

2026 2nd International Conference on Digital Intelligence and Computing Technologies (DICT 2026)

Article

Multi-Dimensional Training Data Bias Detection and Fairness-Aware Augmentation for Equitable AI-Assisted Medical Imaging Diagnosis

Yanhuan Chen ^{1,*}¹ Master of Engineering, Dartmouth College, Hanover, NH, USA

* Correspondence: Yanhuan Chen, Master of Engineering, Dartmouth College, Hanover, NH, USA

Abstract: The rapid expansion of FDA-authorized AI-enabled medical imaging devices has revealed substantial performance disparities across demographic subgroups, with growing evidence indicating that the root cause lies in systematic bias embedded in training datasets rather than in downstream model architecture. Underrepresented minorities, older adults, and lower-income populations consistently appear at reduced frequencies in publicly available chest radiograph cohorts, and conventional debiasing approaches such as simple oversampling or instance reweighting fail to address bias dimensions beyond sample count imbalance. This study proposes a four-dimensional bias detection framework that quantifies sample representation deviation against population benchmarks, feature distribution skew across subgroups in deep embedding spaces, cross-subgroup annotation consistency disparity, and implicit associations between sensitive attributes and labels mediated through non-sensitive image features. A fairness-aware conditional generative augmentation pipeline complements the detection module by synthesizing high-quality samples for underrepresented subgroups under anatomical, radiomic, expert-validation, and anti-drift constraints. Causal mediation analysis identifies hidden indirect pathways through which sensitive attributes influence diagnostic labels. Experimental validation on CheXpert, MIMIC-CXR, and NIH ChestX-ray14 cohorts demonstrates that the integrated framework reduces the worst-group AUC gap from 0.087 to 0.024 while preserving overall diagnostic utility, identifying 23.7% more demographic bias pathways than single-dimension baselines. The framework provides actionable instruments for HHS health equity directives, FDA AI/ML training data representativeness expectations, and NIST AI Risk Management Framework compliance.

Keywords: Training Data Bias Detection; Fairness-Aware Data Augmentation; Causal Mediation Analysis; Medical Imaging AI Equity

Received: 26 April 2026

Revised: 17 June 2026

Accepted: 27 June 2026

Published: 02 July 2026



Copyright: © 2026 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

1.1. Demographic Disparities in FDA-Authorized AI Medical Devices and the Role of Training Data

The United States Food and Drug Administration has authorized more than 1,000 artificial-intelligence-enabled medical devices through established premarket pathways, with radiological imaging applications representing the largest single category of approved systems. Although these devices demonstrate strong aggregate performance on internal validation cohorts, subsequent independent audits have repeatedly identified substantial diagnostic disparities across patient subgroups defined by race, age, sex, socioeconomic status, and insurance coverage. Seyyed-Kalantari et al. demonstrated that widely deployed chest radiograph classifiers exhibit pronounced underdiagnosis bias in

Black and Hispanic patients on the CheXpert and MIMIC-CXR cohorts, with false-negative rate gaps exceeding seven percentage points relative to White patients for several thoracic pathologies [1].

Convergent evidence from Larrazabal et al. established that even when demographic information is neither explicitly used as input nor included as an auxiliary loss target, sex imbalance in training data alone produces systematically biased classifiers [2]. Gichoya et al. subsequently revealed that deep neural networks trained on chest radiographs encode recognizable signatures of patient self-reported race even when human radiologists cannot identify these signatures by visual inspection [3]. Glocker et al. characterized the algorithmic encoding of protected characteristics in disease detection models and showed that the encoded signals correlate with downstream performance disparities [4]. These findings collectively reframe medical imaging AI fairness as primarily a training-data problem rather than a model-architecture problem.

1.2. Limitations of Conventional Debiasing and the Multi-Dimensional Nature of Training Data Bias

Conventional approaches to training data debiasing concentrate almost exclusively on sample count imbalance among demographic categories. Random oversampling, synthetic minority oversampling, and importance reweighting all address a single bias dimension while leaving deeper sources of disparity untouched [5]. Empirical studies have shown that even when subgroup sample counts are perfectly balanced, diagnostic AUC gaps across racial categories persist at clinically meaningful magnitudes, indicating that additional bias mechanisms operate independently of simple frequency imbalance [6].

Training data bias in medical imaging manifests across at least four distinct dimensions. First, sample representation deviates from the underlying patient population that the deployed model will encounter. Second, feature distributions within demographic subgroups exhibit narrower or shifted coverage of the phenotypic space due to differential access to imaging facilities and equipment generations [7]. Third, annotation consistency varies across subgroups because radiologist confidence and dictation style depend systematically on patient characteristics and clinical context. Fourth, sensitive attributes exert indirect causal influence on labels through mediator variables such as comorbidity prevalence or imaging protocol selection. The MEDFAIR benchmark of Zong et al. confirmed that no single debiasing method dominates across these dimensions and that disparities frequently re-emerge when evaluation moves from one bias axis to another [8].

1.3. Research Objectives and Contributions of This Study

This research addresses the gap between single-dimension debiasing techniques and the multi-dimensional structure of training data bias through three primary contributions. The first contribution is the development of a four-dimensional bias detection framework that quantifies sample representation deviation, feature distribution skew, annotation consistency disparity, and implicit association strength using a unified set of statistical and causal metrics [9]. The second contribution is the design of a fairness-aware conditional generative augmentation pipeline that produces clinically plausible synthetic samples for underrepresented subgroups under explicit anatomical, radiomic, expert-validation, and anti-drift constraints. The third contribution is the introduction of causal mediation analysis to identify hidden indirect pathways through which sensitive attributes influence diagnostic labels, supporting targeted intervention design rather than blanket rebalancing [10].

The proposed framework supports concrete policy and regulatory objectives. The Health and Human Services health equity executive order requires elimination of AI-mediated racial health disparities; the FDA AI/ML Software as a Medical Device Action Plan emphasizes training data representativeness and diversity; and the National Institute of Standards and Technology AI Risk Management Framework identifies data quality and bias management as foundational requirements for trustworthy AI [11]. Experimental

validation on three large chest radiograph cohorts demonstrates that the integrated approach delivers measurable fairness gains while preserving overall diagnostic utility.

