

**FAIRNESS-AWARE MULTI-AGENT
REINFORCEMENT LEARNING FOR EQUITABLE
SEPSIS DETECTION ACROSS DEMOGRAPHIC
SUBGROUPS**

Praneesh Khanna¹

University of Maryland, College Park, MD, praneesh.khanna1204@gmail.com

Abstract

Sepsis, a dysregulated host response to infection leading to life-threatening organ dysfunction, is a leading cause of in-hospital mortality in the United States and worldwide, with an estimated one in five deaths globally attributed to the condition.¹ The prognosis for septic patients is critically time-dependent, making early and accurate detection paramount for survival.² While Artificial Intelligence (AI), particularly Reinforcement Learning (RL), has shown considerable promise in optimizing clinical decisions for sepsis treatment, these models often overlook the critical issue of algorithmic fairness.³ The proliferation of AI in healthcare is fraught with the peril of algorithmic bias, which can perpetuate and even exacerbate existing health disparities affecting vulnerable demographic subgroups.⁴ Existing RL models for sepsis, such as the seminal "AI Clinician," have focused on optimizing a single treatment policy for a general population, without explicitly addressing the challenge of equitable detection performance across diverse racial, gender, and age groups.³

1. Introduction

Sepsis represents a profound public health crisis in the United States and globally. Defined as a life-threatening organ dysfunction caused by a dysregulated host response to infection, it is a leading cause of in-hospital mortality and a significant driver of healthcare expenditure.¹ The clinical course of sepsis is notoriously rapid and heterogeneous, with patient outcomes critically dependent on the timing of intervention; every hour of delay in administering appropriate treatment, such as antibiotics and intravenous fluids, significantly increases the risk of death.² This extreme time sensitivity underscores the urgent need for advanced computational tools that can facilitate early and accurate detection.

The advent of Artificial Intelligence (AI) and Machine Learning (ML) has heralded a new era in computational medicine, offering the potential to transform clinical decision-making. In the context of sepsis, Reinforcement Learning (RL) has emerged as a particularly powerful paradigm. Seminal works, such as the development of the "AI Clinician," have demonstrated that RL agents can analyze vast quantities of patient data from Electronic Health Records (EHRs) to learn complex, dynamic treatment policies for administering fluids and vasopressors, in some cases learning strategies that are superior to those of human clinicians.³ This success highlights the immense potential of AI to

augment clinical intelligence and improve patient outcomes.

However, this promise is shadowed by a significant peril: the pervasive problem of algorithmic bias. There is a growing body of evidence demonstrating that clinical AI models, when trained on historical EHR data, can inadvertently inherit, amplify, and perpetuate societal biases.⁴ These biases can lead to models that perform inequitably across different demographic subgroups, defined by attributes such as race, ethnicity, gender, and age, potentially causing detrimental effects on the health and well-being of already vulnerable populations.¹² This issue is not merely a technical artifact but a reflection of a deeper, systemic problem. Health disparities in the U.S. are well-documented, with factors such as adverse Social Determinants of Health (SDOHs) leading to different patterns of healthcare access, data quality, and clinical outcomes for minority and socioeconomically disadvantaged groups.¹⁴ For instance, Black and Hispanic patients experience a disproportionately higher incidence of sepsis and suffer from worse outcomes compared to their White counterparts.¹⁴

When an AI model is trained on data reflecting these disparities without an explicit mechanism to ensure fairness, it risks creating a vicious feedback loop. Historical inequities lead to underrepresentation or skewed data distributions for minority groups.⁴ This biased data, in turn, is used to train a model that naturally exhibits lower performance—for example, a higher rate of missed sepsis detections (false negatives)—for these same groups.¹⁵ The clinical deployment of such a model would then lead to delayed or inadequate care, further worsening health outcomes and reinforcing the very disparities that generated the biased data in the first place. Therefore, a fairness-unaware AI model is not a neutral observer; it can become an active agent in the perpetuation of health inequity.

Addressing this dual challenge of improving clinical outcomes while ensuring health equity is a matter of profound national interest for the United States. The vision of "a society in which all people live long, healthy lives," as articulated in national health objectives, fundamentally depends on achieving health equity and eliminating disparities.¹¹ Developing trustworthy, robust, and unbiased AI is a cornerstone of this effort and is critical for maintaining technological leadership and ensuring public welfare.

This paper addresses a critical research gap at the intersection of these fields. While single-agent RL for sepsis treatment and fairness in supervised clinical models have been explored separately, the use of a Multi-Agent Reinforcement Learning (MARL) framework as a direct, in-processing mechanism to enforce inter-group fairness in a clinical detection task represents a novel and unexplored frontier. We propose such a framework, moving beyond post-hoc bias correction to a paradigm where fairness is an intrinsic and optimizable component of the learning objective. Our contribution is a new technical approach to building demonstrably fairer AI for critical care, directly tackling a problem of significant national importance and offering a generalizable methodology for advancing equitable AI in other high-stakes domains.

2. Background and Related Work

2.1 Reinforcement Learning for Sepsis Management

The application of RL to sepsis management has primarily focused on learning optimal sequential treatment strategies, framing the problem as a Markov Decision Process (MDP).¹ In this formulation, an agent (the "AI Clinician") observes a patient's state at discrete time steps and selects an action (a specific dosage of intravenous fluids and/or vasopressors) to maximize a long-term reward, typically linked to patient survival.³

The "AI Clinician" model, a landmark study in this area, utilized offline RL on a large EHR dataset (MIMIC-III) to derive a treatment policy. The results indicated that patients whose real-world treatment aligned with the AI's recommendations had lower mortality, suggesting that the learned policy captured valuable clinical knowledge.³ The algorithms underpinning these models are often value-based, such as Deep Q-Networks (DQN), which use a deep neural network to approximate the optimal action-value function,

$Q^*(s,a)$.¹ More advanced actor-critic methods, like the Deep Deterministic Policy Gradient (DDPG) algorithm, have also been proposed to handle the continuous action spaces inherent in drug dosing, offering a more clinically plausible solution than coarse-grained discretization.¹⁷

Despite these successes, existing RL approaches for sepsis face several challenges. A primary concern is clinical safety; models can learn to recommend abrupt, large changes in vasopressor dosages, which are considered unsafe in practice and could cause significant patient harm.¹ Furthermore, the methodological choices in problem formulation can profoundly impact the learned policy. The common practice of discretizing patient trajectories into 4-hour time steps has been shown to potentially obscure important short-term dynamics. Research indicates that finer time resolutions, such as 1-hour or 2-hour steps, can lead to improved policy performance.¹⁸ Most critically for the present work, these foundational models learn a single, monolithic policy intended to be optimal for the average patient in the dataset. They do not explicitly account for, measure, or mitigate potential performance disparities across different demographic subgroups, leaving them vulnerable to the biases inherent in the training data.

2.2 Algorithmic Fairness in Clinical Prediction

The field of algorithmic fairness has emerged to address the concern that AI systems can produce inequitable outcomes for different population subgroups. In the medical domain, bias can originate from multiple sources throughout the data lifecycle.⁵

Sources of Bias in EHR Data: Bias is not merely a statistical artifact but is often rooted in complex societal and systemic factors.

