

Interférences multimodales d'IA générative pour relancer l'idéation après fixation : étude comparative

par

Fabrice FISCHER

THÈSE PAR ARTICLES PRÉSENTÉE À
L'ÉCOLE DE TECHNOLOGIE SUPÉRIEURE
COMME EXIGENCE PARTIELLE À L'OBTENTION
DU DOCTORAT EN GÉNIE
Ph. D.

MONTREAL, LE 10 SEPTEMBRE 2026

ÉCOLE DE TECHNOLOGIE SUPÉRIEURE
UNIVERSITÉ DU QUÉBEC



Fabrice Fischer, 2026



Cette [licence Creative Commons CC BY-NC-ND 4.0](https://creativecommons.org/licenses/by-nc-nd/4.0/) signifie qu'il est permis de diffuser, d'imprimer ou de sauvegarder sur un autre support une partie ou la totalité de cette œuvre à condition de créditer l'auteur, que ces utilisations soient faites à des fins non commerciales et que le contenu de l'œuvre n'ait pas été modifié.

PRÉSENTATION DU JURY

CETTE THÈSE A ÉTÉ ÉVALUÉE

PAR UN JURY COMPOSÉ DE :

M. Mickaël Gardoni, directeur de thèse
Département de génie des systèmes à l'École de technologie supérieure

M. Erik Andrew Poirier, président du jury
Département de génie de la construction à l'École de technologie supérieure

M. Michel Rioux, membre du jury
Département de génie des systèmes à l'École de technologie supérieure

M. Laurent Simon, membre externe indépendant
Département d'entrepreneuriat et innovation à HEC Montréal

ELLE A FAIT L'OBJET D'UNE SOUTENANCE DEVANT JURY ET PUBLIC

LE 4 AOÛT 2026

À L'ÉCOLE DE TECHNOLOGIE SUPÉRIEURE

AVANT-PROPOS

Ce doctorat est né d'une conviction simple : lorsqu'un designer se fixe trop tôt sur ses premières idées, l'exploration s'appauvrit. Aider à relancer cette idéation, même modestement, peut améliorer la qualité des solutions. C'est autour de cette question, à l'intersection de l'innovation, du Design Thinking et de l'intelligence artificielle générative, que s'est construit mon parcours doctoral.

Dans une démarche de Design Thinking, l'idéation occupe une place centrale : c'est le moment où s'ouvrent, ou au contraire se referment, les espaces de solution. Lorsque les designers se fixent trop tôt sur leurs idées initiales, l'exploration se rétrécit et les possibilités d'innovation s'amenuisent.

J'ai abordé l'IA générative non comme un substitut à l'idéation humaine, mais comme un moyen de produire des interférences capables de relancer l'idéation. Cette distinction traverse l'ensemble de ma démarche.

Au cours de ce doctorat, j'ai approfondi ma compréhension du Design Thinking, de l'idéation et des usages de l'IA générative. Cet apprentissage s'est nourri de la littérature académique, de mon enseignement du Design Thinking en ingénierie et en gestion, ainsi que de ma pratique auprès de groupes variés dans différents contextes organisationnels.

Les échanges avec mon jury m'ont conduit à préciser le cadre méthodologique du programme doctoral, en articulant la méthodologie de recherche en science de la conception (*Design Science Research Methodology*, DSRM) et l'étude de cas comparée (*Comparative Case Study*, CCS). Ce cadrage méthodologique structure l'ensemble du programme et éclaire la lecture des trois articles.

Ce programme doctoral s'est structuré en trois temps. Il a d'abord porté sur la génération de suggestions verbales orientées comme interférences en situation de fixation. Il s'est ensuite poursuivi par la comparaison d'une interférence textuelle spécifique au projet, d'une interférence image au contenu sémantique similaire et d'une pause neutre dans le contexte de l'innovation des bibliothèques publiques. Enfin, il a élargi cette logique comparative aux modalités texte, mosaïque d'images, mosaïque vidéo et pause témoin, afin d'examiner si la

forme de l'interférence influence différemment la relance de l'idéation et le déplacement conceptuel.

À travers ce programme, j'ai cherché à identifier des interférences produites par l'IA générative qui soient à la fois faciles à produire, peu coûteuses, adaptées au projet et utiles lorsque l'idéation s'essouffle.

Les résultats ne permettent pas de formuler une conclusion universelle, mais ils indiquent que toutes les interférences n'ont pas le même potentiel de défixation. Cette recherche demeure exploratoire, notamment parce que les outils d'IA générative évoluent rapidement et que le programme repose sur un nombre limité d'interférences, de cas et de configurations expérimentales. Cette recherche ne prétend pas proposer une recette universelle, mais elle ouvre une piste méthodologique et pratique pour soutenir l'idéation en situation de fixation par des interférences générées par l'IA.

Ce doctorat a renforcé mon désir de poursuivre un travail de recherche à l'interface de l'enseignement, de la pratique et de la production de connaissances. Je souhaite continuer à contribuer à ce champ par la recherche, l'encadrement et l'enseignement.

REMERCIEMENTS

Tout d’abord, je remercie Mickaël Gardoni, mon directeur de thèse, pour son accompagnement exigeant, bienveillant et profondément humain. Ses conseils ont contribué de manière décisive à la maturation de ma réflexion de chercheur. Je remercie également les membres de mon jury, Erik Andrew Poirier, Michel Rioux et Laurent Simon, pour leurs recommandations claires et structurantes, qui ont contribué à orienter la suite de ce travail.

Sur un plan plus personnel, je remercie Amandine, Antoine, Bianca, Arnaud et Vincent, qui m’ont accompagné, soutenu et encouragé à persévérer tout au long de ces années de recherche. Je pense aussi à mon père, Hervé, sociologue et philosophe, dont les années d’enseignement à la Sorbonne ont contribué à tracer un horizon intellectuel qui m’a marqué. Je garde également en mémoire l’héritage scientifique de mon grand-père, Édouard Fischer, ainsi que celui d’Édouard Piette, grand-père de mon grand-père.

Interférences multimodales générées avec l'IA pour la défixation : étude comparative

Fabrice FISCHER

RÉSUMÉ

En Design Thinking, l'idéation est le moment où les espaces de solution devraient s'ouvrir. Pourtant, après une première vague d'idées, les designers tendent à revenir vers leurs idées initiales. La génération s'essouffle alors. Cette fermeture progressive renvoie au phénomène de fixation et peut limiter la recherche d'une meilleure solution.

Cette thèse part du constat que des interférences peuvent contribuer à la défixation, sans que toutes aient la même capacité de relance. La thèse mobilise l'intelligence artificielle générative (GenAI) comme infrastructure d'idéation, c'est-à-dire comme un moyen de produire des interférences spécifiques au projet sans se substituer au designer. Son objectif est d'explorer, de concevoir et de comparer des interférences produites avec GenAI afin d'évaluer dans quelle mesure elles peuvent relancer l'idéation humaine après fixation.

Le programme doctoral articule la méthodologie de recherche en science de la conception (DSRM) et l'étude de cas comparée (CCS).

Le premier article explore la génération de suggestions textuelles orientées comme interférences à la fixation et les évalue par classement forcé. Le deuxième étudie cette logique sur une problématique d'innovation des bibliothèques publiques, en comparant une interférence textuelle spécifique au projet, une interférence image et une pause neutre. Le troisième positionne GenAI comme une infrastructure d'idéation et compare une interférence textuelle, une mosaïque d'images, une mosaïque vidéo et une pause témoin d'une minute.

Les résultats montrent qu'il est possible de produire des suggestions textuelles différenciées selon leur orientation sémantique. Ils indiquent ensuite que, sur le cas des bibliothèques publiques, les modalités texte et image présentent un signal directionnel positif de relance de l'idéation par rapport à la pause neutre, avec un signal plus marqué pour le texte. Enfin, dans le protocole multimodal étudié, l'interférence textuelle présente le signal directionnel le plus fort de relance et de déplacement conceptuel, suivie de la mosaïque d'images. La pause témoin occupe une position intermédiaire, tandis que la mosaïque vidéo apparaît la moins performante dans ce protocole. Ce résultat doit toutefois être compris comme une configuration empirique observée dans le protocole étudié, indiquant que les médias génératifs plus riches ne sont pas automatiquement plus efficaces, et non comme une hiérarchie universelle des modalités d'interférence.

Pris ensemble, les trois articles montrent que les interférences n'ont pas toutes le même potentiel de relance. Leur efficacité dépend de leur construction sémantique et de leur modalité de présentation. La thèse contribue aux travaux sur la fixation, la défixation et l'idéation assistée par GenAI. Elle évalue ces interférences à l'aide de deux indicateurs complémentaires : la fluence et le FixShift.

Sur le plan pratique, la thèse ouvre une voie pour les praticiens qui souhaitent utiliser GenAI de manière simple, peu coûteuse et non prescriptive lorsque la production d'idées s'essouffle. Les résultats doivent toutefois être interprétés avec prudence, car le programme reste

exploratoire. La thèse se conclut par l'identification d'une piste de recherche et d'action prometteuse : utiliser GenAI pour produire des interférences adaptées au projet afin de soutenir la défixation et de relancer l'idéation sans retirer au designer son rôle central.

Mots-clés : idéation, fixation, défixation, interférences générées par IA, infrastructure d'idéation, GenAI, FixShift

Multimodal Generative AI Interferences for Defixation: A Comparative Case Study

Fabrice FISCHER

ABSTRACT

In Design Thinking, the ideation phase plays a central role. Yet, after an initial wave of relatively unoriginal ideas, designers tend to fixate on their initial ideas and quickly exhaust their generative capacity. This decline, both in the number of ideas and in their diversity, refers to the fixation phenomenon, which hinders the search for a better final solution.

This thesis starts from the observation that interferences can contribute to defixation, although not all of them have the same reactivation potential. Their effect appears to depend on their nature, their delivery modality, and their degree of project specificity. The thesis mobilizes generative artificial intelligence (GenAI) as ideation infrastructure, i.e., as a way to generate project-specific interferences without replacing the human designer. Its objective is to explore, design, and compare these interferences in order to assess the extent to which they reactivate human ideation after fixation.

The doctoral program combines Design Science Research Methodology (DSRM) and Comparative Case Study (CCS).

The first article explores the generation with ChatGPT 4.0 of oriented textual suggestions as fixation interferences and evaluates them through forced ranking. The second article extends this logic to a public-library innovation problem by comparing a project-specific text interference, a project-specific image interference, and a neutral break. The third article positions GenAI as ideation infrastructure and compares, under a shared post-fixation protocol, text interference, image mosaic interference, video mosaic interference, and a one-minute neutral break.

The results show that it is possible to produce textual suggestions differentiated by their semantic orientation. They further indicate that, in the public-library case, text and image interferences show a positive directional relaunch pattern compared with the neutral break, with a stronger signal for text. Finally, in the multimodal protocol studied, text-based interference shows the strongest directional signal of reactivation and conceptual movement, followed by the image mosaic; the neutral break occupies an intermediate position, whereas the video mosaic appears least effective under this protocol. This result should be understood as a situated pattern showing that richer generative media are not automatically more effective, not as a universal hierarchy of interference modalities.

Taken together, the three articles show that not all interferences are equal. Their effectiveness depends on both their semantic construction and their delivery modality. The thesis contributes to the literature on fixation, defixation, and ideation assisted by generative artificial intelligence. It shows how to design project-specific interferences and relate them to an indicator of conceptual shift with respect to initial idea families (FixShift).

From a practical perspective, the thesis opens a path for practitioners who wish to use GenAI in a simple, low-cost, and non-prescriptive way when idea production begins to stall. The results should nevertheless be interpreted with caution, as the program remains exploratory.

The thesis concludes not with a universal recipe, but with the identification of a promising avenue for research and action: using GenAI to produce project-specific interferences that support defixation and reactivate ideation without displacing the designer from their central role.

Keywords: ideation, fixation phenomenon, defixation, GenAI-generated interferences, ideation infrastructure, human-AI creativity, FixShift

TABLE DES MATIÈRES

	Page
INTRODUCTION	1
CHAPITRE 1 REVUE CRITIQUE DE LA LITTÉRATURE.....	5
1.1 Design Thinking et idéation.....	5
1.2 Phénomène de fixation.....	6
1.3 Défixation et interférences.....	7
1.4 Intelligence artificielle générative et idéation.....	8
1.5 Positionnement scientifique de la thèse	9
CHAPITRE 2 CADRE MÉTHODOLOGIQUE DU PROGRAMME DOCTORAL.....	11
2.1 Positionnement méthodologique général.....	11
2.2 Articulation entre la méthodologie de recherche en science de la conception (DSRM) et l'étude de cas comparée (CCS).....	12
2.3 Programme doctoral en trois études.....	13
2.4 Construction des artefacts d'interférence.....	14
2.5 Cas, problèmes et protocoles	15
2.6 Participants, cohortes et sites	16
2.7 Données collectées et mesures.....	17
2.8 Procédures d'analyse et validation.....	18
2.9 Limites méthodologiques transversales	19
2.10 Articulation des chapitres-articles.....	20
CHAPITRE 3 COMPARING DIFFERENT CHATGPT4.0 GENERATED ORIENTED WORD SUGGESTIONS AS DESIGN FIXATION INTERFERENCES FOR THE IDEATION PHASE IN DESIGN THINKING	23
3.1 Introduction.....	24
3.2 The problem.....	25
3.3 Research objectives and hypotheses	26
3.4 Theoretical Background.....	27
3.4.1 Design Thinking.....	27

3.4.2	Word interferences to mitigate the design fixation phenomenon	28
3.4.3	Design Science Research Methodology (DSRM).....	29
3.4.4	Generative Artificial Intelligence (GAI) with its Large Language Models (LLM) to generate lists of words as artifacts of the DSRM	29
3.4.5	Comparative Case Studies (CCS)	29
3.5	Methodology	30
3.6	The creation of the artifacts	31
3.6.1	Artifact method	31
3.6.2	Artifact Instantiation: Identification of examples of DT problems	32
3.6.3	Data & sources	33
3.6.4	Data preparation and collection methods.....	33
3.7	Evaluation	34
3.8	Discussion	34
3.8.1	Theoretical and practical contributions.....	35
3.8.2	Limitations and future research	35
3.9	Conclusion	35
CHAPITRE 4	GENAI AS A PROJECT-SPECIFIC INTERFERENCE GENERATOR FOR IDEATION DEFIXATION	37
4.1	Introduction.....	38
4.1.1	Ideation fixation in innovation work.....	38
4.1.2	GenAI as a facilitator, not ideator.....	39
4.1.3	Project-specific interference for ideation launch	39
4.1.4	Study context and research question.....	40
4.1.5	Contributions.....	41
4.2	Literature Review.....	42
4.2.1	Ideation fixation and defixation in creativity and innovation management	42
4.2.2	GenAI as creativity facilitator rather than ideator	44
4.2.3	Project-specific interferences for defixation.....	46
4.2.4	From prior evidence to the present study.....	48
4.2.5	Expected relationships	49

4.3	Methodology	50
4.3.1	Prior studies and evolution of the protocol	50
4.3.2	Research design	51
4.3.3	Empirical context and ideation task	51
4.3.4	Participants and data collection	52
4.3.5	Experimental procedure	52
4.3.6	Interference modalities.....	53
4.3.7	Measures	54
4.3.8	Data cleaning and exclusion criteria	55
4.3.9	Analysis strategy	55
4.3.10	Ethical and methodological boundaries	56
4.4	Results.....	56
4.4.1	Data cleaning and final analytical sample	56
4.4.2	Baseline pre-fixation fluency	57
4.4.3	Post-fixation fluency by modality and control group	58
4.4.4	Change in fluency after the interference	58
4.4.5	Group-level statistical comparison	59
4.4.6	Summary of findings by research question.....	60
4.5	Discussion.....	61
4.5.1	Interpreting the directional relaunch pattern.....	61
4.5.2	GenAI as interference infrastructure, not ideator	62
4.5.3	Text modality as a low-friction defixation mechanism	63
4.5.4	Why image interference may not automatically outperform text	63
4.5.5	Implications for creativity and innovation management	64
4.5.6	Boundary conditions and limitations	65
4.5.7	Future research.....	66
4.6	Conclusion	66
CHAPITRE 5	GENERATIVE AI AS IDEATION INFRASTRUCTURE: RELAUNCHING HUMAN IDEATION THROUGH TEXT, IMAGE, OR VIDEO INTERFERENCES.....	69
5.1	Introduction.....	70

5.1.1	GenAI in the work of innovating	70
5.1.2	From GenAI as direct ideator to GenAI as ideation infrastructure.....	71
5.1.3	Post-fixation human ideation as the problem.....	72
5.1.4	Research gap: limited cross-modal evidence on GenAI-generated interferences	73
5.1.5	Research questions and contributions.....	74
5.2	Theoretical Background.....	75
5.2.1	GenAI in innovation management	75
5.2.2	Human-AI creativity and the direct-ideator paradox	76
5.2.3	Fixation, defixation, and ideation relaunch.....	77
5.2.4	Interferences as human defixation mechanisms	77
5.2.5	Text, image mosaic, and video mosaic as interference modalities	78
5.2.6	Fluency and FixShift as complementary outcomes	79
5.3	Methodology	79
5.3.1	Research design overview.....	80
5.3.2	Empirical setting: the public-library innovation brief	80
5.3.3	GenAI as ideation infrastructure: interference generation logic.....	81
5.3.4	Text interference, image mosaic interference, video mosaic interference, and neutral break.....	82
5.3.5	Participants, cohorts, and data collection protocol	85
5.3.6	Measure: fluency and FixShift.....	86
5.3.7	AI-assisted FixShift coding and human validation.....	87
5.3.8	Data cleaning and exclusions.....	87
5.4	Results.....	88
5.4.1	Dataset overview.....	89
5.4.2	Post-fixation fluency by interference modality and neutral break.....	89
5.4.3	FixShift by interference modality and neutral break	91
5.4.4	Cross-modal pattern	92
5.4.5	Human validation of FixShift coding	93
5.4.6	Robustness and interpretation checks	94
5.5	Discussion	95

5.5.1	GenAI as ideation infrastructure, not direct ideator.....	95
5.5.2	Defixation as a human-centered innovation-process method	96
5.5.3	Interference modality and the limits of richer media.....	97
5.5.4	Implications for GenAI-assisted innovation management.....	98
5.5.5	Limits and boundaries of the findings	99
5.6	Conclusion	101
CONCLUSION GÉNÉRALE.....		103
6.1	Synthèse des résultats	103
6.2	Contribution de la thèse	104
6.3	Limites	105
6.4	Pistes de recherche.....	106
ANNEXE I	FORMULAIRE TYPE ET VARIANTES D’INTERFÉRENCES	109
ANNEXE II	PROMPTS DE GÉNÉRATION DES INTERFÉRENCES	115
ANNEXE III	STRUCTURE MINIMALE DES DONNÉES ET VARIABLES D’ANALYSE.....	121
LISTE DE RÉFÉRENCES BIBLIOGRAPHIQUES		129

LISTE DES TABLEAUX

	Page
Table 4.1	Post-fixation fluency and directional change by comparison group 57
Table 5.1	Post-fixation interferences and neutral break 84
Table 5.2	Post-fixation fluency by interference modality and neutral break 90
Table 5.3	FixShift by interference modality and neutral break 91

LISTE DES ABRÉVIATIONS, SIGLES ET ACRONYMES

CIM	Creativity and Innovation Management
CCS	étude de cas comparée (<i>Comparative Case Study</i>)
DSRM	méthodologie de recherche en science de la conception (<i>Design Science Research Methodology</i>)
DT	Design Thinking
FixShift	indicateur de déplacement conceptuel par rapport aux familles d'idées initiales
Fluence	indicateur du nombre d'idées produites après l'interférence
FP	phénomène de fixation (<i>Fixation Phenomenon</i>)
GAI	intelligence artificielle générative (<i>Generative Artificial Intelligence</i>)
GenAI	intelligence artificielle générative (<i>Generative Artificial Intelligence</i>)
LLM	grand modèle de langage (<i>Large Language Model</i>)
PLM	gestion du cycle de vie des produits (<i>Product Lifecycle Management</i>)
WI	interférence verbale (<i>Word Interference</i>)

INTRODUCTION

0.1 Contexte

L'innovation constitue un moteur important d'amélioration du bien-être humain. Parmi les méthodologies mobilisées pour soutenir la créativité en ingénierie, le Design Thinking occupe une place importante. En Design Thinking (Brown, 2008), l'idéation constitue la phase centrale de l'ouverture des espaces de solution.

Après une première vague d'idées de solution possibles, le designer tend à s'essouffler rapidement dans la génération de nouvelles idées (Thomas & Tee, 2022). En pratique, il tend à se fixer sur une idée initiale de solution (Youmans & Arciszewski, 2014). L'exploration des solutions possibles s'en trouve limitée (Purcell & Gero, 1996). La qualité de la solution finale peut s'en trouver affectée (Brown & Wyatt, 2009).

0.2 Problématique

Cette fixation, bien documentée dans la littérature en design, contraint l'idéation (Jansson & Smith, 1991). Plusieurs formes de défixation ont fait l'objet de recherches, et la littérature pointe notamment vers des distractions ou interférences susceptibles de relancer l'idéation (Sio, Kotovsky, & Cagan, 2015 ; Vasconcelos & Crilly, 2016). Toutes ces interférences ne présentent toutefois pas le même potentiel de relance de l'idéation (Chan et al., 2011). Il est donc intéressant de comparer les impacts respectifs de ces interférences.

Avec les stratégies obliques de Brian Eno, un jeu de cartes propose des mots ou contraintes utilisés comme interférences lorsque l'idéation se bloque, mais ces suggestions restent fixes et génériques (Eno & Schmidt, 1975). Dans cette thèse, l'intelligence artificielle générative est considérée comme une infrastructure d'idéation capable de produire, à la demande, des interférences spécifiques au projet, conditionnées par le brief et le contexte de conception (Luo, Sarica, & Wood, 2021). Elle n'est pas utilisée comme un substitut au designer. Cependant, avec cette nouvelle capacité de génération d'interférences, la question de leur efficacité relative

demeure centrale (Fang et al., 2025 ; Wadinambiarachchi, Kelly, Pareek, Zhou, & Velloso, 2024).

0.3 Objectif général et questions de recherche

Cette thèse exploratoire porte sur la conception d'interférences générées avec l'appui de l'IA générative et sur la comparaison de leur capacité à relancer l'idéation après fixation, à la fois en nombre d'idées produites et en déplacement conceptuel.

- QR1 : Quelles suggestions verbales orientées semblent les plus prometteuses?
- QR2 : Les interférences textuelles et visuelles spécifiques au projet favorisent-elles davantage la relance de l'idéation qu'une pause neutre?
- QR3 : Dans quelle mesure les modalités texte, mosaïque d'images, mosaïque vidéo et pause témoin influencent-elles la relance post-fixation et le déplacement conceptuel des idées?

0.4 Aperçu méthodologique

Cette recherche s'appuie sur deux méthodologies complémentaires. La méthodologie de recherche en science de la conception (Design Science Research Methodology, DSRM) encadre la conception des artefacts d'interférence (Peffer, Tuunanen, Rothenberger, & Chatterjee, 2007 ; De Sordi, 2021 ; Johannesson & Perjons, 2021). L'étude de cas comparée (Comparative Case Study, CCS) structure la comparaison des impacts relatifs des interférences générées (Bartlett & Vavrus, 2017).

Le programme doctoral s'articule en trois études. La première explore les suggestions verbales orientées. La seconde compare une interférence textuelle spécifique au projet, une interférence image et une pause neutre. La troisième positionne GenAI comme infrastructure d'idéation et compare, dans un protocole post-fixation commun, une interférence textuelle, une mosaïque d'images, une mosaïque vidéo et une pause témoin d'une minute.

Les deuxième et troisième études partagent le brief d'innovation des bibliothèques publiques afin de stabiliser le problème de design, mais elles ne répondent pas au même objectif

comparatif. L'étude 2 compare texte, image et pause neutre à partir de la fluence. L'étude 3 étend cette logique à une comparaison multimodale intégrant la mosaïque vidéo et ajoute le FixShift afin de distinguer la reprise de production d'idées du déplacement conceptuel.

0.5 Structure de la thèse

Après cette introduction, le chapitre 1 présente la revue critique de la littérature. Le chapitre 2 expose le cadre méthodologique du programme doctoral. Les chapitres 3, 4 et 5 regroupent les trois articles. Ils suivent une progression allant des interférences verbales orientées à une comparaison texte–image–pause neutre, puis à une comparaison multimodale positionnant GenAI comme infrastructure d'idéation. La thèse se termine par une conclusion générale, suivie des annexes et de la liste des références bibliographiques.

CHAPITRE 1

REVUE CRITIQUE DE LA LITTÉRATURE

Ce chapitre propose une analyse critique de l'état des connaissances sur la génération et la comparaison de différents types et modalités d'interférences de défixation dans la phase d'idéation du Design Thinking.

1.1 Design Thinking et idéation

L'innovation occupe une place centrale dans les discours sur la création de valeur et le développement de nouvelles solutions. Parmi les méthodologies mobilisées pour soutenir la créativité en conception, le Design Thinking (DT) s'est progressivement imposé comme une démarche reconnue d'innovation structurée (Brown, 2008 ; Brown & Wyatt, 2009). Il se révèle particulièrement pertinent pour des problèmes de design ouverts, complexes et centrés sur l'humain. Bien que sa mise en œuvre soit non linéaire, il est souvent décrit à travers cinq phases principales : l'empathie, la définition, l'idéation, le prototypage et le test.

Les phases d'empathie et de définition structurent l'exploration et la formulation du problème, tandis que le prototypage et le test permettent d'évaluer la pertinence, la faisabilité et l'acceptabilité des solutions envisagées. L'idéation occupe une place centrale dans cette démarche, puisqu'elle constitue le moment où les espaces de solution se déploient réellement. Cette phase s'ouvre sur une dynamique de divergence, avant de converger vers des pistes de solution jugées suffisamment prometteuses pour être prototypées et testées (Brown, 2008 ; Brown & Wyatt, 2009).

L'enjeu de la phase d'idéation n'est pas seulement de produire un grand nombre d'idées, mais aussi de générer des idées suffisamment variées pour explorer plusieurs espaces de solution (Thomas & Tee, 2022). Cette diversité conditionne la possibilité de faire émerger des solutions plus originales, plus pertinentes et, potentiellement, plus susceptibles d'être adoptées par leurs usagers et parties prenantes.

Au début de la phase d'idéation, les designers produisent généralement une première vague d'idées, mais cette génération s'essouffle rapidement, ce qui réduit à la fois la quantité d'idées nouvelles et la diversité des solutions effectivement explorées (Thomas & Tee, 2022). De plus, ces premières idées demeurent souvent proches de solutions déjà connues et ne sont pas nécessairement les plus innovantes.

La phase d'idéation est donc traversée par une tension entre potentiel d'ouverture et essoufflement progressif. Cette tension rend nécessaire une meilleure compréhension des freins à l'idéation, au premier rang desquels figure le phénomène de fixation (Purcell & Gero, 1996).

1.2 Phénomène de fixation

Dans le cadre du Design Thinking, le phénomène de fixation renvoie à la tendance du designer à rester ancré sur une idée initiale ou sur une famille restreinte de solutions. Il constitue un problème central et récurrent de l'idéation en design (Jansson & Smith, 1991). Cet ancrage précoce réduit l'exploration ultérieure des alternatives et correspond à un problème cognitif important dans l'activité de conception (Purcell & Gero, 1996).

Les origines de la fixation sont multiples. L'ancrage peut se former aussi bien à partir d'idées auto-générées par le designer ou son équipe qu'à partir d'exemples fournis de l'extérieur (Leahy, Daly, McKilligan, & Seifert, 2020). La littérature distingue ainsi plusieurs formes de fixation et plusieurs manières de les prévenir (Youmans & Arciszewski, 2014). Il ne s'agit donc pas d'un simple réflexe, mais d'un phénomène à sources multiples, ce qui justifie la recherche d'aides opérationnelles telles que des stimuli, des distractions ou, dans le cadre de cette thèse, des interférences.

La fixation ne se contente pas de réduire le nombre d'idées nouvelles : elle réduit également leur diversité, puisque l'activité cognitive du designer se resserre autour d'un nombre limité d'espaces de solution. Dans ce contexte, il ne suffit pas de relancer la production d'idées; il faut aussi susciter une plus grande diversité des idées produites (Sio et al., 2015). Les interférences se trouvent ainsi prises dans une tension entre potentiel d'inspiration et risque d'ancrage dans un nouvel enfermement cognitif (Vasconcelos & Crilly, 2016).

