

Université Catholique de Louvain

Faculté des sciences économiques, sociales et politiques et de communication

Institut Langage & Communication – Pôle Communication

Beauté des données.

Quelle est l'influence de l'ornementation des visualisations
de données sur la construction de sens de l'utilisateur ?

*Dissertation doctorale présentée en vue de l'obtention du grade de docteur en sciences de l'information
et de la communication, par Tiffany Andry*

Composition du jury

François Lambotte (Promoteur)	Université Catholique de Louvain
Pierre Fastrez (Promoteur)	Université Catholique de Louvain
Suzanne Kieffer	Université Catholique de Louvain
Alexandru Telea	Université d'Utrecht
Christophe Hurter	École Nationale de l'Aviation Civile, Toulouse
Julia Bonaccorsi	Université Lumière Lyon 2 Berges du Rhône

Le vendredi 5 mars 2021

Je tiens à manifester toute ma reconnaissance à ceux qui, de près ou de loin, m'ont aidée et soutenue dans cette incroyable aventure.

Mes promoteurs. Merci à François Lambotte de m'avoir offert les meilleures conditions pour évoluer sereinement lors de mon parcours. Merci aussi pour son ouverture, sa disponibilité, et son amour communicatif pour la recherche. Merci à Pierre Fastrez qui, à plusieurs reprises, autour de discussions, a illuminé mon chemin en m'ouvrant à de nouvelles idées et perspectives.

Les membres de mon comité d'accompagnement. Merci à Suzanne Kieffer pour ses conseils en matière de méthode expérimentale. Merci à Christophe Hurter et Alexandru Telea d'avoir échangé durant de nombreuses heures à propos de l'ornementation des visualisations de données. Merci à mon jury pour sa disponibilité et pour nos riches échanges.

Mes collègues, aka « Le couloir du bonheur ». Merci pour toutes ces heures, tous ces thés, toutes ces tartes autour de discussions bienveillantes et respectueuses. Quelle équipe !

Sur le plan personnel, je tiens à remercier ma famille et la troupe d'Arte Corpo pour leurs encouragements et leur soutien sans faille. Merci à Manon pour son travail de relecture. Je remercie également mes chers amis rencontrés il y a dix ans sur le campus de l'UCLouvain FUCaM Mons. Leur amitié est un refuge et, par leurs tendres mots, ils m'ont parfois apaisée et encouragée sans même le savoir.

Table des matières

Introduction	7
Partie I : État de l'art. Visualiser, orner et créer du sens	13
Chapitre 1 La visualisation de données.....	13
1. Définition	13
2. Explorer ou présenter ?.....	14
3. Quelle classification des « dataviz » ?.....	17
4. Le débat des principes de conception des visualisations de données.....	19
4.1. Les éléments de base du système graphique	19
4.2. Efficacité et excellence graphique.....	22
4.3. Vers un minimalisme assumé : principes de conception	23
Chapitre 2 L'ornementation des visualisations de données	33
1. Ornementation ou <i>embellishment</i> ?	34
2. Retour sur les termes de base : entre esthétique, beauté et design	45
2.1. De l'esthétique.....	46
2.2. De la beauté, liée à l'art dans l'histoire	47
2.3. Du design.....	48
3. Conclusion : comment considérer l'ornementation en dataviz ?	51
Chapitre 3 Aborder la construction de sens en visualisation de données.....	55
1. Les différentes composantes de la construction de sens	55
2. La perception visuelle de l'image	56
3. Les mécanismes cognitifs : coordonner ses représentations mentales	59
4. Le <i>sense-making</i> individuel par Dervin.....	61
5. Les facteurs d'engagement envers la visualisation	63
Chapitre 4 Problématisation et positionnement	66
1. Positionnement épistémologique sur la théorie abordée	66
2. Formulation de la question de recherche.....	71
Partie II : Vers la récolte de données. Expérimenter, interviewer	72
Introduction Une récolte de données en phase avec les composantes de la (...)	73
Chapitre 1 Phase 1 : L'exposition aux visualisations de données (...)	75
1. Les visualisations exposées : constitution du corpus test.....	75
2. Les sujets : choix de l'échantillon	80
3. Matériel	81

3.1.	<i>Mobilier, locaux et logistique</i>	81
3.2.	<i>Technologie utilisée</i>	82
4.	Articulation des outils et déroulement des expérimentations.....	83
4.1.	<i>Déroulement général</i>	83
4.2.	<i>Contrôle des effets d'apprentissage et des effets de séquence</i>	85
4.3.	<i>Instrument de présentation des visualisations</i>	87
4.4.	<i>La récolte des appréciations déclarées</i>	88
Chapitre 2 Phase 2 : L'entretien semi-directif		89
1.	Pourquoi choisir l'entretien semi-directif ?.....	89
1.1.	<i>Le paradigme éactif et le cours d'expérience</i>	90
1.2.	<i>Les méthodes « think-aloud »</i>	91
1.3.	<i>L'entretien semi-directif, méthode traditionnelle de récolte de données en sciences de la communication</i>	93
1.4.	<i>Contribution méthodologique de ces modèles</i>	94
2.	Les appréciations déclarées comme guide d'entretien : une méthode nouvelle ...	96
3.	Matériel et logistique	98
4.	Déroulement général	98
5.	Découvrir Dervin sur le tard : de l'acte manqué à l'opportunité d'analyse	99
Conclusion Justesse méthodologique et positionnement épistémologique		101
1.	<i>Eye tracking</i> et entretien semi-directif : pourquoi une telle méthode est (...)	101
2.	Réflexion épistémologique sur la méthodologie.....	103
Partie III. Les appréciations déclarées. Analyses et interprétations		105
Chapitre 1 Méthode d'analyse des appréciations déclarées.....		105
1.	Présentation des données mobilisées dans l'analyse.....	105
1.1.	<i>Présentation des appréciations déclarées</i>	105
1.2.	<i>Présentation des attributs des visualisations de données du corpus</i>	107
2.	Méthode d'analyse des appréciations déclarées et des attributs	111
Chapitre 2 Analyse statistique des appréciations déclarées <i>par attribut</i>		113
1.	L'attribut principal : le type (« Raw » vs. « Junk »).....	113
2.	Corrélations inter-échelles (Kendall Tau b).....	119
3.	Effets de l'attribut « ornement » sur les appréciations déclarées	122
4.	Effets de l'attribut « Figuration » sur les appréciations déclarées.....	128
5.	Effets de l'attribut « Pictogramme » sur les appréciations déclarées	135
6.	Effet de l'attribut « Lie Factor » sur les appréciations déclarées.....	137
7.	Effet de l'attribut « Data-Ink » sur les appréciations déclarées	140
8.	Effet de l'attribut « Embellished Data-Ink » sur les appréciations déclarées.....	145

Chapitre 3	Interprétation des résultats et discussion	151
1.	Interprétation factuelle	151
2.	Discussion	156
2.1.	<i>Vers une recommandation essentielle en conception de visualisations (...)</i>	156
2.2.	<i>La place d'une telle analyse dans la littérature sur l'ornementation</i>	159
Partie IV : La subjectivité dans la construction de sens. Analyse et (...)		163
Introduction		163
Chapitre 1	Analyse thématique	164
1.	Méthode d'analyse par codage thématique	164
2.	Interprétation résultant du codage thématique	165
2.1.	<i>La perception des visualisations</i>	166
2.2.	<i>Les facteurs d'appréciation des visualisations.....</i>	180
2.3.	<i>Critique de la disposition des éléments graphiques</i>	182
2.4.	<i>Les facteurs humains et émotionnels : des fondamentaux.....</i>	189
3.	Conclusion de l'analyse thématique.....	195
Chapitre 2	Les comblements de discontinuité ou le <i>gap bridging</i>	196
1.	Méthode d'analyse inspirée de la <i>sense-making methodology</i>	197
2.	Résultats de l'analyse et interprétation.....	200
2.1.	<i>Description et répartition des verbings par étape de gap bridging</i>	201
2.2.	<i>Les caractéristiques des visualisations de données, cause de (...)</i>	204
2.3.	<i>Les schémas modèle</i>	212
3.	Apport et pertinence de l'analyse des gap bridgings de visualisations (...)	219
Chapitre 3	Ornementer ou prôner le design émotionnel (Norman, 2013)	221
1.	Rappel théorique	222
2.	Pertinence de ce cadre conceptuel.....	223
3.	Vers la visualisation de données idéale ?	225
4.	La simplicité, encore et toujours	226
Conclusion	L'analyse qualitative, la finesse des détails.....	228
Partie V : Fixations et patterns de lecture. Analyse des <i>gaze data</i>		230
Introduction		231
Chapitre 1	Les <i>gaze data</i> récoltées, mises en perspective dans en <i>Infovis</i> et en (...) ...	232
1.	Les données récoltées et le matériel exploitable	232
2.	La compréhension de graphiques et l' <i>eye tracking</i>	234
Chapitre 2	Analyse des fixations	238
1.	Formulation des hypothèses	238
2.	Méthode d'analyse	240

3. Analyse statistique des fixations	241
3.1. Comparaison des fixations en fonction du type (ornementé ou standard)..	241
3.2. Comparaison des fixations en fonction d'attributs significatifs.....	247
3.3. Exploration des zones d'intérêt.....	254
Chapitre 3 Détection des <i>patterns</i> de lecture : le <i>bundling</i>	257
1. Pertinence d'une nouvelle technique de visualisation	257
2. Analyse des chemins de lecture : observations et commentaire général sur le résultat du <i>bundling</i>	261
2.1. Les cycles ou le « 3-stage model »	261
2.2. De la lecture.....	262
2.3. Les icônes apposées	264
2.4. L'ornementation sur le <i>data-ink</i> et la place des étiquettes	265
2.5. La maximisation du <i>data-ink ratio</i> ... ou de l'ornementation épurée	267
2.6. Les formes inhabituelles.....	268
2.7. Ce qui influence le <i>pattern</i>	269
3. Interprétation des observations	272
Conclusion L'analyse des <i>gaze data</i> : entre complexité, engagement et (...).....	275
 L'influence de l'ornementation des visualisations de données sur la construction de sens. Conclusion	276
1.Contexte de la recherche	277
2.L'influence de l'ornementation sur la perception visuelle et l'impression de (...).....	278
3.L'influence de l'ornementation sur la construction de sens individuelle.....	283
4.Recommandations relatives à l'utilisation de l'ornementation	286
5.Limites de la recherche et perspectives futures.....	289
6.Réflexion théorique et épistémologique : quels apports de la recherche aux SIC ?.....	293
Bibliographie.....	297
Table des figures et tableaux.....	305
<i>Figures</i>	305
<i>Tableaux</i>	307

Introduction

Ces dernières années, la quantité d'information produite et accessible n'a cessé de croître, notamment grâce à l'essor de l'activité socio-numérique. En 2010, le volume de données numériques créé à l'échelle mondiale était de 2 zettaoctets¹. En 2020, il s'élève à 47 zettaoctets. Les usages socio-numériques et l'utilisation de technologies numériques sont si foisonnants qu'on s'attend à une explosion exponentielle des données au cours du temps, avec un volume de données numériques de 175 zettaoctets attendu d'ici à 2025 (Reinsel, Gantz, Rydning, 2020). Toutes ces données permettront des avancées technologiques et sociales dans bien des domaines. Mais, en premier lieu, il s'agit d'information brute. Pour comprendre le message contenu dans les données numériques, il est très commun de les visualiser. Ainsi, les nouvelles technologies permettent l'extraction, le traitement algorithmique et la visualisation des données à des fins informationnelles et communicationnelles et transforment profondément nos sociétés (Cardon, 2015). Dans de nombreuses sphères (entreprises, médias, politique, ...), visualiser les données permet de les interpréter et de les comprendre. Cela permet également de les communiquer à une audience qui pourra comprendre une information claire et condensée : la visualisation de données est le media principal par lequel les gens ont accès aux données (Kennedy et Hill, 2018). Elle se définit comme une représentation visuelle visant à faciliter la compréhension (Kirk, 2016). Les graphiques les plus basiques, appris par tous au cours du parcours scolaire, sont des visualisations de données, tout comme la représentation visuelle d'un réseau identifiant les interconnexions entre des centaines de personnes.

En outre, la visualisation de données acquiert également une place grandissante dans les espaces médiatiques numériques. Dans la presse en ligne, sur les blogs et les réseaux sociaux, elle gagne du terrain, grâce notamment à la puissance visuelle des infographies qui permettent de communiquer l'information rapidement dans un espace numérique très saturé. La visualisation de données se développe donc dans un contexte d'évolution intense. Tout d'abord, notre environnement actuel est imprégné d'infobésité, comprise comme une surcharge d'information (Reibel, Desrocher, 2014). Si ce concept existait déjà dans les années 60, il est plus vrai aujourd'hui que jamais (Sauvajol-Rialland, 2014). Alors que les médias diffusaient auparavant l'information à sens unique avec un faible retour de la part du public, le web 2.0 et les réseaux sociaux sur lesquels l'information est diffusée donnent maintenant aux gens la possibilité de réagir, de répondre et de partager. En effet, on peut parler de "web social", où l'utilisateur est un producteur de contenu qui provoque

¹ 1 zettaoctet (Zo), = 10^{21} octets ;

inévitavelmente une mutation de la communication, rendant l'information disponible partout, en tout temps, tout en laissant la possibilité à l'utilisateur de la commenter. L'information est omniprésente et est souvent poussée vers l'utilisateur qui pourrait être intéressé, grâce aux différents algorithmes filtrant les intérêts des utilisateurs sur le web. Les graphiques ou visualisations d'information sont un bon moyen de communiquer quelque chose rapidement et d'intéresser les gens. Ils sont à la fois un gain de temps et d'effort cognitif. Dans ce contexte de surcharge d'information, ils sont souvent mobilisés. L'infobésité frappe également les organisations et les entreprises, où la quantité de données produites crée une course à l'information qui est devenue vitale pour rester efficace. D'une part, la plupart des entreprises doivent être visibles sur les multiples moyens de communication modernes et, d'autre part, elles doivent collecter des données et des informations pour les analyser et affronter correctement la compétition (Sauvajol-Rialland, 2014). C'est pourquoi, les tableaux de bord d'analyse d'activité, rassemblant diverses visualisations de données, sont très utilisés aujourd'hui pour établir des rapports, prendre des décisions et comprendre globalement cette grande masse de données (Adam & Pomerol, 2008).

Plus encore, nous vivons dans une société où la *datafication* bat son plein : les données sont partout autour de nous, dans une logique de production et d'exploitation quasi-constante. « De plus en plus d'aspects de la vie sociale sont *datafied* : les amitiés, les intérêts et les émotions ont été transformés en données quantifiables sur les médias sociaux et autres plateformes désireuses d'extraire de la valeur des opérations sociales quotidiennes » (Kennedy, Hill, 2018, p.1, notre traduction). Aujourd'hui, nous avons tendance à mettre les choses sous forme de chiffres, de données, puis à les visualiser pour les comprendre, ce qui peut donner une impression de neutralité ou, comme le disent Kennedy et Hill (2017), une vision aseptisée de la réalité. Les données, sous leur forme visualisée, semblent objectives ; il semblerait en effet que « les chiffres parlent d'eux-mêmes ». Leur utilisation médiatique s'explique donc également par une dynamique d'intérêt, de rapidité et de conviction. Pourtant, la visualisation de données qui se veut totalement objective n'existe pas parce qu'elle est le résultat d'étapes qui visent à construire l'interprétation d'un phénomène : les données sont une construction, puisqu'elles sont extraites d'une activité. Elles sont ensuite nettoyées et exploitées dans un but précis : être représentées dans un graphique pour expliquer une histoire, pour communiquer une information. La visualisation de données exprime un message, pas un fait en soi. Elle est donc le résultat de choix visant à communiquer un message.

Pour résumer, la société actuelle est en pleine transformation digitale et, dans ce contexte, la visualisation de données est utile à bien des égards, car

- Elle aide à faire sens de l'activité numérique : les *visual analytics*, visualisations représentant des mesures simples dans des tableaux de bord informatiques, se sont énormément développées ces dernières années ;
- Elle permet de communiquer des données qui sont aujourd'hui de plus en plus accessibles (par exemple grâce au *data crunching* ou encore à l'*open data*) ;
- Elle peut être utilisée en journalisme pour raconter des histoires chiffrées, avec plus de précision (*data journalisme*) ;
- Elle améliore la communication et l'utilisation des données dans les entreprises où les évolutions technologiques permettent de mieux générer, stocker, analyser, visualiser les données et prendre des décisions ;
- Elle permet de se démarquer dans un contexte d'infobésité numérique, grâce à ses propriétés visuelles.

Dès lors, dans ces différents domaines et contextes, comment les concepteurs de toutes ces visualisations de données font-ils pour que leurs visualisations soient les plus efficaces possible ? En vérité, ils font assez peu appel à des référentiels ou à des guides qui existent, principalement en langue anglaise. Ils s'y réfèrent peu parce que d'une part, le temps manque et que, d'autre part, ils ne s'y retrouvent pas forcément. Pourtant, de nombreux principes de conception des visualisations de données existent, notamment parce que la visualisation de données est loin d'être un phénomène neuf. Depuis longtemps, l'homme utilise la schématisation pour organiser sa pensée et la communiquer. William Playfair, en 1786, concevait pour la première fois et de manière aboutie les trois graphiques les plus connus au monde : le diagramme en ligne, le diagramme en barres et le diagramme circulaire (Costigan-Eaves, Macdonald-Ross, 1990). De nombreuses pratiques en visualisation de données ont ensuite foisonné. Ce n'est donc pas la visualisation de données qui est nouvelle, mais plutôt le regard que nous portons sur les données puisque nous les visualisons afin de les interpréter (Cardon, 2012). Transcrire des données statistiques en une forme visuellement compréhensible est un enjeu stimulé par la diversité d'outils, de techniques, de pratiques et de tendances.

De nombreux principes de conception de visualisations de données existent donc. Ils permettent de maîtriser au mieux l'organisation de l'information sur un graphique afin qu'elle puisse être saisie de la façon la plus simple, rapide et claire possible par le lecteur. Bertin (1967) et Tufte (1983) sont les deux pères de la visualisation de données : les travaux de chacun d'entre eux ont nourri le domaine de la conception graphique de visualisations de données. Pour Bertin, l'efficacité de la communication est la priorité principale. Il

explique comment faire des choix raisonnés qu'il veut libres de toute subjectivité grâce à un éventail de possibilités graphiques (Bertin, 1967). L'utilisateur doit donc pouvoir décoder aisément l'information qui est encodée sur la visualisation. Ensuite, les travaux de Tufte traitent de différents principes à appliquer dans les représentations visuelles et les visualisations d'information. En général, il montre comment visualiser des données avec simplicité, clarté et élégance, tout en faisant campagne contre les "*chartjunks*"² et autres mauvaises pratiques de conception qui conduisent à l'obstruction de l'information (Tufte, 1983, 1991, 2003). Les règles proposées par les deux auteurs sur lesquels nous nous concentrerons tout au long de cette thèse tendent vers une proposition forte : au plus simple le graphique sera, au mieux il sera compris. Pour eux, l'efficacité du graphique se mesure par le temps de lecture, la quantité d'encre utilisée ou encore l'espace occupé par la représentation graphique. Ils optent donc pour des designs épurés et minimalistes. Ils supposent ainsi qu'en suivant l'ensemble de leurs principes, l'information encodée sera correctement décodée par les lecteurs. Néanmoins, cette vision de la transmission de l'information pousse au débat dans les sciences de l'information et de la communication : en estimant que le lecteur est un élève modèle qui comprendra l'information telle qu'elle a été encodée, ces auteurs oublient en réalité de les considérer pour ce qu'ils sont, c'est-à-dire des sujets interprétants qui réagissent à l'information en fonction de leur propre expérience et en fonction du contexte de lecture.

Ce manque d'attention envers l'utilisateur se ressent aujourd'hui dans les principes de conception des visualisations de données. En effet, d'un point de vue médiatique, on remarque que les designers ne respectent pas du tout cette tendance minimaliste recommandée par les précurseurs de la visualisation de données. Souvent, sur les espaces numériques, les visualisations de données sont embellies par rapport aux recommandations minimalistes. Couleurs, réarrangements, fantaisies, utilisations d'images ou d'icônes, ... Les façons d'embellir les visualisations de données statiques sont nombreuses dans la pratique médiatique et numérique : dans un contexte d'infobésité numérique, nous ne pouvons nier l'importance de l'aspect visuel pour attirer le regard ou agrémenter un propos. De même, il semble plaisant pour les designers d'ornementer leurs visualisations, tout comme cela semblerait nécessaire pour éveiller la curiosité et l'intérêt des lecteurs. En réalité, cet embellissement, que nous décidons d'appeler « ornementation » au cours de cette thèse, n'est pas considéré de manière péjorative et déroge complètement aux principes et règles de conception qui permettent une lecture parfaite selon Bertin et Tufte.

² Bruit graphique, élément perturbateur d'une représentation graphique, traité dans l'état de l'art.

En effet, dans la littérature anglo-saxonne et dans la communauté *Infovis* (communauté scientifique internationale étudiant la visualisation d'information), les préoccupations quant aux pratiques d'ornementation sont grandissantes. Plusieurs études ont notamment été menées et montrent qu'en fait, l'ornementation peut parfois apporter une valeur ajoutée à la visualisation de données, comme une meilleure mémorisation à court terme et à long terme, l'envie de s'engager envers la visualisation en poursuivant la lecture ou encore de l'amusement. Il y a donc une tension entre l'ornementation, critiquée et décriée par les précurseurs de la visualisation de données, et leurs principes minimalistes. Cette tension est au cœur de cette thèse.

En réalité, tout est une question de construction de sens. En se demandant comment le sens est produit, plutôt qu'en se concentrant sur la conception de la visualisation, nous nous concentrons sur l'utilisateur. Il est l'acteur interprétant, celui qui produit le sens grâce à tous les signes à disposition sur le graphique. Ainsi, l'ornementation influence évidemment le sens produit par les utilisateurs. Au-delà de la simple compréhension des données, il se passe quelque chose pour l'utilisateur qui lit la visualisation, qu'elle soit ornementée ou pas : il va faire sens en fonction du sujet abordé, des données, du contexte de lecture, mais surtout de ses caractéristiques personnelles. Il en va de même pour les visualisations minimalistes : l'utilisateur dépasse la simple compréhension en faisant sens et, que les données soient comprises ou pas, une chose est sûre : l'utilisateur produit toujours du sens à travers son expérience d'utilisation. Se pencher sur la construction de sens opérée par les utilisateurs revient à considérer la visualisation de données comme un objet de communication. Cet objet est assez peu étudié en sciences de l'information et de la communication et nous nous y attelons dès lors en nous concentrant sur l'ornementation de ces visualisations de données.

Alors que différentes études abordent l'ornementation pour ce qu'elle apporte de positif sur des éléments très concrets comme la mémorisation ou l'engagement, nous souhaitons l'étudier de façon plus large. Dans cette thèse de doctorat, nous nous questionnons sur l'influence de l'ornementation des visualisations de données sur la construction de sens des utilisateurs. Nous considérons la construction de sens comme l'assemblage de différentes composantes que sont la perception visuelle, la compréhension et la subjectivité (l'expérience). En effet, lorsqu'un utilisateur voit une visualisation, il la regarde (donc la perçoit visuellement), la comprend (ou pas) et se positionne finalement lui-même par rapport à cette visualisation, en tant que sujet interprétant. Nous avons donc développé une méthode de récolte de données originale, en lien avec cette vision de la construction de sens.

Dans la présente monographie, la recherche se présente comme suit. Tout d'abord, nous proposons dans la première partie un état de l'art sur la visualisation de données, sur le phénomène croissant qu'est l'ornementation ainsi que sur la façon dont nous étudions la construction de sens. Tous ces éléments nous aident à problématiser le sujet qui nous anime et à poser une question de recherche précise. Dans la deuxième partie, nous exposons la méthodologie employée et notre positionnement épistémologique. Ainsi, nous y expliquons que la méthode de récolte de données a été appliquée en laboratoire et repose sur un protocole d'expérimentation utilisant différents outils et s'inspirant lui-même de différentes méthodes issues des sciences sociales et du champ de l'interaction homme-machine. Effectivement, 40 personnes, professionnelles de la communication numérique, se sont rendues au Social Media Lab de l'UCLouvain pour participer à une expérience d'une durée allant de 1h00 à 1h30. Elles ont pu observer un corpus de visualisations minimalistes et ornementées. Les données ont pu être récoltées grâce à une expérimentation *eye tracking* conjointe à une récolte d'appréciations déclarées. Après l'expérimentation, les participants ont pris part à un entretien semi-directif visant à expliquer leur ressenti par rapport à l'expérimentation. Pour ce faire, le protocole d'expérimentation devait être correctement réalisé, afin de disposer du plus grand niveau de contrôle durant les expérimentations. Dans cette deuxième partie, rien n'est laissé au hasard : chaque choix de conception expérimentale y est expliqué. Par la suite, les trois parties suivantes concernent chacune l'analyse d'un type de données issu de l'expérimentation. Effectivement, nous avons récolté des données quantitatives (données ordinales), qualitatives (retranscriptions d'entretiens), et des données de regard (*gaze data*). Évidemment, à chaque type de données sa méthode d'analyse. Ces trois dernières parties commencent chacune par un exposé sur la méthode d'analyse employée avant de poursuivre sur les résultats. La partie III présente une analyse statistique, réalisée sur base d'appréciations déclarées récoltées au moment de l'expérimentation. Puis, la partie IV présente les résultats d'une analyse de contenu purement qualitative. La cinquième partie, quant à elle, présente l'analyse de données de regard récoltées grâce à l'*eye tracking*. Pour conclure, nous expliquerons en quoi chaque analyse apporte sa pierre à l'édifice et permet de répondre à la question de recherche. Grâce à cela, nous pourrions commenter la tension entre les principes minimalistes et la pratique d'ornementation des visualisations de données. Nous mettrons également en avant les apports de cette recherche aux sciences de l'information et de la communication, ainsi que les différentes limites qu'elle contient et les perspectives qu'elle ouvre grâce à l'interdisciplinarité.

Partie I. État de l'art

Visualiser, orner et créer du sens

Chapitre 1 La visualisation de données

Qu'est-ce que la visualisation de données, à quoi sert-elle et sous quelle forme allons-nous l'étudier dans le cadre de cette thèse ? Nous ouvrons cet état de l'art en définissant notre objet d'étude et ses différents objectifs. Pour terminer, nous présentons un exemple de classification des visualisations de données qui existent.

1. Définition

L'expression « visualisation de données » n'était pas courante à la fin des années soixante. Jacques Bertin, cartographe français et auteur de *La Sémiologie Graphique* (1967), manuscrit fondamental en la matière, parle plutôt de représentation graphique. Pour lui, il s'agit de « la transcription, dans le système graphique de signes, d'une pensée, d'une information connue par l'intermédiaire d'un système de signes quelconque » (Bertin, 1967, p 8). Bertin assimile alors représentation graphique et sémiologie (la « science des signes ») pour aboutir à la sémiologie graphique. La bonne utilisation des signes doit éviter de laisser toute possibilité d'interprétation personnelle au lecteur, afin qu'il dispose de l'information la plus « véritable » possible. Bertin avait en tête de transmettre les données sous forme graphique dans la plus grande « objectivité » possible. Il explique ainsi comment mobiliser le système de signes propre à la visualisation de données statistiques, la graphique, afin que le lecteur puisse comprendre des informations sur une schématisation abstraite au travers de la perception visuelle. Il énonce les moyens du système graphique et il explique la théorie de l'image, dans laquelle il met en place plusieurs règles de conception propres aux visualisations de données (Bertin, 1967). Également appelée « visualisation d'information », il s'agit de la présentation visuelle de données abstraites (Sallaberry, 2011).

Aujourd'hui, la visualisation de données peut être définie comme une représentation visuelle assistée par ordinateur, soit statique, soit interactive et mouvante (Card, 2012), dont le principal objectif est « de traduire des données sous une forme visuelle pertinente, simple, didactique et pédagogique » (Brasseur, 2015), dans un souci de lisibilité et de mémorisation motivé par l'apparition de « corrélations visuelles ou des relations d'ordre

entre les données » (Saulnier, Thièvre, Viaud, 2006, p.57). Cette représentation est génératrice d'une « amplification cognitive en termes d'acquisition et d'utilisation de l'information » (Polanco cité par Chauvin, 2006). Pour Hearst, ce serait « la représentation d'information réalisée grâce à la transcription graphique ou spatiale, pour faciliter la comparaison, la reconnaissance de formes, pour changer la détection et d'autres compétences cognitives, le tout en faisant usage du système visuel » (Hearst, 2003, en ligne, notre traduction). Il est également bon de noter que « la visualisation de données désigne tous types de représentations visuelles qui facilitent l'exploitation, l'analyse et la communication de données » (Fredriksson, 2015, p.36). Elle est à distinguer du design d'information, qui s'en différencie mais qui la complète : il « résulte d'un équilibre entre des enjeux relevant de l'exactitude scientifique, de l'esthétique, de l'efficacité cognitive et de l'accessibilité, déterminés par un contexte » (Fredriksson, 2015, p.36). Pour Kirk, il s'agit de « la représentation et la présentation de données pour faciliter la compréhension » (Kirk, 2016, p. 19, notre traduction). Par « représentation », Kirk entend les choix globaux, comme la forme du graphique, par exemple. La présentation concerne les décisions de design qui englobent l'apparence complète du projet.

Pour terminer, une distinction est à mettre en évidence. La visualisation de données peut être comprise de deux façons : le mot « visualisation » peut, d'une part, représenter l'acte de visualiser, de transcrire les informations sous une forme visuelle pertinente. C'est également le mécanisme par lequel l'ordinateur génère l'image. D'autre part, ce même mot peut désigner l'objet final : le graphique, la visualisation. Il est important et utile de la saisir, tant l'interdisciplinarité règne dans le domaine : tous n'expriment pas forcément l'idée du même point de vue. Dans cette recherche, nous aborderons les visualisations de données de type statique, à des fins de présentation. De même, nous considérons la visualisation de données comme un objet d'étude, c'est-à-dire comme un type d'image à étudier.

2. Explorer ou présenter ?

Les objectifs de la visualisation de données peuvent être divers. D'une part, elle peut être utilisée à des fins d'exploration. Dans ce cas, l'exploitation du graphique est plutôt individuelle. L'utilisateur visualise les données afin de trouver l'information qu'elles comportent et de les interpréter. Il est dans une réelle posture d'exploration. Il trouve des informations et génère des idées. Les visualisations de données produites dans un but exploratoire évoluent beaucoup au cours de leur utilisation. Le travail personnel de l'utilisateur sur cette visualisation est donc très important : certains choix de construction et d'utilisation de la représentation sont propres à la personne qui l'utilise (Friendly, 2008).

La seconde raison d'utiliser une représentation graphique est d'avoir pour ambition de présenter des données à des fins d'information. Dans ce cas, il faut décider quelle information on souhaite mettre en lumière et ce, dans un dispositif approprié à l'audience. La façon dont la représentation sera perçue par l'utilisateur doit être dans les préoccupations principales du créateur de la visualisation (Friendly, 2008). Dès lors, il doit définir à quels besoins la représentation graphique doit correspondre. Les questions auxquelles répondent les données exploitées dans une visualisation doivent être définies bien avant que la conception de la visualisation des données ne commence. Un même jeu de données peut en effet répondre à plusieurs questions à la fois. Le sujet de la visualisation de données ainsi que son audience sont des éléments déterminants dans sa présentation (Telea, 2016).

Ainsi, d'une part, il y a la visualisation de données en tant qu'outil d'analyse : il permettra à l'utilisateur d'explorer et d'analyser. D'autre part, il y a la visualisation de données en tant qu'outil de communication. Il comporte un message à transmettre à une audience. Bien entendu, les logiques de conception de ces deux types de visualisations de données sont tout à fait différentes.

Kirk (2016) a poussé plus loin les aspects explicatifs, exhibitoires et exploratoires des visualisations de données lorsqu'elles sont réalisées dans un objectif de communication. Le concepteur doit penser au ton qu'il veut donner à sa visualisation ainsi qu'à l'expérience qu'il veut générer pour l'utilisateur. Par exemple, même en présentant des données à des fins d'information, le concepteur peut aujourd'hui donner un aspect exploratoire à sa visualisation grâce au mouvement et à l'interactivité permis par nos supports de diffusion actuels. Kirk, à travers sa *Purpose Map* (Figure 1), propose différents cas de présentation des données, en se référant à l'intention du concepteur. Sur l'axe horizontal de la carte se trouve l'expérience générée par la visualisation chez l'utilisateur. Il y a ainsi

- Les visualisations explicatives, qui laissent assez peu de place à l'interprétation. L'information est clairement établie et le récepteur est assisté autant que possible dans la compréhension de la visualisation, par des guides, annotations et autres. C'est le meilleur type de visualisation quand on connaît peu de choses à propos de son public (Kirk, 2016, p. 77) ;
- Les visualisations à tendance exhibitoire. Dans ces dernières, le récepteur doit clairement faire un grand travail d'interprétation – et c'est voulu. Il doit être capable de comprendre seul le contenu et le contexte de la visualisation. Il tire sa propre conclusion à propos des données (Kirk, 2016, p. 81).

- Les visualisations exploratoires, qui se focalisent sur l'aide apportées à l'utilisateur tout en lui laissant une marge de manœuvre. Elles visent à faciliter le questionnement et la manipulation des données (Kirk, 2016, p. 79).

Sur la *Purpose Map*, sur l'axe vertical, se pose la question du ton défini dans la visualisation, peu importe l'expérience qu'elle génère. De manière générale, c'est subjectivement qu'il est défini. Il permet au concepteur de juger la meilleure manière de favoriser la lisibilité perceptuelle des données. La notion est assez difficile à comprendre et surtout à appliquer. Kirk définit deux tons :

- Le ton de lisibilité, qui est la meilleure approche quand le but du travail du concepteur est de faciliter la compréhension avec un haut niveau de détail. L'audience n'a pas besoin de stimulation pour être séduite et s'intéresser au dispositif, dès lors (Kirk, 2016, p. 82-83) ;
- Le ton de ressenti, qui met davantage l'accent sur certaines valeurs numériques dans le graphique, sur le sens des relations montrées, etc. Le coup d'œil y est parfois plus important que la compréhension profonde. Le postulat, en jouant sur le ressenti de l'utilisateur, est de dépasser la retranscription des données autrement que dans un tableau, en y apportant une plus-value. La sensation émotionnelle engendrée a la capacité d'influencer l'expérience de l'utilisateur ainsi que ses étapes de compréhension (Kirk, 2016, p. 83-14).

Comme on le constate à travers ces caractéristiques établies par Kirk, l'objectif communicationnel peut parfois varier : montrer brièvement un maximum, décrire une situation en détail ou encore susciter l'émotion sur un chiffre fort.

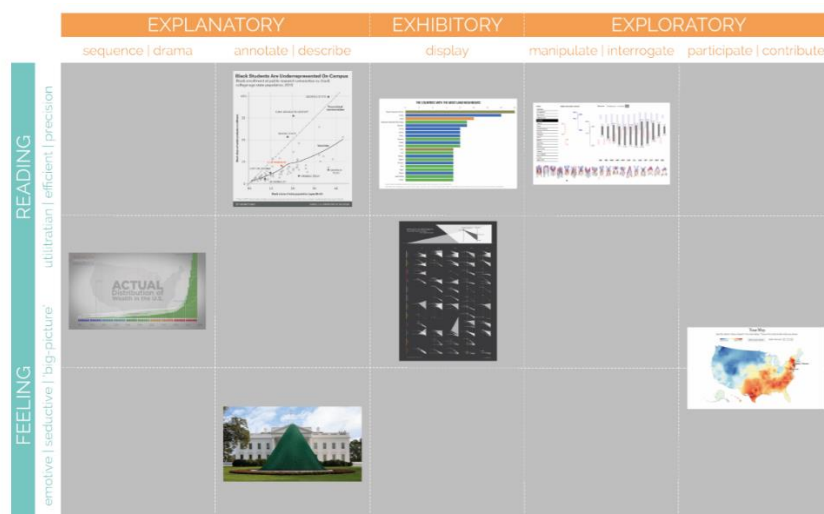


Figure 1 The purpose Map, tirée de Kirk A. (2016), *Data Visualisation. A Handbook for a data driven design*, Los Angeles, Sage, p. 76

Dans le cadre de ce travail de recherche, nous étudierons la visualisation de données comme objet de communication.

3. Quelle classification des « dataviz³ » ?

Nombreuses sont les représentations graphiques. Apparence visuelle, objectif, type de données représentées, ... Elles pourraient ainsi être classées en fonction de nombreuses caractéristiques. Certains auteurs parlent d'ailleurs de la visualisation de données comme d'une « technique » et classent ainsi différentes techniques de visualisations. L'objectif de cette thèse n'est évidemment pas de réaliser une typologie exhaustive des différentes visualisations de données mais il reste intéressant de montrer de quelle manière les classer selon nous. Une classification plus détaillée se trouve en annexe 1.

Afin de classer les visualisations de données, nous nous sommes basée sur l'objectif des différents types de visualisations de données. Nous avons commencé par différencier trois familles de visualisation, comme le fait par ailleurs Bertin (1967) qui les appelait « diagrammes, réseaux et cartes ». Les diagrammes mobilisent en général deux axes (parfois trois dans certains cas) ou deux dimensions du plan. En fonction du message à communiquer (indiquer des tendances, illustrer une temporalité, montrer des statistiques descriptives), les diagrammes descriptifs se subdivisent en différents groupes. Ensuite, nous avons appelé les réseaux « diagrammes nœud-lien », qui sont appelés à montrer des relations dans le plan en deux ou trois dimensions. Nous avons choisi cette terminologie car elle fait davantage appel au rendu visuel du réseau. Effectivement, le réseau, que l'on peut aussi appeler graphe, est un ensemble de nœuds et de liens qui peuvent être visualisés de nombreuses façons différentes, notamment grâce à des *layouts* disponibles dans les programmes de visualisation. Les « diagrammes nœud-lien » font davantage référence à ces *layouts* qu'au graphe en lui-même, qui est en réalité un calcul mathématique. Pour terminer, nous avons identifié les cartes, qui rendent compte d'une localisation spatiale.

Récapitulons donc en entrant davantage dans le détail. Nous distinguons :

- (1) Les diagrammes descriptifs, dont l'objectif est de donner une vue d'ensemble sur une situation donnée. La situation est décrite grâce à des valeurs dans un schéma.
 - a. Les diagrammes de valeurs. Ils regroupent les diagrammes qui, en général, s'appellent *chart* en anglais et qui montrent des tendances, des échelles ou encore la partie d'un tout.
 - b. Les diagrammes de séries chronologiques. Il s'agit de graphiques illustrant des données de type temporel et donc des évolutions dans le temps.

³ Le terme « dataviz » est la contraction de « data visualization ». Il est communément employé pour parler de la visualisation de données.

- c. Les diagrammes de statistiques descriptives. Utilisés dans les statistiques descriptives, ils permettent de montrer différentes notions comme la distribution, les corrélations et autres.
 - d. Les tableaux, que nous classons dans les diagrammes tout comme Bertin (1967). Effectivement, selon lui, un tableau où figurent parfois peu de données est plus lisible. De même, utiliser des variables visuelles simples peuvent aider à la lecture d'un tableau.
- (2) Les diagrammes nœud-lien, qui représentent uniquement des relations entre des entités. Il y a deux types de diagrammes nœud-lien.
- a. Les réseaux, qui représentent des relations à l'aide de nœuds (points) et de liens (arcs ou lignes). Les réseaux sont aussi appelés graphes en français.
 - b. Les arbres, qui sont en réalité des réseaux hiérarchiques. Il y a donc une racine : le réseau est enraciné. Les nœuds ont des parents et des enfants.
- (3) Les cartes, qui permettent ainsi la spatialisation. Il y a
- a. Les cartes géographiques, qui permettent d'indiquer des localisations physiques.
 - b. Les cartes non-géographiques, qui permettent par exemple de montrer une localisation sur un support, comme un écran.

Finalement, Bertin classe ce qu'il appelle « les problèmes graphiques selon » (1) le groupe d'imposition (diagrammes, réseaux, cartes), (2) le nombre de variables visuelles nécessaires, (3) le niveau d'organisation des variables, (4) la longueur des variables et leur implantation (Bertin, 1967, p. 192). Notre propre classification se trouve dans le Tableau 1. Elle se base pour rappel sur le type de message que la visualisation tente de montrer. Elle n'est pas forcément exhaustive et est expliquée en détail dans l'annexe 1.

Tableau 1 Classification des visualisations de données, par l'auteur. Le recensement des visualisations n'est pas exhaustif.

Diagrammes descriptifs	Diagrammes nœud-lien	Cartes
(1) Diagrammes de valeurs • Graphique en barres • Graphique en barres empilées • Diagramme circulaire • Diagramme en anneau • Diagramme en étoile • Coordonnées parallèles	(1) Réseaux - Les tracés de graphes théoriques • Graphe régulier • Random graph • Small World - Layouts (possibilités de visualisation des graphes) • Layout circulaire • Radial axis layout • Force Atlas • Fruchterman-Reingold	(1) Cartes géographiques • Carte à bulles • Choroplèthe • Anamorphose
(2) Diagrammes chronologiques • Graphique en ligne • Graphique en aires • Diagramme de Gantt		(2) Cartes non-géographiques • Carte de chaleur
(3) Diagrammes de statistiques descriptives • Histogramme • Scatter Plot • Diagramme en boîtes à moustaches	(2) Arbres - Layouts (possibilités de visualisation des graphes) • Hierarchical tree • Radial tree • Cone tree	

4. Le débat des principes de conception des visualisations de données

Nombreux sont les principes généraux relatifs à la construction d'une visualisation de données ou permettant de juger de son bon « fonctionnement ». Bien que la littérature soit foisonnante à ce sujet, nous avons choisi de nous concentrer sur deux auteurs en particulier – que nous pourrions également qualifier de « précurseurs » en la matière. Nous nourrissons de temps à autres leurs propos de contributions d'auteurs. Nous avons déjà abordé Bertin, cartographe français et auteur de *La Sémiologie Graphique* (1967), véritable manifeste en la matière. Il a été durant une partie de sa vie contemporain de Tufte, statisticien américain et pionnier dans le domaine de la visualisation d'information. Bien qu'ils ne se soient pas ou peu cités mutuellement, et bien que la présentation de leurs travaux soit très différente, nous constatons une grande similarité sur le fondement de leurs conseils en matière graphique. Ce fondement pose par ailleurs la base de notre questionnement à propos de l'ornementation. Bertin et Tufte conseillent de manière générale de réaliser des graphiques simples, sobres, épurés, où seule l'information a sa place. Alors que le mot « minimalisme » nous vient en tête, d'autres auteurs l'ont verbalisé auparavant (Inbar et al., 2007 ; Hill & al., 2017; Cairo, 2012). Pour saisir le fondement de cette préconisation de minimalisme dans la réalisation des visualisations de données, il faut comprendre ce qu'évoquent les définitions de l'efficacité (Bertin, 1967) et de l'excellence graphique (Tufte, 1983), pensés de manière similaire par les deux précurseurs, avant d'aborder leurs préconisations, porteuses elles aussi de cette intension minimaliste.

Nous ne présenterons pas tous les principes de conception des visualisations de données. Notre objectif ici est de montrer en quoi nos deux auteurs de référence formulent des préconisations minimalistes en termes de design graphique.

4.1. Les éléments de base du système graphique

Avant tout, il est pertinent de définir quels sont les éléments basiques à prendre en compte dans une représentation graphique. Sur base de données et en fonction du besoin et des questions auxquelles répondre, nous communiquons une information. Retenons dès lors que l'information est « le contenu traductible d'une pensée. Il est constitué essentiellement par une ou plusieurs correspondances originales entre un ensemble fini de concepts de variation et un invariant » (Bertin, 1967, p. 9). Ce contenu n'est donc pas à confondre avec le contenant qui, dans notre cas, sera la représentation graphique, elle-même constructible grâce aux moyens du système graphique. La définition de Bertin en ce qui concerne

l'information est celle que nous garderons en tête à chaque fois que nous utiliserons ce terme dans le cadre des visualisations de données.

Dans l'information à transcrire, il y a deux concepts à prendre en compte. Premièrement, il y a l'invariant, qui est la définition invariable valable pour toutes les données. On en a une connaissance précise. Deuxièmement, il y a ce que Bertin appelle « les composantes ». Nous appellerons les composantes « variables », puisqu'il s'agit des éléments variables de l'information. Dans l'exemple de la Figure 2, l'invariant est « nombre de magazines (ordonnées) en fonction du temps (abscisses) ». Les variables sont (1) les quantités en millions d'exemplaires et (2) le temps en années. Les données sont encodées sur les dimensions du plan, qui sont au nombre de deux (x,y sur la Figure 2). Lorsqu'une troisième dimension est à encoder, on utilise les variables visuelles. Aussi appelées « variables rétinienne » par Bertin (1967), il s'agit des éléments sur lesquels il est possible d'encoder de l'information supplémentaire : la couleur, la taille, la forme, le grain, l'orientation, la valeur.



Figure 2 Vente de magazines depuis 2005, graphique à titre d'exemple, par l'auteure

Ensuite, il y existe plusieurs types de données. On peut distinguer les données quantitatives et qualitatives.

Les données quantitatives sont de l'information d'ordre numérique ou statistique. Les représentations graphiques concernées sont donc celles qui nous montrent des nombres et des chiffres qui se rapportent clairement à des quantités. Les données quantitatives se subdivisent ainsi en deux sous-types de données : les données discrètes et les données continues (Howell et al., 2008). Les données discrètes concernent des valeurs numériques qui sont entières et limitées. Il pourra par exemple s'agir du nombre de chanteurs dans une chorale de 20 personnes. Le nombre entier se limite à 20 et il ne peut pas y avoir de valeur intermédiaire entre une ou deux personnes. Par exemple, il est impossible que 3,7 personnes

participent à la chorale. En opposition, les données quantitatives peuvent être infinies ou contenues dans un intervalle sur lequel il sera toujours possible d'opérer des subdivisions. Le poids ou la taille en sont des exemples. Ainsi, imaginons un cuisinier qui dispose d'un kilo de farine à peser dans sa balance. Il pourra diviser sa farine en toutes les valeurs possibles : 256,8 grammes ou 999,87 grammes si ça lui chante. Ces valeurs ne sont pas limitées.

Les données qualitatives sont aussi appelées données catégorielles. Au fond, il peut s'agir de valeurs numériques mais qui se rapportent plutôt à des codes qu'à des quantités. Cela peut donc être des numéros d'habitation ou une taille de vêtement. Même si ces données ne sont pas des quantités, elles peuvent être ordonnées. Il s'agira alors de données ordinales, entre lesquelles il y a une relation d'ordre : « petit, moyen, grand ». Les données qui ne sont pas ordinales sont alors nominales. Par exemple, il n'y a pas de relation d'ordre entre les deux sexes (Homme, Femme). Ces données nominales, de nature catégorique, permettent ainsi de comparer (Ware, 2013).

En ce qui concerne les données textuelles, elles se rapportent tout simplement au langage. Il peut donc s'agir d'hashtag prélevés sur les réseaux sociaux, qu'on pourrait illustrer en un nuage de mots, de commentaires sur un forum ou encore de lettres traditionnelles. Elles nécessitent un traitement d'ordre linguistique et statistique (Kirk, 2016).

Stevens parle davantage de « type de nombres ». Pour lui il y a quatre niveaux de mesures :

- Nominal : c'est la fonction de catégorisation dont nous parlions précédemment.
- Ordinal : c'est une fonction permettant d'ordonner (cf. *supra*).
- Intervalle : cette fonction permet de montrer le creux entre deux valeurs de données. L'intervalle entre deux catégories a toujours la même valeur et un zéro de signifie pas qu'il y a absence de phénomène. Exemple : la température
- Ratio : c'est le plein pouvoir expressif d'un nombre réel. On pourrait par exemple dire « l'objet A est deux fois plus grand que l'objet B ».

(Stevens, 1946, cité par Ware, 2016, p. 27)

4.2.Efficacité et excellence graphique

Bertin définit l'efficacité graphique comme ceci : « Si pour obtenir une réponse correcte et complète à une question donnée, et toutes choses égales, une construction requiert un temps d'observation plus court qu'une autre construction, on dira qu'elle est plus efficace pour cette question » (Bertin, 1967, p.139). Bertin pense l'efficacité comme un critère précis à partir duquel peuvent être classées les constructions graphiques. Sa définition de l'efficacité se base sur la « notion de coût mental de la perception [...] appliquée à la perception visuelle. [...] L'efficacité est liée à la facilité que rencontre le lecteur à chacune des opérations mentales groupées dans ce que l'on appelle la lecture d'un dessin » (Bertin, 1967, p. 139). L'efficacité est plutôt traduite chez Bertin en termes de fonctionnement : un graphique va mieux « fonctionner » qu'un autre. L'auteur est, dans sa vision du terme, très tourné vers le récepteur qui cherche à répondre à une question. Nous retiendrons donc un certain aspect « utilisateur ». Metz commente positivement le principe de Bertin : c'est un critère précis et mesurable qui constitue l'efficacité. « Un graphique est d'autant plus efficace qu'est élevé le nombre des questions auxquelles il permet de répondre par une seule image. » (Metz, 1971, p.162). A ne pas confondre avec l'esthétique, il s'agit de la transmission du plus gros volume d'information possible tout en respectant un certain niveau de rapidité et de netteté (Metz, 1971).

Pour Tufte, « l'excellence graphique consiste à donner au lecteur une vue sur le plus grand nombre d'information, dans le temps le plus court, dans l'espace le plus petit, et avec le moins d'encre possible » (Tufte, 1983, p. 51, notre traduction). Les idées complexes sont ainsi communiquées avec clarté et précision. La fonction majeure du graphique est de *révéler* les données (Tufte, 1983, p.51). Il présente l'excellence graphique comme essayant d'atteindre un certain niveau de qualité, parce que le graphique doit être d'une telle efficience que le récepteur pense d'abord à sa substance (ses données) avant même d'avoir le temps de penser à la méthodologie déployée pour mettre le graphique en place (Tufte, 1983). Chez Tufte, cette idée d'excellence graphique est un peu plus complexe parce qu'elle fait écho à un tas d'autres principes que nous aborderons ensuite : l'intégrité graphique, la maximisation du *data-ink ratio* ou encore l'évitement du *chartjunk* (cf. *infra*). Ces principes visent tous à montrer les données uniquement, dans un espace-temps réduit.

L'efficacité et l'excellence graphique se recoupent sans aucun doute, comme en atteste Kim : « le critère de l'efficacité que Bertin nous présente en mettant l'accent sur la dimension temporelle de l'observation n'est pas loin de l'excellence graphique promulguée par le gourou américain du design graphique » (Kim, 2008.). On constate bel et bien chez les deux auteurs la volonté de représentations graphiques « meilleures », « excellentes », «

efficaces ». Par efficace, on entend de « permettre à l'homme d'intégrer en quelques instants des relations complexes » (Bertin, 1967, p.146) entre plusieurs composantes et ce, « à condition d'utiliser le minimum d'images » (Bertin, 1967, p.146). En ce qui concerne l'excellence graphique, il s'agit de pouvoir communiquer des idées complexes avec clarté, précision et efficience (Tufte, 1983). À ce stade, Bertin et Tufte ne considèrent l'efficacité d'une visualisation de données que du point de vue de la perception visuelle. Leur vision est finalement assez fonctionnelle : il s'agit de faire tourner les rouages d'un système fonctionnant de manière similaire chez tout le monde pour qu'un message puisse être décodé comme il a été conçu. Ce système est la perception visuelle. Néanmoins, comme nous le verrons par la suite, la construction de sens d'une visualisation de données n'est pas si simple et linéaire. Voici ce qui nourrit notre travail de recherche : la construction de sens opérée à partir d'une visualisation de données dépasse la simple perception visuelle et la cognition, alors que ces deux précurseurs les prennent uniquement en compte pour dicter leurs principes de conception qu'ils veulent minimalistes. Tout comme nous, Kirk pense que la définition de l'efficacité peut être élargie à bien des égards : on peut imaginer que la réception d'une visualisation diffère d'une personnalité à l'autre, des compétences de l'individu, etc. Il y a ainsi de nombreuses manières de définir l'efficacité d'une visualisation de données (Kirk, 2016).

4.3. Vers un minimalisme assumé : principes de conception

Plusieurs principes de conception des visualisations de données dictés par Bertin (1967) et Tufte (1983) visent la production de graphiques épurés. Nous les présentons successivement dans cette section en mettant en évidence leurs spécificités et similitudes.

4.3.1. L'intégrité graphique : parce que l'esprit humain est crédule

Comme Cairo le dit pertinemment, « ce que vous concevez n'est jamais ce que votre public finit par interpréter exactement, il devient donc crucial de réduire les risques de mauvaise interprétation » (Cairo, 2016, 1275, notre traduction). Effectivement, bien que la construction de sens ne soit pas uniquement qu'une question de perception visuelle, il s'agit tout de même d'un élément essentiel à contrôler. Nombreux sont les biais qui peuvent perturber la réception de l'information ou, tout simplement, la transmission intégrale des données. Le premier biais, fréquent et évitable, est celui de la distorsion visuelle. Elle intervient lorsque la représentation visuelle des données n'est pas cohérente avec les valeurs numériques (Tufte, 1983, p.55). C'est là qu'intervient le principe d'intégrité graphique. Selon Tufte, il arrive que certains graphiques représentent mal les données, de manière volontaire ou non.

Le tout premier principe de Tufte est celui-ci : la représentation visuelle des données doit être directement proportionnelle à la représentation numérique. La violation de ce principe se vérifie grâce au calcul du « *lie factor* », qui établit le rapport entre la taille de l'effet montré dans le graphique (représentation visuelle) et la taille de l'effet montré dans les données (représentation numérique). La valeur résultant de ce calcul décrit la variation entre l'effet montré par un graphique et l'effet produit par les données. La représentation de nombres physiquement mesurés sur la surface du graphique lui-même doit être directement proportionnelle aux quantités représentées (Tufte, 1983, p.57). Le *lie factor* se calcule ainsi :

$$\text{Lie factor} = \frac{\text{taille de l'effet montré dans le graphique}}{\text{taille de l'effet montré dans les données}}$$

(Tufte, 1983, p.57)

La taille de l'effet se calcule en soustrayant la dernière valeur à la première et en divisant le résultat par la première valeur⁴. Pour assurer l'intégrité du graphique, le *lie factor* doit être compris entre 0,95 et 1,05. Si l'indicateur ne se trouve pas dans cet intervalle, le graphique comporte des distorsions et inexactitudes. Le graphique est alors dans une situation mensongère (mensonge par omission ou intentionnel) (Tufte, 1983, p.57).

De même, afin de satisfaire l'intégrité graphique, le respect des six principes suivants est de mise :

- La représentation de nombres sur le graphique doit être directement **proportionnelle** aux quantités numériques à représenter (*lie factor*) ;
- Des **étiquettes** (labeling) claires et détaillées doivent être utilisées pour éviter l'ambiguïté et la distorsion dans le graphique. Écrire des explications à même le graphique ou préciser quels événements sont "importants" pour les données peuvent également être d'une grande aide ;
- Il est important de montrer les **variations de données** et non les variations de design. En d'autres mots, c'est sur les données que l'accent doit être porté et non le design ;
- Pour les visualisations de données concernant de l'argent, les unités de mesure relatives à la monnaie doivent être **déflatées et standardisées**. Elles valent toujours mieux que des unités nominales ;

⁴ Taille de l'effet = $\frac{(\text{la dernière valeur} - \text{la première valeur})}{\text{la première valeur}}$

- Le nombre de dimensions représentées comme support d'information ne doivent pas excéder le nombre de dimensions réellement contenues dans les données ;
- Les graphiques ne doivent **pas citer les données hors contexte**.

(Tufte, 1983, p.77, notre traduction)

Aux sources du non-respect de l'intégrité graphique, Tufte identifie le manque de compétences quantitatives des concepteurs, l'idée reçue de l'ennui que peuvent provoquer les données statistiques ainsi que l'idée reçue selon laquelle les graphiques sont réalisés pour des personnes incapables de saisir des données statistiques.

4.3.2. La simplicité, révélatrice de sens

La principale caractéristique des principes mis en avant par les deux auteurs est la simplicité esthétique. En effet pour Bertin, il est bon de se demander sur quelle base se construit l'image. Ainsi, avant même d'émettre des principes de design, il revient sur les bases de la graphique en abordant le plan et les différentes variables qu'on peut y inscrire. En effet, c'est en utilisant les dimensions du plan que l'on construit la représentation graphique. Sur le plan, on peut réaliser l'image selon deux ou trois dimensions (la troisième dimension « s'élève » alors du plan). La visualisation de données statique⁵ pourra ainsi admettre trois variables maximum (quatre, dans de rares cas). S'il y a plus de variables à inscrire, il faut penser à rédiger son information en deux représentations visuelles distinctes. En effet, « la perception visuelle n'admet qu'un nombre réduit de variables » (Bertin, 1967, p.9). Ainsi, Bertin émet les principes suivants :

- Dans le plan, l'image est créée sur des **variables homogènes et ordonnées**. Les deux premières variables se concentrent sur les dimensions du plan tandis que la troisième variable est placée en troisième dimension ;
- Pour qu'une seule image seulement puisse contenir l'information à transcrire, les variables doivent être homogènes et l'information doit se rapporter à un concept ordonné. Sinon, il faudra plus d'une image ;
- Si un problème graphique se concentre autour de deux variables visuelles, la construction de base sera le schéma : 
- Si le problème graphique relève de trois variables visuelles, la construction sera différente, respectant le schéma :  Cela signifie que la troisième variable

⁵ Nous parlons ici uniquement de visualisations de données statiques. Les visualisations de données interactives permettent d'inscrire un nombre plus grand de variables

s'élèvera sur la troisième dimension du plan. On utilise alors des variables visuelles pour coder cette information (couleur ou taille par exemple).

(Bertin, 1967, p. 148, 149)

Il est primordial pour Bertin d'être incollable dans la connaissance de ses propres variables. En effet, une représentation graphique doit pouvoir favoriser les comparaisons externes de son observateur, en stimulant sa lecture au plus haut niveau. Pour atteindre ce niveau et cette qualité graphique, Bertin préconise trois principes :

1. « **Construire une information en une image** » (Bertin, 1967, p. 171) ou minimiser les images nécessaires à la compréhension de l'information ;
2. En ce qui concerne les représentations à composantes ordonnables, **simplifier l'image au maximum**, en prenant garde à ne pas réduire le nombre de correspondances ;
3. **Réduire** l'image pour la simplifier lorsqu'elle comporte des composantes ordonnées.

(Bertin, 1967, p.171)

Les propositions (2) et (3) se recoupent parfaitement avec deux principes mis en avant par Tufte.

Le deuxième principe de Bertin, visant à simplifier l'image au maximum, va dans le même sens que le principe « *data-ink ratio* » de Tufte. C'est bien de "l'encre" utilisée sur un graphique dont Tufte se préoccupe ici. L'usage de technologies avancées étant peu répandu à l'époque, on comprend son choix pour ce terme. Pour Tufte, les données deviennent information une fois "dessinées" sur papier, présentes visuellement. Tufte appelle "*data-ink*" le cœur non effaçable d'un graphique, c'est-à-dire les éléments non-redondants, indépendants des nombres, présentes sur papier grâce à l'encre (Tufte, 1983, p. 93). Le *data-ink ratio* se calcule lui aussi, et à l'heure actuelle, on peut le calculer à l'aide des pixels sur ordinateur, tout comme des centimètres sur papier. Voici comment il calcule ce ratio :

$$\text{Data-ink ratio} = \frac{\text{Encre consacrée aux données}}{\text{Total d'encre utilisée dans le graphique}}$$

(Tufte, 1983, p. 93)

L'intérêt d'un tel ratio est de pouvoir comprendre dans quelle mesure l'encre utilisée sur le graphique est bel et bien dédiée aux données. Il y a ainsi dans le graphique le « *data-ink* », c'est-à-dire l'ensemble de l'encre qui se rapporte directement à l'information communiquée, et le « *non data-ink* », concernant dès lors l'encre qui ne concerne pas cette

information (bordures et traits inutiles, fioritures, etc.). Pour Tufte, plus "l'encre" utilisée dans le graphique réfère aux données, mieux c'est.

C'est ainsi que Bertin et Tufte se rejoignent : pour Tufte, il faut ainsi maximiser le *data-ink ratio* (principe de *data-ink ratio maximisation*). Chaque goutte d'encre utilisée a une raison d'être : celle de présenter une nouvelle information. Tufte reconnaît que ce principe peut avoir une grande conséquence sur le design du graphique mais il le juge de bon sens (Tufte, 1983, p.96). Pour augmenter la proportion de *data-ink*, une manipulation simple consiste à effacer, dans les limites du raisonnable, toute l'encre qui ne concerne pas les données (*non data-ink*). De même, il est important d'éviter le *data-ink* qui est redondant. Il vaut mieux effacer cette redondance si elle n'a pas de raison d'être. Ainsi donc, des traits qui ne sont pas en relation pourraient embrouiller le graphique (Tufte, 1983, p.96). Il s'agit d'un principe d'effacement. On va là vers la simplification dont parle Bertin. Bertin dit bien qu'il faut simplifier l'image « en prenant garde à ne pas réduire le nombre de correspondances ». On comprend ainsi que pour lui aussi, il convient de ne garder que l'information nécessaire, se rapportant aux données, sur le graphique.

Ensuite, la troisième proposition de Bertin en termes de simplification est celle-ci : « simplifier l'image par réduction » (Bertin, 1967, p.171.). Cela correspond au « *shrink principle* » sur lequel Tufte insiste. Pour lui, les graphiques peuvent être réduits à leur plus simple expression (Tufte, 1983, p.167). C'est d'ailleurs une façon d'augmenter la densité des données dont les deux auteurs parlent. Lorsqu'il explique le *shrink principle* dans son ouvrage « *The visual display of quantitative information* » (1983), Tufte cite la 214^{ème} page de Bertin dans « *La Sémiologie Graphique* », ventant ainsi les élégants graphiques à petite échelle de cette page, respectant parfaitement son principe. À notre connaissance, c'est l'une des seules fois où Tufte cite son prédécesseur.

Quoi qu'il en soit, les deux auteurs sont sur la même longueur d'ondes : ils prônent la simplification de l'image. Ils en arrivent à cette conclusion car « les fonctions de la représentation graphique forment une chaîne définie par le jeu de la mémoire et qui concourt (...) à la simplification » (Bertin, 1967, p. 166). Puisque la représentation graphique vue comme un objet communicationnel a pour objectif la compréhension du lecteur, sa simplification est obligatoire : « la simplification est créatrice » (Bertin, 1967, p. 166). Bertin appelle le traitement de l'information « simplification logique » car, tout simplement, « le traitement graphique de l'information s'opère par la simplification de l'image » (Bertin, 1967, p. 166).

Pour terminer sur cette idée de simplicité, il convient néanmoins d'ajouter que simplifier la visualisation ne veut pas dire l'appauvrir. Selon Bertin, la simplification est révélatrice

de sens (1967). Il ajoute que « la simplicité des lectures découle du contexte d'une information détaillée et complexe, bien organisée » (Tufte, 1990, p.37, notre traduction). Dans la simplicité et la maximisation du *data-ink ratio*, le souci du détail a donc toute son importance.

Ainsi, pour Cairo, de manière générale, « les graphiques ne doivent pas simplifier les messages. Ils doivent les clarifier, mettre en évidence les tendances, découvrir des modèles et révéler des réalités qui n'étaient pas visibles auparavant » (Cairo, 2012, p. 1286).

4.3.3. Le chartjunk et la mise en garde contre les intentions esthétiques

Poursuivons ensuite avec le *chartjunk*. C'est un concept fortement connu dont Tufte est l'initiateur. Effectivement, il pose question quant à la mémorisation d'une information "mise en scène" et non plus seulement racontée à travers la simple figuration des données. Il s'agit d'une forme de visualisation de données fortement rencontrée, en tout temps, et de fait, critiquée par l'auteur. Le *chartjunk* se définit comme un concentré d'éléments qui ne correspondent pas aux données et d'éléments redondants qui concernent les données, le tout dans un même graphique. Il se manifeste dans des graphiques qui exposent des décorations artistiques ou dans des graphiques conventionnels dans lesquels des éléments inutiles n'ajoutent aucune valeur (Few, 2011). Tufte distingue trois types de *chartjunk*, qu'il nomme de façon assez poétique : l'art optique non intentionnel (*unintentional optical art*), le quadrillage obsolète⁶ (*dreaded grid*), le canard boiteux ou canard d'auto-promotion (*self-promoting duck*) (Tufte, 1983, p. 108).

Par "art optique non-intentionnel", Tufte suggère qu'une représentation graphique peut engendrer un tremblement physiologique dans l'œil du récepteur (et être, à la limite, désagréable). Un mouvement de vibration, appelé « moiré » apparaît alors (Tufte, 1983, p.107).

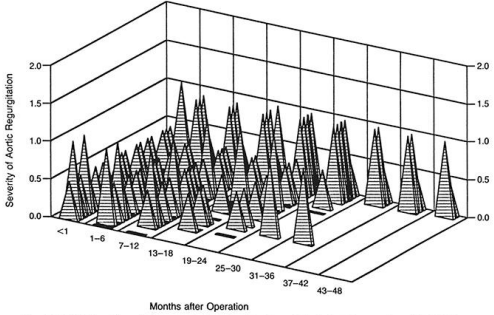
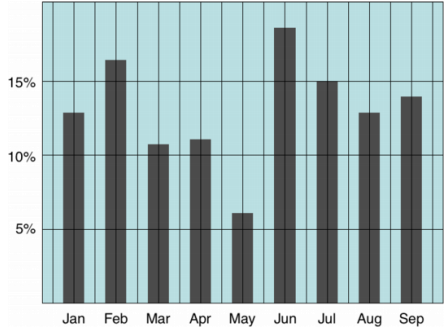
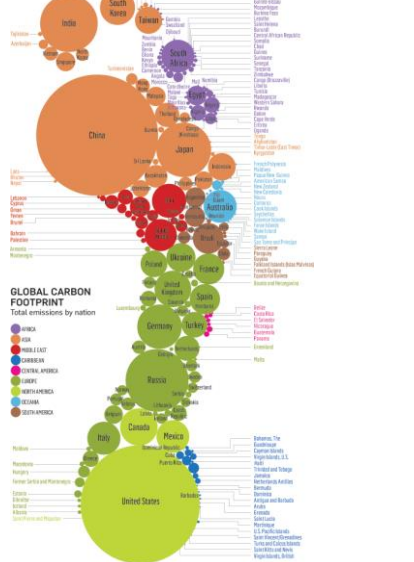
Le quadrillage obsolète, quant à lui, est un élément qui devrait être supprimé lorsqu'il s'avère inutile. A de nombreuses reprises, ce n'est pas le cas. La présence de la grille entre alors en "compétition" avec les données présentées (Tufte, 1983, p. 112). Par contre, si une grille est essentielle et aide à lire, il est primordial de la faire apparaître (Tufte, 1983, p. 116).

Le "canard boiteux" (*self-promoting duck*) est un style de graphique réalisé pour sa forme décorative et pour le design avant tout, et non en fonction des formes qui mettraient les données le plus en valeur. Les mesures de données deviennent alors tout simplement de

⁶ Littéralement « grille redoutée », que préférons traduire par « quadrillage obsolète »

purs éléments de design faisant prévaloir le style sur l'information quantitative (Tufte, 1983, p. 116). De plus, on ajoute souvent à cela une fausse perspective sur la structure des données, ce qui est assez trompeur (Tufte, 1983, p. 118). Différents exemples de ces types de *chartjunk* sont présentés dans le Tableau 2.

Tableau 2 Exemples réels des différents types de chartjunk

Art Optique Non Intentionnel <i>Unintentional Optical Art</i> Outre l’effet de 3D qui est néfaste à la compréhension de ce graphique, il y a un effet de moiré dérangeant pour l’œil au sein même des triangles représentant les données sur ce graphique. L’œil s’y perd. Exemple issu de (Tufte, 1983, p. 109)	 Figure 2. Serial Echocardiographic Assessments of the Severity of Regurgitation in the Pulmonary Autograft in 31 Patients. The numerical grades were assigned according to the severity of regurgitation, as follows: 0, none; 0.5, trivial; 1.0 to 1.5, mild; 2.0, moderate; and 3.0, severe.
Quadrillage obsolète <i>Dreaded grid</i> Outre le fond coloré qui lui aussi est inutile à la compréhension des données, la grille présente derrière les bâtonnets de ce graphique est obsolète. Elle n’aide pas la compréhension et le graphique pourrait très bien exister sans elle. Exemple issu de (Bray T., 2019, http://www.tbray.org/ongoing/data-ink/di1)	
Le Canard ou Graphique d’auto promotion <i>The Duck or Self-Promoting Graph</i> Il est difficile de lire et de saisir en un coup d’œil l’information présentée dans cette visualisation de données. C’est un bel exemple de canard, car l’organisation et la présentation des données a été pensée en fonction de l’apparence visuelle de l’image et non de la pertinence des données. Exemple issu de (Stanford Kay, 2010, http://www.stanfordkaystudio.com/information.html)	

Bertin, lui aussi, laisse peu de place aux efforts artistiques. « [...] Toute variation visuelle apparaît comme significative et l’introduction d’une variation de couleur par exemple, à seule fin esthétique ou décorative est une source d’inefficacité si les différences de couleur

ne correspondent à aucune variable » (Bertin, 1967, p.46). De plus, pour lui, tout effort esthétique dont l'intention est décorative est à éviter. Cela semble cohérent avec l'objectif du système monosémique qu'il utilise pour encoder l'information.

Bertin s'oppose ainsi à l'image polysémique ou pansémique, qui permet l'interprétation parce qu'elle est figurative. « L'image graphique est passée de *l'image morte*, de « l'illustration », à *l'image vivante*, à l'instrument de recherche accessible à tous » (Bertin, 1967, p.8). Il utilise clairement le niveau monosémique de la perception spatiale en ne permettant qu'une seule interprétation. Dans le cas contraire, « l'interprétation est liée au répertoire d'analogies et de hiérarchies de chaque « récepteur » » (Bertin, 1967, p.6). Sachant que ce répertoire est différent d'une personne à l'autre, Bertin veut l'éviter en mobilisant un système de signes arbitraire. Il n'y a aucune place pour l'ornementation dans la visualisation de données selon lui, puisque cela va à l'encontre de l'objectif de la graphique.

Selon Tufte, le *chartjunk* n'atteint pas le but de ses propagateurs. Le graphique ne deviendrait pas plus intéressant et plus attractif parce que des éléments ornementaux y seraient placés. Le *chartjunk* peut ainsi transformer les données et le faire tourner au désastre, selon l'auteur. Le meilleur design provoque l'intrigue et la curiosité et révèle la vraie teneur des données, parfois grâce à un sens narratif, au sens du détail et parfois par une présentation simple et élégante de données intéressantes. Aucune information, découverte ou toute autre substance n'est produite par le *chartjunk* (Tufte, 1983). Ainsi le principe de Tufte est on ne peut plus clair : « oubliez le *chartjunk*, en ce compris les effets de vibration, la grille et le canard » (Tufte, 1983, p. 121). Dans son livre "*Envisionning Information*", Tufte inclut également dans le *chartjunk* les ornements et ajouts esthétiques, plus ludiques. Il pense que les promoteurs du *chartjunk* imaginent les données comme ennuyeuses, et qu'ils ont besoin d'ajouter de la décoration afin de briser l'ennui (Tufte, 1990, p.34). A la lecture de ses écrits, on peut sentir une attitude assez militante, fondamentalement contre les ornements qu'il qualifie de *chartjunk*.

Pour Tufte, il n'y a pas de demi-mesure : une visualisation compréhensible, respectant les principes de conception des visualisations de données dans une certaine mesure et présentant des ornements n'existe pas. Ces ornements sont du *non data-ink* et n'apportent rien à l'information. Est-ce vraiment le cas ? C'est ce que nous étudions dans cette thèse de doctorat⁷ : ce *chartjunk*, ces ornements sont-ils susceptibles d'apporter quelque chose de positif ou négatif à la construction de sens ? Tout le monde n'est pas forcément d'accord

⁷ D'ailleurs, Tufte qualifierait probablement les visualisations que nous étudions comme des canards (*self promoting duck*).

avec cette vision négative des éléments ornementaux. La notion de *chartjunk* pousse fortement au débat. Presque 20 ans après la première apparition du mot « *chartjunk* », une étude intitulée "*Useful Junk? The Effects of Visual Embellishment on Comprehension and Memorability of Charts*" (Bateman & al., 2010) remet le concept en question. Cette étude teste l'influence du *chartjunk* et du graphique « minimaliste » sur la compréhension et sur la mémorisation du récepteur. Elle montre qu'à propos des mêmes données, le *chartjunk* n'est pas forcément moins bien compris qu'un graphique minimaliste. Mais surtout, après 2 à 3 semaines, le *chartjunk* permet une mémorisation de l'histoire racontée par les données de manière nettement significative (Bateman & al., 2010). Ainsi, l'ornementation pourrait avoir un effet positif. Nous aborderons cela plus en détail dans les chapitres suivants.

Il y a donc de la nuance à apporter au discours de Tufte, qui est tout à fait négatif. De plus, les participants comprenaient la valeur du message plus souvent dans les *chartjunk* que dans les graphiques minimalistes et surtout, ils ont souligné le plaisir de les utiliser, ce qui rendait la mémorisation plus facile (Bateman & al., 2010). Few, bien que se montrant impressionné par l'étude, se permet néanmoins une remarque : « les exemples dans cette étude ont consisté seulement en quelques graphiques embellis, créés par une seule personne – Nigel Holmes – et ils n'ont été un succès que dans un but extrêmement précis et limité » (Few S., 2011, p.4, notre traduction). Few établit une critique méthodologique forte à l'encontre de l'article qui, à son avis, fait le buzz. De nombreuses questions se posent encore. Que se passerait-il en n'utilisant pas forcément des cas extrêmes de *chartjunk* ? (Few, 2011). Y aurait-il donc différentes catégories ou niveaux de *chartjunk* ? Ces critiques le laissent penser. L'exploration est donc encore possible dans ce domaine et semble prometteuse. En se focalisant sur l'utilisateur et en choisissant pour angle d'attaque les sciences de la communication, il y a fort à parier que le discours sur le *chartjunk* serait différent.

Cependant, pour Few, Tufte décrit trop légèrement le *chartjunk*. En effet, les traces minimales ou les éléments redondants peuvent être porteurs d'un message. « Les ornements qui représentent les données, même métaphoriquement, peuvent eux-mêmes être considérés comme « *data-ink* ». Les ornements qui ne sont pas des données en eux-mêmes peuvent parfois supporter le dispositif des données de manière utile » (Few S., 2011, p.10, notre traduction). Les ornements graphiques peuvent selon Few supporter l'efficacité de la visualisation de données

- « - en engageant l'intérêt du lecteur ;
- en attirant l'attention sur des éléments qui méritent qu'on insiste ;
- en rendant le message plus facilement mémorisable »

(Few, 2011, p. 10, notre traduction)

Si cela nous pousse à étudier les différents effets de ce que nous appelons plus globalement « l'ornementation »⁸, nous restons cependant prudente : il ne s'agit pas forcément d'une bonne chose dans le cas où les données ne sont pas communiquées avant tout. C'est d'ailleurs ce que Tufte essaie de dire : montrer les données, éviter la distorsion du message et la distorsion visuelle, penser aux données avant de penser au design, ... Un tas de principes qui peuvent facilement être bafoués si le concepteur produit sa visualisation de données comme un objet plaisant et non comme un objet de communication de données précis.

4.3.4. La préconisation d'une esthétique fonctionnelle et élégante

Bien que Bertin et Tufte préconisent le minimalisme en termes de visualisation de données, ils mettent en avant une esthétique en particulier. Pour Bertin, c'est assez simple puisqu'il met en avant l'utilisation d'un système de signes monosémique pour communiquer. Pour lui, l'ornement n'a pas sa place dans le graphique et il ne fait pas plus de commentaire. Néanmoins, c'est un peu plus nuancé chez Tufte qui, lui, parle carrément « d'élégance graphique » (Tufte, 1983, p.177). Pour lui, « l'élégance graphique se trouve souvent dans la simplicité du design et dans la complexité des données » (Tufte, 1983, p. 177). Le pouvoir attractif et visuel du graphique doit ainsi venir du contenu et non d'une intention décorative. La bonne conception est donc celle qui sait conjuguer la simplicité aux jeux de données complexes. Il dit ainsi « l'esthétique est aux artistes ce que l'ornithologie est aux oiseaux » (Newman cité par Tufte, 1983, p. 177). Il n'y a qu'une esthétique valable pour Tufte dès lors : celle qui met en avant la simplicité au maximum. Les autres efforts ne concernent pas la visualisation de données. Nous le verrons par la suite, cette obsession visant à pousser le principe « *less is more* » à son maximum vaudra à Tufte d'être critiqué (Inbar et al., 2007 ; Hill et al., 2017). En préconisant des visuels si épurés et qui semblent particuliers, Tufte en vient à proposer de montrer aux individus des types de visuels avec lesquels ils n'ont jamais été en contact. Cela peut donc être une contrainte : on le voit dans la Figure 3 qui, testée par Inbar et ses collègues dans leur étude en 2007, se voit « boudée » par les lecteurs qui, dès lors, préfèrent des formes de visualisation conformes à leurs habitudes. Dans le chapitre 2, nous allons aborder cette grande question de l'ornementation, que Tufte qualifie de *chartjunk*. La question que

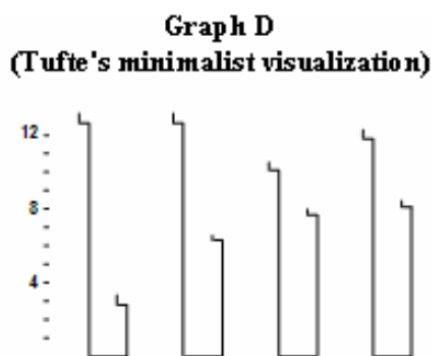


Figure 3 Graphique minimaliste selon Tufte, utilisé dans l'expérimentation de (Inbar et al., 2007)

⁸ Ce qui nous permet ainsi de nous distancer de l'aspect militant du *chartjunk*

soulèvent Inbar et ses pairs est pertinente : et si le minimalisme poussé à l'extrême était inaccessible au citoyen lambda, parce qu'inhabituel ? L'ornementation, quant à elle, est appréciée et semble même avoir des qualités, notamment en termes de mémoration du message.