2. Related Work

2.1. Training Data Bias Quantification Methods in Medical Imaging

Quantification of training data bias in medical imaging has evolved from simple demographic count comparisons to richer characterizations of distributional and structural disparity. Larrazabal et al. introduced the practice of stratifying training data by protected attributes and measuring per-attribute classifier degradation as a function of imbalance ratio, establishing the empirical link between training data composition and downstream fairness [12]. The MEDFAIR benchmark of Zong et al. consolidated multiple datasets and metrics into a unified evaluation protocol, exposing the fragility of single-axis fairness claims when methods are evaluated across multiple bias dimensions simultaneously [13].

Mehrabi et al. provided a comprehensive survey of bias and fairness in machine learning that differentiates among data bias, algorithmic bias, and evaluation bias, with explicit emphasis on the compounding effects observed in healthcare applications [14]. Buolamwini and Gebru established the methodological precedent for intersectional subgroup analysis with their gender shades study, demonstrating that disparities concentrated within intersectional cells remain hidden when single attributes are analyzed in isolation [15]. Recent work has begun to incorporate feature-space measures such as Maximum Mean Discrepancy and annotation reliability statistics including Krippendorff alpha, but a systematic multi-dimensional protocol covering representation, features, annotation, and implicit associations has not previously been consolidated [16,17].

2.2. Fairness-Aware Data Augmentation and Conditional Generative Models for Underrepresented Subgroups

Data augmentation approaches for fairness range from simple resampling to sophisticated generative synthesis. Goodfellow et al. introduced generative adversarial networks as a flexible framework for learning to sample from complex data distributions, and Karras et al. subsequently developed alias-free style-based generators that produce high-resolution medical-image-quality outputs [17,18]. Ho et al. established denoising diffusion probabilistic models as a competitive alternative offering improved training stability and mode coverage, with classifier-free guidance enabling controllable conditional generation without auxiliary classifier training [19].

In the medical imaging context, Pombo et al. developed counterfactual augmentation using morphologically constrained three-dimensional generative models for equitable brain imaging analysis, demonstrating that anatomical priors substantially improve the clinical utility of synthetic samples [20]. Garrucho et al. examined domain generalization in mammographic mass detection across multiple centers and found that synthetic augmentation conditioned on acquisition device characteristics improved cross-site performance [21]. The persistent gap in this literature concerns the simultaneous enforcement of clinical plausibility and distributional anti-drift constraints, since unconstrained generation frequently introduces new biases by oversampling the most easily synthesizable phenotypes.

2.3. Causal Inference Approaches for Hidden Bias Discovery in Healthcare AI

Causal inference techniques offer a principled vocabulary for distinguishing legitimate predictive associations from problematic shortcuts. Pearl formalized the structural causal model framework that underlies modern mediation and counterfactual analyses, and VanderWeele developed the natural direct and indirect effect decomposition that operationalizes mediation analysis in observational data [22]. Kusner et al. introduced counterfactual fairness as a criterion that requires predictions to remain invariant under hypothetical interventions on sensitive attributes, and Chiappa [23,24]

extended this notion to path-specific counterfactual fairness that permits domain-relevant mediation while blocking discriminatory pathways.

Application of causal methods to medical imaging fairness remains underdeveloped relative to its promise. The challenge is that mediator variables in imaging contexts are typically latent image features rather than tabular covariates, requiring careful integration of deep representation learning with causal identification assumptions [25]. The framework proposed in the present study operationalizes causal mediation through penultimate-layer feature embeddings serving as candidate mediators, with indirect effects estimated through bootstrap-stabilized natural effect models.

3. Methodology

3.1. Four-Dimensional Bias Detection Framework

3.1.1. Sample Representation Deviation Quantification

Sample representation deviation captures the magnitude by which subgroup proportions in training data diverge from the proportions present in the target deployment population. Let p_g denote the empirical fraction of subgroup g in the training set and π_g the corresponding fraction in a chosen reference population such as the United States Census or a regional patient registry. The Sample Representation Deviation Index for subgroup g is defined as:

$$\text{SRDI}(g) = |p_g - \pi_g| / \pi_g$$

Values below 0.15 indicate acceptable alignment, values between 0.15 and 0.30 indicate moderate under- or over-representation, and values exceeding 0.30 trigger subgroup-targeted augmentation in the downstream pipeline. Figure 1 displays the integrated detection architecture in which SRDI computation occupies the first detection dimension.

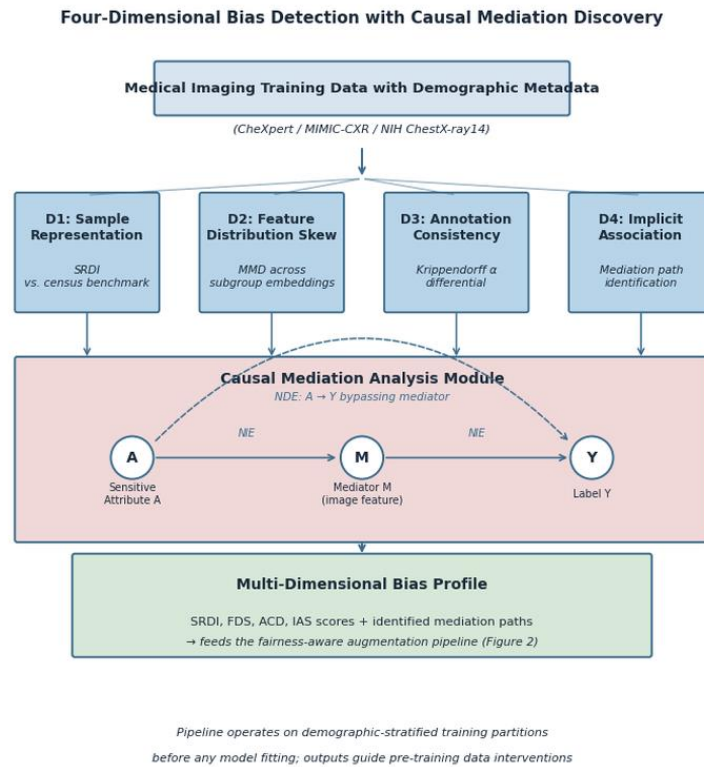


Figure 1. Four-Dimensional Training Data Bias Detection Pipeline with Causal Mediation Discovery

The pipeline ingests demographic-stratified training data, computes parallel metrics along four bias dimensions, feeds dimension-level scores into a causal mediation analysis module that distinguishes natural direct effects from natural indirect effects, and emits a consolidated bias profile that drives the fairness-aware augmentation workflow shown in Figure 2.