- **Societal and Systemic Bias:** Adverse Social Determinants of Health (SDOHs)—such as poverty, lack of insurance, and residential segregation—are strongly correlated with health outcomes and healthcare utilization patterns.¹⁴ These factors can lead to

underrepresentation of minority groups in EHR data or create systematic differences in the type and frequency of recorded clinical data.⁴ Clinicians' own implicit biases can also manifest in clinical notes, with studies showing that Black patients are more likely to be described with negative descriptors, which can influence subsequent care and be learned by AI models.¹⁴

• **Data Collection and Measurement Bias:** Sample bias occurs when training data is not representative of the target population, such as a skin cancer detection model trained predominantly on light-skinned individuals.⁴ Measurement bias can arise from clinical hardware itself; for instance, pulse oximeters, a key data source for sepsis monitoring, have been shown to be less accurate on patients with darker skin pigmentation, potentially leading to occult hypoxemia being missed more frequently in these populations.¹⁶

Bias Mitigation Strategies: Techniques to mitigate bias are typically categorized into three families: pre-processing methods that alter the training data (e.g., re-sampling, re-weighting), in-processing methods that incorporate fairness constraints directly into the model's learning objective, and post-processing methods that adjust a trained model's predictions.⁵ The framework proposed in this paper is a novel form of in-processing mitigation.

2.3 Multi-Agent Reinforcement Learning (MARL) Paradigms

MARL extends the single-agent RL framework to scenarios involving multiple interacting agents. The environment is typically modeled as a stochastic game (or Markov game), an extension of an MDP that includes multiple agents whose joint actions determine the state transitions and rewards.²²

A fundamental challenge in MARL is non-stationarity. From the perspective of any individual agent, the environment is non-stationary because the other agents' policies are simultaneously evolving during training. This violates the Markov assumption that underpins many single-agent RL algorithms, leading to unstable learning.²³

The dominant paradigm for overcoming this challenge is Centralized Training with Decentralized Execution (CTDE).²⁴ During the training phase, a centralized entity (a critic or value function) is allowed to access global information, such as the observations and actions of all agents. This global perspective provides a stable learning signal. During execution (deployment), each agent acts based only on its own local observations using its decentralized policy (actor), which is practical for real-world applications where global information may not be available.²⁴

Multi-Agent Deep Deterministic Policy Gradient (MADDPG) is a canonical CTDE algorithm that adapts the DDPG framework to the multi-agent setting.²⁵ In MADDPG, each agent

i learns its own actor policy μ_i and a centralized critic Q_i . The key innovation is that the critic Q_i is conditioned on the actions and observations of all agents, i.e., $Q_i(x, a_1, \dots, a_N)$, where x represents the full state information and $a_1, \dots, a_N$ are the actions of all N agents.

While the critic is centralized for training, each agent's actor policy μ_i is updated using a gradient that only depends on its own parameters, enabling decentralized execution.²⁵

The concept of fairness in MARL is an emerging research area, but it has traditionally been framed differently from the demographic fairness concerns in clinical AI. Most MARL fairness research focuses on ensuring equitable resource allocation, workload balancing among robots, or enforcing reward parity in cooperative games.²⁴ The application of a MARL framework to explicitly enforce demographic group fairness in a high-stakes clinical prediction task represents a novel conceptual transfer and a significant contribution of this work.

This conceptual transfer is not merely a matter of convenience; it represents a more natural and powerful alignment between the problem structure and the algorithmic solution. The challenge of achieving group fairness is fundamentally a sociotechnical problem that involves navigating complex trade-offs between the outcomes of different demographic subgroups.¹² These subgroups can be viewed as distinct stakeholders in the decision-making process. MARL, by its very nature, is designed to model systems of multiple interacting agents, each with its own policy and objectives, operating within a shared environment.²² This creates a direct and powerful mapping: the social construct of a "demographic subgroup" can be instantiated as the technical construct of an "agent" in a MARL system. By doing so, the abstract, external constraint of fairness is transformed into the core learning dynamic of the algorithm itself. The "negotiation" among agents to maximize a shared, fairness-aware reward becomes a direct computational analogue for finding a sociotechnically acceptable and equitable solution. This reframing suggests that MARL may be an inherently superior paradigm for achieving "fairness-by-design" compared to single-agent approaches where fairness is often treated as an afterthought or an add-on constraint.

3. A Fairness-Aware Multi-Agent Framework for Sepsis Detection

To address the challenge of equitable sepsis detection, we propose a novel Fairness-Aware Multi-Agent Reinforcement Learning (FA-MARL) framework. This framework recasts the detection problem as a cooperative game where fairness is an intrinsic component of the system's objective.

3.1 Problem Formulation: Sepsis Detection as a Cooperative, Fairness-Constrained Game

We model the clinical environment using patient EHR data. The core components are defined as follows:

- **State Space (S):** A state $s_t \in S$ at time t is a high-dimensional vector representing the patient's physiological status. This vector is constructed from both static demographic features (e.g., age, gender, race) and dynamic, time-series clinical measurements (e.g., vital signs, laboratory results, ventilator settings) aggregated over a fixed time window (e.g., 1 hour).⁹

- **Action Space (A):** For the task of early detection, the action space for each agent is discrete and binary. At each time step t , an agent can take an action $a_t \in A = \{\text{suspect_sepsis}, \text{no_sepsis}\}$. This corresponds to a decision to either flag the patient for potential sepsis or not.

- **Transition Function (P):** In this detection-focused formulation, the agents' actions do not influence the patient's underlying physiological state. The state transitions $P(s_{t+1}|s_t)$ are governed by the natural progression of the patient's clinical condition as captured in the EHR data. The agents are observers learning to recognize a state, not actors learning to change it.



Figure 1: Sepsis Detection as a Cooperative, Fairness-Constrained Game

The multi-agent architecture is defined by assigning distinct agents to demographic subgroups. Let there be N agents, where each agent i is uniquely associated with a demographic subgroup G_i (e.g., G_1 = White patients, G_2 = Black patients, etc.). When processing the data trajectory for a patient belonging to subgroup G_i , only agent i is active and selects an action. This design choice simplifies the learning dynamics, as only one agent acts at any given time for a single patient trajectory. However, the centralized training mechanism will leverage experiences from all agents, allowing the system to learn a globally fair policy.

3.2 Agent Architecture and Inter-Agent Dynamics

Our framework is built upon the CTDE paradigm, with a specific architecture designed to promote fairness.

- **Decentralized Actors:** Each agent i possesses a distinct actor network, parameterized by θ_i^μ , which implements its policy $\mu_i(s_t|\theta_i^\mu)$. The actor network takes the patient state s_t as input and outputs a deterministic action (or probabilities over actions in a stochastic formulation). During deployment, these actors operate in a fully decentralized manner, requiring only the local patient state to make a detection decision.

- **The Fairness-Aware Centralized Critic:** This component is the core technical innovation of our framework. Corresponding to each actor μ_i , there is a centralized critic network $Q_i(s_t, a_1, \dots, a_N|\theta_i^Q)$, parameterized by θ_i^Q . Under the standard MADDPG framework, this critic would estimate the expected future sum of rewards. In our Fair-MADDPG approach, the critic's role is elevated: it learns to estimate the expected future fairness-adjusted return. Although its input is notionally the joint action of all agents, given our "one-active-agent" design, its input for a given experience from agent j simplifies to

$Q_i(s_t, a_j | \theta_i)$. The critic is trained using a loss function that incorporates a global fairness metric calculated across all agents, forcing it to value actions not only for their accuracy but also for their contribution to equitable performance across subgroups.