Chapitre 2 L'ornementation des visualisations de données : la théorisation d'un constat général

Ces dernières années, plusieurs auteurs ont constaté l'apparition de pratiques tendant à amener une notion de « beauté », de caractéristiques esthétiques attractives dans les visualisations de données. Face à ce constat, nous avons pour objectif d'étudier « l'embellissement » des visualisations de données. Ce même constat apparaît dans la littérature anglo-saxonne : « *embellished charts* », « *visual embellishment* », ... Néanmoins, de nombreux chercheurs sont frileux lorsqu'il s'agit d'étudier le beau en recherche. La question est là : « comment étudier le beau dans un domaine aussi codifié que la visualisation de données ? ». Tout est une question de termes.

Avec les notions de « *chartjunk* », de *data-ink ratio* ou encore d'élégance graphique, Tufte déconseille tout effort esthétique qui ne se rapporte pas aux données, menant quelque part une guerre face à l'ornementation ayant pour but de plaire avant d'informer. Tout comme Bertin (1967), il conseille des graphismes sobres, épurés, où tout pixel est disposé sur la représentation à des fins utiles. Il faut simplifier le graphique au maximum et enlever tout ce qui est inutile, ce qui s'apparente à une certaine forme de minimalisme. Hill et ses pairs (2017) ou encore Inbar et ses collègues (2007) remettent en question le principe de maximisation du *data-ink ratio*, qu'ils assimilent clairement à une injonction au minimalisme. Tufte n'a jamais réellement étudié les effets de ces principes. C'est pourquoi ces auteurs s'y sont attardés. En réalité, ils n'ont pas étudié la présence d'ornementation dans les visualisations, mais plutôt l'absence de minimalisme dans les graphiques traditionnels. Couleur, grilles, échelles, ... Des éléments qui sont si familiers dans le domaine de *l'Infovis*, et qui seraient superflus selon Tufte. Entre minimalisme, *chartjunk* et ornementation, le brouillard se lève rapidement et il est important de le dissiper. De plus, certains auteurs étudient l'esthétique relative à la visualisation de données (Lau et al., 2007) tandis que d'autres se concentrent sur l'apparence esthétique en tant que style général ou en tant qu'ornementation picturale de ces dernières (Skau et al., 2015 ; Borgo et al., 2012 ; Vande Moere et al., 2012 ; ...).

Dans ce chapitre, nous souhaitons établir la clarté sur la notion d'ornementation. Il semble essentiel de définir ce qui a trait au beau, à l'esthétique et au design. Nous comparerons ainsi diverses études qui traitent de l'« *embellishment* » et nous montrerons en quoi cette notion n'y est pas illustrée de la même manière avant d'en arriver aux questions de vocabulaire. Il s'agira finalement de définir le concept central de cette recherche : « l'ornementation ».

1. Ornementation ou *embellishment* ?

Comme nous le disions, notre attention a été attirée par diverses études qui n'ont pas le même discours à propos du concept qu'elles nomment « *embellishment* » en anglais. D'une même manière, ces études interrogent les préconisations de design minimaliste dans les visualisations de données, telles que propagées par Tufte et Bertin. Cette remise en question vient d'un constat général de l'utilisation croissante d'un potentiel enjolivement dans toutes les sphères d'activité qui utilisent la visualisation d'information : mass-médias, activité socio-numérique, visualisation ambiante, activités internes dans les organisations, ...

A partir de ces études, nous constatons qu'un flou certain gravite autour de la notion d'*embellishment*. Nous identifions pourtant certains points de convergence, que nous détaillons ci-après. Quoi qu'il en soit, l'ornementation se place la plupart du temps dans un contexte d'évolution médiatique et technologique. L'utilisation croissante des infographies dans les médias, ainsi que l'expression d'informations chiffrées sous forme visuelle, est en pleine recrudescence. Ce contexte d'apparition est donc important. Par ailleurs, le terme « *embellishment* » n'est jamais défini clairement ni explicité dans ces articles scientifiques qui sont anglophones. Nous pouvons affirmer qu'il consiste en général en un apport esthétique à une forme standard de visualisation de données (*raw chart* vs. *junk chart*). Cet apport peut consister en un apport pictural, métaphorique ou encore de métamorphose des traits relatifs aux données. Bien qu'il dépasse largement la limite de l'acceptable en termes d'efficacité pour Tufte et Bertin, cet *embellishment* n'est pas seulement un « ajout » à un *raw chart*. Il fait sans aucun doute partie du processus de conception, sans forcément constituer une étape à part entière. Selon nous, « embellissement » n'est pas la traduction correcte du terme anglais « *embellishment* ». La traduction française et non littérale correspond plutôt à l'enjolivure, au sens de la décoration, de l'ornement. Le second sens à la française que l'on pourrait donner à « *embellishment* » est l'exagération d'une histoire, son enjolivement prononcé (Collins Dictionary, 2017, en ligne). Dans ces notions, on décèle bien la volonté « d'ajouter » quelque chose à la forme graphique de base, de la modifier. Mais comme nous venons de le mentionner, il ne s'agit pas d'un ajout, disposé après la visualisation-même de l'image, mais d'un acte de conception. Nous avons donc

choisi de traduire le terme *embellishment* par « ornementation », en prenant conscience qu'elle fait partie intégrante de la conception de la visualisation de données, et qu'il ne s'agit pas d'une étape ultérieure visant à exagérer l'information. Nous nous en tenons donc à la première traduction littérale de « *embellishment* », et évitons ainsi le terme « embellissement », par prudence et afin de ne pas tomber dans les travers du beau en recherche. L'ornementation est donc une pratique en visualisation de données.

Afin de comprendre ce qui motive notre définition de l'ornementation, nous proposons de revenir sur les études qui l'ont inspirée. La plus connue d'entre elles – et sans doute celle qui suscite le plus de réactions – a été menée par Bateman et ses collègues en 2010. Ils y prennent le contrepied de Tufte en montrant que le « *chartjunk* » est moins néfaste qu'on ne le pense. Pour Tufte, le *chartjunk* est une décoration intérieure au graphique qui génère des traits sur l'image (*ink*) non relatifs aux données, et qui donc ne sont pas porteurs d'information utile relative à la donnée (Tufte, 1983). Comme dans les études que nous présenterons ensuite, l'*embellishment* n'est présenté que sous forme d'un constat. A ce propos, Bateman et ses collègues disent : « Beaucoup de designers incluent une grande

variété d'[*embellishments*] visuels dans leurs graphiques, des petites décorations aux grandes images et arrière-plans » (Bateman et al., 2010, p. 2573, notre traduction). Dans cette étude, des visualisations de données ont été présentées à des étudiants, sur écran. Il s'agit de visualisations réalisées par un seul graphiste, Nigel Holmes. Les auteurs les considèrent comme ornementées ; Tufte crierait au *chartjunk*. Le fait qu'elles n'aient été présentées qu'à des étudiants et surtout qu'elles aient été conçues par une seule et même personne est critiqué, notamment par Few (2011), ce qui a donné un peu de plomb dans l'aile à la crédibilité de l'étude. Elle reste tout de même une référence très citée. Les participants ont donc examiné plusieurs graphiques minimalistes et ornementés

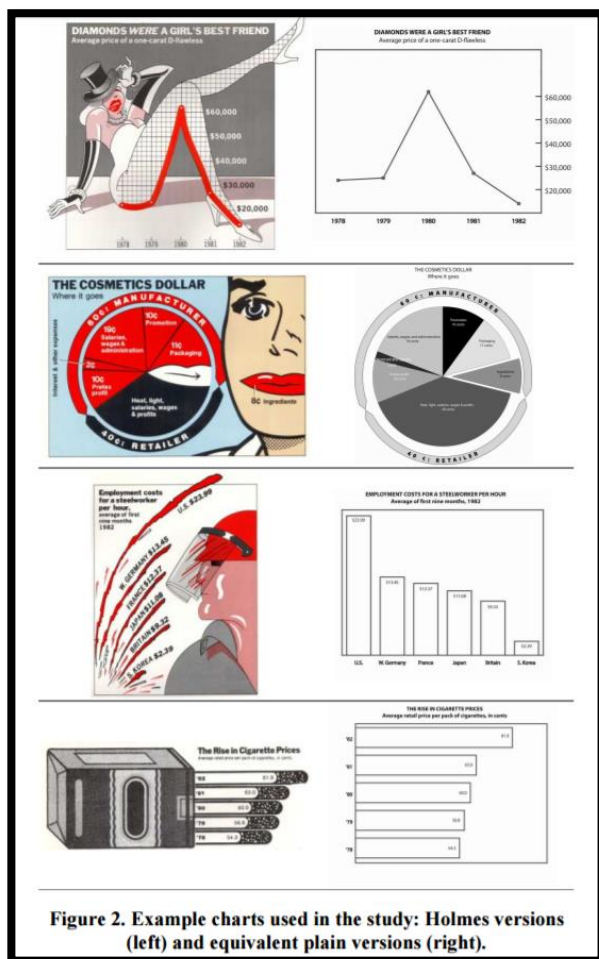


Figure 4 Bateman et al. (2010, p.2576), " Example charts used in the study: Holmes versions (left) and equivalent plain

différents et ont répondu à des questions concernant les sujets abordés par les graphiques. Ensuite, ils ont répondu aux mêmes questions après quelques minutes et après deux semaines. Le résultat est sans appel : le *chartjunk* est une aide à la mémoire. Cependant, puisque les auteurs n'ont pas clairement défini ce qu'est l'*embellishment*, nous nous sommes attardée sur les visuels ayant servi aux expérimentations (Figure 4). Comme on le voit, l'*embellishment* présent dans les visualisations considérées comme *chartjunk* incluent l'ajout d'images et touchent aussi aux traits qui représentent les données : le pixel qui représente la donnée sur le graphique est ornementé (ex. : la jambe, la cigarette, ...).

L'*embellishment* n'est donc pas qu'un ornement, c'est aussi ce qui représente la donnée. En ce qui concerne les traits relatifs aux données, ils ne respectent en aucun cas les règles et principes dictés par Tufte et Bertin. Si certains jugent leur militantisme minimaliste trop strict, il va sans dire que les images ci-contre en sont tout le contraire. L'*embellishment* selon eux est du *chartjunk*.

En 2015, Skau et ses collègues ont étudié l'impact de l'*embellishment* dans leur article « *An evaluation of the impact of visual embellishments in Bar charts* ». Comme le nom de cette étude l'indique, ils ne se sont concentrés que sur les diagrammes en bâtonnets. Bien que les visualisations utilisées soient moins étoffées que les visuels utilisés par Bateman et ses collègues, la modification du *bar chart* de base passe aussi par la modification des traits relatifs aux données. La barre rectangulaire change de forme. C'est d'ailleurs la définition approximative que les auteurs donnent au terme : « [...], les visualisations de données dans l'infographie sont souvent enrichies par l'ajout et la modification du graphique brut⁹ » (Skau et al., 2015, p.1).

Tout d'abord, la notion de *raw chart* (graphique « brut ») suppose qu'il existe une visualisation dénuée de tout artifice, représentant toutes les données de la manière la plus pure possible, dans l'idéal d'une réalité objective. Cette notion résume en deux mots l'idéal de lisibilité selon Bertin et Tufte qui sont dans une vision très objectiviste de l'esthétique. Afin de mener les expérimentations dans le cadre de cette recherche et après avoir parcouru des librairies de visualisations de données en ligne, disponibles à tous, les chercheurs travaillant sur le sujet ont repéré six pratiques d'*embellishment* en ce qui concerne les diagrammes en bâtonnets. Parce qu'ils désiraient détenir un niveau de contrôle élevé lors de leurs expérimentations, ils ne les ont pas directement réalisées avec les exemples trouvés, mais avec des prototypes reproduits sur base de ces exemples. Les *bar charts* « source » ont donc été simplifiés pour les besoins de l'étude. L'*embellishment* concerne ici les traits relatifs aux données. Ainsi il s'agit d'une transformation des traditionnels

⁹ « *Raw chart* »

rectangles qui constituent les *raw bar charts*, comme en atteste la Figure 5 directement issue de l'article. Dans ce cas-ci, l'*embellishment* testé est isolé de toute couleur, dessin, ornement. L'*embellishment* testé ne prend pas en compte les ajouts, les ornements esthétiques, ce qui diffère de l'étude précédente.

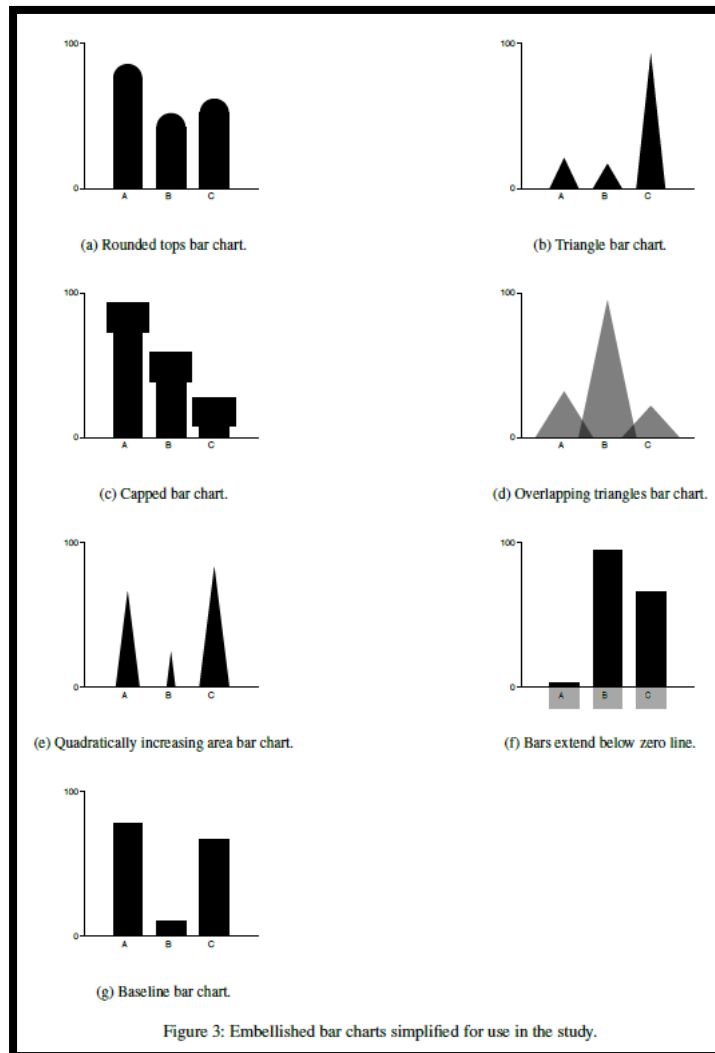


Figure 5 Skau et al. (2015), "Embellished bar charts simplified for use in the study", in *An Evaluation of the Impact of Visual Embellishment in Bar chart*, p.5

Par ailleurs, l'*embellishment* est défini de manière plus précise par Borgo et ses pairs (2012), qui l'associent à une figure de style. « Les [*embellishments*] visuels sont une forme de figures rhétoriques non linguistiques que l'on retrouve fréquemment dans les arts visuels, les arts de la scène, la publicité, les icônes et signes, les symboles culturels, le symbolisme des couleurs, les interfaces utilisateur graphiques, etc. » (Borgo et al., 2012, p. 2759, notre traduction). De fait, les auteurs ont utilisé le terme de « métaphore visuelle » en ce qui concerne les *embellishments* dans les visualisations de données jusqu'à ce qu'ils prennent connaissance des termes « *visual embellishments* » utilisés par Bateman (2010). Les

métaphores visuelles seraient, selon eux, des images qui transmettent un concept via l'utilisation d'un concept secondaire. L'*embellishment* trouverait son utilité dans la métaphore visuelle (Borgo et al., 2012). Des modifications ont été apportées à des visualisations « ordinaires ». C'est ce qu'ils ont récolté dans leur étude. Nous constatons là que l'*embellishment* consiste en l'ajout d'images, pictogrammes, symboles. Les traits propres aux données ne subissent aucun changement tandis que des éléments picturaux sont ajoutés tout autour, derrière et à l'intérieur d'eux (Figure 6).

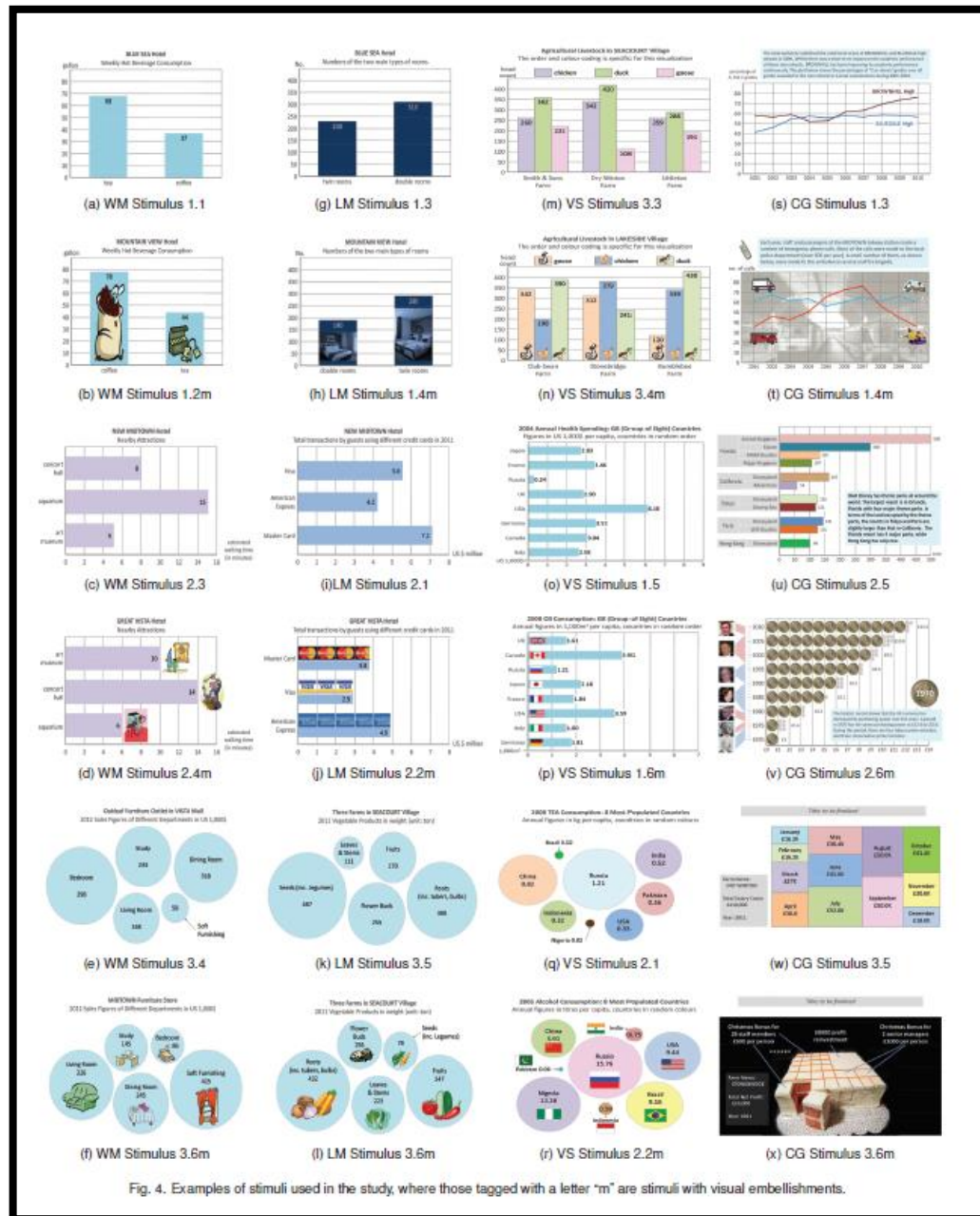


Figure 6 Borgo et al. (2012), "Examples of stimuli in the study" in *An Empirical Study on Using Visual Embellishments in Visualization*, p.2764

L'objectif de Borgo et de ses collègues était d'étudier l'impact de l'*embellishment* sur la mémoire à court terme, à long terme, la recherche visuelle et la compréhension de concepts. Ils concluent que les *embellishments* permettent une meilleure mémorisation, même si le temps de traitement au moment de la lecture de la visualisation est plus long. Leur effet est d'ailleurs très positif sur la vitesse de rappel de la mémoire (les individus se souviennent beaucoup plus rapidement lorsqu'il y a *embellishment*). Évidemment, tous les *embellishments* ne demandent pas la même charge cognitive. En fonction de leur complexité, ils ont un impact sur la performance de l'utilisateur lorsqu'il doit réaliser une recherche ciblée (Borgo et al., 2012).

La prochaine étude sur laquelle nous nous sommes attardée concerne également l'*embellishment* mais considère plutôt le « style » global des visualisations. Il s'agit d'une étude comparative menée en ligne qui porte sur trois visualisations de l'information qui affichent chacune le même ensemble de données mais d'un style esthétique différent. L'*embellishment* n'y est pas réellement défini, mais on retrouve des notions de beauté, d'art et de plaisir. Vande Moere et ses pairs définissent ainsi trois niveaux de « style », dans lesquels l'*embellishment* est graduel. Après avoir répertorié et défini trois styles de visualisation (Figure 7), les chercheurs ont créé leurs propres prototypes, qu'ils ont testé en ligne. Ils ont ainsi différencié les visualisations analytiques, les visualisations de magazine (fréquemment trouvées dans les médias) et les visualisations artistiques (Vande Moere et al., 2012). Dans cet *embellishment* graduel, on retrouve des effets de style relatifs aux données ainsi que des images. Il faut cependant noter que les images sont moins explicites que dans les autres études. L'effort esthétique est plus abstrait que pictural.

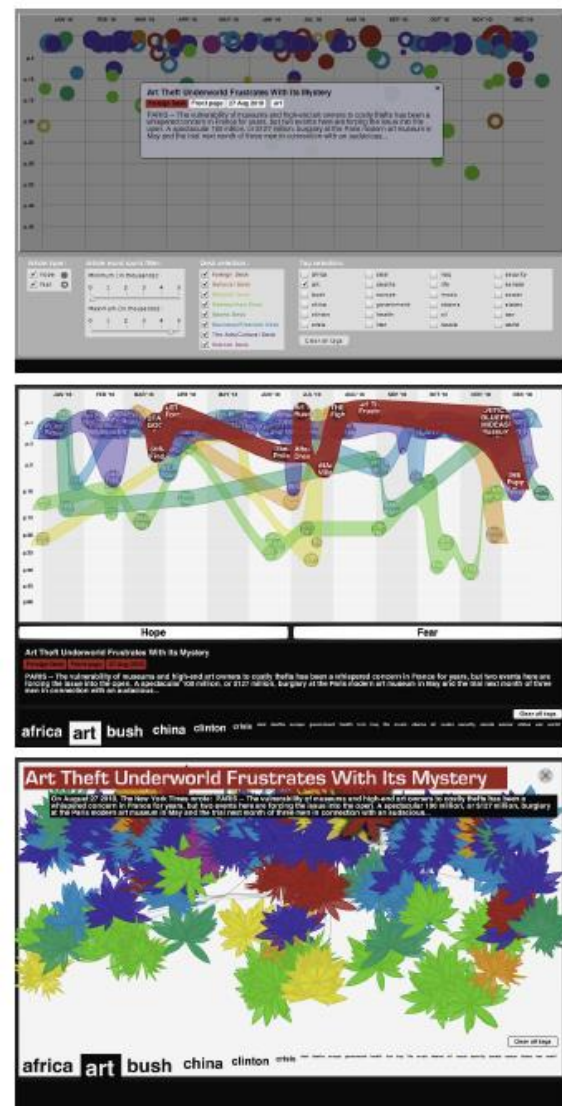


Figure 7 Styles de visualisations repérés dans Vande Moere et al. (2012). Evaluating the effect of style in information visualization. IEEE transactions on visualization and computer graphics, 18(12), 2739–2748.

Par ailleurs, nous avons encore retrouvé la notion d'une mémorabilité accrue par l'*embellishment* dans d'autres études. Les implications de ces recherches sont extrêmement

riches, non seulement du point de vue de la méthode mais également de la contribution théorique. Bien que les résultats obtenus par les chercheurs soient moins précis que prévu, ils ont pu montrer que le style avait peu d'impact sur l'utilisabilité de la visualisation. Par contre, les visualisations de type analytique (première image de la Figure 7) sont celles qui semblent les plus claires, faciles à comprendre, informatives, agréables et fonctionnelles en comparaison à leurs homologues embellies. Ainsi, les auteurs préconisent de n'embellir que lorsque la convivialité est importante dans l'intention de communication (Vande Moere et al., 2012). Les styles n'influencent pas non plus la profondeur de l'analyse, même si les participants jugeaient les visualisations ornementées plus superficielles.

Poursuivons ensuite avec Borkin et ses collègues qui se demandent en 2013 ce qui rend une visualisation de données plus facilement mémorisable. Comme point de départ, ils choisissent le fameux « *chartjunk debate* » initié par Bateman et ses pairs (2010) et critiqué par Few (2011) et citent d'autres études sur les *visual embellishments* des visualisations de données. Les résultats de l'étude de Bateman montraient que la présence de *chartjunk* renforce la mémoire des utilisateurs : il est plus facile pour eux de se souvenir d'une visualisation imagée. Comme nous l'avons déjà souligné auparavant, la définition de l'efficacité des visualisations de données peut grandement être élargie et diversifiée. C'est le cas pour Borkin et ses collègues, qui y incluent la mémorabilité des visualisations de données. Sans surprise, ils incluent ainsi la présence de l'image dans l'efficacité des visualisations de données – pour autant que l'objectif poursuivi par la visualisation ne soit correctement défini (à notre sens, explorer n'est pas communiquer). Ainsi, ces chercheurs ont récolté 410 visualisations de données dans différentes sources et les ont caractérisées selon une taxonomie qu'ils ont réalisée. Ils admettent qu'un tas de typologies peuvent être créées en la matière. Ils ont ainsi classé ces visualisations statiques selon la structure des données, l'encodage visuel et les tâches perceptuelles rendues possibles par ces encodages. Toutes ces visualisations ont été annotées en fonction de leurs propriétés (présence de couleurs, de pictogrammes, etc.). Sur les 410 visualisations de données sélectionnées, 145 sont des exemples extrêmes de minimalisme (bon *data-ink ratio*), 103 sont des exemples de *chartjunk* et 162 sont entre ces deux spectres, avec un *data-ink ratio* moyen. L'expérimentation a été proposée sous forme de jeu à un échantillon d'employés. Chaque travailleur a vu chaque visualisation deux fois au plus, moins s'il réussissait facilement le jeu. Les chercheurs ont ensuite réalisé un indice de mémorabilité afin de pouvoir réaliser des comparaisons. Les auteurs précisent bel et bien qu'ils ont là testé la reconnaissance, la mémorabilité de ces visualisations et non pas la compréhension de ces dernières. Ainsi, s'ils ne l'ont pas appelé en ces termes, c'est de l'influence de l'*embellishment* sur la mémorabilité des visualisations de données dont il s'agit ici (Borkin et al., 2013).

D'ailleurs, Borkin et ses collègues ont pu montrer que des caractéristiques telles que la couleur et l'inclusion d'un objet reconnaissable améliorent la mémorisation. Ils ont également constaté que les visualisations avec un *data-ink ratio* et une densité visuelle assez haute (donc avec beaucoup de « bruit ») sont plus mémorables que les visualisations minimalistes. Ils expliquent aussi que les types de visualisations qu'ils appellent « uniques » (pictogrammes, arbres, réseaux, diagrammes spéciaux, etc.), recueillent des scores de mémorisation plus élevés que les graphiques dits « habituels » (cercles, barres, lignes, ...).

En 2016, Borkin et ses collègues étendent leur travail sur la mémorabilité en menant une étude *eye tracking*. Cet article va plus loin en montrant à quel point les visualisations de données sont reconnues et remémorées en détail lorsqu'elles sont ornementées. Sur base du même corpus que dans leur précédente étude, Borkin et ses collègues ont sélectionné 393 visualisations de données qu'ils ont montrées à 33 citoyens de Boston dans le cadre d'une étude utilisant l'*eye tracking*. Leur travail ne se concentre pas sur des tâches spécifiques, mais davantage sur la localisation des points de fixation détectés par l'*eye tracker* et leur durée. Chaque visualisation a été étiquetée en fonction de caractéristiques et son niveau de redondance du message ou des données a été évalué également. Comme précédemment, il y avait dans ce corpus des visualisations minimalistes ainsi que du *chartjunk*, illustré en images réelles ou en pictogrammes. Par exemple, pour Borkin et ses pairs, un objet reconnaissable sur une visualisation serait un élément de redondance. Comme le montre la Figure 8, directement issue de leur article, les participants ont observé les visualisations durant dix secondes dans un premier temps. Il s'agissait de la phase d'encodage. Les mesures collectées durant cette phase étaient les points de fixation¹⁰ (x,y) et leur durée. Lors de la phase de reconnaissance, 10 minutes plus tard, les participants ont visionné les visualisations. Ils disposaient de 2 secondes par visualisations, ce qui permettait ainsi de trouver grâce à l'*eye tracker* quels étaient les points de « refixation », à

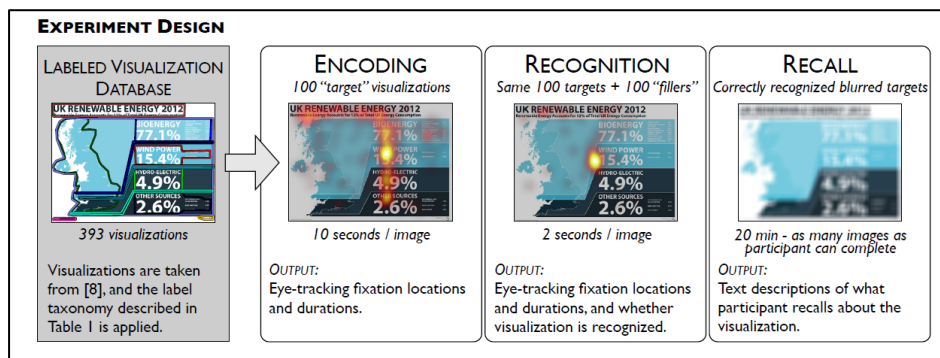


Figure 8 Design expérimental exploité dans Borkin, M. A. et al (2016). *Beyond Memorability: Visualization Recognition and Recall*. IEEE Transactions on Visualization and Computer Graphics, 22(1), 519-528. <https://doi.org/10.1109/TVCG.2015.2467732>

¹⁰ Un point de fixation est l'endroit où s'arrête le regard, généralement compris entre 50 et 200ms.

savoir quels éléments étaient directement observés afin de reconnaître la visualisation de données. Les mêmes mesures ont été prélevées. Enfin, dans la dernière phase, les visualisations, 20 minutes plus tard, étaient suggérées en flou aux participants qui devaient ainsi décrire tout ce dont ils se souvenaient. Les mesures les plus importantes et calculées dans le cadre de cette recherche sont les refixations (combien de fois l'œil revient sur un même élément), l'indice de reconnaissance et la qualité de la description.

En termes de résultats, cette étude a montré que les visualisations qui sont les plus mémorables "en un coup d'œil" sont celles qui reviennent rapidement en mémoire, c'est-à-dire qui nécessitent le moins de mouvements oculaires pour reconnaître la visualisation. Les titres et le texte sont des éléments clés dans une visualisation : les gens passent le plus de temps à regarder le texte dans une visualisation, et plus précisément le titre. L'agencement du texte de manière générale est très important. Les pictogrammes engendrent également une meilleure reconnaissance. La redondance aide également au rappel et à la compréhension. « Les participants se sont souvenus et ont oublié les mêmes visualisations, qu'on leur ait donné 1 ou 10 secondes pour l'encodage, et surtout, les participants se sont mieux souvenus du message principal/contenu des visualisations qui étaient plus mémorables "d'un seul coup d'œil". Ainsi, pour que les gens comprennent une visualisation, "nous devons d'abord les amener à y prêter attention" - la mémorisation est un moyen d'y parvenir. Ensuite, en faisant attention, on peut concevoir et introduire si nécessaire du texte, des pictogrammes et des redondances dans la visualisation afin de faire passer les messages » (Borkin et al., 2016, p. 527, notre traduction).

L'*embellishment* n'est pas défini dans ces deux études mais il se constitue de pictogrammes et d'éléments visuels reconnaissables par l'humain, tels des images. Le *data-ink* est parfois ornementé. Ainsi, on voit les visuels dont parlent Borkin et ses collègues sur la Figure 9. Ils les ordonnent ainsi : « Les dix visualisations les plus mémorables pour chacune des quatre catégories de sources de visualisation : infographiques (en haut à gauche), publications scientifiques (en haut à droite), médias (en bas à gauche) et gouvernement /organisation mondiale (en bas à droite). Dans chaque quadrant, les visualisations sont classées dans l'ordre le plus faiblement mémorisable, du haut à gauche au bas à droite » (Borkin et al., 2013, p. 2314, notre traduction).



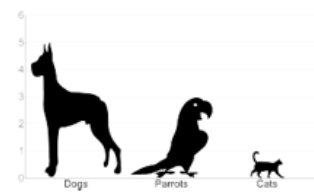
Figure 9 Ornementations étudiées par Borkin et ses collègues : "Top 10 des visualisations les plus mémorables par catégorie", vu dans Borkin, M. A. et al., (2013). What Makes a Visualization Memorable? IEEE Transactions on Visualization and Computer Graphics, 19(12), p. 2314.

Par ailleurs, une autre étude parle également d'ornementation en testant cette fois l'effet des pictogrammes placés dans les visualisations de données sur la mémoire à court terme, la performance et l'engagement de l'individu envers la visualisation. Haroz et ses collègues prétendent ainsi que le pictogramme peut rendre la visualisation de données plus efficace parce qu'il impacte la mémoire, la vitesse de recherche de l'information, l'engagement et les préférences des utilisateurs, tout en admettant que les images superflues sont un élément de distraction. Le pictogramme intelligemment intégré à la visualisation peut donc renforcer ses capacités de transmission du message (Haroz et al., 2015). Selon nous, l'utilisation de pictogrammes relève clairement d'une pratique *d'embellishment*. Ainsi, dans leur étude, ils évaluent les coûts et bénéfices de l'utilisation des « pictogrammes » pour représenter les données (Haroz et al., 2015). Selon eux, la communauté du design d'information s'est fortement inspirée du type « Isotype », ensemble de symboles graphiques extrêmement simples, normalisés et abstraits pour représenter des données socio-scientifiques (Bresnahan, 2011). C'est l'influence de ce style d'ornementation qu'ils observent. En général, le pictogramme fait directement référence à ce à quoi l'on parle. Il s'agit donc d'objets reconnaissables comme en parlent également Borkin et ses pairs (2013, 2016). Haroz et ses pairs ont qualifié différents types d'*embellishments*. Ils se demandent effectivement quels aspects et quelle disposition du pictogramme revêtent le plus d'impact sur l'utilisateur. Haroz et ses collègues les qualifient ainsi : le « stretched pictograph », le « stacked pictograph », le « background » et le « superfluous » (Tableau 3).

Tableau 3 Classification des pictogrammes dans (Haroz et al., 2015)

Type de visuel

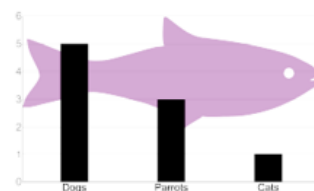
Haroz et al.,
2015



Stretched



Stacked



Background

Pictogramme présent mais complètement
séparé des données Superfluous

Il y a quatre grandes conclusions à tirer de cette étude : « (1) Seuls les pictogrammes intégrés au mapping des données sont bénéfiques (ou du moins ne sont pas préjudiciables). Les pictogrammes et images d'étiquettes superflus sont distrayants et déroutants. (2) La discrétisation d'une barre en une collection de petits éléments (stacked pictograph) améliore l'encodage et le rappel, mais uniquement pour les petites valeurs. (3) Les pictogrammes peuvent aider les gens à mémoriser des informations lors de tâches exigeantes. (4) Les pictogrammes incitent les gens à observer de plus près une visualisation. Nous n'avons trouvé aucune preuve solide que l'utilisation de pictogrammes pour la communication de données nuise à la performance dans nos tâches » (Haroz et al, 2015, p. 9).

Pour terminer, nous constatons ainsi que malgré notre effort de définition, le concept d'ornementation reste flou dans les études dans lesquelles il est abordé, tant les types d'ornementation existants ne sont pas encore distingués de façon claire et uniformisée.

2. Retour sur les termes de base : entre esthétique, beauté et design

Dans les articles que nous avons comparés, nous avons été confrontée à des termes qui se rapprochent et se recourent. Une esthétique « meilleure » serait-elle visée par ces ornements ? Ou veut-on juste que ce soit plus beau ? Ou encore, agit-on par simple souci de « design », parce que c'est dans l'air du temps ? Tous ces mots, souvent utilisés dans de nombreuses disciplines, semblent communs aujourd'hui et pourtant, ils évoquent beaucoup d'éléments déroutants. Un retour aux sources est nécessaire.

Nous pouvons évoquer David McCandless, auteur, designer et data journaliste, dont les livres « *Datavision* » (2011) et les workshops « *Information is beautiful* » sont un succès avéré auprès du grand public. Les workshops portent dans leur titre un mot évocateur : « *beautiful* ». Les visualisations de données peuvent selon lui être belles et utiles à la fois. Il voit dans leur beauté une grande force. C'est un parti pris, à en voir le discours qu'il tient :

« Il semble que nous ayons beaucoup de problèmes d'information dans notre société en ce moment, (...) [comme] le manque de transparence ou simplement l'intérêt. Les informations sont trop intéressantes. Elles ont une qualité magnétique qui m'attire. (...) Et même quand les informations sont terribles, le visuel peut être beau »

(McCandless, en ligne, TED Global 2010).

Finalement, où est l'essentiel ? McCandless, semble avoir choisi : l'information doit être belle. A vrai dire, la beauté est quelque chose de subjectif, débattu en philosophie depuis la nuit des temps. Rappelons-le : Bertin et Tufte, à travers leurs principes, visent l'objectivité des données – bien que nous le questionnions - et ils veulent refléter ce qu'il y a de plus vrai en elles. Ils ne sont évidemment pas contre la beauté mais mettent en garde contre le geste artistique. C'est là que se situe toute la complexité de la beauté en visualisation de données. Dans *The Visual Display of Quantitative Information*, Tufte (1983) titre l'un de ses chapitres « *Aesthetics and Technique in Data Graphical Design* ». Pourtant, l'esthétique est inextricablement liée au beau, à la beauté, tout comme le beau est lié à l'art. Alors, une visualisation de données se devrait-elle d'être élégante comme le dit Tufte, mais pas artistique ? Jamais Tufte ou Bertin ne diraient qu'il est interdit qu'une visualisation soit belle tant qu'elle respecte la véracité des données. Pourtant, il est sûr et certain qu'ils refusent de considérer des visualisations de données comme artistiques. Pour Tufte, si la visualisation est minimaliste et épurée, la beauté apparaîtra d'elle-même. Mais qu'est-ce que l'esthétique et la beauté ?

2.1. De l'esthétique

Pour commencer, l'esthétique est un terme extrêmement mobilisé en visualisation d'information. Si l'on remonte à son étymologie αἰσθητικός (*aisthêtikós*), le sens premier de l'esthétique remonte au sensible. Nous le verrons plus tard, la sensation précède la perception. L'objectif de la pensée esthétique est de se questionner sur la beauté qui se résume elle-même en une chose et non en une idée (Dufrenne, 2017). L'esthétique est descriptive ; elle est philosophie et non science. De plus, l'expérience esthétique est elle-même toute relative à son objet qui lui-même se révèle dans la prise de conscience du sensible par le sujet. En d'autres mots, l'objet esthétique peut se résumer en une transmission d'un émetteur à un récepteur, cette œuvre étant adressée à la perception de l'homme qui le reçoit. L'esthétique peut dès lors être considérée selon deux angles qui sont complémentaires : subjectiviste et objectiviste (Dufrenne, 2017). Dans l'esthétique subjectiviste, le sujet interprétant est tout à fait privilégié. Ainsi, on ne peut nier que la perception d'un objet peut dépendre de la culture, de la vie, du contexte de celui qui l'explore, tout autant que de celui qui la crée. Dans l'esthétique objectiviste, on prend attention aux « conditions psychophysiologiques de l'expérience esthétique plutôt que cette expérience elle-même » (Dufrenne, 2017). Elle se rapproche d'ailleurs du formalisme, qui consiste à se demander ce qu'est réellement l'être esthétique que l'on observe – ce qu'est son être, et non sa valeur. A contrario, les conditions psychosociologiques de la création y sont laissées de côté (Dufrenne, 2017). Au-delà de ces considérations philosophiques à propos de l'esthétique, il nous semble intéressant de la voir rapportée dans un cas concret, tel que l'ont fait Zen et Vanderdonkt dans « *Towards an evaluation of graphical user interfaces aesthetics based on metrics* » (2014), dans une approche objectiviste d'ailleurs. Ainsi, les deux chercheurs se sont demandé s'il était possible de rendre l'esthétique quantifiable au point de l'évaluer grâce à des métriques dans une interface numérique en particulier – un site web par exemple. Si à première vue cela peut sembler très calculé, cela s'explique sans doute par le « standard de beauté » visé par l'esthétique de manière générale. C'est en tout cas ce qui transparaît des citations d'autres auteurs qu'ils citent :

« Il existe un concept d'esthétique universelle considéré comme l'existence d'une norme de beauté partagée entre les cultures, les races et les sexes. S'il ne peut être prouvé pour toutes les caractéristiques des formes, certains aspects tels que la symétrie, la douceur, la brillance et l'équilibre peuvent être considérés comme des principes esthétiques répandus. Les formes qui manifestent ces principes peuvent être appréciées à la fois inter-culturellement et trans-culturellement ». (Van Damme, 1996, cité par Zen et Vanderdonkt, 2014, p.2, notre traduction).

« Les perceptions de l'esthétique par l'utilisateur englobent deux dimensions principales, "l'esthétique classique" et "l'esthétique expressive". Alors que "l'esthétique classique" se réfère aux principes standards de beauté tels que la séquence, l'ordre et la symétrie, "l'esthétique expressive" représente des qualités qui vont au-delà des principes classiques et qui soulignent la créativité et la puissance expressive du designer. "Originalité", "design fascinant" et "effets spéciaux" sont des aspects du design non standard qui font partie de cette catégorie » (Lavie & Tractinsky, 2004, cité par Zen et Vanderdonckt, 2014, p.2, notre traduction).

Il est ainsi possible de faire le lien avec des chercheurs comme Kennedy et Hill (2017, 2018), qui s'intéressent aux individus et à leur contexte environnant lors de la perception d'une visualisation, ce que l'on peut rapprocher de l'esthétique subjectiviste, tandis que Tufte et Bertin sont plutôt dans une perspective objectiviste de l'esthétique. Cette complémentarité trouve son écho en visualisation de données, tout comme le fait que définir une certaine esthétique revienne à rechercher un standard de beauté. Ne se rapporterait-il pas, chez nos deux grands auteurs de référence, au minimalisme qu'ils préconisent tant ? Le style épuré de Tufte et Bertin est un standard de beauté qu'ils préconisent pour la visualisation de données, parce que ce standard de beauté en particulier exerce des effets que les deux auteurs recherchent, ces effets correspondant à leur vision de l'efficacité graphique. L'élégance graphique est une esthétique tout indiquée par Tufte (1983). Dans les faits, c'est sans aucun doute une autre esthétique des données qui se propage au quotidien actuellement.

2.2. De la beauté, liée à l'art dans l'histoire

Avant de parler de beauté, on parlait du « beau » et les réflexions les plus lointaines que l'on connaît à ce sujet remontent à la Grèce Antique. Il est bon de se rappeler qu'avant toute chose, le beau était inextricablement lié au bien. Il était utile à la production du bien et lié au plaisir, il constituait d'ailleurs un moyen d'ascension vers le bien, si l'on se remémore les pensées de Platon et Socrate. Entre 205 et 270, Plotin revient sur l'alliage du plaisir au bien et essaie de caractériser les émotions suscitées par le beau, autrement dit le plaisir esthétique. Les anciens philosophes grecs voyaient le beau dans ce qu'il y avait de plus naturel et voyaient dans la divinité l'accord parfait entre le beau et le bien (Michaud, 2017). Tout de même, cela laisse à réfléchir, une fois transposé à nos thématiques contemporaines. Une belle visualisation donne-t-elle l'illusion d'être de l'ordre du « bien » ? C'est l'apparition du terme esthétique qui opère une séparation entre beau et bien et qui fait plutôt parler de la « beauté ». Par la suite selon Locke, la beauté provoque en

l'homme des idées qui sont affectives, sensibles, qui ne sont d'ailleurs pas la réalité (l'homme produit par ailleurs des idées de « qualité première »). Cela générerait une connaissance confuse, une connaissance autre que la « réalité-même », qui met l'homme dans un état esthétique et sensible (Michaud, 2017). On pourrait considérer cette connaissance confuse comme l'essence-même de l'art, et c'est du point de vue de la visualisation d'information qu'il est possible de le comprendre. Effectivement, Kosara (2007) distingue clairement les visualisations de données artistiques des visualisations pragmatiques. D'une telle manière, une visualisation artistique n'aura pas les mêmes buts et ne pourra être utilisée de la même façon qu'une visualisation pragmatique, même si le flou gravite de temps à autres autour du sujet. Ainsi, Kosara parle du sublime, cette sensation esthétique qui peut être comprise comme ce qui provoque l'émotion, la grandeur ou encore une réponse intellectuelle et émotionnelle profonde (Kosara, 2007, p.4). Le sublime est un critère en esthétique, parmi d'autres comme l'(in)harmonieux, le dissonant ou encore le laid (Michaud, 2017). En *human centered computing*, il s'oppose à l'anti-sublime qui est compris comme user friendly, facile à saisir, qui vise la compréhension immédiate (Sack W., 2004). Ainsi, ce qui différencie les visualisations artistiques des visualisations pragmatiques est cette visée esthétique : les visualisations artistiques relèvent avant tout du sublime et, même si elles sont conçues sur base d'un grand jeu de données et selon des procédés numériques et hautement technologiques, elles ne peuvent prétendre à la même finalité qu'une visualisation de données pragmatique : analyser et comprendre les données avec le moins d'efforts possibles (Kosara, 2007). C'est donc l'anti-sublime qui anime les visualisations de données pragmatiques. Bertin et Tufte conseillent donc quant à la production de visualisations de données pragmatiques et se situent dans l'ordre de l'anti-sublime. C'est pourquoi on constate chez eux un certain rejet quant à l'effort artistique : si l'on suit Locke, il apporte un ajout de signification qui pourrait perturber le message initial que les auteurs veulent transmettre pragmatiquement. Puisque le beau provoque une émotion, une sensation esthétique, il serait pertinent de prendre en compte le ressenti de l'utilisateur quant à la lecture de la visualisation. Cette pensée toute légitime est de l'ordre du design.

2.3. Du design

Les précédentes interrogations quant au beau et à l'esthétique amènent à se questionner davantage sur le design d'information. Cela commence simplement avec le design. Nous éclaircissons ainsi le terme dans cette section et proposons ensuite de nous pencher en particulier sur le design émotionnel (Norman, 2013).

2.3.1. La conception esthétique

Aujourd'hui, le design est fortement associé à l'allure esthétique, belle, « friendly » qu'il évoque alors que dans les faits, il dépasse largement cette idée. Provenant de la langue anglaise, il renvoie à l'idée de conception et de fonctionnalité, mais pas seulement. Dans son étymologie française, « design » rapporte au dessin et au dessein (Colin, 2017). C'est l'alliage d'une image comme conception du plan d'un objet en particulier - avec une notion esthétique liée au dessin – et d'un but bien défini. Lloveria définit très bien comment le design est perçu aujourd'hui : il est doté d'un « modernisme fonctionnel sur le plan esthétique. (...) Le design associe à la valeur informationnelle initiale un réseau de valeurs esthétiques et rhétoriques (le modernisme) » (Lloveria, 2017, p.104). Dans la langue française, « le design s'inscrit dans la terminologie qui lui préexiste : art décoratif, création industrielle et métiers d'art » (Colin, 2017, en ligne).

Il est aussi intéressant de savoir que l'ornement a été considéré comme un « crime » dans la conception des objets de la maison dès le début du XXème siècle, en partie grâce à (ou à cause de) Adolf Loos. A la suite de ce « traumatisme », le design du milieu du siècle s'est efforcé de réintroduire les ornements et de les considérer à l'intérieur de la structure et de la matière-même des objets. En effet, le décoratif n'est pas à séparer complètement de la notion de fonction car « la fonction n'est pas l'usage (...) l'utilisateur « se sert ». La fonction, elle, est attachée à l'objet, c'est lui qui fonctionne » (Colin, 2017, en ligne). Quoi qu'il en soit, le design « renvoie au designer. C'est-à-dire à un mode de conception de l'objet » (Colin, 2017, en ligne). Tout comme le concepteur d'une visualisation, le designer d'un objet n'est jamais seul. Il n'est pas un artiste, qui, lui, réalise son œuvre dans la solitude dans la plupart des cas. Le designer est le professionnel des intentions qui feront en sorte qu'une forme en particulier soit ce qu'elle est. Il est celui qui a l'expérience, qui comprend ce à quoi l'objet doit répondre, le « pourquoi ». Il conçoit : il observe un processus avant d'en être acteur et de le concevoir (Colin, 2017, en ligne). Malgré tout, le design n'est pas à confondre avec le fonctionnalisme. En effet, pour le fonctionnalisme, « la beauté d'une œuvre dépend de son adaptation à sa fonction » (Colin, 2017, en ligne). Dès lors, on pourrait peut-être dire que la pensée de Bertin et Tufte est fonctionnelle, surtout quand on sait que Tufte prétend que la beauté apparaît d'elle-même dans une visualisation qui est « bien faite ».

En réalité, le design apporte quelque chose de supplémentaire ; il accepte l'ornementation au sein même de la conception, l'ornement est dans la fonction. On peut même aller plus loin en disant que « l'esthétique ne se limite pas à la forme visuelle, mais inclut également des aspects flous tels que l'originalité, l'innovation et la nouveauté ainsi que d'autres facteurs plus subjectifs prenant en compte l'expérience de l'utilisateur. L'originalité,

l'innovation et la nouveauté constituent alors, à nos yeux, la *dimension rhétorique* du design d'information » (Lloveria, 2017, p.105). Donc, l'esthétique est utile. Rappelons que le but de la graphique selon Bertin est d'atteindre une monosémie visuelle. En 1981, Bertin rappelle que la notion de préférence (« *Que préférez-vous ?* ») est en rupture totale avec l'objectif de la graphique. Pour lui, c'est une source de confusion. Pourtant, c'est bien ce qui est testé, la plupart du temps, dans les études qui traitent de l'*embellishment*. Dans *Théorie Matricielle de la Graphique* (1981), il met en place un « test », une grille d'analyse qui permet de réaliser sa visualisation de données en se posant les bonnes questions, celles auxquelles le graphique doit répondre. « (...) Ce test souligne qu'on ne regarde pas un graphique ou une carte comme on regarde une peinture ou un signal routier » (Bertin, 1981, p.63). Ainsi, pour lui, la visualisation de données s'éloigne de l'image qui est polysémique, de l'iconicité et du figuratif car ce type de visuels permet une signification propre à chacun, qui se discute, et les réponses qu'un graphique ne peuvent pas se discuter (Bertin, 1981). Or, on ne peut le nier, malgré ses préconisations et malgré la théorie existante sur la graphique ou sur la visualisation d'information, force est de constater que les designers prennent aujourd'hui compte de ce que le récepteur préfère au moment de la conception, souvent sans s'inquiéter de cette volonté de monosémie.

Il y a de nombreuses raisons de vouloir se démarquer dans le contexte communicationnel actuel. On l'a vu plus haut, le designer pense au « pourquoi » de son plan. Actuellement, les technologies, le nombre d'écrans consultés en une journée et la concurrence en ce qui concerne l'information numérique consultée influencent probablement l'attitude du designer, du graphiste, du concepteur de la visualisation finale. « *Que voyez-vous ? Que préférez-vous ?* » (Bertin, 1981, p.63) sont des questions qui ont bel et bien trouvé une place en visualisation d'information aujourd'hui... Si le but des visualisations de données actuelles est tout de même de transmettre l'information liée aux données et de répondre à des questions précises, l'ornementation est une pratique nouvelle qui met parfois à mal les principes de conception de Bertin et Tufte. Ce n'est pas parce qu'on s'attarde sur l'esthétique dans la conception, et donc sur le beau, que cette « beauté » n'a rien d'utile à l'heure actuelle. Il est à supposer que l'individu fasse sens, en prime de l'information contenue dans les données, d'autre chose.

2.3.2. Le design émotionnel

Cela nous mène au design émotionnel selon Norman (2013). Effectivement, Norman a développé sa théorie uniquement pour les objets quotidiens, « *the everyday thing* ». Pourtant, nombreux sont les auteurs qui le citent, même dans le domaine du design d'information qui ne concerne pas les objets tangibles. Pour Norman, il ne faut pas

forcément écarter le beau du design : cela ne doit pas être une conception purement fonctionnelle. Il lie en effet le beau à l'agréable et au plaisir. Ainsi, il prétend que si un objet est agréable et que l'utilisateur ressent du plaisir à l'utiliser, il l'utilisera mieux qu'un objet classique purement fonctionnel. Cet objet attractif fonctionnerait alors mieux. Il lie cela à trois niveaux de traitement cognitif du design : le niveau viscéral, comportemental et réflexif. Ce sont également trois niveaux de design. Le design viscéral est le design qui fait appel à nos réflexes primitifs. A ce niveau, on fait attention aux apparences et à ce qu'elles évoquent pour nous, en établissant clairement ce qui est bon ou mauvais, en ne se basant que sur l'apparence physique de ce qu'il y a en face de nous. Ensuite, le niveau comportemental est lié à l'efficacité d'utilisation mais aussi au plaisir : à la suite de l'utilisation, l'utilisateur va agir en fonction de ce qui fonctionne mieux. C'est finalement assez comparable à l'utilisabilité. Finalement, le niveau réflexif concerne la conscientisation, l'intellectualisation du produit. L'utilisateur va pouvoir réfléchir à propos de son objet, en se positionnant par rapport à lui et en se demandant quelle histoire il peut raconter. C'est à ce niveau que l'utilisateur choisira par exemple un objet moins pratique qu'un autre, mais qui aura plus de signification pour lui. Atteindre un niveau renforce ou inhibe le niveau précédent (Norman, 2013). C'est pourquoi les émotions sont importantes, sachant que le niveau réflexif du design est le plus haut : la simplicité n'est pas une réponse en soi (Norman, 2008). Selon les réflexions de Norman, même si une visualisation de données est moins efficace ou pratique qu'une autre, elle pourrait être correctement utilisée et bien comprise par un utilisateur qui arrive à trouver une émotion positive dans l'ornementation. Nous aurons l'occasion d'exploiter cette théorie par la suite.

3. Conclusion : comment considérer l'ornementation en visualisation de données ?

Puisqu'on ne peut parler d'ornementation sans aborder ces sujets, nous nous sommes interrogée sur l'esthétique, la beauté et le design. Nous avons ainsi compris que ces notions évoquent des éléments fondamentaux en la matière : des ponts se jettent véritablement entre elles toutes. La beauté a trait à l'art et il s'agit de quelque chose qui effraie en recherche, tout comme la catégorisation en sciences effraie les artistes. Comme le dit Lima, « la plupart des artistes sont mal à l'aise face à l'analyse scientifique. Ils pensent qu'elle détruit quelque chose qui touche à l'aspect humain de la créativité. La peur (...) d'un réductionnisme peu subtil (...) est répandue. Le domaine de l'esthétique traite, après tout, de notions émotionnelles, et ce n'est jamais un terrain propice à la recherche » (Lima, 2013, p. 222). Pourtant, de nombreuses raisons nous poussent à croire que ce domaine, finalement délaissé suite à ces réticences, a beaucoup de choses à offrir. Car les émotions induisent

chez les individus une construction de sens qui leur sera propre au moment de l'appropriation de la visualisation de données. Dans une vision subjectiviste, l'esthétique des visualisations de données engendrera des émotions – visées ou non par les concepteurs de ces dernières – qui feront partie des éléments intégrés par les individus au moment de leur construction de sens. De plus, plusieurs auteurs ont questionné la relation entre visualisation de données et art : « *Bridging the gap between art & science* » (Hudelman, 2004), « *The missing link between information visualisation & art* » (Kosara, 2007) sont des exemples d'articles sur le sujet, tandis que d'autres auteurs comme Quispel & Maes, (2014) et Inbar et al. (2007) ne prennent pas à la légère le design émotionnel dans leurs études (Norman, 2013). Il fait effectivement partie intégrante de leur revue de littérature et guide leurs réflexions. Les premiers étudient la perception de la « bonne » visualisation et les deuxièmes étudient les effets issus d'une maximisation du *data-ink ratio*. Ainsi, comme nous l'avons précisé précédemment, Bertin et Tufte préconisent une esthétique épurée et minimaliste en élaborant une suite de principes et de règles. Nous estimons alors que ce qu'ils préconisent est une esthétique à part entière.

Comme nous l'avons précédemment annoncé, nous pensons, tout comme Inbar et ses collègues (2007), que Tufte préconise une esthétique minimaliste pour atteindre la visualisation idéale¹¹. Effectivement, peu de personnes ont l'habitude de consommer ce genre de visuels, bien qu'il soit conseillé comme le plus « efficace » d'un point de vue perceptuel. Ce style minimaliste signifie ainsi un respect total des règles de conception et une absence complète de *chartjunk*, ainsi qu'une maximisation totale du *data-ink ratio* selon Tufte (1983). En contrepied de cela, le *chartjunk* rassemble beaucoup d'éléments. Ainsi, puisque Tufte estime que tout élément non utile et pouvant potentiellement apporter de la distorsion est du *chartjunk*, il s'avère que les graphiques les plus communément réalisés (entreprises, parcours scolaires, ...) et dont nous avons le plus l'habitude seraient du *chartjunk* également, bien que d'un niveau moins important que dans une intention « décorative ». Prenons *Excel* pour exemple qui propose grilles, étiquettes, 3D, couleurs, etc., qui, mal ou trop utilisés, sont du *chartjunk*, et plus précisément de type « *dreaded grid* » ou quadrillage obsolète. Au quotidien, nous ajoutons des éléments obsolètes et effectivement inutiles à la « récolte » de l'information. Pourtant, ils font partie de nos habitudes. C'est ce qu'étudient (Inbar et al., 2007) et Hill et ses collègues (2017) en comparant la différence de perception entre des graphiques « standards » au sens de *chartjunk* « *dreaded grid* » et le style minimaliste préconisé par Tufte. Le résultat de leurs recherches montre que les étudiants participants à l'étude préfèrent et pensent mieux

¹¹ Du moins, de son point de vue objectiviste voulant des visualisations qui transmettent l'information de façon optimale.

comprendre le *chartjunk dreaded grid* que le style minimaliste. La raison n'a pas été testée mais ces chercheurs assimilent cela à un manque d'habitude de consommation des visualisations minimalistes. Quoi qu'il en soit, ces visualisations ne sont pas ornementées d'un point de vue pictural. Nous estimons que l'ornementation est un autre type de *chartjunk*, mais qu'elle ne peut se résumer au troisième type de *chartjunk* selon Tufte, à savoir le « canard boîteux ». Nous insistons ainsi sur l'intention décorative de l'ornementation des visualisations de données et sur l'ajout d'éléments figuratifs de type simple à complexe, disposés à différents endroits du visuel sans forcément provoquer une distorsion du *data-ink*. En somme, en utilisant le mot « ornementation », nous pouvons nous permettre plus de souplesse en considérant toutes les intentions décoratives (formes, couleurs, pictogrammes, dessins, ...) et en nous distançant toutefois de cette vision négative des éléments décoratifs. Si pour Tufte, les ornements font preuve d'un design pensé uniquement en fonction de l'apparence et non de la pertinence et profondeur des données, nous nous permettons de lui imaginer d'autres raisons d'être. La Figure 10 résume notre point de vue.

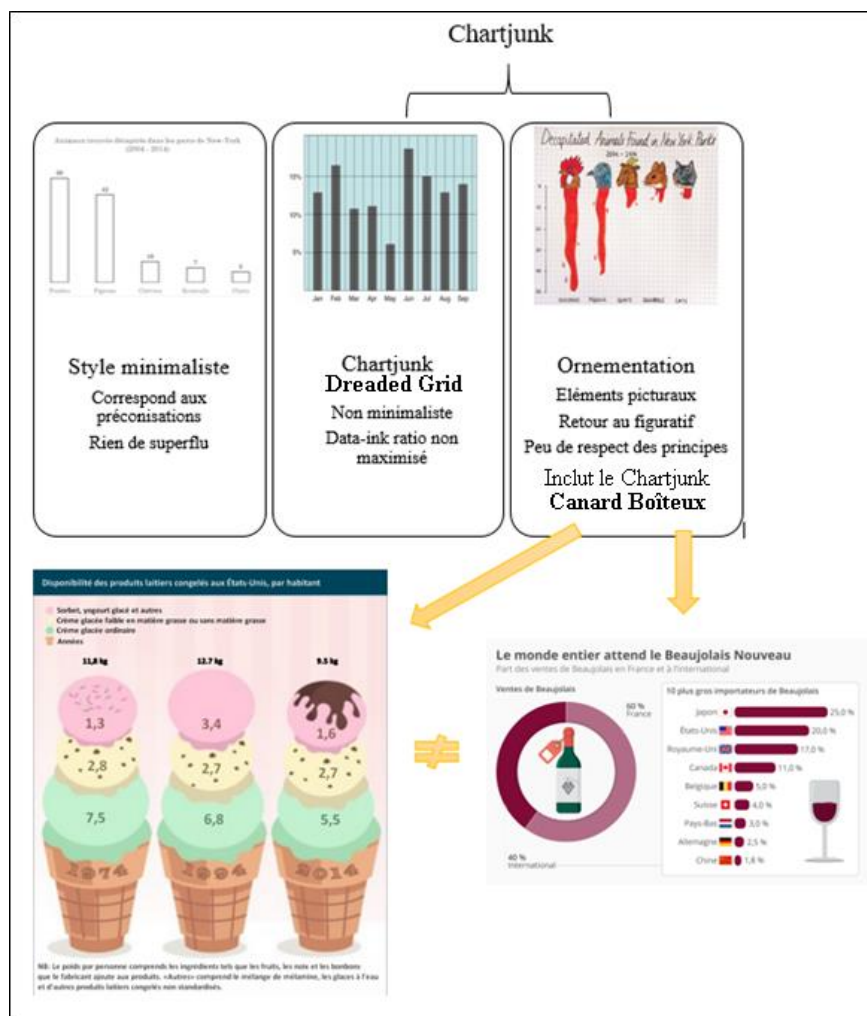


Figure 10 Différencier chartjunk et ornementation, par l'auteure

Par ailleurs, nous avons nous-même identifié dans les articles recensés différentes façons d’ornementer le visuel. Ce recensement n’est évidemment pas exhaustif. Bateman et ses collègues (2010) utilisent de manière similaire les termes *embellishment* et *chartjunk*. L’ornementation se concentrerait donc sur le *data-ink* selon eux. Pour Skau et ses pairs (2015), l’ornementation consiste en la transformation des traits relatifs aux données (*data-ink*) à des fins d’enjolivement. Borgo et ses co-auteurs étudient l’ornementation comme l’utilisation d’images, de dessins et de pictogrammes apposés à un graphique traditionnel ou superposés au *data-ink* sans pour autant en changer la forme. Haroz et ses pairs, quant à eux, étudient les différentes façons d’utiliser le pictogramme (2015). Finalement, dans leurs études de 2013 et 2016, Borkin et son équipe de chercheurs étudient davantage l’ornementation telle que nous l’entendons : ils y incluent des formes reconnaissables par l’œil humain, métaphores, pictogrammes, images communes, dessins, appliqués sur le *data-ink* ou non. Ces différentes considérations de l’*embellishment*, relevées dans toutes ces études, nous seront forcément utiles. Dans le [chapitre 1.2. de la Partie III](#) de cette monographie, nous abordons les attributs nous ayant permis de caractériser les visualisations de données ayant fait l’objet d’expérimentations dans le cadre de cette thèse.

Finalement, entre les notions de sublime, d’artistique, de pragmatique, de fonctionnel, d’émotionnel et bien d’autres qu’ont souligné les réflexions émises lors de ce chapitre, il nous semble possible d’imaginer qu’il puisse exister un type d’ornementation alliant à la fois les principes basiques de conception des visualisations de données et les défis, usages, besoins et envie de notre société contemporaine (Figure 11). Si le *chartjunk* de type « canard boiteux » tel que l’illustre Tufte est effectivement une mauvaise façon de communiquer de manière précise et intégrée, les études de Bateman et ses collègues (2010), de Borkin et ses pairs (2013, 2016) et des autres auteurs cités dans ce chapitre montrent bel et bien que l’ornementation peut apporter des bénéfices, notamment du point de vue de la mémorisation ou de l’intérêt (préférences des individus). Dès lors, nous supposons que ces bénéfices, tout comme l’ornementation elle-même, jouent leur rôle dans la construction de sens.

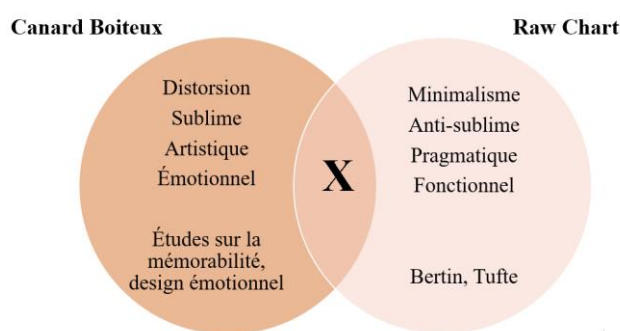


Figure 11 Entre le canard et le raw chart

Chapitre 3 Aborder la construction de sens en visualisation de données

Il y a différentes façons d'étudier le sens en visualisation de données, parce que cette discipline stimule non seulement la perception visuelle mais aussi la cognition, tout en faisant appel à certaines caractéristiques personnelles du lecteur. Ce sont donc les différentes composantes de la construction de sens qui sont abordées dans ce chapitre.

1. Les différentes composantes de la construction de sens

De nombreuses études et concepts existent concernant la construction de sens ou le *sense-making* et ce, dans bien des disciplines. Selon nous, il serait difficile d'affirmer que dans le cadre de notre recherche, une seule vision de la construction de sens est applicable. En nous interrogeant sur les différents éléments pouvant entrer en compte dans la construction de sens de l'utilisateur qui lit une visualisation de données et en gardant en tête que notre étude se base principalement sur l'influence de l'ornementation, nous pouvons mettre en exergue ce qui constitue la construction de sens pour les utilisateurs qui sont au contact des visualisations de données. Effectivement, lorsqu'un individu lit une visualisation de données et crée du sens, il passe à travers différentes étapes transversales (Figure 12) qui sont, selon nous, des « composantes » du *sense-making* :

- (1) La perception visuelle : avant tout, l'individu voit la visualisation de données et prend connaissance de l'information grâce à son système perceptuel.
- (2) La compréhension : la construction de sens est largement étudiée en sciences cognitives, où plusieurs chercheurs prétendent qu'une visualisation de données est une représentation externe au système de pensées de l'individu, lui permettant ainsi de réaliser une « économie cognitive » (Card, 2012). Différents processus cognitifs s'enclenchent alors afin que l'individu atteigne une certaine compréhension de la visualisation de données.
- (3) Le *gap bridging*, qui est un élément de *sense-making* individuel conceptualisé par Dervin (1992, 2003). Il est influencé par
 - a. Les caractéristiques personnelles et émotionnelles, propres à l'individu
 - b. Les caractéristiques propres à la visualisation de données.(a) et (b) constituent ainsi les facteurs humains et émotionnels d'engagement envers les visualisations de données, mis au jour par Kennedy et Hill (2017, 2018).

Il y a de nombreuses manières d'étudier la construction de sens. L'identification de ces composantes découle de notre propre réflexion. La visualisation de données demande à l'individu de voir, regarder, décoder des informations transcrites de manière arbitraire (ou pas, dans le cas de l'ornementation) et de réagir en fonction de sa personnalité, de son environnement, du contexte de transmission de l'information, etc. C'est la raison pour laquelle l'association de ces trois composantes transversales nous semble être une tentative d'étudier le *sense-making* de façon complète. Cette construction théorique nous permet ainsi d'associer différentes théories et concepts afin de tendre vers une étude consistante et complexe du *sense-making* engendré par la visualisation de données.

Notre état de l'art se poursuivra ainsi sur ces différents éléments avant de nous permettre, dans le chapitre suivant, de problématiser concrètement l'objet de recherche qui nous occupe.

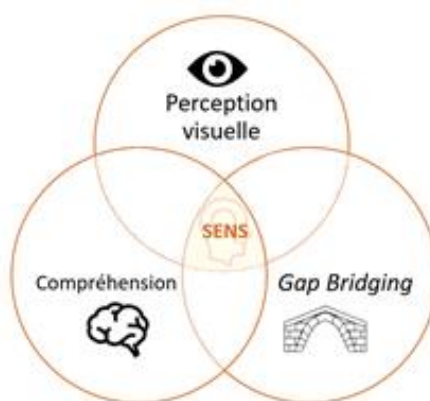


Figure 12 Les composantes du sense-making en visualisation de données, par l'auteure

2. La perception visuelle de l'image

Pour commencer, il nous semble légitime de considérer les visualisations de données statiques comme des images, bien qu'abstraites. Bertin définit le graphique comme l'utilisation d'images transformables (1967). Rappelons néanmoins que l'objectif selon lui est d'utiliser un système de signes arbitraire et monosémique, laissant peu de place à l'interprétation personnelle du graphique, communiquant ainsi un message avec le plus d'objectivité possible¹² (Bertin, 1967). Il n'en reste pas moins que ces supports visuels composés de signes arbitraires (et souvent ornementés) sont des images. La visualisation de données en tant qu'action est reconnue dans la littérature pour produire des images. Qu'elle soit statique ou non, c'est bien cela que construisent les algorithmes dans les programmes de visualisation. « La sortie d'un algorithme de visualisation est exprimée

¹² Dans l'idée bien-sûr qu'un lecteur « modèle » lise la représentation graphique

sous la forme d'une image. Cette image est construite sur tous les ordinateurs au moyen d'un langage graphique » (Baudel, 2002). Sicard parle également de l'image « de science » : « les graphes, graphiques, schémas, tableaux, diagrammes, équations, sont des images, dès lors qu'ils donnent lieu à des figurations mentales. Il serait illusoire de penser que toutes les images de science sont des images immédiates d'un référent réel » (Sicard, 1997, p.47).

Par ailleurs, l'image constitue un objet que nous percevons. Elle fait partie du monde visuel qui nous entoure. Pendant longtemps, Gibson, figure incontournable dans le domaine de la perception visuelle, s'est questionné. « Voir l'environnement, c'est extraire l'information d'un ensemble ambiant de lumière. Que se passe-t-il, alors, quand on voit une image de quelque chose ? » (Gibson, 1978, p.227, notre traduction). Effectivement, l'information dans la lumière ambiante ne consiste pas en des formes et des couleurs ; il s'agit d'invariants. De fait, Gibson se demande si cela implique également que les informations contenues dans une image soient des invariants. Pour lui, une telle supposition est étrange, bien qu'il l'émette : une image n'est-elle pas plutôt composée de formes et de couleurs ? (Gibson, 1978).

Pour Gibson, la vision d'une image découle d'un traitement que le récepteur de l'image effectue sur base de la perception (Gibson, 1978). De fait, l'image est un ensemble...

- de lumières perçues par la pupille de l'œil telles que leur brillance réfère aux différents éléments présents sur la surface ;
- D'angles solides et imbriqués, déterminés par des mesures ou des contrastes d'intensité et de la composition centrale de la lumière ambiante ;
- Considéré comme une structure immobile ;
- D'invariants persistants dans la structure, sans forme et sans nom

(Gibson, 1978, p.228).

Gibson a beaucoup travaillé sur la perception du monde visuel d'un point de vue physiologique et considère donc les images comme quelque chose que le cerveau peut décoder à travers l'œil. Lorsque nous regardons une image, notre système visuel décode les signaux que notre œil lui envoie. Nous sommes donc en possession de petites unités d'information qui sont concernées par la lumière et sa longueur d'ondes (Jeannerod, 2003). Ces unités d'informations peuvent être appelées "données rétiniennes" et c'est à elles que Bertin fait référence en parlant de variables rétiniennes (plus communément appelées variables visuelles par les autres auteurs du champ). "La perception, la représentation que nous nous faisons du monde extérieur résulte d'un processus qui assemble pour chaque

point de l'espace le résultat des traitements opérés dans chaque carte rétinotopique" (Jeannerod, 2003, p. 2). Les cartes rétinotopiques sont inscrites dans le cortex cérébral et conservent des zones de contrastes dans des localisations spatiales déjà perçues auparavant et ce, dès leur arrivée dans notre cortex cérébral visuel (Jeannerod, 2003). Quoi qu'il en soit, toutes les données rétiniennes perçues forment un premier spectre "des propriétés physiques de l'image perceptive" (Jeannerod, 2003, p.2). Néanmoins, pour Gibson, une forme perçue par la rétine est un stimulus basique, qui est aux prémices de la perception. Un dessin ou une image représente également un stimulus, mais la notion d'espace et de profondeur sont à ce moment différentes (Gibson, 1978). Comment pouvons-nous percevoir une perspective dans une image en deux dimensions ? Ce sont les différents ensembles qui constituent l'image qui nous le permettent. Dans le cas de la perspective sur image par exemple, Gibson la présente comme une perspective "artificielle", produite par les composantes de l'image. Elle n'est pas une perspective réelle mais nous en donne l'illusion. Tout comme pour d'autres effets visuels, il parlera « d'illusion de réalité » (Gibson, 1978).

Nous ne nous attarderons pas davantage sur le processus physiologique et les opérations réalisées par l'œil pour percevoir l'image, bien que notre explication soit sommaire. Effectivement, le point d'intérêt dans les travaux de la perception visuelle, du point de vue de la visualisation de données, se situe au niveau des lois de l'illusion et de la perception qui permettent de comprendre comment agencer l'image pour que d'éventuels effets optiques ne troublent pas la perception du lecteur d'une part, et pour profiter de ces effets d'optique d'autre part. Les théories diverses et les modèles préconisés dans le domaine de la perception visuelle permettent d'expliquer "les mécanismes de transmission d'information entre une visualisation et un utilisateur" (Hurter, 2010, p.16).

De plus, nous ne sommes pas la seule à penser que les réflexions de Gibson nourrissent la réflexion. Pour Ware, la théorie de la potentialité de Gibson est riche de sens : à la perception suivrait l'action. Les possibilités perceptives seraient donc les potentialités qui se manifesteraient directement, au moment de la perception. Cette théorie est souvent utilisée à des fins d'information et de décision (Ware, 2013). Cependant, les théories de Gibson ne sont plus toutes neuves et Ware reste critique en mettant certaines de ses limites en avant.

En psychologie, la théorie *Gestalt*, ou théorie de la forme, permet de comprendre, si l'on retourne à ses prémices, pourquoi la perception ne se résout pas seulement au bon fonctionnement de nos sens mais à une réaction qui consisterait à répondre à un signal. Evidemment, les sens sont indispensables au bon fonctionnement de la perception, tout

comme la perception visuelle est indissociable de la vue. La vue serait donc un sens précurrent, c'est-à-dire qui fonctionne « à distance ». De fait, cela nécessite, pour chacun, le développement « d'un espace subjectif » : « la perception étant devenue indissociable de l'action, du fait que l'extension perceptive était liée à une diversification des actes dans le temps accru de la préurrence, le faire de l'organisme a produit des objets pourvus des mêmes caractéristiques que les structures naturelles du percevoir » (Thinès, 2016, en ligne). De cette manière, selon les gestaltistes,

- Les sens ont d'abord induit des réactions à des signaux. Par apprentissage, ces signaux sont devenus, par procédés complexes, des "formes" ;
- Du fait de l'apparition de ces formes, les significations se sont étendues ;
- Le tout mène progressivement vers la représentation et la symbolisation, qui poussent à s'interroger sur les liens entre perception et cognition.

(Thinès, 2016, en ligne)

Il est primordial de comprendre que, d'une part, la perception n'est pas uniquement une activité sensorielle et physiologique : elle est guidée par l'action. D'autre part, la perception a certaines propriétés cognitives non négligeables dans le domaine des visualisations de données. Elle induit ainsi des mécanismes cognitifs importants et nécessaires dans la construction de sens.

3. Les mécanismes cognitifs : coordonner ses représentations mentales

Grâce à la vision, la schématisation aide depuis longtemps l'humain à structurer ses pensées. De par notre expérience dans le monde, chacun d'entre nous possède des représentations mentales internes. Au-delà de la perception visuelle, cela agit sur la cognition. C'est ce que l'on entendait en assurant que « la perception dépasse la sensation ». Effectivement, « Norman (1988, 1993) décrit depuis plusieurs années la cognition en termes de "connaissance dans la tête" et de "connaissance dans le monde" » (Scaife, Rogers, 1996, p. 188, notre traduction). La visualisation de données est une représentation externe à l'esprit humain, lui permettant ainsi d'amplifier sa cognition. « Dans le domaine des sciences cognitives, en général, on s'est efforcé de promouvoir la nécessité d'analyser l'interaction entre les représentations internes et externes » (Scaife, Rogers, 1996, p. 188, notre traduction). C'est dans cette interaction que la construction de sens opérée à partir des visualisations de données intervient. Ainsi, elle améliore le travail cognitif de l'utilisateur. Différentes représentations externes existent ; elles se présentent la plupart du temps sous forme linguistique ou graphique. En interagissant avec une

représentation externe comme la visualisation de données, nous nous préparons à coordonner nos représentations internes et externes (Kirsh, 2010). Nous interagissons donc avec les visualisations de données, ce qui nous permet de créer du sens supplémentaire, de réaliser des actions externes à notre système de pensée « interne » et ce, de manière rapide et efficace. « Les avantages de l'interaction avec une représentation externe sont particulièrement clairs pour les structures [de pensée] complexes. Au fur et à mesure que la complexité d'une spécification linguistique d'une structure visuelle augmente, il devient plus enrichissant de donner un sens à la phrase en construisant un dessin physique et en la regardant que de construire cette forme géométrique dans son esprit et de donner un sens à la phrase de l'intérieur » (Kirsh, 2010, p.3, notre traduction). Ces actions externes ont pour fonction d'ancrer des processus mentaux (internes) sur des processus externes. Autrement dit, le traitement sémantique et pragmatique du problème posé est approfondi (Kirsh, 2010). En bref, les visualisations de données nous permettent de réaliser une réelle économie cognitive : elles diminuent notre charge cognitive en facilitant le travail effectué pour comprendre l'information et la traiter (Zhu, Chen, 2008). La tâche cognitive est alors différente ; le volume de notre traitement cognitif est amplifié (Liu & al., 2008). Un autre élément essentiel de cette « externalisation » est qu'elle permet d'appliquer des mesures concrètes sur des termes abstraits, non visibles. Les représentations externes détiennent ainsi une fonction centrale dans les processus de cognition interne (Scaife & Rogers, 1996). Scaife & Rogers identifient ainsi trois caractéristiques fondamentales de la cognition externe. La visualisation de données permet ainsi

1. Un déchargement de calcul : utiliser des supports externes réduit l'effort cognitif pour traiter l'information.
2. La re-présentation : c'est la façon dont des représentations externes différentes, qui ont la même structure abstraite, rendent la résolution de problèmes plus facile ou plus difficile.
3. La contrainte graphique : les éléments graphiques peuvent limiter les types d'inférences à faire à propos du monde représenté.