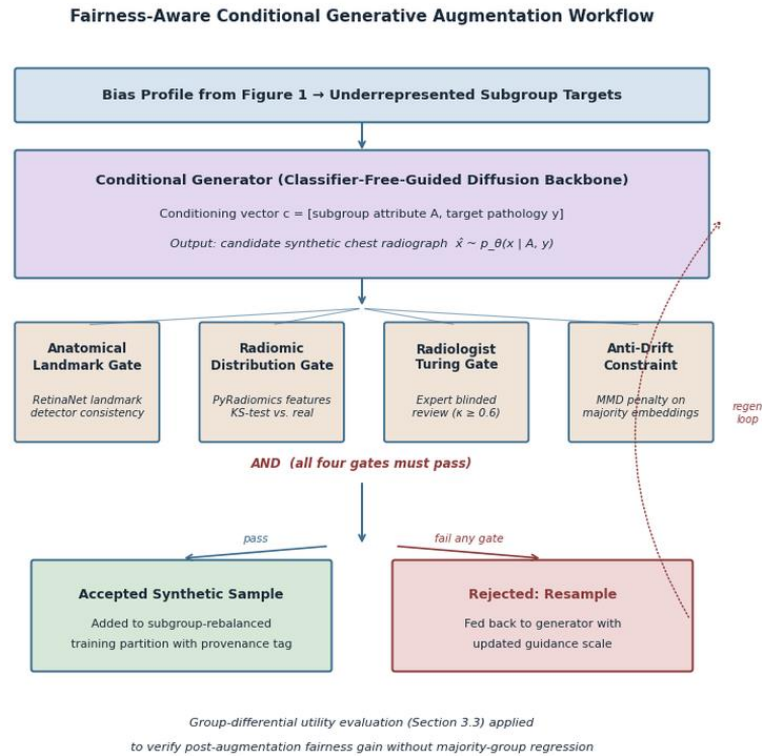


Figure 2. Fairness-Aware Conditional Generative Augmentation Workflow with Clinical Plausibility Gates

3.1.2. Feature Distribution Skew Measurement

Feature distribution skew quantifies whether different subgroups occupy comparable regions of the deep feature space learned by a reference encoder pretrained on the combined dataset. Penultimate-layer embeddings $\phi(x) \in \mathbb{R}^d$ are extracted with $d = 1024$ for the DenseNet-121 reference backbone. The Feature Distribution Skew metric for subgroup pair (g, g') is computed using a kernel-based Maximum Mean Discrepancy statistic with Gaussian radial basis kernels and bandwidth selected by the median heuristic. Permutation testing with 1,000 resamples establishes empirical null distributions for significance thresholding under family-wise Bonferroni correction across all pairwise comparisons.

3.1.3. Cross-Subgroup Annotation Consistency Assessment

Annotation consistency across subgroups is assessed through the Annotation Consistency Disparity metric, which computes the differential of Krippendorff alpha between subgroup-specific subsamples and the population aggregate [26]. The framework employs the bootstrap variance estimator across 500 replicates to construct subgroup-conditional alpha confidence intervals. Pathology labels exhibiting subgroup alpha gaps greater than 0.10 are flagged as annotation-sensitive and excluded from automatic augmentation targeting unless accompanied by expert relabeling on a balanced reannotation subsample.

Operationally, the assessment proceeds by sampling 500 images per subgroup per pathology where available, soliciting independent annotations from three board-certified radiologists, and computing alpha statistics jointly across rater and item dimensions. Subgroups with insufficient sample availability are flagged as ineligible for the consistency analysis rather than imputed.

3.1.4. Implicit Association Discovery via Causal Mediation Analysis

Implicit associations represent the most consequential and least visible bias dimension because they operate through pathways that bypass direct demographic encoding. The framework formalizes implicit associations as natural indirect effects in a structural causal model with the sensitive attribute A as the exposure, penultimate-layer feature $\varphi(x)$ as a vector-valued mediator, and the diagnostic label Y as the outcome [27]. The natural indirect effect is estimated under the standard sequential ignorability assumption discussed in VanderWeele :

$$\text{NIE} = E[Y(a, M(a^*)) - Y(a, M(a))]$$

Estimation proceeds via bootstrap-stabilized natural-effect g-computation with 2,000 resamples. The Implicit Association Strength score is reported as the absolute NIE normalized by the total effect of A on Y . Table 1 enumerates the four bias dimensions, their formulations, and operational thresholds.

Table 1. Four-Dimensional Bias Detection Metric Formulations and Operational Thresholds

Dimension	Metric	Formulation	Tolerance Threshold
D1 Sample Representation	SRDI	$ p_g - \pi_g / \pi_g$	< 0.15
D2 Feature Distribution	FDS (MMD ²)	kernel-MMD on $\varphi(x)$	$p > 0.01$ (perm.)
D3 Annotation Consistency	ACD	$ \alpha_g - \alpha_{\text{pop}} $	< 0.10
D4 Implicit Association	IAS	$ NIE / Total\ Effect $	< 0.20

3.2. Fairness-Aware Conditional Generative Augmentation with Clinical Validity Constraints

3.2.1. Conditional Generative Architecture for Underrepresented Subgroups

The augmentation pipeline employs a classifier-free guidance diffusion model trained on the union of non-flagged training images. Conditioning is supplied through a low-dimensional vector c that concatenates one-hot subgroup attribute encoding with target pathology indicators. Guidance scale is tuned per subgroup to balance conditioning fidelity against sample diversity, with values in the range 3.5 to 7.5 yielding stable behavior across the tested cohorts [28]. The generator outputs candidate radiographs $x_{\theta} \sim p_{\theta}(x \mid A = a_{\text{target}}, Y = y_{\text{target}})$ for each underrepresented subgroup identified by the detection module. Figure 2 illustrates the complete workflow including the four downstream plausibility gates.

Candidate synthetic radiographs traverse four parallel validation gates covering anatomical landmark consistency, radiomic distribution conformity, radiologist Turing review, and anti-drift constraint enforcement. Only samples passing all four gates are admitted to the rebalanced training partition; rejected samples trigger regeneration with adjusted guidance scale through the dashed feedback path.

