

Mapping Market Risk Narratives of Artificial Intelligence: A Social Network Analysis of Public Discourse

Kirana Labbaika*, Vanessa Gaffar, Asep Miftahuddin

Faculty of Economics and Business Education, Universitas Pendidikan Indonesia, Bandung, Indonesia

*Corresponding Email: kirana.labbaika15@upi.edu

Abstract. As AI technologies rapidly progress, privacy and surveillance, cybersecurity, and governance have become increasingly pressing concerns that raise questions about how the narratives of risk and AI influence the market's perception of new technologies. As a component of the public discourse, it is becoming more and more a sign of trust or distrust in technology and a source of legitimization in the digital world. This study utilized a Social Network Analysis (SNA) and Text based analytical method to analyze the perception of the market on AI by analyzing the AI risk narratives in the online public discourse. Data were gathered from X (formerly Twitter) using the SocialX platform over the first quarter of 2026 (1st January-31st March 2026) and found to be 2,883 public posts. The analysis combined trend analysis, sentiment analysis, temporal word analysis, text network analysis, BERTopic topic modeling, zero-shot classification, and SNA. The results showed an attitude of caution, with neutral (57.86%) and negative sentiment (35.93%) dominating the AI risk discourse, suggesting that the public isn't entirely positive about AI. Most prevalent data privacy issues were Data Privacy Risk (43.81%), Surveillance and Control Risk (21.19%) and Cybersecurity Risk (11.24%). Text network analysis further identified data, training, and surveillance as central themes within the discourse. This study contributes theoretically by integrating Framing Theory, Perceived Risk Theory, and market legitimacy perspectives to explain AI risk narratives as indicators of market perception. In practical terms, the results offer insights for trust-building in AI, responsible communication, and technology branding strategies.

Keywords: Artificial Intelligence, Social Network Analysis, Market Perception, Perceived Risk, Technology Marketing, Public Discourse, AI Trust

1 Introduction

Artificial Intelligence (AI) has revolutionized the interaction between organizations, industries, and consumers with the digital technology. AI's image as just a technological improvement has developed into its role as a strategic framework impacting busi-

ness procedures, productivity, consumer decision making and service personalization. AI in tech marketing is seen as a solution to increase productivity, customer satisfaction and cost-efficiency [1, 2]. However, there have been several concerns in the public sphere regarding the risks of AI, such as privacy concerns, digital surveillance, misinformation, algorithmic bias, and job replacement. The uptake of AI in the market isn't just about the benefits; however, it's also about the dangers people feel about using AI [3, 4].

In the digital age, digital public spaces are a key space that will shape the market's perception of technology. The X (formerly Twitter), as a social media, plays a crucial role in sharing opinions, setting the agenda of an issue and disseminating narratives on new technologies. While consumer perception surveys represent a more conventional marketing strategy, digital conversations on social media can leave a trail of behaviors that are the spontaneous reaction of the public to a technology [5–7]. Public opinion of AI can therefore no longer be considered as a mere opinion but as a market signal that mirrors the degree of anxiety, trust, legitimacy and resistance towards AI.

Increased adoption of AI in several industries is making this an ever-relevant phenomenon, but it is not always the case that people trust AI more as it becomes more commonplace. There have been many reports of data breaches, the use of AI-powered surveillance, manipulation of information through generative technologies, and algorithmic discrimination, which have led to growing distrust of AI technology [8–10]. Trust building and responsible governance are as much a problem as technological innovation in the future development of AI. This demonstrates a gap between the ideal condition, AI as a widely accepted and trusted technology, and the empirical condition of intensifying risk narratives that have the potential to undermine market legitimacy for AI [11, 12]. From a theoretical perspective, the perception of risk can play a crucial role in the acceptance of new technologies by consumers. Uncertainty and negative expectations of a product or service can often affect consumer behavior. From an AI perspective, perceived threats can be privacy threats, social threats, psychological threats, or functional threats [13, 14]. The public's strong engagement with such narratives on topics of personal data manipulation or digital surveillance can increase risk perceptions, which in turn might impact trust and market uptake of AI. In other words, public sphere risk narratives can be seen as early signals of the trajectory of market perception formation about AI [4, 15].

There has been significant research about public perception of AI in recent years. The majority of past studies have concentrated on the extent to which individuals accept a technology, have intentions to use it, and feel about the psychological consequences of using AI. Perceived usefulness and ease of use are the variables that are generally used to explain technology acceptance. Furthermore, other studies have shown that trust is an important factor in the acceptance of AI-powered technologies, particularly those that require the sharing of personal information and involve automated decision-making. Furthermore, a few works in the digital marketing field have started identifying the link between risk perception and consumers' attitudes towards intelligent technologies, specifically those that rely on automation and recommendation systems [13, 16].

On the other hand, studies on AI in social media have focused more on sentiment analysis, detection of public opinion, or evaluation of general perceptions of AI technological developments. The conversation around AI in society is one that is both ecstatic and worrying. Nevertheless, most of these studies still use descriptive statistical approaches, sentiment analysis, or individual perception surveys, and have yet to explain how risk narratives are structurally formed within the digital public sphere. In fact, in a highly connected digital communication environment, risk issues do not emerge separately, but are interrelated, forming certain relationship patterns that can collectively influence market perception[17, 18].

Research based on Social Network Analysis (SNA) has begun to be used to understand digital communication structures, information dissemination patterns, and relationships among actors on social media. SNA enables researchers to map relational structures within communication networks to identify central nodes, connectivity patterns, and the formation of discursive communities[19, 20]. In the technological context, several studies have used SNA to map the dissemination of information about digital innovation, including perceptions of AI. However, studies that specifically map the interconnectedness of AI risk narratives as signals of market perception are still relatively limited. Previous research has generally focused on user networks, influential actors (influencers), or the dissemination of technology topics, without linking the structure of risk narratives to marketing implications such as technology legitimacy, trust-building, and barriers to market acceptance[21, 22].

Public responses to AI can also offer a way of side-lining the expectations of the markets regarding new technologies, as they represent social reactions. Risks surrounding technologies in digital environments are routinely amplified, challenged and normalized by repeated interactions between users, institutions and the technology-related communities. This talk is likely to influence the public perception of AI, from it as a tool of productivity to it as a privacy risk or a tool of digital control, etc.[23, 24]. This means that the conversation online can shape the overall market sentiment, either by fostering trust or doubt in AI or opposing the use of AI. Knowing these dynamics becomes more crucial, especially in the case of perceptions that could shape consumer behavior and the acceptance of technology before actual changes are measurable [25–27].

Meanwhile, the growing complexity of discourse on AI necessitates more sophisticated analytical methods that are able to capture the nature of both the content of discourse and the interconnections of concerns within a larger communication structure[28, 29]. Relational analysis of AI risk narratives allows for a more holistic view of risk issues gaining salience, themes converging, and key actors driving narratives. Social Network Analysis offers an opportunity to shift from sentiment analysis to an analysis of structure of digital discourse, as a way to gain more insights into how public concerns around AI can become indicators of marketers' perception and technology legitimacy [30, 31].

From this review, gaps in knowledge are identified that must be sought to be filled. First, past studies tend to isolate the topic of AI-related risk perception from the social media public discourse. Second, most research on AI from an "information behaviour" viewpoint remains based on surveys of consumers' information behaviour and

therefore lacks the ability to capture the "natural" dynamics of the perception of the market which naturally occur when faced with the digital world. Third, SNA research related to AI generally has not explained how the interconnectedness among risk narratives forms market signals regarding the legitimacy of AI technology. Thus, there remains a gap of understanding with respect to the heart of the risks to public discourse, how these risks are interrelated, and the nature of the relationships between these risks with respect to market uptake of AI.

In this study, the novelty lies in the use of public discourse as an indicator of market perception of the technology of AI, in the context of a Social Network Analysis approach. This research, on the contrary, focuses on the structure of the networked narratives of AI risks in the digital public sphere in order to understand how they are collectively and gradually articulated and constructed. This study applies the theoretical frameworks of Framing Theory to gain insight into the framing of AI as a risk object in public discourse, and Perceived Risk Theory to explore how these framings might shape market perceptions of AI. Furthermore, market legitimacy is introduced as an explanation of why the acceptance of AI technology is not a matter of technical skills alone, but also of social legitimacy developed through public communication.

