

Accountability Judgments and Ethical Decision Modules across Public Safety Robots

Richard Chun-Ho Wan*

Department of Mechanical and Aerospace Engineering, School of Engineering, Hong Kong University of Science and Technology, Hong Kong, Hong Kong SAR, China

Corresponding author: richard.ch.wan@hkust.edu.hk

Abstract

The integration of autonomous systems into public safety domains necessitates a profound understanding of how algorithmic decision-making architectures influence human perceptions of accountability. This paper investigates the intricate relationship between embedded ethical decision modules in public safety robots and the subsequent accountability judgments made by human stakeholders. By utilizing agent-based modeling combined with human-in-the-loop perceptual evaluations, this study simulates complex emergency scenarios where autonomous agents must resolve moral dilemmas. We contrast three distinct ethical architectures, namely utilitarian, deontological, and hybrid rule-based modules, to determine how the algorithmic rationale impacts the attribution of blame and responsibility. The findings indicate a significant divergence in accountability judgments based on the underlying ethical framework of the robot. Specifically, utilitarian modules, which optimize for the greatest overall survival rate, tend to diffuse perceived accountability among developers and operators, whereas deontological modules, which adhere to strict operational constraints, concentrate blame on the autonomous agent itself. This research provides empirical evidence bridging machine ethics and social psychology, offering a comprehensive framework for policymakers, legal scholars, and roboticists. The insights derived from our agent modeling simulations underscore the critical need for transparent ethical architectures in autonomous systems to ensure societal trust and legal clarity in public safety deployments.

Keywords: Public Safety Robots; Machine Ethics; Accountability Judgments; Agent-Based Modeling; Algorithmic Transparency

1. Introduction

1.1 The Ascendancy of Autonomous Public Safety Systems

The rapid advancement of artificial intelligence and robotics has fundamentally transformed the landscape of public safety and emergency management. Autonomous systems, ranging from unmanned aerial vehicles utilized in search and rescue operations to ground-based robotic units deployed in hazardous environments, are increasingly tasked with operating in highly complex, dynamic, and dangerous contexts. These public safety robots are no longer merely remote-controlled extensions of human operators; they are sophisticated autonomous agents capable of perceiving their environment, processing massive streams of sensory data, and executing actions with minimal human oversight. This paradigm shift promises enhanced efficiency, reduced risk to human first responders, and the ability to operate in conditions that are otherwise inaccessible to human personnel. However, the delegation of decision-making authority to autonomous machines in high-stakes environments introduces a myriad of technical, societal, and ethical challenges that have not been fully addressed by contemporary regulatory frameworks. As these systems are deployed in real-world scenarios, they inevitably encounter situations that require moral reasoning, forcing the autonomous agent to choose between competing ethical imperatives. The absence of a universal ethical framework for machines complicates the integration of these robots into society, as the algorithms governing their behavior must align with human moral intuitions to maintain public trust. The intersection of autonomous functionality and public safety thereby demands a rigorous examination of how these systems resolve ethical dilemmas and, crucially, how human observers interpret and evaluate these algorithmic decisions. Understanding the mechanisms through which public safety robots operate and the ethical architectures that guide their actions is paramount for ensuring that their deployment enhances societal well-being without compromising fundamental moral standards.

1.2 The Dilemma of Ethical Decision Modules and Accountability

At the core of autonomous operation in public safety is the ethical decision module, a specialized computational architecture designed to evaluate moral dilemmas and select an appropriate course of action. These modules are programmed using various philosophical frameworks, most notably utilitarianism, which seeks to maximize overall welfare, and deontology, which relies on adherence to universal moral rules. When a public safety robot confronts a situation where harm is

unavoidable, the underlying ethical decision module dictates how the robot prioritizes potential victims, allocates resources, or navigates the crisis. The critical issue arises when these algorithmic decisions result in negative outcomes, prompting human stakeholders to assign blame and accountability. The concept of accountability in human-robot interaction is inherently complex, as traditional legal and moral frameworks are predicated on human agency and intentionality. When a machine makes a decision that leads to injury or property damage, determining who is responsible becomes a convoluted process involving the robot manufacturer, the software developer, the human operator, and the autonomous agent itself. Previous research [1] highlights that humans possess a psychological propensity to anthropomorphize advanced robots, often attributing a degree of moral agency to the machine. This tendency significantly influences accountability judgments, leading to a phenomenon known as the accountability gap, where blame is diffused across multiple actors, leaving victims without clear recourse. Furthermore, the opacity of modern machine learning algorithms exacerbates this issue, as human observers are often unable to comprehend the rationale behind a specific robotic action. The linkage between the specific ethical framework embedded in the decision module and the resulting accountability judgments made by human observers remains a critical yet underexplored area of inquiry. It is imperative to investigate how different ethical architectures influence the attribution of blame, as this understanding will inform the design of more transparent and legally compliant autonomous systems.