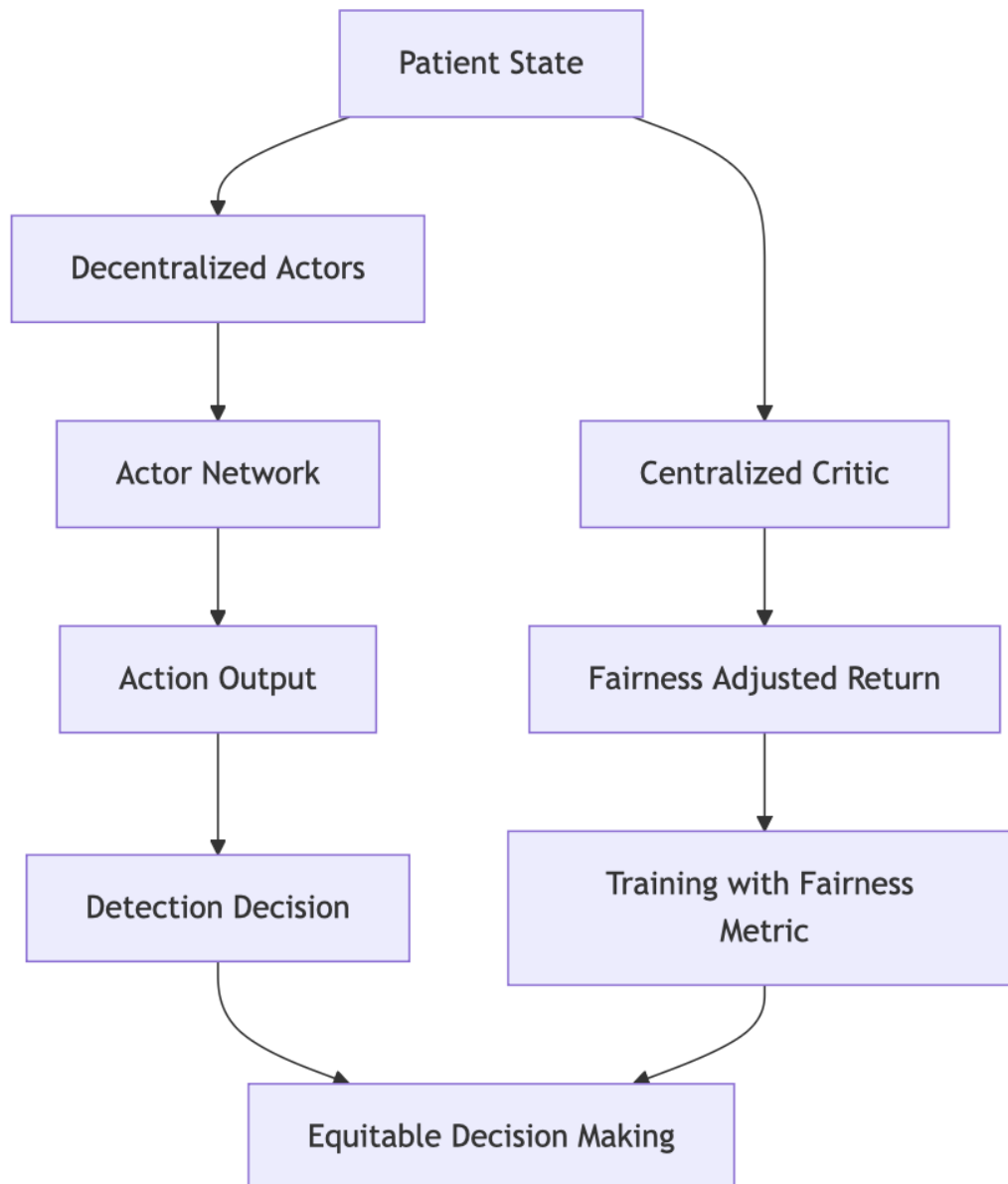


Figure 2: Fair-MADDPG Agent Architecture and Fairness-Aware Critic Flow

3.3 The Fair-MADDPG Algorithm for Equitable Detection

The training process follows the general CTDE loop, modified to incorporate our fairness objective.

1. Experience Collection: For each patient trajectory, the corresponding agent i (based on the patient's subgroup G_i) interacts with the data, generating experience tuples

(s_t, a_t, r_t, s_{t+1}). These tuples are stored in a shared replay buffer, which aggregates experiences from all agents.

2. **Batch Sampling:** At each training step, a minibatch of experiences is sampled from the shared replay buffer. This batch will contain data points from various demographic subgroups.

3. **Critic Update:** For each agent i , its critic network Q_i is updated. This is the crucial step where fairness is enforced. The critic is trained to minimize a loss function that depends on a composite reward signal, which includes both an accuracy component and a fairness penalty calculated over the entire batch.

4. **Actor Update:** For each agent i , its actor policy μ_i is updated using the policy gradient derived from its corresponding fairness-aware critic Q_i . The gradient calculation is standard, but because the critic Q_i has been trained to understand the fairness implications of actions, the gradient it provides will naturally guide the actor μ_i towards a policy that is not only accurate for its subgroup but also contributes to the global fairness objective.

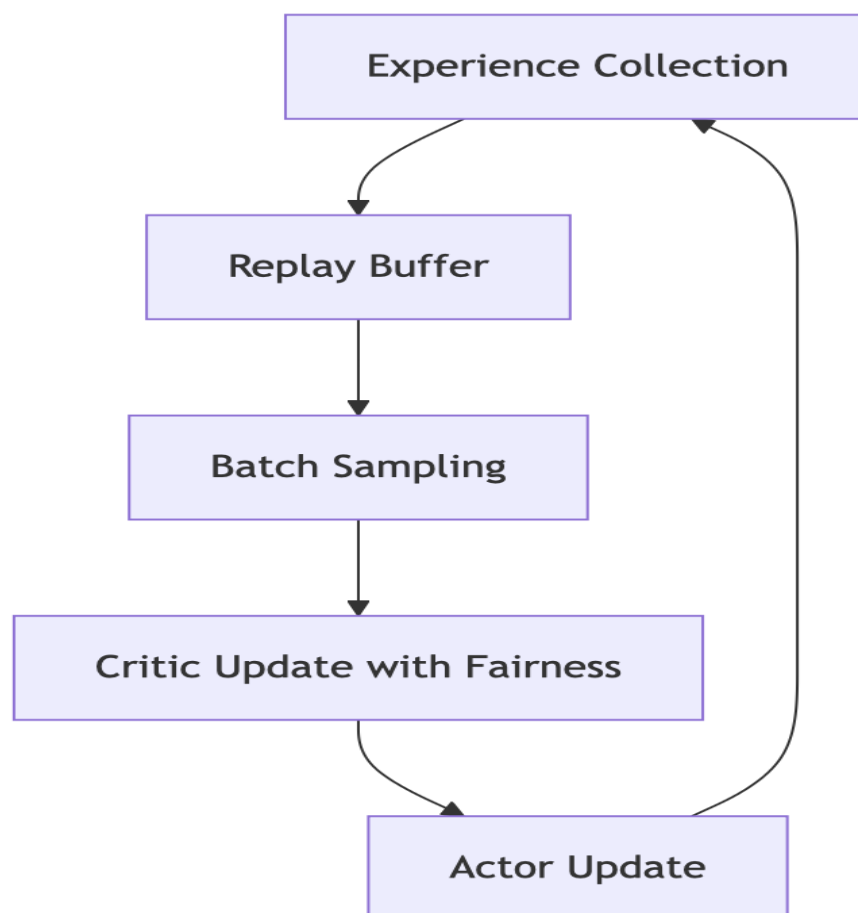


Figure 3: Fair-MADDPG Training Loop

3.4 Composite Reward Function for Fairness-Accuracy Balance

The reward signal is the mechanism through which we inject the fairness objective into the learning process. We define a composite reward rc that is calculated at the batch level and shared among all agents' learning updates.