This study treats online public discourse as a proxy for market perception, rather than general public opinion, on three grounds. First, from a markets-as-conversation and signaling perspective, markets for emerging technologies are constituted not only by transactions but by the circulation of information, expectations, and evaluations among audiences; discourse on X is observable evidence of this collective evaluation, especially for technologies whose value is still being socially negotiated. Second, AI-discourse participants are not a representative population sample but a self-selected, digitally engaged audience, adopters, developers, enterprise users, journalists, policy actors, and opinion leaders, who both shape and are shaped by the market in which AI products compete. Third, digital-marketing research treats user-generated content as behavioral-trace data that can anticipate adoption, brand attitudes, and demand, since expressed concerns and framings precede and condition usage decisions. AI risk discourse is therefore read as an early, aggregate market-perception signal of the trust, legitimacy, and resistance conditions under which AI is received, not as a direct measure of consumer intent. This proxy has explicit boundaries: it indexes an engaged digital public and relates to, but does not confirm, downstream market activity (see Limitations).

The study integrates Framing Theory[32] and Perceived Risk Theory[33] analytically rather than causally. Framing Theory explains how certain AI attributes, privacy, surveillance, security, become salient in discourse, while Perceived Risk Theory holds that greater salience of negative consequences raises perceived risk and conditions technology adoption. The linkage is examined at the aggregate-discourse level, not individual cognition: the presence and comparative salience of specific risk frames via zero-shot classification and topic modeling, their structural centrality via text-network analysis, and their salience over time via temporal word and sentiment analysis. That privacy, surveillance, and cybersecurity frames occupy the most central and frequent positions while co-occurring with neutral-to-negative sentiment provides evidence of association between dominant frames and a cautious risk orientation in

the aggregate signal. The study does not claim individual-level causation from frame exposure to any consumer's risk perception; establishing such causation is left to future research.

Because the analysis interprets discourse through the lens of market legitimacy, it is important to specify which dimension is engaged. Legitimacy comprises three related subtypes: pragmatic (audience self-interest and the technology's practical value or safety), moral (normative judgments about whether the technology and its governance are right), and cognitive (taken-for-grantedness and comprehensibility)[34]. This study focuses on pragmatic and moral legitimacy. Data Privacy and Cybersecurity frames reflect pragmatic concerns, whether AI benefits audiences without harming them, while Surveillance and Control and Governance and Trust frames reflect moral concerns about the acceptability and governance of AI deployment. Cognitive legitimacy is engaged only indirectly, through the high share of neutral, informational discourse and the fragmented topic structure, indicating that AI is still being collectively defined and made comprehensible. Framing the analysis this way shows that the observed patterns reflect audiences' legitimacy stakes, the interests AI serves and the norms of its governance, rather than undifferentiated public concern.

This study aims to identify the dominant AI risk narratives, map the structural interconnectedness of the various risk narratives in the public sphere, and analyse the narratives as signals of a market perception of AI. The results should be applicable to the study of technology marketing, digital consumer behaviour and also provide implications for organisation and AI developers in developing communication strategy, trust building and AI based brand legitimacy.

Accordingly, this study addresses the following research questions. RQ1: What are the dominant risk narratives emerging in public discourse surrounding AI on X/Twitter? RQ2: How are AI risk narratives structurally interconnected within the digital public sphere, and which actors or themes occupy central positions in the discourse network? RQ3: How do public sentiment and thematic patterns reflect market perceptions toward AI-related risks? RQ4: To what extent do dominant AI risk narratives function as market signals associated with trust, legitimacy, and public acceptance of AI technologies? By addressing these questions, this study seeks to provide a more comprehensive understanding of how AI risk discourse evolves in digitally networked environments and how such discourse may shape broader market perceptions toward emerging technologies.

2 Method

The methodology is based on social media analysis, a quantitative method, and the narratives of AI risks in digital public spaces were analyzed based on the conversations on the X (formerly Twitter) social network. To understand how risk narratives are formed, evolve, and influence market perceptions of AI, this study combines several analytical techniques, namely trend analysis, sentiment analysis, sentiment index, word frequency analysis, text network analysis, Social Network Analysis (SNA), temporal word analysis, BERTopic topic modeling, and zero-shot classification. The

methodological combination was chosen due to its capacity to comprehensively explain the structure of interaction, conversational dynamics and the substance of AI risk narratives. The analysis of the network is used to gain insight into the relationship patterns between actors and concepts, and text analysis is used to determine the sentiments, themes and emerging risk perceptions in the public sphere[35, 36]. X (formerly Twitter) data was used through the SocialX social media analytics platform to gather, process, and visualize digital conversation data. The data collection took place from January 1 to March 31, 2026, and involved gathering of public posts concerning AI and the associated issues of risks. The search process included using keywords and hashtags that were related to AI, including terms like AI, artificial intelligence, generative AI, ChatGPT, LLM, privacy concern, user data, personal data, training data, data protection, and others. All the data were preprocessed before analysis to improve the quality and reliability of the data set. In this stage, any duplicate data, hyperlinks, emojis, special characters, symbols and irrelevant text elements were removed. Normalization processes of text such as lowercasing, tokenization, removing stop words, and words without substantive meaning to the research context were then applied.

The preprocessing and analytical techniques reported in this study (data scraping, cleaning of the data, sentiment analysis, temporal word analysis, text network analysis, Social Network Analysis (SNA), BERTopic topic modelling, and zero-shot classification) were done within the SocialX platform, without using separate author-implemented code. The authors set up the data collection query, and interpreted the output and visualization.

It is important to note that this data is limited to a sample of the discourse around AI risks, and not necessarily a representative sample of public opinion. There are 3 scope conditions. First, the data are platform-specific: X/Twitter has a specific demographic and professional profile (most of its users are technologically engaged, professional, and also speak English) and certain features that allow for amplification, which means that it is appropriate to take the results and conclusions of this study with a grain of salt and to refrain from generalizing the findings to other platforms or to non-users of these platforms. Second, the corpus is English-only (100% of posts), meaning that Anglophone discourse with its associated geographic and linguistic biases dominate the corpus and the non-English AI conversation is minimized. Thirdly, retrieval is keyword- and hashtag-driven and risk-oriented; this limits the sample to posts employing the chosen vocabulary. The sample size (2,883 posts from 2,269 unique accounts over a span of three months) is sufficient for the exploratory and structure-mapping goals of this study and the network- and text-analytic methods used, but it is not provided as a probability sample to which inferences about population-level prevalence may be extrapolated. There was no demographic weighting or stratification, because platform data cannot provide reliable demographic attributes. These boundaries are treated explicitly in the interpretation of results and in the Limitations section.

The analysis started with trend analysis which maps the conversation volume dynamics in the observation period. This analysis was employed to derive fluctuations in public attention to the issues of AI risk and to pinpoint the highest discussion intensity periods (spike events). In order to complement the understanding of public percep-

tion, this study then used sentiment analysis techniques to categorize the posts based on whether they were positive, neutral, and negative. The sentiment index then served as a tool to monitor the trend of people's attitude toward AI during the study period.

The next step involved analyzing the content of the discussions using word frequency analysis and the creation of word cloud images to determine the most frequently used words in AI conversations. To account for the changes in issue focus over time, this study used temporal word analysis, in which words that were more frequent during the time period studied were analyzed on the basis of normalization of word counts to specific time periods [36].

Gain insights into the conceptual structure of AI risk narratives, a text network analysis was conducted, in which words are depicted as nodes and co-occurrences as edges. This method enables to identify the semantic relationships between concepts and to create dominant narrative clusters. Furthermore, Social Network Analysis (SNA) was used to create a network of relationships between actors in the digital discourse. These networks consist of nodes corresponding to user accounts and edges representing mentions, replies and reposts between actors. It is applied to detect influential actors, interaction patterns, and community structures in the AI discourse [37].

