

Course:	SKOM12
Term:	Spring 2026
Supervisor	Howard Nothhaft
Examiner	Nils Holmberg

Constructing “AI for Minors”: Corporate Framing and Public Discourse in China and the United States

RUIHE MA & SHURONG WANG

Lund University
Department of communication
Master's thesis



Abstract

This study examines how AI for minors is framed by AI companies and constructed in public discourse in China and the United States. Using Rhetorical Arena Theory and framing theory, the study analyses corporate communication from four AI platforms (ChatGPT, Gemini, Grok, and Doubao) together with public discourse materials, including mainstream media, NGO reports, and online commentary. Reflexive thematic analysis identified five themes in corporate framing and six in public discourse.

The findings show that companies frame AI for minors as a conditional form of access, beneficial but only under specific conditions of parental involvement, safeguards, and responsible use. U.S. companies emphasise developmental empowerment, while Doubao frames protection through functional restriction and regulatory compliance. Public discourse in both contexts challenges corporate framings, but with different intensity: U.S. criticism is systemic and structural, while Chinese criticism focuses on specific functional issues, with public discourse independently constructing AI as a parenting support tool. Tensions between corporate framing and public discourse are present across all dimensions in the U.S., while the Chinese context displays stronger functional coexistence. These patterns are shaped by the regulatory environments, media systems, and communicative dynamics of each rhetorical arena rather than by the technology itself.

Keywords: *AI for minors, strategic communication, corporate framing, public discourse, rhetorical Arena Theory, framing Theory, China, United States*

Acknowledgement

First and foremost, we would like to express our sincere gratitude to our supervisor, Howard Nothhaft, for his guidance throughout this thesis journey. We are grateful for his extensive knowledge and valuable feedback, and even more so for his consistent respect for our ideas. At important moments, he never imposed answers on us, but instead offered thoughtful questions and timely insights that helped guide our thinking forward.

Looking back, it is hard to believe that this thesis began as a conversation in Milan in October 2025. At the time, we were trying to imagine how to create a project that could be both meaningful and interesting to us. Since then, the thesis has grown through countless discussions, revisions, changing timelines, moments of uncertainty, and new ideas. Across three seasons, the project slowly evolved from scattered thoughts into the draft that now exists before us.

The most memorable milestone was undoubtedly the final twelve hours before submission. Even at the very end, we found ourselves discussing the thesis once again, realizing that it could still go one step further. What followed was an entire night of rewriting, revising, questioning, and refining until sunrise. By the end of this night, we felt genuinely proud of what we had created together, and we sincerely hope that those who read this thesis will find something meaningful in it.

We also extend our deepest thanks to each other for the shared patience and understanding. Our partnership formed the perfect synergy to carry this work forward. Although this thesis eventually comes to an end, the experiences and memories we gained along the way will remain with us long after these pages are closed.

With gratitude,

马睿鹤 & 王舒蓉

Table of Contents

1. Introduction.....	5
1.1 Research Area and Background.....	5
1.2 Research Problem and Research Gap.....	9
1.3 Research Aim.....	9
1.4 Research Questions.....	10
1.5 Contribution of the Study.....	11
1.6 Thesis Structure.....	12
2. Literature Review.....	13
2.1 AI and Minors.....	13
2.2 AI Framing, Corporate Communication, and Public Discourse.....	16
2.3 Context Comparison: China and the U.S.....	22
2.4 Synthesis of Literature Review.....	25
3. Theoretical Framework.....	27
3.1 Rhetorical Arena Theory (RAT).....	27
3.2 Framing Theory.....	30
3.3 Synthesis of Theoretical Framework.....	34
4. Methodology.....	36
4.1 Research Approach and Perspective.....	36
4.2 Research Method: Thematic Analysis.....	37
4.3 Research Design.....	38
4.3.1 Sample Selection: Four AI Companies from China and The U.S.....	38
4.3.2 Define “AI for Minors”.....	39
4.3.3 Empirical Material and Analysis Process.....	41
4.4 Reflexivity and Research Ethics.....	45
5. Findings.....	47
5.1 RQ1: Corporate Framing of AI for Minors.....	48
5.1.1 Empowering to Meet Development Needs.....	48
5.1.2 Diffuse Accountability.....	50
5.1.3 Risk Acknowledged But Managed.....	52
5.1.4 Corporate Responsibility Commitment.....	55

5.1.5 Protection Through Functional Restriction.....	57
5.1.6 Synthesis of RQ1 Findings.....	58
5.2 RQ2: How is AI for minors discussed in public discourse?.....	59
5.2.1 Superficial Adaptation.....	59
5.2.2 Contested Accountability.....	61
5.2.3 Multidimensional Harms to Minors.....	62
5.2.4 Needs for Safeguard.....	64
5.2.5 Profit over Protection.....	66
5.2.6 AI as a Parenting Support Tool.....	69
5.2.7 Synthesis of RQ2 Findings.....	70
5.3 RQ3: Alignment, Tensions, and Cross-contextual Differences.....	70
5.3.1 Minor-specific AI: Empowerment, Adaptation, and Parenting Support.....	70
5.3.2 Accountability and Responsibility.....	73
5.3.3 Risk Construction.....	74
5.3.4 Corporate Responsibility and Motivation.....	76
5.3.5 Synthesis of RQ3 Findings.....	77
6. Discussion.....	79
6.1 Legitimacy Construction and Contestation.....	79
6.1.1 The U.S.: Elaborated Framing and Public Challenge.....	79
6.1.2 China: Compliance-centred Framing and Functional Reinforcement.....	81
6.1.3 Comparative Reflection.....	83
6.1.4 Additional Observation: Communicative Style.....	83
6.2 Normative Expectations: What Should AI for Minors Look Like?.....	84
6.2.1 The U.S. Context: Unsettled and Evolving Public Expectations.....	84
6.2.2 The Chinese Context: Relatively Consistent Public Expectations.....	86
7. Conclusion.....	87
7.1 Addressing the Research Questions.....	87
7.2 Contributions to Research and Practice.....	89
7.3 Reflexivity and Future Research.....	90
References.....	92
Appendix.....	106
Appendix A. Thematic Coding Table for RQ1: Corporate Framing of AI for Minors.....	106
Appendix B. Thematic Coding Table for RQ2: Public Discourse on AI for Minors.....	107
Appendix C. Overview of Data Sources.....	110

1. Introduction

1.1 Research Area and Background

Generative artificial intelligence (AI) has become increasingly present in everyday digital life. AI systems are now used for searching information, producing texts and images, supporting learning, entertainment, and social interaction. For minors, these technologies are more than educational tools; they also serve as sources of advice, companionship, creativity, and everyday assistance. This development raises important questions about how AI is made available to younger users, how its risks are understood, and how responsibility for safe use is communicated.

Recent research suggests that generative AI is already widely used by adolescents. For example, Madden et al. (2024) found that 70% of teenagers aged 13 to 18 in the U.S. had used at least one type of generative AI tool. Among these users, 42% reported using AI to avoid boredom, while 18% used AI chatbots for personal advice. At the same time, parental awareness remains limited, with only 37% of parents recognising that their child had used such tools. Although these figures are based on the U.S. context, they point to a wider issue: generative AI is becoming part of young people's everyday information and communication environments, often before adults, institutions, and regulators have established clear norms for its use.

This situation reflects a wider pattern in children's digital media use. Existing research has shown that young people often adopt emerging technologies quickly and integrate them into their daily routines in ways that are not always visible to adults (Livingstone & Stoilova, 2021). Their digital experiences are shaped by both opportunity and risk. On one hand, digital technologies can support learning, participation, creativity, and access to information. On the other hand, they may expose minors to inappropriate content, misinformation, privacy risks, manipulation, or emotional dependency. This balance between opportunity and protection is also

reflected in international child-rights guidance, which emphasises that digital environments should support children's participation and development while also reducing risks and protecting their rights (UNICEF Innocenti, 2025). Recent scholarship on AI and childhood similarly emphasises this dual character, showing that AI may support learning, accessibility, and children's agency while also raising concerns about development, privacy, dependence, and ethical protection (Osório de Barros & Severino Soares, 2025). AI intensifies these concerns because it can respond conversationally, generate personalised content, and simulate forms of social interaction.

Against this background, AI for minors has become an emerging public concern. The issue is not simply whether children and adolescents should use AI. It also concerns how AI systems are designed and explained, and who is expected to manage the risks of use. Several AI companies and platforms have introduced minor-oriented features, parental controls, age-related restrictions, safety settings, and forms of content moderation. These measures are often presented as ways to make AI safer, more appropriate, and more useful for younger users. Yet they also raise questions of responsibility. For example, when companies provide safety features, do they take responsibility for minors' use of AI, or do they shift responsibility to parents, guardians, schools, or users themselves? When companies describe AI as beneficial for learning and creativity, how do they also communicate the possible harms?

Public concerns around AI and minors have also become more visible. Public concern is also visible in media coverage, regulatory discussions, and civil society reports, which have raised concerns about harmful content, misleading information, emotional dependency, privacy, and sexual exploitation. Recent child-safety research has similarly warned that generative AI may amplify risks such as harmful content exposure, sexual grooming, sexual extortion, and other forms of online abuse, while also requiring policy responses that involve both governments and technology companies (NSPCC, 2025). The Internet Watch Foundation, for instance, has reported increasing concerns about AI-generated child sexual abuse material, showing how generative AI can be connected to serious online harm (Internet Watch Foundation, 2024). Regulatory attention has also increased. In the U.S., the Federal Trade Commission has examined how AI chatbots acting as companions may affect children and teenagers, including

how companies evaluate safety, limit risks, and communicate these risks to users and parents (Federal Trade Commission, 2025). In China, regulatory frameworks for generative AI have placed obligations on service providers, including requirements to guide appropriate use and prevent minors from excessive dependence or addiction to generative AI services (Cyberspace Administration of China, 2023). Together, these policies, safety, and child-rights discussions show that AI for minors is being shaped through claims about benefit, protection, accountability, and appropriate use, not through technical design alone.

These developments make AI for minors a communication issue as well as a technological and educational one. Companies do more than release AI products into society. They explain what these products are, who they are for, what benefits they offer, what risks they involve, and who should be responsible for managing those risks. Through external communication, such as product pages, safety explanations, policy statements, help centre materials, and public announcements, companies frame AI for minors in particular ways. They may present it as a tool for empowerment, a controlled learning environment, a manageable risk, or a shared responsibility between platforms, parents, and users.

This places AI for minors within the field of Strategic Communication. From a strategic communication perspective, organisations use communication not simply to transmit information, but use communication purposefully to shape meanings, stakeholder relationships, and legitimacy around their activities (Hallahan et al., 2007). In this case, corporate communication is where companies define the value of AI, explain their role in protecting young users, and respond to public concerns. These communicative choices matter because they influence how AI for minors is understood by parents, educators, regulators, journalists, and the wider public.

At the same time, company communication does not exist in isolation. Public discourse may accept, question, reinterpret, or challenge corporate framings. AI for minors is shaped by how companies describe their products and by how other actors receive and debate those descriptions. The relationship between company communication and public discourse may

involve both alignment and tension. Examining this relationship helps explain how meanings, expectations, and responsibilities around AI for minors are formed.

This study focuses on China and the U.S. as two key contexts for analysing these dynamics. The U.S. plays a central role in the development of generative AI technologies, with companies such as OpenAI, Google, and xAI shaping global AI markets and public debates. Its communication environment is also shaped by market competition, media scrutiny, civil society criticism, litigation, and regulatory debate, all of which affect how technology companies are evaluated (Napoli, 2019). China provides another important context. It is one of the largest markets for digital platform use and AI adoption, and its AI development is closely shaped by state regulation, platform governance, and content control. Chinese authorities have issued specific rules for generative AI services, including obligations concerning content, data, algorithmic accountability, and the protection of minors (Cyberspace Administration of China, 2023).

The comparison between China and the U.S. does not treat the two countries as fixed opposites. It offers a way to examine how the same emerging issue is communicated and debated under different institutional, regulatory, and media conditions. Comparative research has shown that media systems, governance arrangements, and institutional contexts shape how technological risks are framed and understood (Hallin & Mancini, 2004; Just & Latzer, 2017). In the case of AI for minors, these contextual differences may affect how companies present responsibility, how public discourse constructs risk, and how the roles of parents, platforms, users, and regulators are imagined.

Against this background, this thesis approaches AI for minors as a contested communicative issue. It examines what companies say about AI for minors and how these claims relate to broader public discourse. The central concern is how meanings around AI for minors are produced, negotiated, and challenged across corporate communication and public debate.

1.2 Research Problem and Research Gap

Existing research has examined minors' digital media use, online risk, platform governance, and AI regulation. These studies show that minors' engagement with digital technologies is shaped by a tension between opportunity and protection. Digital media may support learning, creativity, and participation, but they may also create risks related to privacy, harmful content, misinformation, and vulnerability (Livingstone & Stoilova, 2021). Research on AI governance has also highlighted issues of accountability, transparency, safety, and regulation in the development and use of emerging technologies (Just & Latzer, 2017; Ulnicane et al., 2021).

However, less attention has been paid to how AI for minors is communicated as a public issue. Existing research often focuses on the risks AI may create, how minors should be protected, or how platforms should be regulated. These perspectives are important, but they do not fully explain how AI companies frame AI for minors in their external communication. They also offer limited insight into how these framings relate to wider public discourse. This gap matters for Strategic Communication because companies provide technologies while also helping to define what these technologies are, what they are for, and who should be responsible for their safe use.

This study addresses this gap by examining AI for minors as a contested communicative issue. It focuses on how AI companies frame AI for minors in their communication with external stakeholders, how AI for minors is constructed in public discourse, and how the two align or diverge across China and the U.S. The analysis is guided by three framing dimensions: how AI for minors is defined, how responsibility is attributed, and what normative or ethical positions are implied. In this way, the study connects debates on AI, minors, and platform responsibility to the field of Strategic Communication.

1.3 Research Aim

The aim of this thesis is to examine how AI for minors is framed by AI companies and constructed in public discourse in China and the U.S. The study focuses on companies'

communication with external stakeholders and on wider public discussions surrounding minors' use of AI.

More specifically, the thesis examines how AI for minors is defined, how responsibility is attributed, and what ethical or normative positions are implied in both company communication and public discourse. By comparing these two forms of communication, the study identifies where corporate framings and public discourse align, where they diverge, and how these patterns differ across the Chinese and U.S. contexts.

The study seeks to understand how meanings, expectations, and responsibilities around AI for minors are formed and contested. It examines AI for minors as a strategic communication issue shaped by organisations, public actors, media discourse, and governance environments, rather than only as a technical or regulatory issue.

1.4 Research Questions

This study is guided by the following research questions:

RQ1: *How do AI companies in China and the U.S. frame AI for minors in their communication with external stakeholders?*

RQ2: *How is AI for minors constructed in public discourse in China and the U.S.?*

RQ3: *What alignments or tensions emerge between AI companies' framing and public discourse on AI for minors, and how do these differ across the two contexts?*

Together, these questions structure the analysis across three levels: company communication, public discourse, and the relationship between the two. RQ1 focuses on how companies present AI for minors through their own public-facing communication. RQ2 examines how AI for minors is discussed and constructed in public discourse. RQ3 brings these two levels together by analysing areas of alignment and tension across the Chinese and U.S. contexts.

1.5 Contribution of the Study

This study contributes to Strategic Communication by examining AI companies as actors that actively shape public meanings around emerging technologies. AI companies provide products and technical services, but they also communicate what these technologies are, what they are for, and how their risks and responsibilities should be understood. This makes them strategic communicators whose public-facing messages influence how emerging technologies are interpreted and legitimised. In the case of AI for minors, this communication is especially important because it involves young users, public trust, and contested questions of protection, autonomy, and responsibility.

The study also contributes to research on AI communication by focusing on minors as a specific user group. Existing discussions of AI often address innovation, safety, regulation, or social impact in broad terms. AI for minors raises more specific concerns about learning, parental involvement, platform responsibility, and age-appropriate protection. This focus shows how AI communication becomes more complex when technological opportunity and minors' protection are closely connected.

In addition, the study brings together company communication and public discourse. Corporate messages do not operate in isolation. They are interpreted, questioned, and debated by other actors in the wider communication environment. Examining both levels allows the study to show how meanings and responsibilities around AI for minors are constructed, negotiated, and contested.

The study also has comparative and practical value. By examining materials from China and the U.S., it considers how similar issues are communicated and debated under different regulatory, institutional, and media conditions, without treating the two contexts as fixed opposites. As AI tools become increasingly accessible to minors, understanding how claims about safety, benefit, and responsibility are communicated can help clarify the broader social expectations surrounding AI for minors.

1.6 Thesis Structure

The thesis is structured as follows. Chapter 2 reviews literature on AI and minors, AI framing and corporate communication, public discourse, and contextual comparison between China and the U.S. It situates the study within existing research and identifies the research gap addressed by this thesis.

Chapter 3 presents the theoretical framework. It focuses on framing theory and Rhetorical Arena Theory, and explains the three analytical dimensions used in this study: how AI for minors is defined, how responsibility is attributed, and what normative or ethical positions are implied in these two different national arenas.

Chapter 4 explains the methodological approach. It introduces reflexive thematic analysis and discusses the research design, data collection, sampling strategy, coding process, and ethical considerations.

Chapter 5 presents the findings. It first examines how AI companies in China and the U.S. frame AI for minors in their external communication. It then analyses how AI for minors is constructed in public discourse across the Chinese and U.S. data. The final part of the chapter addresses RQ3 by bringing the two sets of findings together and identifying areas of alignment and tension between company communication and public discourse.

Chapter 6 discusses the findings in relation to the wider research context. It considers what the findings suggest about legitimacy, public expectations, and AI for minors as a strategic communication issue.

Chapter 7 concludes the thesis by summarising the main findings, outlining the study's contributions, discussing its limitations, and suggesting directions for future research.

2. Literature Review

This chapter reviews existing research relevant to how AI for minors is framed and discussed. It is organised around three areas. The first section examines research on the relationship between AI and minors. The second section reviews studies on corporate framing and public discourse around AI. The third section considers comparative research on the Chinese and U.S. contexts. Together, these three areas identify the research gap that the present study addresses.

2.1 AI and Minors

This section reviews existing research on the relationship between AI and minors, identifying the main research directions and remaining gaps

The most extensively developed body of research examines the risks and harms that AI may pose to minors. This body of work spans multiple dimensions. Studies from parents' perspectives have identified concerns about data collection practices, the dissemination of misinformation, and exposure to unsuitable content. These concerns are supported by broader scholarship suggesting that generative AI models may be unreliable and may produce or spread misinformation and disinformation (Monteith et al., 2024; Zhou et al., 2023). Research on minors' own perspectives also suggests concern about reliance on AI during development. Findings suggest that minors are more likely than adults to attribute human-like social and emotional traits to conversational AI systems (Druga et al., 2018). Research also confirms that minors who are overly dependent on AI may find it more challenging to form real-life social connections (Yu et al., 2025). Scholars have also highlighted that the high level of neural plasticity characterising early life means that both protective and harmful effects of AI can have an outsized impact during this period (Ho & King, 2021). Beyond these specific risks, researchers have raised broader concerns about how AI systems may affect young users at a

vulnerable stage of development, especially because these systems reflect particular values and are increasingly embedded in social interaction (Ruane et al., 2019). AI does more than provide access to content. It actively produces content, responds to individual users, and can simulate forms of social interaction. These characteristics reshape familiar digital risks, including exposure to inappropriate content, emotional dependence, and privacy concerns, while also making them less predictable (Livingstone & Stoilova, 2021). The collective focus of these studies is on the potential risks of minors being exposed to AI technology, including misinformation, disinformation, emotional dependence, developmental harm, and privacy concerns.

The second major research direction focuses on the application of AI in education. Research in this area primarily examines how integrating AI into the learning environment affects students' academic performance. Results show that students' learning outcomes can improve, including in relation to higher-order thinking and learning perception abilities (Wang et al., 2024; Wang & Fan, 2025). Research also explores how minors' interaction with AI can reduce anxiety, enhance learning motivation and self-efficacy, and create a more engaging learning environment (Shao, 2025). Beyond education, researchers have noted that AI can stimulate creativity through personalized experiences, such as customized storytelling and interest-driven dialogues. These studies mainly emphasise the positive educational potential of AI. However, a significant number of students have expressed concerns that over-reliance on AI may weaken their ability to think independently and make responsible decisions (Szmyd & Mitera, 2024). Overall, this research examines AI in several educational roles, including tutoring tools, learning partners, creative assistants, and adaptive systems that respond to individual learners. The impact of these various roles on the autonomy and initiative of minors remains debated, with studies identifying both potential benefits and risks. Nevertheless, systematic research on how AI in minors' education is publicly communicated and evaluated remains limited.

The third research direction focuses on AI governance and ethical frameworks related to minors. Researchers have achieved a deep consensus on the management of AI products. Across AI governance research, scholars have identified five core ethical principles of AI: transparency,

fairness and impartiality, non-harm, accountability, and privacy (Jobin et al., 2019). Building on these general principles, researchers have proposed that AI systems targeting minors should specifically prioritise content filtering, human oversight, transparency regarding non-human entities, and clear accountability mechanisms (Kurian, 2025). However, scholarship in this area also highlights a persistent gap between stated principles and actual practice. Kerr et al. (2020) argue that corporate commitments to ethical AI do not always translate into concrete measures, revealing inconsistencies between declared values and implementation. More broadly, the commercial imperatives driving AI development present a continuous challenge to realising ethical principles in practice (Ruane et al., 2019). In products targeting minors, this gap between principle and practice makes failures of implementation more serious. Existing research has identified this gap, but less is known about how it is communicated, justified, or challenged when AI products are aimed at minors.

A fourth, smaller but relevant body of research examines how minors perceive and make sense of AI. Studies have shown that minors are more inclined than adults to ascribe social and emotional qualities to conversational AI (Druga et al., 2018) and that minors' capacity to distinguish between AI and human interlocutors, while present, is shaped by contextual and situational factors (Xu et al., 2025). This line of research highlights that minors are not passive recipients of AI but active interpreters whose understanding of the technology is still developing. While this direction is not the primary focus of this study, it provides important context for understanding why the way AI is presented and communicated to minors matters. Such communication may shape how minors understand both AI systems and their relationship to the wider world.

The fifth research direction focuses on the role of parents in minors' use of AI. A series of studies have found that parents play an important role in minors' use of AI, but they often lack knowledge about AI and how to provide proper guidance. Research indicates that parents are gatekeepers of digital access and primary protectors of children in the digital environment; their attitudes influence children's trust in or resistance to AI technologies (Livingstone & Blum-Ross, 2020). Studies specifically addressing parental guidance in children's use of generative AI have found that parents employ various strategies to manage their children's use, but often lack

sufficient understanding of how these systems operate, making it difficult to provide effective guidance (Eira et al., 2025). Parents' awareness of their children's actual use of AI also remains limited; research shows that only a small percentage of parents are aware that their children have used generative AI tools (Madden et al., 2024). At the same time, research has revealed an intergenerational trust gap between adolescents and parents, with adolescents often placing greater trust in the accuracy and advisory capabilities of AI than their parents. Without adequate guidance and critical thinking support, this could expose young users to risks (Riolo et al., 2025). These findings are closely relevant to this study, as parental involvement and role positioning are core elements in corporate communication about AI for minors. Although existing research recognises the importance of parents, it says less about how parental responsibility is constructed in corporate framing and public discourse.

