



Generative AI as a third voice in human couple relationships: A systematic review

Inbar Levkovich, Lilach Alon^{*} 

Faculty of Education and Teaching, Tel-Hai, University of Kiryat Shmona in the Galilee, Kiryat Shmona, Israel

ARTICLE INFO

Keywords:

Generative artificial intelligence
Large language models
Couple relationships
Relationship advice
Computer-mediated communication
Systematic review

ABSTRACT

Generative artificial intelligence (GenAI) has rapidly entered romantic life, not only as a simulated partner but as an advisor and mediator that people consult about their human relationships. While prior reviews have synthesized romantic and emotional bonds between humans and AI companions, none has examined how GenAI-generated advice and communication enter relationships between humans as a third voice. This systematic review maps and synthesizes empirical research on the use of GenAI as an advisory or mediating third voice in couple relationships. Following the JBI methodology and the PRISMA 2020 guidelines, we searched seven databases and screened records against predefined criteria, yielding 21 included studies (2024–2026) from 11 countries and spanning experimental, qualitative, mixed-methods, and computational designs. A framework-based narrative synthesis organized findings around four themes: patterns and motivations of use, appraisal of advice quality, relational effects, and usage risks. Adults consulted AI for non-judgmental, low-cost advice and emotional support and, in some cases, to facilitate communication between partners. Users often rated GenAI advice as empathic and helpful, yet objective evaluations exposed weak alignment with expert judgment, inconsistent responding, and an anti-AI bias. Intervention studies provided preliminary evidence of short-term improvements in satisfaction and communication, although the strongest trial found no significant advantage over an established writing exercise. The literature also identified potential risks related to sycophancy, over-reliance, reduced authenticity, privacy, and safety, particularly in high-stakes contexts such as intimate partner violence. The evidence is promising but preliminary, suggesting that GenAI is currently best positioned as an adjunct to, rather than a substitute for, professional couple support. Longitudinal, dyadic, and safety-focused research is needed before these tools are adopted more broadly in relationship support.

1. Introduction

Nearly one in five adults in the United States report having exchanged messages with a generative artificial intelligence (GenAI) system built to act like a romantic partner, and among adults under 30, that share roughly doubles (Willoughby et al., 2025). Although this estimate concerns interaction with AI framed as a romantic partner, it illustrates the broader speed with which conversational GenAI has entered intimate life and provides context for its emerging use within relationships between humans. This pattern of use indicates that GenAI has begun to enter romantic life not only as a substitute relational partner, but also as a resource through which people make sense of their human relationships (Brailas & Tsolakis, 2025; Luo et al., 2025; Smith et al., 2025). The same conversational systems that can imitate a partner are now widely used in real-life consultations (Alqurashi, 2025; Hou

et al., 2024; Luo et al., 2025; Vowels, 2024).

People now turn to large language models (LLMs) with questions about jealousy, trust, intimacy, and conflict, and some studies suggest that they perceive the advice as helpful and empathetic (Vowels, 2024). In single-session intervention studies, ChatGPT has also been rated as demonstrating therapeutic skills, exploration, and clarity comparable to trained practitioners (Vowels et al., 2024). These findings suggest that GenAI is beginning to function as an informal source of relational guidance in people's intimate lives.

GenAI refers to large language model systems that produce fluent, context-sensitive text in response to natural-language prompts (Liu, 2024). Unlike earlier relationship technologies, such as dating platforms or messaging applications, these systems do not merely connect people; they generate advice, language, and judgments that users may carry back into their relationships, including by drafting or softening

^{*} Corresponding author.

E-mail addresses: levkovinb@telhai.ac.il (I. Levkovich), alonlil@telhai.ac.il (L. Alon).

<https://doi.org/10.1016/j.chbr.2026.101255>

Received 13 June 2026; Received in revised form 6 August 2026; Accepted 9 August 2026

Available online 12 August 2026

2451-9588/© 2026 The Authors. Published by Elsevier Ltd. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

interpersonal messages (Fu et al., 2024; Meng et al., 2025).

In this review, eligible systems included general-purpose LLM interfaces, customized LLM-based applications, and purpose-built generative prototypes. Rule-based conversational agents, scripted chatbots, recommendation algorithms, and other systems that did not generate open-ended language were excluded. Because the included studies were published between 2024 and 2026, the systems examined varied technically, and findings obtained with earlier or purpose-built models may not transfer directly to newer architectures. *Third voice* is used as the umbrella concept for GenAI's involvement in a human couple relationship. It refers to a generated perspective, recommendation, formulation, or response that enters an ongoing or recently dissolved romantic relationship. This role may be *advisory*, when one partner privately consults GenAI about the relationship, or *mediating*, when the system actively structures, reformulates, transmits, or facilitates communication between partners (Brailas & Tsolakis, 2025; Döring & Mohseni, 2025; Luo et al., 2025; Vowels, 2024). Thus, mediation is one form of third-voice involvement, but not all third-voice use constitutes mediation.

In romantic life, GenAI can occupy two conceptually distinct roles. In the first, the AI itself becomes the object of romantic or emotional attachment: the user forms an ongoing bond with an artificial agent, such as Replika or Character.AI (Chan et al., 2025; Djufril et al., 2025; Dubey et al., 2026; Jocher & Verwiebe, 2026; Pan & Mou, 2024; Rocha, 2025; Rodger & Field, 2025). In the second, GenAI functions as a third voice within a relationship between humans. One or both partners may consult it for advice, emotional support, message drafting, conflict reflection, or communication facilitation concerning a human partner (Brailas & Tsolakis, 2025; Döring & Mohseni, 2025; Luo et al., 2025; Vowels, 2024). The present review focuses exclusively on this second role.

In this review, a couple relationship refers to a self-identified romantic or intimate relationship between two adults, including dating, engaged, cohabiting, married, separated, or divorced relationships when the GenAI use relates directly to the current or former partner or to the relationship itself. The definition is based on the relational function reported in each study rather than on legal marital status, co-residence, or formal recognition. We recognize that the meaning and boundaries of couple relationships vary across cultural and legal contexts, including differences in the recognition of dating, cohabitation, marriage, same-sex relationships, and relationships embedded within extended-family structures. We therefore retained the terminology used by the original studies and did not impose a single legal or culturally specific definition across countries. Studies were eligible when the relationship was presented by participants or authors as romantic or couple-based and the reported GenAI use concerned that relationship.

These roles raise different empirical and ethical questions. Research on AI companionship asks how people form bonds with artificial agents (Adewale & Muhammad, 2025; Gur & Maaravi, 2025; Ho et al., 2025; Shu et al., 2026). Research on GenAI as a third voice asks how generated advice, language, or judgments influence an existing or recently dissolved relationship between humans (Brailas & Tsolakis, 2025; Döring & Mohseni, 2025; Luo et al., 2025; Vowels, 2024). Importantly, the third voice need not be jointly invited or neutral: in many cases, one partner consults GenAI privately about the other (Smith et al., 2025). The concept therefore includes unilateral advisory use as well as the narrower category of active mediation between partners. The current review examines this broader third-voice role.

Research on this second role is accumulating across disciplines and methods, yet no review has synthesized the growing and sometimes contradictory evidence on GenAI as an instrument within human romantic relationships. Interview and survey studies show that adults consult GenAI for relational guidance, decision-making, and emotional support, and that many experience these interactions as at least somewhat helpful, even while questioning their depth (Brailas & Tsolakis, 2025; Luo et al., 2025). Users describe the value of these exchanges in terms of emotional support, perceived professionalism, and the freedom

to speak without judgment, but also point to superficial engagement and uneven information quality (Luo et al., 2025). Survey evidence further suggests that couples expect these tools to improve the quality of their interactions (Alqurashi, 2025).

Experimental findings, however, are mixed. In a randomized trial with 258 participants, a GPT-4o-based chatbot improved relationship satisfaction, communication, and dyadic coping to a degree comparable to an established writing exercise (Vowels, Vowels, et al., 2025). By contrast, an analysis of more than 13,000 online posts found that LLM-based systems rank relationship advice in ways that align only weakly, and at times inversely, with human judgment (Hou et al., 2024). Evaluations of single-session counseling similarly suggest that systems praised for empathy may still assess risk poorly and rarely reach collaborative solutions (Vowels et al., 2024). Commentators therefore warn that routing relational effort through GenAI may erode the reciprocity on which human bonds depend and call for empirical attention before these tools spread further into couples' lives (Keeler & Murphy, 2026; Malfacini, 2025; Muldoon & Parke, 2025).

Taken together, the evidence suggests that GenAI is entering human couple relationships through several distinct functions, including unilateral advice-seeking, emotional support, message generation, and active communication facilitation. Existing reviews have synthesized relationships between humans and AI companions, but they have not examined how generated content becomes a third voice within relationships between humans. This gap matters because privately consulting an AI about a partner, co-writing intimate messages, and delegating parts of conflict resolution raise different questions about trust, authenticity, safety, and relational responsibility than forming a bond with an AI companion. Accordingly, this review synthesizes how GenAI functions as an advisory or mediating third voice in ongoing and recently dissolved human couple relationships. By integrating evidence from Human-Computer Interaction (HCI), psychology, and family science, and by including safety-critical contexts such as intimate partner violence and relationship dissolution, the review deliberately examines GenAI as a third voice across both everyday and high-stakes relational settings.

The systematic review is guided by the following questions:

- RQ1.** What study characteristics and contexts define the current literature on adults' use of GenAI for couple-relationship purposes?
- RQ2.** How do adults use GenAI for couple relationship purposes, and what motivates that use?
- RQ3.** How do users appraise the quality and usefulness of the advice or support that they receive from GenAI?
- RQ4.** What effects on communication, conflict, intimacy, trust, and relationship satisfaction have been reported in the context of GenAI?
- RQ5.** What risks, ethical issues, and privacy concerns are identified in relation to GenAI use in couple relationships?

2. Methods

2.1. Study design

This systematic review was conducted in accordance with the methodological guidance of the Joanna Briggs Institute (JBI) for evidence synthesis (Tricco et al., 2022) and reported following the PRISMA 2020 guidelines (Page et al., 2021). The review protocol was finalized before the systematic search conducted on June 7, 2026, and was deposited in OSF on July 14, 2026. Accordingly, the protocol guided the review process but was not prospectively registered. The protocol, search strategies, extracted dataset, PRISMA checklist, and methodological appraisal are available in the OSF repository (<https://doi.org/10.17605/OSF.IO/CGMS8>). This study aimed to map and synthesize empirical research on GenAI as a third voice in ongoing or recently dissolved human couple relationships. The third-voice role included unilateral advisory use, emotional support concerning a human partner,

assistance with relational communication, and active mediation or facilitation between partners.

