

Les déterminants de l'acceptation des robots de service

Inventaire documenté des antécédents, arbitrage de validité discriminante et modèle à tester

Dossier de Recherche Appliquée — L3 MHR et L3 MRC en alternance — 2026-2027

ISTHIA — Université Toulouse Jean Jaurès — version 2, 24 septembre 2026

1. Objet, périmètre et méthode

Cette note recense, à partir de la littérature scientifique publiée et vérifiée, les antécédents de l'intention d'accepter un robot de service dans quatre configurations distinctes de l'hôtellerie-restauration : le robot qui nettoie les chambres d'hôtel, le robot serveur en restauration commerciale, le service robotisé en restauration collective, et le robot cuisinier qui prépare le plat devant le convive et le lui remet. Elle en tire un modèle unique à tester, décliné en trois versions correspondant aux trois groupes d'étudiants, et documente explicitement les arbitrages de validité discriminante qui ont conduit à retenir certains construits plutôt que d'autres.

L'inventaire détaillé se trouve dans le classeur Excel qui accompagne cette note : un onglet par configuration, quatorze colonnes par antécédent (définition opérationnelle, recherches qui le mobilisent, contexte et méthode de ces recherches, sens de l'effet, coefficient vérifié lorsqu'il existe, échelle de mesure d'origine, justification contextuelle, statut empirique, risque de recouvrement et décision retenue pour le modèle).

1.1. Règle de vérification des sources

Trois cent neuf références sont mobilisées (128 dans la version 1, 181 ajoutées dans la version 2). Chacune a été contrôlée individuellement auprès de l'API Crossref : identifiant DOI, titre exact, liste d'auteurs, année, revue, volume et pagination. Les entrées bibliographiques au format APA 7 ont été produites directement à partir des métadonnées Crossref, sans ressaisie manuelle, ce qui élimine le risque d'erreur de transcription. Aucune référence dont l'existence n'a pas pu être confirmée n'a été conservée.

Une seconde règle s'applique au contenu : un résultat n'est jamais déduit d'un titre. Lorsque le résumé ou le texte intégral n'a pas pu être consulté, le classeur porte la mention correspondante et aucun sens d'effet n'est avancé. Cette distinction entre existence bibliographique et contenu vérifié est indiquée ligne par ligne.

Une réserve doit être signalée sur les dates. Crossref renvoie fréquemment l'année de mise en ligne, distincte de l'année du numéro ; l'écart atteint un à deux ans pour plusieurs articles du corpus. La bibliographie v2 retient l'année du numéro lorsque Crossref la fournit, et ne retombe sur la date de mise en ligne qu'à défaut — ce qui corrige, par exemple, la date de Henseler, Ringle et Sarstedt (2015).

1.2. Deux absences qui sont des résultats

Deux constats méritent d'être posés d'emblée, parce qu'ils conditionnent tout le reste.

Premièrement, aucune étude empirique publiée ne teste un modèle d'acceptation sur le cas précis d'un robot qui nettoie une chambre d'hôtel occupée. Le seul travail qui s'en approche réellement est celui de Hoang et Tran (2022), publié dans *Tourism Management*, qui porte sur des robots nettoyeurs évalués par des consommateurs dans des lieux touristiques, hôtels et aéroports. Trois autres travaux — Ivanov et Webster (2019), Alma Çallı et al. (2022), Binesh et Baloglu (2023) — mesurent l'intention d'usage ventilée par département hôtelier, le housekeeping y figurant explicitement, mais sans modèle causal propre à cette tâche.

Deuxièmement, aucune étude empirique centrée sur l'utilisateur n'a pu être identifiée en restauration collective. Les requêtes croisant restauration institutionnelle, cantine d'entreprise, restauration universitaire, restauration

scolaire, restauration hospitalière et milieu pénitentiaire avec robot, acceptation, intention, TAM et UTAUT ne renvoient que trois types de résultats : de l'ingénierie informatique sans modèle comportemental, de la recherche nutritionnelle et environnementale sur la restauration scolaire, et des travaux sur le robot d'assistance sociale ou logistique en établissement de soin qui n'abordent pas le service alimentaire du point de vue du convive.

Ces deux absences ne sont pas un défaut de recension : elles délimitent l'espace de contribution du projet, et elles imposent de raisonner par transposition argumentée plutôt que par réplique.

Version 2 — ce qui a changé après évaluation

Cette version intègre les évaluations de deux relecteurs indépendants (Gemini et Mistral) et deux demandes complémentaires : examiner la littérature dont les terrains sont en Asie, et combler les angles morts signalés — préférence humain-robot selon la nature de la tâche, congruence robot-marque, revues francophones. Les deux relecteurs valident le diagnostic de validité discriminante et l'architecture socle-extensions ; leurs objections portent sur quatre points, tous traités ici. Le détail remarque par remarque figure dans l'onglet « Réponse aux évaluateurs » du classeur.

Le modèle G3 est allégé

Les deux relecteurs jugeaient le modèle du groupe restauration collective surchargé, et ils avaient raison : socle, extensions alimentaires, coercition et menace sur l'emploi portaient le nombre de construits latents au-delà de dix. L'arbitrage est le suivant. Le groupe G3 conserve le socle (cinq construits), l'amour perçu dans l'aliment comme médiateur de la manipulation de transparence, la menace sur l'emploi du personnel connu, et la confiance dans le respect des régimes et des allergènes — soit huit antécédents latents. La coercition perçue est rétrogradée en variable de contexte : elle est mesurée par deux ou trois items d'attractivité des alternatives (Lee & Kim, 2022) et utilisée en modérateur. Le dégoût moral n'est plus un construit distinct : Pancer et al. (2025) établissent qu'il est l'expression affective de l'opposition au remplacement de l'emploi, il devient donc la dimension affective de la menace sur l'emploi, dont l'unidimensionnalité sera vérifiée au pré-test.

La variable dépendante est double, dans les trois groupes

Mistral objectait à juste titre que si G3 mesure l'intention d'évitement et les autres groupes l'intention d'accepter, la comparaison multi-groupes porte sur des objets différents. Les deux intentions sont désormais mesurées dans les trois groupes : l'acceptation par les trois items de willingness de Gursoy et al. (2019), l'évitement par les quatre items de Lee et Kim (2026) adaptés au contexte. La littérature justifie deux construits et non un continuum : willingness et objection ne corrélaient qu'à $-0,613$ chez Gursoy et al. (2019), la théorie du double facteur de Cenfetelli (2004) établit que les inhibiteurs ne sont pas l'inverse des facilitateurs, et la théorie de l'évitement des menaces technologiques de Liang et Xue (2009) distingue l'évitement d'une menace de l'adoption d'un bénéfice. La distinction sera vérifiée par analyse factorielle confirmatoire (HTMT $< 0,85$) et la corrélation rapportée par groupe — on attend qu'elle soit plus faible en contexte captif, signe d'ambivalence. H25 porte sur l'acceptation ; l'évitement reçoit des antécédents asymétriques.

