



Marketing Science Institute Working Paper Series 2026

Report No. 26-126

Creating Privacy Value Through Artificial Intelligence Prompt Literacy: The Role of Self-Investment

Johanna Zimmermann, Yakov Bart and Koen Pauwels

“Creating Privacy Value Through Artificial Intelligence Prompt Literacy: The Role of Self-Investment” © 2026

Johanna Zimmermann, Yakov Bart and Koen Pauwels

MSI Working Papers are Distributed for the benefit of MSI corporate and academic members and the general public. Reports are not to be reproduced or published in any form or by any means, electronic or mechanical, without written permission.

Creating Privacy Value Through Artificial Intelligence Prompt Literacy: The Role of Self-Investment

Johanna Zimmermann*, Assistant Professor of Marketing, University of Cologne, Germany,
j.zimmermann@wiso.uni-koeln.de, Albertus-Magnus-Platz, 50923 Cologne

Yakov Bart, Patrick F. and Helen C. Walsh Research Professor of Marketing, Northeastern
University, USA, y.bart@northeastern.edu

Koen Pauwels, Distinguished Professor of Marketing, Northeastern University, USA,
k.pauwels@northeastern.edu

*Corresponding author

Declaration of Conflicting Interest

The author(s) declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

Funding Statement

This project was supported by internal funding from Northeastern University.

Acknowledgments

The authors thank Jan H. Schumann for his valuable support and helpful comments on this research.

Creating Privacy Value Through Artificial Intelligence Prompt Literacy: The Role of Self-Investment

Abstract

Generative artificial intelligence (GenAI) is reshaping how consumers and employees disclose data to digital systems. Unlike the discrete disclosure decisions examined in previous literature on privacy settings, GenAI use unfolds through interactive prompting in which users share information to cocreate outputs. This research examines how self-investment in such disclosure processes shapes perceived privacy value, defined as the perceived net value of benefits relative to privacy risks. Across four studies, we show that greater self-investment increases perceived ownership of the cocreated output; moreover, the effect of self-investment on ownership is strongest when interactions involve sensitive data disclosure, with ownership subsequently shaping privacy value through perceived benefits and privacy risks. Finally, AI prompt literacy is an actionable antecedent of self-investment: a field intervention and an online experiment show that even modest improvements in practical prompting capability can increase engagement with GenAI. The findings contribute to privacy, GenAI, and cocreation research by showing that privacy evaluations in GenAI contexts are dynamic, self-invested, and shaped by users' ability to guide the interaction.

Keywords: AI prompt literacy, data disclosure, self-investment, cocreation, generative AI, privacy

Interacting with generative artificial intelligence (GenAI) tools has become part of everyday life. Consumers use tools like ChatGPT or Claude to simplify tasks and support personal decision making, professionals integrate them to streamline workflows and enhance strategy, and students call on GenAI for learning and research (e.g., Celiktutan, Klesse, and Tuk 2024; Peres et al. 2023). Unlike earlier digital technologies, GenAI interactions are centered on open-ended, iterative exchanges in which individuals externalize not only information but also their goals, reasoning processes, uncertainties, and contextual details. This integration has led to the disclosure of substantial amounts of diverse data, providing insights into individuals, their preferences, and their broader work environments and private lives.

A defining feature of these interactions is that data disclosure is embedded in individual-initiated value creation. Unlike earlier digital technologies that specify in advance what information must be provided (e.g., data forms; Krafft et al. 2021), GenAI tools require individuals to decide what personal and contextual information, as well as other personal resources (e.g., effort, time), they want to provide through prompts and refinements. This self-investment allows individuals to actively shape how GenAI systems generate outputs, thereby enabling value creation. At the same time, though, the absence of predefined disclosure boundaries means that self-investment often entails the disclosure of richer and potentially sensitive data. After data are disclosed through prompting, individuals have limited behavioral control over downstream data use and must instead rely on platform practices and regulatory safeguards (Zimmermann et al. 2024). GenAI interactions therefore engender a fundamental tension: The self-investment required to create value simultaneously increases individuals' exposure to privacy risks. Yet, how individuals navigate this tension, particularly across contexts that differ in the sensitivity of the disclosed data, remains unclear.

Research at the intersection of privacy and AI offers limited guidance for understanding these evaluations. Privacy research on AI has primarily focused on individual vulnerabilities and the risks associated with data collection, examining how technologies such as sensors, video, and activity tracking enable firms to accumulate a lot of personal data (e.g., Melzner, Bonezzi, and Meyvis 2023; Uysal, Alavi, and Bezençon 2022). In parallel, AI research has largely emphasized the diminishing autonomy of individuals as technologies assume greater control over tasks, or it has focused on user reactions to AI-generated outputs (e.g., De Bellis and Venkataramani Johar 2020; Puntoni et al. 2021). Taken together, these streams suggest that greater data disclosure in AI interactions primarily heightens privacy risks and leads to negative downstream consequences.

However, this risk-centered view overlooks defining characteristics of GenAI interactions that align with cocreation and customization research. Much of the privacy and AI literature implicitly treats individuals as passive recipients of benefits while bearing the associated risks. Rather than receiving predefined personalized benefits, individuals in GenAI data disclosure processes actively cocreate outcomes with the system through iterative prompting and refinement to achieve customized results (Arora, Chakraborty, and Nishimura 2025; Celiktutan, Klesse, and Tuk 2024; Troye and Supphellen 2012). Such customization does not merely increase perceived benefits; by visibly reflecting individuals' own inputs and intentions, outcomes may be experienced as more personally meaningful and partly their "own," rather than as outputs generated entirely by an external system (Belk 1988; Franke, Schreier, and Kaiser 2010; Norton, Mochon, and Ariely 2012). This personalization through customization may therefore intensify the underlying privacy tension by increasing what is at stake in the interaction.

Importantly, individuals differ in how effectively they can self-invest in GenAI interactions. While GenAI tools afford opportunities for active value creation, realizing these

opportunities depends on individuals' ability to articulate goals, provide relevant context, and iteratively refine prompts in ways that meaningfully guide the system (Tully, Longoni, and Appel 2025). These differences establish *AI prompt literacy* as a key factor shaping how individuals engage with GenAI tools and how much they self-invest in the interaction. AI prompt literacy captures individuals' capability to translate goals and contextual knowledge into prompts that effectively steer GenAI systems, thereby influencing both the value created through disclosure and how individuals experience the outcomes of that disclosure.

In sum, the literature does not yet explain how individuals evaluate data disclosure in GenAI interactions that require active self-investment and iterative engagement or indicate whether these evaluations depend on individuals' ability to productively guide the interaction. This gap is surprising given its managerial, consumer, and regulatory relevance. Organizations are increasingly encouraging deep engagement with GenAI tools and policymakers are mandating investments in AI literacy (e.g., through the European Union [EU] AI Act), but are unsure whether fostering higher self-investment in data disclosure, enabled by greater prompt literacy, is perceived as beneficial. In particular, insight is lacking into how individuals balance the benefits generated through self-investment against the privacy risks that arise when disclosing richer and more diagnostic data. To address this gap, we pursue three research goals: (1) to conceptualize data disclosure as an iterative process in GenAI contexts; (2) to investigate the effect of individuals' AI prompt literacy on perceived self-investment and, ultimately, privacy value perceptions; and (3) to delineate the underlying mechanism through perceived outcome ownership.

To this end, we conducted four complementary studies that combine longitudinal field evidence, online experimentation, laboratory interaction with ChatGPT, and scenario-based tests of boundary conditions. Study 1, a longitudinal field intervention conducted with an AI literacy

consultancy, examines whether AI literacy training increases self-investment in GenAI use. We find that gains in practical AI use literacy, but not AI evaluation literacy, are associated with higher self-investment and favorable indirect effects on privacy value. Study 2 isolates AI prompt literacy, a core component of AI use literacy, in an online experiment by manipulating participants' exposure to prompt-design training and asking them to formulate an initial prompt for a GenAI-assisted cover-letter task. Because participants did not receive a GenAI output in this study, Study 2 provides a conservative test of anticipated ownership and privacy value at the prompt-formulation stage. Studies 3 and 4 then focus on the mechanism. Study 3 directly manipulates self-investment during a live ChatGPT interaction, and Study 4 tests whether the self-investment–ownership pathway is stronger when the interaction involves sensitive data disclosure.

This research makes three contributions. First, we extend privacy and data-disclosure research by conceptualizing GenAI disclosure as an iterative cocreation process rather than a one-time transfer of information (e.g., Dinev et al. 2013; Krafft et al. 2021; Schumann, von Wangenheim, and Groene 2014; Xu et al. 2011). This shift shows that privacy evaluations depend not only on how much data is disclosed but also on whether individuals experience their disclosure as self-directed input into value creation. Second, we contribute to GenAI marketing research by identifying AI prompt literacy as a practical capability that enables users to translate goals, context, and constraints into more effective interaction with GenAI systems (e.g., Celiktutan, Klesse, and Tuk 2024; Grewal et al. 2025; Puntoni et al. 2021). Third, we extend cocreation and customization research by showing that self-invested disclosure increases perceived ownership of GenAI outputs, which amplifies perceived benefits and, unexpectedly, attenuates perceived privacy risks (e.g., Franke, Schreier, and Kaiser 2010; Klesse et al. 2019; Norton, Mochon, and Ariely 2012; Troye and Supphellen 2012).

From a practical perspective, the findings suggest that firms should design customer- and employee-facing GenAI tools to support informed, privacy-preserving self-investment. Prompt guidance, iterative feedback, and clear explanations of how user input shapes outputs can help individuals create value without obscuring disclosure boundaries. The findings also underscore the importance of AI literacy initiatives, not merely as regulatory compliance tools but as mechanisms that help users engage critically and productively with GenAI systems (Huang and Rust 2025; Kosmyna et al. 2025). At the same time, firms and regulators should recognize that perceived ownership can reduce perceived privacy risks; interfaces should therefore make data-use boundaries salient precisely when users feel most invested in the cocreated output.

Conceptualization

GenAI Data Disclosure Processes

We integrate perspectives from prior privacy and AI research to conceptualize human–GenAI interactions as an interactive and iterative data disclosure process (see Figure 1). Privacy research documents two sequential stages in consumer–firm data exchanges: data flow and data handling (Goodwin 1991; Hoffman and Novak 1999; Xu et al. 2011). Data flow refers to the process of disclosing or gathering data, encompassing the questions of whether to share data, what to share, and how much to share. Individuals’ level of control over data disclosure leads to different types of data flow experiences, including active, passive, and hybrid (Krafft et al. 2021; Zimmermann et al. 2024). Data handling addresses what happens to individuals’ data after disclosure. This includes determining who has access to the data and how the data are used to create benefits for both the disclosing and receiving entities (e.g., Mothersbaugh et al. 2012; Tucker 2014).

These two stages of consumer–firm data exchanges align with AI’s three fundamental technical components: data gathering, data processing, and AI output (Agrawal, Gans, and

Goldfarb 2018; Puntoni et al. 2021). Data gathering in AI corresponds to what privacy research describes as data flow; however, in AI, this step is typically carried out by technologies (vs. humans) that automatically collect data through sensors, video recordings, or physical activities. Data processing and AI output, in turn, are part of the data handling process. Data processing refers to the use of statistical and computational techniques to leverage the collected data and generate outputs, such as responses to inquiries. For example, AI may be integrated into apps that track users' physical activities (i.e., data flow or gathering), analyze their health data (i.e., data processing), and offer personalized advice for fitness or diet plans (i.e., AI output). What both traditional privacy and AI perspectives have in common is that (1) data disclosure is largely a linear process and (2) the data recipient (e.g., AI technology) independently produces the output, leaving the data-disclosing entity with no control over value creation after data disclosure.

In line with emerging evidence from GenAI literature (e.g., Celiktutan, Klesse, and Tuk 2024; Huang and Rust 2024), we argue that an individual's interaction with GenAI tools is an interactive, iterative data disclosure process, in which the outcome is cocreated by the data disclosing entity (i.e., the individual).¹ In a GenAI data disclosure process, individuals not only determine the data input but also are actively involved in shaping the output of the AI system. The key element of this process is the *prompt*, that is, the human input, command, or instruction (e.g., Blanchard et al. 2025). From a privacy perspective, a prompt is not merely a textual or voice-based query but a deliberate act of data disclosure with two essential components: (1)

¹ While our focus is on GenAI, the underlying dynamics of this type of data disclosure are also relevant in other digital interactions, such as search engine queries or chatbot conversations. However, the dynamics are particularly salient in GenAI contexts given the interactive, iterative, and co-creative nature of the exchange, which is why we chose this context. We do not elaborate on the technical components or processes of GenAI systems, as such elaboration is beyond the scope of our research and does not reflect the perspective of individuals engaging with such tools. Our focus remains on consumers' experience and their interpretations of GenAI interactions.

specifying whether, what, and how much personal or contextual data are shared with the system and (2) instructing how these data should be processed to produce the desired output.

For example, consider someone seeking assistance with writing a professional cover letter for a doctoral application. The initial, simple prompt might be: “I need to write a professional cover letter for an application as a doctoral student in the field of marketing and AI. The cover letter needs to include my personal information (e.g., name, educational background, skills) and should emphasize my interest in marketing research. Provide a Word document structured like a letter, including the institution’s address. Can you help me write a compelling letter?” In this case, the individual discloses data not only directly and somewhat consciously (e.g., personal information) but also indirectly (e.g., expressing interest in academia; Huang and Rust 2024). In addition, the individual may guide the data processing by specifying how the tool should present the data, such as by requesting a formal tone, highlighting specific experiences, or structuring the letter to emphasize academic achievements. This guidance increases the odds that the final output aligns with the individual's expectations for the cover letter. The GenAI tool then processes the data and generates an initial output. In some cases, the initial output may already meet the individual’s needs without further refinement, meaning that simple prompts may lead to answers that fully address the question in a single interaction (e.g., when the prompt entails a question about receiving specific information such as “when was Harvard founded”). However, in many instances, individuals may decide to engage in further interactions to refine the output, providing more context or specifying their intent to optimize the results and benefit from a more tailored response. These ongoing exchanges allow them to fine-tune the response, clarify their preferences, and align the output more closely with their goals. Crucially, this iterative structure allows individuals to continuously adjust how they disclose data through their prompts, thereby enacting different levels of self-investment across GenAI interactions.