(Scaife & Rogers, 1996)

Pour Scaife & Rogers, les diagrammes sont donc des représentations externes ayant la fonction de décrire des processus et structures à de hauts niveaux de complexité. L'inférence est ainsi indépendante de la représentation. Il y a une réduction du besoin de chercher et de reconnaître (Scaife & Rogers, 1996), ce qui coïncide avec l'économie cognitive mise en avant par Kirsh (2010). Les éléments perceptuels constituent la plus grande partie du travail à effectuer dans la résolution de problèmes géométriques. Le mécanisme œil – diagramme mène ainsi à des résultats cognitifs et perceptuels qui ont

demandé peu d'efforts (Scaife & Rogers, 1996). La cognition n'est donc pas une étape à posteriori de la perception. Toutes deux sont des composantes fondamentales de la construction de sens.

Par ailleurs, Zhu et Chen parlent de l'extension de la cognition et de l'extension de la mémoire. Effectivement, la visualisation, en plus de nous permettre de comprendre, ancre un souvenir dans notre mémoire. Grâce à ses capacités d'extension de la mémoire et de la cognition, la représentation visuelle d'information est également une aide pour celui qui doit prendre une décision parce qu'elle augmente ses capacités de compréhension et de traitement de l'information (Zhu, Chen, 2008). De plus, la visualisation de données permet de cristalliser la connaissance parce qu'elle rend possible ce qu'ils appellent le « *sense-making* visuel », c'est-à-dire la construction de sens supportée par un dispositif visuel (Chi et al, 1999). On comprend ainsi quelles sont les possibilités d'optimisation des capacités humaines de cognition et de *sense-making* rendues possibles par la visualisation de données.

4. Le *sense-making* individuel par Dervin

En sciences humaines, nombreuses sont les recherches sur le *sense-making*. Néanmoins, nous avons décidé de choisir Brenda Dervin comme auteur de référence. Elle a effectivement développé le concept de *sense-making* individuel dans le paradigme constructiviste. Elle place ainsi la focale sur la manière dont les connaissances sont construites par les personnes. Tout comme nous dans cette recherche, elle considère l'information comme comprise et interprétée du point de vue des individus (Maurel, 2010). Au-delà du simple concept, Dervin fait du *sense-making* une méthodologie à part entière. Elle part de la situation de l'individu et de son expérience subjective. Le *sense-making* est pour elle un comportement interne, à savoir cognitif, comme nous l'avons montré ci-dessus, et externe, où la personne agira dans l'espace et le temps en fonction de son expérience. Cela en fait également un processus communicatif (Maurel, 2010).

Dervin s'oppose à la conception de « sens unique » qui est encore attribuée à la communication, dans une vision « shannonienne » où un émetteur envoie un message à un récepteur qui le décodera sans rétroaction ni recul critique (Jacques, 2016). Selon nous, la visualisation de données est encore trop étudiée sous cet angle à l'heure actuelle. Alors qu'on essaie de comprendre quels principes de conception viseront à mieux transmettre le message et qu'on tente de saisir à quel point l'individu a compris le message envoyé, il y a une chose qui est souvent omise : la dimension subjective de la réception de l'information. Pour Dervin, le *sense-making* est une production dynamique de l'individu en fonction de ce qu'il observe dans son environnement. Cette production se réalise notamment grâce au

comblement de « *gaps* ». C'est d'ailleurs la métaphore « *gap bridging* » que Dervin emploie dans ses travaux (Dervin et al., 2003). Ainsi, il faut imaginer le *sense-making* en trois étapes : la situation initiale, la discontinuité et la phase d'aide ou de résultats, le tout se passant dans un espace-temps précis. « L'individu est appelé à expérimenter différents types de discontinuités dans une situation prenant place dans un contexte précis. L'aide prend la forme de recherche puis d'utilisation d'information, cette information devant être intégrée et comprise pour que l'individu puisse poursuivre la situation initiale » (Maurel, 2010, p. 2). Le « *gap bridging* » est ce qui permet de combler la discontinuité, qui est ce qui apparaît entre « la réalité » et les sens humains (Jacques, 2016). Il s'agit ainsi d'un véritable « pont informationnel » qui consistera à combler le manque grâce à l'expérience, à l'interprétation d'idées, de valeurs, etc. (Dervin, 1992). On se déplace ainsi entre des situations spatio-temporelles dans lesquelles apparaît la discontinuité (elle-même alors constante) et ce, afin de la combler grâce au « *gap bridging* » (Jacques, 2016). Il n'y a ainsi pas d'inertie : tout est en mouvement perpétuel dans cet espace-temps. Le message reçu et décodé par le récepteur ne sera sans doute pas équivalent au message initié par le destinataire. Cette idée est intéressante à exploiter en abordant les visualisations de données.

Dervin insiste ainsi sur l'utilisateur, sur la personne et non sur le système mis en place. Néanmoins, la métaphore de « *gap bridging* » n'en reste pas moins méthodologique puisqu'elle permet selon elle d'analyser les processus mis en place par les individus pour faire sens. La communication individuelle ou collective est ce qui rend le « *gap bridging* » possible : l'utilisateur construit sans cesse sa compréhension de la réalité en comblant des manques. Le *sense-making* assume ainsi que l'utilisateur se situe dans une situation culturelle, historique et autres, en lien avec l'habitus. Pour Dervin, l'individu reste maître de la construction de sens et c'est la production de ce « *gap bridging* » qu'il faut étudier : « L'objectif est donc d'étudier, dans la succession des situations spatio-temporelles, la construction du sens par un individu et de discerner les dimensions statiques (inflexibilité, habitudes, rigidité, stabilité) et les dimensions fluides (flexibilité, hasard, innovation, créativité) » (Dervin et al., 2003, p. 140, cité par Jacques, 2016). Cette méthodologie se concentre ainsi sur les différentes situations de la vie des individus puisqu'elle se focalise sur la construction de sens dans le temps.

Ainsi, nous pouvons nous interroger par rapport à la construction de sens d'un individu qui est confronté à une visualisation de données. Y a-t-il discontinuité lorsqu'une visualisation de données lui est présentée ? À partir de quel moment peut-on dire qu'il y a discontinuité ? L'ornementation des visualisations de données est-elle créatrice de discontinuité, nécessitant ainsi un « *gap bridging* » permettant de faire sens par rapport à une situation encore jamais rencontrée ? Cela pourrait tenir tout son sens dans le cadre de notre

recherche : en vérité, en termes de visualisation de données, il est extrêmement pertinent de se demander quelle situation de discontinuité l'individu tente-t-il de résoudre.

Selon nous, le plus important dans la théorie de Dervin est qu'elle se concentre bel et bien sur l'utilisateur. Elle prend ainsi en compte le vécu des individus, leur expérience, et les éléments qu'ils font entrer en compte lors du « *gap bridging* », tels que « la classe économique, le revenu et l'éducation, [qui] illustrent le type de contraintes structurelles qui limitent la création de nouvelles réponses » (Dervin et al., 2003, p. 275, notre traduction). On ne peut ainsi nier la dimension sociologique qu'elle joint à l'étude de la construction de sens, en prenant donc en compte des éléments propres aux individus. Dans le domaine de la visualisation de données, le rapprochement avec les travaux de Kennedy et Hill (2017, 2018), qui abordent les facteurs humains et émotionnels qui entrent en compte dans la réception des visualisations de données, est ainsi incontournable.

5. Les facteurs d'engagement envers la visualisation

Dans leurs travaux sur la visualisation de données, Kennedy et ses collègues cherchent à comprendre quelle est la place de la visualisation de données dans notre société et ce, à travers plusieurs axes. Ils replacent ainsi la visualisation dans le quotidien du citoyen lambda, « *the everyday life* ». Cette question est fondamentale dans nos travaux de recherche et Kennedy et ses collègues la soulèvent également : au-delà de la conception, qu'est-ce que l'efficacité de la visualisation de données ? Selon eux, les définitions actuelles de l'efficacité doivent être élargies parce que la réception d'une visualisation peut différencier d'une personnalité à l'autre, des compétences de l'individu, de son but et bien d'autres. On comprend dès lors le lien avec le *sense-making* de Dervin, à caractère individuel. Nous remettons donc en question l'aspect minimaliste et déterministe des définitions de Bertin et Tufte qui, au fond, sont fonctionnelles, concernent la réception visuelle de l'information, supposent que le récepteur soit un lecteur modèle, docile et qui ne creusent donc pas la question d'un point de vue sociologique. Leurs règles permettent de transmettre un message de la manière la plus « objective » possible, mais nous savons dès lors que la réception consiste en bien plus qu'un décodage neutre et systématique. Justement, c'est peut-être là que le bât blesse : si ces règles de conception et cette définition de l'efficacité suffisaient, au-delà de leur méconnaissance, pourquoi constate-t-on autant de pratiques d'ornementation des visualisations de données ? De nombreux éléments, propres aux individus d'une part, propres aux visualisations d'autre part, influencent la construction de sens.

Kennedy et Hill (2017, 2018) ont identifié plusieurs facteurs qui influencent la réception d'une visualisation de données. Elles se sont demandé comment les gens interagissent,

appréhendent, et s'engagent avec les données à travers la visualisation. Effectivement, il faut savoir que la visualisation de données est le media principal par lequel le plus de monde entre en contact avec les données (Kennedy, Hill, 2017). Il existe ainsi deux groupes de facteurs qui interfèrent avec l'engagement envers une visualisation : (1) les facteurs humains, sociaux, et propres à la visualisation et (2) les émotions que génèrent les visualisations dans le for intérieur des individus (Kirk, 2016, en ligne ; Kennedy, Hill, 2017, 2018). Nous proposons donc de les détailler. Ils auront par la suite une influence directe sur la méthodologie déployée dans le cadre de cette thèse.

Facteurs humains, sociaux, et propres à la visualisation

- **Le sujet de la visualisation** : une visualisation de données ne peut être présentée indépendamment de son sujet. Si les intérêts du récepteur sont touchés par la visualisation, il aura envie de s'engager. Sinon, il fera moins d'efforts.
- **La source d'où provient la visualisation** : si la visualisation présentée à un individu provient d'un blog en ligne qu'il adore, il y a de fortes chances qu'il s'y intéresse plus. A contrario, si elle vient d'un journal qu'il déteste, il aura des réticences avant même de commencer. Il y a aussi une notion de confiance dans un média qui peut entrer en ligne de compte.
- **Croyances et opinions** : l'individu s'engagera plus envers la visualisation si elle convient à ses croyances et opinions, mais pas seulement. Effectivement, la visualisation peut élargir ses opinions, lui faire découvrir de nouvelles choses ou provoquer des remises en question à propos de ce dont il était convaincu, ce qui constitue un challenge.
- **Le temps** : il s'agit d'une variable importante. Par exemple, si l'individu est dans une situation de travail où il doit agir assez rapidement, son engagement envers la visualisation en sera influencé. Avoir le temps suffisant pour traiter une visualisation est donc primordial.
- **La confiance en soi et ses compétences** : certaines personnes sont plus confiantes que d'autres en ce qui concerne leurs propres capacités à comprendre une visualisation de données. Le doute peut donc donner des motifs de réticence aux personnes qui n'ont pas cette confiance en leurs compétences de lecture et de compréhension de visualisations.
- **Les compétences langagières** : il sera évidemment plus facile pour une personne de saisir une visualisation si elle en comprend la langue, tout comme elle pourrait

s'engager davantage envers une visualisation réalisée en une langue qu'elle comprend mieux qu'une autre.

- **Les compétences mathématiques ou statistiques** : certains graphiques nécessitent ce type de compétences afin de comprendre des notions comme « moyenne » ou « échelle ». Les individus qui ont peu d'affinités avec ces notions peuvent parfois se montrer réticents s'il faut les mobiliser à un niveau élevé.
- **Les compétences de littératie visuelle**, qui aident l'individu à comprendre la signification attachée au visuel.
- **Les compétences informatiques** : elles sont utiles à l'individu quand il s'agit d'interagir avec la visualisation.
- **Les compétences de recul critique** : elles permettent à l'individu de se demander si le graphique est correct, si un élément y a davantage été mis en avant qu'un autre, de questionner l'information, etc.

Facteurs de réaction émotionnelle

- **La visualisation en tant que visualisation** : les individus peuvent tout simplement réagir émotionnellement à une visualisation parce qu'elle en est une. Quelqu'un qui, par nature, n'aime pas les graphiques, pourra se montrer agacé à sa vue par exemple. Les chercheurs du groupe *Seeing Data* ont également montré que les individus pouvaient tout simplement réagir à la visualisation et exprimer une réaction émotionnelle parce que le visuel est plaisant.
- **Les données et le sujet traité** : elles sont inextricablement liées au sujet. Évidemment, les individus peuvent réagir émotionnellement aux données. Imaginons qu'une visualisation traite du nombre de victimes du terrorisme sur une année et que le sujet soit sensible à cette thématique. Il pourrait montrer des émotions comme de l'empathie, par exemple. De même, si un individu déteste le sujet, sa réaction émotionnelle pourrait être de l'ordre de l'énervement. Son intérêt envers la visualisation pourrait être détourné, ou bien renforcé, mais accompagné d'une émotion négative.
- **La source ou le média** : certains médias peuvent sembler plus légitimes que d'autres à un individu, tout comme il peut aimer (ou pas) certaines sources. L'émotion en est donc influencée.

(Kirk, 2016, en ligne; Kennedy, Hill, 2017; Kennedy, Hill, 2018)

En bref, Kennedy et Hill (2017) résument tout ce que nous venons de détailler :

« Les émotions sont provoquées par les données elles-mêmes, le sujet, les endroits dans lesquels les données se situent (médias, sources) et par la perception que les gens ont de leurs propres capacités à comprendre les données et à s'y intéresser. Ces données, principalement appréhendées visuellement, sont presque toujours visuelles et statistiques et, pour les gens ordinaires, elles existent rarement sous une forme perceptuellement disponible en dehors de la visualisation. Cela conduit à un engagement émotionnel avec les données, ce que nous appelons « *feeling numbers* » » (Kennedy, Hill, 2017, p.10, notre traduction). Ces critères confirment également que les émotions ont leur place dans la construction de sens.

Chapitre 4 Problématisation et positionnement

Au cours de cette revue de la littérature, nous avons exploré les différents concepts en lien avec l'ornementation des visualisations de données. Tel un puzzle, nous avons assemblé toutes les pièces nous permettant dès à présent de formuler une problématique et une question de recherche claire et précise. Mais avant cela, réfléchir épistémologiquement aux concepts que nous associons est essentiel.

1. Positionnement épistémologique sur la théorie abordée

Avant tout, pour comprendre la manière dont les individus font sens des visualisations de données, il était nécessaire d'en comprendre les différentes implications théoriques. Pour commencer, nous avons défini la visualisation de données. En connaissant son utilité et les raccourcis cognitifs qu'elle nous offre, il est maintenant aisé de la considérer comme objet d'étude. La connaissance des principaux principes de conception émis par Bertin et Tufte, préconisant le minimalisme en design, est primordiale dans cette thèse. En effet, elle permet de mettre en évidence la tension croissante existant entre cette préconisation et les pratiques grandissantes d'ornementation. En réalité, ces principes ont été pensés pour transmettre un message le plus efficacement possible à un individu et ce, sans doute dans une vision à sens unique de la communication, en supposant que l'utilisateur soit un élève modèle qui comprend la communication telle qu'elle a été encodée. Or, ce n'est pas ainsi que nous considérons le sujet. Si, en théorie, les principes de Bertin et Tufte sont bons pour dépouiller la visualisation de tout biais entravant la compréhension, nous nous distançons de leur vision « shannonienne » de la communication. Notre positionnement épistémologique est constructiviste et interprétatif. Il nous permet de contrôler notre démarche de recherche (Bertacchini, 2009) et d'avancer vers une construction saine de la question de recherche.

Ainsi, pour les constructivistes, c'est par nos représentations du monde et nos connaissances que nous avons accès à la réalité. Cette réalité n'est donc pas atteinte directement, puisque nous nous en faisons tous une interprétation. Il est impossible, pour un constructiviste, d'atteindre une réalité objective (von Glasersfeld, 1994). Justement, alors que dans l'état de l'art, Bertin et Tufte cherchent à atteindre la visualisation de données idéale, régie par des règles strictes, nous nous montrons sceptique. Une visualisation de données objective existerait difficilement puisque qu'elle est le résultat d'étapes qui visent à construire l'interprétation d'un phénomène : la donnée est un construit, elle est extraite d'une activité. Elle est ensuite nettoyée et exploitée dans un but précis : être représentée dans un graphique afin d'expliquer une histoire, de communiquer un message. L'objectivité ne peut donc qu'être remise en question. On peut comprendre aisément cette tentation « d'objectivité absolue », lorsqu'on s'intéresse au *pipeline* que suivent les données avant de devenir une visualisation.

Depuis toujours se cache derrière les chiffres une idée d'objectivité : on leur fait confiance. À l'état brut, ils nous éloignent d'une « connaissance intime » et d'une « croyance personnelle ». Leur impersonnalité leur permet de transmettre l'information de manière standardisée et familière (Kennedy, Hill, 2017, p. 2-3). L'affichage de ces nombres est pourtant un choix humain ; la donnée est construite par l'humain. On parle aujourd'hui de *datafication*, un terme qui, à notre connaissance, n'a pas encore d'équivalent en français. En effet, Kennedy et Hill voient notre monde comme une « *data-driven society* », où une importance de plus en plus grande est accordée aux données. Notre monde actuel est dans une relation profonde avec les nombres. La *datafication* est la mise en nombres de notre vie sociale ; nos amitiés, intérêts et émotions, pourtant qualitatifs de prime abord, sont « *datafied* » (Kennedy, Hill, 2017). Aujourd'hui, la visualisation de données est le canal principal par lequel des individus de toutes sphères sociales confondues ont accès aux données (Kennedy, Hill, 2017, p. 2). Véritable mise en scène, elle est à la fois statistique et visuelle (Kennedy, Hill, 2017). Il est délicat de flanquer l'étiquette « d'objectivité » sur cet élément sociologique. Comme nous l'avons déjà mis en exergue, le produit final qu'est la visualisation de données réalisée à des fins de communication résulte d'un ensemble de choix humains et informatiques. Afin d'arriver à produire la visualisation finale, les données ne sont pas seulement extraites et traduites sous une forme visuelle. De fait, elles subissent de grandes transformations à plusieurs reprises. Ces transformations ont évidemment des implications importantes sur les données (Dasgupta, Kosara, 2012). La visualisation de données est souvent présentée comme un accès transparent aux données alors qu'elle est le point final d'un processus complet (Dasgupta, Kosara, 2012). Les différents traitements informatiques et algorithmiques subis par les données au cours du

processus nécessaire à leur visualisation relèvent évidemment de décisions. Le *pipeline* de visualisation indique ce que la transformation de la donnée en forme visuelle a impliqué tout au long du processus de traitement informatique. La meilleure explication selon nous est celle de Telea, qui, dans son livre « *Data Visualization: Principles and Practice* », illustre parfaitement la situation. En voici un schéma directement extrait (Figure 13).

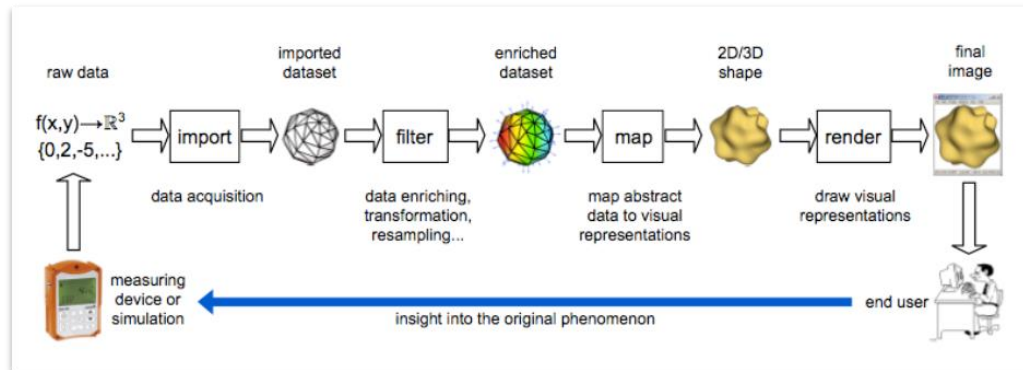


Figure 13 The Visualization Pipeline - A technical View, tirée de A. Telea (2014), *Data Visualization: Principles and Practice*, Boca Raton, Floride, CRC Press

Nous pouvons constater qu'énormément d'étapes et d'éléments entrent en jeu dans la construction d'une visualisation de données et ce, dès la « commande » du projet jusqu'au lancement final. Ainsi, l'idée d'une visualisation objective, miroir pur d'une « réalité » est illusoire car, même si beaucoup d'aspects relèvent aujourd'hui d'actes informatisés, le tout est soumis à des choix humains. La donnée, en elle-même, est par ailleurs un construit qui est un indice de la réalité. La visualisation de données n'illustre pas une réalité objective. Elle en donne une interprétation.

C'est donc bien parce qu'il s'agit d'une interprétation du monde qu'il convient de l'aborder à partir du paradigme constructiviste. Nos deux auteurs de référence, Bertin et Tufte, cherchent à faire respecter des lois de la nature – et notamment de la perception visuelle – pour correspondre à une seule et unique réalité de ce que doit être la visualisation de données idéale. Néanmoins, même en respectant ces « lois », rien ne garantit qu'un destinataire interprète un message tel qu'il a été encodé. Effectivement, Bertin et Tufte partent du principe que le récepteur est un lecteur modèle, docile, et qu'il décodera le message tel qu'il est encodé. Ils sont finalement très déterministes : Bertin parle de « l'information » comme si elle était un miroir parfait de la réalité, tandis que Tufte parle de « vérité » qui émerge des données. Pour eux, l'idée de réalité n'est pas quelque chose qui émerge de l'individu mais de la donnée en elle-même. Malheureusement, cette vision de la réception fait fi des éléments sociologiques, humains, techniques, personnels qui eux aussi entrent en compte dans la réception du message et plus largement dans la construction de sens. Vouloir prendre ces éléments en compte est cohérent avec la vision du *sense-*

making telle qu'exprimée par Dervin, légitimant ainsi les facteurs émotionnels et humains que Kennedy et Hill (2017, 2018) mettent en avant concernant l'engagement d'une personne envers la visualisation de données. En ajoutant une dimension sociologique à notre réflexion, c'est avec plus de souplesse que nous sommes à même de penser et de conceptualiser l'ornementation des visualisations de données. Effectivement, cette pratique grandissante ne peut désormais plus être ignorée ou simplement « bannie » des principes de conception des visualisations de données. En s'inscrivant dans le paradigme constructiviste, nous pensons qu'il est fortuit d'imaginer qu'une voie unique mène vers la visualisation idéale. La raison en est que nos perceptions de la réalité sont intentionnelles et qu'il existe plusieurs constructions mentales de la réalité.

Si nous les remettons en question, nous ne réfutons pas pour autant les principes émis par Bertin et Tufte. Puisque nous concevons la construction de sens comme l'assemblage de diverses composantes que sont la perception visuelle, la compréhension et le *gap bridging*, nous proposons d'admettre que les principes de conception ne peuvent agir que sur deux des trois composantes de la construction de sens, à savoir la perception et la compréhension. Ils ne peuvent néanmoins pas influencer le *gap bridging* réalisé par l'utilisateur : il s'agit d'un paramètre hors de portée du concepteur puisqu'il est le moment d'interprétation personnelle réalisée par l'individu, la visualisation stimulant alors différents facteurs humains et émotionnels. Néanmoins, la seule façon d'influencer cela est d'agir sur les facteurs propres à la visualisation, pouvant déclencher une réaction émotionnelle, telles que la source de la visualisation, la langue employée, etc. Cette dimension sociologique et l'acceptation d'une telle chose nous permettent ainsi de nous inscrire parfaitement dans le paradigme constructiviste, qui accepte ainsi que différentes conceptions de la réalité influencent l'interprétation finale d'une visualisation. Cairo résume très bien notre pensée : il considère que la visualisation de données soit un modèle, « une information plus abondante et de meilleure qualité engendre de meilleurs modèles » (Cairo, 2016, p. 1457, notre traduction). Nous pouvons imaginer ainsi qu'une information de bonne qualité est celle qui respecte les principes de conception. Ainsi, Cairo ajoute « un modèle solidement fondé sur ces méthodes est plus susceptible d'être vrai que faux. Je dis "susceptible" parce que dans cet exercice mental, je suppose que nous ne savons pas ce que "absolument vrai" signifie vraiment. Nous ne pouvons pas. Nous sommes des humains, vous vous souvenez ? » (Cairo, 2016, p. 1457, notre traduction). Selon lui, le modèle (donc la visualisation) se trouve sur un continuum entre l'absolument faux et l'absolument vrai. On ne sait ce qu'est « l'absolument vrai » : cela dépend des personnes. En proposant un modèle, on ne peut donc être sûr que le message sera complètement décodé comme il a été encodé (Cairo, 2016, 1275). Cela correspond bien à notre vision de la visualisation conçue comme une

interprétation des données qui est réinterprétée au moment de la conception. Les données sont ainsi un produit sans cesse réinterprété.

Mais encore, comment se positionner aujourd'hui par rapport à l'ornementation des visualisations de données ? Après avoir souligné que le terme n'était jamais défini dans les recherches qui l'abordent, bien que le considérant toujours de manière similaire (à savoir un « ajout » esthétique par rapport à un *raw chart*, dans l'objectif de rendre le visuel plus beau), nous nous sommes permise de tenter une définition de ce terme¹³, souvent constaté et jamais conceptualisé. En s'inscrivant dans le paradigme constructiviste, nous ne craignons pas d'étudier un objet dont l'appropriation peut être différente en fonction de dimensions subjectives – à savoir ici l'appréciation subjective du « beau ». S'il y a bien un mauvais geste à poser en visualisation de données, c'est celui de penser à la beauté et au design d'une visualisation avant de penser à la consistance de l'information et à la présentation des données. C'est d'ailleurs un principe fondamental de Tufte. L'ornementation est sans aucun doute du *chartjunk* au sens de Tufte, mais nous n'avons pas souhaité adopter ce terme qui nous semble négatif. Comme Few et Edge le suggèrent, les efforts esthétiques pourraient d'ailleurs avoir divers effets bénéfiques sur le destinataire de la visualisation (Few et Edge, 2011). Le terme « ornementation » nous permet d'adopter un positionnement non militant (contrairement au *chartjunk*), tout en acceptant qu'induire une notion de beauté dans les visualisations puisse influencer l'interprétation subjective du destinataire.

L'ornementation n'est pas seulement à concevoir comme quelque chose de potentiellement bénéfique. En effet, le principal danger est que cette pratique, contrairement aux principes de conception préconisés par Bertin et Tufte, est à même d'influencer les trois composantes de la construction de sens. Une visualisation ornementée peut ainsi agir sur la perception visuelle, la cognition et, additionnellement aux principes de conception, sur le *gap bridging* (voir Figure 14). La raison en est que l'ajout d'ornements dépasse volontairement la logique de Bertin pour qui l'utilisation d'un système de signes monosémique limite l'interprétation au décodage des signes arbitraires. En ajoutant des éléments de type polysémique, l'interprétation s'ouvre ainsi également sur des dimensions subjectives et propres à l'individu. Désormais, il convient donc de considérer l'ornementation non seulement comme une force potentielle, comme l'ont démontré certaines études (bienfaits sur la mémorisation, plaisir d'utilisation, etc.) mais également comme un danger au sens du *chartjunk*, quand la beauté serait alors privilégiée au détriment du message.

¹³ Pour rappel, [nous avons défini l'ornementation](#) précédemment.

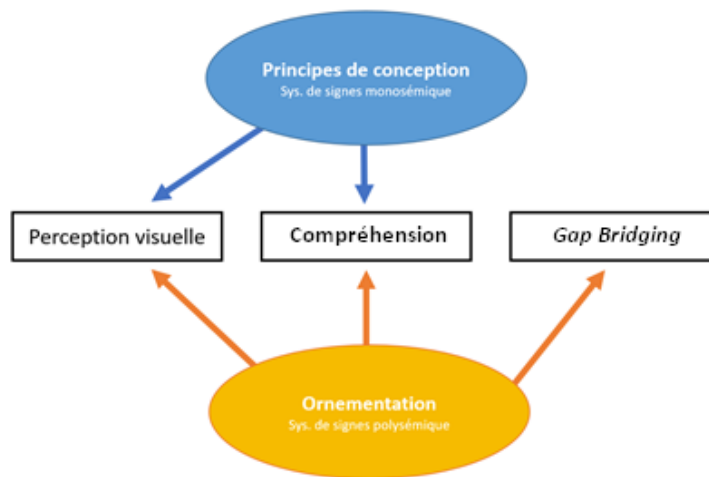


Figure 14 Influence des principes de conception des visualisations de données et de l'ornementation sur les différentes composantes de la construction de sens

2. Formulation de la question de recherche

Désormais, après avoir réfléchi les différentes imbrications théoriques possibles en ce qui concerne les concepts employés et les auteurs mobilisés, nous sommes en mesure de formuler une problématique simple qui est celle de l'influence de l'ornementation des visualisations de données sur la construction de sens de l'utilisateur. Cette problématique s'articule ainsi autour d'une question de recherche précise :

Par rapport à une visualisation de données qui est correcte en regard des principes de conception (que nous appelons « visualisation standard »), qu'est-ce qu'une visualisation de données ornementée induit sur la construction de sens

(1) au niveau de la perception visuelle et de la compréhension ?

(2) au niveau du gap bridging ?

Au-delà de la question de recherche, nous développons des ambitions tendant à proposer des recommandations fortes en ce qui concerne la conception de visualisations de données ornementées. Ainsi, afin de correspondre à l'objectif de présentation d'une visualisation de données et en phase avec les travaux de nos prédécesseurs Bertin, Tufte, Dervin, Kennedy et autres, nous souhaitons découvrir quelles recommandations en termes d'ornementation préconiser afin que toutes les composantes de la construction de sens soient stimulées pour que le message soit perçu de manière fidèle, tout en donnant une plus-value à la visualisation de données. En cela, nous nous concentrerons sur l'objet qui est donné à voir au sujet interprétant.

Partie II. Vers la récolte de données

Expérimenter, interviewer

Introduction Une récolte de données en phase avec les composantes de la construction de sens

Comme nous l'avons exposé dans la première partie, nous concevons la construction de sens comme une interaction constante entre trois composantes que sont la perception visuelle, la compréhension et le *gap bridging*. Afin de répondre à notre question de recherche, il convient de récolter des données permettant d'analyser non seulement la construction de sens du point de vue de la perception et de la compréhension mais aussi d'un point de vue individuel et subjectif. En sciences de l'information et de la communication, déployer une méthodologie permettant d'allier ces composantes qui ne sont pas de même nature est rare. Effectivement, il s'agit de récolter (1) des données quantitatives, relatives à des mesures physiologiques qui concernent la perception visuelle et (2) des données quantitatives et qualitatives, qui touchent à la subjectivité et aux émotions des destinataires des visualisations de données. Un tel choix nous permet de récolter des données à différents niveaux de granularité. En effet, la perception visuelle engendre le regard et plus particulièrement, des trajets oculaires qui sont réalisés par les individus de façon quasi inconsciente. Il est ainsi pertinent de relever des données propres à ce comportement visuel, qui est un fait directement observable. Ensuite, il est également important de relever des indications identiques pour tous les participants. L'objectif, dans ce cas, est de relever des données nous permettant de savoir si un ensemble d'individus note d'une façon similaire une visualisation de données ou pas. « Une telle visualisation de données est-elle claire, est-elle belle ? ». C'est ce genre d'informations que nous chercherons à récolter. Finalement, il sera question d'explorer plus en profondeur les avis et émotions des individus au contact d'une visualisation de données. L'entretien semi-directif est un bon moyen de recueillir de telles informations.

Les trois types de données que nous souhaitons récolter sont complémentaires, parce qu'elles correspondent à différents niveaux de granularité structurés par les composantes de la construction de sens que nous avons mis en avant dans la [première partie](#). La diversité

des méthodes à employer pour collecter ces différentes informations nécessite que nous puissions avoir le contrôle le plus total sur la récolte de données. Ainsi, la récolte de données effectuée dans le cadre de cette recherche s'effectue en laboratoire. Dans le cadre de nos diverses activités de recherche au sein du Social Media Lab de l'UCLouvain, nous avons eu la chance de disposer d'un laboratoire neuf et complètement modulable. Certes, nous n'avons pas saisi pas les données en contexte, comme le permettent les études réalisées en conditions écologiques. Cependant, les conditions de laboratoire nous permettent une flexibilité non négligeable en ce qui concerne le scénario de récolte des données, la fiabilité des données et leur précision. Il s'agit ainsi de saisir de manière complète et originale tous les aspects que revêt la construction de sens en visualisation de données, en incluant la question de l'ornementation dans la réflexion.

Parmi ses outils, le Social Media Lab dispose d'un *eye tracker*¹⁴. Cet appareillage technologique permet d'enregistrer le regard du participant sur un support numérique. Il est ainsi possible de récolter des données de regard (*gaze data*) de manière extrêmement précise, à savoir dans le temps (récolte des données en millisecondes) et dans l'espace (affichage de la localisation du regard à l'écran). Nous avons eu l'opportunité de récolter des données physiologiques très pertinentes par rapport à nos questions de recherches puisqu'elles permettent l'étude de questions relatives à la perception visuelle, composante importante de la construction de sens en visualisation de données. En ce qui concerne l'étude de la construction de sens individuelle et du *gap bridging*, il est selon nous possible de les étudier grâce à des méthodes plus traditionnelles en sciences sociales, comme l'entretien semi-directif. Malheureusement, comme nous le verrons dans le chapitre 5, nous avons découvert la méthodologie de Dervin trop tard pour l'appliquer. S'en inspirer a néanmoins été possible : la récolte de données effectuée dans cette thèse permet de commenter la manière dont les individus comblent la discontinuité, bien que ce ne soit pas la méthode originelle proposée par l'auteure. Nous y reviendrons.

La conception d'un protocole expérimental demande une certaine rigueur ainsi que la justification des choix de conception expérimentale. Quelles caractéristiques influencent le contact d'une personne avec la visualisation de données ? Quelles personnes choisir pour les tests ? Quelles caractéristiques doivent être prises en compte dans les visualisations contenues dans le corpus test ? Le choix de tous ces éléments doit être raisonné. Ils seront présentés dans cette partie, menant ainsi vers le protocole suivi en termes de récolte de données.

¹⁴ Les détails techniques de l'outil sont présentés *infra*.

Finalement, nous avons opté pour des expérimentations en laboratoire. Nous avons mis en place un test d'*eye tracking*, suivi d'un entretien semi-directif basé sur le contenu de l'expérience. Quarante personnes, professionnelles de la communication numérique ou d'un domaine connexe, ont participé à l'expérience, une à une. Cette deuxième partie se présentera comme suit : la première phase de notre méthodologie est l'exposition des participants à un corpus de visualisations de données. Un test d'*eye tracking*, ainsi que la récolte d'appréciations déclarées, ont été menés dans cette phase. Ensuite, suit la seconde phase de notre méthodologie, qui est l'entretien semi-directif. Les sujets ont participé successivement et directement aux deux phases de l'expérience. Pour une meilleure lisibilité et une meilleure justification des choix posés, ces phases sont présentées distinctement dans les chapitres 1 et 2 de cette partie. Pour terminer, nous aborderons les raisons pour lesquelles ces récoltes de données nous permettront de répondre à nos questions de recherche. Par contre, étant donné la diversité des données récoltées, les méthodes d'analyse seront abordées respectivement dans les parties de cette monographie consacrées à leurs analyses et interprétations.

Chapitre 1 Phase 1 : L'exposition aux visualisations de données et l'expérimentation *eye tracking*

Afin de pouvoir mener les expérimentations, il a été primordial de réaliser plusieurs choix définitifs, ayant des implications directes sur les conclusions de cette recherche. Ainsi, nous avons constitué un corpus de visualisations de données qui a été testé auprès d'un échantillon de 40 utilisateurs, scrupuleusement choisis eux aussi pour faire partie de l'expérience. L'expérimentation a aussi été pensée en fonction du matériel à disposition ainsi que de l'information à récolter. Tous ces choix sont justifiés dans ce chapitre.

1. Les visualisations exposées : constitution du corpus test

La première question qui s'est posée lors de la conception des expérimentations fut celle-ci : qu'allons-nous montrer aux participants dans le laboratoire ? Durant quelques semaines, nous avons arpenté le web dans les médias, blogs, journaux nationaux en ligne et réseaux sociaux afin de récolter des visualisations de données qui ne respectaient pas les principes de conception et qui dépassaient le « *chartjunk dreaded grid* », allant jusqu'à une ornementation assumée (voir *supra*). A vrai dire, il n'a pas été difficile de trouver de telles images. Cependant, il fallait, pour pouvoir sélectionner l'image ornementée à des fins expérimentales, que les données soient clairement saisissables. Effectivement, nous avons

décidé de sélectionner des cas réels plutôt que d'inventer des visualisations ornementées. Nous avons tout de même créé des visualisations de données correctes en regard des principes de conception, sur base des mêmes données et du même sujet que les visualisations sélectionnées. En bref, pour une image originale et ornementée, nous en avons créé la visualisation brute, standard (pourtant non pas « minimaliste »¹⁵), tendant ainsi vers ce que préconisent Bertin et Tufte. Le corpus se constitue ainsi de 40 visualisations de données, c'est-à-dire de 20 visualisations ornementées et de 20 visualisations standard aux données correspondantes. De plus, nous avons choisi de nous cantonner au diagramme en bâtonnets et à ses variantes puisqu'il s'agit de la visualisation de données la plus facilement perceptible par l'œil humain (Cleveland & McGill, 1984). Effectivement, une visualisation de données qui aura des propriétés facilitant la perception pré-attentive, et donc nécessitant une tâche perceptuelle élémentaire, mènera vers des appréciations ou jugements plus précis de la part des utilisateurs que d'autres formes graphiques. Le diagramme en bâtonnets nécessite une tâche perceptuelle d'un niveau élémentaire (Cleveland & McGill, 1984). Le choisir permet de se focaliser sur un type de visualisation et de créer des visuels qui sont à la portée de tout l'échantillon. Quoi qu'il en soit, la variable la plus importante du corpus se concentre sur l'ornementation du visuel. Y ajouter des variables relatives aux différents types de graphiques existants aurait apporté de la lourdeur aux expérimentations ainsi qu'une difficulté plus importante en ce qui concerne non seulement la récolte de données qualitative mais aussi l'analyse et comparaison des données qualitatives et quantitatives a posteriori. Par ailleurs, il arrive souvent que les concepteurs de visualisations de données ne choisissent pas le bon type de graphique. Dans la plupart des cas, le *bar chart* est une solution qui se suffit à elle-même, bien que jugée à tort trop peu attrayante visuellement. Dès lors, il arrive que les concepteurs de visualisation de données privilégient d'autres formes de graphiques, traditionnelles ou réinventées, bien moins efficaces d'un point de vue perceptuel. Ce genre d'erreurs, de graphiques ne respectant pas les principes de conception, se retrouvent également dans les expérimentations. Notons que certaines visualisations ornementées du corpus ne sont donc pas des diagrammes en bâtonnets (bubble charts, formes fantaisistes, ...). Elles ont néanmoins été sélectionnées parce que c'est le type de visualisation qu'elles auraient dû être si les principes de conception avaient correctement été fidèles. Ainsi, d'une part, les

¹⁵ Attention toutefois : les visualisations réalisées sont en effet des graphiques « standards » : pas de couleur inutile, étiquettes simplifiées, etc. Cependant, il ne s'agit pas du minimalisme maximal préconisé par Tufte, sans pour autant être du *chartjunk dreaded grid*. Les visuels sont simples mais pourraient être simplifiés d'avantage (vider les barres, enlever les échelles, etc.). Au vu des études de Inbar et al. (2007) et Quispel & Maes (2014), il n'est pas plus mal de ne pas pousser le minimalisme jusqu'à sa forme maximale car les utilisateurs ne sont pas habitués à la lecture de ces visuels.

visualisations standard, réalisées par nos soins, ne sont que des diagrammes en bâtonnets¹⁶. D'autre part, les visualisations récoltées dans les médias sont soit des *bar chart* améliorables, soit des visualisations de données qui auraient dû en être, mais réalisées sous une autre forme. Pour illustrer, la Figure 15 est l'originale récoltée sur *Twitter*. La Figure 16 en est la forme transformée.

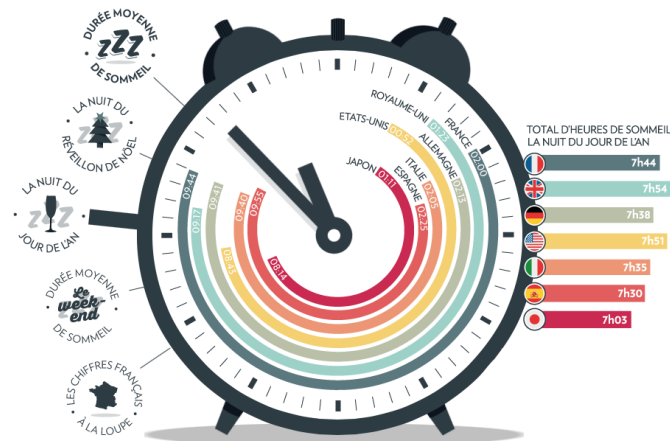


Figure 15 Data Match par Paris Match (2015), « Total d'heures de sommeil la nuit du jour de l'an », adaptée à l'expérimentation par nos soins (retrait de la source), repérée à <http://www.impactvisuel.net/2015/01/06/dataviz-comparer-les-habitudes-du-sommeil.html>

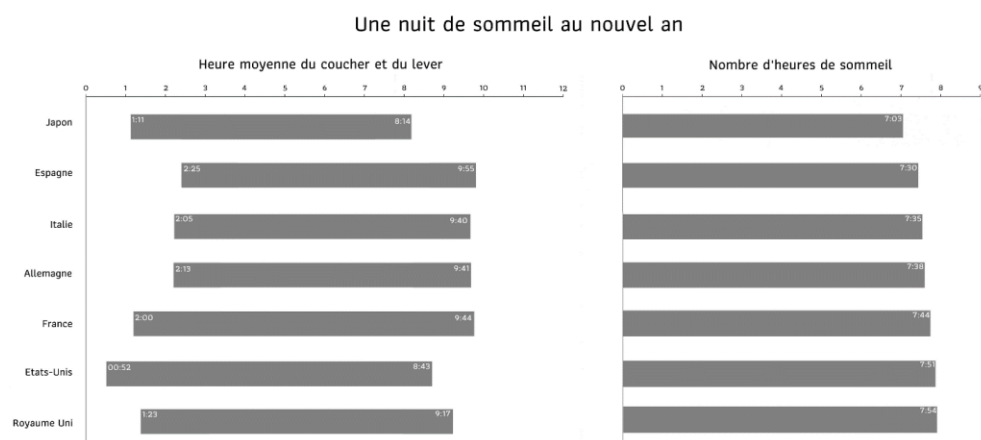


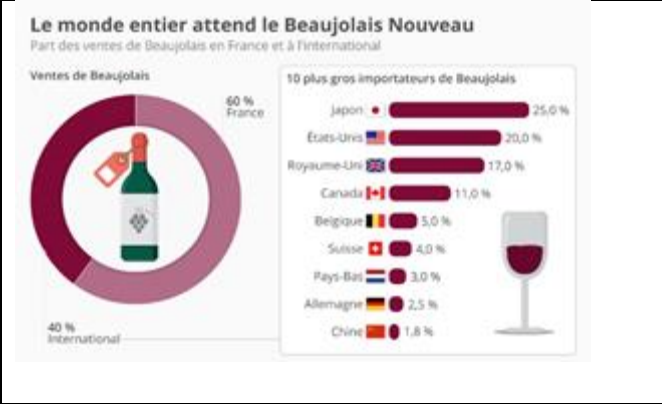
Figure 16 « La nuit de sommeil du 31 décembre », raw visualization inspirée de « Total d'heures de sommeil la nuit du jour de l'an », par l'auteur

De même, en ce qui concerne les images originales, le corpus se compose (1) de visualisations de données dont l'ornementation est un ajout d'images simples ou de pictogrammes sur une visualisation simple, comme une « décoration » (que nous avons appelée « apposition » dans la [Partie I](#) de ce manuscrit) et (2) de visualisations de données dont l'ornementation est concentrée dans le *data-ink*, c'est-à-dire dans les formes et traits du dessin relatifs aux données, comme on peut le voir dans le Tableau 4. Les différentes

¹⁶ À l'exception de deux visualisations qui comportent un diagramme circulaire ne constituant pas l'essentiel de l'information

dispositions d'images et de pictogrammes correspondent également à la classification de Haroz et ses collègues (2015) (stacked, stretched, etc.).

Tableau 4 Exemple de spécificités des visualisations ornementées présentes dans le corpus

	
(1) Ornementation par ajout d'images Apposition d'images simples et pictogrammes Haroz et al. (2015) : « <i>Superfluous pictographs</i> » Figure 17 Statista (2017), « Le monde entier attend le beaujolais nouveau », adaptée à l'expérimentation par nos soins (retrait de la source), repérée à https://fr.statista.com/infographie/11869/le-monde-entier-attend-le-beaujolais-nouveau/	(2) Ornementation appliquée au <i>data-ink</i> Application aux formes relatives à la présentation de données «Haroz et al. (2015) : « <i>Stretched pictographs</i> » Figure 18 Statista (2017), « Le succès du champagne dans le monde », adaptée à l'expérimentation par nos soins (retrait de la source), repérée à https://fr.statista.com/infographie/12229/le-succes-du-champagne-dans-le-monde/

L'ensemble des visualisations choisies correspond aux différents critères identifiés dans notre conclusion du chapitre sur l'ornementation. Tous les visuels se trouvent en annexe 3. Avant de procéder aux expérimentations, nous avons validé les visualisations « standards » en les montrant à nos pairs (expert en littératie médiatique et experte en *infovis*). De plus, les visualisations ont été caractérisées par différents attributs sur base de leur apparence et du respect des principes de conception ou non. Ces attributs n'ayant été utilisés que dans l'analyse, nous les exposerons plus tard.

Nous n'avons pas inclus de visualisations dynamiques ou interactives dans l'expérimentation. Cela aurait nécessité une tâche de la part de l'utilisateur, qui aurait exploré l'image des yeux afin de trouver les endroits cliquables. Il se serait dès lors concentré sur l'utilisabilité du support et non sur son esthétique. Or, c'est bien l'ornementation que nous avons abordé a posteriori lors d'entretiens semi-directifs. Nous avons souhaité que la concentration du participant soit focalisée sur la compréhension du contenu de la visualisation (et non sur la manière dont elle fonctionnerait si elle était interactive) et sur les sentiments et émotions qu'elle pouvait induire, de par son sujet ou

son esthétique. Nous assumons ainsi pleinement le choix de réaliser cette étude en utilisant des visualisations de données statiques uniquement.

De plus, les facteurs humains et émotionnels de Kennedy et Hill (2017, 2018) ont influencé le choix des visualisations mais aussi quelques transformations mineures que nous y avons apporté. Dans le Tableau 5, nous expliquons en quoi ces facteurs ont influencé les choix de visualisations ornementées.

Tableau 5 Facteurs humains et émotionnels selon Kennedy et Hill (2017, 2018): adaptation du corpus test

Le sujet de la visualisation	Les sujets sélectionnés sont simples et légers. Nourriture, séries, nouvel an, dons, ... sont des thèmes accessibles sur lesquels il est facile d'avoir une opinion.
La source d'où provient la visualisation et le média	Nous avons estimé nécessaire d'enlever toute marque d'appartenance aux visualisations présentées aux participants. Marques de journaux, liens URL d'un blog et sources des données ont donc été retirés de ces images afin de ne pas influencer le participant. La focale doit ainsi rester sur les données et leur esthétique.
Le temps	Les visuels étaient à disposition des participants le temps qu'ils ont estimé nécessaire. Ce sont eux qui ont décidé du temps accordé à chaque visualisation.
La confiance en soi et ses compétences	Il est difficile d'avoir un contrôle sur ce facteur. Néanmoins, ce sujet peut être exploré lors de l'analyse des entretiens semi-directifs.
Les compétences langagières	Nous avons traduit la plupart des visualisations en français, puisque la moitié du corpus était en anglais. Seulement deux visualisations n'ont pas été traduites, étant dans un anglais très accessible (niveau A1).
Les compétences mathématiques ou statistiques	Pour des raisons de faisabilité et de précision, nous n'avons pas utilisé de visualisations demandant des connaissances statistiques poussées. Du côté des visualisations standard, seuls des diagrammes en bâtonnets ont été présentés aux participants. Ils sont facilement accessibles et, de plus, il s'agit du type de visualisation de données le plus facilement compréhensible par l'œil humain (Cleveland & McGill, 1984). En ce qui concerne les visualisations ornementées, elles sont simples et accessibles, mais sont susceptibles de générer ce sentiment de besoin de compétences chez l'utilisateur, et cela fait partie du jeu de l'expérimentation.
Les compétences de littératie visuelle	Peu de contrôle à ce sujet dans l'expérimentation. Il est tout de même possible d'explorer la thématique dans l'entretien semi-directif.
Les compétences informatiques	Les visualisations sont statiques. Les participants ne doivent pas manipuler de paramètres informatiques propres à une visualisation dynamique ou interactive.
La visualisation en tant que visualisation	C'est tout l'enjeu des expérimentations : laisser le participant s'exprimer à propos de ce que la visualisation évoque pour lui, autant sur le plan des données que sur le plan de l'attrait visuel.
Les données et le sujet traité	L'intérêt est aussi de voir si le participant aborde en premier lieu le sujet, les données ou l'apparence visuelle de la visualisation lors de l'entretien semi-directif. Laisser parler le participant à propos du sujet de la visualisation (« j'aime – je n'aime pas ce sujet, il m'angoisse, il ne m'évoque rien, ... ») est une clé quant aux éléments constitutifs de la construction de sens.

2. Les sujets : choix de l'échantillon

Ensuite, il a fallu déterminer sur quels individus les expérimentations allaient porter. En effet, le domaine de la visualisation de données est très vaste et interdisciplinaire. Nous ne cherchons pas à généraliser nos résultats. Nous avons réalisé un échantillonnage raisonné (Bertacchini, 2009). De même, il a fallu être à même de définir la taille de l'échantillon pour atteindre un degré de confiance suffisant en ce qui concerne les éléments à en retirer. Il en va de la crédibilité suffisante des éléments relevés (Bertacchini, 2009). Notre échantillon se constitue en grande partie de professionnels de la communication numérique : graphistes, journalistes, chargés de communication numérique, marketing digital, ... Pour apporter un peu de nuance, il était également composé, dans une moindre mesure, de deux *data analysts*, d'un informaticien reconverti dans la communication web et de deux chargés d'événements. Nous n'avons pas tenu compte de l'âge et des caractéristiques démographiques, estimant que, de par leur profession, les personnes constituant l'échantillon étaient averties dans le numérique – et donc susceptibles de croiser de temps en temps des visualisations de données, durant leur temps libre ou leur activité professionnelle.

Dans le cadre de nos fonctions au sein de l'UCLouvain Social Media Lab, nous avons assuré la coordination administrative du certificat universitaire en communication web. Le certificat, qui est un programme de formation continue, vise le perfectionnement de la profession ainsi que la reconversion professionnelle dans le domaine de la communication numérique. Cette tâche nous a ouvert à un réseau de personnes qui travaillent ou sont intéressées par le secteur. Le profil des diplômés de la filière en communication est également adéquat. C'est vers ces deux publics de professionnels que nous avons souhaité nous tourner. Ces personnes proviennent d'un peu partout en Fédération Wallonie-Bruxelles et Bruxelles-Capitale. Environ 142 personnes ont été contactées par invitation (email ou message privé sur LinkedIn) afin d'en recruter 40. Nous n'avons pas souhaité restreindre l'âge des participants¹⁷. Effectivement, il nous semble que le fait d'appartenir à une classe professionnelle particulière et d'y avoir des compétences communes est suffisant. Nous avons également été prudente à propos de leurs caractéristiques physiologiques. Pour des raisons techniques relatives à l'utilisation de l'*eye tracker*, il fallait que leur vue soit correcte. Dans le cas où ils portaient des lunettes, nous avons vérifié la compatibilité de la taille de la monture et des verres au moment de la calibration dans la phase d'expérimentation. Comme nous l'avons précisé, le nombre de personnes choisies

¹⁷ Nous avons cependant conservé l'âge des participants par précaution, bien que nous ne l'utilisons pas dans nos interprétations. Les âges de l'échantillon varient entre 23 et 50 ans (profils professionnels allant de juniors à seniors). À titre d'information, la moyenne d'âge est de 29 ans.

pour participer aux expérimentations s'élève à 40. Nous sommes consciente que cet échantillon n'est pas représentatif de la population étudiée, ce qui n'est pas notre préoccupation première.

3. Matériel

Dans le cadre de nos activités de recherche au Social Media Lab de l'UCLouvain, nous avons pu mener nos expérimentations dans le Usability Lab, composé de plusieurs locaux, de meubles mobiles, ergonomiques, et de matériel technologique de pointe. L'utilisation de ce dispositif a nécessité le développement d'une méthodologie adéquate. C'est pourquoi nous décrivons maintenant le matériel employé.

3.1. Mobilier, locaux et logistique

Nous avons occupé le Usability Lab du Social Media Lab durant la quasi-totalité du mois d'avril 2018. Il se compose de cinq locaux différents, affectés à différentes tâches, et de mobilier ergonomique, adaptable au déroulement des expérimentations (voir Figure 19). Ainsi, nous avons disposé de deux salles indépendantes (utilisées comme salle d'attente et salle de débriefing), d'une régie depuis laquelle nous avons le contrôle des expérimentations, ainsi que de la salle principale dans laquelle les participants ont pris part à l'expérimentation sur un poste d'ordinateur. La salle de débriefing était confortable et tout à fait adéquate au déroulement d'entretiens semi-directifs, de par sa disposition et du matériel dont elle était équipée (fauteuils, grand écran et caméra). L'eyetracker Tobii Pro X3-120 est équipé du programme Tobii Studio, qui permet de créer différentes expérimentations de manière claire et intuitive, et qui nous fournit par la suite des données brutes à analyser ainsi que des données visualisées. Les visualisations faisant partie du corpus test ont cependant été affichées dans une interface web insérée elle-même dans le logiciel (voir *infra*).

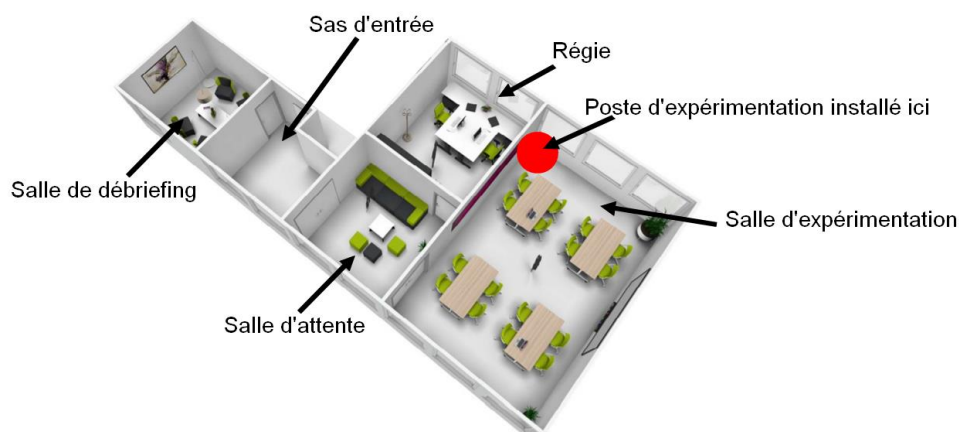


Figure 19 Usability Lab. Plan de C. Peeters annoté par nos soins

3.2. Technologie utilisée

Dans le cadre de cette recherche nous avons récolté des données grâce à une méthodologie axée autour de l'*eye tracking*. L'*eye tracking*, ou oculométrie en français, se base sur l'idée qu'une relation directe existe entre ce que les gens regardent intentionnellement et ce



Figure 20 Tobii Pro (2017), Tobii Pro X3-120, repéré à <https://www.tobii.com/product-listing/tobii-pro-x3-120>

qu'ils sont en train de penser. Il s'agit de l'hypothèse œil-cognition ou *eye-mind* de Just et Carpenter (1976). Ainsi, l'appareillage technologique permet d'enregistrer le regard du participant sur un support donné et fournit ainsi de nombreuses données de regard et métriques, calculées en millisecondes. Nous avons utilisé le dispositif Tobii Pro X3-120¹⁸ (Figure 20). Cet *eye tracker* mesure 32,4 cm de longueur. Il s'agit d'une barre de 1,5 cm de largeur qui se fixe très aisément au bord inférieur de l'écran, ce qui le rend non-intrusif. Les données relatives au regard (plus communément appelés « *gaze data* », qui sont donc les données oculométriques) sont calculées grâce à un *firmware* (micrologiciel) externe. Cela permet de lancer des tests dans un environnement contrôlé, avec un traitement informatique complètement dédié à la récolte de *gaze data*. La fréquence d'échantillonnage (nombre d'échantillons prélevés par unité de temps) s'élève à 120Hz et l'*eye tracker* capte la traçabilité du regard avec une précision de 0,24°. Les contraintes quant à l'utilisation de l'outil ne sont pas nombreuses mais ont nécessité une attention particulière lors de la mise en place du dispositif. Effectivement, le participant n'est pas obligé de rester immobile : il a une certaine liberté de mouvement de la tête, mais tout de même limitée (50cm x 40cm). Il doit également se situer à minimum 50 cm et maximum 90 cm de distance de l'appareil afin qu'il fonctionne correctement. Ainsi, nous avons précisé aux participants qu'ils ne devaient pas faire d'amples mouvements de buste. Une simple chaise, sans roulettes, a été utilisée afin d'éviter tout gigotement. Nous avons utilisé l'outil sous le logiciel Tobii Studio. Il permet de concevoir l'expérience, de récolter les données et de les analyser.

Par précaution, après avoir testé l'outil, nous avons conclu que les individus portant des lunettes pouvaient rencontrer quelques problèmes quant à la captation des *gaze data*. En effet, il est important que la monture et les verres du participant ne soient pas trop épais. L'*eye tracker*, fonctionnant sur une technologie à infrarouge, capte les données grâce aux effets de réverbération dans les pupilles de l'individu. Des faux cils peuvent également gêner. Nous avons tenu compte de ces impératifs lors de la constitution de l'échantillon. Avant de commencer l'expérience, le logiciel procède à une phase de calibration. Il s'agit

¹⁸ Pour voir les détails techniques : <https://www.tobii.com/product-listing/tobii-pro-x3-120/>

d'un processus qui utilise les caractéristiques géométriques des yeux du participant comme base afin d'effectuer par la suite un calcul de point de vue précis et entièrement personnalisé

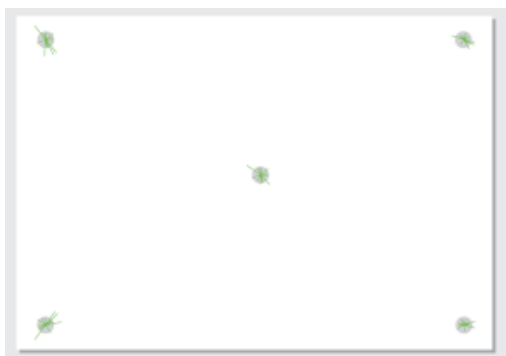


Figure 21 Tobii Pro (2017), "What happens during the eye tracker calibration", repéré à <https://www.tobii.com/learn-and-support/learn/eye-tracking-essentials/what-happens-during-the-eye-tracker-calibration/>

(Tobii, 2017, en ligne). Pour ce faire, le participant doit suivre des yeux une boule de couleur qui se déplace sur l'écran. Cette phase de calibration se fait bien sûr avant que l'expérimentation commence. Lors des expérimentations, la calibration s'est effectuée grâce à cinq points de fixation (voir Figure 21), ce qui dure plus longtemps que pour un nombre de points inférieur, mais qui augmente la précision de la capture.

Nous n'avons pas été présente auprès du participant que lors de la phase de calibration et de l'expérimentation. Dès son installation sur la chaise, nous nous rendions en régie afin de le laisser se concentrer sur l'écran. Nous lui donnions alors les consignes grâce à un micro. Cela a ainsi empêché qu'il ne quitte l'écran des yeux une fois la calibration effectuée et qu'il ne change radicalement de position.

4. Articulation des outils et déroulement des expérimentations

Nous pouvons maintenant faire le point sur la manière dont nous avons articulé le matériel, le corpus test et notre échantillon. Effectivement, comment procéder concrètement à une expérimentation qui tient la route en regard de nos questions de recherche ? Il convient enfin d'en expliquer le déroulement et la place qu'y prennent tous les éléments que nous avons précédemment expliqués. Pour une lecture plus rapide, le scénario d'expérimentation se trouve en annexe 4. Nous prenons néanmoins le temps d'en faire une présentation plus détaillée ici, tandis que l'annexe 4 donne plus de détails pratiques (micro, contrôles ordinateurs, instructions).

4.1. Déroulement général

Tout d'abord, le participant entrait dans le laboratoire et s'installait dans la salle d'attente quelques minutes, le temps de recevoir les détails pratiques relatifs à l'expérience. Effectivement, l'utilisation de l'*eye tracker* nécessite une phase de calibration. L'*eye tracker* détecte la pupille de l'individu et définit ainsi les paramètres propres à la bonne capture du mouvement du regard. De même, la chaise de l'individu ne pouvait pas bouger. Il était évidemment libre de mouvements dans la limite du raisonnable. De trop grands

mouvements de tête, qui feraient quitter le regard de l'écran, ou avancer ou reculer durant l'expérience auraient pu perturber la calibration effectuée préalablement. Informer le participant de tous ces éléments était capital. L'individu s'installait donc dans la salle d'expérimentation et procédait à l'étape de calibration. Il était assisté depuis la régie, afin qu'il puisse rester concentré et garder les yeux rivés sur l'écran. Nous pouvions ainsi donner quelques consignes supplémentaires sans nous inquiéter que le participant ne se retourne pour écouter. Une fois cette étape terminée, l'expérimentation pouvait être lancée.

Vingt visualisations étaient donc présentées à l'individu (dix visualisations standard et dix visualisations ornementées). Nous avons utilisé la méthode de randomisation des images en carré latin (cf. *infra*). Le participant recevait des consignes simples : lire, essayer de comprendre et cliquer lorsque cela était fait. Entre chaque visualisation, un questionnaire (appréciations déclarées) permettait de mentionner à quel point la visualisation a été appréciée, jugée belle, claire, intéressante et comprise. Nous en expliquerons les raisons ensuite.

Une fois cette phase expérimentale terminée, l'individu était invité à se détendre autour d'un café, dans la salle d'attente. Durant ce temps, nous nous occupons de

- Récupérer et consulter les réponses aux appréciations déclarées. Sur une tablette, adjointes au guide de débriefing, elles permettaient de réguler chaque entretien. Le fait de pouvoir les consulter sur un dispositif mobile était d'une grande facilité (voir *infra*).
- Assurer d'un point de vue technique la projection du *gaze replay*¹⁹ dans la salle de débriefing.

Pour que les pensées des individus soient les plus « fraîches » possibles et donc faciles à se remémorer, il convenait d'effectuer ces tâches le plus rapidement possible. Après un instant de pause pour le participant et après l'apprêtement du matériel pour la suite, l'entretien semi-directif pouvait commencer dans la salle prévue à cet effet (phase 2 détaillée dans le [chapitre 2](#)).

¹⁹ Vidéo de l'expérience, montrant le trajet du regard à l'écran

4.2. Contrôle des effets d'apprentissage et des effets de séquence

Quarante personnes ont participé à l'expérience. Elles ont été amenées à voir et essayer de comprendre 20 visualisations de données. Pour rappel, le corpus entier est composé de 40 images (20 visualisations ornementées, issues des médias et 20 visualisations standard correspondantes, créées par nos soins). Ainsi, tous les individus n'ont pas forcément vu les mêmes images mais ils ont observé autant de visualisations standard que de visualisations ornementées. Il est également bon de préciser que les individus qui ont vu une visualisation ornementée n'ont pas pu observer la version standard construite sur base des mêmes données et ce, afin d'éviter les effets d'apprentissage propres aux données et au sujet abordé par la visualisation. Puisque le temps est un facteur important dans la réception d'une visualisation (Kennedy, Hill, 2017, 2018), nous ne voulions pas presser les participants. Ainsi, ils ont disposé du temps qu'ils souhaitaient pour regarder chaque visualisation de données. C'était alors à eux de passer à l'image suivante grâce à une commande (click de la souris).

Ensuite, il était essentiel de nous montrer prudente quant au contrôle des éventuels effets de séquence qui seraient apparus dans les données récoltées si tous les individus visionnaient à chaque fois les visualisations dans le même ordre. Il était donc également important de contrôler les menaces à la validité interne de notre expérience (Campbell & Stanley, 1967). Effectivement, la validité interne est un indicateur important qui permet d'évaluer la fiabilité d'un protocole expérimental. Pour contrôler les diverses menaces et effets indésirables, il était fondamental de randomiser l'ordre dans lequel les visualisations de données ont été présentées à l'individu. Pour cela, nous avons choisi la méthode de randomisation en carré latin. Ce design expérimental est intéressant parce qu'il permet de tenir compte de deux sources de variation minimum ; autrement dit, il permet de contrôler au moins deux sources de nuisance à la validité interne de l'expérience simultanément (Saville & Wood, 1991). Comme son nom l'indique, sur papier, le carré latin consiste à réaliser un carré où le nombre de lignes et de colonnes doit correspondre au nombre de niveaux de traitement. Donc, en ayant quatre traitements, il faut quatre lignes et quatre colonnes pour créer un carré. Cela donne un design où chacun des traitements se trouve dans chaque ligne et dans chaque colonne (Saville & Wood, 1991). Nous n'avons cependant pas réalisé pas un « vrai » carré latin puisque nous ne souhaitons pas que les 40 participants voient 40 visualisations mais uniquement 20 (cf. *supra*). Effectivement, il aurait été laborieux de voir et comprendre 40 images sans que cela ne soit trop long. De même, voir toutes les visualisations de notre corpus aurait pu engendrer des effets d'apprentissage, comme nous l'avons expliqué plus haut. Cependant, la méthode de

randomisation du carré latin reste efficace en ce qui nous concerne, comme en atteste la Figure 22.

<table> <tr><td>1</td><td>2</td><td>3</td><td>4</td></tr> <tr><td>2</td><td>1</td><td>4</td><td>3</td></tr> <tr><td>3</td><td>4</td><td>2</td><td>1</td></tr> <tr><td>4</td><td>3</td><td>1</td><td>2</td></tr> </table>	1	2	3	4	2	1	4	3	3	4	2	1	4	3	1	2	<table> <tr><td>1</td><td>2</td><td>3</td><td>4</td></tr> <tr><td>2</td><td>3</td><td>4</td><td>1</td></tr> <tr><td>3</td><td>4</td><td>1</td><td>2</td></tr> <tr><td>4</td><td>1</td><td>2</td><td>3</td></tr> </table>	1	2	3	4	2	3	4	1	3	4	1	2	4	1	2	3	<table> <tr><td>1</td><td>2</td><td>3</td><td>4</td></tr> <tr><td>2</td><td>4</td><td>1</td><td>3</td></tr> <tr><td>3</td><td>1</td><td>4</td><td>2</td></tr> <tr><td>4</td><td>3</td><td>2</td><td>1</td></tr> </table>	1	2	3	4	2	4	1	3	3	1	4	2	4	3	2	1	<table> <tr><td>1</td><td>2</td><td>3</td><td>4</td></tr> <tr><td>2</td><td>1</td><td>4</td><td>3</td></tr> <tr><td>3</td><td>4</td><td>1</td><td>2</td></tr> <tr><td>4</td><td>3</td><td>2</td><td>1</td></tr> </table>	1	2	3	4	2	1	4	3	3	4	1	2	4	3	2	1
1	2	3	4																																																																
2	1	4	3																																																																
3	4	2	1																																																																
4	3	1	2																																																																
1	2	3	4																																																																
2	3	4	1																																																																
3	4	1	2																																																																
4	1	2	3																																																																
1	2	3	4																																																																
2	4	1	3																																																																
3	1	4	2																																																																
4	3	2	1																																																																
1	2	3	4																																																																
2	1	4	3																																																																
3	4	1	2																																																																
4	3	2	1																																																																

Table 13.8: Standard latin squares for the 4×4 layout.

To obtain a random layout, take the following steps:

1. Use random numbers to choose a standard square.
2. Reorder all the rows except the first of the standard square according to a sequence of random numbers; for example, 3 2 4 for a 4×4 square.
3. Reorder all the columns of the resulting square according to a new random sequence; for example 4 1 3 2 for a 4×4 square.
4. This final square gives you the treatment numbers in your random layout.

Figure 22 Saville & Wood (1991, p.359), Randomization method for a latin square in Latin Square Design. In *Statistical Methods: The Geometric Approach* (p. 340-353). Springer, New York, NY.

Ainsi, en se cantonnant à 20 visualisations par personne, il est possible, en suivant cette méthode, d'obtenir un ordre de présentation des visualisations tel que

1. chaque individu visualise les images dans un ordre différent et
2. chaque individu visualise 10 images standard, créées par nos soins, et 10 images ornementées, issues de cas réels.

Dans ce cas, pour les 40 visualisations, nous avons procédé comme le montre la Figure 23, avec R pour visualisation standard et I pour visualisation ornementée. I1 est donc la version ornementée de R1, par exemple.

Ce qui donne, appliqué aux 40 participants, en leur attribuant 20 visualisations, le résultat figurant dans le Tableau 6.

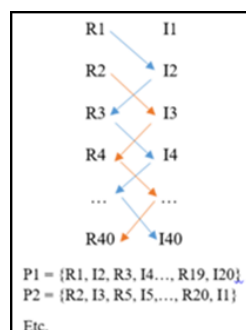


Figure 23 Randomisation de l'ordre des visualisations de données dans l'expérimentation avec la méthode du carré latin

Tableau 6 Randomisation de l'ordre de présentation des visualisations de données en carré latin

P1	R1	I2	R3	I4	R5	I6	R7	I8	R9	I10	R11	I12	R13	I14	R15	I16	R17	I18	R19	I20
P2	R2	I3	R4	I5	R6	I7	R8	I9	R10	I11	R12	I13	R14	I15	R16	I17	R18	I19	R20	I1
P3	R3	I4	R5	I6	R7	I8	R9	I10	R11	I12	R13	I14	R15	I16	R17	I18	R19	I20	R1	I2
P4	R4	I5	R6	I7	R8	I9	R10	I11	R12	I13	R14	I15	R16	I17	R18	I19	R20	I1	R2	I3
P5	R5	I6	R7	I8	R9	I10	R11	I12	R13	I14	R15	I16	R17	I18	R19	I20	R1	I2	R3	I4
P6	R6	I7	R8	I9	R10	I11	R12	I13	R14	I15	R16	I17	R18	I19	R20	I1	R2	I3	R4	I5
P7	R7	I8	R9	I10	R11	I12	R13	I14	R15	I16	R17	I18	R19	I20	R1	I2	R3	I4	R5	I6
P8	R8	I9	R10	I11	R12	I13	R14	I15	R16	I17	R18	I19	R20	I1	R2	I3	R4	I5	R6	I7
P9	R9	I10	R11	I12	R13	I14	R15	I16	R17	I18	R19	I20	R1	I2	R3	I4	R5	I6	R7	I8
P10	R10	I11	R12	I13	R14	I15	R16	I17	R18	I19	R20	I1	R2	I3	R4	I5	R6	I7	R8	I9
P11	R11	I12	R13	I14	R15	I16	R17	I18	R19	I20	R1	I2	R3	I4	R5	I6	R7	I8	R9	I10
P12	R12	I13	R14	I15	R16	I17	R18	I19	R20	I1	R2	I3	R4	I5	R6	I7	R8	I9	R10	I11
P13	R13	I14	R15	I16	R17	I18	R19	I20	R1	I2	R3	I4	R5	I6	R7	I8	R9	I10	R11	I12
P14	R14	I15	R16	I17	R18	I19	R20	I1	R2	I3	R4	I5	R6	I7	R8	I9	R10	I11	R12	I13
P15	R15	I16	R17	I18	R19	I20	R1	I2	R3	I4	R5	I6	R7	I8	R9	I10	R11	I12	R13	I14
P16	R16	I17	R18	I19	R20	I1	R2	I3	R4	I5	R6	I7	R8	I9	R10	I11	R12	I13	R14	I15
P17	R17	I18	R19	I20	R1	I2	R3	I4	R5	I6	R7	I8	R9	I10	R11	I12	R13	I14	R15	I16
P18	R18	I19	R20	I1	R2	I3	R4	I5	R6	I7	R8	I9	R10	I11	R12	I13	R14	I15	R16	I17
P19	R19	I20	R1	I2	R3	I4	R5	I6	R7	I8	R9	I10	R11	I12	R13	I14	R15	I16	R17	I18
P20	R20	I1	R2	I3	R4	I5	R6	I7	R8	I9	R10	I11	R12	I13	R14	I15	R16	I17	R18	I19
P21	I1	R2	I3	R4	I5	R6	I7	R8	I9	R10	I11	R12	I13	R14	I15	R16	I17	R18	I19	R20
P22	I2	R3	I4	R5	I6	R7	I8	R9	I10	R11	I12	R13	I14	R15	I16	R17	I18	R19	I20	R1
P23	I3	R4	I5	R6	I7	R8	I9	R10	I11	R12	I13	R14	I15	R16	I17	R18	I19	R20	I1	R2
P24	I4	R5	I6	R7	I8	R9	I10	R11	I12	R13	I14	R15	I16	R17	I18	R19	I20	R1	I2	R3
P25	I5	R6	I7	R8	I9	R10	I11	R12	I13	R14	I15	R16	I17	R18	I19	R20	I1	R2	I3	R4
P26	I6	R7	I8	R9	I10	R11	I12	R13	I14	R15	I16	R17	I18	R19	I20	R1	I2	R3	I4	R5
P27	I7	R8	I9	R10	I11	R12	I13	R14	I15	R16	I17	R18	I19	R20	I1	R2	I3	R4	I5	R6
P28	I8	R9	I10	R11	I12	R13	I14	R15	I16	R17	I18	R19	I20	R1	I2	R3	I4	R5	I6	R7
P29	I9	R10	I11	R12	I13	R14	I15	R16	I17	R18	I19	R20	I1	R2	I3	R4	I5	R6	I7	R8
P30	I10	R11	I12	R13	I14	R15	I16	R17	I18	R19	I20	R1	I2	R3	I4	R5	I6	R7	I8	R9
P31	I11	R12	I13	R14	I15	R16	I17	R18	I19	R20	I1	R2	I3	R4	I5	R6	I7	R8	I9	R10
P32	I12	R13	I14	R15	I16	R17	I18	R19	I20	R1	I2	R3	I4	R5	I6	R7	I8	R9	I10	R11
P33	I13	R14	I15	R16	I17	R18	I19	R20	I1	R2	I3	R4	I5	R6	I7	R8	I9	R10	I11	R12
P34	I14	R15	I16	R17	I18	R19	I20	R1	I2	R3	I4	R5	I6	R7	I8	R9	I10	R11	I12	R13
P35	I15	R16	I17	R18	I19	R20	I1	R2	I3	R4	I5	R6	I7	R8	I9	R10	I11	R12	I13	R14
P36	I16	R17	I18	R19	I20	R1	I2	R3	I4	R5	I6	R7	I8	R9	I10	R11	I12	R13	I14	R15
P37	I17	R18	I19	R20	I1	R2	I3	R4	I5	R6	I7	R8	I9	R10	I11	R12	I13	R14	I15	R16
P38	I18	R19	I20	R1	I2	R3	I4	R5	I6	R7	I8	R9	I10	R11	I12	R13	I14	R15	I16	R17
P39	I19	R20	I1	R2	I3	R4	I5	R6	I7	R8	I9	R10	I11	R12	I13	R14	I15	R16	I17	R18
P40	I20	R1	I2	R3	I4	R5	I6	R7	I8	R9	I10	R11	I12	R13	I14	R15	I16	R17	I18	R19

Le tableau se lit de façon linéaire et indique ainsi l'ordre dans lequel sont apparues les images pour chaque individu. Ce protocole expérimental est fiable parce qu'il permet de contrôler des facteurs qui, autrement, menaceraient la validité interne des expérimentations, à savoir la maturation, le testing et l'instrumentation (Campbell & Stanley, 1967).

4.3. Instrument de présentation des visualisations

Comment mettre en pratique et élaborer les quarante expérimentations conformément à nos ambitions en termes de randomisation ? Si Tobii Studio est un outil efficace qui permet de concevoir les expérimentations, il y avait deux désavantages à préparer les expérimentations directement sur le logiciel. Premièrement, cela aurait impliqué de concevoir les quarante expérimentations à la main. Ensuite, la version de Tobii Studio disponible à ce moment-là au Social Media Lab ne permettait pas de proposer des échelles de Likert aux participants alors que nous souhaitons récolter des données par ce moyen

également (*infra*). Cependant, il est possible de préparer une expérimentation sur Tobii Studio qui se lance à partir d'une interface web. C'est l'option qui a été choisie.

