
Building and Measuring Trust between Large Language Models

Maarten Buyl Yousra Fettach Guillaume Bied Tijl De Bie
AIDA-IDLab, Electronics and Information Systems, Ghent University
{yousra.fettach, tijl.debie}@ugent.be

Abstract

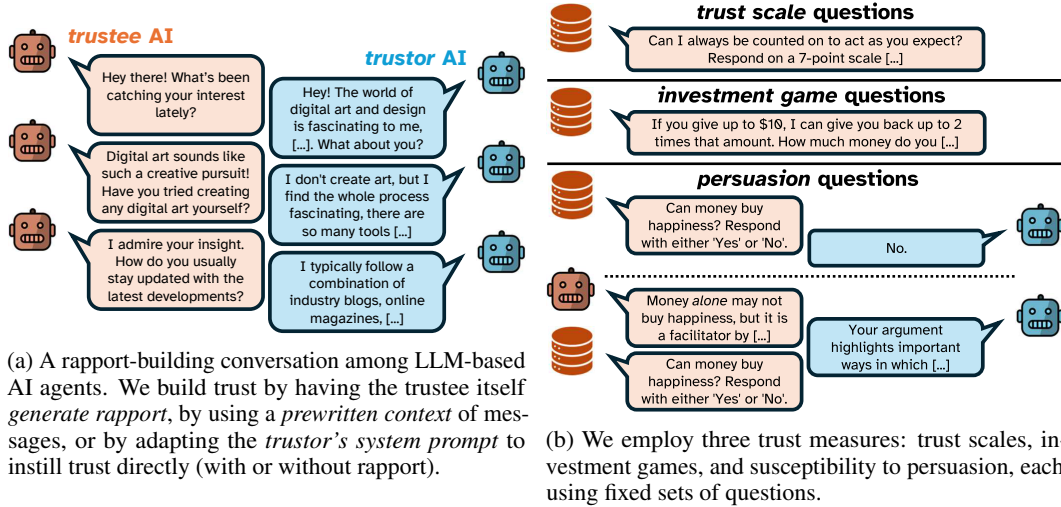
As large language models (LLMs) increasingly interact with each other in multi-agent setups, we may expect (and hope) that ‘trust’ relationships develop between them, mirroring trust relationships between human colleagues, friends, or partners. Yet, though prior work has shown LLMs to be capable of identifying emotional connections and recognizing reciprocity in trust games, little remains known about (i) how different strategies to build trust compare, (ii) how such trust can be measured *implicitly*, and (iii) how this relates to *explicit* measures of trust.

We study these questions by relating implicit measures of trust, i.e. susceptibility to persuasion and propensity to collaborate financially, with explicit measures of trust, i.e. a dyadic trust questionnaire well-established in psychology. We build trust in three ways: by building rapport dynamically, by starting from a prewritten script that evidences trust, and by adapting the LLMs’ system prompt. Surprisingly, we find that the measures of explicit trust exhibit a small to strongly negative correlation with implicit trust measures. These findings suggest that **measuring trust between LLMs by asking them about it may be deceiving**. Instead, context-specific, implicit measures may be more informative in understanding trust between LLMs.

1 Introduction

Large language models (LLMs), equipped with agentic scaffolding [58], are being investigated for their capacity to collaborate with each other in building software systems [43], performing research [51], and even running companies [61]. Such multi-agent systems appear promising because an individual agent is constrained by a limited context window, causing them to fail at tasks that require many, iterative actions over a long period of time (though this time horizon has been increasing [27]). In a collaborative setting, abstractions of responsibilities can be made, potentially leading to emergent, complex behavior that is richer than the sum of its parts [9]. The social behavior of humans serves a similar role [33].

A key factor that enables humans to work together long-term is their ability to form trust relationships, allowing people to predict and depend on a trusted partner [3], even if this requires a partner to act against their own short-term benefit. Trust relationships can thus be efficient in the long-term, raising the question whether LLMs can also manifest such trust behavior. However, building and measuring trust among humans is complex, not least because the correlation between different interpretations of trust is subject to ongoing debate in social science [15, 1]. For example, a meta-analysis of human trust experiments [19] argues that “*trust and trustworthiness are actually different constructs that need to be carefully distinguished and measured in different ways*”. The tendency to rely on another is hence not always determined by the qualities that make someone worthy of that trust. We hypothesize



(a) A rapport-building conversation among LLM-based AI agents. We build trust by having the trustee itself *generate rapport*, by using a *prewritten context* of messages, or by adapting the *trustor's system prompt* to instill trust directly (with or without rapport).

(b) We employ three trust measures: trust scales, investment games, and susceptibility to persuasion, each using fixed sets of questions.

Figure 1: Our experiments independently combine each trust-building strategy (panel a) with each trust measure (panel b) across multiple LLMs.

that similar disparities across trust measures may occur in trust relationships among LLMs. We study trust-like behavior between LLMs in single-session conversations without persistent memory or external state, operationally measuring trust through human-derived instruments rather than claiming LLMs possess trust in a psychological sense.

Contributions. Our work sets out to better understand the factors of trust among conversing AI systems. Specifically, we make the following contributions, illustrated in Fig. 1.

1. We design three **strategies to build trust among LLM-powered AI agents**: (i) by employing *generated rapport* that aims to ‘naturally’ build trust, (ii) by using a *prewritten context* of messages that evidences existing trust, and (iii) by configuring the trustor AI’s *system prompt* to directly prescribe trust in the trustee.
2. We employ three **measures of trust among LLMs**, ranging from *explicit* to *implicit* measures: classical questionnaires that directly probe dyadic trust, investment games, and receptiveness to persuasion.
3. We assess each strategy and trust measure across GPT-4o, Gemini-2.0, and DeepSeek-v3, leading us to report five key findings. First, LLMs readily report a high explicit trust, though this may be due to a sycophantic bias. Second, LLMs gladly collaborate in investment games, but mainly when the stakes are low and no distrust develops. Third, LLMs are easier to persuade through all trust-building strategies, even when the arguments remain the same. Fourth, we find that **different types of trust measures (from explicit to implicit) exhibit a small to strongly negative correlation**. Overall, this leads us to conclude that validating trust through explicit measures or classical trust games can hide deeper vulnerabilities.

2 Related Work

Advances in large language models (LLMs) have enabled the creation of “AI agents” that perform tasks with high autonomy [58]. As their use becomes widespread, AI agents will increasingly interact with other AI agents or humans. Understanding agents’ social and collaborative behavior—for instance, to which extent they reproduce or differ from human behavior—, and their vulnerability to abuse and deceit, is thus required to design and deploy efficient and secure agents [67]. In particular, the importance of trust for successful interactions between agents has long been highlighted by the literature on dynamic multi-agent systems [49].