La prise de conscience du phénomène de fixation par le designer est utile, mais elle demeure insuffisante. Il faut donc agir activement pour relancer l'exploration (Youmans & Arciszewski, 2014). Aussi, l'effet des interférences dépend de la manière dont elles sont introduites (Sio et al., 2015).

Dès lors que la fixation est reconnue comme un problème récurrent de l'idéation, la question centrale devient celle des mécanismes capables de rouvrir le champ de recherche du designer. Dans cette perspective, certaines interférences peuvent être mobilisées comme formes de distraction ou de relance en soutien à la créativité en phase de fixation (Koh & De Lessio, 2018). Cette question se pose désormais également dans des contextes où l'IA générative intervient dans la production de stimuli ou d'interférences, sans pour autant devoir occuper elle-même le rôle d'idéateur (Wadinambiarachchi et al., 2024).

1.3 Défixation et interférences

En phase d'idéation, la défixation peut être comprise comme le processus cognitif et créatif qui permet de dépasser l'ancrage lié au phénomène de fixation. Il ne s'agit pas seulement de produire davantage d'idées, mais de produire des idées suffisamment différentes pour rouvrir d'autres cadrages, d'autres associations et d'autres espaces de solution. L'inspiration peut en effet aussi bien ouvrir qu'enfermer la recherche (Vasconcelos & Crilly, 2016), et la relance de l'idéation dépend d'effets cognitifs différenciés selon les exemples et stimuli mobilisés (Sio et al., 2015).

Une interférence peut être assimilée à une perturbation volontaire introduite lorsque l'idéation s'essouffle, afin de déplacer l'attention hors de la solution dominante ou déjà ancrée. Dans un contexte de design créatif, ces interférences peuvent jouer un rôle utile de relance (Koh & De Lessio, 2018). Elles peuvent notamment agir comme déclencheurs de déplacement conceptuel et de reconfiguration de la recherche (Vasconcelos & Crilly, 2016).

Les interférences peuvent prendre des formes très diverses. Il peut s'agir de mots, d'images, de vidéos, de dessins, d'objets, de sons ou d'autres supports encore. L'important est qu'elles remplissent leur fonction d'interférer avec le phénomène de fixation. Elles peuvent agir par plusieurs mécanismes, notamment le déplacement attentionnel, l'activation d'associations

distantes ou la reconfiguration du problème, en particulier par des analogies et des formes de représentation variées (Chan et al., 2011). L'effet des interférences fondées sur des exemples dépend toutefois de leur structure et de leurs usages (Sio et al., 2015).

Toutes les interférences ne se valent donc pas : leur efficacité dépend de leur contenu sémantique, de leur distance par rapport au problème, de leur pertinence au projet et de leur modalité de présentation. La nature et la représentation de l'interférence influencent sa capacité à soutenir la défixation (Chan et al., 2011), et certaines interférences se révèlent plus utiles que d'autres (Koh & De Lessio, 2018). Dans cette logique de relance indirecte plutôt que prescriptive, les stratégies obliques de Brian Eno constituent une référence importante (Eno & Schmidt, 1975).

Il existe une grande diversité de formes d'interférences, et leur comparaison devient en elle-même un enjeu scientifique. Des recherches récentes explorent déjà la diversité de leurs modalités, par exemple à travers des mots, des dessins ou d'autres supports (Baudoux, Guo, & Goucher-Lambert, 2025). La forme du support influe sur la manière dont l'interférence agit (Chan et al., 2011), et la production de ces interférences fait désormais aussi l'objet de recherches avec l'appui de l'intelligence artificielle générative (Gong, Soomro, Wang, Latif, & Georgiev, 2024).

Les interférences apparaissent ainsi comme une voie prometteuse de défixation, mais il reste encore à préciser quelles interférences, sous quelles formes et sous quelles conditions, soutiennent réellement la défixation.

1.4 Intelligence artificielle générative et idéation

L'un des principaux attraits de l'IA générative est qu'elle modifie les conditions de l'idéation en permettant de produire rapidement des contenus susceptibles de servir d'interférences ou d'appuis à la génération d'idées (Gong et al., 2024 ; Fang et al., 2025).

Toutefois, cette rapidité et ce volume de production ne se traduisent pas nécessairement par une meilleure exploration des solutions ou des espaces de solution. Au contraire, ils peuvent renforcer la convergence, la fixation externe et une forme de dépendance passive du designer

à l'égard des productions générées par le système (Fang et al., 2025 ; Wadinambiarachchi et al., 2024).

Il devient donc important de distinguer deux usages de l'IA générative : d'une part, son usage comme idéateur direct; d'autre part, son usage comme infrastructure de production d'interférences destinées à relancer l'idéation humaine sans fournir directement les solutions (Chen et al., 2025 ; Wadinambiarachchi et al., 2024).

Par ailleurs, l'IA générative peut permettre de produire des interférences spécifiques au projet. (Luo et al., 2021) proposent notamment une logique de distance de connaissance, tandis que (Gong et al., 2024) montrent l'usage de l'IA générative pour produire des interférences textuelles, et que (Fang et al., 2025) en offrent un cadrage d'ensemble.

Enfin, l'IA générative ouvre la possibilité de comparer plusieurs modalités d'interférences, notamment textuelles, images et vidéo. Comme le soulignent (Baudoux et al., 2025), il ne s'agit pas d'un simple changement de support, mais potentiellement d'un changement de mécanisme cognitif.

La littérature progresse rapidement, mais elle reste fragmentée lorsqu'il s'agit de comparer ces modalités d'interférence sous un protocole post-fixation commun, avec des indicateurs capables de distinguer la relance de production d'idées et le déplacement conceptuel. L'IA générative rend donc possible la production automatique d'interférences, mais la question de la comparaison rigoureuse de leur potentiel différencié de défixation demeure ouverte.

1.5 Positionnement scientifique de la thèse

La littérature établit la centralité de l'idéation, documente le phénomène de fixation, suggère le potentiel des interférences et reconnaît un intérêt croissant pour l'IA générative, sans pour autant fournir de cadre comparatif suffisamment intégré (Purcell & Gero, 1996 ; Brown, 2008 ; Koh & De Lessio, 2018 ; Wadinambiarachchi et al., 2024 ; Fang et al., 2025).

Un premier manque de la littérature concerne la comparaison des interférences verbales selon leur logique sémantique et leur capacité à agir face aux limites de connaissance et de générativité (Koh & De Lessio, 2018 ; Thomas & Tee, 2022 ; Fischer & Gardoni, 2025a).

Un deuxième manque concerne la comparaison, dans un même problème d'innovation, entre une interférence textuelle spécifique au projet, une interférence image spécifique au projet et une pause neutre, afin d'évaluer si des interférences générées avec GenAI relancent davantage l'idéation qu'une simple interruption (Chan et al., 2011 ; Koh & De Lessio, 2018 ; Fischer & Gardoni, 2025b).

Un troisième manque concerne la comparaison, sous un protocole post-fixation commun, des modalités texte, mosaïque d'images, mosaïque vidéo et interruption neutre, en mobilisant conjointement la fluence et le FixShift pour distinguer reprise de production d'idées et déplacement conceptuel (Gong et al., 2024 ; Wadinambiarachchi et al., 2024 ; Fang et al., 2025).

La présente thèse répond progressivement à ces trois insuffisances. Elle étudie d'abord des suggestions verbales orientées. Elle compare ensuite texte, image et pause neutre dans le cas des bibliothèques publiques. Elle propose enfin une comparaison multimodale intégrant texte, mosaïque d'images, mosaïque vidéo, interruption neutre, fluence et FixShift. Elle apporte ainsi des contributions théoriques, méthodologiques et pratiques, notamment par l'articulation entre la DSRM et la CCS (Bartlett & Vavrus, 2017 ; Peffers, Tuunanen, & Niehaves, 2018) et par l'appui à une mesure du déplacement conceptuel (Shah, Smith, & Vargas-Hernandez, 2003 ; Zhang, Han, & Ahmed-Kristensen, 2025).

Le positionnement de cette thèse ne consiste donc pas à évaluer la capacité de l'IA générative à remplacer l'idéation humaine, mais à évaluer sa capacité à fonctionner comme infrastructure de production d'interférences capables de soutenir la défixation sans déposséder le designer de son rôle cognitif central.

Pour répondre à ces trois insuffisances, la thèse n'a pas été conçue comme une étude unique, mais comme un programme de recherche en trois études, articulé par un cadre méthodologique spécifique, présenté au chapitre suivant.

CHAPITRE 2

CADRE MÉTHODOLOGIQUE DU PROGRAMME DOCTORAL

Ce chapitre présente le cadre méthodologique du programme doctoral, en explicitant l'articulation entre la conception d'artefacts d'interférence et la comparaison structurée de leurs effets à travers trois études successives.

2.1 Positionnement méthodologique général

Cette thèse s'inscrit dans une recherche orientée vers la conception d'artefacts d'interférence. Dans cette perspective, la théorie ne se limite pas à l'observation de phénomènes existants, mais peut également se développer à partir d'artefacts conçus pour répondre à un problème donné (De Sordi, 2021). Le programme doctoral s'appuie ainsi sur la méthodologie de recherche en science de la conception (*Design Science Research Methodology*, DSRM), qui fournit un cadre pertinent pour concevoir, instancier et évaluer des artefacts de recherche en contexte de conception (Peffer et al., 2018 ; Johannesson & Perjons, 2021).

L'objet central du programme n'est pas l'usage de l'IA générative comme idéateur autonome, mais l'infrastructure d'interférence générée avec l'IA, conçue pour intervenir lorsque l'idéation s'essouffle. Dans cette thèse, l'IA générative est mobilisée comme moyen de produire des interférences spécifiques au projet, conditionnées par l'énoncé du problème, le vocabulaire, le contexte et les contraintes de conception (Luo et al., 2021). Ce positionnement permet de conserver le designer humain au centre du processus d'idéation, tout en utilisant l'IA comme soutien à la défixation.

Les projets de Design Thinking constituent généralement des singularités contextuelles. Ils sont peu nombreux, situés et difficilement répliquables dans une logique strictement expérimentale à grande échelle. Ils relèvent donc plus souvent de comparaisons qualitatives et contextualisées que de protocoles massifs visant des généralisations universelles (Johannesson & Perjons, 2021 ; Mello, 2021).

Dans ce cadre, la thèse vise moins à établir des estimations causales universelles qu'à produire des résultats comparatifs structurés, permettant de dégager des tendances interprétables entre différents types et différentes modalités d'interférences (Bartlett & Vavrus, 2017).

Ce positionnement conduit à articuler deux logiques méthodologiques complémentaires : une logique de conception d'artefact, portée par la DSRM, et une logique de comparaison structurée des effets de cet artefact, portée par l'étude de cas comparée (*Comparative Case Study*, CCS).

2.2 Articulation entre la méthodologie de recherche en science de la conception (DSRM) et l'étude de cas comparée (CCS)

La méthodologie de recherche en science de la conception (*Design Science Research Methodology*, DSRM) et l'étude de cas comparée (*Comparative Case Study*, CCS) ne sont pas juxtaposées dans cette thèse, mais remplissent deux fonctions différentes et complémentaires.

La DSRM est mobilisée pour concevoir l'artefact d'interférence, ses instanciations et ses règles de génération (De Sordi, 2021 ; Peffers et al., 2018). Cette logique de conception structure l'ensemble du programme doctoral : le chapitre 3 porte sur des listes orientées de suggestions verbales, le chapitre 4 sur une comparaison entre interférence textuelle, interférence image et pause neutre, et le chapitre 5 sur des interférences déclinées selon les modalités texte, mosaïque d'images, mosaïque vidéo et pause témoin.

La CCS est, pour sa part, mobilisée pour comparer les instanciations d'un même principe d'interférence, soit par rapport à une condition témoin, soit entre plusieurs modalités ou variantes d'un même artefact (Bartlett & Vavrus, 2017 ; Kaarbo & Beasley, 1999). Cette logique comparative est déjà présente dans le chapitre 3, puis elle se renforce dans les chapitres 4 et 5.

La combinaison entre la DSRM et la CCS est particulièrement adaptée à une recherche dans laquelle il s'agit à la fois de concevoir des solutions et d'en comparer les effets dans des contextes réels et contextualisés (Johannesson & Perjons, 2021 ; Mello, 2021). Elle convient bien à des projets de Design Thinking, qui sont peu nombreux, situés et difficilement répliquables dans une logique strictement expérimentale à grande échelle.

Cette articulation évolue au fil des trois études. Le chapitre 3 correspond à une première exploration qualitative des suggestions orientées. Le chapitre 4 approfondit cette logique par une comparaison entre interférence textuelle, interférence image et pause neutre. Le chapitre 5 en propose enfin une version plus structurée, en comparant texte, mosaïque d'images, mosaïque vidéo et pause témoin selon un protocole post-fixation commun.

Autrement dit, la DSRM porte ici la logique de conception de l'artefact, tandis que la CCS porte la logique de comparaison de ses instanciations. Leur articulation n'est donc pas additive, mais fonctionnelle.

2.3 Programme doctoral en trois études

Le programme doctoral s'organise en trois études qui reflètent une progression logique, à la fois sur le plan scientifique et sur le plan méthodologique.

La première étude constitue une exploration qualitative et comparative de suggestions verbales orientées comme interférences à la fixation. Elle vise à examiner, à partir d'un protocole de classement forcé, quelles formes de suggestions semblent les plus prometteuses pour relancer l'idéation. Cette première étape relève d'une logique exploratoire, fondée sur l'étude de cas comparée (Mello, 2021), tout en s'inscrivant dans une démarche de conception d'artefacts d'interférence (Johannesson & Perjons, 2021). Elle permet aussi d'introduire la question de la limite de générativité dans les processus d'idéation (Thomas & Tee, 2022).

La deuxième étude approfondit cette logique en comparant, dans un même problème d'innovation des bibliothèques publiques, une interférence textuelle spécifique au projet, une interférence image spécifique au projet et une pause neutre. Il ne s'agit plus seulement d'explorer des suggestions verbales orientées, mais d'évaluer si des interférences générées avec GenAI relancent davantage l'idéation qu'une interruption sans contenu sémantique. Cette étude renforce ainsi la dimension comparative du programme, tout en maintenant une logique de conception d'artefacts et d'évaluation contextualisée (De Sordi, 2021 ; Kaarbo & Beasley, 1999).

La troisième étude constitue le volet le plus structuré du programme doctoral. Elle reprend un même principe d'interférence spécifique au projet, mais l'instancie sous forme d'interférence

textuelle, de mosaïque d'images, de mosaïque vidéo et de pause témoin. Elle compare ces conditions dans un protocole post-fixation commun à l'aide de deux indicateurs complémentaires, la fluence et le FixShift, afin de distinguer la relance de production d'idées du déplacement conceptuel entre familles d'idées (Bartlett & Vavrus, 2017 ; Peffers et al., 2018).

Dans son ensemble, le programme doctoral progresse donc d'une comparaison d'orientations sémantiques à une comparaison texte-image-pause neutre, puis à une comparaison multimodale intégrant texte, mosaïque d'images, mosaïque vidéo et pause témoin. Cette progression cumulative conduit naturellement à expliciter, dans la section suivante, la manière dont les artefacts d'interférence sont construits.

2.4 Construction des artefacts d'interférence

Dans cette thèse, l'artefact étudié n'est pas l'idée générée avec l'IA, mais le dispositif d'interférence conçu pour relancer l'idéation lorsque celle-ci s'essouffle. La logique d'artefact est donc centrale : il s'agit de concevoir une forme d'intervention ciblée, de l'instancier dans des formats adaptés au problème traité, puis d'en comparer les effets (Peffers et al., 2018).

Dans la première étude, l'artefact prend la forme de trois listes de suggestions verbales orientées, générées avec ChatGPT 4.0, qui constituent les premières instanciations d'interférences à la fixation.

Dans la deuxième étude, l'artefact est décliné en deux modalités d'interférence spécifiques au projet, texte et image, comparées à une pause neutre.

Dans la troisième étude, un même principe d'interférence spécifique au projet est conservé, mais il est instancié selon plusieurs conditions : une interférence textuelle, une mosaïque d'images, une mosaïque vidéo et une interruption neutre d'une minute. Dans les trois modalités génératives, il ne s'agit pas de fournir directement une solution, mais de produire une interférence susceptible de relancer l'exploration.

Les artefacts sont ainsi construits de manière à éviter la prescription directe de solution, qui risquerait de provoquer un nouvel ancrage de fixation. Ils sont au contraire conçus pour rester

pertinents pour le problème traité tout en maintenant une distance suffisante pour favoriser la défixation (Luo et al., 2021).

Leur construction repose sur une logique de traçabilité explicite, qui inclut les consignes de génération, la logique de sélection, la version du modèle utilisé et la conservation de traces dans le classeur de données ainsi qu'en annexe. La feuille `FixShift_AI_Method_Trace` documente plus particulièrement cette chaîne de traitement.

Une fois les artefacts définis, il reste à préciser dans quels cas, sur quels problèmes et selon quels protocoles ils ont été mobilisés.

2.5 Cas, problèmes et protocoles

Les trois études du programme doctoral ne reposent pas sur exactement les mêmes problèmes de design ni sur les mêmes dispositifs empiriques. La première étude mobilise plusieurs problèmes issus de la littérature sur le Design Thinking, reformulés puis utilisés pour comparer des listes verbales orientées. Elle relève avant tout d'une exploration de principe, visant à examiner la pertinence relative de différentes suggestions comme interférences à la fixation, plutôt qu'à établir une généralisation à partir d'un cas unique.

Les deuxième et troisième études reposent sur le même brief d'innovation des bibliothèques publiques. Elles ne comparent toutefois pas les mêmes conditions. La deuxième étude compare une interférence textuelle, une interférence image et une pause neutre à partir de la fluence. La troisième ajoute la mosaïque vidéo et complète la fluence par le FixShift. Le maintien de ce problème commun permet de comparer les interférences à problème constant et de réduire ainsi le risque de confusion entre le type de problème et le type d'interférence mobilisé. Dans les études 2 et 3, le protocole général conserve une logique similaire : une première phase d'idéation est menée jusqu'à l'essoufflement, puis une interférence ou une pause témoin est introduite avant une nouvelle phase d'idéation.

Dans la troisième étude, une condition témoin est ajoutée sous la forme d'une interruption neutre d'une minute, sans contenu sémantique. Cette condition permet de distinguer l'effet d'une simple pause de celui d'une interférence effectivement produite pour relancer l'idéation.

Le protocole commun ne vise donc pas d’abord une causalité universelle, mais la constitution d’une base comparative homogène entre conditions.

Il reste maintenant à préciser plus finement les participants, les cohortes et les sites dans lesquels ces protocoles ont été mis en œuvre.

2.6 Participants, cohortes et sites

Les participants diffèrent selon les trois études du programme doctoral.

Dans la première étude, la recherche porte sur les préférences exprimées par des praticiens et des participants à l’égard de listes de suggestions, dans une logique qualitative et directionnelle.

La deuxième étude repose sur des expérimentations menées à plusieurs dates et sur plusieurs sites, ce qui permet d’observer les interférences dans des contextes variés, même si les cohortes relèvent toutes d’un niveau de maîtrise. Les expérimentations ont été conduites auprès de quatre cohortes. AIT Group A a participé le 28 novembre 2024, à Bangkok et à distance. AIT Group B a participé le 30 novembre 2024, également à Bangkok et à distance. AIT VN a participé le 9 mars 2025, à Ho Chi Minh-Ville et à distance. HKUST a participé le 12 mars 2025 sur place. Dans cette deuxième étude, toutes les participations étaient volontaires et se déroulaient hors du temps d’enseignement.

Dans la troisième étude, le jeu de données comparatif comprend 138 participants et 983 lignes enregistrées. Il contient 502 pré-idées et 481 lignes post-interférence, dont 2 non-idées exclues du codage, soit 479 idées postérieures effectivement codées. Au niveau des participants retenus pour l’analyse comparative, les effectifs par modalité se répartissent comme suit : 41 pour la condition textuelle, 42 pour la mosaïque d’images, 25 pour la mosaïque vidéo et 22 pour la condition témoin correspondant à une interruption neutre d’une minute. Ces éléments permettent de situer la portée des comparaisons réalisées et d’explicitier la base empirique du programme doctoral.

Il convient maintenant de préciser la nature des données recueillies et les mesures utilisées pour les comparer.

2.7 Données collectées et mesures

Les données recueillies varient selon les trois études du programme doctoral. Dans la première étude, elles comprennent les définitions de problèmes, les listes de suggestions générées et les préférences exprimées par les participants au moyen d'un classement forcé. Dans la deuxième étude, la donnée principale correspond au nombre d'idées nouvelles produites après fixation, à la suite de l'exposition à une interférence textuelle, à une interférence image ou à une pause neutre. Dans la troisième étude, deux mesures complémentaires sont mobilisées : la fluence, qui correspond au nombre d'idées produites après l'interférence, et le FixShift, qui vise à rendre compte du déplacement conceptuel par rapport aux familles d'idées initiales.

La fluence est utilisée ici comme un indicateur de réactivation post-fixation, et non comme une mesure directe de la qualité des idées (Shah et al., 2003).

Le FixShift, pour sa part, ne constitue pas un score général de créativité. Il mesure le degré d'écart conceptuel entre les idées produites après l'interférence et les familles d'idées formulées avant celle-ci, propres à chaque participant. Cette logique de déplacement conceptuel entre familles de solutions s'inscrit dans une perspective où la défixation ne se réduit pas à produire davantage d'idées, mais suppose aussi un changement dans la structure des idées générées (Chan et al., 2011).

Le FixShift repose sur une grille simple à trois niveaux. Un score de 0 correspond au maintien dans la même famille d'idées ou dans la même logique centrale. Un score de 1 correspond à une variation significative à l'intérieur de la même famille. Un score de 2 correspond à l'émergence d'une nouvelle famille d'idées ou d'un nouveau chemin conceptuel. Cette gradation permet de distinguer la simple variation interne d'un déplacement plus net de l'exploration.

Le comptage des idées repose sur une règle volontairement conservatrice : une ligne non vide équivaut à une idée. Ce choix limite l'espace laissé à des interprétations variables, rend le traitement plus visible et renforce la reproductibilité de la démarche. Le traitement des absences et du bruit a également été explicité. Les lignes qui ne correspondent pas à des idées sont exclues du codage FixShift. De même, les réponses post-fixation laissées vides ne sont

pas traitées comme des idées manquantes, mais comme l'absence d'idée produite après l'interférence.

Enfin, pour la mesure de la fluence, les médianes accompagnées des quartiles [Q1, Q3] sont privilégiées comme principaux indicateurs de tendance centrale, plutôt que les moyennes. Ce choix s'explique par la forme des distributions observées, qui sont asymétriques et présentent des valeurs extrêmes.

La crédibilité de ces mesures repose donc moins sur un indicateur unique que sur l'ensemble des procédures d'analyse, de validation et de contrôle mises en place pour limiter les biais d'interprétation.

2.8 Procédures d'analyse et validation

Les procédures d'analyse et de validation diffèrent selon les trois études, tout en conservant une logique comparative commune.

Dans la première étude, le classement forcé a été retenu comme mode d'évaluation directionnel, puisqu'un score absolu aurait supposé une décomposition plus lourde et moins crédible des dimensions d'intérêt.

Dans la deuxième étude, l'analyse repose sur la comparaison du nombre d'idées nouvelles produites après exposition à une interférence textuelle, à une interférence image ou à une pause neutre. Ce choix est cohérent avec un protocole de terrain multi-sites, de faible effectif et à visée directionnelle.

Dans la troisième étude, l'analyse principale est descriptive et comparative plutôt qu'inférentielle et causale. Elle vise à faire émerger des tendances comparatives entre conditions plus qu'à établir des effets causaux universels (Bartlett & Vavrus, 2017).

Le dispositif de la troisième étude n'étant pas pleinement croisé entre les cohortes, les différences observées entre modalités ne peuvent être interprétées comme des effets causaux purs. Pour limiter le risque de pseudo-réplication, les résultats comparatifs sont rapportés au niveau du participant, même lorsque le FixShift est d'abord codé au niveau de chaque idée. Ce choix permet de préserver l'unité d'analyse correspondant à l'expérience réelle du participant, chaque personne n'étant exposée qu'à une seule modalité d'interférence.

Le codage du FixShift assisté par l'IA n'est pas traité comme une vérité de terrain, mais comme un indicateur de substitution nécessitant une validation humaine (Zhang et al., 2025). Cette validation repose sur un sous-échantillon aveugle de 60 idées post-interférence, équilibré à la fois selon les modalités et selon les niveaux de FixShift attribués par l'IA. Les trois évaluateurs ont utilisé la même grille (0,1,2), les mêmes règles de départage, ont travaillé de manière indépendante et n'ont pas eu recours à des outils d'IA. Les résultats de validation indiquent un accord exact de 60,0 %, un accord à une catégorie près de 96,7 %, une différence absolue moyenne de 0,44 par rapport à la moyenne humaine agrégée et un kappa pondéré quadratique de 0,61. Ces métriques doivent être interprétées comme un faisceau convergent d'indices, et non comme une mesure unique et décisive. Dans cette perspective, le coefficient kappa ne signifie pas que l'IA remplace le jugement humain, mais qu'elle fournit un indicateur suffisamment cohérent pour soutenir une lecture comparative directionnelle de l'ensemble du jeu de données.

Enfin, les distributions de comptage observées étant asymétriques et marquées par des valeurs extrêmes, les médianes et les quartiles sont privilégiés pour la mesure de la fluence, tandis que les moyennes et les écarts-types servent surtout à documenter des gradients directionnels entre conditions. La crédibilité de ces mesures repose donc moins sur un indicateur unique que sur l'ensemble des procédures d'analyse, de validation et de contrôle mises en place pour limiter les biais d'interprétation. Dans la première étude, le classement forcé ne constitue pas une mesure directe d'efficacité comportementale post-fixation. Il joue un rôle exploratoire de repérage directionnel des suggestions jugées les plus prometteuses, avant le passage à des comparaisons plus empiriques dans les études suivantes.

Au-delà de cette structure d'analyse, la démarche reste soumise à plusieurs limites transversales qui bornent la portée des résultats.

2.9 Limites méthodologiques transversales

Les projets de Design Thinking constituent généralement des singularités contextuelles, ce qui rend leur répétition à l'identique difficile. Il n'est, par exemple, pas possible de demander à un même designer de recommencer une séance sur un même énoncé de problème dans des

conditions véritablement équivalentes, puisqu'il n'aborde plus ce problème pour la première fois. Cette caractéristique limite la possibilité d'une réplique strictement expérimentale.

À cette première limite s'ajoute le caractère génératif des sorties produites par les grands modèles de langage et, plus largement, par l'intelligence artificielle générative. À consigne identique, les contenus obtenus ne sont pas parfaitement reproductibles, ce qui réduit encore la possibilité d'une reproduction stricte des conditions expérimentales. De plus, la comparaison entre modalités demeure structurellement bornée par des cellules inégales et par un plan qui n'est pas pleinement croisé entre les cohortes, en raison du caractère réel et situé des expérimentations menées.

Dans ce cadre, une cellule vidéo (Vim) exploratoire a été conservée dans le jeu de données élargi, mais exclue de la comparaison principale, parce que sa taille était trop faible pour permettre une interprétation stable. Il s'agissait d'une interférence vidéo unique, distincte de la mosaïque vidéo retenue dans l'analyse comparative principale. Ce choix vise à préserver la cohérence de l'interprétation plutôt qu'à maximiser artificiellement le nombre de conditions comparées.

Les mesures mobilisées présentent elles aussi une portée définie et limitée. La fluence ne constitue pas une mesure de la qualité des idées, mais un indicateur de reprise de la production après fixation. De son côté, le FixShift ne doit pas être interprété comme un score général de créativité, mais comme une mesure du déplacement conceptuel par rapport aux familles d'idées initiales. Les résultats doivent donc être lus comme des preuves comparatives directionnelles, et non comme un classement universel des modalités d'interférence.

Ces limites n'invalident pas la cohérence du programme doctoral. Elles en précisent au contraire la portée et permettent de situer plus exactement le rôle de chaque chapitre-article dans la démonstration d'ensemble.

2.10 Articulation des chapitres-articles

Les trois chapitres-articles jouent des rôles méthodologiques distincts dans le programme doctoral.

Le chapitre 3 correspond à une première étude exploratoire, à la fois qualitative et comparative, qui vise à tester la faisabilité de suggestions verbales orientées comme interférences à la fixation, ainsi que la pertinence d'une évaluation directionnelle par classement forcé.

Le chapitre 4 approfondit cette logique par une étude comparative de terrain portant sur une interférence textuelle spécifique au projet, une interférence image spécifique au projet et une pause neutre, à problème constant. Il permet ainsi d'évaluer, dans un même contexte d'innovation des bibliothèques publiques, si des interférences générées avec GenAI relancent davantage l'idéation qu'une interruption sans contenu sémantique.