Une interface web a été créée par l'informaticien développeur du Social Media Lab et la randomisation en carré latin y a été codée. Ainsi, en configurant au préalable l'expérimentation et en inscrivant le nom et le numéro du participant dans le compte administrateur de notre interface web, nous avons directement accès à l'ordre correspondant à son passage dans l'ensemble des participants. Lorsque l'expérimentation se lançait, l'interface web apparaissait en plein écran avec, premièrement, une page de consignes très simples. Ensuite, sur l'interface, l'expérimentation se passait comme ceci : une visualisation apparaissait, le participant la lisait et après avoir estimé que l'information était clairement saisie, il cliquait et arrivait sur une page avec des échelles de Likert. Après les avoir complétées, il cliquait de nouveau et accédait à la nouvelle visualisation et ainsi de suite, jusqu'à ce que les 20 visualisations aient été observées.

L'*eye tracker* capturait les données oculométriques durant tout le temps d'utilisation de cette interface web. Tobii Studio considère chaque URL comme un contenu à part entière. Le logiciel scinde donc les données et les visualisations (cartes de chaleur, *gaze plot*, etc.) en fonction des URL où se situent les visualisations. Chaque image a ainsi été placée sur une URL ou sous-URL différente, appartenant au même site web. Les données oculométriques ont aussi été captées lors des moments de réponse aux appréciations déclarées. Nous n'en tiendrons néanmoins pas compte dans l'analyse des données. Ainsi, à une image correspondait un lien URL pour lequel toutes les données ont été récoltées. L'utilisation de cette interface web a donc permis une récolte de données supplémentaire pour notre étude (appréciations déclarées) et a augmenté les potentialités de l'outil Tobii.

4.4. La récolte des appréciations déclarées

Les échelles de Likert complétées après le visionnage de chaque visualisation ont permis de récolter des appréciations déclarées, ou *self reported appreciations* dans le champ UX (Albert & Tullis, 2013). De manière quasi instantanée, il est ainsi possible de savoir si l'individu a aimé ou non les visualisations, s'il les comprenait, les trouvait claires ou belles (Figure 24) puisque ces métriques d'appréciation personnelles sont inspirées des facteurs sociaux influençant la réception d'une visualisation (Kennedy, Hill, 2017, 2018). Nous avons estimé que le plus grand intérêt était d'avoir une trace du ressenti directement exploitable dans le guidage de l'entretien semi-directif. Il était évident d'utiliser ces appréciations lors des entretiens pour aborder non pas chaque image avec les individus,

mais celles qui ont suscité des variations importantes d'opinion, ce qui constitue une méthode d'interview inédite (voir *infra*).

J'ai trouvé cette visualisation...		
Pas belle	○ ○ ○ ○ ○	Belle
Pas intéressante	○ ○ ○ ○ ○	Intéressante
Pas claire	○ ○ ○ ○ ○	Claire
Je n'ai pas compris cette visualisation	○ ○ ○ ○ ○	J'ai compris cette visualisation
Je n'ai pas aimé cette visualisation	○ ○ ○ ○ ○	J'ai aimé cette visualisation

Figure 24 Appréciations déclarées utilisées dans les expérimentations

Afin de laisser la concentration de l'individu à un niveau maximal lorsqu'il regardait une visualisation de données, mais aussi pour ne pas perturber les *gaze data*, nous avons choisi de disposer ce questionnaire d'appréciations déclarées après chaque visualisation sur une nouvelle page et non pas sur le même écran que la visualisation de données.

De plus, les résultats de ces appréciations déclarées ont été analysés statistiquement (voir [Partie III](#)). Après avoir visionné les 20 visualisations de données et répondu aux 20 questionnaires d'appréciations déclarées, les participants pouvaient passer à la suite de l'expérience. Après une brève pause, ils changeaient de pièce pour participer à l'entretien semi-directif.

Chapitre 2 Phase 2 : L'entretien semi-directif

Afin d'explorer la composante subjective de la construction de sens, nous avons mené un entretien semi-directif directement après l'expérimentation avec chaque participant. La conduction des entretiens correspond à bon nombre de choix méthodologiques mais aussi de conception expérimentale. Nous les exposons tous dans ce chapitre. En premier lieu, nous présentons les raisons nous ayant poussée à choisir l'entretien semi-directif. Ensuite, nous expliquons comment nous avons utilisé les appréciations déclarées comme guide d'entretien avant de présenter le matériel et la logistique mise en place. Puis, nous expliquons le déroulement général avant de revenir sur la *sense-making methodology* de Dervin (1993), découverte sur le tard.

1. Pourquoi choisir l'entretien semi-directif ?

En *eye tracking*, l'outil à lui seul ne suffit pas : il doit s'utiliser dans le cadre d'une méthode correctement conçue et balisée. Ainsi, ce n'est pas parce qu'on sait où un individu a posé le regard qu'on sait pourquoi il l'a fait. Dans le cadre d'une question de recherche qui

aborde la construction de sens, la récolte de données qualitative est essentielle. Une interrogation s'est vite posée : comment réussir à mener un entretien qui allie la technologie utilisée et ce, en se focalisant sur la construction de sens, le tout en prenant en compte les différents facteurs d'appropriation identifiés en théorie ? Différents courants et différentes méthodes nous ont permis de poser ce choix de manière avisée. Bien qu'aucun d'entre eux n'ait pu être appliqué à notre cas en tant que tel, leurs apports ont été notables. Ils nous ont permis de rester en phase avec notre positionnement épistémologique, tout en étant inspirants. C'est après leur lecture que nous avons été capable d'imaginer notre propre méthode d'entretien.

1.1. Le paradigme éactif et le cours d'expérience

Pour commencer, c'est avec enthousiasme que nous avons pris connaissance du paradigme éactif (Penelaud, 2010) ainsi que du cours de l'action et de l'expérience de Theureau (1992, 2002). Ces derniers ont été largement exploités par Schmitt (2012) qui a analysé la construction de sens des individus lors d'une expérience de visite au musée. La recherche de Schmitt consiste à équiper le sujet de lunettes *eye tracker*²⁰ et de revenir par la suite sur cette expérience avec l'individu, en la revivant avec lui une seconde fois, grâce à la vidéo de sa propre expérience. Ainsi, à posteriori, le sujet devait exprimer, expliquer avec ses mots ce qu'il a compris, le sens que revêtaient pour lui ces moments d'action et d'expérience dans le musée (Schmitt, 2012). Selon le paradigme éactif, l'acteur et son environnement sont dans un système clos autonome dans lequel tout acteur interagit avec ce qui est source de perturbation pour lui (Rouve & Ria, 2008). Les interactions acteur-objets, acteur-acteur, etc. sont des perturbations. L'acteur contribue donc à l'émergence de la chose signifiante. Il n'y a pas d'information « extérieure » selon cette vision mais seulement des perturbations qui provoquent une réorganisation, et qui provoquent donc information, formation à l'intérieur du système (Schmitt, 2012). Le contexte et le sens sont liés, mais le contexte du concepteur n'est pas celui du récepteur. Le destinataire crée le contexte (Schmitt, 2012). Dans notre cas, on peut imaginer que le concepteur d'une visualisation de données la conçoit au mieux pour transmettre son message et que le récepteur en crée le sens. De même, l'intérêt de pratiquer le cours d'expérience est de révéler la connaissance construite par l'individu. Effectivement, il n'aura peut-être pas produit du sens comme l'entend l'émetteur, mais il aura produit une connaissance qui n'en a pas moins de valeur pour lui-même. Cependant, le chercheur ne peut deviner quelle est la connaissance produite par le participant, même s'il possède des données importantes,

²⁰ Le système d'oculométrie est alors disposé sur une monture de lunettes, ce qui assure une grande mobilité

comme les points d'attention captés par l'*eye tracker*. En utilisant des mots et en exprimant ses sensations, le participant test énonce ses connaissances – qui sont toutes relatives et propres à lui-même – et le chercheur est à même de les saisir, bien qu'il doive par la suite les interpréter (Schmitt, 2012).

Malgré la richesse théorique du paradigme énonctif, nous ne pouvons pas nous y inscrire. Effectivement, selon Theureau, le cours d'action ne peut avoir lieu que dans une situation réelle, dans un contexte réel, puisque l'individu réagit par rapport à son environnement, crée du sens à partir de ce qui l'intéresse. Il évince donc les conditions de laboratoire pour les méthodes énonctives (Theureau, 2006). Ce paradigme se base sur l'autopoïèse et l'homéostasie, deux théories issues de la biologie qui décrivent tout système comme capable de se produire, de maintenir son organisation ou de revenir à l'équilibre après une perturbation (Penelaud, 2010). L'analyse de ces éléments en contexte est fondamentale et bien qu'il existe beaucoup de points de convergence entre le paradigme constructiviste et énonctif, nous nous distançons de l'énonction tout en gardant en tête cette autre vision intéressante de la construction de sens qui, de plus, s'articule plutôt bien avec l'idée de *sense-making* de Dervin qui parle d'un comblement de la discontinuité (Dervin, 1992).

Quoi qu'il en soit, c'est bien la méthode de Shmitt (2012) qui a fait germer dans notre réflexion l'idée d'un entretien post-expérimentation.

1.2. Les méthodes « *think-aloud* »

Avec l'idée d'une récolte de données qualitative vient rapidement la question de la temporalité : pourquoi débriefer avec le participant après l'expérience, au lieu de le laisser s'exprimer en direct, pendant l'expérience ? C'est ainsi que nous avons découvert les « *think-aloud methods* », ou, tout simplement, méthodes qui consistent à « réfléchir à voix haute ». Ainsi, cette méthode consiste à dire tout haut ce que l'on pense au plus profond de soi, « à l'intérieur », durant l'expérience. Il est donc demandé au participant de parler, de dire ce qu'il pense le plus possible. Il peut également être demandé au participant qu'il décrive ce qu'il est en train de faire. Toutes nos pensées ne prennent pas forcément une forme verbale. Y mettre des mots ne révèle donc pas toute la substance de leur signification, mais permet d'atteindre des données très complètes et intéressantes malgré tout (Charters, 2003). Par ailleurs, le plus gros avantage est de saisir des pensées qui, autrement, auraient disparu quasi-instantanément. Tout humain perd facilement le fil de ses idées tant le débit de pensées peut parfois s'activer à toute vitesse. Il y a donc là un focus sur la connaissance immédiate. Souvent complétée par une interview rétrospective directement après l'expérience et ce, afin de revenir avec l'individu sur ses dires et d'aller plus loin, la

méthode est très fiable. Plus le temps est court entre l'expérience et cette rétrospection, plus l'interview est fiable (Charters, 2003).

Intéressée dans un premier temps par cette méthode qui aurait pu nous faire gagner un temps considérable, nous nous sommes ensuite ravisée. Effectivement, la méthode « *think-aloud* » ne permet pas de révéler toute la complexité des pensées car l'individu retranscrit « simplement » en mots ce qu'il a ressenti. Ses paroles sont une traduction de ses pensées : à partir de ses pensées, il produit du sens et par la suite il met des mots sur ce sens (Charters, 2003). Néanmoins, les tâches qui demandent une charge cognitive élevée peuvent interférer avec la verbalisation : d'autres processus de réflexion croisent de l'information verbale de différentes origines dans la mémoire du participant qui est en train de « travailler » (Pressley & Afflerbach, 1995, cités par Charters, 2003). En bref, demander à l'individu de réfléchir à voix haute durant une tâche compliquée, qui lui demande une réflexion importante, peut tout simplement embrouiller ses pensées – et donc, in fine, les données du chercheur. Lire, comprendre des visualisations de données et leur attribuer des appréciations nous a semblé être une tâche suffisamment compliquée. Notre crainte était également d'apporter un biais aux *gaze data* : si un utilisateur parle de quelque chose qu'il est en train d'observer, il le regardera plus longtemps, avec plus d'insistance. A l'inverse, la description d'opérations plus introspectives peut provoquer une perturbation du regard (regarder ailleurs, dans le vide, vers le haut...). De même, il cherchera autour de l'élément dont il parle d'autres choses qui pourraient nourrir ses commentaires. Ainsi, les *gaze data* auraient été biaisées et n'auraient pas pu correspondre à une utilisation naturelle du dispositif (Pernice & Nielsen, 2009).

Pour pallier cela, il existe une méthode appelée « *retrospective thinking-aloud method with eye tracking* » (RTE), inspirée de la « *retrospective think-aloud method* » (RTA), qui consiste à inviter l'utilisateur à revoir la vidéo de son expérience et à la commenter, à en dire le plus de choses possible. Dans le cas de la RTE, une dimension supplémentaire est ajoutée : la vidéo de l'expérience contient les actions de l'individu (« *gaze-replay* » de la vidéo captée par l'*eye tracker*) (Elling & al., 2011). Selon Elling & al., il n'y a pas de différence significative entre l'efficacité des RTA ou des RTE. Le mouvement de l'œil en RTE n'a pas de valeur additionnelle (Elling & al., 2011). Effectivement, voir un *gaze replay* est dans la majorité des cas nouveau pour les participants, ce qui peut être intrigant ou amusant, voir distrayant pour eux. De même, un *gaze replay* est très rapide. Ainsi, l'accompagnement de l'individu dans son interprétation est très important, tout comme ralentir la vidéo peut être d'une grande aide, pour le pousser à de meilleurs commentaires (Elling & al., 2011). Les études analysées par Elling et ses collègues concernent dans la plupart des cas des sites web et des actions, tandis que nos expériences

concernaient des images statiques. Nous apportons tout de même une grande attention à leurs préconisations.

1.3.L'entretien semi-directif, méthode traditionnelle de récolte de données en sciences de la communication

Grâce à notre parcours universitaire²¹ et à la recherche menée dans le cadre de notre mémoire²², nous avons acquis des connaissances relatives à la méthode de récolte de données empiriques qu'est l'entretien semi-directif. Nous avons déjà pratiqué l'exercice et étions donc sensibilisée à la manière de conduire de tels entretiens. Songer à ce type de récolte de données nous a donc semblé tout naturel, bien qu'incomplet dans le cadre de la méthodologie usitée dans cette recherche-ci.

Comme son nom l'indique, l'entretien semi-directif laisse une certaine liberté de parole à la personne interrogée. « Le chercheur a établi la liste des thèmes et/ ou des questions qu'il veut aborder. Ils sont repris dans un guide d'entretien qui permettra à l'interviewer de ne rien oublier et ainsi de traiter tous les thèmes prévus avec les personnes interviewées. L'ordre (prévu) dans lequel les thèmes sont abordés ne doit pas forcément être fidèle (...). En principe, on débute toujours par la même question pour tous les entretiens et on clôture l'entretien de manière semblable » (Derèze, 2009, p. 110).

Ainsi, l'entretien semi-directif « réalise un compromis souvent optimal entre la liberté d'expression du répondant et la structure de la recherche. Le répondant s'exprime sur les thèmes qu'il souhaite, et dans son propre langage : la directivité de l'entretien est donc très réduite. Le chercheur en retire deux éléments : (1) des informations sur ce qu'il cherche a priori (les thèmes du guide de l'interviewer) ; et (2) des données auxquelles il n'aurait pas pensé (la surprise venant de la réalité du terrain) » (Romelaer, 2005, p. 104).

Il semblerait que le type d'entretiens menés dans le cadre de cette recherche soit compréhensif (Kaufmann, 1996). En effet, dans notre position interprétative, nous acceptons que le sujet, même s'il n'a pas forcément raison en soi, a ses propres raisons de penser ce qu'il déclare lors de l'entretien. Cela tient tout son sens puisque Kaufmann considère l'individu comme un « concentré du monde social », qui en reflète toutes les contradictions, doutes et ambivalences (Kaufmann, 1996). Le chercheur doit être capable

²¹ Master 120 en information et communication à finalité spécialisée en communication des organisations, institutions et associations, UCLouvain.

²² Étude visant à identifier les facteurs de motivation et d'implication du personnel de terrain dans une entreprise de soins à domicile, sur base de 27 entretiens semi-directifs et six focus groups. Référence : Andry, T (2015), *Le cadre intermédiaire, acteur de la motivation et de l'implication organisationnelle : une approche croisée: réfléchir la communication interne chez Aide & Soins à Domicile Hainaut Oriental*, Mémoire en communication réalisé sous la direction de F. Lambotte, UCLouvain, Mons.

de faire preuve d'empathie et de se mettre dans la peau de l'interlocuteur pour qu'il puisse exposer toutes ses raisons. Cette posture compréhensive permet ainsi d'atteindre une saturation du modèle, au moment où les observations, en se cumulant, confirment ce qui est attendu (Kaufmann, 1996). Cette proposition de Kaufmann est intéressante parce qu'il admet ainsi que l'entretien semi-directif n'est pas une méthode unique ni exacte. Il est fonction d'un artisanat propre au chercheur qui, pour atteindre son objectif de récolte de l'information, bricole sans doute. Kaufmann considère ainsi le participant comme un informateur (Kaufmann, 1996). L'entretien compréhensif, en phase avec la position interprétative des propos des individus, permet de résoudre la difficulté de l'essai de généralisation qui surgit souvent de la récolte de données qualitatives par entretien semi-directif (Kaufmann, 1996).

L'entretien semi-directif, voire compréhensif, semble donc adéquat quant à l'exploration des facteurs humains et émotionnels de Kennedy et Hill (2017, 2018), qui s'apparentent d'ailleurs aux « thèmes » dont parle Derèze (2009), tout en laissant la liberté de parole aux participants. Elle est essentielle dans une recherche sur la construction de sens où seul l'individu peut aborder et expliquer lui-même ce qui l'a aidé à construire ce sens. Le guide d'entretien est alors un outil essentiel, permettant de garder une structure dans la récolte des données. Comme nous le verrons ensuite, des questions seules ne suffisaient pas dans le cadre de notre méthodologie. La technologie employée et l'expérimentation réalisée avant l'entretien mené *de facto* ont eu une incidence sur le guide d'entretien²³ prévu à cet effet.

1.4. Contribution méthodologique de ces modèles

Les différents modèles présentés ont retenu toute notre attention pour différentes raisons. Effectivement, si nous ne pouvons les appliquer en tant que tels, ils ont stimulé notre réflexion quant à des détails méthodologiques d'une importance haute.

Du paradigme énatif, nous avons appris que nul ne peut deviner la connaissance produite par un individu, et qu'il doit l'expliquer lui-même. La recherche de Schmitt (2012), montre qu'un système d'*eye tracking* ne peut se suffire à lui-même dans les recherches qui visent à comprendre les pensées d'un individu, et qu'il doit être absolument combiné à une récolte de données qualitatives, laquelle permettra d'interpréter une connaissance nouvelle énamée par la parole de l'individu. Bien que propre au paradigme énatif, une telle pensée ne nous semble pas en contradiction avec notre positionnement constructiviste et interprétativiste.

²³ Le guide d'entretien utilisé se trouve en annexe 4.

La méthode *think-aloud*, quant à elle, permet de comprendre la puissance de l'expression orale et nous a poussé à confirmer notre choix de débriefing oral à posteriori de l'expérimentation. Bien que des interviews rétrospectives sont souvent menées dans les expériences de ce type, cette méthode peut néanmoins être improductive combinée à une tâche trop compliquée. Une mise en garde a aussi retenu notre attention : une parole n'est pas une pensée, elle en est la traduction et n'en contient donc pas toute la substance. De même, nous avons choisi de faire confiance aux études qui ne confirment pas la plus-value d'une RTE (Elling et al., 2011). Malgré tout, nous avons souhaité utiliser le *gaze replay* dans le cadre de l'entretien, mais sans demander au participant de formuler ses pensées à voix haute.

Après avoir pris connaissance de ces deux modèles de récolte de données, il nous a semblé pertinent de revenir à nos connaissances premières sur l'entretien semi-directif. Ainsi, vu la flexibilité et la souplesse possibles tant dans la conception du guide que dans la conduction de l'entretien, ce type d'entretien nous a semblé tout à fait adaptable à la récolte de données que nous souhaitions mener parce que

- Il laisse libre court à la parole et laisse donc l'opportunité à l'interviewé d'énacter le sens avec une liberté relativement variable ;
- L'ordre des thèmes abordés importe peu : c'est une grande force lorsque le sujet de la discussion concerne 20 visualisations de données au sujet et à l'esthétique différente ;
- Il permet d'aborder tous les sujets à traiter avant de terminer l'entretien, que ce soit de manière inductive (surprise du terrain selon Romelaer (2009)) et déductive (exploration quant aux ressources théoriques identifiées).

La décision s'est donc portée sur la conduction d'un entretien semi-directif a posteriori de l'expérimentation. Il fallait néanmoins l'adapter afin de stimuler la discussion à propos des visualisations d'une part, et en exploitant la force de l'*eye tracking* d'autre part. Ainsi, avec un outil technologique aussi simple qu'une tablette et un écran, nous avons pu mener un type d'entretien inédit : l'entretien semi-directif guidé par les appréciations déclarées.

2. Les appréciations déclarées comme guide d'entretien : une méthode nouvelle

L'objectif des entretiens semi-directifs qui ont directement suivi l'expérimentation était d'éveiller la parole des participants à propos de la façon dont ils ont lu et compris les visualisations de données, en tenant compte de leurs caractéristiques personnelles et en écoutant leur propre interprétation. Nous ne cherchions donc pas uniquement à savoir pourquoi une personne avait posé le regard dans certaines zones de l'écran. L'intérêt était d'évoquer le sens personnel, donnant ainsi de l'importance à l'expression des émotions et à l'habitus de la personne.

Pour mener l'entretien, nous avons préparé des questions à l'avance (voir annexe 4). Elles consistaient à faire parler la personne interrogée à propos de ses propres comportements face aux visualisations de données, des visualisations qu'elle préférait, qu'elle détestait, etc. Ainsi, les questions ouvertes laissaient une certaine souplesse au participant qui, loin de se sentir "interrogé", pouvait parler librement et donner un avis pleinement personnel. Des questions sur les émotions et sur les facteurs de Kennedy et Hill étaient également prévues. Néanmoins, le support nous ayant le plus aidé lors de la conduction des entretiens était la tablette, sur laquelle figuraient les différentes appréciations déclarées complétées par les participants. Directement affichées sous nos yeux grâce à l'interface en ligne conçue par le développeur du Social Media Lab (Figure 25), elles ont permis de mieux structurer les entretiens semi-directifs qui sont au cœur de la démarche de récolte de données subjectives. Ces éléments pouvaient ainsi facilement être mis en évidence : en revenant sur une question à laquelle un individu avait lui-même répondu quasi-spontanément, la discussion sur le sujet se développait très facilement. En voici un exemple :

T²⁴ : Celle-ci par exemple : je vois dans vos réponses que vous la trouvez intéressante mais pas claire. Pourquoi ?

E : Oui, ça me fait penser aux nouveaux graphiques que l'on voit quand on regarde les relations avec les liens sur les réseaux sociaux et tout ça. Je trouve ça toujours intéressant et moderne... mais plus compliqué à comprendre. J'ai l'impression - et je vais peut-être loin - qu'il faut prendre tellement d'attention là-dessus que c'est peut-être comme ça que d'autres idées pourraient surgir. Mais là je parle comme quelqu'un qui devrait traiter la donnée et qui ne chercherait que l'information directe...

²⁴ La lettre « T » nous identifie en tant que chercheure ayant mené l'entretien.

Ainsi, alors que le guide d’entretien permettait de stimuler la discussion sur des thèmes relatifs aux éléments subjectifs que couvrent les facteurs humains et émotionnels de Kennedy et Hill (2017, 2018), les appréciations déclarées permettaient de guider la discussion

- Sur des appréciations personnelles (beauté, clarté, compréhension, appréciation, intérêt) mises en tension avec l’apparence de la visualisation ;
- Sur les visualisations de données ayant été notées de manière unanime (notes maximales (5/5) ou minimales (1/5)) ou avec de grandes tensions entre les appréciations (par exemple : beauté 5/5, clarté 1/5). Cela donnait directement une indication des visualisations méritant une attention particulière lors de l’entretien, ce qui constitue un gain de temps.

Cela nécessitait évidemment une lecture des appréciations au préalable en régie.

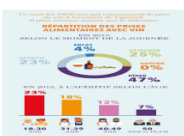
	Cette visualisation est belle	• •
	Cette visualisation est intéressante	• • • •
	Cette visualisation est claire	• • • • •
	J'ai compris cette visualisation	• • • • •
	J'aime cette visualisation	• • • • •
	Cette visualisation est belle	• • • • •
	Cette visualisation est intéressante	• • •
	Cette visualisation est claire	• •
	J'ai compris cette visualisation	• • • • •
	J'aime cette visualisation	• • •

Figure 25 Exemples de résultats des appréciations déclarées avec un intérêt pour le débriefing

De plus, l’interface conçue par le développeur du Social Media Lab était particulièrement adéquate pour l’exercice. Ainsi, il était possible de visualiser un aperçu des réponses, par participant, directement en ligne et quelques secondes à peine après la fin de l’expérimentation, sur une même page. Le téléchargement de fichiers .csv contenant les données est une autre option disponible sur l’interface.

D’un point de vue pratique, lorsque nous abordions des appréciations, nous montrions simultanément la visualisation au participant, afin de stimuler rapidement ses souvenirs. S’il le souhaitait et afin de stimuler davantage la discussion, le participant pouvait demander à voir le *gaze replay*, dont le fonctionnement lui était systématiquement montré en début d’entretien. De plus, lorsqu’il parlait d’une visualisation, elle apparaissait directement sur le grand écran.

3. Matériel et logistique

Pour pouvoir mener les entretiens semi-directifs, nous nous rendions dans la salle de débriefing avec le participant (voir Figure 26). Dans cette pièce, nous disposions :

- D'un écran sur lequel le *gaze replay* ainsi que les images présentées lors de l'expérimentation pouvaient être montrés au participant grâce au logiciel *TightVNC* ;
- De la tablette sur laquelle s'affichaient les appréciations déclarées ;
- Du guide d'entretien papier ;
- D'une caméra et d'un micro.

Les entretiens semi-directifs ont été enregistrés de deux manières : une caméra a enregistré l'entretien dans sa globalité tandis que l'activité sur écran a été enregistrée et fusionnée avec l'enregistrement vocal²⁵. Par précaution, nous avons ainsi deux flux de données ayant enregistré le même entretien. Sur l'un, il est possible de voir le corps et les réactions du participant et sur l'autre, il est possible de voir les visualisations, les *gaze replays* et d'entendre les interventions vocales qui s'y rapportent.

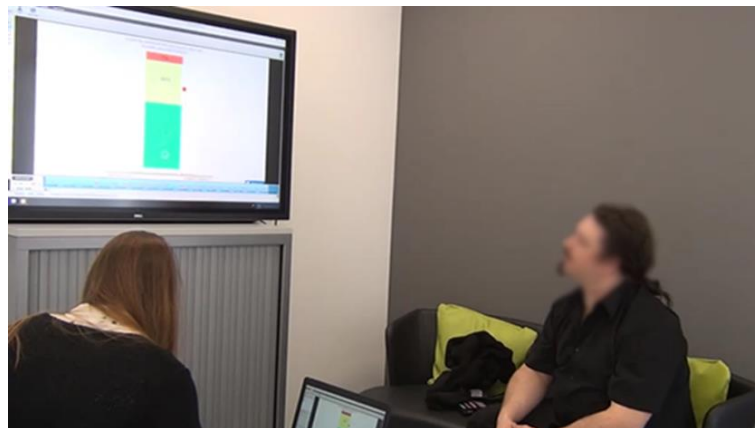


Figure 26 Prise de vue de l'entretien semi-directif avec la caméra sur pied

4. Déroulement général

Une fois le dispositif mis en place et après avoir pris connaissance des appréciations, nous invitons le participant à s'installer dans la salle de débriefing. C'est alors que commençait l'entretien semi-directif, toujours avec la même introduction et la même question : « *Quelle visualisation vous a marqué, positivement ou négativement ?* ».

Initialement, nous avions prévu de projeter le *gaze replay*. Mais les propos de Elling et ses pairs (2011) se sont confirmés : il n'y avait pas de plus-value. Les participants se

²⁵ Logiciel utilisé : OBS Studio.

souvenaient très bien de leurs ressentis et des moments de blocages qu'ils ont pu rencontrer dans le visionnage des visualisations. La projection de l'image seule, sans lancer de gaze replay, suffisait. Ainsi, après 2 expérimentations, nous avons adopté le réflexe d'expliquer aux participants ce qu'était le *gaze replay* et pourquoi il pouvait être utile. Nous leur avons proposé de demander à le voir lorsqu'ils estimaient cela nécessaire. Cela a bien fonctionné : dans les moments de doutes, les participants demandaient à le visionner pour se remémorer leur trajet oculaire. Ce changement opérationnel, réalisé au cours de la série d'expérimentations, n'est pas un problème d'un point de vue méthodologique s'il est réalisé rapidement et de manière minime, selon Romelaer (2005) : « les guides d'entretien peuvent être fixes ou évolutifs. (...) le chercheur conduit un premier entretien, qui lui amène de nouveaux éléments, et en fonction de ce nouveau « morceau de réalité », il ou elle élabore un nouveau guide (...) » (Romelaer, 2005, p.115).

Comme expliqué précédemment, l'entretien était ensuite mené en mobilisant le guide d'entretien et les appréciations déclarées sur la tablette. L'entretien se terminait par la même question pour tout le monde : « *Y a-t-il un sujet que vous souhaitez encore aborder ou quelque chose que vous souhaitez ajouter ?* », avant de finalement expliquer au participant l'objectif de la recherche. Chaque participant a été remercié avec un bon d'achat d'une valeur de 10€.

5. Découvrir Dervin sur le tard : de l'acte manqué à l'opportunité d'analyse

Malheureusement, c'est assez tard dans notre processus de recherche que nous avons découvert la « *sense-making methodology* » de Dervin et ses collègues (2003), que nous avons expliquée plus en détail dans le chapitre 3 de la première partie. Nous avons déjà récolté toutes les données nécessaires à nos travaux de recherche.

Pour rappel, « la méthodologie du *sense-making* place l'individu au centre de la construction du sens : les comportements qu'il développe sont des activités qui peuvent être analysées en termes de construction de ponts entre les interstices de la réalité. La communication est pour Dervin le processus qui permet aux individus de construire des ponts (« *gap bridging* ») par-dessus les interstices afin de faire émerger le sens » (Jacques, 2016, p. 104). De même, « la méthodologie de création de sens exige une approche très souple et contextuelle de l'utilisation de l'information. La détection et la réduction des écarts peuvent se manifester de multiples façons, et elles ne sont pas toujours délibérées, instrumentales, intentionnelles et orientées vers un but » (Savolainen, 2006, p.1123, notre traduction).

Dès lors, pendant nos expérimentations, en présentant 20 visualisations à chaque participant et en sachant que certaines étaient confortables à lire et d'autres moins, nous les avons placés dans une situation de discontinuité. Celle-là même se rapproche de la perturbation du système dont parlent Rouve et Ria dans le paradigme éenactif (2008). Ainsi, l'individu tente par ses diverses actions de revenir à l'équilibre, ce qui le pousse à faire information (formation au « dedans » et donc construction de sens). Dans le paradigme éenactif, on ne peut faire fi du contexte pour étudier un fait. Mais dans l'idée du *sense-making* de Dervin et du *gap bridging*, nous créons, de par l'exposition aux diverses visualisations, différentes situations de discontinuité qui dépendent d'un individu à l'autre. Ces situations les incitent à utiliser des souvenirs, l'expérience ou encore l'habitus pour agir (poursuivre la lecture, l'abandonner, donner une bonne ou mauvaise appréciation à la visualisation, ...). Donc, en incluant dans notre guide d'entretien des questions relatives au « *gap bridging* », nous aurions pu en savoir plus sur la construction de sens des utilisateurs. Sow et Chaudiron, par exemple, utilisent cette méthode dans le contexte d'une recherche sur les pratiques communicationnelles des chercheurs sénégalais (2018). Dans leur méthodologie, ils mettent eux aussi les participants à leurs entretiens semi-directifs dans une situation de discontinuité avant de les interroger. Puis, « il s'agit de s'intéresser donc au temps, à l'espace, au mouvement, à la situation, à la discontinuité (gap), au « bridge », aux résultats (*outcomes*). Trois éléments fondamentaux doivent être retenus selon Dervin. Ceux-ci pour répondre aux questions « (1) qu'est-ce qui vous a arrêté dans votre situation ? Que vous a-t-il manqué dans votre situation ? (2) Quelles questions vous vous posez, quelles confusions, difficultés avez-vous ? (3) Quel type d'aide espérez-vous obtenir ? ». (...) Nous nous focalisons en tant qu'enquêteurs sur le « gap » : la discontinuité propre à la situation et sur les stratégies développées par l'acteur pour obtenir des résultats. » (Sow, Chaudiron, 2018, p. 7).

C'est d'un acte manqué dont il s'agit ici : nous aurions pu nourrir davantage notre guide d'entretien pour explorer le *gap bridging* selon Dervin. Malheureusement, remonter le temps est impossible. En revanche, ce qui ne l'est pas, c'est de chercher, dans nos données, les éléments significatifs quant au *gap bridging* que nos participants ont mis en place. Dans [la Partie IV](#), relative à l'analyse des données qualitatives, nous avons donc identifié dans nos données les stratégies employées par les participants pour réaliser le pontage de la construction de sens. Grâce à cela, notre analyse n'en est que plus nourrie et s'inspire d'une méthodologie connue et validée et adaptée à notre situation.

Conclusion Justesse méthodologique et positionnement épistémologique

Pour conclure, nous expliquons dans un premier temps pourquoi la méthodologie mise en place dans le cadre de cette recherche est adaptée à la récolte de données avant d’aborder, dans un second temps, une réflexion épistémologique sur cette même méthodologie.

1. *Eye tracking* et entretien semi-directif : pourquoi une telle méthode est-elle adaptée ?

Pour rappel, notre méthode de récolte de données se déroule en deux phases : une phase de récolte de deux types de données quantitatives, grâce à l’*eye tracker*, et une phase de récoltes de données qualitatives, grâce à la conduction de l’entretien semi-directif. L’intérêt et l’originalité de la méthode résident dans le fait que ces deux phases ne se produisent pas de façon cloisonnée : elles sont fortement liées grâce à la récolte et à l’exploitation des appréciations déclarées. Nous avons finalement récolté trois types de données différentes :

- (1) Des *gaze data*, données fournies par l’*eye tracker* et relatives au regard des participants ;
- (2) Des appréciations déclarées constituées en échelles de Likert ;
- (3) Des données qualitatives, verbales principalement, que nous avons pris soin de retranscrire mot pour mot (environ 230 pages)

Ces données sont donc nombreuses, variées et surtout complémentaires. Se cantonner aux professionnels de la communication numérique nous semble adéquat d’un point de vue de faisabilité, mais aussi parce qu’ils sont sensibilisés à la transmission d’informations dans le domaine du numérique.

Cette méthode est adaptée à notre problématique parce qu’elle correspond presque parfaitement à la manière dont nous conceptualisons la construction de sens, qui, rappelons-le, est un processus s’enrichissant de diverses composantes que sont la perception visuelle, la compréhension et la subjectivité de la construction de sens.

Rappelons-nous de notre question de recherche :

Par rapport à une visualisation de données qui est correcte en regard des principes de conception (que nous appelons « visualisation standard »), qu’est-ce qu’une visualisation de données ornementée induit sur la construction de sens

- (1) au niveau de la perception visuelle et de la compréhension ?
- (2) au niveau du *gap bridging* ?

La récolte des *gaze data* nous permet aisément de commenter l'influence de l'ornementation sur la perception visuelle, puisque nous pouvons ainsi retracer et comparer chaque trajet oculaire en fonction des caractéristiques des visualisations. De nombreuses métriques oculométriques sont également exploitables (cf. *infra*).

Les appréciations déclarées peuvent être un indicateur au niveau de la compréhension : l'individu pense-t-il avoir compris, la visualisation lui semblait-elle claire ? Il est ainsi aisé de comparer le facteur « beauté²⁶ » (donc « ornementation » dans notre langage scientifique) avec d'autres indicateurs qui montrent ce que l'individu rapporte de sa propre cognition (compris, clair) et qui peut également être influencée par d'autres éléments (intérêt).

Les données qualitatives, quant à elles, permettent d'identifier les facteurs humains et émotionnels de Kennedy et Hill (2017, 2018) mais aussi d'identifier, dans une analyse thématique des données textuelles, ce qui a poussé un individu à adopter un comportement et à identifier son ressenti par rapport à l'apparence visuelle des visualisations. Ces données ouvrent ainsi une porte vers la subjectivité des utilisateurs. En mentionnant leurs difficultés, les individus ont eux-mêmes identifié des moments de discontinuité selon Dervin (Dervin, 1992 ; Dervin et al., 2003) et ont donc bien souvent expliqué leur stratégie de pontage (« *gap bridging* »). Ces données offrent donc de grandes potentialités d'analyse.

De plus, afin de correspondre à l'objectif de présentation d'une visualisation de données et en phase avec les travaux de nos prédécesseurs Bertin, Tufte, Dervin, Kennedy et autres, nous souhaitons découvrir quelles recommandations en termes d'ornementation préconiser pour que toutes les composantes de la construction de sens soient stimulées et ce, afin que le message soit perçu de manière fidèle, tout en donnant une plus-value à la visualisation de données. Ainsi, l'analyse et l'interprétation des données dans le cadre de la question de recherche nous permettra d'émettre ce type de recommandations. Pour y arriver, il est nécessaire de croiser les analyses pour produire des interprétations supplémentaires : par exemple, comment les données qualitatives peuvent-elles expliquer les différents ou trajets oculaires détectés ? C'est ainsi qu'on pourra comprendre l'apport des données qualitatives qui n'expliquent pas seulement les raisons d'arrêt du regard, mais qui les mettent également en perspective par rapport au vécu des individus.

Nos choix auraient pu être différents. Premièrement, notre recherche ne tient pas compte du contexte de réception de la visualisation de données, puisqu'elle est entièrement menée en laboratoire. Cela aurait pu être imaginé. Ensuite, nous assumons de ne nous intéresser

²⁶ Nous avons choisi de parler de « beau » auprès des participants par souci d'accessibilité et de vulgarisation.

qu'à la réception de la visualisation de données. Nous prenons en compte les expériences, émotions et autres facteurs des récepteurs. Néanmoins, nous ne nous intéressons pas – ou très peu – aux logiques de conception des visualisations de données. Nous ne nous intéressons pas aux concepteurs : ils ne sont nullement parties prenantes des expérimentations. Pourtant, on peut s'en douter, la conception d'une visualisation de données peut parfois être influencée par le besoin d'intéresser le lecteur, des logiques économiques dans le cadre des médias, ou encore par une méconnaissance des principes de conception des visualisations de données. Bien qu'intéressants, ces éléments pourraient s'intégrer dans le cadre d'une autre recherche sur la conception de visualisation de données. Pour ces expérimentations, nous partons donc du principe que l'ornementation est une façon de ne pas respecter les principes de conception des visualisations de données.

Pour terminer, bien que notre méthodologie soit complexe, elle débouche sur l'obtention de trois jeux de données différentes, de différente nature également. Ainsi, nos travaux de recherches peuvent s'apparenter à un voyage dont les différentes étapes s'apparentent à la création de méthodes. Pour arriver à destination et aux conclusions finales, il est donc essentiel de développer trois méthodes d'analyse de données, propres aux différents types de données obtenues : les *gaze data*, les appréciations déclarées et les retranscriptions textuelles des entretiens. Ces différentes méthodes d'analyse seront présentées avant chaque partie de cette thèse relative aux différentes données récoltées.

2. Réflexion épistémologique sur la méthodologie

Jusqu'ici, nous avons déjà abordé les réflexions épistémologiques qui nous animent quant à l'état de l'art établi dans le cadre de cette recherche. Nous nous inscrivons pour rappel dans le paradigme constructiviste et nous adoptons une position interprétative. Néanmoins, du point de vue du chercheur, il n'est pas toujours aisé d'assumer complètement cette posture. Le bricolage que nous avons réalisé pour développer notre propre méthodologie, ainsi que les diverses méthodes de récolte de données en soi, ont souvent semé le doute sur notre positionnement épistémologique. Finalement, après avoir développé une telle méthodologie, peut-il rester inchangé ?

Il est assez peu commun, dans les sciences humaines et dans des disciplines telles que la communication, de réaliser des expérimentations nécessitant un *eye tracker*. En vérité, le set de données récoltées est lui-même varié et complexe. De manière générale en sciences, les expérimentations visent à « expliquer », grâce à, notamment, des mesures quantitatives, ce qui est plutôt de l'ordre du post-positivisme. Dans notre cas, nous assumons pleinement

que nos expérimentations visent à comprendre et non à expliquer, le tout dans une démarche qui se veut interprétative. Il est ainsi de notre ressort d'éviter d'employer un vocabulaire trop déterministe.

Chaque partie de l'expérimentation et chaque récolte de données a été conçue en fonction de la façon dont nous considérons la construction de sens. L'*eye tracking* relève ainsi des mesures claires, précises et quantitatives relatives à la perception visuelle : à tel moment, l'œil s'est posé à tel endroit. Ce type d'information est donc très factuel, même si les raisons sous-jacentes au déplacement du regard ne le sont pas.

Pour continuer, il convient de garder en tête que nous sommes le sujet interprétant de toutes les informations exprimées par les participants à travers les données. Ainsi, malgré l'installation de ce dispositif technique, c'est le paradigme dans lequel nous nous inscrivons qui nous permet d'avancer vers une compréhension profonde et interprétative du phénomène. La raison est simple : nous évitons tout simplement de basculer dans une explication absolue afin de tendre vers l'interprétation et la compréhension du phénomène étudié, sans prétention explicative.

Dans la visée constructiviste que nous avons précédemment exposée, nous essayons d'interpréter la réalité perçue des utilisateurs de visualisation de données. Cela permet de comprendre le comportement social des individus en contact avec les visualisations de données – ou, du moins, de l'interpréter (Bertacchini, 2009). De même, au-delà de notre position interprétative, nous nous situons dans une démarche d'abduction, selon laquelle « on ne peut pas affirmer avec certitude qu'une explication constitue la cause réelle d'une observation, l'incertitude pouvant porter sur la plausibilité de l'explication, ou bien concerner la validité de la connaissance permettant l'explication » (Catellin, 2004, p.180). Il n'y a donc là pas de pouvoir prédictif. C'est bien normal, selon nous, dans le cas d'une étude oculométrique visant à découvrir ce que regardera l'utilisateur, sans prévoir de tâche à l'avance (cf. *infra*). L'abduction, qui consiste à restreindre ses idées au plus probable sans les juger certaines, est ce qui nous a mené à penser diverses méthodes d'analyse de données.

Partie III. Les appréciations déclarées

Analyse et interprétations

Chapitre 1 Méthode d'analyse des appréciations déclarées

Avant de plonger dans l'analyse et la présentation des résultats propres aux appréciations déclarées, nous tenons à décrire la méthode d'analyse employée. Pour commencer, nous ferons état de la façon dont se présentent les appréciations déclarées. Ensuite, nous énumérerons une série d'attributs des visualisations de données en fonction desquels nous avons mené l'analyse. En effet, les appréciations déclarées diffèrent en fonction de certaines caractéristiques des visualisations présentées aux participants. Pour terminer, une méthode d'analyse statistique a été élaborée afin d'explorer les appréciations déclarées en fonction des différents attributs des visualisations de données. L'explication de cette méthode clôt ce chapitre.

1. Présentation des données mobilisées dans l'analyse

L'analyse statistique dont cette partie fait l'objet a été réalisée à partir des appréciations déclarées, qui sont des données ordinales, couplées aux attributs des visualisations de données faisant partie du corpus, qui sont des données nominales. Successivement, nous présentons dans cette section ces deux types de données.

1.1. Présentation des appréciations déclarées

Durant l'exposition aux visualisations de données dans les expérimentations menées, nous avons pour rappel récolté des appréciations déclarées grâce à des échelles de Likert. Sur une échelle de 1 à 5 (de « pas du tout d'accord » à « complètement d'accord », en passant par « neutre »), les participants ont tous pu donner leur appréciation à propos des *items* « Belle », « Intéressante », « Claire », « Compris », « J'aime » (voir Figure 27).

J'ai trouvé cette visualisation...				
Pas belle	1	2	3	4 5
Pas intéressante	1	2	3	4 5
Pas claire	1	2	3	4 5
Je n'ai pas compris cette visualisation	1	2	3	4 5
Je n'ai pas aimé cette visualisation	1	2	3	4 5
1 = Pas du tout d'accord	2 = Pas d'accord	3 = Neutre	4 = Plutôt d'accord	5 = Complètement d'accord

Figure 27 Appréciations déclarées utilisées dans les expérimentations

Puisque nous avons 40 participants qui ont chacun pu donner leurs appréciations à propos de 20 visualisations, nous aurions logiquement pu compter sur 800 appréciations par *item*. Malheureusement, nous ne pouvons compter que sur 741 appréciations par *item* suite à deux contraintes techniques. La première est qu'une série de données n'a pu être récoltée pour une participante, nous faisant ainsi perdre 20 appréciations par *item*. La seconde est que nous avons constaté plusieurs semaines après la conduite des expérimentations que la visualisation n°40 (standard) du corpus ne correspondait pas totalement à la visualisation n°20 (ornementée). Il y manque en effet deux bâtonnets, ce qui simplifie le jeu de données et n'en fait plus une parfaite homologue. De fait, nous avons choisi de soustraire ces visualisations de l'analyse. Il nous reste ainsi 741 appréciations déclarées par *item*, que nous pouvons ainsi considérer comme données fiables (Tableau 7)

Tableau 7 Données disponibles pour l'analyse, par item

Self Reported Appreciations Data						
		Belle	Intéressante	Claire	Compris	J'aime
N	Valid	741	741	741	741	741
	Missing	0	0	0	0	0

Pour rappel, l'élément variable central de notre corpus est que les visualisations présentées aux participants sont de type standard (nommé « *Raw* » dans l'analyse statistique *infra*) et de type ornementé (nommé « *Junk* » dans l'analyse statistique *infra*). Les participants ayant vu les versions standard des visualisations n'ont pas visionné la visualisation ornementée, dans un souci d'évitement des effets d'apprentissage. Puisque nous avons utilisé la méthode de randomisation en carré latin pour éviter les effets de séquence, nous ne pouvons pas dire que notre échantillon de participants se subdivise en deux sous-populations, l'une ayant testé les visualisations standard et l'autre ayant testé les visualisations ornementées. Par contre, nous pouvons affirmer que nous détenons deux groupes d'appréciations relatives aux visualisations de type standard (*Raw*) et aux visualisations de type ornementé (*Junk*). On dispose ainsi de 370 appréciations pour le type « *Raw* » et 371 appréciations pour le type « *Junk* » (Tableau 8). Les données issues des appréciations déclarées sont des données de type ordinal.

Tableau 8 Nombre d'appréciations récoltées et analysables par item

Data frequencies by type (Raw, Junk)			Belle	Intéressante	Claire	Compris	J'aime
Type							
Raw	N	Valid	370	370	370	370	370
		Missing	0	0	0	0	0
Junk	N	Valid	371	371	371	371	371
		Missing	0	0	0	0	0

De plus, ces appréciations peuvent être grandement enrichies par des observations quant aux éléments du corpus. Ainsi, l'analyse ne sera pas menée sur base de la population mais sur base des objets qui leur ont été présentés. L'analyse se centre donc sur les caractéristiques des visualisations de données présentées, à savoir les attributs que nous détaillerons ci-après, et non sur les caractéristiques des participants. Bien que nous concentrons sur des appréciations déclaratives et donc personnelles, nous souhaitons dès lors étudier l'impact de l'objet (donc des visualisations présentées) sur ces déclarations transcrites en appréciations déclarées. C'est pourquoi l'analyse s'enrichit de caractéristiques que nous avons attribuées a posteriori aux visualisations faisant partie du corpus test. Nous appelons ces caractéristiques « attributs des visualisations de données ». Notre objectif, en utilisant ces attributs à des fins d'analyse, est d'explorer les données via une granularité de plus en plus précise.

1.2. Présentation des attributs des visualisations de données du corpus

Nous allons maintenant aborder les attributs qui caractérisent les visualisations de données faisant l'objet de notre corpus. Effectivement, à la genèse de nos travaux de recherche, la phase d'expérimentation était prévue comme exploratoire. Au vu de la taille et de la variété des données disponibles après ce cycle d'expérimentations, nous avons réalisé qu'il s'agirait tout simplement de la seule récolte de données dans le cadre de cette thèse de doctorat. De nombreuses hypothèses à falsifier pouvaient effectivement être émises sur cette base. Le corpus test, pour la partie ornementée, se constituait donc de visualisations différentes, récoltées dans les médias et à l'intention esthétique évidente mais néanmoins, ces visualisations n'ont pas été sélectionnées en fonction de leurs attributs dont nous faisons état ici. Nous avons donc caractérisé le corpus a posteriori. C'est pourquoi chaque attribut n'apparaît pas forcément à fréquence égale dans l'ensemble du corpus. De plus, puisque c'est l'ornementation que nous étudions, il va sans dire que les visualisations de type standard (*Raw*) présentent dans la plupart des cas une absence de l'attribut évoqué. Ainsi, il est parfois pertinent dans l'analyse de ne concentrer les hypothèses que sur les appréciations des visualisations de type ornementé. Nous présenterons ces analyses restreintes au moment opportun. Les attributs présentés ci-dessous concernent l'ensemble

du corpus et visent à étudier les spécificités de l'ornementation en explorant différents niveaux de granularité.

- (1) Le **type** : standard ou ornementé, le type est le premier attribut et le plus important de notre corpus. Nous l'avons présenté précédemment.
- (2) L'**ornement** : caractérise l'élément dans lequel l'ornementation se manifeste, à savoir un/ des icône(s) ou une forme inhabituelle de graphe. La plupart des visualisations ornementées choisies présentent des icônes au sens de Peirce (1978). L'icône est un signe présentant une analogie avec l'objet qu'il dénote. L'icône peut, sur ces visualisations, être utilisé pour représenter la donnée (par exemple, une icône représentant une bouteille de champagne qui prend place au lieu d'un simple bâtonnet de diagramme) ou pour l'illustrer, apposé à côté du graphique (c'est le type de figuration, qui constitue le 4^{ème} attribut). Nous avons également repéré des symboles (relation arbitraire entre le signe et l'objet qu'il représente) au sens de Peirce dans ces visualisations mais dans une moindre mesure et toujours en présence d'une icône. Nous avons décidé de ne pas en tenir compte dans nos attributs. Par contre, la présence d'ornementation se manifeste parfois différemment, comme une forme de graphe inhabituelle et fantaisiste (par exemple la visualisation n°10 du corpus). Nous avons donc 314 appréciations concernant des visualisations contenant des icônes et 57 appréciations à propos de visualisations à la forme inhabituelle.

Tableau 9 Fréquences des données pour l'attribut "Ornement"

	Frequency	Percent	Number of visualizations
Aucun	370	49.9	19
Icônes	314	42.4	16
Formes inhabituelles	57	7.7	3
Total	741	100.0	38

- (3) La **figuration** : concerne l'endroit où l'ornement figure dans la visualisation. « Apposé » signifie que l'ornement se trouve juxtaposé à la représentation des données. « Sur le *Data-Ink* » signifie que l'ornement constitue la représentation des données. Cet attribut relève de notre propre classification réalisée dans la revue de littérature sur l'ornementation.

Tableau 10 Fréquences des données pour l'attribut "Figuration"

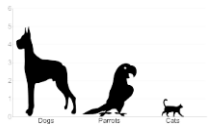
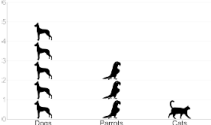
	Frequency	Percent	Number of visualizations
Absence	370	49.9	19
Apposé	118	15.9	6
Sur le data ink	253	34.2	13
Total	741	100.0	38

- (4) Le **pictogramme** : se base sur la caractérisation du pictogramme en visualisation de données selon (Haroz et al., 2015), présenté dans le chapitre de revue de littérature sur l'ornementation. Nous y avons ajouté l'attribut « Simple », qui signifie que le pictogramme ou l'image représente la donnée sans être multiplié (*stacked*) ou sans que sa taille ne soit modifiée (*stretched*). Le pictogramme concerne forcément des visualisations de données ornementées. Nous n'avons donc pas catégorisé les visualisations de données standard en employant cet attribut. De même, les visualisations de données à la forme inhabituelle (attribut : ornement) ne comportent pas de pictogramme. Ainsi, le pictogramme est une façon plus détaillée de considérer l'icône.

Tableau 11 Fréquences des données pour l'attribut "Pictogramme"

	Frequency	Percent	Nb. of visualizations
Stretched	99	13.4	5
Stacked	19	2.6	1
Superfluous	97	13.1	5
Simple	39	5.3	2
Background	0	0	0
Total	254	34.4	13

Rappel des pictogrammes de Haroz et ses collègues (Tableau 3 Classification des pictogrammes dans (Haroz et al., 2015))

Type de visuel	Haroz et al., 2015
	Stretched
	Stacked
	Background
Pictogramme présent mais complètement séparé des données	Superfluous

- (5) Le **lie factor** : cet attribut indique s'il est distorsif ou non. Pour rappel, le *lie factor* est un coefficient qui décrit la variation entre l'effet montré par un graphique et l'effet produit par les données (Tufté, 1983). Un graphique au *lie factor* distorsif montre une distorsion dans les données, ce qui causerait une distorsion dans son

interprétation et sa compréhension selon Tufte (1983). Ainsi, par « Fidèle », nous entendons que le *lie factor* est correct en regard de l'intégrité graphique. Par « Distorsif », nous entendons qu'il y a distorsion dans le graphique.

Tableau 12 Fréquences des données pour l'attribut "Lie Factor"

	Frequency	Percent	Number of visualizations
Distorsif	254	34.3	13
Fidèle	487	65.7	25
Total	741	100.0	38

- (6) Le ***data-ink*** : calcul de la proportion d'encre ou de pixels qui se rapporte directement aux données, selon Tufte (1983). Les attributs sont différenciés en 3 classes de proportions : bas (entre 0 et 50%), moyen (entre 50 et 75%) et élevé (entre 75 et 100%).

Tableau 13 Fréquences des données pour l'attribut "Data-Ink"

	Frequency	Percent	Number of visualizations
Bas	38	5.1	2
Moyen	98	13.2	5
Élevé	605	81.6	31
Total	741	100.0	38

- (7) Le ***data-ink ornementé*** ou *embellished data-ink* : cet attribut montre la proportion de *data-ink* qui est ornementée sur le graphe. La différenciation des 3 classes est identique à la précédente.

Tableau 14 Fréquences des données pour l'attribut "Embellished Data-Ink"

	Frequency	Percent	Number of visualizations
Absence	524	70.7	27
Bas	79	10.7	4
Moyen	80	10.8	4
Élevé	605	81.6	3
Total	741	100.0	38

Le type et le *lie factor* sont les seuls attributs qui constituent des variables de type dichotomique. Ainsi, la variable « Type » contient les possibilités « Raw » et « Junk » (0, 1), tandis que la variable « Lie Factor » contient les possibilités « Distorsif » et « Fidèle » (0, 1). Les autres variables sont catégorielles et laissent donc plusieurs possibilités d'attribution aux visualisations de données. La caractéristique dichotomique ou non de ces

variables tient son importance dans l'analyse. Les données que constituent les attributs sont des données nominales.

2. Méthode d'analyse des appréciations déclarées et des attributs

L'objectif de cette analyse quantitative est de comprendre l'effet de l'ornementation sur les différentes appréciations déclarées des individus et ce, avec des niveaux de granularité différents au sein-même de l'ornementation. Afin de mener cette analyse, nous employons le logiciel d'analyse statistique SPSS. Avant toute chose, ce logiciel nous a permis de générer différents *box plots* représentant la distribution et la répartition des données et ce, pour chaque attribut des visualisations de données. Ces *box plots* nous ont permis de générer différentes hypothèses que nous avons testé par la suite. Comme nous le verrons, les distributions des données pour la quasi-totalité des groupes contenus dans les attributs ne semblent pas suivre une distribution normale. Il est donc nécessaire d'employer des tests statistiques de type non paramétrique pour réaliser l'analyse.

Pour tenter de répondre aux hypothèses relevant de données dichotomiques, c'est-à-dire dont les attributs étaient à deux possibilités (deux variables), nous nous dirigeons initialement vers le test non paramétrique de Mann Whitney U. Néanmoins, en prenant garde aux conditions d'application de ce test statistique, nous avons constaté que le test assume l'homogénéité des variances des distributions qu'il compare. Dans les faits, une telle chose est rarement vérifiée. Mais afin de mener notre analyse avec le plus de justesse possible, nous souhaitons respecter cette condition, dont l'importance a été soulevée par Maxwell & Delaney (1990). Le test de Mann Whitney permet d'identifier si les moyennes des variables représentées sont égales ou pas, tout en assumant une distribution similaire entre les deux groupes. Une façon de vérifier si les distributions sont semblables d'un groupe à l'autre consiste à effectuer un test d'homogénéité des variances, car l'observation de différentes variances statistiquement significatives implique nécessairement différentes formes de distribution (Maxwell & Delaney, 1990). Ainsi, si l'homogénéité des variances est à rejeter, les résultats du test de Mann Whitney ne peuvent être interprétés avec précision.

De même, nous avons souhaité analyser les données en fonction des attributs catégoriels (non dichotomiques) de visualisations de données, notamment en tentant d'évaluer l'effet des attributs sur les appréciations. Le test non paramétrique de Kruskal Wallis est similaire à celui de Mann Whitney et s'applique lorsque l'on souhaite comparer plus de deux variables. Ce test compare les moyennes des rangs des données et permet de rejeter ou non l'hypothèse selon laquelle les moyennes de ces rangs sont égales. Ses conditions

d'application sont identiques à celles de Mann Whitney. Ainsi, en ayant des distributions non normales et une homogénéité des variances non respectée, les résultats de ce test ne sont pas précis eux non plus. Vu la structure-même de nos données, cela ne nous convient pas.

En effet, on pense souvent à tort que les statistiques non paramétriques comme le test de Mann Whitney et le Kruskal Wallis n'ont pas d'hypothèses de départ. Par ailleurs, certains statisticiens pensent que mobiliser ces tests n'est pas un problème si l'on tolère le non-respect de l'homogénéité des variances à un faible niveau. En somme, si les variances ne sont pas homogènes mais que leur score est assez proche de celui qui permettrait d'assumer l'homogénéité des variances, ces auteurs admettent une certaine tolérance (Howell et al., 2008).

Ce n'est néanmoins pas notre cas. Notre design expérimental n'a pas été initialement pensé pour réaliser une analyse statistique. En prenant connaissance des conditions d'application des tests qui pourraient nous servir, nous avons dès lors été dubitative : souvent, les distributions de nos données, en fonction des attributs, ne sont pas normales. La condition d'homogénéité des variances est également rarement respectée. Nous nous sommes ainsi penchée sur les travaux de Gignac (2019), titulaire d'un doctorat en Psychologie, particulièrement intéressé par les méthodes quantitatives qu'il enseigne. Nous nous sommes donc appuyée sur son matériel pédagogique afin de réaliser l'analyse. Nous avons souhaité que l'analyse statistiques soit menée de la manière la plus rigoureuse possible. Ainsi, cette analyse peut sembler fort orthodoxe et peu commune mais nous désirons que le résultat soit le plus respectueux possible des conditions statistiques pour les tests employés. Venant du champ des méthodes qualitatives, il nous est précieux d'accorder du temps et de la rigueur à chaque détail, ce qui transparaît désormais dans notre démarche quantitative.

Il est ainsi possible, selon Gignac (2019), d'adopter une alternative au test de Kruskal Wallis en mobilisant le test de Brown Forsythe. Il est paramétrique mais permet d'assumer que les distributions ne soient pas normales et que l'homogénéité des variances ne soient pas respectées. Attention néanmoins : il s'agit d'un test de comparaison de moyennes. Ainsi, pour pouvoir comparer ces données qui sont de type ordinal, nous choisissons d'en comparer les moyennes des rangs. Le test de Brown Forsythe est donc utilisé dans notre analyse pour comparer les moyennes des rangs. En somme, il s'agit d'utiliser un outil paramétrique pour tendre vers le non paramétrique. Nous comparerons ensuite les

différentes variables entre elles, avant d'appliquer une correction de Bonferroni afin de choisir à quel seuil de signification rejeter nos hypothèses²⁷.

A notre sens, le simple test de Brown Forsythe sur les rangs peut également s'appliquer en alternative du test de Mann Whitney, puisqu'il s'agit du même principe²⁸. De plus, dans le cadre d'une corrélation inter-échelles, nous avons utilisé le test Kendall Tau B, adapté aux données ordinales.

Chapitre 2 Analyse statistique des appréciations déclarées par attribut

Ce chapitre présente l'analyse séquentielle par attribut des visualisations de données. Chaque section illustre la distribution des données qui permet de formuler les hypothèses, suivie des tests statistiques et de la falsification des hypothèses. Nous avons prêté attention aux distributions et à l'homogénéité des variances pour tous les cas traités dans cette analyse. Nous avons donc toujours vérifié si les distributions étaient normales et si l'homogénéité des variances étaient respectées et ce, par attribut. Aucune distribution n'est normale dans l'ensemble de ce jeu de données structuré par type. De temps à autres, l'homogénéité des variances était respectée, mais jamais pour tous les *items* au sein d'un même attribut. Pour plus de cohérence, nous avons choisi de toujours mener le test de Brown Forsythe sur les rangs des variables des attributs étudiés. Le test de Brown Forsythe a toujours été réalisé sur les rangs des *items*.

1. L'attribut principal : le type (« *Raw* » vs. « *Junk* »)

1.1. Statistiques descriptives et génération d'hypothèses

Pour commencer, il convient de faire le point sur les statistiques descriptives à notre disposition. Nous pouvons donc montrer et comparer les statistiques descriptives des visualisations de type standard (*Raw*) versus ornementées (*Junk*), puisqu'il s'agit de l'attribut le plus important de notre corpus. Ces statistiques sont présentées dans le Tableau

²⁷ Effectivement, il faut être attentif avant d'interpréter les résultats des tests : en multipliant les comparaisons, la probabilité qu'un événement rare se produise augmente et, avec lui, la probabilité de rejeter incorrectement H_0 , ce qui constitue une grosse erreur (Gignac, 2019). La correction de Bonferroni permet de compenser ce risque en calculant un nouveau seuil de signification (qui n'est donc plus 0,05). Le calcul est simple : il suffit de diviser le seuil alpha par le nombre de comparaisons effectuées. On obtient ainsi un nouveau seuil auquel comparer la *p-value* afin d'interpréter les résultats.

²⁸ Le test de Welch aurait pu être employé également, en étant utilisé similairement au test de Mann Whitney, plus adapté aux données de type dichotomique. Le test de Brown Forsythe reste néanmoins acceptable ; les résultats ne sont nullement affectés. Il a été utilisé dans une logique d'automatisation de l'analyse.

15. Puisque nous travaillons avec des échelles de Likert et donc des variables ordinales, nous éviterons d'exploiter directement la moyenne de ces variables. S'intéresser à la médiane et au mode de ces variables est dès lors plus pertinent. Dans le Tableau 15, on constate d'ailleurs que les valeurs sont assez similaires. Pour rappel, les codes se rapportent aux appréciations suivantes : (1) pas du tout d'accord, (2) pas d'accord, (3) neutre, (4) plutôt d'accord, (5) complètement d'accord. La médiane et le mode sont donc très élevés pour les *items* « Belle » et « J'aime » tandis qu'ils restent similaires pour les autres *items*. De même, pour les visualisations de type « Junk », les appréciations sont relativement élevées ; pour les visualisations de type « Raw », les *items* « Belle » et « J'aime » ont récolté des appréciations plus basses que les trois autres *items*. Ainsi, les *items* « Belle » et « J'aime » sont les seuls à engendrer des différences importantes entre les types « Raw » et « Junk ».

Tableau 15 Statistiques descriptives des appréciations par groupes de visualisations (Raw, Junk)

		Case Summaries				
Type		Belle	Intéressante	Claire	Compris	J'aime
Raw	N	370	370	370	370	370
	Median	2.00	4.00	4.00	5.00	2.00
	Mode	2	4	5	5	3
Junk	N	371	371	371	371	371
	Median	4.00	4.00	4.00	5.00	4.00
	Mode	4	4	5	5	4

Il convient dès lors d'observer la distribution de ces *items* et de les commenter grâce à des visualisations de la distribution en boîtes à moustaches (ou *box plots*). Ainsi, la Figure 28 montre la distribution des appréciations en fonction des différents *items* pour les visualisations de type « Raw » et la Figure 29 en montre tout autant pour les visualisations de type « Junk ».

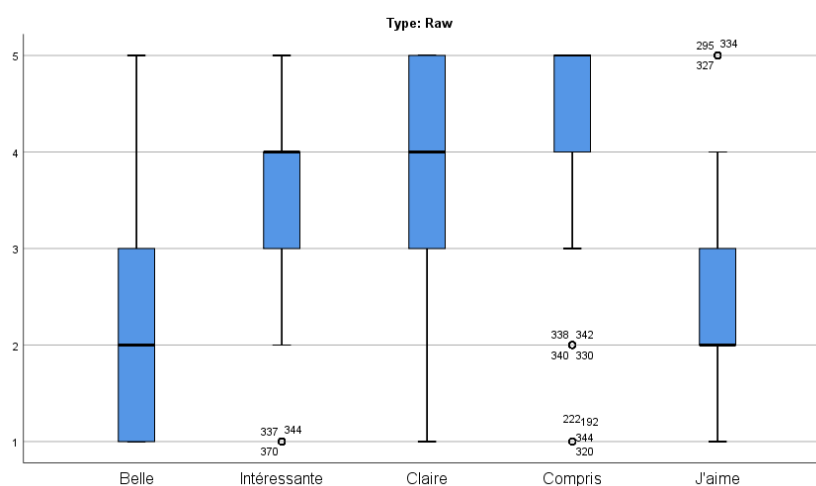


Figure 28 Box plot : distribution des appréciations déclarées en fonction des items pour les visualisations de type « Raw »

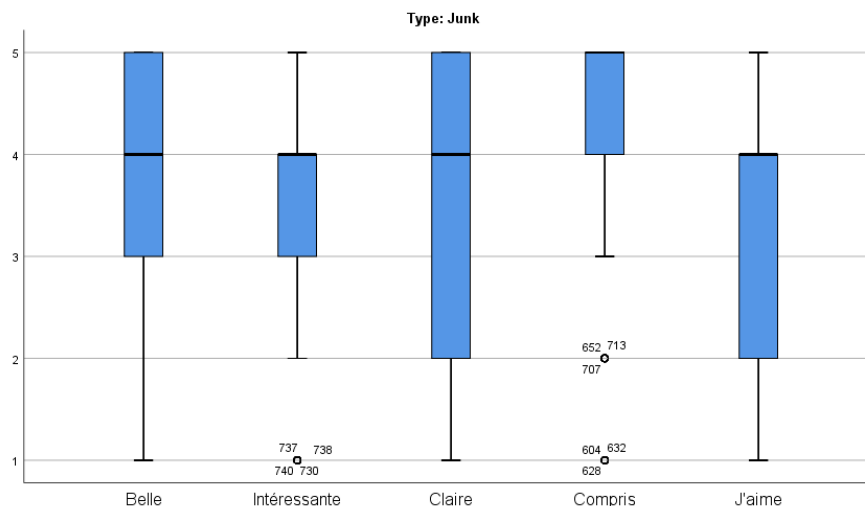


Figure 29 Box plot : distribution des appréciations déclarées en fonction des items pour les visualisations de type « Junk »

Il est aisé de constater que les visualisations de données de type « *Junk* » fournissent des réponses dispersées autour d'appréciations plus élevées que les visualisations de type « *Raw* » en ce qui concerne les items « Belle » et « J'aime ». Pour l'item « Claire » par contre, les données sont davantage dispersées vers des appréciations négatives en ce qui concerne les visualisations de type « *Junk* ». Ce n'est pas le cas pour les visualisations de type « *Raw* », dont la dispersion évolue davantage vers les valeurs positives de l'échelle (même si les deux médianes sont de 4). On pourrait ainsi penser que l'ornementation influencerait négativement l'appréciation de clarté et positivement le jugement esthétique et l'appréciation personnelle. Nous pourrions ainsi constituer les premières hypothèses à tester statistiquement.

Le plus surprenant à l'observation de ces *box plots* est que peu de différences sont à constater sur les autres échelles (« Intéressante » et « Compris »). Nous pourrions ainsi constituer une hypothèse qui concerne l'absence d'effet de l'ornementation sur ces items.

De plus, la plupart des boîtes représentées sont asymétriques (la médiane ne se situe pas au milieu). Pour ces boîtes, cela indique donc une asymétrie de la distribution des données mesurées, ce qui laisse penser que la distribution de toutes ces données ne suit pas une loi normale. Les boîtes symétriques sont celles des items « Belle » et « Claire » pour les visualisations de type « *Junk* » et la boîte de l'item « Belle » pour les visualisations de type « *Raw* ». Les valeurs identifiées autour des points dans les deux graphiques représentent des valeurs aberrantes.

À partir de ces statistiques descriptives, nous sommes donc en mesure d'émettre nos premières **hypothèses** qui sont donc des hypothèses générales :

- H1** L'appréciation de l'*item* « Claire » est significativement plus basse pour les visualisations de type « *Junk* » que pour les visualisations de type « *Raw* ».
- H2** L'appréciation de l'*item* « Belle » est significativement plus haute pour les visualisations de type « *Junk* » que pour les visualisations de type « *Raw* ».
- H3** L'appréciation de l'*item* « J'aime » est significativement plus haute pour les visualisations de type « *Junk* » que pour les visualisations de type « *Raw* ».
- H4** L'attribut « Type » n'a pas d'effet sur l'*item* « Intéressante ».
- H5** L'attribut « Type » n'a pas d'effet sur l'*item* « Compris ».

Ensuite, indépendamment des statistiques descriptives, nous inférons les hypothèses suivantes.

- H6** Les appréciations relatives à l'*item* « Belle » sont fortement corrélées aux appréciations relatives à l'*item* « J'aime ».
- H7** Les appréciations relatives à l'*item* « Belle » sont fortement corrélées aux appréciations relatives à l'*item* « Claire ».
- H8** Les appréciations relatives à l'*item* « Belle » ne sont pas corrélées aux appréciations relatives à l'*item* « Intéressante ».
- H9** Les appréciations relatives à l'*item* « Belle » ne sont pas corrélées pas aux appréciations relatives à l'*item* « Compris ».

Il sera ainsi possible de tenter de valider ces hypothèses grâce à une analyse statistique plus approfondie des données d'échelle (données ordinales) qui, sur base de l'attribut « Type » (*Raw* vs. *Junk*), fournissent une quantité d'information considérable. Les hypothèses H1 à H5 seront testées dans la partie [«1.2. Analyse approfondie des appréciations déclarées en fonction du type »](#), tandis que les hypothèses H6 à H9 seront testées ultérieurement.

1.2. Analyse approfondie des appréciations déclarées en fonction du type

Même si nous n'avons ici que deux variables au sein de l'attribut, nous avons appliqué le test de Brown Forsythe dont le résultat se trouve dans le Tableau 16. À ce stade, puisqu'il ne s'agit que d'une comparaison entre deux groupes (« *Raw* » vs. « *Junk* »), la correction de Bonferroni n'est pas nécessaire.

Tableau 16 Résultat du test de Brown Forsythe pour l'analyse des données en fonction de l'attribut "Type"

		Robust Tests of Equality of Means			
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	499.873	1	729.060	.000
Rank of Intéressante	Brown-Forsythe	.757	1	738.541	.385
Rank of Claire	Brown-Forsythe	6.647	1	737.775	.010
Rank of Compris	Brown-Forsythe	2.679	1	737.159	.102
Rank of J'aime	Brown-Forsythe	79.198	1	721.423	.000

a. Asymptotically F distributed.

Nous considérons ici que

H_0 = la moyenne des rangs des *items* du type « Raw » est égale à la moyenne des rangs des *items* du type « Junk »

Avec pour alternative

La moyenne des rangs des *items* du type « Raw » n'est pas égale à la moyenne des rangs des *items* du type « Junk ».

Pour ce test, notre seuil de signification s'élève à 0,05. Nous pouvons donc le comparer avec nos résultats dans le Tableau 16. Ainsi, pour les rangs de

- « Belle », « Claire » et « J'aime », H_0 est rejetée : les moyennes des rangs de ces *items* ne sont pas égales, ce qui nous pousse à valider l'hypothèse alternative. En somme, pour « Belle », « Claire » et « J'aime », il y a une différence d'appréciations induite par le type.
- « Intéressante » et « Compris » : H_0 ne peut pas être rejetée. On ne peut pas affirmer que le type « Raw » ou « Junk » engendre des différences d'appréciations en ce qui concerne ces *items*.

Nous pouvons donc interpréter les résultats grâce au Tableau 17 (qui pour rappel sont les moyennes des rangs). Dans la suite de nos analyses, nous ne présenterons pas tous les tableaux de moyennes des rangs. Ce tableau-ci nous sert en effet d'exemple, afin d'illustrer le procédé employé.

Tableau 17 Moyenne des rangs des données classées selon l'attribut « Type »

		N	Mean
Rank of Belle	Raw	370	237.98649
	Junk	371	503.65499
	Total	741	371.00000
Rank of Intéressante	Raw	370	364.35541
	Junk	371	377.62668
	Total	741	371.00000
Rank of Claire	Raw	370	390.45811
	Junk	371	351.59434
	Total	741	371.00000
Rank of Compris	Raw	370	381.87703
	Junk	371	360.15229
	Total	741	371.00000
Rank of Jaime	Raw	370	305.97568
	Junk	371	435.84906
	Total	741	371.00000

N'ayant que deux variables à comparer, nous mêlons nos commentaires à la falsification des hypothèses.

1.3. Falsification des hypothèses

H1 L'appréciation de l'*item* « Claire » est significativement plus basse pour les visualisations de type « *Junk* » que pour les visualisations de type « *Raw* ».

➔ Les distributions entre les groupes « *Raw* » et « *Junk* » ne sont pas similaires puisque la moyenne des rangs ne l'est pas non plus.

L'hypothèse **H1 est validée**.

La variable « *Junk* » tend ainsi à générer moins d'appréciations positives pour l'*item* « Claire » que la variable « *Raw* ». Les visualisations ornementées sont donc jugées moins claires que les visualisations standard.

H2 L'appréciation de l'*item* « Belle » est significativement plus haute pour les visualisations de type « *Junk* » que pour les visualisations de type « *Raw* ».

➔ Les distributions entre les variables « *Raw* » et « *Junk* » ne sont pas similaires, selon le résultat du test de Brown Forsythe : les moyennes des rangs sont différentes. L'hypothèse **H2 est validée**.

L'attribut de type « *Junk* » tend ainsi à générer plus d'appréciations positives pour l'*item* « Belle » que l'attribut de type « *Raw* ». Ainsi, les visualisations ornementées sont jugées plus belles que les visualisations standard.

H3 L'appréciation de l'*item* « J'aime » est significativement plus haute pour les visualisations de type « *Junk* » que pour les visualisations de type « *Raw* ».

➔ Les distributions et les moyennes des rangs entre les groupes « *Raw* » et « *Junk* » ne sont pas similaires. L'hypothèse **H3 est validée**.

La variable « *Junk* » tend ainsi à générer plus d'appréciations positives pour l'*item* « J'aime » que la variable « *Raw* ». Les visualisations ornementées sont donc davantage aimées par les participants que les visualisations de type standard.

H4 L'attribut « type » n'a pas d'effet sur l'*item* « Intéressante ».

➔ Effectivement, les variables « *Raw* » et « *Junk* », faisant partie de l'attribut « Type », ne montrent pas de différence significative en ce qui concerne la moyenne des rangs. L'hypothèse **H4 est validée**.

Entre les visualisations ornementées et standard, il n'y a donc pas de différences en ce qui concerne les appréciations relatives à l'intérêt.

H5 L'attribut « Type » n'a pas d'effet sur l'*item* « Compris ».

➔ Effectivement, les variables « *Raw* » et « *Junk* », faisant partie de l'attribut « Type », ne montrent pas de différence significative en ce qui concerne la moyenne des rangs. L'hypothèse **H5 est validée**. Entre les visualisations ornementées et standard, il n'y a pas de différences en ce qui concerne les appréciations relatives à la compréhension.

2. Corrélations inter-échelles (Kendall Tau b)

Afin de répondre aux hypothèses H6 à H9, nous devons étudier les corrélations pouvant éventuellement exister entre nos échelles. Il faut donc utiliser un test de corrélation non paramétrique. C'est pourquoi nous utilisons le test de Kendall Tau b²⁹, qui se base lui aussi sur les rangs. Ainsi, il fournit un coefficient qui représente la concordance entre deux colonnes de données classées. Ce test est généralement utilisé pour étudier la concordance entre des variables ordinales, ce qui est parfait dans notre cas. Non paramétrique, il utilise le rang des variables pour en mesurer la corrélation. Et comme pour tout test de corrélation, nous obtiendrons un coefficient³⁰ entre 0 et 1 (nous ne pouvons obtenir de corrélation

²⁹ Nous n'avons pas utilisé le test de Spearman car il assume que les relations entre les variables soient monotones, ce que nous n'étions pas en mesure de montrer avec le jeu de données. Le test de Kendall Tau b est une probabilité.

³⁰ Ce coefficient est représenté par la lettre grecque Rho (ρ).

négative dans notre cas et il n'y aura donc pas de coefficient négatif). Plus le coefficient tend vers 1, plus la corrélation est forte. La difficulté de l'exercice tient à l'interprétation : on sait s'il y a corrélation ou pas, mais la force de la corrélation tient au commentaire du chercheur. Quoi qu'il en soit, nous partons sur le fait que

$$H_0 : \rho = 0$$

Autrement dit, le coefficient de corrélation n'est pas significativement différent de 0.

Après avoir été réalisé, le test a aisément montré que tous les coefficients de corrélation sont supérieurs à zéro. De plus, le test s'est révélé significatif dans tous les cas : une corrélation existe statistiquement entre tous les *items*³¹. Nous devons rejeter H_0 selon laquelle tous les coefficients ne sont pas significativement différents de zéro, validant donc l'hypothèse alternative selon laquelle tous les coefficients sont différents de zéro. Attention néanmoins : ces coefficients montrent une relation statistique. Nous souhaitons donc les commenter. Ainsi, au-delà de la valeur de 0,5, nous considérons que la corrélation est forte³². Nous considérons la corrélation modérée ou faible sous ce nombre.