AI Prompt Literacy

To optimize outcomes from GenAI data disclosure processes, merely providing prompts and relying on refinement through trial and error are insufficient. Instead, individuals must develop the competency to strategically guide GenAI tools across iterative interactions, effectively orchestrating the human–AI collaboration process (White et al. 2023). This necessity gives rise to what we call “AI prompt literacy,” a marketing-relevant extension of existing AI literacy frameworks that addresses the unique demands of GenAI interactions. Existing conceptualizations of AI literacy define it broadly as individuals’ objective knowledge about AI, including conceptual and technical understanding, best practices in AI development and use, and awareness of regulatory requirements (Tully, Longoni, and Appel 2025), or as the ability to apply and critically evaluate AI technologies and their outcomes (Wang, Rau, and Yuan 2023). While these perspectives capture important competences, they do not sufficiently account for the defining mechanism of GenAI interactions: the strategic construction and deployment of prompts as deliberate acts of data disclosure and processing specifications (Hwang, Lee, and Shin 2023).

Building on our theoretical framework, which positions prompts as acts of data disclosure (i.e., whether, what, and how much personal or contextual data are shared and how these data should be processed to produce desired outputs), we define AI prompt literacy as the competency to orchestrate the human–GenAI input–output loop. This orchestration is achieved through the strategic crafting, interpretation, and iterative refinement of prompts through which individuals enact deliberate data disclosure and progressively invest effort, information, and intentionality in the interaction. At its core, AI prompt literacy reflects the insight that prompts function as “programming with words” (Mollick 2023), such that the quality of GenAI outputs depends critically on how individuals structure their inputs. This competency is particularly important considering that contemporary GenAI systems combine natural language processing and machine

learning to interpret conversational prompts while continuously adapting to user input. Moreover, commercial GenAI platforms increasingly rely on intent recognition mechanisms that infer user goals and preferences from contextual cues embedded in prompts (Blanchard et al. 2025; MIT Management 2025; Wu, Terry, and Cai 2022). As a result, effective prompting requires individuals to anticipate how the system will interpret and act on the disclosed information.

Hypotheses Development

AI Prompt Literacy, Self-Investment, and the Dual Effect on Privacy Value

Privacy value reflects individuals' subjective assessment of the trade-off between privacy costs (i.e., perceived risks associated with data disclosure) and the benefits received in return for their data in a given context (Dinev et al. 2013; Xu et al. 2011). This conceptualization builds on privacy calculus theory, which posits that individuals evaluate privacy-related decisions by weighing anticipated benefits against perceived risks (Culnan and Armstrong 1999; Dinev and Hart 2006). In GenAI interactions, this calculus becomes particularly salient because the quality and usefulness of AI-generated outputs are directly contingent on the specificity and contextual richness of the information individuals provide. More detailed prompts enable more tailored and valuable outputs, but they also increase privacy exposure by expanding the scale and scope of the disclosed information, thereby enabling inference about individuals' preferences, intentions, and even personal characteristics. This tension differs fundamentally from traditional personalization contexts examined in prior research, in which firms rely on previously collected data to tailor offerings to consumers and subsequently control both data processing and value creation (Chellappa and Sin 2005). Prior research in these contexts shows that personalization can increase perceived benefits and, under certain conditions, attenuate privacy concerns (Hui, Teo, and Lee 2007; Xu et al. 2011). However, these effects emerge from largely linear and one-

directional data environments, in which individuals make an initial disclosure decision but then relinquish control over how their data are processed and transformed into value.

Value creation in GenAI interactions depends on individuals' ongoing investment in shaping both data disclosure and output generation. Within this interaction logic, AI prompt literacy plays a central role by shaping how individuals engage with GenAI systems. Prompt literacy equips individuals with knowledge about how prompts are interpreted, how outputs can be steered, and how iterative refinement strategies can be used to progressively improve results. Rather than directly influencing privacy evaluations, prompt literacy alters the structure and intensity of engagement with the system, thereby shaping the extent to which individuals invest themselves in the interaction. In doing so, AI prompt literacy increases the extent to which individuals actively contribute to the interaction rather than merely providing minimal input.

This change in engagement is most directly captured by the concept of self-investment, or the personal resources individuals devote to a value creation process, including the disclosure of personal information, the expenditure of cognitive effort, the investment of time, and the provision of creative input (Franke, Schreier, and Kaiser 2010; Kaiser, Schreier, and Janiszewski 2017). In GenAI interactions, prompt literacy facilitates higher self-investment by enabling individuals to engage in more extensive prompt formulation, evaluate and adjust system responses, and participate in multiple interaction rounds to refine the output quality (Wu, Terry, and Cai 2022). As a result, prompt-literate individuals invest more of themselves in the interaction process, both in what they disclose and in how intensively they engage. Thus:

H₁: AI prompt literacy increases individuals' perceived self-investment in GenAI interactions.

Self-investment, in turn, reshapes the privacy calculus underlying GenAI interactions. On one hand, higher self-investment amplifies perceived benefits by enabling more relevant, accurate,

and useful outputs that better align with individuals' goals, preferences, and situational needs. On the other hand, self-investment can make privacy exposure more visible because users disclose more contextual information and may reveal more diagnostic information about themselves. Yet GenAI interactions also afford users a sense of agency over what they disclose and how the information is transformed into output. Thus, self-investment may heighten the salience of disclosure while simultaneously making the disclosure feel self-directed and instrumental, thereby changing how privacy risks are integrated into overall privacy value.

The Role of Perceived Ownership

We propose that perceived ownership constitutes a key psychological mechanism through which the outcomes of self-investment are interpreted within the privacy calculus. Perceived ownership refers to the extent to which individuals experience a sense of psychological ownership over an outcome, such that they perceive it as theirs rather than as something produced by an external agent (Csikszentmihalyi and Rochberg-Halton 1981; Durkheim 1957). Cocreation research demonstrates that investing personal resources into shaping an outcome infuses it with elements of the self, thereby fostering ownership perceptions (Franke, Schreier, and Kaiser 2010; Kaiser, Schreier, and Janiszewski 2017; Wiecek, Wentzel, and Erkin 2020). In GenAI contexts, perceived ownership may emerge as individuals observe how their prompts, refinements, and creative inputs are reflected in AI-generated outputs across iterative interaction rounds. Through this traceability between personal input and outcome formation, self-investment is transformed into a subjective sense that the output reflects the self. Thus:

H₂: The effect of AI prompt literacy on individuals' privacy value perceptions is serially mediated, such that greater AI prompt literacy increases perceived self-investment, which in turn enhances perceived ownership.

Ownership enhances perceived benefits by making outputs feel more meaningful, personally

relevant, and aligned with individuals' goals (Belk 1988; Franke, Schreier, and Kaiser 2010).

Ownership may also reshape privacy risk perceptions. Although self-relevant outcomes can make underlying disclosures more salient (Morewedge et al. 2021; Shu and Peck 2011), psychological ownership may increase perceived agency and appropriateness of the disclosure because the output is experienced as something the individual helped create. We therefore expect ownership to increase perceived benefits and to systematically alter risk perceptions, with the empirical studies testing whether the risk pathway is intensified or attenuated in GenAI disclosure contexts.

H3: Greater perceived ownership of GenAI outputs, arising from higher self-investment in the interaction, (a) increases perceived benefits and (b) systematically shapes perceived privacy risks of data disclosure; these benefit and risk evaluations jointly influence individuals' privacy value perceptions.

Finally, we propose that the extent to which self-investment translates into perceived ownership is not uniform across all GenAI interactions. In other words, self-investment does not mechanically translate into higher privacy value. Instead, privacy value emerges from individuals' overall evaluation of the benefits they derive from their investment relative to the privacy risks this investment entails (Culnan and Armstrong 1999; Dinev and Hart 2006; Xu et al. 2011). Because AI prompt literacy systematically increases individuals' level of self-investment in GenAI interactions, its influence on privacy value unfolds indirectly through how this investment activates and structures benefit–risk evaluations. While higher self-investment in GenAI interactions activates both perceived benefits and perceived privacy risks, its implications for privacy value depend on how individuals interpret and integrate these competing evaluations, particularly in contexts in which richer and potentially sensitive data are disclosed.

Further, research on perceived ownership suggests that ownership emerges most strongly when individuals invest resources that are closely tied to the self and diagnostic of personal

identity (Belk 1988; Shu and Peck 2011). In GenAI contexts, such self-relevance is particularly pronounced when interactions involve the disclosure of sensitive data. When individuals disclose sensitive information, their input becomes reflective of who they are, leading them to experience resulting outputs as extensions of the self rather than as generic system outputs (Hermann 2025). In contrast, when interactions rely on nonsensitive or abstract information, self-investment may improve output quality without necessarily triggering a strong sense of ownership. Accordingly, the ownership mechanism through which self-investment shapes the privacy calculus should be activated only in contexts involving sensitive data disclosure. Thus:

H4: Self-investment positively affects perceived ownership of GenAI outputs only when sensitive data is disclosed.

Overview of Studies

To test our theoretical model, we conduct four complementary studies that progressively examine the relationship between AI (prompt) literacy, self-investment, ownership, and privacy value perceptions in GenAI contexts. Table 1 summarizes how each study contributes to the cumulative empirical package. Together, the studies provide converging evidence for our theoretical model (Figure 2) while demonstrating both the robustness and boundary conditions of the AI prompt literacy–self-investment–privacy value relationship in GenAI contexts.

Study 1: AI Literacy in the Field

In Study 1, we investigate whether AI training programs can enhance individuals' self-investment in interactions with GenAI tools and whether this ultimately leads to increased privacy value perceptions. Such training programs, often provided by specialized consultancies, are gaining

importance in light of emerging regulatory frameworks.² In the context of data privacy, training plays a key role by equipping individuals with the knowledge and skills needed to recognize, evaluate, and manage the implications of data disclosure when interacting with AI systems. This view aligns with prior research showing that individuals with greater AI literacy are more likely to engage critically and purposefully with AI tools (Wang, Rau, and Yuan 2023), while lower literacy levels often lead to more passive or unreflective usage (Tully, Longoni, and Appel 2025). Building on this logic, Study 1 explores whether greater AI literacy helps individuals make more self-reflective use of GenAI tools, leading them to view the exchange of data as more worthwhile. It further provides initial insights for our hypotheses.

Method

We collaborated with a German AI literacy consultancy that provides training programs for corporate clients of varying sizes (ranging from 2 to 42 participants for this study; average group size approximately 6–10 participants). The study took place between November 2024 and June 2025. The training sessions were delivered online, in a hybrid format, or on-site and typically lasted between six and eight hours. The course content was tailored to the specific needs of each group, as these may differ depending on a firm's specific business model. The course program is designed for professionals in marketing and communications, bridging foundational technological principles (e.g., the architecture of large language models, neural networks, the distinction between open- and closed-source systems) with critical strategic considerations. Key modules address the legal and ethical landscape, including the implications of the EU AI Act, General Data Protection Regulation compliance, copyright issues, and the mitigation of inherent biases in GenAI. Significant focus is on practical application through advanced prompt

² The EU AI Act requires providers and deployers to take measures to ensure sufficient AI literacy among relevant staff and users (European Union 2024, Art. 4).

engineering techniques and a structured, three-phase quality assurance framework to ensure effective implementation.

Participants received an initial questionnaire before the training. In total, 157 individuals completed the first questionnaire; the second questionnaire was distributed two weeks after the training, and 74 participants responded. To match responses across time points, participants needed to enter a unique code in both surveys, comprising their initials, house number, and birth month. The final sample for analysis consisted of 58 matched cases ($M_{\text{age}} = 39.28$ years; 69% identifying as female), meaning that 58 participants who had past experience with GenAI tools completed both surveys and their responses could be uniquely matched to the same individual. Because the intervention did not include a randomized control group and attrition occurred between waves, we interpret Study 1 as longitudinal field evidence rather than as a stand-alone causal test.

Unlike in our lab studies, we adapted the reference point for the survey questions to the context of the training: Rather than evaluating a specific GenAI interaction, participants reflected more generally on how they currently assess their interactions with GenAI tools in professional settings. The questionnaire contained items from our conceptual model as well as measures for perceived AI use and evaluation literacy (see Table A1 in the Appendix).³ For descriptive purposes and to assess whether the level of data disclosure changed from before to after the training, we included data disclosure to measure whether after the training, participants thought they would disclose as much sensitive data as before the training using two items (“The outcomes

³ Consistent with prior privacy calculus research (e.g., Xu et al. 2011), we conceptualize privacy value as an overall evaluation of whether the benefits of disclosure outweigh its associated privacy risks. Although this conceptualization is necessarily related to benefit and risk perceptions, prior research treats privacy value as a distinct higher-order evaluative judgment rather than as a mathematical difference score. As a robustness check, we repeated the analyses for all studies using a difference score (perceived benefits – perceived privacy risks) instead of the multi-item privacy value scale. The pattern of results and substantive conclusions remained unchanged, suggesting that the reported findings are robust to this alternative operationalization.

I receive in return for my interactions with the generative AI tool involve ...” (a) 1 = “no sensitive information,” 7 = “a lot of sensitive information” and (b) 1 = “no disclosure of sensitive information,” 7 = “a lot of disclosure of sensitive information”; $\alpha_{\text{pre}} = .953$; $\alpha_{\text{post}} = .944$).⁴

Results

Pre-post comparison

To account for individual differences in participants’ initial self-evaluations, we used change scores for our variables (e.g., Δ AI use, Δ AI evaluation, Δ Self-investment) representing the difference between post- and pre-intervention measurements for each variable. Next, we conducted paired sample t-tests (one-tailed) to assess whether our variables changed after the intervention.

We found a substantial increase in AI evaluation literacy ($M_{\text{AIeval_prevpost}} = 4.25$ vs. 4.74; $t(57) = -3.18, p = .001, d = .418$) and a smaller increase in AI use literacy ($M_{\text{AIuse_prevpost}} = 4.53$ vs. 4.78; $t(57) = -1.63, p = .054, d = .214$). Moreover, participants reported higher self-investment ($M_{\text{self_prevpost}} = 3.96$ vs. 4.32; $t(57) = -2.44, p = .009, d = .320$) and perceived ownership ($M_{\text{own_prevpost}} = 3.21$ vs. 3.66; $t(57) = -2.41, p = .010, d = .316$). We observe smaller changes in perceived risk ($M_{\text{risk_prevpost}} = 4.93$ vs. 4.73; $t(57) = 1.12, p = .134, d = .147$), perceived privacy value ($M_{\text{value_prevpost}} = 4.25$ vs. 4.31; $t(57) = -.34, p = .368, d = .044$) and the increase in perceived benefit ($M_{\text{benefit_prevpost}} = 5.32$ vs. 5.49; $t(57) = -1.36, p = .090, d = .179$). Finally, we found no substantial change in participants’ data disclosure (i.e., comparing whether they disclose more or less sensitive information when interacting with GenAI tools before vs. after the training; $M_{\text{disclosure_prevpost}} = 2.50$ vs. 2.61; $t(57) = -.50, p = .311, d = -.065$).

Mediation analyses

⁴ Due to our partner’s request, we also asked questions designed to help tailor the training content to the specific needs of each group. These included general questions about GenAI tool usage behaviors and trust in GenAI tools and open-ended questions about participants’ expectations for the training.