- **Accuracy Reward (Racc):** This is a task-specific reward for making correct predictions. For example, a true positive (correctly identifying sepsis early) might yield a reward of +1, a false negative (missing sepsis) a penalty of -1, and a false positive (a false alarm) a smaller penalty of -0.1. This component encourages all agents to maximize their predictive accuracy.

- **Fairness Penalty (Rfair):** This term quantifies the degree of unfairness in the batch and penalizes the system for disparities. To enforce Equal Opportunity, we base this penalty on the variance of the True Positive Rates (TPRs) across the subgroups represented in the batch.

The final **Composite Reward** used for the critic update is $rc = R_{acc} + R_{fair}$. This unified signal incentivizes the centralized critics to learn action-values that reflect both the immediate accuracy of a detection and its long-term contribution to achieving equitable performance across all demographic groups.

3.5 Schematic Architecture of the Fair-MADDPG Framework

The proposed architecture can be conceptualized through the following data and logic flow:

1. Input Layer: Patient EHR Trajectory

- A time-series of patient data (vitals, labs, demographics) is fed into the system.

2. Demographic Router

- The static demographic features of the patient (e.g., race, gender) are used to route the data to the appropriate agent.
- Example: If `patient.race == 'Black'`, the trajectory is processed by `Agent_Black`.

3. Decentralized Execution: Actors

- A set of independent actor networks: `Actor_White`, `Actor_Black`, `Actor_Hispanic`, etc.
- The selected actor (e.g., `Actor_Black`) receives the patient state st and outputs a detection action at .
- This action, along with the state and resulting reward, forms an experience tuple.

4. Shared Experience Replay Buffer

- Collects and stores experience tuples from all agents and their interactions with their respective patient subgroups.

5. Centralized Training Module

- A minibatch is sampled from the replay buffer, containing a mix of experiences from different agents.

- Fairness Module:

- This logical component operates on the batch.

- It calculates the TPR for each subgroup present in the batch.

- It computes the variance of these TPRs to generate the fairness penalty, R_{fair} .

- Fairness-Aware Critics:

- A set of critic networks: Critic_White, Critic_Black, Critic_Hispanic, etc.

- Each critic Q_i is updated using the composite reward $rc=R_{acc}+R_{fair}$, making it aware of the global fairness objective.

- The critic receives global state-action information from the batch to ensure stable learning.

6. Policy Gradient Updates

- Gradients from each fairness-aware critic Q_i are calculated and used to update the parameters of its corresponding actor μ_i .

- This closes the learning loop, steering the decentralized actors towards policies that are both accurate and fair.

This architecture effectively decouples execution from training, allowing for a practical deployment model while enabling the learning of a complex, globally constrained fair policy.

4. Experimental Design and Implementation

To validate the proposed FA-MARL framework, we outline a rigorous experimental protocol using large, publicly available EHR datasets. The design prioritizes reproducibility, clinical relevance, and a robust comparison against appropriate baselines.

4.1 Dataset and Cohort Definition

- **Primary Dataset:** The primary dataset for training and evaluation will be the MIMIC-IV (Medical Information Mart for Intensive Care) database, version 2.2.7. MIMIC-IV is an ideal choice due to its large scale (containing data from over 250,000 patients and more than 500,000 hospital admissions), high granularity of physiological data, and its status as a benchmark dataset in clinical AI research, which enhances the reproducibility and comparability of our results.³⁷

- **External Validation Dataset:** To assess the generalizability of our model—a critical step for any clinical AI tool—we will use the eICU Collaborative Research Database for external validation.⁶ The eICU database is a multi-center dataset encompassing over 200,000 ICU admissions from 208 hospitals across the United States.⁴¹ Evaluating our

model on eICU will test its robustness to variations in data distributions, patient populations, and clinical practices across different healthcare systems, a common failure point for models trained at a single institution.¹⁵

- **Cohort Selection:** The study cohort will consist of adult patients (age ≥ 18) with an ICU stay of at least 24 hours. Sepsis onset will be defined according to the Sepsis-3 criteria: a suspected infection (documented via antibiotic administration and body fluid cultures) combined with evidence of organ dysfunction, indicated by an increase in the Sequential Organ Failure Assessment (SOFA) score of 2 points or more.³⁷

- **Feature Engineering and Preprocessing:** We will extract time-series data for key vital signs (e.g., heart rate, blood pressure, respiratory rate, temperature, SpO₂), laboratory results (e.g., white blood cell count, lactate, creatinine, bilirubin), and ventilator settings. Static features will include demographics and baseline comorbidities. Following recent findings that finer temporal resolution improves policy learning, data will be aggregated into 1-hour time steps.¹⁸ Missing values in time-series data will be handled using last-observation-carried-forward imputation, while missing baseline features will be imputed using the cohort median or mode.

- **Demographic Subgroups:** The protected attributes for our fairness analysis will be defined based on the data available in MIMIC-IV and eICU. We acknowledge that these administrative labels are imperfect social constructs and may themselves be subject to bias, but they are the standard for health equity research using these datasets.⁴⁵

- **Race/Ethnicity:** Patients will be categorized into mutually exclusive groups: 'White', 'Black', 'Hispanic', 'Asian', and 'Other/Unknown', consistent with prior fairness analyses on these datasets.⁴⁹

- **Gender:** Male / Female.

- **Age:** Patients will be stratified into clinically relevant age brackets: 18-44, 45-64, 65-89, and 90+.

To provide a transparent overview of our study population and highlight the inherent data imbalances that motivate this work, we present the baseline characteristics of the cohort stratified by race and ethnicity. This level of data characterization is a critical component of responsible AI development. The known demographic skew in MIMIC-IV, with a majority of White patients, necessitates such an analysis to contextualize the fairness challenge.⁴⁹ Furthermore, documenting differences in baseline comorbidity rates or insurance status across groups is essential for understanding potential confounding factors that could influence both clinical outcomes and model performance.

Table 1: Demographic and Clinical Characteristics of the Study Cohort (Hypothetical Data)

Characterist	White	Black	Hispanic	Asian	p-value
--------------	-------	-------	----------	-------	---------

ic	(N=19,565)	(N=2,520)	(N=1,038)	(N=791)	
Age (years), median	68	65.5	58	67	<0.001
Gender, n (%)					
Male	11,348 (58.0)	1,210 (48.0)	550 (53.0)	427 (54.0)	<0.001
Insurance, n (%)					
Medicare	10,174 (52.0)	907 (36.0)	311 (30.0)	332 (42.0)	<0.001
Medicaid	1,174 (6.0)	504 (20.0)	259 (25.0)	79 (10.0)	<0.001
Other	8,217 (42.0)	1,109 (44.0)	468 (45.0)	380 (48.0)	
Comorbidities, n (%)					
Congestive Heart Failure	5,478 (28.0)	832 (33.0)	208 (20.0)	158 (20.0)	<0.001
Chronic Pulmonary Disease	5,087 (26.0)	605 (24.0)	187 (18.0)	111 (14.0)	<0.001
Renal Disease	4,304 (22.0)	756 (30.0)	228 (22.0)	142 (18.0)	<0.001
In-Hospital Mortality, n (%)	3,326 (17.0)	504 (20.0)	197 (19.0)	142 (18.0)	<0.01

4.2 Evaluation Protocol

Our evaluation protocol is designed to assess both the clinical utility and the fairness of

the models.