Une prémisse corrigée sur la menace sur l'emploi

La relecture approfondie a révélé une erreur d'attribution dans la version 1 : le coefficient non significatif de Molinillo et al. (2022) cité comme preuve de non-significativité de la menace sur l'emploi ne mesure pas l'emploi — leurs « objections à l'usage » mesurent la préférence pour le contact humain, et les auteurs proposent eux-mêmes les pertes d'emploi comme construit à ajouter. La preuve de non-significativité en restauration commerciale vient d'Etemad-Sajadi, Soussan et Schöpfer (2022) : sur 341 répondants suisses exposés à une vidéo de Pepper en service, le remplacement des travailleurs est la préoccupation éthique la plus élevée (moyenne de 5,25 sur 7) mais n'a aucun effet significatif sur l'intention d'utiliser le robot. L'hypothèse H23 est réancrée sur Granulo, Fuchs et Puntoni (2019), qui montrent que la réaction au remplacement

robotique s'inverse selon que la personne remplacée est soi ou autrui, et sur Granulo et al. (2025), qui montrent que la réaction négative aux licenciements est amplifiée par la proximité sociale avec les travailleurs touchés. À la connaissance des recherches menées ici, l'hypothèse que la menace devienne significative lorsque le personnel remplacé est personnellement connu n'a été testée nulle part.

Deux précisions et une référence confirmée

Gemini craignait que les items d'absence de contact humain en G1 captent une misanthropie ou une anxiété sociale générale. Il n'y a pas de construit séparé : le socle mesure le besoin d'interaction humaine avec la même échelle de Dabholkar (1996) dans les trois groupes, et H4b ne prédit que l'inversion de son signe en hébergement. Le mécanisme est désormais documenté par la littérature sur la gêne anticipée et la métaperception (Pitardi, Wirtz, Paluch & Kunz, 2022, 2024 ; Holthöwer & van Doorn, 2023) : le client se sent moins jugé par un robot et le préfère pour les situations qui l'exposent. Enfin, Mistral n'avait pas pu confirmer Topsakal, Yüzbaşıoğlu et Uysal (2026) : la référence existe (International Journal of Consumer Studies, volume 50, numéro 4, Wiley, mise en ligne le 12 juillet 2026, DOI 10.1111/ijcs.70232) ; son résumé n'ayant pas été lu, elle n'est citée que pour son existence.

2. Ce que dit la littérature, configuration par configuration

2.1. Le robot de nettoyage des chambres

La configuration est singulière : le service s'exécute le plus souvent en l'absence du client, sans aucune interaction, et il n'est jugé que sur son résultat. Trois conséquences théoriques en découlent, et chacune est un argument de spécification.

D'abord, l'attente d'effort et la motivation hédonique perdent leur objet. Le client n'utilise pas le robot et n'assiste pas au service : il n'y a ni interface à manipuler, ni spectacle, ni amusement. Or le plaisir anticipé était le chemin le plus fort du modèle 2025-2026. Son retrait doit être assumé et exposé comme un résultat théorique, non subi comme un oubli.

Ensuite, toute la branche relationnelle du modèle d'acceptation des robots de service de Wirtz et al. (2018) devient inopérante, faute d'interaction. Et surtout, Choi et al. (2020) établissent, dans une expérience comparant service humain, service robotisé et service conjoint, que le personnel humain est perçu comme supérieur au robot sur la qualité d'interaction et sur l'environnement physique, mais qu'aucune différence significative n'apparaît sur la qualité de résultat. Dans un service jugé uniquement sur son résultat, le déficit du robot disparaît donc : c'est le meilleur appui théorique disponible pour ce sujet.

Enfin, les résultats de Hoang et Tran (2022) fournissent le squelette causal. La compétence perçue du nettoyeur est plus faible pour un robot que pour un humain (moyennes de 5,39 contre 5,88, $p = 0,01$), et cette moindre compétence dilue l'évaluation de la propreté du lieu puis réduit l'intention de visite. Mais deux modérateurs inversent ce désavantage : lorsque la tâche est répugnante, les consommateurs jugent le robot plus compétent que l'humain, par un mécanisme de préoccupation empathique ; et lorsque le nettoyage est perturbateur, le robot est préféré parce qu'il occasionne une perturbation sociale moindre. La salle de bains d'une chambre d'hôtel est le prototype de la tâche répugnante, et le passage d'une machine y est socialement moins intrusif que celui d'une personne — pas de salutation à échanger, pas de pourboire, pas de gêne.

Deux mécanismes complémentaires méritent d'être retenus. Shin et Kang (2020) montrent que le risque sanitaire perçu médie la relation entre l'interaction attendue et l'intention de réservation, et que la propreté perçue modère cet effet ; le robot, en supprimant l'interaction, réduit donc le risque sanitaire perçu. Magnini et Zehrer (2021) signalent à l'inverse que la présence visible du personnel de nettoyage est elle-même un indice subconscient de propreté : robotiser supprime ce signal heuristique. La tension entre ces deux mécanismes est testable.

Enfin, quatre construits propres à cette configuration n'existent nulle part dans la littérature vérifiée et constituent l'apport possible du groupe : l'intrusion territoriale perçue, le risque perçu pour les effets personnels, la préoccupation de surveillance et de cartographie de l'espace intime — un robot de nettoyage embarque caméras, LiDAR et cartographie tridimensionnelle —, et la préférence positive pour l'absence de contact humain, qui renverse la logique habituelle du besoin d'interaction en faisant de l'absence de contact un bénéfice et non un coût.

2.2. Le robot serveur en restauration commerciale

C'est la configuration la mieux documentée des quatre, et celle où la littérature est la plus contradictoire sur l'anthropomorphisme. Molinillo et al. (2022), sur 645 clients potentiels espagnols, trouvent que les perceptions hédoniques sont le premier déterminant de l'attitude ($\beta = 0,293$) et le premier réducteur d'objections après l'anthropomorphisme ($\beta = -0,203$). Mais Blut et al. (2021), dans une méta-analyse portant sur 11 053 individus et 108 échantillons indépendants, concluent à l'absence de consensus : l'effet de l'anthropomorphisme dépend du type de robot et du type de service. Xie et Lei (2022) établissent un effet non linéaire, en U inversé, de l'anthropomorphisme sur les préoccupations de vie privée. Kim et al. (2019) montrent qu'il accroît la chaleur perçue mais dégrade l'attitude par l'étrangeté. Mende et al. (2019) montrent qu'il provoque un inconfort et une menace pour l'identité humaine qui déclenchent des comportements de consommation compensatoire.

La leçon de spécification est claire : l'anthropomorphisme ne doit jamais être modélisé comme un antécédent exogène linéaire et positif, ce qui était précisément la spécification du modèle 2025-2026.

