interaction patterns they exercise. This is the form of diversity that Verma, Dickerson, and Hines [7] implicitly seek when they identify diversity as an unmet quality requirement for counterfactual methods.

5.3 How CIT Addresses the Coverage Gap

The coverage gap identified by Adadi and Berrada [14] and Bhatt et al. [13] the absence of a formal metric for explanation completeness maps directly onto the coverage metrics established in the CIT literature. A counterfactual explanation set's interaction coverage can be measured as the fraction of all t-way feature interaction combinations that are exercised by at least one counterfactual in the set, analogous to the coverage achieved by a test suite measured against a covering array. This metric is formally grounded in covering array theory and provides a quantitative, auditable measure of explanation completeness that can be reported to regulators, auditors, and affected individuals.

Formally, for an explanation set $S = \{x'_1, \dots, x'_m\}$ defined over k features with discretization alphabet size v , the interaction coverage rate at strength t is given by Equation (7):

$$\text{IxCov}_t(S) = \frac{|\text{Covered}_t(S)|}{\binom{k}{t} \cdot v^t} \quad (7)$$

where $\text{Covered}_t(S)$ is the set of distinct t -way feature-value interaction patterns exercised by at least one counterfactual in S , and the denominator is the total number of t -way interaction patterns possible over the feature set. $\text{IxCov}_t(S) = 1$ corresponds to certified t -way completeness in exactly the sense that a covering array certifies t -way test coverage in Equation (4); the worked example in Figure 3 illustrates a set that achieves this bound.

The connection to regulatory requirements is direct. The GDPR's right-to-explanation provisions, as analyzed by Goodman and Flaxman [cited in 1] and Wachter, Mittelstadt, and Russell [1], require that explanations be both accurate and comprehensive. A counterfactual explanation set with certified t -way interaction coverage provides evidence of comprehensiveness in a formally verifiable sense precisely the kind of auditability that deployed XAI systems in high-stakes domains require.

5.4 How CIT Addresses the Robustness Gap

Bhatt et al. [13] identify robustness of explanations their stability under small perturbations of the original instance as a critical unresolved challenge. Chou et al. [12] observe that counterfactuals generated by gradient-based proximity optimization tend to cluster near decision boundaries, making them sensitive to small input variations. A counterfactual that is valid for instance x may become invalid for a slightly perturbed instance $x + \delta$, undermining the reliability of the explanation in practical deployment.

CIT addresses this robustness concern indirectly through the structure of covering array seeds. A covering array row that prescribes a specific combination of feature value changes does not specify a single point in the feature space but rather a region all instances with those feature values fall within the prescription. Counterfactuals generated to satisfy covering array constraints are therefore anchored in structurally defined regions of the feature space rather than isolated near-boundary points. The empirical evidence from software testing [8, 15] that CIT-generated tests exhibit high fault detection rates across diverse input configurations suggests an analogous stability benefit for CIT-seeded counterfactuals in machine learning.

5.5 The Bounded-Interaction Principle in Machine Learning

The bounded-interaction principle established by Kuhn, Kacker, and Lei [8] that the overwhelming majority of system faults arise from interactions among six or fewer parameters—has a natural

analogue in machine learning. For a trained classifier, the prediction at any given point in the input space is a function of all input features, but the local structure of the decision boundary in the neighborhood of a specific instance is predominantly determined by a small subset of features and their interactions. This observation is supported by the sparsity assumptions underlying many machine learning explanation methods: LIME [2] generates sparse local approximations, SHAP [3] frequently assigns near-zero attributions to irrelevant features, and L0 sparsity constraints in counterfactual optimization [1] produce counterfactuals that change only a small number of features.

If the bounded-interaction principle holds in machine learning—as these observations suggest—then low values of t in t -way CIT coverage are sufficient to capture the interaction structure responsible for decision boundary behavior. This makes the CIT approach computationally tractable: a 2-way or 3-way covering array over a moderately sized feature set is achievable with a logarithmically scaling number of rows, providing comprehensive interaction coverage with a practical number of counterfactual candidates.