Le chapitre 5 constitue le volet le plus structuré du programme doctoral. Il stabilise davantage le protocole, introduit une condition témoin, distingue deux mesures complémentaires, soit la fluence et le FixShift, et ajoute une procédure de validation du codage. Ce chapitre compare, dans un cadre commun, une interférence textuelle, une mosaïque d'images, une mosaïque vidéo et une interruption neutre.

Dans son ensemble, le programme doctoral progresse d'une comparaison qualitative de préférences vers une comparaison texte-image-pause neutre, puis vers une comparaison multimodale appuyée sur une rigueur méthodologique et statistique plus forte. Cette progression donne une cohérence d'ensemble aux trois articles et montre qu'ils ne constituent pas une simple juxtaposition de travaux, mais les étapes successives d'un même programme de recherche.

CHAPITRE 3

COMPARING DIFFERENT CHATGPT4.0 GENERATED ORIENTED WORD SUGGESTIONS AS DESIGN FIXATION INTERFERENCES FOR THE IDEATION PHASE IN DESIGN THINKING

Fabrice Fischer ¹, Mickael Gardoni ^{1,2}

¹ Département de génie des systèmes, École de technologie supérieure,
1100 Notre-Dame Ouest, Montréal, Québec, Canada H3C1K3

² INSA Strasbourg, Strasbourg, France

Article publié dans Sureephong P, Danjou C, Bouras A (eds) Product Lifecycle Management. Leveraging AI, Digital Twins, and Smart Technologies, July 2025
https://doi.org/10.1007/978-3-031-93323-3_16

Abstract. The purpose of this article is to present a comparative case study (CCS) methodology in the generation of lists of words suggestions as interference to the phenomenon of fixation of the design in the ideation phase of a Design Thinking (DT) process. The base line is a random list of suggestions from the Large Language Model (LLM) ChatGPT 4.0. While the more advanced first lists of suggestions are related to knowledge from discipline of the problem and the second list of suggestions is specifically not related to the discipline of the problem. This CCS is qualitative and use the Design Research Science Methodology (DSRM).

Keywords: Design thinking, Ideation, Fixation phenomenon, Word interference, Comparative case studies (CCS), Design research science methodology (DSRM), Large language model (LLM), ChatGPT 4.0, Knowledge limit, Generativity limit

3.1 Introduction

Innovation is at the heart of the discourses on value creation. Innovating could be slow or fast, incremental or disruptive, specific or generic, i.e., it is diverse. Contrary to popular belief that most innovations result from the one-human genius of a Leonardo da Vinci, the innovative designs often result from a structured and systematic innovation methodology conducted by a design team.

There are a few well-established innovations methodologies such as CK, TRIZ, ASIT, which are a bit complex to understand. There is also an apparently easy to comprehend innovation methodology with Design Thinking (DT).

DT is often described as a non-linear five-phase innovation methodology: 1. Empathy, 2. Definition, 3. Ideation, 4. Prototyping, 5. Test. While most individuals have some views on each phase simply based on their names, in practice, DT is complex to implement with a high degree of freedom in the choice of the DT activities and tools, requiring a fair level of expertise and experience to navigate effectively.

In DT, the ideation phase is the phase concerned with generating innovative ideas to the problem identified in phase “1. Empathy” and formalized in phase “2. Definition”. DT presents a fair level of challenges and difficulties, ranging from the design fixation phenomenon (Purcell & Gero, 1996), limitation of knowledge of the design team (Koh & De Lessio, 2018), limitation in experience of the design team, to limitation in generativity (Thomas & Tee, 2022), i.e., variations on the same themes but limited generation of truly new and diverse ideas, and more.

The design fixation phenomenon relates to the designers’ mind fixating on a given solution, thereby limiting their ability to ideate. Interferences contribute to the mitigation of the fixation of the phenomenon.

With the focus on words as interferences, instead of preset words on cards, or a generated words manually or randomly with generative artificial intelligence, could we try to mitigate other limitations of the ideation phase such as the limitation of knowledge, and, the limitation

in generativity? The limit of experience, while possibly not out of reach, seems different as it would require an injection of experience of some sort, and is not part of this study.

So, could we generate a list of words specific to DT projects, which might help to extend the knowledge effectively used by the design team, and, or, could these generated list of words contribute to increase the generativity of the design team while they ideate?

We could generate these lists of word manually or experiment with some digital automation tools. With the more recent advancements and generalization of generative artificial intelligence tools, could we generate a list of words of higher interest to the design team in their ideation effort? How could we generate such a list of words? Would it be of interest to the design team? How could we even qualify and evaluate the notion of “interest” or “value” to a design team, or more cautiously, to a single designer? Can we, as a minimal but sufficient approach, collect the practitioner’s preference for a list over another with regards to limitation of knowledge, and with regards to the limitation in generativity?

3.2 The problem

The problem we address is:

“For a design team at the time of the ideation phase of a DT project, using generative artificial intelligence, could we generate lists of words, as design fixation interferences of higher interest with regards to the limitations of knowledge and the limitation of generativity.”

Specifically, can we generate, with generative artificial intelligence, lists of words to use as design fixation interferences to a design team at the ideation phase of a DT project? Also, can we increase the list of words’ perceivable interest or value with regards to the limitation of knowledge and the limitation of generativity?

This problem comes with many other questions. What Methodology should we use? Which artificial intelligence could we use to generate this list? What do we mean with an interference of “higher interest”? How can we evaluate this “higher interest”?

3.3 Research objectives and hypotheses

The primary research question is:

Can we generate higher interest word suggestions as design fixation interferences for the ideation phase in Design Thinking using ChatGPT4.0?

The secondary research questions are:

How can we generate the lists of word suggestions?

How could we assess that some lists are more interesting than others with regards to the limitation of knowledge and the limitation in generativity?

The specific objective of this study is to:

- generate different lists of words as design fixation interferences in the ideation phase of a DT project.
 - Random words related to the DT project.
 - Words related to knowledge in the discipline of the DT project, with the intent to contribute to the mitigation of the knowledge limitation.
 - Words related to knowledge in a different discipline than the discipline of the DT project, with the intent to contribute to the mitigation of the knowledge limitation as well as the generativity limitation.
- assess if some lists are of relative higher interest to a designer.

The hypothesis are:

- We can generate a random list of words related to a DT project using ChatGPT4.0
- We can generate a more specific list of words related to knowledge related to the discipline of the DT project using ChatGPT4.0
- We can generate a more specific list of words related to knowledge not related to the discipline of the DT project using ChatGPT4.0
- We can assess if a list of words presents a higher interest to a design team versus another list of words by asking practitioners about their preferences.

With this research we expect to contribute to the corpus of knowledge, methodologies and tools available to the designer of a DT project, and possibly other designer facing the design fixation phenomenon.

The scope of this CCS is limited to a few examples while applying the DSRM for the generation of the word lists as artifacts. Its findings will, therefore, be limited to these few examples. It does not intend to offer a generalization of its findings, which may or may not be the topics of other research work.

While respecting our time and resources constraints, this scope is sufficient to illustrate our concepts of the generation of higher interest list of words with generative artificial intelligence while contributing to the mitigation of the limit of knowledge and the limitation of generativity.

3.4 Theoretical Background

This research is at the crossroad the following theoretical topics:

- Design Thinking and its ideation phase.
- Word interferences as a design fixation mitigation prop
- Design Science Research Methodology (DSRM) (De Sordi, 2021)
- Generative Artificial Intelligence (GAI) with its Large Language Models (LLM) to generate lists of words as artifacts of the DSRM
- Comparative Case Studies (CCS) (Kröger, 2021)

3.4.1 Design Thinking

The terms “Design Thinking” are related to innovating and generally encompass different but related dimensions ranging from a concept, an approach, a methodology, activities, tools, up to the project of researching an innovative design solution to a human centric problem or opportunity leveraging all the other dimensions. The DT methodology DT is often described

as a non-linear five-phase innovation methodology: 1. Empathy: where the designers empathize with the user to better understand the problem (or the opportunity) to be addressed with the design solution to innovate.

2. Definition: where the problem or opportunity is further clarified down to a clear formulation of the definition of the problem to address

3. Ideation: where the designers ideate potential design as solutions to the problem and then select the solutions of higher interest

4. Prototyping: where prototypes are being made in preparation for the tests on some characteristics of or the complete products/services

5. Test: the prototypes are tested on some characteristics of or the complete

These phases are non-linear and the approach includes iterating fast between the phases to learn fast and come up with a better design while remaining human centric.

3.4.2 Word interferences to mitigate the design fixation phenomenon

As far as the design fixation is concerned, the phenomenon is that design teams' and individual designers' minds tend to rapidly fixate on a design solution, being a suggested solution or their own idea (Leahy et al., 2020). It is a problem because this fixation phenomenon limits their ability to generate additional design solutions ideas. And previous research has shown that to get to the best final design, a DT team needs to generate many ideas, with the higher number of ideas being positively correlated with the best final design solution.

And, "... my phone just rang, and I lost track of my thoughts ... What? Oh, yes ... " interferences have been proven to help mitigate the fixation phenomenon. All kinds of interferences help, from a pause in the ideation work, to the introduction of words, or the introduction of objects. Brian Eno (Eno & Schmidt, 1975) , with his oblique strategies, designed a box with a collection of cards with words to be used as design prompts, and are effective as fixation interferences.

Now, could some interferences appear to be more potent than others in relaunching the ideation flow of the design teams? A few-minute pause helps a design team that ran out of ideas, but

less so compared to the introduction of words via the Brian Eno cards, or “How would an elephant or a ghost solve the problem, or an object. So, there should be a way to, if not maximize, but at least, enhance the potency of the interferences.

3.4.3 Design Science Research Methodology (DSRM)

The Design Science Research Methodology (DSRM) is concerned with the research of innovative solutions. It includes the development of artifacts with the intent to solve an identified problem.

The artifacts could be constructs, methods, models, and instantiations.

3.4.4 Generative Artificial Intelligence (GAI) with its Large Language Models (LLM) to generate lists of words as artifacts of the DSRM

Generative Artificial Intelligence (GAI) with its Large Language Models (LLM) is quickly becoming mainstream. Based on text or files entered in a prompt, these systems can generate text, images, videos and sounds. ChatGPT 4.0 is one of such LLM systems. It can be used to generate lists of words based on some relevant prompts. We could consider using ChatGPT as a method, and its answers as instantiations of the DSRM.

3.4.5 Comparative Case Studies (CCS)

Comparative Case Studies (CCS) come under various forms. One classic such form consists in establishing a case baseline to compare to, and one or many alternative cases to compare to the baseline.

3.5 Methodology

For this study of comparing potential list of suggestions, the Design Science Research Methodology (DSRM) is used for the creation of the list of suggestions, while the Comparative Case Study (Mello, 2021) methodology is used to compare them.

The Design Science Research Methodology (DSRM) (Johannesson & Perjons, 2021) is particularly pertinent in the context of the creation of artifacts to address a problem. Also, DT projects generally come as singularities with each project having its own context and specificities. DT projects tend to be unique. While some institutions have portfolios of DT project, they are usually only a small number of projects in the low tens and almost never in the hundreds in any given one institution or accessible to a given research team. In that context, the research on DT projects tends to be more qualitative than quantitative. DSRM is well suited for qualitative research. This research is subjected to the same constraints as the other DT project-related research and will be qualitative in nature.

Furthermore, the interest of this research is to assess the “higher interest” nature of some suggestions. The relative adjective “higher” comes with a comparison. The Comparative Case Studies (CCS) methodology is meant to compare one or more alternatives versus a baseline.

The combination of DSRM with CCS offers the combination of the creation of artifacts to address a problem and the comparison of an alternative versus a baseline in the context of qualitative research, which suits well with the research objective and constraints of this research.

The design of the research is to generate three lists of suggestions, a random list as a baseline (A), a list of words related to knowledge in the discipline of the DT project (B), with the intent to contribute to the mitigation of the knowledge limitation, a list of words related to knowledge in a different discipline than the discipline of the DT project (C), and to compare then in terms of a designer interest or value at the time of the ideation phase of a DT project.

The research phases are as follows:

- Define the elements of the research and the context of the comparison: list of suggestions, generate, comparative interest, artifacts.
- Generate the lists of words using ChatGPT4.0 A, B, C.
- Compare the lists and draw learnings applicable to the samples compared.
- Share the findings (this paper).

3.6 The creation of the artifacts

3.6.1 Artifact method

The method is to use ChatGPT to generate lists of ten suggestions composed of three-to-five-words.

3.6.1.1 Generating the list of ten random three-to-five-words related to the DT project

ChatGPT4.0 prompt:

Generate a list of 10 random 3 to 5 related words.

Answer:

xxx

3.6.1.2 Generating a list of ten three-to-five-words of knowledge related to the discipline of the DT project

ChatGPT4.0 prompt:

With the following definition of the problem: XXX

Generate a list of 10 of 3 to 5 words presenting a knowledge related to the discipline of the problem.

With XXX to replace with the definition of the problem

3.6.1.3 Generating a list of 3 to 5 words of knowledge not related to the discipline of the DT project.

ChatGPT4.0 prompt:

With the following definition of the problem: XXX

Generate a list of 10 of 3 to 5 words presenting a knowledge not related to the discipline of the problem.

With XXX to replace with the definition of the problem.

3.6.2 Artifact Instantiation: Identification of examples of DT problems

To identify examples DT problems, we search for DT papers and used identically or similarly worded but shortened versions of the identified problems.

Here are the five problems identified:

Problem #1: How can emergency departments redesign their processes to efficiently manage patient flow, reduce wait times, and improve satisfaction within existing resource constraints?

Problem #2: What innovative pedagogical strategies and technologies can higher education institutions implement to significantly increase student engagement and satisfaction?

Problem #3: How can affordable, sustainable lighting solutions be developed and made accessible to low-income families using local resources and renewable energy?

(Brown & Wyatt, 2009)

Problem #4: How can public libraries innovate their services and spaces to meet evolving community needs and increase user engagement in the digital age?

Problem #5: What innovative, economically and environmentally sustainable water management solutions can be tailored to meet the specific needs of rural communities?

3.6.3 Data & sources

For this experiment we are using the following data:

1. Examples of DT problem definitions: a few examples of known DT problems.

The sources for these examples are academic paper.

2. lists of word suggestions:

- a. a list of ten random three-to-five-words related to the DT project.
- b. a list of ten three-to-five-words of knowledge related to the discipline of the DT project.
- c. a list of ten three-to-five-words of knowledge did not relate to the discipline of the DT project.

These lists are the output of respective prompts, one per list, using ChatGPT 4.0

3. Questionnaire to the practitioners

The authors of this paper prepared the questionnaire.

4. Practitioners' answers to the Questionnaire

The practitioners will fill the questionnaire with their answer and send the filled questionnaire to the author of this paper.

5. Synthesis of the practitioners' answers

The filled questionnaires are the data source for this synthesis.

3.6.4 Data preparation and collection methods

For this experiment we are using the following data:

- List of DT problem definition: we searched into DT academic research papers, for preferably well-cited papers with one DT problem example each. We then summarized and transposed the mentioned or illustrated DT problem in respective problem definition in our own words.
- Lists of word suggestions, we logged into our ChatGPT account, selected the version ChatGPT 4.0 and created a prompt for each list. We then removed the ChatGPT numbering and formatted the list to fit into our Questionnaire tables.

For the practitioners' preferences, we are using the following.

- The questionnaire is prepared by the authors of this research.
- The answers to the questionnaire will come from practitioners.

3.7 Evaluation

We can use the “forced ranking” evaluation method. A scale to measure or evaluate the solution would be great but would require a thorough clarification of the dimensions and a form a grading on each of these dimension. For this study, we will stop at a forced ranking by practitioners to get a directional evaluation of the effectiveness of the artifacts without getting into the evaluation of the intensity of that effectiveness.

3.8 Discussion

This paper presents a methodological approach to navigate qualitatively and directionally the use of GAI to generate higher interest word suggestions as design fixation interferences in a context of a collection of unique occurrences of DT projects.

3.8.1 Theoretical and practical contributions

The combination of DSRM with CCS in a GAI ADT is of interest to DT. While the formulation through this paper might feel a bit heavy to set up and explain, once the methodological framework is in place, its practical use is relatively fast and aligned with the DT bias and appetite for fast learning iteration cycles of ideation-prototype-test.

The DSRM methodology is strong with the contribution of the artifacts as key components to figure-out potentially solutions, which might lead to knowledge. The CCS, with its baseline and a forced ranking of alternatives by the practitioners, allows for a directional intuitive preference in a mostly qualitative context of unique occurrences of DT projects where the contribution to the mitigation of knowledge limitation and generative limitation are difficult to measure. It is practical by construct.

3.8.2 Limitations and future research

To start with, this study is only qualitative. While we would like to believe that more practitioners would answer the questionnaire in a similar way, it can't be confirmed until a quantitative study is conducted to deliver some statistically relevant findings.

Then, the request on our questionnaire is a forced ranking. While it points directionally to a preference, it does not identify the underlying components of that preference, and it does not provide a clear evaluation grading in each of these sub-components.

3.9 Conclusion

This research contributes to the body of knowledge related the words as fixation phenomenon interferences of higher interest related to the mitigation of the knowledge limitation and the generativity limitations. The methodology, combining the DSRM with its method artifact, and

the CCS to establish the DT practitioners' directional preferences are relevant to the qualitative nature DT projects singularities.

The research on using GAI tools to augment the DT process is in its infancy. Here are some research directions of interest. We expect that far more, to all, information from the earlier phases 1. Empathy and 2. Definition, will be consumed into building a better understanding of the problem to inform the GAI generating the list of word interference. Researching and better understand that contextual information within the DT process augmented with GAI could lead to recommendation of even higher interest. As far as the interferences format is concerned, with the fast progression of the sound, image, and video generation, research on these GAI interference formats comparative effectiveness might further enhance the effectiveness of the interferences. And while it might be possible to generate quality sound and visual interferences, would it be or not at the expense of a higher fixation anchoring or a lower generativity? As for the word interferences, the three-to-five-word suggestions length of this paper is restrictive to represent knowledge. What would be the most effective qualitative and quantitative GAI word interference? Is there such a thing as a specific length, or could it be adaptative to the problem concerned and specific to the knowledge introduced? At this point, we know so little about the potential uses and impacts on this fast-evolving Augmented Design Thinking (ADT).

CHAPITRE 4

GENAI AS A PROJECT-SPECIFIC INTERFERENCE GENERATOR FOR IDEATION DEFIXATION

Fabrice Fischer ^{1*}, Mickael Gardoni ^{1,2}

¹ École de technologie supérieure,
1100 Notre-Dame Ouest, Montréal, Québec, Canada H3C1K3
² INSA Strasbourg, Strasbourg, France

Article soumis pour publication, juin 2026

Abstract

Innovation teams rarely struggle to produce a first idea. They more often struggle when their first ideas become too comfortable and the team experiences fixation after an initial burst of ideas. Rather than asking GenAI to produce the ideas, this study gives it the smaller and more managerial role of generating project-specific interferences that help humans restart ideation after fixation. Using a library innovation brief, participants first generated ideas until self-reported fixation and were then exposed to one of three interferences: GenAI-generated text interference, GenAI-generated image interference, or a neutral break. The study evaluates whether participants generate again after fixation by comparing pre- and post-interference idea counts across text, image, and neutral-break groups. By comparing text and image modalities and the neutral-break control group, the study shows how GenAI can support innovation facilitation without becoming the ideator in the room. The findings are interpreted cautiously given the single brief, student sample, and fluency-only outcome, but they provide a practical basis for using GenAI as a defixation and ideation-relaunch tool in innovation workshops and design-thinking contexts.

Keywords: generative AI, creativity management, innovation facilitation, ideation fluency, human-AI creativity, creative stagnation, design thinking

4.1 Introduction

4.1.1 Ideation fixation in innovation work

Innovation projects often fail not because teams lack a first idea, but because they converge around it too quickly. This is especially visible in the front end of innovation, in design thinking workshops, and in early service or product development activities. At that stage, selection is premature. The first job is to keep the solution space open long enough to make selection meaningful.

In practice, however, ideation has a rhythm. The first ideas arrive quickly, then the room starts circling the same territory. Participants reformulate familiar concepts, make minor variations, or report that they have “run out” of ideas. In design research, this situation is commonly associated with fixation, understood as adherence to a limited set of ideas (Jansson & Smith, 1991). Fixation is not limited to one design context. It is a broader cognitive constraint affecting the generation of alternatives (Purcell & Gero, 1996).

For creativity and innovation management, fixation creates a practical problem. Innovation methods invite divergence, but participants remain bounded by their prior knowledge, the framing of the problem, the examples they saw, and the first ideas they produced. This is where defixation becomes a facilitation problem, not only a cognitive concept. It requires a perturbation. The interference must be strong enough to disturb the current path, but not so directive that it becomes the next path. The useful zone lies in between. The interference must be close enough to the problem to matter and distant enough from the first idea family to reopen search.

This study focuses on what happens after participants have produced an initial set of ideas and report that they have reached fixation. The outcome is fluency, defined as the number of ideas generated after the interference. Fluency is not the same as originality, usefulness, or final innovation quality. Still, it is an important first indicator. Before an idea can be evaluated, selected, refined, or implemented, it must first be generated.

4.1.2 GenAI as a facilitator, not ideator

GenAI is now entering the innovation workshop with a very ambitious job description. It is often asked to produce ideas. Large language models and text-to-image systems can generate suggestions, concepts, images, narratives, scenarios, and design alternatives almost instantly. This creates an obvious temptation. If a machine can produce ideas instantly, why not invite it to become the ideator?

This direct use is useful in some situations, but it also creates problems. More output does not necessarily mean broader exploration. AI-generated ideas may be plausible, fluent, and well written, while still remaining close to common patterns. Recent research points toward the same paradox. GenAI can increase the quantity of ideas, but it can also reduce originality, encourage convergence, or create overreliance on machine-generated suggestions (Bellesia, Bellis, & Palombi, 2026 ; Cristofaro & Giardino, 2025 ; Doshi & Hauser, 2024). In design contexts, the AI output can itself become a new source of fixation.

This study adopts a more bounded and facilitation-oriented position. GenAI is not treated as the ideator. The participant remains the ideator. GenAI is used before the second ideation phase to generate project-specific interferences. These interferences are not proposed as solutions. They are not meant to be copied, selected, or evaluated as design ideas. Their role is to disturb the participant's current line of thinking so that human ideation can restart after fixation.

This distinction is central to the paper. The question is not whether AI can generate better ideas than humans. The question is whether GenAI can be used as a facilitator of human creativity. The system produces an interference while the human remains responsible for producing the ideas. This framing is consistent with creativity and innovation management, where the objective is not simply automation, but the design of conditions that help individuals and teams produce, explore, and develop new possibilities.

4.1.3 Project-specific interference for ideation launch

The idea of disturbing fixation is older than GenAI. Designers and innovation facilitators have long used cards, analogies, examples, images, objects, patents, provocations, or breaks to

interrupt habitual thinking. Brian Eno's *Oblique Strategies* (Eno & Schmidt, 1975), for example, are a familiar set of short prompts intended to make the creator step away from the obvious track. The common mechanism is simple. Insert something from outside the current idea family and see whether the search opens again.

Many existing interferences, however, are generic. They are not generated for the specific innovation problem at hand. A generic card may be provocative, but too distant from the project. A patent may be useful in engineering design, but less useful in service innovation or design thinking workshops. A simple pause may help, but it does not introduce new semantic or visual material. The practical challenge is to create interferences that are specific enough to connect with the project, but indirect enough not to prescribe a solution.

GenAI changes the economics of interference. Because GenAI systems can generate text and images from a project brief, facilitators can create interferences on demand. A facilitator can take the problem statement, ask the system to generate short verbal cues or visual material, and introduce those interferences when participants report fixation. The interference becomes project-specific, scalable, inexpensive, and available at the moment when the room needs it.

This study compares two GenAI-generated interference modalities, text and image, with a neutral break. Text is simple, fast, and easy to integrate in most innovation workshops. Image is richer and may activate associations that text does not. At the same time, richer media are not automatically better. Images may attract attention, increase cognitive load, or anchor participants around visual features. For this reason, the relative effect of text and image should be tested rather than assumed. This is the practical promise tested here. It is not about more AI in the workshop, but about better-timed AI.

4.1.4 Study context and research question

The empirical task is based on a public-library innovation brief:

“How can public libraries innovate their services and spaces to meet evolving community needs and increase user engagement in the digital age?”

This brief was selected because it is accessible without specialized technical expertise and broad enough to allow several possible directions, including digital services, community

engagement, physical spaces, education, inclusion, events, technology, partnerships, and user experience. This makes it suitable for studying ideation fixation and defixation in an innovation context.

Participants first generated ideas in response to the brief until they reported that they could no longer generate additional ideas. This self-reported point was used as the operational fixation moment. Participants were then exposed to one of three interference types, namely the text modality, the image modality, or a neutral break, and were asked to generate additional ideas. The study is guided by one research question.

RQ. How do project-specific GenAI-generated text and image interferences affect post-fixation ideation fluency compared with a neutral break?

Two sub-questions follow.

RQ1. Do project-specific GenAI-generated interferences support stronger post-fixation ideation fluency than a neutral break?

RQ2. Do text and image interferences differ in their ability to relaunch ideation after fixation?

The purpose is not to prove that one modality is universally superior. The purpose is to provide comparative evidence on whether and how different project-specific interferences can relaunch ideation after fixation in an innovation task.

4.1.5 Contributions

This study makes three contributions.

First, it contributes to creativity and innovation management by treating defixation as a facilitation problem. Rather than asking whether GenAI can replace human ideation, the study examines whether GenAI can support the conditions under which human ideation continues. This shifts the role of GenAI from idea producer to interference generator.

Second, it contributes to human-AI creativity research by distinguishing between GenAI-generated ideas and GenAI-generated interferences. GenAI-generated ideas may create overreliance or new fixation. GenAI-generated interferences are designed to perturb thinking while keeping the human participant responsible for idea generation.

Third, it contributes practical evidence for innovation facilitators. It compares text interference, image interference, and a neutral break under a common innovation brief. This gives facilitators a more grounded basis for deciding whether and how to introduce GenAI-generated material during ideation workshops. The method is low-cost, scalable, and project-specific, which makes it relevant for design thinking courses, innovation labs, and early-stage innovation projects.

Taken together, these contributions are framed through a methodologically cautious posture that privileges parsimonious interpretation, limited extrapolation, and conservative assumptions. The study does not claim that GenAI is a better ideator than humans, or that text is universally better than image. It asks a focused question with practical consequences. When ideation stalls, can a project-specific GenAI interference help people generate again?

This paper also extends prior work on LLM-generated word interferences for design fixation (Fischer & Gardoni, 2025a, 2025b). Earlier studies established the feasibility of generating problem-specific word interferences and provided preliminary evidence that related word interferences can relaunch ideation after fixation. The present study moves beyond that first stage by comparing GenAI-generated text interference, GenAI-generated image interference, and a neutral break using post-fixation fluency as the main outcome.

4.2 Literature Review

4.2.1 Ideation fixation and defixation in creativity and innovation management

Creativity and innovation management is not only concerned with where ideas come from. It is also concerned with how individuals and teams keep ideation open before evaluation and convergence take over. This is the context in which design thinking is useful for this study. Design thinking has been discussed in CIM not only as a design method, but also as an enacted innovation practice. Carlgren, Rauth, and Elmquist's "Framing Design Thinking: The Concept in Idea and Enactment" (Carlgren, Rauth, & Elmquist, 2016) is useful here because it shows that design thinking is not only a formal process, but a set of practices that are enacted differently across contexts. Mayer and Schwemmler's "The impact of design thinking and its

underlying theoretical mechanisms: A review of the literature” (Mayer & Schwemmler, 2025) further helps to position design thinking as a research object for management, innovation, education, and organizational problem solving, rather than as a purely design-engineering method.

However, design thinking and other innovation methods do not remove cognitive constraints. One of these constraints is fixation. In design research, fixation has been described as an adherence to a limited set of ideas (Jansson & Smith, 1991). It is not restricted to one type of design task, but reflects a broader difficulty in moving away from earlier examples, familiar principles, or first-generated ideas (Purcell & Gero, 1996). Expert designers also recognize fixation as a recurring issue in concept development, shaped by project conditions, design activities, prior experience, and organizational context (Crilly, 2015). For creativity and innovation management, the point is straightforward: innovation methods may invite divergence, but participants can still remain cognitively attached to a narrow part of the solution space.