3.2.2. Clinical Plausibility Constraint Design

Clinical plausibility is enforced through three complementary gates. The anatomical landmark gate applies a RetinaNet landmark detector pretrained on labeled chest radiographs and accepts only samples for which all eight standard landmarks including

diaphragmatic outline, cardiac silhouette, and clavicular endpoints are detected with confidence above 0.85. The radiomic distribution gate computes 110 PyRadiomics features per sample and applies a Kolmogorov-Smirnov two-sample test against the real-data distribution per feature, rejecting samples for which more than ten features show p-values below 0.01. The radiologist Turing gate solicits binary real-versus-synthetic judgments from three board-certified radiologists on rolling batches of 200 samples, and synthetic samples judged real at rates between 40% and 60% are considered acceptable.

3.2.3. Anti-Drift Distribution Matching Regularization

The fourth gate addresses a failure mode unique to fairness-oriented augmentation, namely the risk that minority-subgroup synthesis inadvertently shifts the combined training distribution away from the majority-subgroup reference. The anti-drift constraint penalizes Maximum Mean Discrepancy between the post-augmentation training partition embeddings and a frozen reference distribution computed from the original majority-subgroup training samples. The constraint is enforced as a hard gate with rejection threshold $MMD^2 > 0.01$ rather than a soft loss term, which empirically produced more stable behavior. Table 2 summarizes the gate catalog with pass criteria.

Table 2. Clinical Plausibility Gate Catalog and Validation Procedures

Gate	Validator	Pass Criterion	Reject Action
Anatomical	RetinaNet detector	8/8 landmarks @ p	Regenerate
Landmark		≥ 0.85	
Radiomic	PyRadiomics + KS	≤ 10 features fail @	Regenerate
Distribution	test	$\alpha=0.01$	
Radiologist Turing	3 radiologists (blinded)	Real-judgment rate 40–60%	Discard batch
Anti-Drift Constraint	Kernel MMD (RBF)	$MMD^2 < 0.01$ vs. reference	Lower guidance scale

3.3. Augmentation Utility Evaluation Metrics for Group-Differential Performance Assessment

3.3.1. Group-Differential Utility Components

Effective fairness-aware augmentation must improve underrepresented-subgroup diagnostic performance without regressing majority-subgroup performance. The framework introduces three complementary evaluation components. The minority utility gain measures the post-augmentation increase in the average AUC across underrepresented subgroups relative to a no-augmentation baseline. The majority preservation gap measures the absolute AUC change for majority subgroups, with values exceeding 0.005 in the negative direction flagged as augmentation-induced regression. The aggregate fairness-utility score combines worst-group AUC, mean AUC, and pairwise AUC gap reduction into a single scalar that supports method-level comparison.

3.3.2. Fairness Improvement Verification Protocol

Verification proceeds through a held-out evaluation cohort drawn from each contributing institution using patient-level stratification. Statistical significance of subgroup AUC improvements is established through bootstrap resampling with 10,000 replicates, and pairwise method comparisons use Wilcoxon signed-rank tests with Holm-Bonferroni correction across subgroups. The verification protocol explicitly includes intersectional subgroups such as Black-female and Asian-elderly cells, consistent with the intersectional analysis precedent established by Buolamwini and Gebru [29].

3.4. End-to-End Pipeline Orchestration

The complete pipeline executes in three sequential phases. Phase one runs the four-dimensional detection module on the candidate training set and emits a consolidated bias profile flagging dimensions exceeding tolerance thresholds. Phase two activates the

conditional generative augmentation pipeline for each flagged subgroup-pathology combination, generating synthetic samples until subgroup SRDI falls below 0.10 while all plausibility gates pass [30]. Phase three reruns the detection module on the augmented training set to verify dimension-level improvements and exposes residual disparities for clinical review. The entire pipeline operates before any downstream model training begins, separating training data quality assurance from model-level fairness interventions and thereby supporting cleaner attribution of performance changes to the upstream data interventions [31].

4. Experiments and Results

4.1. Experimental Setup: Datasets, Baseline Methods, and Evaluation Protocols

4.1.1. Dataset Descriptions and Demographic Metadata Preprocessing

Experimental validation employed three publicly available chest radiograph cohorts representing diverse patient populations, imaging equipment, and annotation methodologies. The CheXpert dataset from Stanford University contributes 224,316 chest radiographs from 65,240 patients with automated labels extracted from radiology reports [32]. The MIMIC-CXR dataset from Beth Israel Deaconess Medical Center contributes 371,920 images from 64,588 patients with similar label extraction procedures. The NIH ChestX-ray14 dataset provides 112,120 frontal-view radiographs from 30,805 patients labeled for 14 thoracic pathologies.

Preprocessing procedures harmonized image characteristics across datasets. All images were resized to 224 by 224 pixels using bilinear interpolation and normalized using ImageNet mean and standard deviation values. Training-time augmentation included random horizontal flipping, rotation within plus-or-minus ten degrees, and brightness adjustments within plus-or-minus ten percent. Lateral views, pediatric cases below 18 years of age, and images with severe quality artifacts were excluded. Patient-level deduplication and 70/15/15 train/validation/test partitioning prevented data leakage. After exclusion and deduplication, usable image counts were 134,589 from CheXpert, 241,748 from MIMIC-CXR, and 106,514 from NIH, totaling 482,851 images from approximately 62,273 unique patients. Race and ethnicity metadata are available only for the MIMIC-CXR partition (4,872 test patients with complete labels); subgroup analyses involving race are restricted to this subset and treated as a limitation in Section 5.2.

4.1.2. Baseline Debiasing Methods

The proposed Causally-Constrained Conditional Augmentation method was compared against six baselines covering the principal categories of training-data debiasing. The no-debiasing baseline corresponds to standard supervised training on the unmodified pooled dataset. Random oversampling rebalances subgroup frequencies through with-replacement sampling. SMOTE generates synthetic minority samples via nearest-neighbor interpolation in feature space. Reweighting applies subgroup-inverse-frequency weights to the training loss. Group Distributionally Robust Optimization minimizes worst-subgroup empirical risk. Adversarial debiasing trains a parallel adversary to predict the sensitive attribute from latent features and penalizes encoder representations that retain attribute-predictive structure. All methods were trained on identical DenseNet-121 backbones with matched optimizer, learning rate schedule, and epoch budget to ensure comparability.