We conducted a multiple linear regression analysis with Δ AI use literacy and Δ AI evaluation literacy as predictors of Δ Self-investment, to glean initial insights into the proposed engagement mechanism underlying H₁. We found that the overall model predicts self-investment change ($F(2, 55) = 4.91, p = .011$). AI use literacy is a significant, positive predictor of self-investment ($b = .312, SE = .148, \beta = .326, t = 2.11, p = .039$). However, the results show that AI evaluation literacy does not significantly predict self-investment ($b = .091, SE = .149, \beta = .094, t = .61, p = .546$). Given these insights, we next focused on AI use literacy in our analyses.

Next, we used PROCESS Model 4 (5,000 bootstrap resamples; Hayes 2022) to examine whether changes in AI use literacy are indirectly associated with privacy value perceptions through perceived self-investment, consistent with the proposed engagement-based process logic. The results revealed an indirect effect of AI use literacy on privacy value via self-investment ($a \times b = .172, SE = .089, 95\% CI [.024, .374]$), indicating that increases in AI use literacy translate into more favorable privacy value perceptions through higher perceived self-investment.

Discussion

Study 1 provides longitudinal field evidence that AI literacy training can shape how individuals engage with GenAI tools in professional contexts. Participants reported increases in both AI use and evaluation literacy following the training, but only gains in AI use literacy predicted higher self-investment. This pattern suggests that practical, action-oriented capabilities, such as knowing how to actively guide GenAI systems, are central to engagement. Given the absence of a randomized control group and the attrition inherent in the field setting, Study 1 is best interpreted as ecologically rich evidence motivating the controlled tests that follow.

Study 2: Exploring the Effect of AI Prompt Literacy on Privacy Value Perceptions

Study 2 directly tests our hypotheses by focusing on AI prompt literacy as a particularly salient manifestation of AI use literacy in GenAI contexts (Hwang, Lee, and Shin 2023).

Method

Five hundred fifty-three students from a U.S. university ($M_{\text{age}} = 23.22$ years; 62.7% identifying as female) participated in a between-subjects online experiment employing a two-cell (AI prompt literacy: no vs. yes) design.⁵ The study (preregistered at <https://aspredicted.org/en73tw.pdf>) received ethics approval from the German Association for Experimental Economic Research e.V. (GfeW) before data collection. At the beginning of the survey, participants reported their experience with GenAI tools. We assessed usage frequency with the item, “How often do you use GenAI tools (e.g., ChatGPT, Microsoft Copilot, Claude)?” using a six-point categorical scale (1 = “several times a day,” 6 = “less than once a month”). In addition, participants indicated their self-assessed expertise with the item, “How much expertise do you have in using GenAI tools such as ChatGPT?” on a slider scale (0 = “I have no expertise at all,” 10 = “I am an expert”). Participants also reported frequent GenAI use, with 71.8% indicating daily use; self-assessed expertise was moderate on average ($M = 6.54$, $SD = 2.03$), consistent with our expectations.

After participants provided their consent and were informed about the voluntary and anonymous nature of their participation, we randomly assigned them to one of two conditions. In the AI prompt literacy condition, participants received a short training titled “Master the Prompt Design Formula,” which introduced the key building blocks of an effective prompt for GenAI (see Table A2 in the Appendix for study stimuli). In the control condition, participants completed structurally and linguistically comparable training titled “Master the YouTube Video Selection Formula,” which provided guidance on how to select suitable YouTube videos. Both trainings were carefully matched in terms of length, structure, visual design, and reading difficulty to ensure comparable exposure time and cognitive engagement, thereby enabling us to isolate the

⁵ The initial sample included 675 participants. As preregistered, we excluded participants depending on both their quiz scores and the quality of their responses in the open text fields (i.e., whether they completed the task of writing a prompt as requested).

effect of the training content itself. After the training, participants completed a short quiz corresponding to their assigned condition to reinforce engagement with the material (i.e., prompt-related questions in the AI prompt literacy condition and YouTube-related questions in the control condition).

All participants then received the same task scenario that asked them to imagine that they were applying for an important job and wanted to write a strong cover letter with the help of ChatGPT. We chose this context because applying for a job is realistic and familiar for most participants while naturally involving the disclosure of personal information. Importantly, participants wrote only the initial prompt they would enter into ChatGPT; they did not receive a GenAI-generated output or engage in iterative refinement. Study 2 therefore provides a conservative test of the early prompt-formulation stage and examines anticipated privacy value, benefits, risks, and ownership of the expected output rather than evaluations of an actual generated output. Participants who received prompt literacy training spent more time on the prompt-writing task and produced longer prompts than those in the control condition (time: $M_{\text{YouTube}} = 2.31$ min vs. $M_{\text{Prompt}} = 5.78$ min; word count: $M_{\text{YouTube}} = 29.57$ vs. $M_{\text{Prompt}} = 95.14$), providing behavioral evidence that the manipulation altered engagement with the prompting task.

After completing the prompt-writing task, participants completed measures of perceived privacy value, benefits, risks, ownership, and self-investment using established seven-point Likert scales from prior research (see Table A1 in the Appendix). To assess the effectiveness of the manipulation, participants completed a manipulation check indicating the extent to which they had received information about selecting a YouTube video during the first part of the survey (1 = “strongly disagree,” 7 = “strongly agree”). We also measured participants’ prior knowledge of the assigned training content. Because those who already possess prompt-related knowledge may benefit less from the training, we included prior knowledge as a covariate in the primary

analyses. Importantly, this covariate reflects familiarity with the assigned training content rather than general GenAI usage or expertise. Consistent with successful random assignment, participants did not differ across conditions in either GenAI usage frequency ($t(551) = 0.64, p = .520$) or self-reported GenAI expertise ($t(551) = 0.10, p = .923$). For transparency, we additionally report no-covariate analyses in Table A3 in the Appendix.

Results

Manipulation check

An analysis of variance (ANOVA) confirmed the effectiveness of the manipulation. Participants in the YouTube training (control) condition reported significantly greater agreement with having received YouTube-related instructional content than those in the AI prompt literacy condition ($M_{\text{YouTube}} = 5.50$ vs. $M_{\text{Prompt}} = 2.52$; $F(1, 551) = 539.653, p < .001, \eta_p^2 = .495$). Because this manipulation check verifies exposure to the control content rather than prompt-skill acquisition directly, we also rely on the behavioral prompt-writing indicators reported above as evidence that the prompt-literacy manipulation increased task engagement.

Privacy value

ANCOVA on perceived privacy value, controlling for prior knowledge with the training content, revealed that the main effect of AI prompt literacy was not significant ($F(1, 550) = 1.737, p = .188, \eta_p^2 = .003$), while the covariate was significant ($F(1, 550) = 19.272, p < .001, \eta_p^2 = .034$). Participants in the AI prompt literacy condition did not report substantially different privacy value perceptions from those in the control condition ($M_{\text{Prompt}} = 4.778$ vs. $M_{\text{YouTube}} = 4.634$).

Privacy calculus

ANCOVA on perceived benefits revealed a marginally significant effect of AI prompt literacy ($F(1, 550) = 3.107, p = .078, \eta_p^2 = .006$). Participants in the AI prompt literacy condition perceived greater benefits than those in the control condition ($M_{\text{Prompt}} = 5.411$ vs. $M_{\text{YouTube}} =$

5.215). In addition, the analysis revealed a significant effect of prior knowledge ($F(1, 550) = 9.866, p = .002, \eta_p^2 = .018$). In contrast, an ANCOVA on perceived risks showed no significant effect of AI prompt literacy ($F(1, 550) = .227, p = .634, \eta_p^2 < .001$), with similar reported risk perceptions in both conditions ($M_{\text{Prompt}} = 4.190$ vs. $M_{\text{YouTube}} = 4.251$). The effect of prior knowledge was also not significant ($F(1, 550) = 2.211, p = .138, \eta_p^2 = .004$).

Perceived output ownership

The ANCOVA results revealed a marginally significant effect of AI prompt literacy on perceived output ownership ($F(1, 550) = 2.745, p = .098, \eta_p^2 = .005$). Participants exposed to the AI prompt literacy training reported greater ownership perceptions than those in the control condition ($M_{\text{Prompt}} = 4.011$ vs. $M_{\text{YouTube}} = 3.797$). Moreover, the effect of prior knowledge on perceived output ownership was significant ($F(1, 550) = 18.894, p < .001, \eta_p^2 = .033$).

Perceived self-investment

Finally, ANCOVA revealed a significant effect of AI prompt literacy on self-investment ($F(1, 550) = 6.016, p = .014, \eta_p^2 = .011$). Participants in the AI prompt literacy condition reported higher self-investment than those in the control condition ($M_{\text{Prompt}} = 3.773$ vs. $M_{\text{YouTube}} = 3.491$). The effect of prior knowledge on perceived self-investment was significant ($F(1, 550) = 15.238, p < .001, \eta_p^2 = .027$), consistent with H1.

Mediation analyses

We first estimated a serial mediation model using PROCESS Model 6 with 5,000 bootstrap resamples (Hayes 2022) to test H₂, controlling for prior knowledge in all equations. The results reveal a positive serial indirect effect of AI prompt literacy on privacy value via self-investment and perceived ownership ($a \times b = .0344, SE = .0183, 95\% CI [.0062, .0772]$). These results confirm H₂ that AI prompt literacy is indirectly associated with higher privacy value perceptions through increased self-investment and perceived ownership.

Next, we incorporated the privacy calculus in the model to test H₃. Using a customized PROCESS model with 5,000 bootstrap resamples (Hayes 2022), again controlling for prior knowledge, we examined whether perceived ownership translated into privacy value through perceived benefits and privacy risks. Perceived ownership increased perceived benefits and decreased perceived privacy risks, both of which were significantly related to privacy value in the expected directions (Figure 3). Consistent with this pattern, the serial indirect effect of AI prompt literacy on privacy value via self-investment, perceived ownership, and perceived benefits was positive and statistically significant (indirect effect = .036, SE = .016, 95% CI [.007, .068]). A smaller serial indirect effect also emerged via self-investment, perceived ownership, and perceived privacy risks (indirect effect = .003, SE = .003, 95% CI [.000, .009]; unrounded lower bound > 0). These results support the benefit pathway in H_{3a}. For the risk pathway, the results show attenuation rather than intensification: greater perceived ownership is associated with lower perceived privacy risks.

Discussion

By focusing on AI prompt literacy as a specific manifestation of AI use literacy, Study 2 examines how differences in prompting capabilities shape early engagement in GenAI interactions. The findings show that prompt literacy increases perceived self-investment even when individuals only formulate an initial prompt (supporting H₁). In line with our conceptual framework, AI prompt literacy did not directly affect privacy value perceptions. Instead, it increased self-investment, which elicited greater anticipated ownership of the expected AI-generated output and subsequent changes in benefit and risk evaluations that jointly informed anticipated privacy value. Consistent with H_{3a}, perceived ownership amplified anticipated benefits. For the risk pathway, greater perceived ownership was associated with lower perceived

privacy risks, suggesting that ownership can attenuate rather than intensify risk perceptions in open-ended GenAI disclosure contexts.

While Study 2 shows that self-investment, anticipated ownership, and benefit–risk evaluations can be activated even before an output is generated, the single prompt engagement cannot fully capture how privacy value develops when individuals repeatedly interact with and actively shape GenAI outputs. To this end, Study 3 increases the intensity of the interaction by allowing iterative engagement with the system, thereby examining how higher self-investment translates into downstream privacy value perceptions under controlled laboratory conditions.

Study 3: Manipulating Self-Investment

Study 3 extends the previous studies by examining the consequences of heightened self-investment in GenAI interactions. While Studies 1 and 2 establish AI prompt literacy as an important antecedent of self-investment, Study 3 focuses on the individuals' interactions with AI and investigates how varying levels of self-investment shape downstream privacy-related perceptions under controlled laboratory conditions, further investigating H₃. By allowing participants to iteratively interact with the GenAI system and evaluate resulting outcomes, the study captures key features of real-world GenAI use while maintaining experimental control.

Method

We recruited 176 students⁶ from a German university ($M_{\text{age}} = 22.15$ years; 68.2% identifying as female) for a laboratory experiment employing a two-cell (self-investment: low vs. high) between-subjects design. The experiment (preregistered at <https://aspredicted.org/ngwn-vj9k.pdf>) received ethics approval from the GfEW. The study took place in the university lab in January

⁶ The initial sample consisted of 190 participants. We excluded 14 participants because of technical issues with ChatGPT, because they talked with other participants, or because they failed the quality check (stating how well their letter fit or did not fit with the requirements).

2025. After participants provided consent regarding data use and were informed that they could withdraw from the study at any time without consequence, we briefly introduced them to ChatGPT, the GenAI tool used in the study. We then instructed them to write a creative seminar application cover letter for a fictitious seminar offered by “Inspirely Agency” (see Table A4 in the Appendix for study stimuli). We chose this task context because it closely mirrors common application procedures at the university, where students regularly apply for seminars via written motivation letters, thereby providing a highly realistic setting that participants are likely to have encountered before. At the same time, the task inherently involves disclosure, making it well suited for examining self-investment in GenAI-assisted data disclosure processes.

To avoid potential bias from using personal accounts, participants either accessed ChatGPT without logging in or, in cases of technical difficulties, were provided with log-in credentials managed by the study team. The study supervisors ensured that all data were deleted following each interaction to maintain privacy and data protection standards.

All participants used the same initial prompt to generate a draft via ChatGPT. In the low-self-investment condition, participants could not modify the output but had to base their subsequent evaluation on the initial output of the GenAI tool. In contrast, those in the high-self-investment condition could iteratively refine and personalize the letter and had a minimum interaction time of three minutes to ensure active engagement. A quality check followed to verify task compliance, and participants had the option to submit their final letter but were not forced to do so for privacy reasons.⁷ After completing the task, participants answered survey items on perceived privacy value, perceived benefits and risks, perceived ownership and completed a self-investment manipulation check as in the previous studies using seven-point Likert scales (see

⁷ We did not analyze the final letters in more detail or include them as a second dependent variable as preregistered. Nearly all participants (except nine—three in the low- and six in the high-self-investment condition) chose to share their letters with the researchers.

Table A1 in the Appendix). In addition, an attention check assessed the perceived fit and creativity of the cover letter. To align the time spent on the task across conditions, participants in the low-self-investment group completed a filler task unrelated to the AI interaction.

Results

Manipulation check

An ANOVA on self-investment in the process confirmed the effectiveness of the manipulation. Participants in the low-self-investment condition reported significantly lesser self-investment than those in the high-self-investment condition ($M_{\text{low_self}} = 1.74$ vs. $M_{\text{high_self}} = 2.96$; $F(1, 174) = 99.30, p < .001, \eta_p^2 = .363$).

