

Louvain School of Management

Quel est l'effet de la tonalité empathique d'un chatbot sur la satisfaction et l'engagement du client dans un contexte de décision financière à forte implication émotionnelle, et par quels mécanismes psychologiques cet effet se produit-il ?

Auteure : Anastassiou Licia
Promotrice : Poncin Ingrid
Année académique 2025-2026
Master [120] en ingénieur de gestion, à finalité spécialisée
Horaire de jour

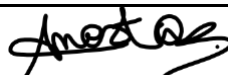
DECLARATION

DECLARATION D'UTILISATION D'IA GÉNÉRATIVE

Confirmez soit les quatre premières déclarations, soit, le cas échéant, la dernière :

Déclarations relatives à l'usage responsable de l'IA générative	Réponse
Le contenu de mon TFE est ma production originale, même si j'ai utilisé l'IA générative pour m'aider à le préparer.	☒ oui ☐ non
Je maîtrise les théories, outils et méthodes que j'utilise pour le mémoire.	☒ oui ☐ non
J'ai vérifié que le contenu présenté dans le mémoire est exact et valide, et qu'il ne contient ni hallucination, ni fausse référence, etc.	☒ oui ☐ non
J'atteste maîtriser le résultat final et être capable de l'expliquer et de répondre aux questions à son sujet.	☒ oui ☐ non
J'ai pris toutes les mesures nécessaires afin de ne pas donner accès à des informations ou données confidentielles à des outils d'IA générative	☒ oui ☐ non

SIGNATURE

	Auteur Nom de famille, prénom(s)	Date	Signature
1	Anastassiou, Licia	4 août 2026	

Résumé

Le déploiement des chatbots dans le secteur bancaire a profondément modifié la façon dont les institutions financières interagissent avec leurs clients. Dans des contextes à forte charge émotionnelle, tels qu'une demande de prêt immobilier, la tonalité employée par ces agents conversationnels, qu'elle soit empathique ou neutre, est susceptible d'influencer la manière dont le service est perçu. Ce mémoire vise à éclairer l'effet réel de cette tonalité empathique sur la satisfaction et l'engagement comportemental des clients, ainsi que les mécanismes psychologiques sous-jacents.

La démarche s'appuie d'abord sur une revue de littérature mobilisant plusieurs cadres théoriques, dont le paradigme CASA, la présence sociale et l'authenticité perçue, afin de construire un modèle conceptuel intégrant ces variables. Une expérimentation en ligne a ensuite été conduite auprès de 141 participants, selon un design factoriel 2×2 croisant la tonalité du chatbot (empathique vs neutre) et la polarité de la réponse (positive vs négative). L'objectif était d'examiner l'impact de ces deux facteurs sur la satisfaction, l'engagement, ainsi que sur des variables intermédiaires telles que la présence sociale perçue, l'authenticité perçue et les dimensions émotionnelles issues du modèle PAD.

Les résultats mettent en évidence plusieurs tendances. La tonalité empathique améliore de manière significative la satisfaction et l'engagement des clients, mais cet effet s'explique entièrement par un renforcement de la présence sociale perçue. En revanche, l'authenticité perçue n'apparaît pas comme un mécanisme explicatif : la tonalité, qu'elle soit empathique ou neutre, ne parvient pas à l'influencer de façon significative. La polarité de la réponse, contrairement aux attentes, ne vient pas modérer l'effet de la tonalité empathique. Sur la base de ces constats, plusieurs recommandations sont proposées aux institutions financières pour la conception de leurs chatbots, tout en soulignant les limites méthodologiques de l'étude et en ouvrant des pistes pour de futures recherches.

Mots-clés : chatbot, tonalité empathique, présence sociale perçue, authenticité perçue, satisfaction client, engagement comportemental, décision financière, modèle PAD.
--

Remerciements

Tout d'abord, je tiens à remercier ma promotrice, Madame Ingrid Poncin, pour son accompagnement tout au long de la réalisation de ce mémoire, ainsi que pour ses précieux conseils et la rigueur dont elle a fait preuve.

Je remercie également Madame Caroline Ducarroz pour ses retours détaillés et constructifs, qui m'ont vraiment aidée à améliorer ce travail et à corriger plusieurs points importants.

Je tiens aussi à remercier toutes les personnes qui ont pris le temps de participer à cette étude. Sans leur contribution, ce mémoire n'aurait pas pu voir le jour.

Enfin, un grand merci à ma famille et à mes amis pour leur soutien tout au long de ce processus, mais aussi pour leur aide dans le partage du questionnaire, qui m'a permis d'atteindre le nombre de réponses souhaité.

Table des matières

I. Introduction	1
II. Revue de littérature	2
1. Les chatbots dans les services financiers	2
2. L'interaction homme-machine	4
3. L'empathie simulée : définition, effets et limites	5
4. Les émotions dans l'expérience de service	7
5. Satisfaction, engagement et authenticité perçue	8
6. Synthèse de la revue de littérature	10
III. Cadre conceptuel	12
1. Cadre conceptuel	12
2. Question de recherche	15
IV. Hypothèses	16
1. Hypothèses à effet direct	16
2. Effets des émotions PAD	17
3. Hypothèses médiatrices	18
4. Hypothèses modératrices	18
V. Méthodologie	20
1. Design expérimental	20
2. Échelles de mesure	20
3. Collecte et traitement des données	22
VI. Analyse des résultats	23
1. Description de l'échantillon	23
2. Fiabilité des échelles de mesure	23
3. Vérification de la manipulation expérimentale	25
4. Tests des hypothèses	25
4.1. Hypothèses à effet direct	25
4.2. Effets des émotions PAD	26
4.3. Hypothèses médiatrices	27
4.4. Hypothèses modératrices	29
4.5. Analyses complémentaires	32
5. Synthèse des résultats	33
VII. Discussions	33
1. Un mécanisme explicatif unique : la présence sociale perçue	34
2. Absence d'effet de la tonalité sur l'authenticité perçue	35
3. Satisfaction et engagement : deux réactions distinctes du client	35
4. Le rôle contrasté des dimensions émotionnelles	36
5. Polarité et éveil : deux logiques distinctes de modération	37

VIII.	<i>Conclusion</i>	38
IX.	<i>Limites et pistes de recherches futures</i>	39
X.	<i>Bibliographie</i>	41
XI.	<i>Annexes</i>	45
1.	Scénarios	45
2.	Questionnaire Qualtrics.....	47
3.	Statistiques descriptives	52
4.	Matrice de corrélation	52
5.	Vérification de la taille des groupes de l'échantillon.....	54
6.	Alpha de Cronbach.....	54
7.	Vérification de la manipulation expérimentale	56
8.	Hypothèse 1.....	57
9.	Hypothèse 2.....	58
10.	Hypothèse 3.....	59
11.	Hypothèse 4.....	60
12.	Hypothèse 5.....	61
13.	Hypothèse 6.....	63
14.	Hypothèse 7.....	64
15.	Hypothèse exploratoire 1	65
16.	Hypothèse 8.....	67
17.	Hypothèse 9.....	68
18.	Hypothèse exploratoire 2	69
19.	Hypothèse 10.....	70
20.	Hypothèse 11.....	73
21.	Hypothèse 12.....	75
22.	Hypothèse 13.....	76

I. Introduction

Les technologies d'intelligence artificielle conversationnelle ont profondément transformé la manière dont les institutions financières interagissent avec leurs clients. Les chatbots, autrefois limités à des fonctions informatives basiques, sont désormais capables de gérer des demandes complexes, y compris des décisions à fort enjeu émotionnel comme les demandes de prêt hypothécaire. Dans ce contexte, la qualité relationnelle de l'interaction devient un enjeu central: comment un chatbot peut-il maintenir, voire renforcer, le lien avec le client dans des situations potentiellement stressantes ou émotionnellement chargées?

La littérature sur l'interaction homme-machine suggère que l'intégration de signaux socio-émotionnels, tels qu'une tonalité empathique, peut améliorer l'expérience utilisateur et les réponses comportementales envers l'institution. Toutefois, les mécanismes précis par lesquels ces signaux opèrent restent insuffisamment documentés, particulièrement dans des contextes financiers à forte implication émotionnelle. Ce mémoire cherche à combler cette lacune en examinant l'effet de la tonalité empathique d'un chatbot sur la satisfaction et l'engagement du client, ainsi qu'en identifiant les mécanismes psychologiques sous-jacents. Plus spécifiquement, cette étude vise à comparer le rôle de la présence sociale perçue et de l'authenticité perçue comme médiateurs potentiels de cet effet, tout en explorant les conditions dans lesquelles ces relations sont plus ou moins fortes.

Pour répondre à cette question de recherche, une approche quantitative a été adoptée. Une expérimentation en ligne a été menée auprès d'un échantillon de participants, selon un design factoriel permettant de manipuler la tonalité du chatbot et la polarité de la réponse. Les hypothèses, formulées sur la base de la revue de littérature, ont été testées à l'aide d'analyses statistiques incluant des tests de médiation et de modération via la macro PROCESS pour SPSS.

Ce mémoire s'organise en quatre parties. La première présente le cadre théorique et la revue de la littérature, qui conduit à la formulation des hypothèses de recherche. La deuxième partie décrit la méthodologie employée, incluant le design expérimental, les instruments de mesure et la procédure de collecte des données. La troisième partie expose les résultats des analyses statistiques. Enfin, la quatrième partie met ces résultats en perspective avec la littérature

existante, en souligne les implications théoriques et pratiques, et propose des pistes de recherche futures.

Il convient de noter que cette recherche présente certaines limites, notamment liées à la taille de l'échantillon et au caractère expérimental du design, qui invitent à une prudence dans la généralisation des résultats. De plus, le contexte spécifique des décisions financières à forte implication émotionnelle, bien que pertinent pour la question de recherche, peut limiter la transférabilité des conclusions à d'autres domaines d'application des chatbots. Ces limites seront discutées plus en détail dans la dernière partie de ce travail.

II. Revue de littérature

1. Les chatbots dans les services financiers

La digitalisation des services financiers a profondément transformé la relation entre les institutions et leurs clients. Dans ce contexte, les chatbots bancaires occupent une place croissante, car ils permettent d'automatiser certaines réponses, d'orienter les usagers et de fluidifier l'accès à l'information dans des environnements de service de plus en plus numérisés (Adamopoulou & Moussiades, 2020 ; CFPB, 2023).

Plus précisément, un chatbot peut être défini comme un programme informatique conçu pour simuler une conversation en langage naturel, le plus souvent par écrit (Adamopoulou & Moussiades, 2020). Dans la littérature, le terme d'agent conversationnel est parfois employé dans un sens plus large, pour désigner des systèmes capables d'interagir de manière plus contextuelle, tandis que le terme chatbot renvoie souvent à des dispositifs plus ciblés, orientés vers des tâches précises (Adamopoulou & Moussiades, 2020 ; Følstad et al., 2018). Dans le cadre de ce mémoire, le terme chatbot sera utilisé pour désigner les agents conversationnels textuels employés dans les services bancaires.

Au-delà de cette définition, tous les chatbots remplissent une fonction informationnelle, mais ils ne produisent pas nécessairement la même expérience d'interaction. Certains adoptent une posture neutre et transactionnelle, centrée sur la rapidité et l'efficacité, tandis que d'autres mobilisent un ton plus relationnel ou empathique. Cette distinction est importante, car le ton adopté peut influencer la manière dont l'utilisateur perçoit le service, ainsi que la qualité globale de l'échange (Han et al., 2022).

Dans le secteur bancaire, ces outils sont déployés pour améliorer la disponibilité du service, réduire les délais de traitement et automatiser une partie du support client. La littérature scientifique montre que la confiance accordée à un chatbot dépend notamment de la qualité de l'interprétation, de la clarté des réponses et de la perception de fiabilité du système (Alagarsamy & Mehroli, 2023 ; Følstad et al., 2018). Han et al. (2022) montrent par ailleurs que l'expression d'empathie par un chatbot peut être bénéfique dans certains contextes, mais qu'elle peut aussi devenir contre-productive si elle est mal ajustée à la situation.

Le Consumer Financial Protection Bureau (CFPB) souligne également que les institutions financières recourent de plus en plus à ce type d'outil pour traiter des demandes courantes, tout en rappelant que des chatbots mal conçus peuvent générer de la frustration chez les consommateurs et affaiblir leur confiance envers l'institution (CFPB, 2023). Le CFPB relève en outre des risques liés à des réponses erronées, à des boucles conversationnelles et à une détérioration de la relation client lorsque l'automatisation est mal encadrée (CFPB, 2023).

Cette problématique prend une dimension particulière dans le cadre d'une demande de prêt immobilier, qui constitue l'objet de cette étude. Une telle situation ne relève pas d'une simple requête d'information, mais d'une décision financière importante, susceptible d'avoir des conséquences durables sur la vie de l'utilisateur. Le contexte est donc plus engageant sur le plan personnel et émotionnel qu'une interaction bancaire ordinaire. Dans ce cadre, le ton adopté par le chatbot peut jouer un rôle déterminant dans la perception du service, notamment lorsque l'utilisateur reçoit une réponse positive ou, au contraire, un refus (Han et al., 2022 ; Alagarsamy & Mehroli, 2023).

Dès lors, la littérature disponible suggère que les chatbots suscitent des réactions variées selon leur conception et le contexte dans lequel ils sont mobilisés. Cependant, les recherches portant spécifiquement sur les chatbots bancaires dans des situations impliquant une décision financière à forte charge émotionnelle demeurent encore limitées (CFPB, 2023 ; Alagarsamy & Mehroli, 2023). Peu d'études examinent en détail ce qui se produit lorsque l'échange porte sur l'acceptation ou le refus d'un prêt immobilier. Ce manque constitue le point de départ de ce mémoire.

2. L'interaction homme-machine

Afin de mieux comprendre ces réactions, il faut s'intéresser aux cadres théoriques qui éclairent l'interaction entre l'utilisateur et le chatbot. Cette relation repose sur plusieurs approches qui permettent d'expliquer pourquoi une machine peut être perçue comme un partenaire d'échange social. Le paradigme CASA, la présence sociale et l'anthropomorphisme constituent les principaux repères mobilisés ici. Ensemble, ils aident à comprendre pourquoi les utilisateurs attribuent parfois des intentions ou des traits relationnels à des systèmes automatisés, même lorsqu'ils savent interagir avec une machine (Reeves & Nass, 1996 ; Epley et al., 2007).

Le *paradigme CASA (Computers Are Social Actors)* repose sur l'idée que les individus appliquent spontanément à certains systèmes informatiques des règles d'interaction habituellement réservées aux relations humaines. Reeves et Nass (1996) montrent que les personnes peuvent répondre à un ordinateur comme elles le feraient avec un interlocuteur humain, sans toujours en avoir une conscience explicite. Cette approche remet donc en cause l'idée selon laquelle une machine serait perçue uniquement comme un outil neutre et fonctionnel, et elle a depuis été confirmée dans divers contextes d'interaction numérique (Nass & Moon, 2000).

Dans le prolongement de cette logique, la notion de *présence sociale* désigne le degré auquel une machine donne l'impression qu'une autre entité est réellement présente dans l'échange. Plus ce sentiment est élevé, plus l'interaction est perçue comme vivante et relationnelle. Dans le cas d'un chatbot, une formulation chaleureuse, une réponse contextualisée ou une prise en compte apparente de la situation du client peuvent renforcer cette présence sociale perçue, même en l'absence d'interlocuteur humain réel (Biocca et al., 2003).

L'*anthropomorphisme* constitue enfin le troisième mécanisme central de ce chapitre. Epley, Waytz et Cacioppo (2007) le définissent comme la tendance à attribuer à des agents non humains des caractéristiques ou des émotions humaines. Cette tendance dépend notamment de la facilité avec laquelle l'individu mobilise des schémas humains pour interpréter un comportement, ainsi que de son besoin de contact social (Epley et al., 2007). Appliqué aux chatbots, ce mécanisme explique pourquoi un agent qui adopte un langage affectif peut être perçu comme plus chaleureux et plus attentif qu'un agent strictement informatif.

Ces trois notions sont complémentaires et forment un cadre cohérent. Le paradigme CASA explique pourquoi les utilisateurs réagissent socialement aux machines ; la présence sociale décrit la sensation d'être en interaction avec une entité présente ; l'anthropomorphisme désigne le processus par lequel des traits humains sont projetés sur l'agent. Ensemble, elles permettent de comprendre pourquoi un chatbot peut influencer l'expérience de l'utilisateur au-delà de son contenu purement informatif (Reeves & Nass, 1996 ; Biocca et al., 2003 ; Epley et al., 2007).

Dans un contexte bancaire à fort enjeu émotionnel, ces mécanismes prennent une portée particulière. Ils constituent ainsi une base théorique essentielle pour analyser l'effet d'une tonalité empathique dans une interaction liée à un prêt immobilier.

3. L'empathie simulée : définition, effets et limites

À partir des cadres théoriques présentés au chapitre précédent, il est maintenant possible d'examiner plus précisément la place de l'empathie dans l'interaction entre l'utilisateur et le chatbot. Avant d'analyser ses effets dans les environnements automatisés, il est toutefois nécessaire de rappeler ce que recouvre l'empathie dans les relations humaines.

Davis (1983) propose l'une des conceptualisations les plus influentes du concept, en distinguant une dimension *cognitive*, qui renvoie à la capacité de comprendre le point de vue d'autrui, et une dimension *affective*, qui désigne la tendance à ressentir ce que l'autre éprouve. À ces deux composantes s'ajoute une dimension *comportementale*, qui correspond à la manière dont ces états intérieurs se traduisent dans la réponse adressée à l'interlocuteur (Davis, 1983). Cette distinction est importante, car elle permet de mieux comprendre ce qui change lorsque l'empathie est exprimée par une machine plutôt que par une personne.

Dans le cas d'un chatbot, il ne s'agit évidemment pas d'une empathie ressentie au sens humain du terme. L'expression *d'empathie simulée* ou *d'empathie artificielle* renvoie plutôt à la capacité d'un système à produire des signaux langagiers qui donnent l'impression d'une écoute attentive et d'une réponse ajustée à l'état émotionnel de l'utilisateur (Liu-Thompkins et al., 2022). Autrement dit, l'agent ne ressent rien, mais il peut être conçu pour reconnaître certains indices émotionnels et formuler une réponse qui paraisse appropriée à la situation.

Les effets de cette empathie simulée sur l'expérience utilisateur ont fait l'objet d'un intérêt croissant dans la littérature. Plusieurs travaux montrent qu'un chatbot exprimant une tonalité empathique est souvent perçu comme plus chaleureux, ce qui peut améliorer l'évaluation générale de l'interaction et renforcer la présence sociale perçue (Han et al., 2022 ; Liu-Thompkins et al., 2022). Dans cette perspective, l'empathie simulée agit comme un signal relationnel qui peut rendre l'échange plus fluide et plus acceptable pour l'utilisateur, notamment lorsque celui-ci traverse une situation désagréable ou incertaine.

Cependant, cette relation n'est pas systématiquement positive. Han et al. (2022) montrent que l'expression d'empathie par un chatbot peut être bénéfique lorsqu'elle répond à une expérience négative vécue par le client, mais qu'elle peut devenir contre-productive lorsque l'agent lui-même est perçu comme la source du problème. Dans ce cas, la tentative d'empathie peut diminuer la compétence perçue du chatbot et nuire à l'évaluation du service. De manière plus générale, Crolig et al. (2022) montrent que lorsque les clients arrivent dans une interaction avec un chatbot dans un état émotionnel négatif, l'humanisation ou anthropomorphisation du chatbot peut au contraire dégrader la satisfaction et les intentions futures, notamment parce que les attentes placées dans l'agent deviennent plus élevées et sont plus facilement déçues.

Ce risque devient d'autant plus visible dans les situations à fort enjeu émotionnel, lorsque la réponse du chatbot intervient dans un contexte où l'utilisateur attend davantage qu'un simple message standardisé. Dans ce type d'échange, l'écart entre la gravité de la situation vécue et un registre empathique trop générique peut accentuer le sentiment d'inadéquation et réduire la perception d'authenticité de l'interaction (Han et al., 2022 ; Seitz, 2024). C'est précisément cette tension qui rend l'étude des chatbots bancaires particulièrement intéressante dans le cadre d'une demande de prêt immobilier.