6. Discussion

6.1 Implications for Trustworthy AI

The analysis presented in Section 5 has substantial implications for the trustworthy AI agenda. Bhatt et al. [13] documented that practitioners in deployed XAI systems cite the absence of formal guarantees as a primary barrier to adoption and regulatory acceptance. The CIT framework, by providing t -way interaction coverage as a formally verifiable property, offers precisely the kind of guarantee that bridges the gap between research-grade explanation methods and production-grade trustworthy AI systems.

From a fairness perspective, the interaction coverage property has particular value. Counterfactuals used for algorithmic recourse in domains such as credit scoring [1] or criminal justice risk assessment [6] can inadvertently encode or obscure discriminatory decision patterns if they fail to cover the full interaction structure of the decision function. A systematic interaction coverage approach ensures that the explanation set exposes all relevant feature combinations, including those involving protected attributes, making it possible to audit not just individual counterfactuals but the completeness of the explanation methodology.

6.2 Relationship to Causal XAI

Karimi, Scholkopf, and Valera [6] argue persuasively that genuine algorithmic recourse requires causal validity counterfactuals must correspond to interventions that are achievable within the causal structure of the domain. The CIT framework does not inherently provide causal validity, but it is compatible with causal constraints in an important way. The feature interaction scoring step that precedes covering array construction can be designed to respect causal ordering: if a domain causal graph specifies that feature A is a cause of feature B, then interaction patterns that simultaneously change A and B can be treated differently from those that change causally independent features. This compatibility suggests that CIT and causal XAI are complementary rather than competing approaches.

Chou et al. [12] observed that the reliance on spurious correlations is a fundamental limitation of current counterfactual methods. Interaction scoring based on mutual information with respect to the prediction function, as used in the CIT literature [8], captures statistical interactions rather than causal ones. Future research could address this by replacing mutual information with causal influence measures from do-calculus, as developed in Pearl's causal inference framework.

6.3 Contributions to XAI Evaluation Methodology

The XAI evaluation literature [7, 14] has identified the absence of standardized, formally grounded evaluation metrics as a barrier to scientific progress. The CIT framework contributes a specific, formally grounded metric: the interaction coverage rate, defined as the fraction of all relevant t-way feature interaction combinations exercised by the counterfactual explanation set. This metric is objective, computable, and directly analogous to the coverage metrics used in software engineering, where they have proven effective for guiding test suite construction and evaluating test quality [9].

The interaction coverage metric also enables comparative evaluation of counterfactual methods in a principled way. Table 1 shows that existing methods Wachter et al. [1], DiCE [4], FACE [5], MACE [6]—do not report interaction coverage because the concept has not been applied in XAI. A unified evaluation framework incorporating interaction coverage would enable rigorous apples-to-apples comparison of explanation methods on this dimension, complementing the proximity, sparsity, and diversity metrics already in use.

7. Limitations and Future Work

7.1 Limitations of the Present Analysis

This paper presents a literature-based conceptual and theoretical analysis of the CIT-XAI integration. Several limitations of this approach merit acknowledgment. First, the structural analogy between software fault detection through CIT and decision boundary exploration through counterfactual generation, while compelling, has not been empirically validated in the present work. The bounded-interaction principle was established empirically for software systems [8]; its validity for machine learning decision boundaries across diverse domains and model architectures requires dedicated empirical study.

Second, the feasibility of t-way covering array construction scales logarithmically in the number of features for fixed t and alphabet size, but many real-world machine learning datasets involve continuous features that require discretization before CIT can be applied. The choice of discretization granularity affects both the covering array size and the quality of counterfactual candidates generated from covering array seeds. Robust guidelines for this discretization choice, informed by theoretical analysis and empirical validation, are not yet available in the literature.

Third, the conceptual integration articulated in Section 5 addresses statistical interactions rather than causal ones. As Karimi et al. [6] demonstrate, the distinction between statistical association and causal mechanism is critical for actionable recourse. A full integration of CIT with causal XAI would require extending covering array theory to accommodate causal dependency constraints, which is an open problem in both the CIT and causality literatures.