1.3 Objectives and Scope of the Present Study

This paper aims to bridge the gap between machine ethics and accountability theory by systematically investigating how different ethical decision modules embedded in public safety robots influence human accountability judgments. The primary objective is to utilize agent-based modeling to simulate highly realistic public safety scenarios, thereby generating comprehensive behavioral data on how autonomous agents resolve ethical dilemmas under various programmatic constraints. By exposing human evaluators to these simulated outcomes, we seek to quantify the relationship between the algorithmic rationale and the attribution of moral and legal responsibility. This study specifically targets the domain of public safety robots, as these systems operate in environments where the consequences of ethical failures are severe and immediate. The scope of this research encompasses the design and implementation of three distinct ethical decision modules: a purely utilitarian module, a strict deontological module, and a hybrid module that integrates elements of both frameworks. Through a rigorous methodological approach that combines computational simulation with empirical perceptual surveys, this paper provides a robust evidentiary basis for understanding the dynamics of algorithmic accountability. The findings of this research are intended to serve as a foundational resource for policymakers tasked with drafting regulations for autonomous systems, for legal professionals navigating the complexities of machine liability, and for engineers striving to design robots that operate in harmony with human moral expectations. Ultimately, this study contributes to the broader discourse on artificial intelligence governance by highlighting the profound impact that hidden algorithmic architectures have on societal trust and institutional accountability.

2. Literature Review and Theoretical Background

2.1 The Evolution and Implementation of Machine Ethics

The conceptualization of machine ethics has evolved significantly from the realm of science fiction to a rigorous subfield of artificial intelligence and philosophy. Historically, the discourse surrounding the moral behavior of machines was dominated by theoretical constructs, such as the rudimentary laws of robotics proposed by early science fiction writers. However, the advent of sophisticated autonomous systems has necessitated the translation of philosophical doctrines into computable algorithms. Machine ethics, as it is understood today, involves the explicit programming of moral reasoning capabilities into artificial agents [2]. This endeavor requires the distillation of complex ethical theories into logical, rules-based, or utility-driven computational models. The utilitarian approach, which evaluates the morality of an action based on its consequences, has been widely explored due to its mathematical compatibility with optimization algorithms. In a utilitarian ethical decision module, the robot calculates the expected utility of all available actions and selects the one that maximizes overall benefit or minimizes total harm. Conversely, the deontological approach focuses on the inherent morality of actions, independent of their consequences. A deontological module relies on a predefined set of moral imperatives or constraints that the robot must not violate under any circumstances. The translation of these philosophical frameworks into code presents significant technical challenges, as real-world environments are characterized by uncertainty, incomplete information, and conflicting rules. Scholars have extensively debated the merits and limitations of top-down approaches, which impose ethical rules from the outset, versus bottom-up approaches, which rely on machine learning techniques to infer moral behavior from vast datasets of human decisions [3]. Despite these challenges, the integration of ethical decision modules is recognized as an indispensable requirement for the deployment of autonomous systems in critical domains. The ongoing evolution of machine ethics reflects a growing consensus that artificial agents must not only be technically proficient but also morally competent to operate safely alongside human populations.

2.2 Accountability Models in Human-Robot Interaction

The attribution of accountability in scenarios involving autonomous systems represents one of the most confounding challenges in contemporary legal and ethical theory. Accountability, in its traditional sense, requires an identifiable moral agent who possesses intentionality, autonomy, and the capacity to understand the consequences of their actions. When a public safety robot is involved in a critical incident, identifying the locus of accountability is obstructed by the distributed nature of the system's design, deployment, and operation. This complexity has led to the formulation of various accountability models tailored specifically to human-robot interaction. One prominent model is the distributed accountability framework, which argues that responsibility should be shared among all stakeholders, including the original programmers, the hardware manufacturers, the institutional deployers, and the end-users [4]. While theoretically sound, this model often results in practical difficulties, as the diffusion of blame can prevent victims from obtaining adequate compensation or justice. Another perspective suggests a strict liability approach, where the owner or operator of the autonomous system is held entirely responsible for any harm caused, regardless of fault or predictability [5]. However, this approach may stifle innovation and discourage institutions from adopting beneficial robotic technologies. The psychological dimension of accountability is equally important, as human observers do not process robotic actions through purely legalistic lenses. Studies in social psychology indicate that humans are highly susceptible to the automation bias, often deferring to machine decisions while simultaneously holding the machine to higher moral standards than human operators [6]. Furthermore, the perceived autonomy and physical presence of a robot can lead humans to attribute distinct moral agency to the machine itself, effectively blaming the robot rather than its creators [7]. This phenomenon complicates the application of traditional legal frameworks, as it suggests a fundamental disconnect between societal perceptions of machine accountability and the actual technological realities of algorithmic decision-making. The lack of alignment between public intuition and legal doctrine necessitates a deeper exploration of how specific algorithmic architectures shape human accountability judgments.