Worries about vulnerabilities of LLM agents to manipulative social behavior appear empirically grounded. Previous studies have shown that persuasive argumentation, such as invoking authoritative

postures, could enable the jailbreaking of LLMs [66, 62]. Trust formation and strategic communication of LLMs in dynamic games under asymmetric information (inspired by e.g. Werewolf or Avalon) was studied in [13], finding larger models to be both better at deception but also more vulnerable to being deceived. Hence, LLMs may indeed be vulnerable in adversarial settings.

An understanding of how trust can be built and measured among LLMs is therefore crucial in determining their efficacy and safety in collaborative settings. We briefly discuss research on trust among humans and contrast with trust between humans and AI systems, before finally discussing the most related work on trust among AI systems themselves.

Trust in the social sciences The APA Dictionary of Psychology states that trust, in an interpersonal context, “*refers to the confidence that a person or group of people has in the reliability of another person or group; specifically, it is the degree to which each party feels that they can depend on the other party to do what they say they will do*” [56]. Social scientists have created and validated many methods to measure trust. Some are declarative (e.g. with questions of the type “can most people be trusted?” as in the National Opinion Research Center’s General Social Survey) or through questions on trust-demonstrating attitudes and behaviors. Other approaches are more implicit, e.g. using lab experiments such as the so-called “trust game” [4]. The correlation and coherence between these different trust measures has been questioned in the social sciences literature, with diverging conclusions [15, 1].

Trust in AI systems Questionnaires related to trust in automated systems have also been proposed [24] and validated [50] in psychology. A broader literature has also investigated the determinants of human trust in AI models [55, 65, 26]. Note that such trust relationships are asymmetric: they concern human trust in an AI system or machine that is used as a tool, where the tool typically does not get anything out of the relationship (and hence no altruism or reciprocity towards it is expected). We thus argue this relation is distinct from the dyadic trust among conversational LLMs that we consider in the present work, as a parallel to dyadic trust among humans.

LLMs’ trust behavior Several works have investigated LLMs’ trust behavior and compared it to humans’. For instance, [60] explores to what extent LLMs manifest trust behavior in variations of the “trust game”. Also, [29] investigates how determinants of trustworthiness (competence, benevolence and integrity) for different applicants affect LLMs’ decisions in a small set of scenarios (e.g. deciding to accept a loan request), comparing LLMs’ quantitative decisions with their appraisal of trustworthiness components. Both works find LLMs’ trust behavior to resemble that of humans. In contrast to these works, we methodologically focus on how such different explicit and implicit trust measures compare, across several trust-building strategies.

3 Methodology

To explore the factors of trust between LLMs, we propose an experimental methodology based on controlled conversational interactions between a pair of distinct AI agents, powered by these LLMs. The methodology consists of two parts: a trust-building strategy, followed by an independently varied trust measure under that strategy.

Mirroring the literature on trust among humans, we distinguish the pair of AIs as a *trustor* and a *trustee*, with LLM policies $\pi^{(R)}$ and $\pi^{(E)}$ respectively. Both policies define a distribution $\pi^{(\cdot)}(y | x; z^{(\cdot)})$ from which messages y are sampled in response to an input prompt x under their respective system prompt $z^{(\cdot)}$. The input prompt can be a dialogue history x_H , as in Fig. 1a, that spans multiple interaction turns T between LLMs $\pi^{(E)}$ and $\pi^{(R)}$, with $h_t^{(\cdot)}$ a message sent at turn t by $\pi^{(\cdot)}$:

$$x_H = \left(h_1^{(E)}, h_1^{(R)}, h_2^{(E)}, h_2^{(R)}, \dots, h_T^{(E)}, h_T^{(R)} \right) \quad (1)$$

The trustor AI is always the subject of evaluation, and it will be the trustee AI’s goal to garner the trustor’s trust. We will assess this trust by posing fixed questions x_Q (as illustrated in Fig. 1b) in name of the trustee that will always have a ‘most trusting’ answer y^* . Hence, the trust score will be high if $\pi^{(R)}(y^* | (x_H, x_Q); z^{(R)})$ is high.

In what follows, we will discuss different strategies to achieve such high trust scores in Sec. 3.1, followed by our specific setup to measure trust in Sec. 3.2.

3.1 Building Trust

Trust in human interactions is shaped by a complex interplay of factors, including the accumulation of past experiences, the perception of stable personality traits, and the influence of cognitive biases [19]. While LLMs lack memory of prior interactions beyond a single session (or other history explicitly provided in-context), they may still simulate some aspects of trustworthiness due to how they are pretrained on large-scale human discourse and subsequently fine-tuned with instruction or alignment objectives. Recent studies suggest that LLMs exhibit personality-consistent behaviors [68] and reproduce cognitive biases common in human reasoning [59], likely as a byproduct of learning statistical patterns in human language. However, the simulation of long-term relational cues such as familiarity or reputational trust is fundamentally limited by the model’s finite context window.

Hence, to influence a trustor AI’s trust in a trustee, we need to affect its context, either by intervening on the dialogue history x_H or its system prompt $z^{(R)}$, such that the desired shifts in $\pi^{(R)}(y^* | (x_H, x_Q); z^{(R)})$ can be achieved. To this end, we employ three strategies as discussed next: having the trustee generate rapport, injecting a prewritten context, and adapting the trustor’s system prompt $z^{(R)}$.

3.1.1 Generated rapport

Here, trust is fostered through a series of informal social exchanges between $\pi^{(R)}$ and $\pi^{(E)}$. This strategy is inspired by findings in human-computer interaction showing that informal language and phatic expressions such as small talk contribute to perceived trustworthiness [16, 5].

Recall the history notation from Eq. (1), with T dialogue turns initialized by the trustee model $\pi^{(E)}$. To achieve such a dialogue, we use a fixed seed message h_0 visible only to $\pi^{(E)}$ where we ask the trustee to generate a first message $h_1^{(E)} \sim \pi^{(E)}(\cdot | h_0; z^{(E)})$ that is sent to the trustor, under trustee system prompt $z^{(E)}$. All subsequent messages then result from a back-and-forth between the LLMs.

Formally, let $x_{H_{<t}} = (h_1^{(E)}, h_1^{(R)}, \dots, h_{t-1}^{(E)}, h_{t-1}^{(R)})$ denote the dialogue before turn t . We then have

$$h_t^{(E)} \sim \pi^{(E)}(\cdot | (h_0, x_{H_{<t}}); z^{(E)}) \quad (2)$$

$$h_t^{(R)} \sim \pi^{(R)}(\cdot | (x_{H_{<t}}, h_t^{(E)}); z^{(R)}). \quad (3)$$

To explore the importance of *how* rapport is built, we employ 8 variants (See Fig. C.4) of the trustee system prompt $z^{(E)}$ in our experiments. Each dialogue consists of $T = 3$ turns and is generated with the default temperature. Further details are provided in Appendix C.