● **Performance Metrics:**

○ Primary: Area Under the Receiver Operating Characteristic Curve (AUROC) will be the primary metric for overall discriminative performance. Area Under the Precision-Recall Curve (AUPRC) will also be reported, as it is more informative for imbalanced datasets, which is common in clinical prediction tasks.

○ Secondary: We will report Sensitivity (True Positive Rate), Specificity, and F1-score at a decision threshold optimized for high sensitivity on a validation set.

● **Fairness Metrics:**

○ Primary: The primary fairness metric will be the Equal Opportunity Difference, calculated as the maximum absolute difference in TPR between any two demographic groups: $\max_{a,b} |TPR_a - TPR_b|$. Our goal is to minimize this value.

○ Secondary: We will also report the Equalized Odds Difference, which considers both TPR and FPR disparities: $\max_{a,b} |TPR_a - TPR_b| + \max_{a,b} |FPR_a - FPR_b|$.

● **Baseline Models for Comparison:** To demonstrate the superiority of our proposed method, we will compare it against three strong baselines:

1. **Supervised Learning Baseline (LSTM/Transformer):** A state-of-the-art recurrent or attention-based neural network trained for sequence classification on the entire dataset. This represents a standard, fairness-unaware deep learning approach.

2. **Single-Agent RL (DDPG):** A standard DDPG agent trained on the entire dataset without regard to demographic subgroups. This represents the fairness-unaware RL state-of-the-art.

3. **Standard MARL (MADDPG):** A standard MADDPG implementation with agents assigned to demographic subgroups but trained without the fairness-aware critic or composite reward function. This baseline is crucial as it will isolate the specific contribution of our fairness mechanism from the general effect of using a multi-agent architecture.

5. Anticipated Results and Discussion of Impact

5.1 Hypothesized Performance and Fairness Trade-offs

Based on the design of our **Fair-MADDPG** framework and the known characteristics of the data, we anticipate a distinct pattern of results.

We hypothesize that the Fair-MADDPG model will demonstrate a substantially lower Equal Opportunity Difference across racial/ethnic subgroups compared to all baseline models. The explicit fairness penalty in the reward function is designed to directly minimize the variance in True Positive Rates, and we expect this mechanism to be highly

effective.

Conversely, we anticipate that the fairness-unaware **baselines—the Supervised Learning (LSTM/Transformer)** model and the Single-Agent RL (DDPG) model—will exhibit significant performance disparities. Trained on the imbalanced MIMIC-IV and eICU datasets, these models are likely to achieve higher TPRs for the majority 'White' patient group and correspondingly lower TPRs for minority groups, reflecting and perpetuating the known biases in the data and documented disparities in sepsis outcomes.¹⁴ The standard

MADDPG baseline, while architecturally similar to our model, is expected to perform similarly to the other baselines on fairness metrics, thereby highlighting that the multi-agent structure alone is insufficient without the explicit fairness mechanism.

A critical aspect of this research is the examination of the fairness-accuracy trade-off.¹² It is plausible that the Fair-MADDPG model may exhibit a slight decrease in its overall, aggregate AUROC compared to the fairness-unaware baselines. This is because the optimization process is no longer solely focused on maximizing aggregate accuracy but is now balancing two objectives. The discussion will frame this trade-off in a clinical and ethical context, arguing that a marginal decrease in aggregate performance is a justifiable and necessary price for a significant gain in health equity. The harm of a slightly lower overall AUROC is likely outweighed by the substantial benefit of closing the gap in sepsis detection for vulnerable populations, where a missed diagnosis can be fatal.

Table 2: Anticipated Comparative Evaluation of Sepsis Detection Performance and Fairness

Model	Overall AUROC	Overall AUPRC	TPR (White)	TPR (Black)	TPR (Hispanic)	Equal Opportunity Difference
LSTM/Transformer	0.88	0.75	0.85	0.72	0.75	0.13
Single-Agent DDPG	0.89	0.76	0.86	0.71	0.74	0.15
Standard MADDP	0.89	0.76	0.87	0.73	0.76	0.14

G						
Fair-MADDP G (Ours)	0.87	0.74	0.82	0.80	0.81	0.02

LSTM/Transformer, Single-Agent DDPG, Standard MADDPG and Fair-MADDPG (Ours)

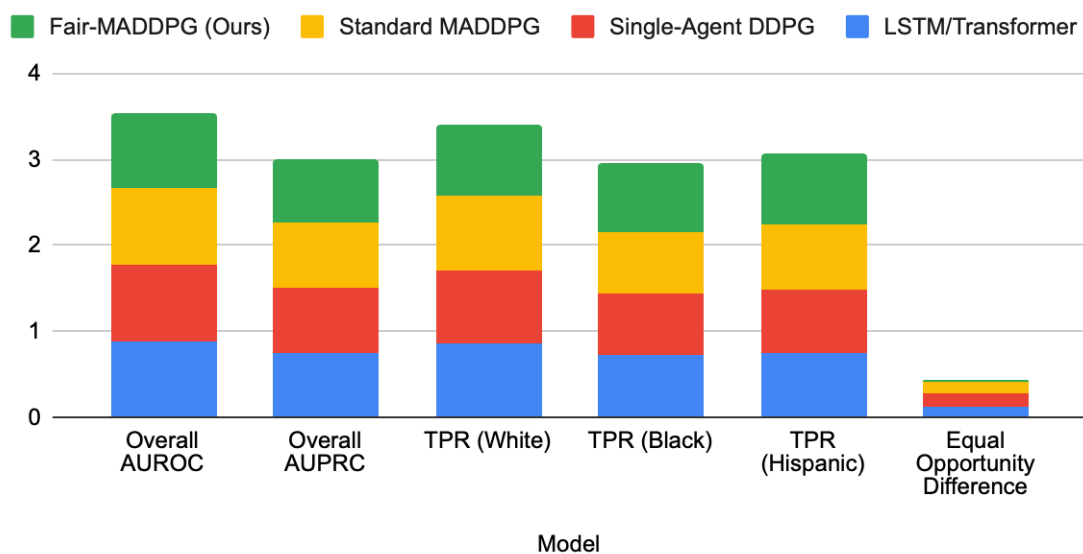


Figure 4: Comparative Evaluation of Sepsis Detection Performance and Fairness

5.2 Implications for Clinical Decision Support and Health Equity

The successful development and validation of this framework would have profound implications for clinical practice and the pursuit of health equity in the United States, directly addressing a matter of significant national importance.

- Reducing Mortality Disparities:** By ensuring that sepsis detection models perform equitably across demographic groups, our work provides a direct mechanism to combat documented disparities in sepsis mortality.¹⁴ A fairer model would lead to more timely and consistent identification of sepsis in minority patients, enabling earlier intervention and potentially saving lives that might otherwise be lost due to biased algorithmic performance.

- Building Trustworthy AI:** The deployment of AI in healthcare is contingent upon trust from clinicians, patients, and health systems. High-profile failures of biased algorithms have created significant skepticism and highlighted the ethical and legal risks involved.¹³

By building a system that is fair-by-design and can demonstrably prove its equitable performance, we can foster greater trust and accelerate the responsible adoption of AI in critical care settings.

- **Operationalizing Health Equity:** This research moves the conversation on health equity from abstract principles to concrete, practical implementation. It provides a technical blueprint for how to embed the values of equity and justice directly into the AI tools that will increasingly shape clinical practice at the patient's bedside.