Ainsi, nous estimons donc les corrélations entre les *items* suivants comme **faibles** :

- « Belle » vs. « Intéressante » : $\rho = 0,282$
- « Belle » vs. « Claire » : $\rho = 0,148$
- « Belle » vs. « Compris » : $\rho = 0,069$

Nous estimons les corrélations entre les *items* suivants comme modérées :

- « Intéressante » vs. « Claire » : $\rho = 0,438$
- « Intéressante » vs. « Compris » : $\rho = 0,379$
- « Claire » vs. « J'aime » : $\rho = 0,437$
- « Compris » vs. « J'aime » : $\rho = 0,302$

Et nous estimons les corrélations entre les *items* suivants comme fortes :

- « Belle » vs. « J'aime » : $\rho = 0,578$
- « Intéressante » vs. « J'aime » : $\rho = 0,513$
- « Claire » vs. « Compris » : $\rho = 0,611$

³¹ Dans tous les cas concernés par ce seuil de significativité : d'une part $p < 0,01$ et d'autre part $p < 0,05$.

³² Les statistiques descriptives présentées en début de chapitre ainsi que l'étude des distributions précédemment menées nous poussent à faire ce choix.

Nous pouvons désormais faire le point sur nos hypothèses générales.

H6 Les appréciations relatives à l'*item* « Belle » sont fortement corrélées aux appréciations relatives à la variable « J'aime ».

➔ L'hypothèse **H6** est **validée**.

H7 Les appréciations relatives à l'*item* « Belle » sont fortement corrélées aux appréciations relatives à l'*item* « Claire ».

➔ S'il existe effectivement une corrélation entre ces deux *items*, elle est très faible ($\rho = 0,148$). Ainsi, l'hypothèse **H7** est **rejetée**.

H8 Les appréciations relatives à l'*item* « Belle » ne sont pas corrélées aux appréciations relatives à l'*item* « Intéressante ».

➔ Nous devons **rejeter** l'hypothèse **H8** parce que la corrélation existe d'un point de vue statistique. Néanmoins, nous devons nuancer le rejet de cette hypothèse par le fait que cette corrélation est très faible ($\rho = 0,282$).

H9 Les appréciations relatives à l'*item* « Belle » ne sont pas corrélées pas aux appréciations relatives à l'*item* « Compris ».

➔ Nous devons **rejeter** l'hypothèse **H9** parce que la corrélation existe d'un point de vue statistique. Néanmoins, nous devons nuancer le rejet de cette hypothèse par le fait que cette corrélation est très faible ($\rho = 0,069$).

Ces dernières hypothèses ont été imaginées parce que ces sujets ont retenu notre attention lors de l'analyse des statistiques descriptives. Néanmoins, il serait appréciable d'apporter un commentaire quant au fait que la corrélation est forte entre les *items* « Intéressante » et « Claire » ainsi que « Claire » et « Compris ». Il sera présent lors de notre interprétation générale. Nous pouvons néanmoins nous rappeler l'importance des corrélations fortes entre

- « Belle » et « J'aime » ;
- « Intéressante » et « J'aime » ;
- « Claire » et « Compris »

3. Effets de l'attribut « Ornement » sur les appréciations déclarées

3.1. Formulation d'hypothèses générales

Le premier attribut que nous comptons analyser en regard des appréciations déclarées est l'ornement. Cet attribut se distingue en trois groupes :

- (1) L'absence d'ornement (nommé « aucun »), qui correspond donc à toutes les visualisations de type « *Raw* »
- (2) La présence d'ornements que nous avons appelés « icônes »
- (3) Les formes inhabituelles (pas de figuration imagée sur la visualisation, mais un *data-ink* qui prend une forme particulière et inédite).

Inévitablement, les groupes (2) et (3) se rapportent à des visualisations de données de type « *Junk* ». Après avoir pris attention aux statistiques descriptives relatives à ces trois groupes, nous estimons que les *box plots* suivants (Figure 30) sont tout à fait propices à la formulation d'hypothèses.

Le *box plot* relatif aux visualisations sans ornement (groupe 1, « absence ») est identique au *box plot* illustré dans le cas de l'attribut « type : *Raw* ». Sa description a donc déjà été réalisée. Néanmoins, la comparaison des *box plots* relatifs aux deux autres groupes est très intéressante. Nous y voyons ainsi que la majorité des appréciations sont positives dans le cas des visualisations du groupe « icônes », tandis que cela est beaucoup plus nuancé dans le cas des visualisations à la forme inhabituelle (groupe 3). Le plus flagrant en ce qui concerne ce groupe est l'*item* « Claire », pour lequel très peu d'appréciations sont positives, étant ainsi affichées comme des *outliers*. De même, alors que nous avons validé **H3** selon laquelle les visualisations de type « *Junk* » sont généralement appréciées plus positivement que les visualisations de type « *Raw* », nous voyons grâce à ce *box plot* que les visualisations *junk* dont l'ornement se concentre sur une forme inhabituelle provoquent davantage d'appréciations négatives. A contrario, toutes les appréciations déclarées sont généralement positives en ce qui concerne la présence d'icônes, avec légèrement plus de réserves pour l'*item* « J'aime ».

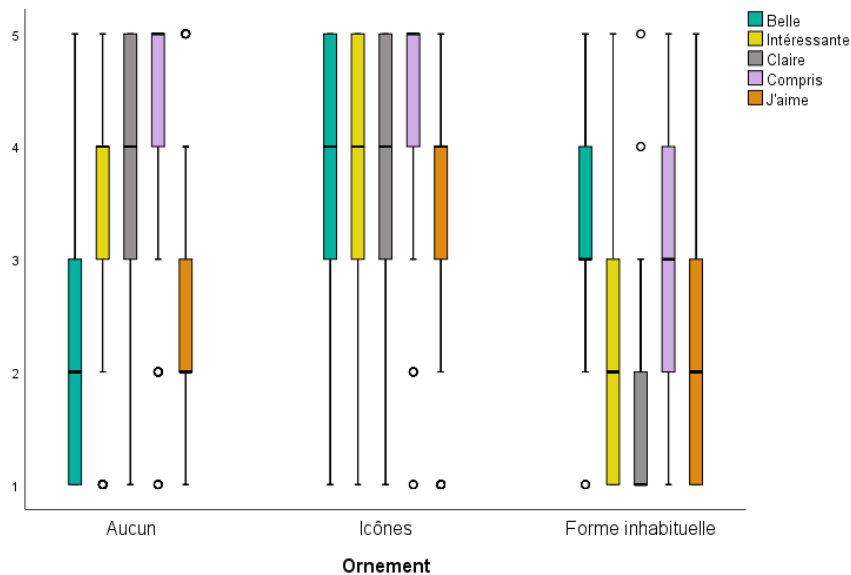


Figure 30 Box Plots illustrant les distributions des groupes de l'attribut « Ornement », pour les 5 items relevés dans les appréciations déclarées

De plus, on pourrait également supposer que les formes inhabituelles tirent l'intérêt des individus pour ces visualisations vers le bas, tout en entachant leur compréhension. De plus, en ce qui concerne les visualisations présentant des icônes, on pourrait supposer, en comparaison avec les visualisations de type « Raw », qu'elles agrandissent l'intérêt des individus pour les visualisations ainsi que leur appréciation personnelle (« J'aime »).

Ainsi, nous en venons aux hypothèses générales suivantes :

- H10** Au sein des visualisations de type « Junk », les appréciations relatives à tous les *items* sont significativement plus hautes pour les visualisations présentant des icônes que pour les visualisations dont la forme est inhabituelle.
- H11** Il n'y a pas de différence significative entre les visualisations sans ornement et les visualisations présentant des icônes en ce qui concerne l'*item* « Clarté ».
- H12** Il n'y a pas de différence significative entre les visualisations sans ornement et les visualisations présentant des icônes en ce qui concerne l'*item* « Compris ».
- H13** Les appréciations relatives aux *items* « Belle », « Intéressante » et « J'aime » sont significativement plus hautes pour les visualisations présentant des icônes que pour les visualisations ne présentant aucun ornement.

H14 Les appréciations relatives aux *items* « Intéressante », « clarté », « Compris » et « J’aime » sont significativement plus hautes pour les visualisations sans ornement que pour les visualisations dont la forme est inhabituelle.

H15 Les appréciations relatives à l’*item* « Belle » sont significativement plus basses pour les visualisations ne présentant aucun ornement que pour les visualisations dont la forme est inhabituelle.

3.2. Analyse approfondie

Nous avons soumis les données au test de Brown Forsythe afin de valider ou non l’hypothèse H0 selon laquelle les moyennes des rangs des variables contenues dans l’attribut « Ornement » sont égales ou non et ce, pour chaque *item*. Les résultats sont dans le Tableau 18. La *p-value* du test est inférieure à 0,05 : nous rejetons l’hypothèse H0 selon laquelle les moyennes des rangs des variables contenues dans l’attribut « Ornement » sont égales. Cela signifie que les appréciations sont différentes en fonction de cet attribut, qui rassemble les variables « Aucun », « Icones » et « Formes Inhabituelles », ce qui laisse entendre que ces variables ont une influence sur les appréciations laissées par les participants.

Tableau 18 Test de Brown Forsythe sur toutes les variables de l’attribut « Ornement »

Robust Tests of Equality of Means		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	209.272	2	159.498	.000
Rank of Intéressante	Brown-Forsythe	23.674	2	243.563	.000
Rank of Claire	Brown-Forsythe	76.913	2	482.139	.000
Rank of Compris	Brown-Forsythe	42.225	2	204.731	.000
Rank of Jaime	Brown-Forsythe	89.358	2	249.700	.000

a. Asymptotically F distributed.

Nous savons ainsi que concernant ces trois variables, les moyennes des rangs ne sont pas identiques. Ce que nous ignorons encore est ce que le test ne nous dit pas : entre quelles variables les moyennes des rangs sont-elles différentes ? Pour le savoir, nous allons de nouveau procéder au test de Brown Forsythe, entre chaque combinaison de variables possibles et pour chaque *item*.

Notre seuil de signification α de référence est 0,05. Nous avons effectué 3 comparaisons pour chaque *item*, ce qui risque donc d’augmenter la probabilité de rejeter incorrectement H0. Nous procédons donc à une correction de Bonferroni pour diminuer cette probabilité. Ainsi, notre seuil de signification ajusté est de 0,016 ($\alpha' = 0,05 / 3$).

Nous pouvons donc dès à présent interpréter les résultats du test de Brown Forsythe selon lequel H_0 signifie l'égalité des moyennes des rangs. H_0 est rejetée au seuil 0,016 et les moyennes effectives de ces rangs nous permettront d'expliquer comment.

1. Entre « Aucun » et « Icônes »

Selon les résultats du

Tableau 19, dont les *p-values* se trouvent dans la dernière colonne, on peut conclure que

- En ce qui concerne les *items* « Belle », « Intéressante » et « J'aime », **H_0 est rejetée** car la *p-value* est inférieure à 0,016. Les moyennes des rangs sont différentes et on peut donc constater un effet de l'icône sur ces items.
 - « Belle » : avec une moyenne des rangs de 237,98 pour « Aucun » et de 515,37 pour « Icônes », on peut conclure de manière significative que les appréciations de l'*item* « Belle » sont plus hautes pour les visualisations présentant des icônes.
 - Intéressante : avec une moyenne des rangs de 364,35 pour « Aucun » et de 407,78 pour « Icônes », on peut conclure de manière significative que les appréciations de l'*item* « Intéressante » sont plus hautes pour les visualisations présentant des icônes.
 - « J'aime » : on constate de nouveau la même chose, avec une moyenne des rangs de 305,97 pour « Aucun » et de 404,73 pour « Icônes ».

En somme, les visualisations de données présentant des icônes sont jugées plus belles, plus intéressantes et plus claires que les visualisations présentant l'absence d'ornement.

- En ce qui concerne les *items* « Claire » et « Compris », **nous ne pouvons rejeter H_0** car la *p-value* est supérieure à 0,016. Nous ne constatons pas d'effet de l'icône ou de l'absence d'ornement sur ces appréciations.

Tableau 19 Test de Brown Forsythe pour les variables « Aucun » et « Icônes » de l'attribut « Ornement », pour tous les items

Robust Tests of Equality of Means					
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	524.087	1	646.363	.000
Rank of Intéressante	Brown-Forsythe	7.894	1	669.988	.005
Rank of Claire	Brown-Forsythe	.136	1	673.747	.712
Rank of Compris	Brown-Forsythe	1.074	1	676.820	.300
Rank of Jaime	Brown-Forsythe	134.758	1	646.647	.000

a. Asymptotically F distributed.

2. Entre « Icônes » et « Forme inhabituelle »

Selon les résultats du Tableau 20, dont les *p-values* se trouvent dans la dernière colonne, on peut conclure que toutes les *p-values* sont significatives et donc inférieures à 0,016 : **nous rejetons donc H0 pour tous les items**. Pour être plus précise, les moyennes des rangs de la variable « Forme Inhabituelle » sont toujours largement inférieures à la moyenne des rangs de la variable « Icônes ». Ainsi, on peut dire que pour tous ces cas, l'icône engendre beaucoup plus d'appréciations positives, pour tous les *items*, que les formes inhabituelles.

Tableau 20 Test de Brown Forsythe pour les variables « Icônes » et « Forme Inhabituelle » de l'attribut « Ornement », pour tous les items

Robust Tests of Equality of Means					
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	7.394	1	70.102	.008
Rank of Intéressante	Brown-Forsythe	47.390	1	77.869	.000
Rank of Claire	Brown-Forsythe	208.619	1	107.277	.000
Rank of Compris	Brown-Forsythe	77.426	1	72.848	.000
Rank of Jaime	Brown-Forsythe	90.303	1	80.953	.000

a. Asymptotically F distributed.

3. Entre « Aucun » et « Forme inhabituelle »

Selon les résultats du Tableau 21, dont les *p-values* se trouvent dans la dernière colonne, on peut conclure que toutes les *p-value* sont significatives et donc inférieures à 0,016 : **nous rejetons donc H0 pour tous les items**. Les moyennes de ces rangs peuvent expliquer certains effets. Ainsi, en ce qui concerne l'*item*

- « Belle » : la moyenne des rangs de « Aucun » s'élève à 237,98 contre 439,07 pour « Forme Inhabituelle ». Les appréciations sont donc davantage positives quand il y a une forme inhabituelle sur la visualisation en comparaison à l'absence d'ornement.
- « Intéressante » : la moyenne des rangs de « Aucun » s'élève à 364,35 contre 211,51 pour « Forme Inhabituelle ». Il semblerait donc que la forme inhabituelle réduit l'intérêt du lecteur en comparaison à une visualisation marquée par l'absence d'ornement.
- « Claire » : la moyenne des rangs de « Aucun » s'élève à 390,45 contre 107,08 pour « Forme Inhabituelle ». Les appréciations sont davantage positives pour les visualisations à l'absence d'ornement. On supposerait donc que les formes inhabituelles rendent les visualisations moins claires.

- « Compris » : la moyenne des rangs de « Aucun » s'élève à 381,87 contre 166,49 pour « Forme Inhabituelle ». La conclusion est très similaire à celle de l'*item* « Claire ».
- « J'aime » : la moyenne des rangs de « Aucun » s'élève à 305,97 contre 221,64 pour « Forme Inhabituelle ». Il semblerait que les individus aiment davantage les visuels marqués par l'absence d'ornementation que ceux marqués par une forme dont ils n'ont pas l'habitude.

Tableau 21 Test de Brown Forsythe pour les variables « Aucun » et « Forme Inhabituelle » de l'attribut « Ornement », pour tous les items

Robust Tests of Equality of Means		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	52.841	1	66.280	.000
Rank of Intéressante	Brown-Forsythe	29.179	1	75.676	.000
Rank of Claire	Brown-Forsythe	203.940	1	104.637	.000
Rank of Compris	Brown-Forsythe	68.691	1	72.709	.000
Rank of Jaime	Brown-Forsythe	10.499	1	74.141	.002

a. Asymptotically F distributed.

3.3. Falsification des hypothèses générales

Les résultats nous permettent ainsi de revenir aux hypothèses générales, d'y répondre et d'interpréter davantage.

H10 Au sein des visualisations de type « Junk », les appréciations relatives à tous les *items* sont significativement plus hautes pour les visualisations présentant des icônes que pour les visualisations dont la forme est inhabituelle.

➔ **L'hypothèse est validée** très simplement grâce aux conclusions de la partie n°2 du point précédent.

H11 Il n'y a pas de différence significative entre les visualisations sans ornement et les visualisations présentant des icônes en ce qui concerne l'*item* « Clarté ».

➔ **L'hypothèse est validée.** Ainsi, les icônes ne semblent pas perturber la clarté des visualisations.

H12 Il n'y a pas de différence significative entre les visualisations sans ornement et les visualisations présentant des icônes en ce qui concerne l'*item* « Compris ».

L'hypothèse est validée. De nouveau, les icônes ne semblent pas perturber la compréhension des visualisations de données, en comparaison aux visualisations ne présentant pas d'ornement.

H13 Les appréciations relatives aux *items* « Belle », « Intéressante » et « J'aime » sont significativement plus hautes pour les visualisations présentant des icônes que pour les visualisations ne présentant aucun ornement.

➔ **L'hypothèse est validée.**

H14 Les appréciations relatives aux *items* « Intéressante », « Clarté », « Compris » et « J'aime » sont significativement plus hautes pour les visualisations sans ornement que pour les visualisations dont la forme est inhabituelle.

➔ **L'hypothèse est validée.**

H15 Les appréciations relatives à l'*item* « Belle » sont significativement plus basses pour les visualisations ne présentant aucun ornement que pour les visualisations dont la forme est inhabituelle.

➔ **L'hypothèse est rejetée.** Contre toute attente, l'*item* « Belle » suscite beaucoup plus d'appréciations positives pour les visualisations à l'absence d'ornement plutôt que pour les visualisations de données dont la forme est inhabituelle.

4. Effets de l'attribut « Figuration » sur les appréciations déclarées

4.1. Formulation d'hypothèses

Nous pouvons désormais aborder l'attribut « figuration » qui se distingue en trois variables :

- (1) Absence : l'absence de figuration correspond aux visualisations de type « *Raw* »
- (2) Apposé : la figuration est séparée du graphe
- (3) Sur le *data-ink* : la figuration se fond dans le graphe.

Nous pouvons dès lors observer les distributions de ces trois groupes de données afin d'établir des hypothèses à vérifier (Figure 31). Ces *box plots* nous montrent des similarités de distribution en ce qui concerne l'intérêt : cet attribut n'aurait donc pas d'effet sur l'appréciation « Intéressante ». Il en va de même pour l'*item* « Compris ». Par contre, en

ce qui concerne les *items* « J’aime » et « Belle », on constate qu’elles obtiennent toutes les deux un score plus élevé pour la variable « apposé » que pour la variable « Sur le *Data-Ink* », tandis qu’il reste assez négatif en ce qui concerne la variable « absence » de figuration. Enfin, on pourrait supposer que la figuration apposée ne perturbe pas la clarté du visuel en comparaison avec un visuel où il y a absence de figuration. Par contre, la figuration « Sur le *Data-Ink* » impacterait négativement cette clarté.

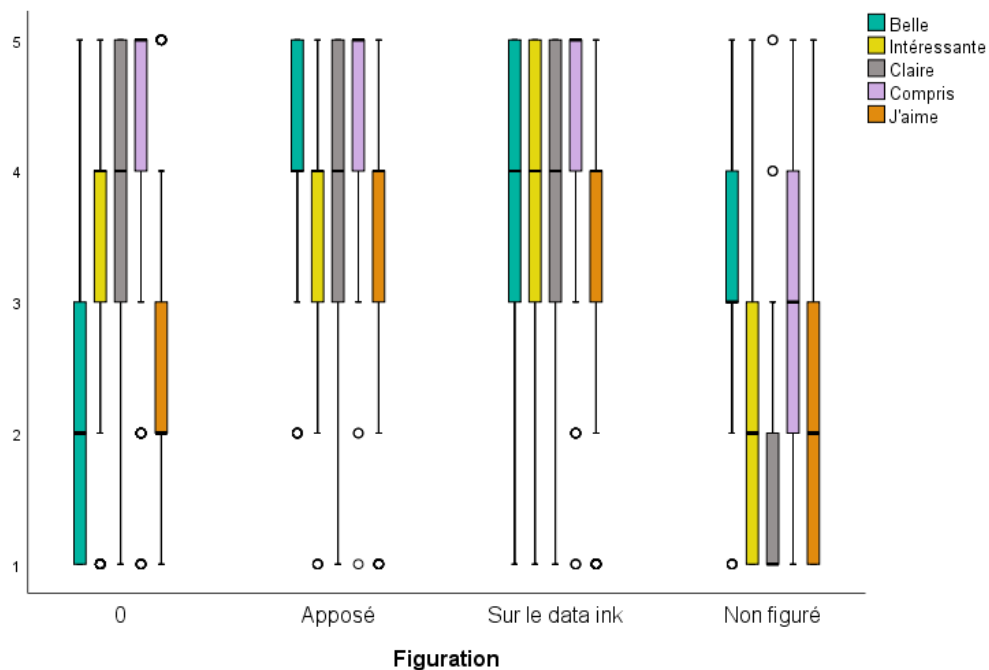


Figure 31 Box Plots représentant les distributions des items pour les variables de l'attribut « Figuration »

Nous pouvons ainsi émettre les hypothèses suivantes.

- H16** La figuration n’a pas d’impact significatif sur les appréciations relatives à l’intérêt.
- H17** La figuration n’a pas d’impact significatif sur les appréciations relatives à la compréhension.
- H18** Les visualisations de données à la figuration apposée obtiennent davantage d’appréciations positives que les visualisations à la figuration sur le *data-ink*, en ce qui concerne l’item « Belle ».
- H19** Les visualisations de données à la figuration apposée obtiennent davantage d’appréciations positives que les visualisations à la figuration sur le *data-ink*, en ce qui concerne l’item « J’aime ».

- H20** En ce qui concerne les *items* « Belle » et « J'aime », les visualisations de données qui sont réalisées sans figuration obtiennent des appréciations plus faibles que les visualisations à la figuration apposée ou sur le *data-ink*.
- H21** Il n'y a pas de différence significative en termes de clarté en ce qui concerne les variables « absence » et « apposé ». Les visualisations de données à la figuration apposée ne perturbent pas la clarté.
- H22** En ce qui concerne l'*item* relatif à la clarté, les visualisations de données dont la figuration se concentre sur le *data-ink* génèrent des appréciations plus négatives que les visualisations à l'absence de figuration ou dont la figuration est apposée.

4.2. Analyse approfondie

Les résultats du test général de Brown Forsythe sur les rangs des *items*, en ce qui concerne l'attribut « Figuration », montrent bien que toutes les moyennes des rangs sont différentes au moins une fois entre les variables de l'attribut exploré excepté en ce qui concerne l'intérêt (Tableau 22). La figuration a donc bel et bien un effet sur les appréciations relevées, excepté pour l'*item* « Intéressante ». Comme on le voit dans le tableau, tous les résultats sont significatifs, avec une *p-value* inférieure à 0,05, excepté pour l'*item* « Intéressante » qui montre une *p-value* de 0,286.

Ainsi, nous rejetons H0 selon laquelle les moyennes des rangs des variables relatives à la figuration sont égales, pour les *items* « Belle », « Claire », « Compris » et « J'aime ». Cette hypothèse peut être conservée en ce qui concerne l'*item* « Intéressante ».

Nous pouvons donc déjà valider l'hypothèse H16 selon laquelle la figuration n'a pas d'impact sur les appréciations relatives à l'intérêt. Nous ne tiendrons donc pas compte de l'intérêt dans la suite de l'analyse relative à cet attribut.

Tableau 22 Test de Brown Forsythe concernant l'attribut "Figuration"

Robust Tests of Equality of Means		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	276.358	2	549.032	.000
Rank of Intéressante	Brown-Forsythe	1.254	2	535.073	.286
Rank of Claire	Brown-Forsythe	5.705	2	532.522	.004
Rank of Compris	Brown-Forsythe	3.728	2	531.419	.025
Rank of Jaime	Brown-Forsythe	42.552	2	490.308	.000

a. Asymptotically F distributed.

Nous pouvons dès lors poursuivre en comparant les moyennes des rangs des variables entre elles, pour savoir quel est l'effet de ces variables sur les appréciations. De nouveau, la correction de Bonferroni nous décide à choisir 0,016 comme seuil de signification ($\alpha' = 0,05/3$) avec lequel comparer nos *p-values*.

1. Entre les variables « Absence » et « Apposé »

H0 selon laquelle les moyennes des rangs sont égales est rejetée pour les *items* « Belle » et « J'aime ». Comme on le voit dans le Tableau 23, la *p-value* est de zéro, ce qui est significatif. En observant les moyennes des rangs, on peut dire que :

- Pour l'*item* « Belle », la moyenne des rangs de la variable « Absence » est de 237,98 contre 553,59 pour la variable « Apposé ». Les visualisations de données dont la figuration est apposée reçoivent donc significativement plus d'appréciations positives en ce qui concerne cet *item* que les visualisations de données à l'absence de figuration.
- Pour l'*item* « J'aime », la moyenne des rangs de la variable « Absence » est de 305,97 contre 474,95 pour la variable « Apposé ». De nouveau, les visualisations de données dont la figuration est apposée reçoivent donc significativement plus d'appréciations positives en ce qui concerne cet *item* que les visualisations de données à l'absence de figuration.

C'est tout l'inverse en ce qui concerne les *items* « Claire » et « Compris » : le résultat au test de Brown Forsythe n'est pas significatif puisque les *p-values* excèdent bien 0,016 (avec respectivement les valeurs de 0,741 et 0,678). Nous ne pouvons donc rejeter H0 selon laquelle les moyennes des rangs sont similaires. Ainsi, il n'y a pas de différence significative en ce qui concerne l'effet de l'une ou l'autre variable sur ces appréciations.

Tableau 23 Test de Brown Forsythe sur les variables "Absence" et "Apposé", contenues dans l'attribut "Figuration"

Robust Tests of Equality of Means					
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	446.600	1	214.935	.000
Rank of Claire	Brown-Forsythe	.109	1	211.247	.741
Rank of Compris	Brown-Forsythe	.173	1	210.217	.678
Rank of Jaime	Brown-Forsythe	71.621	1	189.743	.000

a. Asymptotically F distributed.

2. Entre les variables « Apposé » et « Sur le Data-Ink »

Comme en attestent les *p-values* qui se trouvent dans la dernière colonne du Tableau 24, le résultat est significatif pour les *items* « Belle » et « J'aime ». Ainsi, nous rejetons H_0 selon laquelle les moyennes des rangs de ces deux *items* sont identiques pour ces variables. En observant les moyennes des rangs, on peut dire que :

- En ce qui concerne l'*item* « Belle », la variable « Apposé » reçoit davantage d'appréciations positives que la variable « Sur le Data-Ink » (avec pour moyennes des rangs respectives 553,59 et 480,36).
- En ce qui concerne l'*item* « J'aime », la variable « Apposé » reçoit davantage d'appréciations positives que la variable « Sur le Data-Ink » (avec pour moyennes des rangs respectives 474,95 et 417,60, ce qui n'est somme toute pas un grand effet).

Les *items* « Claire » et « Compris » ne présentent pas de *p-values* significatives. Ainsi, on ne peut rejeter H_0 selon laquelle les moyennes de ces rangs sont similaires et donc, on ne peut dire que la figuration apposée ou sur le *data-ink* ait un réel effet sur ces appréciations.

Tableau 24 Test de Brown Forsythe sur les variables « Apposé » et « Sur le Data-Ink », contenues dans l'attribut « Figuration »

Robust Tests of Equality of Means					
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	18.549	1	291.990	.000
Rank of Claire	Brown-Forsythe	4.651	1	265.826	.032
Rank of Compris	Brown-Forsythe	4.844	1	267.036	.029
Rank of Jaime	Brown-Forsythe	6.535	1	262.276	.011

a. Asymptotically F distributed.

3. Entre les variables « Absence » et « Sur le Data-Ink »

Comme en attestent les *p-values* qui se trouvent dans la dernière colonne du Tableau 25, le résultat est significatif pour les tous les *items* sauf « Compris », qui a une *p-value* de 0,021, dès lors supérieure à notre seuil de signification 0,016. Nous rejetons donc H_0 selon laquelle les moyennes des rangs sont similaires pour tous les *items*, sauf « Compris ». Ainsi, on ne peut affirmer qu'au niveau de l'appréciation relative à la compréhension, il y ait une différence entre l'absence de figuration et la figuration sur le *data-ink*. Pour les autres *items*, nous pouvons dire qu'en ce qui concerne :

- L'*item* « Belle », les visualisations à la figuration concentrée sur le *data-ink* reçoivent davantage d'appréciations positives que les visualisations de données avec absence de figuration (moyennes des rangs respectives : 480,36 contre 237,98).
- L'*item* « Claire », les visualisations avec absence de figuration reçoivent davantage d'appréciations positives que les visualisations de données avec figuration concentrée sur le *data-ink*. La différence entre les moyennes des rangs est significative, même si peu élevée (390,45 contre 336,55).
- L'*item* « J'aime », les visualisations à la figuration concentrée sur le *data-ink* reçoivent beaucoup plus d'appréciations positives que celles manifestant une absence de figuration (moyennes des rangs respectives : 417,30 contre 305,97).

Nous pouvons ainsi de nouveau constater l'intensité de la relation entre les *items* « Belle » et « J'aime », puisque les visualisations qui sont ornementées et dont l'ornementation se concentre sur les traits relatifs aux données obtiennent des scores plus élevés pour ces *items* que les visualisations standard. Par contre, on peut clairement dire que la clarté souffre de cette ornementation, puisque les visualisations de type standard sont jugées bien plus claires.

Tableau 25 Test de Brown Forsythe sur les variables « Absence » et « Sur le Data-Ink », contenues dans l'attribut « Figuration »

		Robust Tests of Equality of Means			
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	308.171	1	478.802	.000
Rank of Claire	Brown-Forsythe	9.758	1	510.317	.002
Rank of Compris	Brown-Forsythe	5.371	1	506.526	.021
Rank of Jaime	Brown-Forsythe	43.875	1	470.009	.000

a. Asymptotically F distributed.

4.3. Falsification des hypothèses générales

Les tests réalisés nous permettent désormais de falsifier nos hypothèses générales et de les commenter.

H16 La figuration n'a pas d'impact significatif sur les appréciations relatives à l'intérêt.

➔ **L'hypothèse est validée** : le type de figuration n'influence nullement les marques d'intérêt manifestées par les participants.

- H17** La figuration n'a pas d'impact significatif sur les appréciations relatives à la compréhension.
- ➔ **L'hypothèse est validée** : tous les tests comparant les variables entre elles se sont montrés non-significatifs. Peu importe donc le type de figuration ou même son absence, la compréhension n'en semble pas impactée, selon ce qu'ont déclaré les participants.
- H18** Les visualisations de données à la figuration apposée obtiennent davantage d'appréciations positives que les visualisations à la figuration sur le *data-ink*, en ce qui concerne l'*item* « Belle ».
- ➔ **L'hypothèse est validée**. Il semblerait donc que les visualisations de données à la figuration apposée soient davantage appréciées par les participants.
- H19** Les visualisations de données à la figuration apposée obtiennent davantage d'appréciations positives que les visualisations à la figuration sur le *data-ink*, en ce qui concerne l'*item* « J'aime ».
- ➔ **L'hypothèse est validée**. Toutefois, le test est significatif de justesse : l'écart entre les moyennes des rangs n'est pas très élevé.
- H20** En ce qui concerne les *items* « Belle » et « J'aime », les visualisations de données qui sont réalisées sans figuration obtiennent des appréciations plus faibles que les visualisations à la figuration apposée ou sur le *data-ink*.
- ➔ **L'hypothèse est validée**.
- H21** Il n'y a pas de différence significative en termes de clarté en ce qui concerne les variables « absence » et « apposé ». Les visualisations de données à la figuration apposée ne perturbent pas la clarté.
- ➔ **L'hypothèse est validée** : le résultat du test n'était pas significatif.
- H22** En ce qui concerne l'*item* relatif à la clarté, les visualisations de données dont la figuration se concentre sur le *data-ink* génèrent des appréciations plus négatives que les visualisations à l'absence de figuration ou dont la figuration est apposée.
- ➔ **L'hypothèse est validée**, même si la différence entre les moyennes des rangs est assez peu élevée, bien que significative.

5. Effets de l'attribut « Pictogramme » sur les appréciations déclarées

5.1. Formulation d'hypothèses

L'observation des *box plots* relatifs aux variables contenues dans l'attribut « Pictogramme » (Figure 32) nous permet d'imaginer différentes hypothèses.

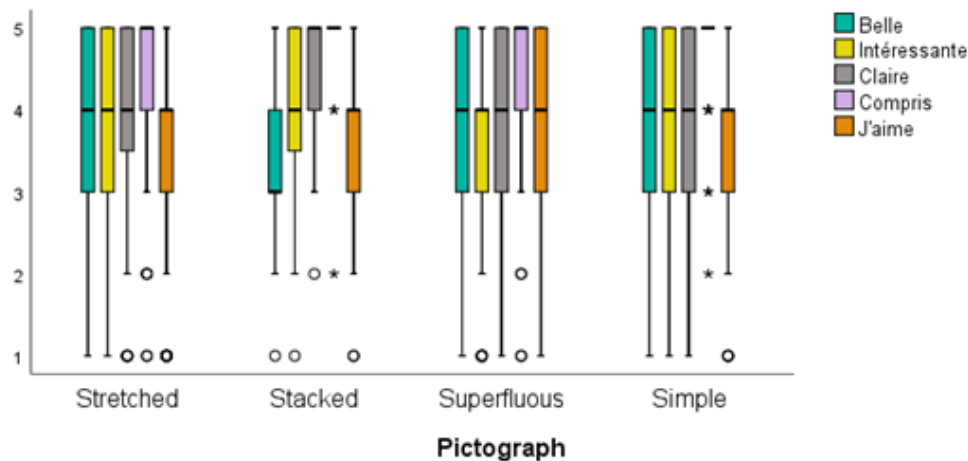


Figure 32 Box Plots des variables relatives à l'attribut "Pictogramme"

De manière générale, on constate que les scores sont assez hauts pour l'*item* « Belle » en ce qui concerne les variables « Stretched », « Superfluous » et « Simple ». Le score est un peu plus bas pour la variable « Stacked ». L'utilisation du pictogramme semble améliorer la clarté, en particulier pour la variable « Stacked ». Le pictogramme semble également avoir une bonne influence sur l'intérêt de manière générale, et cette influence est moindre en ce qui concerne le pictogramme « Superfluous ». Le pictogramme semble avoir une influence positive sur la compréhension, peu importe la forme employée. De même, il semble que le pictogramme « Superfluous » soit préféré des participants. Après ces observations, nous pouvons émettre les hypothèses suivantes :

- H21** En ce qui concerne l'*item* « Belle », il n'y a pas de différence significative entre les variables « Stretched », « Superfluous » et « Simple » de l'attribut « Pictogramme ».
- H22** En ce qui concerne l'*item* « Belle », la variable « Stacked » récolte significativement moins d'appréciations positives que les variables « Stretched », « Superfluous » et « Simple ».

- H23** Le pictogramme améliore significativement les appréciations positives concernant l'intérêt. Autrement dit, les appréciations de la variables « Absent » concernant l'intérêt sont plus basses que les appréciations attribuées aux autres variables.
- H24** Il n'y a pas de différence significative entre les variables « Stacked » et « Simple », en ce qui concerne l'*item* « Compris ».
- H25** Il n'y a pas de différence significative entre les différents types de pictogrammes en ce qui concerne l'*item* « J'aime ».

5.2. Analyse approfondie

Comme on peut le constater dans la dernière colonne du Tableau 26, aucun des résultats du test ne sont significatifs puisque la *p-value* est toujours supérieure à 0,05. Ainsi, pour tous les *items*, nous ne pouvons rejeter H_0 selon laquelle les moyennes des rangs de tous les *items* sont similaires. Cela ne nous permet donc pas d'affirmer que les différents pictogrammes aient un effet sur les appréciations déclarées. Nous avons pour rappel évincé les visualisations de type standard de cette analyse, ainsi que les visualisations ornementées qui ne présentent aucun pictogramme. Elle ne concerne donc que les visualisations ornementées sur lesquelles se trouvent un pictogramme. Ainsi, nous pouvons dire que le type de pictogramme employé n'a pas forcément d'incidence sur les différentes appréciations déclarées.

Notons cependant que ce résultat est probablement influencé par le fait que le nombre de visualisations par catégories de pictogramme est relativement bas pour mener une analyse de ce genre.

Tableau 26 Test de Brown Forsythe sur l'attribut « Pictogramme »

		Robust Tests of Equality of Means			
		Statistic ^a	df1	df2	Sig.
Belle	Brown-Forsythe	1.922	4	167.347	.109
Intéressante	Brown-Forsythe	1.807	4	156.897	.130
Claire	Brown-Forsythe	2.404	4	201.910	.051
Compris	Brown-Forsythe	.504	4	188.676	.733
J'aime	Brown-Forsythe	.743	4	160.632	.564

a. Asymptotically F distributed.

Ainsi, nous ne sommes pas en mesure de valider les hypothèses 21 à 27 concernant l'attribut « Pictogramme ».

6. Effet de l'attribut « *Lie Factor* » sur les appréciations déclarées

6.1. Formulation d'hypothèses

Sur base des *box plots* de la Figure 33, nous pouvons émettre différentes hypothèses relatives au respect du *lie factor*.

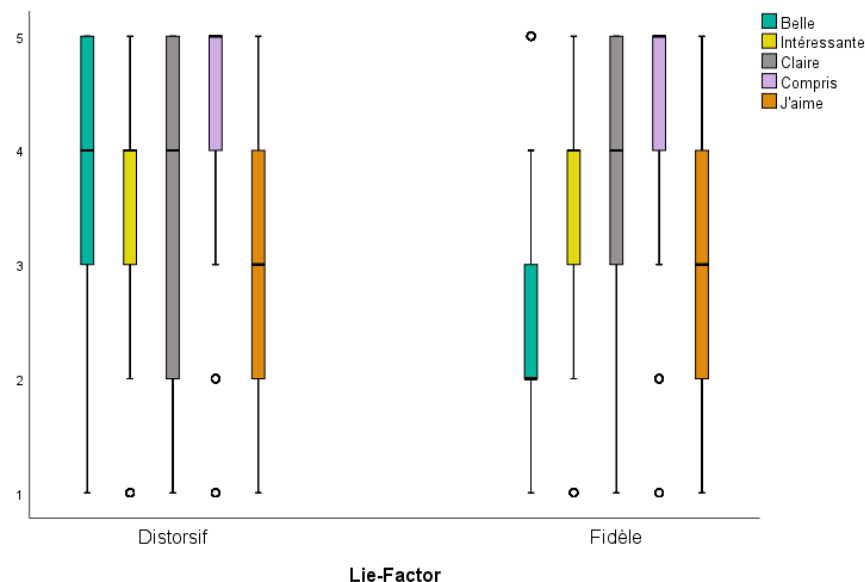


Figure 33 Box Plots présentant les distributions des données en fonction de l'attribut « Lie factor »

Les distributions suivantes supposent que le *lie factor* ait un impact sur deux *items* : la beauté et la clarté. Ainsi, le non-respect du *lie factor* entrainerait des appréciations déclarées plus hautes pour l'*item* « Belle » et plus basses pour l'*item* « Claire », tandis que cela n'aurait pas d'effet significatif sur les autres appréciations déclarées relevées dans le cadre des expérimentations.

H26 Les visualisations de données à l'attribut « Lie Factor Distorsif » suscitent davantage d'appréciations positives en ce qui concerne l'*item* « Belle » que les visualisations de données au « Lie Factor Fidèle ».

H27 Les visualisations de données à l'attribut « Lie Factor Distorsif » suscitent moins d'appréciations positives en ce qui concerne l'*item* « Claire » que les visualisations de données au « Lie Factor Fidèle ».

H28 L'attribut « Lie Factor » n'a aucun effet sur l'*item* « Intéressante ».

H29 L'attribut « Lie Factor » n'a aucun effet sur l'*item* « Compris ».

H30 L'attribut « Lie Factor » n'a aucun effet sur l'*item* « J'aime ».

6.2. Analyse approfondie

Pour cette analyse, nous n'aurons pas besoin d'une correction de Bonferroni puisqu'il n'y a qu'une comparaison à effectuer. Comme le montre la dernière colonne du Tableau 27, la *p-value* est significative pour les *items* « Belle », « Claire », « Compris » et « J'aime ». Effectivement, elle est de 0, ce qui est inférieur à 0,05. Ainsi, on peut dire que pour ces trois *items*, H_0 , supposant l'égalité de la moyenne des rangs est à rejeter. Il y a donc un effet du *lie factor* sur ces *items*, et nous pouvons le commenter :

- « Belle » : la variable « Distorsif » suscite davantage d'appréciations positives que la variable « Fidèle » (moyenne des rangs de 486,97 contre 310,51).
- « Claire » : la variable « Distorsif » suscite significativement moins d'appréciations positives que la variable « Fidèle » (moyennes des rangs de 330,14 contre 392,30).
- « Compris » : la variable « Distorsif » suscite significativement moins d'appréciations positives que la variable « Fidèle » (moyennes des rangs de 346,89 contre 383,57).
- « J'aime » : la variable « Distorsif » suscite davantage d'appréciations positives que la variable « Fidèle » (moyenne des rangs de 410,31 contre 350,49).

En ce qui concerne l'*item* « Intéressante », nous ne pouvons rejeter H_0 et nous ne pouvons donc dire que le *lie factor* a un effet.

Tableau 27 Test de Brown Forsythe sur l'attribut "Lie Factor"

Robust Tests of Equality of Means					
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	148.729	1	551.711	.000
Rank of Intéressante	Brown-Forsythe	.760	1	501.689	.384
Rank of Claire	Brown-Forsythe	14.656	1	474.660	.000
Rank of Compris	Brown-Forsythe	6.490	1	469.588	.011
Rank of Jaime	Brown-Forsythe	13.163	1	474.761	.000

6.2. Falsification des hypothèses

H26 Les visualisations de données à l'attribut « Lie Factor Distorsif » suscitent davantage d'appréciations positives en ce qui concerne l'*item* « Belle » que les visualisations de données au « Lie Factor fidèle ».

➔ L'hypothèse est validée.

H27 Les visualisations de données à l'attribut « Lie Factor Distorsif » suscitent moins d'appréciations positives en ce qui concerne l'*item* « Claire » que les visualisations de données au « Lie Factor Fidèle ».

➔ **L'hypothèse est validée**, ce qui nous laisse conclure que le *lie factor* a un effet néfaste sur la clarté des visualisations, comme le déclarent les participants.

H28 L'attribut « Lie Factor » n'a aucun effet sur l'*item* « Intéressante ».

➔ **L'hypothèse est validée**.

H29 L'attribut « Lie Factor » n'a aucun effet sur l'*item* « Compris ».

➔ **L'hypothèse est rejetée**. Les participants laissent entendre qu'un *lie factor* distorsif entache leur compréhension, puisque les appréciations positives sont significativement plus nombreuses lorsque le *lie factor* est fidèle dans les visualisations de données.

H30 L'attribut « Lie Factor » n'a aucun effet sur l'*item* « J'aime ».

➔ **L'hypothèse est rejetée**. Il semble que les visualisations au *lie factor* distorsif récoltent significativement plus d'appréciations positives concernant l'*item* « J'aime ».

Nous pouvons ainsi dire que la distorsion du graphique a un effet réel sur la perception et les appréciations des participants puisque cela diminue l'appréciation de clarté et de compréhension, tout en augmentant le jugement de beauté et le fait d'aimer le graphique. Ainsi, même en étant conscients que les visualisations de données présentant de la distorsion sont moins claires et compréhensibles que les visualisations de données sans aucune distorsion, les participants les préfèrent tout de même et les jugent plus belles que les visualisations standard et les visualisations ornementées au *lie factor* fidèle.

7. Effet de l'attribut « *Data-Ink* » sur les appréciations déclarées

7.1. Formulation d'hypothèses

Pour cet attribut, le *data-ink* a été scindé en trois classes : bas, moyen et élevé. Les *box plots* présentés dans la Figure 34 nous permettent d'émettre plusieurs hypothèses. On voit effectivement que les visualisations récoltant le plus d'appréciations de beauté sont celles dont le *data-ink* est bas. De même, il semblerait que les visualisations dont le *data-ink* est bas et élevé génèrent des appréciations déclarées plus élevées en ce qui concerne l'intérêt, la clarté et la compréhension que les visualisations dont le *data-ink* est bas. Ce sont par contre les visualisations au *data-ink* moyen qui récoltent le plus d'appréciations positives pour l'item « J'aime ».

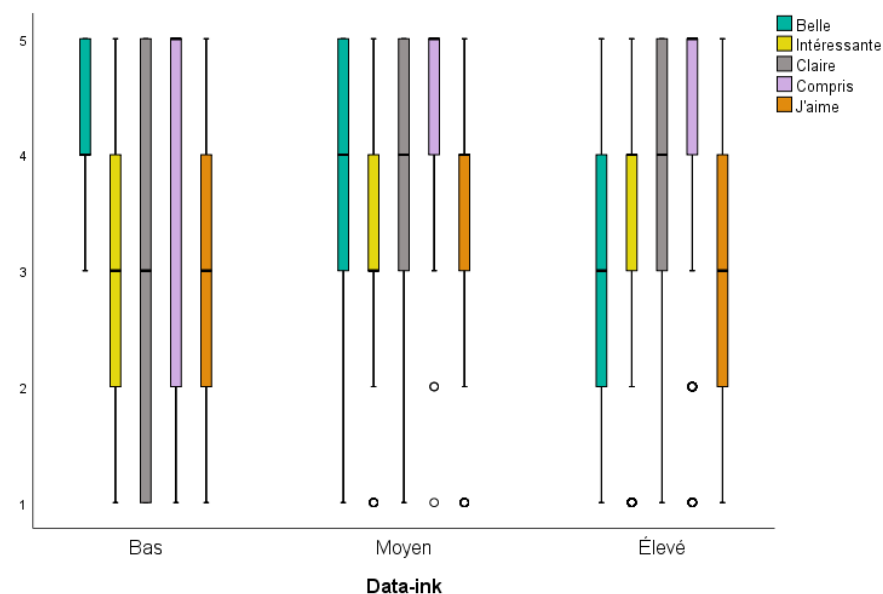


Figure 34 Box Plots représentant les distributions des données en fonction de l'attribut « *Data-Ink* »

- H31** Il n'y a pas de différence significative entre les variables « Moyen » et « Élevé » de l'attribut « *Data-Ink* » en ce qui concerne les items « Intéressante », « Claire » et « Compris ».
- H32** Les visualisations au *data-ink* moyen reçoivent davantage d'appréciations positives pour l'item « Belle » que les autres variables.
- H33** Les visualisations au *data-ink* moyen reçoivent davantage d'appréciations positives pour l'item « J'aime » que les autres variables.
- H34** Les visualisations au *data-ink* bas reçoivent davantage d'appréciations positives pour l'item « Belle » que celles dont le *data-ink* est qualifié de moyen et élevé.

- H35** Les visualisations au *data-ink* bas reçoivent moins d’appréciations positives pour l’*item* « Intéressante » que celles dont le *data-ink* est qualifié de moyen et élevé.
- H36** Les visualisations au *data-ink* bas reçoivent moins d’appréciations positives pour l’*item* « Compris » que celles dont le *data-ink* est qualifié de moyen et élevé.

7.2. Analyse approfondie

Avec un seuil de signification à 0,05, nous pouvons constater que toutes les *p-values* présentes dans le Tableau 28 sont significatives, à l’exception de l’*item* « Intéressante ». Nous rejetons donc H0 selon laquelle les moyennes des rangs sont similaires pour tous les *items*, sauf pour « Intéressante ». Pour savoir quel est l’effet supposé par ce test, nous allons procéder à une comparaison multiple pour les *items* dont la *p-value* est significative (donc inférieure à 0,05). Nous pourrions ainsi amener plus de précisions et de commentaires.

Tableau 28 Test de Brown Forsythe sur l’attribut « Data-Ink »

Robust Tests of Equality of Means					
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	112.813	2	198.347	.000
Rank of Intéressante	Brown-Forsythe	1.132	2	111.993	.326
Rank of Claire	Brown-Forsythe	6.089	2	104.401	.003
Rank of Compris	Brown-Forsythe	3.291	2	87.056	.042
Rank of Jaime	Brown-Forsythe	13.893	2	107.661	.000

a. Asymptotically F distributed.

De nouveau, puisque la comparaison est multiple, nous devons réaliser une correction de Bonferroni. Nous réaliserons trois comparaisons. Notre seuil de signification ajusté est donc de 0,016 ($\alpha' = 0,05/3$).

1. Entre « Bas » et « Moyen »

Comme en atteste la dernière colonne du Tableau 29, seulement deux *p-values* sont significatives. Elles sont relatives aux *items* « Belle » et « Claire ». Elles sont inférieures à 0,016, ce qui nous pousse à rejeter H0 selon laquelle les moyennes des rangs sont similaires. Ces variables de l’attribut « *Data-Ink* » ont donc un effet sur ces deux *items*.

En ce qui concerne l’*item* « Belle », les visualisations dont le *data-ink* est bas ont suscité davantage d’appréciations positives que les visualisations ayant un *data-ink* moyen (moyennes des rangs : 581,85 contre 517,69).

En ce qui concerne l'*item* « Claire », les visualisations dont le *data-ink* est bas ont suscité nettement moins d'appréciations positives que les visualisations ayant un *data-ink* moyen (moyennes des rangs : 284,52 contre 421,30).

En ce qui concerne les autres *items*, à savoir « Intéressante », « Compris » et « J'aime », nous ne pouvons rejeter H0 et donc nous ne pouvons conclure que ces variables ont un effet sur ces appréciations.

Tableau 29 Test de Brown Forsythe sur les variables « Bas » et « Moyen » de l'attribut "Data-Ink"

Robust Tests of Equality of Means					
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	7.451	1	107.436	.007
Rank of Intéressante	Brown-Forsythe	1.591	1	62.794	.212
Rank of Claire	Brown-Forsythe	10.863	1	58.643	.002
Rank of Compris	Brown-Forsythe	5.755	1	52.582	.020
Rank of Jaime	Brown-Forsythe	1.742	1	60.760	.192

a. Asymptotically F distributed.

2. Entre « Moyen » et « Élevé »

Le test de Brown Forsythe se montre significatif à propos de ces variables en ce qui concerne trois items : « Belle », « Claire » et « J'aime » (Tableau 30). Toutes ces *p-values* sont inférieures à 0,016 et nous rejetons donc l'hypothèse d'égalité des moyennes des rangs.

En ce qui concerne l'*item* « Belle », les visualisations dont le *data-ink* est moyen suscitent davantage d'appréciations positives que les visualisations au *data-ink* élevé (moyennes des rangs : 517,69 contre 333,99).

En ce qui concerne l'*item* « Claire », les visualisations dont le *data-ink* est moyen suscitent davantage d'appréciations positives que les visualisations au *data-ink* élevé (moyennes des rangs : 421,30 contre 368,28).

En ce qui concerne l'*item* « J'aime », les visualisations dont le *data-ink* est moyen suscitent davantage d'appréciations positives que les visualisations au *data-ink* élevé (moyennes des rangs : 472,53 contre 351,72).

Pour les autres *items*, à savoir « Intéressante » et « Compris », nous ne pouvons donc estimer qu'il y ait un tel effet puisque le test n'est pas significatif.

Tableau 30 Test de Brown Forsythe sur les variables « Moyen » et « Élevé » de l'attribut « Data-Ink »

Robust Tests of Equality of Means					
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	97.860	1	148.816	.000
Rank of Intéressante	Brown-Forsythe	.002	1	128.170	.965
Rank of Claire	Brown-Forsythe	6.341	1	135.674	.013
Rank of Compris	Brown-Forsythe	2.622	1	138.927	.108
Rank of Jaime	Brown-Forsythe	29.498	1	130.010	.000

a. Asymptotically F distributed.

3. Entre « Bas » et « Élevé »

Cette-fois, nous constatons que seule la *p-value* relative à l'*item* « Belle » est significative, puisqu'inférieure à 0,016 avec une valeur de 0. Pour cet *item*, nous rejetons H0 supposant l'égalité des moyennes des rangs. L'effet est donc à commenter. Les visualisations dont le *data-ink* est bas sont jugées significativement plus belles (davantage d'appréciations positives) que les visualisations au *data-ink* élevé (moyennes des rangs : 531,58 contre 333,95).

En ce qui concerne les autres *items*, à savoir « Intéressante », « Claire », « Compris » et « J'aime », les *p-values* ne sont pas significatives et nous ne pouvons donc pas constater que les variables étudiées dans ce cas aient un effet.

Tableau 31 Test de Brown Forsythe sur les variables « Bas » et « Élevé » de l'attribut "Data-Ink"

Robust Tests of Equality of Means					
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	179.051	1	57.240	.000
Rank of Intéressante	Brown-Forsythe	2.091	1	40.782	.156
Rank of Claire	Brown-Forsythe	4.947	1	40.907	.032
Rank of Compris	Brown-Forsythe	3.230	1	40.045	.080
Rank of Jaime	Brown-Forsythe	2.812	1	40.681	.101

a. Asymptotically F distributed.

7.3. Falsification des hypothèses

Nous pouvons désormais définitivement falsifier et commenter nos hypothèses pour cet attribut.

H31 Il n'y a pas de différence significative entre les variables « Bas » et « Moyen » de l'attribut « *Data-Ink* » en ce qui concerne les *items* « Intéressante », « Claire » et « Compris ».

➔ **L'hypothèse est rejetée.** Elle se vérifie pour les *items* « Intéressante » et « Compris », mais pas pour l'*item* « Claire ». Effectivement, ces variables ont un effet sur la clarté déclarée par les participants : les visualisations au *data-ink* moyen suscitent davantage d'appréciations positives que les visualisations au *data-ink* élevé.

H32 Les visualisations au *data-ink* moyen reçoivent davantage d'appréciations positives pour l'*item* « Belle » que les autres variables.

➔ **L'hypothèse est rejetée** puisque les visualisations dont le *data-ink* est bas ont suscité davantage d'appréciations positives que les visualisations ayant un *data-ink* qualifié de moyen. Par contre, les visualisations dont le *data-ink* est moyen suscitent davantage d'appréciations positives que les visualisations au *data-ink* qualifié d'élevé.

H33 Les visualisations au *data-ink* moyen reçoivent davantage d'appréciations positives pour l'*item* « J'aime » que les autres variables.

➔ **L'hypothèse est rejetée** puisque la comparaison entre les variables « Bas » et « Moyen » n'a pas été significative concernant cet *item*. Par contre, les visualisations au *data-ink* moyen suscitent davantage d'appréciations positives que les visualisations au *data-ink* élevé.

H34 Les visualisations au *data-ink* bas reçoivent davantage d'appréciations positives pour l'*item* « Belle » que celles dont le *data-ink* est qualifié de moyen et élevé.

➔ **L'hypothèse est validée.**

H35 Les visualisations au *data-ink* bas reçoivent moins d'appréciations positives pour l'*item* « Intéressante » que celles dont le *data-ink* est qualifié de moyen et élevé.

➔ **L'hypothèse est rejetée** puisque les tests ne se sont pas révélés significatifs.

H36 Les visualisations au *data-ink* bas reçoivent moins d'appréciations positives pour l'*item* « Compris » que celles dont le *data-ink* est qualifié de moyen et élevé.

➔ **L'hypothèse est rejetée** : les tests ne sont pas significatifs concernant cet *item*.

8. Effet de l'attribut « *Embellished Data-Ink* » sur les appréciations déclarées

8.1. Formulation d'hypothèses

L'attribut « *Embellished Data-Ink* » concerne le *data-ink* qui a été ornementé lors de la phase de conception de la visualisation de données. On le quantifie donc ici. En observant les *box plots* des différentes variables de cet attribut sur la Figure 35, nous pouvons formuler des hypothèses.

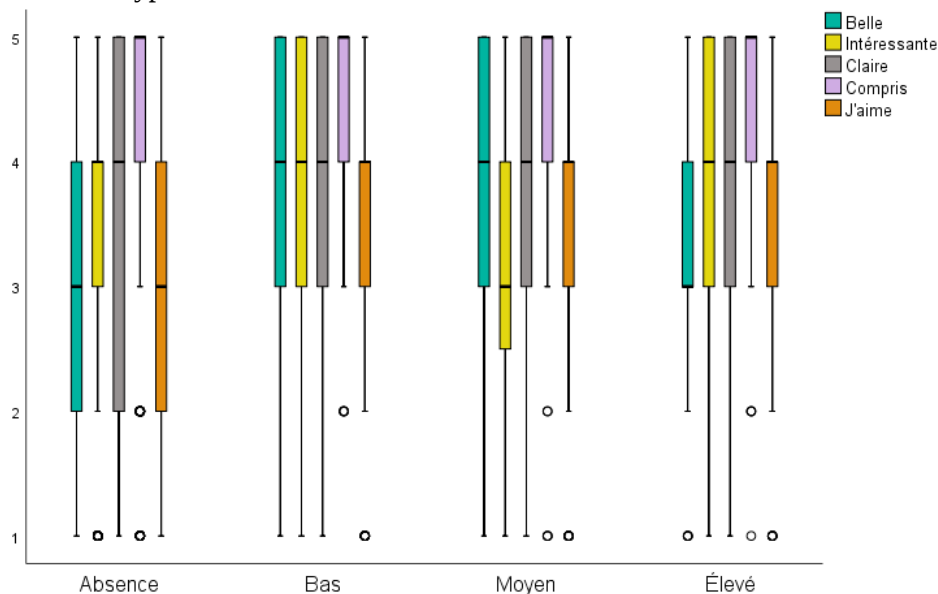


Figure 35 Box Plots représentant les distributions des variables de l'attribut "Embellished Data-Ink"

Avant tout, on voit que cet attribut n'a potentiellement pas d'influence sur l'item « Compris ». Par contre, on voit qu'il a un effet positif sur l'item « J'aime ». Il ne semble pas avoir de différence significative entre les visualisations présentant de l'*embellished data-ink*, bien qu'il semblerait que les visualisations correspondant à la variable « Absence » suscitent moins d'appréciations positives pour cet item. L'intérêt peut être commenté lui-aussi, et il semblerait que les visualisations avec l'attribut « Moyen » suscitent le plus d'appréciations positives en termes de beauté. Nous formulons ainsi les hypothèses suivantes.

- H37** Il n'y a pas de différence significative pour l'item « Belle » en ce qui concerne les variables « Bas » et « Moyen » de « *Embellished Data-Ink* ».
- H38** L'absence de *embellished data-ink* suscite nettement moins d'appréciations positives en termes de beauté que les autres variables.
- H39** L'absence de *embellished data-ink* suscite nettement moins d'appréciations positives en termes de clarté.

- H40** Il n'y a pas de différence significative en termes de clarté en ce qui concerne toutes les variables manifestant la présence de *data-ink*.
- H41** L'absence de *embellished data-ink* suscite nettement moins d'appréciations positives en ce qui concerne l'item « J'aime » que les autres variables.
- H42** L'intérêt récolte significativement moins d'appréciations positives lorsque l'*embellished data-ink* est qualifié de « Moyen ».
- H43** L'intérêt suscite davantage d'appréciations positives lorsqu'il y a présence d'*embellished data-ink*.

8.2. Analyse approfondie

Les tests dont les résultats figurent dans le Tableau 32 nous montre qu'on ne peut rejeter H_0 selon laquelle les moyennes des rangs des *items* « Intéressante », « Claire » et « Compris » sont similaires. Effectivement, leurs *p-values* sont supérieures à 0,05, ce qui nous laisse penser que cet attribut n'a pas d'effet sur ces *items*. Pour la suite, nous ne procéderons à des comparaisons multiples que pour les *items* « Belle » et « J'aime » qui présentent une *p-value* de 0 qui est significative. Nous rejetons donc H_0 pour ces deux *items*. Nous pouvons donc procéder aux comparaisons des variables pour commenter l'effet de cet attribut sur les *items* « Belle » et « J'aime ».

Tableau 32 Test de Brown Forsythe sur l'attribut "Embellished Data-Ink"

Robust Tests of Equality of Means		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	48.725	3	277.191	.000
Rank of Intéressante	Brown-Forsythe	2.411	3	250.635	.067
Rank of Claire	Brown-Forsythe	.457	3	256.423	.713
Rank of Compris	Brown-Forsythe	1.712	3	263.223	.165
Rank of Jaime	Brown-Forsythe	17.456	3	257.234	.000

a. Asymptotically F distributed.

Nous allons procéder à six comparaisons. Notre seuil de signification ajusté selon la correction de Bonferroni sera de 0,012 ($\alpha' = 0,05/4$)

1. Entre « Absence » et « Bas »

Les *p-values* apparaissant dans le Tableau 33 sont inférieures à 0,012 et sont donc significatives. Nous rejetons donc H0 selon laquelle les moyennes des rangs de ces variables pour ces deux *items* sont similaires. Il y a ainsi bien un effet que nous pouvons commenter. Dans les deux cas, les visualisations ayant un *embellished data-ink* bas engendrent davantage d'appréciations positives que les visualisations sans ornementation centrée sur le *data-ink* et ce, pour les deux *items*.

Moyenne des rangs pour « absence », puis « bas » :

- « Belle » : 322,74 contre 503,82
- « J'aime » : 336,77 contre 450,99

Tableau 33 Test de Brown Forsythe pour les variables « Absence » et « bas » de l'attribut "Embellished Data-Ink"

Robust Tests of Equality of Means					
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	74.610	1	115.655	.000
Rank of Jaime	Brown-Forsythe	22.000	1	103.423	.000

a. Asymptotically F distributed.

2. Entre « Absence » et « Moyen »

De nouveau, les *p-values* sont significatives et ne nous permettent pas de valider H0. Les moyennes des rangs ne sont donc pas similaires.

Pour les deux *items*, les visualisations correspondant à la variable « Moyen » récoltent davantage d'appréciations positives que la variable « Absence ». Les moyennes des rangs sont celles-ci : pour « Belle », 322,74 contre 503,82 et pour l'*item* « J'aime », 336,77 contre 450,99.

Tableau 34 Test de Brown Forsythe pour les variables « Absence » et « moyen » de l'attribut "Embellished Data-Ink"

Robust Tests of Equality of Means					
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	77.603	1	116.463	.000
Rank of Jaime	Brown-Forsythe	22.817	1	105.471	.000

a. Asymptotically F distributed.

3. Entre « Absence » et « Élevé »

De nouveau, les *p-values* présentées dans le Tableau 35 sont significatives puisqu'elles s'élèvent à 0. Nous rejetons donc H0 selon laquelle les moyennes des rangs de « Absence » et de « Élevé » sont similaires. Ainsi, les visualisations correspondant à la seconde option suscitent davantage d'appréciations positives que celles correspondant à l'absence d'*embellished data ink* (moyennes des rangs, respectivement : « Belle » avec 322,74 contre 436,15 et « J'aime » avec 336,77 contre 459,83).

Tableau 35 Test de Brown Forsythe pour les variables « Absence » et « Élevé » de l'attribut « Embellished Data-Ink »

		Robust Tests of Equality of Means			
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	22.646	1	76.962	.000
Rank of Jaime	Brown-Forsythe	20.069	1	70.987	.000

a. Asymptotically F distributed.

4. Entre « Bas » et « Moyen »

Cette fois, nous constatons que les *p-values* qui se trouvent dans la dernière colonne du Tableau 36 sont supérieures à 0,012. Ainsi, nous ne rejetons pas H0 selon laquelle les moyennes des rangs sont similaires. Ces variables n'ont donc potentiellement pas d'effet sur les *items* « Belle » et « J'aime ».

Tableau 36 Test de Brown Forsythe pour les variables « Bas » et « Moyen » de l'attribut « Embellished Data-Ink »

		Robust Tests of Equality of Means			
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	.032	1	156.999	.858
Rank of Jaime	Brown-Forsythe	.001	1	156.940	.980

a. Asymptotically F distributed.

5. Entre « Bas » et « Élevé »

De nouveau, nous constatons que les *p-values* sises dans la dernière colonne du Tableau 37 sont supérieures à 0,012. Ainsi, nous ne rejetons pas H0 selon laquelle les moyennes des rangs sont similaires. Ces variables n'ont donc potentiellement pas d'effet sur les *items* « Belle » et « J'aime ».

Tableau 37 Test de Brown Forsythe pour les variables « bas » et « élevé » de l'attribut "Embellished Data-Ink"

		Robust Tests of Equality of Means			
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	5.403	1	123.102	.022
Rank of Jaime	Brown-Forsythe	.066	1	124.194	.798

a. Asymptotically F distributed.

6. Entre « Moyen » et « Élevé »

Nous poursuivons dans la même dynamique puisque les variables « Moyen » et « Élevé » n'ont potentiellement pas d'effet sur les *items* « Belle » et « J'aime ». Nous constatons effectivement que les *p-values* sises dans la dernière colonne du Tableau 38 sont supérieures à 0,012 et nous ne rejetons pas H0 selon laquelle les moyennes des rangs sont similaires.

Tableau 38 Test de Brown Forsythe pour les variables « Moyen » et « Élevé » de l'attribut « Embellished Data-Ink »

		Robust Tests of Equality of Means			
		Statistic ^a	df1	df2	Sig.
Rank of Belle	Brown-Forsythe	6.160	1	124.120	.014
Rank of Jaime	Brown-Forsythe	.055	1	123.748	.815

a. Asymptotically F distributed.

8.3. Falsification des hypothèses

Après avoir mené l'analyse, nous sommes en mesure de répondre aux hypothèses plus générales et de les commenter.

H37 Il n'y a pas de différence significative pour l'*item* « Belle » en ce qui concerne les variables « bas » et « moyen » de « *Embellished Data-Ink* ».

➔ **L'hypothèse est validée.** Nous n'avons effectivement pas constaté d'effet de ces variables sur les *items*. Il n'y a donc presque aucune différence entre elles.

H38 L'absence de *embellished data-ink* suscite nettement moins d'appréciations positives en termes de beauté que les autres variables.

➔ **L'hypothèse est validée.**

H39 L'absence de *embellished data-ink* suscite nettement moins d'appréciations positives en termes de clarté.

➔ **L'hypothèse est rejetée.** Les tests ne nous permettent pas de l'affirmer ; ainsi, l'*embellished data-ink* n'a pas particulièrement d'effet sur la clarté.

H40 Il n'y a pas de différence significative en termes de clarté en ce qui concerne toutes les variables manifestant la présence de *data-ink*.

➔ **L'hypothèse est validée,** mais nous y ajoutons qu'il n'y a pas de différence non plus avec l'absence de *embellished data-ink*.

H41 L'absence de *embellished data-ink* suscite nettement moins d'appréciations positives en ce qui concerne l'*item* « J'aime » que les autres variables.

➔ **L'hypothèse est validée**

H42 L'intérêt récolte significativement moins d'appréciations positives lorsque l'*embellished data-ink* est moyen.

➔ **L'hypothèse est rejetée :** nous ne constatons pas d'effet de l'intérêt sur les *items* en ce qui concerne les variables étudiées ici.

H43 L'intérêt suscite davantage d'appréciations positives lorsqu'il y a présence d'*embellished data-ink*.

➔ **L'hypothèse est rejetée :** nous ne constatons pas d'effet de l'intérêt sur les *items* en ce qui concerne les variables étudiées ici.

Chapitre 3 Interprétation des résultats et discussion

À présent, place à l'interprétation des résultats de l'analyse. Pour commencer, nous réalisons une interprétation factuelle, rendue possible par la falsification des hypothèses abordées par l'analyse. Puis, nous discuterons les apports de l'analyse des appréciations déclarées : peut-on déjà entrevoir une recommandation en termes de design et que penser de l'ornementation des visualisations de données ?

1. Interprétation factuelle

L'analyse statistique des appréciations déclarées nous a permis d'élaborer des hypothèses sur base des distributions des appréciations en fonction des attributs que nous avons assignés aux visualisations de données. Les résultats de cette analyse sont parfois à contrecourant de nos a priori et prouvent leur utilité. Il est maintenant temps de commenter plus largement les réponses à nos hypothèses.

Pour commencer, nous assumons, à la suite de cette analyse, avoir distingué les *items* les uns des autres par une dimension importante, à savoir le fait qu'ils puissent être qualifiés d'affectifs ou réflexifs. Ainsi, nous considérons que les *items* « J'aime » et « Belle » sont des éléments que nous qualifions d'affectifs tandis que les *items* « Intéressante », « Claire » et « Compris » relèvent davantage d'éléments réflexifs. Cela tient tout son sens puisque les appréciations relatives à l'*item* « Belle » sont corrélées aux appréciations relatives à l'*item* « J'aime ». Au vu des déclarations des participants, il semblerait ainsi que la beauté estimée d'une visualisation de données aille de pair avec le fait de l'apprécier ou pas.



Figure 36 Visualisation n°7 du corpus

Mais quand-est-ce qu'une visualisation est jugée « belle » et en fonction de quel attribut ? Quel est l'impact des attributs sur la corrélation entre les *items* « Belle » et « J'aime » ? Nous pouvons dès à présent nous pencher sur cette question. De manière générale, les visualisations de données ornementées sont jugées plus belles et sont favorisées aux

visualisations de type standard. Ainsi, une visualisation ornementée a davantage de chances d'être estimée « Belle » et d'être aimée en fonction de ce critère qu'une visualisation standard. Néanmoins, s'il préfère une visualisation ornementée à une visualisation de type standard, le participant n'est pas dupe de l'apparence de ces dernières. Il a bien des modes de préférences dans les visualisations ornementées qui lui sont présentées et, quasi systématiquement, il préférera les visualisations dont les ornements sont des icônes. Il jugera effectivement moins belles les visualisations de données à la forme inhabituelle, tout en les jugeant moins intéressantes. Ainsi, on peut comprendre que l'individu aime l'ornementation qui reste dans des codes qu'il maîtrise. Il y a bien évidemment différentes façons de placer l'icône : sur le *data-ink* ou non, de façon étirée, superposée, ... Cela est étudié grâce aux autres attributs. Nous pouvons par exemple prendre la visualisation n°7 (Figure 36), très appréciée en termes de beauté, et dont l'ornement se constitue dans les icônes. Cette visualisation de données a l'allure d'un bar chart et l'information est représentée dans les icônes que sont les bouteilles de champagne. Cette visualisation est jugée plus belle, plus intéressante et est plus appréciée que son homologue (visualisation standard, type *Raw*). A contrario, même si les individus préfèrent l'ornementation de manière générale, ils préféreront lire une visualisation standard à une visualisation dont la forme est inhabituelle ou fantaisiste. Ainsi, la visualisation n°10 de notre corpus (Figure 37) était assez peu appréciée en regard de son homologue plus traditionnelle. L'attribut « ornement » est donc très instructif : l'ornementation est appréciée par l'individu, mais pas n'importe comment. L'icône, en général plus simple qu'une forme inhabituelle ou la mise en place d'une métaphore, est très apprécié.

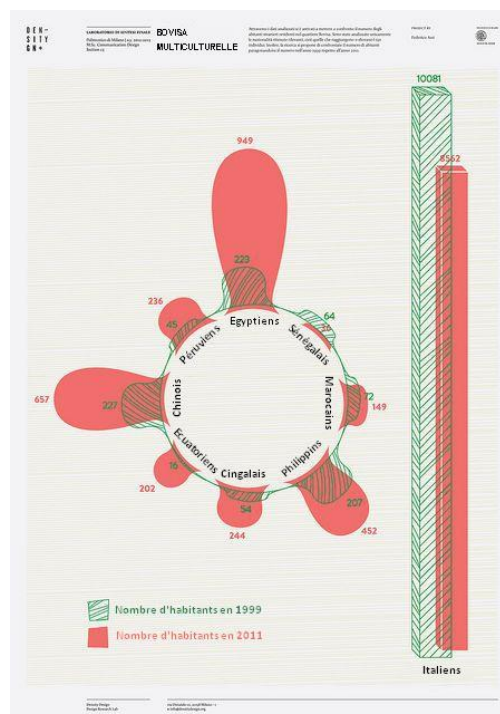


Figure 37 Visualisation n°10 du corpus

Comme nous l'avons mentionné, l'icône peut être placée sur les traits relatifs aux données (*data-ink*). Elle est alors le corps de l'information, elle la représente. L'icône peut également être disposée à côté d'un graphe simple, ce que nous avons qualifié d'« apposé ». Nous avons ainsi pu étudier les appréciations relatives au type de figuration : la figuration, à savoir la manière d'ornementer grâce à l'image, peut se constituer sur le *data-ink* ou de manière apposée. Comme l'a montré l'hypothèse n°18, les visualisations qui sont jugées les plus « Belles » par les participants sont celles dont la figuration est apposée. Elles s'opposent donc à celles qui ont une forme inhabituelle et dont l'icône se constitue en *data-ink*. Lorsque la figuration est apposée, l'ornementation se situe à côté d'une forme de graphe plus conventionnelle. Elle décore, sans constituer l'information principale. C'est ce type de visualisation qui est davantage apprécié. Ainsi, la visualisation n°17 (Figure 38) de notre corpus sera jugée plus belle et sera plus appréciée que la visualisation n°7, où la figuration constitue l'information en étant le *data-ink*. Néanmoins, lorsque l'ornementation se constitue sur les traits relatifs aux données, nous avons vu que les individus préfèrent l'icône à la forme inhabituelle et fantaisiste. Il y a de nouveau une autre façon de connaître les préférences des participants concernant l'icône, car, figuré sur la représentation des données ou non, il peut être représenté différemment. Nous souhaitons alors étudier les différentes appréciations des visualisations de données en fonction de l'attribut « Pictogramme ». Contre toute attente, les types de pictogrammes n'ont pas d'effet sur les appréciations que nous avons relevées. Cela est sans doute dû au fait qu'ils soient peu diversifiés dans notre corpus.

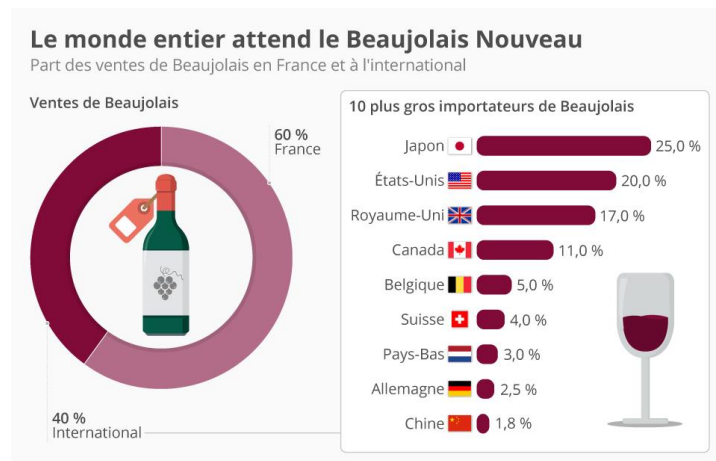


Figure 38 Visualisation n°17 du corpus

En ce qui concerne l'attribut « *Data-Ink* », c'est assez simple. Nous l'avons évalué en pourcentages et nous avons ensuite réalisé des classes en fonction du pourcentage. Moins il y a de *data-ink* sur l'image, plus elle est jugée belle. Dans notre corpus, il y a 20 visualisations où le *data-ink* frôle les 100% (haut). Effectivement, toute l'encre du graphique ne concerne que les données et uniquement les données et en constitue

l'information. En général les appréciations de beauté sont très faibles pour ces visualisations de données.

Nous pouvons maintenant aborder les autres *items* que nous considérons réflexifs (« Clair », « Compris », « Intéressant »). Penchons-nous donc sur l'intérêt, la clarté et la compréhension.

Les tests montrent ainsi que les visualisations de données de type « *Raw* » sont jugées plus claires que les visualisations de données de type « *Junk* ». L'ornementation diminue donc la sensation de clarté, à savoir « l'information saisissable en un coup d'œil » du graphique. Par contre, le type « *Raw* » ou « *Junk* » n'a pas d'incidence sur la compréhension et l'intérêt. C'est du moins ce qu'ont déclaré les participants. Cela montre que les participants établissent une distinction entre la clarté du visuel, que nous considérons comme le caractère intelligible de la visualisation, et la compréhension, qui nécessite un effort d'interprétation.

Lorsque l'ornementation des visualisations de données est réalisée sous forme d'icônes, la clarté et la compréhension ne semblent pas impactées. Par contre, les participants ont jugé que la présence d'icônes sur une visualisation de données la rend plus intéressante. De même, en termes de clarté, de compréhension et d'intérêt, les participants ont déclaré une grande préférence pour l'icône, en comparaison à la forme inhabituelle. Ainsi, nous avons vu précédemment que les individus préfèrent une absence d'ornement à la forme inhabituelle et qu'ils estiment que cette forme de graphique diminue la clarté, la compréhension et l'intérêt des données.

Poursuivons avec le type de figuration qui n'a aucun impact sur l'intérêt et la compréhension des participants, selon leurs déclarations. Néanmoins, le résultat des tests s'est montré différent pour l'*item* « Clarté ». Effectivement, il semblerait que l'ornementation apposée à un graphique plus traditionnel soit plus claire pour les participants qu'une ornementation qui se trouve sur le *data-ink*. A titre d'exemple, la visualisation n°17 de notre corpus est susceptible d'être préférée par les participants parce que les icônes se situent à côté des graphiques conventionnels, en comparaison à la visualisation n°18 qui montre une ornementation concentrée sur le *data-ink* (Figure 39).

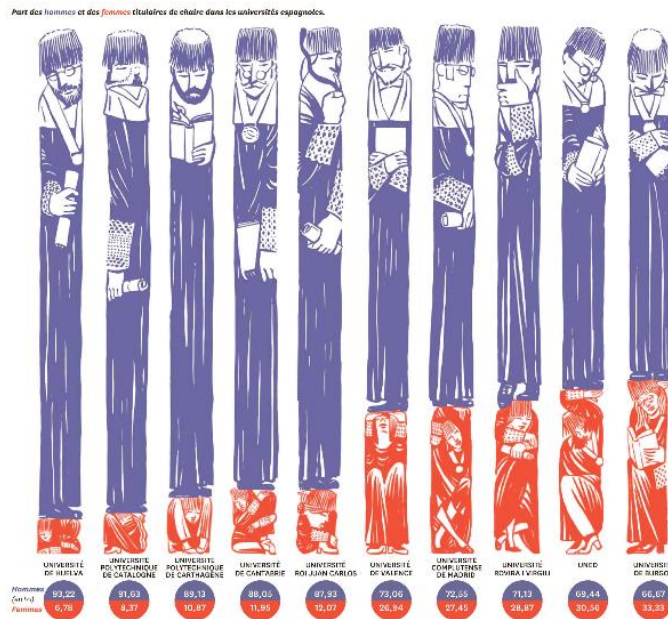


Figure 39 Visualisation n°18 du corpus

Le volume de *data-ink* présent dans le graphe impacte seulement la clarté : les visualisations dont le *data-ink* est moyen semblent plus claires aux participants en comparaisons aux visualisations au haut niveau de *data-ink*. Ainsi, une partie de l'encre qui n'est pas dédiée aux données (donc les ornements) est susceptible de jouer un rôle dans la clarté estimée par les participants. De même, le fait que le *data ink* soit ornementé (« *embellished data-ink* ») n'a pas d'effet significatif sur la compréhension, la clarté et l'intérêt.