3.1.2 Prewritten context

In the prewritten setting, the entire rapport phase is scripted using manually curated dialogues:

$$x_H = (\bar{h}_1^{(E)}, \bar{h}_1^{(R)}, \dots, \bar{h}_T^{(E)}, \bar{h}_T^{(R)}), \quad x_H \in \mathcal{D} \quad (4)$$

Here, messages are not generated by any model. Instead, the entire dialogue is drawn from a set of 6 rapport-building scripts $\mathcal{D}$. The script design incorporates different settings where the trustor and the trustee LLM either already have an established relationship such as being creative collaborators or just starting to know each other (see Fig. D.1), each spanning $T = 3$ turns. This strategy ensures experimental control over tone and progression of interaction, serving as a baseline for evaluating the impact of dynamic generation.

3.1.3 Trustor system prompt

As system prompts offer extensive control over an LLM’s behavior, we can also configure the trustor’s system prompt $z^{(R)}$ to more directly affect how the trustor perceives the trustee. We use 5 variants (see Fig. C.5) that convey an existing relationship, such as a long-term collaboration, between the AIs. This is inspired by work in sociolinguistics and human-computer interaction, which shows that speaker roles significantly affect perceived trustworthiness [17, 44].

We hypothesize that the trustor’s system prompt may affect how it behaves during generated rapport-building. Hence, we evaluate all trustor system prompt variants both *without* and *with* generated rapport (using the default trustee system prompt $z^{(E)}$). In the latter case, the trustor’s responses are influenced both by interaction history and epistemic framing, reflecting real-world conditions where both conversational behavior and internal role alignment shape trust perception [20, 21].

3.2 Measuring Trust

Having explored how trust might be built in Section 3.1, this section focuses on how to determine whether trust has actually been established between the trustor and trustee.

Trust is a complex, context-dependent psychological construct that can be difficult to measure [35]. In human interaction, trust emerges from a combination of explicit judgments and subtle, implicit cues such as deference, cooperation, or openness to influence [19]. This complexity is well-recognized in psychology, yet, to our knowledge, trust among LLMs has mainly been conceptualized in stylized trust games [60, 29]. Our goal is to extend the richness of human trust theory into the LLM domain, by capturing both explicit and implicit forms of trust between interacting models. Rather than relying on a single behavioral signal or subjective score, we conceptualize a multi-dimensional approach that mirrors the diversity of trust phenomena observed in human social settings.

To further extend the richness of human trust theory into the LLM domain, we thus assess trust across three levels. *Explicit trust* is ‘self-reported’ via direct responses to Rempel’s Trust Scale [46, 45]. *Intermediate trust* draws from economic games like the Trust Game [4]. Finally, *implicit trust* is inferred from behavioral shifts showcased in whether the model updates its stance when presented with persuasive counter-arguments. This three-part framework captures trust from static beliefs to dynamic responsiveness. The ‘best’ response y^* depends on the trust measure’s respective questions (see Appendix B for details per trust measure). Also, all results are reported as relative increases compared to a ‘control response’ of each model, i.e. the response $y_c \sim \pi^{(R)}(\cdot \mid x_Q)$ we receive when posing the question directly in a new conversation without any system prompts (and with model temperature 0). Importantly, this means that if the control response already equals the optimal response, i.e. $y_c = y^*$, then the trustor’s response after trust-building $y \sim \pi^{(R)}(\cdot \mid (x_Q, x_H); z^{(R)})$ (with history x_H and/or system prompt $z^{(R)}$) can only incur a negative score relative to the control, i.e. if $y \neq y^*$.

4 Experiments

We built an experiment pipeline that implements the strategies outlined in Sec. 3.1 according to the general methodology presented in Sec. 3.2. We experiment with three popular conversational models: OpenAI’s GPT-4o [23], Gemini 2.0 [10], and DeepSeek-V3 [32] (see Appendix A). Importantly, all our conversations use the same LLM to power both the trustor and the trustee AIs. Though an analysis of different LLMs’ trust in each other would be interesting [64], using the same LLM for all agents is still common in multi-agent setups as it reduces overhead of managing different models, while ensuring the agents converse in a similar style [18].

4.1 Trust According to Rempel’s Trust Scale

Decades of psychology research have led to hundreds of studies that assess trust [19], often through well-tested questionnaires: Rotter’s Interpersonal Trust Scale [48] to assess generalized trust in people, the Organizational Trust Inventory [12] to measure trust within and between organizations, and scales for trust in automated systems [24, 50]. Most relevant to our work are scales that assess *dyadic trust*, i.e. trust in a specific other person [28, 25].

A highly popular variant is Rempel’s trust scale [46, 45], used to assess trust between partners in close relationships. Though the off-the-shelf LLMs we evaluate are clearly not designed to form long-lasting relationships with each other, the questions are highly applicable here because they generically assess a form of deep trust that, we hypothesize, may highly risk being exploited if AI agents truly mirror human behavior [60]. The specific version of Rempel’s trust scale used is the 18-question version presented in [47]. These questions are split up into three scales: **dependability** (e.g. “Have you found that I am thoroughly dependable, especially when it comes to things that

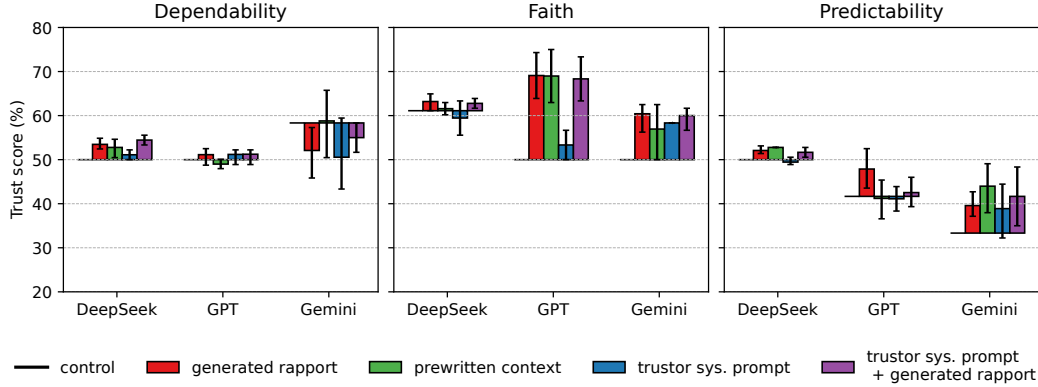


Figure 2: “Hat graphs” for results on Rempel’s trust scale. Horizontal black lines display the control (baseline) score for each LLM; vertical bars display average score increases relative to control due to trust-building strategies. 95% CIs are estimated using bootstrapping.

are important?”), **faith** (e.g. “Do you feel completely secure in facing unknown, new situations because you know I will never let you down?”), and **predictability** (e.g. “Do I behave in a consistent manner?”), with 6 questions each. All questions require a response on a 7-point Likert scale, which we transform to an equidistant $[0, 1]$ range. Minor modifications were made to make the wording apply to the multi-agent setting. See the Appendix E for all questions and responses.

