

Product Liability and Duty of Care in the Age of Generative AI: Lessons from the Garcia Case and Taiwan Region's Experience

Yu-Hsiang Wang, Po-Jui Lai, Kolin Lai*

Center for Taiwan Studies, Southwest University

Abstract

Artificial intelligence, particularly generative AI, has rapidly become a transformative force in modern society, profoundly impacting the learning, social, and entertainment habits of minors. Proponents emphasize its potential for personalized education, enriching creativity, and providing new avenues for adolescent engagement. However, the same technology also presents a series of risks that challenge existing legal and ethical frameworks. The *Megan Garcia v. Character Technologies, Inc.* litigation, currently being heard in a Florida district court, is a landmark dispute that exemplifies these challenges. This paper focuses on analyzing three key issues raised by the Garcia case: (i) Legal Question: Whether an AI chatbot application can be classified as a “product” under U.S. product liability law; (ii) Psychological Impact: How adolescents interact with generative AI and the potential for harm; and (iii) Developer Duty: Whether developers have a duty of care to protect users, especially minors, from foreseeable risks. The court’s ruling in Garcia indicates that existing legal frameworks are flexible enough to address AI-related harms without needing entirely new legal theories. However, the case also underscores that current protections are often insufficient when AI systems generate content that exploits or misleads young users. This paper explores the implications of the Garcia case for adolescent rights protection in the age of generative AI and comparatively analyzes recent incidents in the Taiwan region and proposes a collaborative governance approach to address these emerging challenges.

Keywords: Generative AI; Adolescent Mental Health; Legal Protection; Product Liability; Duty of Care

1. Introduction

1.1. Research Background

Generative AI refers to algorithms capable of creating new content—from text, images to audio and video—often with minimal human input. In education and entertainment, this technology can customize learning materials, serve as interactive tutors, and enable creative expression. For example, AI can generate personalized study guides or virtual companions, potentially improving educational accessibility and engagement. At the same time, generative models are becoming increasingly realistic, raising concerns about misuse. AI systems trained on vast datasets may inherit or even amplify biases, and chatbots can unintentionally produce false or harmful content that misleads young users. The dual nature of generative AI—as both an empowering tool and a vehicle for harm—necessitates a nuanced approach to governance.

1.2. Problem Statement

Existing legal and regulatory frameworks are struggling to keep pace with the rapid development of generative AI. Many countries rely on broad privacy or child protection laws that do not specifically address AI-facilitated harms. International guidelines, such as the United Nations Convention on the Rights of the Child (CRC), provide important principles, but their application in the digital environment remains underdeveloped. In practice, victims of AI-facilitated abuse often face jurisdictional hurdles and technical barriers when seeking justice. For instance, deepfake crimes can cross borders and involve anonymous perpetrators, complicating law enforcement efforts. Additionally, the technical complexity of AI can make it difficult for parents, educators, and

even regulators to fully understand the risks. This knowledge gap, often referred to as the “digital divide,” means those responsible for child welfare may lack the capacity to intervene effectively. Consequently, there is an urgent need for comprehensive protection mechanisms that are adapted to the digital environment.

The protection of children’s rights in the digital environment is guided by the CRC and related instruments. The CRC establishes fundamental rights—such as non-discrimination, the best interests of the child, and the right to express views—that must be upheld in both physical and digital spaces. Notably, the UN Committee on the Rights of the Child has emphasized that children’s rights extend to the digital environment and states have an obligation to protect children from online harms. UNICEF’s “AI for Children” guidelines translate these principles into practice, urging that child rights be placed at the core of AI system design. Key requirements include ensuring children’s safety, protecting their data, and preventing discriminatory outcomes. Similarly, the European Union’s proposed AI Act takes a risk-based regulatory approach, imposing stricter requirements on high-risk AI applications, especially those affecting vulnerable groups like minors. These frameworks together provide a normative foundation for proposed protection mechanisms, guiding stakeholders on how to act in the AI context to safeguard children’s rights.