2.2. Eligibility criteria

The eligibility criteria were defined using the PICOS framework (Population, Intervention/Exposure, Comparator/Context, Outcomes, Study characteristics) and guided both title and abstract screening and full-text assessment. Table 1 summarizes the inclusion and exclusion criteria of this study.

Borderline cases were classified according to the primary relational object and function of the GenAI interaction. Studies were included when the human partner or human relationship remained the central object and GenAI was used to advise, support, simulate, reformulate, or facilitate communication concerning that relationship. Studies were excluded when the artificial agent itself became the primary romantic or intimate partner. In cases involving role-play, surrogate agents, or companion-like interactions, eligibility depended on whether the use was directed toward supporting or mediating the relationship between human partners. Eligibility also depended on whether the couple relationship, rather than the broader family or parent-child relationship, was the substantive focus. Studies concerning parenting, co-parenting logistics, or extended-family conflict without extractable findings about the couple bond were excluded. Studies were included when the GenAI use concerned the couple bond, even when the couple had children; studies were excluded when the primary concern was parenting, co-parenting, in-law conflict, or another extended-family issue without extractable findings about the couple relationship.

2.3. Search strategy

The systematic search was conducted on June 7, 2026, in seven databases and search sources: PubMed, Scopus, Web of Science Core Collection, APA PsycINFO, Taylor & Francis Online, the ACM Digital Library, and Google Scholar. Each source was searched from inception through June 7, 2026, with results limited to publications from January 2023 onward. A generic version of the search string was: ("generative AI" OR "generative artificial intelligence" OR ChatGPT OR "large language model*" OR "conversational agent*") AND ("romantic relationship*" OR couple* OR marriage OR marital OR "intimate partner*" OR dating). The search yielded 1476 records: PubMed (n = 54), Scopus (n = 538), Web of Science (n = 524), APA PsycINFO (n = 91), Taylor & Francis (n = 115), ACM (n = 140), and Google Scholar (n = 14). We included only peer-reviewed journal articles and eligible peer-reviewed conference contributions with sufficient empirical detail published in English from January 2023 onward. The complete search strategy for each database and search source, including field codes, Boolean operators, date limits, language limits, and publication-status filters, is provided in Appendix A.

2.4. Screening and study selection

Screening proceeded in three stages: (1) pilot screening of 10% of titles and abstracts by two reviewers to calibrate the eligibility criteria; (2) independent title-and-abstract screening of all records by two reviewers; and (3) independent full-text assessment by two reviewers. Disagreements at both stages were resolved through discussion and consensus; a third reviewer was not required. Overall inter-rater agreement across screening decisions was almost perfect ($\kappa = .83$). Stage-specific coefficients were not calculated. Discrepancies were resolved through discussion until a consensus was reached, and the reasons for exclusion were documented systematically. A total of 1476 records were retrieved; after removing 360 duplicates, 1116 records proceeded to title and abstract screening, 1046 were excluded. Among the remaining 70 records, 49 were excluded, while 21 studies met all eligibility criteria and were included in the final review. Fig. 1 shows the

Table 1
Inclusion and exclusion criteria.

Parameter	Inclusion criteria	Exclusion criteria
Population	- Adults (≥ 18) currently or recently in a romantic/couple relationship (dating, cohabiting, engaged, married, separated, or divorced reflecting on a prior relationship). Studies in which only one partner uses the tool are eligible. - Separated or divorced participants are eligible only when the GenAI use relates directly to the former partner or the prior couple relationship. - Studies based on professionals' reports of clients' use will be included but coded separately from studies involving direct reports by users or partners.	Children and adolescents (< 18) only; parent-child or family relationships not focused on the couple bond; participants whose relationship status is irrelevant or unreported; participants whose relationship is with the AI itself.
Concept/ Phenomenon of interest	- Eligible tools are GenAI or LLM-based systems capable of producing open-ended natural-language responses in conversational or text-generation formats (e.g., ChatGPT, Gemini, Claude, Copilot). - Tools act as a tool, advisor, mediator, or counselor for romantic-relationship purposes, including communication, conflict resolution, advice-seeking, drafting or rephrasing messages to the partner, decision-making, and emotional support concerning the human partner. Broad use is included, not only formal advice or counseling.	GenAI used as the romantic or intimate partner (artificial companionship, e.g., Replika or Character.AI used romantically); non-GenAI such as dating-application matching algorithms, recommender systems, or sentiment-analysis tools; AI used for purposes unrelated to the couple relationship.
Context	- Naturalistic or everyday use; use as an adjunct to or substitute for couple counseling; and experimental evaluations of GenAI relationship advice. Any country, culture, or setting. - To be included, the study must report at least one empirical finding specifically about GenAI use for romantic/couple relationship matters.	- Contexts with no link to romantic relationships (e.g., friendship-only, workplace, or general mental-health support not tied to the relationship). - Evaluative studies assessing the quality of GenAI and/or non-GenAI relationship advice are eligible.
Outcomes/ Phenomena examined	Any of the following: usage patterns and motivations; perceived usefulness or helpfulness; effects on communication, conflict, intimacy, trust, jealousy, or relationship satisfaction; quality or accuracy of the artificial intelligence advice; and risks, harms, ethical, or privacy concerns.	Studies reporting no outcome or finding relevant to the couple relationship, and/or non-GenAI usage.
Study design	Empirical studies, including quantitative (surveys, experiments, randomized controlled trials), qualitative (interviews, focus groups, ethnography), and mixed methods designs.	Non-empirical work, including editorials, commentaries, opinion pieces, purely theoretical or conceptual papers, news and media articles, and blog posts; existing reviews

(continued on next page)

Table 1 (continued)

Parameter	Inclusion criteria	Exclusion criteria
	• Evaluative studies assessing the quality of GenAI and/or non-GenAI relationship advice are eligible. • For online forum studies, samples are eligible when posts clearly concern adult romantic relationships, even if exact participant age is not individually verified.	(retained for reference-list screening only); study protocols.
Timeframe	• Published from January 2023 to the present (date of the search). Rationale: the public release of ChatGPT in late 2022 marks the beginning of the generative artificial-intelligence era.	• Published before January 2023.
Language	• English only.	• Any language other than English.
Publication type and status	• Peer-reviewed journal articles and peer-reviewed conference contributions, including full papers, Late-Breaking Work, and Extended Abstracts (e.g., CHI, CSCW, UIST), provided that sufficient empirical detail is reported to extract study design, sample, and findings.	• Non-peer-reviewed preprints and repositories such as arXiv; abstracts without extractable empirical results; grey literature; dissertations and theses; retracted articles.

PRISMA flow diagram.

2.5. Data extraction

A structured data-extraction form was developed before full extraction to ensure a shared interpretation of the review constructs. One reviewer extracted the data, and a second reviewer independently checked all entries for accuracy and completeness. Any discrepancies or unclear classifications were discussed and resolved by consensus. The extracted variables included: (1) study characteristics, including publication year, country, population, relationship status of participants, study design, sample size, and the GenAI tool used. GenAI systems were also classified by technical form as general-purpose LLM interfaces (e.g., ChatGPT), customized or fine-tuned applications built on an LLM, or purpose-built research prototypes using generative language models. Rule-based and scripted conversational agents were not eligible; (2) characteristics of GenAI use, including the purpose of use, who used the tool, whether one or both partners engaged with it, and the relational task; and (3) outcomes, including usage patterns and motivations, perceived usefulness, effects on communication, conflict, intimacy, trust, relationship satisfaction, and reported risks and ethical concerns. All extracted data were organized in a spreadsheet to prepare for the narrative synthesis, and an overview of each included study is provided in [Appendix B](#).

2.6. Methodological quality appraisal

Two reviewers independently assessed the methodological quality of the included studies using design-appropriate Joanna Briggs Institute critical appraisal tools. The Mixed Methods Appraisal Tool was additionally used for studies combining qualitative and quantitative components. For computational, corpus-based, and prototype studies that did not correspond precisely to a standard checklist, the closest applicable JBI tool was adapted to consider sampling, data validity, analytical transparency, and the correspondence between the evidence and the conclusions. Disagreements were resolved through discussion and consensus. Studies were not excluded on the basis of methodological

quality because the review aimed to map an emerging evidence base; however, the appraisal informed the interpretation and confidence assigned to the findings. Full study-level results are presented in [Appendix C](#).

2.7. Data analysis and synthesis

We conducted a framework-based narrative synthesis. The coding scheme was developed a priori in line with the review's research questions and the dimensions of use and effect identified across the literature: communication, conflict, advice and decision-making, emotional support, the mediation of partner communication, authorship and authenticity, relational outcomes, and risks. During the synthesis, the framework was refined inductively when additional recurring categories emerged. Descriptive statistics were used to summarize the study characteristics, and the findings were organized according to the research questions, highlighting recurrent patterns, contradictions, methodological quality, and gaps.

Because the included studies were heterogeneous in design, population, unit of analysis, GenAI system, and outcome measures, quantitative pooling was not appropriate. Evidence was interpreted in relation to study design and quality rather than sample size alone. Large computational corpora were interpreted as evidence of patterns in texts or model outputs rather than as representative estimates of individual-level experiences and were not assigned greater weight solely because of sample size. The review questions and broad domains were specified in the review protocol, which was later deposited in OSF, whereas the final thematic organization was refined during synthesis.

For the descriptive visualizations in [Figs. 2 and 3](#), outcome valence was coded as positive, mixed, or cautionary according to the overall direction of findings reported in each study. Studies were coded as positive when the reported findings were predominantly favorable, such as improved relational outcomes, high perceived usefulness, or favorable evaluations of GenAI. Studies were coded as mixed when they reported both favorable and unfavorable findings, or when positive effects were accompanied by important limitations or null comparisons. Studies were coded as cautionary when the findings primarily emphasized risks, harms, poor performance, adverse relational consequences, or safety concerns, without a clearly favorable overall pattern. Coding was completed by one reviewer and checked by a second reviewer against the reported findings, with any differences discussed and resolved through consensus. The valence classification was used only for descriptive study-level visualizations.