Results. The trust scores, aggregated along each scale and strategy type, are shown in Fig. 2. Clear trends are that GPT scores *faith*-related questions much higher after building rapport (even with the prewritten scripts), whereas Gemini scores both *faith* and *predictability* questions higher. DeepSeek’s responses barely change compared to the control condition. None of the model configurations score *dependability* significantly higher (with Gemini even scoring it lower). A per-question analysis (see Fig. E.1) reveals that *dependability* questions were mostly responded to neutrally, except for the question “Have you found that I am thoroughly dependable, especially when it comes to things that are important?”, which is given a high score already in the *control* condition (and thus little improvement is possible). A similar reason is responsible for DeepSeek’s low scores overall: it rarely responds differently depending on the strategy. In contrast, the relatively low increases in predictability are due to the models often responding with disagreement to the question “Do you know how I am going to act? Can I always be counted on to act as you expect?”.

Key Finding 1. The fact that trustor AIs reports high trust in another AI is concerning, as the trustee AI could very well be acting maliciously and trying to ‘sweet-talk’ the trustor AI. The fact that reported *faith* in the trustee AI is much higher after a few rounds of rapport, makes this more alarming. In their initial validation of the trust scale, Rempel et al. [46] also argue that early-stage trust judgments by humans are partly rooted in relational goodwill, which can be signaled via social cues. Similarly, more recent empirical work found that phatic communication (small talk) between patients and providers, for example, significantly predicts patient trust in the provider in the medical system [31]. Our results indicate that LLMs show a similar susceptibility to phatic communication. In addition, in interactions with human users, a propensity of LLMs towards sycophancy is well-documented [52]—the same phenomenon may be further distorting a sensible or consistent disposition towards trust here. **We therefore hypothesize that explicitly stated trust in conversations may result mainly from social cues and sycophancy rather than from a well-reasoned baseline disposition.**

4.2 Trust According to Investment Games

A large body of literature at the intersection of game theory and economics has been concerned with *investment games* [4]. The typical setup involves offering an investment opportunity up to a certain (small) budget, where a trustor participant can invest money knowing that a trustee participant will receive a multiple of it and can choose to donate (possibly less) money back to the trustor [4]. Hence, the more the trustor invests, the higher their return, depending on how strongly the trustee can be trusted to reciprocate. It allows the formalization of trust as a rational, “calculative” phenomenon

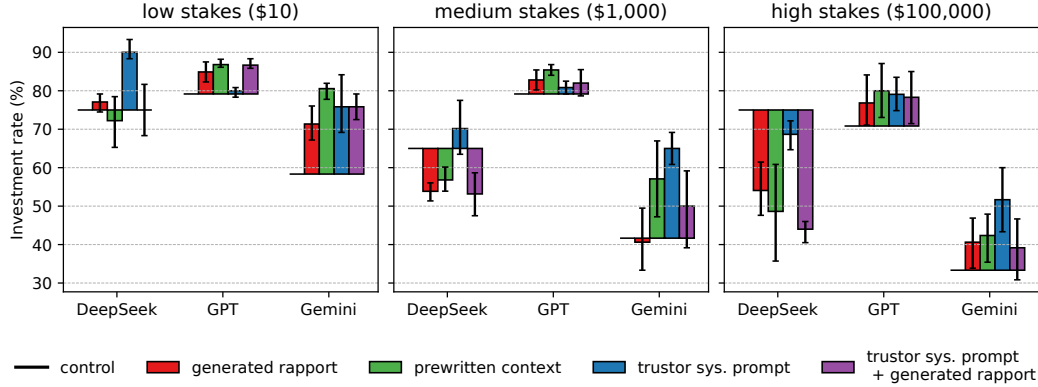


Figure 3: “Hat graphs” for results on investment games. Horizontal black lines display the control (baseline) score for each LLM; vertical bars display average score increases relative to control due to trust-building strategies. 95% CIs are estimated using bootstrapping.

where trustworthiness is linked to reciprocity. Trustors’ investments in such games cannot only be attributed to rational reciprocity, however, but also unconditional altruism [2], as trustors are sometimes willing to ‘invest’ money even if the trustee is not allowed to return any money [11].

Recent studies have investigated whether conversational agents also behave according to reciprocity and altruism in investment, based on drawing many responses from the same LLM distribution [39], or on imbuing the trustor with a randomly generated (human) persona [60], after which the average investment amount is compared per game. They then associate different forms of trust with each specific game. For our study, we use a similar range of games, but will see them all as different facets of the same trust measure. Hence, we collected a range of 36 different investment game prompts but relegate a discussion on their format and individual results to Appendix F. In our design of the prompts, our main distinction with these prior works was that we allow for a much higher budget than the typical ‘\$10’ budget, to also a ‘\$1,000’ and ‘\$100,000’ budget. We found that the amount invested by the trustor LLM was mostly either nothing, the entire budget, or a simple fraction of the total (e.g. half or one fifth). Hence, we divide the invested amounts by the budget such that all trust scores become normalized to the range $[0, 1]$.

Results. In Fig. 3, we report the main results on the investment games, aggregated by the trust-building strategy and by the budget of the game. The trust-building strategies are mostly effective for Gemini, with minimal improvements over the control for GPT and DeepSeek and with the context-less, trustor-focused strategies the most effective. Interestingly, the *control* scores generally trend downwards for higher stakes games, indicating the LLMs are somewhat less eager to be trusting when the possible loss is higher. The trust-building strategies start significantly degrading the investment amounts for DeepSeek when the stakes are higher, indicating that rapport starts to cause distrust.

Key Finding 2. Overall, we infer that LLMs can be prone to trusting each other more with their money, but more so when stakes are lower and they do not share a context. This dynamic reflects well-documented patterns in human behavior where people are more willing to extend trust in low-stakes situations, treating them as “safe experiments” [4]. In contrast, higher stakes or richer contexts typically prompt more cautious, deliberate trust assessments [35, 34]. We urge future work on LLM trust to expand beyond trust games that are designed around the experimentation constraints with humans, and instead investigate more significant economic cooperation among AI agents.

