



ARTICLE

An Edge-Assisted Internet-of-Vehicles Computing Framework for Fair Tail-Risk Allocation in Cooperative Autonomous Driving

Shih-Lin Lin^{*}

Graduate Institute of Vehicle Engineering, National Changhua University of Education, Changhua City, Taiwan

*Corresponding Author: Shih-Lin Lin. Email: lin040@cc.ncue.edu.tw

Received: 25 April 2026; Accepted: 18 June 2026; Published: 23 July 2026

ABSTRACT: Connected automated driving increasingly relies on cooperative perception from onboard sensors, roadside units (RSUs), smart traffic lights, and vehicle-to-everything (V2X) links, but communication uncertainty can concentrate residual risk on vulnerable road users (VRUs). This study proposes an Ethical-Improved risk-allocation objective for edge-assisted Internet of Vehicles (IoV) cooperative autonomous driving. The objective internalizes responsibility as a bounded risk weight, normalizes equality and maximin terms, and adds explicit VRU tail-risk and VRU/Ego ratio penalties. The evaluation is organized into two strictly separated tracks. In the planner-objective track, Ethical-Improved reduces physical collision rate, aggregate harm, inequality, and VRU tail risk relative to Standard, Selfish, and Ethical-Orig baselines. In the communication-aware IoV track, edge-assisted cooperative perception achieves the best within-track collision-proxy and normalized-harm performance under sampled packet delivery ratio (PDR), age of information (AoI), RSU coverage, perception confidence, and edge delay conditions. Physical collision rate and communication-conditioned collision proxy are reported separately and are not numerically compared across tracks. The results support auditable tail-risk regularization as a controlled-simulator design principle for fairer cooperative autonomous driving, while external validation remains necessary before deployment-oriented claims can be made.

KEYWORDS: Internet of vehicles; edge-assisted cooperative perception; vehicle-road cooperation; autonomous driving; fairness-aware trajectory planning; tail-risk regularization; vulnerable road user protection

1 Introduction

As autonomous-driving systems increasingly assume driving decisions, ethical trajectory planning must proactively manage risk across heterogeneous road users, including ego occupants, third parties, and vulnerable road users (VRUs). Traditional binary dilemmas are insufficient; real-world navigation requires continuous, severity-aware risk allocation.

Ethical trajectory planning is increasingly posed as an optimization problem that balances overall harm and distributive justice (e.g., the multi-principle risk-cost framework of Geisslinger et al. [1]). Recent approaches draw on tail-aware risk metrics like Conditional Value-at-Risk to curb extreme outcomes [2] and employ paired statistical tests to compare collision outcomes [3]. Fairness is typically quantified via inequality indices such as the Gini coefficient [4], guided by Rawlsian maximin principles of justice [5]. These technical formulations enable algorithms to internalize ethical concepts, but they must be evaluated in realistic driving contexts. Traditional motion planning and decision-making techniques for automated

driving have been extensively surveyed, covering deterministic and sampling-based planners, optimization-based control, and learning-based methods [6–8]. Complementary surveys focus on prediction and risk assessment, highlighting challenges like occlusions, interactive behavior, and probabilistic forecasting [9]. To address this, scenario-based safety assessment frameworks have been developed, from structured scenario catalogs and coverage metrics to simulation-driven validation toolchains [10]. Recent work also proposes runtime safety monitors using online verification to prevent planned maneuvers from entering provably unsafe states [11]. These advances motivate our focus on measurable, scenario-level risk outcomes that support both planning-time optimization and rigorous post hoc evaluation. While these foundations yield safer trajectories in aggregate, they often use average-case performance metrics that can obscure how residual risk is distributed across different road users. This motivates the first gap: ethical planners need explicit VRU tail-risk control, not only aggregate harm reduction.

Early debates on autonomous vehicle (AV) ethics were dominated by hypothetical “trolley problem” scenarios. Goodall [12] pioneered an explicit algorithmic approach to handle AV crash dilemmas, while others questioned whether encoding “accident algorithms” in software is appropriate [13]. Meanwhile, empirical studies of public moral preferences in unavoidable collision scenarios revealed systematic tensions between what individuals consider acceptable for society vs. for themselves [14]. A large-scale cross-cultural experiment confirmed that such moral judgments vary with societal norms and context [15]. Some research has examined decisions under uncertainty; for example, people often prefer a default (status quo) action rather than actively causing potential harm when outcomes are probabilistic [16]. Conceptual analyses have cautioned that an exclusive focus on rare crash dilemmas can be misleading, arguing that everyday risk management—who bears risk, and how much—is ethically central to AV behavior [17]. Recent literature reviews synthesize these debates and highlight gaps between high-level ethical principles, regulatory guidance, and implementable control objectives [18]. Liu [19] further warns that without careful design, AV decision policies could impose inequitable risk burdens on certain groups, leading to unjust outcomes. User studies further suggest that public acceptance of AVs may hinge on perceived behavioral fairness [20], with different responses to passenger-prioritizing (“selfish”) and utilitarian AV policies [21]. Notably, a recent empirical study found that participants were even willing to accept higher risk to themselves to reduce risk to others [22]. Beyond preference surveys, crash responsibility connects risk allocation to duty-of-care reasoning [23], ethics settings may be embedded before a conflict occurs [24], structured AV ethics procedures are needed [25], moral-judgment models require responsibility, consequence, and intention [26], and residual crash-risk allocation is an independent fairness problem [27]. Policy analyses warn that safety requirements should not be displaced by abstract crash ethics [28]. Well-being perspectives add that automated mobility affects users and non-users beyond immediate collisions [29]. Methodological studies call for transparent ethical decision processes that can be audited after deployment [30]. Sector guidelines translate these principles into practical requirements for automotive artificial intelligence [31]. Control studies show that ethical priorities can be encoded through constraints and vehicle-control costs [32]. This motivates the second gap: responsibility should be encoded as a bounded internal risk weight, not as an external discount or subtractive reward that weakens the planner’s nonzero duty of care.

A third issue concerns scale and calibration. Safety-validation research shows that rare-event reliability requires evidence beyond limited simulation [33]. Risk-aware path planning highlights spatially heterogeneous hazards and context-dependent exposure [34]. Contingency planning prepares alternative maneuvers under multimodal prediction uncertainty [35]. Together with the Gini and maximin principles introduced above, these works show that fairness and worst-off terms can be difficult to interpret across traffic densities, speeds, occlusion patterns, and conflict geometries. This motivates the third gap: equality and maximin

components need normalization before their trade-off weights can be compared across heterogeneous urban scenarios.

Finally, connected and cooperative driving extends ethical planning from an onboard decision problem to an information-conditioned IoV problem. Connected-vehicle surveys show that cooperative perception benefits vulnerable-road-user interactions [36]. Empirical external human-machine interface studies show that intent displays affect pedestrian crossing decisions [37]. Vehicle-to-pedestrian communication reviews summarize how external interfaces support interaction design [38]. However, most ethical objectives do not explicitly condition risk allocation on packet delivery ratio, age of information, roadside-unit coverage, perception confidence, or edge inference delay. This motivates the fourth gap: ethical IoV planners should make communication quality auditable in safety and fairness outcomes. Collectively, these research efforts provide a broad foundation—spanning motion planning, safety validation, ethical theory, and human factors—upon which we build our approach. Our work extends this literature by explicitly targeting worst-case risk inequalities (the “tail risks”) in a multi-principle objective, enabling a quantitative examination of fairness–safety trade-offs that previous frameworks could only qualitatively discuss.

Research Gaps and Contributions

The reviewed literature leaves four gaps that motivate this study. First, existing ethical planners often combine utilitarian, egalitarian, maximin, and responsibility principles without explicitly controlling VRU tail-risk dominance. Second, responsibility is commonly treated as an external discount or subtractive term, which may reduce the planner’s duty of care too aggressively. Third, fairness terms are often sensitive to scenario scale, making cross-scenario interpretation difficult. Fourth, most ethical-planning formulations are not explicitly conditioned on Internet of Vehicles (IoV) communication quality, such as packet delivery ratio (PDR), age of information (AoI), roadside unit (RSU) coverage, perception confidence, and edge inference delay.

To address these gaps, this paper makes four contributions: (i) it reformulates responsibility as a bounded internal risk weight; (ii) it normalizes equality and maximin terms to improve interpretability across heterogeneous scenarios; (iii) it introduces explicit VRU tail-risk and VRU/Ego ratio penalties; and (iv) it extends the planner to a communication-aware IoV setting with edge-assisted and communication-impaired operating conditions.

Building on this literature, this study reformulates the baseline ethical planner as a communication-aware risk-allocation framework for IoV cooperative autonomous driving. Relative to the prior multi-principle formulation [1], the proposed objective introduces three structural changes: responsibility is moved from an external subtractive bonus to an internal bounded weight on aggregate risk; equality and maximin components are normalized so that their trade-off weights remain comparable across heterogeneous scenarios; and two explicit VRU-tail regularizers are introduced to control exceedance above a calibrated threshold and persistent VRU-over-Ego tail dominance. The resulting objective is intended to be numerically stable, auditable, and easier to interpret under matched scenario and communication-state comparisons.

The ego vehicle retains onboard sensing and local receding-horizon planning, while roadside units (RSUs), smart traffic lights, roadside cameras, and V2X message exchange support cooperative perception, occlusion mitigation, and awareness beyond the local sensor field of view [36]. External human-machine interfaces (eHMI) may further support communication between automated vehicles and vulnerable road users [37,38].

The system model contains five interacting layers: (i) the ego vehicle, which performs local sensing, state estimation, and low-latency motion planning; (ii) RSU and smart-infrastructure nodes, which provide occlusion lifting, VRU tracking, and right-of-way context; (iii) vehicle-to-vehicle (V2V), vehicle-to-infrastructure (V2I), and vehicle-to-pedestrian (V2P) links, which disseminate cooperative perception packets and intent messages; (iv) an edge server, which fuses local and infrastructure observations for short-term prediction and group-wise risk aggregation; and (v) a cloud layer, which is limited to long-term model updates and policy calibration rather than hard real-time control.

In the proposed framework, local-only data remain available at every planning cycle, while infrastructure-assisted observations enter through a communication state $z_t = \{a_t^{V2X}, \text{pdr}_t, \text{AoI}_t, c_t^{\text{RSU}}, q_t^{\text{perc}}, d_t^{\text{edge}}\}$. Here a_t^{V2X} denotes link availability, pdr_t the packet delivery ratio, AoI_t the age of information, c_t^{RSU} the current RSU coverage state, q_t^{perc} the perception confidence under occlusion, and d_t^{edge} the edge inference delay. This state is logged together with safety and fairness outputs so that the planner can be audited not only for which road-user group bears risk, but also for which IoV/V2X operating conditions shaped that allocation.