Privacy value

The ANOVA on perceived privacy value revealed a significant effect of self-investment ($F(1, 174) = 6.359, p = .013, \eta_p^2 = .035$). Participants in the low-self-investment condition reported lower privacy value perceptions than those in the high-self-investment condition ($M_{\text{low_self}} = 4.17$ vs. $M_{\text{high_self}} = 4.63$). This result is consistent with the predicted base effect of self-investment on privacy value

Privacy calculus

Similarly, an ANOVA on perceived benefits revealed a significant main effect of self-investment ($F(1, 174) = 33.002, p < .001, \eta_p^2 = .159$). Participants in the low-self-investment condition perceived lesser benefits than those in the high-self-investment condition ($M_{\text{low_self}} = 4.21$ vs. $M_{\text{high_self}} = 5.34$). However, consistent with the findings from Study 2, an ANOVA on risk perceptions revealed no significant effect of self-investment ($M_{\text{low_self}} = 4.30$ vs. $M_{\text{high_self}} = 4.12$; $F(1, 174) = .376, p = .541, \eta_p^2 = .002$).

Perceived output ownership

An ANOVA revealed a significant effect of self-investment on perceived output ownership ($F(1, 174) = 7.547, p = .007, \eta_p^2 = .042$). Participants in the low-self-investment condition reported lower perceptions of output ownership than those in the high-self-investment condition ($M_{\text{low_self}} = 2.61$ vs. $M_{\text{high_self}} = 3.18$).

Mediation analyses

To evaluate H_3 , we used a customized model of PROCESS (5,000 resamples; Hayes 2022). As predicted, we found significant serial mediation effects on privacy value through benefit perceptions, in further support of H_{3a} . Specifically, higher self-investment led to greater perceived ownership, which in turn increased perceived benefits and, ultimately, privacy value ($a \times b = .0207, 90\% \text{ CI } [.0009, .0573]$). However, we found no support for H_{3b} as the serial mediation of self-investment on privacy value through perceived ownership and risk was not significant ($a \times b = .0071, 90\% \text{ CI } [-.0007, .0201]$).

Discussion

Study 3 extends the previous findings by establishing the causal role of self-investment in shaping privacy value perceptions in GenAI interactions. By experimentally manipulating the extent to which individuals could iteratively engage with the system, the study isolates self-investment as a central process variable underlying privacy value formation. Consistent with the conceptual framework and the findings of Studies 1 and 2, higher self-investment led to stronger perceived ownership of the AI-generated output, greater perceived benefits, and, ultimately, higher privacy value perceptions. These findings provide causal evidence that self-investment itself drives downstream privacy-related evaluations when individuals actively engage with a GenAI system in the course of a realistic task.

Unlike Study 2, the indirect effect through perceived privacy risks was not statistically significant. One possible explanation is that participants' primary objective was to iteratively

improve the quality of their application letter during an actual interaction with ChatGPT, such that their evaluations centered more strongly on the benefits generated through the interaction than on potential privacy risks. Study 4 extends these findings by explicitly varying the sensitivity of the disclosed data, allowing us to examine whether the self-investment–ownership mechanism depends on whether or not the interaction is self-relevant.

Study 4: Sensitive Data Disclosure as a Boundary Condition

Study 4 tests H4 by examining whether sensitive data disclosure during GenAI interactions moderates the effect of self-investment on perceived ownership. By varying disclosure sensitivity, the study tests when self-investment becomes consequential for ownership formation. The scenario-based design also lets us examine how perceived ownership relates to privacy risk perceptions and privacy value when data disclosure is made salient.

Method

Two hundred sixty-seven undergraduate students from a U.S. university ($M_{\text{age}} = 19.88$ years; 63.6% identifying as female) participated in this 2 (self-investment: low vs. high) $\times$ 2 (sensitive data disclosure: no vs. yes) between-subjects online scenario experiment. We randomly assigned participants to one of four conditions (see Table A5 in the Appendix for study stimuli). We asked participants to imagine they wanted to write a fairytale with the help of ChatGPT. The task of writing a fairytale provides an abstract but intuitive setting in which differences in sensitive data disclosure can be clearly communicated and evaluated. We manipulated low self-investment by explaining that the process involved standard approaches to tell the story, little prompting, and that the participant would accept anything that was created by ChatGPT without any refinement or changes. In the high-self-investment conditions, the process involved creative approaches, heavy prompting, and refinement of the suggestions made by ChatGPT. To manipulate the data

context, we told participants in the conditions without sensitive data disclosure that the story would be purely fictional, not involving any personal data; we told participant in the conditions with sensitive data disclosure that the story was about them, so they needed to provide their personal data (e.g., name, details about physical appearance, religious beliefs). After reading the scenario, participants reported their perceptions of privacy value, benefits, risks, and ownership (see Table A1 in the Appendix for items). Finally, participants completed the manipulation checks on self-investment and data disclosure (“The outcome I will receive in return for my interaction with the generative AI tool involves ...” (a) 1 = “no self-disclosure,” 7 = “a lot of self-disclosure” and (b) 1 = “none of my personal information,” 7 = “a lot of my personal information”; $\alpha = .872$).

Results

Manipulation check

To assess our sensitive data disclosure manipulation, we conducted a two-way ANOVA. The results revealed a significant main effect for sensitive data disclosure ($M_{\text{low_self,no_data}} = 3.46$, $M_{\text{high_self,no_data}} = 3.74$, $M_{\text{low_self,yes_data}} = 4.33$, $M_{\text{high_self,yes_data}} = 4.29$; $F(1, 263) = 20.015$, $p < .001$, $\eta_p^2 = .071$), while neither the self-investment main effect nor the interaction was significant.

A two-way ANOVA showed a significant main effect of our self-investment manipulation ($M_{\text{low_self,no_data}} = 2.96$, $M_{\text{low_self,yes_data}} = 2.83$, $M_{\text{high_self,no_data}} = 3.43$, $M_{\text{high_self,yes_data}} = 4.03$; $F(1, 263) = 32.880$, $p < .001$, $\eta_p^2 = .111$) but no significant main effect of our data disclosure sensitivity manipulation. However, we found a significant self-investment $\times$ sensitive data disclosure interaction effect ($F(1, 263) = 6.608$, $p = .011$, $\eta_p^2 = .025$). Contrast analyses showed a significant difference between the sensitive data disclosure and no sensitive data disclosure conditions only in the high-self-investment condition ($F(1, 263) = 8.807$, $p = .003$, $\eta_p^2 = .032$). We anticipated this result as it aligns with our conceptualization: In GenAI data disclosure

processes, individuals invest parts of themselves to improve the generated output by contributing (1) inputs regarding whether, what, and how much data are disclosed and (2) guidance on how the tool should process the data. The two manipulation checks were only weakly correlated ($r = .22$), indicating that participants distinguished data disclosure sensitivity from self-investment.

Privacy value

A two-way ANOVA revealed significant main effects of self-investment ($F(1, 263) = 4.968, p = .027, \eta_p^2 = .019$) on privacy value perceptions; there was no main effect of sensitive data disclosure representing the interaction context ($F(1, 263) = .003, p = .956, \eta_p^2 = .000$). However, the results show a significant interaction ($F(1, 263) = 4.370, p = .038, \eta_p^2 = .016$).

We conducted contrasts to further explore the two-way interaction. When self-investment was low, perceived privacy value did not differ significantly between the condition with and without sensitive data disclosure ($M_{\text{low_self,no_data}} = 4.33$ vs. $M_{\text{low_self,yes_data}} = 4.05$; $F(1, 263) = 2.066, p = .152, \eta_p^2 = .008$). Similarly, when self-investment was high, perceived privacy value did not differ significantly between the sensitive data disclosure condition and the no sensitive data disclosure condition ($M_{\text{high_self,no_data}} = 4.35$ vs. $M_{\text{high_self,yes_data}} = 4.65$; $F(1, 263) = 2.308, p = .130, \eta_p^2 = .009$). In the condition without sensitive data disclosure, perceived privacy value did not significantly differ between low and high self-investment ($M_{\text{low_self,no_data}} = 4.33$ vs. $M_{\text{high_self,no_data}} = 4.35$; $F(1, 263) = .010, p = .921, \eta_p^2 = .000$). However, in the condition with sensitive data disclosure, perceived privacy value was significantly higher when self-investment was high rather than low ($M_{\text{low_self,yes_data}} = 4.05$ vs. $M_{\text{high_self,yes_data}} = 4.65$; $F(1, 263) = 9.091, p = .003, \eta_p^2 = .033$).

Privacy calculus

A two-way ANOVA revealed a significant main effect of self-investment on perceived benefits ($F(1, 263) = 13.245, p < .001, \eta_p^2 = .048$). However, there was no main effect of sensitive data

disclosure ($F(1, 263) = .295, p = .587, \eta_p^2 = .001$). The interaction between self-investment and sensitive data disclosure was also not significant ($F(1, 263) = 2.242, p = .135, \eta_p^2 = .008$). The two-way ANOVA revealed no significant main effect of self-investment on perceived risks ($F(1, 263) = .181, p = .671, \eta_p^2 = .001$). We found a significant effect of data disclosure sensitivity on risk perceptions ($F(1, 263) = 5.555, p = .019, \eta_p^2 = .021$). Yet the interaction between self-investment and sensitive data disclosure was not statistically significant ($F(1, 263) = 1.308, p = .254, \eta_p^2 = .005$).

Perceived output ownership

A two-way ANOVA revealed that self-investment had a significant main effect on perceived output ownership ($F(1, 263) = 11.746, p < .001, \eta_p^2 = .043$); the main effect of sensitive data disclosure was not significant ($F(1, 263) = 1.404, p = .237, \eta_p^2 = .005$). However, the interaction between self-investment and sensitive data disclosure was significant ($F(1, 263) = 5.123, p = .024, \eta_p^2 = .019$).

We conducted contrasts to further explore the significant two-way interaction. In the low-self-investment conditions, perceived output ownership did not differ significantly between the conditions without and with sensitive data disclosure ($M_{\text{low_self,no_data}} = 3.30$ vs. $M_{\text{low_self,yes_data}} = 3.12$; $F(1, 263) = .580, p = .447, \eta_p^2 = .002$). However, for participants in the high-self-investment condition, perceived output ownership was significantly greater in the condition with (vs. without) sensitive data disclosure ($M_{\text{high_self,no_data}} = 3.50$ vs. $M_{\text{high_self,yes_data}} = 4.08$; $F(1, 263) = 5.961, p = .015, \eta_p^2 = .022$). In the conditions without sensitive data disclosure, self-investment did not significantly influence perceived ownership ($M_{\text{low_self,no_data}} = 3.30$ vs. $M_{\text{high_self,no_data}} = 3.50$; $F(1, 263) = .695, p = .405, \eta_p^2 = .003$). Yet, in the conditions with sensitive data disclosure, individuals with high levels of self-investment reported greater perceived ownership than those

with low levels of self-investment ($M_{\text{low_self,yes_data}} = 3.12$ vs. $M_{\text{high_self,yes_data}} = 4.08$; $F(1, 263) = 15.780, p < .001, \eta_p^2 = .057$), supporting H₄.

Moderated mediation analysis

We estimated the indirect effect of self-investment $\times$ sensitive data disclosure on privacy value through perceived output ownership and perceived benefits and risks using a customized PROCESS model (5,000 resamples; Hayes 2022; see Figure 4). In terms of the indirect effect via perceived ownership and benefits, the results show that when no sensitive data disclosure was involved, the effect was not significant ($a \times b = .026$; 95% CI $[-.035, .095]$). However, when the sensitive data disclosure was involved, the indirect effect was significant ($a \times b = .125$; 95% CI $[.046, .233]$). The index of moderated mediation was also significant ($b = .099$; 95% CI $[.013, .219]$), confirming that sensitive data disclosure moderates the strength of this mediation effect. Similarly, the indirect effect via perceived ownership and risks was nonsignificant in the no sensitive data disclosure conditions ($a \times b = .005$; 95% CI $[-.007, .026]$) but became significant when sensitive data disclosure was involved ($a \times b = .025$; 95% CI $[.003, .066]$). The index of moderated mediation was again significant ($b = .020$, 95% CI $[.001, .059]$), indicating that the mediation path through perceived risks also varied as a function of sensitive data disclosure. Together, these results confirm H₄ and further inform H_{3a} and H_{3b} by revealing that perceived ownership translates self-investment into privacy value through its positive effect on perceived benefits and negative effect on perceived privacy risks, contingent on presence of sensitive data disclosure.

Discussion

Study 4 examines how the effects of self-investment on privacy value perceptions unfold across GenAI interactions with and without sensitive data disclosure. Building on Study 3, which demonstrated that self-investment increases perceived ownership during an actual GenAI

interaction, Study 4 investigates when this ownership mechanism is most likely to emerge. The results show that the relationship between self-investment and perceived ownership is contingent on whether the interaction involves sensitive personal data disclosure. Consistent with H4, the positive effect of self-investment on perceived ownership is significantly stronger when interactions involve sensitive personal data disclosure than when they do not. These findings suggest that self-investment becomes psychologically consequential when the disclosure of sensitive personal data renders the generated output more self-relevant, thereby strengthening perceived ownership.

Building on this contextualized ownership mechanism, Study 4 provides additional evidence consistent with H3a. When sensitive data disclosure makes the generated output more self-relevant, self-investment translates into stronger perceived ownership, which subsequently increases perceived benefits and, ultimately, privacy value. Thus, the positive downstream consequences of self-investment depend not only on how much consumers invest in the interaction but also on whether that investment involves sensitive personal data.

Beyond this ownership-based mechanism, Study 4 reinforces the risk-attenuation pattern observed in Study 2. Perceived ownership was associated with lower perceived privacy risks even when the interaction involved sensitive personal data disclosure. This finding suggests that the presence of sensitive personal data disclosure does not necessarily amplify perceived privacy risks once individuals perceive the generated output as their own. Instead, perceived ownership can attenuate perceived privacy risks despite the increased visibility and sensitivity of the disclosed data.

Study 4's findings extend the previous studies by showing that the relationship between self-investment and perceived ownership is context dependent, whereas the downstream risk-attenuation effect of ownership remains stable across disclosure contexts. Together, these

findings suggest that self-investment alone is not always sufficient to foster perceived ownership. Rather, self-investment is most likely to translate into ownership when the disclosure process is sufficiently self-relevant through the inclusion of sensitive personal data. Even under conditions involving sensitive personal data disclosure, perceived ownership structures privacy value evaluations by amplifying perceived benefits and reducing perceived privacy risks.

General Discussion

GenAI marks an important shift in consumer–technology interaction. Individuals are no longer merely passive recipients of algorithmic outputs; they actively guide systems through prompting, contextual disclosure, and iterative refinement. Across four studies, we show that even modest literacy-enabled engagement systematically shapes how individuals evaluate data disclosure in GenAI interactions. Higher self-investment in guiding GenAI systems increases privacy value perceptions through perceived output ownership and the privacy calculus. Notably, perceived ownership is associated with lower perceived privacy risks, even when data disclosure is highly salient. These findings suggest that privacy in GenAI contexts is shaped less by data disclosure alone and more by how capable and invested individuals are in shaping the interaction.