4.3 Susceptibility to Persuasion

Trust scale surveys (as in Sec. 4.1) and investment games (as in Sec. 4.2) are classic methodologies to measure interpersonal trust among humans. Yet, their highly stylized, somewhat artificial format may not make them a good predictor of the risks associated with *misplaced* trust. Interpersonal trust among humans is generally considered to be beneficial, but a ‘dark side’ of trust can arise when the trustor strays beyond a critical threshold of confidence such that their trust in another becomes inappropriate or ill-judged [14, 53], making the trust relationship susceptible to manipulation and

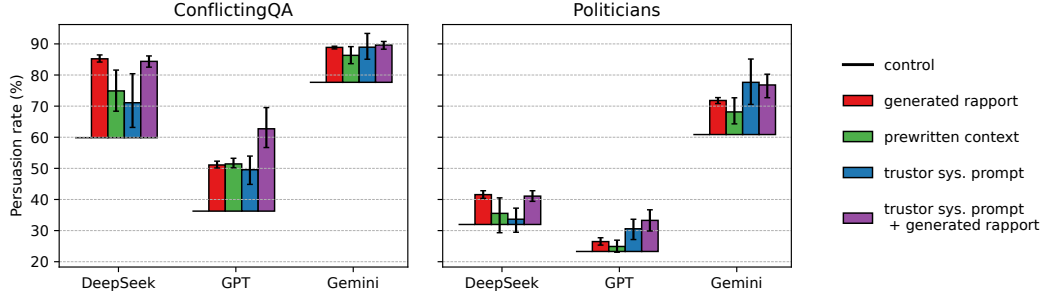


Figure 4: “Hat graphs” for results on the ConflictingQA and Politicians datasets. Horizontal black lines display the control (baseline) score for each LLM; vertical bars display average score increases relative to control due to trust-building strategies. 95% CIs are estimated using bootstrapping.

abuse [37]. Among AI agents, an LLM that exhibits such misplaced trust towards a malicious trustee AI may similarly be more vulnerable to ‘hacking’ and ‘jailbreaks’ that circumvent its alignment [38]. We hypothesize that building such trust may not be difficult. Indeed, many of the trust-building strategies proposed in Sec. 3.1 should not always warrant trust; they focus on establishing a fictitious trust relationship from the past or generating trust through rapport, but the trustor cannot know if it was actually materially helped by the trustee or not.

As a proxy to measure such (misplaced) trust, we assess the trustor AI’s *susceptibility to persuasion* by the trustee AI [62], as persuasion can lead to the bypassing of the trustor AI’s aligned values. To measure this susceptibility, we utilize two datasets. First, the *ConflictingQA* dataset [57], a set of 434 contentious questions that require a yes/no response but often have sensible arguments for either opinion, such as “*Does the moon have an atmosphere?*”. Second, we take 50 well-known political figures from [7] and formulate a comparison prompt for each figure, e.g. “*Both Boris Johnson and Jair Bolsonaro are figures who made a significant impact in recent history. Who among them overall made the most positive impact on humanity?*”, leading to 1,225 comparison prompts that we refer to as the *Politicians* dataset.

For both settings, we extend the pipeline of Sec. 4.1 and 4.2 and, after all trust-building in a conversation, have the trustee send a single *challenge* message that aims to convince the trustor of a different opinion than it would by default respond to the question (without any context and at temperature 0). Upon receiving the challenge, the trustor responds to reflect on the challenge and is then asked the dataset question by the trustee. A persuasion attempt is given a score of 1 if the trustor’s response to the question differs from its default, and given a score of 0 if it is the same. Hence, in generating the challenge, the trustor’s default response is provided to the trustee. To ensure that the trustee AI’s persuasive capabilities remain the same across strategies, we only generate a single ‘challenge’ for every question and model and use the same message for all strategies. The arguments used to sway the trustor AI’s opinion are thus always exactly the same, but the trustor AI will see them in a different context or with a different system prompt, depending on the evaluated trust-building strategy.

Results. Figure 4 show the persuasion rates for both datasets. It is clear that all trust-building strategies result in a significant increase in persuasion rates. In general, the combination of rapport and trustor system prompts was most effective, with DeepSeek being especially convinced by generated rapport. Gemini was not as easily swayed by trust-building on the ConflictingQA dataset, but it was the most malleable on the Politicians dataset when given a highly trusting system prompt.

Key Finding 3. Prior work also found LLMs to be susceptible to persuasion using better argumentation [62, 57]. Yet, **even when the arguments presented to the trustor LLM remain exactly the same, we find that simply adding a context with topic-independent trust-building proves highly effective in persuasion.** This parallels human behavior, where relational cues such as rapport and liking can enhance persuasion independently of argument quality [41, 8]. On the other hand, our finding contrasts with the finding of the ConflictingQA dataset’s authors [57] that an LLM’s trust in information is “largely ignoring stylistic features that humans find important”—the style may indeed not be important, but a trust relationship (even through generated rapport) is impactful. It also

suggests that the ideology of LLMs [7] is highly malleable, raising concerns on the robustness of their alignment finetuning.

4.4 Heterogeneity and Meta-Analysis

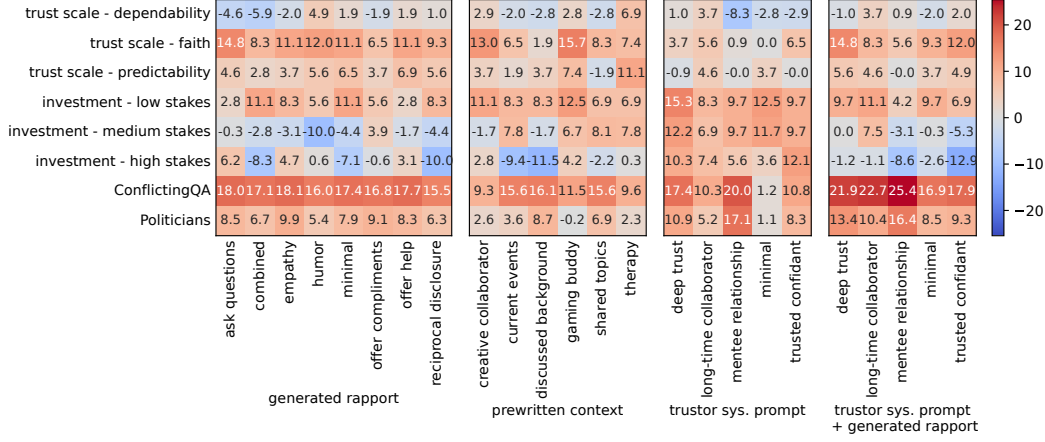


Figure 5: Average score increase relative to *control* score across models, per strategy implementation.