Malheureusement, notre analyse prouve encore, comme cela a déjà été réalisé à de maintes reprises dans la littérature, que le *lie factor* distorsif constitue un danger quant à l'appropriation de l'information. Ainsi, les visualisations de données présentant une distorsion visuelle, comme la visualisation n°14 de notre corpus (Figure 40), sont jugées plus belles et sont davantage appréciées que leurs homologues à la forme standard. En analysant toutes les appréciations en fonction des *items*, nous réalisons que le risque est grand parce que le *lie factor* a un impact considérable sur les appréciations affectives, à savoir les *items* « Belle » et « J'aime ». Les appréciations



Figure 40 Visualisation n°14 du corpus

sont en général assez bonnes, sauf si un attribut que nous avons identifié provoque une baisse du score de l'appréciation. Les appréciations réflexives, quant à elles, sont révélatrices d'un fait important et sont à mettre en balance avec l'aspect affectif. L'hypothèse n°33 a effectivement été validée en ce qui concerne la clarté : de manière générale, les individus ont estimé que les visualisations de données au *lie factor* distorsif étaient moins claires que les autres visualisations qui leur ont été présentées. Pourtant, ils n'estiment pas pour autant que ces visualisations soient moins compréhensibles. Le *lie factor* distorsif n'a pas d'impact sur l'*item* « Compris ». Pour des raisons affectives et parce qu'ils préfèrent les visualisations de données au *lie factor* distorsif, qu'ils jugent plus belles, les participants sont prêts à tolérer un manque de clarté, en pensant que cela n'affecte pas la compréhension. Les participants n'établissent pas de lien clair entre l'intelligibilité visuelle et la compréhension et sont distraits par d'autres facteurs affectifs. Il se pourrait ainsi que les participants pensent avoir compris l'information sans que ce ne soit réellement le cas. Le *lie factor* n'a par contre aucun effet sur l'intérêt déclaré par les participants pour les visualisations présentées.

2. Discussion

L'analyse statistique que nous avons menée nous a permis de tirer des conclusions déterminantes pour la suite de cette thèse. Elles se discutent malgré tout. Effectivement, à ce stade de la recherche, certaines d'entre elles comportent quelques limites importantes, tandis que d'autres confirment et enrichissent les études menées au sujet de l'ornementation. Il est bon de se rappeler que les appréciations sont déclarées par les participants et qu'elles reflètent donc leur propre sensation, qui est tout à fait personnelle. Notre étude ne montre pas si les participants ont compris une visualisation de données ou non. Elle fait état de leur sensation, de leurs impressions. Peut-être ont-ils ainsi la sensation d'avoir compris une visualisation de données sans pour autant que le message initialement encodé par le concepteur n'ait parfaitement été décodé par l'utilisateur. Les appréciations déclarées, récoltées systématiquement et directement après la lecture de chaque visualisation, sont un indicateur quasi-spontané des premières impressions ressenties par les participants.

2.1. Vers une recommandation essentielle en conception de visualisations de données

Le premier constat à tirer concerne la beauté des visualisations de données, inextricablement liée au fait d'aimer une visualisation de données ou pas. Ainsi, il existe peu de cas de figure dans lesquels des individus estiment une visualisation de données belle,

sans l'aimer du tout. La première impression visuelle, la sensation de beauté, joue incontestablement un rôle dans le fait d'apprécier une visualisation ou pas. Souvenons-nous par ailleurs des facteurs humains et émotionnels de Kennedy et Hill (2017, 2018) et plus particulièrement du fait qu'un individu soit prédisposé à s'engager ou pas envers une visualisation de données en fonction de son apparence visuelle. Nous sommes ici confrontés à l'évidence statistique montrant que les individus préfèrent les visualisations qu'ils estiment jolie. Bien que tout à fait subjective, la beauté évoque des critères esthétiques partagés par ces individus issus du même monde professionnel. Nous pouvons sans aucun problème supposer qu'une apparence visuelle plaisante prédispose l'individu à s'engager dans l'interaction avec cette visualisation de données.

Bien que Tufte craigne le *chartjunk* et les efforts artistiques qu'il estime être une distraction, nous constatons à travers les appréciations déclarées de nos participants un certain recul critique. Ils ne préfèrent pas à tout prix les visualisations de données ornementées. Lorsque l'intelligibilité, qui transparait à travers l'*item* « Clarté », se voit fortement impactée, le participant se distancie de la visualisation de données, lui donnant une cote plus basse en termes d'appréciation. Là où les efforts esthétiques sortent de l'habitude, les craintes de Tufte se confirment : l'information est difficilement saisissable et peu claire. L'utilisateur n'est néanmoins pas un récepteur dénué de bon sens : les visualisations à la forme inhabituelle ne sont ni aimées, ni jugées claires, ni même belles ! Les visualisations standard, simples, sont d'ailleurs jugées plus belles. Pour Tufte, c'est de la présentation des données présentées en toute simplicité qu'émerge la beauté. Cela se vérifie en partie. Néanmoins, nous avons pu montrer l'importance des icônes qui, selon les individus, n'impactent pas leur sensation de clarté et de compréhension, tout en augmentant leur sensation de beauté. Nous pouvons ainsi conclure que l'icône prédispose l'individu à s'engager envers la visualisation de données. Mais de nouveau, notre étude tout en granularité montre encore le discernement dont font preuve les participants. L'icône peut être apposée à côté d'un graphique standard ou bien constituer le corps même du graphique (être le *data-ink*). L'icône apposée est préférée à celle qui se constitue en *data-ink* et elle est aussi jugée plus belle. Si l'icône qui constitue le *data-ink* est quant à elle jugée plus belle que l'absence d'icônes, et donc que les graphiques standard, il s'avère qu'elle diminue la clarté de l'information. Il n'y a pourtant pas de différences en termes de clarté en ce qui concerne les graphiques standards ou les graphiques ornementés par des icônes apposées.

Tout ceci nous permet d'entrevoir une recommandation en termes de conception de visualisations de données, bien qu'elle soit à confirmer grâce à la suite de nos analyses. Au vu de nos précédentes conclusions, l'utilisateur est plus enclin à s'engager envers une visualisation qu'il trouve belle et qu'il aimera. L'ornementation favorise ces sensations de

beauté et d'affection, tandis que l'icône apposée ne diminue pas la sensation de clarté ressentie par les participants en comparaison à une visualisation de type standard. Ainsi, en termes d'ornementation, il semblerait que l'utilisation d'icônes apposées à un graphique plutôt standard soit à favoriser afin de stimuler des émotions positives chez l'individu tout en assurant un bon niveau d'intelligibilité.

Selon nous, cette préconisation n'entre pas en conflit avec le concept de *lie factor* émis par Tufte. Que du contraire, elle aurait plutôt tendance à la renforcer. Pour rappel, ce facteur sert à montrer la relation entre l'effet montré par le graphique et l'effet montré par les données. Si l'effet montré par le graphique est plus grand ou plus petit que celui présent dans les données, il n'est pas fidèle et apporte donc une distorsion amenant à tromper l'utilisateur. Le risque évident que nous avons constaté dans cette étude est que, dans le cas de visualisations de données ornementées, les utilisateurs estiment les visualisations de données au *lie factor* distorsif³³ sont moins claires que celles dont le *lie factor* reflète fidèlement la relation entre l'effet montré dans le graphique et l'effet exprimé par les données. Néanmoins, ces visualisations de données au *lie factor* distorsif sont elles aussi jugées belles et sont donc forcément aimées, ce qui prédispose l'utilisateur à s'engager positivement envers la visualisation de données alors même qu'il ne se rend pas compte qu'elle n'est pas claire. Nous constatons que bien souvent, c'est lorsque l'ornementation se concentre sur le *data-ink*, et donc sur les traits relatifs aux données, que le *lie factor* est distorsif. Avant tout, l'œil humain est mauvais juge en termes d'évaluation de quantités représentées par des aires. Ainsi, comment estimer si une icône constituée en *data-ink*, par exemple une boule de glace ou une bouteille de champagne dans notre corpus, puisse représenter fidèlement une quantité ? De même les visualisations dont l'icône se constitue sur le *data-ink* ont été jugées moins claires que les visualisations de données standard ou encore celles dont l'icône est apposée. En préconisant une ornementation utilisant des icônes apposées à un graphique standard, nous misons sur la prudence en évitant les probables distorsions visuelles que peuvent amener les icônes qui se constituent en *data-ink*, en se distançant ainsi d'un *lie factor* pouvant facilement être distorsif. Évidemment, il est indispensable que le graphique standard (type diagramme en bâtonnets par exemple) respecte lui-même la notion de *lie factor* avant même d'y apposer une icône.

Il est vrai, par contre, que notre prescription en termes de design entre en conflit avec le principe de *data-ink ratio maximisation* de Tufte. Notre analyse montre effectivement que les visualisations préférées par les participants sont celles dont le *data-ink ratio* est de niveau moyen. Cela tient tout son sens : les visualisations de données au *data-ink* de bas

³³ Donc dont la relation entre l'effet montré par le graphique et celui montré par les données n'est pas fidèle

niveau sont en général des visuels très chargés, ce qui encombre la clarté. Les visualisations de données au *data-ink* dont le niveau est élevé sont en général soit de type standard donc sans aucune ornementation ou sont des formes inhabituelles ou encore des visualisations de données ornementées dont le *data-ink* est représenté par des icônes. De nouveau, nous comprenons que les visualisations au *data-ink* moyen, qui en général représentent des graphiques simples auxquels sont apposées des icônes, semblent très claires aux participants. Notre analyse consolide donc notre prescription qui est d'employer des icônes apposées en termes d'ornementation. De même, rappelons les études de Inbar et ses collègues (Inbar et al., 2007) et Hill et ses pairs (Hill et al., 2017), selon lesquelles le minimalisme en visualisation de données perturbe grandement l'utilisateur. Ainsi, tout comme nous, ils remettent en question le principe de maximisation du *data-ink ratio*.

2.2. La place d'une telle analyse dans la littérature sur l'ornementation

Nous pouvons déjà, de manière anticipée, replacer cette analyse parmi les études déjà évoquées concernant l'ornementation. C'est l'occasion de mentionner en quoi nos conclusions renforcent certaines contributions passées, en quoi elles s'en distancient et surtout quels sont leurs apports.

Les recherches sur l'ornementation en visualisation de données ont étudié beaucoup de comportements différents de l'utilisateur. Effectivement, que dire de l'attractivité, de la mémoire ou encore de la compréhension d'une visualisation de données ? Les études à ce sujet sont foisonnantes.

Bateman et ses collègues (2010) ont montré que les visualisations ornementées peuvent être comprises de la même façon que les visualisations minimalistes. Ils ont pour cela évalué le niveau de compréhension des participants, ainsi que leur mémorisation : les participants, qui étaient des étudiants, se souvenaient mieux des visualisations ornementées. Borgo et ses collègues (2012) ont eux aussi évolué dans le même sens, en montrant l'impact positif sur la mémoire à long terme. Comme nous l'avons précédemment mis en avant, Borkin et ses pairs, dans deux articles (2013 ; 2016) ont exploré quelles parties d'un graphique sont mémorisables : une haute densité visuelle, un *data-ink ratio* assez bas et surtout des objets reconnaissables par l'œil humain, qui constituent en réalité le cœur de la mémorisation. Ces objets aident l'individu à produire un souvenir, même après dix secondes (Borkin et al., 2016). Puisque nous avons relevé des appréciations, nous pouvons compléter ces études, en apportant une valeur ajoutée.

Nous pouvons néanmoins nuancer les propos de Haroz et son équipe (2015) qui ont quant à eux étudié comment les pictogrammes impactent la mémoire, la performance et

l'engagement des individus auprès des visualisations de données. Nous nous distançons de cette étude sur un point : pour ces chercheurs, le pictogramme « Superfluous », qui s'apparente à ce que nous appelons « icône apposée » est une distraction pour l'utilisateur. Ils disent aussi : « Les pictogrammes ne nuisent pas au spectateur tant qu'ils sont utilisés pour représenter des données. Mais le fait d'inclure une image de fond inutile dans une visualisation semble distraire, et peut détourner l'attention des données » (Haroz et al., 2015, p.9). De même, ils conseillent d'utiliser des pictogrammes pour les tâches exigeantes (lire une barre dans un diagramme en bâtonnets). Cette recommandation est pertinente pour eux : en remplaçant les barres d'un diagramme par des pictogrammes, la mémoire à long terme sera davantage stimulée. Bien que notre corpus n'était pas suffisamment fourni pour pouvoir étudier l'impact des différents pictogrammes tels que classifiés par Haroz et ses pairs, nous constatons tout de même que le ressenti de nos participants diffère de leur recommandation. Effectivement, utiliser un pictogramme en lieu et place des données revient à placer une icône en *data-ink*. Or, l'analyse statistique a clairement montré que les participants ont estimé ces visualisations moins claires en comparaison avec les visualisations dont l'icône est apposée, ce que Haroz appelle « Superfluous ». Si cela stimule certainement la mémoire et l'engagement de l'utilisateur, la prescription d'Haroz et de son équipe risque d'impacter la clarté de la visualisation de données (et ce, même si nous parlons ici d'une clarté déclarée par les participants). Ainsi, si l'icône apposée semble être une distraction selon Haroz et son équipe, il s'avère qu'elle n'impacte pas la clarté en ce qui concerne la lecture des données, tout en étant un indice important de mémorabilité. Certes, l'individu regarde plus d'éléments sur l'image. Nous aurons l'occasion d'étudier cela et de le détailler davantage grâce à l'analyse des *gazes data* ([Partie V](#)). Par contre, Haroz et ses collègues insistent sur le fait que l'utilisation de pictogrammes très simples provoquent l'engagement de l'utilisateur. Nous nous accordons avec ce propos, puisque nous avons constaté davantage de jugement de beauté et d'appréciations positives en ce qui concerne les visualisations de notre corpus présentant des pictogrammes simples.

En réalité, peu de recherches dans le domaine ont relevé des appréciations déclarées, s'intéressant souvent à la mémoire ou à la compréhension des individus. Nous avons davantage souhaité nous concentrer sur le ressenti des participants, et donc sur ce qu'ils pouvaient déclarer. Nous nous rapprochons en cela de l'étude de Quispel & Maes (2014), qui ont étudié comment les professionnels en design graphique et les gens « de la vie de tous les jours »³⁴ notent l'attractivité et la clarté de visualisations de données dont le type de construction (standard ou non standard) et le mode d'expression (pictural ou abstrait) diffèrent. Il y a donc quelques croisements à effectuer entre leur étude et notre analyse. Par

³⁴ Que nous appellerons ensuite « profanes ».

type de construction standard, Quispel et Maes entendent des *pie chart* et des *bar chart*. Par non-standard, ils entendent les autres types de construction. Le mode d'expression pictural correspond à ce que nous qualifions d'ornementation, avec l'utilisation d'images ou pictogrammes, tandis que le mode abstrait concerne simplement les visualisations non-ornementées. Quispel et Maes ont demandé à des étudiants en design graphique de réaliser des visualisations de données en fonction d'un set de données qui leur a été fourni. Ensuite, ils ont testé ces visualisations sur deux échantillons d'utilisateurs (des professionnels en design graphique et de profanes) afin de relever leurs appréciations en ce qui concerne l'attractivité, la clarté et leur appréciation générale. Nous avons résumé leurs résultats dans le Tableau 39, qui ne reflète pas la profondeur de l'étude mais permet de mettre en évidence les nuances entre nos résultats respectifs.

Tableau 39 Brefs résultats de l'étude de Quispel et Maes (2014)

	Professionnels en design		Profanes	
	Type de construction	Mode d'expression	Type de construction	Mode d'expression
Attractivité	Standard > Non standard	Picturales > Abstraites	Non Standard > Standard	Abstraites > Picturales
Clarté	Standard > Non standard	Abstraites > Picturales	Standard > Non standard	Abstraites > Picturales
Généralement	=	=	Standard > Non standard	Abstraites > Picturales

Dans ce tableau, le signe > signifie « est mieux noté que ».

De manière générale, nous retenons de cette étude que le citoyen lambda préfère les visualisations non ornementées, les trouvant donc plus claires et même plus attractives, tandis que les professionnels en design graphique jugent les visualisations ornementées plus attractives que les visualisations non-ornementées, bien que moins claires. Le constat est très clair pour les chercheurs : « Les résultats montrent des différences claires entre les deux groupes cibles de l'étude : les professionnels attribuent une meilleure note à la visualisation picturale et non standard qu'à la visualisation standard et abstraite. Les profanes préfèrent les visualisations standard et abstraites » (Quispel et Maes, 2014, p.113, notre traduction). Il est frappant de constater que les profanes préfèrent les visualisations abstraites, qui correspondent aux visualisations standard de notre corpus (type : « Raw »), tandis que l'inverse se manifeste pour les professionnels en design graphique. L'étude de Quispel et Maes nous permet ainsi de positionner notre échantillon de participants, essentiellement constitué de professionnels de la communication numérique, dont l'objectif quotidien est de communiquer et présenter des informations, ce qui est un point commun avec les professionnels en design graphique sondés dans l'étude de nos collègues. Tout comme ces professionnels, nos participants préfèrent eux aussi les visualisations ornementées. Le croisement de nos analyses respectives montre ainsi l'importance de

l'échantillon participant à la recherche. Quoi qu'il en soit, nous pouvons constater ici que l'ornementation des visualisations de données semble plus importante pour ceux qui ont l'habitude de présenter l'information plutôt que pour les profanes. Dans tous les cas, les visualisations que nous qualifions de standard (c-à-d. type « *Raw* ») sont jugées plus claires que les visualisations ornementées. Cela ne remet néanmoins pas en doute les spécificités issues de notre analyse statistique ainsi que la recommandation que nous avons effectuée dans le paragraphe précédent.

Pour conclure, notre analyse se concentre sur des appréciations rapportées par les participants. Elle montre bien que dans les appréciations réflexives, seule la clarté dépend de l'ornementation. En effet, nous constatons que l'intérêt et la compréhension sont des appréciations qui varient très peu au cours de l'expérimentation. Au final, l'individu pense toujours avoir compris l'information ou, s'il ne la comprend pas, cela ne dépend pas forcément de l'ornementation. De même, nous pensions *a priori* que l'ornementation aurait un grand effet sur l'intérêt des participants et il n'en est rien. Comme nous l'avons déjà mentionné, ces appréciations reflètent le ressenti de l'individu, ce qui est non négligeable à l'heure où le récepteur est de plus en plus considéré comme crucial dans la conception du message permise par la visualisation d'information.

Malgré tout, ceci reste une analyse quantitative qui comporte des limites. Elle ne nous donne ni accès aux émotions des individus et ne nous permet pas de comprendre les facteurs qui les ont poussés à estimer qu'ils ont compris ou non une visualisation. L'analyse nous a néanmoins permis de proposer une recommandation en termes de conception pour les personnes souhaitant orner leurs visualisations de données. Cette recommandation est importante et peut s'avérer déterminante dans le champ de l'*Infovis*. Pourtant, à ce stade, nous avons encore besoin de confirmer ses fondements afin de consolider cette proposition. Il va sans dire que les deux autres jeux de données, à savoir les retranscriptions des entretiens semi-directifs et les *gaze data*, nous permettront d'alimenter la réflexion.

Partie IV La subjectivité dans la construction de sens

Analyse et interprétation des entretiens semi-directifs

Introduction

Cette quatrième partie concerne la subjectivité des participants dans la construction de sens. Que disent-ils, que pensent-ils, que ressentent-ils ? En quoi leurs caractéristiques personnelles entrent-elles en compte dans la construction de sens ? Cette fois, c'est sur le ressenti de l'individu que nous nous concentrons.

Pour commencer, nous présenterons l'analyse thématique menée sur la transcription des entretiens semi-directifs. Basée sur les facteurs humains et émotionnels de Kennedy et Hill (2017, 2018) ainsi que sur un codage thématique de type inductif, elle révélera quels éléments présents dans les visualisations ont influencé la perception des participants et leur appropriation des visualisations de données.

Ensuite, nous proposerons une adaptation de la *sense-making methodology* de Dervin et de ses collègues (2003) afin d'analyser les actions et les stratégies mises en place par les participants lorsqu'ils ont rencontré des problèmes à la lecture des visualisations de données du corpus.

Pour terminer, nous lierons la visualisation de données au concept de design émotionnel de Norman (2013). Il semblerait en effet que les propos des participants illustrent correctement la théorie proposée par le psychologue cognitiviste américain.

Dans ce chapitre, nous présentons l'analyse et les interprétations issues d'un premier codage thématique réalisé sur les transcriptions des interviews. Une fois la méthode présentée, cette première analyse met en exergue les propos les plus importants des participants à propos de la visualisation de données et de l'ornementation en particulier.

1. Méthode d'analyse par codage thématique

Des quarante entretiens menés directement après les expérimentations, nous avons retenu de nombreux éléments au grand potentiel pour notre analyse. Les participants, travaillant tous dans un domaine connexe à la communication numérique, étaient clairement en mesure de distinguer ce qui, à leur avis, leur permettait de comprendre et d'apprécier ou non une visualisation de données. Au-delà de l'expression de leurs constats et ressentis, ils se sont, dans la plupart des cas, montrés capables de synthétiser leur opinion sur l'ornementation des visualisations de données, même si cela n'était pas demandé dans l'exercice. Effectivement, l'ornementation, à considérer désormais comme une pratique de visualisation, fait partie intégrante du quotidien des professionnels de la communication digitale et de ses métiers connexes. Dans certains cas, elle semble d'ailleurs indispensable.

Afin d'explorer en profondeur les données qualitatives issues de l'entretien semi-directif, nous avons mené une analyse thématique sur les retranscriptions des entretiens (mot-à-mot). Comme son nom l'indique, elle suit son cours en suivant plusieurs catégories de thèmes. L'analyse thématique « consiste à lire un corpus, fragment par fragment, pour en définir le contenu en le codant selon des catégories qui peuvent être construites et améliorées au cours de la lecture » (Fallery & Rodhain, 2007, p.9). Propre aux approches constructivistes, elle est particulièrement pertinente dans notre cas puisqu'elle permet d'interpréter un contenu (Fallery & Rodhain, 2007). Les thèmes évoluent donc au cours de l'analyse : ils peuvent émerger avant et pendant l'analyse, et être corrigés après l'analyse avant de procéder à une relecture. Il est possible de consulter les thèmes que nous avons encodés pour cette analyse dans l'annexe n°6. Une partie des thèmes utilisés, que nous appelons « codes » a pu être définie au préalable de l'analyse. Les facteurs humains et émotionnels de Kennedy et Hill (2017, 2018) en faisaient partie, puisqu'il nous semblait logique que les participants citent certains de ces facteurs pour expliquer s'ils se

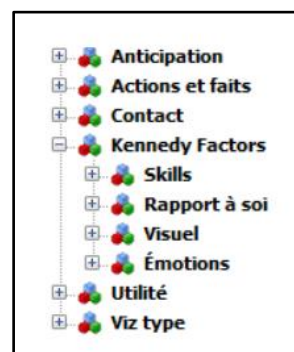


Figure 41 Catégories de codes utilisées dans l'analyse thématique

sentaient engagés ou pas auprès des visualisations. D'autres codes enregistrés au préalable concernaient tout simplement l'identifiant de chaque visualisation testée, afin que nous puissions être en mesure de lire dans un tableau tous les avis émergents à propos d'une visualisation en particulier et de, pourquoi pas, les croiser avec un autre code ou encore d'établir des comparaisons (code : « viz type »). Pour la suite des codes employés, ils ont émergé grâce à des thèmes plus larges mentionnés par les participants. Au cours de l'entretien, les participants ont en effet réagi par rapport à des actions qu'ils ont menées et des faits qu'ils ont constatés lors de la lecture de la visualisation de données (code : « actions et faits »). Ensuite, les participants se sont également exprimés sur les dimensions visuelles et émotionnelles des visualisations ainsi que des compétences à mobiliser pour pouvoir comprendre ces graphiques. Cela correspondait parfaitement aux facteurs humains et émotionnels mis en évidence par Kennedy et Hill (2017, 2018) en ce qui concerne l'engagement d'un individu envers une visualisation de données. En fonction des entretiens, nous nous sommes permise d'enrichir ces facteurs et de les subdiviser en diverses sous-catégories (code : « *Kennedy Factors* »). Puis, l'objectif ou l'utilité des pratiques a également été mentionné par les participants (code : « utilité »). Nous avons aussi codé la fréquence de contact des individus avec les visualisations de données, ainsi que leur capacité à décrire presque à la perfection quel devrait être le visuel alternatif de la visualisation qu'ils ont observée (par exemple, quelle serait l'ornementation d'un graphique standard). Les différentes catégories sous lesquelles nous avons placé nos codes se trouvent dans la Figure 41. L'analyse thématique a été réalisée grâce au logiciel QDA Miner.

2. Interprétation résultant du codage thématique

Le codage thématique permet d'interpréter les différentes opinions tenues par les participants au cours des entretiens semi-directifs. Il permet de voir dans quelle mesure ces opinions convergent ou divergent, en y apportant toute la nuance propre à l'analyse qualitative.

En premier lieu, nous abordons les propos relatifs à la perception des visualisations de données présentées dans le corpus. Quatre grands « thèmes » ont pu être dégagés de ce propos : la sensation d'effort, l'idée de modernité, la surcharge visuelle et, finalement, les facteurs d'incompréhension.

Ensuite, nous nous concentrons sur les facteurs permettant d'apprécier (ou non) les visualisations de données. Ils nous permettent de formuler quelles sont, aux yeux des participants, les forces et les faiblesses de l'ornementation.

Finalement, nous illustrons la pertinence des facteurs humains et émotionnels de Kennedy et Hill (2016) en montrant en quoi ils occupent une place importante dans la construction de sens des participants confrontés au corpus de visualisations de données. Nous concluons ensuite en rappelant les apports de cette analyse thématique.

2.1. La perception des visualisations

En s'exprimant à propos des contenus qu'ils venaient de lire, les quarante participants nous ont fait part de leur propre perception des visualisations de données qui leur ont été exposées. Nous présentons ici par thème les différents propos mentionnés fréquemment et de façon convergente par les participants.

2.1.1. La sensation d'effort

Premièrement, nombreux sont les participants, qui, au contact d'une visualisation de données, s'apprêtent à devoir fournir un effort de compréhension, ou du moins, attendent de la construction graphique qu'elle amoindrisse cet effort lié à la perception de données chiffrées. Cette notion « **d'effort à fournir** » a été mentionnée par les participants 55 fois au cours de la récolte des données – donc près d'une fois par entretien.

La sensation d'effort à fournir de la part de l'individu a été expliquée plusieurs fois par **l'incompréhension d'éléments visuels** disposés sur la visualisation, comme l'échelle, ou encore une illusion d'optique qui perturbe la réception de l'information, comme dans la visualisation n° 4 du corpus. Ensuite, le fait de devoir « **chercher** » **les données** de manière visuelle afin de reconstituer l'information a été mentionné 10 fois. Effectivement, la disposition des éléments tels que titre, légende, échelle, variables, etc. suivrait un ordre logique auquel s'attendent les participants, par convention. Lorsque ce n'était pas le cas, ils étaient conscients de l'effort à fournir. Il s'agirait donc d'un effort supplémentaire à la lecture et à la compréhension intuitive des données. Dans plusieurs cas, en fonction des visualisations de données et des intérêts des participants, cet effort n'était pas forcément vu comme négatif s'il était joint à **une expérience agréable** (voir *infra*).

Un autre élément, mentionné à 9 reprises, intervient dans la sensation d'effort à fournir : la **présence d'une surcharge d'éléments** sur la visualisation de données. Cette surcharge décrite par les participants concerne tantôt les visualisations sur lesquelles apparaissaient beaucoup de données (par exemple la n°23 dans le corpus), tantôt les visualisations chargées d'un point de vue esthétique. Ainsi, la surcharge d'éléments visuels tels que la couleur abusive, les pictogrammes et autres mises en scènes a engendré une sensation d'effort à fournir chez plusieurs consommateurs des visualisations n°14 et 16. En réalité, cela va de pair avec une remarque qui a été émise de façon répétée : le pictogramme ou

l'**image très simple** amoindrirait l'effort à fournir lors de la lecture, ce qui ne serait pas le cas d'images plus complexes, qui alourdisent la visualisation au point de créer un sentiment d'effort. Dans la suite, nous reviendrons sur le caractère simple des ornements et sur les appréciations positives que cela a apporté à plusieurs participants.

Un autre élément influençant la sensation d'effort à fournir est **le caractère attractif** de la visualisation. Ceci dit, en ce qui concerne les visualisations de type standard, il est sûr que l'absence d'éléments imagés ne contribue pas à l'attractivité de la visualisation et, au-delà de cela, ne donne pas « l'envie » au lecteur.

N : Ici, ce n'est pas délimité, il faut vraiment faire un travail pour comprendre où vont les pays par rapport aux animaux, c'est pas bien délimité. Et puis en plus c'est gris... (visualisation n°26).

T : Qu'aurais-tu fait ?

N : J'aurais mis plus d'espace, j'aurais mis la tête d'un chat, la tête d'un chien, d'un poisson, les drapeaux des pays... ça aurait servi à ce qu'on le voit. Là, telle que je vois l'information, je n'ai pas envie de la lire. Dommage parce que le sujet est fun, pourquoi pas.

T : Sur celle-ci, qu'est-ce que tu aimes ? (visualisation n°6)

M : Le graphisme et le sujet. J'ai trouvé ça intéressant de voir la comparaison entre certains pays. Et tous les animaux sont dans des couleurs différentes, ça permet de bien distinguer et ensuite, tu as la forme de l'animal, le symbole de l'animal en lui-même qui est repris et ça ajoute de la clarté en plus je trouve. Je trouvais ça clair. Tout ce qui est autre chose que des bâtonnets, d'office, ça m'a plus intrigué je trouve que ça demande plus de recherche de compréhension quand il y a une image et du coup je reste plus longtemps mais comme c'est quelque chose qui est plus attractif pour l'œil, ça demande plus de temps mais ça ne me dérange pas de passer plus de temps puisque l'image est quelque chose qui paraît moins... ce sont des données mais je n'ai pas l'impression de lire un graphique. Quand ce sont des barres, ça me demande plus d'efforts de rester dessus, regarder les pourcentages, etc.

Ce dernier verbatim illustre assez bien l'avis général quant à l'effort à fournir qu'ajoutent les ornements. La visualisation de données n°6 présente des données simples, en pourcentages, imagées grâce à des pictogrammes simples sur une échelle de pourcentages. Son attractivité a grandement été soulignée par les participants qui, conscients des efforts

supplémentaires à fournir lors de la lecture, s'en réjouissent finalement, soulignant l'aspect agréable que la tâche a alors suscité.

Un autre élément lié à la sensation d'effort à fournir concerne davantage les visualisations de données ornementées. De **nouvelles formes de visualisation** ou bien des présentations originales et uniques, qui correspondent alors à l'attribut « Ornement > Forme inhabituelle » mobilisé dans le cadre de l'analyse statistique, demandent au lecteur un temps d'adaptation et, forcément, de concentration qui lui demande de dépasser ses conventions habituelles. A la vue de ces nouvelles formes, la sensation d'effort peut survenir, tandis que le fait de considérer cet effort comme positif ou négatif dépend d'un participant à l'autre.

W : Oui, c'est toujours mieux d'avoir des informations qui s'écartent des schémas qu'on peut connaître. Donc, le fait que ce soit quelque chose qu'on connaît moins, on doit lui accorder plus de temps parce que... c'est comme ça que je l'interprète. C'est différent et les autres graphiques c'est toujours la même chose ! Là il me faut un temps d'adaptation, de calibrage pour me dire que l'information n'est pas structurée linéairement et que je vais devoir plus chercher que d'habitude. Mais je préfère.

Finalement, le **temps** que prend la lecture d'une visualisation influence également la sensation d'effort à fournir, plus particulièrement sur des visualisations de type standard, pour lesquelles la remarque a été faite une dizaine de fois. Effectivement, si le participant a l'impression d'y passer beaucoup de temps sans que l'expérience ne soit agréable, sa sensation d'effort serait susceptible d'être plus importante.

Nous avons choisi d'utiliser les termes « sensation d'effort » et non tout simplement « effort ». Premièrement, ce qui est ressenti comme un effort varie d'un participant à l'autre. Mais par-dessus tout, il est fondamental de comprendre que tous les individus font des efforts différents lors de la lecture de la visualisation de données, dépendant de leur personnalité, caractère, expérience, mais aussi de leur propre perception de l'effort en lui-même. Ainsi, pour un même effort consenti pour plusieurs personnes (par exemple, la lecture d'un titre), certains n'auront pas eu la sensation de réaliser un effort tandis que ce sera le cas pour d'autres. On parle donc bien de sensation d'effort, du fait d'avoir ressenti la nécessité de se dépasser afin de pouvoir comprendre la visualisation de données. Ces verbatims illustrent bien les nuances dont nous tenons compte :

F : Quand je vois un graphique, oui bien sûr je m'apprête à faire un effort mais je trouve qu'aujourd'hui on fait moins d'effort grâce aux infographies et tout ce qui est possible de faire ... Donc il se peut que je m'attende un peu à ce qu'on me prémâche le boulot avec tout ça. Et donc quand on revient sur un graphique simple, ça peut me sembler compliqué même si ça ne l'est pas vraiment.

T : Tu parles de « facile », est-ce que ça veut dire que ce visuel que tu trouves “beau” te donne l'impression de devoir faire moins d'effort ?

V : Tout à fait. Je pense que le visuel, c'est des dessins donc c'est un peu comme une BD ou un dessin animé et que quand on voit un graphique, peut-être que mon subconscient le lie à l'école, et tout ça. Je le lie à Excel aussi, à une équation, et j'imagine que quelqu'un a fait un travail statistique... Tandis qu'avec une infographie, beaucoup moins.

Ensuite, il est également intéressant de savoir que les participants ont mentionné l'effort lorsqu'ils parlaient d'une visualisation ornementée et ce, à 33 reprises. Ils ont mentionné la sensation d'effort en parlant d'une visualisation de type standard à 10 reprises. Comme nous l'avons indiqué précédemment, la sensation d'effort peut être perçue de manière positive par les participants. Cette sensation – et la perception qu'en ont les participants – tient sans aucun doute un rôle dans la construction de sens. C'est pourquoi nous nous penchons désormais sur les cas dans lesquels les participants estiment que **leur effort est amoindri**.

Il va sans dire que, conformément aux principes de conception des visualisations de données principalement dictés par Bertin et Tufte, **des éléments visuels correctement mobilisés** favorisent la minimisation de l'effort à fournir (couleur utilisée dans un objectif dissociatif, étiquetage correctement réalisé, etc.). Les visuels épurés sont souvent décrits comme clairs, malgré quelques exceptions. Pour une dizaine de personnes, le fait de trouver une visualisation **belle** leur donnait l'impression d'un effort amoindri. Toutefois, rappelons que c'est d'une impression dont les participants nous ont fait part puisque l'analyse statistique basée sur la récolte quasi-instantanée des appréciations déclarées a montré que la corrélation entre le fait de trouver une visualisation « belle » et le fait de la comprendre est très faible. « L'amusement » provoqué par l'ornementation peut ainsi donner la sensation que l'effort soit amoindri sans pour autant en impacter finalement la compréhension.

La **simplicité**, de manière générale, est un gros atout pour les participants. L'exemple parfait est la visualisation n°13 et son homologue n°33 qui parlent toutes deux de

l'augmentation des quantités servies au restaurant. Seulement trois variables à afficher et très peu de données : voilà qui ravit les participants. En un seul coup d'œil l'information est saisie. La simplicité va encore plus loin : dans le visuel n°13, de grandes images simples remplacent les barres du diagramme. Le fait qu'elles constituent le *data-ink* se discute du point de vue des principes, mais cet exemple montre qu'une bonne partie des participants qui ont vu cette visualisation ont l'impression d'être aidés pour plusieurs raisons : l'image est une **allusion au sujet**, ce qui permet à l'individu de rester plongé dans l'univers ce qui, de plus, provoquerait ainsi une **sensation agréable**. Nous pourrions alors dire dans ce cas que l'effort n'est pas réellement amoindri mais est plutôt contextualisé et il s'avère que plusieurs participants parlent d'un « effort soulagé ».

T : Selon tes appréciations, pour celle-ci c'est tout l'inverse : je vois qu'elle est belle à tes yeux, intéressante, claire et que tu as bien compris (visualisation n°13).

C : Oui, je préfère ce genre de représentations, on est dans le graphisme, on voit les objets, ce dont on parle, on doit moins lire, on comprend directement, il y a des couleurs.

A : Le fait que la silhouette des bâtiments nous parle directement. Fatalement on va comprendre sans devoir lire. Et là... Mmmh le chiffre a été représenté par un rond donc on voit qu'on n'a pas besoin de lire le chiffre, on peut voir le rond et se dire que ça représente une proportion, on peut comparer deux ronds directement. Et visuellement, le mauve évoque la nuit... Et la silhouette en mauve foncé dans la nuit c'est très pertinent ainsi que le logo Instagram dans le fond aussi (visualisation n°8).

Plus l'élément décoratif est simple, à l'instar des pictogrammes, qui sont très appréciés, plus cela s'affirme. Ainsi, le fait que l'allure esthétique de la visualisation fasse allusion au sujet ou soit apte à le rappeler au participant provoque une sensation d'effort amoindri.

T : Avez-vous envie de regarder ou de lire l'image ? (visualisation n°32)

V : Moyennement je dirais. Parce que c'est un peu brut comme info si je dois comparer à Game of Thrones, là il y a une espèce de mise en scène mais on est tout de suite mis dans le sujet. On ne doit pas réfléchir Star Wars, les sabres, les couleurs, etc.

T : vous le feriez autrement alors ?

V : Oui, je prendrais la symbolique du sujet.

T : Il n'y a pas de contexte alors ?

V : *Je dirais plutôt qu'il n'y a pas d'enrobage. Ce sont des données brutes et on les comprend c'est tout. Ce n'est pas fun. C'est compréhensible et il y a de la couleur mais ce n'est pas amusant.*

T : *C'est important l'enrobage comme vous dites ?*

V : *Oui, je pense que c'est important, que ça aide, que ça attire l'attention et ça aide à comprendre peut-être plus vite. Mais compréhensible c'est compréhensible voilà, le fait qu'il n'y ait pas l'enrobage ne détruit pas la compréhension non plus...*

Dans ce verbatim, l'aspect « fun » et amusant a été soulevé par le participant. Il a souvent été mentionné et nous y reviendrons. Cet aspect ne semble pas diminuer la sensation d'effort. Cependant, il est important en ce qui concerne la construction de sens.

Pour revenir à la diminution de la sensation d'effort, le fait que la visualisation de données, de par son esthétique, provoque une sensation agréable à la lecture, est aussi un facteur de sensation d'effort diminué. A cet aspect agréable se mêle également le **facteur attractif** de la visualisation.

T : *Si je te montre maintenant celle-ci... Qu'en penses-tu ? (visualisation n°12)*

B : *C'est mieux... C'est plus agréable à regarder mais le sujet est inutile... Mais oui, je vais passer plus de temps sur le graphique s'il est plus facile à regarder.*

T : *ça c'est plus facile ?*

B : *Oui ! Surtout si le sujet est lourd à la base... Le rendre plus attractif en mettant des... Là, c'est plus clair quoi, avec des illustrations et tout ça change complètement, ça change tout ! C'est mieux.*

A : *En général, à travers les 20 images, quand c'était les simples graphiques type Excel en noir et blanc bah ça donne déjà moins envie.*

T : *C'est important de « donner envie » ?*

A : *dans un esprit « com' », on se dit directement qu'on doit attirer tu vois... Pour un nouveau public, tu as envie de visualiser ça autrement.*

T : *Les efforts esthétiques de ce genre y contribuent ?*

A : *Oui, je pense. C'est plus agréable. Même avec un contenu plus lourd, c'est plus sympa.*

En opposition aux aspects agréables, attractifs et aux allusions au sujet produites par les ornements, les participants ont souvent cité l'aspect « brut » ou « scolaire » des

visualisations de type standard qui nécessitent une lecture sérieuse, bien que facile (voir *infra*).

Le Tableau 40, la Figure 42 et la Figure 43 résument les propositions majeures concernant la sensation d'effort. Pour rappel, la sensation d'effort à fournir a été mentionnée 55 fois. Afin d'en expliquer les raisons, nous nous sommes appuyée sur les verbatims dans lesquels le code « effort à fournir » croise d'autres codes. Notons que ces mêmes codes apparaissent dans une bien plus large mesure dans les entretiens, en fonction des contextes et des difficultés rencontrées par les participants (incompréhension 117x, confusion 116x, recherche d'information 43x, surcharge visuelle 43x, aspect rébarbatif 58x). Nous avons agi de façon similaire avec la sensation d'effort amoindri, qui a été mentionnée au nombre de 63 fois. D'une même façon, les codes permettant de contextualiser cette sensation d'effort amoindri ont été mentionnés plus largement dans d'autres contextes dans l'ensemble des entretiens (clarté 212x, ornement 56x, simplicité 60x, sensation agréable 53x, beauté 125x, allusion au sujet 58x, attractivité 110x). Rappelons que l'objectif ici n'est pas quantitatif mais qualitatif : en montrant que certains thèmes sont récurrents une fois croisés à d'autres thèmes, nous justifions la structure adoptée pour notre propre interprétation qui se veut riche en détails.

Tableau 40 Tableau de synthèse sur la sensation d'effort

Sensation	Favorisée par
Effort à fournir	Incompréhension Confusion Recherche d'informations Surcharge visuelle Aspect rébarbatif
Effort amoindri	Beauté Simplicité Allusion au sujet Sensation agréable Attractivité

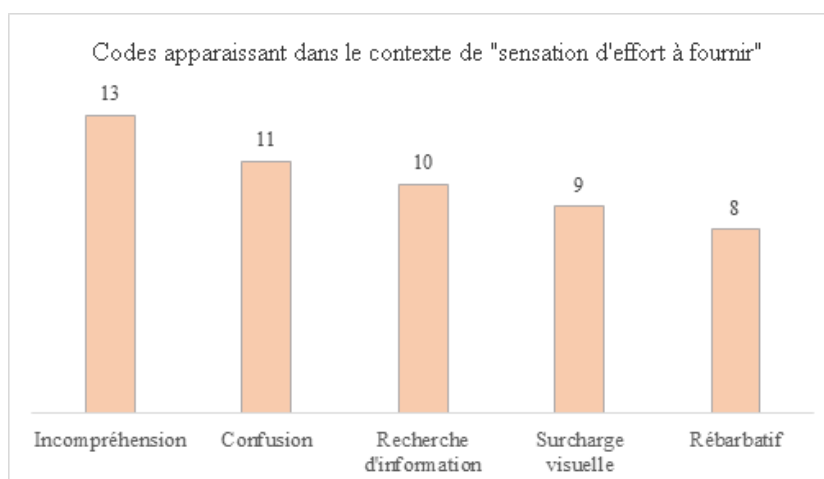


Figure 42 Codes apparaissant dans le contexte de la « sensation d'effort à fournir »

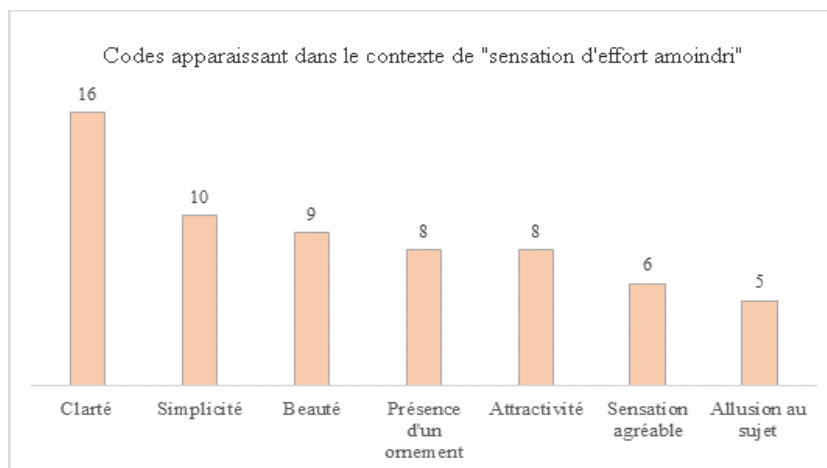


Figure 43 Codes apparaissant dans le contexte de la « sensation d'effort amoindri »

2.1.2. L'idée de modernité

Sans doute influencés par leur métier, qui vise souvent à « présenter de l'information », un grand nombre de participants a indiqué que l'aspect « neuf » ou « vieillot » des visualisations de données impactait leur propre réception de la visualisation de données.

Tout d'abord, pour quelques participants, les diagrammes en bâtonnets de type traditionnel évoquaient en eux un souvenir d'écolier. Effectivement, de tels graphiques sont largement utilisés dans le **cadre scolaire** où c'est par le biais de diagrammes que les élèves établissent un premier contact avec des données de type statistique

Si, de manière générale, les participants estiment que le *bar chart* est quelque chose de simple qui donne une idée des données de manière très claire, il a été décrit comme **vieillot** à 9 reprises.

Selon les personnes qui se sont exprimées sur l'aspect vieillot des visualisations, il y a là de quoi décourager le lecteur, ou, du moins, il n'y a rien qui l'encourage à lire la visualisation de données.

Cependant, les avis peuvent surprendre : si, parce qu'il est utilisé depuis longtemps, le *bar chart* peut parfois se doter d'un air vieillot, ce sont les visualisations ornementées qui seraient les plus « ringardes ». Effectivement, l'esthétique du diagramme en bâtonnets ne change pas ou très peu, tandis que **l'esthétique des visualisations ornementées** est justement soumise à des jugements de goût qui peuvent donc dépendre de tendances. Prenons pour exemple la visualisation n°15, faisant état dans un *bar chart* horizontal des pays les plus généreux en œuvre de charité. Dans les barres du diagramme s'alignent des pictogrammes de billet vert, dans un but décoratif. Il n'y a évidemment pas de vérité unanime, puisqu'on parle ici de jugement de goût. Si certains ont trouvé l'idée appréciable,

cela a constitué une mauvaise expérience pour d'autres qui y voyaient des codes faisant appel à des expériences lointaines, ou à une esthétique qu'ils ont connue par le passé sur des supports numériques. De ce fait, l'esthétique de cette visualisation semble avoir perdu de sa modernité. Il en va de même pour d'autres visualisations ornementées, qui, selon les participants, font penser à « de vieux livres d'histoire » ou « une vieille illustration réalisée à la main ». Bien souvent, c'est à propos des visualisations de données où la part d'ornementation est chargée et complexe que surviennent ces commentaires. Il serait donc opportun de penser que, de nouveau, miser sur une ornementation simple minimise les risques de perdre en modernité.

T : Celle-ci par contre, tu ne la trouvais pas très belle (visualisation n°15)

A : Non c'est archaïque je trouve... On se croirait sur MSN, avec des billets là...

T : Qu'est-ce qui fait MSN ?

A : Je ne sais pas mais c'est raté, simpliste pas moderne... Mais encore une fois j'ai dit que c'était fort clair.

Mais encore, dans le même ordre d'idées, l'originalité des visualisations, citée 17 fois, est un élément crucial quant à la construction de sens opérée par l'utilisateur. Tout d'abord, il y a l'idée que les outils de visualisation sont aujourd'hui nombreux et accessibles à tous. Leurs caractéristiques esthétiques sont connues et vont de pair avec l'idée de modernité.

T : Celle-ci était plutôt intéressante, claire et comprise. Par contre pas spécialement belle et vous n'avez pas spécialement aimé (visualisation n°25).

L : Oui voilà, c'est la présentation la plus traditionnelle des années 80. Avec les outils modernes, on peut la mettre en relief, la mettre en 3D tout en gardant le même format niveau des proportions et autres qui resteraient lisibles. Là on n'est pas dans le professionnel, on n'est plus dans la communication classique où on a besoin d'un visuel rapide clair et précis et qu'on peut passer en revue sur lequel on peut s'attarder si on le souhaite mais qui rapidement s'exprime complètement.

T : Qu'ajouteriez-vous alors ?

L : De nombreuses choses, de la 3D, des couleurs flashy et un peu plus pétant, on peut également s'amuser avec le texte : un petit peu plus propre et un petit peu plus net.

T : Est-ce que ces éléments la contribuent à quelque chose ?

L : Oui, pour moi c'est clair, limpide et net, on voit bien les proportions, elles sont bien délimitées (...). On peut rendre vraiment plus glamour que cette version là même si ce type de graphe me paraît quand même l'un des plus clair.

De plus, puisque ces outils sont accessibles, certains participants ont en tête une idée de « soin » que le concepteur se doit d'avoir à l'égard du lecteur. Une visualisation soignée est originale, moderne, dans l'air du temps.

Quand on voit ça, on se dit que la personne n'a pas vraiment fait d'effort pour nous transmettre l'info quoi... Visuellement, il n'y a pas grand-chose, on dirait qu'elle a fait ça à la va-vite sans s'occuper du fait qu'on aime ou pas.

Ainsi, le fait de devoir « dynamiser » des visualisations de type standard, voire de les « imager », a notamment beaucoup été évoqué (25 fois). Nous reviendrons par la suite sur les commentaires concernant les éléments visuels. Il s'avère tout de même que l'élément le plus cité dans ce besoin de « dynamisation » est la couleur, qui, par ailleurs, est aussi évoquée comme indispensable à l'attraction du regard.

T : Sur celle-ci, tu ajouterais quelque chose ? (visualisation n°29)

J : une couleur autre que le gris : ça a un côté tout de suite attractif parce que le gris ça fait papier journal, ça rapporte à quelque chose de lassant. Il faut rendre la chose didactique et attrayante. Didactique ça allait peut-être mais si les gens étaient attirés par le gris ça se saurait...

T : Celle-ci, tu l'as bien comprise (visualisation n°34).

N : Oui mais il n'y a pas d'images et c'est dommage. On parle de crème glacée donc on aurait pu mettre un pot de glace à la place du rectangle. J'aurais préféré...

T : Pourquoi ?

N : c'est pour qu'il y ait un impact. Si tu mets une image, je vais retenir directement l'information sans faire d'effort alors que là, je ne m'en suis même pas souvenue.

En résumé, les visualisations de données les plus vieillottes ne sont pas simplement les plus traditionnelles (de type standard dans notre expérimentation), mais celles qui sont ornementées et qui ne sont plus « à la mode ». Visiblement, le design d'information connaît lui aussi des tendances. Les participants n'ont pas manqué de pointer l'originalité du visuel, rendue accessible aux concepteurs grâce aux outils de visualisation. Il est donc attendu de ces mêmes concepteurs qu'ils prennent soin de l'utilisateur. Finalement, les bénéfices de l'ajout d'images au sein des visualisations de données ont également été soulignés par les participants.

2.1.3. Une surcharge visuelle qui dérange

Si le fait d'imager les visualisations de données afin de provoquer une allusion au sujet enchante les participants, ils ne sont pas dupes pour autant. Il y a très certainement un juste milieu à trouver dans l'imagerie à mobiliser en termes de visualisation de données. Effectivement, l'ensemble des participants a indiqué que la réception d'une visualisation de données pouvait être perturbée par une **surcharge d'éléments** situés sur la visualisation et ce, à 91 reprises. Imager, oui, mais n'allons pas trop loin...

Ainsi, la surcharge d'éléments concerne diverses situations décrites par les individus :

- Un visuel **trop chargé** et trop décoré du point de vue **esthétique**. A 39 reprises, des visualisations trop chargées, à l'effort esthétique trop insistant, en sont venues à perturber la réception de l'information. Il y a deux exemples parfaits dans le corpus : les visualisations n°1 et 16. Le point commun de ces deux visualisations est que l'intention décorative du concepteur englobe l'entièreté du support, et ce, sans unicité, puisque le pictogramme est utilisé à outrance et sans raison si ce n'est « décorer ».
- Un **nombre de données trop conséquent**, qui effraie l'utilisateur : à 12 reprises, peu importe le type standard ou ornementé de la visualisation, des individus se sont dits irrités par le nombre de données affichées à l'écran. Le meilleur exemple, peu importe son aspect ornementé ou non, se situe dans les visualisations n° 3 et 23, qui illustrent l'espérance de vie de 39 animaux.
- Un **texte trop important** : à 6 reprises, la présence de texte explicatif jugé trop long sur les visualisations ornementées a été dérangeant ;
- **Plus d'une information** dispensée par visualisation : doubles graphiques, information supplémentaire et autres ont dérangé les participants dans leur réception une dizaine de fois.

Comme on le voit, les manières de surcharger une visualisation de données sont nombreuses. Ceci dit, c'est bien la surcharge d'éléments esthétiques qui est le plus souvent citée, malgré les atouts que ces éléments décoratifs peuvent apporter, comme nous avons pu en donner un aperçu dans [le point 2.1.1. de ce chapitre](#).

Cette surcharge, d'où qu'elle provienne, provoque selon 8 participants un balayement de l'écran des yeux : ils se plaignent ainsi de devoir « **chercher** » l'information.

Par ailleurs, le fait de devoir **simplifier une visualisation** de données afin de la rendre plus compréhensible a été mentionné 35 fois. Ce propos concernait à 21 reprises des visualisations ornementées et à 3 reprises seulement des visualisations de type standard. Ainsi, c'est surtout l'esthétique qui est à simplifier et non pas à bannir, selon 13 avis.

L'étiquetage a également été mentionné : mieux réalisé, il permet dans certains cas une simplification évidente. Par ailleurs, il est bon de noter que pour certains intervenants, l'ajout d'esthétique contribue à la perception de simplicité que peut susciter la visualisation. Cela s'explique d'une part à l'accoutumance aux infographies mais aussi d'autre part à l'adoucissement de l'aspect sérieux des visualisations de type standard.

T : Sur celle-ci, on a des barres, qu'est-ce que ça évoque ? (visualisation n°33)

E : C'est très clair, j'ai mis clarté au maximum à chaque fois... Mais là par exemple c'est tout à fait con mais j'aurais peut-être utilisé les frites, mettre un petit paquet et un gros paquet. Un petit hamburger et un gros hamburger. Parce que c'est une donnée relativement simple à comprendre et je trouve que là au contraire... en partant de l'idée que le dessin doit t'aider à comprendre le plus facilement possible, ça donne un aspect pseudo scientifique qui est totalement inutile parce que c'est une bête analyse dans ce cas-ci. Donc je ne vois pas pourquoi on force les gens à peut-être réfléchir sur quelque chose qui en une fraction de seconde avec un petit paquet et un gros paquet de frites serait compris.

2.1.4. Les facteurs d'incompréhension

Les participants ont manifesté à plusieurs reprises un sentiment de **confusion**. Effectivement, il s'agit d'un thème cité 201 fois au cours des entretiens. Nous avons donc considéré de nombreux éléments comme ayant une influence négative, allant de la simple confusion jusqu'à l'**incompréhension** de la représentation graphique. Ainsi, 122 citations mêlant confusion et incompréhension peuvent être comptabilisées en ce qui concerne des visualisations ornementées. En ce qui concerne les visualisations de type standard, la confusion a été mentionnée 57 fois, ce qui est tout de même moindre. Par ailleurs, 22 citations exprimaient un point de vue de manière générale sur le sentiment de confusion que peut produire la lecture d'un graphique.

La thématique de la **clarté**, signifiant qu'une visualisation de données est intelligible au premier coup d'œil, a été mentionnée 211 fois. Dans cet ensemble, 118 des citations que nous avons considérées comme un indice de clarté se rapportent à des visualisations ornementées, contre 81 en ce qui concerne les visualisations de type standard, tandis que 12 citations exprimaient un point de vue de manière générale.

Il convient d'être prudent par rapport à cette information. Effectivement, cela ne signifie pas que les visualisations de type standard sont estimées moins souvent claires que leur égale qui est de type ornementé. Les participants avaient effectivement le choix de parler de certaines visualisations ou pas, et sans obligation d'aborder chaque visualisation.

Comme l'ont mentionné Borkin et ses collègues (Borkin et al., 2013; 2016), les individus se rappellent mieux et plus rapidement d'une visualisation qui est ornementée, surtout si elle l'est avec des objets qui sont reconnaissables pour l'esprit humain (photos, pictogrammes, etc.). Il est donc logique que les participants aient mentionné plus souvent les visualisations ornementées. De plus, cela se vérifie clairement dans le vocabulaire qu'ils emploient : pour parler d'une visualisation ornementée, les participants mentionnent l'ornementation qui la compose (« l'horloge », « les animaux », « les sabres laser », etc.) tandis que sans ornementation, il est plus difficile pour eux de se rappeler des titres et des sujets des visualisations de type standard. Cette observation est parfaitement en phase avec les études de Borkin et ses pairs (2013, 2016).

D'ailleurs, de manière générale, les participants sont unanimes quant à la grande clarté et à la simplicité des visualisations de type plus classique qui, ancrées dans les habitudes, se lisent facilement grâce à différentes conventions.

Non c'est assez clair. Les barres c'est souvent clair mais pas joli. En fait, sur tous les graphiques en barres, j'ai mis que c'était clair.

Le fait que beaucoup de personnes aient estimé des visualisations ornementées comme claires pourrait s'expliquer par le fait que l'ornementation leur a procuré une sensation d'effort amoindri, comme expliqué dans [la section 1.1.](#) de ce chapitre. En effet les visualisations ornementées sont tout autant porteuses d'éléments pouvant semer la confusion. Le discours est nuancé en fonction des caractéristiques de chacun. Nous le verrons ultérieurement.

Quoi qu'il en soit, de nombreux éléments peuvent apporter de la confusion au moment de la lecture d'une visualisation de données.

En premier lieu, le **manque de clarté** engendré par un **effort esthétique** a été mentionné 45 fois. Une nouvelle forme de visualisation, parce qu'elle est encore inconnue, peut contribuer à ce manque de clarté. Mais encore, la disposition inutile ou hasardeuse d'éléments visuels peut également être fatale pour la clarté : une ornementation distrayant la lecture des données ou une légende trop petite ou mal placée n'en sont que des exemples.

« J'ai beaucoup cherché. Ce qui m'a intrigué, c'est cette confusion entre la légende, le rectangle mentionné italien et en dessous on me met le nombre d'habitants en 1999. Et dans le graphe, je vois Péruvien mais alors ça veut dire quoi, que le vert hachuré ce sont des Italiens Péruviens ? J'ai compris la volonté de faire un graphe mais je n'ai pas compris l'issue. Je ne sais pas dire si ce sont des Italiens Chinois ou des Chinois en Italie... l'idée n'est pas mauvaise graphiquement mais c'est du dessiné, c'est du pas

propre, c'est du old school. Il y a un fond, les traits ne sont pas nets et de nouveau, j'ai l'impression qu'on a plutôt travaillé sur l'aspect visuel, l'aspect graphique, sur la visibilité que sur le bon format de graphe. Ce n'est pas parce que c'est nouveau que c'est bon ! » (visualisation n°10)

Le **manque d'information** disposée sur la visualisation a été cité 29 fois comme élément principal de confusion. Un titre mal expliqué ou fantasque, le manque d'un texte explicatif ou encore le manque de guidance d'un point de vue visuel (traits, flèches, couleurs) ont par exemple été cités.

Comme dans le premier point, et c'est logique, le fait de devoir **chercher l'information** sur la visualisation parce qu'elle ne correspond pas à un **ordre intuitif de lecture** engendre de la confusion également (cité 25 fois dans ce cadre-ci).

Le **vocabulaire** a pu poser problème également. Les visualisations n°12 et 18 ainsi que leurs homologues de type standard sont deux bons exemples. Les mots « chaire » ou encore « univers étendu » (dans le cadre de la saga *Star Wars*) ont posé problème aux participants, tout comme la méconnaissance du quartier italien de Bovis (visualisations n°10 et n°30). Ce problème de compétences langagières, déjà soulevé par Kennedy et Hill (2017, 2018), a mené dans bien des cas à l'abandon de la lecture. Un étiquetage précisant ces termes aurait pu éviter le phénomène. **L'étiquetage** a d'ailleurs lui-même été pointé du doigt : absent, mal réalisé ou affiché visuellement hors des conventions habituelles, il a pu, dans une dizaine d'occasions, renforcer l'incompréhension due à un manque d'information.

T : Tu me dis que ce qui n'est pas clair, ce n'est pas la présentation des données mais le titre ? Si j'avais « évolution du repas moyen au fast food » ça irait ?

A : ça ne m'aurait pas suffi... je m'en fous des grammes à la limite. J'aurais plutôt mis une échelle graduée. Pour situer les grammes par rapport à une échelle et de me dire Ok, dans le premier c'est un peu près au rapport de 1 à dans le 2è et dans le 3è, au total, en moyenne, si on veut considérer la boisson dans le repas, on est dans un rapport de 1 à « autant » et sans la boisson, on est dans un rapport de 1 à « autant ».

Le **mauvais encodage des données** a également provoqué l'incompréhension. Il en va de même pour la **difficulté de comparer** des formes sphériques en termes de proportions. Les barres empilées, même de type standard, ont aussi été mentionnées pour la difficulté de comparer des proportions. Quoi qu'il en soit, la difficulté de comparaison a été mentionnée dans d'autres cas, notamment lorsque le nombre de données est trop important ou que les efforts esthétiques distraient cette comparaison. Ainsi, un visuel trop chargé de par son esthétique ou par le nombre de données provoque donc de la confusion (mentionné 22 fois).

Nous abordions en début de chapitre les visualisations qui présentent plus d'une information à la fois, provoquant ainsi une sensation d'effort à fournir plus importante. Pour compléter cela, une douzaine de citations estiment ce phénomène comme facteur d'incompréhension.

Finalement, les problèmes de compréhension dus à une mauvaise maîtrise de l'**optique** ont de nombreuses fois été cités (26 fois). Le parfait exemple est la visualisation n°4.

Celle-là c'était horrible ! On ne distingue parfois même pas certaines couleurs ! ça fait mal aux yeux... Et quand on veut aller plus en profondeur, surtout dans les roses et rouges... J'ai même pas pu distinguer ! Il y aurait dû avoir moins de données et des couleurs distinctes...

Pour résumer, la sensation d'un effort amoindri, que nous avons déjà abordée, contribue à la clarté d'une visualisation, tout comme la disposition des éléments graphiques qui présuppose un ordre de lecture logique. En ce qui concerne les facteurs d'incompréhension, ils sont classés par importance dans la Figure 44.

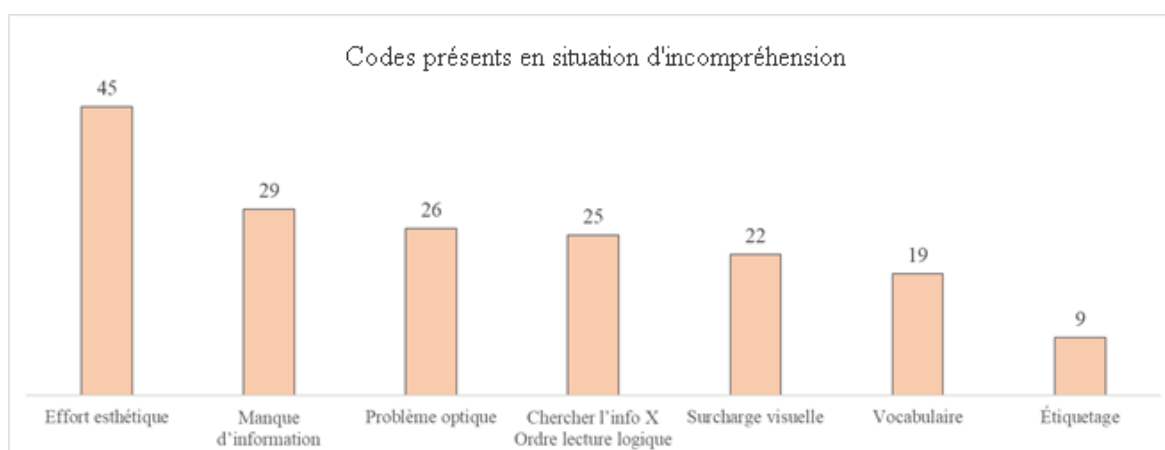


Figure 44 Nombres de codes menant à l'incompréhension / la confusion

2.2. Les facteurs d'appréciation des visualisations

Dès à présent, nous nous concentrons sur les éléments constitutifs des différentes visualisations de données du corpus. Pointés par les participants, ces derniers nous permettront de définir quelles sont les forces et les faiblesses de l'ornementation, c'est-à-dire quels éléments mènent vers une appréciation positive ou négative de la visualisation selon l'utilisateur.

2.2.1. Vers une appréciation positive

L'ensemble des participants a indiqué aimer une visualisation 74 fois en tout. Ces appréciations positives concernaient 61 fois des visualisations ornementées. Elles ont concerné 13 fois des visualisations de type standard. Nombreuses sont les raisons pour lesquelles une visualisation peut être appréciée ou pas. Le critère le plus important est celui de la clarté. Une visualisation bien claire est donc susceptible d'être appréciée selon les entretiens. Souvenons-nous des appréciations déclarées : statistiquement, la corrélation inter-échelles entre « J'aime » et « clair » n'est pas très élevée ($\rho = 0,437$). Les entretiens semi-directifs nous permettent donc de nuancer ce fait. Cela s'est plus souvent produit lorsque l'individu s'exprimait à propos d'une visualisation ornementée. Ainsi, une visualisation de type ornementée semble appréciée si le participant l'estime claire. Le second critère le plus mentionné est la capacité d'évocation, de rappel du sujet que l'ornementation apporte à la visualisation. Cela est très apprécié par les participants. L'ornementation peut parfois être vue comme une aide, un effort diminué, ce qui, de fait, rend la visualisation plus appréciable. La beauté, le visuel, le jugement de goût positif pour l'allure esthétique de la visualisation ont également été très soulignés : une visualisation ornementée qui est claire et « belle » aux yeux du lecteur est fortement appréciée. La simplicité y joue un rôle important : l'ornementation qui alourdit la visualisation avec une complexité visuelle est en général boudée. L'ornementation citée positivement par les participants est donc plutôt simple (pictogrammes ou dessins simples plutôt qu'une métaphore ou image complexe). L'originalité et l'aspect amusant qu'ils suscitent peuvent être un avantage également (cité 11 fois).

En somme, la visualisation ornementée est très appréciée si elle est claire, si l'ornementation est simple et si son esthétique plait. En ce qui concerne les visualisations de type standard, c'est le sujet de la visualisation qui influencera davantage une attitude positive de la part des participants, qui regrettent alors que rien ne soit « mis en scène » pour un sujet qu'ils affectionnent.

En résumé, les forces de l'ornementation sont

- La capacité **de rappel au sujet**
- L'**esthétique** stimulant le jugement de goût
- L'aspect **agréable, amusant**

Et ce, à condition que l'ornementation en elle-même soit **simple** et que la visualisation soit **claire**.

2.2.2. Vers une appréciation négative

Deux critères sont essentiels lorsque les participants donnent une appréciation négative à une visualisation de données : l'incompréhension des données et leur esthétique disgracieuse, souvent reprochée aux graphiques en barre. Cela menait à des réflexions telles que « *Donc je passe* » ou encore « *ce n'est pas gai* ». Ces deux éléments ont été mentionnées plus de dix-huit fois. On le voit, il a été plus simple pour les participants de mentionner les points négatifs des visualisations proposées plutôt que les positifs. Ensuite, les visualisations surchargées d'éléments généraient souvent un sentiment négatif. Qu'il s'agisse du nombre de données ou de l'ornementation trop chargée, il s'avère que souvent, la volonté de simplifier ces graphiques s'est manifestée. Ensuite, le désintérêt pour le sujet présenté a mené une douzaine de fois à l'explication d'une appréciation négative. Cela s'apparente au facteur de Kennedy et Hill (2017, 2018) sur les données et le sujet traité, qui influencent alors émotionnellement l'engagement du participant ou pas. Certains graphiques, de par leur sujet, ont été jugés ennuyeux, mais davantage lorsqu'il s'agissait de diagrammes en bâtonnets. Par contre, la beauté d'une visualisation lui a parfois fait défaut. « *C'est beau mais...* » : si les données sont peu claires ou en trop grand nombre, si cette beauté est exagérée, on glisse dès lors vers une appréciation négative. Certains éléments décoratifs et couleurs présentés dans les visualisations ornementées ont parfois fortement déplu (7 fois), ce qui dépend du jugement de goût personnel sur lequel le concepteur des visualisations a finalement assez peu de contrôle.

Ceci nous permet d'identifier les différentes faiblesses de l'ornementation aux yeux des participants :

- La **surcharge visuelle**, rapidement atteinte, diminue la clarté de la visualisation
- Une esthétique pouvant être démodée ou **déplaisante**
- L'**inutilité** de l'ornementation si les données et l'information sont mal présentées.

2.3. Critique de la disposition des éléments graphiques

Dès à présent, nous allons nous pencher sur les critiques positives ou négatives formulées par les participants en ce qui concerne les différents attributs visuels des visualisations de données. Les quarante participants ont effectivement souligné, de manière positive ou négative, le poids des différents éléments visuels présents sur les visualisations qui leur ont été proposées. Les codes suivants ont donc émergé : [artistique, unité de mesure, conventions visuelles, couleurs, étiquetage, forme ronde, problème optique, légende, ordre de lecture logique, simplicité].

Ainsi, certaines visualisations ornementées ont été qualifiées d'**artistiques** 6 fois. Positivement, cela réveillait l'intérêt des participants. Négativement, l'effort esthétique était réalisé au détriment de l'information.

L'**unité de mesure**, quant à elle, a été mentionnée 14 fois et mentionnée lorsqu'elle générait une appréciation négative. Un mauvais choix d'unité de mesure pouvait engendrer une incompréhension. Une mesure trop vague pouvait aussi donner au récepteur le sentiment d'un manque d'information, l'empêchant de comprendre. Le mauvais choix de l'unité de mesure a été expliqué 8 fois comme facteur d'incompréhension.

Ensuite, les participants se sont dit troublés lorsqu'ils étaient confrontés à des visualisations de données qui sortaient des **conventions visuelles** habituelles (15 fois). Cela provoquait ainsi un dérangement.

T : Pourquoi c'est pas clair ?

A : C'est la manière dont c'est organisé... Le fait que j'ai vu des graphiques qui avant allaient crescendo, là ça va decrescendo. L'échelle est inversée donc j'ai mis un peu plus de temps à la comprendre.

De même, l'allure **épurée** des graphiques préconisée par Tufte sort des **habitudes** des participants, qui estiment ainsi qu'il est plus difficile de comprendre ces visualisations puisqu'ils n'en ont pas l'habitude.

C : Au début je n'ai pas compris ce que voulaient dire les chiffres parce que je cherchais la barre verticale et au final, j'ai vu que c'était un simple graphique « nombre de morts par saison ». (...)

T : Tu as cherché un axe, par habitude ?

C : Oui, quand tu regardes une barre tu as l'habitude d'être référé à l'axe vertical.

Puis, les diagrammes en bâtonnets dont les barres étaient horizontales perturbaient certains utilisateurs qui disaient préférer lire l'information disposée sur des **barres verticales**.

J'avoue que je préfère les barres verticales. Parce qu'en général tu lis le titre, tu descends et tu as toutes les extrémités d'un coup et tu peux voir qui est le plus grand, le plus petit et analyser à ton aise quoi. Bizarrement, ça a toujours été utile pour mieux comprendre et te rendre compte du message à faire passer.

La **couleur** est l'élément visuel le plus cité par les participants : 105 fois au total. Ainsi, de nombreuses fois, l'absence de couleurs dans le cas des visualisations de type standard a été soulignée. Les participants, bien conscients que la couleur n'apporte rien à la compréhension, jugent son utilisation utile pour renforcer l'aspect agréable de la visualisation et pour attiser l'intérêt du récepteur. La couleur a été décrite une dizaine de

fois comme attirante. Cependant, sur les visualisations ornementées, les participants mentionnaient aussi quelques fois le surplus de couleur.

Par contre juste le gris ; mettre une couleur pour que ce soit un peu plus sympa à regarder mais sinon je trouvais ça bien.

T : Celle-là, qu'en penses-tu ?

A : On dirait qu'il y a beaucoup d'information et il n'y en a pas tant que ça, il y a trop de couleurs aussi...

T : C'est trop ? Trop de chichi ?

A : Trop oui ! Après c'est joli mais bon ... Les couleurs oui ça distrait quand ça ne sert à rien en plus là il y en a trop.

Les couleurs sont fun pour un sujet qui n'est pas forcément fun, ça sert plus à attirer l'œil. C'est pas forcément utile... ça dépend pour qui et pour quoi. Après le « jamais » est en rouge, ça attire l'attention sur le mauvais, la plupart du temps c'est plutôt neutre et le toujours est en vert donc les couleurs représentent bien ce que ça signifie au final.

De plus, il va sans dire que les couleurs sont sujettes aux **jugements de goût**. Des visualisations n'ont tout simplement pas été appréciées parce que les participants d'un point de vue personnel n'aimaient pas les couleurs employées.

Parce que ça met de la créativité à l'image et à la façon donc on va répondre à la question. On va peut-être plus s'attarder à regarder une image colorée parce qu'elle est plus jolie qu'une image en noir et blanc. Une image qui est jolie ne veut pas dire qu'elle est intéressante ou claire. La dernière image que tu m'as montrée, elle était jolie mais je n'ai rien compris.

Dans les cas où elle est présente en **surcharge**, la couleur est un facteur de confusion. Cependant, elle apporte la clarté lorsqu'elle entre en cohérence avec le message communiqué.

T : ça c'est plus conventionnel mais tu trouvais ça beau aussi. C'est une question de couleur ?

A : Oui on a compris que mauve c'est la crème glacée, on voit que la part diminue à chaque fois. C'est visuel, coloré et simple donc j'aime bien. J'ai compris rapidement.

Par ailleurs, l'**étiquetage** et l'**absence d'axes** ont aussi largement été commentés. De nouveau, les commentaires des participants entrent en contradiction avec les conventions minimalistes préconisées par Tufte. Une dizaine de fois, les participants ont exprimé qu'il était perturbant de voir des visualisations sans axes gradués et aux étiquettes de données écrites à la fin des barres. L'endroit inadéquat de l'étiquette, sur des visualisations de type standard comme ornementé, a été mentionné, tandis que l'absence d'un chiffre exact a souvent dérangé.

Peut-être que le fait qu'il y ait trop de chiffres, ça m'a embrouillé et j'ai cherché sur l'axe vertical et pas trouvé. Encore une fois, il n'y a pas le chiffre exact. Peut-être que ça dépend de ce qu'on veut représenter, si au Royaume-Uni on importe plus, mais si tu veux le chiffre exact c'est pas clair.

C'est assez redondant... On voit quand même que les 18-30 ans consomment le plus à l'apéritif et c'est tout. Donc on peut enlever le texte, ce qui diminuerait la longueur.

Une quinzaine de fois, la **forme ronde** a été décrite comme plaisante, bien que son problème principal soit de devoir chercher l'information dans les différentes bulles présentes sur la visualisation. Certains expliquent la difficulté de différencier les proportions. D'autres le voient comme une aide malgré la réelle difficulté perceptuelle qu'engendre l'utilisation d'aires.

C : Oui. Je n'avais pas vu du tout ces icônes elles m'ont pas du tout aidé. Les gros ronds m'ont permis de comprendre les jours différents pour lire une série. Après, ce que j'aime bien c'est que c'est des ronds et pas des bâtonnets donc ça change un peu.

T : ça t'a aidé ?

C : Les ronds ? Non, le gros est censé attirer l'œil, mais tu sais pas quoi regarder, tu ne comprends pas comment c'est fait... Moi ça m'aide plus quand c'est classé par ordre croissant et décroissant et là c'est un peu perturbant, c'est ça qui m'a gêné dans ce graphique. Mais niveau compréhension c'est pas dingue quoi...

M : Les tailles de ballon comme ça, je trouve ça sympa et le fait qu'il y ait des icônes pour illustrer des séries je trouve ça sympa aussi. Le fait que ce soit sur la même échelle comparaison c'est bien aussi. Le fait que ça ne soit pas juste 5 saisons de 10 épisodes, non, là, on donne une unité de mesure qui donne 1 jour = 8h et je

crois que je l'ai déjà dit mais le fait que la taille des ballons soit proportionnelle ou plus ou moins à ça c'est très bien.

Par ailleurs, une visualisation en particulier a provoqué la même gêne à chaque participant. Le problème **optique**, effectivement, causé par une succession de points aux couleurs mal choisies, a empêché la lecture de la visualisation n°4. Nous l'avions déjà précisé.

La **légende**, quant à elle, a été mentionnée 9 fois, souvent dans un sens négatif. Les participants, lorsque la légende était sur le côté ou sous le graphique, l'ont estimée mal positionnée. Un ordre de lecture « logique » a souvent été réclamé par ces individus qui estimaient que cela manquait.

T : Celle-là est très mauvaise élève. Pas belle et pas intéressant, à moitié compris...

L : C'est parce que j'ai d'abord regardé à gauche puis je me suis demandé où était la France, et puis j'ai tellement regardé à droite je pense et du coup je me suis dit ok donc ça concerne l'international. On lit de gauche à droite donc je pense que dans l'autre sens ça aurait été mieux.

Ce n'est pas ordonné en fait, en haut on a le groupe des invertébrés mais on commence par les huîtres et puis on a les insectes et on va du plus grand au plus petit et puis pour les autres groupes on va du plus petit au plus grand si je comprends bien et donc ce n'est pas logique.

A en croire les participants, la lecture est facilitée lorsque la légende est placée directement sous un titre qui est facilement compréhensible, afin de saisir les informations en un seul trajet oculaire. Cet ordre de lecture que les participants estiment plus logique a été mentionné 30 fois.

De même, la **simplicité** du visuel a énormément été mentionnée : 60 fois au total. Une douzaine de fois, la simplicité a été expliquée comme un facteur **aidant** l'individu. Une dizaine de fois, elle apporte la clarté, ce qui est cohérent avec les principes de Bertin et Tufte. La simplicité du visuel ne concerne toutefois pas uniquement les visuels de type standard. Elle concerne également les visuels de type ornementé, mais épurés dans le style esthétique qui leur est propre. À 31 occasions, les participants ont souligné la simplicité de la visualisation ornementée en général. Les 25 autres citations appréciant la simplicité du visuel ont été émises à propos de visualisations de type standard. Les 4 dernières citations parlaient du sujet de manière générale, en l'appréciant positivement.