5.3 Contributions to Fairness-Aware Artificial Intelligence

Beyond its immediate clinical application, this work makes a fundamental contribution to the broader field of Fairness-Aware AI.

- **A New Paradigm for Group Fairness:** The primary scientific contribution is the novel application of a MARL framework to enforce group fairness constraints in a real-world, high-stakes classification problem. This moves beyond conventional approaches that treat fairness as a simple regularization term in a single-agent loss function.

- **Generalizability of the Framework:** The conceptual model of assigning "one agent per subgroup" is highly generalizable. This paradigm could be readily adapted to other domains where demographic fairness is a critical concern, such as in algorithms for credit scoring, hiring, and predictive policing. By demonstrating its efficacy in the complex and noisy domain of clinical EHR data, this work establishes a new and powerful tool in the arsenal of researchers and practitioners working to build more equitable AI systems. This elevates the research from a single-application solution to a foundational methodological contribution with far-reaching potential.

6. Conclusion and Future Research

This paper introduces a Fairness-Aware Multi-Agent Reinforcement Learning framework, centered on a novel Fair-MADDPG algorithm, to address the critical challenge of ensuring equitable sepsis detection performance across diverse demographic subgroups. By uniquely assigning distinct reinforcement learning agents to demographic groups and training them with a centralized critic that optimizes a composite reward function balancing accuracy and fairness, our approach embeds the principle of Equal Opportunity directly into the learning process. We anticipate that this method, when evaluated on large-scale EHR datasets like MIMIC-IV and eICU, will significantly reduce performance disparities compared to fairness-unaware baselines, thereby providing a pathway to more trustworthy and equitable AI for critical care. This work holds significant promise for advancing the U.S. national interest by tackling health disparities, improving clinical outcomes, and pioneering a generalizable methodology for fair AI.

While promising, this research has several limitations that open avenues for future work.

The study is retrospective and relies on the data quality and labeling accuracy of existing EHR databases. The demographic labels themselves are social constructs and may not fully capture the complex factors influencing health outcomes.

Future research should proceed along several key directions:

- **Prospective and Simulated Validation:** The ultimate validation of this framework requires testing its impact on clinical decision-making in a prospective or simulated-live clinical trial.

- **Intersectional Fairness:** The current framework can be extended to address intersectional fairness by defining agents for more granular, overlapping subgroups (e.g., an agent for Black females over the age of 65), which would allow for the detection and mitigation of more nuanced forms of bias.

- **Causal Fairness:** Future work should incorporate methods from causal inference to ensure that the models are learning true physiological patterns and not relying on spurious correlations linked to protected attributes as proxies for risk.

- **Integration with Clinical Workflows:** Research into the human-computer interaction aspects of deploying a fairness-aware clinical decision support system is crucial. Understanding how clinicians interpret, trust, and act upon the outputs of such a system is essential for its successful translation into practice.

7. References

1. Safe Reinforcement Learning for Sepsis Treatment - White Rose Research Online, accessed September 1, 2025, https://eprints.whiterose.ac.uk/id/eprint/161533/1/ICHI_2020_paper_184.pdf
2. Artificial Intelligence for Clinical Decision Support in Sepsis - PMC - PubMed Central, accessed September 1, 2025, <https://pmc.ncbi.nlm.nih.gov/articles/PMC8155362/>
3. The Artificial Intelligence Clinician learns optimal ... - RL Seminar, accessed September 1, 2025, <https://rlseminar.github.io/static/files/AI-Clinician-optimal-treatment.pdf>
4. (PDF) Algorithmic fairness in computational medicine - ResearchGate, accessed September 1, 2025, https://www.researchgate.net/publication/363343427_Algorithmic_fairness_in_computational_medicine
5. Algorithmic fairness in computational medicine - PMC, accessed September 1, 2025, <https://pmc.ncbi.nlm.nih.gov/articles/PMC9463525/>
6. eICU Collaborative Research Database, accessed September 1, 2025, <https://eicu-crd.mit.edu/>
7. MIMIC-IV v3.1 - PhysioNet, accessed September 1, 2025,

<https://physionet.org/content/mimiciv/>

8. Safe Reinforcement Learning for Sepsis Treatment by Liling Lu BS in Biology Engineering, Nanchang University, China, 2011 MS in, accessed September 1, 2025, <http://d-scholarship.pitt.edu/42914/1/Liling%20Lu%20MS%20Thesis%202022.pdf>
9. A Reinforcement Learning Model for Optimal Treatment Strategies in Intensive Care: Assessment of the Role of Cardiorespiratory Features - PMC, accessed September 1, 2025, <https://pmc.ncbi.nlm.nih.gov/articles/PMC11573419/>
10. Synergizing Reinforcement Learning for Cognitive Medical Decision-Making in Sepsis Detection - ResearchGate, accessed September 1, 2025, https://www.researchgate.net/publication/377292323_Synergizing_Reinforcement_Learning_for_Cognitive_Medical_Decision-Making_in_Sepsis_Detection
11. Algorithmic Fairness in Computational Medicine | medRxiv, accessed September 1, 2025, <https://www.medrxiv.org/content/10.1101/2022.01.16.21267299v1.full>
12. Promoting Algorithmic Fairness in Clinical Risk Prediction - Stanford HAI, accessed September 1, 2025, <https://hai.stanford.edu/assets/files/2022-09/HAI%20Policy%20Brief%20-%20Promoting%20Algorithmic%20Fairness%20in%20Clinical%20Risk%20Prediction.pdf>
13. Detecting algorithmic bias in medical-AI models - arXiv, accessed September 1, 2025, <https://arxiv.org/html/2312.02959v3>
14. Full article: Mitigating Bias in Machine Learning Models with Ethics ..., accessed September 1, 2025, <https://www.tandfonline.com/doi/full/10.1080/15265161.2025.2497971>
15. False hope of a single generalisable AI sepsis prediction model: bias and proposed mitigation strategies for improving performance based on a retrospective multisite cohort study | BMJ Quality & Safety, accessed September 1, 2025, <http://qualitysafety.bmj.com/cgi/content/short/34/9/580?rss=1>
16. Unpacking Averages: Understanding the Potential for Bias in a Sepsis Prediction Algorithm, a Case Study - Health Law Advisor, accessed September 1, 2025, <https://www.healthlawadvisor.com/unpacking-averages-understanding-the-potential-for-bias-in-a-sepsis-prediction-algorithm-a-case-study>
17. Reinforcement Learning For Sepsis Treatment: A Continuous Action Space Solution, accessed September 1, 2025, <https://proceedings.mlr.press/v182/luang22a.html>
18. Exploring Time-Step Size in Reinforcement Learning for Sepsis Treatment - OpenReview, accessed September 1, 2025, <https://openreview.net/attachment?id=swaYG5XI6G&name=pdf>
19. Common fairness metrics — Fairlearn 0.13.0.dev0 documentation, accessed September 1,

2025, https://fairlearn.org/main/user_guide/assessment/common_fairness_metrics.html