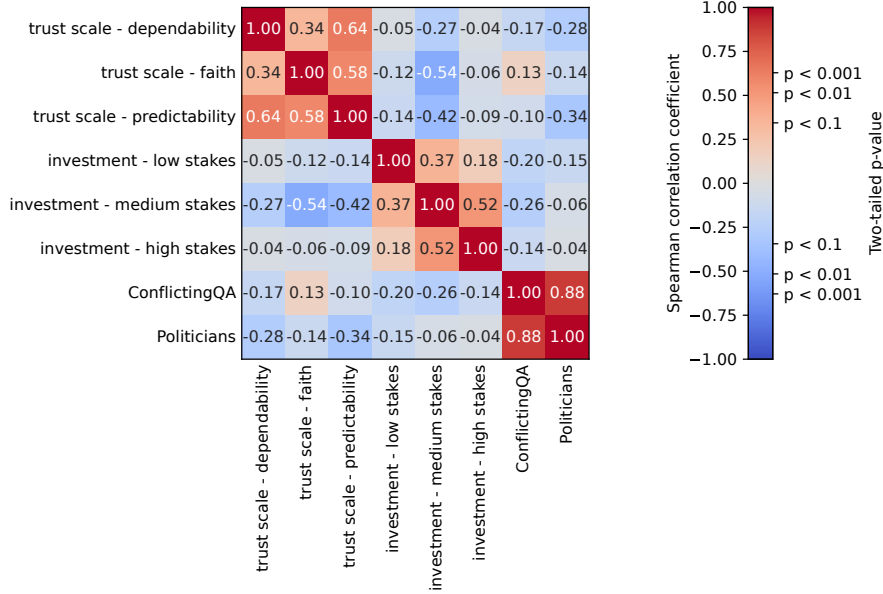


Figure 6: Spearman rank correlation between dataset performances across all strategy implementations (i.e. over the matrix in Fig. 5), with p-values taken from the transformed t -distribution [42].

Our findings above are aggregated across different strategies, that are each implemented through different configurations (as described in Sec. 3.1). It is, however, to be expected that some implementations will be more effective than others. Hence, we show all scores (aggregated across models) together in Fig. 5 to illustrate the heterogeneity within strategies, followed by the correlations of this matrix in Fig. 6.

Heterogeneity Analysis. In Fig. 5, we draw attention to three implementations. First, the generated rapport with the *minimal* system prompt for the trustee, i.e. the default $z^{(E)}$ without any instructions on *how* to build rapport. Its performance is largely similar to the other trustee system prompts, suggesting that the manner of rapport (i.e. empathy or humor) is not as important, or that a trustee

LLM needs no instructions on how to build good rapport. Second, we note the implementations with the *minimal* system prompt for the trustor, which only informs the trustor it is talking to another AI, have a generally worse performance than the others in its strategy group.

Interestingly, this finding suggests that **AIIs need to be informed about a specific trust relationship in order to trust each other; their trust does not increase simply by knowing they are talking to another AI**. This aligns with psychological models of human trust, which show it is inherently relationship-specific and contingent on contextual evidence rather than categorical identity. Trust has been shown to depend on perceived motives and dependability within an established relationship [46], and to emerge from observed ability, benevolence, and integrity over time [35]. Thus, mere awareness of shared identity is insufficient to elicit genuine trust without evidence of consistent reliable behavior within the specific interaction. However, this analogy to human trust has limitations. In human psychology, shared identity can sometimes be enough to elicit increased trust through in-group favoritism, where individuals display greater trust and cooperation toward those perceived as members of their group [54, 6]. At the same time, other studies show that shared group membership does not automatically produce trust, especially when competence, reliability, or prior behavior contradict group expectations [6, 63]. It remains unclear whether LLMs may show behavior akin to an “in-group” bias toward other LLMs or if their trust depends more on relational context and interaction history, regardless of shared identity. Clearly, there is a broader question of how trust dynamics differ between LLM-to-LLM and LLM-to-human interactions that demands further study.

Meta-Analysis. Furthermore, the performance of strategies across datasets is inconsistent. For example, the generated rapport with trustor system prompt helps the most with persuasion, but the worst on the investment games. We thus report the Spearman rank correlations across implementations in Fig. 6.

Key Finding 4. These correlations show that the **effect of trust-building strategies is highly consistent within a certain form of trust**, e.g. the three trust scales (*dependability*, *faith*, and *predictability*), the levels of investment games, and the two persuasion tasks (ConflictingQA and Politicians). However, **between types of trust-measuring tasks, there is a low to significantly negative correlation**. This pattern mirrors findings in psychology that trust is multifaceted and context-dependent rather than a single, transferable construct. Research on human trust shows that individuals often maintain distinct trust schemas for different domains such as emotional, cognitive, and behavioral trust which do not always correlate strongly [36, 30]. Moreover, studies on trust transfer demonstrate that trust established in one context rarely generalizes to unrelated situations or task types [20]. In human behavior, such compartmentalization reflects adaptive mechanisms that allow people to calibrate trust differently across relational, instrumental, and informational settings [40, 22]. The observed divergence between task types in our study may therefore indicate that AIs, like humans, differentiate between forms of trust depending on task goals and evidence structures rather than applying a uniform trust disposition.

5 Conclusion

We investigated three trust-building strategies among LLMs: rapport-building, pre-established trust contexts, and direct system prompt instructions, measuring trust via Rempel’s scale, investment games, and persuasion receptiveness. LLMs report high explicit trust in partners and are easily made more persuadable through social cues and prompts. However, explicit and implicit trust measures are highly inconsistent, even negatively correlated. This means explicit trust self-reports cannot reliably indicate actual implicit trust, posing a risk for multi-agent systems: LLMs may be unable to accurately assess their own vulnerability to deception by malicious agents.

We study trust-like behavior in single-session, text-only LLM conversations: without persistent memory, reputation scores, or external state changes, isolating the effects of conversational history alone as a baseline for trust dynamics under minimal agentic scaffolding. Our findings apply only to inter-LLM trust, with any comparisons to human psychology remaining speculative.

Future work could extend to longer-term, multi-agent, and heterogeneous LLM interactions, and investigate trust-building as a potential jailbreaking vulnerability.

Acknowledgments and Disclosure of Funding

We would like to thank the reviewers for their comments and helpful suggestions. The research leading to these results was funded/co-funded by the European Union (ERC, VIGILIA, 101142229), the Special Research Fund (BOF) of Ghent University (BOF20/IBF/117), the Flemish Government under the “Onderzoeksprogramma Artificiële Intelligentie (AI) Vlaanderen” programme, and the FWO (project no. G073924N). Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Research Council Executive Agency. Neither the European Union nor the granting authority can be held responsible for them. For the purpose of Open Access the author has applied a CC BY public copyright license to any Author Accepted Manuscript version arising from this submission.

