

Metaemotions and Consumer Decision Making in Petcare : A Meta-Care Framework for Human Centred Conversational AI

Kavita Gupta ¹, Dr. Swati Punjani ²

¹PhD Scholar, School of Leadership and Management, Manav Rachna International Institute of Research and Studies

e-mail : kavita_gupta.phd@mriu.ac.in

ORCID - <https://orcid.org/0009-0002-4814-5770>

²Assistant Professor, School of Leadership and Management, Manav Rachna International Institute of Research and Studies.

ORCID <https://orcid.org/0009-0003-2176-5002>

ABSTRACT

Conversational Artificial Intelligence (CAI) systems are increasingly used by consumers as service advisors, guides, and support for emotionally sensitive and care-oriented domains such as pet health, nutrition, behaviour, and welfare. As consumers increasingly rely on these AI-mediated service interactions, it has become an important service marketing concern to understand these systems and their role in shaping emotional experiences. While existing research has primarily evaluated CAI platforms on dimensions such as accuracy, transparency, and ethical risk, considerably less attention has been paid to the role of CAIs in shaping users' emotional experiences, influence on emotional burden, emotional reassurance and their metaemotional responses—how users feel about their own emotions when interacting with Conversational AI.

To address this gap this study examines CAI as human-centred service advisor through a qualitative comparative analysis of five widely used CAI platforms - ChatGPT, Gemini, Perplexity, Copilot and Meta AI, focusing on their responses to identical pet care-related queries. A prompt-controlled research design, using a standardized set of pet care queries was developed to use across four categories: routine care, dietary decisions, health concerns, and ethically ambiguous situations. Responses from each CAI system were systematically analyzed using a two-layer coding framework capturing (1) primary emotional cues embedded in AI responses (e.g., reassurance, urgency, caution, empathy) and (2) metaemotional affordances, including emotional regulation, amplification of anxiety or confidence, guilt normalization, and dependence signalling. Cross-platform comparison was conducted to identify systematic differences in emotional framing, uncertainty communication, and reflective guidance.

Findings demonstrate that CAI platforms create distinct human-centred service experiences by varying in their emotional reassurance, ethical framing, communication of uncertainty and support for users' confidence, anxiety, guilt, and responsibility. The study contributes to service marketing and human-AI interaction literature by introducing the META-CARE (Metaemotional and Ethical Care) Framework, a qualitative evaluation tool for emotionally sensitive and AI-mediated service experiences. Beyond technical performance, the framework highlights the role of CAI in shaping emotional burden, metaemotional affects, consumer trust, and decision confidence. It offers practical guidance for designing more human-centred AI service interactions in care-oriented contexts, including pet care.....

Keywords: Conversational AI, Metaemotions, Human-Centred Service Experience, AI Service Advisor, Consumer Decision-Making, META-CARE Framework, Emotions..

INTRODUCTION

Conversational AI is increasingly replacing or complementing traditional frontline customer support engagements. Consumers now seek advice, reassurance, and decision support through AI – “First -Contact” - even before consulting service centre professionals or online communities in everyday caregiving decisions. Consequently, CAI is evolving from an information

Advances in Consumer Research

retrieval tool to a larger role of being a human-centred service advisor whose quality is judged by its ability to reduce emotional burden, build confidence and support decision-making. In contrast to traditional information retrieval, modern large language model (LLM) or CAI systems produce synthesised recommendations that frequently incorporate normative indicators (“should/should not”), risk evaluations (“urgent/caution”), and emotional framing

(“reassuring/empathetic”). This transition is particularly significant in consumer contexts focused on care, when judgements are ethically complex, emotionally intense, and often made under ambiguity (Grewal et al. 2025). Pet care serves as a prime illustration of this environment. Pet owners frequently request guidance on nutrition, standard welfare practices, symptom analysis, and ethically complex choices including confinement, breeding, behaviour modification, and end-of-life decisions. These actions significantly impact animal wellbeing and influence carers' perceptions of their responsibilities.

Academic and applied evaluation of conversational AI has focused on performance and safety criteria such coherence, relevance, faithfulness/groundedness, and dangerous content protection. Misinformation, false reassurance, delayed veterinarian consultation, and responsibility diffusion are also discussed in veterinary and pet health domains. However, a methodological gap remains: many evaluation approaches either (a) focus on generic conversational quality without anchoring analysis in structured decision domains, or (b) focus on domain risk narratives without providing a replicable, text-analytic protocol for comparing how different AI systems actually construct advice in ethically and emotionally complex situations.

Literature Review

Evaluation of Conversational AI - From Performance Metrics to Qualitative Frameworks

Early evaluation of conversational AI systems was dominated by performance-oriented metrics, including coherence, task completion, relevance, and factual correctness. Such metrics were largely inherited from natural language processing benchmarks and information retrieval traditions, where conversational agents were evaluated primarily as functional systems rather than advisory or relational entities (Radziwill & Benton, 2017). While these approaches provided necessary foundations, they were insufficient for evaluating conversational agents operating in human-centric and care-oriented contexts, where correctness alone does not determine quality or appropriateness.

In response, qualitative and mixed-method evaluation frameworks began to emerge like FAST framework—Fidelity, Accuracy, Safety, and Tone—designed to evaluate deployed AI health-coach dialogues in real clinical environments (Neary et al., 2025). FAST explicitly recognizes that conversational AI must be evaluated not only for informational accuracy but also for safety compliance and emotional tone, thereby legitimizing qualitative judgment as part of AI evaluation. Similarly, CAPE and CAPE-II were developed to evaluate psychotherapy conversational agents across dimensions such as rapport, boundary management, ethical safety, and therapeutic appropriateness (Sobowale et al., 2025). These frameworks show that rubric-based qualitative evaluation is viable and important for AI systems in vulnerable, trusting, and emotionally engaged domains. Beyond domain-specific frameworks, broader multi-layer interpretive models have been proposed to assess human–AI interaction quality which typically examine layers such as system capacity (task effectiveness, usability),

alignment (social presence, warmth), interaction levers (persona framing, onboarding), and perceived outcomes (user satisfaction, trust) (Liu et al., 2024). While conceptually rich, such frameworks remain largely interaction-centric and do not directly operationalize ethical decision-making or emotional self-regulation processes.

This literature makes three key points viz. Qualitative evaluation of conversational AI is important and necessary. Current frameworks are domain-specific or interaction-focused. While emotional tone and safety are acknowledged, deeper emotional processes are still unclear.

Ethical Reasoning and Risk Framing in AI-Mediated Advice

Parallel to evaluation research, scholars have increasingly examined the ethical risks of AI-mediated advice, particularly in healthcare and caregiving domains. Concerns include misdiagnosis, inappropriate reassurance, delayed professional intervention, and responsibility diffusion between AI systems and human decision-makers (Floridi et al., 2018; Jokar et al., 2024). Veterinary and animal-health scholarship has echoed these concerns, warning that AI-generated advice may unintentionally encourage risky self-management or undermine professional care pathways (Jokar et al., 2024).

However, much of this literature remains normative rather than analytical. Even when ethical reasoning is acknowledged, it is rarely operationalized into evaluative parameters that can be systematically applied across AI systems. Recent comparative studies of large language models have highlighted differences in risk framing, disclaimer usage, and evidence signposting, particularly when models respond to health-related queries (Rane et al., 2024). These studies suggest that AI systems embed distinct normative orientations—ranging from precaution-first to reassurance-oriented approaches—but they stop short of examining how such orientations interact with users' emotional experiences or self-judgment.