The actor network was built and plotted in the SocialX platform based on mention, reply and repost relationships between accounts, where the size of the nodes indicates the number of accounts that a single account is connected with (its degree) in the network. The visualization of the platform also helps to break up the graph into segments, which show the existence of different discussions that make up different discussion communities. This enables the identification of structurally prominent actors (those nodes drawn as the biggest and most connected) and community structures, but these are mostly brought to the surface as a visualization; the underlying node-level centrality values (betweenness centrality, eigenvector centrality) and formal community-detection statistics (modularity) are not exported for independent reporting.

In addition, BERTopic topic modeling was employed to delve deeper into the latent themes in the public conversation around AI using sentence embeddings, dimensionality reduction through UMAP, and HDBSCAN clustering algorithms. With the aim of gaining a deeper understanding of the perceptions of risk, the study employed zero-shot classification to classify the posts based on the AI risks without any labeled data. The risk categories were formulated by a literature review of the AI governance field, and modeled after the concept of Perceived Risk theory, covering six areas: Data Privacy Risk, Surveillance and Control Risk, Cybersecurity Risk, Governance and Trust Risk, Training Data and Intellectual Property Risk, and Bias and Fairness Risk. Through this method, the public perception of AI risk as a market sign of technology legitimacy, trust and acceptance can be more easily interpreted in a structured manner [38].

The six risk categories were drawn from the literature in AI governance and the AI trustworthy AI literature, and were a priori defined and then fed into the zero-shot classification module of the SocialX platform that uses a pretrained natural-language-inference model to assign each post to a category without needing labeled training data. The design was selected to produce theoretical interpretability and to assure comparability between the risk dimensions being interested, and to not classify posts

by any category that did not have a corresponding category in the design, the posts were retained under the 'other/unclassified' provision. The classification was therefore performed by the platform's pretrained model, and the validity of the categories is based mainly on the theoretical justifications (face and content validity) provided by the established risk and governance frameworks. We recognise that zero-shot labels are probabilistic and can change depending on the wording of the category descriptions, and that the assignments were not independently re-validated by the authors; this is included in the Limitations section.

The built-in BERTopic module of SocialX was used for topic modeling. Topic modeling was done using the BERTopic module built into SocialX which creates sentence embeddings and reduces the dimensionality of the documents using the UMAP algorithm, followed by clustering using the HDBSCAN algorithm. Since HDBSCAN does not need a fixed number of topics, it generated a great number of very small clusters (164 in total) and documents that had low frequencies were assigned to a noise cluster rather than forced into topics. The analysis in this study relied solely on the output automatically generated by the platform; no clustering steps were configured or executed manually by the authors. Rather than interpreting the raw number of clusters, the output was read at a higher level, with broader thematic families reported below where clusters shared similar meaning. It should be noted that the internal parameters of the platform's pipeline (e.g., the embedding model, UMAP/HDBSCAN settings, and coherence diagnostics) are not fully visible to the user; this reproducibility limitation is acknowledged in the Limitations section.

This study categorized the research process into an integrated workflow comprising data collection, data preprocessing, analytical process and research output to give a clear overview of the analytical sequence. The workflow provides a detailed visual representation of how multiple interconnected techniques were applied to the raw social media content to create structured analytical insights, allowing to systematically explore AI risk narratives, interaction dynamics, public perceptions, and market legitimacy signals in digital discourse.

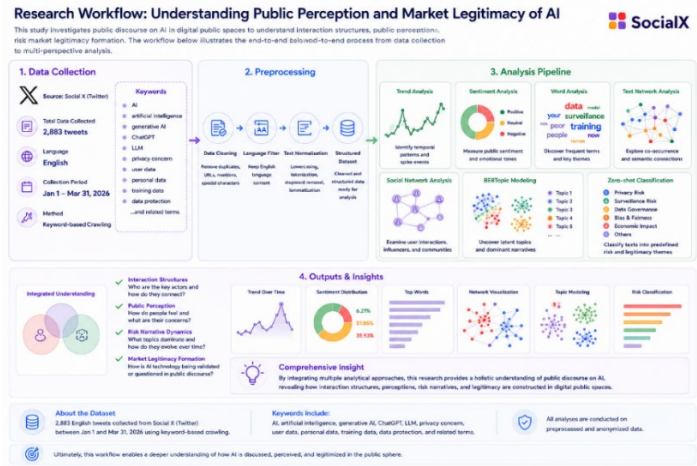


Fig. 1. Research Workflow

Data collection is performed through SocialX, and then, data is preprocessed before performing the various analyses, such as trend analysis, sentiment analysis, word analysis, text network analysis, Social Network Analysis, BERTopic, and zero-shot classification. These stages allow a more comprehensive understanding of the structures of interaction, perception of the public, dynamics of risk narrative, and the processes of building market legitimacy for the AI technology in the public digital sphere.

3 Research Results and Discussion

3.1 Results

Dataset Characteristics.

The research data is collected from platform X (formerly Twitter) through SocialX between January 1–March 31, 2026. The data was gathered with the help of keywords about AI, Large Language Models (LLM), data privacy issues, digital security, surveillance and training data. After a data cleaning process, a total of 2,883 public posts were collected and analyzed.

Before conducting the main analyses, it is important to describe the characteristics of the dataset used in this study. Table 1 provides a summary of the important features of the research data set, including data sources, time of data collection, characteristics of the accounts, language distribution, and types of interactions captured from the public conversation about AI risk.

Table 1. Research Dataset Characteristics

Characteristics	Description / Value
Data source	X (Twitter), via SocialX
Data collection period	1 January – 31 March 2026
Total posts	2,883
Unique user accounts	2,269
Number of attributes	36
Number of languages	1
Dominant language	English , 2.883 posts (100%)
Verified accounts	1.975 posts (68.5%)
Non-verified accounts	908 posts (31.5%)
Type of interaction	Mentions, replies, reposts, quotes
Engagement metrics	Views, likes, replies, reposts, quotes

The research data encompasses 36 attributes such as text information, user metadata, engagement data and inter-account communication relationships. These variables that can be made available include: Posted content (full text), published date, language, mentions, hashtags, account verification, number of followers, and many types of interactions such as likes, replies, reposts, and quotes. Moreover, the dataset also captures relationships between communications like reply relationships and quoted communications, enabling the analysis of content of the communication as well as the interactions between users. The discussions on AI risk appear to have relatively broad

participation with a total of 2,883 posts from 2,269 unique accounts. All posts were detected to be in English (100%). This bounds the dataset to Anglophone discourse and, rather than rendering it globally representative, delimits the scope of the findings to English-language AI risk conversation, particularly issues of data privacy, digital security, surveillance, and the use of training data (see the representativeness note in the Method and the Limitations section). The majority of posts come from verified accounts (68.5%), while the remaining 31.5% are from non-verified accounts. In addition, the dataset also records various forms of interaction such as mentions, replies, reposts, and quotes, which form the basis for constructing communication networks in Social Network Analysis (SNA).

AI Risk Conversation Trends on Social Media.

Trend analysis was conducted to identify the dynamics of conversation volume related to AI risk during the observation period, namely the first quarter of 2026 (January 1–March 31, 2026). The analysis results show that conversation intensity fluctuated throughout the study period with several spike events at certain times. To better understand fluctuations in public attention toward AI risk, Figure 2 presents the conversation volume trend over the observation period.



Fig. 2. AI Risk Conversation Volume Trend

Based on Figure 2, the number of daily conversations shows a relatively fluctuating pattern throughout the observation period. From early January to mid-February, the intensity of conversations tended to remain at a moderate level. Conversations started to ramp up towards the end of February and reached its highest point in mid-March 2026, with more than 400 posts per day. The 7-day moving average pattern suggests that there is an increasing trend in discussion intensity towards the end of the study period. The results indicate that the discourse on the risks of AI evolved over time during the first quarter of 2026.

Structure of the Actor Network in AI Risk Discourse.

Understand the interaction of actors between conversations regarding AI risk, Social Network Analysis (SNA) was performed. The network was constructed based on relationships of mentions, replies, and reposts between users. The overall communication network on the AI platform (Figure 3) demonstrates the interaction structure of actors in AI risk discussions.