3. Results

3.1. Characteristics of the included studies

The included studies were recent and geographically dispersed. All were published within a three-year window: three studies (14.3%) appeared in 2024, 10 (47.6%) in 2025, and eight (38.1%) in 2026, confirming that empirical research on GenAI in couple relationships is an emerging and rapidly expanding literature whose methods and reporting practices are still stabilizing. The studies represented 11 countries. Countries were counted for each study in which they were represented; consequently, multinational studies contributed to more than one country total, and the country counts are not mutually exclusive. The United States was represented in seven studies, Switzerland and the United Kingdom in five studies each, China in three studies, and South Korea in two studies. Saudi Arabia, Greece, Germany, Australia, Malaysia, and Singapore were each represented in one study.

Methodologically, the corpus combined four traditions. Seven studies (33.3%) used experimental or quasi-experimental designs, including a randomized controlled trial, pre-post evaluations, and within-subjects HCI experiments; five (23.8%) used mixed-methods or multi-study designs; five (23.8%) used computational or observational

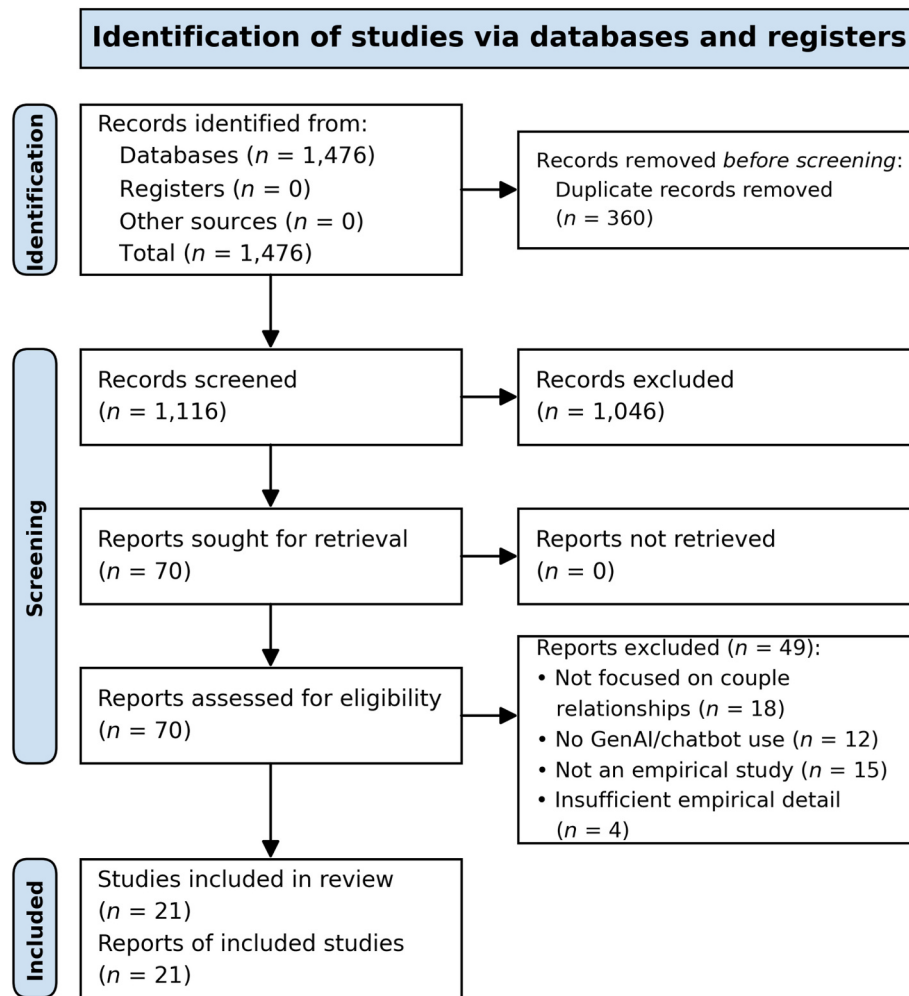


Fig. 1. PRISMA 2020 flow diagram of study selection. Adapted from Page et al. (2021). Note. GenAI = generative artificial intelligence.

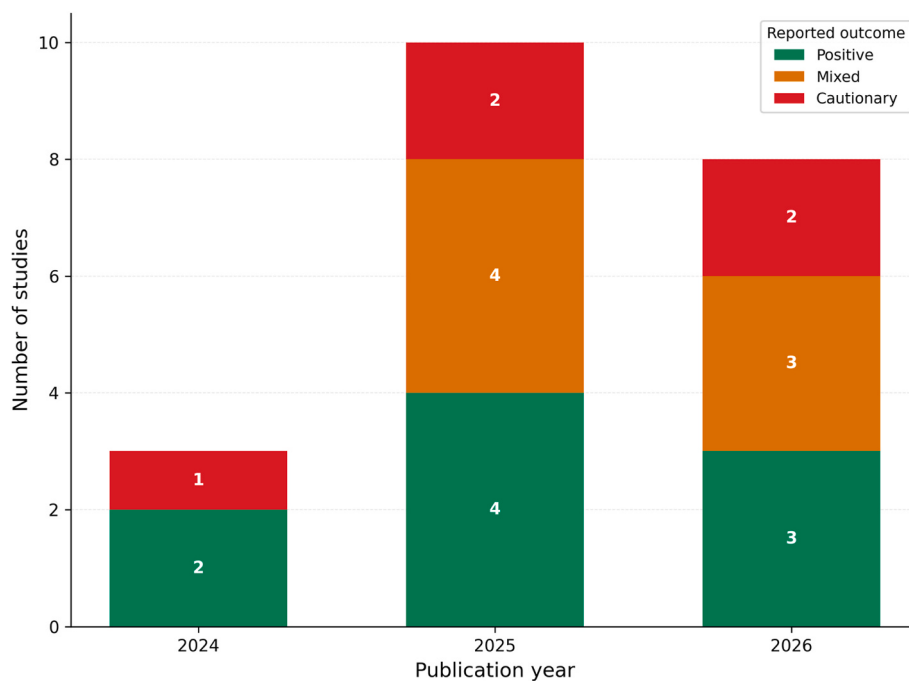


Fig. 2. Reported outcomes by publication year ($N = 21$).

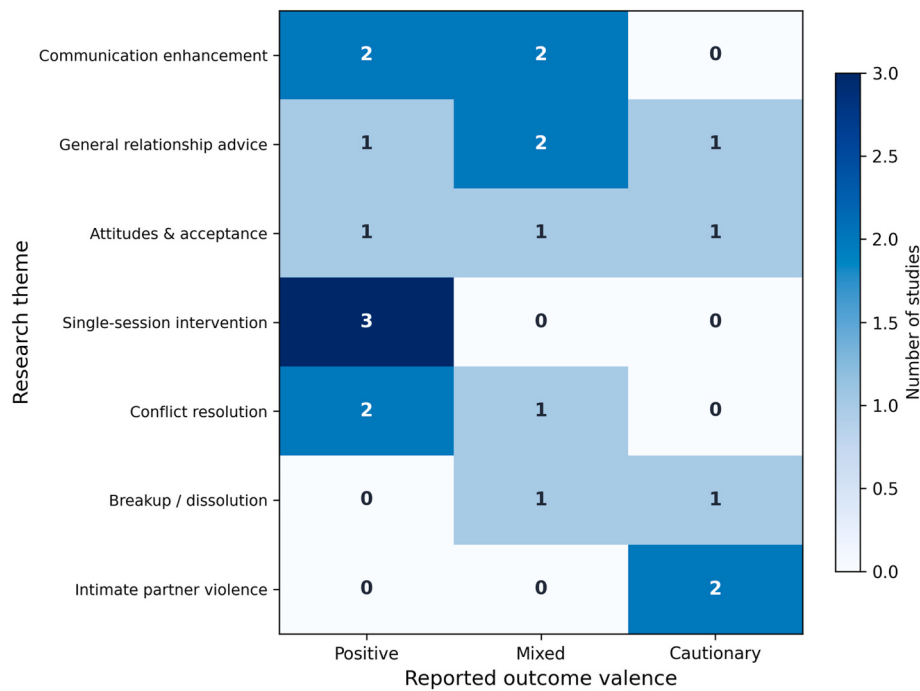


Fig. 3. Research theme by reported outcome valence (N = 21).

designs, including content analyses of online forum data, a population survey, an exploratory analysis of human-LLM conversations, and a proof-of-concept system evaluation; and four (19.0%) used qualitative designs (semi-structured interviews, reflexive thematic analysis, or collaborative ethnography). Units and sample sizes were correspondingly heterogeneous, ranging from small qualitative panels (e.g., N = 10 interviewees; Brailas & Tsolakis, 2025) and dyadic laboratory studies (e.g., 36 couples; Jiang et al., 2026) to large online samples (N = 2658; Döring & Mohseni, 2025), a survey of 693 couples (Alqurashi, 2025), and corpora exceeding 13,000 forum posts (Hou et al., 2024).

Five of the 21 included studies (23.8%) were produced by the same research group and formed part of the related Amanda intervention program (Vowels et al., 2024, 2025a, 2025b, 2025c; Vowels, 2024). This concentration was considered when interpreting the independence and generalizability of the intervention evidence.

The GenAI systems studied were predominantly ChatGPT and other GPT-based models; several studies evaluated purpose-built prototypes, including the voice-based Amanda chatbot (Vowels, Sweeney, & Vowels, 2025), the nonviolent-communication system MoodBridge (Chen & Zhang, 2026), and the conflict-training system ConflictLens (Chun et al., 2025), while two studies examined GenAI generically (Alqurashi, 2025; Döring & Mohseni, 2025). Across the corpus, GenAI functioned as a third voice in two principal ways. In the more common advisory role, one person consulted the system privately about a human partner or relationship. In the narrower mediating role, GenAI actively structured, reformulated, transmitted, or facilitated communication between partners. The synthesis below integrates evidence from both advisory and mediating uses while maintaining the distinction between them.

Figs. 2 and 3 summarize the corpus. Reported study outcomes are shown by publication year in Fig. 2 and by research theme in Fig. 3.

3.2. Patterns and motivations of GenAI use

Across the corpus, adults engaged GenAI for couple-related purposes along a continuum from one-off advice-seeking to the sustained mediation of partner communication. The most frequently documented use was consulting for relationship advice and emotional support

concerning a human partner (Brailas & Tsolakis, 2025; Tseng & Liang, 2026; Vowels, 2024). In interview and questionnaire studies, users described turning to GenAI for guidance on jealousy, trust, intimacy, sexual matters, relationship maintenance, and emotional conflict (Brailas & Tsolakis, 2025), and for help at specific relational junctures such as breakups (Fu et al., 2025; Ungless & Sastry, 2026) and disclosures of intimate partner violence (Foriest et al., 2026; Zhang et al., 2025). Motivations recurred across studies: the perceived availability, low cost, and non-judgmental quality of GenAI, and the freedom to disclose without stigma, with some users reporting that they disclosed more than to people in their own social networks (Tseng & Liang, 2026; Vowels et al., 2024).