Thus, while ethical risk is widely recognized, there remains a lack of structured qualitative methodologies capable of capturing ethical reasoning as a discursive practice within AI-generated guidance.

2.3 Pet Care as a Distinct Care-Ethical Decision Domain

Pet care occupies a unique position within care-oriented decision-making literature. Unlike clinical healthcare, pet care decisions are often made outside formal institutional settings, yet they involve dependent beings, asymmetrical power, and long-term welfare consequences. Research in veterinary ethics has shown that pet owners frequently experience moral distress when making decisions about confinement, treatment escalation, or euthanasia, particularly under conditions of uncertainty (Yildirim, 2025).

Empirical studies further demonstrate that pet caregiving is emotionally complex. Hoummady et al. (2025) identify distinct caregiver profiles in end-of-life decision-making, each associated with different emotional burdens and

support needs. Similarly, Kogan et al. (2023a, 2023b) document disenfranchised guilt among dog and cat owners, linking caregiver guilt to anxiety and depressive symptoms. These findings establish pet care as a domain where **emotional self-evaluation** is not incidental but central to decision-making.

Despite this, pet care remains underrepresented in conversational AI evaluation research. Existing studies that reference animal health primarily focus on accuracy or safety risks, rather than examining how AI systems structure advice in morally ambiguous situations or how they influence caregivers' emotional self-regulation. This gap makes pet care a theoretically and empirically valuable domain for extending qualitative AI evaluation frameworks.

2.4 Meta-Emotions and Their Relevance to AI-Mediated Decision-Making

Meta-emotions—emotions about one's own emotions—have been extensively theorized in psychology as second-order affective processes that shape emotion regulation, moral judgment, and identity (Miceli & Castelfranchi, 2019 ; Greiner et al., 2025). Meta-emotional experiences such as guilt about feeling frustrated, anxiety about uncertainty, or shame about perceived inadequacy influence not only how individuals feel, but how they interpret and act upon those feelings (Norman, 2014).

In caregiving contexts, meta-emotions play a particularly critical role. Caregivers often evaluate themselves through their emotional responses, interpreting guilt as moral failure or anxiety as irresponsibility. In pet care, where social validation is limited and decisions are often private, such meta-emotional processes can strongly shape behavior (Kogan et al., 2023a).

Human–AI interaction research has increasingly acknowledged emotional tone and empathy as important dimensions of conversational AI quality. However, existing work typically treats emotion as a first-order phenomenon (e.g., whether an AI sounds empathetic) rather than examining how AI responses shape users' emotions about their emotions. As a result, meta-emotional processes remain largely invisible in AI evaluation frameworks.

This omission is consequential. When conversational AI systems respond to emotionally charged queries, they may normalize guilt, amplify anxiety, validate moral conflict, or encourage emotional distancing. These effects are not captured by existing evaluation metrics focused on correctness, safety, or tone.

2.5 Research Gap

Although service marketing literature has extensively examined customer experience, emotional labour, and human-centred service design, little is known about how conversational AI creates emotionally supportive service experiences or helps in dealing with metaemotions – feelings about feelings- in care oriented contexts. Existing AI evaluation frameworks rarely examine emotions and/or metaemotions as dimensions of AI mediated service quality. Existing literature reveals a clear methodological and conceptual gap. Qualitative

frameworks for evaluating conversational AI are growing in care-related domains. The ethical risks of AI-mediated advice are known, and pet care is recognized as an emotionally complex caregiving domain. Metaemotions are well-established in psychology and reported among pet caregivers. However, current approaches do not provide a standardized, decision-domain-specific instrument for evaluating conversational AI guidance in pet care, nor do they operationalize metaemotions as empirical features of AI discourse. Conversational AI systems are rarely evaluated as ethical and emotional advisors or service aids, despite functioning as such in practice.

The present study addresses this gap by extending qualitative evaluation of conversational AI to include pet-care decision contexts and by introducing a structured, multi-layer methodology that explicitly captures meta-emotional reasoning within AI-generated guidance.

Research Objectives

To conduct a cross-platform comparative analysis of conversational AI systems (ChatGPT, Gemini, Perplexity, Copilot and Meta AI) to identify similarities and differences in ethical orientation, emotional framing, uncertainty communication, and meta-emotional engagement.

To identify emergent patterns and themes in how CAI systems frame pet-care advice.

To develop a standardized qualitative evaluation instrument for pet-care decision-making.

To develop a human-centred service evaluation framework capable of assessing emotional uplift, burden and metaemotional engagement during CAI mediated service interactions.

Methodology

This study adopts a qualitative, comparative research design grounded in an interpretivist epistemology. Conversational Generative AI (CAI) systems are treated as frontline digital service advisors that produce discursive guidance in response to emotionally and ethically complex human queries. Accordingly, the unit of analysis is the AI-generated response text, examined for its reasoning structure, emotional framing, ethical orientation, and meta-emotional affordances.

The present study employs qualitative analysis to examine how AI systems construct meaning, frame responsibility, and influence emotional self-appraisal in caregiving contexts. This approach is consistent with prior qualitative evaluations of conversational agents in health coaching and psychotherapy, which argue that ethical and emotional dimensions cannot be reliably captured through automated metrics alone.

Five widely used conversational AI systems were selected for comparison based on their market prominence, accessibility, and relevance to consumer-facing advisory use:

ChatGPT (OpenAI)

Gemini (Google)

Perplexity AI

Meta AI

Microsoft Copilot

The inclusion of five systems—rather than two or three—was intentional. Prior comparative studies have typically focused on limited pairwise comparisons (e.g., ChatGPT vs. Gemini), often emphasizing performance or accuracy differences. This study enables a broader cross-platform qualitative comparison, reducing the risk that observed patterns are artefacts of a single system’s design philosophy. Each system was used through its publicly available conversational interface within the same data-collection period to reduce temporal volatility in model behaviour. Untrained, anonymous versions of these instruments were utilised during the research to guarantee an impartial evaluation.

To ensure methodological rigor and comparability, all five AI systems were provided with an identical standardized instruction prompt prior to administering the questionnaire. The use of a standardized instruction prompt is consistent with best practices in qualitative comparison of conversational agents, as it minimizes prompt-induced variance and ensures that differences in responses can be attributed to model behavior rather than task framing (Radziwill & Benton, 2017). The full instruction prompt is reproduced in Annexure A. The prompt explicitly positioned the AI as a care-oriented educational advisor, discouraged alarmist or legalistic disclaimers, and encouraged acknowledgment of uncertainty and emotional difficulty.

A 17-item standardized questionnaire was developed to elicit AI-generated guidance across four distinct pet-care decision domains. The four-domain structure was informed by veterinary ethics literature emphasizing that pet-care decision-making spans preventive care, lifestyle choices, acute risk assessment, and morally complex judgments (Yildirim, 2025; Hoummady et al., 2025). The instrument was designed to reflect realistic, everyday pet-care dilemmas rather than hypothetical edge cases. Questions were deliberately open-ended and reflective, aiming to elicit reasoning depth, ethical justification, emotional sensitivity, and uncertainty management rather than factual recall. The four domains included are:

Routine care practices

Dietary decisions

Health concerns and uncertainty

Ethically ambiguous situations

The full questionnaire is provided in Annexure B.