This study keeps the terms fixation and defixation because they name the mechanism clearly. At the same time, the argument is not limited to engineering design. Creative work is not only an input-output process, it involves cognitive operations that structure what people notice, combine, transform, or ignore (Pinkow, 2023). AI-augmented creative problem solving also requires attention not only to idea production, but to how humans monitor and control their creative cognition (Hoßbach & Isaksen, 2025).

Defixation is the complementary process. It refers to an attempt to move participants away from dominant idea families and to reopen ideation. Prior design and engineering research has studied several defixation strategies, including fixation classification and prevention methods, patent-based distraction, provided and self-generated examples, representational variation, and design-independent incubation interventions (Atilola & Linsey, 2015 ; Koh & De Lessio, 2018 ; Leahy et al., 2020 ; Wu, Zhou, & Zhi, 2023 ; Youmans & Arciszewski, 2014). These references remain important because they define the fixation/defixation mechanism. However, they should not become the only theoretical center of the article. The CIM question is therefore

practical and managerial: when the room gets stuck, what can a facilitator introduce that helps participants generate again without taking the ideator role away from them?

This study uses post-fixation fluency as the first measure of defixation. Fluency is the number of ideas generated after the interference. It is not a complete measure of creativity. It does not establish originality, usefulness, feasibility, or final innovation value. Still, it is a relevant first-order outcome in a defixation study because the first question is whether ideation restarts at all. Before ideas can be selected, improved, combined, or evaluated, participants must produce additional ideas. This study therefore treats fluency as an indicator of ideation relaunch, not as a substitute for full creativity evaluation.

This choice of fluency reflects a parsimonious measurement strategy. It does not pretend that more ideas are automatically better ideas. But in a defixation study, the first test is more basic. After participants say they have run out of ideas, do they generate again?

4.2.2 GenAI as creativity facilitator rather than ideator

GenAI has changed the practical conversation about ideation because it makes idea production cheap, fast, and almost too easy. In product innovation, ideation is positioned as the starting point of the innovation process, with AI increasingly involved across opportunity identification and analysis, idea generation, and idea screening and selection (Pescher & Tellis, 2025). Their review is useful for the present study because it places AI-assisted ideation inside innovation management rather than treating it only as a design or HCI issue. Their conclusions are also cautious: AI can improve the speed, efficiency, and cost of ideation, and it appears to increase the average creativity of generated ideas, but evidence remains conflicting on whether AI improves the generation of top ideas. They also distinguish between idea screening, where AI already appears useful, and idea selection, where human judgment remains important.

This distinction matters for the present research. A first way to use GenAI is to treat it as an ideator. The human writes a prompt, the system generates ideas, and the human selects among them. This approach has practical value, but it changes the creative process in ways that must be managed. GenAI can increase speed and quantity, but more output does not automatically mean broader exploration or better ideas. The objective of the present study is therefore

narrower and different: we do not test whether GenAI can produce ideas for participants. We test whether GenAI can generate project-specific interferences that help participants relaunch their own ideation after fixation.

Recent CIM research provides the conceptual frame for this facilitation view. AI and creativity can be understood as a multilevel phenomenon involving individuals, groups, and organizations (Grilli & Pedota, 2024). Technology can complement creative work not only by improving knowledge search, storage, or transformation, but also by extending the creative domain itself through new symbolic and material elements that workers can recombine in creative tasks (Pedota & Piscitello, 2022). In a GenAI-specific extension of this logic, hybrid creative processes depend on how human agents select, interpret, and use AI-generated possibilities (Pedota, Cicala, & Basti, 2025). This is relevant to the present study because GenAI-generated interferences can be understood as temporary additions to the participant's ideation domain. In the front end of innovation, GenAI can act as a nudge, sparring partner, accelerator, or collaboration enhancer, but its use may also introduce bias, creativity leveling, or reduced human innovation experience (Bellesia et al., 2026).

These CIM publications point toward a facilitation view of GenAI. GenAI is not necessarily most useful when it provides the final idea. It may also be useful when it changes the conditions under which humans generate ideas. This facilitation view is also consistent with work on perspective switching and metacognitive control in AI-augmented creative problem solving (Bordas, Le Masson, & Weil, 2025 ; Hoßbach & Isaksen, 2025). In the present study, GenAI is therefore not treated as the ideator. It is treated as a generator of interferences. The participant remains the ideator.

This distinction also responds to a growing concern in the GenAI creativity literature: AI support may have nonlinear effects. Experimental evidence with product innovators suggests that AI support can have nonlinear effects across no-AI, moderate-AI, and high-AI conditions (Cristofaro & Giardino, 2025). Their key finding is not that “more AI is better.” Instead, they report an inverted-U relationship: moderate AI assistance produced the highest cognitive engagement and the strongest combination of idea quantity, originality, and feasibility, while high AI assistance increased risks of automation bias, reduced originality, and increased

overconfidence. This finding is directly relevant to the present research. It supports a bounded use of GenAI in ideation, where AI is configured to support human cognition without displacing it. This is close to the managerial intuition behind the present study. The problem is not whether AI should be absent or omnipresent. The problem is calibration. How much AI, at what moment, and in what role?

This study adopts that bounded view. GenAI is used before the second ideation phase to generate project-specific interferences. Participants do not evaluate AI ideas, copy AI ideas, or select AI-generated solutions. They are exposed to an interference and then asked to generate their own ideas. This is consistent with a formal model of hybrid creative process that emphasizes the human–GenAI dyad and the active role of human agents in selecting from AI-generated possibilities (Pedota et al., 2025). In the present study, however, the selected AI-generated possibility is not a final idea. It is an interference used to perturb fixation while leaving the human participant in the ideator role.

4.2.3 Project-specific interferences for defixation

Interference is a practical idea before it is a GenAI idea. It rests on the assumption that a designer who is stuck may not need a complete solution. The designer may need a displacement. This can be a word, image, object, example, analogy, or break that shifts attention away from the dominant idea family. Interferences can take the form of words, images, objects, examples, analogies, patents, breaks, or provocations. Brian Eno’s Oblique Strategies are a familiar example of short text provocations intended to disturb habitual thinking. In engineering design, patent documents have been studied as both fixation sources and distractions, while fixation can also arise from provided and self-generated initial examples (Koh & De Lessio, 2018 ; Leahy et al., 2020). Classification and prevention methods have been proposed, and different representations, including computer-aided design, photographs, and sketches, can influence fixation and creativity (Atilola & Linsey, 2015 ; Youmans & Arciszewski, 2014).

Our prior work (Fischer & Gardoni, 2025a) adapted this interference logic to LLM-generated word suggestions. In one study, we proposed a DSRM and comparative case study approach

to generate three types of ChatGPT 4.0 word-interference lists: random words related to the design-thinking project, words related to knowledge in the discipline of the project, and words related to knowledge outside the discipline of the project. The study focused on whether such lists could be generated and whether practitioners perceived some lists as more interesting for addressing knowledge and generativity limits during ideation.

In a subsequent empirical version (Fischer & Gardoni, 2025b) using the public-library innovation brief, we compared the number of post-fixation ideas generated after exposure to the three word-interference lists. That work reported average post-fixation idea counts of 4.60 for random problem-related suggestions, 4.67 for discipline-related knowledge suggestions, and 4.17 for non-discipline-related knowledge suggestions, with the caveat that the valid sample sizes were small. The present study therefore does not repeat the earlier word-list comparison. It uses that work as a foundation for testing whether interference modality and the presence of a neutral control group change the interpretation of GenAI-supported defixation.

These references show that what participants are exposed to matters. However, many traditional interferences are generic, difficult to prepare, or tied to a specific domain. A generic card may be provocative, but it is not generated for the project at hand. A patent may be useful in technical design, but less useful in a service innovation or public-library innovation task. A neutral break may help by interrupting the task, but it does not add project-related material. This creates the opportunity for GenAI-generated interferences. They can be created on demand from the project brief and varied by modality while remaining connected to the same problem.

Project-specificity is central to the present study. A project-specific interference is not a solution. It is tied to the brief, but it should not prescribe what the participant should write. If the interference is too close to a solution, it may create a new fixation. If it is too far from the brief, it may become irrelevant. Our earlier word-interference studies explored this balance by generating lists that were related to the problem, related to the discipline of the problem, or deliberately outside the discipline of the problem. The current study keeps the project-specific logic but simplifies the empirical comparison: instead of comparing three types of word lists, it compares text interference, image interference, and a neutral break.

Text and image interferences may affect defixation through different channels. Text interferences provide verbal cues that can support semantic reframing, category shifts, or new associations. Image interferences provide visual material that may activate different associations, spatial interpretations, or metaphorical links. Generative AI model type and visual stimulus type can interact in shaping design creativity (Ge & Hou, 2025). Their experiment suggests that T2I models can enhance creative performance, but also that AI-generated visual support can increase repetition and fixation risk under some conditions. This is the precise point needed for the present study. Modality matters, but richer or more visual material should not be assumed to be better. Text and image interferences should be compared empirically.

The comparison with a neutral break is also necessary because a break is not a straw man. It may reduce immediate cognitive pressure and may itself help defixation. Without a neutral break, any additional ideas after the interference could simply reflect more time, not better interference. By comparing text interference, image interference, and a neutral break, the present study asks whether project-specific GenAI material adds something beyond interruption alone. That is a deliberately focused question, and a useful one for facilitators.

4.2.4 From prior evidence to the present study

The literature supports four points. First, fixation is a persistent constraint in design and innovation work. It limits the ability of participants to move beyond first ideas, examples, familiar knowledge, or dominant problem framings. Second, defixation can be supported by interferences that perturb attention, association, or problem representation. Third, GenAI can support creativity and innovation, but it can also produce bias, convergence, creativity leveling, overreliance, or reduced diversity, especially when AI-generated outputs become common reference points across participants or teams (Bellesia et al., 2026 ; Cristofaro & Giardino, 2025 ; Doshi & Hauser, 2024). Fourth, text and image interferences may work differently, but the literature does not yet provide enough comparative evidence on project-specific GenAI-generated interferences introduced after fixation and compared with a neutral break.

This is the gap addressed in the present study. We study GenAI not as a direct ideator, but as a facilitator that generates project-specific interferences. We compare text interference, image interference, and a neutral break in the same library innovation task. The dependent variable is post-fixation ideation fluency. This narrow outcome is deliberate. It allows us to test whether ideation is relaunched after fixation without claiming to measure the full creative value of the ideas. Further measures of originality, usefulness, or conceptual distance can extend this work, but the first question is whether participants start generating additional ideas again.

The study is also informed by design science logic because it develops and evaluates an interference, and by comparative case study logic because it compares different interference modalities under a common task structure. These methodological references belong mainly in the Method section. They help explain how the interference is built and compared. They are not the theoretical center of the research. The theoretical center is the CIM question. How can GenAI be configured as a project-specific facilitation mechanism for defixation in innovation work?

4.2.5 Expected relationships

Based on this literature, the study expects project-specific GenAI interferences to support post-fixation ideation fluency more strongly than a neutral break. The logic is direct. A neutral break interrupts the task, but it does not introduce new project-related material. A GenAI-generated interference both interrupts the fixation state and introduces brief verbal or visual material tied to the project. This may help participants activate associations that were not present in their first idea set.

The study also expects text and image interferences to differ in their effect on post-fixation fluency. Text and image are not interchangeable. Text may support semantic reframing and category movement. Image may support visual association and alternative interpretation. However, richer media should not be assumed to be more effective. Images may also attract attention to surface features, increase cognitive load, or create another form of anchoring. For this reason, the study compares the two modalities empirically rather than assuming that one is superior.

This leads to the following research questions:

RQ1. Do project-specific GenAI-generated interferences support stronger post-fixation ideation fluency than a neutral break?

RQ2. Do text and image interferences differ in their ability to relaunch ideation after fixation?

4.3 Methodology

4.3.1 Prior studies and evolution of the protocol

This study builds on two earlier stages of research on LLM-generated word interferences for design fixation. The first stage examined whether ChatGPT 4.0 could generate different types of project-related word suggestions to serve as design fixation interferences. It compared random project-related words, discipline-related knowledge words, and non-discipline-related knowledge words, and evaluated practitioner preferences regarding their perceived usefulness for addressing knowledge and generativity limits (Fischer & Gardoni, 2025a). The second stage applied these word-interference lists to a public-library innovation brief and compared post-fixation idea counts across the three list types (Fischer & Gardoni, 2025b). That study showed that LLM-generated word interferences could relaunch ideation, but it remained limited by small valid samples and the absence of a neutral control group.

The present study is designed as the next step in that research program. It retains the general logic of project-specific GenAI-generated interference after fixation, but changes the comparison in three ways. First, it compares interference modality rather than word-list type. Second, it adds a neutral-break control group to distinguish the effect of project-specific GenAI interference from the effect of interruption alone. Third, it uses post-fixation fluency as the main comparable outcome across text, image, and control group. The present study is therefore a modality and control-group comparison, not a repetition of the earlier word-list studies.

4.3.2 Research design

The study uses a controlled comparative design informed by design science logic and comparative case study logic. The design science component appears in the construction of the project-specific interferences: GenAI is used to generate text and image materials intended to perturb participants' thinking after fixation. The comparative component appears in the comparison of three post-fixation groups under the same ideation task: text interference, image interference, and neutral break.

The unit of analysis is the individual participant. Each participant first generates ideas in response to the same innovation brief, stops when reaching self-reported fixation, receives one assigned group, and then generates additional ideas. The design is not positioned as a full randomized laboratory experiment. It is a controlled comparative study conducted across applied educational settings, with a common task, a common procedure, and group-level comparison. This framing is consistent with the practical constraints of design thinking and innovation workshop research, where repeated exposure to the same brief can contaminate subsequent ideation. Group assignment followed the structure of the data collection sessions rather than a fully randomized laboratory allocation. This is why the study is presented as a controlled comparative study rather than as a randomized experiment.

4.3.3 Empirical context and ideation task

The empirical task used the following public-library innovation brief:

“How can public libraries innovate their services and spaces to meet evolving community needs and increase user engagement in the digital age?”

This brief was selected for three reasons. First, it is accessible to participants without requiring specialized technical knowledge. Second, it is broad enough to support several idea categories, including digital services, community programming, learning spaces, inclusion, partnerships, events, technology, and user experience. Third, it is sufficiently concrete to allow project-specific interferences to be generated from the same problem statement. The library brief

therefore offers a useful context for studying fixation and defixation in early-stage innovation work.

4.3.4 Participants and data collection

Participants were recruited from master-level student groups involved in business, analytics, marketing, or design-thinking-related learning contexts. Participation was voluntary. The study was conducted through an online form that presented the innovation brief, collected pre-interference ideas, introduced the assigned modality, and collected post-interference ideas. The same general procedure was used across cohorts. The form separated the first ideation phase from the post-interference phase so that idea counts could be extracted before and after the fixation point.

The participant pool is appropriate for an exploratory creativity and innovation management study because the task does not require professional library expertise or advanced technical design skills. At the same time, the use of students creates a boundary condition. The results should be interpreted as evidence from an innovation-education and workshop-like context, not as direct evidence from professional design teams. The final sample is reported by modality in the Results section after applying the data-cleaning criteria. This avoids treating recruitment counts as the main methodological result and keeps the Methodology section focused on the design, procedure, and measurement logic.

4.3.5 Experimental procedure

The procedure followed four steps. First, participants received the library innovation brief and were asked to generate as many ideas as possible. Second, participants continued until they reached a point at which they could no longer generate additional ideas. This self-reported point was treated as the operational fixation moment. Third, participants were exposed to one of three comparison groups: GenAI-generated text interference, GenAI-generated image interference, or a neutral break. Fourth, participants were asked to generate additional ideas after the interference.

The protocol was designed to capture a simple sequence. Generate, stop, perturb, generate again. The pre-interference phase established the participant's initial idea set, the self-reported fixation point marked the moment when the participant reported being unable to continue, and the post-interference phase captured whether ideation restarted. The same sequence was used across the three comparison groups so that differences in post-fixation fluency could be interpreted as group-level differences rather than differences in task structure.

4.3.6 Interference modalities

The study compares two GenAI-generated interference modalities and one neutral-break control group.

The same public-library innovation brief was used to generate the GenAI materials for both modalities. The prompt logic was first to generate three variations of project-related verbal cue lists that could suggest alternative directions without proposing complete solutions. These three text-list variations formed the basis of the text interference. The same cue-list logic was also used as the semantic basis for generating the image interference, so that the visual mosaic remained close to the same project-related content while changing the form of presentation from verbal to visual. The text-list variations and image mosaic were generated before data collection using ChatGPT. Because ChatGPT versions may evolve over time, the study treats the generated materials as fixed study materials rather than as outputs that can be reproduced identically from the same prompt. In both GenAI modalities, the material was designed to perturb the participant's current line of thinking while leaving idea generation to the participant. The neutral-break control group received no project-specific GenAI material and continued after a task interruption. The generated materials were kept constant within each modality so that the comparison focused on modality and control-group differences rather than on variation in the interference content.

4.3.6.1 Text interference

In the text modality, participants received one of three GenAI-generated text-list variations derived from the library innovation brief. Each variation consisted of short project-related

verbal cues designed to suggest alternative directions without giving participants ready-made solutions. The purpose was to create a brief semantic perturbation that could help reopen ideation after fixation while keeping idea generation in the hands of the participant.

4.3.6.2 Image interference

In the image modality, participants received a GenAI-generated visual mosaic based on the same project-related cue-list logic used to structure the text interference. The purpose was to keep the semantic content close to the text modality while changing the form of presentation from verbal to visual. The image interference was presented as visual material rather than as a list of solution ideas. Its purpose was to create a visual and associative perturbation that could trigger new interpretations, associations, or directions after fixation. The image material was not evaluated as a design solution. It was used only as an interference.

4.3.6.3 Neutral break

In the neutral-break control group, participants did not receive project-specific GenAI material. The break functioned as an interruption control. This control group is important because time and task interruption may themselves support defixation. The three comparison groups therefore differ on two dimensions: whether the interference is GenAI-generated and whether it is project-specific. The text and image modalities are both GenAI-generated and project-specific, while the neutral break is neither. This structure allows the study to compare project-specific GenAI interference against interruption alone.

4.3.7 Measures

The main dependent variable is post-fixation ideation fluency, operationalized as the number of additional ideas generated after the interference. Pre-interference idea count is recorded mainly to describe baseline idea generation across comparison groups. Because the pre- and post-interference phases are not equivalent in duration or cognitive status, the main outcome remains the number of additional ideas generated after fixation.

The focus on fluency is deliberate. Fluency does not measure originality, feasibility, usefulness, or final innovation value. It is nevertheless appropriate for this study because the research question concerns ideation relaunch. The first question is whether participants generate additional ideas after fixation. Evaluation of idea quality, novelty, or conceptual distance is reserved for future work.

4.3.8 Data cleaning and exclusion criteria

Responses were reviewed before analysis. Responses were excluded when participants did not consent, did not complete the required ideation fields, did not provide usable idea entries, or entered material that did not correspond to idea generation. Duplicate responses and responses that showed clear non-compliance with the task were also removed. Exclusion decisions were applied before modality-level comparison.

The cleaned dataset contains only valid participants with identifiable pre-interference and post-interference idea counts. Exclusion decisions were applied before group-level comparison to avoid changing the analytical sample after observing the results.

4.3.9 Analysis strategy

The analysis proceeds in three steps. First, descriptive statistics are reported for each comparison group, including final sample size, pre-interference idea count, post-fixation idea count, mean, median, standard deviation, and standard error where appropriate. Second, the three comparison groups are compared on post-fixation fluency. Third, non-parametric comparisons are used when idea-count distributions are skewed or when group sizes make parametric assumptions difficult to justify.

The analysis emphasizes cautious interpretation. The study compares group-level patterns and does not claim that the effects are universal across all innovation tasks, participant populations, or GenAI systems. If inferential tests are used, they are interpreted alongside effect sizes and descriptive patterns rather than as the sole basis for conclusions.

4.3.10 Ethical and methodological boundaries

Participation was voluntary, and responses were used only when participants consented to research use. The study used an educational and low-risk ideation task. Participants were not evaluated on the quality of their ideas for grading purposes.

Several methodological boundaries should be noted. The study uses a single innovation brief, a student sample, and a fluency-only outcome. The interference materials depend on the GenAI system, version, prompt logic, and generation date. For reproducibility, the core prompt logic and generated materials are summarized in Section 3.6, while the study remains focused on the comparison of the text modality, the image modality, and the neutral-break control group. The comparison is therefore best understood as evidence from a controlled innovation-task setting, not as a universal test of all possible text and image interferences. These boundaries are consistent with the study's purpose: to provide an initial comparative test of whether project-specific GenAI interferences can relaunch ideation after fixation better than a neutral break.

4.4 Results

4.4.1 Data cleaning and final analytical sample

The Results section reports only the three comparison groups aligned with the research question: the text modality, the image modality, and the neutral-break control group. Other modalities present in the broader data collection are excluded from the main analysis because the present paper focuses on the comparison of project-specific GenAI-generated text and image interferences against interruption alone.

After data cleaning, the final analytical sample included 111 valid participants: 42 in the text modality, 46 in the image modality, and 23 in the neutral-break control group. All retained participants had identifiable pre-interference and post-interference idea counts. Excluded responses were removed before analysis when they lacked consent, did not complete the required ideation fields, contained unusable idea entries, or showed clear non-compliance with the task.

Table 1 reports pre-fixation and post-fixation fluency by comparison group, with emphasis on directional change. Because the comparison groups were not fully crossed across cohorts and the neutral-break control group started with higher pre-fixation fluency, the table should not be read as a simple raw post-interference ranking. Instead, it supports a cautious comparison of directionality. Did idea generation increase or decrease after the interference?

Table 4.1 Post-fixation fluency and directional change by comparison group

Comparison group	n	Pre [Q1, Q3]	Md [Q1, Q3]	Post [Q1, Q3]	Md	Δ total	Δ M (SD; SE)	Direction
Text modality	42	1 [1, 4]		2 [1, 5]		16	0.38 (3.77; 0.58)	Positive
Image modality	46	1 [1, 3]		1 [1, 4]		7	0.15 (3.03; 0.45)	Positive
Neutral-break control	23	1 [1, 8.5]		3 [1, 5]		-12	-0.52 (4.98; 1.04)	Negative

Note. Md = median; Q1 = first quartile; Q3 = third quartile; Δ = post-fixation idea count minus pre-interference idea count; M = mean; SD = standard deviation; SE = standard error. Direction refers only to the sign of the average pre/post change and should not be interpreted as a statistically significant effect.

4.4.2 Baseline pre-fixation fluency

Before the interference, participants varied in the number of ideas generated before self-reported fixation. Mean pre-fixation fluency was 3.24 ideas in the text modality, 2.59 in the image modality, and 5.87 in the neutral-break control group. However, the median pre-fixation count was 1 across all three comparison groups. This indicates skewed distributions, with a small number of high-fluency participants, particularly in the neutral-break control group.

This baseline pattern matters for interpretation. The neutral-break control group started with a higher mean pre-fixation fluency than the two GenAI-generated modalities. Therefore, the post-fixation results should not be read only through raw post-interference means. The analysis

also considers medians, interquartile ranges, and change from pre- to post- interference fluency.

4.4.3 Post-fixation fluency by modality and control group

Post-fixation fluency remained skewed across the three comparison groups. The text modality produced a mean of 3.62 post-fixation ideas and a median of 2. The image modality produced a mean of 2.74 and a median of 1. The neutral-break control group produced a mean of 5.35 and a median of 3.

On raw post-fixation counts, the neutral-break control group therefore shows the highest mean and median. This does not mean that the neutral break was necessarily more effective as an interference, because the same group also had the highest pre-fixation mean. The descriptive pattern suggests that participants who generated more ideas before fixation also tended to generate more ideas after the interference. For this reason, post-fixation fluency is interpreted together with baseline fluency and change scores.

4.4.4 Change in fluency after the interference

Change in fluency was calculated as the number of post-fixation ideas minus the number of pre-interference ideas. This measure is interpreted cautiously because the pre- and post-interference phases are not equivalent. The pre-interference phase captures ideation until self-reported fixation, while the post-interference phase captures additional idea generation after the assigned modality or neutral break. Nevertheless, change scores provide a useful descriptive indication of whether idea generation increased or decreased after the interference. The change scores provide a cautiously encouraging pattern for the GenAI-generated modalities. Mean change was positive in both the text modality and the image modality, while it was negative in the neutral-break control group. The text modality showed the strongest directional increase, followed by the image modality, whereas the neutral break showed a decline. More specifically, the text modality showed a total increase of 16 ideas and an average increase of 0.38 ideas per participant. The image modality showed a total increase of 7 ideas

and an average increase of 0.15 ideas per participant. In contrast, the neutral-break control group showed a total decrease of 12 ideas and an average decrease of 0.52 ideas per participant. Although these average changes are small in absolute terms, they are not negligible relative to a median pre-fixation fluency of 1 idea across the three comparison groups. The difference between the text modality and the neutral-break control group represents a directional gap of 0.90 ideas per participant in mean change. This gap is directionally consistent with a modest ideation-relaunch effect for the text modality, especially because the text and image modalities moved upward on average while the neutral-break control moved downward.

However, this pattern should not be overinterpreted. Median change was 0 across all three comparison groups, and the standard deviations were large relative to the mean changes. The result is therefore best interpreted through a cautious and parsimonious lens: it provides descriptive and directional evidence of possible ideation relaunch in the GenAI-generated modalities, especially the text modality, but not statistically conclusive evidence of modality superiority.

4.4.5 Group-level statistical comparison

Because the comparison groups were not fully crossed across cohorts, the statistical comparison is treated as exploratory rather than confirmatory. The neutral-break control group had higher pre-fixation fluency than the text and image modalities, which means that raw post-fixation differences cannot be interpreted as pure modality effects. This is especially important because participants who generated more ideas before fixation also tended to generate more ideas after the interference.

For this reason, the analysis uses two complementary checks. First, non-parametric Kruskal–Wallis comparisons were used to compare pre-fixation fluency, post-fixation fluency, and change in fluency across the three comparison groups. These tests provide a distribution-sensitive comparison of the observed group differences, but they do not adjust for cohort composition or baseline fluency. Second, an exploratory baseline-adjusted model was examined, predicting post-fixation fluency from pre-fixation fluency and comparison group.

This model accounts for baseline ideation fluency, but it should not be interpreted as removing all cohort-level confounding.

The Kruskal–Wallis comparisons did not indicate statistically significant differences across the three groups for pre-fixation fluency, $H(2) = 1.61$, $p = .447$, post-fixation fluency, $H(2) = 1.98$, $p = .371$, or change in fluency, $H(2) = 2.29$, $p = .318$. In the exploratory baseline-adjusted model predicting post-fixation fluency from pre-fixation fluency and comparison group, pre-fixation fluency was strongly associated with post-fixation fluency, $\beta = 0.77$, $p < .001$, while the overall comparison-group effect was not statistically significant, $F(2, 107) = 0.12$, $p = .884$. This suggests that initial ideation fluency explained more of the observed post-fixation variation than the assigned comparison group.

These results call for a methodologically cautious interpretation. The descriptive change pattern is directionally favorable to the GenAI-generated modalities, especially the text modality, but the inferential checks do not establish a statistically significant comparison-group effect. Given the non-fully-crossed cohort structure and the strong association between pre-fixation and post-fixation fluency, the most defensible interpretation is that the results are consistent with a possible relaunch effect under conservative assumptions, but they are not sufficient to claim that text or image interferences outperform a neutral break.

4.4.6 Summary of findings by research question

For RQ1, the results provide directional but not statistically conclusive evidence that project-specific GenAI-generated modalities may support post-fixation ideation relaunch more than a neutral break. The text and image modalities showed small positive mean changes from pre- to post-interference, while the neutral-break control group showed a negative mean change. However, median change was 0 across all three comparison groups, and the Kruskal–Wallis test on change in fluency was not statistically significant, $H(2) = 2.29$, $p = .318$. The result is therefore encouraging for the use of GenAI-generated interferences, but it is not strong enough to claim that they outperform a neutral break across contexts.

For RQ2, the text modality showed a more favorable descriptive pattern than the image modality. It had a larger positive mean change and a higher post-fixation median. This suggests

that short verbal cues may be easier for participants to integrate into post-fixation ideation than richer visual material. However, the difference between text and image was modest and not statistically significant. The result should therefore be read as a tentative pattern for future testing rather than as evidence that text interference is superior to image interference.

Overall, the results refine the expected relationships while requiring restraint in the strength of the conclusions. The neutral-break control group had the highest raw post-fixation fluency, partly reflecting its higher baseline fluency. In contrast, the text and image modalities showed small positive average changes after the interference, while the neutral-break control group declined. This mixed pattern supports a cautious “glass half full” interpretation. GenAI-generated interferences may help relaunch ideation for some participants, especially in the text modality. The present dataset, however, does not establish a statistically robust or generalizable superiority effect, especially once baseline pre-fixation fluency is considered.