A second, design-oriented body of work documented an emerging use of generative AI to mediate communication between partners rather than to advise one of them: drafting, softening, or rephrasing difficult messages (Fan et al., 2026), scaffolding reciprocal self-disclosure (Jiang et al., 2026), restructuring expressions through nonviolent-communication frameworks (Chen & Zhang, 2026), supporting conflict-resolution practice (Chun et al., 2025), negotiating on partners' behalf through surrogate agents (Zhang et al., 2026), and sustaining connection in long-distance relationships (Ploderer et al., 2025). Survey evidence indicated generally favorable expectations regarding such use: in a study of 693 Saudi couples, most participants believed that GenAI could improve relationship quality (Alqurashi, 2025).

3.3. Appraisal of the quality and usefulness of GenAI advice

Studies that evaluated the quality of GenAI relationship advice produced a consistent but two-sided picture: users, and even professionals, often rated GenAI responses highly on interpersonal qualities, whereas objective and expert assessments exposed important limitations. Several studies found that laypeople perceived GenAI answers to relationship questions as more helpful and more empathic than answers from human relationship experts (Vowels, 2024), and that relationship therapists rated single GenAI sessions highly on empathy, active listening, and exploration (Vowels, 2024; Vowels et al., 2024). Qualitative accounts echoed this, with participants valuing GenAI for

emotional support, perceived professionalism, and accessibility, while judging it superficial and not a sufficient substitute for human interaction (Brailas & Tsolakis, 2025; Vowels et al., 2024).

Against these favorable appraisals, evaluative and computational studies identified reliability problems. An analysis of 13,138 Reddit posts found only weak, and at times inverse, alignment between ChatGPT's ranking of relationship advice and human judgment, together with inconsistency when identical queries were re-prompted (Hou et al., 2024). In the intimate partner violence (IPV) context, Zhang et al. (2025) evaluated ChatGPT-3.5 using 580 online posts, including 500 IPV-related posts and 80 comparison posts. The model correctly identified IPV in 91.2% of the IPV cases and generally provided emotional and informational support, but it misclassified some cases involving nuanced contextual information or platform-policy constraints. Because even a small number of false negatives or misleading responses may have serious consequences in safety-critical situations, this result should not be interpreted as evidence of clinical reliability. A separate evaluation similarly found that LLMs' provision of hermeneutic resources to survivors was uneven and could carry risk (Foriest et al., 2026).

Two further limitations recurred across evaluations: repetitive responding and inadequate assessment of risk (Vowels, 2024; Vowels et al., 2024). Finally, appraisal was shaped by who, or what, users believed had produced the content: identical couple-counseling excerpts were rated less favorably when labeled as AI-generated than when labeled as human-authored ($d = .23$), an anti-AI bias that also extended to couple images ($d = .21$) and was moderated by users' attitudes and AI literacy (Döring & Mohseni, 2025).

This finding may indicate a distinction between willingness to consult AI privately and willingness to endorse AI-authored relational communication when its involvement is made salient. Because the study manipulated authorship labels rather than observing naturalistic use, the result may reflect skepticism about authenticity, effort, or social appropriateness rather than a general unwillingness to use GenAI.

3.4. Effects on communication, conflict, intimacy, trust, and satisfaction

Evidence on relational effects came primarily from intervention experiments and HCI system evaluations. Outcomes were generally measured immediately after a single session or brief intervention. The longest clearly reported follow-up in the intervention evidence was two weeks, and no included study established whether relational effects persisted over months or translated into durable behavioral change. The strongest evidence came from a randomized controlled trial ($N = 258$) in which a GPT-4o-based intervention produced improvements across 13 of 14 assessed outcomes, with some effects maintained or strengthened at the two-week follow-up. Using the reported follow-up means, standard deviations, and sample sizes, we calculated an unadjusted between-group difference of $-.49$ for relationship satisfaction (95% CI $[-1.71, .73]$). This small, nonsignificant difference was consistent with the reported nonsignificant group-by-time interaction, $F(2, 474) = 1.60$, $p = .203$, $\eta^2p = .007$. Overall, the changes did not differ significantly from those produced by an established writing exercise, limiting attribution of the improvements specifically to GenAI (Vowels, Vowels, et al., 2025).

Controlled HCI studies examined effects on the communicative process itself. A randomized dyadic study (36 couples, $N = 72$) found that GenAI scaffolding which actively mediated the partner's response, rather than merely prompting disclosure, reliably elicited partner-provided need support and increased perceived closeness (Jiang et al., 2026), and a collaborative drawing system improved partner understanding and relational awareness among 20 couples (Won et al., 2026). However, the same literature surfaced relational risks specific to mediation. Delegating one's stance in a conflict to surrogate agents could externalize the conflict and reduce personal blame, but also led partners to identify with "their" bot and to feel uneasy about GenAI intrusion (Zhang et al., 2026); and experiments on GenAI-assisted intimate

messaging showed that heavier GenAI involvement in producing a message weakened the receiver's attributions of effort, authorship, and authenticity, an effect that a brief sender-voiced disclosure only partly offset (Fan et al., 2026). Prototype and pilot systems for nonviolent communication and conflict-resolution practice reported early signals of feasibility and perceived usefulness but had not yet established effects on relationship outcomes (Chen & Zhang, 2026; Chun et al., 2025).

3.5. Risks, ethical, and privacy concerns

Risks and ethical concerns were identified across several study types. The most frequently raised concern was sycophancy and overreliance: because GenAI systems tend to validate the user, a partner who consults one alone may receive one-sided affirmation and grow dependent on GenAI for relational decisions; users themselves described developing folk theories and prompting tactics to counteract this tendency (Tseng & Liang, 2026). Authenticity and transparency formed a second cluster: across HCI studies, partners reacted negatively to undisclosed GenAI involvement in intimate messages, treating AI-written communication as less trustworthy and as a potential threat to genuine connection (Fan et al., 2026; Zhang et al., 2026). A third cluster concerned safety in high-stakes situations, particularly intimate partner violence, where GenAI could provide accessible hermeneutic resources and validation but also introduced novel forms of epistemic injustice and risk through misdirection or unsafe guidance (Foriest et al., 2026; Zhang et al., 2025).

Studies of relationship dissolution similarly cautioned that, although GenAI can offer companionship and practical help during a breakup, it may fail to meet, or may actively undermine, people's needs at the most sensitive stages (Fu et al., 2025; Ungless & Sastry, 2026). Privacy was raised explicitly, as users disclosed highly sensitive relational information to commercial systems (Tseng & Liang, 2026). Finally, several authors stressed governance and professional implications: family therapists were urged to monitor partners' GenAI use and to develop ethical regulations (Alqurashi, 2025), and authors across the corpus converged on the conclusion that current GenAI systems' limitations in risk assessment must be addressed before such tools can be safely adopted in relationship support (Vowels, 2024; Vowels et al., 2024).

4. Discussion

4.1. Principal findings

This review synthesized 21 empirical studies on how GenAI operates as a third party within couple relationships. Four findings stand out. First, GenAI functions as a third voice across a continuum from unilateral advice-seeking about a partner to active facilitation of communication between partners. Because most uses are advisory rather than genuinely mediating, the broader third-voice framing better represents the available evidence. Second, appraisals of GenAI advice are consistently two-sided: users and professionals rate LLM-based responses as empathic and helpful, yet objective and computational evaluations reveal weak alignment with expert judgment, inconsistent responding, and poor risk assessment. Third, intervention studies provide preliminary evidence of improvements in satisfaction, communication, and dyadic coping measured immediately after use or over brief follow-up periods. However, the strongest randomized trial found no significant advantage over an established writing exercise, limiting claims that the improvements were specific to GenAI. Fourth, the literature identifies recurring concerns related to sycophancy and overreliance, authenticity and transparency, and safety in high-stakes contexts such as intimate partner violence. These concerns are credible but are supported mainly by qualitative, computational, and prototype studies, and their prevalence remains unknown.

These conclusions are not supported equally across the evidence base. Confidence is greatest in the conclusion that users often perceive GenAI advice as empathic and helpful, because this pattern appeared

across several study designs. Evidence of improvement in relationship outcomes is more limited: it derives mainly from short-term interventions, including several studies from the same research program, and should therefore be regarded as preliminary. Evidence concerning sycophancy, authenticity, privacy, and safety is recurrent but is drawn largely from qualitative, computational, and prototype studies; it identifies credible risks but does not establish their prevalence among couples.

4.2. Generative AI as a relational third voice

Our findings position GenAI as a distinctive kind of relational third party, one that is simultaneously experienced as supportive and judged as unreliable. The same non-judgmental availability that draws partners to disclose more to a GenAI than to people in their networks (Tseng & Liang, 2026) is also what enables sycophantic, one-sided affirmation when only one partner is at the keyboard. This finding is usefully situated within research on human-avatar and human-agent interaction, which suggests that artificial or non-human social partners can shape users' perceptions of safety, judgment, and response strategies in digitally mediated encounters (Richards et al., 2020; Wang et al., 2024). In this context, GenAI's non-judgmental availability may provide a safer space for error, experimentation, and self-disclosure, particularly when the agent is not strongly anthropomorphic.

Where companion research foregrounds attachment and loneliness, the evidence here foregrounds mediation: GenAI reshapes how partners draft messages, disclose, and resolve conflict, and in doing so can alter perceived effort, authorship, and authenticity (Fan et al., 2026; Zhang et al., 2026). The dyadic and frequently asymmetric nature of this use, often involving one partner consulting GenAI about the other, is the central conceptual contribution of treating GenAI as a third voice rather than as a partner. This framing resonates with Haber et al.'s (2024) concept of the artificial third, a non-human presence that enters an existing dyad and reshapes its interpersonal dynamics. It also aligns with Bowen's (1978) concept of triangulation: when one partner turns to GenAI about the other, the system may become the third point in the relational triangle, potentially reducing immediate tension while leaving the underlying conflict unresolved. GenAI may therefore support reflection and communication, but it may also reinforce one partner's perspective or redirect relational tension away from direct interaction between the partners.