Verbatim AI responses were collected without revision or follow-up to preserve analytical integrity. Conversational memory was disabled, and each response was treated as an independent textual unit. Responses were aligned question-wise in a comparative dataset, enabling systematic cross-platform analysis while maintaining traceability between analytical claims and original text.

Analysis followed a three-layer qualitative framework integrating construct validation, thematic patterning, and meta-emotional examination. This approach extends existing qualitative evaluation frameworks (e.g., FAST; CAPE) by explicitly distinguishing between construct presence, discursive pattern formation, and meta-emotional affordances. Responses from five conversational AI systems—ChatGPT, Gemini, Perplexity, Meta AI, and Microsoft Copilot—were analyzed as advisory texts to assess how guidance is reasoned, framed, and emotionally positioned.

The analysis does not seek to rank systems, but to identify convergence and divergence in ethical reasoning and emotional framing. In the first analytical layer, responses were deductively coded against predefined evaluative parameters derived from prior frameworks and adapted to pet-care contexts, with each parameter operationally defined to enhance consistency and auditability (Neary et al., 2025). The resulting framework comprises twelve parameters mapped to pet-care domains and associated meta-emotional constructs, enabling transparent cross-system comparison without reducing qualitative differences to numerical scores. The full parameter definitions are provided in Table 1.

Each AI response was first coded against the pre-defined analytical parameters. Coding focused on presence, absence, and manner of expression, rather than frequency counts. For example, depth of guidance was assessed by whether responses articulated reasoning and trade-offs, not by response length. After parameter coding, results were inductively analysed for recurring themes across models. Only themes that appeared consistently across prompts and models were maintained to ensure robustness rather than anecdotal interpretation. In the third stage, AI responses were reviewed for metaemotional cues, meaning how emotions about other emotions were expressed, such as guilt or anxiety. Metaemotions were coded only when they were clearly evident in the text through direct wording, trade-off framing, or emotional regulation signals. This conservative approach avoids speculative psychological attribution and grounds claims in textual evidence (Table 1).

Table-1: Integrative Framework for Analysing Conversational AI Responses in Pet Care

Table 1: Integrative Framework for Analysing Conversational AI Responses in Pet Care

Parameter	Analytical Definition	Categories	Linked Meta-Emotion Constructs	What Is Compared Across AI Tools
Depth of Guidance	Degree to which advice includes reasoning, trade-offs, and consequences	Routine Care Diet, Health	Responsibility, Competence	Whether responses remain superficial or provide

				reflective, principle-based guidance
Clarity & Actionability	Extent to which advice is implementable and decision-oriented	Routine Care Health	Confidence, Anxiety	How guidance is translated into actionable steps or thresholds
Risk Framing & Precaution	How risks are identified, emphasized, or moderated	Health Concerns Ethics	Anxiety, Fear	Differences in warning intensity and escalation logic
Ethical Reasoning Structure	Moral logic used to justify recommendations	Ethically Ambiguous Situations Health	Guilt, Moral Conflict	Welfare ethics vs harm minimization vs normative rules
Emotional Sensitivity	Acknowledgement of pet-owner emotional states	Ethically Ambiguous Situations Health	Empathy, Emotional Validation	Variation in reassurance and emotional acknowledgement
Meta-Emotional Awareness	Recognition of emotions influencing decisions	Ethically Ambiguous Situations Diet	Guilt, Anxiety, Self-Judgment	Whether AI reflects on emotions about emotions
Normative vs Contextual Framing	Rule-based versus situationally adaptive advice	Routine Care Ethics	Responsibility, Social Pressure	Rigidity versus contextual flexibility
Consistency Across Care Domains	Coherence of principles across all care categories	All Categories	Trust, Reliability	Stability of care logic across contexts
Handling of Uncertainty	Response to ambiguity or incomplete information	Health Concerns Ethics	Anxiety, Ambivalence	How uncertainty is acknowledged and managed
Evidence Signposting	Reference to scientific or professional norms	Health Concerns Diet	Trust, Reassurance	Reliance on veterinary or consensus reasoning
Tone & Moral Positioning	Overall stance adopted in advice delivery	Routine Ethics	Shame, Pride	Judgment-free versus subtly moralizing tone
Standardization Orientation	Implicit push toward uniform care principles	All Categories	Responsibility, Moral Assurance	Promotion of convergent care standards

AI-Mediated Service Experience Analysis and Results

This section presents the findings of the three-layer qualitative analysis conducted on AI-generated responses to the standardized 17-item pet-care questionnaire. Results are organised in alignment with the analytical framework. Comparative observations across five conversational AI systems—ChatGPT, Gemini, Perplexity, Meta AI, and Microsoft Copilot—are reported too.

5.1. Depth of Guidance and Reasoning Structure

Across all five AI systems, responses showed moderate to high depth of guidance, though the style of reasoning varied. ChatGPT and Gemini consistently offered multi-step, principle-based explanations that connected immediate advice to long-term welfare considerations. These systems frequently contextualised guidance within

broader care principles (e.g., cumulative health effects, chronic stress, welfare trade-offs). Perplexity emphasized procedural clarity, structuring guidance around decision thresholds and escalation rules, resulting in narrower but highly actionable reasoning. Meta AI and Copilot provided concise and coherent advice, with sufficient depth for routine and dietary contexts but limited elaboration in ethically ambiguous situations.

Overall, all systems functioned as explanatory advisors rather than purely reactive responders, with differences reflecting how advisory depth framed user responsibility and competence in decision-making. This parameter revealed how advisory depth activates the meta-emotions of responsibility and competence, positioning consumers as moral decision-makers rather than task executors.

5.2 Clarity and Actionability

Clarity and actionability capture the extent to which advice can be readily implemented. Differences were pronounced across models. Clarity and actionability emerged as a high-coverage construct across all systems, particularly in routine care and health-concern domains. Perplexity and Copilot exhibited the strongest emphasis on decision thresholds, especially in health-related questions (e.g., when to seek professional care). Responses frequently used conditional phrasing (“if X persists,” “when symptoms escalate”) that translated uncertainty into action cues. ChatGPT and Gemini balanced clarity with contextual explanation, sometimes sacrificing immediacy for reflective framing. Meta AI provided general action guidance but fewer explicit thresholds. This variation suggests different advisory orientations: threshold-driven guidance versus reflective contextual guidance, rather than absence of clarity. High actionability supported confidence, whereas lower explicitness sustained decision anxiety, demonstrating how emotional outcomes are shaped by advisory structure.

5.3 Risk Framing and Precaution

Risk framing was universally present but differed in intensity and emphasis. Gemini and Perplexity consistently adopted a precaution-first orientation, framing uncertainty itself as a risk factor and emphasizing early intervention. Delays in care were frequently constructed as cumulative or compounding risks. ChatGPT also foregrounded risk but often coupled it with reassurance and emotional normalization, mitigating potential alarmism. Meta AI and Copilot framed risk more conservatively, acknowledging potential harm without sustained escalation language. Across systems, risk was framed less as probabilistic uncertainty and more as moral responsibility to prevent avoidable harm, indicating convergence around welfare-oriented precaution. Risk framing thus functioned as an emotional calibrator influencing anxiety and fear.

