

Creepy Trolley Problem: Algorithms, Inequalities, and Safety as Privilege

Mehmet A. ORHAN, Ph.D.
EM Normandie Business School
Paris Campus, France

Sep 1, 2026

TEACHING NOTE



Creepy Trolley Problem: Algorithms, Inequalities, and Safety as Privilege

Summary

This case study explores the ethical dilemma faced by Alex Hughes, the Senior Programmer at DriveSafe Systems, a leading autonomous vehicle manufacturer. Hughes is approached by a team of certified service mechanics operating in the UAE who propose modifying the company's self-driving algorithms to prioritize the safety of wealthy, famous, and influential individuals through real-time facial recognition, protecting registered clients whether they are drivers, passengers, or pedestrians on the street. The mechanics, who possess the technical capability to hack the system themselves, seek Hughes' insider collaboration to bypass DriveSafe's internal integrity checks, and further request that he assess the feasibility of scaling this service globally. This case study examines the ethical, legal, and societal implications of embedding social hierarchies into life-and-death algorithmic decisions, and challenges students to confront a world where safety itself may become a commodity.

Teaching objectives

By the end of this case study, students should be able to:

1. understand the ethical implications of integrating facial recognition and social evaluation technologies into autonomous vehicle decision-making systems.
2. analyze how algorithmic design choices can encode and amplify existing social inequalities and privileges,
3. evaluate the legal and regulatory landscape governing AI-driven safety systems across different jurisdictions, including the UAE and the European Union, and
4. propose principled responses to coercive insider collaboration requests that balance ethical integrity, corporate responsibility, and professional risk.

Keywords

AI ethics, algorithmic control, dystopia, face recognition, inequalities

Synopsis

This case is designed to help students move beyond abstract ethical theory and engage with the concrete design and commercial pressures facing engineers and managers in the AI and autonomous systems industry.

This case study revolves around Alex Hughes, the Senior Programmer at DriveSafe Systems, one of the world's fastest-growing autonomous vehicle manufacturers. Alex is approached by a team of licensed third-party mechanics, certified DriveSafe service partners based in the UAE, with an unsettling commercial proposal. They suggest modifying DriveSafe's self-driving algorithms so that vehicles continuously scan the faces of everyone in their camera field, pedestrians, cyclists, passengers in other cars, bystanders, and, upon detecting a registered premium client, reprioritize their collision-avoidance logic to protect that individual above all others. In this model, safety travels with the client's face, operating silently and invisibly across an entire fleet of vehicles. The mechanics have already demonstrated that the modification is technically feasible through the software update pathway they use for routine maintenance, but they require Alex's insider knowledge of DriveSafe's audit logs and integrity checks to execute it cleanly. They are not merely asking for permission; they are inviting Alex into a lucrative collaboration while implying that the modification will proceed with or without him. Beyond the UAE, the mechanics harbor global ambitions, asking Alex to map the regulatory and commercial feasibility of a worldwide rollout, turning the local pilot into the foundation of a global ecosystem built on algorithmic privilege. Alex must decide whether to comply, resist, expose, or find a third path, knowing that each choice carries profound professional and moral consequences.

Target Learning Group

This case is appropriate for both undergraduate and postgraduate (MBA/MSc) students. It is particularly suitable for courses in business ethics, AI strategy, technology management, public policy, and organizational behavior. It may also serve as an interdisciplinary entry point for students specializing in computer science, law, or philosophy, who wish to engage with the human and organizational dimensions of AI deployment. The case may also suit in executive education programs for professionals in the automotive, consulting, or technology sectors.

Teaching Objectives:

- Understand the ethical considerations and challenges associated with embedding facial recognition and social evaluation logic into autonomous vehicle safety systems.
- Explore how algorithmic design choices can encode social hierarchies, commodify safety, and produce discriminatory outcomes at scale.
- Examine the legal and regulatory frameworks governing AI safety systems, including the EU AI Act, ISO 26262, anti-discrimination law, and data privacy regulation across different jurisdictions.
- Analyze the dynamics of coercive insider collaboration requests and the ethical obligations of professionals with access to safety-critical systems.
- Evaluate the global scalability of ethically problematic technologies and the role of regulatory arbitrage in enabling their spread.