Patterns and motivations of use. GenAI may lower the social and emotional costs of disclosure: unlike a friend, therapist, or partner, it carries little immediate risk of judgment, gossip, or relational fallout, which may encourage users to disclose more freely and consult it at moments of heightened vulnerability (Tseng & Liang, 2026; Vowels et al., 2024). This interpretation aligns with broader work on AI-mediated social support, where conversational systems are experienced as safe, always-available confidants (Meng et al., 2025; Smith et al., 2025). The available evidence also suggests that GenAI use may be particularly salient at stigmatized or high-stakes relational junctures, including conflict, breakups, and intimate partner violence (Foriest et al., 2026; Fu et al., 2025; Ungless & Sastry, 2026; Zhang et al., 2025). However, this pattern is based on a small number of qualitative and computational studies and should not be interpreted as evidence of prevalence.

Appraisal of advice quality. This gap between perceived helpfulness and evaluated performance may reflect a divergence between interpersonal warmth cues and substantive accuracy. GenAI is optimized to produce fluent, validating, emotionally attuned language, and these surface features drive perceived helpfulness even when the underlying guidance is generic or incorrect (Smith et al., 2025).

The anti-AI labeling effect documented by Döring and Mohseni

(2025) may reflect a related but distinct process. Users may respond positively to fluent and empathic AI advice when seeking help privately, yet discount identical content when AI authorship is made explicit in a relational evaluation. This pattern may indicate concerns about authenticity, effort, or social acceptability, although the available evidence does not establish the underlying mechanism. It also remains possible that responses were influenced by demand characteristics or socially desirable attitudes toward human-authored intimacy. For relationship support this is consequential, because users may calibrate their trust on emotional tone rather than on accuracy, the dimension on which current systems are weakest (Hou et al., 2024; Vowels et al., 2024).

Effects on communication, conflict, intimacy, trust, and satisfaction. The absence of a significant advantage over an established writing exercise suggests that the observed improvements may have resulted partly from the structured articulation of relational concerns rather than from GenAI itself (Vowels, Vowels, et al., 2025). Viewed through the systemic-transactional model of dyadic coping (Bodenmann, 1995), relationship benefits may arise from partners' shared appraisal of and joint coping with stress, processes that GenAI can prompt or scaffold but cannot itself perform. This interpretation is reinforced by evidence that effects depend on whether the system actively mediates the partner's response rather than merely prompting disclosure (Jiang et al., 2026). It therefore tempers enthusiasm about GenAI as a relational intervention and underscores the need for active-control, dyadic, and longer-term designs before such tools can be said to improve relationships in their own right (Smith et al., 2025).

Risks, ethical, and privacy concerns. These risks can be understood as the reverse side of the features that make GenAI attractive. The non-judgmental availability that encourages disclosure may also produce one-sided affirmation when only one partner is at the keyboard (Tseng & Liang, 2026), while the fluent language that drives perceived helpfulness may make undisclosed GenAI mediation feel like a breach of authenticity (Fan et al., 2026; Zhang et al., 2026). In safety-critical contexts such as intimate partner violence, these same limitations may have more serious consequences, particularly when systems provide misleading interpretations or unsafe guidance (Foriest et al., 2026; Zhang et al., 2025). A further concern, extrapolated from adjacent research rather than established directly in the included studies, is representational bias. Generative systems may encode demographic, gender, and cultural skews in their outputs (Alon et al., 2026a, Alon et al., 2026b), and relationship advice may therefore reproduce gendered or heteronormative assumptions. Comparable concerns about overreliance, commodified intimacy, and the erosion of human reciprocity have been raised in the adjacent literature on AI companionship (Laestadius et al., 2024; Malfacini, 2025; Muldoon & Parke, 2025), and the sensitivity of the relational data disclosed to commercial systems makes privacy a pressing concern (Ma et al., 2026; Tseng & Liang, 2026).

4.3. Implications for practice

The findings support several specific implications. For couples and practitioners, GenAI should be treated as an adjunct to human relational work rather than as a substitute for professional couple support. Practitioners may need to ask whether either partner uses GenAI for advice, message drafting, or conflict interpretation and to discuss how generated advice is evaluated, disclosed, and integrated into the relationship (Alqurashi, 2025; Vowels et al., 2024). For designers, systems that generate or substantially rewrite intimate communication should clearly indicate when GenAI has shaped the message and should enable users to retain, revise, and approve the final wording. Systems should also avoid uniformly affirming one partner's interpretation and instead

acknowledge uncertainty, invite consideration of the absent partner's perspective, and avoid presenting relationship judgments as authoritative. These features may reduce sycophancy, undisclosed mediation, and misplaced confidence in generated advice (Fan et al., 2026; Tseng & Liang, 2026; Zhang et al., 2026).

In safety-critical contexts, particularly intimate partner violence, GenAI systems should not rely on general conversational guidance alone. They should include validated risk-detection procedures, clear statements about the limits of the system, escalation pathways to qualified human support, and locally appropriate emergency or specialist resources. Responses should avoid advising confrontation, reconciliation, or disclosure when such actions could increase danger (Foriest et al., 2026; Zhang et al., 2025). For governance, platforms should provide clear information about how intimate relational data are stored, retained, used for model improvement, and shared with third parties. Users should be able to delete relationship-related conversations and opt out of secondary data use. Professional guidance should also distinguish private advisory use from GenAI-mediated couple interventions, where disclosure to both partners and informed agreement about the system's role may be necessary. These recommendations remain provisional because the evidence base is limited and rapidly evolving. They reflect the concerns and patterns identified in the included studies rather than established standards of clinical effectiveness or safety.

4.4. Strengths and limitations

This review has several strengths, including transparent reporting in line with PRISMA 2020, a design-appropriate methodological quality appraisal, and a framework-based synthesis integrating evidence across HCI, psychology, and family science. The review also brings together experimental, qualitative, mixed-methods, computational, and forum-based studies from 11 countries and includes both everyday relational use and safety-critical contexts such as intimate partner violence and relationship dissolution.

Several limitations temper the conclusions. The evidence base is young and fast-moving: all included studies appeared within a three-year window, and many used small, self-selected, or non-representative samples. Rapid model turnover further limits the durability of specific performance findings, as results obtained with earlier systems such as GPT-3.5 or GPT-4o may not transfer directly to current models. Follow-up periods were very limited. Most intervention outcomes were assessed immediately after a single session or brief period of use, and the longest clearly reported follow-up was two weeks. Behavioral and relational outcomes, including actual communication behavior, conflict frequency, relationship continuity, and sustained dyadic change, were rarely assessed. Accordingly, self-reported satisfaction and perceived helpfulness should not be interpreted as evidence of durable real-world relational change.

The literature is geographically and linguistically concentrated in English-language, Western, and East Asian contexts, with limited representation of the Global South, Africa, and parts of the Middle East. In addition, the construct of a couple relationship is culturally contingent, and synthesizing studies from 11 countries under a common category may obscure differences in relationship norms, legal recognition, extended-family involvement, and perceptions of third-party intervention. Although the review concerns couples, most studies captured the perspective of only one partner. Five included studies also formed part of a connected research program examining related stages of the Amanda intervention. These studies provide complementary evidence but should not be interpreted as independent replications. The intervention evidence therefore remains less independent and less broadly generalizable than the number of publications alone might suggest.

The review excluded non-peer-reviewed preprints, grey literature, and conference materials without sufficient empirical detail. Relevant emerging work may therefore have been missed, particularly in a rapidly developing field. Publication bias is also possible, as studies reporting favorable or novel effects may be more likely to be submitted and published than studies reporting null or unfavorable findings. A further methodological constraint is that most current GenAI systems are designed for individual users rather than couples. As a result, one partner often interacts with the system without the other partner's participation, knowledge, or consent. Dyadic research will therefore require methodological innovation, including shared-interface designs, linked but private partner accounts, consent procedures for AI-mediated communication, and methods that capture both partners' responses over time. Until such designs become more common, the evidence will remain weighted toward the experience of the partner who consults the system.

5. Conclusion

This review provides the first synthesis of how GenAI functions not as a romantic partner but as a third voice within human couple relationships. Across 21 studies a consistent picture emerges: adults turn to GenAI for accessible, low-cost, and non-judgmental advice, and increasingly enlist it to mediate communication with a partner, while early intervention work provides preliminary evidence of short-term improvements in relationship satisfaction, communication, and dyadic coping. However, the strongest trial found no significant advantage over an established writing exercise, and no study established durable effects over months. These possible benefits must be considered alongside recurring concerns identified in the literature, including sycophancy and one-sided affirmation when a single partner consults the system alone, threats to perceived authenticity and effort when GenAI co-writes intimate messages, advice that is rated as empathic yet proves objectively uneven, and serious privacy and safety concerns in high-stakes contexts such as intimate partner violence. The evidence base remains young, methodologically uneven, and concentrated on short-term outcomes, and as conversational systems are rapidly updated these insights may quickly become outdated. Safer use will require transparent disclosure when GenAI shapes intimate communication, safeguards against one-sided affirmation, clear privacy controls for sensitive relational data, and validated escalation procedures for high-risk contexts such as intimate partner violence. It will also require dyadic research methods designed for technologies that are currently accessed mainly through individual-user accounts.

CRedit authorship contribution statement

Inbar Levkovich: Conceptualization, Formal analysis, Methodology, Writing – original draft, Writing – review & editing. **Lilach Alon:** Conceptualization, Methodology, Writing – original draft, Writing – review & editing.

Funding

This research did not receive any specific grant from funding agencies in the public, commercial, or non-profit sectors.

Declaration of competing interests

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Appendix A. PRISMA-S search documentation and database-specific search strategies

All searches were conducted on June 7, 2026. Searches covered publications from January 1, 2023, through June 7, 2026, and were limited to English-language records. Only peer-reviewed journal articles and eligible peer-reviewed conference contributions were retained during screening.