In general, existing research on AI and minors primarily examines risk, education, governance, ethics, minors' perceptions, and parental perspectives. It provides essential information for understanding the opportunities and risks that AI presents to minors, the direction of AI governance, and the role of parents in this area. However, they offer limited insight into how AI for minors is strategically communicated, how public actors accept, negotiate, or question these meanings in different contexts. This gap connects the issue of AI and minors to strategic communication and provides the starting point for this study.

2.2 AI Framing, Corporate Communication, and Public Discourse

This section turns from the effects and governance of AI use to the communication through which AI is made meaningful. It reviews research on AI framing, corporate communication, and public discourse, using scholarship from communication studies, science and technology studies, and organisational communication. These studies show how AI is constructed through competing narratives of innovation, risk, ethics, responsibility, and governance, and how these meanings are produced, interpreted, and contested by multiple actors.

A substantial strand of AI framing research has focused on how AI is represented in news media. Content-analytic and review-based studies show that AI enters public debate through recurring frames of innovation, economic progress, risk, ethical concern, and governance (Chuan

et al., 2019; Chuan, 2023). Longitudinal research adds a temporal dimension to this picture, showing that AI framing shifts over time rather than remaining stable. Analysis of three decades of AI coverage in *The New York Times* shows that AI discussion increased sharply after 2009 and remained more optimistic than pessimistic overall, even as specific concerns such as loss of control became more visible (Fast & Horvitz, 2017). More recent longitudinal work suggests that this balance has since shifted, with AI news discourse becoming increasingly critical as broader social and political debates around the technology have developed (Nguyen & Hekman, 2024). These studies show that AI framing is thematically varied and historically shaped by the evolving social and political contexts in which AI is discussed. Research on major American newspapers shows that AI coverage differs across thematic domains and in the relative emphasis placed on benefits and risks (Chuan et al., 2019). Broader research on AI as a contested emerging technology further supports this view by showing that public communication around AI is often shaped by competing interpretations of benefits, risks, ethical concerns, and governance needs (Ulnicane et al., 2021; Sartori & Bocca, 2023). This body of work treats AI less as a fixed category than as a communicative object whose representations vary across contexts, topics, and time. It also shows how media framing helps establish the categories through which AI becomes publicly discussable. However, because this research primarily examines media representation, it provides limited insight into how AI meanings are produced beyond journalism, including by technology companies, users, advocacy groups, and wider publics. This limitation makes it necessary to consider AI framing as a multi-actor process.

A multi-actor perspective treats AI-related meanings as produced across several sites, not only within news media. In this view, AI-related meanings are produced, circulated, and contested by actors with different roles, interests, and institutional positions. Policymakers may frame AI through governance and regulation, companies through innovation and responsibility, experts through technical or ethical concerns, advocacy groups through risk and accountability, and users through personal experience and everyday use. This broader view is consistent with work that treats framing as a collective and contested process, while AI-specific studies show that interpretations of AI vary across communicative settings (Benford & Snow, 2000; Deng & Ahmed, 2025).

Studies of AI discourse have extended media-oriented framing research by showing how AI meanings are shaped across different communicative sites and social positions. Research on AI framing in trending social media news shows that optimism and concern can coexist within the same discourse, with different social groups emphasizing different dimensions of the technology (Deng & Ahmed, 2025). Studies of social media discussions similarly show that ordinary users actively frame the future of AI through their own expectations, emotions, and concerns, rather than simply receiving meanings from media or institutional actors (Ocal & Crowston, 2024). Media organisations, therefore remain important in translating complex technological developments into public narratives, but they operate alongside users, companies, experts, policymakers, and advocacy groups whose interpretations also shape how AI is understood (Chuan, 2023). This perspective moves the review beyond media representation and introduces the two sites most relevant to the present study: corporate communication, where companies formulate public-facing claims about AI, and public discourse, where such claims may be interpreted, questioned, or reworked.

Within this multi-actor environment, corporate AI communication has become a particularly important area of research on AI framing. Technology companies occupy a distinctive position because they communicate the systems they develop and commercialise, rather than commenting on AI from an external position. Corporate AI communication has been examined across materials such as product descriptions, policy statements, safety pages, public commitments, blog posts, and ethics statements (Mao et al., 2025; de Laat, 2021; Zhang et al., 2026). Research on corporate AI communication shows that such materials often frame AI through usefulness, responsibility, safety, trust, and social benefit (Mao et al., 2025; de Laat, 2021; Zhang et al., 2026). In this context, terms such as “Responsible AI,” “Trustworthy AI,” and “AI for Good” are more than descriptive labels; they are normative framings through which companies position AI as beneficial, governable, and socially acceptable (Mao et al., 2025). Research on corporate responsible AI commitments shows how these broad framings are translated into public principles such as fairness, transparency, safety, privacy, accountability, and social benefit, often through website statements or multi-stakeholder initiatives (de Laat, 2021). Studies of corporate AI ethics statements similarly suggest that such texts are

stakeholder-facing communication: they align organisational AI practices with ethical expectations, establish benchmarks for responsible conduct, and signal legitimacy to stakeholders (Zhang et al., 2026). Together, research on corporate AI narratives, responsible AI commitments, and ethics statements suggests that public-facing communication defines expected value, acknowledges selected risks, and makes particular forms of responsibility visible to stakeholders (Mao et al., 2025; de Laat, 2021; Zhang et al., 2026). Corporate AI communication can be read as legitimacy-oriented framing: it enables companies to present themselves as responsible actors within broader debates about AI governance, although whether these claims are accepted depends on external interpretation and public evaluation (Iwanowska, 2025; Salehi et al., 2026).

At the same time, research on ethical AI expectations and digital ethicswashing raises questions about the limits of corporate responsibility narratives. Broad ethical commitments may become performative when they are not connected to concrete practices, institutional accountability, or operational changes (Kerr et al., 2020), while ethical claims may also operate as symbolic or misleading communication when they create an impression of responsibility without substantive action (Schultz et al., 2025). Still, such statements should not be dismissed as empty rhetoric: they may signal organisational commitment and establish normative benchmarks, but they need to be read critically when accountability and implementation remain unclear (de Laat, 2021; Kerr et al., 2020; Schultz et al., 2025). Much of this literature still examines corporate AI communication at a general level. Less attention has been paid to how legitimacy-oriented framings operate when AI is communicated as suitable, safe, or beneficial for minors, or how such framings are interpreted in public discourse.

Corporate AI communication research explains how companies frame AI through claims of usefulness, responsibility, and social benefit. Public discourse research shifts attention to how such claims are interpreted beyond corporate settings. This shift prevents public discourse from being treated as a secondary reaction to corporate messaging. Instead, it positions public communication as a space where claims about usefulness, safety, responsibility, and social benefit are compared with wider social expectations and lived experiences. For the present study, this is relevant because corporate framings of AI for minors need to be examined within a public

arena already shaped by concerns about protection, accountability, and appropriate use. Studies of socio-technical imaginaries and technological expectations approach emerging technologies as objects whose meanings are shaped through promises, narratives, institutional discourse, and public debate, rather than through technical features alone (Jasanoff & Kim, 2019; Borup et al., 2006; Sartori & Bocca, 2023). In this view, public expectations are not simply reactions to technological development, but are collectively produced, circulated, and revised across media coverage, policy debate, corporate messaging, user experience, and everyday discussion (Jasanoff & Kim, 2019; Borup et al., 2006; Sartori & Bocca, 2023).

In AI communication research, these studies position public expectations as collective and shifting, rather than as fixed reactions to technology (Jasanoff & Kim, 2019; Borup et al., 2006; Sartori & Bocca, 2023). Research on public perceptions of AI shows that major technological events can alter these expectations quickly. For example, Wang and Liang's (2024) study of public discourse before and after the launch of ChatGPT found that ChatGPT amplified both public enthusiasm and expectations surrounding AI, as well as concerns and fears about its consequences. Such findings point to the mixed character of public discourse, where optimism and concern may coexist and shift as AI becomes more visible in everyday life, policy debate, and media coverage.

Thus, public discourse scholarship complicates a company-centred view of AI communication. Corporate claims enter an environment already shaped by broader narratives of hope and fear, prior beliefs, institutional debates, and lived experiences with technology (Sartori & Bocca, 2023; Ulnicane et al., 2021). Research on AI trust adds a more individual-level perspective. In an experimental study, Kim and Song (2023) show that message framing and decision ownership shape trust in an AI decision-support context. This finding is relevant here because it indicates that the communication of AI capacities and limitations can affect user evaluation, even though such work focuses on individual trust rather than broader public discourse. However, existing research has paid less attention to how this process unfolds around specific AI issues, including how corporate claims are echoed, modified, or contested in public discourse.

The literature above points to three areas of AI communication that are relevant to this study. Media framing research has established that AI is represented through recurring yet changing frames, with variation across topics and over time (Chuan et al., 2019; Chuan, 2023; Fast & Horvitz, 2017; Nguyen & Hekman, 2024). However, the central issue for this study is less media representation itself than the relationship between corporate communication and public discourse. Corporate AI communication research has shown how companies frame AI through responsibility-oriented narratives, including claims about usefulness, safety, trust, and social benefit (Mao et al., 2025; de Laat, 2021; Zhang et al., 2026). Public discourse and expectation studies, in turn, have examined how AI is interpreted, debated, and contested by different publics (Jasanoff & Kim, 2019; Borup et al., 2006; Wang & Liang, 2024; Sartori & Bocca, 2023). Yet these bodies of work are often examined separately, leaving less understanding of how corporate claims are accepted, modified, or challenged within public discourse, and how such relationships produce alignments or tensions around the same AI issue.

This creates both a conceptual and methodological limitation. In strategic communication, gap analysis has been used to explain discrepancies between organisational claims or performance and societal expectations, including expectation gaps and legitimacy gaps (Kim & Ji, 2018). Earlier work on technology communication similarly suggests that discrepancies between organisational messaging, media representation, and public expectations can emerge through the way technological promises circulate and are interpreted (Kazoleas & Teigen, 2006). More recent research on corporate AI ethics statements also shows that AI-related corporate communication is closely tied to stakeholder expectations and legitimisation cues (Zhang et al., 2026). Methodologically, many studies examine either corporate communication or public discourse, rather than placing the two side by side. This makes it difficult to analyse whether corporate responsibility claims align with public concerns, or whether they are questioned and reworked through public discourse. In the context of AI, such discrepancies are significant because legitimacy depends on what companies claim and on whether these claims are perceived as credible, accountable, and aligned with wider social expectations (Iwanowska, 2025; Salehi et al., 2026).

This gap is especially relevant for AI for minors. As discussed in the previous section, AI for minors involves intensified questions of value, safety, protection, and responsibility. Yet existing AI communication research has paid limited attention to how companies frame AI when minors are involved, or how these framings are evaluated in public discourse. Examining corporate communication and public discourse together makes it possible to identify what companies claim about AI for minors and where these claims are accepted, questioned, or contested. This provides the basis for analysing the alignments and tensions between corporate framing and public discourse in the case of AI for minors.

2.3 Context Comparison: China and the U.S.

This section reviews existing research on AI communication and governance in China and the U.S. It examines the characteristics and differences between the two countries in terms of governance frameworks, media environments, and communication conditions, and discusses the significance of these differences for this study.

First, AI policy discourse provides an important contextual background. A number of scholars have conducted detailed analyses of the two countries' AI strategies (Allen, 2019; Rasser et al., 2019). As a rapidly developing technology, AI has a significant impact on the global economy and international politics, making AI governance an area of active competition among major international players (Chen & Xu, 2025). China and the U.S. share one of the world's most important bilateral relationships, and both countries view AI as a top national priority, having each developed national strategies (China in 2017 and the U.S. in 2019) to guide and regulate the technology's future development (Ding & Kong, 2019; Hine & Floridi, 2024). Comparative analysis of official policy documents reveals that U.S. documents encompass a broader range of concerns and exhibit greater emotional intensity, with competitive dynamics with China implied. Chinese national-level documents, by contrast, are characterised by a consistently positive, development-oriented framing, with particular emphasis on innovation and technological application (Hine & Floridi, 2024). These basic differences in policy discourse form part of the wider rhetorical context in which companies and publics communicate about AI, including its implications for minors.

At the governance and regulatory level, research reveals significant differences in governance frameworks across countries regarding the protection of minors' rights in the digital environment. Studies of China's governance framework describe a multi-layered system consisting of statutory laws and regulations, authoritative guidelines for policy implementation, multi-level technical standards, and corporate self-regulatory measures (Chen & Xu, 2025). This reflects a state-led, multi-stakeholder governance model in which compliance is mandatory, and the state assumes primary responsibility for social protection. Research on the U.S. context shows that the fundamental principle underpinning AI strategy has been innovation. The application of AI has been demonstrated to encourage innovation and economic growth, which in turn benefits society, particularly in terms of commercial development. The U.S. government's approach to its regulatory function is one of relative limitations. This reflects a more permissive regulatory approach, expressed in policy language such as “allowing a thousand flowers to bloom” and “creating space for innovation and growth in AI” (Executive Office of the President et al., 2016; Cath, 2018). In response to growing concerns about the risks associated with minors' use of AI, the U.S. updated the Children's Online Privacy Protection Act (COPPA) in April 2025 to incorporate regulations governing AI services involving minors (Federal Trade Commission, 2025). Meanwhile, China has imposed obligations on AI service providers, including guiding reasonable use and preventing excessive dependence by minors (Cyberspace Administration of China, 2023). These regulatory approaches reflect notably different management and protection logics. Policy research reveals that the U.S. primarily protects minors through rights-based frameworks (such as privacy laws), while China employs a different approach: behavioral restrictions, such as addiction prevention systems and screen time controls (Chen & Xu, 2025). These regulatory differences not only affect the obligations that businesses must fulfill but also the default responsibility frameworks they follow and societal expectations.

Media systems and communicative environments-wise. Research on media systems provides further context for understanding these structural differences. Hallin and Mancini's (2004) comparative media systems theory suggests that media systems are shaped by the wider context of political history, structure, and culture. Research drawing on cultural dimensions has found that in individualistic cultures, AI is often viewed as a potential threat to human autonomy

and privacy, while in collectivist cultures, it is more commonly seen as an extension of the self or a tool for the common good, with greater public support for government regulation (Suerdem & Akkilić, 2021). A close analysis of Western media reveals a consistent emphasis on individual autonomy, risk monitoring, and countering geopolitical threats, while the Chinese media reveals a frequent emphasis on the benefits of AI for economic cooperation, stability, and the creation of heroic figures (Ding & Kong, 2019). Consistent with this, Bareis and Katzenbach (2022) also find that the absence of a stringent regulatory framework has led the media to assume a 'watchdog' role in the U.S., with closer attention to both risks and potentials. Empirical studies of AI-related media discourse confirm that these structural differences produce notably different communicative patterns. A comparative study of the New York Times and the People's Daily found that the two countries focus on different AI domains: the U.S. places significant emphasis on technological expertise and specialised talent, while China prioritises internet applications. Furthermore, the tones diverge: U.S. media often adopts a cautious stance, particularly concerning potential risks and military implications, while Chinese media exhibits a more optimistic and positive perspective (Ding & Kong, 2019). Research on Chinese digital platforms has shown that public discourse about AI on WeChat closely aligns with official narratives and is dominated by industry and political actors, with positive evaluations of AI's economic potential prevailing (Zeng et al., 2020). Scholars have characterised the resulting Chinese media environment as reflecting a technological solutionism narrative, in which AI is depicted as a catalyst for social advancement rather than a source of unmanageable risk (Bareis & Katzenbach, 2022). This aligns with previous research findings. These patterns suggest that corporate participants in both countries face different structural pressures when communicating about AI for minors. Research indicates that U.S. companies need to navigate a public discourse characterized by skepticism and a focus on risk and rights, while Chinese companies need to operate within regulatory frameworks and implement nationally recognized narratives of security and social responsibility (Ding & Kong, 2019; Zeng et al., 2020). These differences in cultural, regulatory, and media environments collectively constitute two distinct communication environments, influencing corporate framing and how corporate messages are received.

Despite these insights, existing research primarily focuses on broader policy discourse, regulatory frameworks, and media environments in China and the U.S. While a few studies have begun to explore how corporate actors construct their visions of AI governance within different national contexts (Mao et al., 2025), this research direction has not yet addressed the interplay between corporate framing and public discourse surrounding specific issues, such as minors' use of AI. Less is known about how different rhetorical contexts, protection logics, media systems, and communicative environments shape the relationship between corporate framing and public expectations, or whether they produce different patterns of alignment and tension around specific issues. This is the gap addressed by the present study.

2.4 Synthesis of Literature Review

This literature review has brought together three bodies of scholarship. The first concerns AI and minors, where existing research has mainly examined risk, education, governance, ethics, minors' perspectives, and parental roles. This literature establishes AI for minors as a sensitive and responsibility-laden issue, but it remains largely focused on effects, harms, design principles, and governance responses rather than on communication.

The second body of literature concerns AI framing, corporate communication, and public discourse. Research in this area shows that AI is not understood through fixed meanings, but through competing frames and narratives shaped by media, companies, users, institutions, and publics. It also shows that companies use responsibility-oriented communication to position AI as useful, safe, trustworthy, and socially beneficial, while public discourse may interpret, negotiate, or challenge these claims.

The third body of literature concerns China–U.S. contextual differences. Existing research shows that AI communication is shaped by different governance frameworks, regulatory logics, media systems, and political and cultural environments. The U.S. context is more strongly associated with rights-based protection, public scrutiny, and risk-oriented debate, while the Chinese context is more closely connected to state-led governance, compliance, social stability, and developmental narratives.

All in all, these studies provide important foundations, but they remain insufficiently connected. Research on AI and minors rarely examines how the issue is strategically communicated. Research on AI framing and public discourse often studies corporate communication and public interpretation separately. Comparative research on China and the U.S. tends to focus on broader AI governance or media discourse rather than on AI for minors specifically. This study addresses these gaps by examining how AI companies frame AI for minors, how public discourse accepts, negotiates, or contests these framings, and how patterns of alignment and tension differ across Chinese and the U.S. contexts.

3. Theoretical Framework

In order to examine how “AI for minors” is strategically framed by technology companies, how it is negotiated in public discourse, and how these dynamics differ across the Chinese and the U.S. contexts, this study uses two complementary theoretical frameworks. Rhetorical Arena Theory (Frandsen & Johansen, 2000) provides the macro-level structure, defining the communicative arena and explaining how multiple voices interact within it. Framing theory (Entman, 1993) provides a micro-level analytical lens for examining how meaning is constructed and what positions are embedded in communicative choices. Rather than applying these frameworks in their original forms alone, this study includes later scholarly developments to make the framework more suitable for the present research. Together, RAT and framing theory provide the structural and interpretive tools needed to address the research questions.

3.1 Rhetorical Arena Theory (RAT)

Dissatisfied with the limitations of existing crisis communication frameworks, and inspired by Image Repair Theory (Benoit, 1995), which centers on an organization's communicative efforts to restore a damaged reputation, and Situational Crisis Communication Theory (Coombs, 1999), which focuses on crisis management to prevent and mitigate reputational threats, Frandsen and Johansen first conceptualized the "Rhetorical Arena" in 2000. As a more expansive theoretical alternative, it highlighted the importance of the rhetorical strategies employed by diverse actors, the textual conventions through which these strategies are conveyed, the agenda-setting and gatekeeping functions of mass media, and the cultural contexts that shape how crises are interpreted and negotiated.

Building on this, Johansen and Frandsen (2007) initially used the term "stakeholders" to describe the range of actors participating in the rhetorical arena. However, recognizing that not all communicative participants hold formal stakeholder status in relation to an organization, they

subsequently refined this terminology, replacing "stakeholders" with "voices" (Frandsen & Johansen, 2022). This shift highlights the uncertainty and complexity of the communicative arena. In situations of crisis, the multiple voices present within the arena collectively contribute to the process of construction. They communicate, negotiate, argue, ignore, or discuss together (Johansen & Frandsen, 2007; Frandsen & Johansen, 2020). It acknowledges that the arena is populated by a fluid and unpredictable range of communicators, whose involvement cannot be fully anticipated or controlled. Thus, RAT is far from linear models of communication. It positions the arena as a dynamic discursive space in which meaning is continuously produced and contested through the interaction of multiple, often competing, voices. Multiple actors speaking simultaneously influence how an issue is understood (Frandsen & Johansen, 2017). RAT posits that these discussions can occur across a variety of channels and settings, which may result in divergent interpretations of the same feature (Frandsen & Johansen, 2022).

It is important to recognize that not all voices within the rhetorical arena carry equal weight. Extending Frandsen and Johansen's framework, Raupp (2019) argues that the distribution of rhetorical resources among participants is inherently asymmetrical. Corporate and mainstream media actors, using their institutional resources, media channels, and strategic communication capabilities, hold a structurally dominant position, serving as primary sources of information and, to a certain extent, shaping the narratives that dominate public discourse (Van Aelst & Walgrave, 2016). While individual voices are growing stronger through digital platforms, this power imbalance remains particularly evident in issues involving minors, whose interests are often articulated on their behalf by parents, advocacy groups, and regulatory bodies. Recognizing this asymmetry is crucial for understanding who speaks in each context, whose voices are amplified, whose are marginalised, and whose are marginalized, and whose interests the prevailing narratives ultimately serve.

Crucially, RAT is an open theoretical framework that encompasses a wide range of controversy types, rather than being confined to conventional crisis scenarios (Frandsen & Johansen, 2022). The issue of minors' use of AI is not a crisis in the traditional sense with a clear triggering event; it is an ongoing, multi-actor debate in which technology companies, regulators, media organizations, parents, and young users all participate as distinct voices. Each party

represents distinct interests and employs different resources and rhetorical strategies. Within the differing cultural and political contexts of China and the U.S., the relative power and positions of these parties also vary, adding an extra layer of complexity to the analysis. This configuration is consistent with the patterns identified in the literature review, which establishes that AI meanings are produced through the interaction of multiple actors with divergent interests, and that corporate communication strategies and public expectations are shaped by distinct institutional and cultural environments. This makes the multi-voiced, non-linear perspective offered by RAT particularly well-suited for the present study.

Using the concept of the rhetorical arena, this study conceptualises the Chinese and the U.S. contexts as two structurally distinct communicative environments, whose defining characteristics have been established through the literature review. In the context of China's regulatory landscape, the rhetorical arena is defined by state oversight. This establishes the boundaries within which all participating voices must operate. AI companies such as Doubao, as corporate entities, are required to adhere to the current national regulations. This creates a clearly defined boundary structure for the 'Chinese rhetorical arena'. Within this structure, the public shows a fundamental trust in the state, which is then extended to enterprises operating under state regulation (Steinhardt, 2012). The power structure here follows a top-down model, although the media and the public have a voice, national regulations are always broadcast through the loudspeakers. In addition, the optimistic cultural attitude towards AI is also reflected in the public, who tend to view AI as a useful tool rather than a threat. This has been shown to affect the 'pattern of interaction', which is consistent with the broader 'technological solutionism' narrative identified in existing literature (Bareis & Katzenbach, 2022). Consequently, an atmosphere of trust and positivity is fostered, reflected in a broadly held optimism that technology can improve everyday life, in which corporate and public voices tend to align rather than contest one another.