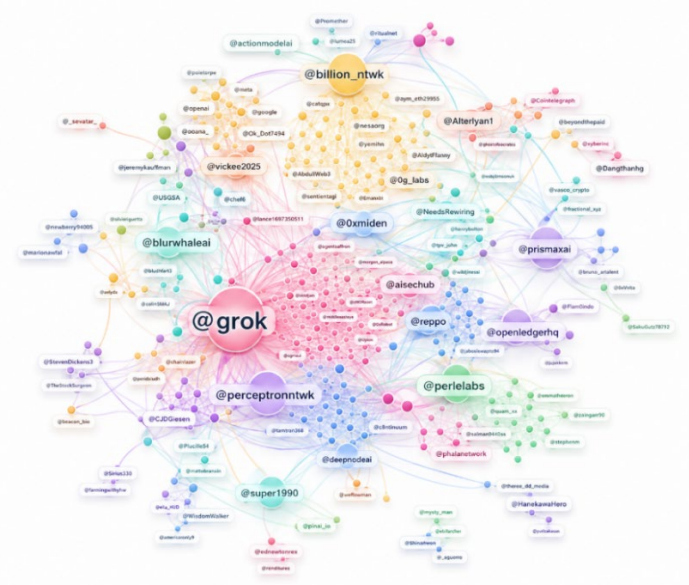


Fig. 3. Social Network Structure of AI Risk Conversation

The visualization results show that the conversation network forms several inter-connected community clusters. The size of the nodes suggests how well interconnected the accounts are in the network. In the network there are several accounts that seem to be at the center, with greater sizes than other accounts. For instance, accounts like @grok, @billion_ntwk, @prismaxai, and @perceptronntwk seem to have moderately high connectivities when it comes to AI risk conversations. The visualization also illustrates that the discussion is not conducted within one uniform community, but is spread across various discussion communities that have distinct interaction patterns.

The graph organizing into a few distinct community clusters suggests that the discourse is structured into various segments, with some accounts, like @grok, @billion_ntwk, @prismaxai and @perceptronntwk, being the most connected and active in the network. The thematic prominence of the privacy-, surveillance-, and data-related risks in the semantic (text) network imply these risk narratives not only influenced the discourse at its margins, but also held significant positions.

Public Sentiment toward AI Risk.

Sentiment analysis was carried out to detect public opinion trends on the topic of AI risk. Emotions were categorized as being positive, neutral, or negative. To gain deeper insights into public attitudes about AI risk, Figure 4 shows the proportion of sentiment categories in the conversations.



Fig. 4. Distribution of Sentiment in AI Risk Conversations

Based on the analysis results, it is shown that neutral sentiment is the dominant sentiment in the conversation, which is 1,668 posts representing 57.86% of the total posts, negative sentiment is 1,036 posts representing 35.93% of the total posts, and positive sentiment is 179 posts representing 6.21% of the total posts. The vast majority of the conversations about AI risk are neutral information or discussion, and the amount of negative sentiments is higher than that of positive.

With the majority of neutral sentiment, it appears AI risk conversation was largely informational and issue-oriented, and not so much emotional responses to AI incidents or developments but rather discussions about data governance, privacy issues, and other technological matters. The relatively high share of negative sentiment, however, suggests that there are significant concerns within the public about the potential impacts of AI adoption, especially concerning privacy, surveillance, and digital security issues. In contrast, the relatively low proportion of positive sentiment indicates that the positive stories about AI and its potential benefits did not dominate the discussion of risks over the observation period.

However, some methodological considerations are appropriate when making sense of the overwhelmingly neutral sentiment. Neutral classification does not necessarily mean that it is cautious or sceptical, but rather the content of the information, the description or the news is devoid of any particular affective connotation. Without the neutral share, therefore, this study does not imply caution. Rather, the characterization of 'cautious acceptance' is based on the pattern of three signals: the very low percentage of positive sentiment (6.21%), the high percentage of negative sentiment

(35.93%), and the high clustering of the discourse around the themes of risk, namely, data privacy, surveillance, and security, with much of the discourse being neutral or informational. In other words, even though the conversation is negative to positive, caution is inferred, and even though the conversation is high on the topics of risk, the large neutral share is interpreted more conservatively as issue-driven exchanges of information without skepticism. A finer-grained stance or emotion detection would be required to distinguish between neutral posts that are actually cautionary instead of descriptive-only posts, but this is noted as an area for improvement.

Sentiment Dynamics Over Time.

A temporal sentiment index analysis was performed to reflect public sentiment towards AI risk, as public sentiment changes over time, the temporal sentiment index in the observation period is shown in Figure 5.

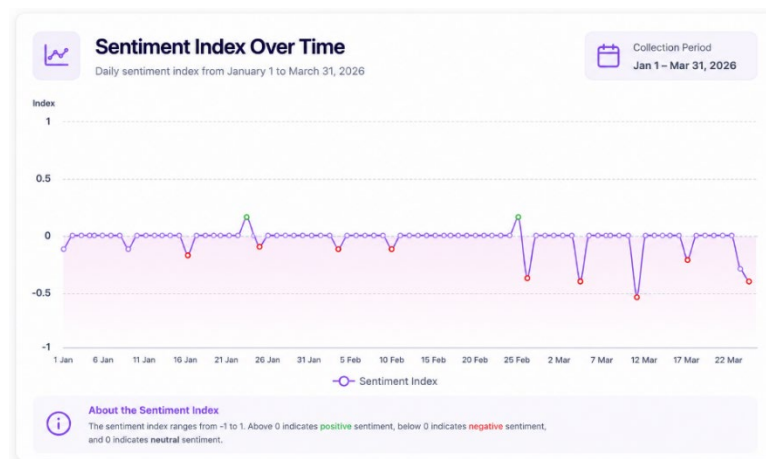


Fig. 5. Sentiment Index Over Time

From the figure 5, the pattern of the difference of the sentiment index values is fluctuating during the observation period. The majority of the index values are in the neutral to negative zone and some of them are above neutral in those periods showing positive sentiment. The trend shows that public attitudes towards the risks of AI are shifting in the first quarter of 2026.

Several brief upturns in sentiment were seen, but the overall picture was that the sentiment around public discussion of AI risks was cautious, not optimistic, and remained largely around neutral and negative scores. Spikes in positive sentiment could be connected with events like the emergence of innovations associated with AI, the introduction of new governance frameworks, or the dissemination of trust-worthy stories and narratives. The positive sentiment highs might be tied to events such as the emergence of AI-related innovations, new governance frameworks, or the dissemination of trust-worthy stories and narratives. The steady rise in neutral and negative sentiment, though, indicated that privacy, surveillance, and data-endangering issues were constant issues in public conversations during the first quarter of 2026.

Word Frequency and Temporal Patterns of Conversation.

Performing the temporal word analysis, it was possible to determine the most prominent terms in the public discussion on AI risk, as well as how the themes of the discussion changed during the observation period. This analysis went beyond sentiment polarity to give insight into the substantive issues that were in the public sphere and that may have influenced market perceptions on AI. To better understand shifts in discussion themes, Figure 6 presents the temporal trend of dominant words in AI risk discourse.



Fig. 6. Temporal Word Frequency Trend

The results showed that the terms “training”, “data” and “surveillance” were used most of the time during the research period. However other words that came up multiple times in lower frequency were “human” and “your.” These are some of the key words that stood out as central issues in the public discourse on AI risk, including data governance, training data, digital monitoring, and the human-AI relationship.

The temporal distribution of words shows a changing pattern that tended to follow the changes in conversation volume. The number of several terms showed the significant changes in the time frame when they appeared, which suggests the change in public attention over the time. For instance, the word training saw the highest rise around mid-February and then it became the most prevalent word throughout the observation period. The trend indicates increasing public awareness of concerns around the origin of AI data, IP rights, and ethical concerns around the development of large models. The continued strength of this term highlights that the discussion around AI was not solely a technical one, but has shifted to questions of transparency and responsible development.