5.4 Ethical Reasoning Structure

Ethical reasoning structure refers to the moral logic used to justify recommendations. Ethical reasoning was evident across all systems, particularly in responses to ethically ambiguous questions (Q13–Q17). Across models, welfare-oriented reasoning dominated. ChatGPT and Gemini articulated ethical trade-offs explicitly, often acknowledging moral conflict and uncertainty. Perplexity framed ethics through practical consequences and decision logic. Meta AI and Copilot adopted a more normative tone, reinforcing responsibility without extended moral reasoning. Ethical justification was typically grounded in harm minimization, proportionality, and long-term quality of life. Explicit ethical reasoning activated guilt and moral ambivalence, while implicit reasoning reduced emotional reflection. Notably, none of the systems relied on rigid rule-based morality; ethical reasoning was predominantly contextual and situational.

5.5 Emotional Sensitivity

Emotional sensitivity captures whether AI responses acknowledge caregiver emotional states. ChatGPT and Gemini regularly acknowledged emotional difficulty.

Meta AI offered supportive but surface-level reassurance. Perplexity and Copilot largely bypassed emotional acknowledgment, focusing instead on informational clarity. Emotional sensitivity shaped perceived empathy but did not necessarily indicate deeper emotional engagement.

5.6 Meta-Emotional Awareness

Meta-emotional awareness refers to whether AI responses engage with emotions about emotions, such as guilt about indulgence or anxiety about uncertainty. ChatGPT consistently demonstrated meta-emotional awareness, reframing guilt and anxiety as indicators of responsible care rather than personal failure. Gemini and Meta AI showed limited implicit engagement occasionally normalizing emotional difficulty without explicitly addressing second-order emotions. Perplexity and Copilot exhibited minimal meta-emotional reflection. This parameter constitutes a key empirical novelty, extending AI evaluation beyond sentiment or tone.

5.7 Normative versus Contextual Framing

This parameter examines whether guidance is rule-based or situationally adaptive. Gemini and Copilot leaned toward normative framing, reinforcing stable standards and obligations. ChatGPT adopted contextual framing, allowing flexibility while maintaining welfare principles. Perplexity and Meta AI occupied intermediate positions. Normative framing amplified social pressure and responsibility, while contextual framing moderated self-judgment.

5.8 Consistency Across Care Domains

Consistency across care domains refers to whether AI systems maintain coherent principles across routine care, diet, health, and ethical contexts. ChatGPT and Gemini exhibited high internal consistency, reinforcing stable welfare-oriented logic across scenarios. Perplexity, Meta AI, and Copilot showed moderate consistency, with shifts in emphasis depending on context. Consistency supported trust and reliability, linking advisory coherence to emotional reassurance.

5.9 Handling of Uncertainty

Handling of uncertainty refers to how AI systems respond to ambiguous or incomplete information. All systems acknowledged uncertainty, but differed in how it was managed. ChatGPT and Gemini frequently made uncertainty explicit and framed it as an inherent feature of caregiving decisions. Perplexity converted uncertainty into decision rules or escalation criteria. Meta AI and Copilot acknowledged uncertainty but often deferred to professional consultation without elaboration. This parameter influenced ambivalence and anxiety regulation.

5.10 Evidence Signposting

Evidence signposting captures references to scientific, veterinary, or professional norms. It shaped trust and reassurance, even in the absence of explicit citations. Evidence signposting was uneven. Gemini and Perplexity most frequently referenced biological, veterinary, or consensus logic (explicit or implicit). ChatGPT occasionally referenced general principles without

citation. Meta AI and Copilot relied more on commonsense reasoning.

5.11 Tone and Moral Positioning

Tone and moral positioning refer to whether advice is delivered in a judgment-free, directive, or subtly moralizing manner. ChatGPT adopted a non-judgmental, reflective tone. Gemini occasionally conveyed moral obligation through duty-oriented language. Perplexity and Copilot maintained neutral, instructional tones. Meta AI emphasized supportive reassurance. Tone influenced experiences of shame, pride, or reassurance, shaping emotional outcomes beyond content.

5.12 Standardization Orientation

Standardization orientation captures whether AI systems promote consistent care principles across contexts. ChatGPT and Gemini implicitly reinforced standardized welfare norms. Perplexity and Copilot emphasized procedural consistency. Meta AI tolerated greater variability in interpretation. This parameter connects advisory logic to moral assurance and responsibility normalization, reinforcing the META-CARE framework's integrative value.

Inductive thematic analysis revealed five dominant cross-platform themes that recurred across questions and systems. Precaution First theme was strongest in Gemini and Perplexity responses but present across all systems. Across health and ethically ambiguous contexts, uncertainty consistently triggered caution rather than reassurance. Even when minimizing alarm, systems framed early intervention as morally preferable to delayed action..

Responsibility amplification theme appeared prominently in ChatGPT and Copilot responses and consistently across ethically ambiguous scenarios. Pet ownership was framed as an active moral role, extending beyond basic provision to ongoing adaptation and vigilance. Responsibility amplification was constructed as continuous rather than episodic.

Several systems explicitly normalized emotional distress, framing guilt, anxiety, or uncertainty as common aspects of responsible caregiving rather than indicators of failure. ChatGPT demonstrated the most explicit emotion-normalization reassurance (theme), while Meta AI and Gemini engaged this theme more implicitly.

Implicit moral hierarchy theme was universal across all five AI systems. A stable moral hierarchy was evident across responses: welfare > convenience > emotional comfort. This hierarchy was rarely stated overtly but consistently structured advice.

Theme of professional authority counters simplistic narratives of AI over-delegation and suggests nuanced positioning of authority. Professional consultation was recommended selectively, particularly when risk escalated or uncertainty persisted. Deferral was framed as supportive rather than substitutive of owner judgment.

The third analytical layer examined metaemotions as observable features of AI-generated responses. Guilt emerged most prominently in dietary decisions and end-of-life scenarios. ChatGPT and Gemini frequently reframed guilt as a signal of care rather than moral failure. Perplexity acknowledged guilt indirectly through discipline framing. Meta AI and Copilot referenced guilt minimally. Anxiety was frequently evoked through risk escalation language, particularly in health contexts. Gemini and Perplexity responses contained stronger urgency cues, while ChatGPT balanced risk with reassurance. Responsibility was the most consistently present meta-emotion across all systems. It was framed as stewardship, obligation, or duty, often reinforcing caregiver identity. Moral conflict appeared explicitly in responses addressing trade-offs between convenience, cost, emotional attachment, and welfare. ChatGPT most explicitly acknowledged “no perfect choice” scenarios, while other systems implied conflict through trade-off logic. Several systems engaged in emotional regulation by normalizing distress, encouraging reflection, or distancing panic from decision-making. This was most explicit in ChatGPT responses.

Across five conversational AI systems, strong convergence was observed in welfare prioritization, precautionary ethics, and responsibility framing. Divergence was primarily stylistic and tonal rather than ethical (Table 2). Critically, meta-emotional processes were empirically observable across AI responses, supporting their inclusion as evaluative constructs rather than interpretive overlays. In sum

All five conversational AI systems demonstrated substantial construct coverage across ethical, emotional, and advisory dimensions.