2 Ethical Risk Objective

The core of any optimization-based motion planner is its cost function. This section defines the risk space and contrasts the baseline multi-principle objective with the proposed formulation.

2.1 Mathematical Definitions and Risk Space

Let $u \in U$ denote a candidate maneuver for the autonomous vehicle (AV) in a given driving scenario. This maneuver u may represent a trajectory segment (a sequence of (x, y, v, t) states) or a discrete control sequence (steering and acceleration inputs) over a finite horizon.

Let $R_i(u)$ denote the predicted risk for road-user group i under maneuver u . Risk is defined here as a normalized harm proxy, typically computed as the integral of collision probability weighted by the expected severity of the collision (often parameterized by delta-V or kinetic energy transfer). We consider three distinct road-user groups:

$$i \in \{\text{Ego}, \text{Third}, \text{VRU}\} \quad (1)$$

where:

- Ego: Refers to the ego vehicle occupants.
- Third: Refers to third-party road users, such as occupants of other vehicles.
- VRU: Refers to vulnerable road users, including pedestrians, cyclists, and motorcyclists, who lack the protective enclosure of a vehicle.

The group-wise risk vector for a single scenario is thus expressed as:

$$\mathbf{R}(u) = [R_{\text{Ego}}(u), R_{\text{Third}}(u), R_{\text{VRU}}(u)]^T \quad (2)$$

This vector $\mathbf{R}(u)$ represents the distribution of group-wise risk induced by maneuver u and serves as the quantity optimized under a specified set of moral principles.

2.2 Baseline Multi-Principle Ethical Objective (Ethical-Orig)

For reference, the “Ethical-Orig” baseline [1] combines Utilitarianism, Egalitarianism, Maximin, and Responsibility into a single scalar risk cost:

$$J_{\text{Risk}}(u) = w_B J_B(u) + w_E J_E(u) + w_M J_M(u) - w_R J_R(u) \quad (3)$$

here, $w_B, w_E, w_M, w_R \geq 0$ are scalar trade-off weights that determine the relative importance of each principle.

1. Utilitarian/Bayesian Term (J_B)

The utilitarian term $J_B(u)$ seeks to minimize the total aggregate risk. It is typically expressed as the average risk across all relevant actors:

$$J_B(u) = \frac{1}{|S_R|} \sum_{i \in S_R} R_i(u) \quad (4)$$

where S_R denotes the set of relevant road users included in the risk evaluation for the scenario, and $|S_R|$ is its cardinality used for scale normalization. Minimizing J_B aligns with the principle of efficiency—reducing the total expected harm regardless of who suffers it.

2. Equality Term (J_E)

The equality term $J_E(u)$ penalizes unequal risk allocation across groups, reflecting egalitarian values. A common instantiation uses pairwise absolute differences, e.g.,

$$J_E(u) \propto \sum_{i,j} |R_i(u) - R_j(u)| \quad (5)$$

This term increases with the disparity among $\{R_i(u)\}$. While theoretically sound, unnormalized equality metrics can lead to “leveling down”, where the system prefers a scenario where everyone faces high risk (low variance) over one where some are safe and others are slightly at risk (high variance).

3. Maximin (Worst-Off) Term (J_M)

Based on Rawlsian principles of justice, specifically the Difference Principle, $J_M(u)$ focuses on the outcome of the least advantaged group. It implements a minimax strategy (minimizing the maximum risk):

$$J_M(u) = \max_i R_i(u) \quad (6)$$

This term discourages maneuvers that impose extreme risk on any single group, serving as a robust safety floor.

4. Responsibility Discount Term (J_R)

The term $J_R(u)$ captures responsibility-sensitive risk discounting. In the baseline formulation Eq. (3), it enters as a subtractive term $-w_R J_R(u)$. The intended effect is to reduce the protection weight for parties deemed responsible for the hazardous situation (e.g., a pedestrian jaywalking). If J_R represents the “risk attributable to responsible parties,” subtracting it lowers the total cost, making the maneuver more attractive.

5. Critique of the Baseline

In practice, the component terms $\{J_B, J_E, J_M, J_R\}$ can exhibit incompatible scales. J_B might be dominated by high-probability low-severity events, while J_M is driven by rare high-severity spikes. This scale mismatch complicates weight tuning, often leading to solutions where one term dominates the others

entirely. Moreover, the subtractive responsibility term introduces incentive distortions. By subtracting $w_R J_R(u)$, the optimizer can effectively “earn points” by finding trajectories where high risk coincides with high responsibility. This reduces the penalty for harm rather than the harm itself, potentially encouraging aggressive behaviors toward non-compliant road users that conflict with the legal doctrine of “last clear chance.”

2.3 Improved Objective: Ethical-Improved

To address scale incompatibility and perverse incentives, the Ethical-Improved objective internalizes responsibility weighting, normalizes fairness terms, and introduces explicit VRU tail-risk regularization. The improved objective is defined as:

$$J_{\text{Risk}}^{\text{Imp}}(u) = w_B \left(\frac{1}{|S_R|} \sum_{i \in S_R} \delta_i(u) R_i(u) \right) + w_E \hat{J}_E(u) + w_M \hat{J}_M(u) + \lambda_T \max\{0, R_{\text{VRU}}(u) - \tau\}^2 + \lambda_\rho \max\{0, R_{\text{VRU}}(u) - \rho R_{\text{Ego}}(u)\}^2 \quad (7)$$

The role and motivation of each component introduced in Eq. (7) are summarized below.

1. Bounded Responsibility Discount Factor (δ_i)

Instead of subtracting responsibility as an external bonus, responsibility is implemented as a bounded discount factor $\delta_i(u)$ applied directly to the risk $R_i(u)$ within the utilitarian sum:

$$\delta_i(u) = \max\{\epsilon, 1 - r_i(u)\}, \epsilon > 0 \quad (8)$$

where $r_i(u) \in [0, 1]$ denotes the responsibility score assigned to group i under maneuver u .

- $r_i(u) = 1$ indicates full responsibility (e.g., a pedestrian jumping into traffic).
- $r_i(u) = 0$ indicates no responsibility (e.g., a pedestrian on a green walk signal).
- ϵ is a strictly positive lower bound (e.g., $\epsilon = 0.05$).

This formulation discounts the risk assigned to responsible parties but prevents the weight from ever reaching zero. The lower bound ϵ is crucial; it ensures that even for a fully responsible agent, the AV retains a non-zero incentive to avoid collision. This aligns with ethical frameworks that value human life intrinsically, regardless of fault, and legal standards that require minimizing harm even to negligent actors. By embedding responsibility inside the utilitarian aggregation, the subtractive “bonus” distortion is removed: cost reduction requires reducing weighted risk rather than exploiting responsibility assignments.

2. Normalized Fairness Terms ($\hat{J}_E, \hat{J}_M$)

To ensure the trade-off weights w_E and w_M behave consistently across varying traffic densities and speeds, normalized fairness terms $\hat{J}_E(u)$ and $\hat{J}_M(u)$ are used, scaled to comparable magnitudes. A standard min-max normalization is:

$$\hat{J}(u) = \frac{J(u) - J_{\min}}{J_{\max} - J_{\min} + \eta} \quad (9)$$

where $J_{\min}$ and $J_{\max}$ are estimated from a representative set of training scenarios, and $\eta > 0$ is a small constant for numerical stability ($\sim 10^{-12}$). Alternatively, quantile normalization can be used to provide robustness against outliers. This normalization makes the weights w_E and w_M more interpretable as tuning parameters across heterogeneous scenarios.

3. Explicit VRU Tail Regularization (λ_T)

This term addresses the “tyranny of the average” by explicitly penalizing extreme VRU risk via a squared hinge loss:

$$\text{Penalty}_T = \lambda_T \max\{0, R_{\text{VRU}}(u) - \tau\}^2 \quad (10)$$

where the hinge operator is defined as $h(x) = \max\{0, x\}$. The squared hinge form yields a quadratic penalty once the threshold is exceeded and provides a continuous gradient at the activation point, which is advantageous for gradient-based solvers.

- Threshold τ : τ is a tolerance threshold, typically set to a high quantile of the VRU risk distribution (e.g., $q \in [0.90, 0.99]$).
- Mechanism: The penalty is zero when $R_{\text{VRU}}(u) \leq \tau$ and grows super-linearly when $R_{\text{VRU}}(u) > \tau$, thereby acting as a soft constraint against extreme VRU outcomes.

4. Protective Ratio Cap (λ_ρ)

The final term imposes a relational constraint between VRU and ego risk:

$$\text{Penalty}_\rho = \lambda_\rho \max\{0, R_{\text{VRU}}(u) - \rho R_{\text{Ego}}(u)\}^2 \quad (11)$$

where $\rho > 0$ is typically set near unity (e.g., $\rho \approx 1$).

This term discourages solutions in which VRU risk systematically exceeds ego risk, i.e., $R_{\text{VRU}}(u) > \rho R_{\text{Ego}}(u)$. It structurally prevents the planner from achieving ego safety by disproportionately offloading tail risk onto VRUs, encoding a protective stance toward vulnerable parties.

Overall, Eq. (7) preserves the interpretability of the baseline multi-principle structure while explicitly targeting tail outcomes. Responsibility is incorporated as an internal weighting (avoiding subtractive incentives), fairness terms are scale-aligned via normalization, and the hinge- and ratio-based penalties suppress extreme VRU risk and persistent VRU-over-ego tail dominance.

5. Parameter Interpretation and Calibration

The parameters in the Ethical-Improved objective are treated as policy-calibration parameters rather than universally optimal constants. The tail threshold τ defines the risk level above which VRU exposure is treated as an ethically salient tail event. For the planner-objective results reported in Section 4, τ is fixed at the empirical $q_{0.95}$ VRU-risk quantile. The sensitivity analysis in Section 4.8 evaluates $q_{0.90}$, $q_{0.95}$, $q_{0.975}$, and $q_{0.99}$. The lower-bound parameter ϵ prevents responsibility weighting from eliminating the planner’s duty of care, while ρ controls the acceptable VRU/Ego tail-risk ratio. The numerical stabilizer η is used only to avoid division by zero in normalized fairness terms and does not encode an ethical preference.