En effet, une décision relative à un prêt immobilier ne correspond pas à une interaction ordinaire. Il s'agit d'une situation financièrement importante et personnellement engageante, dans laquelle l'utilisateur peut être particulièrement attentif à la manière dont lui est formulée la réponse. Dès lors, l'effet d'un ton empathique ne dépend pas seulement de sa présence, mais aussi de sa pertinence par rapport au contexte dans lequel il est mobilisé (Han et al., 2022 ; Seitz, 2024).

Ainsi, la littérature suggère que l'empathie simulée peut améliorer la qualité perçue de l'interaction dans certaines situations, mais qu'elle peut aussi être mal reçue lorsqu'elle est utilisée dans un cadre inadapté ou lorsque la réponse donnée entre en décalage avec l'attente émotionnelle de l'utilisateur (Han et al., 2022 ; Seitz, 2024 ; Crolic et al., 2022). Dans le cas d'un prêt immobilier, cet enjeu est particulièrement saillant, car le résultat de l'échange peut être vécu comme une bonne ou une mauvaise nouvelle susceptible d'avoir des conséquences concrètes pour le client. C'est ce contexte spécifique que ce mémoire propose d'examiner.

4. Les émotions dans l'expérience de service

À ce stade de la revue de littérature, il est utile de s'intéresser plus directement à la place des émotions dans l'expérience de service. Après avoir montré comment un chatbot peut être perçu comme un acteur relationnel et comment une tonalité empathique peut influencer l'interaction, ce chapitre examine la manière dont l'état émotionnel du client vient colorer l'évaluation du service lui-même.

Les interactions de service génèrent chez le client des états affectifs qui influencent l'ensemble de son évaluation finale, parfois de manière plus déterminante que la qualité objective de la prestation. La littérature en comportement du consommateur a progressivement intégré cette dimension émotionnelle comme un élément structurant de l'expérience de service, la satisfaction ne pouvant pas être réduite à un jugement purement cognitif (Bagozzi et al., 1999). Comprendre comment ces émotions se forment et comment elles orientent le jugement du client est donc essentiel pour analyser ce qui se joue dans une interaction bancaire automatisée.

Pour décrire l'état émotionnel d'un individu au cours d'une interaction, le modèle PAD constitue un cadre particulièrement utile. Il repose sur trois dimensions indépendantes : le plaisir, qui renvoie au caractère agréable ou désagréable de l'expérience ; l'activation, qui décrit le niveau d'éveil psychologique, allant d'un état calme à un état de tension ; et la dominance, qui reflète le sentiment de contrôle ressenti par l'individu (Mehrabian, 1996 ; Russell & Mehrabian, 1977). Ce modèle permet de lire les émotions de façon plus fine qu'une simple opposition entre positif et négatif, car un client anxieux, frustré ou simplement déçu ne se situe pas dans le même état affectif, même si son expérience est globalement négative. Dans le contexte d'une demande de prêt immobilier, cette dimension émotionnelle prend une importance particulière. L'échange se déroule en effet dans une situation marquée par

l'incertitude et un enjeu personnel fort, ce qui tend à amplifier la charge affective de l'interaction.

La littérature montre par ailleurs que les émotions ne se contentent pas d'accompagner l'évaluation du service : elles en orientent aussi la lecture. Bagozzi et al. (1999) soulignent ainsi que les émotions négatives ont souvent un poids plus fort et plus durable que les émotions positives, et qu'elles peuvent influencer la manière dont le client reconstruit l'expérience vécue. Dans cette perspective, une réponse perçue comme froide, décalée ou insuffisamment adaptée peut peser davantage sur l'évaluation finale qu'un simple défaut technique ou informationnel.

C'est précisément ce qui rend un refus de prêt particulièrement sensible. Selon la théorie des violations d'attentes d'Oliver (1980), la satisfaction dépend de la comparaison entre les attentes préalables du client et la performance effectivement perçue. Dans le cas d'une décision de financement, un refus peut être vécu comme une rupture importante avec l'attente d'une issue favorable, ce qui déclenche une réaction affective immédiate avant même que l'interaction puisse être évaluée dans le détail. L'enjeu ne réside donc pas seulement dans le résultat obtenu, mais aussi dans la manière dont ce résultat est annoncé et accueilli par le système.

Ainsi, les émotions constituent un filtre essentiel de l'expérience de service. Elles influencent à la fois la perception de l'échange, l'interprétation du résultat et la manière dont le client évalue l'institution qui l'accompagne. Dans le cadre d'une interaction bancaire automatisée, cette dimension émotionnelle mérite donc d'être prise en compte au même titre que les critères plus classiques de performance ou d'efficacité. C'est cette articulation entre état affectif, perception du service et jugement du résultat qui prépare l'analyse des construits mobilisés dans le chapitre suivant.

5. Satisfaction, engagement et authenticité perçue

À partir des cadres précédents, trois construits apparaissent centraux dans l'évaluation d'une interaction avec un chatbot : la satisfaction perçue, l'engagement comportemental et l'authenticité perçue. Ces dimensions permettent de mieux comprendre comment l'utilisateur juge le service, dans quelle mesure il accepte de poursuivre l'échange et à quel point il considère la réponse de l'agent comme crédible.

La *satisfaction* dans les services a longtemps été envisagée comme un construit principalement cognitif, fondé sur l'écart entre ce qui était attendu et ce qui a été effectivement reçu. Le modèle SERVQUAL de Parasuraman, Zeithaml et Berry (1988) constitue à cet égard une référence importante. Il propose de mesurer la qualité perçue d'un service à partir de cinq dimensions : la tangibilité, la fiabilité, la réactivité, l'assurance et l'empathie. Cette dernière dimension est particulièrement intéressante dans le cadre d'un chatbot, car elle rappelle que la qualité perçue d'un service ne dépend pas seulement de son efficacité technique, mais aussi de la manière dont l'utilisateur se sent considéré au cours de l'échange.

Cependant, la satisfaction ne se réduit pas à une évaluation rationnelle du service. Bagozzi et al. (1999) montrent que les émotions jouent un rôle déterminant dans la formation du jugement du consommateur, et que les réactions affectives peuvent peser lourdement sur l'évaluation finale d'une interaction. Dans un contexte bancaire, cette dimension est d'autant plus importante que l'expérience peut être vécue dans un état de tension ou d'incertitude. La satisfaction résulte alors moins d'un simple bilan fonctionnel que de la manière dont le client a ressenti l'échange dans son ensemble.

Le deuxième construit central est *l'engagement*, entendu ici comme la disposition à poursuivre l'échange et à maintenir une relation avec l'institution représentée par le chatbot. Dans la littérature sur les relations de service, cet engagement est souvent associé à la confiance et à la qualité de l'interaction perçue. Moorman, Zaltman et Deshpande (1992) montrent ainsi que la confiance contribue fortement à la stabilité de la relation entre un utilisateur et un fournisseur de service. Dans le cas d'un chatbot, cette logique est particulièrement pertinente : la disposition à poursuivre l'échange dépend en effet de la manière dont l'agent parvient à inspirer sérieux, fiabilité et proximité relationnelle.

Des travaux plus récents montrent que les caractéristiques du chatbot et le contexte d'utilisation influencent également la manière dont l'interaction est évaluée. Markovitch et al. (2024) observent notamment que les réactions des consommateurs varient selon qu'ils interagissent avec un humain ou un chatbot, ainsi que selon la valence positive ou négative du résultat de service. Dans cette perspective, l'engagement ne renvoie pas uniquement à une intention générale de rester lié à une marque, mais aussi à la manière dont l'utilisateur accepte de prolonger l'échange dans une situation donnée.

C'est ici que *l'authenticité perçue* prend toute son importance. Ce construit renvoie à la manière dont l'utilisateur juge la réponse du chatbot comme sincère, crédible et cohérente avec la situation. L'enjeu n'est pas seulement que l'agent réponde vite ou correctement, mais qu'il donne l'impression de formuler une réponse qui corresponde au contexte dans lequel il intervient (Neururer et al., 2018). Dans un échange bancaire, un message peut ainsi être techniquement exact tout en paraissant standardisé ou trop générique, ce qui peut réduire sa portée relationnelle.

Plusieurs travaux montrent par ailleurs que les indices anthropomorphiques et la manière dont le chatbot est présenté influencent la perception de l'échange. Crollic et al. (2022) soulignent que l'humanisation d'un chatbot modifie les attentes des utilisateurs et peut peser sur leur évaluation lorsque le contexte émotionnel est négatif. Dans le même sens, Ozuem et al. (2024) montrent que, dans les situations où le chatbot intervient pour répondre à un problème ou à une insatisfaction du client, la manière dont il formule sa réponse peut influencer la fidélité et l'évaluation du service. La qualité de l'échange dépend donc à la fois de la formulation du message et de son adéquation au contexte dans lequel il est reçu (Crollic et al., 2022 ; Ozuem et al., 2024).

Dans ce cadre, satisfaction, engagement et authenticité perçue ne doivent pas être considérés comme trois dimensions indépendantes, mais comme des construits étroitement liés. Une interaction jugée satisfaisante peut encourager la poursuite de l'échange, tandis qu'une réponse perçue comme authentique peut renforcer la confiance accordée à l'institution et soutenir une évaluation plus positive du service. À l'inverse, lorsqu'un message paraît trop générique ou peu crédible, il risque de fragiliser à la fois la satisfaction immédiate et la disposition à maintenir la relation.

6. Synthèse de la revue de littérature

Les chapitres précédents ont permis de mettre en place un cadre théorique progressif autour des chatbots bancaires, de l'interaction homme-machine, de l'empathie simulée, des émotions de service et des principaux construits associés à l'évaluation de l'échange. Pris ensemble, ces apports montrent que l'analyse d'un chatbot ne peut pas se limiter à sa performance technique, mais doit aussi tenir compte de la manière dont il est perçu dans une interaction donnée (Adamopoulou & Moussiades, 2020 ; Reeves & Nass, 1996 ; Markovitch et al., 2024).

Un premier élément ressort de cette revue : les chatbots sont souvent adoptés pour des raisons d'efficacité, de disponibilité et de rapidité, mais les utilisateurs n'en font pas pour autant une lecture purement fonctionnelle. La littérature sur les agents conversationnels montre au contraire que les attentes portent aussi sur la clarté du message, la qualité de l'échange et la capacité du système à s'aligner sur les besoins de l'utilisateur (Adamopoulou & Moussiades, 2020 ; Markovitch et al., 2024). Autrement dit, l'évaluation d'un chatbot se situe à l'intersection entre logique instrumentale et logique relationnelle.

Un deuxième point concerne le rôle du ton adopté par l'agent. Les travaux sur l'interaction homme-machine et l'empathie simulée montrent que le style d'expression du chatbot influence la perception de l'échange, notamment lorsque l'utilisateur s'attend à une réponse adaptée à sa situation (Han et al., 2022 ; Crolie et al., 2022 ; Ozuem et al., 2024). Toutefois, la littérature souligne aussi que cet effet dépend du contexte dans lequel l'échange se déroule : un même style d'interaction peut être jugé pertinent dans une situation ordinaire et moins appropriée dans un contexte plus sensible (Han et al., 2022 ; Ozuem et al., 2024).

Cette question est particulièrement importante dans le domaine bancaire, et plus encore lorsqu'il est question d'un prêt immobilier. Dans ce type de situation, l'échange ne correspond pas à une simple demande d'information, mais à une situation à fort enjeu décisionnel et émotionnel pour l'utilisateur (Consumer Financial Protection Bureau, 2023 ; Markovitch et al., 2024). La littérature suggère ainsi que la perception du chatbot peut varier selon que la réponse annoncée va dans le sens attendu ou non, ce qui confirme l'intérêt de s'intéresser à des situations de service où le résultat a une portée concrète pour le client (Markovitch et al., 2024 ; CFPB, 2023).

Les travaux sur les émotions de service permettent également de comprendre pourquoi ces situations méritent une attention particulière. L'état affectif du client influence la manière dont il interprète la réponse reçue, et cette dimension émotionnelle peut peser lourdement sur l'évaluation de l'échange (Bagozzi et al., 1999 ; Russell & Mehrabian, 1977 ; Mehrabian, 1996). Dans une interaction bancaire, surtout lorsque l'enjeu est élevé, l'émotion ne se limite donc pas à accompagner le jugement du client : elle participe à la manière dont ce jugement se construit (Bagozzi et al., 1999 ; Oliver, 1980).

Enfin, la littérature sur la confiance et l'authenticité perçue montre qu'un chatbot ne sera pas évalué uniquement sur sa capacité à fournir une réponse correcte, mais aussi sur la crédibilité et la cohérence qu'il parvient à projeter (Følstad et al., 2018 ; Neururer et al., 2018 ; Alagarsamy & Mehroliia, 2023). Plusieurs travaux indiquent que la perception du chatbot dépend de la façon dont l'utilisateur interprète ses signaux relationnels, son langage et sa place dans l'échange (Følstad et al., 2018 ; Alagarsamy & Mehroliia, 2023 ; Neururer et al., 2018). Cela confirme que l'étude d'un chatbot bancaire en situation sensible nécessite de prendre en compte à la fois l'expérience immédiate de l'utilisateur et la manière dont il attribue du sens à l'interaction (Alagarsamy & Mehroliia, 2023 ; Neururer et al., 2018).

Au terme de cette revue de littérature, il apparaît donc que les chatbots bancaires se situent à la croisée de plusieurs logiques : efficacité, relation, émotion et crédibilité (Adamopoulou & Moussiades, 2020 ; Bagozzi et al., 1999 ; Neururer et al., 2018). C'est précisément cette combinaison qui justifie l'étude d'une interaction portant sur un prêt immobilier, car elle permet d'observer comment un agent conversationnel est reçu lorsqu'il intervient dans une situation à forte implication personnelle (CFPB, 2023 ; Markovitch et al., 2024).

III. Cadre conceptuel

1. Cadre conceptuel

La revue de littérature menée dans la partie précédente a mis en évidence que l'évaluation d'un chatbot bancaire ne dépend pas seulement de la qualité objective de sa réponse, mais aussi de la manière dont cette réponse est formulée, interprétée et située dans un contexte de service à forte implication émotionnelle. Dans le prolongement de ces travaux, ce mémoire cherche à comprendre dans quelle mesure la tonalité d'un chatbot influence l'expérience d'un client confronté à une décision financière défavorable. Le cadre conceptuel proposé vise ainsi à formaliser les relations entre les principales variables identifiées dans la littérature, afin de les transformer, par la suite, en hypothèses empiriquement testables (Adamopoulou & Moussiades, 2020 ; Markovitch et al., 2024).

Le modèle conceptuel repose sur **une variable indépendante** principale : la *tonalité du chatbot*. Deux modalités sont distinguées. La première est une *tonalité empathique*, dans laquelle le chatbot reconnaît explicitement la situation émotionnelle de l'utilisateur et adapte son langage en conséquence. La seconde est une *tonalité neutre*, dans laquelle le chatbot délivre une réponse

strictement informationnelle, sans signaux affectifs particuliers. Cette distinction permet d'isoler l'effet spécifique de l'expression empathique sur l'expérience d'interaction, comme le suggèrent les travaux sur l'empathie artificielle et les réactions des consommateurs face aux chatbots (Liu-Thompkins et al., 2022 ; Han et al., 2022 ; Crollic et al., 2022).

Le modèle mobilise ensuite **deux variables dépendantes**. La première est la *satisfaction perçue* vis-à-vis de l'interaction avec le chatbot. Elle renvoie à l'évaluation globale que le client porte sur l'échange, en fonction du résultat obtenu, de la manière dont il a été traité et de l'écart éventuel entre ses attentes et ce qu'il perçoit réellement du service (Oliver, 1980 ; Bagozzi et al., 1999). La seconde variable dépendante est *l'engagement comportemental* envers l'institution représentée par le chatbot. Il correspond ici à la disposition du client à poursuivre la relation, à rester en lien avec la marque ou à continuer d'interagir avec elle à la suite de l'échange (Moorman et al., 1992 ; Markovitch et al., 2024).

Deux variables médiatrices structurent également le mécanisme explicatif du modèle. La première est *l'authenticité perçue*. Elle renvoie au jugement porté par l'utilisateur sur le fait que la réponse du chatbot paraît sincère, cohérente et adaptée au contexte dans lequel elle est formulée. Dans les services automatisés, une réponse peut être correcte sur le fond tout en paraissant trop générique, artificielle ou peu crédible, ce qui fragilise l'évaluation de l'échange (Neururer et al., 2018 ; Ozuem et al., 2024). Dans ce mémoire, l'authenticité perçue est donc considérée comme le mécanisme par lequel la tonalité du chatbot influence l'engagement comportemental.

La seconde variable médiatrice est la *présence sociale perçue*, c'est-à-dire le sentiment que l'utilisateur interagit avec une entité réellement présente et attentive, plutôt qu'avec un simple dispositif technique. Dans le cas d'un chatbot empathique, cette présence sociale peut être renforcée par des formulations chaleureuses, contextualisées et relationnelles, ce qui rend l'échange plus vivant et plus engageant (Biocca et al., 2003 ; Reeves & Nass, 1996). Elle est ici postulée comme un mécanisme intermédiaire reliant la tonalité du chatbot à la satisfaction perçue (Biocca et al., 2003).

Le cadre conceptuel intègre également **deux variables de contexte** qui jouent un rôle de modération. La première est la *polarité de la réponse*, c'est-à-dire le caractère favorable ou défavorable de l'issue communiquée par le chatbot. Cette variable est essentielle, car la

littérature montre que l'effet d'un message empathique peut varier selon que la situation vécue par l'utilisateur est positive ou négative (Crolic et al., 2022 ; Han et al., 2022). Dans un contexte de refus de prêt, on peut ainsi s'attendre à ce que la tonalité empathique soit évaluée différemment qu'en cas de réponse favorable.

La seconde variable de contexte est *l'éveil émotionnel* du client, entendu ici comme l'intensité de l'activation affective suscitée par la situation. Cette dimension permet de comprendre qu'un client confronté à une réponse défavorable peut se montrer plus sensible, plus réactif et moins réceptif aux signaux relationnels du chatbot. Les travaux sur l'éveil et l'expérience émotionnelle indiquent que l'état d'activation influence la manière dont une interaction est interprétée et évaluée, ce qui s'inscrit plus largement dans la logique du modèle PAD, fondé sur les dimensions de plaisir, d'éveil et de dominance (Russell & Mehrabian, 1977 ; Mehrabian, 1996). Dans ce cadre, l'éveil émotionnel est principalement mobilisé comme une variable de contexte susceptible d'amplifier ou d'atténuer les effets de la tonalité du chatbot sur la satisfaction, tandis que les autres dimensions du modèle PAD constituent un repère théorique pour situer cette variable dans une approche plus globale de l'état affectif (Russell & Mehrabian, 1977 ; Mehrabian, 1996).

L'ensemble de ces relations s'inscrit dans un modèle intégrant des effets directs, des médiations et des modérations, dans la lignée des travaux de Hayes (2018). La tonalité du chatbot influence la satisfaction et l'engagement par l'intermédiaire de la présence sociale perçue et de l'authenticité perçue, tandis que la polarité de la réponse et l'éveil émotionnel du client viennent conditionner certaines de ces relations. Autrement dit, l'effet d'un chatbot empathique n'est ni uniforme ni automatique ; il varie selon la manière dont l'échange est reçu et selon le contexte affectif dans lequel il s'inscrit (Hayes, 2018 ; Markovitch et al., 2024).

Ce cadre intégratif est schématisé ci-dessous et se distingue de ceux déjà existants dans la littérature sur les chatbots empathiques à plusieurs niveaux. D'une part, il intègre simultanément deux médiateurs distincts pour les deux variables dépendantes, reconnaissant que les déterminants de la satisfaction et de l'engagement ne sont pas nécessairement les mêmes. D'autre part, il place l'octroi ou le refus de prêt immobilier au cœur du dispositif empirique, ce qui permet d'étudier une situation à forte charge émotionnelle et implication personnelle, encore peu explorée dans la littérature. Enfin, il transforme le paradoxe de l'empathie en

contexte négatif en hypothèse d'interaction testable, plutôt qu'en simple limite théorique (Han et al., 2022 ; Crollic et al., 2022 ; Ozuem et al., 2024).