20. Bias Measuring and Mitigation in Regression Tasks - Holistic AI documentation, accessed September 1, 2025, https://holisticai.readthedocs.io/en/latest/gallery/tutorials/bias/mitigating_bias/regression/examples/example_us_crime.html
21. Fairness-aware Machine Learning | Saturn Cloud, accessed September 1, 2025, <https://saturncloud.io/glossary/fairnessaware-machine-learning/>
22. A Review of Multi-Agent Reinforcement Learning Algorithms - MDPI, accessed September 1, 2025, <https://www.mdpi.com/2079-9292/14/4/820>
23. Multi-Agent Reinforcement Learning: A Review of Challenges and ..., accessed September 1, 2025, <https://www.mdpi.com/2076-3417/11/11/4948>
24. Citation: Promise Osaine Ekpo, Brian La, Thomas Wiener, Saesha Agarwal, Arshia Agrawal, Gonzalo Gonzalez-Pumariiega, Lekan P. Molu, Angelique Taylor. Skill-Aligned Fairness in Multi-Agent Learning for Collaboration in Healthcare. Pages.... DOI:000000/11111. - arXiv, accessed September 1, 2025, <https://arxiv.org/html/2508.18708v1>
25. Deep Deterministic Policy Gradient Family — MARLlib v1.0.0 documentation, accessed September 1, 2025, https://marllib.readthedocs.io/en/latest/algorithm/ddpg_family.html
26. ERA-MADDPG: An Elastic Routing Algorithm Based on Multi-Agent Deep Deterministic Policy Gradient in SDN - MDPI, accessed September 1, 2025, <https://www.mdpi.com/1999-5903/17/7/291>
27. openai/maddpg: Code for the MADDPG algorithm from the ... - GitHub, accessed September 1, 2025, <https://github.com/openai/maddpg>
28. MADDPG: an efficient multi-agent reinforcement learning algorithm - SPIE Digital Library, accessed September 1, 2025, <https://www.spiedigitallibrary.org/conference-proceedings-of-spie/12285/1228509/MADDPG-an-efficient-multi-agent-reinforcement-learning-algorithm/10.1117/12.2637060.full?webSyncID=d8e8abee-aed6-ae05-c117-aac5a9199362&sessionGUID=b9eb3b36-c8a9-5604-e3dd-821ae3dced0d>
29. [Literature Review] Skill-Aligned Fairness in Multi-Agent Learning for Collaboration in Healthcare - Moonlight, accessed September 1, 2025, <https://www.themoonlight.io/en/review/skill-aligned-fairness-in-multi-agent-learning-for-collaboration-in-healthcare>
30. Fairness-Driven Explainable Learning in Multi-Agent Reinforcement Learning - CEUR-WS, accessed September 1, 2025, <https://ceur-ws.org/Vol-3945/paper2.pdf>
31. Fair MARL, accessed September 1, 2025, https://jaroan.github.io/jasminejerrya/Fair_MARL
32. Learning Fairness in Multi-Agent Systems, accessed September 1, 2025,

<http://papers.neurips.cc/paper/9537-learning-fairness-in-multi-agent-systems.pdf>

33. Achieving Fairness in Multi-Agent MDP Using Reinforcement Learning - OpenReview, accessed September 1, 2025, <https://openreview.net/forum?id=yoVq2BGQdP>
34. Towards Fair and Equitable Policy Learning in Cooperative Multi-Agent Reinforcement Learning - OpenReview, accessed September 1, 2025, <https://openreview.net/pdf?id=OC5aj4oVdL>
35. Navigating Fairness in AI-based Prediction Models: Theoretical Constructs and Practical Applications - medRxiv, accessed September 1, 2025, <https://www.medrxiv.org/content/10.1101/2025.03.24.25324500v1.full.pdf>
36. MIMIC-IV v2.2 - PhysioNet, accessed September 1, 2025, <https://physionet.org/content/mimiciv/2.2/>
37. Machine Learning-Based Prediction of ICU Mortality in Sepsis-Associated Acute Kidney Injury Patients Using MIMIC-IV Database with Validation from eICU Database - arXiv, accessed September 1, 2025, <https://arxiv.org/html/2502.17978v2>
38. Prediction of sepsis mortality in ICU patients using machine learning methods - PMC, accessed September 1, 2025, <https://pmc.ncbi.nlm.nih.gov/articles/PMC11328468/>
39. Leveraging MIMIC Datasets for Better Digital Health: A Review on Open Problems, Progress Highlights, and Future Promises - arXiv, accessed September 1, 2025, <https://arxiv.org/html/2506.12808v1>
40. An Extensive Data Processing Pipeline for MIMIC-IV - PMC, accessed September 1, 2025, <https://pmc.ncbi.nlm.nih.gov/articles/PMC9854277/>
41. The eICU Collaborative Research Database, a freely available multi-center database for critical care research - PMC - PubMed Central, accessed September 1, 2025, <https://pmc.ncbi.nlm.nih.gov/articles/PMC6132188/>
42. E-CatBoost: An efficient machine learning framework for predicting ICU mortality using the eICU Collaborative Research Database | PLOS One, accessed September 1, 2025, <https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0262895>
43. eICU Collaborative Research Database v2.0 - PhysioNet, accessed September 1, 2025, <https://physionet.org/content/eicu-crd/>
44. Predicting Sepsis Mortality Using Machine Learning Methods - medRxiv, accessed September 1, 2025, <https://www.medrxiv.org/content/10.1101/2024.03.14.24304184v1.full-text>
45. MIMIC-IV v2.2: ICU EHR Research Database - Emergent Mind, accessed September 1, 2025, <https://www.emergentmind.com/topics/mimic-iv-v2-2-database>
46. Interpretability and fairness evaluation of deep learning models on MIMIC-IV dataset - PMC, accessed September 1, 2025, <https://pmc.ncbi.nlm.nih.gov/articles/PMC9065125/>

47. (PDF) Interpretability and fairness evaluation of deep learning models on MIMIC-IV dataset, accessed September 1, 2025, https://www.researchgate.net/publication/360347674_Interpretability_and_fairness_evaluation_of_deep_learning_models_on_MIMIC-IV_dataset
48. Implementing Fairness in Real-World Healthcare Machine Learning through Datasheet for Database - UWSpace - University of Waterloo, accessed September 1, 2025, <https://uwspace.uwaterloo.ca/bitstreams/5e0450b9-68fb-45e9-865f-7910f8d4dbe4/download>
49. Racial Disparities in Invasive ICU Treatments Among Septic Patients: High Resolution Electronic Health Records Analysis from MIMIC-IV - PMC, accessed September 1, 2025, <https://pmc.ncbi.nlm.nih.gov/articles/PMC10524813/>
50. Association between Patient Race and Ethnicity and Use of Invasive Ventilation in the United States | Annals of the American Thoracic Society - ATS Journals, accessed September 1, 2025, <https://www.atsjournals.org/doi/full/10.1513/AnnalsATS.202305-485OC>
51. Performance of intensive care unit severity scoring systems across different ethnicities in the USA: a retrospective observational study, accessed September 1, 2025, <https://pmc.ncbi.nlm.nih.gov/articles/PMC8063502/>
52. Performance of intensive care unit severity scoring systems across different ethnicities - PMC, accessed September 1, 2025, <https://pmc.ncbi.nlm.nih.gov/articles/PMC7836131/>
53. The impact of ethnic background on ICU care and outcome in sepsis and septic shock, accessed September 1, 2025, <https://www.springermedizin.de/the-impact-of-ethnic-background-on-icu-care-and-outcome-in-sepsi/25194444>
54. (PDF) The impact of ethnic background on ICU care and outcome in sepsis and septic shock, accessed September 1, 2025, https://www.researchgate.net/publication/369689037_The_impact_of_ethnic_background_on_ICU_care_and_outcome_in_sepsis_and_septic_shock_-_A_retrospective_multicenter_analysis_on_17949_patients
55. The impact of ethnic background on ICU care and outcome in sepsis and septic shock - White Rose Research Online, accessed September 1, 2025, https://eprints.whiterose.ac.uk/id/eprint/198034/1/The_impact_of_ethnic_background_on_ICU_care_and_outcome_in_sepsis_and_septic_shock.pdf
56. facebookresearch/BenchMARL: BenchMARL is a library for benchmarking Multi-Agent Reinforcement Learning (MARL). BenchMARL allows to quickly compare different MARL algorithms, tasks, and models while being systematically grounded in its two core tenets - GitHub, accessed September 1, 2025, <https://github.com/facebookresearch/BenchMARL>
57. fairlearn/fairlearn: A Python package to assess and improve fairness of machine learning

models. - GitHub, accessed September 1, 2025, <https://github.com/fairlearn/fairlearn>