4.5 Discussion

4.5.1 Interpreting the directional relaunch pattern

The results should be interpreted as encouraging within a deliberately conservative analytical frame, rather than confirmatory. The study does not show statistically significant superiority of the GenAI-generated modalities over the neutral-break control group. However, the descriptive direction is consistent with the theoretical logic of project-specific interference. The text modality and the image modality both produced positive average change from pre- to post-interference, while the neutral-break control group produced a negative average change. This matters because the median pre-fixation fluency was only 1 idea across the comparison groups. In that context, a mean increase of 0.38 ideas in the text modality and 0.15 ideas in the image modality remains small in absolute terms, but still relevant as a directional relaunch signal.

The most useful interpretation is therefore not that GenAI “works better” than a break. The safer interpretation is that project-specific GenAI-generated interferences may help some participants restart ideation after fixation, especially when the interference is textual and easy

to process. This distinction is important for creativity and innovation management. In applied ideation settings, facilitators often need small, low-cost mechanisms to help participants continue generating ideas. A small directional relaunch signal can still be practically interesting if it is inexpensive, scalable, and easy to integrate into a workshop.

The neutral-break result is also informative. The neutral-break control group had higher raw post-fixation fluency, but it also had stronger baseline fluency. The baseline-adjusted model confirms that pre-fixation fluency was strongly associated with post-fixation fluency. For that reason, raw post-fixation counts alone are not sufficient to evaluate the interference. The directional change pattern gives a more useful reading: GenAI-generated modalities moved upward on average, while the neutral break moved downward. This is not definitive evidence, but it supports further investigation.

4.5.2 GenAI as interference infrastructure, not ideator

The study's main contribution is to reposition GenAI as an interference infrastructure rather than as an ideator. Much of the practical discussion around GenAI in innovation focuses on whether the system can produce ideas, concepts, or solutions. That framing is useful but incomplete. It places the AI output at the center of the creative act and risks shifting the human role toward prompt writing and selection. The present study takes a different view. GenAI is used before the second ideation phase to generate material that perturbs the participant's thinking. The participant remains responsible for idea generation.

This view extends the argument developed in our prior work on LLM-generated word interferences. Earlier studies showed that ChatGPT could generate project-specific word lists and that these lists could be used to relaunch ideation after fixation, although the studies were limited by small samples and the absence of a neutral control group. The current study keeps the same interference logic but moves it into a stronger comparison: text modality, image modality, and neutral break. This helps clarify whether the effect is only a consequence of receiving more time, or whether project-specific material may add something beyond interruption.

For CIM, the implication is that GenAI does not need to replace human ideation to be valuable. It can also support the management of the creative process. In this role, GenAI becomes a facilitation tool: it helps create prompts, cues, images, or provocations that a human facilitator can introduce at a specific moment. This is closer to innovation facilitation than to automation. The value lies not in the AI's final answer, but in its ability to generate timely, project-specific perturbations.

4.5.3 Text modality as a low-friction defixation mechanism

The text modality showed the most favorable descriptive pattern. It had the largest positive total change, the largest positive mean change, and a higher post-fixation median than the image modality. This does not prove that text is superior, but it suggests that short verbal cues may be a particularly usable form of post-fixation interference.

One possible explanation is cognitive economy. After participants report fixation, they may not have much remaining attention or working memory available. A short text interference can be read quickly, interpreted flexibly, and translated into a new idea without requiring the participant to decode a rich visual scene. Text also leaves more room for the participant's own interpretation. It provides a semantic direction without imposing a detailed representation.

This is consistent with the logic of word interferences in the previous paper. The purpose of the interference is not to offer a solution, but to create a small displacement in the participant's mental path. In that sense, a text cue may be strong enough to trigger a new association but weak enough to avoid replacing the participant's ideation process. This makes text interference attractive for innovation facilitators. It is fast to generate, easy to display, easy to adapt to a brief, and does not require visual production quality.

4.5.4 Why image interference may not automatically outperform text

The image modality remains theoretically interesting, but the results do not suggest that richer media automatically create stronger ideation relaunch. Images may activate visual associations that text does not, but they may also introduce surface-feature anchoring, cognitive load, or a

new fixation point. A participant may focus on specific visual elements in the image rather than using the image as an open-ended prompt. This is particularly relevant in a post-fixation moment, where the participant may already be cognitively constrained.

The weaker descriptive pattern for the image modality should therefore not be read as a failure of visual interference. It suggests that visual material may require more careful design. The level of abstraction, realism, density, and distance from the original brief may matter. An image that is too literal may anchor participants. An image that is too abstract may be difficult to use. The current study used image interference as a broad modality, but future work should examine different types of visual interferences more precisely.

For practice, this means that facilitators should not assume that “more vivid” or “more creative-looking” GenAI output is automatically better. Image interferences may work best when they are deliberately designed as ambiguous, open, and non-solutional. The facilitation challenge is not only to generate images, but to generate images that perturb without prescribing.

4.5.5 Implications for creativity and innovation management

The first implication is that defixation can be treated as a facilitation problem. Innovation workshops often invest heavily in opening prompts, but less attention is given to the moment when the room slows down. The present study suggests that the fixation moment can be managed through lightweight interferences. The managerial question is therefore not only how to start ideation, but how to restart it.

The second implication is that project-specificity matters. Generic creativity cards, random prompts, or breaks can be useful, but they are not necessarily connected to the innovation problem. GenAI makes it possible to generate interferences from the project brief itself. This gives facilitators a way to maintain relevance while still introducing distance from the participant’s current idea set. The interference is neither fully generic nor a direct solution; it occupies an intermediate space that may support relaunch.

The third implication is that GenAI can be used in a bounded way. The results support a conservative managerial stance. Do not ask GenAI to replace the ideator. Use it to support the ideation environment. This approach reduces the risk of overreliance on AI-generated ideas

and keeps human agency central. It also aligns with the idea that AI can be useful when it is embedded into a process with clear roles, timing, and constraints.

The fourth implication is practical. Text interferences may be the easiest starting point for facilitators. They require little preparation, can be generated quickly, can be adapted to different briefs, and can be inserted into design thinking or innovation workshops without changing the whole method. Image interferences may also be valuable, but they require more attention to visual design and possible anchoring effects.

4.5.6 Boundary conditions and limitations

Several limitations shape the interpretation of the results. First, the comparison groups were not fully crossed across cohorts. This means the analysis cannot fully separate modality effects from cohort or session effects. The baseline-adjusted model accounts for pre-fixation fluency and shows that baseline fluency is strongly associated with post-fixation fluency, but it does not remove all cohort-level confounding.

Second, the study uses a student sample and a single public-library innovation brief. This context is appropriate for an exploratory ideation study because the task is accessible and does not require specialist knowledge. However, the results may differ with professional innovation teams, technical design tasks, entrepreneurial opportunity development, or longer workshops. Third, the study focuses only on fluency. This is appropriate because the research question concerns ideation relaunch, but it is not a complete creativity measure. A larger number of ideas does not necessarily mean more original, useful, or feasible ideas. Future studies should evaluate idea quality, novelty, usefulness, and conceptual distance.

Fourth, the interference depends on the GenAI system, prompt logic, and generated materials used at the time of the study. Different models, prompt structures, or image styles may produce different outcomes. The core prompt logic and generated materials are summarized in Section 3.6 to support interpretation of the intervention without shifting the article toward a technical prompt-engineering appendix. Future research should test whether the same interference logic holds across different GenAI systems, prompt designs, and visual styles.

4.5.7 Future research

Future research should first replicate the study with a fully crossed and randomized design. Each cohort should ideally include participants assigned to text, image, and neutral-break control groups under the same session conditions. This would make it easier to separate modality effects from cohort effects.

Second, future work should move beyond fluency. The next step is to examine whether the additional ideas generated after interference are original, useful, feasible, or distant from the first idea set. This would show whether GenAI-generated interferences merely increase quantity or also improve the quality and diversity of ideation.

Third, future studies should compare different types of text and image interferences. Text can vary in abstraction, distance, framing, and specificity. Images can vary in realism, ambiguity, density, and metaphorical content. Understanding these design parameters would help facilitators generate better interferences.

Finally, future research should study how facilitators actually use these interferences in live innovation workshops. The same GenAI-generated material may work differently depending on how it is introduced, timed, and discussed. The central question is therefore not only whether GenAI can generate useful interferences, but how human facilitators can embed them effectively into creative work.

4.6 Conclusion

This study examined whether GenAI can be used not as a direct ideator, but as a project-specific interference generator to relaunch human ideation after fixation. Using a public-library innovation brief, the study compared two GenAI-generated interference modalities, text and image, with a neutral-break control group. The outcome was post-fixation ideation fluency, measured through pre- and post-interference idea counts.

The findings are cautiously encouraging. The text and image modalities both showed positive average change after the interference, while the neutral-break control group showed a negative average change. The text modality showed the most favorable descriptive pattern, with a total

increase of 16 ideas and an average increase of 0.38 ideas per participant. This directional pattern is meaningful relative to the low median pre-fixation fluency, but it is not statistically conclusive. The comparison groups were not fully crossed across cohorts, the neutral-break control group had higher baseline fluency, and the statistical checks did not establish a robust comparison-group effect after accounting for pre-fixation fluency. The results therefore support a methodologically cautious interpretation: GenAI-generated interferences may help some participants relaunch ideation after fixation, especially when the interference is textual, but the present study does not establish that GenAI interferences outperform a neutral break across contexts.

The main contribution of the paper is conceptual and practical. Conceptually, it distinguishes GenAI-generated ideas from GenAI-generated interferences. In the first case, GenAI becomes a source of candidate solutions. In the second, GenAI becomes part of the facilitation infrastructure that helps humans continue ideating. This distinction matters for creativity and innovation management because it keeps human agency at the center of the creative process while still using GenAI to support the conditions of ideation. Practically, the study suggests that short, project-specific text interferences may offer a low-cost and scalable way to support ideation relaunch in innovation workshops, design-thinking courses, and early-stage innovation projects.

The study also clarifies the limits of this first empirical test. The study used a single brief, a student sample, and a fluency-only outcome. It did not evaluate originality, usefulness, feasibility, or conceptual distance. Future research should therefore test the approach in fully crossed and randomized designs, with larger and more diverse samples, multiple innovation briefs, and richer creativity outcomes. Such work can examine not only whether participants generate more ideas after GenAI-generated interferences, but whether those ideas are more original, useful, and distant from the initial fixation space.

Overall, the study supports a restrained but useful conclusion: GenAI does not need to be the genius in the room to contribute to innovation work. It can also be used as a timely, project-specific defixation tool that perturbs thinking and helps human participants restart ideation

when they reach a creative block. The practical lesson is simple: sometimes the value of GenAI is not to provide the answer, but to help people produce one more idea.

CHAPITRE 5

GENERATIVE AI AS IDEATION INFRASTRUCTURE: RELAUNCHING HUMAN IDEATION THROUGH TEXT, IMAGE, OR VIDEO INTERFERENCES

Fabrice Fischer ¹, Mickael Gardoni ^{1,2}

¹ Département de génie des systèmes, École de technologie supérieure,
1100 Notre-Dame Ouest, Montréal, Québec, Canada H3C1K3

² INSA Strasbourg, Strasbourg, France

Article soumis pour publication, juin 2026

Abstract

Generative Artificial Intelligence (GenAI) is often announced as transformative, yet we are only beginning to understand what it changes in the work of innovating. Much of the expectation is that GenAI will act as a direct ideator, thereby taking the human out of the ideation loop. Against that expectation, and from a decisively human-ideation-centric perspective, this study examines GenAI as an ideation infrastructure that generates project-specific interferences to help the human ideator relaunch ideation after fixation.

Using a public-library innovation brief, participants first generated ideas until they reported having run out of ideas and reached self-reported fixation. To help them come up with more and different ideas, they were then exposed to one of three GenAI-generated interference modalities: text interference, image mosaic interference, or video mosaic interference. Other participants took a neutral break.

The study evaluates post-fixation ideation through fluency and FixShift, a measure of how far post-interference ideas depart from participants' pre-interference idea families. The descriptive and exploratory statistical results indicate moderate directional differences across the interferences and the neutral break. In this use case, text shows the strongest directional relaunch pattern, followed by the image mosaic, the neutral break, and the video mosaic, suggesting that richer media are not automatically more effective.

Together, the study positions GenAI-assisted defixation as a human-centered innovation-process method and provides comparative evidence on how text, image, and video interferences affect post-fixation fluency and conceptual movement, with practical implications for GenAI-assisted innovation management.

Keywords: Generative AI, human-AI creativity, ideation defixation, innovation management, interference modalities, design fixation, FixShift

5.1 Introduction

5.1.1 GenAI in the work of innovating

Generative Artificial Intelligence is entering the work of innovating with a very ambitious job description. It is expected to accelerate idea generation, support opportunity exploration, visualize possibilities, assist product and service development, synthesize user information, and change how organizations create value. This expectation is now visible across innovation-management research, where AI and GenAI are studied in relation to innovation management, new product development, front-end ideation, employee creativity, improvisation, and human-AI creative work (Eisenreich, Just, Gimenez-Jimenez, & Füller, 2024 ; Pescher & Tellis, 2025 ; Roberts & Candi, 2024 ; Tang, Li, Yao, Zhang, & Zhu, 2026 ; Vallé, Seitz, Kanbach, & Schuhmacher, 2026 ; Ma, Ge, Li, & Wang, 2026).

Yet innovating is not only the production of an output list. It is also the human work of noticing, hesitating, reframing, rejecting the comfortable, connecting distant elements, and producing one more idea when the first family of ideas has become too dominant. This is where the direct-ideator view of GenAI becomes too narrow. GenAI can generate many ideas quickly, but innovation also depends on how humans remain active in framing, judging, redirecting, and sometimes refusing the apparent logic of the problem (Grilli & Pedota, 2024 ; Pedota et al., 2025 ; Raj, Srivastava, & Behera, 2026).

This paper focuses on that second question. It examines GenAI in a narrower and more human-centered role: not as a direct ideator, but as an ideation infrastructure able to generate project-specific interferences. These interferences are introduced when participants report that they have run out of ideas. Their role is not to provide the solution. Their role is to disturb the current ideation path so that the human ideator can relaunch ideation after fixation.

5.1.2 From GenAI as direct ideator to GenAI as ideation infrastructure

The direct-ideator view of GenAI is understandable. A system that can generate fluent text, images, and videos in seconds looks immediately useful for innovation work. It can produce ideas, variants, representations, scenarios, and prototypes at a speed that human teams cannot match. This makes GenAI attractive for the front end of innovation, where teams are expected to explore possibilities before selecting and developing a few of them (Pescher & Tellis, 2025 ; Bell, Pescher, Tellis, & Füller, 2024 ; Bellesia et al., 2026).

But more output does not automatically mean better ideation. GenAI-generated ideas can be plausible, fluent, and well presented while still remaining close to dominant patterns, common solutions, or the logic already present in the prompt. Recent work points to this double-edged effect. GenAI can support creativity, ideation, and improvisation, but it can also encourage convergence, overreliance, and reduced diversity if the human starts working around machine-generated suggestions instead of moving beyond them (Doshi & Hauser, 2024 ; Eisenreich et al., 2024 ; Ma et al., 2026 ; Wadinambiarachchi et al., 2024).

This is the direct-ideator paradox. GenAI can help produce ideas, but it can also narrow what humans continue to imagine. The system can be powerful inside the box it is given. That box is made of training data, prompt wording, provided context, and the patterns the model can recombine. This is useful. It is also a limit. Innovation sometimes begins when a human decides that the given frame is wrong, incomplete, or too comfortable.

This study takes another route. GenAI is not asked to produce the ideas to be evaluated. It is asked to produce interferences that can be introduced after human ideation stalls. The difference is important. A GenAI-generated idea competes for the role of solution. A GenAI-generated interference has a smaller and more precise function: it perturbs thinking without

prescribing the next answer. The participant remains responsible for producing, interpreting, selecting, and developing the ideas.

This paper describes that role as GenAI as ideation infrastructure. The infrastructure is not valuable because it is creative instead of the human. It is valuable because it can generate project-specific, non-prescriptive material at the moment when the human ideator needs a disturbance. In this role, GenAI becomes less a substitute for ideation and more a timed support for human defixation.

5.1.3 Post-fixation human ideation as the problem

Innovation teams rarely struggle to produce a first idea. They more often struggle when the first ideas become too comfortable. After an initial burst, the room starts circling the same territory. Participants reformulate familiar concepts, make minor variations, or report that they have “run out” of ideas. In design cognition, this situation is commonly associated with fixation, understood as adherence to a limited set of ideas (Jansson & Smith, 1991 ; Purcell & Gero, 1996).

For innovation management, fixation is not only a cognitive label. It is a practical problem in the process of innovating. The front end of innovation requires divergence before convergence. But divergence becomes difficult when prior knowledge, first ideas, the wording of the brief, and external examples occupy too much of the ideation space. The problem is not that participants have no intelligence or no creativity. The problem is that their ideation has become locally stable too early.

Defixation therefore becomes a managerial and methodological problem: what can be introduced when human ideation stalls, without taking the ideator role away from the human? A neutral break can interrupt the flow, but it does not add new semantic or visual material. A text interference can provide short verbal cues. An image mosaic interference can provide static visual perturbation. A video mosaic interference can add dynamic visual material, but it may also capture attention or increase cognitive load. These alternatives should not be assumed equivalent.

This paper studies that post-fixation moment. Participants first generated ideas for a public-library innovation brief until they reported having run out of ideas. They were then exposed to one of three GenAI-generated interference modalities, text interference, image mosaic interference, or video mosaic interference. Other participants took a neutral break. The question is whether these interferences help humans relaunch ideation, and whether they help humans move away from their own pre-interference idea families.

5.1.4 Research gap: limited cross-modal evidence on GenAI-generated interferences

The literature now gives several reasons to take GenAI seriously in innovation work. Studies examine AI in innovation management, AI in new product development, GenAI-supported ideation, human-AI creativity, AI affordances, improvisation, and AI as a non-human innovation intermediary (Just, 2024 ; Roberts & Candi, 2024 ; Eisenreich et al., 2024 ; Pedota et al., 2025 ; Vallé et al., 2026 ; Tang et al., 2026 ; Pescher & Tellis, 2025). In parallel, design and creativity research has long studied fixation, defixation, examples, analogies, distractions, incubation, and stimuli as ways to reopen ideation (Jansson & Smith, 1991 ; Leahy et al., 2020 ; Purcell & Gero, 1996 ; Sio et al., 2015 ; Vasconcelos & Crilly, 2016 ; Youmans & Arciszewski, 2014).

However, three gaps remain important for the present paper.

First, much of the GenAI discussion still treats GenAI as a producer of ideas. This is useful, but it leaves underdeveloped the more bounded role studied here: GenAI as an infrastructure that generates interferences while the human remains the ideator. This distinction matters because AI-human creativity is not only a question of automation. It is also a question of augmentation, agency, and the distribution of roles between machine output and human interpretation (Grilli & Pedota, 2024 ; Pedota et al., 2025 ; Raj et al., 2026).

Second, the role of interference modality remains insufficiently understood. Text, image, and video do not disturb thinking through the same cognitive channel. Text may be easier to reinterpret. Images may activate associations that words do not. Video may provide richer movement and atmosphere, but may also reduce the space left for human interpretation. Prior

work on analogical examples, design fixation, generative AI stimuli, and visual AI tools suggests that representational form matters, but it does not yet provide a clear post-fixation comparison of GenAI-generated text, image mosaic, and video mosaic interferences under one shared protocol (Chan et al., 2011 ; Atilola & Linsey, 2015 ; Gong et al., 2024 ; Fang et al., 2025 ; Wadinambiarachchi et al., 2024 ; Ge & Hou, 2025).

Third, existing studies are difficult to compare because they often differ in task, timing, stimulus construction, participant context, and measured outcome. Some studies examine GenAI as direct ideator. Others examine text stimuli, visual stimuli, AI-supported design, or idea evaluation. What is still missing is a common post-fixation protocol comparing text interference, image mosaic interference, video mosaic interference, and neutral break within the same innovation brief, using both a quantity measure and a measure of conceptual movement (Shah et al., 2003 ; Nelson, Wilson, Rosen, & Yen, 2009 ; Orwig, Diez, Vannini, Beaty, & Sepulcre, 2021).

This paper addresses that gap by comparing three project-specific GenAI-generated interference modalities and a neutral break after self-reported fixation. Previous studies in this domain have investigated the use of targeted, AI-produced stimuli to overcome creative blocks, demonstrating initial viability for this method (Fischer & Gardoni, 2025a, 2025b). It evaluates not only whether participants generate again, but also whether their post-interference ideas depart from the idea families they had already produced.

5.1.5 Research questions and contributions

The paper is guided by three research questions.

RQ1. How can GenAI be used as ideation infrastructure to generate project-specific interferences for human ideation defixation?

RQ2. How do text interference, image mosaic interference, video mosaic interference, and neutral break differ in their ability to relaunch post-fixation ideation, as measured through fluency?

RQ3. How do text interference, image mosaic interference, video mosaic interference, and neutral break differ in FixShift, understood as the departure of post-interference ideas from participants' pre-interference idea families?

Together, these questions separate the method, the relaunch effect, and the conceptual movement effect.

The paper makes three contributions.

First, it contributes a human-centered methodology for using GenAI in innovation work. Instead of treating GenAI as the ideator, the paper treats it as an ideation infrastructure that generates timed, project-specific interferences. This shifts the managerial question from “Can GenAI produce ideas?” to “Can GenAI help humans produce again when ideation stalls?”

Second, it provides comparative evidence on three GenAI-generated interference modalities and a neutral break under a shared post-fixation protocol. In doing so, the paper does not assume that more media richness is better. It tests whether text interference, image mosaic interference, and video mosaic interference differ in their post-fixation effects.

Third, it contributes FixShift as a complementary way to evaluate defixation. Fluency captures whether ideation relaunches. FixShift captures whether post-interference ideas move away from the participant's pre-interference idea families. This distinction matters because generating more ideas and generating different ideas are related, but they are not the same.

Taken together, the paper positions GenAI-assisted defixation as a manageable innovation-process method. The claim is not that GenAI should replace human ideation, or that one interference modality is universally superior. The claim is more limited and more practical: when human ideation stalls, project-specific GenAI-generated interferences may help create the right disturbance at the right moment, so that humans can produce the next idea.

5.2 Theoretical Background

5.2.1 GenAI in innovation management

GenAI is becoming part of the work of innovating. It supports opportunity exploration, idea generation, visualization, user insight synthesis, and new product development. This is why it

is attractive for innovation management: it promises speed, volume, accessibility, and flexibility at the front end of innovation, where the objective is to open possibilities before convergence begins (Bell et al., 2024 ; Fetrati, Hansen, & Akhavan, 2022 ; Ma et al., 2026 ; Pescher & Tellis, 2025 ; Roberts & Candi, 2024 ; Tang et al., 2026 ; Vallé et al., 2026).

This literature makes GenAI relevant beyond a technical design context. The open question is how to use it without moving the human out of the ideator role. This paper studies GenAI as innovation facilitation rather than automation: a way to generate a temporary, project-specific disturbance when human ideation stalls, while the human ideator still produces the ideas.

5.2.2 Human-AI creativity and the direct-ideator paradox

The direct-ideator view of GenAI is powerful because it is easy to understand. The system generates output. The human receives it. The process looks faster and more productive. Yet more output does not automatically mean better ideation. GenAI can produce fluent and attractive ideas that remain close to common patterns, reproduce familiar solutions, or become a new fixation source (Doshi & Hauser, 2024 ; Eisenreich et al., 2024 ; Wadinambiarachchi et al., 2024). Recent work on human-AI creativity points to the same tension: the value of GenAI depends on how the human and the system share the creative task (Grilli & Pedota, 2024 ; Hoßbach & Isaksen, 2025 ; Pedota et al., 2025 ; Raj et al., 2026 ; Tang et al., 2026).

This is where machine recombination and human reframing differ. GenAI is powerful inside the box it is given: training data, prompt wording, provided context, and recombinable patterns. Innovation often requires noticing that the box itself is wrong, incomplete, or too comfortable. The role of GenAI here is not to think outside the box for the participant. It is to disturb the box enough for the participant to think again.

The present study responds to the direct-ideator paradox with a more bounded use of GenAI. The participant does not ask GenAI for solutions. The participant does not choose among AI ideas. The participant receives an interference and then generates ideas personally. GenAI prepares the disturbance. The human still does the ideation.

This distinction is central to the paper. The study is not asking whether GenAI is creative. It is asking whether GenAI can help humans become generative again after fixation.

5.2.3 Fixation, defixation, and ideation relaunch

Innovation teams rarely struggle to produce a first idea. They often struggle after the first ideas have arrived and become comfortable. Participants repeat them, reword them, or make small variations. Design research names this difficulty fixation, commonly understood as adherence to a limited set of ideas (Jansson & Smith, 1991 ; Purcell & Gero, 1996). Fixation can come from examples, problem framing, prior knowledge, or earlier ideas, and examples can act as both inspiration and fixation sources (Leahy et al., 2020 ; Sio et al., 2015 ; Vasconcelos & Crilly, 2016).

Defixation attempts to reopen that movement. It is not simply “more ideas,” because more ideas can still be repetitions. Defixation means that ideation starts again and moves away from the dominant idea family. This is why the present paper evaluates both fluency and FixShift. The practical question is simple. What can be introduced when human ideation stalls, without taking the ideator role away from the human? A neutral break interrupts the flow, but adds no project-specific material. An interference disturbs the current path without prescribing the next idea. Prior work has examined examples, analogies, patents, incubation, and representational changes for that purpose (Atilola & Linsey, 2015 ; Chan et al., 2011 ; Koh & De Lessio, 2018 ; Leahy et al., 2020 ; Youmans & Arciszewski, 2014). GenAI changes the practical setting because the interference can be generated from the project brief, at low cost, and at the moment of fixation.

5.2.4 Interferences as human defixation mechanisms

An interference is a deliberate perturbation introduced after ideation stalls. It is not the solution and not an idea to copy. It is a disturbance meant to interrupt persistence around the dominant idea path.

Designers and creators have long used cards, analogies, examples, images, patents, provocations, and breaks to escape habitual thinking. Brian Eno’s Oblique Strategies are a familiar example (Eno & Schmidt, 1975). In engineering design, patents, examples, analogies, and representational formats can inspire, but they can also pull attention toward surface

features or familiar solution principles (Atilola & Linsey, 2015 ; Chan et al., 2011 ; Koh & De Lessio, 2018 ; Srinivasan et al., 2018 ; Vasconcelos & Crilly, 2016).

This is why the word interference is useful. It emphasizes the tactical role of the material. The interference should be strong enough to disturb, but not so strong that it becomes the next fixation. Project-specific GenAI creates a useful middle ground: the interference can speak the language of the brief without answering the brief, bringing new cues while leaving the ideation work to the human.

5.2.5 Text, image mosaic, and video mosaic as interference modalities

The modality of the interference should matter. Text, image mosaic, and video mosaic do not ask the same work from the human ideator. They do not use the same channel, create the same cognitive load, or leave the same interpretive space.

Text interference is the lightest modality in this study. It provides verbal cues that can be read quickly and reinterpreted easily. A short phrase can suggest a direction without imposing a full representation, which may support semantic reframing and new associations (Gong et al., 2024).

Image mosaic interference introduces visual material. It may activate associations that words do not, but it can also anchor attention on visible surface features. Recent work on GenAI-enhanced conceptual design and visual stimuli supports the relevance of visual AI outputs, while also pointing to the need to evaluate their effects on fixation and creativity (Baudoux et al., 2025 ; Fang et al., 2025 ; Ge & Hou, 2025 ; Sattelle & Ortiz, 2025 ; Wadinambiarachchi et al., 2024 ; Zhu & Luo, 2025).

Video mosaic interference adds movement, sequence, and atmosphere. It may feel richer and more engaging, but richer does not mean easier to use after fixation. A video mosaic may create too many cues at once or capture attention instead of freeing it. The neutral break is different again. It is a non-semantic interruption baseline that helps separate the value of interruption alone from the value of project-specific material.

The cross-modal comparison is therefore part of the theory. If GenAI is used as ideation infrastructure, the form of the generated interference matters. The infrastructure is not only the model. It is also the modality through which the human encounters the disturbance.

5.2.6 Fluency and FixShift as complementary outcomes

A defixation study needs to know whether ideation restarts and whether it moves. Fluency answers the first question. It counts how many ideas participants generate after the interference or the neutral break. This is basic, but not trivial. If participants report that they have run out of ideas, generating again is already meaningful. Fluency is also one of the classic measures used to evaluate ideation effectiveness, even though it should not be confused with originality, usefulness, feasibility, or final innovation value (Nelson et al., 2009 ; Shah et al., 2003).