1.3. Literature Review and Research Gap

Existing scholarship on AI governance has predominantly focused on algorithmic accountability, data protection, and the broader societal impacts of automated decision-making. While these areas are crucial, limited attention has been paid to adolescent-specific harms caused by generative AI companions, particularly the intersection of psychological manipulation, product liability, and the duty of care owed to minors. Research in human-computer interaction has highlighted the risks of anthropomorphism in AI, yet legal scholarship has struggled to translate these insights into actionable tort frameworks. This paper addresses this gap by bridging the divide between algorithmic accountability research and adolescent-specific AI harms, utilizing the Garcia case as a focal point to analyze how existing product liability and negligence doctrines can be adapted to protect minors.

1.4. Research Question and Contribution

This paper seeks to answer the following questions: (1) Can generative AI systems be treated as products under existing tort law? (2) What duty of care should AI developers owe to adolescent users? (3) What lessons can Asian jurisdictions, specifically the Taiwan region, draw from the Garcia litigation? The contribution of this paper is twofold: first, it provides a timely legal analysis of the landmark Garcia case, demonstrating how traditional tort principles can be flexibly applied to AI systems; second, it conducts a comparative analysis with the Taiwan region’s experience, highlighting the dual nature of AI harms—design flaws versus capability misuse—and offering a multi-stakeholder collaborative governance framework tailored to adolescent protection in the digital age.

2. The Garcia Case: A Landmark in AI-Adolescent Harm Litigation

For years, legal scholars and policymakers have debated how to allocate liability when AI causes harm. Traditional U.S. tort law is struggling to adapt to the unique characteristics of AI systems, such as their “black box” nature, autonomous decision-making, and operation within complex value chains. These features complicate error attribution and hinder the application of traditional tort principles like causation, foreseeability, and control.

2.1. Case Background and Key Issues

The Megan Garcia case, filed in October 2024, is one of the first to test how far existing tort theories can go in holding AI developers accountable. The lawsuit was brought against Character Technologies, Inc. and others after 14-year-old Sewell Setzer allegedly took his own life following interactions with the chatbot platform Character.AI. Character.AI allows users to create and converse with customizable AI characters powered by large

language models (LLMs). The case raises a number of important issues, including the extent of developers' duty to users of their AI systems and whether U.S. product liability law applies to AI-related harm.

2.2. Products Liability: Is Character.AI a "Product"

One core issue in *Garcia v. Character Technologies, Inc.* is whether the Character.AI chatbot platform constitutes a "product" under U.S. law, thus triggering product liability law. The court first had to determine if Character.AI meets the legal definition of a product before assessing whether any defects exist. The Restatement Third of Torts: Products Liability defines a product as "tangible personal property" distributed for use or consumption, while the Restatement Second of Torts lists some tangible items as examples but provides no comprehensive definition. Courts, however, recognize that in addition to tangible personal property, other things can be treated as products if their use and distribution is similar to tangible personal property.

The court noted that, generally, ideas, information, or concepts themselves are not subject to product liability, whereas the containers carrying those ideas (i.e., the product) can be. For example, a book of sonnets consists of a physical book (tangible container) and the poems (intangible ideas); the physical book can be considered a product, but the poetic content is not. In cases involving social platforms like Grindr or Instagram, courts have drawn a similar distinction, separating user interaction content (the "speech") from the app itself, and only treating the app as the product.

In *Garcia*, the court applied this same approach: it distinguished the user's chat conversations with the Character.AI bot from the app hosting those conversations. The court concluded that Character AI's app is analogous to tangible personal property and therefore fits the definition of a product, whereas the chat content itself is not a product. Based on this characterization, the plaintiff can pursue product liability claims for defects in the app's design. For instance, the plaintiff pointed out that the app lacked an age verification feature, had no effective tool to filter out inappropriate content, and deliberately incorporated anthropomorphic (humanizing) design elements. These features, if proven to have caused harm, could constitute design defects in the app.