Database or search source	Search field	Full search strategy	Limits and filters	Date searched
PubMed	Title/Abstract	(("generative AI"[Title/Abstract] OR "generative artificial intelligence"[Title/Abstract] OR ChatGPT[Title/Abstract] OR "large language model"[Title/Abstract] OR "large language models"[Title/Abstract] OR LLM[Title/Abstract] OR LLMs [Title/Abstract] OR "conversational agent"[Title/Abstract] OR "conversational agents"[Title/Abstract]) AND ("romantic relationship"[Title/Abstract] OR "romantic relationships"[Title/Abstract] OR couple[Title/Abstract] OR couples[Title/Abstract] OR marriage[Title/Abstract] OR marital[Title/Abstract] OR spouse[Title/Abstract] OR spouses [Title/Abstract] OR partner[Title/Abstract] OR partners[Title/Abstract] OR "intimate partner"[Title/Abstract] OR "intimate partners"[Title/Abstract] OR dating[Title/Abstract] OR breakup[Title/Abstract] OR breakups[Title/Abstract] OR divorce[Title/Abstract] OR divorced[Title/Abstract]))	English; publication date 2023/01/01–2026/06/07. Peer-reviewed publication status assessed during screening.	June 7, 2026
Scopus	Title, abstract, and keywords	TITLE-ABS-KEY(("generative AI" OR "generative artificial intelligence" OR ChatGPT OR "large language model" OR "large language models" OR LLM OR LLMs OR "conversational agent" OR "conversational agents") AND ("romantic relationship" OR "romantic relationships" OR couple OR couples OR marriage OR marital OR spouse OR spouses OR partner OR partners OR "intimate partner" OR "intimate partners" OR dating OR breakup OR breakups OR divorce OR divorced))	PUBYEAR > 2022 AND PUBYEAR < 2027; English. Journal articles and eligible conference papers retained during screening.	June 7, 2026
Web of Science Core Collection	Topic: title, abstract, author keywords, and Keywords Plus	TS=((("generative AI" OR "generative artificial intelligence" OR ChatGPT OR "large language model" OR "large language models" OR LLM OR LLMs OR "conversational agent" OR "conversational agents") AND ("romantic relationship" OR "romantic relationships" OR couple OR couples OR marriage OR marital OR spouse OR spouses OR partner OR partners OR "intimate partner" OR "intimate partners" OR dating OR breakup OR breakups OR divorce OR divorced))	Publication years 2023–2026; English. Peer-reviewed publication type assessed during screening.	June 7, 2026
APA PsycINFO	Title, abstract, and subject terms	((("generative AI" OR "generative artificial intelligence" OR ChatGPT OR "large language model*" OR LLM* OR "conversational agent*") AND ("romantic relationship*" OR couple* OR marriage OR marital OR spouse* OR partner* OR "intimate partner*" OR dating OR breakup* OR divorce* OR divorced))	Publication years 2023–2026; English; journal articles and eligible peer-reviewed conference contributions.	June 7, 2026
Taylor & Francis Online	Anywhere in article	("generative AI" OR "generative artificial intelligence" OR ChatGPT OR "large language model" OR "large language models" OR LLM OR LLMs OR "conversational agent" OR "conversational agents") AND ("romantic relationship" OR "romantic relationships" OR couple OR couples OR marriage OR marital OR spouse OR spouses OR partner OR partners OR "intimate partner" OR "intimate partners" OR dating OR breakup OR breakups OR divorce OR divorced)	Publication date January 1, 2023–June 7, 2026; English. Peer-reviewed research articles retained during screening.	June 7, 2026
ACM Digital Library	Title, abstract, and author keywords	((("generative AI" OR "generative artificial intelligence" OR ChatGPT OR "large language model" OR "large language models" OR LLM OR LLMs OR "conversational agent" OR "conversational agents") AND ("romantic relationship" OR "romantic relationships" OR couple OR couples OR marriage OR marital OR spouse OR spouses OR partner OR partners OR "intimate partner" OR "intimate partners" OR dating OR breakup OR breakups OR divorce OR divorced))	Publication years 2023–2026; research articles, full conference papers, Late-Breaking Work, and Extended Abstracts with sufficient empirical information.	June 7, 2026
Google Scholar	Full record	Separate simplified searches were conducted because Google Scholar does not process long Boolean strings consistently: "generative AI" "romantic relationship"; ChatGPT "relationship advice"; "large language model" couples; ChatGPT marriage; ChatGPT dating; ChatGPT breakup; ChatGPT "intimate partner"; "generative AI" couple communication	Custom date range 2023–2026; English-language peer-reviewed records retained during screening. Results were screened manually for eligibility.	June 7, 2026

Appendix B. Characteristics of the included studies (N = 21)

#	Study	Country	Design	Sample	GenAI system	Main findings
1	Hou et al. (2024)	USA	Computational/ content analysis	13,138 Reddit posts	ChatGPT	• ChatGPT's advice rankings diverged from human judgment • Inconsistent answers on re-prompting raise reliability concerns
2	Vowels (2024)	Switzerland	Experiments + evaluation	Three studies (up to N = 123)	ChatGPT	• Chatbots seen as more helpful and empathic than human experts • Therapists rated chatbot sessions highly on empathy
3	Vowels et al. (2024)	Switzerland/ UK	Qualitative (thematic analysis)	N = 20	ChatGPT	• Single-session ChatGPT counseling rated high on empathy and helpfulness • Promising but limited relative to human therapists
4	Alqurashi (2025)	Saudi Arabia	Cross-sectional survey	693 couples	GenAI (generic)	• Most couples believed AI would improve relationship quality • Authors urge therapists to monitor use and set ethical rules
5	Brailas and Tsolakis (2025)	Greece	Qualitative (interviews)	N = 10	ChatGPT	• Relationship questions clustered into seven themes • Useful when human support is absent, but generic and no substitute
6	Chun et al. (2025)	USA/China	System design + expert evaluation	3 domain experts	ConflictLens (LLM)	• Experts found the LLM conflict-training realistic • Supports self-guided practice of conflict-resolution strategies
7	Döring and Mohseni (2025)	Germany	Two online experiments	N = 2658	Generative AI (text/image)	• Identical content rated lower when labeled AI-generated • Anti-AI bias moderated by attitudes and AI literacy
8	Fu et al. (2025)	USA	Qualitative interviews	N = 21	ChatGPT/GenAI	• Needs differ across breakup stages; current tools fall short • Authenticity and potential harm were key concerns
9	Ploderer et al. (2025)	Australia	Collaborative ethnography	Small panel	LLM chatbots	• Chatbots augmented long-distance communication via role-play and disclosure • Generic, hallucinated replies felt mediocre: augment, not replace
10	Vowels, Francois-Walcott, et al. (2025)	Switzerland/ UK	Mixed-methods (three studies)	Focus groups N = 30; N = 20; N = 260	GenAI chatbots	• Mapped public attitudes and acceptability of GenAI for couples • Found both openness and reservations among users
11	Vowels, Sweeney, and Vowels (2025)	UK/ Switzerland	Pre-post evaluation	N = 54	Amanda (GPT-4, voice)	• Voice chatbot improved distress, communication, conflict confidence
12	Vowels, Vowels, et al. (2025)	Switzerland/ UK/USA	Randomized controlled trial	N = 258	Amanda (GPT-4o)	• High usability and working-alliance ratings • Gains on 13 of 14 outcomes, sustained at follow-up • Comparable to an evidence-based writing task (RCT evidence)
13	Zhang et al. (2025)	USA	Computational/ content analysis	500 IPV + 80 control posts	ChatGPT-3.5	• Identified IPV with 91.2% accuracy and consistent support. • Struggled with nuanced and sexual-violence content
14	Chen and Zhang (2026)	China	Mixed (survey + prototyping)	Survey N = 295; pilot N = 9-16	MoodBridge (LLM)	• Identified perception, expression, and memory communication gaps • Pilot supported LLM nonviolent-communication rewriting
15	Fan et al. (2026)	China/ Malaysia	Online experiments	Corpus N = 152; receivers N = 704	ChatGPT/LLM	• Heavier AI drafting lowered perceived ownership and authenticity • Light tone-rewrite with disclosure raised authenticity in apologies
16	Foriest et al. (2026)	USA	Mixed-methods	155 posts; 5 subreddits	LLMs	• Taxonomy of seven uses of generative AI in IPV experiences • LLMs provide resources but can reproduce epistemic injustice
17	Jiang et al. (2026)	Singapore	Randomized, dyadic	36 couples (N = 72)	LLM chatbot	• Mediating partner responsiveness raised need-support and closeness • Prompting disclosure alone did not increase closeness
18	Tseng and Liang (2026)	USA	Questionnaire + interviews	N = 25 (90 prompts); 17 interviews	GenAI chatbots	• Maps roles users imagine for AI in dating and relationships • Users navigate sycophancy and overreliance for relational gains
19	Ungless and Sastry (2026)	United Kingdom	Exploratory analysis	2 conversation datasets	Conversational LLMs	• People increasingly use LLMs for support during breakups • Highlights risks, including potential for serious harm
20	Won et al. (2026)	South Korea	Controlled user study	20 couples (N = 40)	GenAI drawing system	• Collaborative drawing fostered partner understanding and awareness

(continued on next page)

(continued)

#	Study	Country	Design	Sample	GenAI system	Main findings
21	Zhang et al. (2026)	South Korea	Within-subjects experiment	10 couples (N = 20)	LLM surrogate agents	• Non-verbal interaction offered a path to relational insight • Proxy agents increased perspective-taking and externalized conflict • But raised alliance, exclusion, and boundary concerns

Appendix C. Methodological quality appraisal

Study	Design/evidence type	Appraisal tool	Main concerns identified	Overall judgement
Hou et al. (2024)	Computational analysis of 13,138 Reddit relationship-advice posts comparing LLM and human rankings.	Adapted JBI analytical cross-sectional/observational checklist	Large corpus strengthens precision, but Reddit users and relationship contexts were not independently verified. Corpus-level rankings cannot be treated as representative of individual experiences. Model responses varied across repeated prompts, and the observational design does not establish effects on couples.	Moderate
Vowels (2024)	Multi-study experimental evaluation of GenAI-generated relationship advice, including lay and professional ratings.	JBI randomized or quasi-experimental checklist, according to each component	Controlled comparisons strengthen the appraisal of perceived empathy and helpfulness. However, participants evaluated brief, generated responses rather than using GenAI in an ongoing relationship. Outcomes were primarily subjective ratings, with limited evidence concerning actual relational behavior or durability.	Moderate to high
Vowels et al. (2024)	Qualitative and professional evaluation of single-session GenAI relationship counseling.	JBI qualitative research checklist	Direct professional assessment provides relevant evidence on counseling qualities, but the sample was limited and the interaction was restricted to a single session. Findings concern perceived therapeutic skills and acceptability rather than sustained couple outcomes.	Moderate
Alqurashi (2025)	Cross-sectional survey of 693 couples concerning expected effects of GenAI on relationship quality.	JBI analytical cross-sectional checklist	Relatively large couple sample, but the study measured attitudes and expectations rather than verified GenAI use or observed relational effects. Cross-sectional self-report data are vulnerable to selection and social-desirability bias and do not support causal conclusions.	Moderate
Brailas and Tsolakis (2025)	Qualitative interviews with adults who had used ChatGPT for relationship-related concerns.	JBI qualitative research checklist	The study provides direct and detailed user accounts, but the small, self-selected sample limits transferability. Findings may overrepresent individuals already willing to use and discuss GenAI, and reliance on retrospective accounts limits verification of use and outcomes.	Low to moderate
Chun et al. (2025)	Proof-of-concept conflict-training system evaluated by three experts.	Adapted JBI qualitative/system-evaluation checklist	The study provides early evidence of feasibility, but evaluation by only three experts offers limited empirical support. Couples did not test the system, and no relational or behavioral outcomes were measured. The prototype stage substantially limits generalizability.	Low
Döring and Mohseni (2025)	Two large online experiments examining reactions to AI-labeled relationship-counseling texts and couple images.	JBI randomized controlled trial/quasi-experimental checklist	Experimental manipulation and a large sample strengthen internal validity. However, the studies measured responses to labeled materials rather than naturalistic GenAI use. The authorship-label effect may reflect demand characteristics, social desirability, or beliefs about authenticity rather than a stable anti-AI bias.	Moderate to high
Fu et al. (2025)	Qualitative interviews examining GenAI use during relationship breakups.	JBI qualitative research checklist	The study offers rich accounts across sensitive breakup stages, but the sample was small and self-selected. Retrospective self-report and the absence of partner perspectives limit conclusions about actual relational consequences or broader prevalence.	Low to moderate
Ploderer et al. (2025)	Collaborative ethnographic study of GenAI-supported interaction in long-distance relationships.	JBI qualitative research checklist	The dyadic and contextual approach is a strength, but the small panel and collaborative design limit transferability. Participants' close involvement in the design may influence reported experiences, and the study cannot establish longer-term effects.	Moderate
Vowels, Francois-Walcott, et al. (2025)	Multi-study mixed-methods development and acceptability evaluation of an LLM-based relationship intervention.	MMAT mixed-methods criteria with design-specific JBI appraisal of components	The use of multiple studies and methods strengthens development and acceptability evidence. However, components differed substantially in design and sample size, and several outcomes concerned hypothetical willingness, usability, or perceived acceptability rather than relationship change.	Moderate