The same applies to the word “data” which had numerous peaks during the observation period and has been and will continue to be on the agenda as it comes to the use of personal data, data governance and possible breaches of privacy. This recurrence of the term comes at a time when data collection, consent and the use of personal information in AI training and deployment are growing concerns for the public. The prominence of these data-related discussions in a market context suggests that

views on privacy and information security could be key factors influencing public confidence in AI technologies.

The term “surveillance” had a more variable temporal trend, getting more attention in late February and early March. The pattern is that AI was often defined as a means of monitoring or a tool for social control, and not just as a productivity tool. As surveillance language becomes more visible, people might be more aware of topics like “behavioural monitoring”, “automated surveillance” and “institutional oversight” related to surveillance. To complement the temporal analysis, a word cloud of dominant terms in the discourse of AI risks is presented in Figure 7.



Fig. 7. Word Cloud of Dominant Terms in AI Risk Discourse

The word cloud also supports this conclusion by highlighting the importance of training data, surveillance, privacy, human interaction, and data governance in the context of AI risk. These findings, combined, indicate that the debate on risks of AI was not disjointed, ubiquitous, or tangential, but rather focussed on a fairly consistent suite of issues around control, accountability, and responsible use. So, as far as market perception goes, the prominence of these themes suggests that when it comes to the acceptance of AI technologies, companies may find that transparency, privacy protection, and ethical governance in communication and implementation become a growing factor.

Topic Modeling in AI Risk Discourse.

The BERTopic approach to topic modeling was used to find the key thematic clusters in the discussion of AI risk. Figure 8 illustrates the BERTopic modelling results, which can help to understand the thematic structure of AI risk discourse.

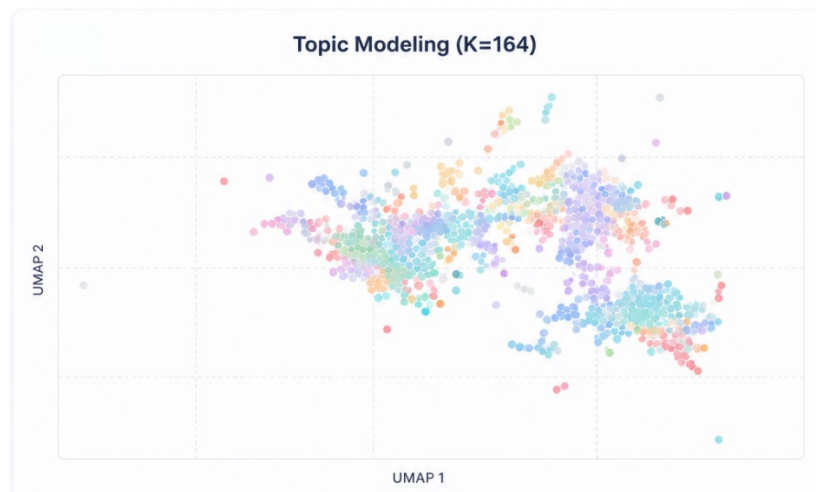


Fig. 8. Topic Modeling Results (BERTopic)

The modeling results reveal a number of topic groups emerging in a space of two dimensions (UMAP projection). The visualization shows a number of topic clusters, some of which are semantically close, others are not. There are some topic groups that seem to overlap in various aspects of data privacy, surveillance, training data, security and AI governance. The distribution of topics shows that the discussions about risks of AI are emerging in different interrelated issue groups.

The BERTopic model yielded 164 clusters of topics, suggesting a very multidimensional and fragmented structure of the discussion around AI risks. The distribution of topics, however, was not even; in fact, only a few topics had relatively high percentages of conversation, and many topics had relatively low percentages of conversation. The pattern indicates that there was a focus on a few main concerns in the public conversation about the risks of AI, while many other topics were more niche or context-specific. The noise cluster is also relatively large, suggesting that a significant portion of posts were not strongly connected to one dominant topic for the conversation, and thus the conversation is relatively heterogeneous.

The actual cluster count is high, and there is a large noise cluster demonstrating the heterogeneity of open online conversation; the interpretive grouping of clusters that are in close proximity to each other into higher-level thematic families (as opposed to clusters based on the raw number of clusters) is one of the ways to prevent over-interpretation of algorithmically generated micro-topics as substantive themes.

The different ratios of topics also indicate that there was no single narrative concerning the risks of AI, but rather a series of interrelated topics that developed concurrently. Semantic clusters also frequently included references to things such as data privacy, surveillance, training data, governance, and cybersecurity - suggesting these topics were recurring thematic anchors in public discourse. In a market perception sense, this fragmentation implies that mistrust and lack of legitimacy towards AI can originate from a mix of concerns, not only one specific issue, further highlighting the

significance of communicating multidimensionally on AI branding and governance issues.

Conceptual Relationships Between Words.

To investigate the relationships between the words that frequently co-occurred in AI risk-related conversations, text network analysis was used. In order to better show the conceptual connections between dominant terms, the text network structure of AI risk discourse is presented in Figure 9.

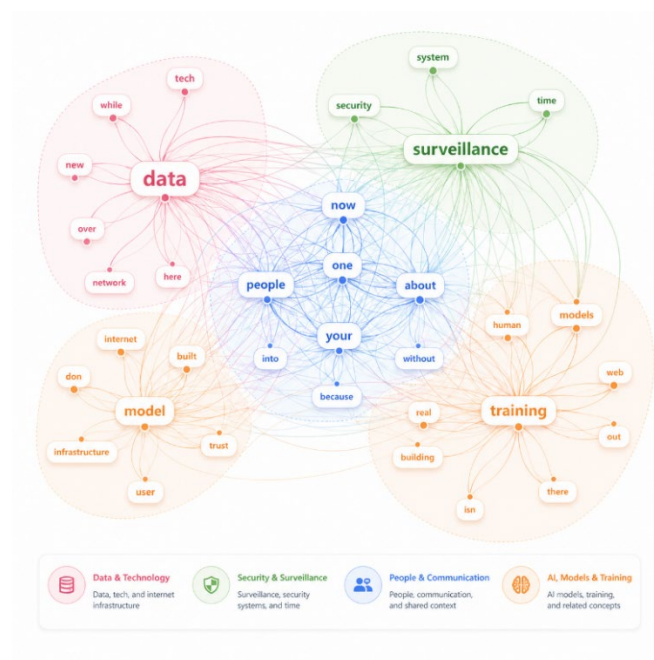


Fig. 9. Text Network Analysis

The analysis results indicate that the word "surveillance" is in the middle of the word network along with "data," "training," and "model." The discourse of AI risk is conceptually clustered through several groups of words, which suggest interrelations between the words. The network visualization shows relationships between words that create multiple discussion groups, such as groups focused on data and technology, surveillance and security, public communication and AI models and training data.

The terms used were often concentrated in the middle of the image, indicating a strong emphasis on the discourse of "data," "surveillance," "training," and "model," which strongly signaled concerns in the domain of information governance, surveillance practices, and technological development. The relationality between these concepts shows that public concerns were not isolated issues, but rather intermeshed stories of privacy, security, technical infrastructure and societal impacts of AI. These semantic clusters from a discourse perspective suggest that the public's perception of

the risk of AI was influenced by repetitive reinforcing connections among data use, system transparency, and the effects of AI deployment.

AI Risk Classification Using Zero-Shot Classification.

To analyze posts for a set of predefined categories of AI risk, a zero-shot classification analysis was used. The classification findings are shown in predefined risk categories in Figure 10 for better understanding of the distribution of perceived risks of AI.