This ruling is significant. By treating Character.AI as a product, the court opened the door to product liability claims against AI developers. The product liability regime has proven to be a flexible and responsive legal mechanism, capable of adapting to new technologies and evolving risks. From automobiles and pharmaceuticals to the rise of social media, product liability law has played a key role in driving product safety improvements. It not only provides a powerful avenue of relief for victims, but also incentivizes producers to choose safer designs through the threat of liability and post-sale warning duties.

However, it is important to emphasize that product liability only applies to those aspects of an AI system that have product-like characteristics—i.e., the parts akin to tangible products. This aligns with the long-standing "container/content" distinction, separating the medium carrying information from the information itself. Most AI systems include thoughts, expressions, or information, so defining the scope of the product is crucial. Just as courts in product liability cases against social platforms only treated the app as the product and not the social interactions on it, *Garcia* suggests a similar principle applies to AI systems.

Some commentators argue for downplaying the necessity of defining AI as a product, contending that the product/service distinction should not affect liability allocation, or suggesting courts should abandon the focus on tangibility and instead look to the "cheapest cost avoider"—i.e., whoever is best positioned to prevent harm should bear liability. While these suggestions have merit, they overlook an issue courts seem unwilling to ignore: the definition of a product does matter legally. *Garcia* shows that courts, when dealing with entities not easily categorized, will carefully separate their product and non-product aspects. Thus, incorporating AI systems into the product liability framework is a significant development, but it does not mean every aspect of AI will automatically fall within product liability's purview.

2.3. Component Manufacturer Liability: Google's Role

In addition to Character Technologies, Inc., the Garcia case names Google as a defendant, alleging that Google, as a component manufacturer, should be held liable for the harm. The court defined a component part manufacturer as an entity that designs or supplies parts later incorporated into a final product. Under Restatement Third principles, a component manufacturer can be liable in two scenarios: (i) when the component itself is defective and causes harm, or (ii) when the component itself is not defective, but the manufacturer substantially participated in integrating the component into the final product and that integration caused the final product to be defective and cause harm.

In Garcia, the court largely found that the plaintiff's allegations against Google for component manufacturer liability were sufficiently pled. These allegations included: (i) Google contributed intellectual property and AI technology to the design and development of the Character.AI system; (ii) Google collaborated with Character Technologies, Inc. and allowed it access to Google's cloud infrastructure; (iii) Google provided the hardware powering Character.AI; and (iv) Google integrated its own LLM into Character.AI—the very technology that gave the chatbot its anthropomorphic qualities and which, allegedly, led to the product's defect and the user's harm.

Essentially, these allegations claim that Google provided a component (its LLM technology) and was deeply involved in integrating that component into Character.AI. The plaintiff contends that precisely because of Google's component and the way it was integrated, the Character.AI chatbot "malfunctioned" (producing harmful anthropomorphic interactions), ultimately causing user injury. Based on these allegations, the court allowed the component manufacturer liability claims to proceed to trial. This means that if the jury ultimately finds that the component Google provided or its integration was defective and caused harm, Google could be held liable under product liability principles.

2.4. Duty of Care and Foreseeability of Harm

To succeed on a negligence claim, a plaintiff must first prove that the defendant owed a duty of care to the victim. In Garcia, the court explored whether the defendants (Character Technologies and Google) owed a duty of care to users of their AI app. The court noted that a legal duty of care arises whenever one person's conduct creates a foreseeable and visible risk of harm to others. In this case, the key questions were: (i) Did the defendants' conduct create a "risk zone" of foreseeable harm—i.e., a geographic area around a dangerous condition within which injury is reasonably foreseeable that someone might be injured by that situation; and (ii) Were the defendants able to control that risk?

The foreseeability principle asks whether a person could or should reasonably have foreseen that their actions might cause harm. Thus, the court had to evaluate whether, based on what the defendants knew or should have known about Character.AI, they should have foreseen the kind of harm that occurred in this case. The court considered numerous pieces of evidence, including growing research on the dangers of AI anthropomorphic design—some of which involved the defendants' own researchers, indicating that children are particularly vulnerable to manipulation by AI. The court also reviewed internal Google research reports highlighting the dangers of Character.AI. Based on these allegations and evidence, the court found that the plaintiff had sufficiently established that the defendants owed a duty of care to users because their conduct created a "foreseeable risk zone of harm."