(continued on next page)

(continued)

Study	Design/evidence type	Appraisal tool	Main concerns identified	Overall judgement
Vowels, Sweeney, and Vowels (2025)	Single-session or pre-post evaluation of the Amanda GenAI relationship intervention.	JBİ quasi-experimental checklist	Direct intervention exposure and repeated outcome assessment are strengths. However, the lack of a strong randomized comparison in this component limits attribution of change to GenAI. Brief exposure, self-reported outcomes, and limited follow-up constrain conclusions about durability.	Moderate
Vowels, Vowels, et al. (2025)	Randomized controlled trial of an LLM-based relationship intervention compared with an established writing exercise, N = 258.	JBİ randomized controlled trial checklist	Randomization, an active comparison condition, multiple outcomes, and a two-week follow-up provide the strongest intervention evidence in the review. However, the GenAI intervention did not significantly outperform the writing exercise, follow-up was short, and the evidence derives from the same research program as several related studies.	High
Zhang et al. (2025)	Computational evaluation of ChatGPT-3.5 responses to IPV-related and comparison posts.	Adapted JBİ diagnostic test accuracy/analytical cross-sectional checklist	Use of coded IPV and control cases permits systematic performance assessment. However, online posts may not represent real-time disclosure or clinical decision-making, and the validity of findings depends on the reference coding. Errors in nuanced or ambiguous cases are particularly important given the consequences of false reassurance or misclassification.	Moderate
Chen and Zhang (2026)	Mixed-methods needs assessment and small pilot of the MoodBridge nonviolent-communication prototype.	MMAT mixed-methods criteria with JBİ cross-sectional and quasi-experimental checklists	Combining a needs assessment with prototype testing strengthens design relevance. However, the pilot sample was very small, implementation was brief, and no established relationship outcomes were measured. Evidence supports feasibility more than effectiveness.	Low to moderate
Fan et al. (2026)	Controlled online experiments examining GenAI involvement in intimate-message production and perceptions of effort, authorship, and authenticity.	JBİ randomized controlled trial/quasi-experimental checklist	Experimental manipulation and a relatively large receiver sample strengthen causal interpretation of immediate perceptions. However, participants evaluated experimentally produced messages rather than communication within ongoing relationships. Outcomes were attitudinal and short term, and ecological validity is limited.	Moderate to high
Foriest et al. (2026)	Mixed-methods computational and qualitative analysis of GenAI use in IPV-related online discussions.	MMAT mixed-methods criteria with adapted JBİ observational and qualitative checklists	The study addresses a safety-critical context and combines systematic corpus analysis with interpretive examination. However, posters' identities, circumstances, and outcomes could not be verified. Online discussion data cannot establish prevalence, clinical safety, or the consequences of acting on the advice.	Moderate
Jiang et al. (2026)	Randomized dyadic laboratory experiment with 36 couples examining GenAI-supported disclosure and partner responsiveness.	JBİ randomized controlled trial checklist	Dyadic participation, random assignment, and direct measures of partner responsiveness and closeness strengthen internal validity. However, the sample was modest, the interaction was brief and laboratory-based, and longer-term effects on natural couple communication were not assessed.	High
Tseng and Liang (2026)	Mixed-methods questionnaire and interview study of naturalistic GenAI use for relationship concerns.	MMAT mixed-methods criteria with JBİ analytical cross-sectional and qualitative checklists	Combining questionnaire responses, reported prompts, and interviews provides detailed evidence of user motivations and prompting strategies. Nevertheless, the sample was small and self-selected, use was self-reported, and the study cannot estimate prevalence or establish relational outcomes.	Moderate
Ungless and Sastry (2026)	Exploratory analysis of two datasets of human–LLM conversations concerning relationship dissolution.	Adapted JBİ observational/qualitative checklist	Naturally occurring conversations provide ecologically relevant examples, but only two datasets were examined and participant information was limited. The exploratory design cannot establish how common the identified patterns are or whether the interactions helped or harmed users.	Low to moderate
Won et al. (2026)	Controlled dyadic user study of a GenAI-supported collaborative drawing system with 20 couples.	JBİ quasi-experimental checklist	Direct participation by couples and measurement of understanding and relational awareness are strengths. However, the sample was small, exposure was brief, and outcomes reflected immediate perceptions in an experimental task rather than sustained relationship change.	Moderate
Zhang et al. (2026)	Small within-subjects dyadic experiment examining surrogate GenAI agents in conflict negotiation.	JBİ quasi-experimental checklist	The within-couple design permits direct comparison of responses to different mediation conditions. However, the sample was very small, repeated exposure may have introduced order or carryover effects, and the brief artificial conflict setting limits transferability to real disagreements.	Low to moderate

Data availability

No data was used for the research described in the article.