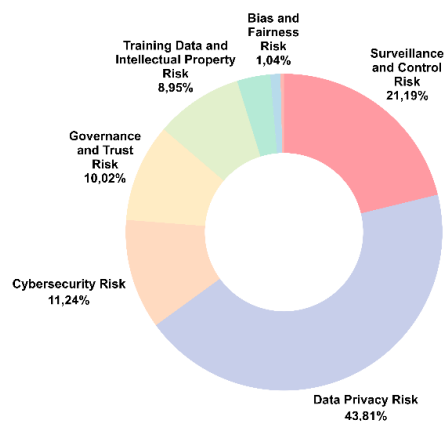


Fig. 10. Distribution of AI Risk Categories

As seen in the classification results, the Data Privacy Risk category is the most predominant and covers 43.81% of the conversation, followed by Surveillance and Control Risk category at 21.19% and Cybersecurity Risk at 11.24%. Governance and Trust Risk was marked at 10.02 percent and Training Data and Intellectual Property Risk was 8.95 percent. Bias and Fairness Risk have the lowest percentage with 1.04%. The results of this distribution, which reveal a shift in the emphasis of AI-related risk discussion in public discourse over the study period, suggest that the critical issues of AI risks in public discourse have shifted.

It is important to note that the issues discussed in public debates were mainly related to the use and privacy of personal data in the context of AI. The issue of Data Privacy Risk was the most prominent on the list of public debates. AI also was often linked to surveillance and the power of technology, as reflected in the relatively high rate of Surveillance and Control Risk. The public's concerns were not only on privacy, however, as the Cybersecurity, Governance and Trust category, combined with the Training Data and Intellectual Property Risk category, indicates. In contrast, the relatively small percentage of Bias and Fairness Risk indicates that issues of algorithmic fairness were not as commonly discussed during the observation period as issues related to privacy and digital control.

Summary of Research Findings.

The overall results show a vibrant discourse on AI dangers on the X (formerly Twitter) in the first quarter of 2026 across various domains. Time series of these conversations showed varying conversation patterns over time, different actors and different levels of connectivity, and contained a higher proportion of neutral and negative sentiment. Topic analysis, word network analysis, and risk classification revealed that the topics on AI risk included these subject areas and covered data privacy, surveillance, cyber security, training data, governance, and trust in different amounts.

3.2 Discussion

AI Risk Narratives as Market Signals.

The findings indicate that public opinion in the first quarter of 2026 over the risks of AI on the X (formerly Twitter) can be interpreted as a market signal towards the technology of AI. Changes in the volume of discussion, a predominance of a neutral-negative sentiment and the significant concern with privacy, surveillance, and cyber-security suggest that public acceptance of AI occurs in an environment of uncertainty and caution. The digital public sphere can be understood as a space for the negotiation of expectations, concerns and legitimacy of technological innovation from the technology marketing perspective. Therefore, social media posts are not just a user's spontaneous reaction, but also a sign of new technology's market perception.

The results indicate that public discourse can be used as a proxy of market perception of AI in early stages of technological legitimacy, particularly in the case of emerging technologies, which have not yet reached full social legitimacy. Trend analysis reveals that there were substantial rises in the level of discussion of AI risk in particular periods, particularly around the middle of March 2026. The increase indicates that AI's visibility is changing and evolving across different scenarios. The move towards digital consumer behaviour introduces additional risk conversation, suggesting that there is an increasing awareness of social and technological risks in the market from the use of AI.

Moreover, the tone is largely neutral and negative, suggesting that while there are positive views of AI, they are not the only ones. Instead, the public conversation is characterized by a "cautious acceptance" in which technology is seen as beneficial, but with a "care" given to the possible dangers that it may pose. It is important because of technology adoption; the degree of acceptance of the technology in the market also depends on the risks that are perceived as relevant to the technology.

Data Privacy Risk as a Barrier to Market Acceptance of AI.

One of the most noticeable findings of this study is the percentage of "Data Privacy Risk" (43.81%) that is above other categories. That's a reflection of the centrality of privacy issues in public debates on digital public sphere and AI. Increased emphasis on data privacy means that consumers are more likely to think about the potential risks they could face when they think about new technologies, as Perceived Risk Theory would predict. However, privacy concerns, particularly privacy around the use of

personal information, collection without consent and ambiguity regarding data governance, can be significant challenges for the adoption of technology.

This can also be verified as per the text network analysis which shows that concepts such as data, training, and surveillance are the central concepts in the AI discourse's conceptual network. It also implies that people don't only consider the subject of AI but also data management. At the same time, AI is regarded as a technology which is tightly linked with the utilization of the user information and the large data-driven model training processes.

In marketing speak, privacy stories win, meaning that building trust in AI is likely to be achieved through its perceived security and control over the data. So, besides technology, AI is also evaluated based on the companies' ability to keep user data secure. The discovery suggests that there is a need to increase the visibility of transparency, privacy and data accountability aspects in the context of AI-driven brand communication strategies. Features such as privacy protection, explainable AI, and transparency about data usage may prove to be turning points for gaining market legitimacy for AI technology.

Surveillance Narratives and the Formation of Trust in AI.

The Surveillance and Control Risk category accounted for 21.19% of the proportion, other than the privacy concern, there was a significant narration. This finding reveals part of the public discourse on AI as it could be used in a monitoring and control, or 'digital surveillance' way, as well as a technology to be used for productivity. The findings are consistent with the Framing Theory, which indicates that the framing of the discussion of AI in the public sphere affects the perception of AI. The public often sees AI in the context of aspects such as surveillance, monitoring, or control, which bring up more social and psychological concerns.

The implications of this are significant as it is how the public thinks about AI that will influence market trust. Increased concerns regarding institutional control, privacy issues, or misuse of technology can affect the attitude toward technology's acceptance if it turns into a surveillance tool. The results highlight the importance of focusing on the positioning of AI as an “empowering technology” and not a “monitoring technology” in the context of technology marketing. AI's role as an augmentation tool for humans, a productivity aid, and a decision-making support tool will be more resonant than as a tool of dominance and of surveillance.

Cybersecurity, Governance, and Market Legitimacy toward AI.

The study also finds that Cybersecurity Risk (11.24% of total) and Governance and Trust Risk (10.02% of total) are important elements of the public discussion on AI. Although these two types of concerns may not be the first that come to mind when discussing privacy, they indicate that the sense of security and institutional management is also influencing how the legitimacy of AI is perceived.

The new technologies need to be socially legitimized, meaning that 'publics' need to view the technology as safe, trustworthy and legitimate. As cybersecurity and gov-

ernance regularly crop up in the public discourse, it's also a sign that the public is asking questions about the regulation and governance of AI, as well as its capabilities.

BERTopic and text network analysis results revealed that there is a relation between some of the issues of data, security, training and surveillance. The interconnection suggests that the perception of the AI risk is shifting from a set of individual problems to a whole system of interconnected problems.

The findings in the field of marketing indicate that gaining market legitimacy for AI may require a combination of security, regulation, and responsible communication about AI. Organizations involved in developing AI systems may have to make an effort to reaffirm the idea of a secure system, explain the transparency of the algorithms, and ensure that there are human oversight mechanisms in place to boost public trust.

In addition to these general guidelines, the outcomes are translated into concrete recommendations to organizations developing and promoting AI, tailored to the risk. These recommendations are associated to the various risk categories that were identified in the discourse:

Data privacy (dominant concern): Assure that data is only used for specific purposes; publish data-use notices, third-party, "nutrition labels" data-privacy measures; offer granular consent and opt-out from training data use; process data on the device or anonymize data-processing; where feasible, disclose data-retention and deletion policies; all of this can be supported by independent privacy certifications and third-party data privacy audits.

Avoidance of dark patterns as a way to frame the situation, explicit commitments against secondary use of behavioural data, surveillance and control; counter surveillance framing through positioning and design, user facing controls, transparency of monitoring features can decrease control related concern.

Security Maturity: demonstrate signal security using industry standards (e.g., ISO/IEC 27001, SOC 2), vulnerability-disclosure and bug-bounty programs, clear incident response and public disclosure of model & data security protection measures.

Governance and trust: Ensure institutional legitimacy with written principles of responsible AI, model documentation and system documentation (model cards, system cards), mechanisms for human oversight and escalation, and incorporate new regulations and frameworks (e.g., EU AI Act, NIST AI Risk Management Framework etc.); independent oversight or ethics review increases credibility.

Training data and IP: communicate data-provenance practices, disclose training data sourcing and licensing, and offer opt-out from, and attribution and compensation for, creators.

Timing of stakeholder engagement/communication: engagement of identified users, journalists and policy actors in discourse network and focus on 'spikes' of risk narratives – focus on building trust-related communication at times of peaks of public attention.