Pet-care guidance was consistently framed through welfare-first ethical reasoning rather than rule-based prescriptions.

Meta-emotions—particularly responsibility, guilt, and anxiety—were present and structured in AI responses, with variation in explicitness across systems.

Differences between AI systems reflected advisory style, not absence of ethical or emotional engagement.

Table 2 : Comparative evaluation outcome across Conversational AI models

Evaluation Parameter	ChatGPT	Gemini	Perplexity	Meta AI	Copilot
Depth of Guidance	High; reflective, principle-based reasoning with trade-offs	High; grounded in welfare and prevention logic	Moderate; structured, procedural reasoning	Low–Moderate; functionally sufficient, limited elaboration	Moderate; task-focused, limited ethical depth

Clarity & Actionability	High; clear with guidance contextual explanation	Moderate; clear but less step-wise	High; explicit steps, thresholds, escalation points	Moderate; general actionable cues	High; concise, step-oriented guidance
Risk Framing & Precaution	Balanced; caution with reassurance	Strong precautionary bias	Escalation-driven via decision rules	Conservative; avoids alarmism	Conservative; risk-averse framing
Ethical Reasoning Structure	Contextual; explicit acknowledgment of moral trade-offs	Normative; duty-based welfare emphasis	Implicit; embedded within procedural logic	Implicit; welfare prioritised without elaboration	Normative; rule-based welfare framing
Emotional Sensitivity	High; consistent emotional acknowledgment	Moderate; selective emotional validation	Low; minimal emotional reference	Moderate; supportive but surface-level	Low; emotionally neutral tone
Meta-Emotional Awareness	High; reframes guilt/anxiety as responsible care	Low–Moderate; implicit normalization	Low; largely absent	Low; limited reflection on emotions	Low; absent
Normative vs Contextual Framing	Contextual; situational flexibility within welfare bounds	Normative; consistent rule-based standards	Mixed; conditional logic	Mixed; situational but under-theorized	Normative; standardized guidance
Consistency Across Care Domains	High; stable principles across all contexts	High; consistent welfare logic	Moderate; context-specific emphasis	Moderate; shifts across domains	Moderate; functional consistency
Handling of Uncertainty	Reflective; uncertainty normalized and monitored	Precautionary; uncertainty prompts escalation	Resolute; uncertainty resolved via rules	Acknowledged but weakly managed	Resolved via professional deferral
Evidence Signposting	Moderate; welfare-based reasoning	Moderate–High; references to veterinary consensus	High; biological or professional logic invoked	Low; commonsense reasoning	Low–Moderate; implicit professional norms
Tone & Moral Positioning	Non-judgmental, reflective	Mildly moralizing via duty framing	Neutral, instructional	Supportive, reassuring	Neutral, directive
Standardization Orientation	Implicit standardization via ethical consistency	Strong standardization via welfare norms	Procedural standardization	Flexible interpretation	Procedural standardization

Designing Human-Centred AI Service Experiences

The findings of this study show that conversational AI should be seen as both an information retrieval system and as a human-centred service advisor shaping consumers' attitudes, behaviours, perceptions and emotional experiences throughout the service journey. In care-related contexts such as pet care, consumers usually seek reassurance, direction, and emotional support besides data and facts. As a result, the quality of AI-mediated service experiences depends on both the accuracy of recommendations and effectiveness of AI systems in recognising emotional uncertainty, communicate risk, and aid confident decision-making.

The comparative analysis reflects that all five conversational AI systems gave broadly appropriate guidance, they differed in their ability to deliver emotionally helpful service experiences. These differences suggest that in the development and evaluation of conversational AI systems emotionally intelligent service design should become an important consideration. Drawing upon the insights from this study, six design principles are proposed for creating higher and better human-centred AI service experiences.

6.1 Conversational AI as Ethical and Emotional Advisors in Pet-Care Contexts

The findings of this study demonstrate that conversational AI systems—ChatGPT, Gemini, Perplexity, Meta AI, and Microsoft Copilot—function not merely as informational tools but as frontline digital service advisors embedded in ethical and emotional reasoning. Across all five systems, responses consistently extended beyond factual instruction to include moral positioning, precautionary logic, and affective framing. This confirms that, in pet-care contexts, conversational AI participates in value-laden guidance rather than neutral information delivery.

This observation aligns with earlier qualitative evaluations of conversational agents in health coaching and psychotherapy, which highlight the importance of tone, safety, and rapport (Neary et al., 2025; Sobowale et al., 2025). However, the present findings extend this literature by showing that ethical reasoning and emotional framing are not peripheral features but structural components of AI-generated pet-care advice, particularly in situations involving uncertainty and moral trade-offs.

6.2 Risk Framing, Uncertainty, and Responsibility Amplification

A dominant pattern across all five AI systems was the precaution-first framing of uncertainty. Uncertainty was rarely presented as a neutral informational gap; instead, it was framed as a cue for increased vigilance and moral responsibility. This approach mirrors veterinary ethics literature that prioritizes harm minimization under uncertainty (Yıldırım, 2025), suggesting that conversational AI systems implicitly adopt welfare-oriented ethical heuristics.

Simultaneously, responsibility was consistently amplified rather than diffused. Ownership was framed as an ongoing moral commitment requiring attentiveness, adaptation, and sometimes sacrifice. This finding is particularly important in light of concerns that AI-mediated advice may encourage over-reliance or abdication of human responsibility (Jokar et al., 2024). Instead, the data indicate that conversational AI systems tend to reinforce, rather than erode, caregiver accountability—though with varying degrees of emotional calibration.

6.3 Meta-Emotions as a Central but Under-Evaluated Dimension

A key contribution of this study lies in empirically demonstrating that metaemotions are observable and structured within AI-generated pet-care guidance. Guilt, anxiety/fear, moral conflict, and emotional regulation cues were embedded in reasoning patterns, reassurance strategies, and ethical justifications across systems.

Importantly, systems differed in how explicitly they engaged with meta-emotions. ChatGPT most consistently acknowledged and normalized emotional distress, distinguishing guilt from moral failure and framing anxiety as a natural consequence of care. Gemini and Perplexity engaged meta-emotions more implicitly, often through risk escalation or procedural clarity. Meta AI and Copilot maintained supportive tone but exhibited limited explicit meta-emotional reflection.

These differences are theoretically significant. Psychological research indicates that meta-emotions influence self-efficacy, moral self-concept, and decision *Advances in Consumer Research*

confidence (Miceli & Castelfranchi, 2019; Norman, 2014). When conversational AI systems normalize guilt or contextualize anxiety, they may support healthier emotional regulation; when urgency is emphasized without such scaffolding, anxiety may be inadvertently amplified. Existing AI evaluation frameworks do not capture this distinction, underscoring the need for meta-emotion-sensitive assessment.

Overall, these findings suggest that future conversational AI systems should be evaluated not only on technical performance, but also on their capacity to create emotionally supportive, ethically responsible, and trustworthy service experiences. The proposed META-CARE Framework offers a practical foundation for assessing these dimensions and may assist researchers, AI developers, and service organisations in designing conversational AI systems that address consumers' emotional as well as informational needs.