In contrast, the U.S. rhetorical arena is characterized by a market-driven dynamic in which corporate actors, such as ChatGPT, Gemini, and Grok, are competing with one another. This has boosted the dynamism and innovation of marketing, but has also facilitated the formation of a competitive communicative environment. Unlike the boundary-setting function of

state regulation in the Chinese context, the U.S. rhetorical arena is relatively broad and flexible, with its boundaries continuously negotiated among diverse voices. It is shaped by regulatory agencies, non-governmental organisations, public media, and individual users. These actors participate as watchdogs, advocates, or critics, raising concerns about potential risks, rights violations, and corporate accountability. Nevertheless, power asymmetry persists within this arena. Corporate and mainstream media actors retain significant structural advantages in terms of media access, strategic communication resources, and agenda-setting capacity, even amid the more open and contested discursive environment. Different voices negotiate responsibilities, challenge corporate narratives, and debate what constitutes ethically acceptable AI development for minors.

The present study employs the RAT of Johansen and Frandsen (2022), conceptualising China and the U.S. as two rhetorical arenas in which multiple voices interact under unequal conditions of power “to identify, describe, and explain patterns within the multiple communication processes taking place” inside of the “AI for minors” field (Frandsen & Johansen, 2017; Raupp, 2019). In doing so, RAT serves as the overarching structural lens: it guides the analysis of how corporate voices construct public framings of AI for minors (RQ1), how public voices respond to and contest these framings across media and online platforms (RQ2), and how the interaction between these voices differs across the two arenas (RQ3). RAT provides this study with a structured yet flexible lens for mapping the communicative dynamics of each arena. The framing analysis that follows is grounded in this arena-based understanding of communication.

3.2 Framing Theory

While RAT maps the arena and its participants, framing theory explains how these actors construct meaning within it. It focuses on how they selectively construct and promote particular meanings of AI for minors. The literature review identifies several patterns that make this focus necessary. Corporate actors tend to frame AI in normative terms, highlighting responsibility, trust, and social benefit rather than technical features. At the same time, public expectations of ethical AI are not fixed but shaped through communication. A recurring gap has been identified

between how companies frame AI and what publics expect. This gap has consequences for how AI technologies are evaluated and legitimized in the public sphere. These findings converge on a shared concern: meaning construction in AI communication is neither neutral nor unified; it depends on the positions, interests, and contexts of different actors (Ulnicane et al., 2021; Rotolo et al., 2015). Framing theory provides a way to analyse this process systematically.

Framing refers to the selective emphasis of certain aspects of an issue while others are downplayed. This process shapes how problems are defined, how causes are assigned, and what responses are seen as appropriate (Entman, 1993). Gamson and Modigliani (1989) similarly conceptualize frames as interpretive packages that give meaning to an issue by linking particular ideas, values, causes, and solutions. This means that a frame is not the same as a topic or theme. A text may mention safety, education, or responsibility, but the analytical question is how these elements are combined into a particular interpretation of AI for minors. In this study, framing is treated as an active communicative practice. It is used to examine how meanings of AI for minors are produced, circulated, and contested. For example, a company may describe its AI product as a "safe learning companion," while a news article frames AI for minors as a threat to cognitive development. These are not neutral descriptions. They guide attention, assign responsibility, and shape what actions appear justified.

This process is especially important in the context of AI for minors, where meanings are contested, and stakes are high. The technology can be framed as a tool for learning, creativity, and equal access to knowledge, or as a source of risk, dependency, and ethical harm. These interpretations are not mutually exclusive, and they do not come from the technology itself. As the literature review establishes, they are produced through the interaction of multiple actors (Benford & Snow, 2000; Deng & Ahmed, 2025). Corporate communication tends to emphasize safety features, responsible design, and developmental benefits, while public discourse is more mixed and sometimes contradictory. As shown by Wang and Liang (2024), it is shaped by broader narratives of hope and fear, prior beliefs, and ongoing debates. The coexistence of these competing frames means that no single interpretation dominates (Chuan, 2023). Framing should be understood as an ongoing process of negotiation rather than a fixed representational outcome.

Within this multi-actor environment, the relationship between corporate framing and public discourse is critical. The literature review shows that companies do not simply respond to existing public expectations. They actively seek to shape them, constructing normative imaginaries of AI, and positioning their technologies within broader societal and regulatory debates (Mao et al., 2025). Yet these constructions are not received passively. As Kerr et al. (2020) demonstrate, public expectations of ethical AI are produced through interactions among policymakers, industry actors, and the public. They may not align with corporate messaging. This creates a relationship in which corporate frames and public expectations interact, and sometimes conflict. This study focuses on that interaction. It examines which frames are used and how they relate to one another across different communicative contexts.

Framing is also shaped by socio-political context. The meanings that different actors attach to AI for minors depend on institutional settings, governance structures, and media systems. This is why the comparison between China and the U.S. matters: the same technology may be framed differently across contexts, and the relationship between corporate communication and public discourse may also vary (Malmberg & Trondal, 2021). In China, public discourse about AI often aligns closely with official and industry narratives (Zeng et al., 2020). While in the U.S., it is more diverse and often more critical. These contextual differences influence both how AI is framed and how those frames are received and interpreted.

To operationalize framing theory as an analytical tool, this study identifies three dimensions through which the construction of AI for minors is examined. Based on Entman's understanding of framing as involving problem definition, causal interpretation, moral evaluation, and treatment recommendation, this study operationalises framing through three analytical dimensions. These dimensions do not mechanically reproduce Entman's four functions, but adapt the logic of framing theory to the specific context of AI for minors. They also came from patterns identified in the literature review, allowing the analysis to remain both theoretically grounded and empirically sensitive. They allow for systematic comparison between corporate communication and public discourse, and across the two national contexts.

a) The first dimension is functional orientation. It examines whether AI for minors is presented mainly as a tool for learning or development, or as a source of risk, harm, and ethical concern. This reflects the central tension between benefit and risk identified in existing research on AI framing (Chuan, 2023; Nguyen & Hekman, 2024). Instead of treating benefit and risk as fixed categories, this dimension focuses on how different actors may prioritize one side of this tension in their communication.

It examines how emphasis is placed on either benefits or risks, and how this emphasis may vary across different communicative contexts. In this way, it helps clarify how the role of AI for minors is defined in practice. This dimension also considers how these framings relate to context-specific concerns. Prior research suggests that AI discourse may be linked to different priorities across socio-political contexts. In some cases, emphasis is placed on education, development, and collective progress, while in others, greater attention is given to privacy, individual rights, and protection from harm. These variations provide a basis for examining how functional orientation is constructed differently across contexts.

b) The second dimension is responsibility attribution. It examines how responsibility for the safe and appropriate use of AI for minors is distributed among actors such as technology companies, parents, educators, users, and regulatory bodies. This dimension is particularly relevant to RQ1, as corporate communication strategies may assign responsibility in specific ways. These choices shape how companies present their role in managing risks and ensuring safety.

It is also central to RQ3. Differences between how responsibility is framed and how it is expected within public discourse may indicate areas of tension. Responsibility may be attributed to companies, users, or regulatory institutions. These variations reflect different assumptions about who should protect minors in digital environments. Analysing responsibility attribution makes it possible to identify where alignment may exist and where gaps may emerge between corporate communication and public expectations. It also provides a way to examine how these differences relate to broader concerns about trust, accountability, and legitimacy in AI (Iwanowska, 2025; Salehi et al., 2026).

c) The third dimension is normative positioning. It examines how appropriate and inappropriate uses of AI for minors are defined. This includes what is encouraged, restricted, or problematized in different forms of communication. In other words, it focuses on the implicit value judgements embedded in different framings about what minors' relationship with AI should look like. Beyond risk and benefit, this dimension captures the norms that shape expectations of how AI should be used (Borup et al., 2006; Jasanoff & Kim, 2019).

Different forms of communication may promote or question particular uses of AI, contributing to ongoing debates about appropriate boundaries. These boundaries are not fixed. They are shaped through interpretation and negotiation. This dimension also reflects broader assumptions about childhood and technology. It captures how different actors construct desirable forms of interaction between minors and AI, such as ideas about autonomy, dependence, and development. This helps reveal the values that underpin different framings and how they vary across contexts.

These three dimensions provide a structured way to analyse how meanings of AI for minors are constructed across actors and contexts. They make it possible to examine how corporate framing relates to public expectations, and to identify where alignment or tension may emerge. Applying these dimensions across China and the U.S. allows the study to capture both differences in framing content and differences in how meanings are negotiated within distinct socio-political environments. This establishes a clear link between the theoretical framework and the empirical analysis that follows.

3.3 Synthesis of Theoretical Framework

The combination of RAT and framing theory provides a multi-dimensional framework for analyzing the construction of AI for minors. RAT treats China and the U.S. as distinct rhetorical arenas, mapping the range of actors involved and the power asymmetries that shape their interactions. Framing theory provides the tools to examine how meaning is constructed by different actors. The three analytical dimensions, namely functional orientation, responsibility attribution, and normative positioning, link the theoretical framework to the empirical analysis.

They enable a systematic comparison of how corporate actors frame their technologies and how public discourse negotiates or challenges those frames.

Together, this framework provides the basis for identifying patterns of alignment and tension between corporate communication and public expectations across the two contexts.

4. Methodology

This chapter presents the methodological approach adopted in this study. Given the exploratory nature of the research questions and the limited pre-existing frameworks in the emerging field of AI applications for minors, it was challenging to develop a priori testable hypotheses. For this reason, a qualitative research approach was adopted to explore and interpret the meanings constructed in this field (Flick, 2023).

This study is grounded in the interpretivist paradigm, which assumes that reality is socially constructed through discourse and meaning-making processes (Prasad, 2017). Thematic analysis and an inductive approach were used to identify, analyse, and report patterns in the collected data (Braun & Clarke, 2006; Thomas, 2003). The following part describes the research design, including the selection of sample companies, the geographic scope, and the criteria and procedures for data collection. The remainder of this section discusses considerations regarding reflexivity and research ethics.

4.1 Research Approach and Perspective

A qualitative research method was applied in this study to better understand how the meaning of “AI for minors” is negotiated and constructed through communication processes between AI companies and external stakeholders.

In the field of strategic communication research, emphasis is placed on the 'actors', 'influence', and 'effects' in communication processes (Falkheimer & Heide, 2018). The interpretivist tradition believes that individuals actively construct meanings through their interactions with objects, events, and social actors. According to this perspective, reality is socially constructed and shaped by subjective interpretations in different social contexts (Prasad, 2017). AI companies are not neutral actors; they are strategic communicators. The public is not

passive receivers, they are active negotiators. This standpoint directs the study toward the process of constructing meaning.

This interpretivist orientation aligns with the theoretical framework of this study. Framing theory treats meaning as the product of selective emphasis, and Rhetorical Arena Theory examines how multiple voices interact within communicative environments. Both frameworks assume that meanings are constructed through communication rather than given in advance. The methodological approach, therefore, focuses on identifying how different actors frame AI for minors and how these framings relate to each other across contexts.

To operationalise this approach, the study employs reflexive thematic analysis (Braun & Clarke, 2019) as the primary analytical method. Thematic analysis is used to identify recurring patterns in corporate communication and public discourse, which are then interpreted through the lens of framing theory. The three analytical dimensions established in the theoretical framework, functional orientation, responsibility attribution, and normative positioning, serve as sensitising concepts that guide the interpretation of the data without predetermining the thematic structure. In this way, the themes identified through inductive analysis can be understood as frames: structured ways of defining issues, attributing responsibility, and positioning normative expectations around AI for minors.

4.2 Research Method: Thematic Analysis

To analyse the data, this study employs reflexive thematic analysis as outlined by Braun and Clarke (2019). Thematic analysis is particularly appropriate for this study, as it can facilitate the description of the patterns in the data. Moreover, it can be regarded as a constructionist method, which examines the ways in which events, realities, meanings, experiences, and identities are produced through a range of discourses operating within society. (Braun & Clarke, 2006). This is consistent with the interpretivist perspective of this study. While the analysis remains open to emergent patterns within the material, the coding process is theoretically informed by the three analytical dimensions established in the theoretical framework. These dimensions function as sensitising concepts rather than fixed coding categories.

This study followed the six phases of analysis outlined by Braun & Clarke (2006): 1) familiarizing ourselves with the data; 2) generating initial codes; 3) searching for themes; 4) reviewing themes; 5) defining and naming themes; 6) producing the report. Analysis involves a constant process of iteration across the entire data set.

The credibility of the analysis was supported by the involvement of two researchers with different backgrounds. According to Braun and Clarke (2019), involving two researchers supports a more collaborative and reflexive analytic process. Given the cross-national material used in this study, the coding process involves continuous dialogue between the two researchers to ensure that specific words are translated consistently within their respective contexts. For instance, the English term 'safeguard' may correspond to three distinct Chinese concepts: "保障" (to guarantee functionality), "保护" (to prevent harm), and "维护" (to maintain stability). This study prioritises translation to preserve these context-specific implications. Furthermore, assumptions and researcher positioning are inherent components of qualitative research; this study acknowledges that differences in interpretation are valuable, as they require critical reflection and deeper engagement with the data, and therefore allow reasonable interpretive differences to remain (Braun & Clarke, 2019).

4.3 Research Design

4.3.1 Sample Selection: Four AI Companies from China and The U.S.

This study uses purposive sampling to select cases that are particularly relevant to its aims. China and the U.S. were chosen as comparative contexts because they represent two of the world's most influential AI markets.

The rapid development of Chinese AI models, such as DeepSeek, has challenged America's ability to maintain its AI leadership, as DeepSeek debuted a powerful open-source AI model in early 2025. This highlights the importance of examining these two contexts (Yuan, 2025). Further evidence of this concentration can be found in LLM Stats (n.d.). According to the Code Arena ranking system (based on human evaluation across multiple coding tasks), 59 of the top 60 large language models originate from either China or the U.S. Among these, 35 U.S.

models are developed by four major companies: OpenAI's ChatGPT, Google's Gemini, xAI's Grok, and Anthropic's Claude. While 24 Chinese models are developed by eight companies. This distribution reinforces the dominance of these two countries in the current AI landscape.

The present study selected four influential companies from the Chinese and U.S. markets that are the most influential, widely used, and representative, and that offer an "AI for minors" feature. ChatGPT, Gemini, Grok, and Doubao. In the U.S., the three dominant players in the consumer chatbot market, ChatGPT, Gemini, and Grok, have each implemented dedicated youth-safety frameworks or "Minor Modes." According to Reuters (Singh & Bajwa, 2026), these platforms collectively command the market, with ChatGPT holding 52.9%, Gemini 29.4%, and Grok 17.8% of the U.S. chatbot share. The inclusion of these high-reach platforms ensures that the analyzed communication strategies represent the industry's mainstream discourse. To provide a meaningful cross-national contrast, this study incorporates Doubao, a leading AI application developed by ByteDance (the parent company of TikTok). Data from QuestMobile (2026) indicates that Doubao has emerged as China's most widely adopted AI platform, surpassing competitors like DeepSeek with 84 million daily active users (DAU) at launch and a peak of 145 million during the Chinese New Year. As Doubao has also established a mandatory "Minor Mode" (未成年人模式) under Chinese regulatory requirements, while other Chinese AI systems have not yet been equipped with features for minors, it serves as a critical case for examining how AI safety is communicated within the unique Chinese socio-political context.

China and the U.S. have distinctive media systems and socio-political contexts (Chen & Rowlands, 2022). These differences provide a valuable basis for examining how public trust is shaped and contested across different environments surrounding AI for minors and youth safety. In addition, the selection of these cases ensures analytical feasibility, considering the languages accessible to the researchers.

4.3.2 Define "AI for Minors"

As the selected AI companies either did not use the same term or lacked a term to describe their AI features designed for children, this study uses the term 'AI for minors'. Given that the four platforms define age limits differently (e.g., 13–18 or under 18), this study adopts a

broader, unified age range (under 18) as the analytical framework to ensure consistency and comparability across cases. In this context, “minors” refers to individuals under the age of 18.

These four conversational AIs allow human users to communicate with an automated system using natural language. The interaction may be speech and/or text-based (Ruane, Birhane & Ventresque, 2019). The following section provides a brief overview of these four conversational AIs and their features designed for minors, to offer essential background information. This overview is based on information retrieved in April 2026.

ChatGPT was launched by OpenAI in 2022. ChatGPT is for individuals 13 years of age and older (OpenAI, 2025b). Starting in August 2025, the company introduced dedicated teen safety features to better protect younger users. Since then, a series of measures have been gradually implemented, including age estimation, under-18 safety policies (content moderation), a default safe experience for users under 18, and parental controls. ChatGPT uses automatic age estimation, or user-provided age information, or users add family members to their account (e.g., 'my child' or 'my parents or guardian'), to adjust the level of protection (OpenAI, 2025b).

Gemini is an AI tool developed by the multinational technology company Google. In late 2025, Google published a set of guidelines outlining enhanced protective measures for minors in the Google Gemini Apps suite. These guidelines included the implementation of content filters and the Family Link feature, which, as per Google's instructions, requires parents to enable access for their children under the age of 13 (or the relevant age in the user's country). The primary objective of this measure is to empower parents to supervise their children's online activities. Gemini positions itself as an educational assistant for minors (Google, 2025).

Grok, a product of Elon Musk's xAI and integrated into the X (formerly Twitter) platform, employs a distinctive approach. Its reliance on data from platform X, which prioritises freedom of expression and more dynamic social interactions, gives rise to concerns regarding informational bias (de Carvalho Souza & Weigang, 2025). Grok is a software application that allows adolescents aged between 13 and 17 years, provided that they have obtained the consent of their parents or guardians. It appears that Grok has a general control mechanism in place for all users, with no specific “features for minors” description available on the official site (xAI

LLC, 2025). The “Kids Mode” can be turned on by younger users or their parents in the Grok app. Grok has also launched a virtual companion assistant for children. It is optional to set a password to turn it off. In July 2025, the company's founder, Elon Musk, announced plans on X to release a product named "Baby Grok," which was designed with children in mind (Musk, 2025). However, as of March 2026, this product had not yet been officially released.

Doubao (豆包) is an AI assistant developed by ByteDance. In late 2025, Doubao implemented its dedicated 'Minor Mode' (未成年人模式). Guardians can control whether this mode is enabled for their children under 18. They can also choose to set a separate password to manage this mode. Once enabled, Doubao will provide more age-appropriate content and restrict or disable features such as chatting with AI agents outside of Doubao, displaying related videos in responses, browsing third-party websites, and AI creation, according to the guardian's preferences. The activation process is integrated into the user interface via a standardised pathway: [Settings] > [Minor Mode] > [Enable].

4.3.3 Empirical Material and Analysis Process

This study relies on two main types of empirical material, enabling a discourse comparison between companies' self-presentation and public interpretation. To ensure the quality and rigor of the data, this study adheres to the four quality criteria proposed by Scott (1990) and synthesized by Flick (2023): authenticity, credibility, representativeness, and meaning. The data were collected from reliable sources.

The material for research question 1 consists of publicly available texts produced by AI companies that offer or promote age-regulation tools, such as “Minor Mode” features. These materials are from their official websites. They are suitable for analysis because they illustrate how organizations strategically present and frame their roles and responsibilities regarding minors.

ChatGPT first mentioned its AI for minors in August 2025 (OpenAI, 2025a). Since that time, the company has released several upgraded features. For the purposes of this study, three documents were selected: the Teen Safety Blueprint published in November 2025, and two

parental guidance documents — *Tips for Talking to Your Teen About AI* and *A Family Guide to Help Teens Use AI Responsibly*. The Teen Safety Blueprint is a comprehensive document developed by OpenAI that clearly explains the reasons for introducing the minor-friendly feature, the measures taken to protect and optimise the experience for minors, and the company's vision for the future (OpenAI, 2025b). The two parental guidance documents supplement this with practical, operational information on how parents can discuss AI use with their children and manage their engagement with the platform (OpenAI, 2025c; OpenAI, 2025d). After the publication of the Blueprint, OpenAI did not release any new comprehensive documents, only feature upgrade notices that were already outlined in it. Together, these three documents provide a solid basis for analysis.

Gemini, Grok, and Doubao each have only one official statement regarding their AI for minors, so these were selected as the analysis material. Both Gemini and Doubao provide a full page guiding guardians on how to ensure proper use of their AI tools by minors, while Grok includes a single statement in its Terms of Service explicitly specifying the minimum user age and the key considerations users should be aware of (xAI LLC, 2025).

This study conducted a reflexive thematic analysis of the five corporate documents identified above, guided by the three analytical dimensions established in the theoretical framework: functional orientation, responsibility attribution, and normative positioning. The analytical process is illustrated in Figure 1.

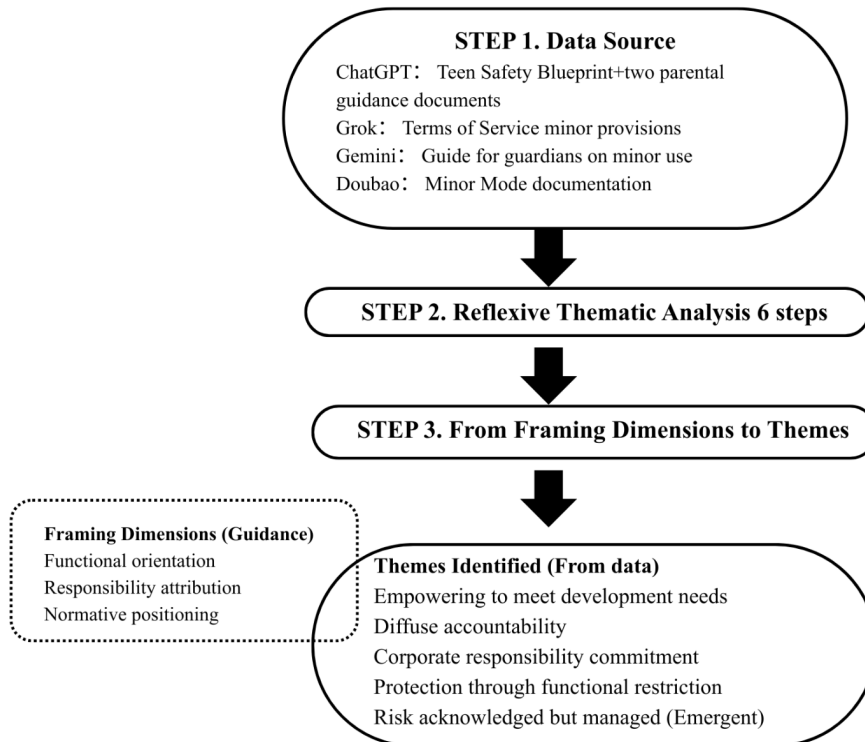


Figure 1: RQ1 Data Analysis Process

To capture the multiplicity of voices theorised within the rhetorical arena framework, this study examines a diverse range of online sources and analyses the communicative dynamics within each arena (Frandsen & Johansen, 2022). The second type of data consists of: 1) mainstream media news articles; 2) non-governmental organisations' reports and public discourse; 3) independent commentary; and 4) professional technology websites covering AI applications for minors. As RAT acknowledges, not all voices carry equal weight within the arena (Raupp, 2019). The materials collected in this study represent those voices that have gained visibility in the public sphere, rather than a comprehensive account of all public opinion. These data are used to address Research Question 2. These materials provide insight into publicly circulating concerns, evaluations, and interpretations surrounding AI for minors. For companies based in the U.S., the data were primarily collected from major news platforms such as The New York Times, The Guardian, and Los Angeles Times, as well as from non-governmental organisations such as Common Sense Media and FairPlay, which focuses on

supporting children's well-being in the digital environment. Opinion was also gathered from experts and industrial workers via independent commentary websites and professional technology websites. For the Chinese company, the data were primarily collected from mainstream news platforms, supplemented by articles from WeChat, one of the most widely used social platforms in China. The materials were collected by searching keywords in the format of "[AI model name] + minors/kids/teens/ kids mode, etc." Despite the personalisation of search engine results, the impact of personalisation on the findings was mitigated by broadening the scope of the keyword search and including a larger number of results.