In the simulator, the responsibility score $r_i(u)$ is assigned from scenario-level right-of-way ambiguity, occlusion exposure, and crossing-intent uncertainty. It is not used to absolve the planner from avoiding harm; instead, it adjusts the internal weighting of predicted risk through $\delta_i(u) = \max\{\epsilon, 1 - r_i(u)\}$. Thus, even road users with higher responsibility retain a nonzero safety weight.

2.4 Communication-Aware Extension for Internet-of-Vehicles Conditions

To support IoV cooperative driving, the improved objective is extended from a purely trajectory-centric formulation to a communication-aware objective $J_{\text{Risk}}^{\text{Imp}}(u; z_t)$, where z_t records V2X availability, packet delivery ratio, age of information, RSU coverage, occlusion-aware perception confidence, and edge inference

delay. The communication state does not replace local sensing; instead, it adjusts the normalized group-wise risk proxy and adds a penalty for stale, missing, or delayed cooperative perception. In occluded scenes, the VRU tail threshold and VRU/Ego ratio cap can be tightened so that tail penalties activate earlier when infrastructure support is unreliable.

The communication-aware objective is specified explicitly as Eq. (12), where the communication penalty is defined in Eq. (13).

$$J_{\text{IoV}}(u, z_t) = J_{\text{Improved}}(u; \tilde{R}(z_t)) + \lambda_c P_{\text{comm}}(z_t) \quad (12)$$

where

$$P_{\text{comm}}(z_t) = w_p(1 - PDR) + w_a \max(0, AoI - AoI_{\text{max}}) + w_r(1 - RSU_{\text{cov}}) + w_q(1 - q_{\text{perc}}) + w_d \max(0, d_{\text{edge}} - d_{\text{max}}) \quad (13)$$

here, a lower packet delivery ratio (PDR), a larger age of information (AoI), unavailable or partial RSU coverage (RSU_{cov}), lower perception confidence q_{perc} , and a larger edge inference delay d_{edge} all increase $P_{\text{comm}}(z_t)$. These communication degradations also inflate the communication-conditioned risk proxy $\tilde{R}_i$ for occluded VRU states, so stale or missing cooperative perception makes the planner more conservative and activates the VRU tail-risk and VRU/Ego ratio penalties earlier, while local sensing remains available at every planning cycle.

3 Evaluation Protocol

To rigorously validate the proposed objective, we established a comprehensive simulation-based evaluation protocol. The goal was to compare the ethical performance of different planners under identical, high-stress conditions.

3.1 Methods Compared

We compare four distinct decision-making methods, differing only in the objective function minimized during trajectory planning:

1. Standard: a safety-oriented baseline that minimizes the unregularized risk surrogate over the planning horizon. It should be interpreted as a planner without explicit fairness or tail regularization, not as an oracle for minimizing realized harm totals.
2. Selfish: an ego-favoring baseline that assigns greater weight to ego risk than to third-party or VRU risk, representing a consumer-protection preference rather than a distributive-justice objective.
3. Ethical-Orig: the baseline multi-principle ethical method that minimizes Eq. (3). This serves as the principal ethical baseline used for comparison in this study.
4. Ethical-Improved: the proposed method minimizes Eq. (7), featuring bounded responsibility, normalized fairness, and tail regularization.

3.2 Scenario Design and Logged Variables

We used a paired-scenario evaluation methodology. All methods were evaluated on the same scenario set so that differences in outcomes can be attributed to decision logic and communication conditions rather than to stochastic variation in the environment.

- Sample Size: We evaluated $N = 2000$ paired scenarios.

Scenario proportions: The 2000 scenarios were sampled as 25% unprotected left turns, 25% occluded crossings, 25% right-of-way ambiguity, and 25% heterogeneous crowds (500 paired cases per class).

- Scenario Composition: The scenarios were curated to include diverse urban driving situations, including:
 - Unprotected Left Turns: The AV must navigate a gap in oncoming traffic while watching for pedestrians on the crosswalk.
 - Occluded Crossings: A pedestrian emerges from behind an obstruction (e.g., a parked truck), forcing a sudden reaction.
 - Intersections with Right-of-Way Ambiguity: Scenarios where responsibility is dynamic or unclear.
 - Heterogeneous Crowds: Mixes of cyclists, pedestrians, and other vehicles.

For each scenario n and each method, we logged the following telemetry:

- Group-Wise Risk Vector: $\mathbf{R}_n = [R_{\text{Ego}}^{(n)}, R_{\text{Third}}^{(n)}, R_{\text{VRU}}^{(n)}]^T$, representing the estimated risk levels for the ego vehicle, third-party road users, and vulnerable road users, respectively.
- Collision Indicator: $C \in \{0, 1\}$, where $C = 1$ indicates a physical collision occurred.
- Travel Time (TT): The time required to complete the scenario segment (e.g., time-to-goal), serving as a proxy for efficiency.
- Harm Proxies: Cumulative harm was computed as an ex post simulator aggregate based on impact-severity measurements such as impact-velocity squared. The same harm proxy was used across all planner baselines.

For the IoV extension, each paired scenario also logs end-to-end (E2E) latency, packet delivery ratio, message-drop sensitivity, age of information (AoI), communication overhead, edge delay, RSU coverage, and occlusion-specific safety metrics. A lightweight communication-aware simulator samples baseline-specific communication states, evaluates candidate maneuvers {brake, yield, balanced, proceed, evade}, and reports both ethical and IoV-system metrics under matched scenario seeds.

The 2000 paired scenarios were generated using a lightweight in-house Python-based traffic-risk simulator rather than CARLA or SUMO. The simulator implements a discrete maneuver set {brake, yield, balanced, proceed, evade}, scenario-dependent occlusion, density, speed, ambiguity, and VRU-presence factors, and group-wise risk estimation for ego occupants, third-party road users, and VRUs. Pedestrians and cyclists in heterogeneous-crowd scenarios are modeled as mixed VRU agents with sampled initial positions, speeds, occlusion exposure, and crossing intent. Their motion follows rule-based constant-velocity or gap-acceptance trajectories with scenario-dependent uncertainty. Fully interactive social-force or game-theoretic pedestrian responses are not modeled and are therefore listed as a limitation. Table 1 summarizes the scenario-generator configuration and heterogeneous VRU modeling assumptions.

Table 1: Scenario-generator and heterogeneous vulnerable-road-user (VRU) modeling details.

Component	Implementation	Purpose/Limitation
Simulator platform	Lightweight in-house Python-based traffic-risk simulator	Clarifies that the scenario bank was not generated in CARLA or SUMO.
Scenario families	Unprotected left turns; occluded crossings; right-of-way ambiguity; heterogeneous crowds	Four scenario classes are sampled under common paired seeds. 25% each ($n = 500$).

(Continued)

Table 1 (continued)

Component	Implementation	Purpose/Limitation
Decision space	{brake, yield, balanced, proceed, evade}	Discrete maneuver set used by the candidate planner.
Scenario variables	Density, occlusion, ambiguity, speed, VRU presence	Variables control hazard level and group-wise risk.
Mixed VRU crowd	Pedestrians and cyclists with sampled speed, crossing intent, and occlusion exposure.	Represents heterogeneous VRU conditions.
VRU motion model	Rule-based constant-velocity or gap-acceptance trajectories with uncertainty	Reactive social-force or game-theoretic behavior is not modeled.

3.3 Tail-Focused Reporting and Fairness Diagnostics

Given our hypothesis that averages obscure ethical failures, we adopted a Tail-Focused Reporting strategy. We focus on the “Top-K” outcomes—the worst scenarios in the dataset.

Let K denote the tail set size. We set $K = 100$. For each group i , we extract the top- K risk values from the set $\{R_i^{(n)}\}_{n=1}^N$. The distribution of these top-100 values characterizes the extreme outcomes (the “tail”).

Beyond tail visualization, we computed four fairness diagnostics:

1. Per-Scenario Gini Coefficient ($G^{(n)}$)

The Gini coefficient is widely used in economics to measure inequality. Here, it is adapted to quantify risk inequality across the three road-user groups. For a specific scenario n , the Gini coefficient is calculated from the risk vector $\mathbf{R}^{(n)}$:

$$G^{(n)} = \frac{\sum_{a=1}^m \sum_{b=1}^m |R_a^{(n)} - R_b^{(n)}|}{2m \sum_{a=1}^m R_a^{(n)} + \epsilon_G}, m = 3 \quad (14)$$

where $m = 3$ is the number of groups, and $\epsilon_G > 0$ is a small stabilizer. Interpretation: $G^{(n)} = 0$ implies perfect equality (all groups face identical risk). $G^{(n)} \approx 1$ implies maximal inequality (one group bears all the risk). A lower mean Gini indicates a method that consistently distributes burden equitably.

2. Worst-Off Risk ($R_{\max}^{(n)}$)

This metric tracks the magnitude of risk faced by the least fortunate group in a scenario:

$$R_{\max}^{(n)} = \max_{i \in \{\text{Ego}, \text{Third}, \text{VRU}\}} R_i^{(n)} \quad (15)$$

Reporting the distribution of $R_{\max}^{(n)}$ —including its mean and tail statistics—enables evaluation of adherence to the Rawlsian maximin principle, which prioritizes minimizing the risk experienced by the worst-off group.

3. Worst-Off Group Distribution ($g^{(n)}$)

We identify which group is the worst-off in each scenario:

$$g^{(n)} = \arg \max_{i \in \{\text{Ego}, \text{Third}, \text{VRU}\}} R_i^{(n)} \quad (16)$$

We report the empirical frequencies of $g^{(n)}$ across the 2000 scenarios.

Significance: If $g^{(n)} = \text{VRU}$ in a disproportionately high percentage of cases (e.g., >33%), it suggests a systemic bias against vulnerable users. Reducing this frequency is a key ethical objective.

4. VRU/Ego Tail Ratios

To quantify the specific relationship between VRU and Ego risks in the tail, we compute quantile ratios:

$$\text{Ratio}_p = \frac{Q_p(R_{\text{VRU}}^{(n)})}{Q_p(R_{\text{Ego}}^{(n)}) + \epsilon_R}, p \in \{0.95, 0.99\} \quad (17)$$

and the Top-K Median Ratio:

$$\text{Ratio}_{\text{TopK}} = \frac{\text{median}(\text{TopK}\{R_{\text{VRU}}^{(n)}\})}{\text{median}(\text{TopK}\{R_{\text{Ego}}^{(n)}\}) + \epsilon_R} \quad (18)$$

where $\epsilon_R > 0$ is a numerical stabilizer. Ratios below unity (Ratio < 1.0) indicate that VRU tail risk does not exceed Ego tail risk, consistent with stronger VRU protection.