Trois construits sont, eux, propres au service en salle et bien établis. La chaleur perçue, dimension du modèle du contenu des stéréotypes, nourrit la valeur relationnelle attendue (Belanche et al., 2021) — avec un revers documenté par Choi et al. (2021) : une chaleur attendue plus élevée chez l'humanoïde produit davantage d'insatisfaction après un échec de processus. La sécurité physique perçue, mesurée par Jang et Lee (2020) et reprise par Molinillo et al. (2022) avec un β de 0,178 sur l'attitude, n'a d'équivalent dans aucune des autres configurations : le robot circule avec des plats chauds parmi les convives. L'authenticité, enfin, est le construit le plus sensible pour un terrain français : Song et al. (2022) montrent que le recours au robot aux niveaux cœur et facilitant du produit dégrade l'authenticité de service et de marque, et Cheng et al. (2026) établissent une dissociation décisive — le serveur humain nourrit l'authenticité morale, le robot nourrit l'authenticité du plat. Confondre les deux revient à annuler deux effets de signe opposé.

Un dernier résultat structure le terrain : Hu et Min (2025) montrent, sur deux expériences, que les restaurants de cuisine locale obtiennent de meilleures notes d'authenticité, de qualité, de congruence et d'intention de fréquentation avec un service intégralement humain, tandis que les établissements de restauration rapide et à thème futuriste obtiennent une congruence comparable quelle que soit la configuration. La légitimité du robot serveur dépend donc du positionnement de l'établissement, ce qui en fait un modérateur manipulable plutôt qu'une constante.

2.3. Le service robotisé en restauration collective

Le terrain étant vierge, la question devient : que faut-il emprunter, et à quoi ? La réponse tient en une observation. Les modèles dominants de l'acceptation technologique — TAM, UTAUT2, AIDUA, modèle d'acceptation des robots de service — présupposent tous un consommateur libre de ses choix dans un marché concurrentiel. Le convive d'une cantine d'entreprise, d'un restaurant universitaire, d'un hôpital ou d'un établissement pénitentiaire n'adopte pas : il subit. Ces modèles ne sont donc pas seulement incomplets dans ce contexte, ils y sont mal spécifiés.

Le gisement pertinent est la littérature sur les technologies en libre-service en usage contraint, que la recherche sur les robots ignore presque entièrement. Lee et Kim (2026) en fournissent la transposition la plus directe : sur 389 personnes de 65 ans et plus, placées devant des scénarios de borne unique sans alternative humaine, ils montrent que la coercition perçue — définie comme la conscience subjective d'un choix structurellement

contraint — augmente la perte d'autonomie perçue et l'affect négatif, qui prédisent tous deux l'intention d'évitement. Le résultat le plus instructif est que l'affect négatif domine l'évaluation cognitive : sous contrainte, l'émotion l'emporte sur le calcul. En amont, Reinders et al. (2008), Liu (2012), Reinders et al. (2015) et Feng et al. (2018) constituent le socle théorique de l'usage forcé et de la réactance.

Trois conséquences pratiques en découlent. D'abord, la variable dépendante doit changer : en contexte captif, l'intention d'utiliser a peu de variance, et c'est l'intention d'évitement — le report vers le sandwich, la gamelle ou l'extérieur — qui est comportementalement interprétable. Ensuite, la valeur perçue devient difficilement mesurable : en restauration subventionnée ou gratuite, le sacrifice monétaire tend vers zéro et le construit s'effondre sur l'utilité perçue. Deux hypothèses concurrentes s'ouvrent d'ailleurs à ce sujet, sans que la littérature permette de trancher : la gratuité abaisse-t-elle les exigences, ou fait-elle du repas un avantage social acquis dont la dégradation par la robotisation est lue comme une atteinte au contrat implicite ? Enfin, la captivité n'est pas binaire mais graduée — cantine d'entreprise en centre-ville, restaurant universitaire isolé, hôpital, établissement médico-social, détention — ce qui offre un plan quasi expérimental naturel que la restauration commerciale ne peut pas produire.

Deux construits spécifiques méritent d'être développés. La confiance institutionnelle d'abord : dans tout le corpus vérifié, la confiance porte sur le robot, alors qu'en restauration collective l'usager n'a pas choisi le prestataire. L'objet de confiance se déplace vers le décideur — l'employeur, l'université, l'hôpital, la collectivité — et la question devient : a-t-il robotisé pour mon confort ou pour réduire sa masse salariale ? La typologie de McKnight et al. (2002), qui sépare la confiance institutionnelle des croyances de confiance, fournit l'ancrage. La menace perçue sur l'emploi ensuite : aucune étude vérifiée ne la mesure comme antécédent du côté du client, et Molinillo et al. (2022) trouvent même un effet non significatif des objections à l'usage sur l'acceptation ($\beta = -0,061$, $p = 0,062$). Mais en restauration collective le convive connaît personnellement l'agent menacé, le croise hors du service, partage parfois son employeur. L'objection éthique cesse d'être abstraite, et l'hypothèse que ce construit devienne significatif lorsque la personne remplacée est identifiée et socialement proche est, à mon sens, le résultat le plus publiable que le projet puisse produire.

2.4. Le robot cuisinier qui cuisine devant le convive

Cette configuration croise trois littératures, et c'est ce croisement qui la rend intéressante.

La première est celle de la psychologie du consommateur alimentaire. Pancer et al. (2025), sur quatre expériences, établissent que l'automatisation de la préparation réduit l'amour perçu dans l'aliment, ce qui dégrade le goût anticipé, l'évaluation globale et la disposition à payer. Surtout, ils montrent que le dégoût observé est d'origine morale — une opposition au remplacement de l'emploi — et non hygiénique, et qu'une communication centrée sur le bénéfice client le normalise. C'est un résultat contre-intuitif dont la portée méthodologique est grande : tester une hypothèse de contagion ou de souillure sans distinguer la cause attribuée revient à mesurer l'indignation morale en croyant mesurer la répulsion sensorielle.

La deuxième est celle de la néophobie. Ross et al. (2022), sur 1 045 répondants irlandais, trouvent que la néophobie technologique alimentaire, mesurée par l'échelle de Cox et Evans (2008), exerce l'effet total négatif le plus fort du modèle ($\beta = -0,369$, dont $-0,167$ en direct et $-0,202$ en indirect), et que la confiance en la science en est l'antidote le plus puissant ($\beta = -0,445$). C'est un levier de communication directement actionnable.

La troisième est celle de la transparence opérationnelle, et c'est là que se trouve l'angle mort le plus intéressant. Buell, Kim et Tsay (2017) établissent, sur deux expériences de terrain et deux expériences de laboratoire menées en restauration, que rendre le processus visible augmente la qualité perçue de 22,2 % et réduit le temps de traitement perçu de 19,2 %, par un mécanisme d'effort perçu : le client qui observe le travail l'apprécie davantage. Or aucune étude vérifiée n'a transposé ce mécanisme à un agent robotique. Et le mécanisme suppose un agent capable d'effort méritoire, condition qu'un robot ne remplit pas. L'hypothèse que la transparence opérationnelle ne produise aucun effet, voire un effet négatif, lorsque le cuisinier est une machine est à la fois la plus originale et la plus défendable du dossier.