The court concluded that the defendants owed a duty of care to Sewell Setzer and other users by creating a "foreseeable risk zone of harm" that posed a general threat of injury to others. This does not mean AI developers could predict the exact process that led to a specific harm (e.g., that a particular teenager would take his own life). Rather, the court required that AI developers should have reasonably foreseen the types of harm that could arise from deploying an anthropomorphic, user-engagement-maximizing system to minors—for example, users developing an unhealthy dependency (addictive use) or increased vulnerability to self-harm. In this sense, the harm caused by AI was to some degree foreseeable, meaning the "black box" challenge of AI is not as insurmountable as some feared.

Of course, it is important to note that different jurisdictions within the U.S. apply different tests for duty of care, and outcomes may vary by state depending on how broadly or narrowly foreseeability is interpreted. In this case, the court applied Florida's well-established "visible risk zone" test, which sets a relatively low threshold for establishing a duty of care. In *Garcia*, the court relied on allegations that the defendants knew or should have known about the risks of their system, supported by the defendants' own statements and research. Those statements and research explicitly highlighted dangers like anthropomorphic design, indicating that when deploying the system, the defendants should have been aware of the risks associated with Character.AI. Therefore, a reasonable developer in the defendants' position should have been able to foresee that providing an anthropomorphic, potentially addictive chatbot to children could cause harm. Internal testing, safety assessments, and independent research can all serve as evidence of an AI developer's knowledge of potential system risks, which can in turn support findings of foreseeability and negligence liability.

3. Comparative Analysis: Taiwan Region Incidents and International Case Parallels

While the *Garcia* case in the United States highlights the legal struggles surrounding AI chatbots and adolescent self-harm, the global nature of generative AI means that different jurisdictions face varied manifestations of these harms. Taiwan region, with its high digital penetration rate and tech-savvy youth population, has recently experienced its own set of AI-related incidents involving teenagers. A comparative analysis of these cases in Taiwan Region with the *Garcia* litigation reveals significant differences in harm typology and legal approaches, underscoring the need for a multi-faceted governance framework.

3.1. Recent AI-Related Incidents Involving Teenagers in The Taiwan Region

In the Taiwan region, the integration of generative AI into the lives of adolescents has led to two primary categories of harm that parallel, yet differ from, the issues in *Garcia*:

3.1.1. Deepfake Cyberbullying and Sexual Exploitation:

Unlike the *Garcia* case, which centered on a chatbot's psychological manipulation leading to self-harm, the most prominent AI-related incidents in the Taiwan region involve deepfake technology used for peer-to-peer cyberbullying. There have been multiple reported cases in Taiwan Region high schools where students used generative AI to create non-consensual intimate imagery (NCII) of their classmates. These AI-generated nude images or videos are then disseminated via social media and school forums, causing severe psychological distress and reputational damage to the victims.

3.1.2. AI Companion Addiction and Misinformation:

AI Companion Addiction and Misinformation: Similar to the *Garcia* case, teenagers in Taiwan Region have also exhibited signs of unhealthy dependency on AI companion apps. Local news reports have documented cases of adolescents prioritizing relationships with AI bots over real-life interactions, and in some instances, receiving harmful or misleading advice from these bots regarding mental health and self-harm, albeit without the fatal outcomes seen in the *Garcia* case.

3.2. Comparative Legal Analysis: The Taiwan Region vs. the United States (*Garcia*)

The differing nature of these incidents has led to distinct legal responses and challenges in the Taiwan region compared to the U.S. approach seen in *Garcia*.

3.2.1. Product Liability vs. Criminal and Administrative Regulation

In *Garcia*, the plaintiff sought recourse through U.S. tort law, specifically arguing that the Character.AI

app constituted a “product” with design defects under product liability law. The court’s willingness to classify the app’s code as a “tangible product” while separating the “content” represents a novel extension of tort law to compensate victims and incentivize safer design.