Conclusion and Recommendations

This study demonstrates that major conversational AI systems operate as frontline digital service advisors in pet-care decision-making by consistently incorporating ethical reasoning, precautionary risk framing, and responsibility attribution into their guidance. Crucially, the analysis reveals that metaemotions—such as guilt, anxiety, and moral conflict—are embedded within AI-generated guidance and vary in explicitness across systems. These divergences have important implications for user trust, emotional regulation, and caregiver identity, underscoring the need for evaluation frameworks that extend beyond accuracy and safety to account for the metaemotional impact of AI-mediated advice.

Drawing directly from the empirical findings, this study proposes the META-CARE (Metaemotional and Ethical CARE) Framework as an integrative qualitative evaluation model for conversational AI systems operating in care-oriented domains.

META-CARE is grounded in three principles:

Conversational AI systems participate in ethical framing, even when not explicitly designed as moral agents.

Emotional and metaemotional processes shape how advice is interpreted, trusted, and acted upon.

Care decisions require context-sensitive judgment, not universal prescriptions.

The framework consists of three analytically distinct but interrelated layers:

Layer 1: Advisory Constructs

The base of the META-CARE framework represents Advisory Constructs, which form the functional foundation of conversational AI guidance. This layer captures how advice is delivered and whether it is usable for consumer decision-making. It integrates parameters such as depth of guidance, clarity and actionability, risk framing, handling of uncertainty, and evidence signposting. This layer addresses questions such as whether the AI provides sufficient reasoning, whether advice can be acted upon, how uncertainty is

communicated, and whether guidance is consistent across care contexts. In marketing terms, this layer aligns with perceived usefulness, decision support quality, and trust formation. Without robust advisory constructs, higher-order ethical or emotional engagement lacks operational grounding. The analysis shows that advisory constructs influence not only usability but also consumer confidence and perceived competence. For example, deeper guidance and explicit reasoning positioned consumers as reflective decision-makers, while procedural clarity reduced decisional friction.

Layer 2: Ethical–Discursive Orientation

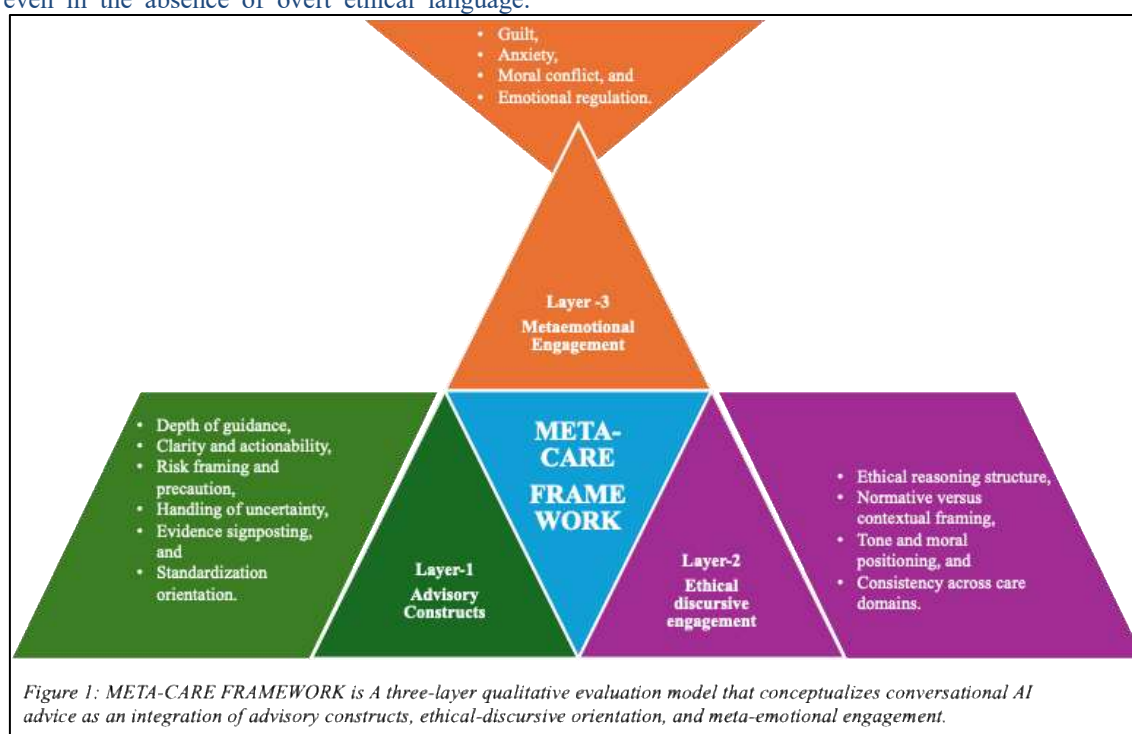
This layer captures the moral logic embedded in AI advice, drawing from parameters such as ethical reasoning structure, normative versus contextual framing, tone and moral positioning, and consistency across care domains. The analysis shows that conversational AI systems inevitably adopt ethical positions—whether explicitly or implicitly—by prioritizing welfare, escalating risk, or framing responsibility. Across models, welfare-oriented ethics were dominant, but the manner of ethical articulation varied—from explicit acknowledgment of moral conflict to implicit harm minimization. These differences shaped how responsibility and obligation were internalized by the user. This layer explains how AI advice shapes moral responsibility, obligation, and social norms, even in the absence of overt ethical language.

From a consumer research perspective, this layer reflects how AI systems function as moral intermediaries, influencing judgments about “right” care practices.

Layer 3: Metaemotional Engagement

This layer represents the most novel contribution of the framework. It integrates emotional sensitivity, meta-emotional awareness, and standardization orientation to examine how AI advice engages with metaemotions. The analysis demonstrates that metaemotional engagement is uneven across models, with only some systems actively reframing negative emotions into signals of responsible care. This layer moves AI evaluation beyond empathy or sentiment to examine emotional self-regulation and moral reassurance. Rather than asking whether AI sounds empathetic, this layer asks whether AI advice actively shapes emotional self-appraisal—for example, by reframing guilt as responsible care or normalizing anxiety under uncertainty. This layer is particularly relevant for marketing and service research, as it explains how AI-mediated advice influences emotional well-being, moral reassurance, and long-term trust.

META-CARE is not intended as a scoring tool but as a qualitative audit and design diagnostic framework, adaptable to other caregiving domains such as elder care, parenting, or chronic illness management (Figure 1).



Theoretical and Practical Implications

The findings of this study carry important implications across stakeholder groups. From a theoretical perspective, this research advances marketing and AI scholarship by integrating meta-emotions into the evaluation of conversational AI, repositioning AI-generated advice as ethically and emotionally structured discourse rather than neutral information delivery. The findings clarify both the strengths and limitations of existing qualitative evaluation

frameworks. Frameworks such as FAST and CAPE have successfully foregrounded safety, tone, and ethical boundaries in care-oriented AI interactions. However, these frameworks do not explicitly operationalize ethical trade-off reasoning under uncertainty or meta-emotional engagement as evaluative dimensions. Rather than positioning prior frameworks as insufficient, the present study demonstrates the need for contextual extension. For marketers and service providers in care-oriented categories, the results demonstrate that conversational AI interfaces function as extensions of brand ethics, shaping

consumer trust, perceived responsibility, and emotional reassurance; applying the META-CARE framework can help align AI-mediated interactions with brand values and long-term relationship goals. For pet parents, the study highlights how AI advice structures not only decisions but emotional self-appraisal, influencing guilt, anxiety, and moral confidence—underscoring the need for transparent, context-sensitive, and emotionally balanced guidance. For AI designers and developers, the findings indicate that technical accuracy alone is insufficient; design choices related to advisory depth, uncertainty handling, ethical framing, and meta-emotional engagement materially affect user experience and well-being, making META-CARE a practical diagnostic and design lens. For policymakers and regulators, the study suggests that governance frameworks focused solely on safety and correctness overlook the ethical and emotional consequences of AI advice in sensitive domains, pointing to the need for evaluation criteria that incorporate responsibility attribution and emotional regulation.