Branding: construct brand legitimacy through "trustworthy-by-design" communication that emphasizes transparency, accountability and safety rather than capability

claims; this evidence suggests that acceptance is achieved through these rather than through capability claims.

Toward a Cautious AI Market Acceptance.

Overall, the findings suggest that the market is not fully rejecting AI yet, it still hasn't shown unconditional acceptance. Public debate around AI is shifting towards a range of 'cautious acceptance' due to the prevalence of neutral and even negative sentiment, along with the increased visibility of privacy and surveillance concerns and the rise of security and governance issues.

The situation implies that the adoption of AI in the market is likely to be slow, the level of which will be affected by the ability of organisations to minimize the increasing risk perception of the public at large. That is, the effectiveness of the use of AI doesn't seem to depend solely on the quality of its technology, but also on how companies can create trust, legitimacy, and communication that is in keeping with public expectations.

3.3 Limitations and Future Research

There were several limitations of this study that limited interpretation of the results. First, with respect to the discourse–market linkage, the analysis reflects the perception of an engaged digital public within one single platform, and does not directly assess the attitudes and actions of consumers, investors, and organizational decision makers; the use of a discourse proxy for the market, therefore, is meant to be indicative but not confirmatory, and the correspondence of the discourse with downstream adoption has not yet been validated. Second, in terms of representativeness, the corpus is platform-specific (X/Twitter); English-only; retrieved over a 3-month window according to risk-oriented keywords; it is by construction geographically, linguistically and demographically biased and the results should not be interpreted as prevalence estimates at the population level. Third, the study is exploratory and correlational in nature, meaning that it uncovers relationships between the dominant risk frames and a cautious sentiment orientation, but does not test causal processes at the level of the individual from risk framing to risk perception, and this would require experimental or longitudinal designs. Fourth, sentiment classification is coarse-grained; a large neutral share does not necessarily mean "cautious" and finer-grained stance or emotion detection is required to distinguish between "descriptive" and "cautious".

Fifth, and most important, the sampling frame is deliberately risk-focused, causing a selection bias: the design deliberately aims to sample posts that lend themselves to a narrative of a negative role for AI, while at the same time under-representing positive AI stories, such as innovation, productivity, efficiency, and economic opportunity, that are also part of public conversation. This low share of positive sentiment (6.21%) should therefore be interpreted as a property of a risk-oriented sample, rather than indicative of a lack of or reduced presence of positive narratives in the broader AI conversation: the result characterizes the structure of AI risk discourse specifically, and would overstate the negativity if seen as representative of the structure of the

market conversation. Future work should involve balanced sampling including benefit/opportunity framed discourse, multi-platform and multi-lingual data, and more fine-grained stance detection, as well as testing the predictive ability of such signals, by triangulating them with consumer-behavior surveys and/or market indicators.

Lastly, all data were collected in the modular pipeline of SocialX and all computations were made there. However, this gives an integrated workflow, but the internal parameters of some modules (e.g., embedding and classification models, UMAP/HDBSCAN parameters, network-centrality computations) are not fully exposed to the user, and several diagnostic metrics (e.g., topic-coherence scores, independent inter-rater validation of the zero-shot labels, formal centrality and modularity values) could not be reported. This is a limitation on transparency and reproducibility: findings based on platform-internal computation should be interpreted accordingly, and future work should complement platform outputs with open, fully documented computation pipelines, and independent validation of platform's automated classifications, such as manual classifications of a labeled subsample.

4 Conclusion

The purpose of this study is to gain insight into the narratives of AI risks circulating in the public discourse and how they affect market perceptions of AI technology by conducting a Social Network Analysis (SNA) and analyzing public conversations using natural language processing (NLP) on X (formerly Twitter) during the first quarter of 2026. The findings indicate that AI risk discourse unfolded in a dynamic manner, with multiple actors and the emergence of discussion patterns focused on issues around data privacy and surveillance, cybersecurity, and technology governance.

In relation to RQ1, the results suggest that the discourse of AI risks was overwhelmingly centered around Data Privacy Risk, followed by Surveillance and Control Risk and Cybersecurity Risk. For RQ2, in both SNA and text network analysis, there was a cohesive structure of discourse that connected the influential actors and themes like data, surveillance, training, and models. For RQ3, the sentiment analysis and thematic analysis indicate that the public's sentiment toward AI was not positive but rather negative, with a noticeable prevalence of negative sentiments. These results suggest that the perception of AI's benefits and risks in the markets is not solely a matter of what benefits are expected and how, but also of how the risks are discussed in the public sphere. As for RQ4, the results also indicate that AI risk narratives serve as signals that influence trust, legitimacy, and public acceptance of AI technologies.

This study makes a theoretical contribution to the literature on technology marketing by providing evidence that the public discourse can be considered a market signal in the development of perceptions of emerging technologies. The combined insights of Framing Theory, Perceived Risk Theory, and market legitimacy indicates that the negotiation of adopting AI is driven by concerns about data privacy, digital surveillance, technology security, and institutional governance. In a practical sense, the study has implications for organisations and AI developers to improve their communication

strategies to include transparency, data protection, system security, and responsible AI communication to ensure public trust in AI technologies. Further studies are recommended in using a digital discourse analysis to consumer behavior survey to gain a better understanding of the relationship between risk perception and intention to adopt AI.

5 Acknowledgements

This article used DeepL AI Translator as a translation tool from Indonesian to English, and ChatGPT 5.5 to improve the readability of the content, making the language easier to understand.

Reference

1. Alotaibi, R.T.T.: Risk network analysis and simulation research in municipal engineering projects contracted by China in Saudi Arabia. *Sci. Rep.* 16, 870 (2026). <https://doi.org/10.1038/s41598-025-32219-z>.
2. Ayub, A.: Exploring Human-AI Communication in the Digital Era: A Need to Develop New Theories for Human-AI Communication. *International Social Sciences and Education Journal.* 2, 30–36 (2025). <https://doi.org/10.61424/issej.v2i2.181>.
3. Budic, M.: AI and Us: Ethical Concerns, Public Knowledge and Public Attitudes on Artificial Intelligence. In: *Proceedings of the 2022 AAAI/ACM Conference on AI, Ethics, and Society.* pp. 892–892. ACM, New York, NY, USA (2022). <https://doi.org/10.1145/3514094.3539518>.
4. Cho, H., Hooi, R.: Risk perceptions and trust mechanisms related to everyday AI. In: *Research Handbook on Artificial Intelligence and Communication.* pp. 163–175. Edward Elgar Publishing (2023). <https://doi.org/10.4337/9781803920306.00019>.
5. Dupuis, M., Long, S.: X Marks the Spot: An Examination of X and Its Transformation from Twitter. *European Conference on Social Media.* 12, 38–46 (2025). <https://doi.org/10.34190/ecsm.12.1.3568>.
6. Dwork, C., Minow, M.: Distrust of Artificial Intelligence: Sources & Responses from Computer Science & Law. *Daedalus.* 151, 309–321 (2022). https://doi.org/10.1162/daed_a_01918.
7. Fildor, D., Škrinjaric, B., Budak, J.: Anatomy of AI Risk Perception: Differentiating Privacy, Misinformation, and Technical Failures. *Int. J. Hum. Comput. Interact.* 1–21 (2026). <https://doi.org/10.1080/10447318.2025.2610439>.
8. Friemel, T.N.: Social Network Analysis. In: *The International Encyclopedia of Communication Research Methods.* pp. 1–14. Wiley (2017). <https://doi.org/10.1002/9781118901731.iecrm0235>.
9. Gonçalves, A.R., Pinto, D.C., Rita, P., Pires, T.: Artificial Intelligence and Its Ethical Implications for Marketing. *Emerging Science Journal.* 7, 313–327 (2023). <https://doi.org/10.28991/ESJ-2023-07-02-01>.
10. HORVATH, K.: CONSUMER ATTITUDES TOWARD ARTIFICIAL INTELLIGENCE: A COMPARATIVE ANALYSIS OF MEASUREMENT SCALES. *The Annals of the University of Oradea Economic Sciences.* (2024). [https://doi.org/10.47535/1991AUOES33\(2\)030](https://doi.org/10.47535/1991AUOES33(2)030).