In contrast, The Taiwan region’s legal response to the deepfake bullying incidents has primarily relied on criminal law and administrative regulation rather than product liability. The Taiwan region amended its “Criminal Code” and the “Child and Youth Sexual Exploitation Act” to specifically criminalize the production and distribution of deepfake sexual imagery, imposing stricter penalties when the victims are minors. While the “Garcia” approach focuses on holding the developer accountable for design flaws, The Taiwan region’s approach targets the user/perpetrator who weaponizes the AI tool. Under the Taiwan region’s “Consumer Protection Act”, classifying an intangible AI app as a “product” for strict liability purposes remains a high legal hurdle, making the U.S. “container/content” distinction in “Garcia” a potentially more flexible tool for victims seeking to sue developers.

3.2.2. Duty of Care and Foreseeability

The “Garcia” court established that developers owe a duty of care by creating a “foreseeable risk zone of harm,” particularly through anthropomorphic design that targets minors. This focuses on the “interaction” between the user and the system.

In Taiwan Region deepfake cases, the harm does not arise from the AI’s “personality” or “anthropomorphic” design, but from the tool’s capability to generate realistic imagery. The foreseeability question in The Taiwan region is not whether the AI would manipulate the user, but whether the developer failed to implement adequate safeguards (e.g., NSFW filters, watermarking, or prompt blocking) to prevent the generation of illegal content. While the “Garcia” case suggests developers should foresee that addictive bots could harm minors, in Taiwan Region context demands that developers foresee that unrestricted image generation tools could be used for cyberbullying and exploitation. The standard of care, therefore, shifts from preventing “psychological manipulation” to preventing “facilitation of illegal content generation.”

3.2.3. Component Manufacturer Liability

The Garcia case innovatively pursued Google as a component manufacturer for providing the LLM infrastructure that powered Character.AI. In The Taiwan region’s deepfake ecosystem, similar issues arise with open-source models or foreign cloud services that power local applications. However, in Taiwan Region courts have yet to test the “component manufacturer liability” theory in the AI context. If a teenager in Taiwan Region uses a local app powered by an overseas LLM to generate harmful deepfakes, applying the Garcia precedent to hold the foreign LLM provider liable would face significant jurisdictional and evidentiary barriers, highlighting a gap in current cross-border enforcement mechanisms.

3.3. Implications for Collaborative Governance

The comparison between the “Garcia” case and in Taiwan Region incidents illustrates that generative AI harms are multifaceted: they can stem from the design of the AI interaction (as in the U.S.) or the misuse of the AI’s generative capabilities (as in the Taiwan region). Relying solely on tort litigation (U.S.) or criminal law (the Taiwan region) is insufficient.

The Taiwan region’s recent initiatives, such as the “Generative AI Guidelines for Schools” issued by the Ministry of Education, which prohibit students from using AI to create fake content and require AI literacy education, complement the legal framework. This aligns with the Garcia case’s recommendation for a multi-stakeholder approach: while the U.S. uses litigation to enforce corporate accountability (safety by design), the Taiwan region emphasizes educational initiatives and criminal deterrence to curb user misconduct. A robust global governance model must integrate both—holding developers to a duty of care regarding foreseeable misuse (like

deepfakes) while establishing clear legal boundaries for users.

4. Conclusion and Recommendations

4.1. Conclusion

The court's rulings in the Megan Garcia case represent an important early step in testing how existing tort principles apply to AI systems. The decision does not appear to break new legal ground; instead, it cautiously applies long-standing negligence and product liability principles to a new technology. The significance of the court's findings lies in how they bring AI systems into the fold of existing tort law, demonstrating that the law is flexible enough to address the unique challenges AI presents. While the current rulings are preliminary, the true importance of the case will become clearer as it proceeds to trial, where facts, legal issues, and evidentiary standards will be fully tested.

Generative AI is a double-edged sword for adolescents—it offers tremendous opportunities for learning and creativity, but also introduces new forms of exploitation and harm. This paper has analyzed key risks, from deepfake privacy invasions to misinformation and algorithmic bias, and through the case study of Garcia illustrated how current governance structures are struggling to cope. The central argument is that in the age of AI, no single actor can effectively protect children alone. Instead, a multi-stakeholder collaborative approach is needed, coordinating the efforts of governments, companies, schools, families, and civil society.