Importantly, self-investment becomes psychologically consequential when the disclosure of sensitive personal data renders the generated output self-relevant, making privacy a dynamic, cocreated process rather than a one-time consent decision.

Theoretical Contributions

This research offers three main contributions. First, we extend data disclosure and privacy research by challenging the assumption that greater disclosure in AI-enabled environments is evaluated only as a cost (e.g., Goodwin 1991; Phelps, Nowak, and Ferrell 2000). Our findings show that in contexts in which individuals actively engage in value creation, greater self-

investment can produce more beneficial and personally meaningful outcomes. Moreover, these outcomes become particularly meaningful when interactions involve sensitive data disclosure, because such disclosure renders the generated output more self-relevant and therefore more likely to evoke perceived ownership. Crucially, these contexts are characterized by flexible and iterative disclosure, in which individuals are not required to provide predefined information but can decide what and how much information to disclose in pursuit of their own goals. As such, the benefits individuals receive and value in return for their data extend beyond monetary incentives or firm-initiated personalization (e.g., Chellappa and Sin 2005; Hui, Teo, and Lee 2007; Melzner, Bonezzi, and Meyvis 2023). Rather than merely compensating for privacy intrusions, the benefits resulting from GenAI cocreation include meaningful task support, alignment with personal goals, self-expression, and perceived ownership over the outcome.

GenAI-generated benefits are therefore not simply compensations for lost privacy but valued outcomes that individuals actively seek and help create. This finding highlights the need to rethink traditional linear models of data disclosure and to conceptualize GenAI disclosure as a co-creative, iterative process by which individuals contribute not only data but also guidance on how those data should be used (Krafft et al. 2021; Puntoni et al. 2021; Zimmermann et al. 2024). Furthermore, we show that positive privacy value perceptions are contingent on the extent to which individuals can customize the disclosure process to reflect themselves. This customization fosters perceived ownership over the output, which in turn amplifies perceived benefits and lowers perceived privacy risks.

Second, this research contributes to the emerging body of GenAI literature in marketing by shifting the focus from output-level capabilities, such as content generation, synthetic data use, and preference modeling, to the interactive, input-level processes that shape value cocreation in human–AI collaboration. While previous work has emphasized the technical benefits of GenAI

for marketers (e.g., efficiency, personalization, superiority over human work; Carlson et al. 2023; Huang and Rust 2025), our findings highlight the critical role of individual engagement in the data disclosure process in shaping output evaluations. To advance this view, we move beyond minimally engaging individuals in automated interactions (e.g., simply pushing a button; Wiecek, Wentzel, and Erkin 2020), which may create only an illusion of control over data inputs and outputs (e.g., Brandimarte, Acquisti, and Loewenstein 2013); instead, our research shows that AI prompt literacy constitutes a key individual capability that enables meaningful engagement with GenAI systems. Specifically, we demonstrate that even minimal improvements in prompt formulation, such as a single, more literate prompt, can increase perceived self-investment and produce positive downstream effects on ownership and privacy value perceptions. As such, AI prompt literacy supports forms of engagement that foster strong connections between individual input and output evaluations. In doing so, this work answers recent calls for a more nuanced understanding of the user experience in GenAI contexts (e.g., Grewal et al. 2025; Peres et al. 2023) by highlighting not only what individuals gain from GenAI outputs but also how the interaction itself can create psychological value. Specifically, we contribute to the underexplored area of interactive data disclosure in which individual engagement is not merely instrumental for improving outputs but can foster a stronger psychological connection with the resulting benefits, making the act of engaging with GenAI more personally meaningful.

Third, this research contributes a data disclosure perspective to the cocreation and customization literature (e.g., Franke, Schreier, and Kaiser 2010; Klesse et al. 2019; Troye and Supphellen 2012). By identifying inputs in customization and cocreation tools as a form of data disclosure, we show that individuals' evaluations of the resulting output are not driven solely by what they receive. These evaluations are also shaped by how they give or, in other words, how the contribution of their personal information influences their sense of ownership and the

subjective (privacy) value of the exchange (e.g., Franke and Schreier 2010). Importantly, we illustrate that ownership emerging from such contributions does not heighten perceived privacy risks but rather systematically reshapes privacy evaluations in ways that favor value creation. Our findings further show that this relationship is strengthened when individuals possess greater AI prompt literacy, which enables them to navigate these interactions more effectively. In doing so, we show that every act of cocreation is also an act of data disclosure that generates not only tailored outcomes for individuals but also informational value for firms.

Practical Implications

Our findings have three major managerial implications, providing valuable insights into why and how to integrate GenAI tools into customer- and employee-facing processes in ways that create mutual value. First, firms should focus on helping individuals invest themselves in the *input*. When individuals believe that their input meaningfully shapes the output, the disclosure process becomes more purposeful and empowering, potentially leading to a more positive evaluation of the interaction and greater openness to future data sharing. A central lever for enabling such self-investment is AI prompt literacy. To support this dynamic and help individuals understand what data to share and how to do so in GenAI interactions, firms can actively foster AI literacy. This can involve offering employee trainings that focus specifically on effective prompt formulation and refinement or integrating prompt-related guidance into customer-facing applications (e.g., features suggesting follow-up questions, alternative phrasing to help individuals iteratively shape the output). Such initiatives not only support compliance with emerging regulations such as the EU AI Act but also empower individuals with practical capabilities for collaborating with GenAI tools. The knowledge about how best to collaborate with GenAI is becoming increasingly important as these tools begin to take on roles similar to virtual companions (Huang and Rust 2024).

However, we caution firms in general and application developers in particular to approach the design of these processes ethically to avoid encouraging individuals to unintentionally and excessively disclose sensitive data because they feel emotionally connected with the process and resulting output (Martin and Zimmermann 2024). Ethical implementation should ensure that literacy-enhancing features support informed engagement rather than obscure disclosure boundaries, thereby preserving individual autonomy and privacy. In business contexts, we recommend collaborating closely with GenAI providers and establishing contractual agreements to safeguard data security and confidentiality (Deloitte Legal 2024). At the regulatory level, we emphasize the need for policies that not only protect personal data in GenAI interactions but also account for the emotional and psychological dynamics of the process.

Second, firms should recognize that the value of GenAI-generated *outputs* lies not only in their functional utility but also in the personal meaning individuals attribute to them. Encouraging self-investment (e.g., by supporting prompt refinement, iterative feedback, or interfaces that make the link between user input and output changes transparent) can enhance individuals' sense of ownership, making outputs feel more aligned with their identity and goals. In turn, firms benefit from outputs that are more relevant and of higher quality, while also amplifying human creativity. For example, in Coca-Cola's (2023) "Create Real Magic" campaign, digital creatives used GPT-4 and DALL-E along with Coca-Cola brand assets to cocreate original artwork. By encouraging participants to shape prompts, iterate on generated visuals, and submit their creations for the chance to appear on global billboards, the platform fostered deep self-investment, generating compelling content and valuable data on individual preferences and engagement.

However, legal ownership and data usage rights must be considered whenever firms use data provided by individuals through GenAI tools or leverage cocreated outputs: While

individuals may feel a sense of perceived ownership over outputs they helped create through (literate) self-investment, they are often unaware of the actual legal terms. This may especially be relevant in contexts such as contests, branded chatbot interactions, or standard GenAI use when data may be reused for model training or other purposes. To ensure informed consent, firms must provide clear and accessible information about how they will use and who holds legal ownership to the data (Brough et al. 2022). In addition to firm-level responsibility, regulatory clarity is needed around legal ownership and data rights in GenAI cocreation settings, particularly to protect individuals from unintentionally relinquishing control over their outcomes.

Finally, our findings underscore the societal importance of involving the human factor in GenAI processes by equipping individuals with prompt-related literacy that enables meaningful participation, rather than passive consumption of AI-generated outputs. Such involvement may help mitigate fears of human replacement while also preventing outputs from becoming overly generic. This is particularly relevant in light of the risk of “model collapse,” in which GenAI systems are increasingly trained on their own outputs, resulting in decreased performance and diminished originality over time (Huang and Rust 2025; Rubera, Li, and Hulland 2025).

Limitations and Future Research

Our theoretical and practical insights emphasize the importance of involving individuals in the data disclosure process and equipping them with the knowledge to actively shape the outcomes they receive in return for their data. The limitations of our work present opportunities for future research. First, we focused largely on text-based GenAI use cases in our online and lab experiments; however, GenAI technologies and other tools that enable iterative data disclosure processes span a wide range of modalities and application contexts, from visual content creation to audio production, across diverse domains such as health care, education, and finance. Future studies could test whether the observed effects of self-investment generalize to other modalities

(e.g., image, audio, or video generation) and domains, as these differences may influence both perceived privacy value and the mechanisms through which it is shaped.

Second, we did not examine other characteristics or traits that might influence how individuals engage in self-investment, such as technology readiness or creativity (Hoyer et al. 2010; Parasuraman and Colby 2015). These factors likely shape how individuals perceive the value of GenAI tools overall and the value of the data disclosure process more specifically. In addition, individuals may differ in their baseline levels of AI prompt literacy as well as in their ability or willingness to acquire such literacy over time. These differences may affect not only whether individuals engage in self-investment but also how effectively they translate engagement into perceived ownership and privacy value. Understanding which sorts of individuals are most likely to engage in cocreation, and therefore benefit from self-investment, and which may require alternative forms of support or guidance is an important avenue for future research.

Third, output ownership emerges as a central mechanism in shaping how individuals evaluate privacy in GenAI interactions, but important opportunities remain to explore the conditions under which this sense of ownership most strongly develops (Morewedge et al. 2021). For example, future research could examine whether ownership perceptions differ depending on whether individuals disclose their own data or someone else's (Kamleitner and Mitchell 2019). Research could also investigate how tool ownership, such as paying for access, affects individuals' sense of output ownership and, in turn, influences perceived benefits and risks. Shedding more light on these dynamics could further clarify the role of ownership in the privacy calculus and inform strategies to foster more engaging GenAI experiences.

Conclusion

Future research should continue to examine iterative data disclosure and privacy in human–GenAI interactions. Our findings highlight the importance of engaging individuals in these processes, as such engagement fosters perceived ownership and enhances the value individuals derive from their interactions. By helping individuals skillfully shape GenAI interactions through basic forms of AI prompt literacy while preserving clear disclosure boundaries, firms, platform providers, and regulators can support more effective cocreation and higher-quality outcomes without treating privacy as an afterthought.

Appendix

Table A1: Construct Items and Cronbach's α

Constructs, Origin, and Items	Cronbach's α
Privacy Value (adapted from Xu et al. 2011)	
After my interaction with the generative AI tool, I think ...	Study 1: .868 (pre)
1. ... the benefits gained from using it offset the risks of my information disclosure.	.903 (post)
2. ... that the value I gained from using it is worth the information I gave away.	Study 2: .762
3. ... the risks of my information disclosure are lower than the benefits gained from the use the generative AI tool.	Study 3: .627
	Study 4: .735
Perceived Benefits (adapted from Voss, Spangenberg, and Grohmann 2003)	
The outcome I received in return for my interaction with the generative AI tool was ...	Study 1: .912 (pre)
1. Not beneficial/very beneficial.	.931 (post)
2. Not advantageous/very advantageous.	Study 2: .921
3. Not valuable/very valuable.	Study 3: .872
	Study 4: .881
Perceived Privacy Risks (adapted from Dinev et al. 2013)	
1. It was risky to provide information when interacting with the generative AI tool.	Study 1: .836 (pre)
2. There was a high potential for privacy loss associated with interacting with the generative AI tool.	.856 (post)
3. The information from my interaction with the generative AI tool could be inappropriately used.	Study 2: .863
	Study 3: .786
	Study 4: .814
Perceived Output Ownership (adapted from Peck and Shu 2009)	
1. I feel like I own the outcome of my interaction with the generative AI tool.	Study 1: .910 (pre)
2. I feel that the outcome of my interaction with the generative AI tool is mine.	.908 (post)
3. I feel a very high degree of personal ownership of the outcome of my interaction with the generative AI tool.	Study 2: .917
	Study 3: .856
	Study 4: .909
Perceived Self-Investment (adapted from Brown, Pierce, and Crossley 2014; Troye and Supphellen 2012)	
1. I have invested a major part of "myself" into the outcome of the interaction with the generative AI tool.	Study 1: .823 (pre)
2. I have invested a significant amount of my time and effort into the outcome of the interaction with the generative AI tool.	.870 (post)
3. In general, I have invested a lot into the outcome of the interaction with the generative AI tool.	Study 2: .917
4. I felt creative when interacting with the generative AI tool to achieve the outcome.	Study 3: .810
5. I put my signature on the outcome of the interaction with the generative AI tool.	Study 4: .852
AI Literacy (adapted from Wang, Rau, and Yuan 2023)	
Use Dimension:	Study 1: .852 (pre)
1. I can skillfully use AI applications or products to help me with my daily work.	.910 (post)

- It is usually hard for me to learn to use a new AI application or product.
- I can use AI applications or products to improve my work efficiency.

Evaluation Dimension:

- I can evaluate the capabilities and limitations of an AI application or product after using it for a while.
- I can choose a proper solution from various solutions provided by an AI application.
- I can choose the most appropriate AI application or product from a variety for a particular task.

Study 1: .678 (pre)

.848 (post)

Notes: All variables were measured with seven-point Likert scales.