Through the keyword search strategy described above, an initial pool of 89 materials was collected. The two researchers then conducted multiple rounds of detailed screening based on Scott's (1990) four quality criteria, resulting in the final dataset for RQ2, which contains 32 materials. Following this, the materials were analysed using the six phases of reflexive thematic analysis (Braun & Clarke, 2006), from which the themes were developed. The full selection and analysis process is illustrated in Figure 2.

The thematic coding tables for RQ1 and RQ2 are included in Appendix A and Appendix B, respectively. A complete overview of all data sources used in this study is provided in Appendix C.

Given the rapid release of AI guidance documents, new policy documents, and new public discourse may have been published after the researchers completed their search. To minimize this risk, the researchers continuously monitored relevant discourse while conducting our data analysis, up through April 23, 2026. All data were collected from open-access sources, including company websites and public media platforms. Therefore, no special permission was required.

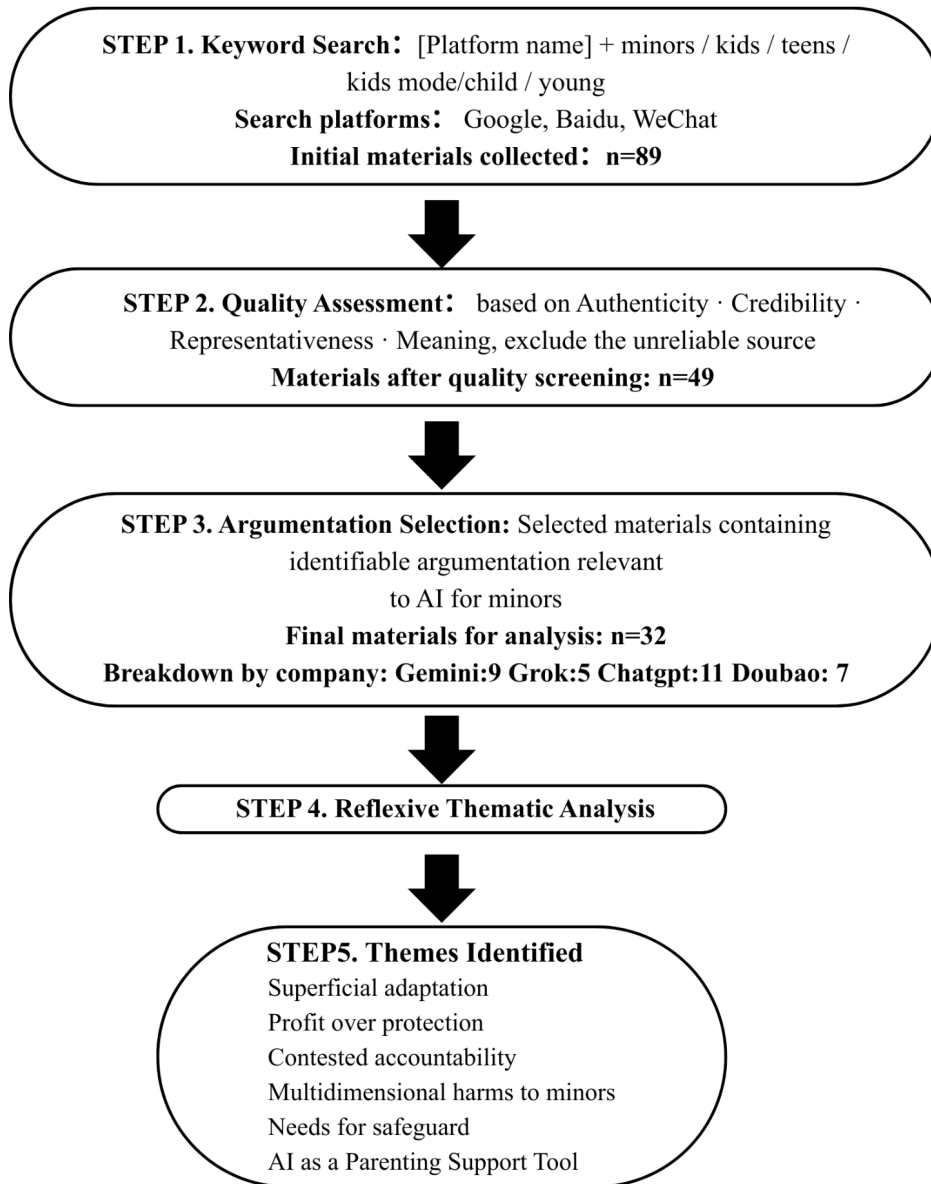


Figure 2: RQ2 Data Selection and Analysis Process

4.4 Reflexivity and Research Ethics

This study acknowledges several limitations. First, the geographical sampling is asymmetrical, comprising three U.S.-based companies and one Chinese influential one. However, due to limited data availability, at the time of this research, Doubao was the only major

Chinese generative AI platform with a specific documented "Minor Mode". Therefore, this study adopts an exploratory approach rather than a strictly balanced comparative design. Second, the analysis is based solely on publicly available communication materials, without interviews with users or employees. Consequently, this study focuses on public discourse and may not capture internal perspectives or individual users' interpretations. Finally, while this study defines "minors" within the broad 0–18 age range to ensure cross-national comparability, it lacks sufficient detail to examine specific stages of development (e.g., early childhood vs. adolescence). The focus remains on the features and narratives designed for minors. Despite these limitations, this study provides an overview of the current state of discussion regarding AI models for minors in China and the U.S. and lays the groundwork for future research on this topic.

In terms of research ethics, this study is based exclusively on publicly accessible documents and does not involve human subjects or private data. The study was not subject to formal institutional ethical review. In this study, ethical conduct, authenticity, and accuracy were considered throughout the research process. To support credibility, the study maintains transparency in its research design, data selection, coding process, and analytical decisions. The research process, the formation of the research question, the rationale behind the selection of specific methods, the procedural steps, and the decisions that influenced the production of data and results are all concepts that can be comprehended in the broadest sense (Flick, 2023).

5. Findings

Through repeated reading and reflexive thematic analysis of the empirical materials, this study identified five themes in response to RQ1 and six themes in response to RQ2. Figure 3 provides an overview of how these themes are distributed across the two rhetorical arenas. The following sections present the findings in the order of the research questions: RQ1 examines how AI companies frame AI for minors, RQ2 examines how AI for minors is discussed in public discourse, and RQ3 brings the two together to analyse patterns of alignment and tension across the U.S. and Chinese contexts.

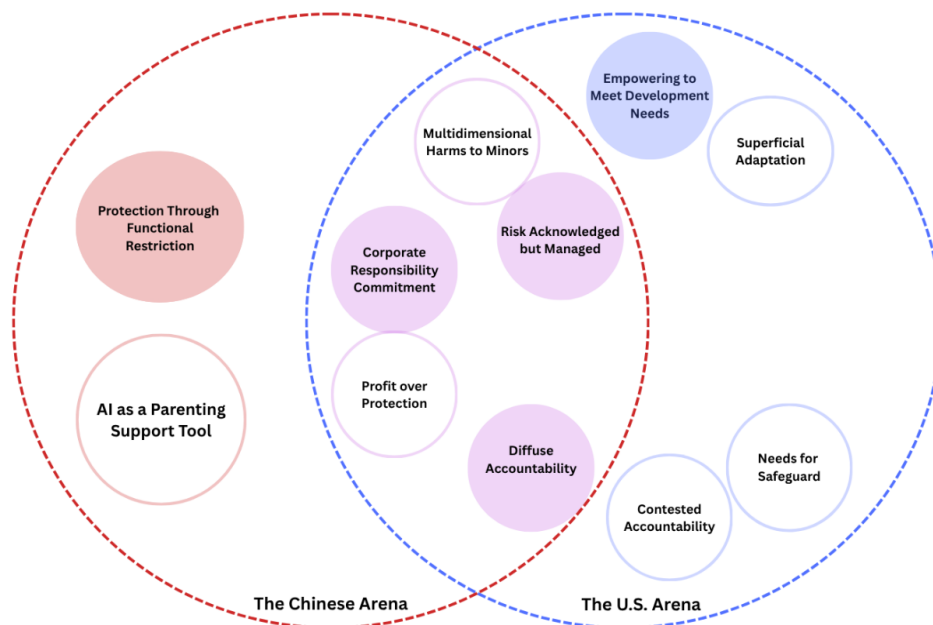


Figure 3: Thematic map: Corporate framing and public discourse themes across the U.S. and Chinese arenas

5.1 RQ1: Corporate Framing of AI for Minors

Following the thematic analysis, five themes were identified that capture how AI companies construct AI for minors in their external communication. These themes concern the developmental value attributed to AI, the distribution of responsibility, the management of risk, corporate responsibility commitments, and protection through functional restriction. Across the cases, these themes show how companies define what AI can offer minors, how risks should be managed, how responsibility is distributed, and what counts as appropriate use.

5.1.1 Empowering to Meet Development Needs

In the U.S.-based corporate materials examined, AI for minors is positioned as a means of meeting developmental needs beyond simply managing risk. Minor users are constructed as a group with distinct needs that AI can actively serve. The overall orientation is therefore both protective and enabling: engagement with AI is framed as something that should be supported under appropriate conditions.

The most explicit articulation of this logic appears in ChatGPT's communication, where access to AI is directly framed as a right. The platform states that "access to safe and trustworthy AI should be seen as a right, being a revolutionary technology that will help people unlock their potential and shape their future" (OpenAI, 2025b). Rather than presenting access as optional, this framing positions AI as something young people should be able to engage with. Exclusion from AI is implicitly positioned as undesirable, while access is linked to young people's future development. This is reinforced by a second formulation that situates teens as "the first generation to come of age in the Intelligence Age," for whom access to AI "at home, at school, and as they prepare to join the workforce" is framed as both timely and necessary (OpenAI, 2025b). The developmental rationale here is explicitly forward-looking: AI is not only presented as useful currently but also as relevant to young people's preparation for adult life.

This construction is further developed in how ChatGPT frames the relationship between age and experience. Rather than positioning youth as a reason to restrict engagement, the platform uses it as a reason to tailor it: "Teens are growing up with AI but aren't grown-up yet.

We believe ChatGPT should meet them where they are" (OpenAI, 2025b). The phrase "meet them where they are" positions teenagers as users whose needs should be responded to directly, rather than managed only around restriction. The platform goes further to argue that the teen version should be "a Teen Edition, rather than a castrated regular version" (OpenAI, 2025b). This explicitly rejects a reduction-based approach to minor-specific design and positions a purpose-built teen experience as a better fit for young users.

Gemini develops the developmental logic in a different direction, foregrounding the creative and educational potential of AI engagement rather than rights-based language. The platform describes the technology as enabling children to "learn and be creative," illustrating this through a range of activities: asking questions about the natural world, writing stories, composing songs, brainstorming for art projects, and working through academic problems with guided support (Google, 2025). Compared with ChatGPT, the emphasis here is less on access as a right and more on AI as a practical support for learning and exploration. This frames AI as more than something children are permitted to use; it becomes something that can enrich their intellectual and creative lives. The inclusion of creativity alongside academic utility extends the developmental rationale beyond instrumental learning goals to broader capacities for expression and imagination. The tone here is inviting rather than regulatory, constructing AI engagement as an exploratory activity.

Parental control features are also framed within this empowerment logic. Both ChatGPT and Grok present personalised parental control as part of their approach for minor users, framing it as a mechanism that enables appropriate access rather than one that simply restricts it. ChatGPT describes parental controls as providing "the tools and flexibility they need to support their teen's development in a way that works best for their family," listing customisable options including model behaviour rules, privacy settings, blackout hours, and feed adjustments (OpenAI, 2025a). Grok similarly frames its controls as enabling parents to "select the appropriate features and limitations for their needs" (xAI LLC, 2025). In this part of the corporate framing, parental involvement supports the possibility of access. It is not presented only as a barrier, but as a way of making AI use more suitable for individual family circumstances.

5.1.2 Diffuse Accountability

Across the corporate communications examined, responsibility for the safe and appropriate use of AI by minors is distributed across multiple parties, including parents, minors themselves, and the platforms. This theme captures how accountability becomes diffuse in corporate communication: companies present themselves as providing tools, safeguards, settings, and guidance, while much of the practical governance of minors' AI use is located with families and users.

In all four platforms, parents are consistently constructed as the primary governance actor. Parental control features appear in ChatGPT, Gemini, Grok, and Doubao alike, but in each case, these features are framed as more than product functionality; they become a mechanism through which families assume active governance over minors' AI use. Gemini instructs parents to "explore together" with their children and to "supervise their use" as children learn how to engage with generative AI (Google, 2025). Grok states that parents and guardians "must agree to our Terms of Service" on behalf of their teenage children and urges them to "exercise care in monitoring the use of Grok by their teenagers" (xAI LLC, 2025). ChatGPT provides an extensive parental control toolkit with a dedicated guide on "tips for talking to your teen about AI" (OpenAI, 2025d). These examples position parental engagement as more than technical oversight. Parents are constructed as active participants in shaping the minor's experience, and their involvement is framed as necessary rather than optional.

Beyond parental supervision, the data also constructs minors themselves as agents who bear some responsibility for their own conduct. Teenagers and children are not positioned solely as subjects of protection but as users capable of learning responsible engagement. Gemini explicitly addresses parents with instructions to "remind your child not to enter personal information" and to "help your child think critically about Gemini Apps responses" (Google, 2025). This constructs the child as someone who can be taught to recognise limits, evaluate responses, and participate in safer use. ChatGPT's teen-facing communication similarly frames young users as users that should be informed, guided, and prepared for AI engagement rather than simply shielded from it (OpenAI, 2025c). In this way, responsibility is extended beyond

parents to include minors themselves, which further diffuses accountability away from the platform alone.

The same pattern appears in how platforms position themselves in relation to responsibility. Across all four companies, the platform's role is framed less as a guarantor of outcomes than as a provider of tools, settings, and guidance. ChatGPT frames its parental controls as providing families with "tools and flexibility" (OpenAI, 2025c), while Grok describes its settings as enabling parents to "select the appropriate features and limitations for their needs" (xAI LLC, 2025). In both cases, the platform provides the infrastructure for governance, but the actual exercise of that governance is located with the family. This distinction is important because it allows companies to present themselves as supportive while still placing basic responsibility outside the platform.

This tool-provider framing is reinforced by how platforms address the boundaries of their own protective role. Both Gemini and Grok acknowledge that platform-level protections have limits, with Gemini noting that children "may encounter content you don't want them to see" despite existing filters (Google, 2025), and Grok similarly indicating that outputs "could" be inappropriate depending on how features are used (xAI LLC, 2025). These acknowledgements do not position the platform as fully responsible for preventing all risks. Instead, they reinforce parental supervision as a necessary complement to platform tools. The communicative logic is that platform protections reduce but do not eliminate risk, while the remaining responsibility is positioned with parents and guardians.

ChatGPT adopts a related but distinct strategy. It frames its protective measures as continuously evolving and research-informed rather than complete or guaranteed. The platform describes its teen protections as "informed by conversations with policymakers" and external experts, and commits to ongoing collaboration with "schools, teachers, and researchers" to improve the teen experience over time (OpenAI, 2025b). In this theme, the language does more than present ChatGPT as responsible; it frames responsibility as procedural and ongoing. The platform does not claim that current protections are final. Instead, responsibility is distributed across tools, expert input, future improvements, and collaborative processes. This keeps the

platform within the responsibility chain, but without positioning it as the sole actor responsible for all outcomes.

The most explicit articulation of diffuse accountability appears in Doubao's communication, where the limits of platform responsibility are stated directly rather than implied through framing. Doubao's minor mode documentation states that "the platform cannot and will not change the guardian's settings," and that guardians "shall bear the relevant consequences that may arise" from enabling minor mode (ByteDance, 2025). This locates the consequences of restriction with the family rather than the platform. A further clause specifies that "the activation of minor mode does not constitute an exemption or reduction of liability for account usage and related conduct" (ByteDance, 2025), making clear that responsibility remains with the user regardless of the protective measures in place. Most explicitly, the platform states that it "is not responsible for events that occurred in the real world" (ByteDance, 2025), a formulation that explicitly delimits the boundary of platform accountability.

A difference emerges between the U.S.-based platforms and Doubao in how accountability is distributed. In ChatGPT, Gemini, and Grok, responsibility is distributed mainly through communicative framing: parents are invited, encouraged, and equipped to take ownership of the minor's experience, and the platform's role is framed as supportive. The diffusion of responsibility is softer, embedded in the framing of parental control, flexibility, and guidance. In Doubao, the shift is more explicit and contractual. Responsibility is formally assigned to the guardian through the terms of the minor's mode agreement, and the platform's non-responsibility is stated as a legal condition rather than implied through communicative tone. Both approaches produce diffuse accountability, but through substantially different mechanisms: one embedded in guidance and support, the other articulated through formal terms and conditions.

5.1.3 Risk Acknowledged But Managed

Across the corporate communications examined, risk is neither denied nor treated as a reason to restrict access. Instead, platforms acknowledge that minors may encounter harmful, inappropriate, or unreliable AI outputs, while framing these risks as conditions that can be

identified, contained, and responded to through existing measures. In this sense, risk is presented as real, but not as uncontrollable. This theme, therefore captures a central balance: companies recognise the need for protection, but they also present risk as something that can be managed without undermining minors' access to AI.

The most visible expression of this logic is the way platforms pair risk acknowledgement with concrete protective responses. ChatGPT describes a range of detection and intervention mechanisms already in place, including "safeguards that direct users to real-world resources if we detect they have expressed suicidal intent" and systems to "escalate risks of physical harm to others for human reviewers who can take action" (OpenAI, 2025b). The platform also reports submitting more than 75,000 cybertips to the National Center for Missing and Exploited Children in the first half of 2025, presenting this as evidence of active and measurable protection (OpenAI, 2025b). These examples give risk management a concrete and institutional form: risks are recognised and attached to detection systems, reporting mechanisms, and human review processes. Gemini similarly describes content filters designed to "identify and block harmful content, avoiding certain responses that violate our policy guidelines like violence, sexually explicit material, and harassment" (Google, 2025). In both cases, risk acknowledgement is immediately followed by what the platform is already doing. This constructs risk as a known condition that can be acted upon, rather than as an open-ended threat.

However, how risk is defined varies across platforms. ChatGPT frames risk as systemic and continuously monitored. Threats are categorised, tracked, and subject to ongoing policy development informed by external experts and policymakers (OpenAI, 2025b). The platform also commits to testing protective policies before deployment and monitoring them after, positioning risk management as a continuous institutional process rather than a fixed solution (OpenAI, 2025b). This makes risk appear as something embedded in the operation of the platform, requiring repeated assessment rather than one-off moderation.

Gemini and Grok take a different approach. Both acknowledge that filters are imperfect and that inappropriate content may occasionally appear, but neither elaborates on the nature or scale of that risk beyond noting its possibility (Google, 2025; xAI LLC, 2025). Risk in this

framing is real but limited. It is presented as something that may occur despite safeguards, rather than as a systemic problem requiring a detailed institutional response. Doubao's construction is different again. Its minor mode documentation focuses on usage addiction and real-world consequences rather than content safety, advising guardians to ensure minors avoid "becoming overly immersed in virtual conversations" and to help verify information obtained from the platform before it informs important decisions (ByteDance, 2025). Here, risk is framed less as exposure to harmful content and more as behavioural or practical harm arising from how the tool is used. The concern is what the AI might generate and how minors might rely on it in everyday life.

A further aspect of this theme concerns how platforms set their default conditions for minor users. ChatGPT states that where a user's age cannot be determined, the platform "will default to the U18 experience," adding that "this might hinder some adults' experience using free products, but we are prioritising kid and teen safety" (OpenAI, 2025b). By naming this trade-off directly, ChatGPT frames protection as a deliberate design choice. Safety is placed before the possibility that some adult users may receive a more restricted experience.

Doubao operates differently. Its documentation specifies that if a guardian disables minor mode, the minor "can use all of Doubao's features" and data handling reverts to standard adult terms (ByteDance, 2025). Protection here is not the baseline but an optional condition activated by guardian choice. The difference is then visible in the starting point of each platform. In ChatGPT's framing, risk is assumed to be present until age is confirmed, so the experience would be safer if applied by default. In Doubao's framing, full functionality remains available unless a guardian activates minor mode.

These default configurations show different ways of managing risk in practice. ChatGPT constructs protection as the baseline condition when the user's age might be under 18, while Doubao positions protection as a user-activated mode. This creates different conditions for minors' engagement with AI: one begins from precaution, while the other begins from access unless restriction is chosen.

5.1.4 Corporate Responsibility Commitment

Across the corporate communications examined, companies vary in how explicitly they construct themselves as responsible actors in relation to AI for minors. This theme concerns how responsibility is presented through specific safeguards and broader forms of self-positioning. In some cases, this self-positioning is highly developed and supported through references to leadership, expert advice, and external collaboration. In others, it is much less visible. This unevenness is important because it shows that corporate responsibility is not framed with the same intensity across all platforms.

The most developed form of this responsibility framing appears in ChatGPT's communication. Here, safety is presented not only as a product concern, but as a matter of principle. ChatGPT's Teen Safety Blueprint opens with a declaration attributed to the company's CEO: "for teens, we prioritize safety ahead of privacy and freedom" (OpenAI, 2025b). By placing this statement in the CEO's voice, the communication gives prioritisation a leadership dimension. It is presented as a company-level stance rather than simply as a technical feature. The platform also positions itself within a broader industry context, describing its CSAM prevention infrastructure as "industry-leading" and calling on other AI companies to "implement commensurate protections for teens" (OpenAI, 2025b). This presents ChatGPT as taking responsibility for its own product and positioning itself as a standard-setter within the AI field.