3.4 Metric Definitions and Track Separation

To avoid ambiguity, this study separates the planner-objective evaluation track from the communication-aware IoV evaluation track. The planner-objective track reports physical collision rate, expected total harm, Gini coefficient, worst-off risk, and VRU tail risk. The IoV track reports communication-conditioned collision proxy, normalized total harm, latency, AoI, dropout index, and occluded-scene VRU risk. These metrics are used for within-track comparison only and should not be numerically compared across tracks. Table 2 provides compact definitions of the collision-related and harm-related metrics used in the planner-objective and IoV evaluation tracks.

Table 2: Compact definitions of collision-related and harm-related metrics.

Metric	Track	Definition	Use
Physical collision indicator C_n	Planner-objective	Binary event; $C_n = 1$ if simulator collision geometry indicates contact.	Physical collision event.
Collision rate	Planner-objective	$(1/N) \sum_n C_n$.	Compare planning baselines.
Collision proxy $\tilde{C}_n$	IoV track	Surrogate likelihood from selected maneuver risk vector and communication state.	Compare connectivity settings only.

(Continued)

Table 2 (continued)

Metric	Track	Definition	Use
Total harm	Planner-objective	Ex post harm over paired physical-simulation scenarios.	Physical-simulation harm.
Normalized total harm	IoV track	Communication-conditioned surrogate harm.	Within-track IoV comparison.
VRU p_{95}/p_{99}	Both, separately	Tail quantiles within each track.	Do not compare across tracks.

4 Results

Across the present $N = 2000$ paired simulation scenarios, Ethical-Improved shows the most favorable combined pattern of safety, distributive fairness, and VRU tail protection among the four compared planner baselines (Figs. 1–6; Tables 3–8). These findings are specific to the common simulator, scenario bank, and fixed parameter setting; therefore, they support compatibility between fairness regularization and safety in this experimental setting rather than a universal claim that both objectives always improve simultaneously.

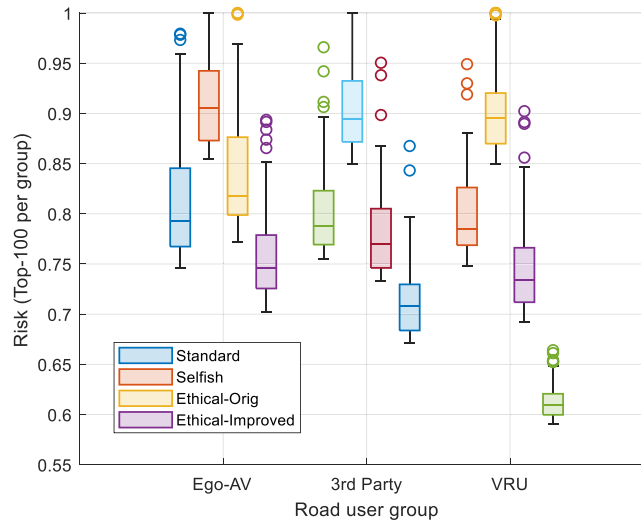


Figure 1: Top-100 risk distribution per road-user group (ego autonomous vehicle (Ego-AV), Third-party, and vulnerable road user (VRU)). Distributions are computed from the $K = 100$ highest risks per group across $N = 2000$ paired scenarios for each method; Ethical-Improved lowers the VRU upper tail while keeping Ego risk comparable.

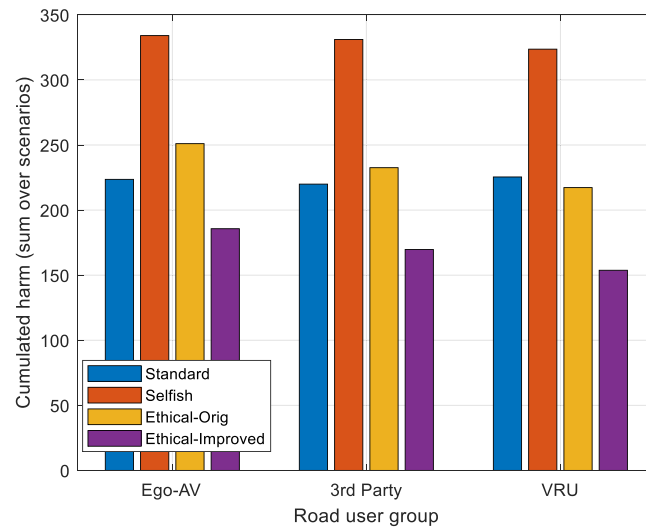


Figure 2: Cumulative personal harm (sum over scenarios) per road user group. Ethical-Improved yields the lowest cumulative harm across all three groups, indicating improved overall safety.

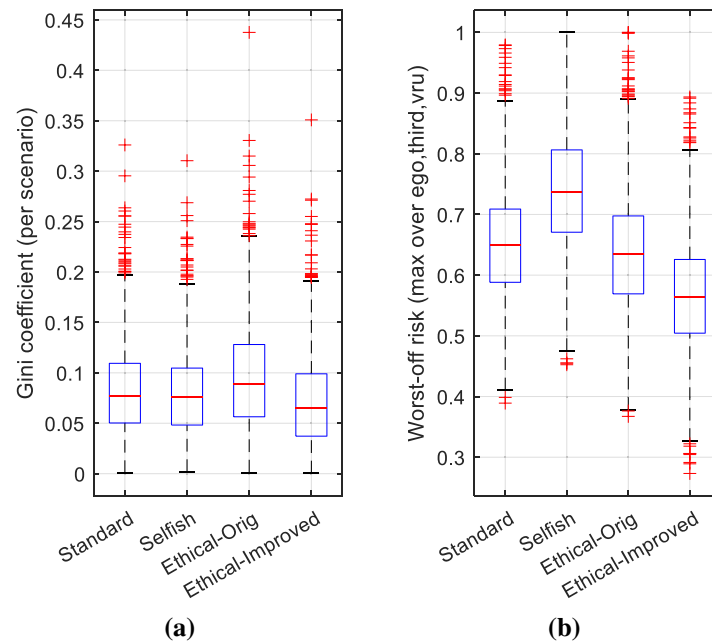


Figure 3: Inequality diagnostics across scenarios. (a) Per-scenario Gini coefficient. (b) Worst-off risk (max across groups). Ethical-Improved shifts both distributions downward, indicating reduced inequality and lower worst-case outcomes.

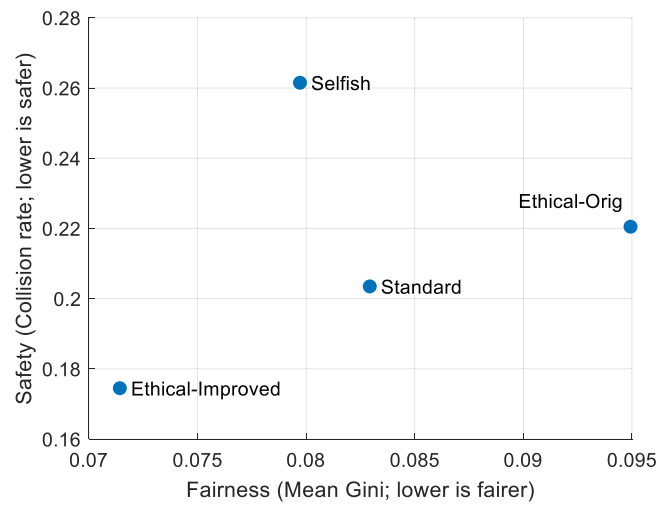


Figure 4: Fairness-safety trade-off. *X*-axis: mean Gini (lower is fairer). *Y*-axis: collision rate (lower is safer). Ethical-Improved occupies the lower-left region, combining improved fairness and safety with a quantified travel-time cost.

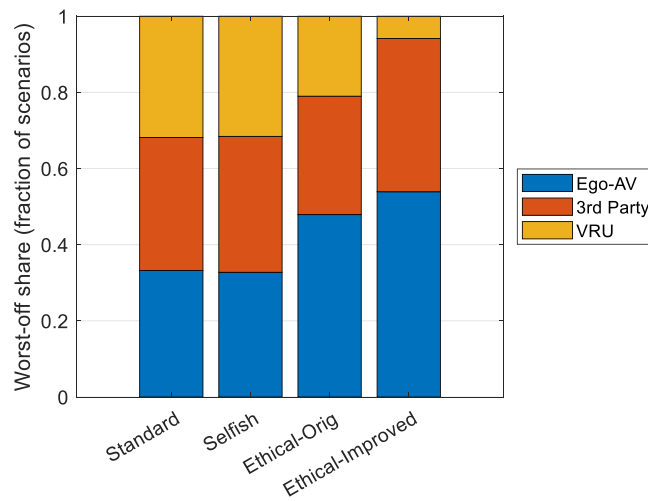


Figure 5: Worst-off group distribution: fraction of scenarios in which each group attains the maximum risk. Ethical-Improved substantially reduces the probability that VRU is worst-off.

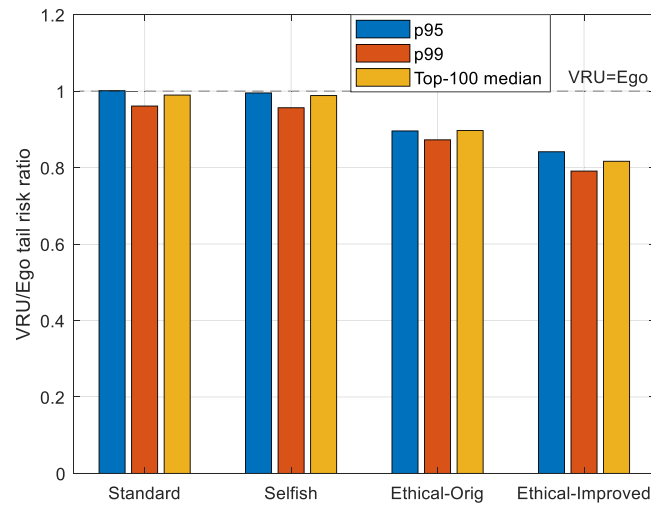


Figure 6: VRU/Ego tail ratios at p_{95} , p_{99} , and top-100 median. Ratios below 1 indicate stronger VRU tail protection relative to Ego; Ethical-Improved achieves the lowest ratios across summaries.

Table 3: Safety and efficiency metrics.