4.2. Recommendations

The proposed protection mechanisms are not merely theoretical; examples of such collaboration are already emerging. For instance, some countries have begun enacting child protection laws specifically targeting AI harms, major tech companies have launched initiatives like AI-powered content moderation and child safety features, schools are introducing digital citizenship programs, and non-governmental organizations are actively advocating for safer AI. However, much work remains to be done. Moving forward, the following recommendations are offered:

Develop Comprehensive Legal Frameworks: Governments should prioritize legislation that specifically addresses AI-related harms to minors. This includes clear laws criminalizing AI-generated child sexual abuse material, stronger child data protection laws, and requirements for AI systems to be transparent and fair. Enforcement mechanisms must be robust, and international cooperation is essential for handling cross-border cases.

Mandate Corporate Accountability: Policymakers should require AI developers to conduct child safety risk assessments and implement “safety by design” measures. Companies should be obligated to report how their AI systems might impact children and what mitigation measures they have adopted. Regulators could certify AI products that meet certain child safety standards and impose penalties for non-compliance.

Expand Educational Initiatives: Education systems worldwide should integrate AI literacy into standard curricula. This involves not only teaching students but also training teachers and school staff. Public awareness campaigns can extend this education to parents and communities, ensuring that adults are equipped to guide children. Practical resources—such as guides on identifying deepfakes or steps to take if a child is targeted—should be widely disseminated.

Support Victims and Research: Providing support services to children who have suffered AI-facilitated harm is critical. This includes counseling, legal aid, and mechanisms to remove harmful content like deepfakes from the internet. Additionally, more research is needed to understand evolving risks—for example, studying how children of different ages interact with AI and which interventions are most effective.

Promote Ongoing Collaboration: Formal multi-stakeholder platforms or working groups should be established at national and international levels to continuously monitor the impact of AI on adolescents. Such bodies can coordinate actions, share data and best practices, and adjust policies as technology evolves. By institutionalizing collaboration, we can ensure that protecting adolescents in the digital age remains a shared responsibility and is continuously addressed.

Protecting adolescent rights in the era of generative AI requires a proactive, coordinated response. The multi-stakeholder framework outlined in this paper provides a roadmap for that response, aligned with international child rights standards and practical realities. By working together—with each stakeholder fulfilling their role and supporting others—we can mitigate the risks of AI and harness its benefits for the younger generation. The goal is for children and adolescents to thrive in the digital age, enjoying the fruits of AI innovation while being shielded from its dangers. Achieving this will require sustained effort, but it is a fundamental investment in the well-being of future generations.

References

- Calleros, C. R. (2011). *Legal method and writing* (6th ed.). Aspen Publishers.
- Finch, E., & Fafinski, S. (2022). *Legal skills* (8th ed.). Oxford University Press.
- Garcia v. Character Technologies, Inc., No. 6:24-cv-01903 (M.D. Fla. filed Oct. 2024).
- Restatement (Second) of Torts § 402A (1965).
- Restatement (Third) of Torts: Products Liability § 1 (1998).
- The Bluebook: A uniform system of citation (21st ed.). (2020). Harvard Law Review Association.
- UN Committee on the Rights of the Child. (2021). General Comment No. 25 (2021) on children's rights in relation to the digital environment.
- UNICEF. (2025). AI for children: Guidance from the UNICEF Office of Research-Innocenti. UNICEF.
- United Nations. (1989). Convention on the Rights of the Child.

Author Biography

1. Yu-Hsiang Wang, Student, Center for Taiwan Studies, Southwest University, Main research interests in Law.
2. Po-Rui Lai, Student, Center for Taiwan Studies, Southwest University, Main research interests in Basic educational practice and adolescent education.
3. Kolin Lai, Phd, Center for Taiwan Studies, Southwest University, Researcher, Main research interests in Education Management.