Un dernier résultat doit être connu, parce qu'il contredit l'argument marketing le plus répandu. Byrd et al. (2024) montrent expérimentalement que les consommateurs préfèrent un serveur humain masqué à un robot propre et désinfecté : le robot n'est pas le dispositif de réassurance sanitaire que l'on suppose.

Deux tensions restent à trancher, et le projet en a les moyens. La première porte sur l'anthropomorphisme : Xiao et Zhao (2022) trouvent qu'il améliore la prédiction de qualité du plat, Lee et Kwak (2025) qu'il réduit l'appréciation attendue et la volonté d'essayer. L'hypothèse de réconciliation la plus plausible est que l'anthropomorphisme aide tant qu'il n'y a pas contact avec l'aliment, et nuit dès qu'il y en a — des mains humanoïdes plongées dans la nourriture activant une contagion symbolique au sens de Rozin et al. (1992). La seconde porte sur la confiance dans le respect des régimes alimentaires et des allergènes : la littérature consommateur est ici inexistante, alors que la question est centrale en restauration collective.

2.5. Ce que fonde la nature de la tâche : le socle théorique de l'adéquation robot-tâche

Les relecteurs demandaient sur quoi reposait l'adéquation perçue robot-tâche, second construit du socle. La réponse est une littérature abondante et convergente, qui manquait à la version 1. Huang et Rust (2018) et Wirtz et al. (2018) établissent que le remplacement par l'IA s'opère au niveau de la tâche, des intelligences mécaniques vers les intelligences empathiques. Castelo, Bos et Lehmann (2019), sur quatre études de laboratoire et deux études de terrain totalisant plus de 56 000 observations, montrent que les algorithmes sont moins utilisés pour les tâches subjectives et que l'objectivité perçue de la tâche est malléable. Longoni et Cian (2022) documentent l'effet « word-of-machine » : l'IA est jugée plus compétente en contexte utilitaire, moins en contexte hédonique. Liu, Wang et Wu (2026), sur 7 428 avis Ctrip et trois expériences, montrent qu'un service fonctionnel est mieux évalué lorsqu'il est rendu par un robot, et un service émotionnel lorsqu'il est rendu par un humain. Yang, Wang, Song et Ma (2024) fondent exactement un « robot-environment fit » par type de tâche sur la théorie person-environment fit. Enfin, la méta-analyse de Peng et al. (2026) — 46 expériences, 27 612 participants — fixe l'ordre de grandeur : une préférence pour l'humain faible en moyenne ($d = -0,128$) et entièrement conditionnelle au contexte, plus forte dans le luxe et les contextes domestiques, inversée dans les situations embarrassantes.

Ces résultats justifient un construit distinct de la compétence perçue : la compétence répond à la question « le robot en est-il capable ? », l'adéquation à la question « est-ce à lui de le faire ? ». Un robot compétent peut être inadéquat pour une tâche empathique, et un robot peu compétent adéquat pour une tâche embarrassante. Choi et al. (2020) en donnent la preuve empirique la plus nette : la qualité de résultat est jugée identique entre humain et robot, la qualité d'interaction ne l'est pas.

2.6. Congruence robot-marque et positionnement

Second angle mort comblé. Choi, Liu et Choi (2022) montrent que les robots à haut contact provoquent des réactions négatives pour une marque de personnalité sincère, non pour une marque excitante, par un fit perçu faible. Guan et al. (2025) répliquent en hôtellerie : la congruence entre le robot et la personnalité de marque nourrit l'amour de marque, la confiance médiatisant. Horng et al. (2025, 2026) introduisent deux congruences distinctes — luxury-technology fit et task-technology fit — qui médiatisent vers l'intention, à Taïwan comme aux États-Unis. Ma et al. (2025) documentent un « désavantage du robot » dans le luxe, par désir de supériorité : le robot ne procure pas le sentiment d'être servi en tant que client de statut. Chan et Tung (2019) trouvaient déjà un effet comportemental du robot positif en hôtellerie économique et de milieu de gamme, nul dans le luxe. Le positionnement de l'établissement est donc un modérateur robuste, à manipuler dans le scénario du groupe G2.

2.7. Ce qu'apportent les terrains asiatiques

Soixante références dont les terrains sont en Corée, en Chine, au Japon, à Taïwan, en Malaisie ou en Inde ont été vérifiées. Elles apportent ce que les terrains occidentaux ont rarement : des clients réellement servis par un robot. La série coréenne de Hwang, Kim et Choe compare systématiquement environ 300 clients servis par un

robot à 300 clients servis par un humain ; Zhang et al. (2026) interrogent 433 clients de restaurants robotisés chinois ; Gong, Guan et Huan (2025) interrogent 269 clients sur site d'un restaurant à robot chef de Guangzhou ; Sthapit et al. (2024) étudient 429 séjours réels au FlyZoo Hotel d'Alibaba ; Fan, Gao et Han (2022) mènent une expérimentation de terrain de six mois dans 35 hôtels. S'y ajoutent des analyses d'avis massives : 4 107 avis Qunar sur 68 hôtels (Chang et al., 2022), Dianping pour Haidilao (Ma et al., 2023), 1 569 commentaires sur Henn-na Hotel (Yu, 2020).

Sur usage réel, la hiérarchie des antécédents devient très utilitaire : l'efficacité perçue représente 55 % des mentions dans les avis Qunar, la qualité de la nourriture et l'innovation technique expliquent la recommandation tandis que la qualité de service du robot ne l'explique pas (Li, Yuan & Zhao, 2025), et l'influence sociale disparaît ou n'agit qu'indirectement (Chang et al., 2022 ; Gong et al., 2026) alors qu'elle est centrale dans les modèles testés par scénario. Sur l'anthropomorphisme, les commentaires japonais sur Henn-na rejettent massivement les humanoïdes (73,5 % négatifs), l'acceptation est maximale pour un anthropomorphisme moyen (Jia, Chung & Hwang, 2021), et le robot d'apparence animale maximise émotions et chaleur (Cheng & Hwang, 2025) : le format « chat-robot » BellaBot qui domine en Asie n'est pas un hasard. Pour le robot en chambre, le contrôle perçu renforce l'acceptation (Song et al., 2024), les craintes de sécurité nocturne, de caméras et de piratage apparaissent spontanément (Yu, 2020), et la co-production est un antécédent de valeur (Tian & Liu, 2024). Sur le robot cuisinier, la nourriture prime, la compétence du robot chef conditionne l'évaluation alimentaire, et l'attitude anti-robot atténue les effets positifs (Gong et al., 2025) ; le besoin d'interaction humaine reste significatif même sur usage réel (Sung & Jeon, 2020). En revanche, aucune étude client asiatique vérifiable ne porte sur la restauration collective : le vide est le même qu'en Occident.