Method	Collision Rate	Total Harm	Ego p_{95}	VRU p_{95}	Ego p_{99}	VRU p_{99}	Average Time (s)
Standard	0.204	668.980	0.746	0.747	0.861	0.828	42.395
Selfish	0.262	988.740	0.854	0.850	0.971	0.928	32.477
Ethical-Orig	0.221	700.880	0.772	0.692	0.887	0.774	42.528
Ethical-Improved	0.175	509.160	0.702	0.591	0.793	0.627	49.481

Table 4: Fairness metrics.

Method	Mean Gini	Gini p_{95}	Mean Max Risk	VRU Risk Share	Worst-Off Group
Standard	0.083	0.162	0.652	0.332	Third-party
Selfish	0.080	0.156	0.741	0.332	Third-party
Ethical-Orig	0.095	0.185	0.638	0.308	Ego-AV
Ethical-Improved	0.071	0.147	0.567	0.300	Ego-AV

Table 5: Worst-off distribution and VRU/Ego tail ratios.

Method	Worst Ego	Worst Third-Party	Worst VRU	VRU/Ego p_{95}	VRU/Ego p_{99}	VRU/Ego Top-100
Standard	0.333	0.349	0.319	1.001	0.961	0.990
Selfish	0.328	0.357	0.316	0.995	0.956	0.988
Ethical-Orig	0.479	0.311	0.210	0.896	0.873	0.897
Ethical-Improved	0.539	0.403	0.059	0.841	0.791	0.817

Table 6: Risk redistribution and efficiency cost of Ethical-Improved.

Metric	Standard	Ethical-Orig	Ethical-Improved	Interpretation
Collision rate	0.204	0.221	0.175	Lower physical collision frequency.
Total harm	668.980	700.880	509.160	Lower aggregate harm.
Mean max risk	0.652	0.638	0.567	Lower worst-off severity.
Ego p_{95}/p_{99}	0.746/0.861	0.772/0.887	0.702/0.793	Ego absolute tail risk decreases.
VRU p_{95}/p_{99}	0.747/0.828	0.692/0.774	0.591/0.627	VRU tail risk decreases substantially.
Ego worst-off freq.	0.333	0.479	0.539	Higher relative risk ranking.
VRU worst-off freq.	0.319	0.210	0.059	VRU tail dominance is suppressed.
Average travel time	42.395 s	42.528 s	49.481 s	16.7% increase vs. Standard.

Table 7: Selected paired significance tests for Ethical-Improved (negative mean differences indicate reductions vs. baseline).

Comparison	Metric	Mean Per-Scenario Difference	95% Confidence Interval	<i>p</i>	Test
Ethical-Improved vs. Standard	Collision	−0.029	[−0.053, −0.005]	0.022	McNemar exact
Ethical-Improved vs. Standard	Total harm	−0.080	[−0.118, −0.040]	<0.001	Permutation
Ethical-Improved vs. Standard	Gini	−0.012	[−0.013, −0.010]	<0.001	Permutation
Ethical-Improved vs. Standard	Worst-off risk	−0.085	[−0.088, −0.083]	<0.001	Permutation
Ethical-Improved vs. Standard	VRU risk	−0.109	[−0.112, −0.105]	<0.001	Permutation
Ethical-Improved vs. Selfish	Collision	−0.087	[−0.113, −0.061]	<0.001	McNemar exact
Ethical-Improved vs. Selfish	Total harm	−0.240	[−0.285, −0.196]	<0.001	Permutation
Ethical-Improved vs. Selfish	Gini	−0.008	[−0.010, −0.007]	<0.001	Permutation

(Continued)

Table 7 (continued)

Comparison	Metric	Mean Per-Scenario Difference	95% Confidence Interval	<i>p</i>	Test
Ethical-Improved vs. Ethical-Orig	Collision	−0.046	[−0.071, −0.021]	<0.001	McNemar exact
Ethical-Improved vs. Ethical-Orig	Gini	−0.024	[−0.024, −0.023]	<0.001	Permutation
Ethical-Improved vs. Ethical-Orig	VRU risk	−0.049	[−0.053, −0.046]	<0.001	Permutation

Table 8: Paired robustness checks, raw paired effects, standardized effect sizes, and confidence intervals for Ethical-Improved vs. Ethical-Orig.

Metric	Test	<i>p</i>	Raw Paired Effect	Standardized Effect Size	95% CI
Collision	McNemar exact	<0.001	Paired risk difference = −0.046	Cohen's $h = -0.116$; Cohen's d not applicable	Risk difference CI [−0.071, −0.021]; h CI [−0.183, −0.052]
Total harm	Permutation + Wilcoxon	<0.001	Mean difference = −0.0959; Relative reduction = 27.4%	Cohen's $d_z = -0.192$	Raw CI [−0.1188, −0.0748]; d_z CI [−0.212, −0.172]
Gini	Permutation + Wilcoxon	<0.001	Mean difference = −0.024	Cohen's $d_z = -2.104$	Raw CI [−0.024, −0.023]; d_z CI [−2.236, −1.988]
Worst-off risk	Permutation + Wilcoxon	<0.001	Mean reduction = −0.071	Cohen's $d_z = -0.468$	Raw CI [−0.078, −0.064]; d_z CI [−0.492, −0.445]
VRU risk (mean)	Permutation + Wilcoxon	<0.001	Mean difference = −0.049	Cohen's $d_z = -0.614$	Raw CI [−0.053, −0.046]; d_z CI [−0.643, −0.585]
VRU p_{95}/p_{99} risk	Bootstrap/paired quantile comparison	<0.001	0.692/0.774 → 0.591/0.627	Quantile shift; d_z not defined	p_{95} shift CI [−0.118, −0.088]; p_{99} shift CI [−0.155, −0.136]
Travel time	Permutation + Wilcoxon	<0.001	+6.953 s efficiency cost	Cohen's $d_z = +0.674$	Raw CI [+6.508, +7.417] s; d_z CI [+0.644, +0.706]

Note: d_z denotes Cohen's standardized effect size for matched continuous outcomes; collision outcomes are binary and are therefore summarized by paired risk difference and Cohen's h .

4.1 Tail Risk Suppression

Fig. 1 visualizes the distribution of extreme risks using the top-100 risk values per group. Relative to the baselines, Ethical-Improved compresses the upper tail of VRU risk, consistent with the reductions in VRU p_{95}/p_{99} in Table 3. Importantly, Ego tail risk does not increase; Ego p_{95}/p_{99} also decreases to 0.702/0.793, indicating that the tail-regularized objective improves safety without trading VRU protection for greater ego exposure.

For the planner-objective results in Tables 3–8, the Ethical-Improved objective used a static VRU tail threshold calibrated at $\tau = q_{0.95}$. Communication-state-dependent threshold tightening is reserved for the communication-aware extension and policy-level calibration; the robustness of $q_{0.90}$ – $q_{0.99}$ choices is examined in Section 4.8.

4.2 Safety Performance: Collision Rates and Total Harm

Table 3 reports aggregate safety and efficiency metrics. Ethical-Improved attains the lowest collision rate (0.175), improving on Standard (0.204), Selfish (0.262), and Ethical-Orig (0.221), and also yields the lowest total harm (509.16 vs. 668.98, 988.74, and 700.88, respectively). These gains coincide with reduced tail risks for both Ego and VRU; notably, VRU tail risk declines to 0.591 (p_{95}) and 0.627 (p_{99}). Fig. 2 shows the cumulative personal harm across scenarios for each road-user group.

4.3 Fairness and Distributive Justice

Beyond aggregate safety, we evaluate distributive fairness using inequality and worst-off diagnostics that explicitly capture ethically salient tail outcomes.

1. Gini Coefficient (Inequality)

Table 4 summarizes the fairness metrics. Ethical-Improved achieves the lowest mean Gini (0.071) and the lowest 95th-percentile Gini (0.147), indicating a more even allocation of risk across groups than Standard (0.083) and Ethical-Orig (0.095). It also reduces the mean worst-off risk to 0.567, suggesting a systematic reduction in severe outcomes.

The compression of the VRU upper tail in Fig. 1 indicates that the proposed objective does not merely shift the median risk distribution but specifically targets rare high-risk VRU outcomes. This pattern is consistent with the squared hinge tail penalty, which becomes active only when VRU risk exceeds the calibrated threshold. The simultaneous reduction in Ego p_{95}/p_{99} risk suggests that the improvement is not achieved by transferring tail risk from VRUs to ego occupants, but by encouraging earlier conservative maneuvers in high-uncertainty conflict zones.

Fig. 2 should be interpreted as an aggregate harm decomposition rather than a fairness metric alone. The lower cumulative harm across all groups indicates that the proposed regularizers suppress severe conflict outcomes that contribute disproportionately to total harm. This supports the interpretation that fairness regularization improves safety in the present scenario bank because tail events are also high-harm events.

Fig. 3 supports these results at the distribution level: the per-scenario Gini distribution shifts left under Ethical-Improved (Fig. 3a), and the worst-off risk distribution also shifts downward (Fig. 3b). The Selfish baseline has a modest mean-Gini value but much worse collision and harm results, showing that inequality metrics must be read together with safety and tail diagnostics.

2. Worst-Off Group Frequencies

Table 5 reports the worst-off-group distribution and VRU/Ego tail ratios. Ethical-Improved reduces the probability that VRUs are the worst-off group to 0.059, compared with 0.319 for Standard, 0.316 for Selfish,

and 0.210 for Ethical-Orig. Concurrently, the VRU/Ego tail ratios decline to 0.841 (p_{95}) and 0.791 (p_{99}), reflecting stronger relative protection for VRUs in extreme cases.

The lower-left position of Ethical-Improved in Fig. 4 shows that, in this simulator, fairness and collision reduction are not in conflict. However, this point should not be interpreted as a Pareto-free improvement because Tables 3 and 6 show a clear travel-time cost. The figure therefore supports a constrained policy interpretation: the planner improves safety and distributive fairness, but mobility managers must decide whether the associated delay is acceptable for a given urban context.

Fig. 5 provides a complementary view of worst-off-group frequencies. The stacked distribution shows that Ethical-Improved nearly eliminates VRUs as the maximum-risk group across scenarios. In contrast, the increased share of Ego as the worst-off group is accompanied by a lower overall worst-off severity (Table 4).

3. Tail Ratios (VRU vs. Ego)

Fig. 6 summarizes VRU-to-Ego tail-risk ratios at p_{95} , p_{99} , and the top-100 median. Ethical-Improved achieves ratios below 1 across all three summaries (0.841, 0.791, and 0.817), indicating that even in the tail the VRU risk remains below the Ego risk.