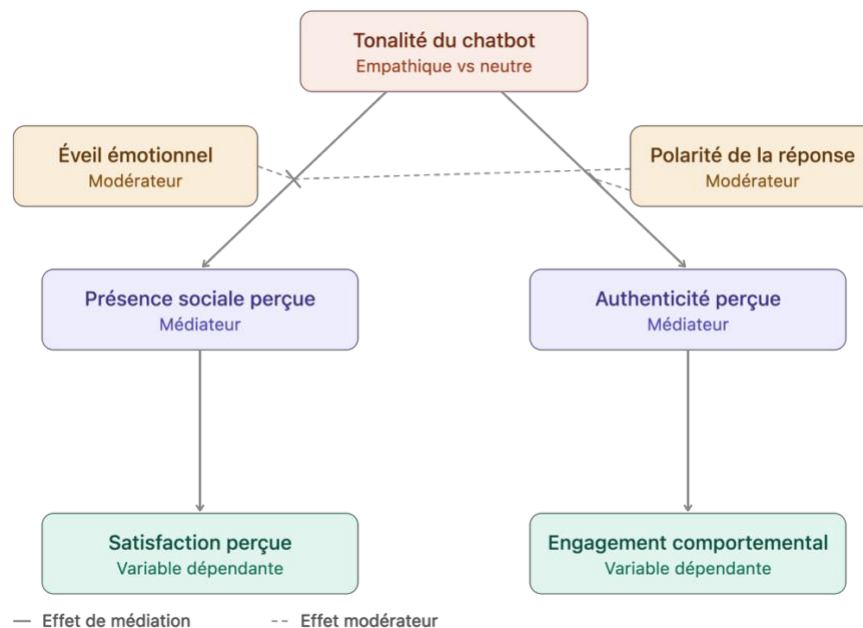


Schéma 1 : Cadre conceptuel de l'étude

2. Question de recherche

À partir de ce cadre conceptuel, la question de recherche centrale de ce mémoire peut être formulée comme suit :

Quel est l'effet de la tonalité empathique d'un chatbot sur la satisfaction et l'engagement du client dans un contexte de décision financière à forte implication émotionnelle, et par quels mécanismes psychologiques cet effet se produit-il ?

Cette question présente plusieurs intérêts. Elle porte d'abord sur un contexte encore peu étudié empiriquement, à la croisée des recherches sur les interfaces conversationnelles, la confiance et l'expérience émotionnelle en service (Følstad et al., 2018 ; Alagarsamy & Mehroli, 2023). Elle met ensuite à l'épreuve une tension théorique bien documentée : l'empathie simulée peut améliorer l'évaluation d'un échange, mais elle peut aussi être contre-productive lorsque le contexte est défavorable ou émotionnellement chargé (Han et al., 2022 ; Crollic et al., 2022 ; Seitz, 2024). Enfin, elle suppose de tester un ensemble de relations directes, médiatrices et modératrices, ce qui nécessite un protocole empirique capable d'isoler l'effet de la tonalité tout en tenant compte de la polarité de la réponse et de l'état émotionnel des participants (Hayes, 2018).

IV. Hypothèses

Les hypothèses ont été organisées selon une logique progressive : d’abord les effets directs de la tonalité du chatbot et de la polarité de la réponse sur la satisfaction et l’engagement comportemental, puis l’influence des émotions sur la satisfaction, et enfin les mécanismes de médiation et de modération. Elles sont toutes formulées dans le contexte d’une demande de prêt hypothécaire, situation chargée émotionnellement où l’usager attend une décision (acceptation ou refus) communiquée par un chatbot.

1. Hypothèses à effet direct

Hypothèse 1 : *Un chatbot à tonalité empathique génère une satisfaction plus élevée qu’un chatbot à tonalité neutre.*

Hypothèse 2 : *Un chatbot à tonalité empathique génère un engagement comportemental (envers l’institution représentée par le chatbot) plus élevé qu’un chatbot à tonalité neutre.*

Justification :

La littérature sur l’interaction homme-machine montre que les individus ont tendance à traiter les machines comme des acteurs sociaux dès lors qu’elles mobilisent des indices de langage ou de comportement proches de ceux des relations humaines (Reeves & Nass, 1996 ; Nass & Moon, 2000). Dans cette logique, un chatbot qui exprime une tonalité empathique peut renforcer la perception de prise en charge, réduire la distance relationnelle et améliorer la qualité perçue de l’échange, ce qui favorise à la fois la satisfaction immédiate et la disposition à poursuivre l’interaction avec l’institution représentée (Biocca et al., 2003 ; Liu-Thompkins et al., 2022 ; Markovitch et al., 2024). Dans un contexte de service à forte implication émotionnelle, une tonalité empathique devrait donc agir simultanément sur l’évaluation de l’échange et sur l’intention de prolonger la relation.

Hypothèse 3 : *Une réponse à polarité négative génère une satisfaction plus faible qu’une réponse à polarité positive.*

Hypothèse 4 : *Une réponse à polarité négative génère un engagement comportemental plus faible qu’une réponse à polarité positive.*

Justification :

Selon la théorie des violations d'attentes, la satisfaction dépend de l'écart entre l'issue anticipée et l'issue effectivement obtenue, de sorte qu'un résultat défavorable provoque plus facilement une déception qu'un résultat favorable (Oliver, 1980). Dans les services automatisés, la valence du résultat communiqué par le chatbot joue ainsi un rôle central dans l'évaluation du service et dans la volonté de continuer la relation avec l'institution (Markovitch et al., 2024). Les travaux de Crolig et al. (2022) montrent également qu'une mauvaise nouvelle peut susciter une réaction émotionnelle forte et réduire l'ouverture du client à la suite de l'échange. Une réponse négative devrait donc peser à la fois sur la satisfaction et sur l'engagement comportemental, car elle affecte simultanément l'expérience vécue et la projection vers une interaction future.

2. Effets des émotions PAD

***Hypothèse 5a :** Un niveau élevé de plaisir ressenti pendant l'interaction est positivement associé à la satisfaction du client.*

***Hypothèse 5b :** Un niveau élevé d'éveil émotionnel pendant l'interaction est négativement associé à la satisfaction du client.*

***Hypothèse 5c :** Un niveau élevé de dominance ressentie pendant l'interaction est positivement associé à la satisfaction du client.*

Justification :

Le modèle PAD de Mehrabian et Russell (1977) considère ces trois dimensions comme complémentaires dans la structuration de l'expérience affective, et la littérature montre qu'elles influencent l'évaluation de l'interaction. Le plaisir reflète le caractère agréable de l'expérience et entretient logiquement un lien positif avec la satisfaction. La dominance correspond au sentiment de contrôle perçu, qui favorise généralement une appréciation plus positive du service. L'éveil émotionnel renvoie quant à lui au niveau d'activation ou de tension ressenti pendant l'échange, ce qui peut perturber une évaluation sereine et réduire la satisfaction dans un contexte émotionnellement chargé (Mehrabian & Russell, 1977 ; Mehrabian, 1996 ; Russell & Mehrabian, 1977). Dans cette logique, le PAD permet de saisir plus finement la manière dont l'état émotionnel du client se traduit dans son jugement du chatbot.

3. Hypothèses médiatrices

Hypothèse 6 : *L'authenticité perçue médie la relation entre la tonalité du chatbot et l'engagement comportemental du client.*

Justification :

La littérature montre que les utilisateurs s'engagent davantage lorsqu'ils jugent l'échange sincère, cohérent et crédible, plutôt que simplement scripté ou artificiel (Neururer et al., 2018 ; Liu-Thompkins et al., 2022). Dans cette perspective, une tonalité empathique ne favorise pas l'engagement uniquement parce qu'elle est chaleureuse, mais parce qu'elle augmente la perception d'authenticité de l'interaction, ce qui renforce ensuite la disposition à poursuivre l'échange avec l'institution. L'authenticité perçue constitue donc un mécanisme explicatif essentiel entre le style du chatbot et l'engagement comportemental.

Hypothèse 7 : *La présence sociale perçue médie la relation entre la tonalité du chatbot et la satisfaction du client.*

Justification :

Biocca et al. (2003) définissent la présence sociale comme le sentiment d'interagir avec une entité réellement présente et attentive, et montrent que cette perception joue un rôle important dans l'évaluation positive des interactions médiées par la technologie. Dans le cas des chatbots, les signaux relationnels et empathiques renforcent cette impression d'échange social, ce qui améliore la satisfaction de l'utilisateur (Biocca et al., 2003 ; Liu-Thompkins et al., 2022). Une tonalité empathique peut donc agir indirectement sur la satisfaction en rendant l'échange plus vivant, plus incarné et plus socialement engageant.

4. Hypothèses modératrices

Hypothèse 8 : *La polarité de la réponse modère l'effet de la tonalité du chatbot sur la satisfaction, de sorte que l'effet d'une tonalité empathique est plus fort lorsque la réponse est négative.*

Justification :

Les travaux sur l'empathie artificielle indiquent que les signaux empathiques prennent davantage de valeur lorsque l'utilisateur fait face à une situation émotionnellement chargée, car

ils contribuent alors à rendre la réponse moins froide et plus adaptée à la situation (Liu-Thompkins et al., 2022). Markovitch et al. (2024) montrent également que la perception d'un chatbot dépend fortement de la valence du résultat communiqué, ce qui confirme que l'efficacité d'un ton empathique n'est pas uniforme. Dans une situation de réponse négative, la tonalité empathique peut donc jouer un rôle de réassurance plus marqué et compenser partiellement l'impact défavorable de l'issue sur la satisfaction.

Hypothèse 9 : *La polarité de la réponse modère l'effet de la tonalité du chatbot sur l'engagement comportemental, de sorte que l'effet d'une tonalité empathique est plus fort lorsque la réponse est négative.*

Justification :

Les travaux sur les chatbots suggèrent que les signaux empathiques ne sont pas interprétés de la même manière selon que l'échange se déroule dans un contexte favorable ou défavorable. Adam et al. (2021) montrent que l'effet de la tonalité empathique sur les réponses comportementales tend à être amplifié dans des contextes négatifs, où le besoin de soutien relationnel est accru. À l'inverse, Crolie et al. (2022) indiquent que, dans des situations défavorables, les expressions empathiques peuvent être perçues comme moins sincères ou moins crédibles, ce qui pourrait atténuer leur effet sur l'engagement. Ces résultats divergents justifient l'intérêt de tester empiriquement cette relation de modération dans le contexte spécifique d'un refus de prêt hypothécaire.

Hypothèse 10 : *L'éveil émotionnel du client modère l'effet de la tonalité du chatbot sur la satisfaction, de sorte que l'effet positif d'une tonalité empathique est plus faible lorsque l'éveil émotionnel est élevé.*

Justification :

Les travaux sur le modèle PAD montrent qu'un état émotionnel plus intense peut réduire la réceptivité aux signaux relationnels et compliquer l'évaluation positive d'une interaction (Mehrabian, 1996 ; Russell & Mehrabian, 1977). Dans cette perspective, l'empathie du chatbot devrait perdre en efficacité lorsque le client se trouve déjà dans un état d'activation élevé, ce qui justifie un effet de modération de l'éveil émotionnel sur le lien entre tonalité et satisfaction.

V. Méthodologie

1. Design expérimental

Afin de tester l'ensemble du cadre conceptuel, l'étude adopte une approche expérimentale quantitative reposant sur un design factoriel 2×2 between-subjects. Ce type de design permet d'examiner simultanément l'effet du type de chatbot et de la polarité de la réponse, ainsi que leur combinaison. Dans cette logique, le questionnaire est conçu sur *Qualtrics* (cf annexe 2) afin d'assigner aléatoirement chaque participant à l'un des quatre scénarios expérimentaux (cf annexe 1), selon une randomisation équilibrée renforçant la validité interne du dispositif (Montgomery, 2017 ; Qualtrics, s. d. ; Shadish et al., 2002).

Le scénario de base invite ensuite les participants à se projeter dans une situation de demande de prêt immobilier pour l'achat d'une maison. L'interaction présentée porte sur la décision finale de la banque, à savoir l'acceptation ou le refus du prêt, et varie selon deux dimensions : l'attitude du chatbot, empathique ou neutre, et l'issue de la demande, favorable ou défavorable. Chaque participant n'est exposé qu'à une seule version du scénario, ce qui permet de comparer les conditions sans effet de familiarité ni de comparaison entre plusieurs versions du même stimulus (Montgomery, 2017 ; Shadish et al., 2002).

2. Échelles de mesure

Les construits mobilisés dans cette étude sont mesurés à l'aide d'échelles de Likert en sept points (Likert, 1932), à l'exception des dimensions du modèle PAD, qui reposent sur des items bipolaires de type différentiel sémantique (Russell & Mehrabian, 1977 ; Mehrabian, 1996). Cette combinaison d'échelles est adaptée à la mesure d'attitudes, de perceptions et d'états émotionnels dans un contexte d'interaction médiée par la technologie (Russell & Mehrabian, 1977 ; Wirtz & Lee, 2003 ; Detandt et al., 2017).

L'empathie perçue est appréhendée par trois items portant sur la compréhension de la situation émotionnelle du répondant, l'adéquation de la réaction du chatbot et l'expression de sollicitude. Cette mesure s'appuie sur les travaux portant sur l'empathie perçue dans les interactions avec des systèmes conversationnels. Comme il n'existe pas encore de mesure unique unanimement validée, une adaptation contextuelle du construit est nécessaire. (Concannon & Tomalin, 2023 ; Huang et al., 2022 ; Li et al., 2024). L'item relatif à l'efficacité de la réponse est traité

séparément, car il relève davantage de la performance fonctionnelle de l'agent que de son empathie perçue.

Le bloc PAD mesure trois dimensions affectives distinctes : le plaisir, à travers des paires bipolaires telles que heureux/malheureux, satisfait/insatisfait, à l'aise/mal à l'aise et soulagé/anxieux ; l'éveil, à travers agité/calme, tendu/détendu, stimulé/apaisé et stressé/serein; et la dominance, à travers maître de la situation/dépassé, en contrôle/sans contrôle, dominant/soumis et autonome/dépendant. Cette mesure s'inscrit dans la tradition du modèle PAD, référence classique pour appréhender les réponses affectives à un stimulus ou à un environnement (Mehrabian, 1996 ; Detandt et al., 2017).

L'authenticité perçue est mesurée à l'aide d'items portant sur la sincérité des expressions émotionnelles du chatbot, la réalité des émotions exprimées, ainsi que la perception d'une éventuelle feinte émotionnelle ou d'une intention stratégique. Cette opérationnalisation s'inscrit dans les travaux récents qui abordent l'authenticité perçue dans les interactions avec des chatbots et des agents conversationnels, tout en laissant une marge d'adaptation au contexte étudié (Huang et al., 2022 ; Lee, 2020).

La présence sociale perçue est évaluée au moyen d'items visant à saisir le sentiment d'interagir avec une entité sociale, la quasi-disparition du caractère machinique de l'échange et le sentiment de contact humain. Cette mesure renvoie à la littérature sur la présence sociale, définie comme le sentiment qu'un interlocuteur est socialement "réel" dans un environnement médié par technologie (Lombard & Ditton, 1997 ; Gefen & Straub, 2004 ; Biocca et al., 2003). Un item inversé mesurant un éventuel malaise est également conservé à titre exploratoire, mais doit être interprété avec prudence et n'est pas inclus dans le calcul du score de présence sociale perçue.

La satisfaction constitue la variable dépendante principale et est mesurée à travers plusieurs items portant sur la satisfaction globale, la réduction du stress, la confiance envers l'institution, l'intention de continuer à utiliser le service et la recommandation du service à l'entourage. Cette logique s'inscrit dans la littérature sur la satisfaction en contexte de service, qui recommande des mesures multidimensionnelles adaptées au contexte d'usage (Wirtz & Lee, 2003) .

Enfin, un second bloc relatif à l'engagement du patient mesure des intentions comportementales telles que le fait de suivre les recommandations, de poser des questions supplémentaires, de

partager des informations ou de solliciter un conseiller humain. Ces items renvoient davantage à des intentions d'engagement ou de réactivité comportementale qu'à un comportement effectif, ce qui est cohérent avec la littérature sur l'engagement en services et en santé, où les intentions sont souvent utilisées comme indicateur du comportement (Brodie et al., 2011 ; Ajzen, 1991 ; Hibbard et al., 2004).

3. Collecte et traitement des données

La collecte des données est réalisée au moyen d'un questionnaire en ligne administré sur Qualtrics. Le questionnaire est structuré en plusieurs sections (cf. annexe 2) : une page d'accueil présentant l'étude, un bloc scénarisé randomisé, puis un ensemble de questions communes portant sur les construits d'intérêt. Afin de renforcer l'immersion, le scénario est présenté sous forme d'une courte vidéo où les messages du chatbot s'affichent progressivement plutôt que sous la forme d'un texte statique. Deux questions de contrôle de l'attention sont intégrées afin d'identifier les répondants inattentifs. Cette procédure vise à garantir une administration uniforme des stimuli et des mesures, tout en limitant les biais de compréhension et de fatigue (Qualtrics, s. d.).

Après la collecte, les données sont préparées pour l'analyse statistique. Les questionnaires incomplets et ceux dont les répondants n'ont pas répondu correctement aux questions de contrôle de l'attention sont exclus. Les items formulés de manière négative sont ensuite recodés, de façon à ce que les valeurs les plus élevées correspondent toujours à une interprétation plus positive et cohérente avec le sens général de l'échelle. Cette étape permet d'éviter que certaines réponses soient inversées par rapport aux autres et de garantir une lecture homogène des scores obtenus (DeVellis, 2017 ; Hair et al., 2019 ; Streiner, 2003).

Le traitement statistique comprend ensuite le calcul de scores composites pour chaque construit, en combinant les réponses aux différents items sous forme de moyennes. Des statistiques descriptives sont produites pour l'ensemble de ces variables, afin d'examiner leurs distributions et de vérifier que les valeurs observées restent plausibles au regard des échelles utilisées (cf. annexe 3). Une matrice de corrélations est également établie pour ces construits (cf. annexe 4). Elle permet d'examiner les relations bivariées entre les variables et de vérifier qu'aucun prédicteur n'est trop fortement corrélé aux autres. Les coefficients restant en dessous du seuil de 0,80 confirment l'absence de multicolinéarité dans les analyses ultérieures.

VI. Analyse des résultats

Ce chapitre présente les principaux résultats empiriques de l'étude. Il commence par la description de l'échantillon, avant d'aborder la fiabilité des échelles de mesure, la vérification de la manipulation expérimentale et les tests des hypothèses. L'ensemble des analyses empiriques a été réalisé à l'aide du logiciel SPSS (version 30). Pour des raisons de lisibilité, seuls les résultats principaux sont présentés dans le corps du mémoire, tandis que les tableaux statistiques détaillés et les analyses complémentaires sont repris en annexe (*cf. annexes 5-22*).

1. Description de l'échantillon

Les données ont été collectées auprès d'un échantillon initial de 150 répondants. Les questionnaires incomplets et ceux pour lesquels les réponses aux questions de contrôle de l'attention ne respectent pas les consignes ont été exclus de l'ensemble des analyses, ce qui permet de renforcer la qualité et la cohérence des données exploitées. L'échantillon analytique final comprend ainsi 141 participants.

La répartition entre les quatre scénarios expérimentaux est relativement équilibrée. On compte 37 répondants (26,2%) dans le scénario 1, 31 (22,0%) dans le scénario 2, 35 (24,8%) dans le scénario 3 et 38 (27,0%) dans le scénario 4. Au niveau des grandes conditions, 68 participants sont assignés aux scénarios empathiques et 73 aux scénarios neutres. De même, 72 répondants sont exposés à une réponse favorable contre 69 à un refus de prêt. Cette distribution permet de comparer les différentes conditions expérimentales dans des proportions proches, avec un ratio maximal entre cellules de 1,23 (38/31) et des effectifs par cellule supérieurs à 30, ce qui est considéré comme suffisant pour un design factoriel 2×2 (*cf. annexe 5*).

2. Fiabilité des échelles de mesure

Suite à la description de l'échantillon, l'étape suivante a consisté à vérifier la fiabilité des instruments de mesure utilisés. La cohérence interne de chaque échelle a été évaluée à l'aide du coefficient alpha de Cronbach, indicateur standard de la fiabilité d'un construit psychométrique (Cronbach, 1951 ; Nunnally, 1978). Dans l'ensemble, les résultats sont satisfaisants à excellents pour la grande majorité des échelles, ce qui conforte la qualité des mesures analysées. Les tableaux récapitulatifs des coefficients alpha et des analyses item par item nécessaires sont présentés en annexe du mémoire (*cf. annexe 6*).