FixShift addresses the second question. It measures how far post-interference ideas depart from the participant's own pre-interference idea families. The focus is not only whether the participant generates again. The focus is whether the participant moves away from what was already produced. This logic is close to work that treats semantic distance and conceptual movement as important signals of creativity, even if FixShift remains a simpler and more task-specific measure (Orwig et al., 2021).

Together, fluency and FixShift capture two different sides of post-fixation ideation. They make it possible to examine whether text interference, image mosaic interference, video mosaic interference, and neutral break differ not only in quantity, but also in how they help humans move beyond earlier idea families. The best interference is not necessarily the richest one. It is the one the human can use to produce the next idea, and possibly a different one.

5.3 Methodology

This section explains how GenAI was used as ideation infrastructure, how the interferences were generated, how participants experienced the post-fixation protocol, and how fluency and FixShift were measured. The methodology has a dual purpose. It documents a possible method for using GenAI-generated interferences when human ideation stalls. It also supports a

comparison of text interference, image mosaic interference, video mosaic interference, and neutral break under a shared empirical protocol.

The study is not designed as a test of GenAI as a direct ideator. Participants did not receive AI-generated solutions to select or improve. They first generated their own ideas. Only after they reported fixation did they receive an interference modality or a neutral break. The second ideation phase was then used to examine whether human ideation relaunched and whether the new ideas moved away from the participant's earlier idea families.

5.3.1 Research design overview

The study uses a DSR-informed comparative design. Design Science Research Methodology supports the development and evaluation of a practical artifact: a method for generating project-specific GenAI interferences when ideation stalls (Peffer et al., 2018, 2007). Comparative Case Study logic supports the comparison of text interference, image mosaic interference, video mosaic interference, and neutral break under the same task brief, fixation trigger, and two-phase ideation structure (Bartlett & Vavrus, 2017).

The study therefore has two linked purposes. The methodological purpose is to document how GenAI can be used as ideation infrastructure. The empirical purpose is to examine whether the form of the post-fixation alternative matters once participants report fixation.

The evidence is interpreted as comparative pattern evidence. The study was conducted across several cohorts and sites, but it was not a fully crossed randomized laboratory experiment. Each participant experienced only one interference modality or the neutral break. For that reason, the results are not presented as universal causal estimates. They are presented as patterned evidence under one shared post-fixation protocol.

5.3.2 Empirical setting: the public-library innovation brief

All participants worked on the same public-library innovation brief:

How can public libraries innovate their services and spaces to meet evolving community needs and increase user engagement in the digital age?

This brief was selected because it is accessible, innovation-oriented, and broad enough to create fixation risk. Participants do not need specialized technical expertise to understand it, yet the brief can generate several idea families, including services, spaces, digital access, community engagement, education, inclusion, partnerships, and user experience. It also produces obvious early ideas, which makes it suitable for studying what happens after the first burst of ideation.

The brief was kept identical across the study. This was important because the paper compares the form of the post-fixation interference, not different innovation problems.

5.3.3 GenAI as ideation infrastructure: interference generation logic

The GenAI artifact is not an idea generator for participants. It is an interference generator. Its purpose is to create material that can be introduced after human ideation stalls.

The generation logic followed four principles.

First, the interference had to be project-specific. It had to be generated from the public-library brief, not from a generic creativity prompt. This is important because a project-specific interference remains connected to the task without directly solving it.

Second, the interference had to be non-prescriptive. It could suggest directions, associations, atmospheres, or possible domains of thought, but it should not provide a ready-made library solution. The participant should not be invited to copy the interference.

Third, the interference had to preserve the human ideator's role. GenAI prepares the disturbance. The participant produces the ideas.

Fourth, the interferences had to be comparable enough for empirical evaluation. The text interference was generated from the public-library innovation brief. The image mosaic and video mosaic were then generated using the project-specific text-interference list as part of their prompt, while preserving the same broad intent: perturb ideation after fixation without prescribing the next idea.

In the present study, the GenAI input was deliberately restricted to a brief-based stimulus chain. The text interference was generated from the public-library innovation brief. The image mosaic and video mosaic were then generated from the project-specific text-interference list,

so that the visual and dynamic interferences remained anchored in the same semantic field without becoming direct solution proposals. This choice kept the comparison focused on interference modality rather than on richer upstream design-research inputs. In a live design-thinking project, the input could be richer. GenAI could synthesize empathy-phase material, interview notes, observations, personas, journey maps, stakeholder maps, user complaints, service data, or community feedback. This would make the interference infrastructure more situated, but also harder to compare experimentally. The present study therefore tests a brief-based version first.

The GenAI-generated interferences were created with ChatGPT tools developed by OpenAI and available during the study period. The study does not evaluate the creativity, quality, or performance of a specific model version or prompt. The model, prompt, and generated outputs are treated as implementation instances of a broader GenAI-assisted defixation method. Because generative outputs vary across model versions, prompt wording, interface conditions, and repeated generations, the relevant methodological object is the interference-generation logic: project-specificity, non-prescription, preservation of the human ideator's role, and cross-modal comparability. The exact prompts are therefore not treated as the object of analysis; they are procedural instantiations of the interference-generation logic and are expected to vary across uses of the method.

This is the central methodological move of the paper. GenAI is not used at the beginning to flood participants with ideas. It is used at a precise moment, after participants report that they have run out of ideas. The method therefore tests GenAI as timed ideation infrastructure.

5.3.4 Text interference, image mosaic interference, video mosaic interference, and neutral break

The study compares three GenAI-generated interference modalities with a neutral break.

1. Text interference

The text interference consisted of a project-specific word or short-phrase list generated from the public-library brief. The text was designed to offer semantic cues without giving direct

solutions. Its role was to create a low-friction verbal perturbation. Participants could read it quickly, reinterpret it, and use it as a starting point for their own additional ideas.

2. Image mosaic interference

The image mosaic interference consisted of a single visual panel composed of multiple project-specific images. The aim was to create a static visual perturbation. The mosaic offered visual associations while avoiding a sequential experience. It was designed to stay close enough to the public-library brief to matter, but indirect enough not to become a proposed solution.

3. Video mosaic interference

The video mosaic interference consisted of a short project-specific dynamic visual interference. It extended the visual logic of the image mosaic into movement, sequence, and atmosphere. Its role was to test whether a richer and more dynamic interference would help participants relaunch ideation after fixation. It was not treated as automatically superior because richer media may also create more cue load.

4. Neutral break

The neutral break consisted of a timed one-minute pause without added semantic or visual material. Its role was not to act as a GenAI-generated interference. Its role was to provide an interruption baseline. It helps answer a practical question: is it enough to stop for one minute, or does project-specific interference material add something beyond interruption?

Table 5.1 Post-fixation interferences and neutral break

Post-fixation alternative	What the participant saw	Generation or selection logic	Exposure mode	Intended role
Text interference	Project-specific word list	Generated from the public-library innovation brief and filtered to avoid direct solution prescription	Static text displayed in the form	Semantic verbal interference
Image mosaic interference	Project-specific image mosaic	Generated using the project-specific text-interference list as part of the prompt and assembled as one visual panel	Static image displayed in the form	Semantic visual interference
Video mosaic interference	Project-specific video mosaic	Generated using the project-specific text-interference list as part of the prompt and the same non-prescriptive intent	Short video displayed in the form	Semantic dynamic visual interference
Neutral break	Neutral digital clock from 60 to 0 seconds	No semantic content added	Timed one-minute clock displayed in the form	Non-semantic interruption baseline

Table 5.1 summarizes the three GenAI-generated interference modalities and the neutral break. It clarifies what participants saw, how each post-fixation alternative was generated or selected, and what role it played in the protocol.

5.3.5 Participants, cohorts, and data collection protocol

Participants completed the study individually through a web form. They all received the same public-library innovation brief. They were asked to enter one idea per line.

Participants were assigned by class group to one post-fixation alternative: text interference, image mosaic interference, video mosaic interference, or neutral break. The assignment was not randomized at the individual level. Class groups were assigned in no particular order, with the practical objective of obtaining participants for each interference modality and for the neutral break. This allocation structure supports comparative pattern analysis, but it does not support strong causal claims at the level of a fully randomized experiment.

Participation was voluntary, responses were analyzed in anonymized form, and the study used a low-risk educational ideation task.

The protocol had two ideation phases. In the first phase, participants generated ideas in response to the library brief. They continued until they reported having run out of ideas. This self-reported moment was used as the operational fixation breakpoint. The study does not claim to measure fixation neurologically or clinically. It uses the participant's own report that ideation has stalled as the practical trigger for the interference.

After this fixation breakpoint, participants were exposed to one of the following post-fixation alternatives: text interference, image mosaic interference, video mosaic interference, or neutral break. The exposure was self-paced within a timed screen. Participants could inspect the text interference, image mosaic interference, or video mosaic interference while a digital clock counted down from 60 seconds to 0. Participants assigned to the neutral break saw the same countdown without added semantic or visual material.

After the countdown, participants continued to the second ideation phase and generated additional ideas in response to the same public-library innovation brief. The second set of ideas is the post-fixation material. It is used to calculate post-fixation fluency and FixShift. This structure makes the participant's own first idea set the reference point for evaluating what happens after the interference or neutral break.

The experiments were conducted across multiple cohorts and sites under the same broad protocol. This gives the study a more realistic empirical base than a single classroom exercise. It also creates limits. The design is multi-site and comparative, but not fully crossed across all cohorts. The analysis therefore focuses on participant-level comparison and directional patterns rather than strong causal claims.

5.3.6 Measure: fluency and FixShift

The study evaluates two post-fixation outcomes: fluency and FixShift.

Fluency measures whether ideation relaunched after fixation. It is counted as the number of post-fixation ideas generated by each participant. The paper also reports pre-fixation idea counts and post-fixation changes when useful. Fluency is not treated as a complete measure of creativity. It does not establish originality, usefulness, feasibility, or final innovation value. It answers a more basic question: after participants said they had run out of ideas, did they generate again?

FixShift measures conceptual movement. It captures how far each post-interference idea departs from the participant's own pre-interference idea families. This matters because more ideas are not necessarily different ideas. A participant can generate additional ideas that remain local variations of the same first family. In this paper, defixation requires more than volume. It also requires some movement away from the earlier idea path.

FixShift is scored on a 0 to 2 scale:

- 0: Same idea family or same core solution logic as the participant's pre-interference ideas
- 1: Partial shift or meaningful variation, but not a clearly new idea family
- 2: Clearly new idea family or distinct conceptual departure

Each post-fixation idea is compared with that participant's own pre-interference idea set. The comparison is not made against a pooled idea corpus. This is important because fixation is participant-specific. What is a new idea family for one participant may not be new for another. Using both fluency and FixShift allows the paper to separate two questions. Fluency asks whether ideation restarted. FixShift asks whether ideation moved to new idea families.

5.3.7 AI-assisted FixShift coding and human validation

FixShift coding was applied to the full post-fixation idea dataset through an AI-assisted coding workflow. The reason was practical. FixShift requires comparing each post-fixation idea with the participant's own pre-interference idea families. This is detailed, repetitive, and difficult to scale manually across the full dataset.

The AI-assisted coding used the same 0 to 2 rubric for all post-fixation ideas. Each idea was interpreted in relation to the participant's own pre-interference idea set. The AI-assisted score was not treated as unquestioned ground truth. It was treated as a scalable proxy that required human validation.

To validate the coding, a blind validation subset of 60 post-fixation ideas was created. The subset was balanced across interference modalities, neutral break, and AI-assigned FixShift levels. Three human raters used the same FixShift rubric. Their judgments were then aggregated and compared with the AI-assisted scores.

Validation was reported with several indicators:

1. exact agreement;
2. agreement within one category;
3. mean absolute difference from the aggregated human mean;
4. quadratic weighted kappa.

This validation logic is important for credibility. It shows that FixShift is not only a black-box metric. It also keeps the interpretation restrained. AI-assisted coding makes full-dataset FixShift possible, but the resulting scores should still be read as comparative pattern evidence.

5.3.8 Data cleaning and exclusions

The analysis uses one idea line as the unit for fluency counting. This is a conservative rule. It keeps the procedure transparent, but it also means that some complex ideas may be counted as one line and some multi-line entries may require judgment.

Blank post-fixation phases were treated as zero post-fixation ideas rather than missing data. This is important because a blank post-fixation phase means that ideation did not relaunch after the interference or neutral break.

Non-idea lines were removed from idea-level coding. In the current dataset, two placeholder non-ideas were excluded from post-fixation idea coding. After this cleaning, the dataset contained 981 coded idea lines, split into 502 pre-interference ideas and 479 post-interference ideas. The comparative dataset contains 138 participants.

One exploratory video-related cell was retained in the broader dataset but excluded from the comparative analysis because it was too small for stable interpretation. This paper therefore keeps the video mosaic interference, but not the VIM-related exploratory material. The main comparison is text interference, image mosaic interference, video mosaic interference, and neutral break.

The final analysis is reported mainly at the participant level. This is necessary because each participant experienced only one interference modality or the neutral break. For fluency, participant-level reporting avoids over-reading unequal idea counts. For FixShift, raw coding is done at the idea level, but results are aggregated per participant. This avoids treating multiple post-fixation ideas from the same participant as independent observations.

These choices define the limits of the method. The study supports comparison under a shared protocol, but it does not claim universal superiority of one interference modality. It provides directional, patterned evidence on how different GenAI-generated interferences and the neutral break relate to post-fixation fluency and conceptual movement.

5.4 Results

This section reports the empirical pattern across the three GenAI-generated interference modalities and the neutral break. The results are presented in four steps. First, the dataset is summarized. Second, post-fixation fluency is compared across text interference, image mosaic interference, video mosaic interference, and neutral break. Third, FixShift is compared across the same post-fixation alternatives. Fourth, the human validation of AI-assisted FixShift coding is reported.

The purpose of this section is not to claim universal superiority of one modality. The purpose is to show what happened in this use case, under one public-library brief and one shared post-fixation protocol.

The analysis is descriptive and comparative. It reports participant-level patterns rather than causal estimates. Because the empirical structure is multi-site but not fully crossed, the results should be read as directional evidence under one shared post-fixation protocol.

5.4.1 Dataset overview

The comparative dataset contains 138 participants and 983 recorded idea lines. After excluding two non-idea lines from the post-interference phase, 981 coded idea lines remained. These coded idea lines are split into 502 pre-interference ideas and 479 post-interference ideas.

The fluency table uses the participant-level one-line counting rule. It reports 481 post-fixation lines before the two non-idea exclusions used for FixShift coding. This distinction should be kept explicit. Fluency is based on participant-level line counts. FixShift is based on coded post-interference ideas after removing non-ideas.

One exploratory video-related cell was retained in the broader dataset but excluded from the comparative analysis. It was too small for stable interpretation. The main comparison therefore includes text interference, image mosaic interference, video mosaic interference, and neutral break.

Because the empirical design is multi-site but not fully crossed, the results are interpreted as comparative pattern evidence. They should not be read as strict causal estimates.

5.4.2 Post-fixation fluency by interference modality and neutral break

Fluency is used here as a relaunch indicator. It tells us whether participants generated again after they had reported fixation. It does not measure idea quality, originality, feasibility, or usefulness.

Table 2 reports pre-interference and post-interference idea counts by post-fixation alternative. Because the distributions are skewed, medians and quartiles are useful to read together with mean delta.

Table 5.2 Post-fixation fluency by interference modality and neutral break

Post-fixation alternative	N participants	Total pre ideas	Total post lines	Pre Median [Q1, Q3]	Post Median [Q1, Q3]	Delta mean (SD)
Text interference	42	136	152	1 [1, 4]	2 [1, 5]	0.38 (3.77)
Image mosaic interference	46	119	126	1 [1, 3]	1 [1, 4]	0.15 (3.03)
Video mosaic interference	27	112	80	3 [1, 5]	2 [1, 3.5]	-1.19 (4.86)
Neutral break	23	135	123	1 [1, 8.5]	3 [1, 5]	-0.52 (4.98)

Notes: Fluency uses the one-line counting rule. Total post lines are reported before removing two non-idea lines used in FixShift coding. Delta mean is calculated as post-fixation fluency minus pre-interference fluency. Because participants were assigned by class groups and not individually randomized, these fluency differences are interpreted as directional comparative evidence.

The descriptive fluency pattern is moderate and should be interpreted cautiously. Text interference is the only post-fixation alternative where the total post lines clearly exceed the total pre ideas. It also has the strongest positive mean delta, at +0.38. Image mosaic interference remains close to neutral, with a small positive mean delta of +0.15. The neutral break shows a

negative mean delta of -0.52. The video mosaic interference shows the weakest fluency pattern, with a mean delta of -1.19.

In this use case, the relaunch pattern is therefore strongest for text interference, weaker but still positive for image mosaic interference, negative for the neutral break, and weakest for video mosaic interference.

5.4.3 FixShift by interference modality and neutral break

Fluency only tells us whether ideation restarted. It does not tell us whether participants moved away from their earlier idea families. FixShift addresses this second question.

Table 3 reports participant-level FixShift summaries. A higher mean FixShift indicates that post-interference ideas depart more from the participant's own pre-interference idea families. A higher percentage of FixShift = 2 indicates a higher share of ideas that belong to clearly new idea families.

Table 5.3 FixShift by interference modality and neutral break

Post-fixation alternative	N participants with post ideas	Mean FixShift (SD) across participants	% FixShift=2 per participant (SD)
Text interference	41	1.45 (0.59)	57.3% (42.1%)
Image mosaic interference	42	1.33 (0.66)	49.1% (43.1%)
Video mosaic interference	25	1.00 (0.69)	32.0% (37.5%)
Neutral break	22	1.20 (0.58)	39.5% (38.0%)

Notes: FixShift is calculated only for participants with at least one coded post-interference idea. Scores are first coded at idea level, then aggregated at participant level. Because FixShift is aggregated at participant level, the table avoids treating multiple post-interference ideas from the same participant as independent observations.

The FixShift results follow the same broad descriptive order as the fluency results, with one important difference. The neutral break sits above the video mosaic interference for conceptual

movement. Text interference shows the strongest descriptive pattern, with a mean FixShift of 1.45 and 57.3% of post-interference ideas scored as FixShift = 2. Image mosaic interference follows, with a mean FixShift of 1.33 and 49.1% FixShift = 2. The neutral break follows, with a mean FixShift of 1.20 and 39.5% FixShift = 2. Video mosaic interference is lowest, with a mean FixShift of 1.00 and 32.0% FixShift = 2.

This matters because the pattern is not only about producing more lines. Text interference is also associated with the strongest conceptual movement away from participants' earlier idea families.

5.4.4 Cross-modal pattern

Taken together, the fluency and FixShift results suggest that the form of the interference matters after fixation has been reported.

In this use case, the overall pattern is:

1. Text interference shows the strongest directional relaunch pattern and the strongest FixShift pattern.
2. Image mosaic interference remains positive, but weaker than text.
3. Neutral break occupies an intermediate position, especially for FixShift.
4. Video mosaic interference appears least effective under the present protocol.

The study does not claim that text is universally superior, that video mosaic is generally ineffective, or that FixShift captures full creativity. It tests whether a brief-based GenAI interference infrastructure produces directional differences in one post-fixation ideation setting.

This cautious reading matters. In this public-library brief and under this post-fixation protocol, richer media did not automatically produce stronger defixation.

This is the main empirical bridge to the Discussion. The result is not simply “text wins.” The more useful result is that GenAI-generated interferences do not have the same effect depending on how the human ideator encounters them.

5.4.5 Human validation of FixShift coding

Because FixShift was applied to the full post-interference idea dataset through an AI-assisted workflow, a human validation step was needed. The objective was not to prove that AI-assisted coding is equivalent to full human judgment. The objective was more practical: to assess whether the AI-assisted FixShift scores were sufficiently aligned with human judgment to support comparative pattern analysis.

A blind validation subset of 60 post-interference ideas was created. The subset was balanced across text interference, image mosaic interference, video mosaic interference, neutral break, and AI-assigned FixShift levels. Three human raters independently applied the same 0 to 2 FixShift rubric. Their judgments were then aggregated and compared with the AI-assisted FixShift scores.

The validation results show moderate but usable alignment. Exact agreement between the AI-assisted scores and the rounded aggregated human judgment reached 60.0%. Agreement within one category reached 96.7%. The mean absolute difference from the aggregated human mean was 0.44 on the 0 to 2 scale. The quadratic weighted kappa was 0.61.

These results support the use of AI-assisted FixShift as a practical full-dataset proxy. The alignment is not perfect, and it should not be presented as such. But it is strong enough for the purpose of this paper, which is to compare directional patterns across interferences and the neutral break. The high within-one-category agreement is especially important because most disagreements remain small on the FixShift scale.

This validation step also defines the boundary of the claim. FixShift is useful here because it makes full-dataset comparison possible. It helps evaluate whether post-interference ideas move away from each participant's pre-interference idea families. At the same time, the AI-assisted scores should be read as comparative evidence, not as exact human-equivalent judgment for every individual idea.

5.4.6 Robustness and interpretation checks

Because fluency delta and mean FixShift distributions were not assumed to be normal, the comparison was checked with nonparametric exploratory tests. Kruskal-Wallis tests were used to assess whether participant-level fluency delta and participant-level mean FixShift differed across text interference, image mosaic interference, video mosaic interference, and neutral break. These checks are used to support interpretation of directional patterns, not to claim universal causal superiority of one interference modality.

For fluency delta, the exploratory Kruskal-Wallis test did not show a statistically significant omnibus difference across post-fixation alternatives, $H(3) = 4.97$, $p = 0.174$. A formula-based baseline-adjusted regression predicting post-fixation fluency from pre-fixation fluency and post-fixation alternative led to the same interpretation. Adding the post-fixation alternatives did not significantly improve model fit, $F(3, 133) = 0.71$, $p = 0.549$. This reinforces the view that the fluency evidence should be read as descriptive rather than confirmatory.

For FixShift, the Kruskal-Wallis test approached the conventional threshold but remained exploratory, $H(3) = 7.55$, $p = 0.056$. Exploratory pairwise comparisons summarized with Cliff's delta showed a negligible difference between text interference and image mosaic interference, $\delta = 0.09$, and small-to-moderate directional differences between text interference and neutral break, $\delta = 0.25$, and between text interference and video mosaic interference, $\delta = 0.37$. Other directional differences were smaller: image mosaic interference versus neutral break, $\delta = 0.15$; image mosaic interference versus video mosaic interference, $\delta = 0.26$; and neutral break versus video mosaic interference, $\delta = 0.15$. The largest directional contrast is therefore between text interference and video mosaic interference on FixShift.

Taken together, these checks are consistent with the descriptive ordering reported above, especially for FixShift, but they do not make the study confirmatory. The evidence remains directional and comparative. It supports the argument that interference modality matters in this use case, while preserving the boundary that the study is not a fully randomized causal test of modality superiority.

5.5 Discussion

The purpose of this study was not to ask whether GenAI can generate better ideas than humans. It was to test a different role for GenAI in innovation work. In this role, GenAI generates project-specific interferences when human ideation stalls. The human remains the ideator.

The results are consistent with that shift in perspective. The evidence is not a universal ranking of text, image, video, and break. Under the present public-library brief, text interference shows the strongest directional pattern on both fluency and FixShift, followed by image mosaic interference, neutral break, and video mosaic interference.

This pattern matters because it says something simple about GenAI-assisted innovation management. More AI, more media, and more richness are not automatically better. The value of GenAI may depend on whether the interference arrives in a usable form, at the right moment, for a human who is trying to ideate again.

5.5.1 GenAI as ideation infrastructure, not direct ideator

The first contribution is to reposition GenAI in the ideation process. In much of the practical discussion, GenAI is treated as a direct ideator: the human asks, the system proposes, and the human selects.

That use is legitimate, but it is not the role studied here. In this study, GenAI is not asked to produce the ideas that should solve the problem. It is asked to produce interferences. These interferences are introduced after participants report that they have run out of ideas. Their role is not to replace human ideation. Their role is to disturb the current path so that human ideation can restart.

This distinction is important for innovation management research because it moves the paper away from the usual automation question. The question is not whether GenAI can replace the ideator. The question is whether GenAI can support the innovation process by creating useful interference at the moment when ideation slows. This is close to the broader movement from automation to augmentation in AI research, but with a specific focus on ideation defixation (Pedota et al., 2025; Raj et al., 2026).

The point is not that humans are always outside the box and AI is always inside it. That would be too simple. The point is more practical. GenAI works from what it has learned and from what it is given. It can recombine, translate, summarize, visualize, and vary that material at remarkable speed. But innovation sometimes begins when a human decides that the given frame is not enough. In this study, GenAI is useful because it creates a disturbance inside the ideation process. The human still has to use that disturbance to leave the first idea family.

In that sense, GenAI becomes an ideation infrastructure. It is not the creative agent in the room. It is part of the method that helps the human ideator continue. The infrastructure is useful because it can generate project-specific, non-prescriptive material quickly. It can do so from the same brief, with different modalities, and at low cost. This gives innovation managers and facilitators a new way to manage the moment of fixation. It also complements recent work showing that GenAI's value in innovation depends less on simple output generation than on how it is embedded into human creative processes, work practices, and innovation routines (Chu, Baxter, & Liu, 2025 ; Roberts & Candi, 2024 ; Tang et al., 2026).

5.5.2 Defixation as a human-centered innovation-process method

The second contribution is to treat defixation as a manageable innovation-process method. Fixation is often described as a cognitive phenomenon. That is useful, but it is not enough for practice. In an innovation workshop, the problem is very concrete. People have generated ideas. Then the room slows down. The ideas become similar. Participants say they have nothing more to add.

At that point, a facilitator needs a practical move. A neutral break is one possible action. It interrupts the work, but it does not add new material. A GenAI-generated interference is another possible action. It interrupts the ideation path while staying connected to the project brief.

The method tested here is simple. Let humans generate first. Wait until they report fixation. Then introduce a project-specific interference and ask them to ideate again. The method does not remove human agency. GenAI prepares the disturbance. Humans do the ideation.

This is why the contribution is not only the result pattern. The paper also proposes a possible way to use GenAI in innovation work. GenAI-assisted defixation can become a timed, project-specific, low-cost method for relaunching human ideation. It is not a complete innovation process. It is a tactical move inside the process.

5.5.3 Interference modality and the limits of richer media

The results suggest that interference modality matters. Text interference, image mosaic interference, and video mosaic interference do not perturb the human ideator in the same way. They do not ask for the same type of attention. They do not leave the same amount of interpretive work to the participant.

Text interference appears to be the most effective in this use case. One possible explanation is its cognitive economy. A short text interference can be read quickly. It can be interpreted flexibly. It can suggest a new direction without imposing a full representation. It gives the human ideator something to work with, but it does not do too much of the work.

Image mosaic interference offers a different type of perturbation. It can open visual associations and suggest atmospheres, objects, places, or uses that text may not activate as easily. In the present data, it remains positive, but weaker than text. This may be because images are richer, but also more specific. A visual element can help. It can also anchor attention on surface features.

The video mosaic interference is the most useful caution in the results. It is richer than text and image. It brings movement, sequence, atmosphere, and more cues. Yet it does not produce the strongest relaunch pattern. In this use case, it appears weaker than the simpler alternatives. This suggests that the question is not only what GenAI can generate. The question is what the human ideator can use after fixation.

The neutral break also matters. It reminds us that interruption alone can have value. However, it does not offer project-specific material. This is why the comparison with text interference and image mosaic interference is useful. It helps separate a simple pause from a pause plus a project-specific perturbation.

The broader message is not that video mosaic interferences are generally weaker. That would be too strong. The message is more useful for innovation management: richer material is not automatically better. After fixation, the participant may need a small disturbance, not a rich experience. The goal is not to admire the interference. The goal is to use it to produce the next idea.

5.5.4 Implications for GenAI-assisted innovation management

The practical implication is not that every innovation team should use text interference and avoid video. That would be too strong. The implication is more careful.

First, innovation managers should not treat GenAI only as a direct idea generator. GenAI can also help prepare the moment where human ideation needs to restart. This extends recent work on AI in innovation management by showing a small but practical role for GenAI inside the ideation process: not producing the solution, but preparing a disturbance that helps human ideation continue (Roberts & Candi, 2024; Pedota et al., 2025; Pescher & Tellis, 2025).

Second, facilitators should treat interference modality as a design choice. Prior work on creativity-support architectures shows that the way creative support is structured matters for creative output, especially when task knowledge varies (Ozturk, Han, & Ren, 2026). Text interference, image mosaic interference, and video mosaic interference do not do the same work. A text interference may be more appropriate when the goal is a fast, low-friction relaunch. An image mosaic interference may be more useful when the group needs visual association. A video mosaic interference may be useful in other settings, but it may require more time, framing, or facilitation than the present protocol allowed.