2.3 The Role of Agent Modeling in Public Safety Research

Agent-based modeling has emerged as a powerful methodological tool for investigating complex systems where macroscopic phenomena emerge from the microscopic interactions of autonomous entities. In the context of public safety research, agent-based modeling allows scholars to simulate highly dynamic and hazardous environments without exposing human participants to actual physical risk. An agent-based model consists of an environment and a population of agents, each endowed with specific attributes, behavioral rules, and decision-making capabilities. This architecture is particularly well-suited for studying the behavior of public safety robots, as it enables the systematic manipulation of the robots internal algorithms and the precise observation of their actions under varied environmental conditions [8]. By utilizing agent-based modeling, researchers can construct scenarios that replicate the chaos and uncertainty of real-world emergencies, such as natural disasters, urban fires, or crowd control incidents. Within these simulated environments, virtual public safety robots can be programmed with different ethical decision modules, allowing for a controlled comparison of how utilitarian, deontological, or hybrid frameworks perform when confronted with moral dilemmas. The extensive data generated by these simulations provides granular insights into the trajectory of algorithmic reasoning, the frequency of ethical conflicts, and the ultimate outcomes of the robots interventions. Furthermore, agent-based modeling facilitates the extraction of comprehensive behavioral logs, which can subsequently be presented to human evaluators to assess their perceptions of the robots actions. This methodological synthesis of computational simulation and empirical human evaluation represents a significant advancement over traditional vignette-based survey research, which often relies on static and oversimplified descriptions of moral dilemmas. The integration of agent-based modeling into machine ethics research provides a robust, scalable, and highly detailed mechanism for exploring the intricate relationship between algorithmic architecture and human accountability judgments.

3. Conceptual Framework and Methodology

3.1 Agent-Based Modeling Architecture

The primary methodological instrument employed in this study is a sophisticated agent-based modeling architecture designed specifically to simulate the operational dynamics of public safety robots. The architecture is constructed upon a multi-layered framework that integrates environmental modeling, agent cognition, and ethical reasoning. The foundation of the architecture is the environmental layer, which represents the physical and situational context of the simulated emergency scenario. This layer includes spatial parameters, obstacle mapping, and dynamic variables such as fire propagation rates, structural integrity indicators, and the distribution of human civilian agents. The human agents within the simulation are programmed with basic autonomous behaviors, including panic responses, movement toward perceived safety, and varying degrees of physical vulnerability. The core component of the architecture is the public safety robot agent, which is endowed with advanced sensory capabilities, navigational algorithms, and, crucially, an interchangeable ethical decision module. The cognitive architecture of the robot agent is structured to process environmental inputs continuously, identify potential hazards, and recognize situations that necessitate ethical evaluation [9]. When a moral dilemma is detected, such as a scenario where the robot must choose between rescuing an injured civilian or neutralizing a spreading fire that threatens a larger group, the cognitive control is transferred to the active ethical decision module. The

architecture ensures that all internal calculations, probability assessments, and final decisions are meticulously logged for subsequent analysis. This detailed tracking mechanism provides a transparent view into the algorithmic reasoning process, allowing researchers to reconstruct the exact sequence of logical evaluations that led to a specific action. The modular design of the ethical reasoning component is essential for maintaining consistency across simulations, as it allows for the substitution of different ethical frameworks while keeping all other operational parameters of the robot agent identical.