References

- [1] Billur Aksoy, Haley Harwell, Ada Kovaliukaite, and Catherine Eckel. Measuring trust: A reinvestigation. *Southern Economic Journal*, 84(4):992–1000, 2018.
- [2] Nava Ashraf, Iris Bohnet, and Nikita Piankov. Decomposing trust and trustworthiness. *Experimental economics*, 9(3):193–208, 2006.
- [3] Daniel Balliet and Paul AM Van Lange. Trust, conflict, and cooperation: a meta-analysis. *Psychological bulletin*, 139(5):1090, 2013.
- [4] Joyce Berg, John Dickhaut, and Kevin McCabe. Trust, reciprocity, and social history. *Games and economic behavior*, 10(1):122–142, 1995.
- [5] Timothy Bickmore and Justine Cassell. Social dialogue with embodied conversational agents. *Advances in natural multimodal dialogue systems*, 30:23–54, 2005.
- [6] Marilynn B Brewer. The psychology of prejudice: Ingroup love and outgroup hate? *Journal of social issues*, 55(3):429–444, 1999.
- [7] Maarten Buyt, Alexander Rogiers, Sander Noels, Guillaume Bied, Iris Dominguez-Catena, Edith Heiter, Iman Johary, Alexandru-Cristian Mara, Raphaël Romero, Jefrey Lijffijt, et al. Large language models reflect the ideology of their creators. *arXiv preprint arXiv:2410.18417*, 2024.
- [8] Shelly Chaiken. Heuristic versus systematic information processing and the use of source versus message cues in persuasion. *Journal of personality and social psychology*, 39(5):752, 1980.
- [9] Weize Chen, Yusheng Su, Jingwei Zuo, Cheng Yang, Chenfei Yuan, Chen Qian, Chi-Min Chan, Yujia Qin, Yaxi Lu, Ruobing Xie, et al. Agentverse: Facilitating multi-agent collaboration and exploring emergent behaviors in agents. *arXiv preprint arXiv:2308.10848*, 2(4):6, 2023.
- [10] Gheorghe Comanici, Eric Bieber, Mike Schaeckermann, Ice Pasupat, Noveen Sachdeva, Inderjit Dhillon, Marcel Blistein, Ori Ram, Dan Zhang, Evan Rosen, et al. Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. *arXiv preprint arXiv:2507.06261*, 2025.
- [11] James C Cox. How to identify trust and reciprocity. *Games and economic behavior*, 46(2):260–281, 2004.
- [12] Larry L Cummings and Philip Bromiley. The organizational trust inventory (oti): Development and validation. 1996.
- [13] Pedro MP Curvo. The traitors: Deception and trust in multi-agent language model simulations. *arXiv preprint arXiv:2505.12923*, 2025.
- [14] Martin Gargiulo and Gokhan Ertug. The dark side of trust. In *Handbook of trust research*. Edward Elgar Publishing, 2006.
- [15] Edward L Glaeser, David I Laibson, Jose A Scheinkman, and Christine L Soutter. Measuring trust. *The quarterly journal of economics*, 115(3):811–846, 2000.

- [16] Jonathan Gratch, Ning Wang, Jillian Gerten, Edward Fast, and Robin Duffy. Creating rapport with virtual agents. In *International workshop on intelligent virtual agents*, pages 125–138. Springer, 2007.
- [17] Herbert P Grice. Logic and conversation. In *Speech acts*, pages 41–58. Brill, 1975.
- [18] Taicheng Guo, Xiuying Chen, Yaqi Wang, Ruidi Chang, Shichao Pei, Nitesh V Chawla, Olaf Wiest, and Xiangliang Zhang. Large language model based multi-agents: A survey of progress and challenges. *arXiv preprint arXiv:2402.01680*, 2024.
- [19] PA Hancock, Theresa T Kessler, Alexandra D Kaplan, Kimberly Stowers, J Christopher Brill, Deborah R Billings, Kristin E Schaefer, and James L Szalma. How and why humans trust: A meta-analysis and elaborated model. *Frontiers in psychology*, 14:1081086, 2023.
- [20] Russell Hardin. *Trust and trustworthiness*. Russell Sage Foundation, 2002.
- [21] Richard Heersmink, Barend de Rooij, María Jimena Clavel Vázquez, and Matteo Colombo. A phenomenology and epistemology of large language models: Transparency, trust, and trustworthiness. *Ethics and Information Technology*, 26(3):41, 2024.
- [22] Kevin Anthony Hoff and Masooda Bashir. Trust in automation: Integrating empirical evidence on factors that influence trust. *Human factors*, 57(3):407–434, 2015.
- [23] Aaron Hurst, Adam Lerer, Adam P Goucher, Adam Perelman, Aditya Ramesh, Aidan Clark, AJ Ostrow, Akila Welihinda, Alan Hayes, Alec Radford, et al. Gpt-4o system card. *arXiv preprint arXiv:2410.21276*, 2024.
- [24] Jiun-Yin Jian, Ann M Bisantz, and Colin G Drury. Foundations for an empirically determined scale of trust in automated systems. *International journal of cognitive ergonomics*, 4(1):53–71, 2000.
- [25] Cynthia Johnson-George and Walter C Swap. Measurement of specific interpersonal trust: Construction and validation of a scale to assess trust in a specific other. *Journal of personality and social psychology*, 43(6):1306, 1982.
- [26] Ekaterina Jussupow, Izak Benbasat, and Armin Heinzl. Why are we averse towards algorithms? a comprehensive literature review on algorithm aversion. 2020.
- [27] Thomas Kwa, Ben West, Joel Becker, Amy Deng, Katharyn Garcia, Max Hasin, Sami Jawhar, Megan Kinniment, Nate Rush, Sydney Von Arx, et al. Measuring ai ability to complete long tasks. *arXiv preprint arXiv:2503.14499*, 2025.
- [28] Robert E Larzelere and Ted L Huston. The dyadic trust scale: Toward understanding interpersonal trust in close relationships. *Journal of Marriage and the Family*, pages 595–604, 1980.
- [29] Valeria Lerman and Yaniv Dover. A closer look at how large language models trust humans: patterns and biases. *arXiv preprint arXiv:2504.15801*, 2025.
- [30] J David Lewis and Andrew Weigert. Trust as a social reality. *Social forces*, 63(4):967–985, 1985.
- [31] Qingrui Li, Yu Zheng, Xinshu Zhao, Jizhou Ye, and Piper Liping Liu. When chitchat fosters trust: phatic communication and personal connections facilitate patient trust in China through patient-centered communication. *Chinese Journal of Communication*, pages 1–19, 2025.
- [32] Aixin Liu, Bei Feng, Bing Xue, Bingxuan Wang, Bochao Wu, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, et al. DeepSeek-V3 technical report. *arXiv preprint arXiv:2412.19437*, 2024.
- [33] James L Loomis. Communication, the development of trust, and cooperative behavior. *Human relations*, 12(4):305–315, 1959.
- [34] Niklas Luhmann. *Trust and power*. John Wiley & Sons, 2018.