Discussion Questions (With Further Questions Not Included in the Case):

1. Ethical Frameworks

Apply at least two ethical frameworks (e.g., utilitarianism, deontology, virtue ethics, or contractarianism) to Alex's situation. Does each framework lead to the same conclusion? Where do they diverge, and why does that divergence matter for how we design AI systems?

2. Safety as a Commodity

The case draws an analogy between priority boarding at an airport and priority safety on the road. Is this analogy valid? Where does it hold, and where does it break down? Can safety ever be ethically commodified, and if so, under what conditions and limits?

3. Regulatory and Legal Dimensions

Would the proposed algorithm modification be legal in the UAE? In the European Union? Consider anti-discrimination law, product liability frameworks for autonomous vehicles, data protection regulation (GDPR), and the EU AI Act's classification of high-risk AI systems. How does regulatory fragmentation across jurisdictions complicate Alex's decision — and does it create an ethical obligation to act even where law is silent?

4. Algorithmic Bias and Social Hierarchy

The proposal relies on facial recognition and behavioural assessment to estimate a person's 'social value.' What are the known accuracy limitations and demographic biases of such technologies? How would those biases compound the ethical problems of the proposed system — particularly for individuals who are misidentified, unregistered, or from underrepresented groups?

5. Stakeholder Analysis

Identify all key stakeholders in this case: Alex, DriveSafe Systems, the mechanics, premium clients, pedestrians and bystanders, governments, and the general public. Map each stakeholder's interests, power, and legitimacy. Who is most vulnerable? Whose interests are least visible in the mechanics' framing of the proposal?

6. The Coercion Dimension (Further Question for Graduate/Executive Programs)

The mechanics make clear that they have the technical capability to proceed without Alex's cooperation. Does this change the nature of Alex's ethical responsibility? Does implicit coercion reduce moral culpability if he complies, or does it increase his obligation to refuse and report? How should organisations design governance systems to protect employees in Alex's position?

7. Global Scalability and Regulatory Arbitrage (Further Question for More Comprehensive, Global Discussion)

The mechanics ask Alex to assess the global feasibility of the service, identifying permissive jurisdictions, wealthy markets, and infrastructure requirements. What are the ethical implications of deliberately targeting weaker regulatory environments for a service that would be prohibited elsewhere? Does the global ambition of the proposal change your moral assessment of it?

8. What Should Alex Do?

Formulate a concrete recommendation for Alex. Should he refuse entirely, propose a counter-offer, escalate internally, blow the whistle to a regulator, or something else? Justify your answer by reference to both ethical reasoning and practical strategy. What are the risks of your recommended course of action, and how would you manage them?

Additional Activities

Divide students into groups and ask them to research one of the following real-world parallels: (a) the MIT Moral Machine experiment and its findings on cross-cultural variation in trolley-problem intuitions; (b) documented cases of biased facial recognition in law enforcement or employment contexts; or (c) existing autonomous vehicle safety regulations in three different jurisdictions of their choice. Each group should present their findings and connect them explicitly to the dilemmas faced by Alex Hughes.

Students may also be asked to draft a one-page internal memo as if they were Alex, addressed to DriveSafe's Chief Ethics Officer, setting out the nature of the approach, the risks to the company, and a recommended course of action. This exercise develops both ethical reasoning and professional communication skills.

Role-play Activity:

Assign students the following roles and ask them to engage in a structured negotiation or debate lasting 20 minutes:

- Alex Hughes as the Senior Programmer, DriveSafe Systems
- Lead Mechanic representing the UAE service partner team
- DriveSafe's Chief Legal Officer
- A representative of a civil liberties or AI ethics organisation
- A prospective premium client (ultra-high-net-worth individual)
- A pedestrian advocacy group spokesperson

Each role-player must argue from their character's position while remaining open to challenge. After the negotiation, the class should debrief: whose arguments were most persuasive, and why? What does that reveal about the power dynamics embedded in the scenario?

Guest Speaker Session:

Instructors may consider inviting a professional from one of the following fields to enrich class discussion: AI ethics and responsible AI deployment; autonomous vehicle regulation or road safety policy; data protection law (particularly in cross-border contexts); or algorithmic auditing. Students should be encouraged to prepare questions in advance, drawing on the case and their research.