Fig. 5 indicates a redistribution of worst-off status away from VRUs. This is ethically desirable only when interpreted together with absolute risk levels: the increased Ego worst-off frequency does not imply higher Ego tail exposure, because Ego p_{95}/p_{99} and mean max risk also decrease. Thus, the figure supports the claim that VRU tail dominance is suppressed without creating an ego-sacrifice policy.

The fact that all VRU/Ego tail ratios remain below 1 under Ethical-Improved indicates that the ratio cap affects the most ethically sensitive part of the distribution rather than only mean outcomes. This provides direct evidence for the protective-ratio mechanism in Eq. (7).

4.4 The Efficiency Trade-Off

Fig. 4 depicts the fairness-safety trade-off by plotting the mean Gini against collision rate. Ethical-Improved lies in the lower-left region, achieving both lower inequality (Table 4) and lower collision rate (Table 3) than the baselines. This result indicates that fairness regularization can be compatible with safety improvement, although it comes with a measurable efficiency cost.

Risk Redistribution and Efficiency Cost

Ethical-Improved improves safety and tail-risk allocation, but it also imposes a measurable efficiency cost. Average travel time increases from 42.395 s under Standard to 49.481 s under Ethical-Improved, corresponding to a 16.7% increase relative to Standard. Relative to Ethical-Orig, travel time increases from 42.528 to 49.481 s, corresponding to a 16.4% increase. Therefore, the efficiency cost should not be interpreted as negligible; it reflects the operational cost of earlier yielding, more conservative conflict handling, and stronger VRU tail-risk protection in high-risk urban scenarios.

4.5 Statistical Significance

Table 7 reports paired statistical tests over the $N = 2000$ matched scenarios. For transparency, the p -values shown are the raw paired-test values for the prespecified comparisons, together with mean differences and 95% confidence intervals. Because 11 comparisons are reported in Table 7, Holm familywise correction was additionally applied; all listed comparisons remain significant at $\alpha = 0.05$ after correction. The statistical evidence should nevertheless be interpreted together with effect sizes and directional consistency.

Table 8 summarizes the paired robustness checks, effect sizes, and confidence intervals for the primary comparison between Ethical-Improved and Ethical-Orig. In addition to p values and confidence intervals,

effect magnitudes are reported to avoid interpreting statistical significance alone. Binary collision outcomes are evaluated with McNemar’s exact test and paired risk difference. Continuous paired outcomes are evaluated with permutation tests and Wilcoxon signed-rank robustness checks. For continuous paired outcomes, standardized effect size was computed as Cohen’s d_z , defined as $d_z = \frac{\bar{D}}{s_D}$, where $D_n = X_n(Ethical - Improved) - X_n(Ethical - Orig)$ for matched scenario n . Bootstrap 95% confidence intervals were computed over paired scenario indices. For binary collision outcomes, Cohen’s d is not directly applicable; therefore, McNemar’s paired risk difference and Cohen’s h are reported.

4.6 Communication-Aware Internet-of-Vehicles Results

This subsection introduces a second evaluation track that is distinct from the planner-objective comparisons reported in Tables 3–8. Here, the communication-aware objective is evaluated under five connectivity/perception settings—local-only, V2X, RSU-assisted, edge-assisted, and communication-impaired—using common scenario seeds so that the reported differences reflect IoV information quality rather than scenario mismatch. Accordingly, Table 9 summarizes the within-track connected-baseline results, Table 10 documents the communication-state sampling rules, and Table 11 reports the corresponding communication and occlusion metrics. These IoV-track tables should not be directly compared numerically with the planner-baseline results in Tables 3–8.

Table 9: Connectivity-baseline safety and fairness metrics under communication-aware IoV evaluation (within-track comparison only; not directly comparable with Tables 3–8).

Connectivity Setting	Collision Proxy	Normalized Total Harm	Mean Gini	VRU p_{95} Risk	VRU/Ego p_{95} Ratio	Avg. Travel Time (s)
Local	0.266	415.817	0.049	0.912	0.961	53.447
V2X	0.257	387.580	0.048	0.854	0.920	53.665
Roadside unit (RSU)	0.232	309.573	0.047	0.719	0.874	52.269
Edge	0.225	291.391	0.046	0.684	0.853	51.722
Comm.-impaired	0.272	429.883	0.051	0.894	0.911	54.251

Note: The metrics in this table are reported for connectivity/perception baselines under the sampled IoV communication states. “Collision Proxy” and “Normalized Total Harm” are used only for within-table comparison in the communication-aware evaluation track and should not be numerically compared with the “Collision Rate” and ex post “Total Harm” reported in Tables 3–8.

Table 10: Communication-state sampling rules used in the IoV evaluation.

Setting	V2X	Packet Delivery Ratio (PDR)	Age of Information (AoI)	Roadside Unit (RSU)	q_{perc}	Edge/Transmission (Tx) Delay	End-to-End (E2E) Latency
Local	0.00	N/A	5 ± 2.75 ms	0.00	0.62	0/18 ms	34.041 ms
V2X	0.88	0.92 ± 0.025	42 ± 8.30 ms	0.15	0.74	0/36 ms	52.043 ms
Roadside unit (RSU)	0.93	0.95 ± 0.025	46 ± 8.90 ms	0.80	0.83	0/42 ms	58.026 ms
Edge	0.98	0.975 ± 0.025	44 ± 8.60 ms	0.95	0.96	9/44 ms	68.810 ms
Comm.-impaired	0.62	0.74 ± 0.06	135 ± 22.25 ms	0.35	0.57	28/84 ms	127.753 ms

Note: PDR = packet delivery ratio; AoI = age of information; RSU = roadside unit; Tx = transmission; E2E = end-to-end.

Table 11: Connectivity-baseline communication and occlusion metrics under communication-aware IoV evaluation. “Comm.-impaired” denotes communication-impaired operation under degraded V2X/RSU conditions.

Case	Latency (ms)	AoI (ms)	Communication Overhead (kB)	Edge Delay (ms)	Drop Rate	Occluded Collision Proxy	Occluded VRU p_{95} Risk
Local	34.041	5.055	0.000	0.241	0.000	0.299	1.003
V2X	52.043	42.093	14.944	0.242	0.039	0.288	0.939
Roadside unit (RSU)	58.026	46.003	23.662	0.226	0.034	0.253	0.754
Edge	68.810	44.043	32.380	9.010	0.032	0.244	0.726
Comm.-impaired	127.753	134.826	11.208	27.953	0.053	0.301	0.984

Table 9 summarizes the IoV-system results, Table 10 documents the communication-state sampling rules, and Table 11 reports the connectivity-baseline communication and occlusion metrics. Edge-assisted cooperative perception achieves the strongest overall trade-off among the connected baselines, yielding the lowest collision proxy (0.225), normalized total harm (291.391), mean Gini (0.046), and VRU p_{95} risk (0.684). Although edge assistance adds nonzero latency, it improves dropout robustness, lowers the occluded-scenario collision proxy to 0.244, and reduces occluded-scene VRU p_{95} risk to 0.726.

The communication-state sampling rules are reported in Table 10, and the connectivity-baseline communication and occlusion metrics are reported in Table 11. “Comm.-impaired” denotes communication-impaired operation under degraded V2X/RSU conditions.

4.7 Communication-State Sampling and Definition of the Communication-Impaired Baseline

The communication-impaired baseline, denoted as “Comm.-impaired” in Tables 9–11, is not a separate planning objective. It uses the same communication-aware Ethical-Improved objective as the other IoV settings but samples z_t from a degraded V2X/RSU operating regime. This regime is characterized by low V2X availability, reduced PDR, high AoI, partial RSU outage, reduced perception confidence, increased edge delay, and higher transmission latency.

The high latency of the communication-impaired condition results from $t_{tx} = 84$ ms, edge-delay mean of 28 ms, high AoI, and unstable RSU coverage. Consequently, “Comm.-impaired” represents stale, lossy, and partially unavailable cooperative perception rather than a different ethical planner.

4.8 One-at-a-Time Sensitivity Analysis

To evaluate parameter robustness, we conducted a compact one-at-a-time sensitivity analysis for the Ethical-Improved planner under the edge-assisted cooperative perception setting. All runs used the same scenario bank, communication-state samples, estimated-risk noise, and collision random numbers, so that differences reflect parameter perturbations rather than random resampling. The tail-risk threshold τ was calibrated from the empirical VRU-risk distribution of the reference yielding maneuver, yielding τ values of 0.639, 0.681, 0.720, and 0.776 for $q_{0.90}$, $q_{0.95}$, $q_{0.975}$, and $q_{0.99}$, respectively.

The sensitivity metrics follow the communication-aware IoV track definitions. Table 12 is based on an independent sensitivity-run reference case with common random numbers; therefore, its baseline values are close to, but not identical to, the Edge row in Table 9.

Table 12: Compact one-at-a-time sensitivity analysis of Ethical-Improved. Sweep rows report the minimum and maximum values across the tested settings.

Parameter	Tested Settings	Safety Metrics	Tail-Risk Metrics	Fairness/ Efficiency
Baseline	$\tau = q_{0.95}; \varepsilon = 0.05;$ $\lambda_T = 1.0; \lambda_\rho = 1.0;$ $\rho = 1.0$	Collision proxy: 0.224	VRU p_{95}/p_{99} : 0.670/0.752	Gini: 0.048
		Normalized harm: 283.959	Ego p_{95}/p_{99} : 0.787/0.874; ratio: 0.850	Average travel time: 51.054 s
τ threshold	$q_{0.90}-q_{0.99}$	Collision proxy: 0.224	VRU p_{95}/p_{99} : 0.665– 0.690/0.744–0.790	Gini: 0.048
		Normalized harm: 283.450–285.596	Ego p_{95}/p_{99} : 0.783– 0.788/0.869–0.876; ratio: 0.844–0.881	Average travel time: 50.936–51.097 s
ε	0.01–0.10	Collision proxy: 0.224	VRU p_{95}/p_{99} : 0.668–0.670/0.751–0.752	Gini: 0.048
		Normalized harm: 283.858–283.959	Ego p_{95}/p_{99} : 0.787/0.874; ratio: 0.849–0.850	Average travel time: 51.054–51.061 s
λ_T	0–1.0	Collision proxy: 0.224	VRU p_{95}/p_{99} : 0.670–0.695/0.752–0.812	Gini: 0.047–0.048
		Normalized harm: 283.959–285.952	Ego p_{95}/p_{99} : 0.783– 0.787/0.869–0.874; ratio: 0.850–0.888	Average travel time: 50.917–51.054 s
λ_ρ	0–1.0	Collision proxy: 0.224–0.225	VRU p_{95}/p_{99} : 0.670– 0.679/0.752–0.764	Gini: 0.048
		Normalized harm: 283.959–285.056	Ego p_{95}/p_{99} : 0.786–0.787/0.874; ratio: 0.850–0.864	Average travel time: 50.934–51.054 s
ρ	0.8–1.2	Collision proxy: 0.222–0.225	VRU p_{95}/p_{99} : 0.652– 0.679/0.730–0.764	Gini: 0.048–0.052
		Normalized harm: 281.167–284.966	Ego p_{95}/p_{99} : 0.786– 0.788/0.874–0.876; ratio: 0.827–0.864	Average travel time: 50.942–51.608 s
Normalization	Normalized vs. raw	Collision proxy: 0.222–0.224	VRU p_{95}/p_{99} : 0.657–0.670/0.744–0.752	Gini: 0.048–0.054
		Normalized harm: 280.932–283.959	Ego p_{95}/p_{99} : 0.787–0.788/0.874; ratio: 0.834–0.850	Average travel time: 51.054–51.875 s

Note: The empirical τ thresholds for $q_{0.90}$, $q_{0.95}$, $q_{0.975}$, and $q_{0.99}$ were 0.639, 0.681, 0.720, and 0.776, respectively. Safety metrics report collision proxy and normalized harm; the VRU/Ego ratio is computed at p_{95} .