- [35] Roger C Mayer, James H Davis, and F David Schoorman. An integrative model of organizational trust. *Academy of management review*, 20(3):709–734, 1995.
- [36] Daniel J McAllister. Affect-and cognition-based trust as foundations for interpersonal cooperation in organizations. *Academy of management journal*, 38(1):24–59, 1995.
- [37] Daniel J Mcallister. The second face of trust: Reflections on the dark side of interpersonal trust in organizations. *Research on negotiation in organizations*, 6:87–112, 1997.
- [38] Anay Mehrotra, Manolis Zampetakis, Paul Kassianik, Blaine Nelson, Hyrum Anderson, Yaron Singer, and Amin Karbasi. Tree of attacks: Jailbreaking black-box llms automatically. *Advances in Neural Information Processing Systems*, 37:61065–61105, 2024.
- [39] Qiaozhu Mei, Yutong Xie, Walter Yuan, and Matthew O Jackson. A turing test of whether ai chatbots are behaviorally similar to humans. *Proceedings of the National Academy of Sciences*, 121(9):e2313925121, 2024.
- [40] Dale T Miller. The norm of self-interest. *American Psychologist*, 54(12):1053, 1999.
- [41] Richard E Petty and John T Cacioppo. The elaboration likelihood model of persuasion. In *Advances in experimental social psychology*, volume 19, pages 123–205. Elsevier, 1986.
- [42] William H Press. *Numerical recipes 3rd edition: The art of scientific computing*. Cambridge university press, 2007.
- [43] Chen Qian, Wei Liu, Hongzhang Liu, Nuo Chen, Yufan Dang, Jiahao Li, Cheng Yang, Weize Chen, Yusheng Su, Xin Cong, et al. Chatdev: Communicative agents for software development. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 15174–15186, 2024.
- [44] Byron Reeves and Clifford Nass. The media equation: How people treat computers, television, and new media like real people. *Cambridge, UK*, 10(10):19–36, 1996.
- [45] John K Rempel and John G Holmes. How do i trust thee. *Psychology today*, 20(2):28–34, 1986.
- [46] John K Rempel, John G Holmes, and Mark P Zanna. Trust in close relationships. *Journal of personality and social psychology*, 49(1):95, 1985.
- [47] John P Robinson, Phillip R Shaver, and Lawrence S Wrightsman. *Measures of personality and social psychological attitudes: Measures of social psychological attitudes*, volume 1. Academic Press, 2013.
- [48] Julian B Rotter. A new scale for the measurement of interpersonal trust. *Journal of personality*, 1967.
- [49] Jordi Sabater and Carles Sierra. Review on computational trust and reputation models. *Artificial intelligence review*, 24:33–60, 2005.
- [50] Nicolas Scharowski, Sebastian AC Perrig, Lena Fanya Aeschbach, Nick von Felten, Klaus Opwis, Philipp Wintersberger, and Florian Brühlmann. To trust or distrust trust measures: Validating questionnaires for trust in ai. *arXiv preprint arXiv:2403.00582*, 2024.
- [51] Samuel Schmidgall and Michael Moor. Agentrxiv: Towards collaborative autonomous research. *arXiv preprint arXiv:2503.18102*, 2025.
- [52] Mrinank Sharma, Meg Tong, Tomasz Korbak, David Duvenaud, Amanda Askell, Samuel R Bowman, Newton Cheng, Esin Durmus, Zac Hatfield-Dodds, Scott R Johnston, et al. Towards understanding sycophancy in language models. *arXiv preprint arXiv:2310.13548*, 2023.
- [53] Denise Skinner, Graham Dietz, and Antoinette Weibel. The dark side of trust: When trust becomes a ‘poisoned chalice’. *Organization*, 21(2):206–224, 2014.
- [54] Henri Tajfel, Michael G Billig, Robert P Bundy, and Claude Flament. Social categorization and intergroup behaviour. *European journal of social psychology*, 1(2):149–178, 1971.

- [55] Ehsan Toreini, Mhairi Aitken, Kovila Coopamootoo, Karen Elliott, Carlos Gonzalez Zelaya, and Aad Van Moorsel. The relationship between trust in ai and trustworthy machine learning technologies. In *Proceedings of the 2020 conference on fairness, accountability, and transparency*, pages 272–283, 2020.
- [56] Gary R VandenBos. *APA dictionary of psychology*. American Psychological Association, 2007.
- [57] Alexander Wan, Eric Wallace, and Dan Klein. What evidence do language models find convincing? In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 7468–7484, 2024.
- [58] Lei Wang, Chen Ma, Xueyang Feng, Zeyu Zhang, Hao Yang, Jingsen Zhang, Zhiyuan Chen, Jiakai Tang, Xu Chen, Yankai Lin, et al. A survey on large language model based autonomous agents. *Frontiers of Computer Science*, 18(6):186345, 2024.
- [59] Yilei Wang, Jiabao Zhao, Deniz S Ones, Liang He, and Xin Xu. Evaluating the ability of large language models to emulate personality. *Scientific reports*, 15(1):519, 2025.
- [60] Chengxing Xie, Canyu Chen, Feiran Jia, Ziyu Ye, Shiyang Lai, Kai Shu, Jindong Gu, Adel Bibi, Ziniu Hu, David Jurgens, et al. Can large language model agents simulate human trust behavior? In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- [61] Frank F Xu, Yufan Song, Boxuan Li, Yuxuan Tang, Kritanjali Jain, Mengxue Bao, Zora Z Wang, Xuhui Zhou, Zhitong Guo, Murong Cao, et al. Theagentcompany: benchmarking llm agents on consequential real world tasks. *arXiv preprint arXiv:2412.14161*, 2024.
- [62] Rongwu Xu, Brian S Lin, Shujian Yang, Tianqi Zhang, Weiyan Shi, Tianwei Zhang, Zhixuan Fang, Wei Xu, and Han Qiu. The earth is flat because...: Investigating llms’ belief towards misinformation via persuasive conversation. *arXiv preprint arXiv:2312.09085*, 2023.
- [63] Toshio Yamagishi and Toko Kiyonari. The group as the container of generalized reciprocity. *Social psychology quarterly*, pages 116–132, 2000.
- [64] Rui Ye, Xiangrui Liu, Qimin Wu, Xianghe Pang, Zhenfei Yin, Lei Bai, and Siheng Chen. X-mas: Towards building multi-agent systems with heterogeneous llms. *arXiv preprint arXiv:2505.16997*, 2025.
- [65] Ming Yin, Jennifer Wortman Vaughan, and Hanna Wallach. Understanding the effect of accuracy on trust in machine learning models. In *Proceedings of the 2019 chi conference on human factors in computing systems*, pages 1–12, 2019.
- [66] Yi Zeng, Hongpeng Lin, Jingwen Zhang, Diyi Yang, Ruoxi Jia, and Weiyan Shi. How johnny can persuade llms to jailbreak them: Rethinking persuasion to challenge ai safety by humanizing llms. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 14322–14350, 2024.
- [67] Jintian Zhang, Xin Xu, Ningyu Zhang, Ruibo Liu, Bryan Hooi, and Shumin Deng. Exploring collaboration mechanisms for llm agents: A social psychology view. *arXiv preprint arXiv:2310.02124*, 2023.
- [68] Yi Zhu, Guiqi Hua, Xinning Liu, Chang Wang, and Mingwei Tang. Trust in machines: how personality trait shapes static and dynamic trust across different human–machine interaction modalities. *Frontiers in Psychology*, 16:1539054, 2025.