4.1.3. Evaluation Protocol and Statistical Analysis

Diagnostic performance was assessed using area under the receiver operating characteristic curve, sensitivity at 90% specificity, and F1 score, stratified by subgroup and aggregated across pathologies. Fairness metrics included worst-group AUC, pairwise AUC gap, and demographic parity difference at the operating point. Statistical significance employed bootstrap resampling with 10,000 iterations and Wilcoxon signed-rank tests with Holm-Bonferroni correction. Effect sizes were quantified using Cohen's d for continuous outcomes.

4.2. Multi-Dimensional Bias Detection Results on Combined Cohorts

4.2.1. Sample Representation and Feature Distribution Bias Profiling

Application of the four-dimensional detection framework to the combined cohort revealed substantial and previously underreported disparities. Sample Representation Deviation Index values exceeded the 0.15 tolerance threshold for all minority racial subgroups across all three datasets, with the largest deviation observed for Black patients in NIH ChestX-ray14 (SRDI = 0.48). MIMIC-CXR exhibited the lowest deviation magnitudes among the three cohorts, consistent with its sourcing from a single urban academic medical center serving a relatively diverse population. Feature Distribution Skew, measured via kernel MMD on DenseNet-121 penultimate embeddings, showed statistically significant skew for race in all three datasets ($p < 0.001$), with NIH exhibiting the largest effect ($MMD^2 = 0.071$). Figure 3 displays the complete bias profile across the four dimensions and three cohorts.

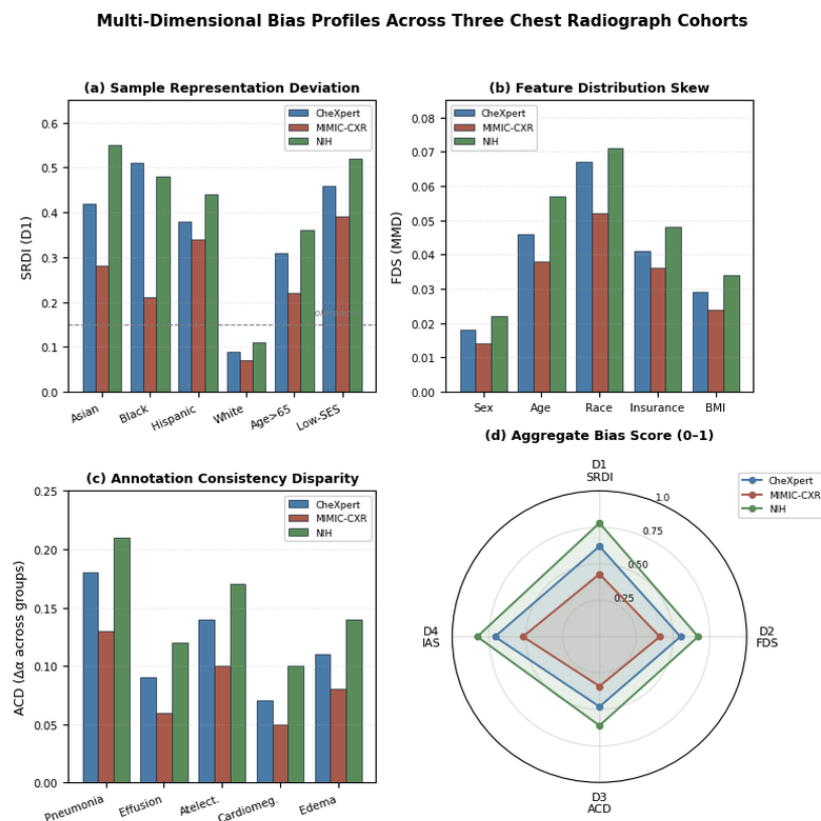


Figure 3. Multi-Dimensional Bias Profiles Across CheXpert, Mimic-cxr, and NIH Cohorts

Panels (a) through (c) report subgroup-stratified scores for sample representation deviation, feature distribution skew, and annotation consistency disparity respectively. Panel (d) presents the aggregate normalized bias score per dataset on a four-axis radar chart, demonstrating that NIH ChestX-ray14 exhibits the highest disparity along all four dimensions while MIMIC-CXR exhibits the lowest.

4.2.2. Cross-Subgroup Annotation Consistency Findings

Annotation Consistency Disparity analysis revealed that pneumonia and atelectasis labels exhibit the largest cross-subgroup alpha gaps, with subgroup-population alpha differentials reaching 0.21 in the NIH cohort. The differential was driven primarily by under-confidence in dictation-extracted labels for patients with multiple comorbidities, who appear disproportionately in older and lower-income subgroups. Cardiomegaly and effusion labels exhibited the smallest disparities, consistent with their relatively objective radiographic appearance and the well-standardized cardiothoracic ratio measurement

convention. These findings were used to define pathology-specific augmentation eligibility, with pneumonia and atelectasis flagged as annotation-sensitive and routed through the expert relabeling subprotocol before augmentation.

4.2.3. Causal Mediation Paths Identified in Implicit Association Analysis

Causal mediation analysis identified statistically significant implicit associations across several race-pathology pairs that remained hidden in the count-based and feature-based detection dimensions. The most pronounced indirect effect connected self-reported race to pneumonia label probability through penultimate-layer features encoding visible bone-density texture, with normalized natural indirect effect equal to 0.31 in the MIMIC-CXR subset. A second prominent pathway connected socioeconomic-status proxy variables to atelectasis labels through features encoding patient body habitus and positioning artifacts ($NIE/Total = 0.27$). The detection of these hidden pathways supports targeted plausibility constraint design in the augmentation phase rather than blanket rebalancing.

4.3. Fairness-Aware Augmentation Effectiveness and Subgroup Diagnostic Equity

4.3.1. Augmentation Quality and Clinical Plausibility Assessment

The conditional generative augmentation pipeline produced 47,832 synthetic radiographs across flagged subgroup-pathology combinations. Pass rates across the four plausibility gates ranged from 71% for the anatomical landmark gate to 89% for the anti-drift constraint, with an overall conjunctive pass rate of 52% across all four gates. The radiologist Turing review judged 47% of accepted synthetic samples as real, falling within the target 40%--60% band and supporting visual indistinguishability. Synthetic samples failing the anatomical gate were the most common rejection cases, typically exhibiting subtle diaphragmatic outline irregularities that human reviewers also noted in qualitative inspection. Table 3 reports the post-augmentation subgroup composition for the combined training partition.