A second recurring element is the framing of responsibility as continuous and evolving rather than complete. ChatGPT commits to ensuring that "kids and teen safeguards will be open and guided by research," to supporting independent research, and to seeking "ongoing feedback from teens, parents, educators, and experts" as it builds and improves (OpenAI, 2025b). Well-being features are described as "informed by the latest research, guided by external experts, and updated as the industry learns more about best practices" (OpenAI, 2025b). This language presents protection as work in progress. Responsibility here is not a condition that has been achieved, but an ongoing process of improvement. This framing also leaves space for uncertainty: the platform can acknowledge that safeguards need development while still presenting itself as actively engaged in that development.

A third element is the use of external relationships to support responsibility claims. ChatGPT names specific external actors as sources of guidance, including "policymakers, including state attorneys general," experts who have joined an "Expert Council on Well-Being and AI," and ongoing collaboration with "schools, teachers, and researchers" (OpenAI, 2025b). The platform also commits to working with "child safety advocacy groups" and "parent, teacher, and community organizations" (OpenAI, 2025b). By naming these relationships, the platform presents corporate responsibility as externally informed rather than purely self-declared. The company is not simply claiming to care about teen safety; it positions that care as informed, accountable, and subject to external scrutiny. It is attached to networks of expertise, policy input, and community engagement.

Doubao's responsibility framing operates through a different register. Rather than expert collaboration or research-informed development, Doubao grounds its responsibility in regulatory compliance. The platform commits to "continuously optimising the functional experience of minor mode" and to providing "higher quality and healthier content" for minors. It also states that feature adjustments will be made "based on laws and regulations," and that data practices in minor mode are governed by formal legal requirements (ByteDance, 2025). Responsibility here is constructed as adherence to an external legal framework. The platform presents itself as compliant and improving, but the source of authority is regulatory rather than collaborative. Compared with ChatGPT, the emphasis is less on leading or shaping best practice and more on operating within established rules.

Gemini and Grok offer a more limited version of this responsibility framing. Neither platform uses much language of commitment, continuous improvement, or external endorsement in relation to minor users. Their communications describe what protections exist, but they do not construct a broader narrative of corporate responsibility around those protections. This absence becomes visible when compared to the explicit self-positioning present in ChatGPT and to a lesser extent, Doubao. The point is not that Gemini and Grok lack safeguards, but that their communication does not make corporate responsibility a central theme in the same way.

These differences show that corporate responsibility is constructed through different communicative routes. ChatGPT builds responsibility through expert networks, research commitments, and industry leadership claims. Doubao builds it through legal compliance and formal regulatory adherence. Gemini and Grok give this question less explicit attention. Together, these differences suggest that responsibility is not simply present or absent, but framed with different degrees of explicitness and through different sources of authority.

5.1.5 Protection Through Functional Restriction

In Doubao's corporate communication, protection for minors is framed primarily through functional restriction. Unlike the empowerment-oriented framing found in the U.S.-based corporate materials examined in this study, Doubao does not mainly present AI for minors as a tool for learning, creativity, or open-ended development. Instead, it constructs minor use as something that should take place within a more limited and guided version of the platform. The platform states that minor mode is designed to guide minors toward positive and healthy use, protect their personal information, and provide guardians with more management tools (ByteDance, 2025). Protection is thus more than a general safety claim; it is built into the conditions and boundaries of use.

This restriction-based logic first appears in how Doubao configures access. Children under 14 require explicit guardian consent to use the general mode, while users aged 14 to 18 are given more choice but are still recommended to use minor mode (ByteDance, 2025). This distinction recognises differences between younger children and teenagers, but it does not position older minors as ordinary adult-like users. Even for teenagers, the protected mode is framed as the more suitable setting. In this sense, protection is not constructed through excluding minors from AI, but through directing them toward a more controlled version of the service.

The clearest expression of this theme is the reduction of platform affordances. When minor mode is enabled, Doubao states that some functions will be limited or disabled, including access to external agents, third-party web links, related video recommendations, AI creation tools, and some real-time call scenarios (ByteDance, 2025). The significance of these restrictions is not simply that certain features are unavailable. Doubao defines the minor version through a

narrower range of possible actions and encounters. A safer AI environment is produced by limiting what minors can access, generate, and interact with.

The theme also appears in how Doubao frames appropriate use. The platform links minor mode to “positive and healthy” use and states that minors should be guided to understand the limits of AI, treat generated content cautiously, and verify information before relying on it in important or sensitive situations (ByteDance, 2025). This does not only concern what the platform blocks. It also concerns what kind of AI the platform presents as acceptable. Appropriate use is framed as cautious, limited, and directed away from over-reliance. Protection, therefore works through reduced functions and by defining the proper boundaries of engagement.

Overall, Doubao constructs AI for minors as a bounded version of the platform. The minor mode is made protective through three connected forms of limitation: access is configured by age, platform affordances are reduced, and appropriate use is framed as cautious and healthy. This theme is distinct from the broader framing of risk as manageable because it shows how suitability for minors is constructed through narrowing the AI experience itself. In Doubao’s framing, AI becomes more appropriate for minors because it offers less, exposes them to less, and defines acceptable use within a more controlled environment.

5.1.6 Synthesis of RQ1 Findings

The findings show that AI companies frame AI for minors as both beneficial and in need of control. Across the dataset, AI is positioned as a developmental resource that can support minors’ learning, creativity, and future development. At the same time, companies do not present minors’ AI use as risk-free. Instead, they frame risks as manageable through parental involvement, technical safeguards, policies, and platform settings. Companies also construct their own role in different ways, including through explicit commitments, external validation, regulatory compliance, or limited articulation of responsibility.

Overall, the RQ1 findings show that corporate communication constructs AI for minors as a governed form of access rather than as a freely available technology. The key difference lies in how this governance is justified. In the U.S.-based materials examined, governance often

works to enable access: safeguards, parental controls, and age-sensitive design are used to make AI participation acceptable. In Doubao's communication, governance works more through limitation: the platform becomes appropriate for minors because it offers less and defines a narrower range of acceptable use. This difference matters because it shows that companies do not define appropriate AI use for minors in the same way. This variation provides the basis for examining how public discourse later responds to these corporate constructions of AI for minors.

5.2 RQ2: How is AI for minors discussed in public discourse?

Following the thematic analysis, six themes were identified in public discourse surrounding AI for minors: Superficial Adaptation, Contested Accountability, Multidimensional Harms to Minors, Profit over Protection, Needs for Safeguarding, and AI as a Parenting Support Tool. Compared with corporate communication, public discourse is generally more critical and questioning, although the intensity and focus differ across the U.S. and Chinese contexts. In the American material, criticism is more systemic and directed toward developmental adequacy, corporate accountability, product design, safeguarding, and rights-related concerns. In the Chinese material, concerns are present but more often focus on specific functional issues and practical use, while a visible strand of discourse positively constructs AI as a tool for parenting and learning. These patterns of public discourse can be understood through the lens of RAT as the arena's responsive voices contesting, negotiating, and at times reinforcing the framings put forward by corporate actors, while recognising that actors have unequal communicative power across contexts. The findings are presented below through six themes.

5.2.1 Superficial Adaptation

The first theme identified was that existing AI products do not adequately address the developmental needs of minors. Several commentators have noted that technology companies frame their products as specifically designed for minors, claiming to cater to minors' developmental needs; yet in practice, they often appear to make superficial modifications to products that were originally designed for adult users. This critique is chiefly directed at ChatGPT and Gemini.

A study by Common Sense Media (2025a, 2025b) on Gemini and ChatGPT points out that both platforms lack age-appropriate adjustments for underage users. Specifically, they use the same response process when interacting with minors of different ages. For example, a 5-year-old and a 12-year-old might receive the same response when interacting with the AI, which may not be the most suitable for either child (Common Sense Media, 2025a). Furthermore, the Gemini U13 model was designed to provide help and detailed answers, consistent with the design logic of the adult version. However, content that is “useful” to adults is not always suitable for children: the system’s emphasis on “usefulness” leads to a large amount of detailed information in response to children’s questions, often resulting in text that is too complex and lengthy for younger users. Common Sense Media (2025a) recommends that parents and guardians consider using alternative digital content designed specifically for children, such as curated educational websites, libraries, and expert-reviewed media, which are more in line with the developmental needs of minors. The current discussion highlights a broader design problem: neither platform adequately considers the cognitive and emotional needs of users across different age groups, resulting in product design that is misaligned with the developmental stages of its target users. Gemini has also faced further criticism for its adult-oriented design logic, which deepens this mismatch.

In contrast, Grok and Doubao are largely absent from this specific public discourse. In the U.S. arena, while Grok faces significant criticism, this is primarily focused on harmful content and profit-driven business models rather than developmental adaptation. This may reflect either limited investment in minor-specific features or a shift in public attention toward more pressing concerns. In the Chinese arena, Doubao's minor mode is essentially a restricted version of the original system, as stated in the official product introduction. However, this superficial adaptation has not triggered substantial criticism or sustained public discussion, which may reflect the general public opinion in China that if a product actually matches its advertising, it doesn't need discussion, or that the Chinese public accepts the restrictions made for safety reasons, or that compliance with relevant regulations is sufficient to prove their applicability.

Public discourse in the U.S. arena consistently identifies a gap between companies' claims of developmental responsiveness and the actual design logic of their products. This

adaptation is criticised as superficial because it continues to follow the design logic of adult-oriented products. Grok and Doubao remain largely absent from this discussion, though for different reasons across the two arenas. This pattern suggests that public focus and forms of discussion differ across the two arenas. In the U.S., advocacy organizations and the media feel a responsibility to scrutinize and comment on the design of minors' products by technology companies, creating a more concentrated critical discourse. In China, however, the focus of public discussion does not seem to be on this aspect, so similar challenges have not emerged.

5.2.2 Contested Accountability

In public discussions, the focus on responsibility has primarily been on Gemini, while discussions surrounding ChatGPT, Grok, and Doubao have been relatively less frequent. In both the U.S. and China, the public emphasizes that AI cannot replace parental companionship, especially considering the risks of addiction and over-reliance. However, in China, this does not necessarily constitute a direct challenge to companies attempting to deflect or shift responsibility. The Chinese public views parental involvement as a natural and necessary supplement to supervising and protecting minors' use of AI, and as an inherent obligation (陆玖商业评论, 2025). In the U.S., the debate on responsibility is more direct and complex.

Fairplay and EPIC (2025), two U.S.-based child advocacy organizations, formally urged the Federal Trade Commission (FTC) to investigate Google's potential violations of the Children's Online Privacy Protection Act (COPPA) in relation to its Gemini AI companion for children under 13. Their letter explicitly states that "Google's announcement disclaimed any responsibility to mitigate these harms or implement safeguards around its own product. Instead, Google unfairly shifts all responsibility onto parents." Both organizations argue that any company choosing to sell products to minors bears a direct responsibility to ensure the safety and suitability of those products, and that Google's attempt to shift this responsibility to parents and guardians fails to meet this standard. This structural criticism resonated in broader media commentary: a Forbes contributor noted that generative AI's shift of management responsibility to parents has become a disturbing but increasingly prevalent reality (Spehar, 2025). Robbie Torney makes a similar point, arguing that the parental controls currently used by several AI

companies are "difficult to set up" and "shift the responsibility back to parents." (Hill, 2025). This type of public discourse reflects a wider concern in the U.S. material that if a company provides services to minors, the primary responsibility lies with that technology company. Any attempt to shift or deflect responsibility is seen as an evasion of accountability and is considered unacceptable.

These findings reveal different ways in which responsibility for minors' AI use is debated in the two national arenas. In the U.S., the debate primarily revolves around structural issues: advocacy groups and the media explicitly question the legitimacy of shifting responsibility, arguing that companies selling products to minors have an on-delegable responsibility to ensure minors' safety, a responsibility that cannot be transferred through default settings or disclaimers. In China, both businesses and the public recognize the irreplaceable responsibility of parents in guiding children's use of AI, which does not constitute a direct challenge to attempts by businesses to shift responsibility. Parental responsibility is seen as a necessary complement to, rather than a replacement for, corporate and regulatory obligations. These differing views on the attribution and allocation of responsibility reflect different assumptions about responsibility, shaped by the political and cultural contexts of the two countries.

5.2.3 Multidimensional Harms to Minors

Despite the widespread recognition of AI's potential benefits in educational contexts, endorsed at the policy level in the U.S. (The White House, 2025) and echoed in parental accounts of AI's capacity to foster creativity and personalised learning experiences (Wong, 2025), public discourse across both arenas reveals sustained concern about the harms AI may pose to minors. This theme generates the most sustained and wide-ranging debate in the data, spanning emotional, cognitive, informational, and privacy-related concerns.

In the U.S. arena, the most prominent concern centres on emotional attachment. A broad coalition led by Fairplay and EPIC (2025) has campaigned against Google's rollout of the Gemini AI companion to young children, arguing that minors struggle to distinguish between AI chatbots and humans, making them particularly susceptible to misplaced trust. This concern is evident in discussions surrounding all three AI platforms in the U.S. CNN (Duffy, 2025) and The

Guardian (Wong, 2025) both reported that teenagers using ChatGPT perceived the AI's responses to their emotional expressions as highly human, leading them to sometimes believe they were building a genuine relationship with the AI. Similarly, Grok's companion features were found to position the AI as a more empathetic, approachable, and trustworthy partner than the real people in the user's life (Common Sense Media, 2026). AI platforms are found to attempt to cultivate a genuine connection, and teenagers have limited resistance to this temptation. Critics argue that AI platforms may cultivate a sense of genuine connection, while teenagers may have limited capacity to resist this appeal. UNICEF, the United Nations' children's agency, and other children's groups have noted that the A.I. systems could confuse, misinform, and manipulate young children who may have difficulty understanding that the chatbots are not human (Singer, 2025). Mathilde Cerioli, chief scientist at Everyone.AI, a nonprofit focused on developing ethically sound AI for teenagers, points out that teenagers with limited social experience and a greater sense of loneliness are more easily drawn to chatbots. "They are already facing more challenging situations, and chatbots will further exacerbate their predicament," she says (Hill, 2026). As public critics have pointed out, the danger lies in the fact that minors may trust these systems as much as they would trust friends, but AI fundamentally lacks the ability to understand the consequences of its responses or a minor's context, including the minor's mental health history, cultural background, and family situation (Common Sense Media, 2025b). Instead of being based on developmental psychology or personalized understanding, these systems rely on algorithmic predictions to generate statistically probable responses that may be completely unsuitable for vulnerable young users. This increases the likelihood of minors being harmed in such unequal relationships.

Public discourse is not uniformly critical on this point. Testing reported by Common Sense Media (2025a) indicates that Gemini does, in certain instances, refuse inappropriate interactions and remind users that it is not human, suggesting that some baseline safeguards are functional, even if critics argue they remain insufficient and inconsistently applied.

Beyond emotional attachment, public discourse in the U.S. arena identifies a cluster of cognitive and informational harms. Critics argue that AI use among minors risks diminishing the diversity of creative expression, undermining the development of critical thinking skills, and

exposing young users to misinformation and reinforced stereotypes. Since AI systems mirror the biases present in society, minors may encounter the same kinds of discrimination they face offline, potentially contributing to a distorted worldview (Duffy, 2025; Common Sense Media, 2025b; Spehar, 2025). Furthermore, privacy protection is another recurring concern, particularly in relation to Gemini. Advocacy groups have called on Google to implement additional safeguards to ensure that data collected through Gemini is not repurposed for internal uses or shared with third parties without appropriate consent (Fairplay & EPIC, 2025).

In China, public concerns about emotional attachment and exposure to misinformation also exist. Doubao's short video function has sparked controversy. Many parents worry that minors will become addicted to watching videos, comparing Doubao's short videos to a built-in Douyin (周小白, 2025). Criticism of Doubao's short video function reflects anxieties about entertainment replacing productive use and the potential for minors to become addicted; the public also expresses concerns about AI easily generating misinformation that minors may have difficulty identifying (张格格, 2025; 一颗青木, 2026). Although the Chinese public also acknowledges similar concerns about the potential harm of minors using AI, the discussions are mostly focused on specific functions and are less detailed and sustained than those in the U.S. Across these concerns, public discourse links harm to the relational and design conditions of AI systems. Human-like interaction, limited contextual understanding, informational unreliability, and prolonged engagement are presented as features that can make minors especially vulnerable.

5.2.4 Needs for Safeguard

In the U.S. arena, public discourse identifies the harms associated with AI for minors and diagnoses the structural vulnerabilities in existing safeguard mechanisms and proposes standards for improvement. In the Chinese arena, no comparable public debate around safeguard adequacy is evident in the data. The absence of this discourse is consistent with the structural characteristics of the Chinese rhetorical arena identified in the theoretical framework, where state regulation defines the boundaries of acceptable practice and compliance itself functions as the dominant measure of safety, leaving limited discursive space for public questioning of whether existing safeguards are genuinely sufficient.

Criticism of ChatGPT's safeguarding capacity centres on three interconnected failures. First, the platform has been criticised for weak age verification mechanisms and for not making protective features the default setting. As Toni Brunton-Douglas, a senior policy officer, argues: "vulnerable children could still be left exposed. Tech companies must not view child safety as an afterthought. It's time to make protection the default" (Booth, 2025). Second, testing by Common Sense Media (2025b) revealed that ChatGPT repeatedly failed to recognize concerning behavioural patterns that a human familiar with the teenager would readily detect, highlighting the platform's inability to apply contextual or relational understanding to safeguarding. Adam Ryan, a California teenager who died in April 2025, managed to bypass ChatGPT's security protocols by asserting that his intention was to utilise the information for the purpose of literary composition (Regalado, 2025). Third, even when potentially harmful interactions are identified, the platform lacks adequate mechanisms for crisis intervention or escalation prevention. These vulnerabilities are compounded by evidence that ChatGPT's safeguarding features become less reliable over extended interactions (Duffy, 2025). In addition, Jared Moore, a researcher at Stanford University, has highlighted that ChatGPT makes broad yet ambiguous promises without providing any mechanism for evaluation (Hill, 2025).

Similar structural weaknesses are identified across Gemini and Grok. Grok provides no strict age verification for adult content, meaning teenagers can access inappropriate material without restriction (Burgess & Varner, 2026). Unlike a human interlocutor who would register context clues, such as references to school, parents, high school relationships, or being underage, Grok continues to provide access to inappropriate content and risky advice even when such signals are present (Common Sense Media, 2026). Gemini faces parallel criticisms, including a failure to redirect minors to trusted adults, an opaque human review process, unclear internal data controls, and insufficient disclosure of minor-safety metrics (Common Sense Media, 2025a). These discussions suggest that these three U.S. platforms share common problems: fragile and inefficient protection mechanisms and unclear risk assessment mechanisms. They cannot accurately determine a user's age based on context, nor can they take appropriate measures when necessary. The effectiveness of protection mechanisms gradually weakens over prolonged interaction, and the public has no way of knowing how the platforms assess risks and

take protective measures. In addition to calls for improved corporate design, some U.S. public voices also call for stronger government regulation of AI platforms targeting minors. This reflects the view that industry self-regulation is insufficient to address the challenges and that regulatory intervention is necessary to enforce effective protection standards (Spehar, 2025; Pederson, 2025).

Public opinion in the U.S. has developed a set of expectations beyond regulatory compliance: AI companies should incorporate contextual intelligence, strengthen security safeguards, make protection the default rather than the exception, disclose risk assessment mechanisms, and demand that the government establish a robust accountability mechanism. However, in China, such standards have not yet emerged in the public sphere, and there is limited public discussion of safeguard adequacy. This indicates a difference in the two countries' assumptions regarding the responsibility for defining and enforcing the safety boundaries of AI systems for minors.

5.2.5 Profit over Protection

This theme can be found in public opinion in both China and the U.S., where the public questions how AI companies can balance commercial interests with claims about protecting minors.

In the U.S. arena, Andrew McStay, a professor of technology and sociology at Bangor University, argues that there is a fundamental problem with exposing minors to technologies based on large language models: "Parents need to realize that these technologies are not designed with children's best interests in mind" (Wong, 2025). A CNN Business article similarly argues that all AI tools are designed to be supportive and agreeable (Duffy, 2025), raising concerns about consequences such as AI companies prioritizing engagement over safety and chatbots encouraging continued use. This theme is reflected in public discussions regarding ChatGPT, Grok, and Gemini.

A Los Angeles Times article focused on the "freedom," deregulation, and empowerment of minors advocated by OpenAI co-founder Sam Altman, arguing that these were merely

superficial moral rhetoric designed to mask the platform's pursuit of accelerated growth and higher profits (Pederson, 2025). A Common Sense Media (2025b) assessment report noted that ChatGPT frequently added prompts at the end of replies to encourage further user interaction. For example, "Would you like me to guide you through a relaxation exercise right now?" This practice repeatedly pulls teenagers back into the conversation rather than guiding them to engage with trusted adults, especially when discussing mental health issues. The public interprets the platform's emphasis on providing unrestricted access for minors while encouraging continuous interaction as reflecting the profit-driven nature of AI platforms, inconsistent with their claimed minor-first principle. Stanford University researcher Jared Moore, who studied how ChatGPT addresses mental health crises, pointed out that the simplest approach when encountering disturbing conversations is to end them directly (Hill, 2025).

Common Sense Media (2026) points out that Grok, through its companion and gamified interactive features, encourages users to engage in continuous and potentially compulsive interactions and actively guides users to continue using the app through push notifications. Furthermore, instead of recognizing warning signs and guiding vulnerable teenagers to seek external support resources, Grok's companion features may reinforce users' emotional dependence during interactions. These measures are interpreted as encouraging continued use of Grok, rather than protecting minors. In addition to interaction issues, Grok has also been criticized for facilitating the creation of unconsented sexual imagery. Through its integration with X, users can generate unconsented digital images and edit others' photos, a practice that spread globally around New Year's Day 2026 (Common Sense Media, 2026). Following strong public outcry, xAI switched its image generation function to a paid subscription model, a decision widely interpreted as profit-driven rather than aimed at curbing the surge in image abuse (Grose, 2026). Critics, therefore, frame the response as commercially driven, with public controversy becoming part of the platform's visibility rather than a reason for stronger protection.

Gemini has been criticized for its default enabling mechanism. As one parent put it, "I received an email like this last week. Note that it didn't ask if I wanted my child to use Gemini; instead, it warned me that if I didn't want my child to use it, I had to opt out. This is

unreasonable” (Spehar, 2025). One commentator argued that Google’s decision to promote the Gemini AI chatbot to teenagers after receiving only one email from a parent was irresponsible (Christ, 2025). Furthermore, Robbie Torney, head of the Common Sense Media AI project, believes that Gemini’s parental controls are merely a stopgap measure, not a true solution for ensuring minors’ AI safety in the long run (Hill, 2025). In addition, commentators argue that Google’s Gemini AI for kids could lead to a generation becoming habitually dependent on Google’s AI ecosystem (TechRadar, 2025). Cultivating usage habits among minors also reflects concern over long-term commercial interests.