Future Research Directions

Future research may empirically test how different metaemotional framings influence user trust, decision confidence, and long-term caregiving behavior. Quantitative validation of META-CARE dimensions and application of the framework to other caregiving domains represent further promising directions.

Limitations

The study is limited to analysis of AI-generated responses and does not measure downstream user behavior or emotional outcomes. Additionally, responses were analyzed in isolation rather than within extended conversational exchanges.

Ethical Statement and Data Availability

This study analyzed AI-generated textual responses and did not involve human participants, personal data, or animal subjects. Accordingly, formal ethical approval was not required. Nevertheless, ethical sensitivity was central to the study design, given its focus on caregiving decisions, emotional influence, and moral framing.

All data analyzed in this study consist of verbatim AI-generated responses collected using publicly accessible conversational AI interfaces. The aligned comparative dataset and coding matrices are available as supplementary material upon reasonable request to the corresponding author. No human or animal participants involved as no primary survey done.

Disclosure

No potential conflict of interest was reported by the author(s).

Funding

No fundings or financial support received for this study by the author(s).

Contribution Statement

Authors are involved in the conception and design, or analysis and interpretation of the data; the drafting of the paper; revising it critically for intellectual content; and the final approval of the version to be published; and all authors agree to be accountable for all aspects of the work..

REFERENCES

1. Araujo, T. (2018). Living up to the chatbot hype: The influence of anthropomorphic design cues and communicative agency framing on conversational agent and company perceptions. *Computers in Human Behavior*, 85, 183–189. <https://doi.org/10.1016/j.chb.2018.03.051>
2. Ashfaq, M., Yun, J., Yu, S., & Loureiro, S. M. C. (2020). I, chatbot: Modeling the determinants of users' satisfaction and continuance intention of AI-powered service agents. *Telematics and Informatics*, 54, Article 101473. <https://doi.org/10.1016/j.tele.2020.101473>
3. Bitner, M. J., Ostrom, A. L., & Morgan, F. N. (2008). Service blueprinting: A practical technique for service innovation. *California Management Review*, 50(3), 66–94. <https://doi.org/10.2307/41166446>
4. Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101. <https://doi.org/10.1191/1478088706qp063oa>
5. Braun, V., & Clarke, V. (2021). One size fits all? What counts as quality practice in (reflexive) thematic analysis? *Qualitative Research in Psychology*, 18(3), 328–352. <https://doi.org/10.1080/14780887.2020.1769238>
6. Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion. *Journal of Marketing Research*, 56(5), 809–825. <https://doi.org/10.1177/0022243719851788>
7. Chaturvedi, R., Verma, S., Das, R., & Dwivedi, Y. K. (2023). Social companionship with artificial intelligence: Recent trends and future avenues. *Technological Forecasting and Social Change*, 193, Article 122634. <https://doi.org/10.1016/j.techfore.2023.122634>
8. Davenport, T., Guha, A., Grewal, D., & Bressgott, T. (2020). How artificial intelligence will change the future of marketing. *Journal of the Academy of Marketing Science*, 48(1), 24–42. <https://doi.org/10.1007/s11747-019-00696-0>
9. De Keyser, A., Köcher, S., Alkire, L., Verbeeck, C., & Kandampully, J. (2019). Frontline service technology infusion: Conceptual archetypes and future research directions. *Journal of Service Management*, 30(1), 156–183. <https://doi.org/10.1108/JOSM-03-2018-0082>
10. Dietvorst, B. J., & Bartels, D. M. (2022). Consumers object to algorithms making morally relevant tradeoffs because of algorithms' consequentialist decision

strategies. *Journal of Consumer Psychology*, 32(3), 406–424. <https://doi.org/10.1002/jcpy.1266>

11. Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>

12. Dietvorst, B. J., Simmons, J. P., & Massey, C. (2018). Overcoming algorithm aversion: People will use imperfect algorithms if they can (even slightly) modify them. *Management Science*, 64(3), 1155–1170. <https://doi.org/10.1287/mnsc.2016.2643>

13. Etzioni, A., & Etzioni, O. (2017). Incorporating ethics into artificial intelligence. *The Journal of Ethics*, 21(4), 403–418. <https://doi.org/10.1007/s10892-017-9252-2>

14. Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People—An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>

15. Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660. <https://doi.org/10.5465/annals.2018.0057>

16. Gnewuch, U., Morana, S., Adam, M. T. P., & Maedche, A. (2022). Opposing effects of response time in human–chatbot interaction: The moderating role of prior experience. *Business & Information Systems Engineering*, 64, 773–791. <https://doi.org/10.1007/s12599-022-00755-x>

17. Gottman, J. M., Katz, L. F., & Hooven, C. (1996). Parental meta-emotion philosophy and the emotional life of families: Theoretical models and preliminary data. *Journal of Family Psychology*, 10(3), 243–268. <https://doi.org/10.1037/0893-3200.10.3.243>

18. Greiner, D., & Lemoine, J.-F. (2025). Bridging the gap: User expectations for conversational AI services with consideration of user expertise. *Journal of Services Marketing*, 39(2), 76–94. <https://doi.org/10.1108/JSM-02-2024-0056>

19. Grewal, D., Satomino, C. B., Davenport, T., & Guha, A. (2025). How generative AI is shaping the future of marketing. *Journal of the Academy of Marketing Science*, 53, 702–722. <https://doi.org/10.1007/s11747-024-01064-3>

20. Gursoy, D., Chi, O. H., Lu, L., & Nunkoo, R. (2019). Consumers' acceptance of artificially intelligent (AI) device use in service delivery. *International Journal of Information Management*, 49, 157–169. <https://doi.org/10.1016/j.ijinfomgt.2019.03.008>

21. Houmady, S., Chaise, L., Guillot, M., & Rebout, N. (2025). All pet owners are not the same: End-of-life caregiver expectations and profiles. *Topics in Advances in Consumer Research*

Companion Animal Medicine, 65, Article 100960. <https://doi.org/10.1016/j.tcam.2025.100960>

22. Huang, M.-H., & Rust, R. T. (2018). Artificial intelligence in service. *Journal of Service Research*, 21(2), 155–172. <https://doi.org/10.1177/1094670517752459>

23. Huang, M.-H., & Rust, R. T. (2021). A strategic framework for artificial intelligence in marketing. *Journal of the Academy of Marketing Science*, 49(1), 30–50. <https://doi.org/10.1007/s11747-020-00749-9>

24. Jokar, M., Abdous, A., & Rahmanian, V. (2024). AI chatbots in pet health care: Opportunities and challenges for owners. *Veterinary Medicine and Science*, 10(3), Article e1464. <https://doi.org/10.1002/vms3.1464>