7.2 Future Research Directions

Five primary directions for future research emerge from the analysis. The most immediate is empirical validation: a systematic experimental study comparing counterfactual explanation methods with and without CIT-inspired interaction coverage on benchmark datasets, measuring not just traditional metrics (proximity, sparsity, diversity) but also the interaction coverage rate defined in Section 5.3. Such a study would test the hypothesis that the bounded-interaction principle generalizes to machine learning decision boundaries and would establish the practical value of the CIT-XAI integration.

Second, the development of adaptive interaction strength selection algorithms that automatically determine the appropriate value of t for a given dataset and model would improve the practical usability of the framework. Current CIT practice typically uses fixed t values selected based on available test budgets [9]; for XAI, a data-driven approach that identifies the empirical interaction strength of the model's decision function would be more principled.

Third, the integration of CIT with causal XAI frameworks [6] represents a rich theoretical challenge. Covering arrays over causally constrained parameter spaces where some combinations of parameter values are causally impossible would require extensions to covering array theory that have not been studied. This problem connects to the literature on constrained covering arrays in software testing [9] and could be advanced through collaboration between the causal inference and CIT communities.

Fourth, human-centered evaluation studies are needed to determine whether the structural diversity guaranteed by CIT-based interaction coverage actually improves human understanding and decision-making. Bhatt et al. [13] emphasize that the ultimate criterion for XAI quality is whether explanations improve human decisions; structural coverage properties are proxies for this goal. Controlled user studies with domain experts in medicine, finance, and law would test whether interaction-complete explanation sets lead to better decision quality than proximity-optimal single counterfactuals.

Fifth, extension of the CIT-XAI framework to deep learning models operating on image, text, and graph data represents an important challenge. The discretization of continuous feature spaces that enables CIT application to tabular data does not translate directly to high-dimensional structured data. Concepts from structural testing of neural networks [11] and multi-level covering arrays in CIT [9] may provide relevant starting points for this extension.

8. Conclusion

This paper has presented a comprehensive, literature-grounded analysis of the research foundations for integrating Combinatorial Interaction Testing with counterfactual explainable machine learning. Drawing on authoritative published research from the XAI literature [1-7, 12-14], the software testing literature [8, 9, 15, 16], and the machine learning testing literature [10, 11], we have established three principal contributions.

First, we have documented and systematically organized eight specific research gaps in the existing counterfactual XAI literature gaps that are acknowledged or implied by published studies but have not been collectively recognized as components of a unified problem. These gaps include the single-feature perturbation bias pervading existing CF methods, the absence of formal interaction coverage guarantees, the insufficiency of pairwise distance diversity measures, and the lack of a cross-domain connection between software testing methodology and XAI evaluation.

Second, we have articulated a detailed structural analogy between the software fault detection problem addressed by CIT [8, 9] and the decision boundary exploration problem addressed by counterfactual XAI [1, 4, 5, 6, 7]. This analogy is not merely metaphorical: the mathematical structure of covering arrays maps directly onto the problem of ensuring that a counterfactual explanation set covers all relevant t-way feature interaction patterns at a model's decision boundary. The bounded-interaction principle established empirically by Kuhn, Kacker, and Lei [8] provides strong a priori justification for the tractability of this approach.

Third, we have shown how CIT principles specifically address the documented gaps in diversity [4, 7], coverage [14, 13], robustness [1, 12, 13], and cross-domain integration [9-11]. The interaction coverage rate, naturally derived from CIT's formal coverage guarantees, provides an auditable, formally grounded metric for explanation completeness that complements existing evaluation criteria in the XAI literature [7, 14].

The integration of CIT with counterfactual XAI represents a disciplinary cross-pollination with substantial potential. Software engineering has invested decades in developing rigorous, empirically validated testing methodologies; explainable machine learning has identified precisely the kinds of structural coverage problems that these methodologies were designed to address. As machine learning systems continue to proliferate in high-stakes, regulated domains, the principled engineering discipline offered by CIT provides a pathway toward explanation methods that satisfy not just the intuitive criteria of proximity and plausibility but the formal criteria of completeness and verifiability that trustworthy AI demands.