In the Chinese context, criticism of AI companies pursuing commercial interests also exists. Some parents have reported that when their children use Doubao for academic research, the platform displays video links below the answers. Clicking the links causes the videos to play in full screen and refresh automatically, effectively diverting children from learning to entertainment, and there is no option to turn off short video recommendations (周小白, 2025; 张格格, 2025). Industry insiders believe this design reflects Doubao's strategy of enriching content formats to maintain user engagement and compete with rivals like DeepSeek, prioritizing commercial interests over the protection of minors (周小白, 2025).

In both arenas, public opinion has criticized AI companies for prioritizing commercial interests over their stated commitment to protecting minors. However, the nature and scope of this criticism differ across contexts. In the U.S., concerns extend beyond specific product features to include structural issues and parental rights: critics question whether children's default access to AI platforms infringes on parents' rights to informed consent, and whether the design logic of these systems fundamentally violates the best interests of minors. In the Chinese arena, criticism is comparatively more contained, focusing on specific functional concerns. While rights-adjacent arguments occasionally surface, they remain peripheral rather than central to the discourse. The theme points to a structural conflict between engagement-driven platform incentives and child protection. In public discourse, features that encourage prolonged use, emotional attachment, opt-out access, or entertainment diversion show how minors’ welfare can become secondary to retention, growth, and commercial visibility.

5.2.6 AI as a Parenting Support Tool

A distinctive theme in Chinese public discourse is the positive construction of AI as a practical tool for parenting and children's learning. The theme is mainly found in the Chinese material. Although the American material sometimes mentions AI's educational potential, it does not construct AI as a parenting support tool in the same sustained way. Importantly, this discourse presents AI as a supplement to parenting, rather than as a replacement for parents.

Public discussion presents Doubao as a versatile educational assistant with an extensive knowledge base and the capacity for customised role-based conversations. Parents describe using the platform to help children practise English conversation, explore historical knowledge through dialogue, check homework, and answer academic questions. Beyond educational functions, Doubao's consistent and patient communication style is valued by parents who find managing children's learning emotionally taxing, with some framing the platform as a helpful complement that allows parents a degree of relief from the pressures of supervising children's study (陆玖商业评论, 2025). A substantial part of the coverage frames Doubao positively in this regard, emphasising its role in supporting parenting and facilitating children's academic development (陆玖商业评论, 2025). This positive framing is further reflected in instructional content such as practical guides on how to make the most of the platform, which circulate widely in the Chinese public sphere, positioning Doubao primarily as a tool for improving everyday learning and family life (云分, 2026).

This theme is important because it represents a form of public meaning construction that operates independently of corporate framing. Doubao's own communication does not actively position the platform as a parenting support tool, yet public discourse has constructed this role on its own terms. This pattern is consistent with the technological solutionism narrative identified in the literature review, in which AI is viewed primarily as a practical tool for social advancement and everyday improvement rather than as a source of unmanageable risk (Bareis & Katzenbach, 2022).

5.2.7 Synthesis of RQ2 Findings

In both contexts, public discourse does not unconditionally accept the interpretations of companies' use of AI for minors. In both arenas, concerns exist about tech companies prioritizing profits over the safety of minors, the potential for misinformation, and the risk of emotional dependence.

However, the nature and intensity of public discourse differ between the two arenas. In the U.S., criticism is often systemic and structural, questioning corporate responsibility, design flaws, and potential violations of the rights of minors and guardians, while also offering directions for improvement for companies and government agencies. The U.S. public opinion is characterized by a complex interplay of diverse critical voices. In China, while criticism exists, it is relatively restrained, primarily focusing on specific functional limitations, without directly questioning the legitimacy of the companies. Simultaneously, a prominent feature of Chinese public discourse is the active construction of AI as a parenting aid, with the public actively viewing platforms like Doubao as practical resources for children's learning and family life. This meaning is constructed largely outside Doubao's own corporate framing.

5.3 RQ3: Alignment, Tensions, and Cross-contextual Differences

This section addresses RQ3 by comparing the themes identified in RQ1 and RQ2. It examines how corporate framing and public discourse align, diverge, or remain only loosely connected across the U.S. and Chinese contexts. The comparison is organised around four recurring issues: how AI for minors is constructed as a developmental issue, how responsibility is attributed, how risk and safeguards are discussed, and how corporate responsibility is evaluated. By comparing these themes across the two contexts, this section shows how AI for minors becomes a shared concern while also revealing different forms of alignment and tension between corporate framing and public discourse.

5.3.1 Minor-specific AI: Empowerment, Adaptation, and Parenting Support

In the U.S.: Empowering or superficial?

This subsection brings together the RQ1 theme of *Empowering to meet development needs* with the RQ2 themes of *Superficial Adaptation* and *Multidimensional Harms to Minors*. The tension is most visible in the U.S. context, while a different pattern emerges in the Chinese arena, as discussed later in this subsection.

The comparison shows a broad alignment: both corporate framing and public discourse in the U.S. treat AI for minors as a distinct issue requiring more than ordinary product access. Across the material, minors are not constructed as ordinary adult users. Corporate framing presents minors' AI use as requiring specific conditions, such as parental supervision, platform settings, safety tools, restricted access, or guidance around appropriate use. Public discourse similarly treats minors as a user group requiring special attention, especially where AI use is linked to cognitive, emotional, and developmental vulnerability. In this broad sense, both sides recognise that AI for minors requires a different form of attention than AI for adult users.

The tension lies in how this special attention is understood. In the U.S. context, corporate framing constructs AI for minors mainly through an empowerment logic. AI is presented as something that can support learning, creativity, exploration, and future readiness when used under appropriate conditions. Age-specific design and safeguards are therefore used to legitimise minors' access to AI: the issue is not whether minors should be excluded from AI, but how their access can be guided, adjusted, and made safer.

The U.S. public discourse does not fully accept this empowering framing. The theme of *Superficial Adaptation* questions whether current minor-facing AI products are genuinely designed around minors' developmental needs, or whether they remain modified versions of adult systems. This critique is reinforced by the theme of *Multidimensional Harms to Minors*, where public discourse links AI use to concerns such as emotional dependency, blurred AI-human boundaries, manipulation, and weakened critical or creative development. These concerns challenge the assumption that guided access is sufficient to make AI developmentally beneficial. From this perspective, the material suggests that safeguards and minor-specific settings may indicate recognition of minors as a special user group, but are not always treated as sufficient evidence of substantive age-appropriate design.

In China: Restriction and positive construction

This subsection brings together the RQ1 theme of *Protection through Functional Restriction* with the RQ2 themes of AI as a Parenting Support Tool.

In Doubao's corporate framing, AI for minors is constructed primarily through functional restriction. The minor mode is defined by what is limited or disabled, and appropriate use is framed as cautious, bounded, and directed away from over-reliance. Protection is built into the conditions of use rather than expressed through claims of developmental empowerment. Doubao does not actively position its product as a learning tool, a creative companion, or a resource for minors' development. The framing centres on what the platform removes or controls, not on what minors can gain.

Public discourse, however, does not challenge this restrictive framing, nor does it simply echo it. Instead, Chinese public discourse constructs its own positive narrative around AI for minors, one that is largely absent from Doubao's corporate communication. Public voices frame AI as a practical parenting support tool, highlighting its capacity to assist with homework, facilitate language practice, answer children's questions, and offer parents a degree of relief from the pressures of supervising learning. This positive construction operates independently of corporate framing: it is not a response to Doubao's claims, but a meaning that the public has produced on its own terms.

The relationship between corporate framing and public discourse in this area can be characterised as a soft alignment: not a direct convergence in which both sides articulate the same position, but a functional coexistence in which corporate restriction and public positive construction operate without mutual contradiction. Doubao limits the platform; the public finds value within and beyond those limits. Neither side directly engages with the other's framing, yet their positions do not conflict. This pattern differs from the U.S. context, where corporate empowerment claims are actively challenged by public discourse. In the Chinese arena, the absence of direct contestation produces a quieter but different form of alignment.

5.3.2 Accountability and Responsibility

This subsection compares the RQ1 theme of *Diffuse Accountability* with the RQ2 theme of *Contested Accountability*. The U.S. context shows both alignment and tension: responsibility is treated as central, but its distribution is contested. In the Chinese context, the relationship is mainly aligned, as parental or guardian responsibility is naturalised rather than problematised.

The broad alignment is that both corporate framing and public discourse recognise that AI for minors involves more than one responsible actor. Corporate framing presents minors' AI use as requiring the participation of guardians, users, platforms, and sometimes wider institutional actors. Public discourse also recognises that guardians have an important role in guiding minors' AI use. The point of disagreement is not whether parental involvement is necessary, but how responsibility is distributed and whether this distribution strengthens or weakens corporate accountability.

Within corporate framing, responsibility is constructed as shared across actors and mechanisms. Companies are positioned as providers of tools, safeguards, settings, age-specific features, and guidance, while guardians are positioned as supervisors who activate, monitor, or manage minors' AI use. This framing makes AI for minors appear governable through a combination of platform design, parental oversight, and user conduct. It is precisely this distribution that becomes contested in the U.S. context.

The U.S. public discourse challenges this distribution when it appears to shift responsibility away from companies. The RQ2 findings show that public criticism does not necessarily reject parental involvement. Instead, it questions whether parents are being made responsible for risks shaped by corporate design choices, default settings, or insufficient safeguards. Parents and users may be asked to manage many of the risks, while companies retain control over product design, default settings, data practices, and the scope of safeguards. Public discourse questions whether warnings, parental controls, opt-out mechanisms, or general guidance are sufficient. These measures may support parental involvement, but they do not replace corporate responsibility for making AI products safe and developmentally appropriate for minors.

In the Chinese context, this tension is less visible. Parental or guardian responsibility is more often naturalised than contested. Public discourse tends to treat guardian involvement as a necessary complement to minors' AI use, particularly in relation to supervision, guidance, and the prevention of overuse. As a result, the distribution of responsibility between platforms and guardians is less likely to be interpreted as responsibility-shifting. Instead, it appears closer to a shared governance logic in which parental supervision is treated as a normal condition of minors' AI use. This produces stronger alignment between corporate framing and public discourse than in the U.S. context.

The relationship between the two themes centres on the meaning of shared responsibility. In the U.S. context, distributed responsibility is partly aligned with public concern for minors' protection, but it is also challenged as a possible way of reducing corporate accountability. In the Chinese context, the relationship is more clearly aligned. Guardian responsibility is not problematised as responsibility-shifting, but treated as a necessary and legitimate part of minors' AI use. Public discourse reinforces rather than challenges the corporate framing of shared responsibility.

5.3.3 Risk Construction

This subsection compares the RQ1 theme of *Risk acknowledged but managed* with the RQ2 themes of *Multidimensional Harms to Minors* and *Needs for Safeguarding*. Both alignment and tension are present in the two contexts, but they appear differently.

The alignment is shared across both contexts: both corporate framing and public discourse recognise that AI for minors involves risk. Companies do not present minors' AI use as entirely risk-free, and public discourse also treats protection as necessary. Both sets of findings construct safeguards as central to AI for minors. The tension lies not in whether risk exists, but in how risk is understood and how it can be controlled.

Corporate framing in the two contexts shows a degree of similarity. In both, risk is presented as identifiable and manageable, something that can be addressed through safety mechanisms, platform rules, age-specific settings, parental controls, or restricted modes. AI for

minors is not framed as harmless, but as something that can be made acceptable through structured safeguards. Within this shared logic, however, the emphasis differs. In the U.S. context, companies stress that their risk management processes are scientific, systematic, and consistent. In the Chinese context, the focus of risk shifts from what the platform might generate to how minors might use it, with Doubao concentrating on behavioural risks such as preventing over-immersion.

The public discourse patterns produce different forms of tension in the two contexts. In the U.S. arena, public discourse constructs risk in a broader and less controllable way. Concerns extend beyond isolated safety incidents to include emotional dependency, blurred AI-human boundaries, misinformation, harmful content, overuse, privacy violations, and possible effects on creativity and critical thinking. These concerns make risk appear broader than the categories directly addressed in corporate communication. This is developed further in the theme of Needs for Safeguard, where public discourse moves beyond general calls for protection and questions the reliability of the safeguards companies already claim to provide. Safeguards are treated as fragile, incomplete, easy to bypass, or unable to respond to context-specific signs of vulnerability. From this perspective, filters, settings, and parental controls are not necessarily sufficient if the system cannot recognise prolonged emotional dependence, escalating distress, or age-related vulnerability. Corporate communication presents safeguards as evidence of control; U.S. public discourse challenges whether such control is reliable in practice.

In the Chinese arena, public discourse also raises concerns about multidimensional harms, including misinformation and emotional attachment. However, these discussions concentrate primarily on specific functional design issues, such as the short video feature leading to excessive entertainment and potential addiction among minors. corporate framing and public discourse in the Chinese context share a relatively consistent definition of risk, with both sides concentrating on behavioural concerns such as addiction and over-immersion rather than content safety or systemic harm. Systematic diagnosis of safeguard adequacy and proposals for structural improvement are largely absent.

In summary, both contexts display alignment in acknowledging the existence of risk, and corporate framing in both contexts shares a similar logic of presenting risk as manageable, though with slightly different emphases. However, the depth and scope of public discourse produce different tensions. In the U.S. context, public voices mount a multidimensional challenge to both the definition of risk and the effectiveness of corporate safeguards. In the Chinese context, tension exists but remains at the level of functional criticism, without escalating into the kind of structural challenge to corporate legitimacy observed in the U.S. arena.

5.3.4 Corporate Responsibility and Motivation

This subsection compares the RQ1 theme of *Corporate Responsibility Commitment* with the RQ2 theme of *Profit over Protection*. Both alignment and tension are present in the two contexts, but they appear with different intensities.

The shared alignment across both contexts is that corporate communication and public discourse converge on a common normative principle: when it comes to minors as a vulnerable group, safety should be prioritised. Both sides agree that corporate responsibility toward minors matters.

However, how companies construct their responsibility commitment and how public discourse challenges it differ across the two contexts. In the U.S. context, AI companies consistently position themselves as responsible actors, collaborating with external experts, establishing advisory councils, and committing to ongoing improvement informed by research. ChatGPT in particular constructs an image of industry leadership by combining CEO-level statements, expert advisory councils, and engagement with policymakers. In the Chinese context, Doubao's source of authority is the external legal framework. It presents itself as a compliant actor rather than an industry leader, but similarly commits to contributing to minors' well-being.

The tension lies not in the principle itself, but in how it is evaluated in practice. The theme of *Profit over Protection* is present in public discourse in both contexts. In the U.S. arena, critics argue that corporate responsibility rhetoric coexists with design choices that prioritise engagement over protection: minimal age verification, opt-out rather than opt-in default settings,

push notifications that draw users back, conversational systems designed to be persistently agreeable, follow-up prompts that sustain interaction, and companion features that cultivate emotional attachment. From the perspective of public discourse, these patterns suggest that commercial interests remain central to product design, and that responsibility commitments function more as legitimising strategies than as effective limits to corporate behaviour. In the Chinese context, public discourse has also raised concerns about whether Doubao's short video feature reflects commercially driven logic, and expressed dissatisfaction with commercial interests infiltrating learning functions, but these criticisms have not formed a sustained or large-scale challenge.

The relationship between these two themes centres on the credibility of corporate responsibility claims. Both sides share the normative position that safety should come first. However, the gap between stated commitment and perceived practice constitutes a recurring source of tension, reflecting the legitimacy gap between corporate framing and public expectations identified in the literature review.

In summary, both contexts show public expectations of corporate responsibility and public questioning of whether companies are driven by commercial interests. However, the framing strategies and the intensity of public challenge differ. In the U.S. context, companies construct an image of industry leadership and commit to sustained, research-informed product improvement, while U.S. public discourse forms a cross-source critical narrative questioning whether commercial interests undermine these promises. In the Chinese context, Doubao positions itself as a legal compliance actor and commits to improving its product with minors' wellbeing in mind, while Chinese public discourse, though critical of commercial logic in specific instances, does not extend to questioning the legitimacy of corporate involvement in providing AI to minors.

5.3.5 Synthesis of RQ3 Findings

Across both contexts, corporate framing and public discourse share a common baseline of agreement. Both sides recognise AI for minors as an issue requiring special attention, acknowledge that risks are present in minors' AI use, affirm the necessity of parental

involvement, and agree that companies should take responsibility and continuously improve their products. However, tensions exist alongside these alignments, and they appear differently across the two contexts.

In the U.S. arena, tension is present across all four thematic comparisons: companies' empowerment claims are challenged as superficial adaptation, their diffusion of responsibility is contested as evasion, their risk framing is expanded and questioned by public discourse, and their responsibility commitments are scrutinised as potentially undermined by commercial interests.

In the Chinese arena, the pattern is more varied. The first comparison about empowerment and adaptation produces neither alignment nor tension in the traditional sense. Doubao does not actively claim developmental empowerment, and public discourse does not challenge the restrictive approach. Instead, public discourse itself constructs AI as a tool capable of supporting minors' learning and development, a framing that is absent from corporate communication. The second comparison about accountability produces alignment, with public discourse accepting the assignment of responsibility to guardians as a natural arrangement. Tension in the Chinese context is concentrated in the latter two comparisons: risk construction and corporate motivation. In both cases, criticism exists but remains focused on specific functional issues and does not escalate into structural challenges to corporate legitimacy.

These findings suggest that while the basic normative premises around AI for minors are shared across contexts, the depth, scope, and nature of tensions between corporate framing and public discourse are shaped by the communicative environment in which they develop.

6. Discussion

6.1 Legitimacy Construction and Contestation

Legitimacy is not something companies can secure through communication alone. As discussed in the literature review, meanings around AI are shaped through interaction among multiple actors, including companies, media, regulators, advocacy groups, and users (Chuan, 2023; Chuan et al., 2019). This is especially important in the case of AI for minors, where the technology is associated with both future participation and heightened vulnerability. Companies, therefore, need to make AI appear useful and innovative, but also appropriate, protected, and socially acceptable for young users. This section examines how legitimacy around AI for minors is constructed and contested in the U.S. and China.

6.1.1 The U.S.: Elaborated Framing and Public Challenge

In the U.S. context, companies construct legitimacy through a relatively complete and multi-layered framing strategy. AI for minors is not presented only as a risky technology requiring control, but also as a developmental resource connected to learning, creativity, exploration, personalised support, and future participation in an AI-shaped society. At times, this framing approaches a form of quasi-entitlement: minors are not explicitly said to have a legal right to AI, but exclusion from AI is made to appear undesirable because it may limit their future opportunities. In this sense, corporate communication makes access itself part of the legitimacy claim.

However, this access is not framed as unconditional. It is made acceptable through a wider governance vocabulary, including safeguards, parental controls, expert involvement, and responsibility commitments. Governance functions as an enabling mechanism: it allows companies to present minors' AI use as supervised rather than uncontrolled. This extends the

literature on corporate AI framing, where companies often use broad normative labels such as “Responsible AI,” “Trustworthy AI,” or “AI for Good” to position themselves within ethical and societal debates (Mao et al., 2025). In the case of minors, this legitimacy work becomes more specific. AI is framed as responsibly provided, but also as something young users should not be entirely excluded from.

This elaborate legitimacy frame, however, does not produce stronger public acceptance. In the U.S. public discourse examined here, corporate claims are met with a sharp and sceptical response. Public actors do not necessarily reject AI for minors in principle, but they question whether companies have earned the authority to define what safe and appropriate AI use for minors should mean. The legitimacy gap arises from doubt over whether these claims are credible, sufficient, and consistent with actual product design.

This gap can be understood in several ways. First, public discourse treats legal compliance, parental controls, and technical safeguards as minimum conditions rather than sufficient proof of responsibility. Following rules or offering safety tools does not automatically demonstrate that companies have adequately protected minors. Second, corporate commitments are evaluated against product features that may encourage emotional attachment, prolonged interaction, or dependence. When these features appear to conflict with child protection, public discourse reads them as evidence that corporate responsibility remains unstable. Third, the sharpness and distrust visible in public discourse may also reflect a stronger accountability culture in which media, advocacy groups, and civil society actors are expected to question corporate narratives rather than accept them as self-evident.

This reflects the contested and multi-actor nature of framing, in which meanings are negotiated and challenged among different voices rather than controlled by a single organisation (Benford & Snow, 2000). In this sense, public contestation becomes part of the legitimacy process: companies construct legitimacy through claims about opportunity, protection, and responsibility, while public discourse tests whether such claims remain credible in practice.

6.1.2 China: Compliance-centred Framing and Functional Reinforcement

The Chinese company materials show a more bounded legitimacy logic. Doubao constructs AI for minors as acceptable through restriction, supervision, and compliance. The minor mode defines appropriateness by limiting certain functions, placing use under guardian management, and framing appropriate engagement as cautious and healthy. This gives corporate legitimacy a narrower basis than in the U.S. materials. AI for minors is made acceptable through controlled conditions of use, rather than through an expansive claim about children's access to future-oriented AI participation.

This bounded framing also gives the company's legitimacy claim a degree of stability. Doubao does not place heavy emphasis on AI as a developmental entitlement or an open-ended resource for minors. The platform is presented as something that can be made suitable when its functions are limited, and its use is placed within legal and guardian-based arrangements. Regulatory compliance therefore works as a visible legitimacy cue. It signals that the product belongs within a recognised governance framework, while functional restriction signals that children's use has been adjusted to fit the risks associated with minors.

Public discourse in the Chinese materials responds to this framing with limited confrontation. Concerns exist, especially around overuse, misinformation, entertainment diversion, and parental guidance. However, these concerns usually focus on whether AI remains within appropriate boundaries. When Doubao is criticised for short-video recommendations or entertainment-oriented design, the issue is framed as a deviation from proper use. AI for minors is expected to support learning and family guidance; concerns emerge when it pulls children toward distraction, excessive engagement, or platform entertainment.

The theme of AI as a parenting support tool is significant because it gives public discourse a constructive role in the legitimacy process. Chinese public discourse does more than evaluate whether Doubao's restrictions are adequate. It also assigns AI a socially acceptable place within family life. AI becomes legitimate when it is imagined as a tool that helps parents manage learning, answer children's questions, support homework, or reduce the emotional burden of tutoring. This shifts the meaning of AI for minors from a direct child–technology

relationship to a parent-mediated support relation. The product is accepted less as a child-facing technology and more as a tool that can assist parental supervision.

This also reshapes how responsibility is understood. In the U.S. materials, parental involvement often becomes a site of contestation because it can appear to reduce corporate accountability. In the Chinese materials, parental involvement is more easily incorporated into a positive model of appropriate use. AI is expected to assist parents, while parents remain the legitimate supervisors of children's engagement with the tool. This makes the distribution of responsibility appear less like burden-shifting and more like a normal arrangement for governing minors' technology use. The public construction of AI as a parenting support tool, therefore reinforces the corporate emphasis on guardian involvement, even though this meaning is not strongly developed in Doubao's own communication.