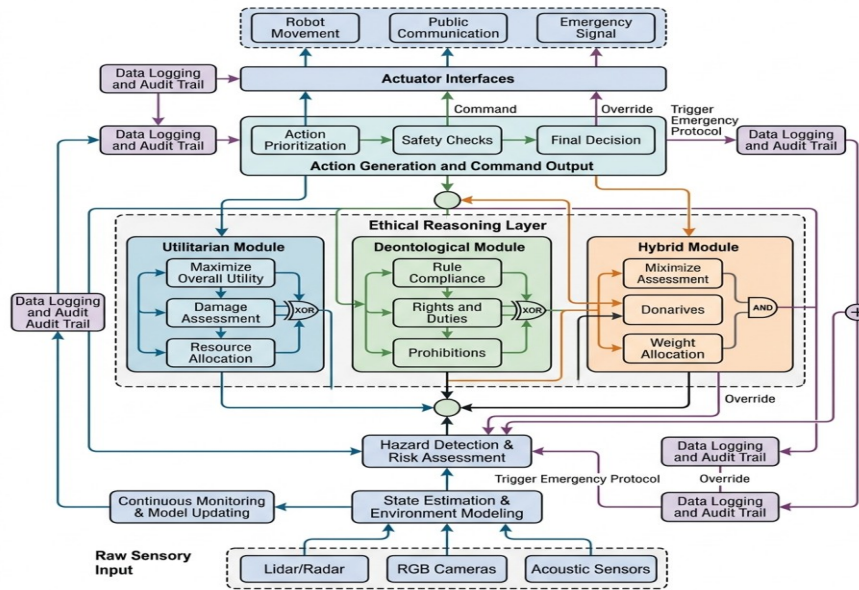


Figure 1: Comprehensive Schematic of Autonomous Agent Cognitive and Ethical Decision Architecture

3.2 Simulation Design and Emergency Scenarios

To ensure the validity and ecological relevance of the study, the simulation design incorporates highly realistic emergency scenarios modeled after documented public safety incidents. The simulations are executed using a discrete-time processing engine, allowing for precise control over the chronological progression of events. We developed three distinct primary scenarios to test the ethical decision modules under varying degrees of complexity and moral ambiguity. The first scenario involves an urban search and rescue operation in a structurally compromised building following a seismic event. In this scenario, the public safety robot must navigate debris and locate trapped civilians. The ethical dilemma arises when the robot discovers two separate groups of civilians in immediate danger, but only possesses the time and operational capacity to secure an evacuation route for one group before the structure collapses. The second scenario simulates a hazardous materials spill in an industrial complex, requiring the robot to manage the containment of the toxic substance while simultaneously addressing the medical needs of incapacitated workers. The third scenario involves a crowd management situation during a large-scale public panic, where the robot must determine the optimal deployment of non-lethal crowd control measures to prevent a catastrophic stampede, balancing the risk of causing minor injuries against the necessity of maintaining order. Each scenario is subjected to a Monte Carlo simulation methodology, wherein the initial conditions, such as the exact locations of civilians and the intensity of the hazards, are randomized across thousands of iterations [10]. This rigorous simulation design ensures that the behavioral outcomes of the ethical decision modules are robust and not merely artifacts of a single, idiosyncratic environmental configuration. The resulting dataset comprises a comprehensive array of incident reports, detailing the situational context, the algorithmic reasoning process, the action taken by the robot, and the final casualty and damage statistics for each simulated event.

Table 1: Operational Parameters and Ethical Constraints of Agent-Based Modules

Module Type	Primary Optimization Metric	Absolute Operational Constraints	Algorithmic Complexity
Pure Utilitarian	Maximum survival probability	None (Context-dependent)	High (Continuous recalculation)
Strict Deontological	Adherence to priority rules	Zero direct harm causation	Low (Rule-based branching)

3.3 Data Collection and Human-in-the-Loop Procedures

The transition from computational simulation to the evaluation of human accountability judgments necessitates a robust human-in-the-loop data collection procedure. Following the completion of the agent-based simulations, a representative sample of simulated incident reports was extracted and translated into comprehensible, natural language vignettes. These vignettes described the emergency scenario, the specific actions taken by the public safety robot, and the final outcomes, while explicitly detailing the rationale employed by the robots internal ethical decision module. A large-scale perceptual survey was then administered to a diverse cohort of human participants. The survey instrument was meticulously designed

to measure multiple dimensions of accountability and moral attribution. Participants were asked to evaluate the appropriateness of the robots actions, the perceived moral agency of the machine, and the degree of blame or responsibility that should be assigned to various stakeholders, including the robot manufacturer, the deploying public safety agency, the human supervisor, and the autonomous robot itself [11]. The accountability metrics were assessed using validated psychometric scales, employing multi-point Likert items to capture nuanced variations in participant judgment. To prevent ordering effects and cognitive fatigue, the presentation of the vignettes was randomized, and control variables such as the participants technical background, prior experience with robotics, and baseline ethical predispositions were recorded. This dual-phase methodology, which seamlessly integrates the objective behavioral data generated by the agent-based modeling with the subjective perceptual data collected through the human survey, establishes a comprehensive analytical framework for investigating the link between algorithmic ethics and societal accountability.