Key Takeaways

- Algorithmic decision-making in safety-critical systems carries profound ethical weight that cannot be resolved by technical optimization alone.
- Facial recognition and social evaluation technologies, when embedded in autonomous systems, risk translating existing social inequalities into physical harm.
- Safety is a fundamental human right, not a commodity; proposals to commodify it through premium algorithmic protection represent a categorical ethical departure from existing norms.
- The global reach of autonomous vehicle platforms means that local ethical failures can rapidly become global ones; regulatory arbitrage is an ethical problem, not merely a legal one.

- Professionals with access to safety-critical systems carry heightened ethical obligations, particularly when approached with coercive or clandestine modification requests.
- Organizations must build governance structures, including clear whistleblowing pathways, ethics review boards, and algorithmic audit functions, to protect employees in Alex's position and prevent such proposals from progressing informally.

Notes to Instructors

This case study provides a rich opportunity for students to move beyond abstract ethical theory and engage with a scenario that is technically plausible, commercially motivated, and morally complex. Instructors should be attentive to several dynamics that commonly arise in classroom discussion.

First, some students may initially frame the proposal as a straightforward market transaction, premium service, willing buyers, and require prompting to recognize the structural asymmetry between registered clients and unregistered bystanders. The latter have no contractual relationship with DriveSafe, no awareness that they are being scanned, and no capacity to opt out. Instructors should press students on what ethical frameworks, if any, could legitimize differential protection for persons outside any commercial relationship.

Second, the coercion dimension that the mechanics can proceed without Alex tends to generate lively debate about complicity and moral responsibility under duress. Instructors should resist allowing students to use coercion as a full exculpatory argument. The question of what Alex is obliged to do, as distinct from what he is legally required to do, should be kept clearly in view.

Third, the global ambition of the proposal raises important questions about regulatory arbitrage that students in management programs may find particularly relevant. Instructors should encourage students to consider whether a strategy that is legal in one jurisdiction but prohibited in another is therefore ethical and what this implies for global technology governance.

Finally, instructors should note that this case is based on a fictional scenario, though all underlying technologies described are real and currently available. Students should be encouraged to engage with the scenario on its own terms rather than seeking to resolve it by reference to what DriveSafe 'would really do.'

Delivery Method

Before Assigning the Case (30 Mins):

Instructors should begin by introducing students to the original trolley problem through a brief reading or video (see Background Material). Students benefit from encountering the philosophical foundations before engaging with the autonomous vehicle context. A useful warm-up question is: 'If a self-driving car must choose between two unavoidable collisions, how should it be programmed to decide?' This surfaces existing intuitions without prematurely framing the ethical terrain.

For undergraduate classes, the 8-minute MIT Moral Machine video provides an accessible and visually engaging entry point. For postgraduate and executive audiences, the Nature article by Awad et al. (2018) on the Moral Machine experiment provides a richer empirical foundation. Instructors should avoid assigning detailed critical commentary on

autonomous vehicle ethics before the case discussion, as this may foreclose student reasoning prematurely.

Detailed Analyses and Solutions to Discussion Questions (60-90 mins)

Allow 10 minutes for case reading within the classroom if the case has not been pre-assigned. After reading, ask students to identify the three most important ethical issues before opening group discussion. Ensure that students correctly identify all layers of the dilemma: the facial recognition premium model, the coercive invitation to collaborate, and the global expansion ambition. Cases where students conflate these layers should be corrected early, as each raises distinct ethical and legal questions.

The following structure is recommended for a 90-minute session:

- Opening (10 min): Show-of-hands intuition poll. 'Would you pull the lever?' then 'What if pulling the lever protected only passengers who had subscribed to a premium safety plan?' Use the gap between responses to open discussion.
- Small Group Work (20 min): Assign Discussion Questions 1–5 to separate groups. Each group appoints a spokesperson.
- Plenary Debrief (35 min): Groups present. Instructor draws out the tension between utilitarian and Rawlsian answers (Q1 and Q5), challenges the airport analogy (Q2), surfaces jurisdictional complexity (Q3), and raises the coercion question (Q6).
- Decision Point (20 min): Students individually draft a 200-word recommendation memo as Alex. Two or three volunteers share their recommendations. Instructor closes by noting that both action and inaction carry moral weight — and that this is the defining feature of a genuine ethical dilemma.
- Wrap-Up (5 min): Connect to live regulatory debates, EU AI Act, ISO 26262, UAE AV pilot programs. Assign a 1,000-word reflective essay for assessment (Optional).