L'interprétation a jusqu'ici porté sur les différents éléments visuels constituant la visualisation et sur la manière dont les participants les ont considérés. Cependant, à travers les données qualitatives récoltées, il est possible de se questionner à propos du rapport à l'image des participants, dans le contexte de l'appropriation d'une visualisation de données. Effectivement, les participants ont énormément commenté la présence ou l'absence d'images, de « dessins » et de pictogrammes.

La présence ou absence d'**images** a été mentionnée 39 fois. Pour certains, il arrivait que les visualisations de données manquent d'une « **mise en contexte imagée** ».

E : C'est gris, il n'y a pas de mise en valeur du titre, des infos mais je ne sais pas s'il en faut plus parce que au moins on a l'information très vite. C'est étonnant il y a d'habitude un axe sur le côté et là c'est en fait beaucoup plus clair d'avoir les morts directement au-dessus de la barre. Mais c'est fluctuant...

T : Tu aurais changé quelque chose ?

E : J'aurais vu un décor ou quelque chose qui se rapporte à la série, un fond peut-être.

A 25 reprises, les participants expriment la nécessité de **dynamiser**, d'imager les visualisations de données standard. La plupart du temps, l'objectif exprimé est de rendre l'expérience du lecteur agréable.

B : Je l'ai trouvé claire mais... pas efficace.

T : Efficace pour toi, c'est ... ?

B : Il faut vraiment faire un peu d'effort pour s'intéresser et aller en profondeur mais il peut être rendu attractif. Le sujet est intéressant, ce sont des informations que je n'ai pas l'habitude de voir mais à la fin on s'ennuie... Et la couleur c'est aussi ennuyeux

T : Les couleurs changent quelque chose ?

B : Oui et on peut facilement trouver des illustrations de quelqu'un qui regarde un film à la télévision, de films en ligne ou... On sait clairement marquer la différence en bas avec des illustrations et ça prendrait un minimum de forme...

L'envie d'illustrer des visualisations de type standard a été exprimée une dizaine de fois, le tout pour marquer l'esprit d'une part mais aussi pour faire allusion au sujet de la visualisation. De même, il est important pour le participant que l'image constitue un bon compromis, comme on l'a vu précédemment, entre **simplicité et esthétique**, voire même « décoration », dans un but attractif. De plus, les images, selon les participants, constituent

une aide à la mémorisation, ce qui constituerait une économie dans la lecture également. La cohérence entre le texte et l'image est soulignée et est inévitable. De nouveau, cela a été mentionné par Borkin et ses pairs (2013, 2016).

Au-delà de « l'image » ou du « dessin », les participants ont énormément mentionné le **pictogramme** (65 fois au total) qui, comme nous l'avons expliqué précédemment, est extrêmement apprécié pour de nombreuses raisons : la simplicité et le style épuré du pictogramme en tant qu'élément de communication, son accessibilité ainsi que la capacité de rappel, d'allusion au sujet que peuvent provoquer ces pictogrammes. Une dizaine de fois, le pictogramme a été expliqué comme un élément amenant de la clarté à la visualisation de données, voire une aide réelle à la compréhension (mentionné une dizaine de fois également). Le pictogramme est donc un élément amenant clarté, aidant la mémorisation et qui attire également l'intérêt du lecteur et ce, peu importe la forme que prend le pictogramme, appelé *pictograph* dans le travail d'Haroz et ses collègues (2015) qui, pour rappel, ont étudié la réception de différents types de pictogrammes (cf. [supra](#))³⁵.

G : C'est pas sexy, le sujet peut être intéressant pour ceux qui sont fan d'animaux mais c'est pas un truc qui m'excite. Après c'est intéressant de voir la différence. Je me suis demandée s'il y avait vraiment 66% des argentins qui avaient un chien. La barre du dessus est en trop et ça manque de couleurs. Ici au lieu des mots par contre ça aurait pu être des pictogrammes.

T : le visuel t'a rebuté ?

G : C'est triste quoi, ça manque de couleur, c'est vintage (rires). Ça dépend dans un rapport hyper sérieux on a plein d'outils aujourd'hui pour faire des graphiques plus sympas que ça, c'est juste ça... Pour faciliter la lecture.

En résumé, les participants ont été capables de critiquer les différents éléments visuels composant les visualisations de données. En qualifiant certaines visualisations d'artistiques, les participants ont insisté sur leur attrait mais également sur le potentiel danger pour la visualisation de ne pas se consacrer pleinement à l'information. Un mauvais choix de mesure peut également mener à l'incompréhension, tandis que sortir des conventions visuelles et des habitudes de lecture des participants peut porter à confusion. Dans les visuels épurés, l'absence d'axes, d'images ou encore la disposition des étiquettes de données peuvent perturber les lecteurs qui n'y sont pas habitués. De plus, même lorsqu'elles s'avèrent inutiles, les couleurs semblent en réalité indispensables aux utilisateurs pour leur aspect agréable. Attention par contre de ne pas basculer dans la

³⁵ Voir la [section 2 du chapitre 2 de la Partie I.](#)

surcharge de couleurs, menant encore une fois vers la confusion. Finalement, la présence d'images simples (et de pictogrammes en particulier) est hautement appréciée lorsque la visualisation de données allie simplicité, esthétique et mise en contexte.

2.4. Les facteurs humains et émotionnels : des fondamentaux

A présent, il est opportun d'aborder les facteurs humains et émotionnels qui sont entrés en compte dans l'appropriation des visualisations de données présentées aux participants. Effectivement, après avoir étudié les opinions des participants, s'attarder sur différents facteurs qui leur sont propres nous permettra d'avancer plus en profondeur dans l'analyse. Il va de soi que les codes thématiques relatifs à cette analyse sont principalement inspirés des facteurs humains et émotionnels de Kennedy et Hill (2017, 2018). Nous nous sommes permise de les adapter en fonction de l'analyse. De plus, tous les facteurs ne sont pas forcément apparus dans le discours des participants.

Dans l'analyse, nous avons donc pris en compte ces éléments, directement inspirés des facteurs humains et émotionnels de Kennedy et Hill (2017, 2018) :

- Les compétences, dont
 - o Les compétences de langage ;
 - o Les compétences mathématiques et statistiques ;
- Les éléments qui ont un rapport à soi, comme
 - o La confiance en soi ;
 - o L'intérêt du sujet ;
 - o L'implication personnelle dans le sujet ;
- L'incidence du visuel
→ a été traité précédemment ; nous n'y reviendrons pas.
- Les émotions et les sensations, comme
 - o Appréciations positives et négatives (traité précédemment, nous n'y reviendrons pas)
 - o Sensation d'amusement ;
 - o Sensation agréable ;
 - o Aspect sérieux, scolaire, professionnel ;
 - o Sensation de beau, de laid ;
 - o L'envie de lire ;
 - o La sensation d'ennui.

Tous les facteurs ne sont pas forcément apparus dans le discours des participants. Nous avons uniquement traité ceux qui ont émergé de la discussion.

2.4.1. Analyse des facteurs humains

Pour commencer, les participants ont évoqué leurs **compétences de langage** à 19 reprises. Effectivement, rien d'étonnant puisque deux visualisations de données en particulier ont posé problème et ce, qu'elles aient été présentées sous une forme standard ou ornementée. Certains participants ont eux-mêmes mentionné la gêne occasionnée par ce problème de vocabulaire, tandis que d'autres ont tout simplement expliqué ne pas comprendre. Quoi qu'il en soit, ce souci a constitué une réelle entrave à la compréhension et ce, dans tous les cas. Les visualisations concernées sont les n°12, 18, 32 et 38. La plupart des participants n'ont pas compris ce qu'était une « chaire » à l'université et n'étaient pas familiarisés au vocabulaire relatif à *Star Wars*. Évidemment, ce vocabulaire est lié aux connaissances ou à la culture. La méconnaissance du vocabulaire employé a souvent mené à l'abandon de la lecture de la visualisation.

C : Voilà, soit mon vocabulaire n'est pas assez développé et je ne connais pas, soit c'est un mot qui aurait dû être expliqué parce que c'est spécifique et personne ne connaît. J'ai été honnête, j'aurais pu faire comme si j'avais compris... Mais non, c'est clair dans le sens où on voit qu'il y a plus d'hommes que de femmes, mais utiliser un mot si spécifique... Soit tu le mets dans un programme que tout le monde connaît et comme ça t'es sûr de ton message, mais ici je ne savais pas ce que c'était donc il aurait peut-être fallu une définition de « chaire » avec un astérisque en dessous... Quelque chose qui puisse faire comprendre, si tu mets ça dans une revue pour tout le monde, tout le monde ne sait pas ce que c'est donc pour moi c'est pas clair à cause de ça.

Ensuite, les **compétences mathématiques** et statistiques ont également été mentionnées, mais dans une moindre mesure (une dizaine de fois). Ce manque de compétences – ou du moins les difficultés que ces compétences engendrent, a constitué plusieurs fois une mauvaise expérience pour les participants.

T : Tu n'appréciais pas les graphiques en barre ?

E : Bah disons que c'est le graphique classique qu'on utilise dans les travaux quoi, celui que t'as fait en mathématiques et qui te laisse un goût amer.

J'ai mis du temps, j'ai regardé 36x la légende et l'explication du haut... Il a fallu que je réfléchisse, j'ai trouvé ça intéressant mais... j'en sais rien, j'ai l'impression que soit tu mets les infos verticalement ou bien sur l'axe. En tout cas, dans mon cas, je suis pas très chiffré.

En ce qui concerne les éléments relatifs au **rapport à soi**, ils incluent selon nous différents facteurs comme la confiance en soi, l'intérêt porté au sujet de la visualisation et l'implication personnelle dans le sujet de la visualisation.

La confiance en soi, à mobiliser au moment de la lecture du graphique, a été mentionnée une douzaine de fois. De manière générale, il s'agit plutôt d'un manque de confiance en soi, qui génère alors une appréciation négative de la part du participant. Le manque de confiance en soi peut effectivement mener à de l'incompréhension. Néanmoins, quelques fois, dans une bien moindre mesure, la profession du participant (marketing, employé dans une agence, ...), dans le domaine du visuel, mène justement à une confiance en soi qui lui permet d'apporter un regard critique sur la visualisation.

Il me manque une info pour dire de quoi on parle. En fait ce qui me dérange... oui c'est ça le mot, « dérange », c'est que j'ai eu l'impression de dire que j'ai bien compris, mais que je n'ai rien pour me rassurer et me dire que j'ai en effet bien compris. C'est plus un problème de confiance en moi. J'ai toujours des doutes donc s'il n'y a pas quelque chose dans ce visuel qui va me convaincre qu'il n'y a pas de doute à avoir et bien je vais me sentir perdue et je ne vais pas arrêter de me demander si je suis sûre de ce que j'ai compris.

Par ailleurs, l'intérêt que le participant éprouve pour le sujet de la visualisation est extrêmement important dans la compréhension, l'appréhension et la construction de sens de l'utilisateur. **L'intérêt du sujet** pour le participant est apparu 111 fois dans les différentes conversations. L'intérêt du sujet a été mentionné 57 fois lorsqu'une visualisation de type standard était concernée, contre 52 fois lorsqu'il s'agissait de visualisations ornementées. Le type de visualisation a finalement peu d'importance dans ce cadre.

Ainsi, lors de la lecture d'une visualisation de type standard, si l'esthétique simple et sans fioritures déplait au participant, il la dépasse afin de prendre connaissance de l'information si celle-ci l'intéresse. Par contre, si le sujet n'intéresse pas le participant, il ne se dépasse pas pour saisir l'information, ce qu'il ferait peut-être si le visuel était ornementé. Ce constat est important car il répond à un vide laissé par l'analyse statistique : effectivement, l'*item* « Intéressant » récolté dans le cadre des appréciations déclarées semblait peu précis et n'a rien révélé d'important au cours de cette analyse. Par contre, l'intérêt du sujet analysé de façon qualitative et selon ce qu'entendent Kennedy et Hill (2017, 2018) permet de mieux saisir la subtilité de l'intérêt du sujet pour les participants.

Je l'avais bien aimée de par le sujet. Ça m'a vraiment plu. Je suis fan de voyages et d'Insta. Ça me parle à fond... Parce que l'esthétique encore une fois est pas dingue mais le sujet m'a emmené.

Je n'aime pas parce qu'elle est banale. Ce n'est pas trop compliqué me semble-t-il. La présentation est moche. Les données apparaissent clairement et sont intéressantes. Elle est clairement compréhensible mais elle n'est pas terrible.

Elle est en gris, on aurait dit que c'est fait exprès pour qu'on ne la trouve pas belle. En plus, c'est pas intéressant. On ne se focalise pas vraiment... c'est toujours un peu moins plaisant de regarder une infographie moins jolie. Enfin, l'information peut-être tout aussi intéressante mais c'est plus sympa de la regarder dans une infographie qui est chouette.

Pour terminer, nous avons constaté plusieurs cas dans lesquels l'**implication personnelle** du participant dans la visualisation de données l'incite à davantage s'approprier la visualisation. A 41 reprises, le fait de se sentir concerné ou non par la visualisation a été soulevé. Le participant pouvait, par exemple, rechercher son propre pays ou son animal préféré dans une visualisation qui lui était présentée. Le fait d'être concerné par la visualisation a incité l'engagement. Au contraire, le participant qui se cherche dans la visualisation sans s'y trouver pourra finalement marquer un désintérêt.

T : Qu'est-ce qui t'a vraiment plu ?

L : C'est la couleur, c'est de voir que certains bâtiments que j'ai vus sont des lieux phares. Et je me suis par exemple étonnée de la Sagrada Familia qui était moins visitée que d'autres choses.

C'est chouette, il y a des choses que j'ai déjà visitées donc j'ai été surprise de voir quand même le nombre et le chiffre final.

J'ai compris que ça parlait d'alcool mais je n'avais pas vu que ça parlait forcément de vin, c'est en lisant le titre que je m'en suis rendu compte. Je la trouve bien faite. On décompose en détail avec les 60 % des parts de vente qui sont en France et dans les 40 %, on voit en détaillé tous les importateurs. Je trouvais ça sympa et je me demandais où en serait la Belgique. En fait, dans toutes les visualisations où il y avait des pays, je me demandais où en était la Belgique, mon pays.

2.4.2. Analyse des facteurs émotionnels

Pour cette partie, nous avons distingué les éléments relatifs aux émotions mais également aux sensations, qui précèdent l'émotion.

Avant tout, la **sensation agréable** a été mentionnée 53 fois au total par les participants. La couleur, les pictogrammes et les différents efforts esthétiques renforcent cette sensation, pour autant que cette esthétique plaise au participant. L'effort réalisé par le concepteur est ainsi grandement apprécié. Une sensation de clarté et d'effort amoindri en découlent très souvent. De même, l'aspect agréable attise l'intérêt. Nous ne nous étendrons pas plus sur le sujet qui a été traité précédemment dans cette section.

Les participants ont également mentionné une **sensation d'amusement** au contact de certaines visualisations de données. Dans le cas contraire, il arrivait qu'ils réclament cet amusement. De plus, ce qui peut générer une appréciation positive et de l'amusement peut également apporter de la confusion. Là où l'effort esthétique apporte l'originalité qui amuse, il arrive que des éléments visuels essentiels à la transmission et à la réception de l'information soient transformés, jusqu'à rendre confus le participant. Prenons pour exemple la visualisation n°19, de type ornementé, qui montre la proportion d'animaux retrouvés décapités dans les parcs de New York. Des coulées de sang rouge, pointant vers le bas et classées par ordre décroissant, étaient présentées aux participants. Alors que certains d'entre eux la trouvaient « parlante », « choc » ou « fun », d'autres regrettaient que la forme de la coulée les empêche de saisir le chiffre exact, puisque l'extrémité de la coulée, similaire à une barre, était arrondie et donc ne permettait pas clairement de saisir l'information. Cela aurait pu être rapidement contrôlé, en ajoutant par exemple une étiquette de données.

Il n'y a pas de délimitation précise pour la goutte de sang et il faut réfléchir où elle s'arrête. (...) n'est pas hyper clair. Et en plus on a une échelle de 0 à 50 ok mais c'est 50 quoi ? C'est en unité ou en dizaines je n'en sais rien.

J'ai beaucoup aimé parce que c'est clair en fait, il n'y avait même pas besoin de mettre le nom des animaux en dessous et puis au niveau esthétique, que ce soit fait en dessin à la main ça m'a plus capté que quelque chose qui est fait à l'ordinateur, ça change. C'est original.

Comme nous l'avons indiqué dans le premier chapitre, la sensation d'amusement est inextricablement liée à la **sensation agréable** ainsi qu'à l'**attisement de l'intérêt** du participant. La sensation d'amusement peut être stimulée par le sujet de la visualisation ou par son esthétique. Cependant, lorsque le sujet de la visualisation est amusant aux yeux du

participant mais que la visualisation est de type standard, l'absence d'esthétique « ludique » est souvent regrettée. De fait, tout est lié : cela a évidemment un rapport avec l'univers dans lequel s'insère la visualisation.

C'est marrant, c'est assez esthétique et de prime abord je pense que ça colle à la charte graphique de la série. Ça m'aide à contextualiser, je comprends de quoi ça parle ouais.

Les participants ont également mentionné le ressenti d'une visualisation à l'**aspect sérieux**, scolaire et professionnel. Cela concernait davantage des visualisations de type standard. Les visualisations ornementées quant à elles peuvent sembler professionnelles si l'ornementation employée est simple, avec des pictogrammes. Par ailleurs, les visualisations de type standard peuvent donner l'impression au participant que le concepteur n'a pas souhaité accorder de soin à la visualisation. Cette remarque a été formulée plusieurs fois.

Ça a une connotation sérieuse, scolaire et puis, pour les gens, un graphique c'est ça. Maintenant, on a des outils qui permettent d'évoluer donc pourquoi ne pas en profiter ?

Certaines visualisations peuvent également donner aux participants une **sensation d'ennui** (à 58 reprises). La visualisation de type standard peut générer une sensation d'effort à fournir. L'ennui peut parfois mener à l'incompréhension, l'abandon ou ne pas générer d'intérêt.

Je pense que si tu prends ça sur un feed Facebook, c'est pas le genre de choses qui va attirer le clic d'onglet ! Si tu mets un petit paquet de frites, un gros paquet de frites, de la couleur, c'est fun fact ou un truc dans le genre, ça va tout de suite attirer l'œil et l'intérêt.

Oui, c'est ennuyeux et on n'a pas envie d'aller voir si on ne sait pas de quoi on va parler. C'est un graphique... Voilà et je ne vais pas faire l'effort d'aller voir de quoi ça parle.

Pour terminer, les dernières sensations qu'ont ressenties les participants étaient des **sensations esthétiques** (le beau, le laid). La beauté des visualisations a été soulignée 125 fois. Cette sensation prédispose évidemment positivement le participant à s'approprier la visualisation de données. La beauté a été soulignée sur des visualisations qui semblaient claires aux participants et ce, 37 fois. A 29 reprises, la visualisation qui semble jolie génère pourtant de l'incompréhension. Le fait de trouver une visualisation laide n'encourage pas son appropriation. A 25 reprises, une visualisation considérée laide restait claire malgré

tout. A seulement 10 reprises, une visualisation de données a été jugée laide et incompréhensible, ce qui est dérisoire.

On constate ainsi que les différents facteurs humains et émotionnels de Kennedy et Hill (2017, 2018) se retrouvent pleinement dans le discours des participants. Leur existence ou le fait qu'ils soient stimulés influence l'appropriation des utilisateurs et permet de comprendre dans quelle mesure l'ornementation peut répondre à certains problèmes posés par ces facteurs.

3. Conclusion de l'analyse thématique

Les conclusions de l'analyse thématique sont nombreuses et donnent accès à des pistes ayant été peu explorées auparavant dans la littérature, ou du moins assez peu d'un point de vue aussi qualitatif.

Pour commencer, la sensation d'effort est fondamentale dans l'appropriation des visualisations de données. Certains éléments visuels peuvent amoindrir ou renforcer la sensation d'effort à fournir, en stimulant également d'autres sensations comme la sensation agréable, d'amusement ou encore d'attractivité. L'ornementation permet d'augmenter l'attractivité de la visualisation et de donner au lecteur l'envie de la lire. Néanmoins, toutes les ornements ne mènent pas forcément à ce bénéfice : parfois, l'esthétique peut sembler dépassée ou peu moderne, tandis que le simple diagramme en bâtonnets semble rester « traditionnel » sans pour autant être vieillot.

L'ornementation peut parfois mener à une surcharge visuelle dont l'effet est dès lors négatif. D'autres surcharges existent : trop de données, trop de texte ou encore plus d'une information par visualisation jouent en défaveur de la visualisation également.

Par ailleurs, de nombreux facteurs mènent à l'incompréhension. Nous en avons retenu plusieurs : un effort esthétique utilisé à mauvais escient, le manque d'information, les problèmes optiques, le vocabulaire, etc. Nous avons également pu identifier les forces et faiblesses de la visualisation. Parmi les forces, nous avons noté la capacité de rappel au sujet, l'esthétique stimulant un jugement de goût positif ou encore l'aspect agréable et amusant et les aspects positifs qui en découlent. Parmi les faiblesses, nous retenons que rien ne garantit qu'une esthétique ne soit pas déplaisante, que le danger de la surcharge visuelle reste présent et que, surtout, rien ne vaudra jamais une information correctement présentée et ce, indépendamment de tout ornement.

L'interprétation des discours en fonction de divers facteurs humains et émotionnels de Kennedy et Hill (2017, 2018) montre qu'il y aura toujours une part d'incertitude concernant

l'utilisateur. Effectivement, le concepteur n'a aucun contrôle sur le sentiment de confiance en soi ou sur les compétences mathématiques de son destinataire par exemple. Néanmoins, montrer que le rôle de ces facteurs n'est pas hypothétique mais réel permet de se préparer à toute éventualité.

Finalement, la conclusion la plus importante de cette analyse thématique concerne la simplicité visuelle des visualisations de données. Effectivement, nombreux sont les atouts de l'ornementation. Malgré tout, il semblerait que la préconisation de Tufte, « *less is more* », s'applique également en la matière. La simplicité des ornements (images, pictogrammes, apparence visuelle dans sa globalité) a largement été soulignée. Toutes les ornements ne s'équivalent pas : il semblerait que les ornements les plus simples et épurés apportent bien plus de bénéfices que leurs homologues surchargés.

Une prochaine piste dans le cadre de cette thèse sera donc de comparer les différents trajets oculaires des participants en fonction du type d'ornementation (épurée ou chargée). Nous y viendrons dans la [Partie V](#). Pour le moment, en complément de la présente analyse, nous allons nous concentrer sur les différentes stratégies mises en place par les participants lorsqu'ils ont été confrontés à la difficulté liée à la présence d'ornementation.

Chapitre 2 Les comblements de discontinuité ou le *gap bridging* dans la théorie du *sense-making*

Ce deuxième chapitre concernant la subjectivité dans la construction de sens se concentre sur les actions des participants au cours de leur processus de lecture. Comment les participants font-ils sens de l'information, en particulier lorsqu'ils sont confrontés à une situation ne leur permettant pas une lecture logique et intuitive ? Souvent, il arrive que l'ornementation brise la caractéristique préattentive des visualisations de données. Grâce à une adaptation de la *sense-making methodology* de Dervin et ses collègues (2003), nous nous penchons plus en détail sur les stratégies mises en place par les individus pour faire sens des visualisations de données dont l'ornementation ne permet pas une appropriation quasi-spontanée de l'information. Nous évaluerons les bénéfices et les inconvénients apportés par les différents pontages opérés dans le cadre de ce que Dervin et ses pairs appellent le *gap bridging*.

1. Méthode d'analyse inspirée de la *sense-making methodology*

Comme nous l'avons précédemment exposé dans la [Partie II](#), c'est sur le tard que nous avons découvert la méthodologie du *sense-making* de Dervin. Nos données étaient déjà récoltées. Nous n'avons donc pas mené nos entretiens en posant exactement les questions préconisées par Dervin dans sa méthodologie. Néanmoins, les discours des participants sont naturellement porteurs de difficultés qu'ils ont rencontrées, ainsi que de stratégies mises en place pour pallier ces problèmes. En nous inspirant de la méthodologie de Dervin appliquée dans différentes études (Lezon Rivière & Ihadjadene, 2014 ; Reinhard & Dervin, 2015 ; Spurgin, 2006), nous sommes en mesure de nous concentrer sur le pontage réalisé par les individus en situation de discontinuité. En utilisant la méthode du *verbing*, nous pouvons ainsi nous concentrer sur les actions des individus plutôt que sur les visualisations en elles-mêmes. Le *verbing*, qui pousse ainsi à employer des verbes plutôt que des noms afin d'aborder l'action, est pratiqué lors de l'entretien et de l'analyse. En entretien, il s'agit de poser certaines questions amenant le participant à expliquer ce qu'il a ressenti et comment il a agi. Lors de l'analyse, il s'agit de coder les verbatims clés à l'aide de verbes, afin de définir l'action décrite par le participant (Dervin, 1998).

Nous avons des raisons de croire que nos entretiens peuvent être analysés selon la méthode de *verbing*. Lezon Rivière et Ihadjadene décrivent l'entretien idéal de *sense-making methodology* comme ceci : d'abord la situation de l'individu est abordée, en étant un « récit d'expériences comportant des contraintes, des barrières dans le présent et en connexion avec le temps passé » (Lezon Rivière et Ihadjadene, 2014, p.13). Nous nous sommes assurée de poser aux participants des questions en lien avec l'affinité envers les visualisations de données, les habitudes et les différentes contraintes rencontrées à la lecture de chaque visualisation. Ensuite, l'entretien idéal devrait aborder la situation de discontinuité (*gap*) « exprimée par les questions, les interrogations, les confusions, les inconnu(e)s et les craintes de l'acteur » (Lezon Rivière et Ihadjadene, 2014, p.13), avant d'aborder les pontages réalisés par les participants pour combler les discontinuités. Ces pontages sont constitués « à l'aide des idées, de la cognition, des pensées, des attitudes, des croyances, des valeurs, des sentiments, des émotions, des intuitions, des souvenirs, des histoires ou encore des récits » (Lezon Rivière et Ihadjadene, 2014, p.13). Les participants ont souvent proposé une description de leur propre appropriation des visualisations proposées. Le verbatim suivant est représentatif de la manière dont, généralement, les participants ont expliqué comment ils faisaient sens de ces visualisations. On y retrouve le problème (*gap*), la stratégie permettant d'y remédier (pontage) ainsi que la situation finale.

T : Et donc pour rebondir sur la couleur nous avons ceci.

O : Oh lala ça m'a fait mal aux yeux, j'ai quand même essayé de faire un effort mais les couleurs se mélangent et pour distinguer « soupe », « salade », je dois avouer que... Non ce n'est pas le genre de graphiques qui me plaît... Autant ça peut paraître simple dans le sens où ce ne sont que des couleurs avec un nom, mais visuellement, je n'arrivais pas à faire la distinction et ça m'a... J'ai vite passé !

Dans le verbatim précédent, la discontinuité apparaît comme ceci : la participante perçoit difficilement les données ; elle les voit mal. En pontage, elle mentionne qu'elle a fait un effort en insistant. Finalement, elle explique qu'elle a abandonné, ce qui constitue la situation finale. Au passage, on apprend également que ce genre de graphiques lui déplait. Il est possible d'affecter un verbe à chaque situation : mal voir (*gap*), faire un effort (pontage) et abandonner (situation finale). Étant donné que les entretiens ont porté sur une expérience menée en laboratoire et donc indépendamment de tout contexte, il est difficile d'établir une situation initiale qui est différente pour chaque visualisation. La situation initiale est presque la même dans tous les cas : les individus sont venus au laboratoire en sachant qu'ils participeraient à une étude sur la visualisation de données. Ils ont tous des affinités avec la communication numérique et, évidemment, des caractéristiques personnelles propres et singulières. Contrairement aux études que nous avons mentionnées (Lezon Riviere & Ihadjadene, 2014 ; Reinhard & Dervin, 2015 ; Spurgin, 2006), nous ne disposons pas de situations initiales sur lesquelles pratiquer l'analyse grâce au *verbing*. Néanmoins, décrire la discontinuité, le pontage et la situation finale pour chaque verbatim concernant la description d'une appropriation de visualisation et y attribuer des *verbings* pourra nous permettre de définir différents schémas de *sense-making* opéré par les participants. La description de chaque *verbing* employé se trouve en annexe 7. Comme nous l'avons tout juste expliqué, nous avons procédé de cette manière (voir Tableau 41) avec chaque verbatim possible, en réalisant l'analyse grâce à une grille disponible en annexe 8.

Tableau 41 Exemple d'analyse d'un verbatim (participante n°10) propre au gap bridging d'une visualisation de données

N°	Verbatim	Discontinuité	Pontage	Situation finale
4	<i>T : Et donc pour rebondir sur la couleur nous avons ceci</i> <i>O : Oh lala ça m'a fait mal aux yeux, j'ai quand même essayé de faire un effort mais les couleurs se mélangent et pour distinguer « soupe », « salade », je dois avouer que..... Non c'est pas le genre de graphiques qui me plait... Autant ça peut paraître simple dans le sens où ce ne sont que des couleurs avec un nom, mais visuellement, je n'arrivais pas à faire la distinction et ça m'a... J'ai vite passé !</i>	Douleur aux yeux, difficulté de distinguer l'information Verbing : mal voir	Fait un effort pour distinguer les couleurs, se force et se concentre Verbing : se concentrer, insister, faire un effort	N'aime pas, ne poursuit pas ses tentatives de compréhension Verbing : ne pas aimer, abandonner

L'analyse pratiquée comme telle comporte plusieurs limites. Étant donné que les entretiens n'ont tout de même pas été menés dans l'optique de la *sense-making methodology*, il est fréquent dans l'analyse que certains verbatims soient incomplets, ne nous permettant pas d'identifier la discontinuité, le pontage ainsi que la situation finale. Néanmoins, nous sommes assurée qu'ils soient suffisamment nombreux afin que l'analyse soit consistante. À propos des visualisations ornementées, il y a ainsi 169 verbatims exploitables en tout. Nous n'avons pas mené l'analyse du *gap bridging* sur tous les témoignages concernant les visualisations de type standard.

En effet, nous avons davantage orienté l'analyse vers les visualisations ornementées. Très rapidement, après avoir pratiqué l'analyse sur les visualisations standard également, nous avons conclu que les « schémas » de *gap bridging* étaient presque toujours similaires. Il était très fréquent que, d'une part, nous ne détectons pas de discontinuité et que, d'autre part, les discontinuités détectées et leurs comblements soient semblables. Le plus grand constat de l'analyse des visualisations de type standard est que, bien souvent, le manque de couleur et l'aspect gris et austère engendrent une discontinuité due au visuel qui n'est pas attirant. Le pontage se fait généralement par une concentration et un intérêt pour le sujet (ou un désintérêt, par ailleurs) avant d'aboutir à une situation finale tendant vers la compréhension du sujet, mais sans pour autant que le participant ne l'apprécie (en général, la situation finale est « aimer » ou « ne pas aimer » en fonction du sujet et des affinités de

la personne). Par ailleurs, il nous a semblé évident de nous concentrer davantage sur les visualisations ornementées en ce qui concerne cette analyse puisque ce sont les discontinuités causées par l'ornementation que nous étudions.

Avant de présenter les résultats de l'analyse, nous tenons également à préciser que nous nous sommes concentrée sur les discontinuités « négatives », qui représentaient un problème à résoudre pour les participants. Effectivement, selon Dervin, un *gap* n'est pas forcément négatif ; ce n'est pas seulement un problème. La surprise ou l'étonnement peuvent être des discontinuités positives. Diverses études ont montré les bienfaits de l'ornementation, comme la mémorisation ou l'engagement. De même, notre précédente analyse thématique a largement montré les points positifs de l'ornementation. Nous n'avons donc pas comptabilisé ces discontinuités positives et nous nous sommes concentrée sur celles qui, uniquement, représentaient un problème de compréhension et d'appropriation pour l'individu. Cela nous a permis d'étudier quelles stratégies étaient mises en place par les participants lorsque l'ornementation rendait la compréhension difficile. Cela a aussi permis d'évaluer si ces stratégies ont mené vers des situations finales qui étaient des réussites ou des échecs d'appropriation des visualisations (par exemple, en *verbing* de situation finale : comprendre vs. abandonner). Ainsi, plusieurs visualisations ornementées du corpus n'ont pas généré de discontinuité (et donc de problème de compréhension ou d'appropriation) auprès de plusieurs participants.

2. Résultats de l'analyse et interprétation

Sur l'ensemble des *verbings* générés, nous avons identifié 169 schémas de *gap bridging*. Nous appelons « schéma » l'enchaînement logique d'une discontinuité, d'un pontage et de la situation finale, comme l'illustre la Figure 45. Le schéma montre donc en un verbe ou deux comment la personne a comblé la discontinuité afin d'arriver à la situation finale.

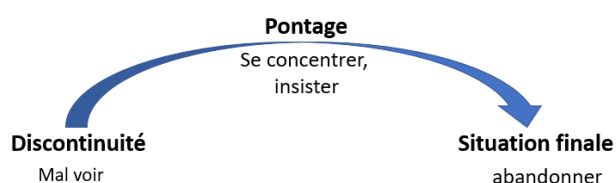


Figure 45 Exemple de schéma de *gap bridging* caractérisé par les *verbings* (participante 10, visualisation 4).

Au sein des 169 schémas réalisés, nous avons identifié neuf modèles, à savoir des schémas qui se répètent très fréquemment au sein des différentes expériences vécues par les participants. La particularité de ces schémas est qu'ils ne sont pas forcément propres à un

type de visualisation. Les schémas décrivent les actions des participants, indépendamment des caractéristiques des visualisations de données. Avant de présenter les différents schémas « modèles » que nous avons identifiés, il nous semble pertinent d'aborder sommairement les *verbings* que nous avons détectés dans l'ensemble de l'analyse, puis de montrer les similitudes de discontinuités causées par certaines caractéristiques des visualisations du corpus.

2.1. Description et répartition des *verbings* par étape de *gap bridging*

Dans cette section, nous présentons quelles actions sont menées dans le cadre de la discontinuité, du pontage et de la situation finale. Ces actions sont menées par les participants et se répartissent au sein de différents *verbings* identifiés au cours de l'analyse. Nous présenterons successivement la discontinuité, le pontage et la situation finale, dans leur ensemble.

2.1.1. La discontinuité dans l'ensemble

Pour commencer, concentrons-nous sur les différents *verbings* identifiés en tant que discontinuités. La Figure 46 et le Tableau 42 en donnent une idée de proportion. On voit ainsi que la discontinuité négative qui survient le plus souvent est « ne pas comprendre ». Effectivement, le participant ne comprenait pas du tout l'information au premier coup d'œil. S'ensuivent les *verbings* « déranger » (le participant se dit dérangé par la disposition de l'information ou le sujet) et « perturber » (le participant se dit perturbé malgré le fait qu'il comprenne). Dans une moindre mesure, les *verbings* « douter » et « identifier un problème informationnel » apparaissent. Par « douter », on entend que le participant ne peut pas affirmer que la visualisation soit compréhensible en un coup d'œil. Les problèmes informationnels détectés peuvent être des incohérences au sein même des données, une surcharge ou un manque d'information.



Figure 46 Nuage de mots représentant les différentes discontinuités rencontrées par les participants

Tableau 42 Verbinges des discontinuités

Code	Count
Ne pas comprendre	56
Déranger	39
Perturber	38
Identifier un Pb Info	20
Douter	14
Se questionner	13
Mal voir	10
Comparer	7
Décourager	5
Quantifier	4
Ennuyer	2
Heurter	2
Intriguer	2
Sortir de l'habitude	1

2.1.2. Le pontage dans l'ensemble

Les stratégies mises en place pour pallier les problèmes rencontrés sont très diversifiées (Tableau 43 et Figure 47). La relecture du visuel est très fréquente, tout comme la recherche précise d'informations dans le graphique. Les *verbings* identifiés ensuite de manière moins fréquente mais tout de même importante sont ceux-ci : « faire un effort » (le participant explique qu'il se dépasse pour comprendre la visualisation), « se concentrer » (le participant prend consciemment un temps de concentration) ou encore « faire appel à son expérience » (le participant se souvient de situations similaires ou fait appel à ses connaissances). On voit ainsi que l'action posée par les participants concerne tantôt la visualisation (relire, comparer, regarder, ...), tantôt leur propre engagement dans l'expérience de lecture (chercher, commenter, insister, ...).

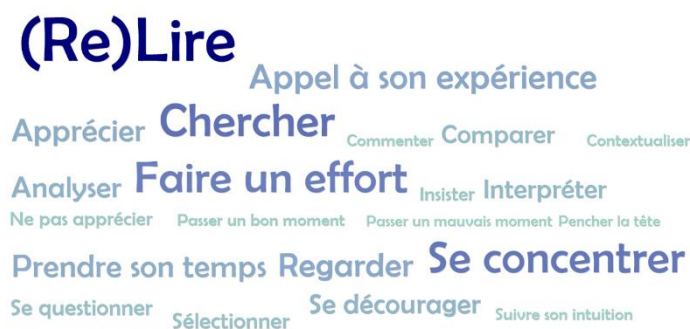


Figure 47 Nuage de mots représentant les différents pontages rencontrés par les participants

Tableau 43 Verbing des pontages

Code	Count
(Re)Lire	34
Chercher	24
Se concentrer	23
Faire un effort	21
Appel à son expérience/ ses connaissances	16
Regarder	15
Apprécier	14
Prendre son temps	13
Analyser	12
Se décourager	11
Interpréter	10
Comparer P	9
Se questionner P	7
Sélectionner	7
Insister	5
Ne pas apprécier P	5
Commenter	3
Contextualiser P	3
Passer un bon moment P	3
Suivre son intuition	3
Passer un mauvais moment	2
Pencher la tête	2

2.1.3. Les situations finales dans l'ensemble

Comme on le voit sur la Figure 48 et le Tableau 44, les situations finales les plus fréquentes sont caractérisées par le *verbing* « comprendre » et « prescrire ». Par « prescrire », nous entendons que les participants, eux-mêmes professionnels de communication, expliquent ce qu'ils auraient mis en place s'ils étaient à la place du concepteur de la visualisation observée afin que le message soit transmis de façon plus efficace. Les critiques pleuvent également, négatives en général. Dans une moindre mesure viennent les *verbings* « ne pas comprendre », « critiquer », « conclure » ou encore « apprécier ». Ces *verbings* sont à la fois positifs et négatifs, ce qui n'empêche qu'ils soient utilisés ensemble (par exemple : « comprendre »



Figure 48 Nuage de mots représentant les différentes situations finales atteintes par les participants

Tableau 44 Verbins des situations finales

Code	Count
Comprendre	39
Prescrire	34
Critiquer	24
Apprécier SF	21
Ne pas comprendre	21
Conclure	20
Ne pas apprécier SF	19
Abandonner	16
Douter SF	16
Prendre du recul	13
Échouer	9
S'énervé	5
Ne pas faire confiance	4
Contextualiser	3
Interpréter SF	3
S'interroger SF	3
Déranger SF	2
Passer un bon moment	2
Passer un mauvais moment SF	2

2.2. Les caractéristiques des visualisations de données, cause de pontages hétérogènes

Comme nous l'avons précédemment exposé, tous les verbatims concernant les différentes visualisations de données, identifiés dans les entretiens semi-directifs ne comportaient pas forcément de discontinuité et nous n'avons pas tenu compte des discontinuités positives. De même, toutes les visualisations ne provoquent pas forcément le même nombre de schémas de *gap bridging* et donc de discontinuités négatives. Dans notre analyse, l'apparition d'un schéma est causée par la détection d'une discontinuité négative. Certaines caractéristiques des visualisations, qui sont autres que les attributs exposés précédemment, influencent l'apparition des discontinuités sans pour autant influencer le pontage mis en place par les participants. Cela est bien normal et c'est tout l'intérêt de la *sense-making methodology*, qui se concentre sur l'individu et sur les actions qu'il mène pour faire sens. Cela montre également que les individus font sens des visualisations de données en fonction de leur propre expérience, même si c'est bien la présentation visuelle qui peut occasionner des discontinuités. Ainsi, on pourrait comprendre l'intention minimaliste de Bertin et Tufte dans la conception des visualisations de données : en étant le plus simple possible visuellement, on tenterait donc de réduire au maximum l'apparition de

discontinuités. De même, l'analyse de la construction de sens par le biais de la *sense-making methodology* montre aussi qu'il est fortuit d'imaginer que tous les individus décodent les visualisations de données de la même façon : en mettant en avant les actions commises par les individus pour établir un pontage par-delà la discontinuité, elle montre que les différents contextes et expériences des individus influencent indéniablement les actions et donc la construction de sens. La Figure 49 montre la répartition des différents schémas de *gap bridging* identifiés pour les visualisations de type ornementé.

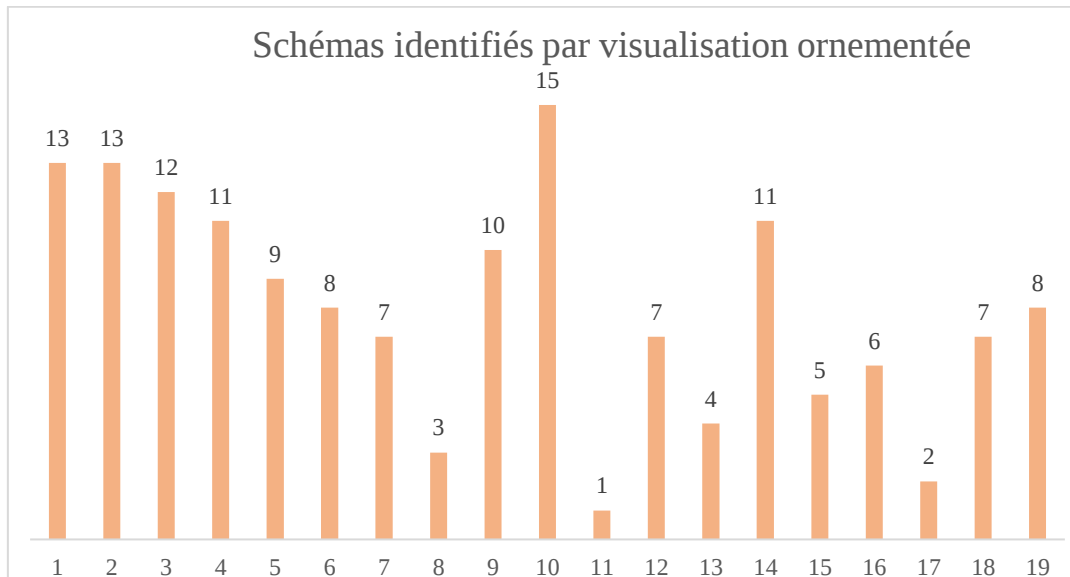


Figure 49 Nombre de schémas de *gap bridging* identifiés par visualisation de données ornementée

Rappelons que chaque visualisation ornementée a été vue par 20 personnes. Il y a donc un maximum de 20 verbatims concernant chaque visualisation ornementée, au sein desquels nous avons identifié différents schémas de *gap bridging*. On constate ainsi sur la Figure 49 que certaines visualisations ont plus souvent occasionné des discontinuités négatives que d'autres. La détection des schémas de *gap bridging* est donc un indicateur de la complexité de la visualisation. La complexité n'est pas forcément visuelle : elle peut aussi concerner la compréhension du texte ou encore la difficulté du sujet abordé. Nous avons identifié des similitudes entre les visualisations qui occasionnent un nombre similaire de schémas de *gap bridging* :

- Les visualisations n°8, 11, 13, 15 et 17 (Figure 50) occasionnent chacune entre 1 et 5 schémas de *gap bridging*, ce qui est relativement faible. Ces visualisations ont toutes un élément en commun : elles sont d'une simplicité visuelle évidente. L'ornement est en général très simple et épuré, les icônes utilisées sont très proches du pictogramme.

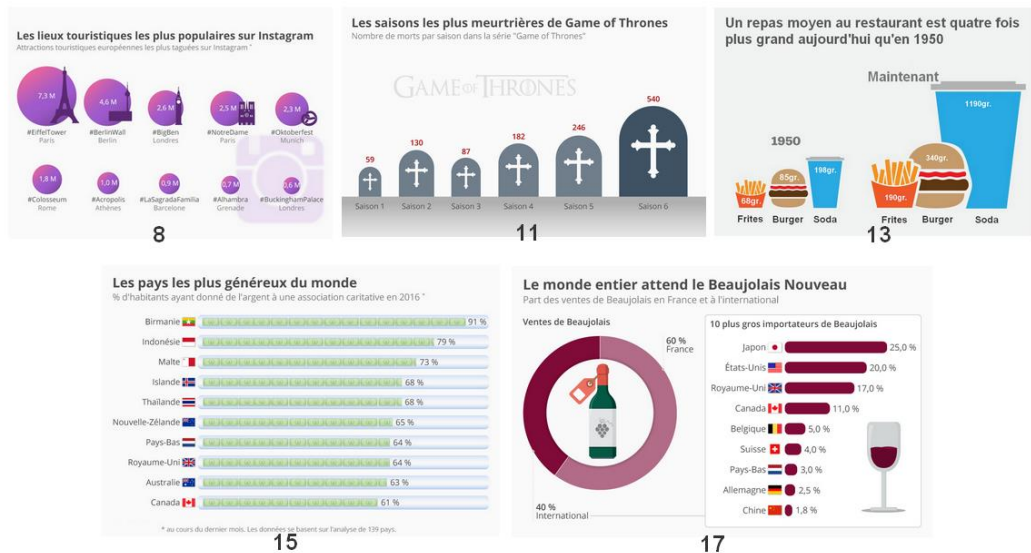


Figure 50 Miniatures des visualisations n°8, 11, 13, 15, 17

- Les visualisations n°5, 6, 7, 9, 12, 16, 18 et 19 occasionnent chacune entre 5 et 10 schémas de *gap bridging*. Leur complexité est plus grande et il apparaît que
 - o Les visualisations n°5, 7, 12, 16 et 18 ne sont pas toujours directement comprises du point de vue textuel (discontinuité : « ne pas comprendre le texte »). Sur la Figure 51, le texte qui pose un problème est encadré en rouge³⁶. Ainsi, la discontinuité est causée par une mauvaise compréhension de l'information textuelle et non de l'ornementation. Les visualisations de type standard et qui comportent le même texte causeraient donc le même type de discontinuités. Néanmoins, il arrive que les choix de design relatifs à l'ornementation altèrent le texte (exemple : visualisation n°5, absence de titre qui cause une distorsion de l'information). Il est fréquent que la situation finale relative au pontage de ces discontinuités soit négative : les participants comprennent l'intention informative en soi mais ne peuvent affirmer comprendre complètement le sujet abordé par la visualisation.

³⁶ Texte problématique : Visualisation n°5 : absence de titre // Visualisation n°7 : Chiffre d'affaire à l'export // Visualisation n°12 : « En cas d'overlaps entre la trilogie originelle et l'univers étendu » // Visualisation n°16 : « regarder des films en ligne vs. acheter l'accès de films en ligne » // Visualisation n°18 : « titulaire d'une chaire ».

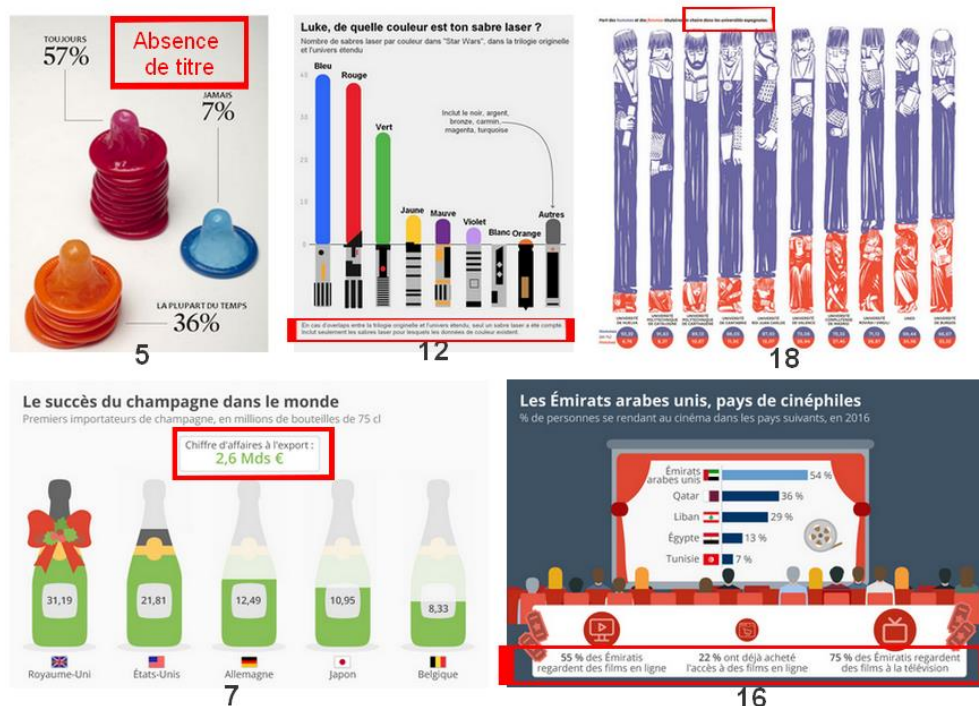


Figure 51 Miniatures des visualisations n°5, 12, 18, 7 et 16.

- Les visualisations n°18 et 19 (Figure 52) proposent un visuel et/ ou un sujet qui peut heurter la sensibilité des participants. Nous sommes ici dans la situation inverse : c'est le visuel choisi, et donc clairement l'ornementation, qui oriente les pontages réalisés dans le cadre des discontinuités rencontrés par les participants. La visualisation n°18 présente un déséquilibre entre les hommes et les femmes. Les femmes semblent « écrasées » par les hommes dans ce visuel. Bon nombre de participants s'est dit troublé par cette présentation, tantôt estimée comme « exagérée », « orientée » ou encore « une réalité parlante ». De fait, la discontinuité causée est plutôt d'ordre affectif et les pontages réalisés par les participants dépendent de leur expérience (le fait d'être préoccupé par la thématique de l'égalité des genres ou pas, par exemple). Une situation similaire se produit avec la visualisation n°19, qui présente sous forme dessinée des têtes d'animaux dont la coulée de sang représente les données. Causant de nouveau une discontinuité d'origine plutôt affective, les pontages divergent de nouveau en fonction de l'expérience et même du niveau d'humour de chacun : certains estiment cela drôle, d'autres qualifient ce visuel de « glauque ». La situation finale en est influencée, étant positive lorsque la visualisation séduit grâce à son originalité et négative lorsque justement, son style audacieux déplaît et ce, même si, dans tous les cas

l'information est comprise. Ces types de schémas de *gap bridging* nous montrent bel et bien qu'au-delà de la compréhension se joue une construction de sens relative à l'habitus de l'individu.



Figure 52 Miniatures des visualisations 18 et 19

- Les visualisations n°1, 2, 3, 4, 10 et 14 occasionnent chacune entre 10 et 15 schémas de *gap bridging*, ce qui est assez élevé en regard des verbatims disponibles. Différents problèmes, bien relatifs à l'ornementation, ont été soulevés :
 - Les formes inhabituelles (Figure 53), que nous avons précédemment considéré comme un attribut des visualisations du corpus, sont les éléments qui causent le plus de discontinuités. Elles provoquent l'étonnement, la nouveauté et ne permettent pas une présentation claire de l'information : la discontinuité consiste tout simplement à une incompréhension au premier coup d'œil et elle est rarement comblée. La situation finale reste la plupart du temps une situation d'incompréhension. Cela conforte les résultats qui ont précédemment émergé à propos des visualisations qui comportent l'attribut « forme inhabituelles » : estimées d'une faible beauté et d'une faible clarté, les appréciations déclarées ont montré qu'elles ne sont pas aimées des participants, ce qui se vérifie de nouveau ici. Notons un grand consensus en ce qui concerne les schémas de *gap bridging* de la visualisation n°4 : le visuel proposé occasionne un dérangement optique et gratte l'œil du lecteur, ce qui constitue une discontinuité quasi-systématique. Le pontage le plus souvent réalisé est d'insister ou de tenter de compter les points sur la visualisation. Il se joint

généralement d'agacement avant de se muter en abandon en situation finale.

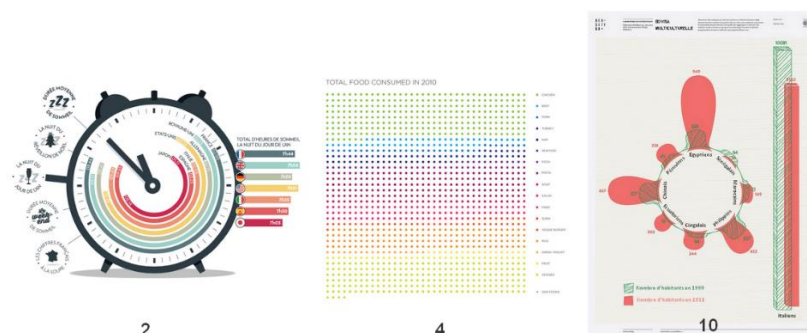


Figure 53 Miniatures des visualisations n°2, 4, 10

De plus, au-delà de tous ces éléments, nous ajoutons qu'une discontinuité est récurrente pour les visualisations n°1, 3, 6, 9 et 14 (Figure 54). En effet, c'est une surcharge visuelle qui cause une discontinuité : les participants doutent dès lors de ce qu'ils lisent. Ainsi, ils ont identifié des éléments qui chargent le visuel et en diminuent sa puissance informative. Sur la visualisation n°6 par exemple, l'ajout de lignes supposées aider à la lecture des nombres alourdit l'information. Il en va de même pour l'utilisation abusive de texte et de couleurs sur la visualisation n°9 ou sur la visualisation n°14. Nous sommes ici devant une des rares occasions où le pontage est similaire entre les participants face à cette discontinuité. En général, ils relisent et commentent l'information en pontage, pour en arriver en situation finale à produire une recommandation. Celle-ci est quasi-systématique : alléger le visuel. De manière générale, ce schéma de *gap bridging* est récurrent lorsqu'on se trouve en situation de surcharge visuelle.



Figure 54 Miniatures des visualisations n°1, 3, 6, 9 et 14

Puis, nous avons souhaité nous pencher sur les attributs des visualisations de données afin de voir si, au-delà des discontinuités occasionnées, ils génèrent une certaine homogénéité dans les pontages réalisés par les participants. Mais de manière générale, les attributs des données influencent peu les différents schémas de *gap bridging*, à l'exception de l'attribut « ornement », pour lequel les formes habituelles ont de nouveau montré leur inefficacité en termes de construction de sens. Nous avons aussi observé que les visualisations ornementées dont le *lie factor* est fidèle génèrent très peu de schémas dont la situation finale est un échec de compréhension. Les situations finales restant en situation d'incompréhension sont plus fréquentes lorsque le *lie factor* est distorsif, ce qui nuance les résultats de l'analyse statistique. Au-delà de cela, les autres attributs ne semblent pas influencer les schémas de *gap bridging*, même si d'autres caractéristiques similaires entre les visualisations de données du corpus, comme présentées ci-dessus, peuvent avoir cet effet.

En effet, nous venons de montrer dans cette section que certaines caractéristiques communes des visualisations de données du corpus – et pas forcément les attributs – peuvent occasionner des discontinuités similaires chez les participants. Sauf quelques exceptions, nous constatons que malgré tout, les schémas de *gap bridging* occasionnés ne sont pas identiques d'un participant à l'autre : si les discontinuités se ressemblent, les pontages sont quant à eux très hétérogènes. Cela est porteur de sens : ce sont les gens qui, en agissant, construisent le sens. En ayant tous des expériences, vécus et habitudes différentes, ils ne réalisent pas les pontages de la même façon. C'est d'ailleurs tout l'intérêt d'une analyse comme celle-ci : si, sommairement, nous avons indiqué quels étaient les éléments de similitude, la grille d'analyse complète (disponible en annexe 8) montre la grande variété des pontages et des situations finales pour les schémas existants.

À propos des exceptions, avant d'aborder les schémas modèles qui peuvent s'appliquer à toute visualisation du corpus ou presque, nous avons identifié plusieurs schémas qui s'appliquent à une visualisation précise ou à un cas précis. Ainsi :

- À propos de la visualisation n°4, le schéma de *gap bridging* est récurrent et quasi identique pour les participants qui ont trouvé la visualisation problématique. En effet, sur la Figure 55, on constate que l'individu, d'un point de vue visuel, perçoit très mal l'information qui est encodée sur la visualisation. La visualisation provoque donc un dérangement optique. Malgré une relecture et malgré le fait que les participants, en général, insistent afin de comprendre l'information en se forçant à la lire, la situation finale la plus fréquente est l'abandon. Les participants expliquent aussi pourquoi ils

abandonnent la lecture de la visualisation en formulant certaines critiques sur les éléments les ayant poussés à abandonner l'appropriation de l'information.

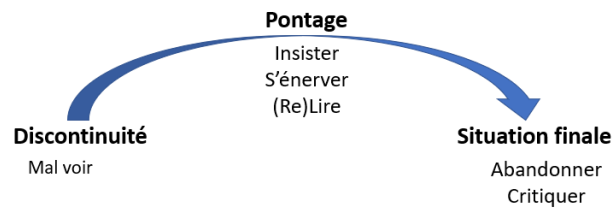


Figure 55 Schéma de gap bridging relatif à la visualisation n°4

- À propos de la surcharge visuelle, comme indiqué ci-dessus, les participants qui qualifient une visualisation de « chargée visuellement » en viennent souvent à réaliser un schéma de *gap bridging* similaire. En général, ce qui leur permet de faire sens est le commentaire qu'ils sont capables d'élaborer au moment de la relecture de la visualisation, ce qui constitue un pontage important. Ils font sens de la visualisation en imaginant comment elle pourrait être allégée d'un point de vue visuel. C'est pourquoi la situation finale se solde par une recommandation (« prescrire »).

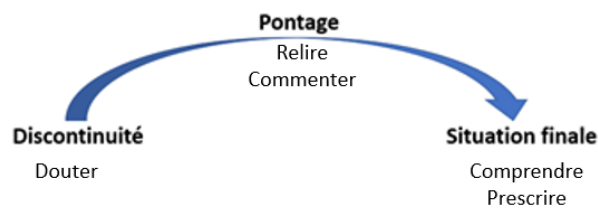


Figure 56 Schéma de gap bridging relatif aux visualisations chargées visuellement

- À propos de la dimension affective de la discontinuité, elle influence le pontage soit positivement, soit négativement, comme nous avons pu l'illustrer précédemment avec les visualisations n°18 et 19. Ainsi, il s'avère que le sujet de la visualisation couplé à un choix de design orienté puisse influencer sur le pontage réalisé par l'utilisateur qui, en fonction de sa personnalité, de ses habitudes et expériences, va aboutir à une situation finale positive ou négative. En définitive, ce genre de pontages, qui consistent simplement à se laisser toucher affectivement – de façon positive ou pas – survient logiquement lorsque la discontinuité est elle-même de type affectif. Ainsi, en choisissant d'élaborer une communication de type « choc » ou engagée, il arrive de perdre le lecteur ou, au contraire, de le pousser à s'engager. Cela dépend donc bel et bien du lecteur, qui peut réaliser le pontage de façon positive ou négative, et non pas du caractère interpellant de la visualisation. Ainsi, comme on le voit sur la Figure 57, la discontinuité est de l'ordre du trouble, de l'interpellation, et peut déboucher sur une situation positive (faire écho aux expériences du

lecteur, passer un moment agréable) qui aboutit à une situation finale positive elle aussi, ou sur une situation négative où, en passant un moment désagréable, le lecteur n'apprécie pas l'information et ce, même si, dans les deux cas de figure, les données en soi sont comprises par l'utilisateur.

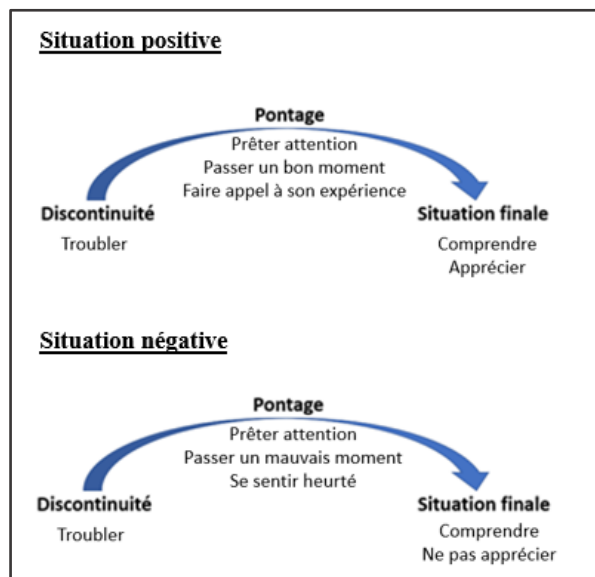


Figure 57 Schéma de gap bridging relatif à une discontinuité de type affectif

2.3. Les schémas modèles

Au sein de l'ensemble des schémas encodés dans le cadre de cette analyse (au nombre de 169), nous en avons identifié neuf que nous qualifions de « modèles ». Ces schémas modèles ont donc tendance à se répéter au cours de l'analyse et ce, peu importe la visualisation concernée ainsi que ces caractéristiques. Évidemment, certaines visualisations tendront davantage vers l'émergence d'un schéma modèle plutôt qu'un autre, comme nous l'avons précédemment montré, mais de manière générale, ces modèles sont susceptibles de s'appliquer à n'importe quelle visualisation du corpus, en fonction de la personne et de son expérience.

Ainsi, pour une visualisation, plusieurs schémas sont possibles, en fonction de la personne qui la lit. Nous avons recensé ceux qui apparaissent le plus souvent. Pour ce faire, un travail textuel a été réalisé : dans chaque transcription d'entretien, nous avons sélectionné les verbatims relatifs à chaque visualisation. Nous étions donc par exemple en mesure de lire tous les verbatims relatifs à la visualisation ornementée n°1, puis n°2 et ainsi de suite. Grâce à notre grille d'analyse (annexe 8), nous avons pu identifier les *verbings* relatifs à la discontinuité, au pontage et à la situation finale, par personne et par visualisation. Ensuite, par visualisation, nous avons sélectionné les schémas de *gap bridging* les plus fréquents,

avant de les comparer avec les schémas relatifs aux autres visualisations afin de sélectionner les 9 modèles que nous présentons désormais. Comme nous l'avons précisé auparavant, les discontinuités que nous avons retenues sont celles qui sont indicatrices d'un problème pour le participant. Les discontinuités se ressemblent donc souvent. Par contre, il y a deux issues possibles dans les situations finales identifiées : soit elles sont positives (le participant comprend, apprécie, ...), soit elles sont négatives (il ne comprend pas, il échoue, ...). Nous allons donc classer les neuf modèles en fonction du succès ou non du participant lors de la situation finale. Rappelons que la situation finale était exprimée par l'individu lui-même dans les verbatims avant que nous n'en ayons fait un *verbing*. Classer les schémas modèles en fonction de la situation finale permet de se concentrer sur le pontage. Rappelons néanmoins que chaque situation est propre à l'individu qui la vit selon Dervin (1998). En réalité, il y a eu autant de situations différentes que de discontinuités identifiées. Néanmoins, nous avons tenté de rassembler les situations qui se ressemblent en les caractérisant d'un ou plusieurs *verbing(s)* à chaque fois sans pour autant réexpliquer chaque contexte. Notre objectif est de détecter les stratégies menant à une situation finale positive ou négative de manière générale.

2.3.1. Les situations finales positives

La Figure 58 illustre le premier schéma modèle identifié. Il était fréquent que les participants se questionnent à propos de la visualisation qu'ils venaient de lire.

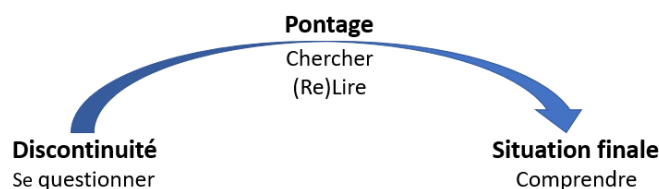


Figure 58 Schéma modèle n°1 de gap bridging à situation finale positive

Le questionnement faisant l'objet de la discontinuité avait deux sources :

- La présentation de l'information sur le graphique, qui pouvait ne pas être claire :
« On peut se demander pourquoi c'est une fois un homme, une fois une femme. Ce n'est pas explicité mais on comprend que peu importe »
 (visualisation n°6).

Dans ce verbatim, la participante interroge le fait que les ornements utilisés (des petits bonshommes) soient des hommes et des femmes, alors que le sexe n'est pas une variable représentée par les données.

- Une interrogation à propos du sujet présenté par la visualisation de données
« J'ai un *apriori* négatif en disant que si les italiens ont diminué et les étrangers ont augmenté, c'est quoi ce truc, ce propos, est-ce que c'est pour soutenir une lutte anti-émigration ? » (visualisation n°10)

Pour les personnes se trouvant dans ces situations, le pontage s'est effectué de deux façon :

- En cherchant l'information qui n'était pas saillante au premier coup d'œil sur le graphique ;
- En relisant tout simplement la visualisation

Ces deux actions ont généralement permis aux participants de trouver les éléments de réponse à leur questionnement initial, menant vers la compréhension en situation finale.

La Figure 59 amène un peu plus de nuance au fait de se questionner en discontinuité. Effectivement, ce questionnement peut être motivé par le doute. L'hésitation fait douter le participant qui se pose alors des questions :

« Ce sont des prix fixes et je n'arrive pas à dire si le 492 c'est Johannesburg ou Dar Es Salaam et c'est ce qui fait que je me suis posé des questions et qui a perturbé ma compréhension » (visualisation n°1)

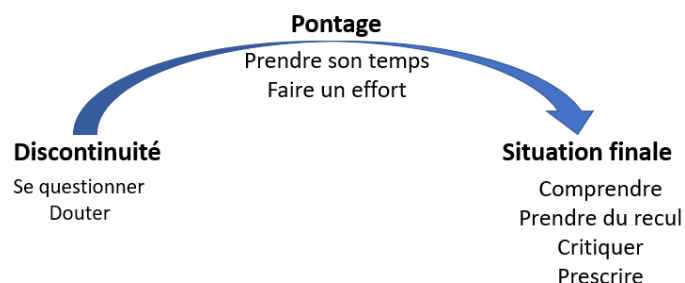


Figure 59 Schéma modèle n°2 de gap bridging à situation finale positive

En guise de pontage, ces personnes ont pris leur temps et ont décrit l'action de « faire un effort ». Par ces deux actions, les participants montrent qu'ils dépassent une simple lecture intuitive. Ils font le choix de s'arrêter et de regarder la visualisation plus en profondeur, ce qu'ils décrivent comme un effort :

« J'ai quand même pris le temps donc au final j'ai compris, en prenant mon temps... Si tu prends vraiment le temps tu arrives à comprendre mais si tu passes ça dans un livre vite fait tu comprends rien quoi » (visualisation n°7).

« C'est très joli à voir mais il faut faire un effort pour comprendre l'information qu'il y a dans la visualisation » (visualisation n°2).

Néanmoins, que l'issue soit positive ou négative, la situation finale s'accompagne souvent de critiques ou d'une prise de recul de la part des participants lorsque le pontage à réaliser demande un effort qui dépasse la simple relecture. En critiquant, les participants décrivent ce qui leur a posé problème et déplu dans la visualisation. Lorsqu'ils prescrivent, ils expliquent comment ils auraient mieux compris la visualisation si, dans l'idéal, elle avait été présentée comme ils la décrivent :

« Ça aurait été mieux de tout mettre sur une même ligne, ça aurait encore été plus clair (...). Ça aurait été sympa plus en longueur » (visualisation n°8)

La Figure 60 montre une nouvelle situation de discontinuité : le participant se décrit comme perturbé. La perturbation peut être causée par l'organisation de l'information en elle-même ou par l'ornement présent sur la représentation. Il se met également à douter ; il ne comprend pas l'information au premier coup d'œil. Dès lors, les stratégies de pontage menant vers une compréhension sont plus nombreuses.

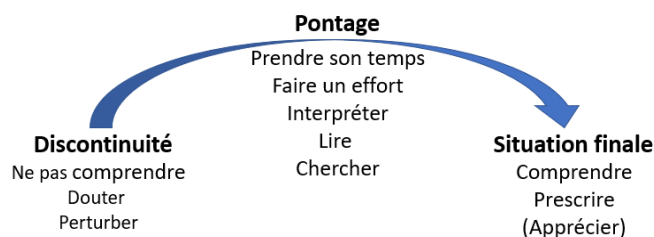


Figure 60 Schéma modèle n°3 de gap bridging à situation finale positive

Comme précédemment, le participant prend son temps et fait un effort. Il relit également l'information et cherche visuellement les éléments lui permettant d'interpréter l'information. Au-delà du simple décodage d'information, le participant se met dans la peau d'un interprète. On constate de nouveau que les participants expliquent comment ils auraient aimé voir l'information présentée.

Sur la Figure 61, les participants se disent perturbés et dérangés par des éléments visuels présents dans la représentation graphique. Néanmoins, ils ne sont pas dans une situation d'incompréhension, mais les éléments visuels occasionnant un dérangement les amènent à douter de l'information.

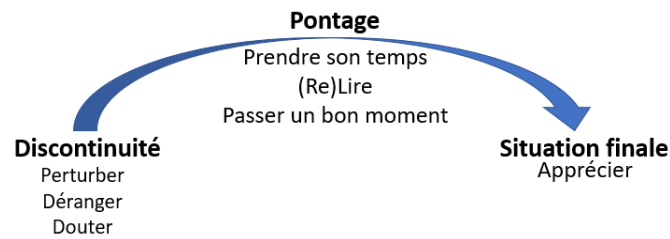


Figure 61 Schéma modèle n°4 de gap bridging à situation finale positive

Les pontages opérés sont ceux-ci : prendre le temps de saisir l'information et la relire. Néanmoins, le visuel est plaisant et les participants en viennent à passer un bon moment :

Elle m'a marqué, il y a des trucs partout et quand on retire les informations qui sont répétées plusieurs fois, et que à ce moment-là, on comprend que c'est à partir du moment auquel ils se couchent et se lèvent, c'est un comparateur fun quoi, c'est ludique (visualisation n°2).

Ainsi, malgré la difficulté occasionnée par l'ornementation, les participants en viennent à apprécier la visualisation plutôt que de simplement la comprendre.

Si on avait juste mis 18-30 ans et rien d'autre j'aurais lu et compris tout de suite mais je pense que les photos sont importantes parce qu'on voit plus ou moins la tranche d'âge dans laquelle la personne se situe, j'ai bien aimé celle-ci (visualisation n°9).

Finalement, le dernier schéma menant à une situation finale positive montre que les participants sont parfois confrontés à la discontinuité parce qu'ils détectent un problème informationnel dans la visualisation, comme une mauvaise couleur employée selon eux ou, tout simplement, un « mauvais » design :

Je me suis dit qu'il manquait des données... ça pourrait être une belle image de pub mais il manque des infos, un message derrière... (visualisation n°5).

2.3.2. Les situations finales négatives

Sur le schéma de la Figure 62, on constate que des éléments visuels peuvent parfois déranger et perturber les participants. Il leur arrive également de douter de l'information qu'ils décodent.

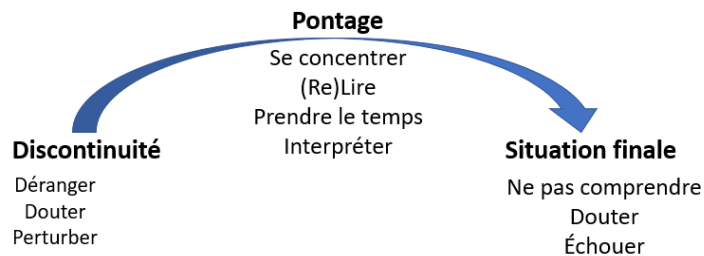


Figure 62 Schéma modèle n°5 de gap bridging à situation finale négative

Les pontages réalisés pour pallier ces problèmes sont relatifs à la personne, qui souhaite dès lors se concentrer sur la visualisation, l'interpréter en prenant le temps de contempler les détails afin d'augmenter les chances de compréhension. Néanmoins, il arrive que ces stratégies n'aboutissent pas : le visuel n'a pas toujours pu aider les participants malgré les pontages mis en place. Ainsi, les participants ont eux-mêmes indiqué comme situation finale qu'ils n'ont pas compris l'information, qu'ils échouaient dans l'appropriation de la visualisation de données et que, surtout, le doute apparu en discontinuité n'était pas levé en situation finale.

J'émet toujours des doutes donc s'il n'y a pas quelque chose dans ce visuel qui va me convaincre il n'y a pas de doute à avoir, je vais me sentir perdue. Je ne vais pas arrêter de me demander si je suis sûre de ce que j'ai compris (visualisation n°5).

Sur la Figure 63, le schéma montre qu'une discontinuité présentant l'incompréhension peut rester non résolue malgré les pontages mis en place comme le fait de faire un effort, de relire l'information, ou encore de chercher les détails. Il arrive parfois également que le participant fasse appel à sa propre expérience ou à ses souvenirs afin d'arriver à comprendre certains choix de conception.

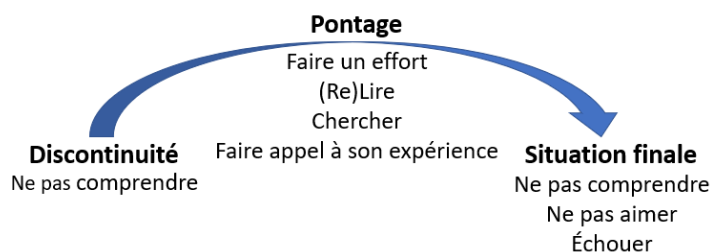


Figure 63 Schéma modèle n°6 de gap bridging à situation finale négative

Cela se vérifie effectivement dans le verbatim suivant.

Je sais pas si je suis fatiguée mais sur le coup j'ai trouvé ça tellement improbable parce que j'ai déjà vu Star Wars (visualisation n°12)

Le fait que la visualisation n'est pas en cohésion avec l'expérience ou les souvenirs du participant ne mène pas seulement au fait qu'il échoue dans la compréhension de la visualisation : en plus de cela, le participant n'apprécie pas la visualisation en situation finale.

Le schéma qui apparaît sur la Figure 64 montre que des éléments visuels perturbent les participants et les amènent à douter de l'information. Certains participants dans cette situation décidaient de prendre le temps de lire l'information correctement mais ressentaient un malaise à ce moment-là.

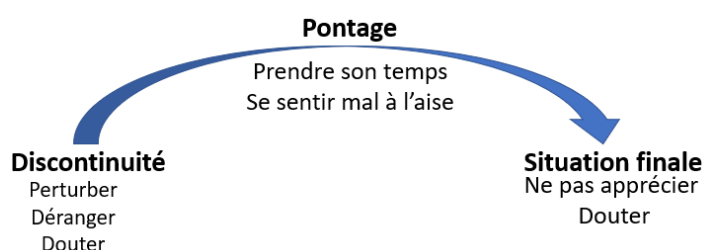


Figure 64 Schéma modèle n°7 de gap bridging à situation finale négative

La situation finale qui en est issue, suite au malaise, est le fait de pas apprécier la visualisation de données. De plus, les doutes émis lors de la discontinuité ne sont toujours pas comblés.

Et alors bêtement, ce côté qui a été fait pour choquer... ça donne pas envie de s'y attarder. L'image a tellement une connotation négative que... non... ou j'ai un trop petit cœur j'en sais rien (visualisation n°19)

Le schéma n°10 (Figure 65) montre une situation dans laquelle le dérangement constituant une continuité débouche sur la critique et sur le fait de ne pas aimer une visualisation malgré le fait que les participants vivant ce schéma comprennent l'information d'un point de vue purement fonctionnel.

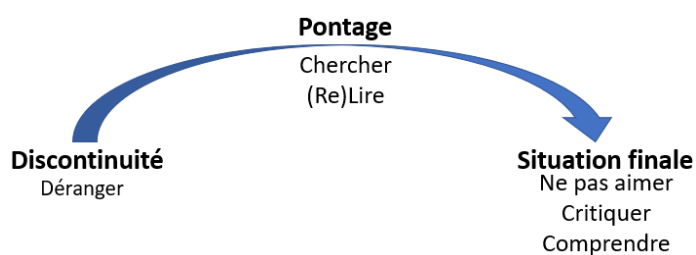


Figure 65 Schéma modèle n°8 de gap bridging à situation finale négative

Parfois, les participants, qui ne comprennent pas l'information en discontinuité, relisent la visualisation mais n'apprécient pas cette expérience de lecture. Le fait de ne pas passer un bon moment lors de la lecture n'encourage pas les participants à poursuivre et ils abandonnent alors la lecture de la visualisation.

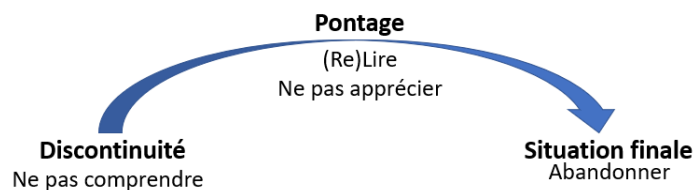


Figure 66 Schéma modèle n°9 de gap bridging à situation finale négative

Le fait de ne pas apprécier la lecture de la visualisation peut aussi se manifester par des sentiments négatifs comme l'ennui ou l'agacement.

Je trouvais que c'était du chipotage dans tous les sens... j'ai lu vite fait et j'ai passé...

3. Apport et pertinence de l'analyse des gap bridgings de visualisations ornementées

Dans cette analyse, nous ne nous sommes pas penchée sur les discontinuités rencontrées par les lecteurs de visualisations standard. Les raisons en sont simples :

- Les discontinuités sont en réalité très peu fréquentes ;
- Les discontinuités qui apparaissent ont généralement trait à l'information plutôt qu'au visuel (mauvaise compréhension du titre ou de l'unité de mesure) ou bien concernent un jugement de goût (exemple : ne pas apprécier les barres grises du diagramme en bâtonnets) ;
- Le constat est le même que dans l'analyse statistique : les visualisations sont claires et bien comprises (peu de discontinuités) mais de manière générale, elles ne sont pas appréciées d'un point de vue esthétique (discontinuités apparentes). Dans cette situation, en guise de pontage, le participant se concentre en général sur le sujet de la visualisation qui, souvent comprise en situation finale, est appréciée ou pas en fonction du sujet.