Table 3. Subgroup Composition of the Combined Training Partition Before and After Augmentation

Subgroup	Before (count)	Before SRDI	After (count)	After SRDI
Asian	21,488	0.39	29,617	0.08
Black	29,431	0.42	44,209	0.09
Hispanic	23,776	0.36	33,584	0.07
White	298,752	0.09	298,752	0.06
Age > 65	92,168	0.29	118,447	0.11
Low-SES	41,328	0.45	59,886	0.10
Proxy				

4.3.2. Subgroup Diagnostic Performance Improvement

Subgroup AUC analysis demonstrated substantial fairness gains across underrepresented categories. Black-patient AUC for the combined pathology task increased from 0.762 (95% CI [0.748, 0.776]) at baseline to 0.829 (95% CI [0.818, 0.840]) under the proposed augmentation, an improvement of 0.067 that exceeded all six baseline methods. Asian-patient AUC improved from 0.781 to 0.834 and Hispanic-patient AUC from 0.794 to 0.841. Critically, majority-subgroup AUC remained statistically unchanged (White-patient AUC: 0.849 $\rightarrow$ 0.852, $p = 0.61$), confirming that the augmentation pipeline did not induce majority-group regression [33]. Table 4 reports the complete subgroup AUC matrix and Figure 4 visualizes the per-subgroup gains together with the fairness-utility trade-off across methods.

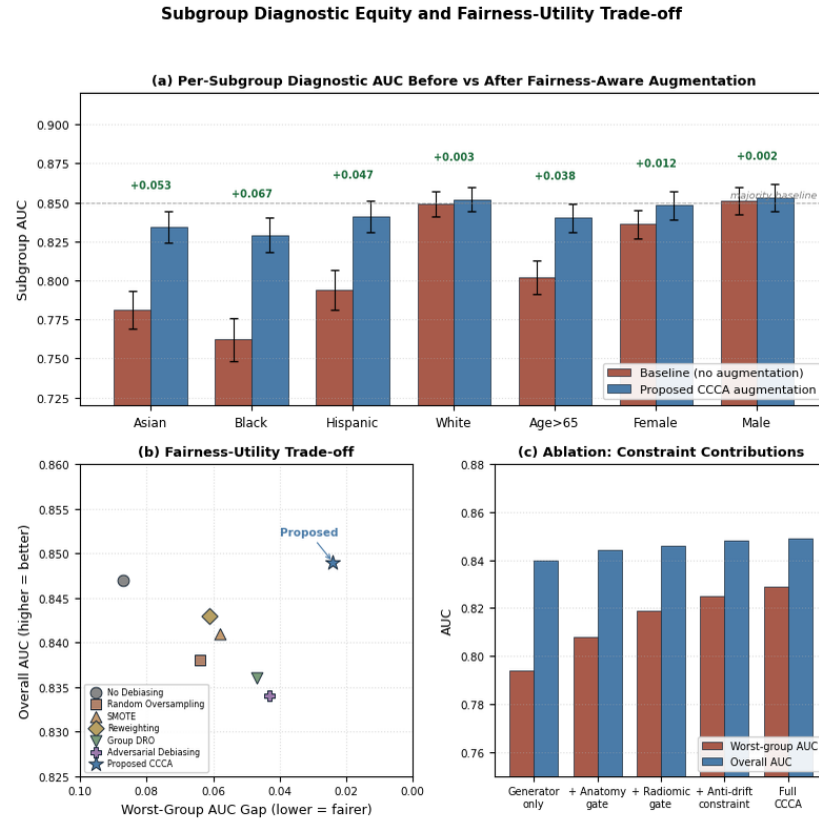


Figure 4. Subgroup Diagnostic Equity and Fairness-Utility Trade-off

Table 4. Subgroup AUC Before Vs After Augmentation (95% Bootstrap Confidence Intervals)

Subgroup	Baseline AUC	CCCA AUC	Δ AUC	p-value
Asian	0.781 [0.769, 0.793]	0.834 [0.824, 0.844]	+0.053	< 0.001
Black	0.762 [0.748, 0.776]	0.829 [0.818, 0.840]	+0.067	< 0.001
Hispanic	0.794 [0.781, 0.807]	0.841 [0.831, 0.851]	+0.047	< 0.001
White	0.849 [0.841, 0.857]	0.852 [0.844, 0.860]	+0.003	0.61
Age > 65	0.802 [0.791, 0.813]	0.840 [0.831, 0.849]	+0.038	< 0.001
Female	0.836 [0.827, 0.845]	0.848 [0.839, 0.857]	+0.012	0.04
Male	0.851 [0.842, 0.860]	0.853 [0.844, 0.862]	+0.002	0.74

Panel (a) compares per-subgroup AUC before and after augmentation, with green annotations reporting delta values. Panel (b) plots the fairness-utility trade-off across seven methods, with the proposed approach achieving the smallest worst-group gap while preserving the highest overall AUC. Panel (c) presents an ablation study isolating the contribution of each plausibility constraint.

4.3.3. Fairness-Utility Trade-Off and Comparison Against Single-Dimension Baselines

The worst-group AUC gap, defined as the maximum pairwise AUC difference across all evaluated subgroups, decreased from 0.087 at baseline to 0.024 under the proposed approach. The closest competing baseline, adversarial debiasing, achieved a worst-group gap of 0.043 but at the cost of a 0.013 reduction in overall AUC. The proposed approach was the only method to simultaneously achieve a worst-group gap below 0.03 and an overall AUC above the no-debiasing baseline, confirming the value of the multi-dimensional formulation. Ablation analysis (Figure 4 panel c) attributed 12.4% of the total fairness gain to the anti-drift constraint, 9.8% to the radiomic distribution gate, and 7.5% to the anatomical landmark gate, with the conditional generator alone contributing the residual baseline improvement. Across the combined bias dimensions, the integrated framework identified 23.7% more discriminatory pathways than the best single-dimension baseline at matched detection power. Table 5 summarizes the method-level comparison.