Trois précautions de transposition. Les répondants asiatiques ont une familiarité élevée avec les robots : les effets de nouveauté seront plus forts en France, et l'influence sociale peut y peser davantage, comme dans les études par scénario. Beaucoup d'études chinoises sont des scénarios et les échantillons coréens des panels en ligne : vérifier l'usage réel avant de transposer un coefficient. Enfin, une seule étude vérifiable teste les valeurs culturelles comme modérateurs (Chi, Chi, Gursoy & Nunkoo, 2023, États-Unis contre Chine) : distance hiérarchique et évitement de l'incertitude sont à réestimer sur un terrain français, non à importer.

2.8. La littérature francophone

Une recherche systématique a été menée par requêtes Crossref filtrées par revue (Décisions Marketing, Recherche et Applications en Marketing, Revue Française de Gestion, Management & Avenir, Téoros, Mondes du tourisme, Systèmes d'Information et Management, Gestion 2000, entre autres), sur HAL et sur theses.fr, de 2012 à 2026. Le constat est net : le corpus francophone à comité de lecture sur l'acceptation des robots de service par le client en hôtellerie-restauration est quasi inexistant. Trois articles francophones portent sur les robots, aucun en hôtellerie-restauration : Goudey et Bonnin (2016) montrent qu'un anthropomorphisme partiel aide les familiers du smartphone et nuit aux autres ; Passebois-Ducros et Flacandji (2022) étudient un robot à Lascaux IV ; Yang, Garnier et Djelassi (2025) montrent qu'un robot qui imite l'approche d'un vendeur accroît l'intrusivité perçue — résultat directement transposable au robot serveur qui s'approche de la table. Un seul article francophone croise IA, tourisme et client (Garcia et al., 2021, sur un chatbot touristique).

Cette absence est compensée de trois manières. D'abord par les laboratoires français publiant en anglais : Meyer-Waarden et Cloarec (2026, TSM Toulouse) sur l'éthicité perçue des robots sociaux en restaurant contre à domicile, le groupe Hasan et de Kervenoael (Tourism Management, 2020 ; JHTT, 2024), et El Halabi et Trendel (2024, Journal of Service Research), qui montrent que nommer un robot humanoïde réduit l'étrangeté perçue — un levier simple pour le malaise perçu du socle. Ensuite par les cadres francophones transposables sur les technologies en libre-service (Feenstra & Glérant-Glikson, 2017 ; Lao & Vlad, 2018 ; Ba & Vignon, 2013). Enfin par l'ancrage local en sociologie de l'alimentation : le CERTOP (Poulain ; Laporte & Poulain, 2014, sur la synchronisation sociale du repas d'entreprise) et la thèse en cours de David Rodriguez, dont la netnographie de réactions aux robots dans des restaurants français documente précisément l'effet de

contamination du plat, le spectre de la suppression d'emplois et la politisation du repas — sur la plateforme expérimentale Ovalie de l'université. C'est un contact évident pour le projet.

3. Existe-t-il un travail fondateur ?

La question mérite une réponse mesurée plutôt qu'une impression. Deux indicateurs ont été calculés. Le premier est la notoriété absolue, à savoir le nombre total de citations rapporté par l'API Semantic Scholar en septembre 2026, pour trente-huit candidats. Le second est le taux de co-citation à l'intérieur du domaine : pour trente articles empiriques récents du corpus, les références déposées auprès de Crossref ont été téléchargées — 2 558 références au total, dont 2 015 portant un DOI exploitable — et l'on a compté dans combien de ces trente bibliographies chaque candidat apparaît.

La distinction entre les deux indicateurs n'est pas cosmétique. La notoriété absolue favorise les références matricielles de la recherche en systèmes d'information, qui ne disent rien de spécifique sur les robots de service : Davis (1989) totalise 71 107 citations et Venkatesh et al. (2003) 46 187, mais ces travaux situent le champ sans le différencier. Le taux de co-citation identifie au contraire le point de passage obligé de la littérature spécifique.

La réponse est alors nette. Le travail fondateur du domaine est l'article de Wirtz, Patterson, Kunz, Gruber, Lu, Paluch et Martins publié en 2018 dans le *Journal of Service Management* sous le titre « Brave new world: Service robots in the frontline ». Il est cité dans quinze des trente bibliographies examinées, soit la moitié du corpus, et cette présence traverse les quatre configurations — nettoyage, service en salle, cuisine, méta-analyses. Avec 1 751 citations, il n'est pourtant pas le plus cité dans l'absolu : sa force est d'être incontournable. Il fournit le vocabulaire du champ, la typologie des tâches — cognitives-analytiques, cognitives-intuitives, émotionnelles-sociales — et le modèle d'acceptation à trois branches, fonctionnelle, socio-émotionnelle et relationnelle, que la plupart des auteurs reprennent ou amendent.

Deux piliers complètent ce socle. L'échelle de volonté d'intégration des robots de service de Lu, Cai et Gursoy (2019), présente dans douze bibliographies sur trente, est l'instrument de mesure standard de la variable dépendante du domaine. Et le modèle d'acceptation des dispositifs d'intelligence artificielle de Gursoy, Chi, Lu et Nunkoo (2019), présent dans huit bibliographies sur trente mais fort de 1 351 citations, est le modèle théorique le plus répliqué, y compris hors hôtellerie — Li et al. (2024) l'appliquent aux robots médicaux. C'est aussi celui dont le projet 2025-2026 était issu, ce qui en fait le point de comparaison naturel du travail 2026-2027.

Un dernier chiffre mérite d'être signalé, parce qu'il concerne directement le problème que cette note cherche à résoudre. **Le critère HTMT de Henseler, Ringle et Sarstedt (2015) n'est cité que dans quatre des trente bibliographies examinées, et le critère de Fornell et Larcker (1981) dans dix.** Autrement dit, deux tiers des études publiées dans ce domaine ne documentent leur validité discriminante par aucun de ces deux critères. Documenter sérieusement la sienne serait déjà, en soi, une contribution.

Une réserve technique s'impose sur ce second indicateur : Crossref ne recense que les références effectivement déposées par l'éditeur avec un DOI. Une référence citée sans DOI n'apparaît pas dans le décompte. Les taux ci-dessus sont donc des planchers, et Gursoy et al. (2019) est en particulier sous-estimé.

4. Pourquoi la matrice HTMT a posé problème, et comment ne pas recommencer

Le modèle 2025-2026 a dû être nettoyé en raison de recouvrements entre construits. Ce n'est pas un accident local : c'est une caractéristique structurelle de cette littérature, et il est possible de le montrer sur les travaux publiés eux-mêmes.

Trois exemples suffisent. Chez Molinillo et al. (2022), l'intention de recommander prédit la volonté d'accepter avec un bêta de 0,781 et un R^2 de 0,736 : un coefficient de cette magnitude entre deux intentions est un signal de recouvrement, pas un résultat substantiel. Chez Marghany et al. (2025), l'attitude prédit l'intention à 0,565 pendant que l'effet direct de l'attente de performance devient non significatif et négatif ($\beta = -0,036$) : configuration typique d'un recouvrement entre attitude et intention autant que d'une médiation véritable. Et dans la méta-analyse de Begum et al. (2025), la présence sociale perçue est qualifiée de plus mauvais prédicteur alors que sa corrélation avec l'intention atteint 0,594, les auteurs suspectant une relation confondue par des variables externes — ce qui est exactement la signature d'un problème de validité discriminante non diagnostiqué.