11. Seyed Mohammad Seyed Hosseini: Modeling Systemic Communication Risks in Global Financial Media: A Review of Networked Information Flows Across Stock, Forex, and Cryptocurrency Markets. *Review of Communication Research*. 11, 215–224 (2023). <https://doi.org/10.52152/RCR.V11.109>.
12. Tasriqul Islam, Sadia Afrin, Neda Zand: AI in Public Governance: Ensuring Rights and Innovation in Non-High-Risk AI Systems in the United States. *European Journal of Technology*. 8, 17–27 (2024). <https://doi.org/10.47672/ejt.2577>.
13. K, R., Izhar, R., Bhatti, S.N.: Examining the Role of AI Experience in Shaping Consumer Attitudes and Adoption Behavior in Social Media Marketing. In: 2025 International Conference on Innovation in Artificial Intelligence and Internet of Things (AIIT). pp. 1–7. IEEE (2025). <https://doi.org/10.1109/AIIT63112.2025.11082958>.
14. Khan, L.: Social Network Analysis. In: *The Data Analytics Advantage*. pp. 130–156. Oxford University Press New York (2025). <https://doi.org/10.1093/oso/9780197814222.003.0007>.
15. Krieger, J.B., Boudier, F., Wibral, M., Almeida, R.J.: A systematic literature review on risk perception of Artificial Narrow Intelligence. *J. Risk Res.* 28, 1111–1129 (2025). <https://doi.org/10.1080/13669877.2024.2350725>.
16. Li, Z., Lyu, M., Li, L., Huang, J.: Decoding Global AI Risk Perception Evolution: Intercultural drivers and Public Discourse Patterns in Algorithmic Societies, (2025). <https://doi.org/10.21203/rs.3.rs-6468789/v1>.
17. Carrió-Pastor, M.L., Perinán-Pascual, C. eds: *Research Trends in Discourse Analysis and Language Acquisition in the Artificial Intelligence Era*. Peter Lang Verlag (2026). <https://doi.org/10.3726/b23335>.
18. Mammadov, A.: From natural (human to human) discourse to artificial intelligence-assisted (human to artificial intelligence) discourse. *International Review of Pragmatics*. 18, 116–130 (2026). <https://doi.org/10.1163/18773109-01801006>.
19. Moon, W.-K., Wei, X., Overton, H., Kim, J.K.: Between Innovation and Caution: How Consumers' Risk Perception Shapes AI Product Decisions. *Journal of Current Issues & Research in Advertising*. 47, 157–180 (2026). <https://doi.org/10.1080/10641734.2025.2468474>.
20. Murphy, J., Link, M.W., Childs, J.H., Tesfaye, C.L., Dean, E., Stern, M., Pasek, J., Cohen, J., Callegaro, M., Harwood, P.: *Social Media in Public Opinion Research: Executive Summary of the Aapor Task Force on Emerging Technologies in Public Opinion Research*. *Public Opin. Q.* 78, 788–794 (2014). <https://doi.org/10.1093/poq/nfu053>.
21. N. O. Sadiku, M., Padi, D., O. Sadiku, J.: Artificial Intelligence in Marketing. *International Journal of Engineering Research and Advanced Technology*. 11, 01–13 (2025). <https://doi.org/10.31695/IJERAT.2025.8.1>.
22. Neri, H., Cozman, F.: The role of experts in the public perception of risk of artificial intelligence. *AI Soc.* 35, 663–673 (2020). <https://doi.org/10.1007/s00146-019-00924-9>.
23. Neyazi, T.A., Ng, S.W.T., Hobbs, M., Yue, A.: Understanding user interactions and perceptions of AI risk in Singapore. *Big Data Soc.* 10, (2023). <https://doi.org/10.1177/20539517231213823>.
24. Surabhi NV, Arun Ajith K: Artificial Intelligence's Detrimental Effects on Digital Marketing via Social Media Platforms and Case Studies. *International Journal of Advanced Research in Science, Communication and Technology*. 524–532 (2025). <https://doi.org/10.48175/IJARSCT-23157>.
25. Okoidigun, O.G., Emagbetere, E.: Trusting Artificial Intelligence: A Qualitative Exploration of Public Perception and Acceptance of the Risks and Benefits. *Open J. Soc. Sci.* 13, 64–80 (2025). <https://doi.org/10.4236/jss.2025.133006>.

26. Priyadharshini, Ms.K., Arulraj, Dr.S.: ARTIFICIAL INTELLIGENCE (AI) IN MARKETING. *INTERANTIONAL JOURNAL OF SCIENTIFIC RESEARCH IN ENGINEERING AND MANAGEMENT*. 08, 1–3 (2024). <https://doi.org/10.55041/IJSREM29127>.
27. Said, N., Potinteu, A.E., Brich, I.R., Buder, J., Schumm, H., Huff, M.: An Artificial Intelligence Perspective: How Knowledge and Confidence Shape Risk and Opportunity Perception, (2022). <https://doi.org/10.31234/osf.io/5zvha>.
28. Samuel, J., Khanna, T., Esguerra, J., Sundar, S., Pelaez, A., Bhuyan, S.S.: The Rise of Artificial Intelligence Phobia! Unveiling News-Driven Spread of AI Fear Sentiment Using ML, NLP, and LLMs. *IEEE Access*. 13, 125944–125969 (2025). <https://doi.org/10.1109/ACCESS.2025.3588179>.
29. Schwarz, A.: The mediated amplification of societal risk and risk governance of artificial intelligence: technological risk frames on YouTube and their impact before and after ChatGPT. *J. Risk Res.* 1–23 (2024). <https://doi.org/10.1080/13669877.2024.2437629>.
30. Shi, J., Kapucu, N., Zhu, Z., Guo, X., Haupt, B.: Assessing Risk Communication in Social Media for Crisis Prevention: A Social Network Analysis of Microblog. *J. Homel. Secur. Emerg. Manag.* 14, (2017). <https://doi.org/10.1515/jhsem-2016-0058>.
31. Skarżyńska, E., Paliszkievicz, J., Dąbrowski, I., Mendel, M.: Risks, failures, and ethical dilemmas of AI technologies and trust. In: *Trust in Generative Artificial Intelligence*. pp. 3–11. Routledge, New York (2025). <https://doi.org/10.4324/9781003586937-2>.
32. Goffman, Erving., Berger, B.M.: *Frame analysis : an essay on the organization of experience*. Northeastern University Press (1986).
33. Peter, J.P., Ryan, M.J.: An Investigation of Perceived Risk at the Brand Level. *Journal of Marketing Research*. 13, 184–188 (1976). <https://doi.org/10.1177/002224377601300210>.
34. Suchman, M.C.: *Managing Legitimacy: Strategic and Institutional Approaches*. *Academy of Management Review*. 20, 571–610 (1995). <https://doi.org/10.5465/amr.1995.9508080331>.
35. Vulpe, S.-N., Rughiniș, R., Țurcanu, D., Rosner, D.: AI and cybersecurity: a risk society perspective. *Front. Comput. Sci.* 6, (2024). <https://doi.org/10.3389/fcomp.2024.1462250>.
36. Wajid, M.A., Camacho-Zuñiga, C.: Mapping WUN expert discourse on responsible and ethical AI: a multinational expert network analysis. *Front. Commun. (Lausanne)*. 10, (2025). <https://doi.org/10.3389/fcomm.2025.1689751>.
37. Ying, S., Shen, Z., Xu, Z.: What Factors Affect Consumers' Trust in AI? *Advances in Economics, Management and Political Sciences*. 260, 127–132 (2026). <https://doi.org/10.54254/2754-1169/2026.BJ31948>.
38. Zhang, B.: Public Opinion toward Artificial Intelligence. In: *The Oxford Handbook of AI Governance*. pp. 553–571. Oxford University Press (2023). <https://doi.org/10.1093/oxfordhb/9780197579329.013.36>.