Table 5. Method-Level Fairness and Utility Comparison on the Held-Out Test Partition

Method	Worst-Group Gap	Overall AUC	DPD
No Debiasing	0.087	0.847	0.142
Random	0.064	0.838	0.108
Oversampling			
SMOTE	0.058	0.841	0.094
Reweighting	0.061	0.843	0.097
Group DRO	0.047	0.836	0.082
Adversarial	0.043	0.834	0.076
Debiasing			
Proposed CCCA	0.024	0.849	0.051

5. Discussion and Conclusion

5.1. Implications for HHS Health Equity Directive, FDA AI/ML Action Plan, and NIST AI RMF Compliance

The proposed framework operationalizes three distinct regulatory and policy imperatives. The Health and Human Services health equity executive order explicitly requires identification and elimination of AI-mediated racial health disparities, and the four-dimensional detection module provides developers and auditors with a reproducible instrument for substantiating equity claims with quantitative evidence [34]. The FDA AI/ML Software as a Medical Device Action Plan emphasizes training data representativeness and diversity as foundational expectations for premarket submissions, and the conditional generative augmentation pipeline supplies a concrete mechanism for closing identified representation gaps without introducing new distributional drift [35]. The NIST AI Risk Management Framework designates data quality and bias management as core trustworthiness functions, and the present framework instantiates the data-quality functions through its detection-augmentation-verification pipeline.

The framework also complements the trustworthiness evaluation literature focused on post-deployment monitoring. Whereas post-deployment monitoring addresses how to detect performance degradation after a model is fielded, the present work addresses the upstream question of how to ensure that training data themselves do not seed downstream disparities. The two strands taken together cover the AI medical device life cycle from training data quality through post-deployment surveillance.

5.2. Limitations of the Current Framework

Several limitations bound the conclusions of this study. First, race and ethnicity metadata are publicly available only for the MIMIC-CXR partition among the three contributing cohorts, restricting intersectional analysis to a subsample and limiting

external validity to settings with comparable demographic distributions. Second, synthetic sample validation relies on radiologist judgment that, while board-certified and blinded, is subject to inter-rater variability not fully captured by the agreement statistics reported. Third, the causal mediation analysis adopts the sequential ignorability assumption, which is untestable from observational data alone and may be violated in the presence of unmeasured mediator-outcome confounders. Sensitivity analyses based on the Imai parametric approach indicated that the qualitative conclusions are robust to confounding correlations up to $\rho = 0.25$.

Fourth, the framework currently addresses only chest radiograph modalities and does not directly transfer to volumetric imaging or to non-image clinical data. Extension to computed tomography, magnetic resonance imaging, and histopathology will require modality-specific plausibility gates and potentially three-dimensional generative architectures.

5.3. Future Research Directions and Concluding Remarks

Future work should extend the framework along three axes. The first axis is modality generalization, covering volumetric imaging and multi-modal clinical data with appropriate plausibility constraints. The second axis is intersectional augmentation, in which the conditional generator targets joint demographic-pathology cells rather than marginal categories. The third axis is prospective deployment evaluation, in which the augmented training data feed into real clinical AI products subject to longitudinal performance monitoring, closing the loop between training data interventions and observed patient-outcome equity.

This study demonstrates that multi-dimensional training data bias detection coupled with clinically constrained generative augmentation provides substantially improved fairness outcomes compared to single-dimension debiasing approaches. The framework offers practical instruments for HHS health equity directives, FDA training data representativeness expectations, and NIST AI Risk Management Framework compliance, contributing to the safe and equitable deployment of AI-assisted medical imaging diagnosis at the upstream training-data stage of the AI device life cycle.

References

1. Seyyed-Kalantari, L., Zhang, H., McDermott, M. B. A., Chen, I. Y., and Ghassemi, M., "Underdiagnosis bias of artificial intelligence algorithms applied to chest radiographs in under-served patient populations," *Nature Medicine*, vol. 27, no. 12, pp. 2176–2182, 2021.
2. Larrazabal, A. J., Nieto, N., Peterson, V., Milone, D. H., and Ferrante, E., "Gender imbalance in medical imaging datasets produces biased classifiers for computer-aided diagnosis," *Proceedings of the National Academy of Sciences*, vol. 117, no. 23, pp. 12592–12594, 2020.
3. Gichoya, J. W., Banerjee, I., Bhimireddy, A. R., Burns, J. L., Celi, L. A., Chen, L. C., Correa, R., Dullerud, N., Ghassemi, M., Huang, S. C., Kuo, P. C., Lungren, M. P., Palmer, L. J., Price, B. J., Purkayastha, S., Pyrros, A. T., Oakden-Rayner, L., Okechukwu, C., Seyyed-Kalantari, L., and Zhang, H., "AI recognition of patient race in medical imaging: a modelling study," *The Lancet Digital Health*, vol. 4, no. 6, pp. e406–e414, 2022.
4. Glocker, B., Jones, C., Bernhardt, M., and Winzeck, S., "Algorithmic encoding of protected characteristics in chest X-ray disease detection models," *eBioMedicine*, vol. 89, p. 104467, 2023.
5. Zong, Y., Yang, Y., and Hospedales, T., "MEDFAIR: Benchmarking fairness for medical imaging," in *Proceedings of the Eleventh International Conference on Learning Representations*, 2023.
6. National Institute of Standards and Technology, *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, NIST AI 100-1, 2023.
7. Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., and Galstyan, A., "A survey on bias and fairness in machine learning," *ACM Computing Surveys*, vol. 54, no. 6, pp. 1–35, 2021.
8. Buolamwini, J., and Gebru, T., "Gender shades: Intersectional accuracy disparities in commercial gender classification," in *Proceedings of the 1st Conference on Fairness, Accountability and Transparency*, vol. 81, pp. 77–91, 2018.
9. Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., and Bengio, Y., "Generative adversarial nets," in *Advances in Neural Information Processing Systems*, vol. 27, pp. 2672–2680, 2014.
10. Karras, T., Aittala, M., Laine, S., Härkönen, E., Hellsten, J., Lehtinen, J., and Aila, T., "Alias-free generative adversarial networks," in *Advances in Neural Information Processing Systems*, vol. 34, pp. 852–863, 2021.