L'inventaire regroupe les construits recensés en treize familles conceptuelles, c'est-à-dire treize ensembles de termes qui mesurent la même chose sous des noms différents. La famille de la performance en compte douze : utilité perçue, attente de performance, efficacité, fonctionnalité, consistance, assurance de service, qualité de résultat, intelligence perçue, bénéfices perçus, valeur fonctionnelle, compétence perçue et facette « performance » de la confiance. La famille de l'humanité du robot en compte neuf, dont l'anthropomorphisme et la présence sociale automatisée, qui partagent littéralement les mêmes items chez la plupart des auteurs.

De ce diagnostic découlent dix règles de spécification, détaillées dans l'onglet correspondant du classeur. Les quatre plus importantes sont les suivantes.

- Un seul construit mesuré par famille conceptuelle. Retenir la compétence perçue signifie abandonner explicitement les onze autres membres de sa famille, et traiter l'assurance de service ou l'intelligence perçue comme des indicateurs de ce construit et non comme des construits distincts.
- Jamais attitude et intention dans le même modèle. L'attitude spécifique envers le robot est supprimée et les antécédents pointent directement vers l'intention ; l'attitude générale envers les robots est conservée, mais comme variable de contrôle exogène, parce qu'Ivanov et Webster (2019) établissent que c'est le meilleur prédicteur unique de la volonté d'utiliser un robot et que ne pas la contrôler gonflerait artificiellement tous les construits spécifiques.
- Manipuler plutôt que mesurer ce qui peut l'être. L'apparence du robot est manipulée par deux versions de la vidéo stimulus, mécanoïde et anthropomorphe, au lieu d'être mesurée. Cette seule décision retire du modèle de mesure toute la famille de l'humanité du robot, qui est la première source de recouvrement du domaine — et qui était l'antécédent exogène principal du modèle 2025-2026.
- Séparer strictement les traits et les états. Un trait — anxiété technologique, néophobie, aversion aux germes — est un modérateur ou une covariable ; un état — malaise, propreté perçue, risque perçu — est un antécédent ou un médiateur. La confusion des deux est une source classique de HTMT élevé, et elle est fréquente dans ce corpus.

Les six autres règles portent sur l'unicité de la variable dépendante, l'ancrage explicite de chaque batterie d'items sur son objet, le plafonnement à six ou huit construits latents, le pré-test cognitif avant collecte plutôt que la purification après coup, la publication du HTMT avec intervalles de confiance bootstrap et un seuil conservateur de 0,85, et la préférence pour les chaînes causales sur les faisceaux de prédicteurs parallèles.

5. Le modèle proposé

Le modèle repose sur un socle commun mesuré à l'identique dans les trois groupes, ce qui autorise une comparaison multi-groupes, et sur deux à six extensions propres à chaque configuration. Chaque construit du socle est l'unique représentant de sa famille conceptuelle.

5.1. Le socle commun

Construit	Rôle	Pourquoi celui-là
Compétence	Antécédent — route instrumentale	Unique représentant de la famille de la performance ;

Construit	Rôle	Pourquoi celui-là
perçue du robot		en C1, il prend la forme validée de la compétence perçue du nettoyeur (Hoang & Tran, 2022).
Adéquation perçue robot-tâche	Antécédent — route normative	Construit le plus discriminant et le plus original. Il demande « est-ce à lui de le faire ? » là où la compétence demande « en est-il capable ? ».
Malaise perçu	Médiateur affectif unique	Remplace les émotions génériques de l'AIDUA. État réactif au stimulus, distinct de l'anxiété technologique qui est un trait.
Besoin d'interaction humaine	Antécédent — route sociale	Le seul construit dont le signe attendu diffère selon le contexte : négatif en restauration, positif en hébergement où l'absence de contact est un bénéfice.
Plaisir anticipé	Antécédent — route hédonique (G2 et G3)	Conservé du modèle 2025-2026, mais délibérément absent du modèle G1 : un service invisible n'offre ni spectacle ni amusement.
Intention d'accepter et intention d'éviter	Deux variables dépendantes distinctes (v2)	Acceptation (Gursoy et al., 2019, 3 items) et évitement (Lee & Kim, 2026, 4 items) mesurés dans les trois groupes ; distinction vérifiée par AFC, HTMT < 0,85.
Apparence du robot	Variable manipulée	Mécanoïde contre anthropomorphe, dans la vidéo stimulus. Supprime du modèle de mesure la famille conceptuelle la plus problématique.
Attitude générale envers les robots	Contrôle exogène	Covariable obligatoire : sans elle, les construits spécifiques capteraient une simple technophilie.

5.2. Les extensions par groupe

Le groupe hébergement ajoute la propreté anticipée de la chambre en médiateur, l'intrusion territoriale perçue, la préoccupation de surveillance et le risque pour les effets personnels, avec l'aversion aux germes en modérateur et le caractère répugnant de la tâche en manipulation. Le groupe restauration commerciale ajoute la chaleur perçue, la sécurité physique perçue, l'authenticité morale du service et l'influence sociale, avec la congruence au positionnement de l'établissement en manipulation. Le groupe restauration collective, allégé en version 2, ajoute la transparence opérationnelle en manipulation, l'amour perçu dans l'aliment en médiateur, la menace sur l'emploi du personnel connu (dimensions cognitive et morale) orientée vers l'évitement, et la confiance dans le respect des régimes et des allergènes, avec la captivité (attractivité des alternatives) et la néophobie technologique alimentaire en modérateurs. Dans les trois groupes, l'intention d'accepter et l'intention d'éviter sont mesurées séparément.

Modèle intégré de l'acceptation d'un robot de service en hôtellerie-restauration Socle commun aux trois contextes et extensions contextuelles - DRA L3 MHR / L3 MRC 2026-2027 - version 2 (24/09/2026)