Table 12 reports a compact range summary. Across the tested perturbations, collision proxy remains within 0.222–0.225, mean Gini remains within 0.047–0.054, and average travel time remains within 50.917–51.875 s. The τ sweep shows the expected safety-efficiency pattern: stricter thresholds reduce VRU p_{95}/p_{99} risk but slightly increase travel time. Fig. 7 visualizes this τ -threshold sensitivity for VRU p_{95} risk and average

travel time. The ρ sweep provides the clearest ratio-control effect, with $\rho = 0.8$ yielding the lowest VRU/Ego p_{95} ratio at the cost of higher travel time.

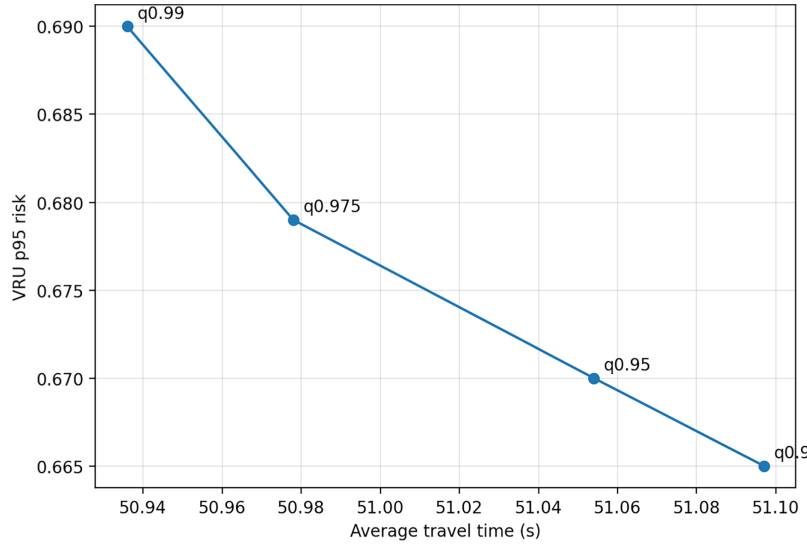


Figure 7: Sensitivity of VRU p_{95} risk and travel time across τ thresholds. Stricter τ thresholds reduce VRU tail risk at a modest travel-time cost.

5 Discussion

The results provide three main implications for ethical IoV risk allocation and should be interpreted as relative objective behavior within a controlled simulator. The comparison between Ethical-Orig and Ethical-Improved shows that adding equality and maximin terms to a weighted sum is insufficient when component scales are misaligned. The squared hinge loss acts as a soft constraint on extreme VRU outcomes, while normalized fairness terms make the trade-off weights more interpretable across heterogeneous scenarios. The bounded responsibility factor supports treating responsibility as an internal risk weight rather than as an external subtractive bonus. The formulation can discount risk for agents who contributed to a hazard while still preserving a nonzero duty-of-care incentive to avoid collisions whenever feasible. Within this simulator, fairness regularization is compatible with improved safety because severe tail events contribute disproportionately to cumulative harm. However, this should not be read as a universal law. The weights are policy parameters, the scenario bank is simulator-generated, and future work should perform component ablations, broader multivariate calibration sweeps, packet-level network co-simulation, public-benchmark transfer, and real roadside perception validation. The compact one-at-a-time sensitivity analysis in [Section 4.8](#) shows that the main safety, VRU-tail-risk, fairness, and travel-time trends remain stable across tested perturbations of τ , ϵ , λ_T , λ_ρ , ρ , and normalization.

5.1 Ego Worst-Off Frequency and Policy-Calibrated Deployment

The increase in Ego worst-off frequency under Ethical-Improved requires explicit ethical interpretation. Ethical-Improved increases the frequency with which Ego is the maximum-risk group to 0.539, while reducing the VRU worst-off frequency to 0.059. This does not mean that the proposed planner requires ego occupants to self-sacrifice for VRUs. Rather, it indicates that after disproportionate VRU tail dominance is constrained, the Ego more often becomes the largest remaining risk component in the group-wise ranking. This ranking change should be interpreted together with absolute risk levels: Ego p_{95}/p_{99} risks decrease from 0.772/0.887 under Ethical-Orig to 0.702/0.793 under Ethical-Improved, and mean max risk decreases from

0.638 to 0.567. Thus, the result reflects a change in relative risk ranking rather than an increase in absolute Ego tail exposure.

Nevertheless, a VRU-protective planner should not be interpreted as authorizing unlimited risk transfer to ego occupants. Practical deployment should include an explicit ego safety floor, hard collision-avoidance constraints, maximum allowable ego-risk thresholds, and legal duty-of-care requirements toward passengers. From a public-acceptance perspective, limited altruistic yielding may be acceptable when it reduces severe VRU harm, but users are unlikely to accept a system that substantially increases passenger danger. The proposed method should therefore be understood as a policy-calibrated ethical planner rather than a universal moral rule.

5.2 Mechanism Relative to the Baseline Ethical Planner

The results suggest that Ethical-Improved improves both safety and fairness by suppressing severe tail events rather than merely redistributing risk. Compared with Ethical-Orig and the baseline ethical planner in [1], the proposed formulation changes three mechanisms: responsibility is internalized as a bounded risk weight, fairness terms are normalized to reduce scale sensitivity, and VRU tail-risk dominance is penalized directly. The observed reductions in both VRU and Ego p_{95}/p_{99} risk are consistent with earlier conservative conflict handling, such as braking or yielding before the vehicle enters a high-uncertainty conflict zone.

5.3 Practical Implications of the Travel-Time Cost

The 16.7% travel-time increase has practical implications for deployment. At the individual-vehicle level, an additional 7.086 s per high-risk segment may be acceptable in school zones, unsignalized crossings, or dense VRU areas, where severe VRU harm reduction is a priority. At the traffic-system level, however, repeated conservative yielding could reduce intersection throughput, increase queue formation, and create secondary delays during peak demand. User acceptability may therefore depend on whether the delay is transparent, predictable, and limited to high-risk contexts. Mobility managers may need to calibrate the tail-risk threshold and ratio penalty by road type, time of day, VRU density, and congestion state rather than deploy fixed ethical weights across all contexts. The present study does not quantify network-level throughput or queue spillback, so this operational impact should be evaluated in future microscopic and macroscopic traffic simulations before deployment.

5.4 Prioritized Limitations and Deployment Priorities

No real-world dataset or hardware-in-the-loop experiment was used in the present study. The current validation is limited to controlled paired simulations generated by an in-house lightweight traffic-risk simulator. Therefore, the results should be interpreted as objective-level and mechanism-level evidence rather than deployment-ready validation.

These limitations are prioritized according to their impact on deployment validity. The most important limitation is that the results are obtained from an in-house lightweight simulator; therefore, the conclusions should be interpreted as objective-level and mechanism-level evidence rather than deployment-ready validation. The second limitation is the rule-based VRU model, which does not fully capture interactive pedestrian or cyclist responses. The third limitation is the lightweight communication model, which should be replaced or complemented by packet-level network co-simulation before deployment claims are made. To reduce the limitation of data availability, the generated scenario seeds, aggregate risk logs, metric tables, and analysis scripts will be released through the GitHub repository specified in the Data Availability statement. [Table 13](#) summarizes the prioritized limitations and corresponding deployment-validation priorities.

Table 13 : Prioritized limitations and deployment priorities.

Priority	Limitation	Effect on the Conclusion	Required Next Step
1	In-house lightweight simulator	Limits external validity	Validate in CARLA, SUMO, or public benchmarks
2	Rule-based VRU behavior	May underestimate interactive pedestrian/cyclist responses	Add social-force or game-theoretic VRU agents
3	Lightweight communication model	Does not capture packet-level network dynamics	Add ns-3, C-V2X, or 5G NR-V2X co-simulation
4	One-at-a-time sensitivity only	Does not prove global parameter optimality	Add multivariate calibration
5	Limited deployment calibration	Policy weights may vary by jurisdiction	Add stakeholder-informed calibration

5.5 Scalability Considerations

The proposed framework is locally scalable at the ego-planner level because the optimization is performed over a fixed, discrete set of maneuvers and a small number of risk groups, rather than over all individual agents in the network. The edge/RSU layer supplies communication-conditioned perception inputs, while the ego vehicle solves the risk-allocation objective at each planning cycle. However, large-scale IoV deployment with dense traffic and multiple RSUs was not evaluated in this study. Such deployment would require runtime profiling, multi-RSU coordination, packet-level V2X co-simulation, and stress testing under high vehicle and VRU density. [Table 14](#) summarizes the validation and scalability issues that must be addressed before deployment-ready claims can be made.

Table 14: Validation and scalability clarification for deployment readiness.

Concern	Present Status in This Study	Deployment Implication/Required Next Step
Real-world dataset/hardware-in-the-loop validation	No real-world dataset or hardware-in-the-loop experiment was used. Validation is limited to controlled paired simulations generated by an in-house lightweight traffic-risk simulator.	Interpret the results as objective-level and mechanism-level evidence. Before deployment claims, validate with real-world trajectory/perception data, public benchmarks, and hardware-in-the-loop experiments.
Dense traffic/multiple-RSU scalability	The optimization is locally scoped to the ego planner, a fixed discrete set of maneuvers, and a small group-wise risk vector. Dense traffic and multiple-RSU deployment were not evaluated in this study.	Conduct runtime profiling, multi-RSU coordination tests, packet-level V2X co-simulation, and high-density vehicle/VRU stress testing.