References

- [1] G. Warren, R. M. J. Byrne, and M. T. Keane, "Categorical and continuous features in counterfactual explanations of AI systems," in Proc. 28th ACM Conference on Intelligent User Interfaces (IUI '23), 2023, pp. 171-187.
- [2] J. Shin, "Feasibility of local interpretable model-agnostic explanations (LIME) algorithm as an effective and interpretable feature selection method: Comparative fNIRS study," Biomedical Engineering Letters, vol. 13, no. 4, pp. 689-703, 2023.
- [3] L. Ter-Minassian, S. Ghalebikesabi, K. Diaz-Ordaz, and C. Holmes, "Challenges and opportunities of Shapley values in a clinical context," arXiv preprint arXiv:2306.14698, 2023.

- [4] M. Wang, D. Wang, W. Wu, S. Feng, and Y. Zhang, "T-COL: Generating counterfactual explanations for general user preferences on variable machine learning systems," arXiv preprint arXiv:2309.16146, 2023.
- [5] E. A. Small, J. N. Clark, C. J. McWilliams, K. Sokol, J. Chan, F. D. Salim, and R. Santos-Rodriguez, "Counterfactual explanations via locally-guided sequential algorithmic recourse," arXiv preprint arXiv:2309.04211, 2023.
- [6] G. De Toni, B. Lepri, and A. Passerini, "Synthesizing explainable counterfactual policies for algorithmic recourse with program synthesis," *Machine Learning*, vol. 112, pp. 1389-1409, 2023.
- [7] T. Laugel, A. Jeyasothy, M.-J. Lesot, C. Marsala, and M. Detyniecki, "Achieving diversity in counterfactual explanations: A review and discussion," in *Proc. 2023 ACM Conference on Fairness, Accountability, and Transparency (FAccT '23)*, 2023, pp. 1859-1869.
- [8] D. R. Kuhn, D. R. Wallace, and A. M. Gallo, Jr., "Software fault interactions and implications for software testing," *IEEE Transactions on Software Engineering*, vol. 30, no. 6, pp. 418-421, June 2004.
- [9] Y. Le Traon and T. Xie, "Combinatorial testing and machine learning for automated test generation," *Software Testing, Verification and Reliability*, vol. 33, Art. e1846, 2023.
- [10] F. Tambon, V. Majdinasab, A. Nikanjam, F. Khomh, and G. Antoniol, "Mutation testing of deep reinforcement learning based on real faults," in *Proc. IEEE Conference on Software Testing, Verification and Validation (ICST)*, 2023, pp. 188-198.
- [11] J. Yu, S. Duan, and X. Ye, "A white-box testing for deep neural networks based on neuron coverage," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 34, no. 11, pp. 9185-9197, Nov. 2023.
- [12] G. Carloni, A. Berti, and S. Colantonio, "The role of causality in explainable artificial intelligence," arXiv preprint arXiv:2309.09901, 2023.
- [13] Y. Xuan, E. Small, K. Sokol, D. Hettiachchi, and M. Sanderson, "Comprehension is a double-edged sword: Over-interpreting unspecified information in intelligible machine learning explanations," arXiv preprint arXiv:2309.08438, 2023.
- [14] S. Ali, T. Abuhmed, S. El-Sappagh, K. Muhammad, J. M. Alonso-Moral, R. Confalonieri, R. Guidotti, J. Del Ser, N. Diaz-Rodriguez, and F. Herrera, "Explainable Artificial Intelligence

- (XAI): What we know and what is left to attain trustworthy artificial intelligence," Information Fusion, vol. 99, Art. 101805, 2023.
- [15] A. A. Hassan, S. Abdullah, K. Z. Zamli, and R. Razali, "Q-learning whale optimization algorithm for test suite generation with constraints support," Neural Computing and Applications, vol. 35, no. 34, pp. 24069-24090, 2023.
- [16] X. Guo, X. Song, J.-T. Zhou, F. Wang, K. Tang, and Z. Wang, "A memetic algorithm for high-strength covering array generation," IET Software, vol. 17, no. 4, pp. 538-553, 2023.
- [17] F. Kluck, Y. Li, J. Tao, and F. Wotawa, "An empirical comparison of combinatorial testing and search-based testing in the context of automated and autonomous driving systems," Information and Software Technology, vol. 160, Art. 107225, 2023.