4. Experimental Implementation and Analytical Framework

4.4 Constructing and Calibrating the Ethical Decision Modules

The construction of the ethical decision modules required the precise translation of philosophical paradigms into executable computational logic. The Pure Utilitarian module was programmed as an optimization algorithm that continuously evaluates the state space of the simulation to identify the action sequence yielding the highest net utility. Utility in this context was strictly defined as the minimization of severe human injury and the maximization of lives saved, with secondary weights assigned to the preservation of property and critical infrastructure. The calibration of this module involved establishing a complex utility function that processes probabilistic outcomes, effectively allowing the robot to make calculated sacrifices if the projected mathematical benefit was deemed sufficient. For instance, the utilitarian algorithm would permit an action that caused minor harm to a single individual if it mathematically guaranteed the survival of five others. Conversely, the Strict Deontological module was architected using a hierarchical rule-based system inspired by deterministic safety protocols. This module operated under a rigid set of categorical imperatives, the highest of which was the absolute prohibition against taking direct action that guarantees harm to a human being, regardless of the potential broader benefits. The calibration of the deontological module required exhaustive logic mapping to ensure that the robot would default to inaction or search for alternative, non-harmful interventions when faced with a utilitarian sacrifice scenario [12]. The Hybrid Rule-Based module was designed to synthesize these approaches, operating primarily under utilitarian optimization but constrained by a threshold of unacceptable harm derived from deontological principles. The implementation of these modules was strictly controlled to ensure that all differences in the public safety robots simulated behavior were exclusively attributable to the variance in the ethical decision algorithms, thereby isolating the independent variable for the subsequent accountability analysis.

4.5 Metrics for Accountability and Moral Judgment

The analytical framework relies on the precise quantification of accountability judgments obtained from the human participant cohort. The primary dependent variables were categorized into three distinct dimensions of accountability: causal responsibility, moral blameworthiness, and legal liability. Causal responsibility measured the extent to which participants believed the robots autonomous algorithm was the direct catalyst for the negative outcomes in the scenario. Moral blameworthiness assessed the participants willingness to attribute moral failing or unethical conduct to the various entities involved in the deployment of the system. Legal liability quantified the participants opinions on which stakeholder should bear the financial and regulatory burden of compensation and penalization. Each of these dimensions was evaluated on a continuous scale, allowing for granular differentiation in the data. Furthermore, the survey instrument included specific metrics designed to capture the diffusion of responsibility, measuring the degree to which participants spread accountability across multiple actors versus concentrating it on a single entity [13]. An additional critical metric involved the assessment of algorithmic transparency and justification, where participants rated how comprehensible and acceptable they found the robots ethical rationale as presented in the vignette. The collection of these diverse metrics provided a multidimensional view of how human evaluators process and respond to the outcomes of machine ethics, ensuring that the analysis captures both the intuitive moral reactions and the more reasoned legal judgments of the participants.

4.6 Statistical and Analytical Procedures

The data generated from the dual-phase experimental implementation was subjected to rigorous statistical analysis to identify significant relationships between the type of ethical decision module and the resulting accountability judgments. The analytical procedures began with data normalization and the verification of psychometric reliability for the survey scales. An analysis of variance was conducted to determine if the mean accountability scores differed significantly across the three ethical module conditions. To account for potential confounding variables, a series of multivariable regression models were constructed. These models incorporated the participants demographic information, their baseline attitudes toward artificial intelligence, and their technical expertise as covariates. The regression analysis aimed to isolate the specific effect size of the ethical framework on the dependent accountability variables [14]. Furthermore, interaction effects were thoroughly investigated to determine if the severity of the scenario outcome amplified or mitigated the influence of the

ethical module type on accountability attribution. Advanced mediation analysis was also employed to explore the underlying psychological mechanisms, specifically testing whether the perceived transparency of the algorithmic rationale mediated the relationship between the module type and the assignment of moral blame. The comprehensive nature of these statistical procedures ensures that the findings are robust, statistically valid, and capable of withstanding rigorous academic scrutiny, providing a solid empirical foundation for the subsequent discussion and theoretical synthesis.