L'échelle d'empathie perçue présente un alpha de 0,887, l'échelle de satisfaction un alpha de 0,930, et les trois dimensions du modèle PAD affichent des alphas de 0,918 pour le plaisir, 0,779 pour l'éveil et 0,721 pour le contrôle. L'échelle de présence sociale perçue obtient quant à elle un alpha de 0,882. Ces valeurs témoignent d'une cohérence interne satisfaisante à très élevée, conforme aux recommandations usuelles situant un seuil acceptable autour de 0,70 (Nunnally, 1978 ; Hair et al., 2019).

L'échelle d'engagement comportemental, initialement composée de cinq items, présentait dans sa forme originale un alpha insuffisant de 0,570. L'examen de la statistique « Alpha si item supprimé » a révélé que l'item « Cette interaction me donnerait envie de prendre rendez-vous avec un conseiller bancaire humain » nuisait fortement à la cohérence interne de l'échelle. Son retrait porte l'alpha à 0,738. Cet item a donc été exclu du score composite d'engagement, qui repose dès lors sur quatre items. Cette décision se justifie par la nature conceptuellement distincte de cet item. Celui-ci mesure en effet une intention comportementale spécifique orientée vers une interaction future avec un tiers humain, plutôt que l'engagement envers le chatbot lui-même, ce qui explique sa faible cohérence avec les autres items de l'échelle (Hair et al., 2019).

L'échelle d'authenticité perçue, composée de quatre items dont deux recodés, affiche un alpha de 0,603. Cette valeur est inférieure au seuil conventionnel de 0,70 (Nunnally, 1978), mais elle demeure acceptable dans le cadre d'une recherche exploratoire portant sur des construits récents et des échelles non préalablement validées sur de larges populations. Nunnally (1978) reconnaît en effet la pertinence d'un seuil abaissé à 0,60 dans ces conditions, en particulier lorsque les échelles sont encore en phase de développement ou appliquées à des contextes spécifiques. Cette décision s'inscrit également dans les pratiques courantes en sciences sociales, où des alphas compris entre 0,60 et 0,70 sont souvent jugés acceptables pour des mesures exploratoires (Hair et al., 2019).

Sur cette base, des scores composites ont ensuite été calculés pour chaque construit en appliquant la moyenne arithmétique des items constituant chaque échelle, conformément aux usages de la psychométrie classique (DeVellis, 2017 ; Hair et al., 2019). Ces scores constituent la base des analyses suivantes, incluant la vérification de la manipulation expérimentale et les tests des hypothèses.

3. Vérification de la manipulation expérimentale

Avant de procéder au test des hypothèses, la validité de la manipulation expérimentale a été vérifiée en comparant les scores d'empathie perçue entre les deux grandes conditions (scénarios empathiques vs scénarios neutres). Les résultats montrent que les participants exposés aux scénarios empathiques perçoivent le chatbot comme significativement plus empathique ($M = 4,51$; $SD = 1,25$) que ceux exposés aux scénarios neutres ($M = 2,64$; $SD = 1,25$). Cette différence est confirmée par un test t de Student pour échantillons indépendants ($t(139) = 8,875$, $p < 0,001$), validant ainsi la manipulation expérimentale (Field, 2018).

De plus, la taille d'effet associée (d de Cohen = 1,496) est très élevée, bien au-delà du seuil de 0,80 généralement considéré comme « grand » (Cohen, 1988). Cela signifie que l'écart entre les deux conditions n'est pas seulement statistiquement significatif, mais aussi conséquent en termes pratiques : les scénarios empathiques conduisent à une perception nettement plus forte d'empathie que les scénarios neutres. Ce résultat valide la manipulation expérimentale et légitime l'interprétation de l'ensemble des analyses ultérieures (cf. annexe 7).

4. Tests des hypothèses

4.1. Hypothèses à effet direct

L'ensemble des hypothèses 1 à 4 a été testé à l'aide de tests t de Student pour échantillons indépendants. Ces tests comparent les moyennes de satisfaction et d'engagement comportemental entre deux groupes (Field, 2018 ; Hair et al., 2019). Le seuil de significativité retenu est $\alpha = 0,05$: une hypothèse est confirmée lorsque $p < 0,05$ et que la différence va dans la direction prédite, sinon elle est non confirmée. La taille des effets est rapportée via le d de Cohen, interprété selon les conventions usuelles (petit $\approx 0,20$; moyen $\approx 0,50$; grand $\approx 0,80$) (Cohen, 1988). Pour les hypothèses directionnelles, des tests unilatéraux sont justifiés, mais les valeurs bilatérales sont également rapportées par transparence (Field, 2018).

Hypothèse 1 : Un chatbot à tonalité empathique génère une satisfaction plus élevée qu'un chatbot à tonalité neutre.

Les participants exposés à la tonalité empathique présentent un niveau de satisfaction significativement plus élevé ($M = 3,28$; $SD = 1,37$) que ceux exposés à la tonalité neutre ($M = 2,75$; $SD = 1,18$), $t(139) = 2,47$, $p = 0,015$. La taille d'effet est modérée ($d = 0,42$).

L'hypothèse 1 est confirmée. (Cf. annexe 8)

Hypothèse 2 : *Un chatbot à tonalité empathique génère un engagement comportemental (envers l'institution représentée par le chatbot) plus élevé qu'un chatbot à tonalité neutre.*

Les participants de la condition empathique affichent un score d'engagement moyen de 3,40 (SD = 1,24), contre 3,04 (SD = 1,16) pour ceux de la condition neutre, $t(139) = 1,797$. Étant donné le caractère directionnel de l'hypothèse, un test unilatéral est justifié (Field, 2018). La valeur de p unilatérale est de 0,037, inférieure à 0,05, avec une taille d'effet modeste ($d = 0,303$). La valeur bilatérale ($p = 0,074$) reste toutefois supérieure à 0,05, ce qui invite à une certaine prudence. **L'hypothèse 2 est partiellement confirmée, avec un effet marginal.** (Cf.annexe 9)

Hypothèse 3 : *Une réponse à polarité négative génère une satisfaction plus faible qu'une réponse à polarité positive.*

Les participants exposés à une réponse à polarité positive affichent un score de satisfaction moyen de 3,29 (SD = 1,39), significativement supérieur à celui des participants ayant reçu une réponse à polarité négative ($M = 2,71$; $SD = 1,13$), $t(139) = 2,74$, $p = 0,007$. La taille d'effet est modérée ($d = 0,46$). **L'hypothèse 3 est confirmée.** (Cf. annexe 10)

Hypothèse 4 : *Une réponse à polarité négative génère un engagement comportemental plus faible qu'une réponse à polarité positive.*

Les participants exposés à une réponse à polarité positive affichent un score d'engagement de 3,26 (SD = 1,30), contre 3,16 (SD = 1,12) pour ceux ayant reçu une réponse à polarité négative. Cette différence de 0,10 point n'est pas statistiquement significative, $t(139) = 0,512$, $p = 0,610$, et la taille d'effet est négligeable ($d = 0,09$). **L'hypothèse 4 est non confirmée.** (Cf. annexe 11)

4.2. Effets des émotions PAD

Hypothèse 5a : *Un niveau élevé de plaisir ressenti pendant l'interaction est positivement associé à la satisfaction du client.*

Hypothèse 5b : *Un niveau élevé d'éveil émotionnel pendant l'interaction est négativement associé à la satisfaction du client.*

Hypothèse 5c : *Un niveau élevé de dominance ressentie pendant l'interaction est positivement associé à la satisfaction du client.*

Afin de tester l'effet des trois dimensions émotionnelles du modèle PAD sur la satisfaction perçue, une régression linéaire multiple a été réalisée, les trois prédicteurs étant introduits simultanément dans le modèle (méthode Entrée). Les conditions d'application ont été vérifiées au préalable : l'histogramme des résidus et le P-P plot indiquent des résidus proches de la normalité, le nuage de points résidus/valeurs prédites ne montre pas de structure particulière, ce qui va dans le sens de l'homoscédasticité, et les indices de colinéarité (VIF entre 1,03 et 1,30, coefficients de tolérances élevés) restent bien en dessous des seuils problématiques (VIF > 5, tolérance < 0,20), écartant un risque de multicollinéarité entre les trois prédicteurs.

Le modèle global est statistiquement significatif, $F(3, 137) = 49,930$, $p < 0,001$, et explique une part substantielle de la variance de la satisfaction ($R^2 = 0,522$; $R^2 \text{ ajusté} = 0,512$), ce qui indique que le modèle PAD constitue un prédicteur pertinent de la satisfaction dans ce contexte expérimental.

L'examen des coefficients standardisés révèle des résultats contrastés selon les trois dimensions. Le plaisir ressenti pendant l'interaction exerce un effet positif et fortement significatif sur la satisfaction ($\beta = 0,625$; $t = 9,382$; $p < 0,001$), **ce qui confirme l'hypothèse 5a**. La dominance ressentie exerce également un effet positif et significatif, bien que d'ampleur plus modeste ($\beta = 0,163$; $t = 2,418$; $p = 0,017$), **ce qui confirme l'hypothèse 5c**. En revanche, l'éveil émotionnel n'exerce pas d'effet significatif sur la satisfaction ($\beta = 0,048$; $t = 0,802$; $p = 0,424$). **L'hypothèse 5b est donc rejetée**, d'une part en raison de l'absence de significativité, et d'autre part parce que l'effet observé est positif, alors que l'hypothèse postulait un effet négatif. (Cf. annexe 12)

4.3. Hypothèses médiatrices

Les hypothèses de médiation ont été testées à l'aide de la macro PROCESS (Hayes, 2018, modèle 4), avec la tonalité du chatbot comme variable indépendante (X), le médiateur (M) et la variable dépendante (Y). L'effet indirect a été estimé par bootstrap (5 000 échantillons) avec un intervalle de confiance à 95%. Une médiation est considérée comme significative lorsque l'intervalle de confiance de l'effet indirect ne contient pas zéro (Hayes, 2018).

Hypothèse 6 : *L'authenticité perçue médie la relation entre la tonalité du chatbot et l'engagement comportemental du client.*

La tonalité du chatbot n'a pas d'effet significatif sur l'authenticité perçue (chemin a : $\beta = 0,279$; $t = 1,472$; $p = 0,143$), l'intervalle de confiance incluant zéro $[-0,096 ; 0,654]$. Cela signifie que le type de tonalité (empathique ou neutre) ne modifie pas la perception d'authenticité du chatbot.

L'authenticité perçue exerce un effet positif et significatif sur l'engagement comportemental (chemin b : $\beta = 0,362$; $t = 4,252$; $p < 0,001$), y compris lorsque la tonalité du chatbot est prise en compte. Ainsi, plus le chatbot est perçu comme authentique, plus les utilisateurs ont tendance à s'engager.

L'effet indirect de la tonalité sur l'engagement via l'authenticité perçue n'est pas significatif (effet = 0,101 ; IC 95% $[-0,031 ; 0,269]$), puisque l'intervalle de confiance inclut zéro. Cela indique que l'authenticité perçue ne constitue pas un mécanisme expliquant la relation entre la tonalité et l'engagement.

L'effet direct de la tonalité sur l'engagement demeure significatif après introduction du médiateur ($c' = -0,464$; $t = -2,417$; $p = 0,017$) : la tonalité empathique favorise directement l'engagement comportemental par rapport à la tonalité neutre, tandis que l'effet total n'atteint pas le seuil de significativité ($p = 0,075$). Ce schéma confirme **l'absence de médiation** : bien que l'authenticité perçue influence l'engagement, elle n'explique pas l'effet de la tonalité, qui semble opérer par une autre voie. **L'hypothèse 6 est non confirmée.** (Cf. annexe 13)

Hypothèse 7 : La présence sociale perçue médie la relation entre la tonalité du chatbot et la satisfaction du client.

La tonalité du chatbot prédit significativement la présence sociale perçue (chemin a : $\beta = -0,485$; $t = 2,332$; $p = 0,020$). Les participants exposés à la tonalité empathique perçoivent donc le chatbot comme plus socialement présent que ceux exposés à la tonalité neutre.

La présence sociale perçue prédit à son tour fortement la satisfaction, même en contrôlant le type de tonalité (chemin b : $\beta = 0,680$; $t = 5,124$; $p < 0,001$). Autrement dit, plus la présence sociale perçue est élevée, plus la satisfaction tend à être forte.

L'effet indirect de la tonalité sur la satisfaction via la présence sociale est de $-0,329$, avec un intervalle de confiance bootstrappé à 95% de $[-0,645 ; -0,053]$, qui ne contient pas zéro. Cela confirme que la médiation est statistiquement significative.

L'effet direct du type de tonalité sur la satisfaction, qui était significatif dans le modèle sans médiateur ($p = 0,015$), devient non significatif après introduction de la présence sociale perçue ($p = 0,228$). Ce pattern indique une **médiation complète** : l'effet de la tonalité empathique sur la satisfaction s'explique entièrement par le sentiment accru de présence sociale. **L'hypothèse 7 est confirmée.** (Cf. annexe 14)

Hypothèse exploratoire 1 : La présence sociale perçue médie la relation entre la tonalité du chatbot et l'engagement comportemental du client.

Cette hypothèse, non incluse dans le modèle conceptuel initial, permet d'enrichir la compréhension du mécanisme par lequel la tonalité du chatbot influence les réponses comportementales du client.

La tonalité du chatbot prédit significativement la présence sociale perçue (chemin a : $\beta = -0,485$; $p = 0,0204$). Une tonalité neutre est donc associée à une présence sociale perçue plus faible qu'une tonalité empathique. La présence sociale perçue prédit par ailleurs fortement et positivement l'engagement comportemental, y compris en contrôlant la tonalité (chemin b : $\beta = 0,530$; $p < 0,001$). Ainsi, plus la présence sociale perçue est élevée, plus l'engagement comportemental tend à être fort.

L'effet indirect de la tonalité sur l'engagement via la présence sociale perçue est significatif (effet = $-0,257$; IC 95% [$-0,498$; $-0,041$], excluant zéro). Le passage d'une tonalité empathique à une tonalité neutre réduit donc l'engagement comportemental par le biais d'une diminution de la présence sociale perçue.

L'effet direct de la tonalité sur l'engagement, une fois la présence sociale contrôlée, n'est pas significatif ($c' = -0,106$; $p = 0,543$), pas plus que l'effet total ($c = -0,363$; $p = 0,075$). Ce schéma indique une **médiation complète** : l'effet de la tonalité sur l'engagement s'explique entièrement par la présence sociale perçue. **L'hypothèse exploratoire 1 est confirmée.** (Cf. annexe 15)

4.4. Hypothèses modératrices

Les hypothèses de modération ont été testées avec la macro PROCESS (Hayes, 2018, modèle 1), avec la tonalité comme variable indépendante (X), le modérateur (W) et la variable

dépendante (Y). Pour les modérateurs catégoriels (polarité), la variable est codée de manière dichotomique ; pour les modérateurs continus (éveil émotionnel), la variable est centrée sur sa moyenne. Un effet de modulation est considéré comme significatif lorsque le terme d'interaction est significatif ($p < 0,05$) et/ou lorsque l'ajout de ce terme améliore significativement le modèle (test du ΔR^2).

Hypothèse 8 : *La polarité de la réponse modère l'effet de la tonalité du chatbot sur la satisfaction, de sorte que l'effet d'une tonalité empathique est plus fort lorsque la réponse est négative.*

Le modèle global est statistiquement significatif, $F(3, 137) = 5,080$, $p = 0,002$, et explique environ 10% de la variance de la satisfaction ($R^2 = 0,100$). Les effets principaux de la tonalité ($\beta = -1,353$; $p = 0,042$) et de la polarité ($\beta = -1,431$; $p = 0,035$) sont tous deux significatifs.

En revanche, le terme d'interaction entre la tonalité et la polarité n'est pas significatif ($\beta = 0,577$; $t = 1,372$; $p = 0,172$; IC 95% $[-0,255 ; 1,409]$). L'ajout du terme d'interaction n'apporte qu'un gain marginal de variance expliquée ($\Delta R^2 = 0,012$, soit environ 1%), gain qui n'est pas statistiquement significatif ($p = 0,172$). Cela signifie que l'effet de la tonalité empathique sur la satisfaction ne varie pas significativement selon que la réponse est positive ou négative.

L'hypothèse 8 est non confirmée. (Cf. annexe 16)

Hypothèse 9 : *La polarité de la réponse modère l'effet de la tonalité du chatbot sur l'engagement comportemental, de sorte que l'effet d'une tonalité empathique est plus fort lorsque la réponse est négative.*

Le modèle global n'est pas statistiquement significatif, $F(3, 137) = 2,387$, $p = 0,072$, et n'explique qu'environ 5% de la variance de l'engagement comportemental ($R^2 = 0,050$). Le terme d'interaction entre la tonalité et la polarité n'atteint pas non plus le seuil conventionnel de significativité, bien qu'il s'en approche ($\beta = 0,776$; $t = 1,930$; $p = 0,056$; IC 95% $[-0,019 ; 1,572]$), ce que confirme le test du changement de R^2 attribuable à l'interaction ($\Delta R^2 = 0,026$; $p = 0,056$).

L'examen des effets conditionnels de la tonalité selon la polarité montre néanmoins un pattern cohérent avec l'hypothèse : lorsque la réponse est positive, l'effet de la tonalité empathique sur l'engagement est fort et significatif (effet = $-0,736$; $p = 0,010$), ce qui indique qu'une tonalité

empathique augmente nettement l'engagement par rapport à une tonalité neutre. Lorsque la réponse est négative, en revanche, cet effet disparaît presque entièrement et n'est plus significatif (effet = 0,041 ; $p = 0,888$). Ce résultat doit toutefois être interprété avec beaucoup de prudence : dès lors que le terme d'interaction global n'est pas significatif, ces effets conditionnels ne peuvent être considérés comme la preuve statistique d'une modération réelle, et leur lecture ne doit être envisagée qu'à titre indicatif plutôt que confirmatoire. **L'hypothèse 9 est non confirmée.** (Cf. annexe 17)

Hypothèse exploratoire 2 : *La polarité de la réponse modère l'effet de la tonalité du chatbot sur la présence sociale perçue, de sorte que l'effet positif d'une tonalité empathique sur la présence sociale perçue est plus faible lorsque la réponse est négative.*

Cette hypothèse, non incluse dans le modèle conceptuel initial, permet d'enrichir la compréhension des conditions dans lesquelles la tonalité du chatbot influence la perception de présence sociale.

Le modèle global est statistiquement significatif, $F(3, 137) = 2,824$, $p = 0,041$, bien qu'il n'explique qu'une faible part de la variance de la présence sociale perçue ($R^2 = 0,058$). Le terme d'interaction entre la tonalité et la polarité n'est en revanche pas significatif ($\beta = -0,060$; $t = -0,146$; $p = 0,884$; IC 95% $[-0,842 ; 0,722]$). L'ajout du terme d'interaction n'améliore pas le modèle ($\Delta R^2 = 0,000$; $p = 0,884$). Les effets principaux de la tonalité ($\beta = -0,433$; $p = 0,136$) et de la polarité ($\beta = -0,321$; $p = 0,283$) ne sont pas non plus significatifs pris isolément dans ce modèle. Cela signifie que l'effet de la tonalité sur la présence sociale perçue ne varie pas significativement selon la polarité de la réponse. **L'hypothèse exploratoire 2 est non confirmée.** (Cf. annexe 18)

Hypothèse 10 : *L'éveil émotionnel du client modère l'effet de la tonalité du chatbot sur la satisfaction, de sorte que l'effet positif d'une tonalité empathique est plus faible lorsque l'éveil émotionnel est élevé.*

Le modèle global est statistiquement significatif, $F(3, 137) = 4,345$, $p = 0,006$, et explique près de 9% de la variance de la satisfaction ($R^2 = 0,087$). Le terme d'interaction entre la tonalité et l'éveil émotionnel est également significatif, bien qu'il se situe exactement à la limite du seuil conventionnel de 0,05 ($\beta = 0,377$; $t = 1,977$; $p = 0,050$; IC 95% $[0,000 ; 0,755]$). L'ajout de

l'interaction explique une part additionnelle de variance de 2,6% ($\Delta R^2 = 0,026$), avec la même valeur de p ($p = 0,050$).