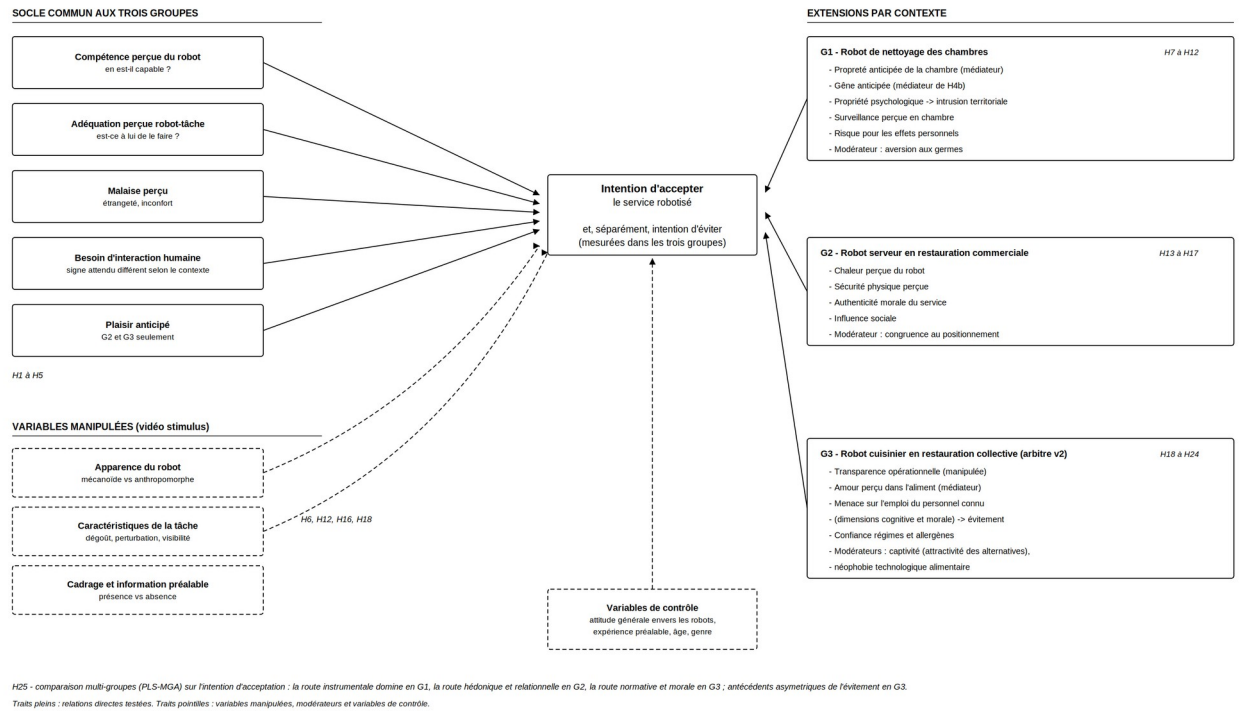


Figure 1 — Modèle intégré : socle commun, variables manipulées et extensions contextuelles.

Les trois modèles de groupe

Un socle commun identique, des extensions propres à chaque configuration de service - version 2 (24/09/2026)

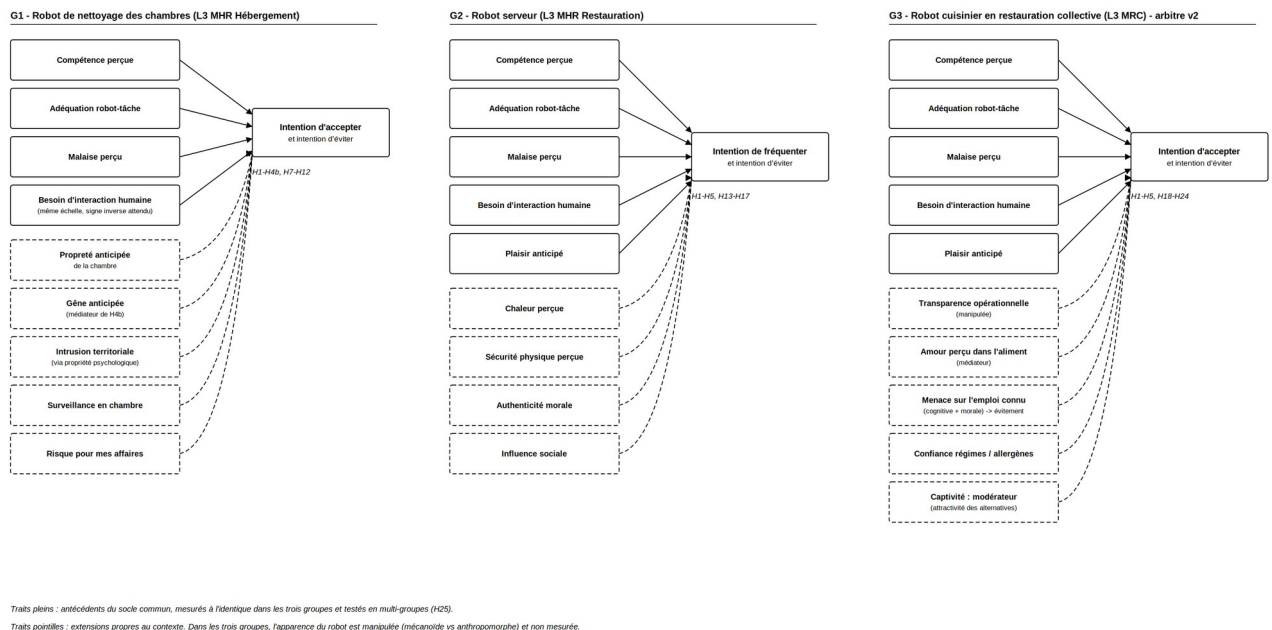


Figure 2 — Les trois modèles de groupe.

6. Les hypothèses

Vingt-cinq hypothèses sont proposées (H20 ayant été fusionnée dans H23 et une hypothèse de mesure ajoutée) : le socle commun, les extensions par groupe et une hypothèse transversale de comparaison. Le tableau ci-dessous les reprend sous forme abrégée ; l'énoncé complet et l'ancrage bibliographique de chacune figurent dans l'onglet « Modèle et hypothèses » du classeur.

Code	Énoncé abrégé	Ancrage principal
H1	Compétence perçue → intention (+)	Hoang & Tran (2022) ; Begum et al. (2025)
H2 (v2)	Adéquation robot-tâche → intention (+), au-delà de la compétence ; plus élevée pour les tâches objectives, plus faible pour les tâches émotionnelles ou symboliques	Huang & Rust (2018) ; Castelo et al. (2019) ; Longoni & Cian (2022) ; Liu et al. (2026) ; Peng et al. (2026)
H3	Malaise perçu → intention (-)	Mende et al. (2019) ; Baek et al. (2025)
H4a	Besoin d'interaction humaine → intention (-) en G2 et G3	Dabholkar (1996) ; Shin & Jeong (2020)
H4b (v2)	Besoin d'interaction humaine → intention (+) en G1, médiation par la gêne anticipée	Pitardi et al. (2022, 2024) ; Holthöwer & van Doorn (2023) — hypothèse forte
H5	Plaisir anticipé → intention (+) en G2 et G3 ; sans objet en G1	Gursoy et al. (2019) ; Molinillo et al. (2022)
H6	L'effet de l'apparence dépend du contact avec l'aliment	Xiao & Zhao (2022) contre Lee & Kwak (2025) — réconciliation
H7	Compétence → propreté anticipée → intention (médiation)	Hoang & Tran (2022) ; Vos et al. (2019)
H8 (v2)	Intrusion territoriale → intention (-), d'autant plus forte que la propriété psychologique de la chambre est élevée	Construit nouveau ; Jiang, Balaji & Lyu (2026)
H9	Surveillance perçue en chambre → intention (-)	Jia et al. (2024) ; Song B. et al. (2024)
H10	Risque pour les effets personnels → intention (-)	Construit nouveau ; Park (2020)
H11	L'aversion aux germes modère compétence → intention	Duncan et al. (2009) ; Shin & Kang (2020)
H12	Le caractère répugnant de la tâche inverse le désavantage du robot	Hoang & Tran (2022) — réplification
H13	Chaleur perçue → intention (+), distincte de la compétence	Belanche et al. (2021) ; Choi et al. (2021)
H14	Sécurité physique perçue → intention (+)	Jang & Lee (2020) ; Molinillo et al. (2022)
H15	Authenticité morale → intention (+) ; le robot la dégrade	Song X. et al. (2022) ; Cheng et al. (2026)
H16	La congruence au positionnement modère l'adéquation perçue	Hu & Min (2025) ; Kao & Huang (2023)
H17	Influence sociale → intention (+), propre à la consommation publique	Chi et al. (2022) ; Zemke et al. (2025)
H18	La transparence opérationnelle n'opère pas pour un agent robotique	Buell, Kim & Tsay (2017) — jamais testé, hypothèse centrale