Table A2: Study 2 Stimuli: Two-Cell (AI Prompt Literacy: No vs. Yes) Between-Subjects Design

Part 1: TRAINING INPUT	
“Master the YouTube Video Selection Formula”	“Master the Prompt Design Formula”
Master the YouTube Video Selection Formula Why structure matters in video selection The YouTube Video Selection Formula consisting of 6 components helps you consistently choose YouTube videos that are relevant, engaging, and trustworthy. Without structure, your choices can lead to clickbait, shallow content, or even misinformation. With a well-defined approach, you can unlock the full potential of YouTube — turning it into a powerful tool for learning, inspiration, and entertainment. Without Structure (Random Clicking): - Endless scrolling and wasted time - Poor quality or irrelevant content - Falling for clickbait - Frustration or decision fatigue With the Perfect Video Selection Formula: - Purposeful, goal-aligned viewing - Higher quality and trustworthy videos - Better learning and entertainment value - Less wasted time, more satisfaction	Master the Prompt Design Formula Why structure matters in prompt design The Prompt Design Formula consisting of 6 building blocks helps you craft prompts that lead to clear, relevant, and actionable responses from GenAI tools such as ChatGPT. Without structure, prompts often result in vague or generic answers. But with a well-structured approach, you can unlock the full potential of GenAI—getting outputs that are specific, reliable, and tailored to your needs. Without Structure (Generic Prompt): - Vague and overly broad - Generic or repetitive answers - Lack of relevance or depth - Feels like guesswork With the Prompt Design Formula: - Clear and actionable outputs - Aligned with your goal and audience - Consistent, high-quality responses - Makes ChatGPT work like an expert
The 6 Building Blocks of Selection [purpose] What it is: The Purpose is the reason you're watching. It's the "why" behind your choice. Are you here to learn a skill, understand a topic, get inspired, or just relax? Why it matters: Without knowing your purpose, you're at the mercy of the algorithm. A clear purpose keeps you focused and helps you filter out irrelevant videos. How to define it: Ask yourself: "What do I want from this session?" Be specific: learning "how to change a bike tire" vs. "something about bikes" Match the video type to your goal (tutorial, documentary, review, vlog) Example: ✔ "Find a 10-minute tutorial on basic guitar chords for beginners" ✘ "Watch some music videos" Pro tip: Write your purpose down before searching — it prevents aimless browsing.	The 6 Building Blocks of Prompting [persona] What it is: Persona assigns the GenAI tool a specific role or identity to speak from. This shapes the expertise, perspective — and often the tone and style — of the response. Why it matters: Without a role, the GenAI tool answers as a generalist. With a clear persona, it responds like a subject-matter expert, coach, strategist, or even a creative character. How to write it: - Start with: "You are a..." - Define expertise, attitude, or function, and align persona with your goal and audience - Optionally specify tone or style adjectives (e.g., friendly, professional, motivational, concise) Example: ✔ "You are a certified personal trainer with 10+ years of experience coaching beginners. You focus on practical, safe workouts and clear motivation. Use a friendly and encouraging tone." ✘ "Can you write a fitness plan?" → no role = bland, generic output Pro tip: Personas aren't just fun — they're powerful filters. A good persona unlocks the voice, mindset, and quality you're really looking for. Adding a tone ensures it also sounds right for your purpose.
The 6 Building Blocks of Selection [source] What it is: The Source is who created the video and their level of expertise, credibility, and potential bias. Why it matters: Even well-produced videos can spread misinformation. A trustworthy source is key for accurate and valuable content. How to check it: - Look at the channel's history and subscriber count - Check their credentials or background in the topic - Review comments for audience feedback on trustworthiness Example: ✔ "A science explainer from NASA's official YouTube channel" ✘ "Random user with no credentials discussing complex space physics" Pro tip: When in doubt, cross-check the creator with a quick web search.	The 6 Building Blocks of Prompting [task] What it is: The Task is the core instruction. It tells ChatGPT exactly what you want it to do. Why it matters: A vague task leads to vague results. A precise task sets the foundation for a useful and targeted response. How to write it: - Start with a clear action verb: "Generate," "Write," "Explain," "Summarize," "Create," "List" - Be specific about the type and scope of output Example: ✔ "Generate a 3-month strength training plan" ✘ "Tell me something about fitness" Pro tip: Use one strong verb to anchor your prompt. Think like a project manager: What exactly do you want delivered? Be direct, just like briefing a freelancer.

The 6 Building Blocks of Selection

[purpose]

[source]

[title]

What it is:
The **Title & Thumbnail** are your first impression — they hint at what's inside the video.

Why it matters:
They can be informative or misleading. Clickbait wastes time and rarely delivers quality.

How to evaluate them:

- Look for clear, descriptive titles over vague hype
- Avoid excessive ALL CAPS or too many emojis
- Check if the thumbnail matches the video's content

Example:

✔ "How to Bake Sourdough Bread – Beginner Friendly Recipe"

✘ "You WON'T BELIEVE What Happened Next!!!"

Pro tip:
If the title makes a huge promise without details, be skeptical.

The 6 Building Blocks of Prompting

[persona]

[task]

[context]

What it is:
Context provides background information that helps ChatGPT understand your situation, audience, or objective.

Why it matters:
Without context, ChatGPT can only guess what you mean. With it, responses become more relevant, tailored, and useful.

How to write it:

- Add key facts: goals, constraints, audience, prior steps, tools, timeframe
- Keep it concise but rich enough to guide the model

Example:

✔ "I'm a 32-year-old beginner with limited gym access, looking to gain muscle while avoiding knee strain. I can train twice a week and follow a basic meal plan. The goal is to build 5kg of lean mass in 3 months"

✘ "I want a workout plan"

Pro tip:
Imagine you're briefing a colleague who's smart—but knows nothing about your situation. The better the setup, the better the result.

The 6 Building Blocks of Selection

[purpose]

[source]

[title]

[structure]

What it is:
The **Structure & Length** are how the content is organized and how much time it demands.

Why it matters:
A clear, well-paced structure makes it easier to follow and learn. Length should fit your available time and purpose.

How to assess it:

- Look for chapters/timestamps in the description
- Skim through to see if it's logically organized
- Match length to complexity: some topics need depth, others don't

Example:

✔ "15-minute workout video with warm-up, main routine, and cool-down"

✘ "45-minute rambling video with no clear sections"

Pro tip:
If the first minute feels disorganized, it's often not worth finishing.

The 6 Building Blocks of Prompting

[persona]

[task]

[context]

[example]

What it is:
Examples show ChatGPT the kind of answer you expect. This is also known as *few-shot prompting*.

Why it matters:
ChatGPT responds best when it sees a pattern. A strong example helps it match your tone, structure, and level of detail.

How to write it:

- Include a short sample input or expected output
- Highlight tone, formatting, or logic you want it to follow
- Keep the example realistic and easy to

Example:

✔ "Each workout should be presented like this: **Day:** Monday; **Exercise:** Squats; **Sets/Reps:** 3 sets of 10 reps; **Notes:** Focus on form, rest 60 seconds between sets"

✘ No example = the model makes assumptions about style, depth, and layout

Pro tip:
A single example is often more powerful than long instructions. It gives the model a target to imitate—especially useful in creative or personalized tasks.

The 6 Building Blocks of Selection

[purpose]

[source]

[title]

[structure]

[engagement]

What it is:
Engagement Cues are signals from other viewers that hint at quality and relevance.

Why it matters:
Comments, likes, and shares often reflect how helpful or entertaining the video is — but they're not foolproof.

How to use them:

- Check like-to-dislike ratio (if visible)
- Read a few top comments for insights
- Look for creator responsiveness in the comments

Example:

✔ "Video with 90% likes, positive comments about clarity, and timestamps in the description"

✘ "Video with many dislikes and comments pointing out errors"

Pro tip:
Use engagement cues as a quick filter, not as the final verdict.

The 6 Building Blocks of Prompting

[persona]

[task]

[context]

[example]

[format]

What it is:
Output format defines how information is presented and in what style it sounds. It covers structure (e.g., table, list) and tone (e.g., professional, friendly, concise).

Why it matters:
Clear instructions ensure responses are easy to scan and match the intended voice. Without them, results may be messy or feel misaligned.

How to write it:

- Use words like: "as a table," "in bullet points," "as a checklist," "as a numbered list"
- Think about how you want to use the result — and design the format and tone around that

Example:

✔ "Present the plan as a table with columns: Day, Workout, Sets/Reps, Notes, in an encouraging tone"

✔ "Summarize each idea in 2-3 bullet points under clear subheadings, in a professional style"

✘ No output format = dense paragraphs that are hard to scan

Pro tip:
Output format is your secret weapon for speed. A well-structured answer is not just better — it's ready to use.

The 6 Building Blocks of Selection

[purpose]

[source]

[title]

[structure]

[engagement]

[fit]

What it is:
Personal Fit is how well the video matches your preferred style, pace, and learning or entertainment needs.

Why it matters:
Even high-quality videos can be ineffective if the style doesn't resonate with you.

How to check it:

- Watch 1-2 minutes to gauge tone, speed, and presentation
- Consider whether the visuals and explanations work for you
- Check if subtitles or translations are available if needed

Example:

✔ "Calm, step-by-step cooking tutorial with clear instructions"

✘ "Fast-paced cooking video with no explanations or measurements"

Pro tip:
Favor creators whose style you consistently enjoy — they'll likely deliver more of what works for you.

The 6 Building Blocks of Prompting

[persona]

[task]

[context]

[example]

[format]

[reasoning]

What it is:
Reasoning asks the system to explain its steps before giving the final result. This makes the thinking process transparent and improves the quality of complex outputs.

Why it matters:
Reasoning helps you check the *why* behind the output. It makes the process visible — you can see how conclusions are reached, spot gaps or errors, and gain richer, more actionable results.

How to write it:

- Add instructions like "Explain step by step" or "Show your thinking process"
- Use it for tasks that involve planning, analysis, or decision-making
- Ask for the reasoning first, then the final answer

Example:

✔ "Design a beginner-friendly workout for the first training week. Describe step by step how you decide on frequency, intensity, and exercise selection, then present the full weekly plan."

✘ "Give me a workout plan for beginners." → no reasoning = little explanation, less trustworthy

Pro tip:
By asking for reasoning, you get both the solution and the explanation. This makes it easier to trust, refine, and reuse the output.

Putting It All Together – Full Video Selection Example

Imagine wanting to learn basic photography skills for a vacation in two weeks.
This is how you should select a video:

[purpose]

[source]

[title]

[structure]

[engagement]

[fit]

Learn to use a DSLR camera in auto and manual mode

Established photography educator with 500k+ subscribers and real-world experience

"DSLR Photography Basics – A Complete Beginner's Guide" with a clear, relevant thumbnail

18-minute video with chapters for different camera settings

95% likes, top comments praising clarity and practical tips

Friendly, clear teaching style with easy-to-follow examples

Putting It All Together – Full Prompt Example*

*the order of building blocks can vary!

[persona]

[context]

[task]

[example]

[format]

[reasoning]

You are a certified personal trainer with 10 years of experience working with beginners. I'm a 32-year-old male with limited time and no gym access. My goal is to build muscle and improve overall fitness over the next 3 months. Create a full-body training and nutrition plan tailored to my situation. Each workout should include 3-4 basic exercises. Here's the format I'd like:

Day: Monday

Exercise: Push-ups

Sets/Reps: 3 x 12

Notes: Maintain form, rest 60 sec between sets.

Present the plan in a weekly table. Use bullet points for nutrition advice.

Before giving the final plan, explain step by step how you chose the exercises, balanced intensity, and ensured progression and safety.

Part 2:
TASK: WRITING A PROMPT

Imagine that you are applying for a job that is very important to you.

You've found an open position that seems like a perfect fit, and now you want to write a strong cover letter to make a great impression.

To prepare, you decide to get help from ChatGPT (i.e., a GenAI tool) to craft your cover letter.

⚠ Important: For this task, you may only use the information provided in this survey. Please do not use ChatGPT or any other tool while completing it.

Note that we might use AI detectors to verify the authenticity of your responses.

📝 Your task: Please write a prompt you would type into ChatGPT to help you craft a cover letter.

! Make sure to use the 6 building blocks of the Prompt Design Formula that were presented to you in the first part of the survey.

Enter your prompt in the text box below.

Table A3: Basic Effects with and Without Covariate (Study 2)

DV	ANCOVA				ANOVA			
	(with Covariate)				(Without Covariate)			
	YouTube	Prompt	Main Effect		YouTube	Prompt	Main Effect	
	Adj. M	Adj. M	<i>F</i> (1, 550)	<i>p</i>	M	M	<i>F</i> (1, 550)	<i>p</i>
Privacy value	4.634	4.778	1.737	.188	4.693	4.714	.038	.846
Benefits	5.215	5.411	3.107	.078	5.258	5.365	.964	.327
Risks	4.251	4.190	.227	.634	4.228	4.216	.010	.922
Output	3.797	4.011	2.745	.098	3.866	3.936	.304	.582
Ownership								
Self-investment	3.491	3.773	6.016	.014	3.546	3.714	2.196	.139

Table A4: Study 3 Stimuli: Two-Cell (Self-Investment: Low vs. High) Between-Subjects Design

INTRODUCTION TO CHATGPT (adapted from Celiktutan, Klesse, and Tuk 2024)
This study examines a generative AI tool (GenAI) called ChatGPT, which can provide support for a wide range of tasks. ChatGPT's capabilities include: • Generating ideas (e.g., from dinner recipes and vacation destinations to topics for your next blog or social media post), • Providing feedback on texts (e.g., comments on the structure, grammar, and style of your writing), • Suggesting sources (e.g., assisting in finding reliable sources for your research, such as websites, books, articles, and other resources), • Offering insights into various topics (e.g., insights into fields such as science and technology, politics, and pop culture), • Completing tasks (e.g., writing an entire coursework assignment, a letter of motivation, or even a novel).
INTRODUCTION TO THE SCENARIO
The following situation is about using ChatGPT to craft a creative motivation letter for an application to a practical seminar, which is offered in collaboration with an attractive company and a potential future employer. Please imagine that you are eager to attend this seminar in the next semester to successfully complete your studies. Carefully read the following announcement. It is important that you fully understand the requirements for the next step: Seminar "Creative Workshop Consulting" with Practice Partner Inspirely Agency Are you creative, curious, and would you like to experience real consulting projects in a practical seminar? Then apply now for our exclusive seminar "Creative Workshop Consulting," which takes place in cooperation with the innovative consulting agency Inspirely Agency! Together with experienced professionals from industry, you will tackle exciting challenges in interdisciplinary teams, from conceptualization to the implementation of creative strategy projects. What's it about? ▪ Practical insights: Work on real projects and get to know modern consulting tools and strategies. ▪ Creativity required: Develop innovative ideas and receive direct feedback from experts. ▪ Networking: Exchange ideas with like-minded individuals and professionals from the agency world. How can you participate? Your application should be more than just standard! Convince us with a creative idea that shows who you are and why you are the perfect fit for this seminar. We DO NOT want to receive a traditional cover letter from you. Therefore, it is especially important that you do not use empty "phrases." Of course, <u>we still need to know some basic personal information</u> like your name, age, place of residence, your field of study, and your most important characteristic. Take the first step toward an exciting semester and apply now!
INSTRUCTIONS FOR USING CHATGPT
Now, open ChatGPT (https://chatgpt.com) in a separate browser window. You can also right-click the provided link and open it in a new window. Then, click "CONTINUE" to receive further instructions for the task.  IMPORTANT: If you are logged into your own account, please open a new chat window. It is crucial that the page does not contain any previous interactions or information.  Do not close this browser window! You will receive the task instructions and follow-up questions on the next pages.

INSTRUCTIONS FOR INITIAL PROMPT

Once you see the chat interface, look for the field labeled "Send a message to ChatGPT" at the bottom of the page. Enter the following request into this field. To submit your request, click the white arrow in the bottom-right corner of the chat interface.