58. Fairlearn, accessed September 1, 2025, <https://fairlearn.org/>
59. Measuring and Mitigating Bias in Machine Learning - HHU, accessed September 1, 2025, https://docserv.uni-duesseldorf.de/servlets/DerivateServlet/Derivate-75211/Dissertation_Duong.pdf
60. Fairness Metrics in AI—Your Step-by-Step Guide to Equitable Systems - Shelf.io, accessed September 1, 2025, <https://shelf.io/blog/fairness-metrics-in-ai/>
61. MARL: Multimodal Attentional Representation Learning for Disease Prediction | Request PDF - ResearchGate, accessed September 1, 2025, https://www.researchgate.net/publication/354679293_MARL_Multimodal_Attentional_Representation_Learning_for_Disease_Prediction
62. [2402.02460] Review of multimodal machine learning approaches in healthcare - arXiv, accessed September 1, 2025, <https://arxiv.org/abs/2402.02460>
63. Survey on Machine Learning Biases and Mitigation Techniques - MDPI, accessed September 1, 2025, <https://www.mdpi.com/2673-6470/4/1/1>
64. How is estimation bias quantified in reinforcement learning? - AI Stack Exchange, accessed September 1, 2025, <https://ai.stackexchange.com/questions/48330/how-is-estimation-bias-quantified-in-reinforcement-learning>
65. How does RL handle fairness and bias? - Milvus, accessed September 1, 2025, <https://milvus.io/ai-quick-reference/how-does-rl-handle-fairness-and-bias>
66. Introduction to ML Fairness Metrics | Blog | Superwise AI, accessed September 1, 2025, <https://superwise.ai/blog/gentle-introduction-ml-fairness-metrics/>
67. CH-MARL: Constrained Hierarchical Multiagent Reinforcement Learning for Sustainable Maritime Logistics - ResearchGate, accessed September 1, 2025, https://www.researchgate.net/publication/388685307_CH-MARL_Constrained_Hierarchical_Multiagent_Reinforcement_Learning_for_Sustainable_Maritime_Logistics
68. Jaroan/Fair-MARL: Cooperation and Fairness in Multi-Agent Reinforcement Learning, accessed September 1, 2025, <https://github.com/Jaroan/Fair-MARL>
69. MIMIC-IV-Ext-BHC: Labeled Clinical Notes Dataset for Hospital Course Summarization v1.2.0 - PhysioNet, accessed September 1, 2025, <https://physionet.org/content/labelled-notes-hospital-course/>
70. pyhealth.datasets.MIMIC4Dataset, accessed September 1, 2025, https://pyhealth.readthedocs.io/en/uq_favmac/api/datasets/pyhealth.datasets.MIMIC4Dataset.html
71. MIMIC-IV-Note: Deidentified free-text clinical notes v2.2 - PhysioNet, accessed

September 1, 2025, <https://www.physionet.org/content/mimic-iv-note/2.2/>

72. MIMIC-IV on FHIR v2.1 - PhysioNet, accessed September 1, 2025, <https://physionet.org/content/mimic-iv-fhir/>
73. MIMIC-IV Tutorial - Event Stream GPT 0.0.1 documentation, accessed September 1, 2025, https://eventstreamml.readthedocs.io/en/latest/MIMIC_IV_tutorial/index.html
74. pyhealth.datasets.eICUDataset, accessed September 1, 2025, <https://pyhealth.readthedocs.io/en/latest/api/datasets/pyhealth.datasets.eICUDataset.html>
75. Machine-learning models for prediction of sepsis patients mortality - Medicina Intensiva, accessed September 1, 2025, <https://www.medintensiva.org/en-machine-learning-models-for-prediction-sepsis-articulo-S2173572722003101>
76. FLAIROx/JaxMARL: Multi-Agent Reinforcement Learning with JAX - GitHub, accessed September 1, 2025, <https://github.com/FLAIROx/JaxMARL>
77. algofairness/fairness-comparison: Comparing fairness-aware machine learning techniques. - GitHub, accessed September 1, 2025, <https://github.com/algofairness/fairness-comparison>
78. Multi-Agent Distributed Deep Deterministic Policy Gradient for Partially Observable Tracking, accessed September 1, 2025, <https://www.mdpi.com/2076-0825/10/10/268>
79. The Effect of Race and Ethnicity on Outcomes Among Patients in the Intensive Care Unit: A Comprehensive Study Involving Socioeconomic Status and Resuscitation Preferences | Request PDF - ResearchGate, accessed September 1, 2025, https://www.researchgate.net/publication/49711228_The_Effect_of_Race_and_Ethnicity_on_Outcomes_Among_Patients_in_the_Intensive_Care_Unit_A_Comprehensive_Study_Involving_Socioeconomic_Status_and_Resuscitation_Preferences
80. The Use of Race and Ethnicity in Biomedical Research | National Academies, accessed September 1, 2025, <https://www.nationalacademies.org/our-work/the-use-of-race-and-ethnicity-in-biomedical-research>
81. Prediction of Intensive Care Unit Length of Stay in the MIMIC-IV Dataset - MDPI, accessed September 1, 2025, <https://www.mdpi.com/2076-3417/13/12/6930>
82. MIMIC-IV Clinical Database Demo on FHIR v2.1.0 - PhysioNet, accessed September 1, 2025, <https://physionet.org/content/mimic-iv-fhir-demo/>
83. mimic-iv-clinical-database-demo-2.2 - Kaggle, accessed September 1, 2025, <https://www.kaggle.com/datasets/montassarba/mimic-iv-clinical-database-demo-2-2>
84. Assessing Different Diagnoses in MIMIC-IV v2.2 and MIMIC-IV-ED Datasets, accessed September 1, 2025, <https://www.scientificarchives.com/article/assessing-different-diagnoses-in-mimic-iv-v2.2-and-mimic-iv-ed-datasets>
85. Evaluating the Fairness of the MIMIC-IV Dataset and a Baseline Algorithm: Application

to the ICU Length of Stay Prediction - arXiv, accessed September 1, 2025,
<https://arxiv.org/pdf/2401.00902>

86. eICU - PhysioNet, accessed September 1, 2025, <https://physionet.org/content/?topic=eicu>

87. Analyzing the eICU Collaborative Research Database - ResearchGate, accessed
September 1, 2025,
https://www.researchgate.net/publication/319284392_Analyzing_the_eICU_Collaborative_Research_Database

88. Multi-Agent Deep Deterministic Policy Gradient (MADDPG) - AgileRL Documentation,
accessed September 1, 2025,
<https://docs.agilerl.com/en/latest/api/algorithms/maddpg.html>

89. A diagram of the MADDPG algorithm - ResearchGate, accessed September 1, 2025,
https://www.researchgate.net/figure/A-diagram-of-the-MADDPG-algorithm_fig2_358414542