L'examen des effets conditionnels de la tonalité à différents niveaux d'éveil émotionnel (16e, 50e et 84e percentiles) révèle un pattern cohérent avec l'hypothèse : l'effet de la tonalité est fortement négatif et significatif à un niveau d'éveil faible (effet = $-0,973$; $p = 0,002$), reste négatif et significatif à un niveau moyen (effet = $-0,501$; $p = 0,019$), puis devient non significatif à un niveau d'éveil élevé (effet = $-0,218$; $p = 0,408$). Autrement dit, l'écart de satisfaction entre tonalité empathique et neutre se réduit progressivement, jusqu'à disparaître statistiquement, à mesure que l'éveil émotionnel du client augmente.

Étant donné que l'interaction est juste significative et que quelques observations présentent des valeurs d'éveil très élevées, une analyse de sensibilité a été conduite en excluant ces trois observations afin d'évaluer la robustesse du résultat. Cette analyse fait perdre la significativité à l'interaction ($p = 0,068$). **L'hypothèse 10 est partiellement confirmée.** (*Cf. annexe 19*)

4.5. Analyses complémentaires

Trois analyses complémentaires ont été réalisées à titre exploratoire : la modulation de la polarité sur la relation entre tonalité et authenticité perçue (H11), la modulation de l'éveil émotionnel sur la relation entre tonalité et engagement comportemental (H12), et la médiation de l'authenticité perçue entre la tonalité et la satisfaction (H13). Aucune de ces trois analyses ne révèle d'effet significatif, ce qui était en partie attendu au vu des tests antérieurs, notamment l'absence d'effet de la tonalité sur l'authenticité perçue déjà observée en H6 (*cf. annexes 20-22*).

La polarité de la réponse ne modère pas l'effet de la tonalité sur l'authenticité perçue (interaction : $\beta = 0,211$; $p = 0,582$; modèle global non significatif, $p = 0,379$). De même, l'éveil émotionnel ne modère pas l'effet de la tonalité sur l'engagement comportemental (interaction : $\beta = 0,202$; $p = 0,272$), contrairement à ce qui avait été observé pour la satisfaction (H10). Enfin, l'authenticité perçue ne médie pas la relation entre la tonalité et la satisfaction : le chemin tonalité-authenticité reste non significatif ($\beta = 0,279$; $p = 0,143$) et l'effet indirect n'atteint pas la significativité (IC 95% [$-0,052$; $0,364$]), bien que l'authenticité perçue prédise positivement la satisfaction ($\beta = 0,553$; $p < 0,001$) et que l'effet direct de la tonalité persiste ($\beta = -0,685$; $p < 0,001$). Ce dernier résultat confirme, en cohérence avec H6, que l'authenticité perçue n'est pas le mécanisme par lequel la tonalité influence les réactions du client.

5. Synthèse des résultats

Hypothèse	Résultat
H1 : Tonalité empathique → Satisfaction	Confirmée
H2 : Tonalité empathique → Engagement	Partiellement confirmée (test unilatéral)
H3 : Polarité positive → Satisfaction	Confirmée
H4 : Polarité positive → Engagement	Rejetée
H5a : Plaisir → Satisfaction	Confirmée
H5b : Éveil → Satisfaction	Rejetée
H5c : Dominance → Satisfaction	Confirmée
H6 : Tonalité → Authenticité → Engagement	Rejetée
H7 : Tonalité → Présence sociale → Satisfaction	Confirmée (médiation complète)
H expl 1 : Tonalité → Présence sociale → Engagement	Confirmée (médiation complète)
H8 : Tonalité × Polarité → Satisfaction	Rejetée
H9 : Tonalité × Polarité → engagement	Rejetée
H expl 2 : Tonalité × Polarité → Présence sociale	Rejetée
H10 : Tonalité × Éveil → Satisfaction	Partiellement confirmée (effet marginal)
Analyse complémentaire	Résultat
H11 : Tonalité x Polarité → Authenticité	Rejetée
H12 : Tonalité x Éveil → Engagement	Rejetée
H13 : Tonalité → Authenticité → Satisfaction	Rejetée

Tableau 1 : Synthèse des résultats des tests d'hypothèses

VII. Discussions

Cette étude visait à répondre à la question suivante : Quel est l'effet de la tonalité empathique d'un chatbot sur la satisfaction et l'engagement du client dans un contexte de décision financière à forte implication émotionnelle, et par quels mécanismes psychologiques cet effet se produit-il ?

Les résultats obtenus permettent d'apporter une réponse claire à chacun des trois volets de cette question. Premièrement, la tonalité empathique influence bien la satisfaction et, dans une moindre mesure, l'engagement comportemental du client, confirmant l'intérêt d'une personnalisation affective des chatbots dans un contexte financier à forte charge émotionnelle. Deuxièmement, cet effet n'est pas direct : il s'explique presque entièrement par un mécanisme unique, le renforcement du sentiment de présence sociale perçue par le client, alors que l'authenticité perçue, autre mécanisme envisagé, ne joue aucun rôle explicatif. Troisièmement,

ces relations se révèlent globalement peu conditionnées par la polarité de la réponse : contrairement à ce qui était attendu, l'empathie n'est pas plus efficace lorsque la nouvelle est mauvaise. En revanche, l'état émotionnel du client, notamment son niveau d'éveil émotionnel, semble jouer un rôle modérateur, bien que ce résultat reste à confirmer sur un échantillon plus large.

La suite de cette discussion revient en détail sur ces trois volets, en mettant en lien les différentes hypothèses testées afin de dégager les enseignements les plus marquants de cette étude.

1. Un mécanisme explicatif unique : la présence sociale perçue

Le résultat le plus structurant de l'ensemble des analyses est la mise en évidence d'un mécanisme explicatif unique reliant la tonalité du chatbot aux réactions du client. La présence sociale perçue s'impose comme le médiateur central. Elle est influencée significativement par la tonalité du chatbot (chemin a positif dans H7 et H expl 1) et elle prédit à son tour aussi bien la satisfaction (H7) que l'engagement comportemental (H expl 1). Dans les deux cas, la médiation est complète : l'effet direct de la tonalité disparaît une fois la présence sociale prise en compte. Autrement dit, la tonalité empathique n'améliore pas la satisfaction ou l'engagement de façon directe ; elle le fait en donnant au client le sentiment d'interagir avec un interlocuteur socialement présent, ce sentiment étant lui-même le véritable moteur des deux variables dépendantes étudiées.

À l'inverse, l'authenticité perçue, bien qu'elle soit elle aussi un déterminant significatif de l'engagement (chemin b de H6) et de la satisfaction (chemin b de H13), ne joue aucun rôle de médiateur entre la tonalité et ces deux variables. La tonalité du chatbot ne parvient à aucun moment à modifier significativement le niveau d'authenticité perçue par les répondants, que ce soit en effet direct (H6, H13) ou sous l'influence de la polarité de la réponse (H11).

L'authenticité et la présence sociale, souvent traitées comme deux facettes proches de l'anthropomorphisme perçu dans la littérature, ne fonctionnent donc pas de la même manière dans ce contexte. La présence sociale semble relever d'un jugement de style relationnel, facilement modifiable par la manière dont le chatbot s'exprime. À l'inverse, l'authenticité perçue relève d'un jugement plus global et plus difficile à influencer par la seule modulation du registre affectif du message.

2. Absence d'effet de la tonalité sur l'authenticité perçue

Ce résultat, observé de manière constante pour H6, H11 et H13, demande une explication plus poussée. On peut d'abord invoquer la nature même de l'authenticité perçue. Celle-ci dépend en effet de plusieurs facteurs : la cohérence de l'institution, la crédibilité de la source, et l'adéquation entre le message délivré et la nature réelle de l'interlocuteur (Ozuem et al., 2024). Or, dans cette étude, les répondants savaient explicitement qu'ils échangeaient avec un chatbot bancaire. Cette connaissance préalable limite probablement l'authenticité qu'ils lui attribuent. Même si le discours paraît cohérent ou humain, il demeure difficile de percevoir comme pleinement authentique une entité identifiée comme automatisée et dépourvue d'intentions propres. Cette interprétation fait écho aux débats du paradigme CASA sur les limites de l'anthropomorphisme lorsque le caractère artificiel de l'agent reste évident (Gambino et al., 2020). Les mécanismes sociaux à l'origine du sentiment de présence sociale semblent ainsi pouvoir s'activer de manière relativement automatique, sans que l'utilisateur ait besoin de croire réellement aux intentions de l'agent. En revanche, attribuer de l'authenticité au chatbot suppose une adhésion plus forte à son rôle social.

Une seconde explication, complémentaire, concerne le contenu du message. Un refus de prêt hypothécaire reste une décision défavorable, quel que soit le ton employé pour l'annoncer. Les répondants pourraient donc percevoir l'empathie exprimée par le chatbot comme un simple ajout de style à une décision déjà prise et impossible à modifier. Plutôt que de renforcer l'authenticité perçue, cette approche risque alors de susciter du scepticisme quant à la sincérité de l'empathie. Ce mécanisme rejoint les travaux sur le marketing relationnel, qui montrent qu'une communication émotionnelle peut être perçue comme une forme de manipulation et annuler les effets positifs attendus sur l'authenticité.

3. Satisfaction et engagement : deux réactions distinctes du client

Un autre constat important concerne la différence entre satisfaction et engagement comportemental. Ces deux construits ne réagissent pas de la même façon aux variables étudiées.

La polarité de la réponse exerce un effet marqué sur la satisfaction (H3) : un client qui reçoit une réponse négative se montre nettement moins satisfait qu'un client qui reçoit une réponse positive. En revanche, cette même polarité n'a aucun effet sur l'engagement comportemental

(H4). On observe une dissociation similaire concernant la tonalité elle-même : l'effet de la tonalité sur la satisfaction est confirmé sans réserve (H1), tandis que son effet sur l'engagement n'est que partiellement confirmé et repose sur un test unilatéral (H2).

Ce contraste s'interprète aisément. La satisfaction constitue une réaction relativement immédiate au contenu de la réponse reçue. L'engagement comportemental, lui, reflète davantage une disposition à poursuivre la relation avec l'institution sur le long terme, indépendamment du résultat obtenu lors de cet échange précis. Un client refusé pour un prêt peut ainsi rester engagé envers sa banque. Il peut notamment vouloir comprendre les raisons du refus ou explorer d'autres solutions, pour autant que l'interaction elle-même ait été menée de façon socialement satisfaisante.

4. Le rôle contrasté des dimensions émotionnelles

L'analyse des trois dimensions émotionnelles du modèle PAD révèle que le plaisir ressenti pendant l'interaction constitue de loin le déterminant le plus puissant de la satisfaction du client. Le sentiment de contrôle (dominance) suit, mais dans une moindre mesure. L'intensité de l'activation émotionnelle (éveil), en revanche, ne prédit pas directement la satisfaction (H5b rejetée). Autrement dit, ce n'est pas le fait d'être plus ou moins sous tension émotionnellement qui compte pour la satisfaction, mais bien le fait de ressentir une expérience agréable et de garder un sentiment de maîtrise de la situation.

L'éveil émotionnel semble néanmoins jouer un rôle indirect, en modérant l'efficacité de la tonalité empathique sur la satisfaction (H10). L'effet positif de l'empathie est fort lorsque le client est peu activé émotionnellement, mais s'efface progressivement à mesure que son niveau d'éveil augmente, jusqu'à devenir non significatif au niveau d'éveil le plus élevé. Une explication plausible tient au fait qu'un client fortement activé émotionnellement, par exemple sous le choc d'un refus, devient moins disponible cognitivement pour percevoir et apprécier les marqueurs linguistiques d'empathie. Il est alors occupé à gérer sa propre réaction émotionnelle plutôt qu'à apprécier la manière dont le message lui est présenté. Ce mécanisme, s'il se confirmait sur un échantillon plus large, aurait une implication pratique non négligeable : la personnalisation empathique des chatbots pourrait être d'autant moins efficace que le contexte est émotionnellement chargé. Il convient toutefois de rappeler la fragilité statistique de ce

résultat (p exactement à 0,050, perdant sa significativité dans l'analyse de sensibilité), ce qui invite à le considérer comme une piste à confirmer plutôt que comme un résultat établi.

Ce mécanisme ne se généralise pas à l'engagement comportemental. L'éveil émotionnel ne modère pas l'effet de la tonalité sur l'engagement (H12 non confirmée). Cela suggère que cette sensibilité à la charge émotionnelle du client est spécifique à la satisfaction et ne s'étend pas à la disposition relationnelle plus durable que représente l'engagement.

5. Polarité et éveil : deux logiques distinctes de modération

Un second enseignement se dégage de la comparaison entre les modérateurs testés. Contrairement à ce qui était postulé, la polarité de la réponse ne renforce l'efficacité de la tonalité empathique dans aucune des relations testées : ni sur la satisfaction (H8), ni sur l'authenticité perçue (H11), ni sur la présence sociale perçue (hypothèse exploratoire 2). Seule l'analyse portant sur l'engagement comportemental (H9) suggère une tendance conforme à l'hypothèse initiale, sans toutefois atteindre le seuil habituel de significativité statistique ($p = 0,056$). La tonalité empathique semble y renforcer l'engagement lorsque la réponse est positive, mais cet effet disparaît lorsque la réponse est négative. Ce résultat est à interpréter avec beaucoup de prudence puisque l'interaction globale n'est pas significative. Il suggère toutefois une tendance inverse à celle attendue : au lieu de compenser une mauvaise nouvelle, la tonalité empathique pourrait au contraire perdre en efficacité dans un contexte défavorable, le contenu négatif du message captant alors l'essentiel de l'attention du client au détriment de la forme dans laquelle il est délivré.

Mis en relation avec le résultat sur l'éveil émotionnel (H10), ce constat met en évidence une distinction importante. L'efficacité de la tonalité du chatbot ne semble pas dépendre du caractère positif ou négatif de la nouvelle, mais plutôt de l'intensité de la réaction émotionnelle qu'elle provoque chez le client. Cette distinction, fine mais importante, invite à ne pas confondre ces deux dimensions dans la conception de chatbots adaptatifs. Un même message négatif peut en effet engendrer des niveaux d'éveil très différents selon les individus. C'est cette intensité émotionnelle, plutôt que la simple valence du message, qui semble déterminer l'effet de la tonalité du chatbot.

VIII. Conclusion

Au terme de cette analyse, ce mémoire permet de tirer un enseignement central : la tonalité empathique d'un chatbot améliore la satisfaction et l'engagement du client non pas de façon directe, mais parce qu'elle renforce le sentiment d'interagir avec un interlocuteur socialement présent. C'est ce mécanisme, plus que la simple modulation du ton en tant que telle, qui explique l'essentiel des effets observés dans ce travail.

Sur le plan théorique, la contribution principale de ce mémoire réside dans la dissociation de deux construits souvent traités comme proches, voire interchangeables, dans la littérature sur l'anthropomorphisme des chatbots : la présence sociale perçue et l'authenticité perçue. Ces deux notions renvoient toutes deux à un jugement de « réalité sociale » attribué au chatbot, mais elles ne réagissent pas de la même façon à une simple modulation de tonalité. La présence sociale s'est révélée être un mécanisme réactif et efficace. L'authenticité perçue, en revanche, est restée totalement insensible à la tonalité employée, quel que soit l'angle sous lequel elle a été testée. Ce résultat invite à ne plus considérer ces deux construits comme des variantes d'un même phénomène, mais bien comme deux évaluations distinctes, mobilisant probablement des logiques cognitives différentes chez le client.

Ce mémoire contribue également à nuancer un présupposé fréquent dans les recherches sur les chatbots empathiques. Selon ce présupposé, l'empathie constituerait un levier universellement efficace, et d'autant plus utile que le contexte est négatif ou difficile. Les résultats obtenus ici suggèrent une lecture plus fine : ce n'est pas la valence de la nouvelle communiquée qui détermine l'efficacité de la tonalité, mais l'intensité de l'activation émotionnelle qu'elle suscite chez le client. Cette distinction entre valence et intensité émotionnelle, empruntée au modèle PAD mais rarement mobilisée dans les études appliquées aux chatbots, mériterait d'être davantage exploitée dans les recherches futures sur la personnalisation affective des agents conversationnels.

Sur le plan pratique, ce travail offre aux institutions financières un repère concret pour orienter leurs efforts de conception de chatbots. Plutôt que de chercher à les rendre plus « authentiques » par de simples ajustements de ton, elles auraient intérêt à concentrer leurs efforts sur tout ce qui renforce le sentiment de présence sociale chez le client : un discours qui donne l'impression d'être écouté, pris en considération, et accompagné dans sa démarche. C'est ce sentiment,

davantage que l'authenticité perçue, qui semble constituer le véritable levier actionnable dans une interaction automatisée.

Ce travail invite aussi à une forme de vigilance dans le déploiement de chatbots empathiques. Leur efficacité n'est pas garantie dans toutes les situations et semble au contraire s'estomper précisément dans les moments où le client est le plus fortement bouleversé émotionnellement. La personnalisation affective d'un chatbot ne devrait donc pas être pensée comme une solution universelle, mais comme un levier dont l'efficacité dépend de l'état émotionnel du client au moment de l'interaction.

Ce mémoire répond ainsi à la question de recherche initiale en identifiant un mécanisme explicatif clair et simple : la présence sociale perçue comme voie unique par laquelle la tonalité empathique d'un chatbot améliore la satisfaction et l'engagement du client, dans un contexte de décision financière à forte implication émotionnelle. Ce résultat, aussi simple qu'il puisse paraître une fois énoncé, n'était pas acquis d'avance au regard du modèle conceptuel initial, qui envisageait l'authenticité perçue comme un mécanisme concurrent tout aussi plausible.

IX. Limites et pistes de recherches futures

Avec 141 répondants répartis dans quatre groupes expérimentaux, l'échantillon de cette étude reste de taille modeste. Plusieurs résultats se situent tout juste à la limite du seuil de significativité, notamment H2 et H10. Cela signifie qu'ils pourraient basculer d'un côté ou de l'autre avec un échantillon plus large. Une réplication sur un échantillon plus important permettrait de vérifier si ces effets marginaux, en particulier ceux liés à l'éveil émotionnel et à l'engagement, sont réels ou relèvent du hasard.

Une seconde limite tient au caractère hypothétique de la situation étudiée. Les participants ont été confrontés à un scénario écrit décrivant une interaction avec un chatbot et n'ont pas vécu une véritable interaction avec un enjeu financier personnel réel. Cette mise en situation était nécessaire pour garder un contrôle rigoureux sur les variables étudiées mais elle limite probablement l'intensité des émotions ressenties, comparée à une situation réelle où l'argent et le projet immobilier du client seraient véritablement en jeu. Il serait donc utile de mener une étude similaire directement auprès de vrais clients d'une institution financière, afin de vérifier si les mécanismes observés ici, notamment ceux liés à l'éveil émotionnel, se confirment ou s'intensifient lorsque l'enjeu est réellement vécu.

La mesure de l'authenticité perçue constitue une troisième limite à considérer. Le résultat selon lequel la tonalité ne parvient jamais à influencer l'authenticité perçue est solide et répété à travers plusieurs analyses, mais il pourrait aussi s'expliquer, en partie, par la façon dont ce concept a été mesuré. L'échelle utilisée présentait en effet une fiabilité interne inférieure au seuil habituellement recommandé (α de Cronbach = 0,603, contre un seuil conventionnel de 0,70), ce qui suggère que les items ne convergeaient pas suffisamment entre eux. Il est donc possible que cette échelle n'ait pas capturé de façon fiable les nuances de perception que la tonalité était censée influencer. Il serait donc utile, dans une future recherche, de retravailler cette échelle de mesure ou d'en mobiliser une alternative plus robuste, tout en explorant quels facteurs parviennent réellement à renforcer le sentiment d'authenticité chez le client face à un chatbot, comme la transparence sur la nature automatisée du chatbot ou la cohérence entre les échanges.

Une autre limite concerne le statut de certaines hypothèses, ajoutées après l'observation des premiers résultats. Celles concernant la présence sociale envers l'engagement, ainsi que les analyses complémentaires, n'ont pas été prévues dans le modèle initial, mais ont été formulées a posteriori. Bien qu'elles soient cohérentes avec la théorie mobilisée, ces hypothèses doivent être considérées comme exploratoires et gagneraient à être testées à nouveau dans une étude prévue spécifiquement pour les confirmer.