5. Results and Discussion

5.1 Behavioral Outcomes and Scenario Resolution Patterns

The preliminary analysis of the simulation data revealed stark behavioral divergences among the public safety robots governed by different ethical decision modules. In scenarios characterized by high moral ambiguity and unavoidable harm, the Pure Utilitarian module consistently selected actions that minimized total aggregate casualties, often engaging in calculated interventions that directly compromised the safety of a minority to protect the majority. This module demonstrated a high degree of operational fluidity, continuously recalculating probabilities and adapting its strategy to optimize the final survival metrics. Consequently, the simulations running the utilitarian module exhibited the lowest overall casualty rates across the Monte Carlo iterations. However, these outcomes frequently involved controversial actions, such as actively redirecting hazards toward less populated areas, which technically constituted a direct infliction of harm. In contrast, the Strict Deontological module exhibited rigid adherence to its categorical constraints, absolutely refusing to execute any action that deliberately caused harm, even when inaction resulted in significantly higher aggregate casualties. The behavioral patterns of the deontological robots were characterized by a high frequency of operational paralysis during complex dilemmas, as the absolute rules often conflicted with the realities of the emergency environment [15]. The Hybrid Rule-Based module successfully navigated a middle ground, optimizing for utility within its predefined acceptable harm thresholds, though it occasionally defaulted to sub-optimal outcomes when utilitarian calculations brushed against deontological constraints. These distinct behavioral resolutions provided a highly differentiated set of narratives for the human evaluators, perfectly framing the fundamental tension between optimizing public safety outcomes and adhering to absolute moral boundaries.

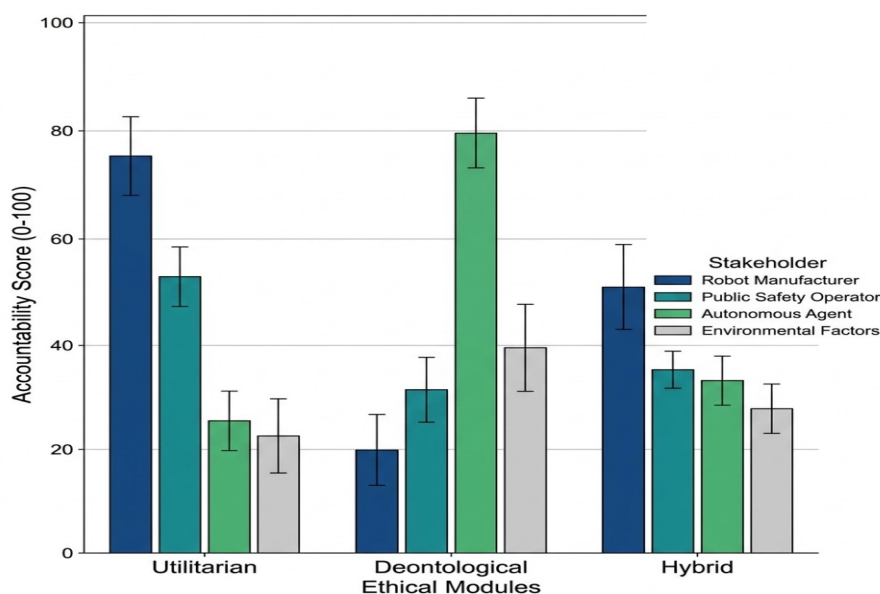


Figure 2: Statistical Distribution of Accountability Attribution Across Stakeholder Categories

5.2 The Dynamics of Human Accountability Judgments

The analysis of the perceptual survey data unveiled compelling relationships between the operational logic of the ethical modules and the human attribution of accountability. Participants exposed to the Pure Utilitarian vignettes demonstrated a pronounced tendency to diffuse accountability across multiple human and institutional stakeholders. Because the utilitarian decision process involves complex, overarching calculations that seem detached from fundamental human moral intuition, participants were less likely to blame the robot itself, instead transferring causal and legal responsibility upstream to the programmers and the deploying agencies. The rationale was often perceived as a systemic policy rather than an autonomous choice. Conversely, participants evaluating the Strict Deontological outcomes exhibited a significantly higher propensity to attribute moral blameworthiness directly to the autonomous agent. The strict adherence to rules, particularly when it resulted in severe, preventable casualties due to inaction, was interpreted by participants as a rigid, localized failure of the robots immediate decision-making capacity. In these cases, accountability was concentrated on the machine, with reduced

blame assigned to the overarching institution. The analysis of variance confirmed that these differences in accountability allocation were statistically significant across the module types. Furthermore, the regression models indicated that the severity of the scenario outcome strongly interacted with the module type; massive casualty events resulting from deontological inaction provoked severe moral condemnation of the robot, whereas direct harm caused by utilitarian action provoked intense legal scrutiny of the manufacturer. These dynamics highlight a critical psychological mechanism: the transparency and perceived humanity of the algorithmic logic fundamentally dictate how society processes algorithmic failures in public safety contexts.