Discussion Questions

1. Ethical Frameworks

Apply at least two ethical frameworks (e.g., utilitarianism, deontology, virtue ethics, or contractarianism) to Alex's situation. Does each framework lead to the same conclusion? Where do they diverge, and why does that divergence matter?

2. Safety as a Commodity

The case draws an analogy between priority boarding and priority safety. Is this analogy valid? Where does it hold, and where does it break down? Can safety ever be ethically commodified, and if so, under what conditions?

3. Regulatory and Legal Dimensions

Would the proposed algorithm modification be legal in the UAE? In the European Union? Consider anti-discrimination law, product liability law (particularly as applied to autonomous vehicles), and relevant AI regulation such as the EU AI Act. How does regulatory fragmentation across jurisdictions complicate Alex's decision?

4. Algorithmic Bias and Social Hierarchy

The case states that facial recognition and behavioral assessment technologies already exist and could be used to estimate a person's "social value." What are the known accuracy limitations and demographic biases of such technologies? How would those biases compound the ethical problems of the proposed system?

5. Stakeholder Analysis

Identify all key stakeholders in this case: Alex, DriveSafe Systems, the mechanics, potential premium clients, pedestrians, governments, and the public. Map each stakeholder's interests, power, and legitimacy. Who is most vulnerable, and whose interests are most likely to be ignored?

6. What Should Alex Do?

Formulate a concrete recommendation for Alex. Should he refuse entirely, propose a counter-offer, escalate internally, or report to a regulator? Justify your answer by reference to both ethical reasoning and practical strategy. What are the risks of your recommended course of action, and how would you manage them?

Answers to Discussion Questions:

Background: The Core Ethical Structure

Before addressing individual questions, instructors should ensure students understand the ethical architecture of the case. Three distinct but related dilemmas are present: (1) whether it is ethical to embed social hierarchy into life-and-death algorithmic decisions at all; (2) whether Alex's insider knowledge gives him special obligations that an ordinary bystander would not have; and (3) whether a service that is commercially attractive in a permissive jurisdiction becomes ethical by virtue of its legality there. Students who conflate these three questions will produce weaker analyses.

1. Ethical Frameworks

A utilitarian analysis initially appears to support the mechanics' proposal: if premium clients are statistically rare, and if the system protects them with high accuracy, total harm may be only marginally increased. Instructors should press students on this reasoning. The utilitarian calculus is complicated by: (a) the accuracy limitations of facial recognition, which introduce unpredictable misclassification; (b) the systemic effect of normalizing identity-based safety differentials, which may produce broader social harms; and (c) the interests of unregistered bystanders who are not parties to any agreement and whose exposure to increased risk is entirely non-consensual.

A deontological analysis, drawing on Kant's categorical imperative, reaches a clear prohibition: the proposal treats unregistered individuals merely as means to the commercial ends of registered clients, violating the fundamental principle of equal human dignity. Students should be encouraged to articulate why the identity of the person matters morally when the physical consequences of a collision are independent of social status.

A Rawlsian analysis, applying the veil of ignorance, is particularly powerful here. Behind the veil, no rational agent would choose a society where their physical safety in a public space depends on their wealth or fame, since they might find themselves on the unprotected side of that distinction. This framework connects the individual dilemma to the broader question of what kind of technological society we are building.

2. Safety as a Commodity

The airport analogy is useful precisely because it reveals where the ethical line lies. Priority boarding is a convenience, not a right; missing a boarding queue causes inconvenience, not injury. The proposed safety premium, by contrast, does not merely accelerate a benefit; it transfers physical risk from one person to another. The correct analogy is not priority boarding but a lifeboat with limited capacity, where the allocation of places is determined by wealth. Most

students will recognize that this crosses a categorical ethical threshold, even those who are broadly comfortable with market-based differentiation.

Instructors should push students to articulate where the line between legitimate service differentiation and ethically impermissible risk transfer lies. This produces useful discussion about the distinction between positional goods (goods whose value depends on others not having them) and fundamental rights.