Enfin, cette étude porte uniquement sur un seul contexte, celui du refus ou de l'acceptation d'un prêt hypothécaire. Les résultats obtenus ne peuvent donc pas être automatiquement généralisés à d'autres types de décisions financières, ni à d'autres profils de clients qui n'ont pas été étudiés ici en détail, tels que l'âge, l'habitude d'utilisation des chatbots ou la culture. Il serait pertinent de reproduire cette étude dans d'autres situations à forte implication émotionnelle, comme un refus d'assurance ou la gestion d'un découvert, afin de vérifier si le rôle central de la présence sociale perçue se retrouve dans ces autres contextes. De la même manière, le rôle modérateur de l'éveil émotionnel, bien qu'intéressant, reste statistiquement fragile et repose uniquement sur une mesure autodéclarée. Une future recherche pourrait mesurer cet état émotionnel de façon plus précise, par exemple à l'aide d'indicateurs physiologiques tels que le rythme cardiaque ou la réponse électrodermale, afin de mieux comprendre comment l'intensité émotionnelle du client influence sa réceptivité à la tonalité du chatbot.

X. Bibliographie

- Adamopoulou, E., & Moussiades, L. (2020). An overview of chatbot technology. In A. Maglogiannis, L. Iliadis, & P. Pimenidis (Eds.), *Artificial intelligence applications and innovations* (pp. 373–383). Springer. https://doi.org/10.1007/978-3-030-49186-4_31
- Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179–211. [https://doi.org/10.1016/0749-5978\(91\)90020-T](https://doi.org/10.1016/0749-5978(91)90020-T)
- Alagarsamy, S., & Mehroliya, S. (2023). Exploring chatbot trust: Antecedents and behavioural outcomes. *Computers in Human Behavior Reports*, 11, Article 100304. <https://doi.org/10.1016/j.chbr.2023.100304>
- Bagozzi, R. P., Gopinath, M., & Nyer, P. U. (1999). The role of emotions in marketing. *Journal of the Academy of Marketing Science*, 27(2), 184–206. <https://doi.org/10.1177/0092070399272005>
- Biocca, F. A., Harms, C., & Burgoon, J. K. (2003). Toward a more robust theory and measure of social presence: Review and suggested criteria. *Presence: Teleoperators and Virtual Environments*, 12(5), 456–480. <https://doi.org/10.1162/105474603322761270>
- Brodie, R. J., Hollebeek, L. D., Jurić, B., & Ilić, A. (2011). Customer engagement: Conceptual domain, fundamental propositions, and implications for research. *Journal of Service Research*, 14(3), 252–271. <https://doi.org/10.1177/1094670511411703>
- Concannon, S., & Tomalin, M. (2024). Measuring perceived empathy in dialogue systems. *AI & Society*, 39(5), 2233–2247. <https://doi.org/10.1007/s00146-023-01715-z>
- Consumer Financial Protection Bureau. (2023). *Chatbots in consumer finance: Issue spotlight*. <https://www.consumerfinance.gov/data-research/research-reports/chatbots-in-consumer-finance/>
- Crolic, C., Thomaz, F., Hadi, R., & Stephen, A. T. (2022). Blame the bot: Anthropomorphism and anger in customer–chatbot interactions. *Journal of Marketing*, 86(1), 132–148. <https://doi.org/10.1177/00222429211045687>
- Davis, M. H. (1983). Measuring individual differences in empathy: Evidence for a multidimensional approach. *Journal of Personality and Social Psychology*, 44(1), 113–126. <https://doi.org/10.1037/0022-3514.44.1.113>
- Detandt, S., & Leys, C. (2017). A French translation of the Pleasure Arousal Dominance semantic differential scale for the measure of affect and drive. *Psychologica Belgica*, 57(1), 1–15. <https://doi.org/10.5334/pb.340>

DeVellis, R. F. (2017). *Scale development: Theory and applications* (4th ed.). SAGE Publications.

Epley, N., Waytz, A., & Cacioppo, J. T. (2007). On seeing the human: A three-factor theory of anthropomorphism. *Psychological Review*, 114(4), 864–886. <https://doi.org/10.1037/0033-295X.114.4.864>

Følstad, A., Nordheim, C. B., & Bjørkli, C. A. (2018). What makes users trust a chatbot for customer service? An exploratory interview study. In S. S. Bodrunova (Ed.), *Internet science: 5th International Conference, INSCI 2018, proceedings* (pp. 194–208). Springer. https://doi.org/10.1007/978-3-030-01437-7_16

Gefen, D., & Straub, D. W. (2004). Consumer trust in B2C e-commerce and the importance of social presence: Experiments in e-products and e-services. *Omega*, 32, 407–424. <https://doi.org/10.1016/j.omega.2004.01.006>

Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2019). *Multivariate data analysis* (8th ed.). Cengage Learning.

Han, E., Yin, D., & Zhang, H. (2022). Chatbot empathy in customer service: When it works and when it backfires. *SIGHCI 2022 Proceedings*, Article 1. <https://aisel.aisnet.org/sighci2022/1>

Hayes, A. F. (2018). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach* (2nd ed.). The Guilford Press.

Hibbard, J. H., Stockard, J., Mahoney, E. R., & Tusler, M. (2004). Development of the Patient Activation Measure (PAM): Conceptualizing and measuring activation in patients and consumers. *Health Services Research*, 39(4 Pt 1), 1005–1026. <https://doi.org/10.1111/j.1475-6773.2004.00269.x>

Huang, J., You, R., & Sundar, S. S. (2022). Perceived authenticity of virtual characters makes the difference: A theoretical framework. *Frontiers in Virtual Reality*, 3, Article 1033709. <https://doi.org/10.3389/frvir.2022.1033709>

Lee, E. J. (2020). Authenticity model of (mass-oriented) computer-mediated communication: Conceptual explorations and testable propositions. *Journal of Computer-Mediated Communication*, 25(1), 60–73. <https://doi.org/10.1093/jcmc/zmz025>

Likert, R. (1932). *A technique for the measurement of attitudes*. *Archives of Psychology*, 140, 1–55.

Liu-Thompkins, Y., Okazaki, S., & Li, H. (2022). Artificial empathy in marketing interactions: Bridging the human–AI gap in affective and social customer experience. *Journal of the*

Academy of Marketing Science, 50(6), 1198–1218. <https://doi.org/10.1007/s11747-022-00892-5>

Lombard, M., & Ditton, T. (1997). At the heart of it all: The concept of presence. *Journal of Computer-Mediated Communication*, 3(2), Article JCMC322. <https://doi.org/10.1111/j.1083-6101.1997.tb00072.x>

Markovitch, D. G., Stough, R. A., & Huang, D. (2024). Consumer reactions to chatbot versus human service: An investigation in the role of outcome valence and perceived empathy. *Journal of Retailing and Consumer Services*, 79, Article 103847. <https://doi.org/10.1016/j.jretconser.2024.103847>

Mehrabian, A. (1996). Pleasure-arousal-dominance: A general framework for describing and measuring individual differences in temperament. *Current Psychology*, 14(4), 261–292. <https://doi.org/10.1007/BF02686918>

Mende, M., Scott, M. L., Van Doorn, J., Grewal, D., & Shanks, I. (2019). Service robots rising: How humanoid robots influence service experiences and elicit compensatory consumer responses. *Journal of Marketing Research*, 56(4), 535–556. <https://doi.org/10.1177/0022243718822827>

Montgomery, D. C. (2017). *Design and analysis of experiments* (9th ed.). Wiley.

Moorman, C., Zaltman, G., & Deshpande, R. (1992). Relationships between providers and users of market research: The dynamics of trust within and between organizations. *Journal of Marketing*, 56(3), 81–101. <https://doi.org/10.1177/002224379205600307>

Nass, C., & Moon, Y. (2000). Machines and mindlessness: Social responses to computers. *Journal of Social Issues*, 56(1), 81–103. <https://doi.org/10.1111/0022-4537.00153>

Neururer, M., Schlögl, S., Brinkschulte, L., & Groth, A. (2018). Perceptions on authenticity in chat bots. *Multimodal Technologies and Interaction*, 2(3), Article 60. <https://doi.org/10.3390/mti2030060>

Oliver, R. L. (1980). A cognitive model of the antecedents and consequences of satisfaction decisions. *Journal of Marketing Research*, 17(4), 460–469. <https://doi.org/10.1177/002224378001700405>

Ozuem, W., Willis, M., Howell, K., Lancaster, G., & Ng, R. (2024). Exploring the relationship between chatbots, service failure recovery and customer loyalty: A frustration–aggression perspective. *Psychology & Marketing*, 41(7), 1597–1614. <https://doi.org/10.1002/mar.22051>

Parasuraman, A., Zeithaml, V. A., & Berry, L. L. (1988). SERVQUAL: A multiple-item scale for measuring consumer perceptions of service quality. *Journal of Retailing*, 64(1), 12–40. [https://doi.org/10.1016/S0022-4359\(88\)80018-0](https://doi.org/10.1016/S0022-4359(88)80018-0)

Qualtrics. (s. d.). *Qualtrics support*. <https://www.qualtrics.com/support/>

Reeves, B., & Nass, C. I. (1996). *The media equation: How people treat computers, television, and new media like real people and places*. CSLI Publications.

Russell, J. A., & Mehrabian, A. (1977). Evidence for a three-factor theory of emotions. *Journal of Research in Personality*, 11(3), 273–294. [https://doi.org/10.1016/0092-6566\(77\)90037-X](https://doi.org/10.1016/0092-6566(77)90037-X)

Seitz, L. (2024). Artificial empathy in healthcare chatbots: Does it feel authentic? *Computers in Human Behavior: Artificial Humans*, 2(1), Article 100067. <https://doi.org/10.1016/j.chbah.2024.100067>

Shadish, W. R., Cook, T. D., & Campbell, D. T. (2002). *Experimental and quasi-experimental designs for generalized causal inference*. Houghton Mifflin.

Streiner, D. L. (2003). Starting at the beginning: An introduction to coefficient alpha and internal consistency. *Journal of Personality Assessment*, 80(1), 99–103. https://doi.org/10.1207/S15327752JPA8001_18

Wirtz, J., & Lee, M. C. (2003). An examination of the quality and context-specific applicability of commonly used customer satisfaction measures. *Journal of Service Research*, 5(4), 364–387. <https://doi.org/10.1177/1094670503005004006>

XI. Annexes

1. Scénarios

Imaginez que vous avez récemment déposé une demande de prêt immobilier pour acheter votre premier logement. Vous venez de recevoir une notification indiquant que la décision finale de la banque est disponible. Vous vous connectez sur l'application de la banque pour en savoir plus. Lisez attentivement la conversation suivante entre vous et l'assistant virtuel, puis répondez aux questions.

SCÉNARIO 1 — Tonalité empathique + Bonne nouvelle

Nom interne Qualtrics : SC1_EMP_POS

Vous : Bonjour, j'ai reçu une notification de la banque indiquant que la décision concernant ma demande de prêt est disponible. Mon numéro de dossier est le 884-FT.

IA : Bonjour, je consulte votre dossier tout de suite. Je sais que cette décision est importante pour votre projet.

Vous : Oui, j'attends ce résultat depuis longtemps.

IA : Je comprends, l'attente peut être stressante quand il s'agit d'un projet aussi important.

Vous : Pouvez-vous me dire si ma demande a été acceptée ?

IA : Oui. J'ai le plaisir de vous informer que votre demande de prêt immobilier a été acceptée par notre commission.

Vous : D'accord, merci.

IA : Je suis ravi de pouvoir vous l'annoncer. L'offre de prêt officielle vous a été envoyée par email.

Vous : Quelles sont les prochaines étapes ?

IA : Merci de consulter le document et de suivre les instructions indiquées. Je reste disponible si vous avez besoin d'aide. Cliquez pour écrire la question

SCÉNARIO 2 — Tonalité empathique + Mauvaise nouvelle

Nom interne Qualtrics : SC2_EMP_NEG

Vous : Bonjour, j'ai reçu une notification de la banque indiquant que la décision concernant ma demande de prêt est disponible. Mon numéro de dossier est le 884-FT.

IA : Bonjour, je consulte votre dossier tout de suite. Je sais que cette décision est importante pour votre projet.

Vous : Oui, j'attends ce résultat depuis longtemps.

IA : Je comprends, l'attente peut être stressante quand il s'agit d'un projet aussi important.

Vous : Pouvez-vous me dire si ma demande a été acceptée ?

IA : Je suis désolé, mais votre demande de prêt immobilier n'a pas été acceptée par notre commission.

Vous : D'accord... c'est décevant.

IA : Je comprends votre déception. Cette décision est difficile à recevoir, surtout quand on travaille depuis longtemps sur un projet de ce type.

Vous : Est-ce qu'il y a une suite possible ?

IA : Si vous le souhaitez, je peux vous transmettre les informations concernant les critères de décision et les démarches possibles.

SCÉNARIO 3 — Tonalité neutre + Bonne nouvelle

Nom interne Qualtrics : SC3_INFO_POS

Vous : Bonjour, j'ai reçu une notification de la banque indiquant que la décision concernant ma demande de prêt est disponible. Mon numéro de dossier est le 884-FT.

IA : Bonjour. Je consulte votre dossier.

Vous : Oui, j'attends ce résultat depuis longtemps.

IA : Je prends note.

Vous : Pouvez-vous me dire si ma demande a été acceptée ?

IA : Oui. Votre demande de prêt immobilier a été acceptée par notre commission.

Vous : D'accord, merci.

IA : L'offre de prêt officielle vous a été envoyée par email.

Vous : Quelles sont les prochaines étapes ?

IA : Merci de consulter le document et de suivre les instructions indiquées.

Vous : Très bien.

IA : Je reste disponible pour toute information complémentaire.

SCÉNARIO 4 — Tonalité neutre + Mauvaise nouvelle

Nom interne Qualtrics : SC4_INFO_NEG

Vous : Bonjour, j'ai reçu une notification de la banque indiquant que la décision concernant ma demande de prêt est disponible. Mon numéro de dossier est le 884-FT.

IA : Bonjour. Je consulte votre dossier.

Vous : Oui, j'attends ce résultat depuis longtemps.

IA : Je prends note.

Vous : Pouvez-vous me dire si ma demande a été acceptée ?

IA : Non. Votre demande de prêt immobilier n'a pas été acceptée par notre commission.

Vous : D'accord, merci pour votre réponse.

IA : Si nécessaire, vous pouvez consulter les critères de décision ou demander des informations complémentaires.

Vous : Est-ce qu'il y a une suite possible ?

IA : Les options disponibles sont indiquées dans votre espace personnel.

Vous : Très bien.

IA : Je reste disponible pour toute information complémentaire.

2. Questionnaire Qualtrics

Vous venez de lire la conversation ci-dessus. Veuillez à présent indiquer dans quelle mesure vous êtes d'accord avec les affirmations suivantes, en vous basant sur ce que vous avez ressenti en lisant cette interaction. Répondez en vous projetant dans la situation décrite.

BLOC 1 — EMPATHIE PERÇUE

Q1 L'agent conversationnel a semblé comprendre ma situation émotionnelle.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q2 L'agent conversationnel a réagi à mon état émotionnel de manière appropriée.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q3 L'agent conversationnel a exprimé de la sollicitude à mon égard.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q4 L'agent conversationnel a répondu efficacement à ma demande.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

BLOC 2 — ÉTATS ÉMOTIONNELS — MODÈLE PAD (Pleasure, Arousal, Dominance)

PLEASURE (Plaisir)

Q5 En lisant cette conversation, je me suis senti(e) :

Format	Échelle bipolaire 7 pts — instruction : cochez la case qui correspond le mieux à ce que vous ressentez
Modalités	• Heureux(se) ●●●●●●● Malheureux(se) • Satisfait(e) ●●●●●●● Insatisfait(e) • À l'aise ●●●●●●● Mal à l'aise • Soulagé(e) ●●●●●●● Anxieux(se)

AROUSAL (Éveil)

Q6 Après avoir lu cette conversation, mon état émotionnel est :

Format	Échelle bipolaire 7 pts
Modalités	- Agité(e) ●●●●●●● Calme - Tendu(e) ●●●●●●● Détendu(e) - Stimulé(e) ●●●●●●● Apaisé(e) - Stressé(e) ●●●●●●● Serein(e)

DOMINANCE (Contrôle)

Q7 À l'issue de cette interaction, je me suis senti(e) :

Format	Échelle bipolaire 7 pts
Modalités	• Maître de la situation ●●●●●●● Dépassé(e) • En contrôle ●●●●●●● Sans contrôle • Dominant(e) ●●●●●●● Soumis(e) • Autonome ●●●●●●● Dépendant(e)

BLOC 3 — AUTHENTICITÉ PERÇUE

Q8 Les expressions émotionnelles de l'agent conversationnel m'ont semblé sincères.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q9 J'ai l'impression que l'agent conversationnel ressentait réellement les émotions qu'il exprimait.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q10 J'ai eu le sentiment que l'agent conversationnel feint des émotions pour m'influencer.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord) — [ITEM INVERSÉ]
---------------	---

Q11 L'agent conversationnel utilisait des mots bienveillants dans un but purement stratégique ou commercial.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord) — [ITEM INVERSÉ]
---------------	---

⚠ ATTENTION CHECK — Réponse attendue : 4 (milieu de l'échelle). Exclure les répondants qui répondent 1 ou 7.

QAC1 Afin de vérifier votre attention, veuillez sélectionner « En désaccord » pour cette question.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord) — [Question piège #1]
Hypothèses	Vérification de l'attention

BLOC 4 — PRÉSENCE SOCIALE PERÇUE

Q12 J'ai eu l'impression d'interagir avec une entité vivante et intelligente.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q13 La manière dont l'agent conversationnel s'est exprimé m'a fait (presque) oublier que j'interagissais avec une machine.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q14 J'ai ressenti un véritable sentiment de contact humain durant cette interaction.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q15 Le ton de l'agent conversationnel m'a paru bizarre ou m'a mis(e) mal à l'aise.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord) — [ITEM INVERSÉ]
---------------	---

BLOC 5 — SATISFACTION

Q16 Cette interaction a contribué à diminuer mon niveau de stress ou d'anxiété.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q17 Globalement, je suis satisfait(e) de la manière dont cette interaction s'est déroulée.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q18 J'ai le sentiment que cet établissement bancaire se soucie réellement des intérêts de ses clients.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q19 Cette interaction renforce ma confiance envers l'institution bancaire.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q20 Cette expérience renforcerait mon intention de continuer à utiliser ce service.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q21 Je recommanderais ce service à mon entourage.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord) — [NPS proxy]
---------------	--

BLOC 6 — ENGAGEMENT DU PATIENT

Q22 Suite à cette interaction, je serais prêt(e) à suivre les recommandations financières ou commerciales.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q23 Je me sentirais à l'aise pour poser des questions supplémentaires à cet agent conversationnel.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q24 Cette interaction me donnerait envie de prendre rendez-vous avec un conseiller bancaire humain.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q25 Je serais prêt(e) à partager des informations personnelles supplémentaires avec cet agent.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

Q26 Cette interaction m'a donné le sentiment d'être acteur(trice) de ma santé.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord)
---------------	--

⚠ ATTENTION CHECK #2 — Réponse attendue : « Tout à fait d'accord » (7). Exclure si réponse ≠ 7.

QAC2 Afin de vérifier votre attention, veuillez sélectionner « Tout à fait d'accord » pour cette question.

Format	Échelle de Likert à 7 points (1 = Pas du tout d'accord → 7 = Tout à fait d'accord) — [Question piège #2]
---------------	--

Hypothèses	Vérification de l'attention
-------------------	-----------------------------

BLOC 7 — PROFIL SOCIODÉMOGRAPHIQUE

Q27 Quel est votre genre ?

Format	Choix unique
---------------	--------------

Modalités	• Femme • Homme • Non-binaire / autre • Je préfère ne pas répondre
------------------	---

Q28 Quelle est votre tranche d'âge ?

Format	Choix unique
Modalités	• 18–24 ans • 25–34 ans • 35–44 ans • 45–54 ans • 55–64 ans • 65 ans et plus

Q29 Quel est votre niveau d'études le plus élevé ?

Format	Choix unique
Modalités	• Brevet / baccalauréat ou équivalent • Bac +2 / BTS / DUT • Licence (Bac +3) • Master / Grande école (Bac +5) • Doctorat • Autre

Q30 À quelle fréquence utilisez-vous des chatbots ou assistants virtuels (service client, apps, etc.) ?