Code	Énoncé abrégé	Ancrage principal
H19	L'amour perçu dans l'aliment médie l'effet de l'automatisation	Pancer et al. (2025) ; Fuchs et al. (2015)
H21	Confiance régimes et allergènes → intention (+)	Construit nouveau ; Huang et al. (2023)
H22 (v2)	La captivité (faible attractivité des alternatives) modère l'effet de la menace sur l'emploi et du malaise sur l'évitement	Lee & Kim (2022) ; Lee & Kim (2026) ; Kim, Ng & Kim (2009)
H23 (v2)	Menace sur l'emploi du personnel connu (cognitive + morale) → évitement (+), amplifiée par la proximité relationnelle ; n.s. en commercial	Etemad-Sajadi et al. (2022) ; Granulo et al. (2019, 2025) ; Pancer et al. (2025) — inédite
H24	La néophobie technologique alimentaire modère l'ensemble	Cox & Evans (2008) ; Ross et al. (2022)
H25 (v2)	Sur l'acceptation, la structure des déterminants diffère entre configurations (PLS-MGA) ; sur l'évitement, antécédents asymétriques en G3	Chi et al. (2022) ; Hu & Min (2025) ; Claudy et al. (2015)
VD (v2)	Acceptation et évitement sont deux construits distincts (HTMT < 0,85), moins corrélés en contexte captif	Gursoy et al. (2019) ; Cenfetelli (2004) ; Liang & Xue (2009)

7. Protocole recommandé

La collecte reste séparée par groupe de contexte, conformément à l'arbitrage déjà retenu : chaque groupe interroge sa propre population sur son propre scénario, sans affectation aléatoire des répondants aux trois configurations. À raison d'une vingtaine de répondants par étudiant et de données mises en commun, chaque groupe devrait réunir entre 250 et 300 observations, ce qui est confortable pour un modèle de sept construits latents estimé en moindres carrés partiels.

Deux vidéos stimulus par groupe sont nécessaires, l'une présentant un robot d'aspect délibérément mécanique, l'autre un robot d'aspect anthropomorphe, le reste du scénario étant strictement identique. Ce plan à deux conditions suffit à tester H6 et à retirer du modèle de mesure la famille conceptuelle la plus encombrante.

Trois précautions de conception méritent d'être signalées aux étudiants dès la première séance de rédaction du questionnaire. La première concerne la traduction : les items de coercition perçue portent sur l'absence structurelle d'alternative et non sur la difficulté d'usage, or en français « je suis obligé de » glisse facilement vers « j'ai du mal à ». La deuxième concerne l'ancrage : chaque batterie d'items doit nommer explicitement son objet — ce robot, cet établissement, mon employeur, mes affaires, mon régime alimentaire — faute de quoi deux construits distincts fusionneront dans l'esprit du répondant. La troisième concerne la distinction entre la préoccupation égocentrée et la préoccupation altruiste : « je perds quelque chose » et « la personne qui me sert perd son poste » sont deux construits différents que les répondants confondront si les items ne les séparent pas explicitement.

Enfin, un pré-test cognitif sur une dizaine de répondants, mené avant la collecte, suffit généralement à détecter les items qui fusionnent. C'est un investissement d'une demi-journée qui évite le nettoyage a posteriori — lequel, rappelons-le, revient à redéfinir le construit sans le dire.

8. Limites et points restant à trancher

Quatre points appellent une décision ou une vérification complémentaire.

- L'échelle SRPC de Jia, Chi et Lu (2024), qui est l'instrument le plus prometteur pour mesurer la préoccupation de surveillance en chambre, n'a pas pu être consultée en texte intégral : ses dimensions et ses items doivent être récupérés avant toute mobilisation.

- L'arbitrage du modèle G3 est désormais tranché (huit antécédents latents, coercition en variable de contexte, dégoût moral fusionné dans la menace sur l'emploi). Il reste à vérifier au pré-test l'unidimensionnalité du construit « menace sur l'emploi du personnel connu » avec ses items cognitifs et moraux ; si elle n'est pas obtenue, la dimension morale sera traitée comme médiateur en chaîne (menace → indignation → évitement) plutôt que comme construit parallèle.
- Deux publications portent exactement sur le contexte du groupe G3 — Min, Martinez et Stringam (2025, IJHM ; 2026, Journal of Foodservice Business Research) sur les robots de livraison en restauration sous contrat — mais leurs résumés n'ont pas pu être consultés. Elles sont à lire en priorité via l'accès institutionnel.
- La saillance de la menace sanitaire, qui inversait la préférence humain-robot au plus fort de la pandémie (Kim et al., 2021), a probablement décliné depuis 2023. La traiter en variable de contrôle situationnelle plutôt qu'en antécédent est prudent — et vérifier empiriquement cette décroissance est en soi une contribution modeste mais propre.
- Les indicateurs bibliométriques de la section 3 (citations Semantic Scholar, co-citations Crossref) sont, comme le note Mistral, des arguments d'aide à la décision et non des résultats reproductibles par les étudiants : ils sont présentés comme tels.
- Les effets du genre sont contradictoires dans la littérature vérifiée : Parvez et al. (2024) et Ayyildiz et al. (2022) trouvent les femmes plus favorables, Binesh et Baloglu (2023) et Ivanov et al. (2018) trouvent l'inverse. Le genre doit donc rester une variable de contrôle et ne pas faire l'objet d'une hypothèse.

Une dernière remarque, qui déborde le cadre méthodologique. La littérature vérifiée est dominée par la Corée du Sud, la Chine, Taïwan, les États-Unis, la Thaïlande, l'Espagne et Singapour. La recherche systématique de la version 2 le confirme : aucune étude quantitative à comité de lecture sur un échantillon français, la seule trace empirique française étant une netnographie en cours au CERTOP. C'est une contribution disponible, et elle ne coûte rien de plus que la collecte déjà prévue.