6 Conclusions

This work presents a responsibility-aware tail-risk objective for edge-assisted IoV cooperative autonomous driving. The formulation internalizes responsibility as a bounded risk weight, normalizes fairness terms, and adds explicit VRU tail-risk and VRU/Ego ratio penalties. Under the fixed 2000-scenario simulator protocol, the method improves the combined safety-fairness profile relative to Standard, Selfish, and Ethical-Orig baselines. The one-at-a-time sensitivity analysis further shows that these conclusions are qualitatively robust across the tested values of τ , ϵ , λ_T , λ_ρ , and ρ , as well as across normalization settings. The current evidence remains limited by simulator-generated scenarios, a one-at-a-time robustness analysis rather than a full multivariate calibration study, and a lightweight communication model. Before deployment, claims should be tested with systematic ablations, jurisdiction-aware calibration, packet-level network co-simulation, public connected-driving benchmarks, and real roadside perception datasets. Overall, the manuscript should be read as an auditable objective-design framework linking ethical risk allocation with V2X availability, information freshness, RSU coverage, occlusion uncertainty, and edge inference latency. The limitations further clarify that no real-world dataset or hardware-in-the-loop experiment was used in the present study and that dense-traffic, multi-RSU scalability remains an open deployment-validation step.

Acknowledgement: Not applicable.

Funding Statement: This work was supported in part by the National Science and Technology Council, Taiwan, under Grant NSTC 114–2221-E-018–003.

Availability of Data and Materials: The generated scenario seeds, simulation logs, aggregated metric tables, and analysis scripts are available in a public GitHub repository at <https://github.com/lin040/iov-tail-risk-fairness.git>.

Ethics Approval: Not applicable.

Conflicts of Interest: The authors declare no conflicts of interest.

References

1. Geisslinger M, Poszler F, Lienkamp M. An ethical trajectory planning algorithm for autonomous vehicles. *Nat Mach Intell.* 2023;5(2):137–44. doi:10.1038/s42256-022-00607-z.
2. Rockafellar RT, Uryasev S. Optimization of conditional value-at-risk. *J Risk.* 2000;2(3):21–41. doi:10.21314/jor.2000.038.
3. McNemar Q. Note on the sampling error of the difference between correlated proportions or percentages. *Psychometrika.* 1947;12(2):153–7. doi:10.1007/BF02295996.
4. Ceriani L, Verme P. The origins of the Gini index: extracts from *Variabilità e Mutabilità*, (1912) by Corrado Gini. *J Econ Inequal.* 2012;10(3):421–43. doi:10.1007/s10888-011-9188-x.
5. Rawls J. *A theory of justice*. Cambridge, MA, USA: Belknap Press of Harvard University Press; 1971.
6. Paden B, Čáp M, Yong SZ, Yershov D, Frazzoli E. A survey of motion planning and control techniques for self-driving urban vehicles. *IEEE Trans Intell Veh.* 2016;1(1):33–55. doi:10.1109/TIV.2016.2578706.
7. Schwarting W, Alonso-Mora J, Rus D. Planning and decision-making for autonomous vehicles. *Annu Rev Control Robot Auton Syst.* 2018;1(1):187–210. doi:10.1146/annurev-control-060117-105157.
8. Katrakazas C, Quddus M, Chen WH, Deka L. Real-time motion planning methods for autonomous on-road driving: state-of-the-art and future research directions. *Transp Res Part C Emerg Technol.* 2015;60:416–42. doi:10.1016/j.trc.2015.09.011.
9. Lefèvre S, Vasquez D, Laugier C. A survey on motion prediction and risk assessment for intelligent vehicles. *ROBOMECH J.* 2014;1(1):1. doi:10.1186/s40648-014-0001-z.
10. Riedmaier S, Ponn T, Ludwig D, Schick B, Diermeyer F. Survey on scenario-based safety assessment of automated vehicles. *IEEE Access.* 2020;8:87456–77. doi:10.1109/ACCESS.2020.2993730.

11. Pek C, Manzinger S, Koschi M, Althoff M. Using online verification to prevent autonomous vehicles from causing accidents. *Nat Mach Intell*. 2020;2(9):518–28. doi:10.1038/s42256-020-0225-y.
12. Goodall NJ. Ethical decision making during automated vehicle crashes. *Transp Res Rec J Transp Res Board*. 2014;2424(1):58–65. doi:10.3141/2424-07.
13. Nyholm S, Smids J. The ethics of accident-algorithms for self-driving cars: an applied trolley problem? *Ethical Theory Moral Pract*. 2016;19(5):1275–89. doi:10.1007/s10677-016-9745-2.
14. Bonnefon JF, Shariff A, Rahwan I. The social dilemma of autonomous vehicles. *Science*. 2016;352(6293):1573–6. doi:10.1126/science.aaf2654.
15. Awad E, Dsouza S, Kim R, Schulz J, Henrich J, Shariff A, et al. The moral machine experiment. *Nature*. 2018;563(7729):59–64. doi:10.1038/s41586-018-0637-6.
16. Meder B, Fleischhut N, Krumnau NC, Waldmann MR. How should autonomous cars drive? A preference for defaults in moral judgments under risk and uncertainty. *Risk Anal*. 2019;39(2):295–314. doi:10.1111/risa.13178.
17. Sparrow R, Howard M. When human beings are like drunk robots: driverless vehicles, ethics, and the future of transport. *Transp Res Part C Emerg Technol*. 2017;80(suppl 3):206–15. doi:10.1016/j.trc.2017.04.014.
18. Martinho A, Herber N, Kroesen M, Chorus C. Ethical issues in focus by the autonomous vehicles industry. *Transp Rev*. 2021;41(5):556–77. doi:10.1080/01441647.2020.1862355.
19. Liu HY. Irresponsibilities, inequalities and injustice for autonomous vehicles. *Ethics Inf Technol*. 2017;19(3):193–207. doi:10.1007/s10676-017-9436-2.
20. Gill T. Ethical dilemmas are really important to potential adopters of autonomous vehicles. *Ethics Inf Technol*. 2021;23(4):657–73. doi:10.1007/s10676-021-09605-y.
21. Liu P, Liu J. Selfish or utilitarian automated vehicles? Deontological evaluation and public acceptance. *Int J Hum*. 2021;37(13):1231–42. doi:10.1080/10447318.2021.1876357.
22. Krügel S, Uhl M. The risk ethics of autonomous vehicles: an empirical approach. *Sci Rep*. 2024;14(1):960. doi:10.1038/s41598-024-51313-2.
23. Hevelke A, Nida-Rümelin J. Responsibility for crashes of autonomous vehicles: an ethical analysis. *Sci Eng Ethics*. 2015;21(3):619–30. doi:10.1007/s11948-014-9565-5.
24. Gogoll J, Müller JF. Autonomous cars: in favor of a mandatory ethics setting. *Sci Eng Ethics*. 2017;23(3):681–700. doi:10.1007/s11948-016-9806-x.
25. Evans K, de Moura N, Chauvier S, Chatila R, Dogan E. Ethical decision making in autonomous vehicles: the AV ethics project. *Sci Eng Ethics*. 2020;26(6):3285–312. doi:10.1007/s11948-020-00272-8.
26. Dubljević V. Toward implementing the ADC model of moral judgment in autonomous vehicles. *Sci Eng Ethics*. 2020;26(5):2461–72. doi:10.1007/s11948-020-00242-0.
27. Kauppinen A. Who should bear the risk when self-driving vehicles crash? *J Appl Philos*. 2021;38(4):630–45. doi:10.1111/japp.12490.
28. Lundgren B. Safety requirements vs. crashing ethically: what matters most for policies on autonomous vehicles. *AI SOCIETY*. 2021;36(2):405–15. doi:10.1007/s00146-020-00964-6.
29. Ferdman A. Bowling alone in the autonomous vehicle: the ethics of well-being in the driverless car. *AI SOCIETY*. 2024;39(3):1171–83. doi:10.1007/s00146-022-01565-1.
30. Gros C, Werkhoven P, Kester L, Martens M. A methodology for ethical decision-making in automated vehicles. *AI SOCIETY*. 2025;40(8):6245–56. doi:10.1007/s00146-025-02370-2.
31. Lütge C, Poszler F, Acosta AJ, Danks D, Gottehrer G, Mihet-Popa L, et al. AI4People: ethical guidelines for the automotive sector-fundamental requirements and practical recommendations. *Int J Technoethics*. 2021;12(1):101–25. doi:10.4018/ijt.20210101.0a2.
32. Thornton SM, Pan S, Erlien SM, Gerdes JC. Incorporating ethical considerations into automated vehicle control. *IEEE Trans Intell Transp Syst*. 2017;18(6):1429–39. doi:10.1109/TITS.2016.2609339.
33. Kalra N, Paddock SM. Driving to safety: how many miles of driving would it take to demonstrate autonomous vehicle reliability? *Transp Res Part A Policy Pract*. 2016;94(4):182–93. doi:10.1016/j.tra.2016.09.010.
34. Ryan C, Murphy F, Mullins M. Spatial risk modelling of behavioural hotspots: risk-aware path planning for autonomous vehicles. *Transp Res Part A Policy Pract*. 2020;134:152–63. doi:10.1016/j.tra.2020.01.024.

35. Mustafa KA, Ornia DJ, Kober J, Alonso-Mora J. RACP: risk-aware contingency planning with multi-modal predictions. *IEEE Trans Intell Veh.* 2025;10(1):228–43. doi:10.1109/TIV.2024.3411530.
36. Reyes-Muñoz A, Guerrero-Ibáñez J. Vulnerable road users and connected autonomous vehicles interaction: a survey. *Sensors.* 2022;22(12):4614. doi:10.3390/s22124614.
37. de Clercq K, Dietrich A, Núñez Velasco JP, de Winter J, Happee R. External human-machine interfaces on automated vehicles: effects on pedestrian crossing decisions. *Hum Factors.* 2019;61(8):1353–70. doi:10.1177/0018720819836343.
38. Rouchitsas A, Alm H. External human-machine interfaces for autonomous vehicle-to-pedestrian communication: a review of empirical work. *Front Psychol.* 2019;10:2757. doi:10.3389/fpsyg.2019.02757.