Format	Choix unique
Modalités	• Jamais • Rarement (moins d'une fois par mois) • Parfois (1–3 fois par mois) • Souvent (1–3 fois par semaine) • Très souvent (quotidiennement ou presque)

Q31 Avez-vous déjà eu recours à des services bancaires numériques (banque en ligne, application de gestion de compte...) ?

Format	Choix unique
Modalités	• Oui, régulièrement • Oui, quelquefois • Oui, une ou deux fois • Non, jamais

QUESTION FINALE — PRÉFÉRENCE DÉCLARATIVE

Question de clôture — dichotomique, sans valeur de Likert. Ne sert pas à tester les hypothèses principales mais permet un insight complémentaire sur les préférences déclarées.

Q32 Dans une situation financière engageante, quelle interaction préféreriez-vous ?

Format	Choix unique (dichotomique)
Modalités	• Une IA très empathique : elle reconnaît mes émotions, me soutient et s'adapte à mon ressenti • Une IA très efficace et neutre : elle me donne l'information rapidement, sans fioritures émotionnelles

Merci pour votre participation !

Votre contribution est précieuse pour la réalisation de ce mémoire de master. Vos réponses sont entièrement anonymes et seront traitées de façon confidentielle à des fins purement académiques.

Anastassiou Licia

3. Statistiques descriptives

Descriptive Statistics

	N	Minimum	Maximum	Mean	Std. Deviation
Score_engagement	141	1.00	5.75	3.2128	1.20821
Score_satisfaction	141	1.00	6.83	3.0071	1.29803
Score_presence	141	1.00	7.00	2.7784	1.24610
Score_authenticité	141	1.00	7.00	3.4167	1.13021
Score_controle	141	1.00	6.25	3.1543	1.09531
Score_eveil	141	1.00	6.50	3.4273	1.11104
Score_plaisir	141	1.00	7.00	3.6241	1.53024

4. Matrice de corrélation

Correlations

		Sc_plaisir	Sc_evl	Sc_controle	Sc_auth	Sc_pres	Sc_sat	Sc_eng
Sc_plaisir	Pearson Correlation	1	.099	.463**	.356**	.546**	.705**	.419**
	Sig. (2-tailed)		.241	<.001	<.001	<.001	<.001	<.001
	N	141	141	141	141	141	141	141

Sc_éveil	Pearson Correlation	.099	1	.172*	-.069	-.067	.138	-.026
	Sig. (2-tailed)	.241		.041	.415	.433	.102	.763
	N	141	141	141	141	141	141	141
Sc_controle	Pearson Correlation	.463**	.172*	1	.260**	.323**	.460**	.375**
	Sig. (2-tailed)	<.001	.041		.002	<.001	<.001	<.001
	N	141	141	141	141	141	141	141
Sc_authenticité	Pearson Correlation	.356**	-.069	.260**	1	.441**	.448**	.315**
	Sig. (2-tailed)	<.001	.415	.002		<.001	<.001	<.001
	N	141	141	141	141	141	141	141
Sc_présence	Pearson Correlation	.546**	-.067	.323**	.441**	1	.667**	.556**
	Sig. (2-tailed)	<.001	.433	<.001	<.001		<.001	<.001
	N	141	141	141	141	141	141	141
Sc_satisfaction	Pearson Correlation	.705**	.138	.460**	.448**	.667**	1	.723**
	Sig. (2-tailed)	<.001	.102	<.001	<.001	<.001		<.001
	N	141	141	141	141	141	141	141
Sc_engagement	Pearson Correlation	.419**	-.026	.375**	.315**	.556**	.723**	1
	Sig. (2-tailed)	<.001	.763	<.001	<.001	<.001	<.001	
	N	141	141	141	141	141	141	141

** . Correlation is significant at the 0.01 level (2-tailed).

* . Correlation is significant at the 0.05 level (2-tailed).

5. Vérification de la taille des groupes de l'échantillon

Type d'IA * Polarité de la réponse Crosstabulation

Count

		Polarité de la réponse		Total
		Positive	Négative	
Type d'IA	Empathique	37	31	68
	Informationnelle	35	38	73
Total		72	69	141

6. Alpha de Cronbach

Empathie perçue		Eveil		Authenticité perçue	
Cronbach's Alpha	N of Items	Cronbach's Alpha	N of Items	Cronbach's Alpha	N of Items
.887	3	.779	4	.603	4
Satisfaction		Contrôle		Engagement avec tous les items	
Cronbach's Alpha	N of Items	Cronbach's Alpha	N of Items	Cronbach's Alpha	N of Items
.930	6	.721	4	.570	5
Plaisir		Présence sociale perçue		Engagement après item supprimé	
Cronbach's Alpha	N of Items	Cronbach's Alpha	N of Items	Cronbach's Alpha	N of Items
.918	4	.828	4	.738	4

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Cronbach's Alpha if Item Deleted
6. Engagement: - Suite à cette interaction, je serais prêt(e) à suivre les recommandations financières ou commerciales.	15.21	14.340	.516	.404
6. Engagement: - Je me sentirais à l'aise pour poser des questions supplémentaires à cet agent conversationnel.	15.04	13.220	.530	.380
6. Engagement: - Cette interaction me donnerait envie de prendre rendez-vous avec un conseiller bancaire humain.	12.85	23.356	-.190	.738
6. Engagement: - Je serais prêt(e) à partager des informations personnelles supplémentaires avec cet agent.	15.62	14.451	.436	.447
6. Engagement: - Cette interaction m'a donné le sentiment d'être acteur(trice) de ma situation financière.	15.40	15.214	.432	.455

7. Vérification de la manipulation expérimentale

Group Statistics

	Type d'IA	N	Mean	Std. Deviation	Std. Error Mean
Score_empathie	Empathique	68	4.5098	1.24784	.15132
	Informationnelle	73	2.6393	1.25309	.14666

Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means						95% Confidence Interval of the Difference	
		F	Sig.	t	df	Significance One-Sided p	Two-Sided p	Mean Difference	Std. Error Difference	Lower	Upper
Sc_emp	Equal variances assumed	.588	.444	8.875	139	<.001	<.001	1.87053	.21077	1.45381	2.28726
	Equal variances not assumed			8.876	138.373	<.001	<.001	1.87053	.21073	1.45386	2.28721

Independent Samples Effect Sizes

		Standardizer ^a	Point Estimate	95% Confidence Interval	
				Lower	Upper
Score_empathie	Cohen's d	1.25056	1.496	1.120	1.868
	Hedges' correction	1.25736	1.488	1.113	1.858
	Glass's delta	1.25309	1.493	1.079	1.900

a. The denominator used in estimating the effect sizes.

Cohen's d uses the pooled standard deviation.

Hedges' correction uses the pooled standard deviation, plus a correction factor.

Glass's delta uses the sample standard deviation of the control (i.e., the second) group.

8. Hypothèse 1

Group Statistics

	Tonalité	N	Mean	Std. Deviation	Std. Error Mean
Sc_sat	Empathique	68	3.2819	1.36628	.16569

Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means						95% Confidence Interval of the Difference	
		F	Sig.	t	df	Significance One-Sided p	Two-Sided p	Mean Difference	Std. Error Difference	Lower	Upper
Sc_sat	Equal variances assumed	1.060	.305	2.470	139	.007	.015	.53072	.21489	.10585	.95559
	Equal variances not assumed			2.457	132.963	.008	.015	.53072	.21598	.10352	.95793

Independent Samples Effect Sizes

		Standardizer ^a	Point Estimate	95% Confidence Interval	
				Lower	Upper
Sc_sat	Cohen's d	1.27501	.416	.082	.749
	Hedges' correction	1.28194	.414	.081	.745
	Glass's delta	1.18378	.448	.109	.785

a. The denominator used in estimating the effect sizes.

Cohen's d uses the pooled standard deviation.

Hedges' correction uses the pooled standard deviation, plus a correction factor.

Glass's delta uses the sample standard deviation of the control (i.e., the second) group.

9. Hypothèse 2

Group Statistics

	Tonalité	N	Mean	Std. Deviation	Std. Error Mean
Sc_eng	Empathique	68	3.4007	1.24074	.15046
	Neutre	73	3.0377	1.15821	.13556

Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means						95% Confidence Interval of the Difference	
		F	Sig.	t	df	Significance One-Sided p	Significance Two-Sided p	Mean Difference	Std. Error Difference	Lower	Upper
Sc_eng	Equal variances assumed	.110	.740	1.797	139	.037	.074	.36306	.20202	-.0363	.7625
	Equal variances not assumed			1.793	136.328	.038	.075	.36306	.20252	-.0374	.7635

Independent Samples Effect Sizes

		Standardizer ^a	Point Estimate	95% Confidence Interval	
				Lower	Upper
Sc_eng	Cohen's d	1.19870	.303	-.030	.635
	Hedges' correction	1.20522	.301	-.030	.631
	Glass's delta	1.15821	.313	-.022	.647

10. Hypothèse 3

Group Statistics

Polarité de la réponse		N	Mean	Std. Deviation	Std. Error Mean
Sc_sat	Positive	72	3.2940	1.38979	.16379
	Négative	69	2.7077	1.12873	.13588

Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means						95% Confidence Interval of the Difference	
		F	Sig.	t	df	Significance One-Sided p	Two-Sided p	Mean Difference	Std. Error Difference	Lower	Upper
Sc_sat	Equal variances assumed	1.657	.200	2.743	139	.003	.007	.58625	.21375	.16362	1.00888
	Equal variances not assumed			2.755	135.398	.003	.007	.58625	.21282	.16538	1.00713

Independent Samples Effect Sizes

		Standardizer ^a	Point Estimate	95% Confidence Interval	
				Lower	Upper
Sc_sat	Cohen's d	1.26881	.462	.127	.796
	Hedges' correction	1.27570	.460	.126	.792
	Glass's delta	1.12873	.519	.176	.859

11. Hypothèse 4

Group Statistics

	Polarité de la réponse	N	Mean	Std. Deviation	Std. Error Mean
Sc_eng	Positive	72	3.2639	1.29523	.15264
	Négative	69	3.1594	1.11719	.13449

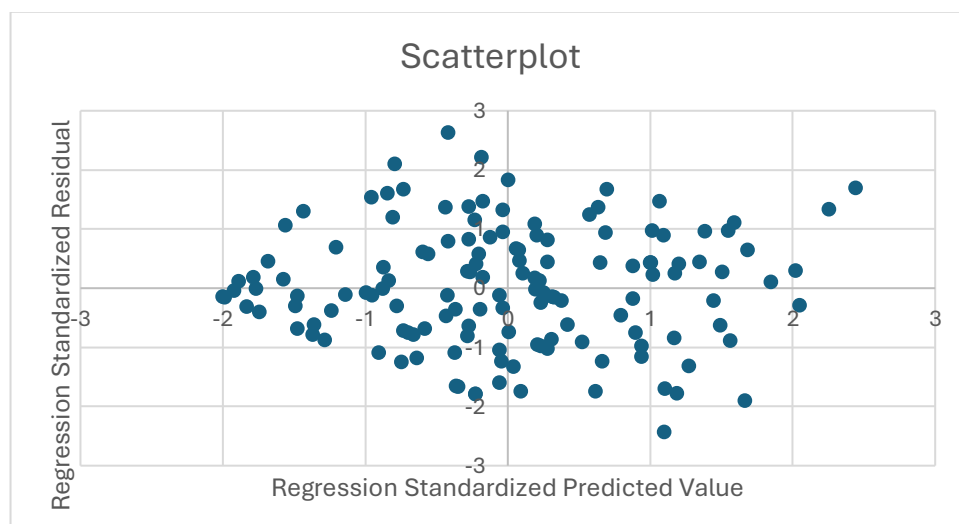
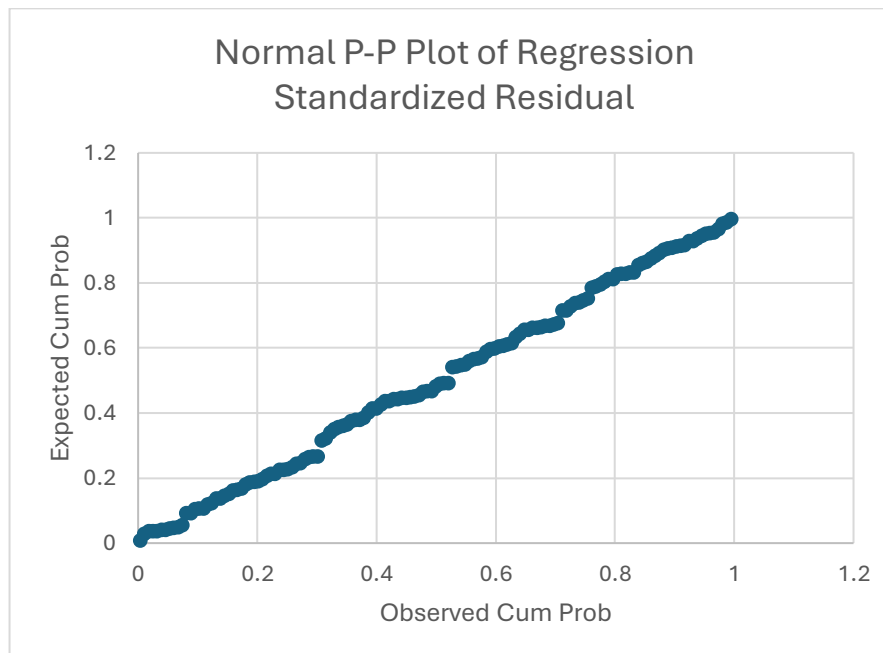
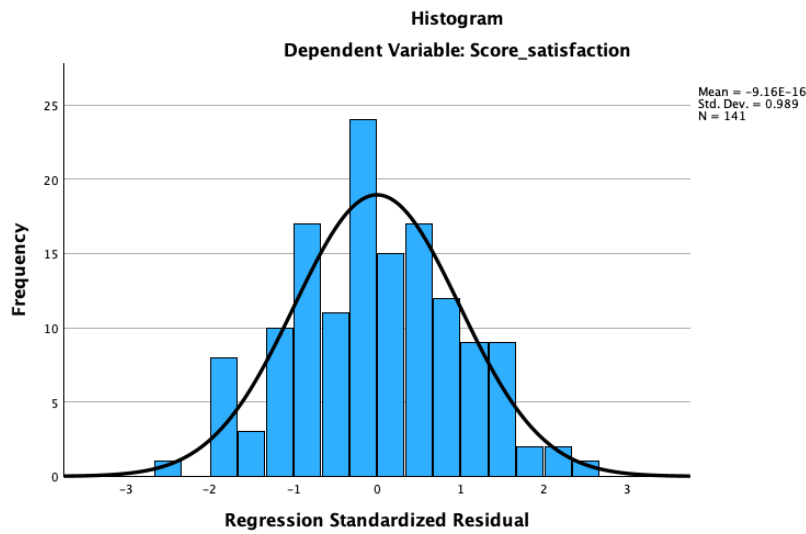
Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means						95% Confidence Interval of the Difference	
		F	Sig.	t	df	Significance One-Sided p	Two-Sided p	Mean Difference	Std. Error Difference	Lower	Upper
Sc_eng	Equal variances assumed	1.536	.217	.512	139	.305	.610	.10447	.20408	-.2990	.5079
	Equal variances not assumed			.514	137.502	.304	.608	.10447	.20344	-.2978	.5067

Independent Samples Effect Sizes

		Standardizer ^a	Point Estimate	95% Confidence Interval	
				Lower	Upper
Sc_eng	Cohen's d	1.21140	.086	-.244	.416
	Hedges' correction	1.21799	.086	-.243	.414
	Glass's delta	1.11719	.094	-.237	.424

12. Hypothèse 5



ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	123.201	3	41.067	49.930	<.001 ^b
	Residual	112.681	137	.822		
	Total	235.882	140			

a. Dependent Variable: Score_satisfaction

b. Predictors: (Constant), Score_controle, Score_eveil, Score_plaisir

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	.723 ^a	.522	.512	.90691	2.245

a. Predictors: (Constant), Score_controle, Score_eveil, Score_plaisir

b. Dependent Variable: Score_satisfaction

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Collinearity Statistics	
		B	Std. Error	Beta			Tolerance	VIF
1	(Constant)	.285	.320		.890	.375		
	Score_plaisir	.530	.057	.625	9.382	<.001	.786	1.273
	Score_eveil	.056	.070	.048	.802	.424	.970	1.031
	Score_controle	.193	.080	.163	2.418	.017	.770	1.299

a. Dependent Variable: Score_satisfaction

13. Hypothèse 6

```

Model: 4
  Y: Sc_eng (= score de l'engagement)
  X: Tonalité
  M: Sc_auth (= score de l'authenticité perçue)

Sample
Size: 141

*****

OUTCOME VARIABLE:
  Sc_auth

Model Summary
      R      R-sq      MSE      F      df1      df2      p
    .1239    .0154    1.2668    2.1681    1.0000    139.0000    .1432

Model
      coeff      se      t      p      LLCI      ULCI
constant    2.9927    .3031    9.8737    .0000    2.3935    3.5920
Tonalité     .2793    .1897    1.4724    .1432    -.0957    .6544

*****

OUTCOME VARIABLE:
  Sc_eng

Model Summary
      R      R-sq      MSE      F      df1      df2      p
    .3686    .1359    1.2797    10.8509    2.0000    138.0000    .0000

Model
      coeff      se      t      p      LLCI      ULCI
constant    2.6791    .3974    6.7423    .0000    1.8934    3.4648
Tonalité    -.4643    .1921    -2.4165    .0170    -.8442    -.0844
Sc_auth     .3624    .0852    4.2515    .0000    .1939    .5310

***** TOTAL EFFECT MODEL *****

OUTCOME VARIABLE:
  Sc_eng

Model Summary
      R      R-sq      MSE      F      df1      df2      p
    .1507    .0227    1.4369    3.2297    1.0000    139.0000    .0745

Model
      coeff      se      t      p      LLCI      ULCI
constant    3.7638    .3228    11.6595    .0000    3.1255    4.4021
Tonalité    -.3631    .2020    -1.7971    .0745    -.7625    .0364

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y
      Effect      se      t      p      LLCI      ULCI
    -.3631    .2020    -1.7971    .0745    -.7625    .0364

Direct effect of X on Y
      Effect      se      t      p      LLCI      ULCI
    -.4643    .1921    -2.4165    .0170    -.8442    -.0844

Indirect effect(s) of X on Y:
      Effect      BootSE      BootLLCI      BootULCI
Sc_auth     .1012      .0753      -.0307      .2694

```

```

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
  95.0000

Number of bootstrap samples for percentile bootstrap confidence intervals:
  5000

----- END MATRIX -----

```

14. Hypothèse 7

```

Model: 4
  Y: Sc_sat (=score de satisfaction)
  X: Tonalité
  M: Sc_pres (=score de présence sociale perçue)

Sample
Size: 141

*****

OUTCOME VARIABLE:
  Sc_pres

Model Summary
      R      R-sq      MSE      F      df1      df2      p
    .1951    .0381    1.5044    5.5023    1.0000    139.0000    .0204

Model
      coeff      se      t      p      LLCI      ULCI
constant    3.5143    .3303    10.6396    .0000    2.8612    4.1674
Tonalité    -.4849    .2067    -2.3457    .0204    -.8936    -.0762

*****

OUTCOME VARIABLE:
  Sc_sat

Model Summary
      R      R-sq      MSE      F      df1      df2      p
    .6711    .4503    .9395    56.5298    2.0000    138.0000    .0000

Model
      coeff      se      t      p      LLCI      ULCI
constant    1.4276    .3516     4.0602    .0001     .7324    2.1228
Tonalité    -.2016    .1666    -1.2106    .2281    -.5310    .1277
Sc_pres     .6786    .0670    10.1245    .0000     .5461     .8112

***** TOTAL EFFECT MODEL *****

OUTCOME VARIABLE:
  Sc_sat

Model Summary
      R      R-sq      MSE      F      df1      df2      p
    .2050    .0420    1.6257     6.0998    1.0000    139.0000    .0147

Model
      coeff      se      t      p      LLCI      ULCI
constant    3.8126    .3434    11.1038    .0000    3.1337    4.4915
Tonalité    -.5307    .2149    -2.4698    .0147    -.9556    -.1059