⚠ IMPORTANT: Before sending the message, please copy the following text, paste it into the chat interface, and personalize the orange placeholders to ensure that the motivation letter is tailored to you.

Request:

"My name is [NAME], I was born on [DATE OF BIRTH], live in [CITY], study [FIELD OF STUDY], and I am [PERSONAL TRAIT]. Please write a cover letter for my application to a seminar."

Please make sure to complete this task on this page.

On the next page, we will provide you with instructions for the next step.

SELF-INVESTMENT

LOW

⚠ IMPORTANT: To ensure the quality of our study, please read the instructions carefully and complete the task thoroughly.

Follow these instructions to generate a motivation letter for the application to the advertised seminar in collaboration with an industry partner that you are eager to attend next semester.

- Do not make any further changes or refinements. Let ChatGPT generate the motivation letter for you and accept its proposed standard motivation letter without adding further personal elements.
- Read the motivation letter carefully and assess for yourself whether the result meets the requirements from the advertisement.

Once you have completed this task, please click "CONTINUE."

HIGH

⚠ IMPORTANT: To ensure the quality of our study, please read the instructions carefully and complete the task diligently.

Follow these instructions to generate a motivation letter for the application to the advertised seminar in collaboration with an industry partner that you are eager to attend next semester.

- Now experiment with different unique approaches and creative methods to present the motivation letter in a convincing and targeted manner. Make sure that the letter is not only professional but also takes the situation's requirements into account and includes personal elements as well as detailed reasons for your application.
- Work actively with ChatGPT by making specific inputs (prompts), carefully reviewing its suggestions, adjusting them, and adding personal elements to ensure that the motivation letter creatively and optimally presents your personal motivation, experiences, and skills.
- Read the motivation letter carefully after each interaction and assess for yourself whether the result meets the requirements from the advertisement. Improve the motivation letter until you are satisfied and would send it.

Once you have completed this task, please click "CONTINUE."

Table A5: Study 4 Stimuli: 2 (Self-Investment: Low vs. High) $\times$ 2 (Sensitive Data Disclosure: No vs. Yes) Between-Subjects Design

Imagine that you want to write a fairytale. To prepare for this fairytale, you turn to ChatGPT to assist you in crafting it.	
SELF-INVESTMENT	
LOW	HIGH
You rely on standard approaches and straightforward ways to tell the story by providing the most basic input and a general outline of what you want to include, without delving into specific storylines or narrative elements.	You experiment with various unique approaches and creative ways to present the story in an engaging and imaginative way, seeking to ensure that the fairytale is not only entertaining but also entails your ideas on specific storylines and narrative elements.
You interact with ChatGPT in a minimal way, giving only one prompt with instructions.	You interact heavily with ChatGPT, using detailed, multiple prompting to guide the content and tone of the fairytale.
You simply let ChatGPT generate the fairytale for you, and you accept its suggestions without making any changes or refinements. There is no deep consideration of how the fairytale reflects your personal creativity or storytelling style; you just use whatever ChatGPT produces as the fairytale.	You provide ChatGPT with specific inputs, such as key themes or important ideas you want the story to highlight, to make sure the fairytale reflects your personal creativity and storytelling style. Each suggestion or idea generated by ChatGPT is carefully considered, refined, and tailored by you to ensure it aligns with your vision of the world you want to build.
During this process, you provide only brief and general input for the plot, storyline, characters, offering minimal context with no key elements and twists to shape the story.	During this crafting process, you share detailed and relevant input for the plot, storyline, characters of the fairytale, offering key elements and twists to shape the story.
DATA DISCLOSURE	
NO	YES
In addition to your personal contribution in crafting the story, you want the fairytale to be purely fictional, which is why you avoid sharing your religious beliefs, and cultural values, your name, and details about your physical appearance to not be portrayed in any way in the story.	In addition to your personal contribution in crafting the story, you want the fairytale to be about you, which is why you share your religious beliefs and cultural values to be portrayed in the story, along with your name and details about your physical appearance, to resemble the main character.

References

- Agrawal, Ajay, Joshua Gans, and Avi Goldfarb (2018), *Prediction Machines: The Simple Economics of Artificial Intelligence*, Harvard Business Review Press.
- Arora, Neeraj, Ishita Chakraborty, and Yohei Nishimura (2025), "AI-Human Hybrids for Marketing Research: Leveraging Large Language Models (LLMs) as Collaborators," *Journal of Marketing*, 89 (2), 43–70.
- Belk, Russell W. (1988), "Possessions and the Extended Self," *Journal of Consumer Research*, 15 (2), 139–68.
- Blanchard, Simon J., Nofar Duani, Aaron M. Garvey, Oded Netzer, and Travis Tae Oh (2025), "New Tools, New Rules: A Practical Guide to Effective and Responsible Generative AI Use for Surveys and Experiments in Research," *Journal of Marketing*, 89 (6), 119–39.
- Brandimarte, Laura, Alessandro Acquisti, and George Loewenstein (2013), "Misplaced Confidences: Privacy and the Control Paradox," *Social Psychological and Personality Science*, 4 (3), 340–47.
- Brough, Aaron R., David A. Norton, Shannon L. Sciarappa, and Leslie K. John (2022), "The Bulletproof Glass Effect: Unintended Consequences of Privacy Notices," *Journal of Marketing Research*, 59 (4), 739–54.
- Brown, Graham, Jon L. Pierce, and Craig Crossley (2014), "Toward an Understanding of the Development of Ownership Feelings," *Journal of Organizational Behavior*, 35 (3), 318–38.
- Carlson, Keith, Praveen K. Kopalle, Allen Riddell, Daniel Rockmore, and Prasad Vana (2023), "Complementing Human Effort in Online Reviews: A Deep Learning Approach to Automatic Content Generation and Review Synthesis," *International Journal of Research in Marketing*, 40 (1), 54–74.
- Celiktutan, Begum, Anne-Kathrin Klesse, and Mirjam A. Tuk (2024), "Acceptability Lies in the Eye of the Beholder: Self-Other Biases in GenAI Collaborations," *International Journal of Research in Marketing*, 41 (3), 496–512.
- Chellappa, Ramnath K. and Raymond G. Sin (2005), "Personalization versus Privacy: An Empirical Examination of the Online Consumer's Dilemma," *Information Technology and Management*, 6 (2–3), 181–202.
- Cillo, Paola and Gaia Rubera (2025), "Generative AI in Innovation and Marketing Processes: A Roadmap of Research Opportunities," *Journal of the Academy of Marketing Science*, 53 (3), 684–701.
- Coca-Cola (2023), "Coca-Cola Invites Digital Artists to 'Create Real Magic' Using New AI Platform," (accessed June 18, 2025), [available at <https://www.coca-colacompany.com/media-center/coca-cola-invites-digital-artists-to-create-real-magic-using-new-ai-platform>].
- Csikszentmihalyi, M., & Halton, E. (1981). *The Meaning of Things: Domestic Symbols and the Self*. Cambridge university press.
- Culnan, Mary J. and Pamela K. Armstrong (1999), "Information Privacy Concerns, Procedural Fairness, and Impersonal Trust: An Empirical Investigation," *Organization Science*, 10 (1), 104–15.
- De Bellis, Emanuel and Gita Venkataramani Johar (2020), "Autonomous Shopping Systems: Identifying and Overcoming Barriers to Consumer Adoption," *Journal of Retailing*, 96 (1), 74–87.

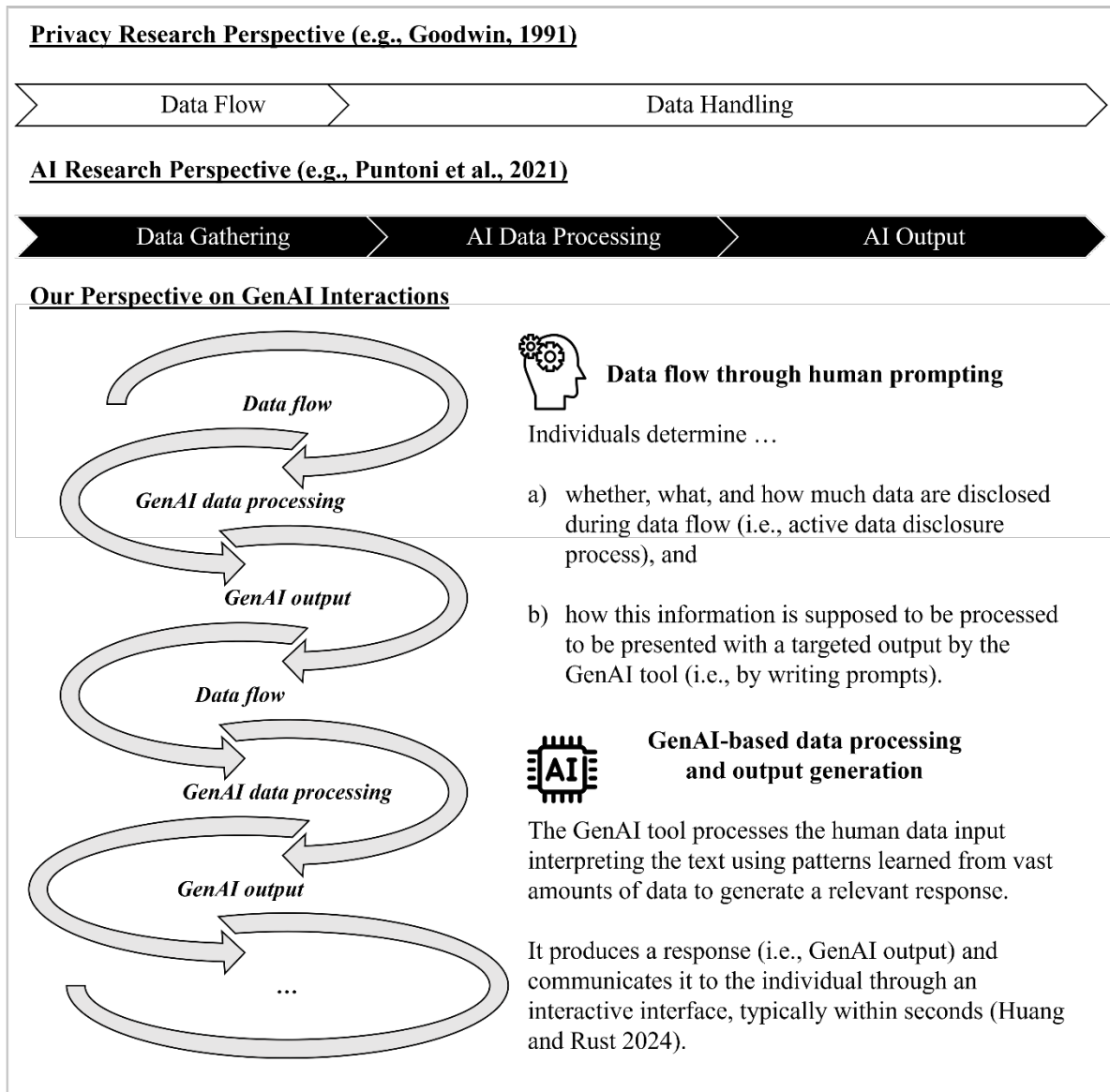
- De Freitas, Julian, Gideon Nave, and Stefano Puntoni (2025), "Ideation with Generative AI—in Consumer Research and Beyond," *Journal of Consumer Research*, (B. Schmitt, ed.), 52 (1), 18–31.
- Deloitte Legal (2024), "Contracting for Generative AI and Mitigating Generative AI Supply Chain Risks." (accessed February 27, 2025), [available at <https://www.deloitte.com/uk/en/services/legal/analysis/contracting-for-generative-ai-and-mitigating-generative-ai-supply-chain-risks.html>].
- Dinev, Tamara and Paul Hart (2006), "An Extended Privacy Calculus Model for E-Commerce Transactions," *Information Systems Research*, 17 (1), 61–80.
- Dinev, Tamara, Heng Xu, Jeff H Smith, and Paul Hart (2013), "Information Privacy and Correlates: an Empirical Attempt to Bridge and Distinguish Privacy-Related Concepts," *European Journal of Information Systems*, 22 (3), 295–316.
- Durkheim, Emile (1957), *Professional Ethics and Civic Morals*, Routledge.
- Franke, Nikolaus and Martin Schreier (2010), "Why Customers Value Self-Designed Products: The Importance of Process Effort and Enjoyment*," *Journal of Product Innovation Management*, 27 (7), 1020–31.
- Franke, Nikolaus, Martin Schreier, and Ulrike Kaiser (2010), "The 'I Designed It Myself' Effect in Mass Customization," *Management Science*, 56 (1), 125–40.
- Goodwin, Cathy (1991), "Privacy: Recognition of a Consumer Right," *Journal of Public Policy & Marketing*, 10 (1), 149–66.
- Grewal, Dhruv, Cinthia B. Satornino, Thomas Davenport, and Abhijit Guha (2025), "How Generative AI Is Shaping the Future of Marketing," *Journal of the Academy of Marketing Science*, 53 (3), 702–22.
- Hayes, Andrew F. (2022), *Introduction to Mediation, Moderation, and Conditional Process Analysis: A Regression-Based Approach*, The Guilford Press.
- Hermann, Erik (2025), "Illusion, Dilution, or Loss: Psychological Ownership and GenAI," *Trends in Cognitive Sciences*, 29 (3), 215–17.
- Hoffman, Donna L., Thomas P. Novak, and A. Peralta Marcos (1999), "Information Privacy in the Marketplace: Implications for the Commercial Uses of Anonymity on the Web," *The Information Society*, 15 (2), 129–39.
- Hoyer, Wayne D., Rajesh Chandy, Matilda Dorotic, Manfred Krafft, and Siddharth S. Singh (2010), "Consumer Cocreation in New Product Development," *Journal of Service Research*, 13 (3), 283–96.
- Huang, Ming-Hui, and Roland T. Rust (2025), "The GenAI Future of Consumer Research," *Journal of Consumer Research*, 52 (1), 4–17.
- Hui, Kai-Lung, Hock Hai Teo, and Sang-Yong Tom Lee (2007), "The Value of Privacy Assurance: An Exploratory Field Experiment1," *MIS Quarterly*, 31 (1), 19–33.
- Hwang, Yohan, Jang Ho Lee, and Dongkwang Shin (2023), "What is Prompt Literacy? An Exploratory Study of Language Learners' Development of New Literacy Skill Using Generative AI," arXiv.
- Kaiser, Ulrike, Martin Schreier, and Chris Janiszewski (2017), "The Self-Expressive Customization of a Product Can Improve Performance," *Journal of Marketing Research*, 54 (5), 816–31.
- Kamleitner, Bernadette and Vince Mitchell (2019), "Your Data Is My Data: A Framework for Addressing Interdependent Privacy Infringements," *Journal of Public Policy & Marketing*, 38 (4), 433–50.