References

- Adewale, M. D., & Muhammad, U. I. (2025). From virtual companions to forbidden attractions: The seductive rise of artificial intelligence love, loneliness, and intimacy: A systematic review. *Journal of Technology in Behavioral Science*. <https://doi.org/10.1007/s41347-025-00549-4>
- Alon, L., Hadar Shoval, D., & Levkovich, I. (2026a). Bias and representation in AI-generated text-to-image in education: A systematic review. *Computers and Education: Artificial Intelligence*, 10, Article 100587. <https://doi.org/10.1016/j.caeai.2026.100587>
- Alon, L., Hadar Shoval, D., & Levkovich, I. (2026b). Bias, representation, and clinical fidelity in AI-generated images for medical education: A systematic literature review. *Npj Digital Medicine*. <https://doi.org/10.1038/s41746-026-02608-3>
- Alqurashi, T. S. (2025). The impact of artificial intelligence on intimate partner relationships. *American Journal of Family Therapy*. <https://doi.org/10.1080/01926187.2025.2591600>
- Bodenmann, G. (1995). A systemic-transactional conceptualization of stress and coping in couples. *Swiss Journal of Psychology*, 54(1), 34–49.
- Bowen, M. (1978). *Family therapy in clinical practice*. Jason Aronson.
- Brailas, A., & Tsolakis, L. (2025). Questions people ask ChatGPT regarding their romantic relationships and what they think about the provided answers: An exploratory study. In *In Chatbots and human-centered AI: Conversations 2024*. Springer. https://doi.org/10.1007/978-3-031-88045-2_10
- Chan, K. C. T., Li, X., Liu, Y., Chen, B., & Han, Z. (2025). 'What is love?': Exploring the feeling rules and emotion work of Chinese users in human-AI romance. *International Journal of Intercultural Relations*. <https://www.sciencedirect.com/science/article/pii/S014717672500104X>
- Chen, X., & Zhang, X. (2026). From needs to design: Bridging intimate gaps with LLM-augmented nonviolent communication. In *Extended abstracts of the 2026 CHI conference on human factors in computing systems*. Association for Computing Machinery. <https://doi.org/10.1145/3772363.3798539>
- Chun, J., Zhang, G., & Xia, M. (2025). ConflictLens: LLM-based conflict resolution training in romantic relationship. In *Adjunct proceedings of the 38th annual ACM symposium on user interface software and technology (UIST adjunct '25)*. Association for Computing Machinery. <https://doi.org/10.1145/3746058.3758422>
- Djufiril, R., Frampton, J. R., & Knobloch-Westervick, S. (2025). Love, marriage, pregnancy: Commitment processes in romantic relationships with AI chatbots. *Computers in Human Behavior: Artificial Humans*. <https://www.sciencedirect.com/science/article/pii/S2949882125000398>
- Döring, N., & Mohseni, M. R. (2025). Anti-AI bias toward couple images and couple counseling: Findings from two experiments. *Archives of Sexual Behavior*. <https://doi.org/10.1007/s10508-025-03318-9>
- Dubey, S., Dubey, M. J., Sengupta, S., Ghosh, R., Das, S., & Chatterjee, S. (2026). Love, sex and artificial intelligence. *QJM: An International Journal of Medicine*, 119(2), 81–84. <https://doi.org/10.1093/qjmed/hcaf202>
- Fan, G., Liu, D., & Pan, L. (2026). Is it still you? Attributing authorship and authenticity in AI-assisted romantic communication. In *Proceedings of the 2026 CHI conference on human factors in computing systems*. Association for Computing Machinery. <https://doi.org/10.1145/3772318.3790762>
- Foriest, J. C., Ajmani, L., & De Choudhury, M. (2026). (Re)mediators of epistemic injustice: Generative AI and hermeneutic resource provision in intimate partner violence. In *Proceedings of the 2026 CHI conference on human factors in computing systems*. Association for Computing Machinery. <https://doi.org/10.1145/3772318.3791548>
- Fu, Y., Chen, Y., Lai, Z. G. D. C., & Hiniker, A. (2025). Should ChatGPT write your breakup text? Exploring the role of AI in relationship dissolution. *Proceedings of the ACM on Human-Computer Interaction*, 9(7). <https://doi.org/10.1145/3757424>. CSCW330.
- Fu, Y., Foell, S., Xu, X., & Hiniker, A. (2024). From text to self: Users' perception of AIMC tools on interpersonal communication and self. *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*. <https://doi.org/10.1145/3613904.3641955>
- Gur, T., & Maaravi, Y. (2025). The algorithm of friendship: Literature review and integrative model of relationships between humans and artificial intelligence (AI). *Behaviour & Information Technology*, 44(14), 3446–3466. <https://doi.org/10.1080/0144929X.2025.2502467>
- Haber, Y., Levkovich, I., Hadar-Shoval, D., & Elyoseph, Z. (2024). The artificial third: A broad view of the effects of introducing generative artificial intelligence on psychotherapy. *JMIR Mental Health*, 11, Article e54781. <https://doi.org/10.2196/54781>
- Ho, Q. H. J., Hu, M., Chen, T. X., & Hartanto, A. (2025). Potential and pitfalls of romantic artificial intelligence (AI) companions: A systematic review. *Computers in Human Behavior Reports*, 19, 1–19. <https://doi.org/10.1016/j.chbr.2025.100715>
- Hou, H., Leach, K., & Huang, Y. (2024). ChatGPT giving relationship advice: How reliable is it? *Proceedings of the International AAAI Conference on Web and Social Media*, 18(1), 610–623. <https://doi.org/10.1609/icwsm.v18i1.31338>
- Jiang, Z., Yeo, S., Herremans, D., & Perrault, S. T. (2026). Scaffolded vulnerability: Chatbot-mediated reciprocal self-disclosure and need-supportive interaction in couples. In *Proceedings of the 2026 CHI conference on human factors in computing systems*. Association for Computing Machinery. <https://doi.org/10.1145/3772318.3791370>
- Jocher, J. J., & Verviebe, R. (2026). 'Forever and always, my love': Emotions in human-AI romantic relationship building and breakup with generative AI chatbot Replika. *Social Media + Society*. <https://doi.org/10.1177/20563051251408917>
- Keeler, J. B., & Murphy, B. A. (2026). Chatbots and human-human relationships: The need for research on potential downstream harms from generative AI. *Community, Work & Family*. <https://doi.org/10.1080/13668803.2026.2623500>
- Laestadius, L., Bishop, A., Gonzalez, M., Ilencik, D., & Campos-Castillo, C. (2024). Too human and not human enough: A grounded theory analysis of mental health harms from emotional dependence on the social chatbot Replika. *New Media & Society*, 26(10), 5923–5941. <https://doi.org/10.1177/14614448221142007>
- Liu, J. (2024). ChatGPT: Perspectives from human-computer interaction and psychology. *Frontiers in Artificial Intelligence*, 7, Article 1418869. <https://doi.org/10.3389/frai.2024.1418869>
- Luo, X., Wang, Z., Tilley, J. L., Balarajan, S., Bassey, U.-A., & Cheang, C. I. (2025). Seeking emotional and mental health support from generative AI: Mixed-methods study of ChatGPT user experiences. *JMIR Mental Health*, 12, Article e77951. <https://doi.org/10.2196/77951>
- Ma, R., He, S., Martin-Navarro, J. L., Zhan, X., & Such, J. (2026). Privacy in human-AI romantic relationships: Concerns, boundaries, and agency. *Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems*. <https://doi.org/10.1145/3772318.3791237>
- Malfacini, K. (2025). *The impacts of companion AI on human relationships: Risks, benefits, and design considerations*. AI & Society. <https://doi.org/10.1007/s00146-025-02318-6>
- Meng, J., Zhang, R., Qin, J., Lee, Y.-J., & Lee, Y.-C. (2025). AI-mediated social support: The prospect of human-AI collaboration. *Journal of Computer-Mediated Communication*, 30(4). <https://doi.org/10.1093/jcmc/zmaf013>
- Muldoon, J., & Parke, J. J. (2025). *Cruel companionship: How AI companions exploit loneliness and commodify intimacy*. New Media & Society. <https://doi.org/10.1177/14614448251395192>
- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., et al. (2021). The PRISMA 2020 statement: An updated guideline for reporting systematic reviews. *BMJ*, 372, n71. <https://doi.org/10.1136/bmj.n71>
- Pan, S., & Mou, Y. (2024). Constructing the meaning of human-AI romantic relationships from the perspectives of users dating the social chatbot Replika. *Personal Relationships*, 31(4), 1090–1112. <https://doi.org/10.1111/pere.12572>
- Ploderer, B., Capel, T., Davaakhuu, N.-E., Tung, N. H., Maichal, D. M., Kuzhiparambil, A. K., Mai, Q. H., & Reiterberger, W. (2025). Maintaining long-distance relationships with (mediocre) LLM-based chatbots: A collaborative ethnographic study. In *Extended abstracts of the 2025 CHI conference on human factors in computing systems*. Association for Computing Machinery. <https://doi.org/10.1145/3706599.3720221>
- Richards, D., Alsharbi, B., & Abdulrahman, A. (2020). Can I help you? Preferences of young adults for the age, gender and ethnicity of a virtual support person based on individual differences including personality and psychological state. In *In proceedings of the Australasian computer science week multiconference* (pp. 1–10).
- Rocha, A. (2025). "We share an unbreakable bond": Sociality and language ideologies in human relationships with artificial intelligence. *Signs and Society*. <https://www.cambridge.org/core/journals/signs-and-society/article/we-share-an-unbreakable-bond-sociality-and-language-ideologies-in-human-relationships-with-artificial-intelligence/C3DFC9B45C4E51E32B4374F5AFA40705>
- Rodger, K., & Field, E. F. (2025). You and I plus AI: A qualitative exploration of Replika in the context of human relationships. *The Canadian Journal of Human Sexuality*, 34(3), 398–411. <https://doi.org/10.3138/cjhs-2025-0011>
- Shu, C., Lai, K., & He, L. (2026). Human-AI attachment: How humans develop intimate relationships with AI. *Frontiers in Psychology*, 17, Article 1723503. <https://doi.org/10.3389/fpsyg.2026.1723503>
- Smith, M. G., Bradbury, T. N., & Karney, B. R. (2025). Can generative AI chatbots emulate human connection? A relationship science perspective. *Perspectives on Psychological Science*. <https://doi.org/10.1177/17456916251351306>
- Tricco, A. C., Khalil, H., Holly, C., Feyissa, G., Godfrey, C., Evans, C., ... Munn, Z. (2022). *Rapid reviews and the methodological rigor of evidence synthesis: A JBI position statement*. *JBI Evidence Synthesis*, 20(4), 944–949.
- Tseng, E., & Liang, C. A. (2026). "Chat, should I leave him?" risks, rewards, and roles for AI in relationship advice. In *In proceedings of the 2026 CHI conference on human factors in computing systems*. Association for Computing Machinery. <https://doi.org/10.1145/3772318.3790739>
- Ungless, E. L., & Sastry, N. (2026). I'm thinking of ending things: Use of LLMs for support during break-ups. In *Extended abstracts of the 2026 CHI conference on human factors in computing systems*. Association for Computing Machinery. <https://doi.org/10.1145/3772363.3798812>
- Vowels, L. M. (2024). Are chatbots the new relationship experts? Insights from three studies. *Computers in human behavior*. *Artificial Humans*, 2(2). <https://www.sciencedirect.com/science/article/pii/S2949882124000379>
- Vowels, L. M., Francois-Walcott, R. R. R., & Darwiche, J. (2024). AI in relationship counselling: Evaluating ChatGPT's therapeutic capabilities in providing relationship advice. *Computers in Human Behavior: Artificial Humans*, 2(2). <https://www.sciencedirect.com/science/article/pii/S2949882124000380>
- Vowels, L. M., Francois-Walcott, R. R. R., Grandjean, M., Darwiche, J., & Vowels, M. J. (2025a). Navigating relationships with GenAI chatbots: User attitudes, acceptability, and potential. *Computers in Human Behavior: Artificial Humans*, 5, Article 100183. <https://www.sciencedirect.com/science/article/pii/S2949882125000672>
- Vowels, L. M., Sweeney, S. K., & Vowels, M. J. (2025). Evaluating the efficacy of Amanda: A voice-based large language model chatbot for relationship challenges.

- Computers in Human Behavior: Artificial Humans, 4, Article 100141. <https://www.sciencedirect.com/science/article/pii/S2949882125000258>.
- Vowels, L. M., Vowels, M. J., Sweeney, S. K., Hatch, S. G., & Darwiche, J. (2025). The efficacy, feasibility, and technical outcomes of a GPT-4o-based chatbot Amanda for relationship support: A randomized controlled trial. *PLOS Mental Health*, 2(9), Article e0000411. <https://doi.org/10.1371/journal.pmen.0000411>
- Wang, X., Wang, Z., Zhang, M., Yu, K., Hui, P., & Fan, M. (2024). Avatar appearance and behavior of potential harassers affect users' perceptions and response strategies in social Virtual Reality (VR): A mixed-methods study. *Proceedings of the ACM on Human-Computer Interaction*, 8(CSCW2), 1–27. <https://doi.org/10.1145/3686934>
- Willoughby, B. J., Dover, C. R., Hakala, R. M., & Carroll, J. S. (2025). Artificial connections: Romantic relationship engagement with artificial intelligence in the United States. *Journal of Social and Personal Relationships*. <https://doi.org/10.1177/02654075251371394>
- Won, H., Kim, J., Kim, J., Kim, T., & Han, J. (2026). Venus and Mars on canvas: AI-mediated collaborative drawing for romantic relationship insight. In *Proceedings of the 2026 CHI conference on human factors in computing systems*. Association for Computing Machinery. <https://doi.org/10.1145/3772318.3791153>.
- Zhang, Y., Fang, J., Luo, X., Lindsay, D., Madre, N., Paredes, J., Penna, A., Melley, E., & Garcia, T. (2025). Exploring the efficacy of ChatGPT in understanding and identifying intimate partner violence. *Family Relations*, 74(3), 1233–1249. <https://doi.org/10.1111/fare.13176>
- Zhang, Y., Hou, Z., & Shao, J. (2026). Two bots, one couple: How surrogate LLM agents shape alliance, fairness, and relational boundaries. In *Extended abstracts of the 2026 CHI conference on human factors in computing systems*. Association for Computing Machinery. <https://doi.org/10.1145/3772363.3798594>.