3. Regulatory and Legal Dimensions

The EU AI Act (2024) classifies safety-critical autonomous systems as high-risk, requiring conformity assessment, transparency obligations, and prohibition of biometric identification systems used for real-time public surveillance in certain contexts. The proposed modification would likely violate multiple provisions, including those prohibiting AI systems that exploit personal characteristics to produce discriminatory outcomes. ISO 26262, the functional safety standard for road vehicles, explicitly prioritizes minimizing total harm and does not permit identity-based weighting of safety decisions.

In the UAE, the regulatory environment for autonomous vehicles is more permissive but rapidly evolving. Dubai's Roads and Transport Authority has issued AV framework guidelines that do not currently address identity-based safety prioritization, creating a genuine regulatory gap. Instructors should note that the absence of explicit prohibition is not the same as permission, and that international human rights norms, including the right to life and the prohibition of discrimination, apply regardless of domestic regulatory silence.

GDPR analysis is also relevant: the biometric data of pedestrians and bystanders would be processed without consent, almost certainly in violation of Article 9 GDPR if any European data subjects are involved. This extraterritorial reach is an important teaching point for students who may assume that non-EU deployments fall outside European law.

4. Algorithmic Bias and Social Hierarchy

Facial recognition systems, even the most advanced commercially available systems, display significantly higher error rates for individuals with darker skin tones, women, and older adults (Buolamwini & Gebru, 2018). In a safety-critical context, these inaccuracies may mean that a registered client is misidentified and unprotected, or that an unregistered individual is falsely identified as a client and receives protection at the expense of others. The system's claimed purpose is therefore undermined by the technical limitations of the technology on which it depends.

Beyond accuracy, the very concept of estimating 'social value' from attire, gait, or demographic attributes encodes existing social biases into physical safety outcomes. Students should be encouraged to identify whose definitions of social value are embedded in such a system, and whose are not.

5. Stakeholder Analysis

The most frequently overlooked stakeholder group is pedestrians and bystanders, individuals who have no relationship with DriveSafe, no awareness of the system, and no recourse if harmed by its decisions. Their interests are entirely absent from the mechanics' commercial framing. A robust stakeholder analysis should also distinguish between the interests of DriveSafe as a corporation (reputational and liability risk), DriveSafe's shareholders (commercial risk), and the broader public (safety and trust in autonomous systems). These interests do not necessarily align, and students should be asked to map where conflicts arise.

6. The Coercion Dimension

This question tends to generate the most intense classroom debate. The key distinction instructors should maintain is between moral culpability (which may be reduced by coercion) and moral obligation (which is not). Even if Alex's culpability for the modification is reduced if it proceeds without him, he retains an obligation to act, to escalate internally, to document the approach, or to report to a relevant authority. Students who argue that coercion relieves Alex of all responsibility should be pressed on whether the same logic would apply in other professional contexts: a bank employee coerced into facilitating fraud, for example.

7. Global Scalability and Regulatory Arbitrage

The global expansion request escalates the moral stakes considerably. Regulatory arbitrage, deliberately deploying a technology in jurisdictions where it would be prohibited elsewhere, is increasingly recognized as an ethical failure in its own right, not merely a legal strategy. Instructors should ask students whether the mechanics' global mapping request, which essentially asks Alex to identify the weakest regulatory environments, is compatible with any credible corporate ethics framework.

8. What Should Alex Do?

There is no single correct answer, but students should be discouraged from treating refusal as a costless default. Refusing without escalating leaves the threat in place; the mechanics retain their capability and may find another collaborator. A complete response to the dilemma requires Alex to: (a) refuse to collaborate; (b) document the approach in writing immediately; (c) escalate to DriveSafe's legal, ethics, and executive leadership; and (d) consider whether external disclosure, to a regulator, to DriveSafe's board, or to the public, is warranted given the severity of the risk. Students should also consider Alex's personal legal exposure: knowingly possessing information about a plan to hack a safety-critical system without disclosing it may itself carry liability in some jurisdictions.

Experience of Using This Case in Teaching

In early implementations of this case, a significant number of students initially focused on the business opportunity presented by the mechanics' proposal, framing the dilemma primarily as a question of risk management rather than ethics. This pattern suggests that instructors should invest time at the outset in establishing that the case is not asking students to optimize a commercial decision, but to evaluate whether a commercial decision should be made at all.