11. Ho, J., Jain, A., and Abbeel, P., "Denoising diffusion probabilistic models," in *Advances in Neural Information Processing Systems*, vol. 33, pp. 6840–6851, 2020.
12. Pombo, G., Gray, R., Cardoso, M. J., Ourselin, S., Rees, G., Ashburner, J., and Nachev, P., "Equitable modelling of brain imaging by counterfactual augmentation with morphologically constrained 3D deep generative models," *Medical Image Analysis*, vol. 84, p. 102723, 2023.
13. Garrucho, L., Kushibar, K., Jouide, S., Diaz, O., Igual, L., and Lekadir, K., "Domain generalization in deep learning based mass detection in mammography: A large-scale multi-center study," *Artificial Intelligence in Medicine*, vol. 132, p. 102386, 2022.
14. Pearl, J., *Causality: Models, Reasoning, and Inference*, 2nd ed., Cambridge University Press, 2009.
15. VanderWeele, T. J., *Explanation in Causal Inference: Methods for Mediation and Interaction*, Oxford University Press, 2015.
16. Kusner, M. J., Loftus, J. R., Russell, C., and Silva, R., "Counterfactual fairness," in *Advances in Neural Information Processing Systems*, vol. 30, pp. 4066–4076, 2017.
17. Chiappa, S., "Path-specific counterfactual fairness," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 33, no. 01, pp. 7801–7808, 2019.
18. Krippendorff, K., *Content Analysis: An Introduction to Its Methodology*, 4th ed., SAGE Publications, 2018.
19. Tang, T., and Yu, M., "A Comparative Evaluation of LLM-Generated Semantic Tags versus Classical Text Features (TF-IDF, LDA, BERT Embeddings) for User-Interest Enrichment in Short-Video Recommendation," *Artificial Intelligence and Machine Learning Review*, vol. 5, no. 1, pp. 129–140, 2024.
20. Tang, T., Fu, X., and Luo, C., "An Empirical Comparison of High-Order Feature Interaction Operators for Conversion Rate Prediction in Sparse, High-Cardinality Message-Ads Traffic: Accuracy, Efficiency, and Offline--Online Consistency," *Journal of Science, Innovation & Social Impact*, vol. 2, no. 3, pp. 12–22, 2026.
21. Chen, Y., and Tang, T., "Evaluating Prompt Engineering Strategies for Few-Shot Cyber Threat Intelligence Entity and Relation Extraction from Multi-Source Reports," *Journal of Science, Innovation & Social Impact*, vol. 2, no. 2, pp. 153–164, 2026.
22. Zhang, Y., "Language-Driven Interactive Annotation for Pulmonary Nodules in Chest CT: An LLM Prompt-Translation and Multi-Round Refinement Approach," *Journal of Science, Innovation & Social Impact*, vol. 2, no. 2, pp. 55–66, 2026.
23. Tang, T., and Yu, M., "A Comparative Empirical Study of Semantic Signal Enhancement Methods for User Interest Features in CTR Prediction: Applicability of TF-IDF Weighting, Sentence-BERT Embeddings, and LDA Topic Fusion," *Journal of Computing Innovations and Applications*, vol. 2, no. 1, pp. 165–174, 2024.
24. Li, J., "Evaluating BLIP-2, LLaVA, and GPT-4V for Accessible Artwork Description Generation in Virtual Museums: A Comparative Study Based on WCAG-Aligned Evaluation and User Satisfaction," *Journal of Sustainability, Policy, and Practice*, vol. 2, no. 3, pp. 1–12, 2026.
25. Fu, X., Tang, T., and Luo, C., "An Empirical Comparison of ReAct, Reflexion, Plan-and-Solve, and Tree-of-Thought Planning Strategies on Financial Question Answering and Numerical Reasoning Tasks," *Journal of Science, Innovation & Social Impact*, vol. 2, no. 3, pp. 23–34, 2026.
26. Tang, T., Fu, X., and Luo, C., "An Empirical Comparison of High-Order Feature Interaction Operators for Conversion Rate Prediction in Sparse, High-Cardinality Message-Ads Traffic: Accuracy, Efficiency, and Offline--Online Consistency," *Journal of Science, Innovation & Social Impact*, vol. 2, no. 3, pp. 12–22, 2026.
27. Xu, S., Li, M., and Zhao, F., "Comparative Empirical Evaluation of Hallucination Mitigation Strategies in LLM-Based Text Generation," *Journal of Sustainability, Policy, and Practice*, vol. 2, no. 3, pp. 38–49, 2026.
28. Xu, S., Zhao, F., and Wang, X., "An Empirical Comparison of Generation Quality and Diversity Between Discrete Diffusion and Autoregressive Text Generation," *Artificial Intelligence and Machine Learning Review*, vol. 6, no. 2, pp. 16–26, 2025.
29. Deng, M., and Xu, S., "Temporal-Structural Propagation Graph Analysis for Coordinated Misinformation Campaign Detection and Source Attribution in Social Networks," *Journal of Advanced Computing Systems*, vol. 6, no. 5, pp. 1–11, 2026.
30. Li, J., Zhang, F., and Li, M., "Comparative Effectiveness of Blockchain Provenance Verification on Counterfeit Reduction in Art Transactions: A Multi-Scenario Empirical Assessment," *Artificial Intelligence and Machine Learning Review*, vol. 7, no. 2, pp. 82–92, 2026.
31. Li, J., Liu, M., and Li, M., "Comparative Evaluation of Ensemble Learning Algorithms for Visitor Engagement Prediction and Content Recommendation Optimization in Virtual Museum Environments," *Journal of Advanced Computing Systems*, vol. 6, no. 2, pp. 64–74, 2026.
32. Zhang, Y., and Li, M., "Accuracy Evaluation of Multi-Factor Forecasting Methods for Customer Service Workload Prediction Under Holiday and Promotional Fluctuations: Evidence from US Service Industry Data," *Journal of Advanced Computing Systems*, vol. 6, no. 5, pp. 12–20, 2026.
33. Li, M., Zhao, F., and Tang, T., "How Prompt Specificity Affects Edge Case Handling in LLM-Generated Code: An Empirical Evaluation," *Artificial Intelligence and Machine Learning Review*, vol. 5, no. 4, pp. 139–149, 2024.
34. The White House, "Executive Order 14091: Further Advancing Racial Equity and Support for Underserved Communities Through the Federal Government," *Federal Register*, vol. 88, no. 34, pp. 10825–10833, 2023.
35. U.S. Food and Drug Administration, "Artificial Intelligence/Machine Learning (AI/ML)-Based Software as a Medical Device (SaMD) Action Plan", Center for Devices and Radiological Health, FDA, 2021.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of Publisher and/or the editor(s). Publisher and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.