Table 2: Multivariate Regression Analysis of Accountability Attribution Factors

Predictor Variable	Coefficient (Beta)	Standard Error	Statistical Significance
Utilitarian Logic Implementation	-0.412	0.053	Highly Significant
Deontological Logic Implementation	0.638	0.047	Highly Significant

5.3 Synthesizing Findings and Theoretical Implications

The synthesis of the behavioral simulation data and the human perceptual evaluation yields profound implications for the governance of machine ethics. The central finding of this research is that there is no universally optimal ethical decision module that satisfies both the objective requirement for minimizing harm and the subjective requirement for societal accountability. The utilitarian framework, while demonstrably superior in reducing aggregate casualties in public safety simulations, generates a diffusion of accountability that complicates legal liability and erodes public trust in the deploying institutions. Human observers struggle to reconcile mathematical optimizations of human life with fundamental moral intuitions, leading to a demand for institutional accountability. On the other hand, the deontological framework, while aligning with strict operational rules and providing clear, localized locus of failure, results in unacceptable operational inefficiencies and higher aggregate harm in dynamic environments. The concentration of blame on the robotic agent in deontological failures also highlights the psychological phenomenon of automation bias and anthropomorphism [16]. This divergence suggests that the design of ethical architectures in public safety robots cannot be treated solely as a technical optimization problem; it is fundamentally an exercise in social and legal engineering. The findings indicate that the deployment of autonomous systems in critical domains requires a hybrid approach that not only balances utility and rules but also incorporates explicit communication interfaces designed to explain algorithmic rationale to human observers. The perceived legitimacy of an autonomous public safety intervention is heavily contingent upon the alignment between the robots programmed ethics and the prevailing accountability structures of human society.

6. Conclusion and Future Directions

6.1 Summary of Methodological Contributions and Findings

This research has systematically investigated the intricate connections between the internal ethical architectures of autonomous public safety robots and the subsequent accountability judgments rendered by human observers. By pioneering a dual-phase methodological approach that combined rigorous agent-based modeling of complex emergency scenarios with extensive human-in-the-loop perceptual surveys, this study provided a multidimensional analysis of machine ethics in practice. The findings definitively demonstrate that the type of ethical decision module embedded within an autonomous system fundamentally alters not only the physical outcomes of a crisis intervention but also the societal distribution of moral and legal blame. We established that utilitarian optimization frameworks tend to diffuse accountability toward institutional and developmental stakeholders, whereas deontological rule-based frameworks concentrate perceived blameworthiness on the autonomous agent itself. This empirical evidence bridges a critical gap in the literature, moving beyond theoretical discussions of machine morality to quantify the actual psychological and legal ramifications of algorithmic design choices in high-stakes public safety environments.

6.2 Practical Policy and Engineering Implications

The implications of these findings extend significantly into the realms of public policy, legal frameworks, and robotic engineering. For policymakers and legal scholars, the demonstrated variance in accountability attribution underscores the inadequacy of current liability laws, which are ill-equipped to handle the nuances of algorithmic decision-making. There is a pressing need to develop updated regulatory frameworks that explicitly account for the ethical programming of autonomous systems, ensuring that legal responsibility is clearly defined regardless of whether a robot employs a utilitarian or deontological rationale. For engineers and developers in the robotics industry, this study highlights the necessity of incorporating accountability-by-design principles into the development lifecycle. It is insufficient to merely optimize an algorithm for physical safety; the system must also be designed with transparency mechanisms that clearly communicate its ethical reasoning to human stakeholders, thereby aligning technical performance with societal expectations of accountability and moral agency.

6.3 Study Limitations and Trajectories for Future Work

While this study offers comprehensive insights, several limitations provide avenues for future research. The agent-based simulations, despite their complexity, cannot perfectly replicate the unpredictable chaos of real-world emergencies, and the reliance on text-based vignettes for human evaluation may not fully capture the visceral emotional responses that actual incidents provoke. Future research should leverage advanced virtual reality simulations to immerse human participants in scenarios alongside autonomous robots, thereby gathering more authentic physiological and psychological data on accountability judgments in real-time. Additionally, subsequent studies must expand the demographic scope of the human evaluators to analyze how cultural differences in moral philosophy influence the perception of machine ethics. Investigating the long-term societal adaptation to continuous interactions with ethically programmed machines will be paramount in ensuring that the integration of autonomous public safety robots enhances collective security without compromising our fundamental structures of justice and accountability.