The coercion dimension, the revelation that the mechanics can proceed without Alex, has consistently generated the most animated discussion, and instructors should be prepared to spend additional time here. Students sometimes use coercion as a complete exculpatory argument, and gentle but firm challenge is needed to redirect attention toward Alex's positive obligations.

Overall, this case has been received enthusiastically by students who appreciate its combination of philosophical depth, technical grounding, and narrative tension. The global expansion subplot has proven particularly productive for students interested in international business and regulatory strategy, surfacing questions about the ethics of market entry that generalise well beyond the autonomous vehicle context.

United Nations Sustainable Development Goal 10: Reducing inequalities

This case study directly engages SDG 10, which aims to reduce inequality within and among countries. The proposed algorithmic system would institutionalize differential physical safety based on wealth and social status, a concrete mechanism for deepening inequality in one of

the most fundamental human interests: the right to life and bodily integrity. Students should be encouraged to connect the case to SDG 10's call for universal access to basic protections, regardless of economic or social status, and to ask what it means for a society when life-protecting systems are subject to market logic.

What Happened Next

The scenario described in this case is fictional, but the underlying technologies and commercial pressures it depicts are real. Several developments in the autonomous vehicle industry are directly relevant to the dilemmas Alex faces.

Several major autonomous vehicle developers, including Waymo and Cruise, have published public safety commitments that explicitly reject identity-based decision-making, affirming that their systems treat all road users equally regardless of personal characteristics. These commitments reflect both ethical conviction and awareness of the catastrophic reputational and legal consequences of the alternative. The fictional DriveSafe's situation, in which an insider is invited to undermine these commitments covertly, illustrates the governance gap between public commitment and private temptation that all technology companies must actively manage.

In 2023, the National Highway Traffic Safety Administration (NHTSA) in the United States opened investigations into multiple autonomous vehicle incidents, signalling increasing regulatory appetite for scrutiny of AI-driven safety decisions. The direction of travel in major markets is unmistakably toward greater transparency, accountability, and prohibition of discriminatory algorithmic logic, making the mechanics' proposal not only currently unethical but almost certainly commercially unviable in the medium term.

Conclusion

Alex Hughes's dilemma exemplifies one of the most consequential questions of the algorithmic age: when machines make decisions that determine who is protected and who is endangered, what values do we embed in their code, and who decides? The trolley problem, long a philosophical thought experiment, is now a live engineering challenge with real commercial pressures, real regulatory gaps, and real people in the path of real vehicles. This case does not offer a comfortable resolution. Both action and inaction carry moral weight, and the ethical frameworks that guide us do not always converge. What the case does offer is a set of tools for reasoning carefully about a world in which safety can be priced, faces can be ranked, and the consequences of algorithmic design extend far beyond the engineer who writes the code. Students who engage seriously with these questions will be better equipped, as professionals, as citizens, and as human beings, to navigate the ethical terrain of an increasingly autonomous world.

Background Material

Awad, E., Dsouza, S., Kim, R., Schulz, J., Henrich, J., Shariff, A., Bonnefon, J.-F., & Rahwan, I. (2018). The Moral Machine experiment. *Nature*, 563(7729), 59–64.

<https://doi.org/10.1038/s41586-018-0637-6>

European Parliament. (2024). *EU AI Act: First regulation on artificial intelligence*. European Parliament News.

<https://www.europarl.europa.eu/news/en/headlines/society/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence>

The Moral Machine Game. MIT <https://www.moralmachine.net/>

Thomson, J. J. (1985). The trolley problem. *Yale Law Journal*, 94(6), 1395–1415.

<https://doi.org/10.2307/796133>

Further Reading

Barocas, S., Hardt, M., & Narayanan, A. (2023). *Fairness and machine learning: Limitations and opportunities*. MIT Press. <https://fairmlbook.org>

Buolamwini, J., & Gebru, T. (2018). Gender shades: Intersectional accuracy disparities in commercial gender classification. *Proceedings of Machine Learning Research*, 81, 1–15.

<https://proceedings.mlr.press/v81/buolamwini18a.html>

ISO 26262. (2018). *Road vehicles, Functional safety*. International Organization for Standardization.

Noble, S. U. (2018). *Algorithms of oppression: How search engines reinforce racism*. New York University Press.

Rawls, J. (1971). *A theory of justice*. Harvard University Press.

Waymo. (2023). *Waymo safety report*. <https://waymo.com/safety/>