This pattern can be understood as a form of bounded usefulness. AI for minors is acceptable when it is useful, supervised, educationally oriented, and contained. Its legitimacy weakens when these boundaries are crossed, especially when learning functions become connected to entertainment, commercial retention, or overuse. This explains why criticism exists in the Chinese material but does not usually become a broader challenge to corporate authority. Public discourse questions specific design choices, while the broader idea of AI as a useful family and learning tool remains relatively intact.

This finding connects with earlier research suggesting that Chinese AI discourse often gives more weight to application, stability, and social advancement, while criticism tends to be expressed through governance-compatible language (Ding & Kong, 2019). It also complicates the technological solutionism perspective discussed in the literature review, where AI is often imagined as a practical tool for solving social or everyday problems (Bareis & Katzenbach, 2022). In the Chinese materials examined here, this solution-oriented view is conditional. AI is welcomed as useful only when its use remains bounded by education, family supervision, and regulatory order. Legitimacy is maintained through a combination of formal control and practical value: the product appears acceptable because it is both governable and useful within an approved social role.

6.1.3 Comparative Reflection

This contrast shows that the relationship between corporate framing and legitimacy is not linear. Elaborate corporate framing does not automatically produce stronger legitimacy, and narrower compliance-based framing does not necessarily produce weaker legitimacy. In the U.S. materials, companies make stronger claims about opportunity, protection, expertise, and responsibility. These claims also create clearer standards for public scrutiny. The more companies define AI for minors as beneficial and responsibly governed, the more public actors can question whether these claims are credible in practice.

In the Chinese materials, corporate framing is more limited, but it fits more closely with visible expectations around compliance, supervision, and proper use. Legitimacy is supported by a narrower set of criteria: the product should be regulated, supervised, useful, and kept within appropriate boundaries. Public discourse tends to evaluate whether AI remains within this acceptable role, rather than challenging the company's authority to provide AI for minors more broadly.

The main difference lies in what each public arena treats as sufficient grounds for legitimacy. In the U.S. context, safety claims require further proof and are tested through public scrutiny. In the Chinese context, compliance and bounded usefulness carry stronger legitimating weight. Public discourse, therefore, plays different roles across the two settings: in the U.S., it functions more as a counter-frame that challenges corporate authority; in China, it more often works as a supplementary frame that adds practical and educational value while raising boundary-focused concerns. This comparison also suggests that media and public discourse actively shape whether corporate claims become credible, contested, or reinforced (Chuan, 2023; Ocal & Crowston, 2024).

6.1.4 Additional Observation: Communicative Style

A further observation concerns the communicative format and style of corporate communication. Although this study focused on the argumentative content of corporate framing, the materials also show notable stylistic differences across the two contexts.

Doubao's communication adopts a contractual and legalised format. The minor mode documentation reads like formal terms and conditions, positioning the platform as a service governed by rules, obligations, and liability boundaries. This style reinforces the compliance-centred character of its legitimacy framing, making the platform appear already situated within a formal order of rules and obligations. By contrast, the U.S. companies use more narrative and identity-oriented communication, making corporate responsibility appear more visibly performed and therefore more open to public evaluation. ChatGPT positions itself through leadership statements and public commitments, Gemini adopts an inviting learning-oriented tone, and Grok's communication remains comparatively minimal.

These differences matter because rhetorical strategies include how messages are formatted and delivered within the rhetorical arena (Frandsen & Johansen, 2000). From a strategic communication perspective, communicative form can shape how audiences perceive credibility, authority, and organisational intent (Falkheimer & Heide, 2018). This study does not examine how these stylistic differences affect public interpretation, since the analysis focused primarily on thematic content. Future research could examine how the form of corporate communication influences the public meaning of AI for minors across different cultural and regulatory contexts.

6.2 Normative Expectations: What Should AI for Minors Look Like?

The findings of this study reveal that the alignment and tensions between corporate framing and public discourse, as well as the public normative expectations regarding AI for minors, have not yet fully taken shape. This uncertainty is particularly pronounced in the U.S. context, while the Chinese context presents a more stabilised pattern.

6.2.1 The U.S. Context: Unsettled and Evolving Public Expectations

In the U.S. public discourse, normative expectations around AI for minors show signs of being contradictory and unsettled across several dimensions.

First, the U.S. public's demands regarding the use of AI by minors are inherently contradictory. At the macro level, the public has widely criticised the potential harms of AI to

minors and called for stronger restrictions and protections. At the same time, U.S. policy actively promotes the integration of AI into education, and the public recognises its potential value in learning and development. As a result, public discourse contains two competing expectations: minors should be protected from the potential harms of AI, but they also should not fall behind in the AI era. The question of how to balance these two demands has not yet reached a consensus between the government and the public. Similar tensions also appear in discussions of specific platform designs. For instance, the public is calling for greater parental authority to monitor their minors' interactions with AI, while simultaneously emphasizing the importance of respecting minors' privacy and condemning platforms for diverting responsibility by involving parents. These tensions do not necessarily reflect inconsistency within public opinion. Instead, they point to practical dilemmas that remain unresolved. The public has expectations regarding the safety and feature-richness of products, yet lacks a clear understanding of how to achieve this balance in specific designs. This may reflect the difficulty of forming stable social expectations around emerging technologies in their early stages. The evolution of such expectations is a gradual process that unfolds through the collaborative deliberations of multiple stakeholders. In the case of minors, these expectations become even harder to define because debates around children, technology, education, and responsibility are already highly contested in the U.S. public sphere.

Second, a lack of consensus exists within the U.S. public regarding social norms of accountability. On the one hand, public opinion demands that corporations assume primary, inalienable responsibility; on the other hand, the public also calls for stronger government regulation while emphasising the necessity of parental involvement. This suggests that issues such as the allocation of responsibility for the protection of minors using AI and the identification of the ultimate bearer of responsibility have not been clearly or consistently defined.

Third, the normative expectations of the U.S. public are not static but continuously shaped through interactions with corporate behaviour. The findings of this study suggest a correlation between public criticism of profit-driven design strategies and demands for enhanced protection. The public perceives corporate engagement-driven design strategies as both a "source of the problem" and a "reason for stronger protection." As a company's approach becomes

increasingly profit-driven, the public's protective demands intensify. This constitutes a cyclical tension. When commercial motives become more visible, calls for stronger protection also increase. This signifies that the public's normative expectations are perpetually recalibrated and elevated in accordance with corporate conduct, rather than being anchored to a pre-existing, uniform standard.

These observations suggest that in the U.S., a coherent framework for public normative expectations regarding minors' use of AI has not yet been established. These expectations are internally contradictory, lack consensus on accountability, and are constantly reshaped through interactions with corporate behaviour. This reflects the broader difficulty of forming stable social expectations around emerging technologies that involve vulnerable users and rapidly evolving forms of communication.

6.2.2 The Chinese Context: Relatively Consistent Public Expectations

By contrast, within the materials collected for this study, Chinese public normative expectations display a relatively high degree of consistency. AI is widely regarded as a learning and parenting aid in the education of minors, and parental supervision is seen as a natural arrangement. Limited versions of AI for children are readily accepted by the public. This contrasts sharply with the debates about rights and responsibilities observed in the U.S. However, whether this consistency indicates that Chinese public normative expectations are more mature, or whether the structure of the public discursive environment discourages challenges to the legitimacy of compliant companies, meaning that potential contradictions have not yet fully surfaced, is a question that this study cannot answer. Future research could further examine how different discursive environments shape the formation of public normative expectations around emerging technologies and vulnerable user groups.

7. Conclusion

This thesis began with a relatively simple concern: generative AI is becoming more powerful and more present in young people's everyday lives, while minors remain a particularly vulnerable group of users. Minors have begun to use AI more widely for learning, creativity, entertainment, advice, and companionship. As this use becomes more common, cases and concerns related to minors' safety have also made AI for minors a more visible public issue.

Existing discussions have paid attention to AI risks, minors' digital safety, and platform regulation, but less is known about how AI companies themselves frame AI for minors, and how these framings relate to wider public discourse. This thesis addresses this gap by treating AI for minors as a strategic communication issue, rather than only a technical, educational, or regulatory one. It examines how meanings around AI for minors are produced through corporate communication, and how these meanings are accepted, negotiated, or challenged by public actors.

By comparing China and the U.S., the thesis further examines how similar concerns around minors and AI take shape under different institutional, regulatory, and media conditions. The comparison is used to understand how different communication environments shape the relationship between corporate framing and public discourse, while avoiding a simple opposition between the two contexts. The following section returns to the three research questions and summarises the main findings of the study.

7.1 Addressing the Research Questions

To answer the first research question, AI companies frame AI for minors through a balancing logic: minors' access to AI is made legitimate only when it is presented as supervised, restricted, and manageable. Across the corporate materials, AI is not constructed as simply

beneficial or simply dangerous. It is presented as a technology that may support young users, but only under specific conditions of parental involvement, platform safeguards, age-related restrictions, and responsible use. This framing allows companies to acknowledge minors' vulnerability without rejecting minors' access to AI altogether. It also distributes responsibility across several actors. Parents are positioned as supervisors, minors are encouraged to learn responsible use, and companies present themselves as providers of tools, settings, guidance, and safeguards. However, this balance differs across contexts. The U.S.-based platforms generally place more emphasis on enabled access, developmental value, and guided use, while Doubao's communication is more strongly oriented toward restriction, regulatory compliance, and guardian responsibility. Corporate communication constructs AI for minors as a conditional form of access: possible and potentially valuable, but only when managed through multiple actors and mechanisms.

Regarding the second research question, public discourse constructs AI for minors as a contested and still unsettled issue. It does not simply reject minors' use of AI, but questions the conditions under which such use can be considered safe, appropriate, and responsible. Across the material, public voices challenge the developmental fit of existing products, the commercial logic behind engagement-oriented design, the distribution of accountability, and the adequacy of current safeguards. These concerns give public discourse a more critical and interrogative role than corporate communication. However, this role is not performed in the same way across the two contexts. In the U.S. arena, public discourse develops into a more systemic critique of corporate responsibility, design failures, commercial incentives, and multidimensional harms to minors. In the Chinese arena, public discourse is more mixed and contained: concerns are raised around overuse, entertainment displacement, misinformation, and functional design issues, while AI is also frequently presented as a useful tool for learning, parenting, and academic support. Overall, public discourse raises the standard for what AI for minors should involve, but the U.S. arena turns this into a more explicit challenge to corporate legitimacy, while the Chinese arena keeps the critique closer to practical use and functional improvement.

Finally, to answer the third research question, the study found that alignment and tension between corporate framing and public discourse are shaped by the communication environment

in which they appear. At a general level, both corporate communication and public discourse recognise that AI for minors requires special attention, involves risk, and depends on some combination of platform responsibility and parental involvement. However, this shared baseline did not translate into the same relationship between businesses and public opinion in both countries. The U.S. public discourse more actively countered corporate communications. Claims by businesses regarding empowerment, shared responsibility, manageable risks, and corporate commitment were questioned due to concerns about suitability, shifting responsibility, weak safeguards, and commercial motives. In China, corporate framing and public discourse exhibited a stronger functional coexistence; tensions remained, but AI was generally perceived as a learning and parenting support rather than a risk to be controlled. The findings suggest that the relationship between corporate framing and public discourse is shaped by specific communication environments.

Taken together, these findings show that AI for minors becomes socially acceptable through specific forms of communication, governance, and public expectation. Companies attempt to legitimise minors' AI use by framing it as valuable, manageable, and responsible, while public discourse tests, modifies, or reinforces these claims in context-specific ways. The U.S. arena shows a more contested legitimacy process, where expectations around safety, accountability, privacy, parental authority, and age-appropriate design remain unsettled. The Chinese arena shows stronger functional coexistence, where compliance, supervision, restriction, and practical usefulness support a more consistent expectation of AI as a learning and parenting aid. The thesis shows that what counts as safe, responsible, and appropriate AI for minors is shaped through communication within specific public arenas.

7.2 Contributions to Research and Practice

This study contributes to research in strategic communication by showing how AI companies take part in shaping public meaning around emerging technologies. In the case of AI for minors, corporate communication does not only describe product features or safety measures. It also defines what kind of access is appropriate for minors, how responsibility should be distributed, and how risks should be understood. By comparing China and the U.S., the study

further shows that corporate framing does not operate in isolation. Similar concerns around AI, safety, and minors produce different patterns of acceptance, negotiation, and contestation depending on the communicative environment. This contributes to cross-contextual communication research by demonstrating how institutional settings, media systems, and public expectations shape the relationship between corporate framing and public discourse.

At the practical level, the study offers insights for AI companies, policymakers, and child-safety stakeholders. For companies, the findings suggest that responsibility claims, safeguards, and parental controls are unlikely to secure public trust on their own. Public actors may also evaluate whether these claims are consistent with product design, default settings, commercial incentives, and actual user experience. For policymakers and child-safety stakeholders, the study highlights the need to consider how responsibility is distributed among platforms, parents, minors, and regulators. AI for minors requires technical safeguards, but also clearer expectations around age-appropriate design, default protection, transparency, parental involvement, and the commercial logic of engagement-driven systems.

7.3 Reflexivity and Future Research

Finally, it is worth reflecting on the fact that many findings in this study align with common narratives about China and the U.S., such as limited public discussion space in China and systematic public criticism in the U.S. Both researchers were aware of these stereotypes and sought to minimise their influence through collaborative coding and ongoing dialogue. Whether the findings reflect genuine structural differences, or have been partly shaped by pre-existing cognitive frameworks cannot be fully resolved. The collected materials, while equivalent in nature, already reflect different media operating logics, which may contribute to the observed differences. Nevertheless, the findings offer a more nuanced picture than stereotypes suggest: the Chinese context is not devoid of critical voices, and the American context is not characterised solely by contestation. These nuances are precisely what stereotypes fail to capture. These reflections should also be understood alongside the methodological boundaries of the study, including its reliance on publicly available materials, the different media environments from which the materials were collected, and the broad age category of minors.

Future research could test the robustness of these findings by incorporating more diverse data sources, such as user interviews, internal corporate perspectives, or cross-platform comparisons. In addition, several open questions raised in the literature review and discussion of this study merit further exploration, including how different discursive environments shape the formation of public normative expectations, how AI for minors and AI for adults differ in corporate communication and public perception, whether the communicative format of corporate framing influences public interpretation, and how misalignment between public expectations and regulatory standards is produced.

References

- Af Malmborg, F., & Trondal, J. (2021). Discursive framing and organizational venues: mechanisms of artificial intelligence policy adoption. *International Review of Administrative Sciences*, 89(1), 39-58.
- Allen, G. C. (2019). *Understanding China's AI strategy: Clues to Chinese strategic thinking on artificial intelligence and national security*. Center for a New American Security.
<https://www.cnas.org/publications/reports/understanding-chinas-ai-strategy>
- Bareis, J., & Katzenbach, C. (2022). Talking AI into Being: The Narratives and Imaginaries of National AI Strategies and Their Performative Politics. *Science, Technology, & Human Values*, 47(5), 855-881
- Benford, R. D., & Snow, D. A. (2000). Framing processes and social movements: An overview and assessment. *Annual review of sociology*, 26(2000), 611-639.
- Benoit, W.L. (1995): Accounts, Excuses, and Apologies: *A Theory of Image Restoration Strategies*. State University of New York Press, New York.
- Borup, M., Brown, N., Konrad, K., & Van Lente, H. (2006). The sociology of expectations in science and technology. *Technology analysis & strategic management*, 18(3-4), 285-298.
- Booth, R. (2025, September 2). Parents could get alerts if children show acute distress while using ChatGPT. *The Guardian*.
<https://www.theguardian.com/technology/2025/sep/02/parents-could-get-alerts-if-children-show-acute-distress-while-using-chatgpt>

- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative research in psychology*, 3(2), 77-101.
- Braun, V., & Clarke, V. (2019). Reflecting on reflexive thematic analysis. *Qualitative research in sport, exercise and health*, 11(4), 589-597.
- Burgess, M., & Varner, M. (2026, January 7). Grok is generating sexual content far more graphic than what's on X. *WIRED*.
<https://www.wired.com/story/grok-is-generating-sexual-content-far-more-graphic-than-whats-on-x/>
- ByteDance. (2025, August 26). 未成年人使用须知.
https://lf9-cdn-tos.draftstatic.com/obj/ies-hotsoon-draft/grace_legal/youthmode.html
- Cath, C. (2018). Governing artificial intelligence: ethical, legal and technical opportunities and challenges. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 376(2133).
- Chuan, C. H., Tsai, W. H. S., & Cho, S. Y. (2019, January). Framing artificial intelligence in American newspapers. In *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society* (pp. 339-344).
- Chuan, C. H. (2023). A critical review of news framing of artificial intelligence. *Handbook of critical studies of artificial intelligence*, 266-276.
- Chen, X., & Xu, L. (2025). State, society, and market: Interpreting the norms and dynamics of China's AI governance. *Computer Law & Security Review*, 59, 106206.
<https://doi.org/10.1016/j.clsr.2025.106206>
- Chen, Y., & Rowlands, I. H. (2022). The socio-political context of energy storage transition: Insights from a media analysis of Chinese newspapers. *Energy Research & Social Science*, 84, 102348.

- Christ, D. (2025, May 3). Gemini AI for kids: Why parents must slam the brakes. *Dan Christ*. <https://www.danchrist.com/gemini-ai-for-kids-parents-slam-the-brakes/>
- Coombs, W.T. (1999). Ongoing Crisis Communication. *Planning, Managing, and Responding*. Sage, London.
- Common Sense Media. (2025a, September 5). *AI risk assessment: Gemini (under age 13)*. <https://www.common Sense Media.org/sites/default/files/featured-content/files/csm-ai-risk-assessment-gemini-u13-09052025.pdf>
- Common Sense Media. (2025b, October 23). *ChatGPT: AI rating*. <https://www.common Sense Media.org/ai-ratings/chatgpt-5>
- Common Sense Media. (2026, January 22). *Grok. Common Sense Media AI ratings*. <https://www.common Sense Media.org/ai-ratings/grok>
- Cyberspace Administration of China. (2023, July 13). 生成式人工智能服务管理暂行办法 [Interim measures for the management of generative artificial intelligence services]. http://www.cac.gov.cn/2023-07/13/c_1690898327029107.htm
- Deng, R., & Ahmed, S. (2025). Perceptions and paradigms: An analysis of AI framing in trending social media news. *Technology in Society*, 81, 102858.
- Ding, H., & Kong, Y. (2019). Constructing Artificial Intelligence in the US and China: A Cross-Cultural, Corpus-Assisted Study. *China Media Research*, 15(1).
- de Carvalho Souza, M. E., & Weigang, L. (2025). Grok, Gemini, ChatGPT and DeepSeek: Comparison and applications in conversational artificial intelligence. *Inteligencia Artificial*, 2(1).
- de Laat, P. B. (2021). Companies committed to responsible AI: From principles towards implementation and regulation?. *Philosophy & technology*, 34(4), 1135-1193.

- Druga, S., Williams, R., Park, H. W., & Breazeal, C. (2018, June). How smart are the smart toys? Children and parents' agent interaction and intelligence attribution. In *Proceedings of the 17th ACM conference on interaction design and children* (pp. 231-240).
- Duffy, C. (2025, August 26). Parents of 16-year-old sue OpenAI, claiming ChatGPT advised on his suicide. *CNN*.
<https://edition.cnn.com/2025/08/26/tech/openai-chatgpt-teen-suicide-lawsuit>
- Eira, M., Rasouli, A., & Charisi, V. (2025). Parents' Perceptions About the Use of Generative AI Systems by Adolescents. In *Proceedings of the 24th Interaction Design and Children (IDC '25)*, 927–931. <https://doi.org/10.1145/3713043.3731508>
- Entman, R. M. (1993). Framing: Towards clarification of a fractured paradigm. *McQuail's reader in mass communication theory*, 390, 397.
- Executive Office of the President, National Science and Technology Council, Committee on Technology. (2016). *Preparing for the future of artificial intelligence*.
https://www.whitehouse.gov/sites/default/files/whitehouse_files/microsites/ostp/NSTC/preparing_for_the_future_of_ai.pdf
- Fairplay & EPIC. (2025, May 21). *Letter to Google re: Gemini rollout for children under 13* [PDF]. *Fairplay for Kids*.
https://fairplayforkids.org/wp-content/uploads/2025/05/Letter-to-Google-re-Gemini-EPIC-and-Fairplay_5.21.25.pdf
- Falkheimer, J., & Heide, M. (2018). What is strategic communication?. *Strategic Communication: An Introduction* (pp. 56-83). Routledge.
<https://doi-org.ludwig.lub.lu.se/10.4324/9781315621555-4>
- Fast, E., & Horvitz, E. (2017, February). Long-term trends in the public perception of artificial intelligence. In *Proceedings of the AAAI conference on artificial intelligence* (Vol. 31, No. 1).

Federal Trade Commission. (2025, September 11). *FTC launches inquiry into AI chatbots acting as companions*.

<https://www.ftc.gov/news-events/news/press-releases/2025/09/ftc-launches-inquiry-ai-chatbots-acting-companions>

Flick, U. (2023). *An introduction to qualitative research*. SAGE.

Frandsen, F., & Johansen, W. (2000). Retorik og krisekommunikation [Rhetoric and crisis communication]. *Rhetorica Scandinavica*, 14, 45–63.

Frandsen, F., & Johansen, W. (2017). Organizational crisis communication: A multivocal approach. *SAGE*

Frandsen, F., & Johansen, W. (2020). Arenas and voices in organizational crisis communication: how far have we come. In: Crisis Communication. *Handbooks of Communication Science* (ed. F. Frandsen and W. Johansen), 195–212. Berlin: De Gruyter Mouton.

Frandsen, F., & Johansen, W. (2022). Rhetorical Arena Theory. In *The Handbook of Crisis Communication* (eds W.T. Coombs and S.J. Holladay).

<https://doi-org.ludwig.lub.lu.se/10.1002/9781119678953.ch12>

Gamson, W. A., & Modigliani, A. (1989). Media discourse and public opinion on nuclear power: A constructionist approach. *American journal of sociology*, 95(1), 1-37.

Google. (2025, September). *Guide your child's Gemini Apps experience*.

<https://support.google.com/gemini/answer/16109150?hl=en>

Grose, J. (2026, April 8). *Deepfake nudes are haunting America's teens*. *The New York Times*.

Hallahan, K., Holtzhausen, D., Van Ruler, B., Verčič, D., & Sriramesh, K. (2007).

Defining strategic communication. *International journal of strategic communication*, 1(1), 3-35.