25. Kogan, L. R., Bussolari, C., Currin-McCulloch, J., Packman, W., & Erdman, P. (2022). Disenfranchised guilt—Pet owners' burden. *Animals*, 12(13), Article 1690. <https://doi.org/10.3390/ani12131690>

26. Kogan, L. R., Bussolari, C., Currin-McCulloch, J., Packman, W., & Erdman, P. (2023). Dog owners: Disenfranchised guilt and related depression and anxiety. *Human-Animal Interactions*, 2023, 1–16. <https://doi.org/10.1079/hai.2023.0016>

27. Kogan, L. R., Currin-McCulloch, J., Bussolari, C., & Packman, W. (2023). Cat owners' disenfranchised guilt and its predictive value on owners' depression and anxiety. *Human-Animal Interactions*, 2023, Article 0044. <https://doi.org/10.1079/hai.2023.0044>

28. Larivière, B., Bowen, D., Andreassen, T. W., Kunz, W., Sirianni, N. J., Voss, C., Wunderlich, N. V., & De Keyser, A. (2017). “Service encounter 2.0”: An investigation into the roles of technology, employees and customers. *Journal of Business Research*, 79, 238–246. <https://doi.org/10.1016/j.jbusres.2017.03.008>

29. Lemon, K. N., & Verhoef, P. C. (2016). Understanding customer experience throughout the customer journey. *Journal of Marketing*, 80(6), 69–96. <https://doi.org/10.1509/jm.15.0420>

30. Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103. <https://doi.org/10.1016/j.obhdp.2018.12.005>

31. Longoni, C., Bonezzi, A., & Morewedge, C. K. (2019). Resistance to medical artificial intelligence. *Journal of Consumer Research*, 46(4), 629–650. <https://doi.org/10.1093/jcr/ucz013>

32. Luo, X., Tong, S., Fang, Z., & Qu, Z. (2019). Frontiers: Machines versus humans: The impact of artificial intelligence chatbot disclosure on customer purchases. *Marketing Science*, 38(6), 937–947. <https://doi.org/10.1287/mksc.2019.1192>

33. Mariani, M. M., Hashemi, N., & Wirtz, J. (2023). Artificial intelligence empowered conversational agents: A systematic literature review and research agenda. *Journal of Business Research*, 161, Article 113838.

<https://doi.org/10.1016/j.jbusres.2023.113838>

34. Martin, K. (2019). Ethical implications and accountability of algorithms. *Journal of Business Ethics*, 160(4), 835–850. <https://doi.org/10.1007/s10551-018-3921-3>
35. Mayer, J. D., & Gaschke, Y. N. (1988). The experience and meta-experience of mood. *Journal of Personality and Social Psychology*, 55(1), 102–111. <https://doi.org/10.1037/0022-3514.55.1.102>
36. Miceli, M., & Castelfranchi, C. (2019). Meta-emotions and the complexity of human emotional life. *Cognitive Systems Research*, 55, 220–233. <https://doi.org/10.1016/j.cogsys.2019.02.001>
37. Moor, J. H. (2006). The nature, importance, and difficulty of machine ethics. *IEEE Intelligent Systems*, 21(4), 18–21. <https://doi.org/10.1109/MIS.2006.80>
38. Neary, M., Fulton, E., Rogers, V., Wilson, J., Griffiths, Z., Chuttani, R., & Sacher, P. M. (2025). Think FAST: A novel framework to evaluate fidelity, accuracy, safety, and tone in conversational AI health coach dialogues. *Frontiers in Digital Health*, 7, Article 1460236. <https://doi.org/10.3389/fdgh.2025.1460236>
39. Norman, E., & Furnes, B. (2016). The concept of metaemotion: What is there to learn from research on metacognition? *Emotion Review*, 8(2), 187–193. <https://doi.org/10.1177/1754073914552913>
40. Ostrom, A. L., Parasuraman, A., Bowen, D. E., Patrício, L., & Voss, C. A. (2015). Service research priorities in a rapidly changing context. *Journal of Service Research*, 18(2), 127–159. <https://doi.org/10.1177/1094670515576315>
41. Patrício, L., Fisk, R. P., Falcão e Cunha, J., & Constantine, L. (2011). Multilevel service design: From customer value constellation to service experience blueprinting. *Journal of Service Research*, 14(2), 180–200. <https://doi.org/10.1177/1094670511401901>
42. Puntoni, S., Reczek, R. W., Giesler, M., & Botti, S. (2021). Consumers and artificial intelligence: An experiential perspective. *Journal of Marketing*, 85(1), 131–151. <https://doi.org/10.1177/0022242920953847>
43. Radziwill, N. M., & Benton, M. C. (2017). Evaluating quality of chatbots and intelligent conversational agents. *Software Quality Professional*, 19(3), 25–36. <https://doi.org/10.48550/arXiv.1704.04579>
44. Rane, N., Choudhary, S., & Rane, J. (2024). Gemini versus ChatGPT: Applications, performance, architecture, capabilities, and implementation. *Journal of Applied Artificial Intelligence*, 5(1), 69–93. <https://doi.org/10.48185/jaai.v5i1.1052>
45. Sobowale, K., & Humphrey, D. K. (2025). Evaluating the quality of psychotherapy conversational agents: Framework development and cross-sectional study. *JMIR Formative Research*, 9, Article e65605. <https://doi.org/10.2196/65605>
46. Spitznagel, M. B., Marchitelli, B., Gardner, M., & Carlson, M. D. (2020). Euthanasia from the veterinary client's perspective: Psychosocial contributors to euthanasia decision making. *Veterinary Clinics of North America: Small Animal Practice*, 50(3), 591–605. <https://doi.org/10.1016/j.cvsm.2019.12.008>
47. Sundar, S. S. (2020). Rise of machine agency: A framework for studying the psychology of human–AI interaction. *Journal of Computer-Mediated Communication*, 25(1), 74–88. <https://doi.org/10.1093/jcmc/zmz026>
48. van Doorn, J., Mende, M., Noble, S. M., Hulland, J., Ostrom, A. L., Grewal, D., & Petersen, J. A. (2017). Domo arigato Mr. Roboto: Emergence of automated social presence in organizational frontlines and customers' service experiences. *Journal of Service Research*, 20(1), 43–58. <https://doi.org/10.1177/1094670516679272>
49. Vargo, S. L., & Lusch, R. F. (2008). Service-dominant logic: Continuing the evolution. *Journal of the Academy of Marketing Science*, 36(1), 1–10. <https://doi.org/10.1007/s11747-007-0069-6>
50. Wirtz, J., Patterson, P. G., Kunz, W. H., Gruber, T., Lu, V. N., Paluch, S., & Martins, A. (2018). Brave new world: Service robots in the frontline. *Journal of Service Management*, 29(5), 907–931. <https://doi.org/10.1108/JOSM-04-2018-0119>
51. Xu, Y., Shieh, C.-H., van Esch, P., & Ling, I.-L. (2020). AI customer service: Task complexity, problem-solving ability, and usage intention. *Australasian Marketing Journal*, 28(4), 189–199. <https://doi.org/10.1016/j.ausmj.2020.03.005>
52. Yıldırım, M. (2025). Veterinary ethics in practice: Euthanasia decision making for companion and street dogs in Istanbul. *Animals*, 15(17), Article 2585. <https://doi.org/10.3390/ani15172585>