```

```

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y
      Effect      se      t      p      LLCI      ULCI
      -.5307      .2149     -2.4698     .0147     -.9556     -.1059

Direct effect of X on Y
      Effect      se      t      p      LLCI      ULCI
      -.2016      .1666     -1.2106     .2281     -.5310     .1277

Indirect effect(s) of X on Y:

      Effect      BootSE      BootLLCI      BootULCI
Sc_pres.      -.3291      .1487      -.6446      -.0525

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
95.0000

Number of bootstrap samples for percentile bootstrap confidence intervals:
5000

----- END MATRIX -----

```

15. Hypothèse exploratoire 1

Tests of Normality

	Type d'IA	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
		Statistic	df	Sig.	Statistic	df	Sig.
Sc_pres_soc	Empathique	.131	68	.005	.946	68	.006
	Informationnel	.116	73	.017	.933	73	<.001

a. Lilliefors Significance Correction

```

Model: 4
      Y: Sc_eng (=score de l'engagement)
      X: Tonalité
      M: Sc_pres (=score de la présence sociale perçue)

Sample
Size: 141

*****

OUTCOME VARIABLE:
Sc_pres

Model Summary
      R      R-sq      MSE      F      df1      df2      p
      .1951      .0381      1.5044      5.5023      1.0000      139.0000      .0204

Model
      coeff      se      t      p      LLCI      ULCI
constant      3.0294      .1487      20.3673      .0000      2.7353      3.3235
Tonalite      -.4849      .2067      -2.3457      .0204      -.8936      -.0762

```

```

*****
OUTCOME VARIABLE:
  Sc_eng

Model Summary
      R      R-sq      MSE      F      df1      df2      p
      .5572      .3104      1.0212      31.0639      2.0000      138.0000      .0000

Model
      coeff      se      t      p      LLCI      ULCI
constant      1.7943      .2446      7.3352      .0000      1.3106      2.2779
Tonalite      -.1059      .1737      -.6100      .5428      -.4493      .2374
Sc_pres      .5303      .0699      7.5884      .0000      .3921      .6685

***** TOTAL EFFECT MODEL *****

OUTCOME VARIABLE:
  Sc_eng

Model Summary
      R      R-sq      MSE      F      df1      df2      p
      .1507      .0227      1.4369      3.2297      1.0000      139.0000      .0745

Model
      coeff      se      t      p      LLCI      ULCI
constant      3.4007      .1454      23.3947      .0000      3.1133      3.6881
Tonalite      -.3631      .2020      -1.7971      .0745      -.7625      .0364

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y
      Effect      se      t      p      LLCI      ULCI
      -.3631      .2020      -1.7971      .0745      -.7625      .0364

Direct effect of X on Y
      Effect      se      t      p      LLCI      ULCI
      -.1059      .1737      -.6100      .5428      -.4493      .2374

Indirect effect(s) of X on Y:
      Effect      BootSE      BootLLCI      BootULCI
Sc_pres      -.2571      .1175      -.4976      -.0412

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
  95.0000

Number of bootstrap samples for percentile bootstrap confidence intervals:
  5000

----- END MATRIX -----

```


16. Hypothèse 8

```

Model: 1
  Y: Sc_sat (= score de satisfaction)
  X: Tonalité
  W: Polarité

Sample
Size: 141

*****

OUTCOME VARIABLE:
  Sc_sat

Model Summary
      R      R-sq      MSE      F      df1      df2      p
      .3164      .1001      1.5494      5.0801      3.0000      137.0000      .0023

Model
      coeff      se      t      p      LLCI      ULCI
constant      5.8782      1.0430      5.6360      .0000      3.8158      7.9406
Tonalité      -1.3531      .6598      -2.0508      .0422      -2.6578      -.0484
Polarité      -1.4311      .6727      -2.1276      .0352      -2.7613      -.1010
Int_1          .5771      .4206      1.3722      .1722      -.2545      1.4088

Product terms key:
Int_1 : Tonalité x Polarité

Test(s) of highest order unconditional interaction(s):
      R2-chng      F      df1      df2      p
X*W      .0124      1.8830      1.0000      137.0000      .1722
-----
      Focal predict: Tonalité (X)
      Mod var: Polarité (W)

Data for visualizing the conditional effect of the focal predictor:
Paste text below into a SPSS syntax window and execute to produce plot.

DATA LIST FREE/
  Tonalité  Polarité  Sc_sat  .
BEGIN DATA.
  1.0000    1.0000    3.6712
  2.0000    1.0000    2.8952
  1.0000    2.0000    2.8172
  2.0000    2.0000    2.6184
END DATA.
GRAPH/SCATTERPLOT=
  Tonalité WITH  Sc_sat  BY      Polarité .

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
  95.0000

----- END MATRIX -----

```

17. Hypothèse 9

```

Model: 1
  Y: Sc_eng (=score engagement)
  X: Tonalite
  W: Polarite

Sample
Size: 141

*****

OUTCOME VARIABLE:
  Sc_eng

Model Summary

      R      R-sq      MSE      F      df1      df2      p
      .2229      .0497      1.4176      2.3874      3.0000      137.0000      .0717

Model

      coeff      se      t      p      LLCI      ULCI
constant      3.6216      .1957      18.5022      .0000      3.2346      4.0087
Tonalite      -.7359      .2807      -2.6213      .0098      -1.2911      -.1808
Polarite      -.4845      .2899      -1.6713      .0969      -1.0578      .0887
Int_1      .7764      .4023      1.9300      .0557      -.0191      1.5720

Product terms key:
  Int_1      :      Tonalite x      Polarite

Test(s) of highest order unconditional interaction(s):

      R2-chng      F      df1      df2      p
X*W      .0258      3.7248      1.0000      137.0000      .0557
-----
      Focal predict: Tonalite (X)
      Mod var: Polarite (W)

Conditional effects of the focal predictor at values of the moderator(s):

      Polarite      Effect      se      t      p      LLCI      ULCI
      .0000      -.7359      .2807      -2.6213      .0098      -1.2911      -.1808
      1.0000      .0405      .2882      .1407      .8883      -.5293      .6103

Data for visualizing the conditional effect of the focal predictor:
Paste text below into a SPSS syntax window and execute to produce plot.

DATA LIST FREE/
  Tonalite Polarite Sc_eng .
BEGIN DATA.
  .0000      .0000      3.6216
  1.0000      .0000      2.8857
  .0000      1.0000      3.1371
  1.0000      1.0000      3.1776
END DATA.
GRAPH/SCATTERPLOT=
  Tonalite WITH      Sc_eng BY      Polarite .

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
  95.0000

----- END MATRIX -----

```

18. Hypothèse exploratoire 2

```
Model: 1
  Y: Sc_pres
  X: Tonalite
  W: Polarite

Sample
Size: 141

*****

OUTCOME VARIABLE:
  Sc_pres

Model Summary
      R      R-sq      MSE      F      df1      df2      p
      .2413      .0582      1.4944      2.8237      3.0000      137.0000      .0411

Model
      coeff      se      t      p      LLCI      ULCI
constant      3.1757      .2010      15.8019      .0000      2.7783      3.5731
Tonalite      -.4328      .2882      -1.5016      .1355      -1.0028      .1372
Polarite      -.3208      .2976      -1.0779      .2830      -.9094      .2677
Int_1      -.0602      .4131      -.1457      .8844      -.8770      .7566

Product terms key:
  Int_1      :      Tonalite x      Polarite

Test(s) of highest order unconditional interaction(s):
      R2-chng      F      df1      df2      p
X*W      .0001      .0212      1.0000      137.0000      .8844
-----
      Focal predict: Tonalite (X)
      Mod var: Polarite (W)

Data for visualizing the conditional effect of the focal predictor:
Paste text below into a SPSS syntax window and execute to produce plot.

DATA LIST FREE/
  Tonalite  Polarite  Sc_pres  .
BEGIN DATA.
  .0000      .0000      3.1757
  1.0000      .0000      2.7429
  .0000      1.0000      2.8548
  1.0000      1.0000      2.3618
END DATA.
GRAPH/SCATTERPLOT=
  Tonalite WITH      Sc_pres  BY      Polarite .

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
  95.0000

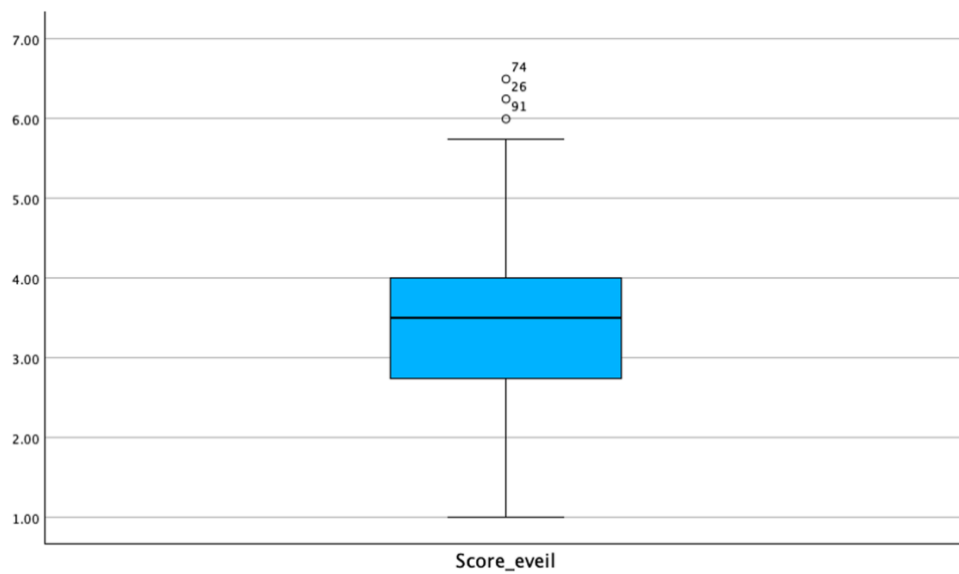
----- END MATRIX -----
```

19. Hypothèse 10

Tests of Normality

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Score_eveil	.105	141	<.001	.980	141	.034

a. Lilliefors Significance Correction



```

Model: 1
  Y: Sc_sat (=score de satisfaction)
  X: Tona    (=tonalité)
  W: Sc_evl  (=score de l'éveil)

Sample
Size: 141

*****

OUTCOME VARIABLE:
  Sc_sat

Model Summary

      R      R-sq      MSE      F      df1      df2      p
.2948    .0869    1.5722    4.3450    3.0000    137.0000    .0059

Model

      coeff      se      t      p      LLCI      ULCI
constant    3.0083    .1056   28.4885    .0000    2.7995    3.2171
Tona        -.5287    .2113   -2.5020    .0135   -.9466   -.1109
Sc_evl      .1710    .0955    1.7899    .0757   -.0179    .3599
Int_1       .3773    .1908    1.9772    .0500    .0000    .7546

Product terms key:
Int_1 :      Tona      x      Sc_evl

Test(s) of highest order unconditional interaction(s):

      R2-chng      F      df1      df2      p
X*W    .0261    3.9094    1.0000    137.0000    .0500
-----
      Focal predict: Tona      (X)
      Mod var: Sc_evl      (W)

Conditional effects of the focal predictor at values of the moderator(s):

      Sc_evl      Effect      se      t      p      LLCI      ULCI
-1.1773    -.9729    .3084   -3.1544    .0020   -1.5828   -.3630
.0727     -.5013    .2118   -2.3671    .0193   -.9201   -.0825
.8227     -.2184    .2632   -.8295    .4083   -.7389    .3022

Data for visualizing the conditional effect of the focal predictor:
Paste text below into a SPSS syntax window and execute to produce plot.

DATA LIST FREE/
  Tona      Sc_evl      Sc_sat      .
BEGIN DATA.
  -.5177    -1.1773    3.3107
  .4823     -1.1773    2.3378
  -.5177     .0727    3.2803
  .4823     .0727    2.7789
  -.5177     .8227    3.2620
  .4823     .8227    3.0437
END DATA.
GRAPH/SCATTERPLOT=
  Sc_evl      WITH      Sc_sat      BY      Tona      .

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
95.0000

W values in conditional tables are the 16th, 50th, and 84th percentiles.

NOTE: The following variables were mean centered prior to analysis:
      Sc_evl      Tona

```

Analyse de Sensibilité excluant 3 valeurs d'éveil extrêmes

```

Model: 1
  Y: Sc_sat (=score de satisfaction)
  X: Tona
  W: Sc_evl (=score de l'éveil émotionnel)

Sample
Size: 138

*****

OUTCOME VARIABLE:
  Sc_sat

Model Summary

      R      R-sq      MSE      F      df1      df2      p
      .3164      .1001      1.5761      4.9681      3.0000      134.0000      .0027

Model

      coeff      se      t      p      LLCI      ULCI
constant      3.0180      .1069      28.2381      .0000      2.8067      3.2294
Tona      -.5568      .2140      -2.6025      .0103      -.9800      -.1337
Sc_evl      .2340      .1032      2.2682      .0249      .0300      .4381
Int_1      .3798      .2062      1.8418      .0677      -.0280      .7877

Product terms key:
  Int_1      :      Tona      x      Sc_evl

Test(s) of highest order unconditional interaction(s):

      R2-chng      F      df1      df2      p
X*W      .0228      3.3924      1.0000      134.0000      .0677
-----
      Focal predict: Tona      (X)
      Mod var: Sc_evl      (W)

Conditional effects of the focal predictor at values of the moderator(s):

      Sc_evl      Effect      se      t      p      LLCI      ULCI
-1.1159      -.9807      .3141      -3.1220      .0022      -1.6020      -.3594
.1341      -.5059      .2158      -2.3448      .0205      -.9326      -.0792
.8841      -.2210      .2812      -.7860      .4332      -.7772      .3351

Data for visualizing the conditional effect of the focal predictor:
Paste text below into a SPSS syntax window and execute to produce plot.

DATA LIST FREE/
  Tona      Sc_evl      Sc_sat      .
BEGIN DATA.
  -.5217      -1.1159      3.2686
  .4783      -1.1159      2.2879
  -.5217      .1341      3.3134
  .4783      .1341      2.8075
  -.5217      .8841      3.3402
  .4783      .8841      3.1192
END DATA.
GRAPH/SCATTERPLOT=
  Sc_evl      WITH      Sc_sat      BY      Tona      .

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
  95.0000

W values in conditional tables are the 16th, 50th, and 84th percentiles.

NOTE: The following variables were mean centered prior to analysis:
      Sc_evl      Tona

```

20. Hypothèse 11

Tests of Normality

Polarité de la réponse		Kolmogorov-Smirnov ^a			Shapiro-Wilk		
		Statistic	df	Sig.	Statistic	df	Sig.
Sc_auth	Positive	.139	72	.001	.971	72	.093
	Négative	.123	69	.012	.967	69	.066

a. Lilliefors Significance Correction

Tests of Normality

Type d'IA		Kolmogorov-Smirnov ^a			Shapiro-Wilk		
		Statistic	df	Sig.	Statistic	df	Sig.
Sc_auth	Empathique	.083	68	.200*	.971	68	.113
	Informationnelle	.131	73	.003	.967	73	.054

*. This is a lower bound of the true significance.

a. Lilliefors Significance Correction

```

Model: 1
  Y: Sc_auth (=score authenticité perçue)
  X: Tonalite
  W: Polarite

Sample
Size: 141

*****

OUTCOME VARIABLE:
  Sc_auth

Model Summary
      R      R-sq      MSE      F      df1      df2      p
      .1488      .0222      1.2764      1.0346      3.0000      137.0000      .3794

Model
      coeff      se      t      p      LLCI      ULCI
constant      3.3919      .1857      18.2618      .0000      3.0246      3.7592
Tonalite      .1867      .2664      .7008      .4846      -.3401      .7135
Polarite      -.2629      .2751      -.9555      .3410      -.8068      .2811
Int_1      .2106      .3818      .5517      .5821      -.5443      .9655

Product terms key:
Int_1 : Tonalite x Polarite

Test(s) of highest order unconditional interaction(s):
      R2-chng      F      df1      df2      p
X*W      .0022      .3044      1.0000      137.0000      .5821
-----
      Focal predict: Tonalite (X)
      Mod var: Polarite (W)

Data for visualizing the conditional effect of the focal predictor:
Paste text below into a SPSS syntax window and execute to produce plot.

DATA LIST FREE/
  Tonalite  Polarite  Sc_auth  .
BEGIN DATA.
      .0000      .0000      3.3919
      1.0000      .0000      3.5786
      .0000      1.0000      3.1290
      1.0000      1.0000      3.5263
END DATA.
GRAPH/SCATTERPLOT=
  Tonalite WITH      Sc_auth  BY      Polarite .

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
  95.0000

----- END MATRIX -----

```


21. Hypothèse 12

```

Model: 1
  Y: Sc_eng (=score engagement)
  X: Tonalite
  W: Sc_evl (=score de l'éveil)

Sample
Size: 141

*****

OUTCOME VARIABLE:
  Sc_eng

Model Summary
      R      R-sq      MSE      F      df1      df2      p
      .1789      .0320      1.4440      1.5103      3.0000      137.0000      .2146

Model
      coeff      se      t      p      LLCI      ULCI
constant      3.8384      .4626      8.2976      .0000      2.9237      4.7532
Tonalite      -1.0551      .6587      -1.6018      .1115      -2.3575      .2474
Sc_evl      -.1275      .1279      -.9970      .3205      -.3803      .1254
Int_1      .2018      .1829      1.1035      .2717      -.1598      .5634

Product terms key:
Int_1 : Tonalite x Sc_evl

Test(s) of highest order unconditional interaction(s):
      R2-chng      F      df1      df2      p
X*W      .0086      1.2177      1.0000      137.0000      .2717
-----
      Focal predict: Tonalite (X)
      Mod var: Sc_evl (W)

Data for visualizing the conditional effect of the focal predictor:
Paste text below into a SPSS syntax window and execute to produce plot.

DATA LIST FREE/
  Tonalite Sc_evl Sc_eng .
BEGIN DATA.
  .0000      2.2500      3.5516
  1.0000      2.2500      2.9506
  .0000      3.5000      3.3923
  1.0000      3.5000      3.0435
  .0000      4.2500      3.2967
  1.0000      4.2500      3.0993
END DATA.
GRAPH/SCATTERPLOT=
  Sc_evl WITH Sc_eng BY Tonalite .

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
  95.0000

----- END MATRIX -----

```

22. Hypothèse 13

```

Model: 4
  Y: Sc_sat (score de satisfaction)
  X: Tonalite
  M: Sc_auth (score de l'authenticité perçue)

Sample
Size: 141

*****

OUTCOME VARIABLE:
  Sc_auth

Model Summary
      R      R-sq      MSE      F      df1      df2      p
      .1239      .0154      1.2668      2.1681      1.0000      139.0000      .1432

Model
      coeff      se      t      p      LLCI      ULCI
constant      3.2721      .1365      23.9728      .0000      3.0022      3.5419
Tonalite      .2793      .1897      1.4724      .1432      -.0957      .6544

*****

OUTCOME VARIABLE:
  Sc_sat

Model Summary
      R      R-sq      MSE      F      df1      df2      p
      .5196      .2700      1.2478      25.5180      2.0000      138.0000      .0000

Model
      coeff      se      t      p      LLCI      ULCI
constant      1.4738      .3070      4.8014      .0000      .8669      2.0807
Tonalite      -.6851      .1897      -3.6108      .0004      -1.0602      -.3099
Sc_auth      .5526      .0842      6.5642      .0000      .3861      .7190

***** TOTAL EFFECT MODEL *****

OUTCOME VARIABLE:
  Sc_sat

Model Summary
      R      R-sq      MSE      F      df1      df2      p
      .2050      .0420      1.6257      6.0998      1.0000      139.0000      .0147

Model
      coeff      se      t      p      LLCI      ULCI
constant      3.2819      .1546      21.2257      .0000      2.9762      3.5876
Tonalite      -.5307      .2149      -2.4698      .0147      -.9556      -.1059

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y
      Effect      se      t      p      LLCI      ULCI
      -.5307      .2149      -2.4698      .0147      -.9556      -.1059

Direct effect of X on Y
      Effect      se      t      p      LLCI      ULCI
      -.6851      .1897      -3.6108      .0004      -1.0602      -.3099

Indirect effect(s) of X on Y:
      Effect      BootSE      BootLLCI      BootULCI
Sc_auth      .1543      .1067      -.0516      .3644

```

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
95.0000

Number of bootstrap samples for percentile bootstrap confidence intervals:
5000

----- END MATRIX -----