Third, the neutral break should not be dismissed. Sometimes people need interruption before they need more material. In practice, the useful question may be one of sequence. A facilitator may first use a break, then introduce an interference. Or the facilitator may use a text interference when the group still has energy, and reserve richer media for later exploration.

Fourth, managers should evaluate GenAI-assisted ideation with more than idea count. Fluency matters because ideas must exist before they can be selected. But fluency does not say whether the new ideas move away from the old idea families. FixShift helps address that missing part.

It gives a way to ask whether ideation only restarted, or whether it also moved. This is consistent with broader work on ideation metrics and semantic distance, where quantity and conceptual movement are treated as different dimensions of creative output (Shah et al., 2003; Nelson et al., 2009; Orwig et al., 2021).

A further implication concerns the input to the interference. This study uses a brief-based stimulus chain because it gives a clean comparison across text interference, image mosaic interference, and video mosaic interference. In practice, innovation teams often have more material before ideation begins. They may have empathy interviews, field observations, user complaints, stakeholder maps, journey maps, service data, or community feedback. GenAI is especially useful at synthesizing such material and generating from information provided to it. Future GenAI-assisted defixation methods could therefore move from brief-based interferences to empathy-enriched interferences. The brief-based version increases comparability; the empathy-enriched version may increase relevance. That may make the disturbance more human-centered and more project-specific, provided the human team remains responsible for interpreting the material and producing the ideas.

The broader message is practical. GenAI-assisted innovation management should not be reduced to asking better prompts for better AI ideas. This is also consistent with recent work treating GenAI as part of a guided design or creativity-support system, rather than only as an output generator (Abrusci, Dabaghi, D’Urso, & Sciarrone, 2025 ; Fang et al., 2025). It should also include the design of moments, cues, and interferences that help humans continue their own ideation.

5.5.5 Limits and boundaries of the findings

These findings should be read with restraint. The study uses one innovation brief, a public-library case. The participants are Master’s-level students rather than professional innovation teams. The empirical structure is multi-site, but not fully crossed. Each participant experienced only one interference modality or the neutral break. The results are therefore comparative pattern evidence, not strict causal proof.

The measures also define the boundaries of the paper. Fluency captures whether participants generated again after fixation. It does not establish idea quality, originality, usefulness, feasibility, or implementation value. FixShift adds a measure of conceptual movement, but the full-dataset coding is AI-assisted. The human validation supports its use as a practical proxy, but it does not make it equivalent to full human judgment for every idea.

The generated interferences are also tied to the GenAI tools, prompts, and visual choices used in the study. Different prompts, different models, different image styles, or different video designs may produce different results. This is especially important for the video mosaic. The present finding may say less about video as a general modality than about how this video mosaic functioned under this timing and this brief.

The input structure is also a boundary. The present study uses a brief-based stimulus chain: the text interference was generated from the public-library brief, and the image and video mosaics were generated from the project-specific text-interference list. This keeps the comparison clean, but it also simplifies what happens in real design-thinking projects. In practice, ideation usually follows empathy and definition work. The team may have interviews, observations, personas, journey maps, stakeholder tensions, user complaints, and other material. A richer GenAI-assisted defixation method could use that material as input. The present study does not test that richer version. It tests a first brief-based version.

Future research should test the method with professional teams, different briefs, larger and more balanced samples, repeated use across ideation moments, and alternative text, image, and video designs. The next step is not only to ask which interference modality is best. It is to understand what kind of interference is useful for which fixation moment, with which participants, and for which innovation task. Future work should also test empathy-enriched interferences generated from interviews, observations, personas, journey maps, user feedback, or service data. This would move the method from project-specific to insight-specific.

The central lesson remains simple. GenAI does not need to be the ideator in the room to matter. Sometimes its value is to create the right interference, at the right moment, so that humans can produce the next idea.

5.6 Conclusion

This paper contributes at three connected levels: a method, a comparison, and a shift in the role assigned to GenAI in innovation ideation. First, it proposes a practical method for using GenAI as ideation infrastructure: not to produce the ideas, but to generate project-specific interferences that can be introduced when human ideation stalls. In that role, GenAI does not replace the ideator. It creates a temporary perturbation, while the human remains responsible for producing, interpreting, selecting, and developing the ideas.

Second, the paper provides comparative evidence that the interference modality matters for both post-fixation fluency and FixShift. In this use case, text interference shows the strongest directional relaunch pattern, followed by the image mosaic interference, the neutral break, and the video mosaic interference. This pattern matters because it suggests that relaunching ideation is not only a question of adding more media, more signal, or more AI. A short text interference may sometimes be more usable than a richer visual or dynamic interference because it perturbs without overloading, and because it leaves more interpretive work to the human ideator.

Third, the paper contributes to GenAI-assisted innovation management by turning defixation into a manageable innovation-process method. The managerial question is not only how to start ideation, but also how to restart it when the room, the team, or the individual reaches a creative block. Project-specific GenAI interferences offer one possible answer: they are low-cost, scalable, brief-specific, and can be introduced at the moment when ideation slows without asking GenAI to become the source of the solution.

The findings should be read with restraint. They do not establish a universal ranking of text, image, video, and break. The evidence comes from one innovation brief, a Master's-level participant sample, a multi-site structure that is not fully crossed, and a protocol where each participant experienced only one interference modality or the neutral break. FixShift also relies on AI-assisted full-dataset coding, even though that coding was checked against a balanced human validation subset. The study also uses a constrained brief-based stimulus chain as GenAI input. That was useful for comparison, but it is not the richest possible version of the

method. Future research should therefore test the method with professional teams, multiple briefs, larger and more fully crossed samples, richer outcome measures, including originality, usefulness, feasibility, and creativity assessment (Sarkar & Chakrabarti, 2011), and empathy-enriched interferences built from interviews, observations, personas, journey maps, user feedback, or service data.

The practical lesson is simple: GenAI does not need to be the ideator in the room to matter. Sometimes its value is to create the right interference, at the right moment, so that humans can produce the next idea.

CHAPITRE 6

CONCLUSION GÉNÉRALE

Cette thèse part du constat que l'idéation s'essouffle après une première production d'idées, notamment parce qu'elle se heurte au phénomène de fixation, qui réduit l'exploration des espaces de solution. Elle étudie l'IA générative non comme un substitut au designer, mais comme une infrastructure d'idéation capable de produire des interférences de défixation spécifiques au projet. La recherche s'articule autour de trois études successives, allant des suggestions verbales orientées à une comparaison texte-image-pause neutre, puis à une comparaison multimodale texte-mosaïque d'images-mosaïque vidéo-pause témoin.

6.1 Synthèse des résultats

Les trois études ne sont pas juxtaposées et ne constituent pas trois analyses redondantes d'un même effet. Elles construisent progressivement un même raisonnement : d'abord vérifier que des suggestions générées avec GenAI peuvent jouer un rôle d'interférence, ensuite comparer texte, image et pause neutre sur la fluence dans le cas des bibliothèques publiques, puis étendre cette comparaison à un protocole multimodal intégrant la mosaïque vidéo et le FixShift. Le maintien du brief des bibliothèques publiques dans les études 2 et 3 permet de stabiliser le problème de design, tandis que les modalités des interférences comparées et les mesures mobilisées évoluent.

Le premier article montre qu'il est possible de générer, à partir d'un problème donné, des suggestions textuelles orientées comme interférences à la fixation et d'en comparer la valeur perçue au regard des enjeux de connaissance et de générativité.

Le deuxième article montre, sur le cas des bibliothèques publiques, que les interférences textuelles et image spécifiques au projet présentent un signal directionnel positif de relance de l'idéation par rapport à une pause neutre, avec un signal plus marqué pour le texte. Cette deuxième étude doit être lue comme une comparaison pratique intermédiaire, utile pour

distinguer l'effet de deux modalités d'interférence générées avec GenAI d'une interruption sans contenu sémantique, mais encore insuffisante pour soutenir des inférences fortes.

Le troisième article montre que, dans le protocole comparatif étudié, l'interférence textuelle présente le signal directionnel le plus fort de relance et de déplacement conceptuel, suivie de la mosaïque d'images. La pause témoin occupe une position intermédiaire, notamment pour le FixShift, tandis que la mosaïque vidéo apparaît la moins efficace dans ce protocole. Ce résultat doit toutefois être compris comme une configuration empirique observée, qui indique que les médias génératifs plus riches ne sont pas automatiquement plus utiles, et non comme une hiérarchie universelle des modalités d'interférence.

Pris ensemble, les trois articles montrent que les interférences n'ont pas toutes le même potentiel de relance. Leur effet dépend à la fois de leur construction sémantique, de leur forme de présentation et du moment où elles sont introduites dans l'idéation.

La thèse ne débouche donc pas sur une solution unique, mais sur une compréhension plus fine des conditions dans lesquelles des interférences peuvent relancer l'idéation après fixation.

6.2 Contribution de la thèse

D'un point de vue théorique, la contribution centrale de la thèse est de déplacer le rôle attribué à GenAI. L'enjeu n'est pas de montrer que l'IA générative peut produire les idées à la place du designer. L'enjeu est de montrer qu'elle peut soutenir les conditions de l'idéation humaine, en produisant des interférences capables de rouvrir l'exploration lorsque celle-ci se referme. La thèse contribue ainsi aux travaux sur la fixation, la défixation et l'idéation assistée par GenAI. Elle apporte un cadre comparatif qui distingue la nature de l'interférence de sa modalité de présentation.

Sur le plan méthodologique, elle articule la DSRM et la CCS dans un programme doctoral progressif, plutôt que dans une seule étude ponctuelle. Elle montre comment concevoir des interférences spécifiques au projet, les instancier dans plusieurs formats et les comparer de manière structurée. Elle introduit la combinaison de deux mesures complémentaires : la fluence et le FixShift. La fluence rend compte de la reprise de production d'idées. Le FixShift rend compte du déplacement conceptuel par rapport aux familles d'idées initiales d'un même

participant. Cette combinaison ne permet pas de mesurer directement la qualité globale des idées ni la créativité au sens large. Elle permet plus modestement de distinguer une reprise de production d'un déplacement conceptuel post-interférence. Par ailleurs, la validation humaine du codage assisté par IA permet de soutenir l'usage d'un proxy pratique sur l'ensemble du jeu de données, tout en maintenant une interprétation prudente.

Sur le plan pratique, la thèse ne propose pas un « meilleur format universel », mais un cadre d'usage raisonné pour produire, sélectionner et comparer des interférences de défixation adaptées au projet. Concrètement, la thèse suggère qu'un praticien peut mobiliser des interférences générées à partir du brief lorsque l'idéation commence à tourner autour des mêmes idées. Les résultats invitent à commencer par des interférences simples, notamment textuelles ou image. Les modalités plus riches, comme la vidéo, peuvent ensuite être explorées, mais leur effet ne doit pas être supposé supérieur simplement parce qu'elles sont plus riches.

6.3 Limites

La thèse demeure exploratoire. Le programme repose sur un nombre limité d'interférences, de cas et de configurations expérimentales, ce qui borne la portée de généralisation des résultats. La comparaison multimodale repose sur une structure multi-site et non pleinement croisée; elle ne permet donc pas d'attribuer les écarts observés à des effets causaux purs de la modalité. Elle permet plutôt d'identifier des tendances comparatives à l'échelle du protocole retenu. Le brief unique sur les bibliothèques publiques stabilise la comparaison des interférences, mais limite le transfert direct des résultats à d'autres problèmes de design. Enfin, les participants sont majoritairement des étudiants de niveau maîtrise, et non des designers professionnels en contexte industriel.

Concernant les mesures, le FixShift appliqué à l'ensemble des données repose sur un codage assisté par IA, validé humainement, mais non équivalent à une vérité de terrain. La fluence ne doit pas être confondue avec la qualité des idées, pas plus que FixShift ne doit être interprété comme un score général de créativité.

Enfin, l'évolution rapide des outils d'IA générative rend toute conclusion temporellement située.

Ces limites n'invalident pas la cohérence du programme, mais elles imposent de lire les résultats comme des preuves comparatives directionnelles, situées et interprétatives, plutôt que comme des prescriptions universelles.

6.4 Pistes de recherche

La première piste de recherche consiste à étendre la comparaison à des cohortes plus larges et plus pleinement croisées, afin de réduire davantage les risques de confusion. Il faudrait également répliquer les études sur d'autres problèmes de design, au-delà du brief sur les bibliothèques publiques, afin de tester la robustesse des tendances observées. Il serait également utile d'intégrer des designers professionnels et des contextes industriels réels.

Une autre piste de recherche importante concerne la réplication de la validation humaine de FixShift sur des sous-échantillons plus larges et sur d'autres jeux de données. Il serait pertinent d'examiner si le pattern observé entre texte, image, vidéo et interruption neutre se maintient sous d'autres règles de durée et d'autres protocoles post-fixation.

Il existe d'autres formes d'interférences spécifiques au projet à explorer, y compris d'autres combinaisons multimodales ou des formes audio et visuelles plus ciblées.

Une piste plus fine concerne la forme optimale des interférences verbales, par exemple leur longueur, leur densité sémantique ou leur caractère adaptatif au problème.

Il serait utile d'intégrer plus finement les informations issues des phases amont du Design Thinking, notamment l'empathie et la définition, afin de conditionner encore plus précisément les interférences générées.

L'enjeu n'est pas d'automatiser la créativité, mais de développer des manières plus fines de soutenir l'idéation humaine lorsqu'elle se fixe ou s'essouffle.

En définitive, cette thèse ne propose pas une recette universelle de défixation, mais une manière de concevoir et de comparer des interférences capables de relancer l'idéation lorsqu'elle s'essouffle. Elle montre que l'IA générative peut aider le designer à

produire davantage d'idées et à explorer d'autres directions, dès lors qu'elle est mobilisée comme infrastructure d'idéation plutôt que comme substitut créatif. GenAI peut ainsi agir indirectement : produire, au bon moment, l'interférence située qui aide le designer à rouvrir l'espace des possibles.

ANNEXE I

FORMULAIRE TYPE ET VARIANTES D'INTERFÉRENCES

I.1 Objet de l'annexe

Cette annexe présente la structure des formulaires mobilisés dans les trois études. Elle décrit aussi les principales variantes d'interférences utilisées dans le programme doctoral. Son objectif est de documenter la collecte des réponses et l'introduction des interférences au moment où l'idéation s'essouffle. Elle ne reproduit pas l'intégralité des formulaires bruts, mais fournit une version représentative de leur logique de fonctionnement, en cohérence avec les protocoles décrits dans les chapitres 3, 4 et 5. Les formulaires diffèrent selon les études, mais conservent un principe commun : une première phase d'idéation, l'identification d'un point d'essoufflement, l'introduction d'une interférence, puis une seconde phase d'idéation.

I.2 Formulaire type de l'étude 1

La première étude repose sur un questionnaire de préférence destiné à comparer trois listes de suggestions verbales orientées. Le participant ne produit pas lui-même d'idées à partir d'un problème donné. Il évalue plutôt, pour plusieurs problèmes de Design Thinking, l'intérêt relatif de trois listes A, B et C. Ces listes visent deux limites de l'idéation : la limite de connaissance et la limite de générativité. Les trois listes comparées sont :

- Liste A : suggestions aléatoires liées au problème;
- Liste B : suggestions liées à des connaissances relevant de la discipline du problème;
- Liste C : suggestions liées à des connaissances extérieures à la discipline du problème.

Consigne type

Pour chacun des problèmes ci-dessous, veuillez classer les listes A, B et C par ordre décroissant d'intérêt pour atténuer la limite de connaissance, puis la limite de

générativité. Indiquez, pour chaque problème, la liste que vous jugez la plus intéressante, la suivante, puis la moins intéressante.

Problèmes utilisés dans l'étude 1

1. How can emergency departments redesign their processes to efficiently manage patient flow, reduce wait times, and improve satisfaction within existing resource constraints?
2. What innovative pedagogical strategies and technologies can higher education institutions implement to significantly increase student engagement and satisfaction?
3. How can affordable, sustainable lighting solutions be developed and made accessible to low-income families using local resources and renewable energy?
4. How can public libraries innovate their services and spaces to meet evolving community needs and increase user engagement in the digital age?
5. What innovative, economically and environmentally sustainable water management solutions can be tailored to meet the specific needs of rural communities?

Gabarit de réponse

Problème 1 — Limite de connaissance : A / B / C

Problème 1 — Limite de générativité : A / B / C

Problème 2 — Limite de connaissance : A / B / C

Problème 2 — Limite de générativité : A / B / C

Problème 3 — Limite de connaissance : A / B / C

Problème 3 — Limite de générativité : A / B / C

Problème 4 — Limite de connaissance : A / B / C

Problème 4 — Limite de générativité : A / B / C

Problème 5 — Limite de connaissance : A / B / C

Problème 5 — Limite de générativité : A / B / C

I.3 Formulaire type des études 2 et 3

Les deuxième et troisième études reposent sur un formulaire individuel construit selon une logique commune. Tous les participants reçoivent d'abord le même énoncé de problème, puis produisent une première série d'idées. Lorsqu'ils déclarent avoir épuisé leurs idées, une interférence leur est présentée. Ils poursuivent ensuite l'idéation dans une seconde phase. Cette structure permet de maintenir un protocole homogène entre conditions, tout en faisant varier uniquement la nature ou la modalité de l'interférence introduite.

Le formulaire complet comprenait également des sections d'identification, de consentement et de navigation propres à l'outil de collecte utilisé. Ces éléments ne sont pas reproduits intégralement dans la présente annexe, car ils relèvent davantage de la gestion éthique et administrative de la collecte que de la logique méthodologique centrale de la thèse. Seule la structure fonctionnelle utile à la compréhension du protocole est retenue ici.

Structure fonctionnelle retenue

- présentation du problème de design;
- première phase d'idéation;
- signalement de l'essoufflement;
- exposition à une seule condition d'interférence;
- deuxième phase d'idéation;
- soumission finale.

Extrait représentatif

Problème de design

How can public libraries innovate their services and spaces to meet evolving community needs and increase user engagement in the digital age?

Première phase

Veillez écrire toutes vos idées, une idée par ligne, jusqu'à ce que vous n'ayez plus de nouvelles idées.

Déclenchement

Lorsque, et seulement lorsque, vous avez épuisé vos idées, passez à la section suivante.

Deuxième phase

Après exposition à une seule interférence, veuillez écrire toutes vos nouvelles idées, une idée par ligne, jusqu'à épuisement.

I.4 Variantes d'interférences utilisées dans les études 2 et 3

Les variantes d'interférences diffèrent selon les études, mais conservent une même logique : intervenir après l'essoufflement de l'idéation sans prescrire directement une solution. Dans l'étude 2, chaque participant est exposé à une seule des trois perturbation suivantes :

- une interférence textuelle spécifique au projet;
- une interférence image spécifique au projet;
- une pause neutre sans contenu sémantique.

Dans l'étude 3, chaque participant est exposé à une seule des quatre conditions suivantes :

- une interférence textuelle spécifique au projet;
- une mosaïque d'images spécifique au projet;
- une mosaïque vidéo spécifique au projet;
- une interruption neutre d'une minute sans contenu sémantique.

L'étude 2 compare donc deux modalités générées avec GenAI à une interruption neutre, tandis que l'étude 3 étend cette logique à une comparaison multimodale incluant la vidéo et le FixShift.

À des fins d'illustration, les interférences des études 2 et 3 sont décrites ici selon leur fonction expérimentale plutôt que selon l'intégralité de leurs contenus. L'étude 2 compare une interférence textuelle spécifique au projet, une interférence image spécifique au projet et une pause neutre. L'étude 3 étend cette logique à une interférence textuelle, une mosaïque d'images, une mosaïque vidéo et une interruption neutre d'une minute.

I.5 Fonction méthodologique du formulaire

Dans les trois études, le formulaire ne sert pas seulement à collecter des réponses. Il constitue une partie intégrante du dispositif expérimental, puisqu'il organise la séquence pré-interférence, le signalement du point d'essoufflement, l'exposition à une condition précise, puis la collecte post-interférence. À ce titre, il participe directement à la comparabilité des conditions et à la cohérence méthodologique du programme doctoral.

ANNEXE II

PROMPTS DE GÉNÉRATION DES INTERFÉRENCES

II.1 Objet de l'annexe

Cette annexe présente les principales consignes de génération utilisées pour produire les interférences dans les trois études du programme doctoral. Lorsque les formulations exactes sont documentées dans les articles, elles sont reproduites ici. Lorsque le manuscrit documente surtout la logique de génération sans conserver mot à mot l'invite initiale, l'annexe restitue cette logique sous une forme synthétique et fidèle. L'objectif est de rendre explicite la manière dont les artefacts d'interférence ont été générés, filtrés et instanciés.

II.2 Prompts de l'étude 1 — suggestions verbales orientées

La première étude vise à générer trois listes de suggestions verbales, chacune composée de dix suggestions de trois à cinq mots. Le problème de Design Thinking est inséré dans l'invite lorsque cela est nécessaire.

Prompt A : suggestions aléatoires liées au problème

- Generate a list of 10 random 3 to 5 related words.

Prompt B : connaissances liées à la discipline du problème

- With the following definition of the problem: XXX

Generate a list of 10 of 3 to 5 words presenting a knowledge related to the discipline of the problem.

Prompt C : connaissances extérieures à la discipline du problème

- With the following definition of the problem: XXX

Generate a list of 10 of 3 to 5 words presenting a knowledge not related to the discipline of the problem.

Note d'instanciation

Dans ces trois prompts, la variable XXX est remplacée par l'une des définitions de problème retenues pour l'étude 1. Les listes obtenues sont ensuite nettoyées de leur numérotation automatique et intégrées au questionnaire de classement forcé.

II.3 Problèmes utilisés comme base de génération dans l'étude 1

Les cinq définitions de problème utilisées dans l'étude 1 sont les suivantes :

1. How can emergency departments redesign their processes to efficiently manage patient flow, reduce wait times, and improve satisfaction within existing resource constraints?
2. What innovative pedagogical strategies and technologies can higher education institutions implement to significantly increase student engagement and satisfaction?
3. How can affordable, sustainable lighting solutions be developed and made accessible to low-income families using local resources and renewable energy?
4. How can public libraries innovate their services and spaces to meet evolving community needs and increase user engagement in the digital age?
5. What innovative, economically and environmentally sustainable water management solutions can be tailored to meet the specific needs of rural communities?

II.4 Logique de génération et d'instanciation de l'étude 2 : texte, image et pause neutre

La deuxième étude reprend le problème commun de l'innovation des bibliothèques publiques. Elle compare deux interférences générées avec GenAI, soit une interférence textuelle spécifique au projet et une interférence image spécifique au projet, à une pause neutre sans contenu sémantique. Les matériaux textuels et visuels sont produits avant la collecte à partir du même brief, puis conservés comme matériaux fixes de l'étude. Ce choix permet de comparer les conditions de manière stable, sans introduire de variation supplémentaire liée à une génération en temps réel.

Énoncé de problème utilisé

How can public libraries innovate their services and spaces to meet evolving community needs and increase user engagement in the digital age?

Condition texte : logique de génération

À partir du brief sur les bibliothèques publiques, générer de courts éléments verbaux spécifiques au projet. Ces éléments doivent suggérer des directions alternatives sans proposer de solutions complètes. Leur fonction est de produire une perturbation sémantique brève, susceptible de relancer l'idéation après fixation, tout en maintenant la génération d'idées du côté du participant.

Condition image : logique de génération

À partir du même brief et de la même logique de contenu que la condition texte, générer une interférence image spécifique au projet. L'image doit rester liée au champ sémantique du brief, mais changer la forme de présentation de l'interférence, en passant du verbal au visuel. Sa fonction est de produire une perturbation image et associative, sans représenter une solution finale et sans prescrire une réponse attendue.

Condition témoin : logique d'instanciation

Prévoir une pause neutre sans contenu sémantique ni visuel ajouté. Cette condition ne constitue pas une interférence générée avec GenAI. Elle sert de point de comparaison pour distinguer l'effet d'une simple interruption de celui d'une interférence spécifique au projet.

II.5 Logique de génération et d'instanciation de l'étude 3 : texte, mosaïque d'images, mosaïque vidéo et pause neutre

La troisième étude reprend le même brief sur l'innovation des bibliothèques publiques, mais étend la comparaison à quatre conditions post-fixation : une interférence textuelle, une mosaïque d'images, une mosaïque vidéo et une pause neutre. Cette étude positionne GenAI comme une infrastructure d'idéation. GenAI n'est pas utilisé pour produire les idées à la place des participants, mais pour produire des interférences spécifiques au projet, introduites après le signalement de l'essoufflement de l'idéation.

Les trois interférences générées avec GenAI respectent une même logique : elles doivent être spécifiques au projet, non prescriptives et comparables entre elles. La liste textuelle est générée à partir du brief. La mosaïque d'images et la mosaïque vidéo sont ensuite générées à partir de cette même base textuelle, afin de maintenir un champ sémantique commun entre les modalités tout en variant leur forme de présentation. La pause neutre, quant à elle, ne contient aucun contenu sémantique ou visuel ajouté.

Énoncé de problème utilisé

How can public libraries innovate their services and spaces to meet evolving community needs and increase user engagement in the digital age?

Condition texte : logique de génération

À partir du brief sur les bibliothèques publiques, générer une liste de mots ou de courtes expressions spécifiques au projet. Cette liste doit offrir des indices sémantiques liés au problème, sans fournir de solution prête à l'emploi. Sa fonction est de créer une perturbation verbale à faible friction, que le participant peut lire rapidement, réinterpréter et utiliser comme point de départ pour produire ses propres idées.

Condition mosaïque d'images : logique de génération

À partir de la liste textuelle spécifique au projet, générer une mosaïque d'images spécifique au brief sur les bibliothèques publiques. La mosaïque doit constituer un seul panneau visuel. Elle doit rester suffisamment proche du champ sémantique du brief pour être pertinente, mais suffisamment indirecte pour ne pas devenir une proposition de solution. Sa fonction est de produire une perturbation image statique susceptible d'activer de nouvelles associations.

Condition mosaïque vidéo : logique de génération

À partir de la même liste textuelle spécifique au projet et de la même intention non prescriptive, générer une courte mosaïque vidéo. Cette condition prolonge la logique visuelle de la mosaïque d'images en y ajoutant mouvement, séquence et atmosphère. Sa fonction est de tester si une interférence plus riche et plus

dynamique peut relancer l'idéation après fixation, sans supposer que la richesse médiatique constitue nécessairement un avantage.

Condition témoin : logique d'instanciation

Prévoir une interruption neutre d'une minute, matérialisée par une horloge numérique de 60 à 0 secondes. Cette condition ne contient aucun contenu sémantique ou visuel ajouté. Elle sert de base de comparaison non sémantique pour évaluer si les interférences spécifiques au projet apportent quelque chose de plus qu'une simple pause.

II.6 Contraintes communes de génération

Dans les trois études, les interférences générées partagent un ensemble de contraintes communes :

- elles doivent être liées au problème ou au contexte de manière suffisante pour rester cognitivement pertinentes;
- elles doivent rester indirectes afin de ne pas prescrire directement une solution;
- elles doivent être comparables entre conditions;
- elles doivent pouvoir être intégrées à un protocole de relance post-fixation.

Le manuscrit précise aussi que, du fait de la nature probabiliste des outils génératifs, le contenu exact peut varier d'une génération à l'autre, tandis que la logique de génération, les contraintes et les règles d'instanciation demeurent documentables et reproductibles.

ANNEXE III

STRUCTURE MINIMALE DES DONNÉES ET VARIABLES D'ANALYSE

III.1 Objet de l'annexe

Cette annexe présente la structure des jeux de données, des variables et des procédures de validation mobilisés dans le programme doctoral. Elle vise à rendre explicite la chaîne de traitement des données utilisée dans les trois études, sans reproduire l'intégralité du classeur de travail. Elle documente plus particulièrement les données retenues pour les analyses de fluence et de FixShift, les règles de nettoyage appliquées, ainsi que la validation humaine du codage assisté par l'IA. Cette annexe complète le cadre méthodologique présenté au chapitre 2 et les résultats détaillés dans les chapitres-articles.

III.2 Jeux de données mobilisés et logique de traitement

Le traitement des données repose sur un classeur de travail structuré en plusieurs feuilles correspondant à différentes étapes de préparation, de synthèse et de validation. Toutes les feuilles du classeur ne sont pas mobilisées au même niveau dans la thèse. Certaines servent au nettoyage intermédiaire, d'autres à la synthèse des résultats, et d'autres encore à la validation humaine du codage. Dans la logique de cette annexe, seules les feuilles effectivement utiles à la compréhension des analyses sont présentées.