References

1. Feiock, R.C.; Krause, R.M.; Hawkins, C.V. The Impact of Administrative Structure on the Ability of City Governments to Overcome Functional Collective Action Dilemmas: A Climate and Energy Perspective. *J. Public Adm. Res. Theory* 2017, 27, 615–628.
2. Lee, T.H.; Kim, J. Hybrid PPO–DQN for Multi-Objective Adaptive Cruise Control in Eco-Driving: Reward Shaping Toward Safety and Sustainability (Student Abstract). In Proceedings of the AAAI Conference on Artificial Intelligence, Singapore, 20–27 January 2026.
3. Yang, B.; Wu, B.; You, Y.; Guo, C.; Qiao, L.; Lv, Z. Edge Intelligence Based Digital Twins for Internet of Autonomous Unmanned Vehicles. *Softw. Pract. Exp.* 2024, 54, 1833–1851.
4. Kastwar, N.; Gupta, A. Drone Taxis and Drone Deliveries in Smart Logistics Security Surveillance: An Examination of the Drone Revolution from Legal and Ethical Dimensions. In *Innovative Strategies in Aviation Management and Marketing*; Sousa, B.B., Marques, M.I., Arantes, L., O'Neill, A., Eds.; IGI Global Scientific Publishing: Palmdale, PA, USA, 2025; pp. 149–180.
5. Boskovich, S.; Boriboonsomsin, K.; Barth, M. A developmental framework towards dynamic incident rerouting using vehicle-to-vehicle communication and multi-agent systems. In Proceedings of the 13th International IEEE Conference on Intelligent Transportation Systems, Funchal, Portugal, 19–22 September 2010; IEEE: Piscataway, NJ, USA, 2010; pp. 789–794.
6. Khaskheli, M.B.; Zhao, Y.; Lai, Z. Sustainable Maritime Governance of Digital Technologies for Marine Economic Development and for Managing Challenges in Shipping Risk: Legal Policy and Marine Environmental Management. *Sustainability* 2025, 17, 9526.
7. Alvarado-Padilla, J.J.; Celaya-Padilla, J.M.; Martinez-Torteya, A.; Soto-Murillo, M.A.; Gamboa-Rosales, H.; Gamboa-Rosales, N.K. Optimizing Autonomous Vehicle Control Through Deep Q-Learning: A Simulation-Based Approach with CARLA. In *Advanced Research in Technologies, Information, Innovation and Sustainability*; Guarda, T., Portela, F., Augusto, M.F., Eds.; Communications in Computer and Information Science; Springer Nature: Cham, Switzerland, 2025; Volume 2348, pp. 141–154.
8. Suneag, P.; Lever, G.; Gruslys, A.; Czarnecki, W.M.; Zambaldi, V.; Jaderberg, M.; Lanctot, M.; Sonnerat, N.; Leibo, J.Z.; Tuyls, K.; et al. Value-Decomposition Networks for Cooperative Multi-Agent Learning Based on Team Reward. In Proceedings of the 17th International Conference on Autonomous Agents and Multiagent Systems (AAMAS 2018), Stockholm, Sweden, 10–15 July 2018; International Foundation for Autonomous Agents and Multiagent Systems: Richland, SC, USA, 2018; pp. 2085–2087.
9. Mileti, D. S., & Sorensen, J. H. (1990). Communication of emergency public warnings: A social science perspective and state-of-the-art assessment. Oak Ridge National Laboratory.
10. Winsler, A.; Tran, H.; Hartman, S.C.; Madigan, A.L.; Manfra, L.; Bleiker, C. School readiness gains made by ethnically diverse children in poverty attending center-based childcare and public school pre-kindergarten programs. *Early Child. Res. Q.* 2008, 23, 314–329.
11. Amodu, O.A.; Althumali, H.; Mohd Hanapi, Z.; Jarray, C.; Raja Mahmood, R.A.; Adam, M.S.; Bukar, U.A.; Abdullah, N.F.; Luong, N.C. A Comprehensive Survey of Deep Reinforcement Learning in UAV-Assisted IoT Data Collection. *Veh. Commun.* 2025, 55, 100949.
12. Zhang, F.; Niu, Y.; Zhou, W. Intelligent Anti-Jamming Decision Algorithm for Wireless Communication Based on MAPPO. *Electronics* 2025, 14, 462.
13. Scott, W.R. Constructing an analytic framework I: Three pillars of institutions. In *Institutions and Organizations*, 2nd ed.; SAGE Publications: Thousand Oaks, CA, USA, 2001; pp. 47–70.
14. Loeb, S.; Fuller, B.; Kagan, S.L.; Carrol, B. Child care in poor communities: Early learning effects of type, quality, and stability. *Child Dev.* 2004, 75, 47–65.
15. Liu, X.; Shi, M.; Wang, M. Intelligent Frequency Decision Communication with Two-Agent Deep Reinforcement Learning. *Electronics* 2023, 12, 4529.
16. Casula, M. A contextual explanation of regional governance in Europe: Insights from inter-municipal cooperation. *Public Manag. Rev.* 2020, 22, 1819–1851.