- Hallin, D. C., & Mancini, P. (2004). *Comparing media systems: Three models of media and politics*. Cambridge University Press.
- Hill, K. (2025, September 2). OpenAI plans to add safeguards to ChatGPT for teens and others in distress. *The New York Times*.
<https://www.nytimes.com/2025/09/02/technology/personaltech/chatgpt-parental-controls-openai.html>
- Hill, K. (2026, April 4). What teens are doing with those role-playing chatbots. *The New York Times*.
<https://www.nytimes.com/2026/04/04/technology/ai-chatbots-teen-roleplay.html>
- Hine, E., & Floridi, L. (2024). Artificial intelligence with American values and Chinese characteristics: A comparative analysis of American and Chinese governmental AI policies. *AI & Society*, 39, 257–278. <https://doi.org/10.1007/s00146-022-01499->
- Ho, T. C., & King, L. S. (2021). Mechanisms of neuroplasticity linking early adversity to depression: Developmental considerations. *Translational Psychiatry*, 11, 517.
<https://doi.org/10.1038/s41398-021-01639-6>
- Internet Watch Foundation. (2024). AI-generated child sexual abuse material: An emerging threat (Public update report). Internet Watch Foundation.
https://admin.iwf.org.uk/media/nadlcb1z/iwf-ai-csam-report_update-public-jul24v13.pdf
- Iwanowska, B. (2025). Algorithmic Legitimacy and the Digital State: Rethinking Governance, Trust, and Accountability in the Age of AI. *AI, Law, Politics*, 1(2), 130-149.
- Jasanoff, S., & Kim, S. H. (Eds.). (2019). Dreamscapes of modernity: Sociotechnical imaginaries and the fabrication of power. *University of Chicago Press*.
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389–399. <https://doi.org/10.1038/s42256-019-0088-2>

- Johansen, W., & Frandsen, F. (2007). *Krisekommunikation: Når virksomhedens image og omdømme er truet*. Samfundslitteratur.
- Just, N., & Latzer, M. (2017). Governance by algorithms: reality construction by algorithmic selection on the Internet. *Media, culture & society*, 39(2), 238-258.
- Kazoleas, D., & Teigen, L. G. (2006). The technology-image expectancy gap: A new theory of public relations. *Public relations theory II*, 415-433.
- Kerr, A., Barry, M., & Kelleher, J. D. (2020). Expectations of artificial intelligence and the performativity of ethics: Implications for communication governance. *Big Data & Society*, 7(1), 2053951720915939.
- Kim, S., & Ji, Y. (2018). Gap analysis. *The international encyclopedia of strategic communication*, 8, 1-6.
- Kim, T., & Song, H. (2023). Communicating the limitations of AI: the effect of message framing and ownership on trust in artificial intelligence. *International Journal of Human-Computer Interaction*, 39(4), 790-800.
- Kurian, N. (2025). “No, Alexa, no!”: Designing child-safe AI and protecting children from the risks of the “empathy gap” in large language models. *Learning, Media and Technology*, 50(4), 621–634. <https://doi.org/10.1080/17439884.2024.2367052>
- Livingstone, S. M., & Blum-Ross, A. (2020). *Parenting for a digital future: How hopes and fears about technology shape children's lives*. Oxford University Press.
- Livingstone, S., & Stoilova, M. (2021). *The 4Cs: Classifying online risk to children*.
- LLM Stats. (n.d.). *LLM Stats*. <https://llm-stats.com/>
- Madden, M., Calvin, A., Hasse, A., & Lenhart, A. (2024). *The dawn of the AI era: Teens, parents, and the adoption of generative AI at home and school (Report)*. Common Sense Media.

- Mao, Y., Richter, V., & Katzenbach, C. (2025). Strategising imaginaries: How corporate actors in China, Germany and the US shape AI governance. *Big Data & Society*, 12(4), 20539517251400727.
- Monteith, S., Glenn, T., Geddes, J. R., Whybrow, P. C., Achtyes, E., & Bauer, M. (2024). Artificial intelligence and increasing misinformation. *The British Journal of Psychiatry*, 224(2), 33-35.
- Musk, E. [@elonmusk]. (2025, July 20). *[We're going to make Baby Grok @xAI, an app dedicated to kid-friendly content] [Post]. X (formerly Twitter).*
<https://x.com/elonmusk/status/1946763642231500856>
- Napoli, P. M. (2019). Social media and the public interest: Media regulation in the disinformation age. *Columbia university press*.
- Nguyen, D., & Hekman, E. (2024). The news framing of artificial intelligence: a critical exploration of how media discourses make sense of automation. *AI & society*, 39(2), 437-451.
- NSPCC. (2025). *Viewing generative AI and children's safety in the round*. NSPCC Learning.
- Ocal, A., & Crowston, K. (2024). Framing and feelings on social media: the futures of work and intelligent machines. *Information Technology & People*, 37(7), 2462-2488.
- OpenAI. (2025a, August 26). *Helping people when they need it most*.
<https://openai.com/index/helping-people-when-they-need-it-most/>
- OpenAI. (2025b, November 6). *Introducing the Teen Safety Blueprint*.
<https://openai.com/index/introducing-the-teen-safety-blueprint/>
- OpenAI. (2025c). *A family guide to help teens use AI responsibly*.
<https://cdn.openai.com/pdf/a-family-guide-to-help-teens-use-ai-responsibly.pdf>

- OpenAI. (2025d). *Tips for talking to your teen about AI*.
<https://cdn.openai.com/pdf/tips-for-talking-to-your-teen-about-ai.pdf>
- Osório de Barros, G., & Severino Soares, O. (2025). AI and the next generation: Protecting childhood in the digital age. In *Environmental, social, governance and digital transformation in organizations* (pp. 103-133). Cham: Springer Nature Switzerland.
- Pederson, J. (2025, October 20). Sam Altman's teen suicide problem. *Los Angeles Times*. <https://www.latimes.com/opinion/story/2025-10-20/sam-altmans-teen-suicide>
- Prasad, P. (2017). *Crafting qualitative research: Beyond positivist traditions*. Routledge.
- QuestMobile研究院. (2026, March 3). 2025年AI应用层发展核心报告:原生App巨头筑墙, 新锐破局, AI核心入口雏形初现, 入口争夺战持续 [2025 AI Application Layer Development Core Report: Native App Giants Build Walls, Newcomers Break Through, Core AI Entry Point Emerges, and the Battle for Entry Points Continues]. *QuestMobile*.
<https://www.questmobile.com.cn/research/report/2028739596590878721>
- Raupp, J. (2019). Crisis communication in the rhetorical arena. *Public Relations Review*, 45(4), 101768. <https://doi.org/10.1016/j.pubrev.2019.04.002>
- Rasser M, Lamberth M, Riikonen A, Guo C, Horowitz M, Scharre P (2019) *The American AI Century: a blueprint for action*. Center for a New American Security.
<https://www.cnas.org/publications/reports/the-american-ai-century-a-blueprint-for-action>
- Regalado, F. (2025, September 30). What we know about ChatGPT's new parental controls. *The New York Times*.
<https://www.nytimes.com/2025/09/30/technology/chatgpt-teen-parental-controls-openai.html>

- Riolo, M., Di Nuovo, S., & Maganuco, N. R. (2025). Emotional Intelligence and Adolescents' Use of Artificial Intelligence: A Parent–Adolescent Study. *Behavioral Sciences*, 15(8), 1115. <https://www.mdpi.com/2076-328X/15/8/1142>
- Rotolo, D., Hicks, D., & Martin, B. R. (2015). What is an emerging technology?. *Research policy*, 44(10), 1827-1843.
- Ruane, E., Birhane, A., & Ventresque, A. (2019). Conversational AI: Social and Ethical Considerations. *AICS*, 2563, 104-115.
- Sartori, L., & Bocca, G. (2023). Minding the gap (s): public perceptions of AI and socio-technical imaginaries. *AI & society*, 38(2), 443-458.
- Salehi, A., Motaghi, H., Hashemi, M. S., & Fahim, Y. (2026). Public Trust and Legitimacy in the AI Era: A Systematic Review of Public Media Organizations. *Journal of Communication and Management*, 5(01), 72-84.
- Schultz, M. D., Conti, L. G., & Seele, P. (2025). Digital ethicswashing: A systematic review and a process-perception-outcome framework. *AI and Ethics*, 5(2), 805-818.
- Scott, J. (1990). A matter of record: Documentary sources in social research. *Polity Press*.
- Shao, S. (2025). The role of AI tools on EFL students' motivation, self-efficacy, and anxiety: Through the lens of control-value theory. *Learning and Motivation*, 91, 102154. <https://doi.org/10.1016/j.lmot.2025.102154>
- Singer, N. (2025, May 2). Google plans to roll out its A.I. chatbot to children under 13. *The New York Times*.
<https://www.nytimes.com/2025/05/02/technology/google-gemini-ai-chatbot-kids.html?searchResultPosition=2>

- Singh, J., & Bajwa, A. (2026, February 13). Musk's AI chatbot Grok's U.S. market share jumps amid sexualized images backlash, data shows. *Reuters*.
<https://www.reuters.com/business/media-telecom/musks-ai-chatbot-groks-us-market-share-jumps-amid-sexualized-images-backlash-2026-02-13/>
- Spehar, D. (2025, May 12). Google's Gemini AI kids edition is here: What it means for parents. *Forbes*.
<https://www.forbes.com/sites/dianaspehar/2025/05/12/googles-gemini-ai-kids-edition-is-here-what-it-means-for-parents/>
- Steinhardt, H. C. (2012). How is High Trust in China Possible? Comparing the Origins of Generalized Trust in Three Chinese Societies. *Political Studies*, 60(2), 434-454.
- Suerdem, A., Akkilic, S. (2021). Cultural Differences in Media Framing of AI. In: Schiele, B., Liu, X., Bauer, M.W. (eds) *Science Cultures in a Diverse World: Knowing, Sharing, Caring*. Springer, Singapore. https://doi.org/10.1007/978-981-16-5379-7_10
- Szmyd, K., & Mitera, E. (2024). The impact of artificial intelligence on the development of critical thinking skills in students. *European Research Studies Journal*, 27(2), 1022-1039.
- TechRadar. (2025). Google is working on a Gemini AI app for kids.
<https://www.techradar.com/computing/artificial-intelligence/google-is-working-on-a-gemini-ai-app-for-kids>
- The White House. (2025, April 23). *Advancing artificial intelligence education for American youth [Presidential action]*.
<https://www.whitehouse.gov/presidential-actions/2025/04/advancing-artificial-intelligence-education-for-american-youth/>
- Thomas, D. R. (2003, August). *A general inductive approach for qualitative data analysis*.

- Ulnicane, I., Knight, W., Leach, T., Stahl, B. C., & Wanjiku, W. G. (2021). Framing governance for a contested emerging technology: insights from AI policy. *Policy and Society*, 40(2), 158-177.
- UNICEF Innocenti. (2025). *Guidance on AI and children*. UNICEF Innocenti – Global Office of Research and Foresight.
- Van Aelst, P., & Walgrave, S. (2016). Information and arena: The dual function of the news media for political elites. *Journal of Communication*, 66(3), 496-518.
- Wang, J., & Fan, W. (2025). The effect of ChatGPT on students' learning performance, learning perception, and higher-order thinking: Insights from a meta-analysis. *Humanities and Social Sciences Communications*, 12, 621. <https://doi.org/10.1057/s41599-025-04787-y>
- Wang, S., & Liang, Z. (2024). What does the public think about artificial intelligence? An investigation of technological frames in different technological context. *Government Information Quarterly*, 41(2), 101939.
- Wang Y, Liu W, Yu X, Li B, Wang Q (2024) The impact of virtual technology on students' creativity: a meta-analysis. *Comput Educ* 215. <https://doi.org/10.1016/j.compedu.2024.105044>
- Wong, J. C. (2025, October 2). My son genuinely believed it was real: Parents are letting little kids play with AI. Are they wrong? *The Guardian*. <https://www.theguardian.com/technology/ng-interactive/2025/oct/02/ai-children-parenting-creativity>
- xAI LLC. (2025, November 4). *Terms of service*. <https://x.ai/legal/terms-of-service/>
- Xu, Y., Thomas, T., Yu, C.-L., & Pan, E. Z. (2025). What makes children perceive or not perceive minds in generative AI? *Computers in Human Behavior: Artificial Humans*, 4, 100135. <https://doi.org/10.1016/j.chbah.2025.100135>

- Yuan, J. E. (2025). DeepSeek and the Shifting AI Landscape in China and the US. *AsiaPacific Issues*, 29(174).
- Yu, Y., Sharma, T., Hu, M., Wang, J., & Wang, Y. (2025). Exploring parent-child perceptions on safety in generative AI: Concerns, mitigation strategies, and design implications. In *IEEE Symposium on Security and Privacy 2025* (pp. 2735–2752).
<https://doi.org/10.1109/SP61157.2025.00090>
- Zeng, J., Chan, C. H., & Schäfer, M. S. (2020). Contested Chinese dreams of AI? Public discourse about artificial intelligence on WeChat and People's Daily Online. *Information, Communication & Society*, 25(3), 319-340.
- Zhang, J., Ki, E. J., & Cai, Q. (2026). Legitimizing AI: ethical AI in corporate ethics statements. *Corporate Communications: An International Journal*, 31(1), 21-36.
- Zhou, J., Zhang, Y., Luo, Q., Parker, A. G., & De Choudhury, M. (2023, April). Synthetic lies: Understanding ai-generated misinformation and evaluating algorithmic and human solutions. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems* (pp. 1–20).
- 陆玖商业评论. (2025, December 15). 这届家长，为何用豆包看娃？[Why are parents using Doubao to watch their children?]. 雪球.
- 一颗青木. (2026, April 23). 不要太迷信豆包，不然会出事. [Don't be too superstitious about Doubao, or something bad might happen]. 远方青木微信公众号.
<https://mp.weixin.qq.com/s/-4teUvN4pdKl3X16V7lV9Q>
- 云分. (2026, March 1). 90%的人都浪费了豆包等AI工具的真正用法，这篇文章让你看完效率翻倍！[90% of people are wasting the true potential of AI tools like Doubao. This article will double your efficiency after you read it!] 云分职场笔记微信公众号. <https://mp.weixin.qq.com/s/4QiREZddWM7Enl7FC-5pAQ>
- 张格格. (2025, August 25). AI学习还是无限刷视频？豆包回应视频推送无法关闭争议 [AI learning or endless video scrolling? Doubao responds to controversy over

unclosable video recommendations]. 凤凰网科技

<https://tech.ifeng.com/c/8m6i6dU9OGb>

周小白. (2025, August 28). 豆包里的短视频关不掉引争议, 上线未成年人模式靠谱吗? [The inability to turn off short videos on Doubao has sparked controversy; is launching a mode for minors a reliable solution?]. *TechWeb*微信公众号.

<https://mp.weixin.qq.com/s/vHNwkpdr1DOi7ZKCWsiGtA>

Appendix

Appendix A. Thematic Coding Table for RQ1: Corporate Framing of AI for Minors

THEME	ChatGPT codes	Gemini codes	Grok codes	Doubao codes
Empowering to Meet Development Needs	AI access as developmental entitlement	treat teens differently from adults	personalised parental control	
	child-impact awareness	AI as learning and creativity support		
	teen-specific design logic			
	treat teens differently from adults			
	personalised parental control			
Diffuse Accountability	parental dialogue as safeguard	parental gatekeeping of access	parental gatekeeping of access	parental gatekeeping of access
	ongoing parental guidance	parental co-use and supervision	parental supervision as safeguard	parental supervision as safeguard
	parental monitoring of emotional risks	parental instruction of responsible use		guardian responsibility for outcomes
				parental instruction of responsible use
				parental control over minors' data
				disclaimer of real-world outcome responsibility
				parental activation of Kids Mode
				adult-status fallback outside Kids Mode

Risk Acknowledged but Managed	abnormal activity monitoring	possible Inappropriate content	possible Inappropriate content	minor privacy protection claim
	U18-specific policy need	anthropomorphic interaction risk		addiction / overuse risk
	Expert- and policy-informed protections	factual issues sometimes		factual issues sometimes
	default safety protections for minors	Risk managed through filtering		
	factual issues sometimes			
Corporate Responsibility Commitment	safety prioritised over autonomy and privacy			iterative well-being improvement
	iterative well-being improvement			
	platform duty to protect minors			
Protection through Functional Restriction				Safety through restriction
				Recommended protective mode for teens
				Guidance as behavioural regulation

Appendix B. Thematic Coding Table for RQ2: Public Discourse on AI for Minors

THEME	ChatGPT codes	Gemini codes	Grok codes	Doubao codes
Superficial Adaptation	developmental and education support expectation	lack of age-appropriate adaptation		
	ignore needs of age-segment	adult-oriented adaptation logic		
		inferior to child-specific platforms		

		developmental-stage mismatch		
Profit over Protection	agreeable design to drive engagement	default enrolment into AI features	profited from nonconsensual imagery.	technology improvement expectation
	freedom prioritised over safety		push-notification re-engagement	Prioritized profit over minor protection
	encourages continued engagement		compulsive companion interaction	low-barrier access
	child best interests sidelined		conversation extension over teen support	
			profit-oriented dependency design	
			X integration enabling non-consensual sexual imagery	
			controversy leveraged for attention	
			unhealthy attachment and explicit-content promotion	
Multidimensional Harms to Minors	uncertain developmental effects	difficulty recognizing artificial agents	impaired critical thinking development	AI perceived as beneficial for learning (positive)
	decrease diversity of creative expression	blurred boundary with real life	distorted worldview risk	AI easing parenting burden (positive)
	emotional attachment risk	acting as emotional support provider	unhealthy attachment and explicit-content promotion	easy emotional attachment to AI
	risks of intimacy with AI	manipulable child trust	companion displacement of real relationships	entertainment displacement through short-video features
	algorithmic support without personalised understanding	psychological dependence on AI	reality erosion through X integration	
	illusory AI understanding	inappropriate interaction refusal (positive)	uncensored erotic roleplay	
	indirect romantic content exposure	bias-speech avoidance (positive)	delusion amplification and unsafe affirmation	
	stereotype reinforcement	easy access to inappropriate material	unsafe and inappropriate content	
	misinformation risk	inconsistent safety	misinformation and	

		filtering	conspiracy-theory risk	
	insufficient crisis intervention	possible misinformation exposure	stereotype and bias reinforcement	
		habitual-behaviour influence	gendered harassment and exclusion	
		manipulation concerns		
		vulnerability to misinformation		
		shortcut-taking by teenagers		
		child-disproportionate environmental burden		
		unresolved voice-data sensitivity		
		privacy concern		
Needs for Safeguard	safeguard degradation in long interactions	AI should framed as non-human	weak age verification for adult content	unable to choose whether receive short video recommendations
	stronger age-check demand	failure to redirect to adults	ineffective existing safeguards	parental irreplaceability claim
	default-by-design protection	insufficient parental supervision	self-reported age circumvention	
	lack of safeguarding measures	weak safeguarding structures	failure to infer teen status from context	
	clearer roleplay and storytelling boundaries	unclear human review process	misinformation risk requiring safeguards	
	failure to detect mental-health warning patterns	Unclear internal data controls	safety-boundary erosion over time	
	lack of crisis-intervention mechanisms	insufficient disclosure in child safety metrics.		
	regulatory intervention demand	school-system integration concerns		
	lack of intervention transparency	responsibility shifted to parents		
		AI companies as primary responsible		

		actors		
Contested Accountability		avoidance of direct accountability		AI can't really replace the parents
		responsibility shifted to parents		
		AI companies as primary responsible actors		
AI as a Parenting Support Tool				customised learning support
				AI as accessible knowledge resource
				emotionally stable interaction
				AI-assisted child supervision

Appendix C. Overview of Data Sources

Date	Source	Author	Title & Link	Company
5/2/2025	NEW YORK TIMES	Natasha Singer	Google Plans to Roll Out Its A.I. Chatbot to Children Under 13	Gemini
5/2/2025	Mashable	Timothy Beck Werth	Google will put Gemini AI in the hands of kids under 13	
5/2/2025	Techradar	Eric Hal Schwartz	Google is working on a Gemini AI app for kids	
5/3/2025	PCMag	Will McCurdy	Your Kids Can Now Use Google's Gemini AI	
5/3/2025	Dan Christ	Dan Christ	Gemini AI for Kids: Why Parents Must Slam the Brakes	
5/12/2025	Forbes	Diana Spehar	Google's Gemini AI Kids Edition Is Here: What It Means For Parents	
5/21/2025	Fairplay	Fairplay	Advocates and child development experts urge Google to halt Gemini AI chatbot rollout to young children, request FTC investigation	
5/21/2025	EPIC&Fairplay	EPIC&Fairplay	Re: Potential COPPA Violations in Google's Rollout of AI Chatbot Gemini to Children	
9/5/2025	Common Sense Media	Common Sense Media	Common Sense Media AI Risk Assessment: Gemini Under 13	Grok
1/7/2026	WIRED	Matt Burgess & Maddy Varner	Grok Is Generating Sexual Content Far More Graphic Than What's on X	

1/9/2026	the Guardian	Moira Donegan	Grok is undressing women and children. Don't expect the US to take action	
1/22/2026	Common Sense Media	Common Sense Media	Grok Common Sense Media	
1/27/2026	Education Week	Alyson Klein	'Grok' Chatbot Is Bad for Kids, Review Finds	
4/8/2026	NEW YORK TIMES	Jessica Grose	Deepfake Nudes Are Haunting America's Teens	
8/26/2025	NEW YORK TIMES	Kashmir Hill	A Teen Was Suicidal. ChatGPT Was the Friend He Confided In.	ChatGPT
8/27/2025	CNN Business	Clare Duffy	Parents of 16-year-old sue OpenAI, claiming ChatGPT advised on his suicide	
9/2/2025	NEW YORK TIMES	Kashmir Hill	OpenAI Plans to Add Safeguards to ChatGPT for Teens and Others in Distress	ChatGPT
9/2/2025	the Guardian	Robert Booth	Parents could get alerts if children show acute distress while using ChatGPT	
9/30/2025	NEW YORK TIMES	Francesca Regalado	What We Know About ChatGPT's New Parental Controls	
10/2/2025	the Guardian	Julia Carrie Wong	'My son genuinely believed it was real': Parents are letting little kids play with AI. Are they wrong?	
10/20/2025	L.A. Times	Joshua Pederson	Sam Altman's terrible reason for letting ChatGPT talk to teens about suicide	
10/23/2025	Common Sense Media	Common Sense Media	ChatGPT Common Sense Media	
1/22/2026	Mashable	Rebecca Ruiz	How ChatGPT ends up in children's toys	
2/24/2026	NEW YORK TIMES	Natasha Singer	More Than Half of Teens Use Chatbots for Schoolwork, Survey Finds	
4/4/2026	NEW YORK TIMES	Kashmir Hill	What Teens Are Doing With Those Role-Playing Chatbots	Doubao
8/25/2025	凤凰网科技	张格格	AI学习还是无限刷视频？豆包回应视频推送无法关闭争议	
8/26/2025	澎湃新闻	范佳来	未成年人如何健康使用AI？字节豆包率先上线未成年人模式	
8/28/2025	TechWeb	周小白	豆包里的短视频关不掉引争议，上线未成年人模式靠谱吗？	
12/15/2025	雪球	陆玖商业评论	这届家长，为何用豆包看娃？	
3/10/2026	中文科技资讯	慢放	AI教育喊了十年，豆包做了什么？	
3/31/2026	云分职场笔记	云分	90%的人都浪费了豆包等AI工具的真正用法，这篇文章让你看完效率翻倍！	
4/23/2026	远方青木	一颗青木	不要太迷信豆包，不然会出事	