Le traitement général suit une chaîne simple. Les données brutes sont d'abord recueillies sous forme de réponses pré-interférence et post-interférence. Elles sont ensuite nettoyées, structurées et agrégées selon le niveau pertinent d'analyse.

Pour la fluence, les résultats sont principalement synthétisés au niveau des conditions expérimentales et des participants.

Pour le FixShift, le codage est d'abord réalisé au niveau de l'idée postérieure, puis réagrégué au niveau du participant afin d'éviter de traiter plusieurs idées provenant d'une même personne comme des observations indépendantes.

Cette distinction entre niveau idée et niveau participant est centrale pour l'interprétation des résultats.

Le classeur de données anonymisées comprend un ensemble de feuilles correspondant à des fonctions distinctes dans la chaîne de traitement. On y trouve notamment :

- des feuilles de collecte brute, regroupant les réponses pré-interférence et post-interférence;
- des feuilles de nettoyage intermédiaire, utilisées pour la préparation, les exclusions et la structuration des données;
- des feuilles de synthèse de la fluence, utilisées pour les effectifs, les comptages et les résumés par condition;
- des feuilles de synthèse du FixShift, utilisées pour l'agrégation au niveau participant;
- des feuilles de traçabilité du codage assisté par l'IA;
- des feuilles de validation humaine, comprenant le sous-échantillon, les consignes de codage et la compilation des jugements.

Dans la logique de cette annexe, seules les feuilles effectivement utiles à la compréhension des analyses sont retenues. Les feuilles intermédiaires de nettoyage ne sont pas reproduites intégralement, mais leur fonction est signalée lorsque cela est nécessaire à la compréhension du traitement des données.

III.3 Cohortes, sites et effectifs retenus

Les données du programme doctoral ne proviennent pas d'un seul contexte homogène. Dans la deuxième étude, les expérimentations ont été menées à plusieurs dates et sur plusieurs sites, ce qui permet d'observer les interférences dans des contextes variés, même si les cohortes relèvent toutes d'un niveau de maîtrise. Les expérimentations ont été conduites auprès de AIT Group A le 28 novembre 2024 à Bangkok et à distance, de AIT Group B le 30 novembre 2024 à Bangkok et à distance, de AIT VN le 9 mars 2025 à Ho Chi Minh-Ville et à distance, ainsi que de HKUST le 12 mars 2025 sur place. Dans cette deuxième étude, toutes les participations étaient volontaires et se déroulaient hors du temps d'enseignement.

Dans la troisième étude, le jeu de données comparatif comprend 138 participants et 983 lignes enregistrées. Après exclusion de deux non-idées dans la phase post-interférence, 981 lignes codées demeurent, réparties en 502 pré-idées et 479 post-idées. Au niveau des participants retenus pour les comparaisons de FixShift, les effectifs par modalité sont de 41 pour la condition textuelle, 42 pour la mosaïque d’images, 25 pour la mosaïque vidéo et 22 pour la condition témoin correspondant à une interruption neutre d’une minute. Pour la fluence, les effectifs par modalité sont plus élevés, soit 42 pour la condition textuelle, 46 pour l’image, 27 pour la vidéo et 23 pour la condition témoin. La différence entre les effectifs de la fluence et ceux du FixShift au niveau participant s’explique par le fait que le FixShift n’est calculé que pour les participants ayant au moins une idée postérieure exploitable pour le codage.

Tableau-A III-1 Cohortes, sites, dates et modalités d’exécution des études empiriques

Études concernées	Cohorte / site (étude 2) et code de cohorte (étude 3)	Date(s)	Modalités d’exécution
Études 2 et 3	AIT Group A – AIT SoM, Bangkok (U1-C1)	28 novembre 2024	À distance et sur site
Études 2 et 3	AIT Group B – AIT SoM, Bangkok (U1-C1)	30 novembre 2024	À distance et sur site
Études 2 et 3	AIT VN – Ho Chi Minh-Ville (U3-C4)	9 mars 2025	Texte et vidéo unique exploratoire
Études 2 et 3	HKUST (U4-C5)	12 mars 2025	Texte, image et vidéo unique exploratoire
Étude 3	U3-C6	11 septembre 2025	Image et condition témoin
Étude 3	U1-C7	27 novembre 2025	Image, vidéo mosaïque et condition témoin
Étude 3	U1-C7	29 novembre 2025	Image, vidéo mosaïque et condition témoin
Étude 3	U4-C8	23 février 2026	Image, vidéo mosaïque et condition témoin
Étude 3	U4-C8	25 février 2026	Image et condition témoin
Étude 3	U4-C8	2 mars 2026	Image et condition témoin

Note. Certaines cohortes ont contribué à la fois à l'étude 2 et à l'étude 3. Dans ces cas, le Tableau-A III-1 présente le libellé de site de l'étude 2 ainsi que le code de cohorte utilisé dans l'étude 3. Lorsqu'aucun libellé institutionnel explicite n'est retenu dans le manuscrit pour une cohorte propre à l'étude 3, seul le code de cohorte est conservé.

Les cohortes des deuxième et troisième études ont été mobilisées dans des contextes variés, tout en relevant majoritairement d'un niveau de maîtrise. Cette diversité de mise en œuvre permet d'observer les interférences dans plusieurs contextes pédagogiques, tout en maintenant un cadre de comparaison commun.

III.4 Variables principales et niveaux d'analyse

Les variables mobilisées dans la thèse ne se situent pas toutes au même niveau d'analyse. Certaines relèvent de la ligne de réponse, d'autres de l'idée, d'autres encore du participant ou de la condition expérimentale. Cette distinction est importante, car elle conditionne la manière dont les résultats peuvent être agrégés et interprétés.

Dans les études 1 et 2, les variables principales concernent respectivement les préférences exprimées au moyen du classement forcé et le nombre d'idées nouvelles produites après exposition à une interférence textuelle, à une interférence image ou à une pause neutre. Dans l'étude 3, deux familles de variables sont mobilisées. La première porte sur la fluence, entendue comme le nombre d'idées produites après l'interférence. La seconde porte sur le FixShift, qui mesure le degré de déplacement conceptuel entre les idées produites après l'interférence et les familles d'idées formulées avant celle-ci, propres à chaque participant. La fluence est donc un indicateur de réactivation, tandis que FixShift rend compte d'un changement plus structurel dans la recherche d'idées.

Tableau-A III-2 Variables principales et niveau d'analyse

Variable principale	Nom présent dans le classeur	Niveau d'analyse	Définition courte
Participant	Participant	Participant	Identifiant anonymisé du participant
Cohorte / site	Uni & Course	Participant / cohorte	Code de cohorte, de site ou de groupe d'appartenance
Condition expérimentale	Experiment	Participant / condition	Code de la condition expérimentale associée au participant
Modalité d'interférence	Modality	Participant / condition	Modalité reçue par le participant : texte, image, vidéo ou condition témoin
Pré-idées	Pre-fixation ou PRE_Idea_SET	Participant	Ensemble des idées formulées avant l'interférence
Post-idée	After interference, POST_IDEA ou Idea Txt	Idée	Idée formulée après l'interférence
Nombre d'idées pré	Pre_count	Participant	Nombre d'idées formulées avant l'interférence
Nombre d'idées post	Post_count	Participant	Nombre d'idées formulées après l'interférence
Écart pré / post	Delta_Count	Participant	Différence entre le nombre d'idées postérieures et le nombre d'idées antérieures
Donnée pré manquante	Missing_Pre	Participant	Indicateur signalant l'absence de contenu exploitable avant l'interférence
Nombre d'idées post codées	N_post_ideas	Participant	Nombre d'idées postérieures effectivement retenues pour le codage FixShift
Score FixShift assisté par l'IA	AI_FixShift_Score	Idée	Score de déplacement conceptuel attribué à une idée postérieure sur l'échelle 0–1–2
Justification du codage assisté par l'IA	AI_Rationale	Idée	Justification textuelle courte associée au score FixShift attribué par l'IA
Indicateur de non-idée	Non_idea_Flag	Ligne / idée	Marqueur utilisé pour exclure les lignes qui ne correspondent pas à des idées

Variable principale	Nom présent dans le classeur	Niveau d'analyse	Définition courte
Moyenne du FixShift	Mean_FixShift	Participant	Moyenne des scores FixShift attribués aux idées postérieures d'un même participant
Part des idées à FixShift = 2	%FixShift2	Participant	Proportion d'idées postérieures appartenant à une nouvelle famille d'idées
Identifiant du sous-échantillon humain	HumanSample_ID	Validation humaine	Identifiant interne du sous-échantillon utilisé pour la validation humaine
Identifiant aveugle	Blind_ID	Validation humaine	Identifiant anonymisé utilisé pour le codage humain en aveugle
Score humain final	Final_Human_Score	Validation humaine	Score agrégé final attribué par les évaluateurs humains

Les variables présentées dans le Tableau-A III-2 ne se situent pas toutes au même niveau d'analyse. Certaines relèvent de la ligne ou de l'idée, d'autres du participant, et d'autres encore du sous-échantillon de validation humaine. Cette distinction est importante, car elle conditionne la manière dont les résultats sont agrégés et interprétés. En particulier, la fluence est principalement analysée au niveau du participant et de la condition, tandis que le FixShift est d'abord attribué à chaque idée postérieure avant d'être réagrégué au niveau du participant. Les indicateurs agrégés utilisés dans les analyses comparatives par modalité ne constituent pas des variables élémentaires du jeu de données, mais des résumés calculés à partir des variables principales décrites ci-dessus. Pour la fluence, ces résumés comprennent notamment le nombre de participants, le total d'idées pré et postérieures, les médianes avec quartiles et l'écart moyen. Pour le FixShift, ils comprennent le nombre de participants disposant d'idées postérieures exploitables, la moyenne du FixShift par participant et la part des idées notées 2.

III.5 Règles de comptage, de nettoyage et de codage

Le comptage des idées repose sur une règle volontairement conservatrice : une ligne non vide correspond à une idée. Ce choix réduit l'espace laissé à des interprétations variables, rend le traitement plus visible et favorise la reproductibilité de la démarche. Il implique toutefois que

certaines idées formulées sur plusieurs lignes puissent être insuffisamment comptées, ce qui constitue une limite assumée de la mesure de fluence.

Le traitement des absences et du bruit a également été explicité. Les lignes qui ne correspondent pas à des idées sont exclues du codage FixShift. De même, les réponses postérieures laissées vides ne sont pas traitées comme des données manquantes, mais comme l'absence d'idée produite après l'interférence. Le codage FixShift repose sur une grille à trois niveaux

- 0 : même famille d'idées ou même logique centrale
- 1 : variation significative dans la même famille
- 2 : nouvelle famille d'idées ou nouveau chemin conceptuel

Cette gradation permet de distinguer la simple variation interne d'un déplacement conceptuel plus net, sans confondre défixation et créativité générale.

III.6 Validation humaine du codage FixShift

Le codage FixShift assisté par l'IA n'est pas traité comme une vérité de terrain, mais comme un indicateur de substitution nécessitant une validation humaine. Cette validation repose sur un sous-échantillon aveugle de 60 idées postérieures, équilibré selon les modalités et selon les niveaux de FixShift attribués par l'IA. Trois évaluateurs humains ont appliqué la même grille 0–1–2, à partir de consignes communes et sans recourir à des outils d'IA.

L'objectif n'était pas de montrer que l'IA remplace le jugement humain, mais d'évaluer si elle fournit un codage suffisamment cohérent pour permettre une lecture comparative directionnelle sur l'ensemble du jeu de données.

Les indicateurs de validation doivent être interprétés comme un faisceau convergent d'indices plutôt que comme une mesure unique. Dans le sous-échantillon de validation :

- l'accord exact atteint 60,0 %;
- l'accord à une catégorie près atteint 96,7 %;
- la différence absolue moyenne par rapport à la moyenne humaine agrégée est de 0,44;
- le coefficient kappa pondéré quadratique est de 0,61.

Pris ensemble, ces résultats soutiennent l'usage du codage assisté par l'IA comme indicateur pratique à l'échelle du jeu complet, tout en imposant une interprétation prudente des différences observées entre modalités.

III.7 Limites de structure des données

La structure des données impose plusieurs limites de lecture. Les modalités ne sont pas pleinement croisées entre les cohortes, ce qui empêche de présenter les différences observées comme des effets causaux purs. La fluence est rapportée au niveau du participant, alors que le FixShift est d'abord codé au niveau de l'idée, puis réagrégué au niveau du participant, ce qui explique que les effectifs diffèrent d'un tableau à l'autre. Enfin, une cellule vidéo exploratoire a été conservée dans le jeu de données élargi, mais exclue de la comparaison principale parce que sa taille était trop faible pour permettre une interprétation stable. Pour ces raisons, les résultats doivent être lus comme des tendances comparatives structurées plutôt que comme des estimations universelles.

LISTE DE RÉFÉRENCES BIBLIOGRAPHIQUES

- Abrusci, L., Dabaghi, K., D'Urso, S., & Sciarrone, F. (2025). *AI4Design*: A generative AI-based system to improve creativity in design—A field evaluation. *Computers and Education: Artificial Intelligence*, 8, 100401. <https://doi.org/10.1016/j.caeai.2025.100401>
- Atilola, O., & Linsey, J. (2015). Representing analogies to influence fixation and creativity: A study comparing computer-aided design, photographs, and sketches. *Artificial Intelligence for Engineering Design, Analysis and Manufacturing*, 29(2), 161-171. <https://doi.org/10.1017/S0890060415000049>
- Bartlett, L., & Vavrus, F. (2017). Comparative Case Studies: An Innovative Approach. *Nordic Journal of Comparative and International Education (NJCIE)*, 1(1). <https://doi.org/10.7577/njcie.1929>
- Baudoux, G., Guo, C., & Goucher-Lambert, K. (2025). Multimodal generative AI for conceptual design: enabling text-based and sketch-based human-AI conversations. Dans *Proceedings of the Design Society: International Conference on Engineering Design (ICED25)*. <https://doi.org/10.1017/pds.2025.10264>
- Bell, J. J., Pescher, C., Tellis, G. J., & Füller, J. (2024). Can AI Help in Ideation? A Theory-Based Model for Idea Screening in Crowdsourcing Contests. *Marketing Science*, 43(1), 54-72. <https://doi.org/10.1287/mksc.2023.1434>
- Bellesia, F., Bellis, P., & Palombi, G. (2026). Algorithms in the Realm of Fuzziness. How Generative Artificial Intelligence Is Changing the Front End of Innovation. *Creativity and Innovation Management*, caim.70048. <https://doi.org/10.1111/caim.70048>
- Bordas, A., Le Masson, P., & Weil, B. (2025). Switching Perspectives: Enhancing Generative Artificial Intelligence's Creativity, by Humans, With Design. *Creativity and Innovation Management*, 34(4), 1013-1033. <https://doi.org/10.1111/caim.70013>
- Brown, T. (2008). Design Thinking. *HARVARD BUSINESS REVIEW*, 86(6), 84-95.
- Brown, T., & Wyatt, J. (2009). Design Thinking for Social Innovation. *Stanford Social Innovation Review*, 8(1), 31-35. <https://doi.org/10.48558/58Z7-3J85>
- Carlgren, L., Rauth, I., & Elmquist, M. (2016). Framing Design Thinking: The Concept in Idea and Enactment. *Creativity and Innovation Management*, 25(1), 38-57. <https://doi.org/10.1111/caim.12153>
- Chan, J., Fu, K., Schunn, C., Cagan, J., Wood, K., & Kotovsky, K. (2011). On the Benefits and Pitfalls of Analogies for Innovative Design: Ideation Performance Based on Analogical Distance, Commonness, and Modality of Examples. *Journal of Mechanical Design*, 133(8), 081004. <https://doi.org/10.1115/1.4004396>

- Chen, L., Song, Y., Guo, J., Sun, L., Childs, P., & Yin, Y. (2025). How generative AI supports human in conceptual design. *Design Science*, 11, e9. <https://doi.org/10.1017/dsj.2025.2>
- Chu, W., Baxter, D., & Liu, Y. (2025). Exploring the impacts of generative AI on artistic innovation routines. *Technovation*, 143, 103209. <https://doi.org/10.1016/j.technovation.2025.103209>
- Crilly, N. (2015). Fixation and creativity in concept development: The attitudes and practices of expert designers. *Design Studies*, 38, 54-91. <https://doi.org/10.1016/j.destud.2015.01.002>
- Cristofaro, M., & Giardino, P. L. (2025). Human–AI synergy: finding cognitive balance in idea generation for product innovation. *European Journal of Innovation Management*, 1-23. <https://doi.org/10.1108/EJIM-03-2025-0312>
- De Sordi, J. O. (2021). *Design Science Research Methodology: Theory Development from Artifacts*. Cham : Springer International Publishing. <https://doi.org/10.1007/978-3-030-82156-2>
- Doshi, A. R., & Hauser, O. P. (2024). Generative AI enhances individual creativity but reduces the collective diversity of novel content. *Science Advances*, 10(28), eadn5290. <https://doi.org/10.1126/sciadv.adn5290>
- Eisenreich, A., Just, J., Gimenez-Jimenez, D., & Füller, J. (2024). Revolution or inflated expectations? Exploring the impact of generative AI on ideation in a practical sustainability context. *Technovation*, 138, 103123. <https://doi.org/10.1016/j.technovation.2024.103123>
- Eno, B., & Schmidt, P. (1975). *Oblique strategies: over one hundred worthwhile dilemmas* (5th rev. ed). (S.l.) : [The Authors].
- Fang, C., Zhu, Y., Fang, L., Long, Y., Lin, H., Cong, Y., & Wang, S. J. (2025). Generative AI-enhanced human-AI collaborative conceptual design: A systematic literature review. *Design Studies*, 97, 101300. <https://doi.org/10.1016/j.destud.2025.101300>
- Fetrati, M. A., Hansen, D., & Akhavan, P. (2022). How to manage creativity in organizations: Connecting the literature on organizational creativity through bibliometric research. *Technovation*, 115, 102473. <https://doi.org/10.1016/j.technovation.2022.102473>
- Fischer, F., & Gardoni, M. (2025a). Comparing Different ChatGPT4.0 Generated Oriented Word Suggestions as Design Fixation Interferences for the Ideation Phase in Design Thinking. Dans P. Sureephong, C. Danjou, & A. Bouras (Éds), *Product Lifecycle Management. Leveraging AI, Digital Twins, and Smart Technologies* (pp. 207-219). Cham : Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-93323-3_16

- Fischer, F., & Gardoni, M. (2025b, 11 juin). Comparison of three types of LLM-generated word interferences for the fixation phenomenon in design thinking: case study of libraries innovation [Poster presentation]. *Poster*. Repéré à <https://www.irdo.si/irdo2025/posters/78.pdf>
- Ge, W., & Hou, G. (2025). The effects of generative AI model type and visual stimuli type on design creativity. *Artificial Intelligence for Engineering Design, Analysis and Manufacturing*, 39, 1-13. <https://doi.org/10.1017/S0890060425100061>
- Gong, Z., Soomro, S. A., Wang, M., Latif, U. K., & Georgiev, G. V. (2024). Text Stimuli Created by Generative AI in Ideation: An Exploratory Study. Dans *Proceedings of NordDesign 2024* (pp. 496-503). The Design Society. <https://doi.org/10.35199/NORDDESIGN2024.53>
- Grilli, L., & Pedota, M. (2024). Creativity and artificial intelligence: A multilevel perspective. *Creativity and Innovation Management*, 33(2), 234-247. <https://doi.org/10.1111/caim.12580>
- Hoßbach, C., & Isaksen, S. G. (2025). AI-Augmented Approaches to Creative Problem-Solving: A Metacognitive Perspective. *Creativity and Innovation Management*, 34(4), 854-869. <https://doi.org/10.1111/caim.70003>
- Jansson, D. G., & Smith, S. M. (1991). Design fixation. *Design Studies*, 12(1), 3-11. [https://doi.org/10.1016/0142-694X\(91\)90003-F](https://doi.org/10.1016/0142-694X(91)90003-F)
- Johannesson, P., & Perjons, E. (2021). *An Introduction to Design Science*. Cham : Springer International Publishing. <https://doi.org/10.1007/978-3-030-78132-3>
- Just, J. (2024). Natural language processing for innovation search – Reviewing an emerging non-human innovation intermediary. *Technovation*, 129, 102883. <https://doi.org/10.1016/j.technovation.2023.102883>
- Kaarbo, J., & Beasley, R. K. (1999). A Practical Guide to the Comparative Case Study Method in Political Psychology. *Political Psychology*, 20(2), 369-391. <https://doi.org/10.1111/0162-895X.00149>
- Koh, E. C. Y., & De Lessio, M. P. (2018). Fixation and distraction in creative design: the repercussions of reviewing patent documents to avoid infringement. *Research in Engineering Design*, 29(3), 351-366. <https://doi.org/10.1007/s00163-018-0290-y>
- Kröger, M. (2021). *Studying Complex Interactions and Outcomes Through Qualitative Comparative Analysis: A Practical Guide to Comparative Case Studies and Ethnographic Data Analysis*. London : Routledge. <https://doi.org/10.4324/9781003095101>

- Leahy, K., Daly, S. R., McKilligan, S., & Seifert, C. M. (2020). Design Fixation From Initial Examples: Provided Versus Self-Generated Ideas. *Journal of Mechanical Design*, 142(10), 101402. <https://doi.org/10.1115/1.4046446>
- Luo, J., Sarica, S., & Wood, K. L. (2021). Guiding data-driven design ideation by knowledge distance. *Knowledge-Based Systems*, 218, 106873. <https://doi.org/10.1016/j.knosys.2021.106873>
- Ma, L., Ge, L., Li, Y., & Wang, Y. (2026). When employees meet gen AI: The double-edged effects of AI attitude on creativity. *Technovation*, 152, 103495. <https://doi.org/10.1016/j.technovation.2026.103495>
- Mayer, S., & Schwemmler, M. (2025). The impact of design thinking and its underlying theoretical mechanisms: A review of the literature. *Creativity and Innovation Management*, 34(1), 78-110. <https://doi.org/10.1111/caim.12626>
- Mello, P. A. (2021). *Qualitative comparative analysis: an introduction to research design and application*. Washington, DC : Georgetown University Press.
- Nelson, B. A., Wilson, J. O., Rosen, D., & Yen, J. (2009). Refined metrics for measuring ideation effectiveness. *Design Studies*, 30(6), 737-743. <https://doi.org/10.1016/j.destud.2009.07.002>
- Orwig, W., Diez, I., Vannini, P., Beaty, R., & Sepulcre, J. (2021). Creative Connections: Computational Semantic Distance Captures Individual Creativity and Resting-State Functional Connectivity. *Journal of Cognitive Neuroscience*, 33(3), 499-509. https://doi.org/10.1162/jocn_a_01658
- Ozturk, P., Han, Y., & Ren, J. (2026). Crowdsourcing creativity: Support architectures and task-knowledge intensity. *Technovation*, 150, 103393. <https://doi.org/10.1016/j.technovation.2025.103393>
- Pedota, M., Cicala, F., & Basti, A. (2025). Human agents, generative AI, and innovation: A formal model of hybrid creative process. *Technovation*, 148, 103323. <https://doi.org/10.1016/j.technovation.2025.103323>
- Pedota, M., & Piscitello, L. (2022). A new perspective on technology-driven creativity enhancement in the Fourth Industrial Revolution. *Creativity and Innovation Management*, 31(1), 109-122. <https://doi.org/10.1111/caim.12468>
- Peffer, K., Tuunanen, T., & Niehaves, B. (2018). Design science research genres: introduction to the special issue on exemplars and criteria for applicable design science research. *European Journal of Information Systems*, 27(2), 129-139. <https://doi.org/10.1080/0960085X.2018.1458066>
- Peffer, K., Tuunanen, T., Rothenberger, M. A., & Chatterjee, S. (2007). A Design Science Research Methodology for Information Systems Research. *Journal of Management Information Systems*, 24(3), 45-77. <https://doi.org/10.2753/MIS0742-122240302>

- Pescher, C., & Tellis, G. J. (2025). The Role of Artificial Intelligence in the Ideation Process. *Journal of Product Innovation Management*, 42(5), 958-982. <https://doi.org/10.1111/jpim.12791>
- Pinkow, F. (2023). Creative cognition: A multidisciplinary and integrative framework of creative thinking. *Creativity and Innovation Management*, 32(3), 472-492. <https://doi.org/10.1111/caim.12541>
- Purcell, A. T., & Gero, J. S. (1996). Design and other types of fixation. *Design Studies*, 17(4), 363-383. [https://doi.org/10.1016/S0142-694X\(96\)00023-3](https://doi.org/10.1016/S0142-694X(96)00023-3)
- Raj, R., Srivastava, A. K., & Behera, R. (2026). Automation — Augmentation of Artificial Intelligence and Human Intelligence. *Technovation*, 152, 103462. <https://doi.org/10.1016/j.technovation.2025.103462>
- Roberts, D. L., & Candi, M. (2024). Artificial intelligence and innovation management: Charting the evolving landscape. *Technovation*, 136, 103081. <https://doi.org/10.1016/j.technovation.2024.103081>
- Sarkar, P., & Chakrabarti, A. (2011). Assessing design creativity. *Design Studies*, 32(4), 348-383. <https://doi.org/10.1016/j.destud.2011.01.002>
- Sattele, V., & Ortiz, J. C. (2025). AI as an element to overcome creative fixation in design teams. Dans *Proceedings of the Design Society: International Conference on Engineering Design (ICED25)*. <https://doi.org/10.1017/pds.2025.10055>
- Shah, J. J., Smith, S. M., & Vargas-Hernandez, N. (2003). Metrics for measuring ideation effectiveness. *Design Studies*, 24(2), 111-134. [https://doi.org/10.1016/S0142-694X\(02\)00034-0](https://doi.org/10.1016/S0142-694X(02)00034-0)
- Sio, U. N., Kotovsky, K., & Cagan, J. (2015). Fixation or inspiration? A meta-analytic review of the role of examples on design processes. *Design Studies*, 39, 70-99. <https://doi.org/10.1016/j.destud.2015.04.004>
- Srinivasan, V., Song, B., Luo, J., Subburaj, K., Elara, M. R., Blessing, L., & Wood, K. (2018). Does Analogical Distance Affect Performance of Ideation? *Journal of Mechanical Design*, 140(7), 071101. <https://doi.org/10.1115/1.4040165>
- Tang, F., Li, B., Yao, F., Zhang, Z., & Zhu, R. (2026). Generative AI usage and improvisation capability: The mediating role of technological affordances and the moderating role of task uncertainty. *Technovation*, 153, 103535. <https://doi.org/10.1016/j.technovation.2026.103535>
- Thomas, L. D. W., & Tee, R. (2022). Generativity: A systematic review and conceptual framework. *International Journal of Management Reviews*, 24(2), 255-278. <https://doi.org/10.1111/ijmr.12277>

- Vallé, R., Seitz, J., Kanbach, D. K., & Schuhmacher, M. C. (2026). The use of artificial intelligence in new product development: A systematic literature review, conceptual framework, and future research agenda. *Technovation*, 153, 103541. <https://doi.org/10.1016/j.technovation.2026.103541>
- Vasconcelos, L. A., & Crilly, N. (2016). Inspiration and fixation: Questions, methods, findings, and challenges. *Design Studies*, 42, 1-32. <https://doi.org/10.1016/j.destud.2015.11.001>
- Wadinambiarachchi, S., Kelly, R. M., Pareek, S., Zhou, Q., & Velloso, E. (2024). The Effects of Generative AI on Design Fixation and Divergent Thinking. Dans *Proceedings of the CHI Conference on Human Factors in Computing Systems* (pp. 1-18). Honolulu, HI, USA : ACM. <https://doi.org/10.1145/3613904.3642919>
- Wu, Y., Zhou, C., & Zhi, J. (2023). Alleviating design fixation with Alternative Uses Task: The role of an integrated and design-independent intervention in promoting creative incubation. *Thinking Skills and Creativity*, 50, 101406. <https://doi.org/10.1016/j.tsc.2023.101406>
- Youmans, R. J., & Arciszewski, T. (2014). Design fixation: Classifications and modern methods of prevention. *Artificial Intelligence for Engineering Design, Analysis and Manufacturing*, 28(2), 129-137. <https://doi.org/10.1017/S0890060414000043>
- Zhang, J., Han, J., & Ahmed-Kristensen, S. (2025). Exploring the use of LLMs to evaluate design creativity. *Proceedings of the Design Society*, 5, 1773-1782. <https://doi.org/10.1017/pds.2025.10191>
- Zhu, Q., & Luo, J. (2025). Artificial creativity in design: a theory-based framework for implementation. *Proceedings of the Design Society*, 5, 621-630. <https://doi.org/10.1017/pds.2025.10076>