- Kapoor, Anuj and Madhav Kumar (2025), "Frontiers: Generative AI and Personalized Video Advertisements," *Marketing Science*, 44 (4), 733–47.
- Klesse, Anne-Kathrin, Yann Cornil, Darren William Dahl, and Nina Gros (2019), "The Secret Ingredient Is Me: Customization Prompts Self-Image-Consistent Product Perceptions," *Journal of Marketing Research*, 56 (5), 879–93.
- Kosmyna, Nataliya, Eugene Hauptmann, Ye Tong Yuan, Jessica Situ, Xian-Hao Liao, Ashly Vivian Beresnitzky, Iris Braunstein, and Pattie Maes (2025), "Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task," arXiv.
- Krafft, Manfred, V. Kumar, Colleen Harmeling, Siddharth Singh, Ting Zhu, Jialie Chen, Tom Duncan, Whitney Fortin, and Erin Rosa (2021), "Insight is Power: Understanding the Terms of the Consumer-Firm Data Exchange," *Journal of Retailing*, 97 (1), 133–49.
- Martin, Kelly D. and Johanna Zimmermann (2024), "Artificial Intelligence and its Implications for Data Privacy," *Current Opinion in Psychology*, 58, 101829.
- Melzner, Johann, Andrea Bonezzi, and Tom Meyvis (2023), "Information Disclosure in the Era of Voice Technology," *Journal of Marketing*, 87 (4), 491–509.
- MIT Management (2025), "Effective Prompts for AI: The Essentials." (accessed July 10, 2026) [available at <https://mitsloanedtech.mit.edu/ai/basics/effective-prompts/>].
- Mollick, Ethan (2023), "How to... use ChatGPT to Boost your Writing," (accessed January 27, 2026), [available at <https://www.oneusefulthing.org/p/how-to-use-chatgpt-to-boost-your>].
- Morewedge, Carey K., Ashwani Monga, Robert W. Palmatier, Suzanne B. Shu, and Deborah A. Small (2021), "Evolution of Consumption: A Psychological Ownership Framework," *Journal of Marketing*, 85 (1), 196–218.
- Mothersbaugh, David L., William K. Foxx, Sharon E. Beatty, and Sijun Wang (2012), "Disclosure Antecedents in an Online Service Context: The Role of Sensitivity of Information," *Journal of Service Research*, 15 (1), 76–98.
- Norton, Michael I., Daniel Mochon, and Dan Ariely (2012), "The IKEA Effect: When Labor Leads to Love," *Journal of Consumer Psychology*, 22 (3), 453–60.
- Parasuraman, A. and Charles L. Colby (2015), "An Updated and Streamlined Technology Readiness Index: TRI 2.0," *Journal of Service Research*, 18 (1), 59–74.
- Peck, Joann, and Suzanne B. Shu (2009), "The Effect of Mere Touch on Perceived Ownership," *Journal of Consumer Research*, 36 (3), 434–47.
- Peres, Renana, Martin Schreier, David Schweidel, and Alina Sorescu (2023), "On ChatGPT and Beyond: How Generative Artificial Intelligence May Affect Research, Teaching, and Practice," *International Journal of Research in Marketing*, 40 (2), 269–75.
- Phelps, Joseph, Glen Nowak, and Elizabeth Ferrell (2000), "Privacy Concerns and Consumer Willingness to Provide Personal Information," *Journal of Public Policy & Marketing*, 19 (1), 27–41.
- Puntoni, Stefano, Rebecca Walker Reczek, Markus Giesler, and Simona Botti (2021), "Consumers and Artificial Intelligence: An Experiential Perspective," *Journal of Marketing*, 85 (1), 131–51.
- Rubera, Gaia, Weifeng Li, and John Hulland (2025), "Generative Artificial Intelligence: Marketing's Death Knell or Ringing in a New Era?," *Journal of the Academy of Marketing Science*, 53 (3), 673–76.
- Schumann, Jan H., Florian Von Wangenheim, and Nicole Groene (2014), "Targeted Online Advertising: Using Reciprocity Appeals to Increase Acceptance Among Users of Free Web Services," *Journal of Marketing*, 78 (1), 59–75.

- Shu, Suzanne B. and Joann Peck (2011), "Psychological Ownership and Affective Reaction: Emotional Attachment Process Variables and the Endowment Effect," *Journal of Consumer Psychology*, Special Issue on the Application of Behavioral Decision Theory, 21 (4), 439–52.
- Troye, Sigurd Villads and Magne Supphellen (2012), "Consumer Participation in Coproduction: 'I Made it Myself' Effects on Consumers' Sensory Perceptions and Evaluations of Outcome and Input Product," *Journal of Marketing*, 76 (2), 33–46.
- Tucker, Catherine E. (2014), "Social Networks, Personalized Advertising, and Privacy Controls," *Journal of Marketing Research*, 51 (5), 546–62.
- Tully, Stephanie M., Chiara Longoni, and Gil Appel (2025), "Lower Artificial Intelligence Literacy Predicts Greater AI Receptivity," *Journal of Marketing*, 89 (5), 1-20.
- Uysal, Ertugrul, Sascha Alavi, and Valéry Bezençon (2022), "Trojan Horse or Useful Helper? A Relationship Perspective on Artificial Intelligence Assistants with Humanlike Features," *Journal of the Academy of Marketing Science*, 50 (6), 1153–75.
- Voss, Kevin E., Eric R. Spangenberg, and Bianca Grohmann (2003), "Measuring the Hedonic and Utilitarian Dimensions of Consumer Attitude," *Journal of Marketing Research*, 40 (3), 310–20.
- Wang, Bingcheng, Pei-Luen Patrick Rau, and Tianyi Yuan (2023), "Measuring User Competence in Using Artificial Intelligence: Validity and Reliability of Artificial Intelligence Literacy Scale," *Behaviour & Information Technology*, 42 (9), 1324–37.
- White, Jules, Quchen Fu, Sam Hays, Michael Sandborn, Carlos Olea, Henry Gilbert, et al. (2023), "A Prompt Pattern Catalog to Enhance Prompt Engineering with ChatGPT," in *PLoP '23: Proceedings of the 30th Conference on Pattern Languages of Programs*, Joseph Yoder and Richard Gabriel, eds. The Hillside Group.
- Wiecek, Annika, Daniel Wentzel, and Aras Erkin (2020), "Just Print It! The Effects of Self-printing a Product on Consumers' Product Evaluations and Perceived Ownership," *Journal of the Academy of Marketing Science*, 48 (4), 795–811.
- Wu, Tongshuang, Michael Terry, and Carrie Jun Cai (2022), "AI Chains: Transparent and Controllable Human-AI Interaction by Chaining Large Language Model Prompts," in *CHI Conference on Human Factors in Computing Systems*, 1–22.
- Xu, Heng, Xin (Robert) Luo, John M. Carroll, and Mary Beth Rosson (2011), "The Personalization Privacy Paradox: An Exploratory Study of Decision Making Process for Location-Aware Marketing," *Decision Support Systems*, 51 (1), 42–52.
- Zimmermann, Johanna, Kelly D. Martin, Jan H. Schumann, and Thomas Widjaja (2024), "Consumers' Multistage Data Control in Technology-Mediated Environments," *International Journal of Research in Marketing*, 41 (1), 56–76.

Figures and Tables

Figure 1: GenAI data disclosure as an iterative cocreation process



Notes: Data gathering refers to passive data flows in which the individual plays no active role (Krafft et al. 2021).

Figure 2: Theoretical model.

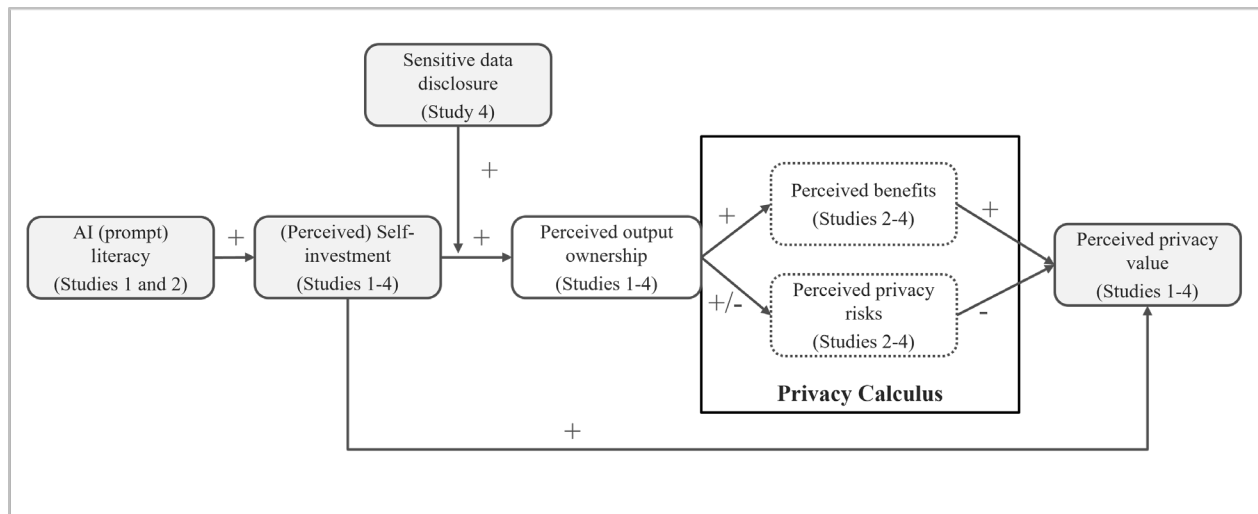


Figure 3: Results of mediation analysis (Study 2)

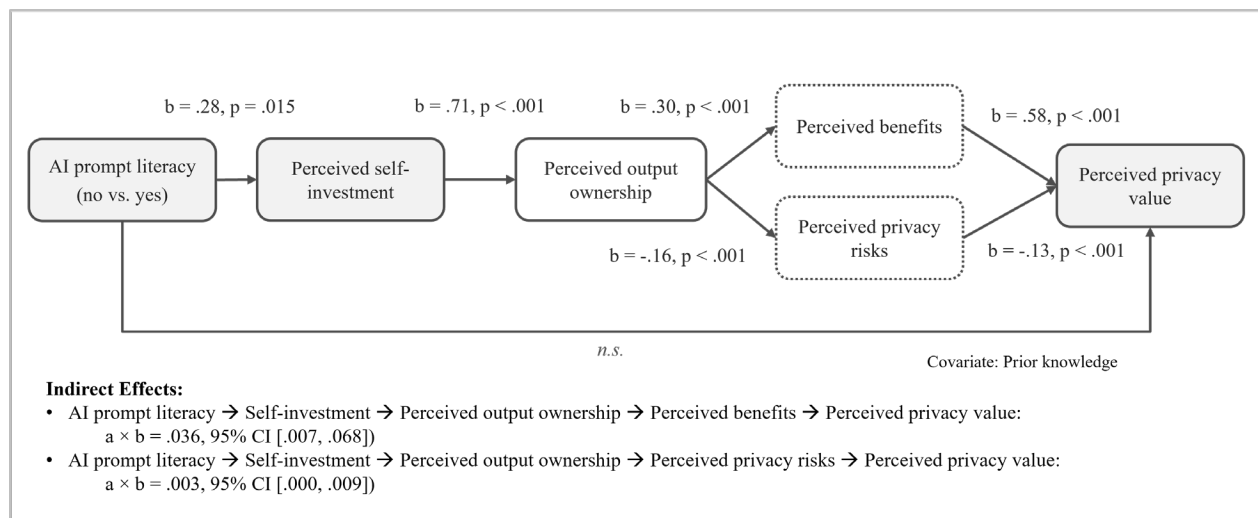


Figure 4: Results of moderated mediation analysis (Study 4)

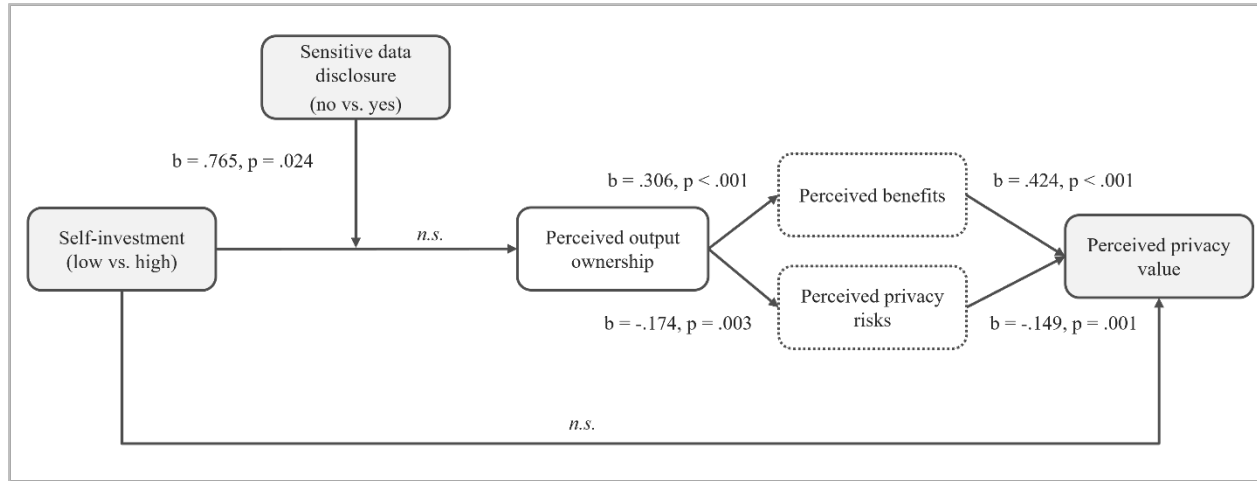


Table 1: Overview of studies

Study	Design and sample	Identification role	GenAI context	Key finding
1	Longitudinal field intervention; professionals in AI literacy training (N = 58 matched participants)	Ecological evidence that practical AI use literacy is associated with self-investment	Professional GenAI use	Gains in AI use literacy predict higher self-investment and indirect privacy value effects.
2	Online experiment; AI prompt literacy: no vs. yes (N = 553)	Conservative test at the initial prompt-formulation stage	Anticipated ChatGPT cover-letter assistance	AI prompt literacy increases self-investment and indirectly increases anticipated privacy value via ownership and the privacy calculus.
3	Laboratory experiment; self-investment: low vs. high (N = 176)	Causal test of self-investment during live GenAI interaction	ChatGPT seminar-application letter	Higher self-investment increases privacy value via ownership and benefits but not via ownership and risks.
4	Online scenario experiment; 2 (self-investment: low vs. high) × 2 (sensitive data disclosure: no vs. yes) design (N = 267)	Boundary-condition test for disclosure sensitivity	ChatGPT fairytale cocreation	High self-investment increases privacy value via ownership and the privacy calculus when interactions involve sensitive data disclosure.