Force est de constater que les visualisations de type ornementé provoquent plus de discontinuités négatives³⁷ que leurs homologues au type standard. Néanmoins, les

³⁷ Pour rappel, nous n'avons pris en compte que les discontinuités négatives, qui constituent un problème. Une discontinuité peut être positive (exemple : surprise, émoi artistique, ...). L'objectif

participants prennent souvent du recul ou sont capables de critiquer le visuel leur ayant été présenté (*verbings* « prendre du recul » et « critiquer » présents 30 fois en situation finale). Nous constatons également que le pontage le plus souvent mis en place, peu importe la discontinuité rencontrée, est de relire la visualisation de données et d’y chercher les éléments apportant des réponses. Cela peut être un succès ou un échec. Les différents schémas de *gap bridging* montrent également l’importance des facteurs humains et émotionnels de Kennedy et Hill (2017, 2018) dans l’appropriation d’une visualisation de données : il est fréquent que l’appel aux expériences et aux souvenirs survienne dans le pontage et qu’en situation finale, cet habitus ait contribué à la compréhension du participant et à son affection pour la visualisation. Une autre situation finale survenant souvent est que le participant soit en désaccord avec la visualisation rencontrée : son expérience ou ses connaissances sont en conflit avec ce que montre la visualisation de données. De même, malgré une compréhension évidente, il n’apprécie pas la visualisation ou conserve des doutes quant à l’information présentée.

C’est donc ainsi que le sens se crée : par de multiples actions, débouchant sur le succès ou l’échec de la compréhension ou de l’appropriation. L’aspect affectif, appelé plutôt engagement dans la discipline *Infovis*, reste important et contribue largement à la construction de sens. Évidemment, loin de nous l’idée de confondre *sens* et *efficacité*. Malgré tout, la façon dont les individus font sens d’une visualisation de données serait à tenir en compte dans une nouvelle définition de l’efficacité (Bertin, 1967) ou de l’excellence graphique (Tufte, 1983). Les différentes discontinuités engendrent souvent un temps d’appropriation de la visualisation qui est plus long. Bien qu’il existe des échecs, le sens s’en trouve plus riche en situation finale lorsqu’il s’agit de succès : l’ornementation provoque davantage d’appréciations positives et stimule la critique positive ou négative. Notons également qu’en tant que professionnels de la communication numérique, les participants ont donné beaucoup de recommandations ou prescriptions en situation finale, en expliquant comment ils auraient préféré voir apparaître les données dans la visualisation.

Pour terminer, les visualisations ornementées ne provoquent pas automatiquement des discontinuités négatives. Comme nous l’avons montré dans la section 2.2., certaines caractéristiques des visualisations, comme le fait d’être très chargées visuellement ou bien d’être caractérisée par l’attribut « forme inhabituelle » engendrent davantage de discontinuités négatives, tandis que les visualisations ornementées mais d’apparence simple en génèrent assez peu. De toute façon, ce qui influence le pontage est contenu dans la personnalité, le vécu, l’expérience et les habitudes du participant qui, à travers tout cela,

était de comprendre les stratégies en place lorsque l’ornementation des visualisations de données pose problème.

agit pour faire sens. C'est ainsi qu'en caractérisant ses actions par des *verbings*, nous sommes à même de comprendre quelles stratégies de construction de sens permettent aux individus d'aboutir à une situation finale positive ou négative. Nous avons tenté de rassembler les schémas similaires et de les interpréter dans l'ensemble. L'exercice est difficile et périlleux : la *sense-making methodology* ne vise pas la généralisation des résultats et c'est tout à fait normal, puisque chaque individu vit ses propres discontinuités et pontages en fonction de son expérience personnelle. Dans les conditions de situations initiales similaires qui ont été mises en place dans le cadre des expérimentations en laboratoire, la mise en commun d'éléments probants a heureusement été possible, mais aurait pu être plus rigoureuse encore puisque les entretiens n'ont pas été menés conformément à la *sense-making methodology* de Dervin et ses collègues (Dervin, 1992 ; Dervin et al., 2003). Malgré tout, en complément de notre précédente analyse thématique, les présents résultats permettent d'en apprendre sur la dimension agissante de la construction de sens en se focalisant sur l'individu. Toujours dans la volonté de se concentrer sur l'individu et sur ce qui anime sa construction de sens, nous abordons le design émotionnel selon Norman (2013) dans le chapitre suivant. Nous y constaterons d'ailleurs un point de convergence avec le concept de *gap bridging*.

Chapitre 3 Ornementer ou prôner le design émotionnel (Norman, 2013)

Les deux analyses précédentes ont montré à quel point les caractéristiques des individus mais également leurs propres actions au moment de la réception d'une visualisation pouvaient influencer leur construction de sens. De même, nos données qualitatives montrent que les participants sont capables de commenter *l'utilisabilité* des visualisations de données et surtout de les critiquer, de se positionner par rapport à elles, comme nous a montré l'analyse selon la *sense-making methodology* (Dervin et al., 2003). Ces faits laissent entendre que le design émotionnel de Norman (2013), qui se compose de différentes étapes de la consommation d'un objet physique, pourrait également s'appliquer à la visualisation de données, en laissant ainsi sa chance à l'ornementation.

Ce sujet est traité plus en profondeur dans un article paru en 2020 dans la revue « *Communiquer. Revue de communication sociale et publique* »³⁸. Inspiré de cet article, ce

³⁸ Andry T. (2020), Visualisation de données et design émotionnel peuvent-ils se conjuguer ? *Communiquer. Revue de communication sociale et publique*, 28, p.53-71.

chapitre présente sommairement les raisons pour lesquelles design émotionnel et visualisation de données peuvent être associés. Nous proposons également un parallèle avec la notion de *gap bridging*.

1. Rappel théorique

Depuis plusieurs années maintenant, l'approche théorique du design émotionnel, conceptualisée par Norman (2013), pousse au débat. En effet, le psychologue cognitiviste émet un principe simple : la simplicité n'est pas toujours la bonne solution en termes de design (Norman, 2008). À ce stade, Norman parle de design d'objets de « la vie de tous les jours » et non de visualisations de données. Nous nous sommes alors questionnée : et si le design émotionnel s'appliquait également à la dataviz ?

Le design émotionnel est une clé intéressante en ce qui concerne la compréhension du poids des émotions dans la consommation d'un objet qui, par principe, se devrait d'être conçu dans la plus grande simplicité. Pour rappel, le mot « design » va bien plus loin qu'une idée de modernité, alliant alors conception, compréhensibilité, contrôle de l'utilisateur et plaisir (Norman, 2008) notamment lié à l'esthétique. En termes de design, pour Norman, les choses attrayantes fonctionnent mieux. « La raison en est qu'elles génèrent un plus grand plaisir chez l'utilisateur qui, dès lors, s'applique davantage dans l'utilisation de l'objet. La cognition et les émotions vont de pair ; les affects ont dès lors un rôle prépondérant dans l'activité vécue (Norman, 2013) » (Andry, 2020). Pour lui, les émotions sont « un système de jugement attribuant des valeurs aux objets et aux événements » (Cahour *dans* Norman, 2013). Les émotions sont issues du sens que chacun crée par rapport à son propre habitus jusqu'à ce qu'une action corporelle soit déclenchée (exemple : sudation en cas de stress).

« Norman estime que le design parfait, qui mêle attraction, compréhension, utilisabilité et plaisir d'utilisation, est un alliage de trois niveaux de design différents interconnectés entre eux et dans lesquels les émotions interviennent avec différents degrés d'importance. Il y a donc :

- Le niveau viscéral, au cours duquel l'individu répond à la vision immédiate avant que ne se mette en route le système réflexif. Ce niveau est préconscient et les apparences l'emportent ;
- Le niveau comportemental qui est le niveau des usages et de l'expérience qui est faite. L'utilisateur évalue les fonctions, la performance et la facilité d'utilisation d'un objet ;
- Le niveau réflexif, qui est le seul niveau de la conscience éveillée de l'utilisateur à propos de l'objet. C'est le niveau de l'interprétation, de la compréhension et du raisonnement. Les pensées et émotions produisent des effets qui y prennent toute leur

ampleur. L'utilisateur intellectualise l'objet et l'intègre en fonction de caractéristiques qui lui sont propres et dont dépendent ses émotions : sa culture, ses expériences, allant jusqu'au sens identitaire. »

(Andry, 2020)

Les deux premiers niveaux de design ne produisent pas d'interprétation consciente, mais des conflits émotionnels peuvent exister entre les différents niveaux. Imaginons une visualisation de données très ornementée et complexe à déchiffrer. Au niveau viscéral, elle pourrait sembler très belle à un individu qui se sentirait dès lors attiré. Son émotion pourrait être positive tout en étant négative au niveau comportemental car il peine à la comprendre malgré son attrait esthétique. Au niveau réflexif, il pourrait se mettre lui-même en perspective par rapport à cette expérience de lecture et conclure « je n'aime pas les chiffres et je ne suis pas motivé à faire un effort ».

2. Pertinence de ce cadre conceptuel

Nous avons de fortes raisons de penser que les trois niveaux de design selon Norman (2013) ont été stimulés dans le laboratoire lors de nos expérimentations. Le niveau réflexif est le seul niveau où l'individu replace l'objet et la perception de lui-même dans le temps. Par notre action de mener un entretien semi-directif, nous l'avons stimulé de manière peu naturelle. Par contre, nous avons peu de marge de manœuvre en ce qui concerne l'immédiateté de la réception des niveaux viscéral et comportemental. Il reste néanmoins possible de revenir sur ces deux premiers niveaux de design, grâce aux commentaires réalisés par les participants : en utilisant des mots et en exprimant leurs sensations, ils ont pu reformuler leurs pensées premières et les communiquer, ce qui a constitué un matériel intéressant à exploiter.

Ainsi, nous pouvons affirmer que les entretiens semi-directifs couvrent beaucoup d'aspects propres aux trois niveaux de design selon Norman (2013).

De même, il est possible d'établir un parallèle entre le *gap bridging* et les différents niveaux de design émotionnel. En effet, les *verbings* identifiés dans le cadre de l'analyse précédente, s'inspirant de la *sense-making methodology* de Dervin et ses collègues (2003) nous permet de l'illustrer. Par exemple, certaines discontinuités sont relatives au niveau viscéral, par exemple lorsqu'un élément visuel dérange le participant mais qu'il ne sait pas expliquer pourquoi, ou tout simplement lorsqu'il sait directement qu'il n'aime pas les bâtonnets gris des visualisations de type standard (*verbings* : déranger, perturber, etc.). Dans le pontage et dans la situation finale, de nombreux exemples montrent que les participants jugent de l'utilisabilité des visualisations de données (niveau comportemental) lorsqu'ils décrivent

par exemple pourquoi une visualisation ne « fonctionne » pas bien à leurs yeux et qu'ils expliquent les actions qu'ils ont mis en place pour qu'elle « fonctionne » malgré tout. Pour terminer, lorsqu'ils situent la visualisation en fonction de leur propre expérience en situation finale, c'est le niveau de design réflexif qui est stimulé. Les niveaux de design comportemental et réflexif sont donc fortement stimulés, à la fois lorsque le participant crée le pontage, mais également lorsqu'il aboutit à la situation finale. Le Tableau 45 nous permet d'illustrer ce parallèle par un exemple. Il concerne le verbatim d'un participant, à propos de la visualisation n°19. On y voit que les trois niveaux de design ne sont pas satisfaits, notamment parce que le visuel fait appel à une expérience professionnelle du participant, allant à l'encontre de ce qui est proposé dans la visualisation. On constate que la situation de *gap* correspond au niveau viscéral : l'emploi de la couleur rouge et l'allusion au sang répulse le participant dans un premier temps. Au niveau comportemental, il est capable d'expliquer que la visualisation est intéressante tout en posant des mots, en conceptualisant ce qui le dérange dans ce visuel. Cela correspond à un pontage : il comble la discontinuité en passant au-dessus de ce qui le dérange pour comprendre l'information. Finalement, au niveau réflexif, on comprend que le participant se situe lui-même par rapport à la visualisation : il n'apprécie pas voir la couleur rouge dans les représentations graphiques car il l'assimile à un ensemble de connotations négatives, de par sa pratique professionnelle.

Tableau 45 Parallélisme entre le design émotionnel et une situation de gap bridging, à propos de la visualisation n°19

Verbatim	Gap et niveau viscéral	Pontage et niveau comportemental	SF et niveau réflexif
<i>A : Ce qui me dérangeait le plus c'était l'effet goutte de sang, après le dessin ça va et il est super clair. C'est intéressant de se dire que des chèvres ont été décapitées à New York (rires) mais c'est pas un visuel que je présenterais chez un client.</i> <i>T : Et tu l'apprécies ?</i> <i>A : ça va, juste le sang. Dans mon agence média on évite le rouge parce que ça fait toujours penser à l'école, l'époque où on se faisait corriger, que c'était pas bon... C'est bête mais on a tous pris l'habitude. Quand il y a trop de rouge ça me gêne.</i>	Verbing : déranger, être mal à l'aise Viscéral : répulsion de la couleur rouge, effet sang	Verbing : apprécier, intéresser, commenter Comportemental : est intéressée par le sujet mais détecte précisément la raison du dérangement	Verbing : comprendre, déranger Réflexif : ressent de la distance par rapport à ce visuel car évite l'utilisation de la couleur rouge dans sa pratique professionnelle

Dans cet exemple, nous constatons donc que les trois niveaux de design ne sont pas correctement stimulés pour que l'information soit communiquée au participant en accord avec son expérience et sa personnalité, dépassant un simple aspect de fonctionnalité. L'effort est là, mais c'est un échec. On le constate aussi grâce aux *verbings* mis en avant par la méthode de Dervin et ses collègues (2003).

À l'inverse, nous pensons qu'en stimulant positivement les émotions des destinataires, il serait possible, en imaginant que les principes de conception des visualisations de données soient correctement respectés, de conjuguer design émotionnel et visualisation de données pour tendre vers la communication de données numériques dépassant le simple aspect fonctionnel. Les bénéfices de l'ornementation, dès lors, pourraient être exploités afin de stimuler les sensations positives identifiées lors de l'analyse thématique (sensation agréable, amusement, attractivité, ...), tendant donc vers des émotions positives prédisposant l'individu à s'engager auprès de la visualisation. Cette méthode du design émotionnel pourrait également permettre de prévenir l'échec d'utilisabilité et de compréhension, comme dans notre précédent exemple. L'exemple employé dans la section suivante illustre une situation positive.

3. Vers la visualisation de données idéale ?

« Pour Norman, les différents niveaux de design sont interconnectés entre eux et suscitent des émotions différentes à chaque niveau. Un état émotionnel positif aide, prédispose l'individu à trouver comment s'approprier un objet et à contourner sa complexité. Un état émotionnel négatif rend cela difficile (Norman, 2013) » (Andry, 2020). En visualisation de données, la représentation « idéale » combinerait donc les trois niveaux de design. Dans nos expérimentations, la visualisation ornementée ayant récolté le plus d'appréciations positives, en ce qui concerne les trois niveaux de design, concerne le *Beaujolais Nouveau* (visualisation n°17, Figure 67).

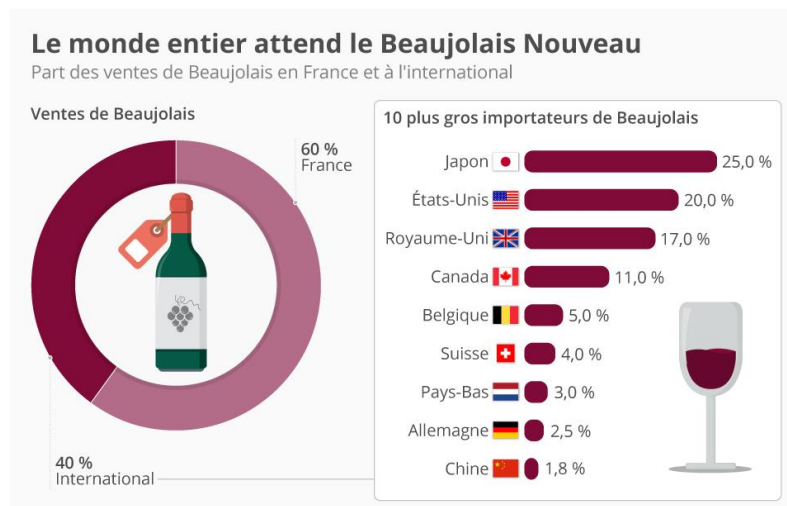


Figure 67 Élément du corpus : visualisation ornementée et adaptée à l'expérimentation (n°17).

Au niveau viscéral, son esthétique est appréciable, les pictogrammes sont très appréciés par les participants, tout comme la couleur, qui est présente pertinemment, en rappelant la lie de vin, mais sans surcharge. En soi, les représentations graphiques utilisées sont simples et claires. « Au niveau comportemental, les pictogrammes sont une allusion réelle au sujet : le vin. La couleur bordeaux ou lie de vin fait tout autant allusion au sujet. La simplicité de l'effort visuel est très appréciée. Au niveau réflexif, la visualisation est décrite comme claire, compréhensible, belle et efficace. Certains participants sont friands de vin, d'autres moins. Ces caractéristiques propres à soi dépendent d'un individu à l'autre mais, de manière générale, la combinaison parfaite entre esthétique et efficacité a été soulignée. Les trois niveaux de design sont atteints et génèrent en général une émotion positive » (Andry, 2020). Le Tableau 46 montre dans quelle mesure les différents niveaux de design se retrouvent dans les verbatims.

Tableau 46 Expression orale des niveaux de design émotionnel pour la visualisation ornementée n°17

Niveau viscéral	<i>C'est super clair, il y a le titre mais aussi la bouteille de beaujolais et le verre de vin ! C'est joli à regarder.</i>
Niveau comportemental	<i>Les icônes sont bien et elle est vite compréhensible. Toutes les informations sont cohérentes et pertinentes et les couleurs rappellent le vin et c'est bien.</i>
Niveau réflexif	<i>Je la trouve bien faite. On décompose en détail avec les 60 % de parts de vente qui sont en France et dans les 40 %, on voit en détaillé tous les importateurs. Je trouvais ça sympa et je me demandais où on serait nous, la Belgique. En fait, dans toutes les visualisations où il y avait des pays, je me demandais où en était la Belgique, mon pays !</i>

Les participants semblent ainsi conquis par l'alliage de la sensation de beauté, de la simplicité et des représentations graphiques plutôt traditionnelles.

4. La simplicité, encore et toujours

Le design émotionnel est d'un grand intérêt : même s'il montre l'importance de l'apparence de l'objet et de la première impression qu'il laisse à l'utilisateur (beau, laid), il ne laisse pas de côté son utilisabilité et ses capacités à provoquer une remise en question chez le destinataire. Cela montre effectivement que toute ornementation n'est pas bonne à utiliser et, en montrant l'exemple de la visualisation n°17 ayant recueilli le plus de commentaires positifs aux trois niveaux de design, force est de constater que de nouveau, c'est une ornementation simple et épurée qui est à recommander.

Cela tient d'ailleurs tout son sens : l'ornementation rendrait la visualisation de données plus agréable et donc plus plaisante à lire, ce qui en améliorerait l'efficacité si l'on se rappelle de la maxime de Norman selon laquelle les choses attrayantes fonctionnent mieux. En réalité, ces choses fonctionnent mieux non pas parce qu'elles sont mieux conçues, mais parce qu'elles prédisposent l'individu à les utiliser, en stimulant positivement leurs émotions. Le plaisir d'utilisation est pleinement pris en compte. Mais pour atteindre les trois niveaux de design et donc la parfaite forme de design émotionnel, il faut que la visualisation de données soit irréprochable au niveau comportemental, après avoir satisfait le niveau viscéral en termes d'apparence visuelle. En d'autres termes, il faut qu'elle « fonctionne » bien. De nombreux principes de conception existent déjà en la matière. Celui que nous avons le plus exploité dans le cadre de ce travail est la maximisation du *data-ink ratio* jusqu'à l'aboutissement d'un visuel minimaliste et épuré. Malgré l'aversion de Tufte pour le *chartjunk* (1983), nous proposons non pas d'opposer ces principes de conception au concept d'ornementation, mais plutôt de les allier. Ainsi, l'ornementation satisfaisant les niveaux viscéral et comportemental serait celle qui n'entraverait pas l'utilisabilité de la visualisation, tout en étant elle-même simple et épurée (puisque, rappelons-le, la surcharge visuelle est un danger de l'ornementation qui endommage la compréhension de l'information). Le niveau réflexif serait sans doute le plus compliqué à satisfaire. Malgré tout, une visualisation de données est un message conçu à destination d'un public cible. La transmission de la visualisation au bon public cible nécessitera *de facto* un intérêt ou un recul critique plus grand de la part d'un public sensibilisé au message qu'on lui communique.

Conclusion L'analyse qualitative, la finesse des détails

L'analyse qualitative, menée en différentes étapes (analyse thématique suivie de la méthode de *verbing*, le tout mis en lien avec la théorie du design émotionnel), a soulevé de nombreux détails et subtilités permettant de commenter la manière dont les individus font sens des visualisations de données standard et ornementées.

Ainsi, les participants parlent d'une sensation d'effort. En général, l'ornementation amoindrit la sensation d'effort, à certaines conditions. De même, ornementation ne rime pas forcément avec modernité puisque, réalisée à partir de tendances, elle peut parfois sembler passée de mode. Les participants sont également capables de décrire différents facteurs d'incompréhension dans les visualisations, qu'elles soient ornementées ou non. Les forces et faiblesses de l'ornementation ont pu être décrites du point de vue des participants. L'analyse thématique a montré l'importance des émotions des individus, de la clarté et de la simplicité du visuel. Les facteurs humains et émotionnels de Kennedy et Hill (2017, 2018) peuvent facilement être repérés dans les entretiens et montrent bien que le concepteur de la visualisation ne contrôle pas du tout la réception de la visualisation. Il y a donc bien une part d'interprétation personnelle, contribuant à l'émergence du *sense-making*.

Nous nous sommes ensuite concentrée sur le comblement des discontinuités (*gap bridging* selon Dervin (Dervin, 1992 ; Dervin et al., 2003) en se focalisait sur les actions des participants et pas sur les visualisations en soi. Uniquement menée sur les verbatims concernant les visualisations ornementées, il s'avère que les attributs n'interviennent pas systématiquement dans la façon dont les individus effectuent le pontage, car plusieurs schémas modèles existent malgré la différence d'attributs et ce, même si certaines caractéristiques des visualisations de données favorisent parfois des discontinuités similaires entre les participants. On constate ainsi différents pontages menant à des réussites et des échecs de la construction de sens. Lorsque le pontage se solde en un échec, il est fréquent que les individus abandonnent, ne comprennent pas ou ressentent un sentiment négatif.

De plus, nous avons exposé que le design émotionnel s'applique à la visualisation de données. Les trois niveaux de design à savoir les niveaux viscéral, comportemental et réflexifs, une fois combinés, peuvent aboutir à une visualisation de données dont l'expérience d'utilisation est un succès. Nous avons constaté qu'une ornementation simple, mobilisant des icônes apposées à un graphique traditionnel, correspond aux trois types de design. Évidemment cela est une piste : on pourrait se demander, lors de la conception

d'une visualisation, en quoi la visualisation finale rencontre les différents niveaux de design émotionnel. En revanche, nous nuancions les propos de Norman qui dit que la simplicité n'est pas une réponse en termes de design. Si l'ornementation s'éloigne de la fameuse simplicité dont parle Norman, nous lui donnons raison. Néanmoins, lorsque l'ornementation apparaît, c'est bien la simplification de cette dernière qui permet de combiner les différents niveaux de design. Ainsi, ornementation et simplicité finissent par s'accorder une fois de plus.

Aussi, les participants ont expliqué que la clarté des visualisations impacte le fait qu'ils aiment ou non la visualisation. Or, dans l'analyse statistique, nous avons identifié une corrélation très faible entre les *items* « Clair » et « J'aime ». Nous supposons dès lors que la prise de recul potentiellement réalisée par les participants entre le moment de l'expérimentation et l'entretien semi-directif a pu influencer leurs propos.

Par contre, les participants ont précisé que sortir des conventions visuelles habituelles était difficile pour eux et diminuait la qualité de leur expérience de lecture. Cela conforte l'analyse quantitative, dans laquelle nous avons constaté que les formes inhabituelles recueillent très peu d'appréciations positives.

D'autres éléments, encore non traités, sont apparus dans l'analyse, comme la disposition de la légende ou encore l'absence de l'axe y qui peut se révéler perturbant pour les participants. À travers l'analyse des *gaze data*, qui suit en [Partie V](#), nous élaborerons différents commentaires à propos des éléments soulevés conjointement dans l'analyse quantitative et qualitative. Nous pourrions ainsi identifier si ce que déclarent oralement les participants dans les entretiens se vérifie réellement d'un point de vue spontané, à travers leur comportement physiologique capturé dans les données de regard.

Partie V. Fixations et patterns de lecture

Analyse des *gaze data*

Introduction

Les deux analyses précédentes ont permis de tirer diverses conclusions quant à la réception de l'ornementation des visualisations de données. La plus importante d'entre toutes concerne la simplicité de ces visualisations ornementées. A priori, à ce stade, nous pouvons avancer qu'une pratique d'ornementation en particulier peut être tout à fait bénéfique et ce, sans perturber aucunement la réception de l'information et sa compréhension. Cette ornementation consiste en la disposition d'icônes simples, « à côté » d'un graphique traditionnel.

Après avoir analysé les appréciations déclarées des participants ainsi que leur discours, nous sommes désormais en mesure de procéder à la dernière analyse des données récoltées, à savoir les données *eye tracking*, aussi appelées « *gaze data* ». Grâce à l'avancement de notre recherche, cette analyse pourra être orientée en fonction des conclusions nouvelles, apparues au cours de ce travail.

Les données de regard sont d'une grande richesse : elles nous permettent de déterminer, pour chaque individu, les endroits de la visualisation sur lesquels son regard s'est arrêté, avec une extrême précision. Après avoir étudié comment les individus font personnellement sens des visualisations de données, nous avons désormais accès à des données représentatives d'un phénomène physiologique et spontané : la vision.

Avant tout, nous présenterons dans cette cinquième partie les données obtenues ainsi que quelques notions élémentaires en *eye tracking*. Ceci permettra de saisir leur importance. Puis, nous proposerons une série d'hypothèses, elles-mêmes basées sur nos précédentes analyses. Pour y répondre, nous traiterons les *gaze data* en menant une analyse statistique sur différentes métriques concernant les fixations (voir *infra*) enregistrées au cours de l'expérimentation. Ensuite, nous étudierons les spécificités des chemins de lecture communs aux participants. Pour cela, nous utiliserons une technique de visualisation des *gaze data* inédite. Ces différentes visualisations permettent de repérer les différents *patterns* de lecture en observant ainsi le rôle de l'ornementation dans le comportement visuel des lecteurs.

Cette analyse nous permettra ainsi de conclure en renforçant la recommandation qui s’est dessinée au cours de notre travail en ce qui concerne l’ornementation des visualisations de données.

Chapitre 1 Les *gaze data* récoltées, mises en perspective en *Infovis* et en *eye tracking*

À l’issue de l’expérimentation, nous avons récupéré les données de regard des participants, aussi appelées *gaze data*, fournies par le logiciel Tobii Studio. Pour commencer, nous expliquons comment se présentent ces données. Ensuite, nous proposons une focale sur la littérature existant à propos de la visualisation d’information et l’*eye tracking*, permettant ainsi d’entrevoir quelles pistes d’analyse s’ouvrent dans le cadre de notre recherche.

1. Les données récoltées et le matériel exploitable

Les expérimentations menées en laboratoire nous ont permis de récolter les *gaze data* de quarante participants, pour trente-huit visualisations de données³⁹. Pour chaque participant et pour chaque visualisation, nous détenons une série de mesures récolées en millisecondes.

Le mouvement des yeux des participants est très informatif à propos de la façon dont ils interagissent avec différents stimuli visuels. L’hypothèse « œil-cognition » de Just et Carpenter (1976) sous-entend que les éléments entrant dans le champ visuel d’un observateur sont directement soumis à un traitement cognitif. Effectivement, l’étude des points de fixation oculaire (voir *infra*) prend ses racines en sciences cognitives, dès 1898 (Just & Carpenter, 1976). L’évolution des connaissances et de la technologie ayant suivi son cours, il est possible aujourd’hui de comprendre ce que le mouvement des yeux révèle à propos du contenu de visualisations d’information (Bylinskii et al., 2017). Différentes métriques existent dès lors et permettent d’interpréter les différentes données issues des expérimentations oculométriques.

Selon l’étude des mouvements de l’œil, la focalisation de l’œil sur un objet est appelée **fixation**. Ainsi, lorsque l’œil cherche un objet et que le regard s’arrête, une fixation se produit, durant généralement entre 200 et 300 ms en fonction de la tâche à effectuer (par exemple : points de fixation de 225ms pour une lecture silencieuse, 275 ms pour une recherche visuelle, etc.). Les mouvements très rapides de l’œil, réalisés entre les fixations s’appellent les **saccades**. Durant le mouvement d’une saccade, il n’y a pas de nouvelle

³⁹ Pour rappel, le corpus est composé de quarante visualisations, soustrait de deux visualisations erronées. Nous prenons donc en compte dix-neuf visualisations par participant, et non vingt.

obtention d'information, tant les yeux bougent rapidement entre les différents stimuli visuels. L'œil perçoit alors un flou et rien de plus ; c'est ce que l'on appelle « la suppression saccadique ». La direction des saccades est influencée par différents points d'intérêts (Rayner, 1998).

Le logiciel Tobii Studio fournit ces différentes métriques, ainsi que d'autres variantes, en secondes et en millisecondes. Il produit aussi différentes visualisations de ces métriques, permettant de se rendre compte visuellement des mouvements oculaires des participants. Il propose ainsi

- Le gaze plot : un enchaînement de nœuds (points de fixation) et de liens (saccades), numérotés, qui permettent de suivre dans quel ordre et où sont apparus les points de fixation. La Figure 68 en est un exemple.

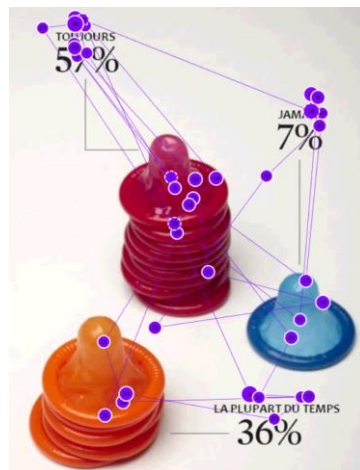


Figure 68 Exemple d'un gaze plot pour un participant, visualisation n°5 du corpus

- La carte de chaleur : montre grâce à son intensité la durée des fixations (vert : moins intense – rouge : plus intense) de manière localisée, comme sur la Figure 69. Elle peut également montrer le nombre de fixations, au besoin.

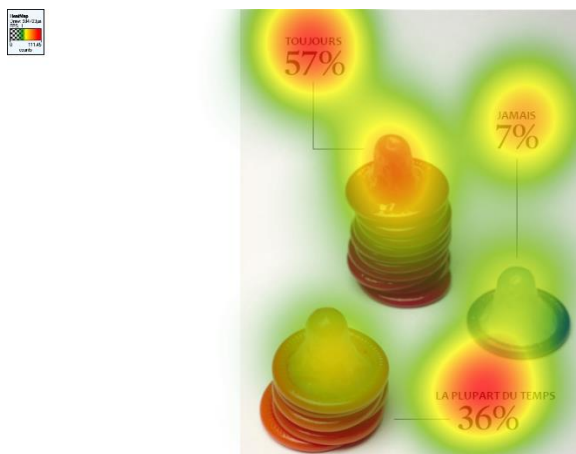


Figure 69 Exemple de carte de chaleur pour la visualisation n°5 du corpus

- Le tout, animé : la représentation animée de ces visualisations permet de comprendre de façon intuitive et temporelle les mouvements oculaires des participants.

Tous ces outils sont disponibles en agrégeant le résultat pour tous les utilisateurs. Malheureusement, dès que les résultats des vingt utilisateurs pour une image sont affichés, l'information devient peu perceptible. Nous avons eu la chance de participer au développement d'un format de visualisation supplémentaire, permettant de montrer les chemins de lecture agrégés pour tous les participants. Il s'agit du *bundling*, largement expliqué dans le chapitre 3 de cette partie. Grâce à cette technique, nous pourrions déterminer si l'ornementation a un impact sur les chemins de lecture des participants.

2. La compréhension de graphiques et l'*eye tracking*

L'utilisation de l'*eye tracking* dans les études portant sur la visualisation de données est très fréquente. L'originalité de notre recherche tient au fait qu'elle se déroule dans la discipline des sciences humaines et que son aspect qualitatif soit très développé. Par ailleurs, les études précédemment menées par d'autres chercheurs à propos de la compréhension de graphes et de l'interprétation des métriques oculométriques nous sont tout à fait bénéfiques.

En effet, Carpenter et Shah ont développé un modèle des processus perceptuels et conceptuels se déroulant lors de la compréhension de graphiques (1998). Pour ce faire, elles ont mené une étude oculométrique au cours de laquelle elles ont prêté attention aux éléments sur lesquels les individus se concentraient visuellement au cours d'une tâche, tout en tentant de déceler certains *patterns*⁴⁰ de lecture. Ainsi, selon Carpenter et Shah, la compréhension de graphiques est un processus incrémental qui suit un modèle intégratif (*integrative model*) en impliquant la reconnaissance de formes, leur interprétation quantitative et qualitative ainsi que leur intégration dans l'esprit de l'utilisateur. Ce modèle dépasse le simple *pattern recognition model* selon lequel l'individu reconnaît des formes et y fait correspondre un ensemble de relations conceptuelles. Le modèle intégratif contient une réelle part interprétative qui se répercute dans les trajets oculaires : si la complexité du graphique augmente, le cycle des différents processus s'intensifierait en fonction des informations supplémentaires à encoder et à mettre en relation. Il est donc normal que l'œil s'arrête plusieurs fois sur les mêmes éléments (fixations et re-fixations), en suivant ces

⁴⁰ Dans la suite de ce chapitre, nous utiliserons les mots « chemin » ou « *pattern* » de lecture pour désigner la même chose, à savoir des modèles de trajets oculaires.

différents cycles de décodage, d'interprétation et d'intégration. C'est pourquoi Carpenter et Shah (1998) montrent que, par exemple, le temps passé par les individus à décoder les différents axes et étiquettes du graphique excède le temps passé à regarder tout simplement la forme du graphique. En 2011, Körner réalise une étude qui renforce les propos de Carpenter et Shah (1998). Il parle effectivement d'un *3-stage model* (modèle à trois étapes) lors de la lecture de graphes complexes. Au cours de cette lecture, l'individu passe à travers deux étapes de recherche puis réalise une étape de raisonnement. Au cours des deux premières étapes, la recherche visuelle est presque identique : les différents « nœuds » cibles, permettant la compréhension⁴¹ sont localisés dans le graphe. Dans la troisième étape, les relations entre ces nœuds sont visuellement établies, menant à l'interprétation et à la compréhension. Selon Körner, si, durant la compréhension, les deux premières étapes servent à la recherche d'information, les *patterns* de lecture, à savoir les différentes zones où se situent les points de fixation, ne devraient pas différer d'un participant à l'autre (Körner, 2011). Dans la troisième étape du modèle, il y a davantage de mouvements oculaires sortant de ces zones, additionnels, permettant l'interprétation. Ainsi, au cours de notre propre analyse, nous tenterons de montrer qu'il existe effectivement différents *patterns* de lecture similaires entre les quarante personnes ayant participé aux expérimentations dans le cadre de notre recherche. Si ces *patterns* existent, ils peuvent nous donner l'opportunité de comparer ce qui diffère entre les *patterns* de visualisations standard et ornementées. Le chapitre 3 y est entièrement consacré.

De plus, comme nous l'avons précédemment mentionné, les métriques oculométriques permettent également d'interpréter certaines informations à propos des visualisations de données. L'*eye tracking* peut être utilisé pour expliquer la performance d'une tâche réalisée par une personne. Dans le cadre de nos expérimentations, la seule tâche demandée à l'individu était de lire et de tenter de comprendre les visualisations qui leur ont été présentées. Le mouvement des yeux est une fenêtre sur la compréhension de processus cognitifs prenant place au moment où l'individu tente de réaliser cette tâche (Bylinskii et al., 2017). Bylinskii et ses collègues ont produit une revue de littérature pertinente à propos des différentes métriques à utiliser lors de la réalisation d'analyses *collectives*, à savoir, l'analyse de plusieurs visualisations, visualisées par plusieurs individus. Ces chercheurs proposent ainsi des interprétations pour différentes métriques, à utiliser dans le cas d'évaluation à grande échelle. Dans leur propre exemple, les chercheurs proposent un corpus de 393 visualisations. Même si l'ampleur de notre étude est plus restreinte, l'interprétation de ces métriques semble justifiée dans notre cas en particulier. Ces

⁴¹ Ces nœuds correspondent en fait à des zones dans le graphique où se situent les informations indispensables à la compréhension

métriques servent à évaluer l'efficacité de certains designs de visualisations de données, ce qui correspond à nos objectifs. De même, analyser l'importance de certaines zones d'intérêt⁴² sera extrêmement pertinent dans notre cas. Nous n'analyserons par ailleurs que des zones d'intérêt relatives à l'ornementation apposée (icônes disposées à côté d'un graphique traditionnel) afin de vérifier définitivement si oui ou non, il s'agit d'une distraction pour l'utilisateur. Pour le reste, nous observerons certaines métriques afin de commenter les différences entre visualisations ornementées et standard ou même entre différents attributs des visualisations ornementées en ce qui concerne l'efficacité, l'engagement ou même la complexité de ces dernières. Pour ce faire, le nombre de fixations et leur durée, par visualisation, suffisent amplement à produire ces interprétations cohérentes. En nous appuyant sur la revue de Bylinskii et ses collègues (2017), nous menons dans le chapitre suivant une analyse statistique basée sur ces métriques. Étant donné que les saccades ne sont pas porteuses d'information, si ce n'est les points d'intérêt attirant le regard, nous ne les avons pas pris en compte dans l'analyse. De plus, la seconde analyse menée dans le chapitre 3 nous permet de nous en passer (cf. *infra*). Le travail de Poole et Ball (2006), même s'il n'est pas consacré à la visualisation d'information mais à l'efficacité de sites web en général, est également éclairant dans ce domaine : comme Bylinskii et ses pairs, ces deux chercheurs commentent les différentes interprétations possibles des mesures récoltées grâce à l'*eye tracking*.

Le Tableau 47 montre notre sélection de métriques mises en avant par Bylinskii et ses pairs, ainsi que par Poole et Ball (2006) en rassemblant leurs significations et les différentes interprétations possibles. Ces métriques sont d'ailleurs à la source des hypothèses qui se trouvent dans le chapitre 2. Les résultats de l'analyse qualitative, menée dans la [Partie IV](#) nous guidera vers l'interprétation la plus correcte des métriques choisies.

⁴² Produire une zone d'intérêt consiste à sélectionner une zone en particulier sur l'image. Ainsi, on peut étudier les métriques oculométriques sur cette zone en particulier.

Tableau 47 Sélection de métriques oculométriques, signification et interprétation. Inspiré de Bylinskii et al. (2017) et de Poole et Ball (2006)

Métriques	Interprétations possibles
Nombre total de fixations sur la visualisation	Un plus grand nombre de fixations indique une efficacité moindre OU Plus d'engagement
Durée des fixations ➔ Aussi possible : durée moyenne des fixations	Une durée de fixation plus longue indique une difficulté à extraire l'information OU Plus d'engagement
Nombre de fixations sur les zones d'intérêt	Plus de fixations sur une zone en particulier indique qu'elle est plus importante OU Plus perceptible que les autres zones sur le graphique
Temps de vision sur les zones d'intérêt	Si le temps de vision est plus grand, le contenu informationnel de cette zone est plus élevé OU Cette zone est complexe OU Elle pousse à l'engagement
<i>Time to first fixation</i> sur une zone d'intérêt	Est révélateur des propriétés de captation d'attention de cette zone.

L'analyse se déroulera en deux temps, dans les deux chapitres suivants. Premièrement, les données numériques relatives aux fixations subiront une analyse statistique afin de répondre à une série d'hypothèses formulées grâce aux résultats des Parties [III](#) et [IV](#) (Chapitre 2). Ensuite, l'exploitation d'une technique de visualisation (*bundling*) nous permettra de définir des *patterns* de lecture tout en décrivant le rôle de l'ornementation dans l'émergence de ces derniers (Chapitre 3).

Dans ce chapitre, nous formulons des hypothèses sur base des analyses précédentes et des interprétations possibles des métriques oculométriques. Nous expliquons ensuite notre méthode avant de présenter l'analyse en elle-même et les différentes interprétations qui en sont issues.

1. Formulation des hypothèses

Pour commencer, résumons brièvement quelques éléments de l'analyse des appréciations déclarées et des entretiens semi-directifs. En ce qui concerne les appréciations corrélées, les *items* « Belle » et « J'aime » sont corrélés : le fait de trouver une visualisation jolie peut être un facteur d'engagement. Les visualisations ornementées sont donc bien souvent préférées aux visualisations standard, de manière générale. Les visualisations ornementées qui présentent des icônes sont souvent appréciées, tandis que les visualisations ornementées à la forme inhabituelle sont boudées des participants. En ce qui concerne la figuration des icônes, les participants la trouvent plus belle et plus claire lorsqu'elle est apposée aux données plutôt que figurant sur les traits relatifs aux données. Les visualisations au niveau moyen de *data-ink* semblent plus claires que celles dont le niveau de *data-ink* est élevé. Puis, en ce qui concerne les conclusions des entretiens semi-directifs, elles sont très nombreuses. Nous retenons néanmoins de nombreux atouts de l'ornementation : l'expérience agréable, la sensation d'effort amoindri, l'attractivité du visuel et surtout une simplicité qui reçoit toutes les louanges des participants qui sont conscients des effets négatifs de la surcharge visuelle. L'analyse des *gap bridgings* nous a également montré la dimension agissante de la construction de sens, inévitablement influencée par l'expérience personnelle des participants, ainsi que par les différentes émotions suscitées par les visualisations de données. Cet aspect apparaît un peu moins dans la présente analyse.

Quoi qu'il en soit, à partir de ces différents éléments, nous souhaitons orienter les hypothèses en fonction de différentes thématiques : la complexité visuelle, l'utilisation d'icônes à la figuration apposée et l'éventuelle distraction causée par l'ornementation et l'impact de tous ces éléments sur les chemins de lecture. Pour répondre à ces hypothèses, nous réaliserons une analyse statistique des métriques relatives aux différents points et durées de fixation récoltés durant l'expérimentation. Les hypothèses sont classées ci-dessous en fonction d'une thématique qui peut en influencer la méthode d'analyse (cf. *infra*).

Comparaison entre le type ornementé le type standard

- H1 La durée moyenne des fixations est significativement plus importante pour les visualisations ornementées que pour les visualisations standard.
- H2 Le nombre de fixations est significativement plus grand pour les visualisations ornementées que pour les visualisations de type standard.

Comparaisons basées sur des attributs significatifs

- H3 Le nombre de fixations est significativement plus élevé sur les visualisations ornementées à la forme inhabituelle que sur les visualisations ornementées dont la figuration se constitue en icônes.
- H4 Le nombre de fixations est significativement plus élevé sur les visualisations dont la figuration se constitue sur les traits relatifs aux données (*data-ink*) que sur les visualisations dont la figuration est apposée aux données.
- H5 Au sein des visualisations ornementées, la comparaison à une situation dans laquelle le *lie factor* est fidèle, la durée moyenne des fixations est significativement plus élevée lorsque le *lie factor* est distorsif.
- H6 Au sein des visualisations ornementées,
- un niveau de *data-ink* élevé induit un temps de vision significativement plus bas qu'un niveau de *data-ink* bas ou moyen ;
 - un niveau de *data-ink* moyen induit un temps de vision significativement plus bas qu'un niveau de *data-ink* bas.

Analyse des zones d'intérêt : de la distraction causée par les icônes apposées

- H7 Les ornements apposés captent rapidement l'attention : ils sont regardés avant la lecture du graphique.
- H8 Les ornements apposés ne sont pas des éléments importants : le nombre de fixations est faible dans ces zones d'intérêt par rapport au nombre moyen de fixations.
- H9 Les ornements apposés ne sont pas d'une grande complexité visuelle : la durée moyenne des fixations dans les zones d'intérêt est relativement basse par rapport à la durée moyenne des fixations pour les visualisations dans leur ensemble.

2. Méthode d'analyse

À l'issue des expérimentations, nous avons récolté différentes données numériques. À la suite d'un traitement léger, nous les avons manipulées afin qu'elles soient présentées dans un tableur en présentant les informations comme le montre le Tableau 48. Comme on le voit sur la ligne « *Feuille 1* », le tableau affiche les métriques qui sont directement interprétables selon Bylinskii et ses collègues (2017) et Poole et Bale (1996). Nous avons également calculé le « Temps de Vision » sur base des différents *TimeStamps*⁴³ fournis par Tobii Studio dans l'extraction de données.

Tableau 48 Présentation des métriques oculométriques dans le tableur pour analyse

Feuille 1 : Fixations					
ID_Participant	ID_Visualisation	Nombre de fixations	Durée moyenne des fixations	Temps de vision	
Feuille 2 : AOIs. Métriques par visualisations					
ID_Participant	Nombre de fixations	Durée moyenne des fixations	Time to first Fixation	Visit Duration	Fixations before

Afin d'extraire et d'obtenir les données relatives aux zones d'intérêt, nous avons tracé ces zones dans le logiciel Tobii Studio. Comme le prévoient les hypothèses, nous n'avons tracé des zones d'intérêt que sur les icônes constituant une figuration apposée. La Figure 70 montre en vert et en orange un exemple de zones d'intérêt pour lesquelles les données ont pu être extraites. Une visualisation peut contenir plusieurs zones d'intérêt. La deuxième ligne du Tableau 48 (*Feuille 2*), montre les différentes métriques extraites par zone d'intérêt. En supplément, nous avons eu l'opportunité d'extraire les métriques « *Visit Duration* » et « *Fixations before* », proposées par Tobii Studio, que nous expliquerons au moment opportun.

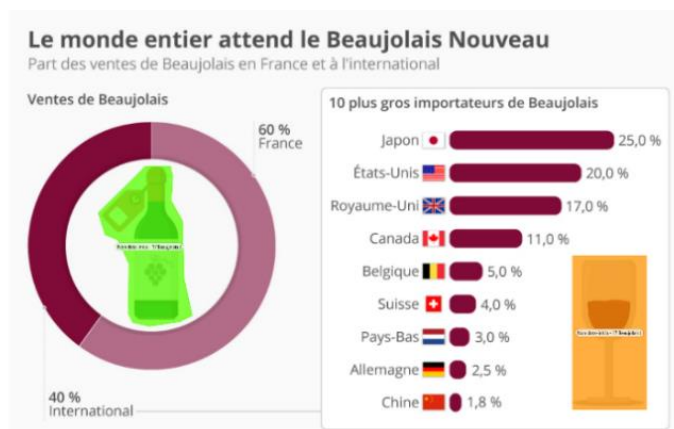


Figure 70 Visualisation des zones d'intérêt pour la visualisation n°17 du corpus

⁴³ Le timestamp désigne le nombre de secondes écoulées depuis le 1^{er} janvier 1970 à minuit UTC précise. Il s'agit d'un horodatage informatique.

Nous avons dessiné des zones d'intérêt sur dix visualisations ornementées du corpus. Effectivement, même si l'ornementation de certaines visualisations du corpus se concentre sur le *data-ink*, il arrivait qu'une icône apposée soit présente sur l'image. Nous avons pris ces cas en compte dans notre analyse. La visualisation des zones d'intérêt se trouve en annexe 10. Dix-neuf zones d'intérêt sont concernées au total.

En ce qui concerne les hypothèses relatives aux fixations dans leur globalité, nous avons mené une analyse statistique en mobilisant principalement les tests de Welch et Brown Forsythe. En revanche, nous avons répondu aux hypothèses concernant les zones d'intérêt, en commentant les données à notre disposition.

3. Analyse statistique des fixations

Différents tests statistiques ont donc été menés sur les métriques extraites du logiciel Tobii Studio. Ces tests permettent d'interpréter le comportement visuel des participants. Nous analysons en premier lieu les données en fonction du type (standard, ornementé), puis en fonction de quelques attributs des visualisations du corpus avant d'explorer les zones d'intérêt.

3.1. Comparaison des fixations en fonction du type (ornementé ou standard)

3.1.1. Analyse des données

Pour commencer, jetons simplement un œil au nombre moyen de fixations par visualisation. Un simple diagramme en bâtonnets nous permet de réaliser une première entrée en matière. Sur la Figure 71, on constate qu'il semblerait que, de manière générale, les visualisations de type ornementé (*Junk* sur la figure) comportent davantage de points de fixation que les visualisations de type standard (*Raw* sur la figure). Sur la figure, la visualisation n°21 est l'homologue standard de la version ornementée n°1. Les barres ont été disposées de manière à faciliter la comparaison entre deux homologues⁴⁴.

⁴⁴ [Voir la version verticale](#)

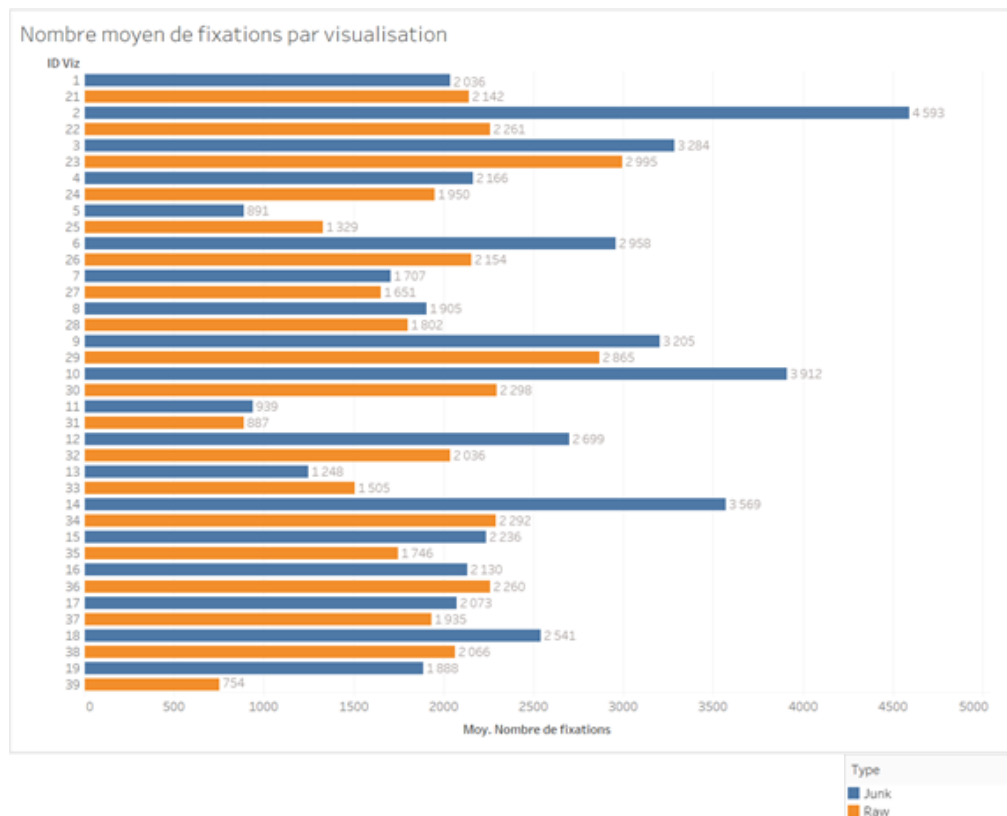


Figure 71 Nombre moyen de fixation par visualisation (agrégation des résultats de tous les participants)

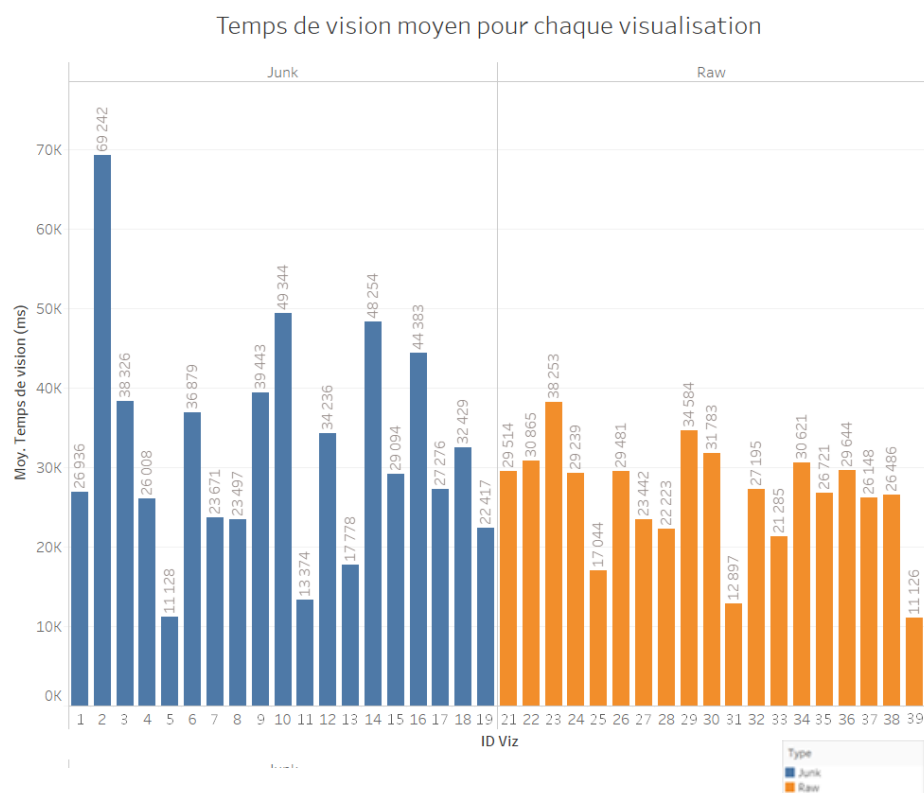


Figure 72 Temps de vision moyen pour chaque visualisation (agrégation des résultats de tous les participants)

Dans presque tous les cas, les visualisations ornementées semblent générer davantage de points de fixation que leurs homologues, mais il y a des exceptions ainsi que de grands écarts. Cette analyse nous permettra de les commenter. Sur la Figure 72, le constat est similaire : le temps de vision est généralement plus élevé pour les visualisations ornementées que pour les visualisations de type standard. Ce que montrent ces deux diagrammes pourrait sembler logique : il y a davantage d'éléments à regarder sur une visualisation ornementée que sur une visualisation de type standard, ce qui pourrait expliquer le nombre de points de fixation et un temps de vision plus grand.

Même si en apparence, cela semble être le cas, peut-on dire que la différence en termes de nombre moyen de fixation est réellement significative entre les visualisations de type ornementé et les visualisations de type standard ? Et surtout, comment interpréter cette différence ? Commençons par répondre aux deux premières hypothèses.

Ainsi, selon **H1**, « la durée moyenne des fixations est significativement plus importante pour les visualisations ornementées que pour les visualisations standard ». Justement, nous n'avons pas encore observé la durée moyenne des fixations. Cette mesure calcule combien de temps durent, en moyenne, les fixations réalisées par un participant sur une visualisation. Nous avons donc calculé la durée moyenne des fixations par visualisation, en agrégeant le résultat pour tous les participants. Cette métrique permet effectivement de déceler soit la complexité, soit l'engagement.

Pour savoir si la durée moyenne des fixations est significativement plus importante pour les visualisations ornementées que pour les visualisations de standard, nous avons besoin de réaliser un test de comparaison de moyennes. En cohérence avec notre précédente analyse statistique ([Partie III](#)), nous vérifions, dès que cela est possible, si nos données suivent une loi normale et si l'homogénéité des variances est respectée. Ces deux conditions nous permettront de choisir le test statistique à appliquer, dans tous les cas de figure de cette analyse.

Selon le Tableau 49, nous devons rejeter H_0 selon laquelle les données suivent une loi normale : la *p-value* est inférieure à 0,05 et est significative ($p = .000$). Par contre, le Tableau 50 nous montre qu'il est impossible de rejeter l'hypothèse selon laquelle les variances de nos données sont homogènes, puisque le test est significatif (colonne « Sig. »).

Tableau 49 Test de normalité de la durée moyenne des fixations, par type (Raw, Junk)

Tests of Normality				
		Kolmogorov-Smirnov ^a		
Type		Statistic	df	Sig.
Raw	Durée moyenne des fixations	.194	381	.000
Junk	Durée moyenne des fixations	.328	378	.000

a. Lilliefors Significance Correction

Tableau 50 Test d'homogénéité des variances pour la durée moyenne des fixations, par type (Raw, Junk)

Test of Homogeneity of Variances					
		Levene Statistic	df1	df2	Sig.
Durée moyenne des fixations	Based on Mean	D2.166	1	757	.141
	Based on Median	1.734	1	757	.188

Puisque nos données ne suivent pas une loi normale, nous choisissons un test non paramétrique, à savoir Mann Whitney U qui assume l'homogénéité des variances. Les raisons pour lesquelles nous prenons attention à l'homogénéité des variances et aux conditions d'application des tests employés sont expliquées plus en détail dans la [méthode d'analyse de la Partie III](#). Nous n'y reviendrons donc pas ici.

Ainsi, selon le test de Mann Whitney U, les moyennes des rangs de « la durée moyenne des fixations » sont similaires en fonction du type (*Raw*, *Junk*).

$$H_0 : \mu_{\text{rangs Raw}} = \mu_{\text{rangs Junk}}$$

Tableau 51 Test de Mann Whitney U pour la durée des fixations en fonction du type (Raw, Junk)

Test Statistics ^a	
Durée moyenne des fixations	
Mann-Whitney U	59644.500
Wilcoxon W	132415.500
Z	-4.094
Asymp. Sig. (2-tailed)	.000

a. Grouping Variable: Type

Selon le Tableau 51, le test de Mann Whitney U est significatif puisque la *p-value* est inférieure à 0,05. Nous rejetons donc H0 selon laquelle les moyennes sont similaires. La moyenne des rangs est de 347, 55 concernant les visualisations standard et de 412, 71 concernant les visualisations de type ornementé. Attention toutefois : la taille de l'effet est de 22%, ce qui signifie que 22% de la variabilité de ces données est due au type. Nous interprétons cela par le fait que les visualisations de données de type ornementé sont elles-mêmes variées et peuvent être classées selon différents attributs, engendrant donc différents résultats concernant les fixations. C'est pourquoi une série d'hypothèses est prévue en ce qui concerne les attributs de ces visualisations.

Quoi qu'il en soit, **l'hypothèse H1 est validée** : la durée moyenne des fixations est significativement plus importante pour les visualisations ornementées que pour les visualisations standard. L'interprétation suivra en fin de paragraphe.

Penchons-nous sur **H2** désormais. Pour rappel, elle suppose que le nombre de fixations est significativement plus grand pour les visualisations ornementées que pour les visualisations de type standard. Les données concernant le nombre de fixations, par type (*Raw*, *Junk*) ne suivent pas une loi normale et ne respectent pas l'homogénéité des variances, selon le Tableau 52 et le Tableau 53. Puisque le test non paramétrique de Mann Whitney U est censé assumer l'homogénéité des variances, nous allons procéder à un test de Welch (qui est paramétrique) sur les rangs du nombre de fixations.

Tableau 52 Test de normalité du nombre de fixations par type (Raw, Junk)

		Kolmogorov-Smirnov ^a		
Type		Statistic	df	Sig.
Raw	Nombre de fixations	.106	381	.000
Junk	Nombre de fixations	.124	378	.000

a. Lilliefors Significance Correction

Tableau 53 Test de Levene pour le nombre de fixations

		Test of Homogeneity of Variances			
		Levene Statistic	df1	df2	Sig.
Nombre de fixations	Based on Mean	18.986	1	757	.000
	Based on Median	14.075	1	757	.000

Nous avons donc calculé les rangs de la variable « nombre de fixations ». En appliquant le test de Welch sur ces rangs, H0 suppose que, pour le nombre de fixations, les moyennes des rangs soient similaires en ce qui concerne le type.

$$H0 = \mu_{\text{rangs Raw}} = \mu_{\text{rangs Junk}}$$

Nous constatons que le test est significatif : la *p-value* est inférieure à 0,05 (Tableau 54). H_0 est rejetée : les moyennes des rangs ne sont pas similaires. La moyenne des rangs est de 350, 68 pour le type standard et de 409, 54 pour le type ornementé. Ainsi, **H2** est validée : nous affirmons que le nombre de fixations est significativement plus grand pour les visualisations ornementées que pour les visualisations de type standard.

Tableau 54 Test de Welch sur les moyennes des rangs du nombre de fixations en fonction du type (Raw,Junk)

Robust Tests of Equality of Means

Rank of Nombre de fixations

	Statistic ^a	df1	df2	Sig.
Welch	13.905	1	754.162	.000
Brown-Forsythe	13.905	1	754.162	.000

a. Asymptotically F distributed.

3.1.2. Interprétation des résultats de l'analyse des fixations en fonction du type

Selon les deux hypothèses précédemment validées, le nombre de fixations et leur durée moyenne est significativement plus grand pour les visualisations de données ornementées. Ces métriques peuvent chacune s'expliquer de deux façons, à en croire Poole et Ball (1996) et Bylinskii et ses collègues (2017).

La durée moyenne des fixations qui est plus haute pour les visualisations de type ornementé peut s'expliquer par le fait que ces représentations soient plus complexes visuellement ou plus engageantes. Quant au nombre de fixations plus élevé, il peut montrer que la tâche de compréhension était plus difficile à réaliser lorsque les visualisations de données étaient ornementées ou bien que celles-ci provoquent davantage d'engagement. En réalité, nous ne devons pas choisir quelle interprétation appliquer de façon générale à notre cas. Toutes ces interprétations sont valables, en fonction des visualisations ornementées et de leur particularité. Ainsi, grâce à la précédente analyse des appréciations déclarées, nous sommes en mesure de proposer quelques explications.

De manière générale, les visualisations de données ornementées sont toujours estimées plus belles et sont favorisées par les participants, ce qui montre qu'elles sont engageantes. De même, l'ornementation diminue significativement la sensation de clarté par rapport à un visuel standard, mais en la gardant à un niveau satisfaisant à certaines conditions. Cela peut tout de même expliquer un nombre de points de fixations sensiblement plus haut pour des visualisations ornementées, en comparaison à des visualisations de type standard. Toutefois, une exception se marque clairement pour un type de visualisations ornementées : celles dont l'ornement se concentre autour d'une forme

inhabituelle et qui ne présente pas d'icônes, comme les visualisations n°2, 4 et 10. Les appréciations déclarées concernant la clarté de ces visualisations de données sont très faibles. Pour ces visualisations, le nombre et la durée des fixations s'expliqueraient donc plutôt par de la complexité et de la difficulté à déceler l'information. C'est pourquoi l'exploration d'une hypothèse est prévue plus tard pour en analyser exactement le nombre de fixations et leur durée moyenne.

L'analyse des appréciations déclarées a aussi montré que les icônes sont très appréciées, mais qu'elles diminuent quelque peu la clarté de la visualisation si elles se constituent sur le *data-ink*. À ce stade, difficile donc de déterminer s'il peut s'agir d'engagement ou de complexité. Nous le déterminerons plus tard grâce aux hypothèses.

Finalement, rappelons-nous qu'il n'y a pas de différence entre une visualisation ornementée et une visualisation standard en ce qui concerne la clarté, lorsque la figuration de l'icône est apposée. Dans ce cas, il y a fort à parier que concernant les visualisations telles que la n°17 du corpus, le nombre de fixations plus haut puisse exprimer de l'engagement. Nous tenterons également de le vérifier ensuite.

En tentant de répondre aux hypothèses suivantes, nous apporterons plus d'éléments tangibles permettant de soutenir ces suppositions.

3.2. Comparaison des fixations en fonction d'attributs significatifs

3.2.1. Analyse des données

Afin d'analyser plus en profondeur l'incidence des attributs sur les données de fixations, nous allons tenter de répondre aux hypothèses H3 à H6.

Selon l'hypothèse **H3**, le nombre de fixations est significativement plus élevé sur les visualisations ornementées à la forme inhabituelle que sur les visualisations ornementées dont la figuration se constitue en icônes. Sans surprise, les données ne suivent pas une loi normale et l'homogénéité des variances n'est pas respectée⁴⁵. Afin de vérifier l'hypothèse, nous comparerons les moyennes des rangs du nombre de fixations en fonction de l'attribut « ornement » en appliquant un test de Welch. Ainsi,

$$H0 : \mu_{\text{rangs icônes}} = \mu_{\text{rangs formes inhabituelles}}$$

⁴⁵ Étant donné que cette tâche est répétitive et que les hypothèses H1 et H2 ont pu servir d'exemple, nous ne présenterons plus les tableaux attestant de la normalité et de l'homogénéité des variances.

Le résultat du test de Welch, montré dans le Tableau 55, est significatif : la *p-value* est inférieure à 0,05. Ainsi, nous rejetons H0 selon laquelle les moyennes des rangs des nombres de fixations de l'ornement « icônes » et de l'ornement « formes inhabituelle » sont similaires. D'ailleurs, la différence entre ces moyennes des rangs est grande : la moyenne des rangs est de 388,12 pour les icônes et de 523,08 pour les formes inhabituelles. Cela montre effectivement que la difficulté de compréhension de l'information est physiologiquement plus grande lorsque les visualisations de données ornementées présentées aux individus sont des formes inhabituelles et fantaisistes. L'hypothèse **H3** est donc validée.

Tableau 55 Test de Welch sur le nombre de fixations, en fonction de l'attribut « Ornement », pour les visualisations ornementées uniquement

Robust Tests of Equality of Means

Rank of Nombre de fixations

	Statistic ^a	df1	df2	Sig.
Welch	21.228	1	86.477	.000

a. Asymptotically F distributed.

L'hypothèse suivante concerne la figuration de l'ornementation. Plus précisément, nous nous concentrons sur le fait que l'icône soit apposée ou qu'elle constitue le *data-ink*.

H4 Le nombre de fixations est significativement plus élevé sur les visualisations dont la figuration se constitue sur les traits relatifs aux données (*data-ink*) que sur les visualisations dont la figuration est apposée aux données.

De nouveau, les données de la variable « nombre de fixations » ne suivent pas une loi normale et ne respectent pas l'homogénéité des variances lorsqu'elles sont considérées du point de vue de l'attribut « figuration » (*apposé, sur le data-ink*). Nous conduisons dès lors un test de Welch sur les rangs de cette variable. L'hypothèse H0 suppose toujours la similarité des moyennes des rangs.

Tableau 56 Test de Welch sur les rangs du nombre de fixations, en fonction de l'attribut « Figuration » (*Apposé, Sur le Data-Ink*)

Robust Tests of Equality of Means

Rank of Nombre de fixations

	Statistic ^a	df1	df2	Sig.
Welch	8.696	1	292.334	.003
Brown-Forsythe	8.696	1	292.334	.003

a. Asymptotically F distributed.

La dernière colonne du Tableau 56 nous montre que le test est significatif, avec une *p-value* de 0,003 qui est donc inférieure à 0,05. Nous rejetons l'hypothèse de similarité des moyennes des rangs. Il y a donc bien une différence significative entre le nombre moyen de fixations pour ces deux groupes de visualisations de données ornementées. Néanmoins, la moyenne des rangs des visualisations dont la figuration est *apposée* est de 432,15, contre 362,15 pour les visualisations dont la figuration se constitue sur le *data-ink*. Ainsi, la tendance que nous constatons est inverse à ce que propose l'hypothèse. Il y a donc bien une différence significative concernant les nombres moyens de fixation entre les deux groupes de visualisation de données ornementées, mais inverse à ce que nous supposions.

L'hypothèse **H4 est rejetée**. Le nombre de fixations n'est pas significativement plus élevé sur les visualisations dont la figuration se constitue sur les traits relatifs aux données (*data-ink*) que sur les visualisations dont la figuration est apposée aux données. C'est tout à fait l'inverse. Nous en interpréterons les raisons par la suite.

Ensuite, selon l'hypothèse suivante, **H5**, au sein des visualisations ornementées, en comparaison à une situation dans laquelle le *lie factor* est fidèle, la durée moyenne des fixations est significativement plus élevée lorsque le *lie factor* est distorsif. Comme précédemment, les données dans ce cas ne suivent pas une loi de normale et l'homogénéité des variances n'est pas respectée. Le test de Welch sur les rangs du nombre de fixations n'est pas significatif. Effectivement, on constate dans la dernière colonne du Tableau 57 que la *p-value* est égale à 0,191 (donc $p > 0,05$).

Tableau 57 Test de Welch sur le nombre de fixations, en fonction du *lie factor* (Fidèle, Distorsif)

Robust Tests of Equality of Means				
Rank of Nombre de fixations				
	Statistic ^a	df1	df2	Sig.
Welch	1.714	1	472.424	.191
Brown-Forsythe	1.714	1	472.424	.191

a. Asymptotically F distributed.

Ainsi, nous ne pouvons pas rejeter H0 selon laquelle les moyennes des rangs des groupes « Fidèle » et « Distorsif » sont similaires. Cela nous pousse donc à affirmer que le *lie factor*, qu'il soit distorsif ou fidèle, ne semble pas induire de différences significatives en ce qui concerne le nombre de fixations des données ornementées. **L'hypothèse H5 est rejetée** : en comparaison à une situation dans laquelle le *lie factor* est fidèle, la durée moyenne des fixations n'est pas significativement plus élevée lorsque le *lie factor* est distorsif.

Pour terminer, l'hypothèse **H6** suppose qu'au sein des visualisations ornementées,

- un niveau de *data-ink* élevé induit un temps de vision significativement plus bas qu'un niveau de *data-ink* bas ou moyen ;
- un niveau de *data-ink* moyen induit un temps de vision significativement plus bas qu'un niveau de *data-ink* bas.

Les différentes données en jeu dans cette hypothèse ne suivent pas une loi normale et ne respectent pas l'homogénéité des variances. Un test de Brown Forsythe est donc réalisé sur les rangs du temps de vision. En comparant les moyennes des rangs, ce test suppose que les différentes moyennes soient similaires.

Tableau 58 Test de Brown Forsythe sur les moyennes des rangs du Temps de vision, en fonction du niveau de data-ink

Robust Tests of Equality of Means

Rank of Temps de vision (ms)

	Statistic ^a	<u>df1</u>	<u>df2</u>	Sig.
Brown-Forsythe	7.182	2	166.923	.001

a. Asymptotically F distributed.

La dernière colonne du Tableau 58 nous permet de rejeter H0 selon laquelle les moyennes des rangs du temps de vision sont similaires lorsqu'elles sont considérées en fonction du niveau de *data-ink*. Il y a donc des différences significatives entre ces moyennes. Puisqu'il y a trois groupes de données, nous avons reproduit successivement le test entre les groupes, deux à deux (élevé et bas, élevé et moyen, moyen et bas). Puisque nous multiplions les comparaisons de moyennes, nous ajustons notre seuil de signification en appliquant la correction de Bonferroni suivante : $\alpha' = 0,05 / 3$ (seuil de signification $\alpha = 0,05$ sur le nombre de comparaisons). Notre nouveau seuil de signification est 0,016.

Entre des niveaux de data-ink bas et élevé

La *p-value* est significative ($p = 0,005$ dans le Tableau 59). L'hypothèse H0 de la similarité des moyennes des rangs du temps de vision est rejetée.

Les moyennes des rangs s'élèvent à 505,79 pour un *data-ink* bas contre 400,48 pour un *data-ink* élevé. La différence est significative et le temps de vision est plus bas pour une visualisation ornementée au *data-ink* élevé, en comparaison à une visualisation ornementée au *data-ink* bas.

Tableau 59 Test de Brown Forsythe sur le temps de vision en fonction du niveau de data-ink (bas, élevé)

Robust Tests of Equality of Means

Rank of Temps de vision (ms)

	Statistic ^a	df1	df2	Sig.
Brown-Forsythe	8.652	1	56.313	.005

a. Asymptotically F distributed.

Entre des niveaux de data-ink moyen et bas

La *p-value* est significative ($p = .000$ dans le Tableau 60). L'hypothèse H_0 de la similarité des moyennes des rangs du temps de vision est rejetée.

Les moyennes des rangs s'élèvent à 505,79 pour un *data-ink* bas contre 351,50 pour un *data-ink* moyen. La différence est significative et le temps de vision est plus bas pour une visualisation ornementée au *data-ink* moyen, en comparaison à une visualisation ornementée au *data-ink* bas.

Tableau 60 Test de Brown Forsythe sur le temps de vision en fonction du niveau de data-ink (moyen, bas)

Robust Tests of Equality of Means

Rank of Temps de vision (ms)

	Statistic ^a	df1	df2	Sig.
Brown-Forsythe	14.796	1	81.171	.000

a. Asymptotically F distributed.

Entre des niveaux de data-ink moyen et élevé

La *p-value*, s'élevant à 0,073 (Tableau 61), n'est pas significative ($p > 0,05$). L'hypothèse H_0 de la similarité des moyennes des rangs du temps de vision ne peut pas être rejetée. Ainsi, il n'y a pas de différence significative entre les temps de vision des visualisations ornementées dont le niveau de *data-ink* est moyen ou élevé.

Tableau 61 Test de Brown Forsythe sur le temps de vision en fonction du niveau de data-ink (moyen, élevé)

Robust Tests of Equality of Means

Rank of Temps de vision (ms)

	Statistic ^a	df1	df2	Sig.
Brown-Forsythe	3.257	1	177.151	.073

a. Asymptotically F distributed.

Pour conclure, **l'hypothèse H6 ne peut être validée que partiellement**. Effectivement, au sein des visualisations ornementées, un niveau de *data-ink* élevé induit un temps de vision significativement plus bas qu'un niveau de *data-ink* bas et un niveau de *data-ink* moyen induit un temps de vision significativement plus bas qu'un niveau de *data-ink* bas. Par contre, il n'y a pas de différence significative en termes de temps de vision entre les visualisations de données aux niveaux de *data-ink* moyen et élevé.

3.2.2. Interprétation des résultats de l'analyse des fixations en fonction des attributs
Selon H3, qui a été validée, les visualisations ornementées à la forme inhabituelle génèrent bien plus de points de fixation que les visualisations ornementées grâce à des icônes. Puisque les appréciations déclarées ont montré que ces visualisations sont jugées moins belles que les visualisations standard, nous n'interprétons pas cela par de l'engagement. De même, l'appréciation de clarté était très basse pour ces visualisations et les différents verbatims les concernant nous poussent à interpréter ce grand nombre de points de fixations comme de la difficulté. Nous savons désormais que la difficulté de ces visualisations n'est pas seulement une interprétation des utilisateurs : elle existe aussi physiologiquement.

Oui j'ai vraiment rigolé, c'est vraiment le pire parce qu'il y a plein de couleurs et il y a trois grands pans d'information, gauche, dans le rond et droite. Et en fait, il n'y a rien qui te dit dans l'image où regarder en premier donc fatalement tu vas un peu partout et au final, tu as total d'heures de sommeil tu comprends quand même. (...) Mais dans l'horloge là pfff franchement, ils ne commencent et finissent pas au même endroit... Vu qu'ils ne commencent pas au même endroit on ne sait pas comparer.

Ensuite, à notre grand étonnement, nous avons rejeté l'hypothèse H4. Le nombre de fixations n'est pas significativement plus élevé sur les visualisations dont la figuration se constitue sur les traits relatifs aux données (*data-ink*) que sur les visualisations dont la figuration est apposée aux données. Avec cette hypothèse, nous espérions rejoindre le résultat de l'analyse des appréciations déclarées en montrant que les visualisations dont l'ornementation se constitue sur le *data-ink* semblent plus complexes visuellement.

Moyenne du nombre de fixations pour l'attribut "Figuration" (apposé, sur le data-ink)

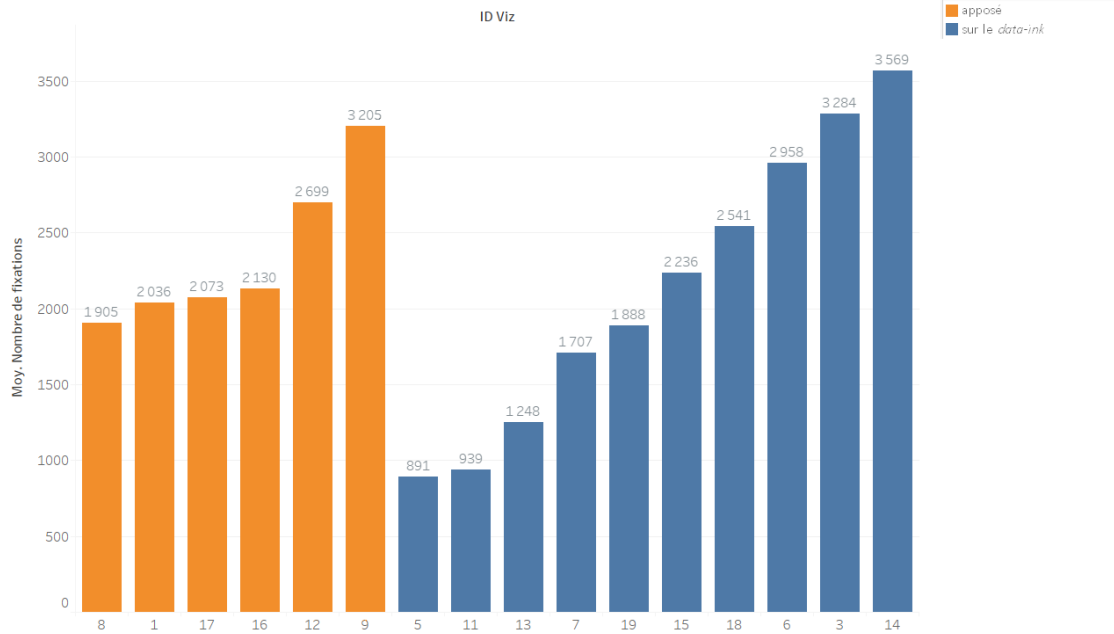


Figure 73 Moyenne du nombre de fixations pour l'attribut « Figuration », par visualisation

La Figure 73 nous montre que le nombre de fixations est globalement plus grand pour les visualisations à l'ornementation apposée. Avant tout, l'analyse des entretiens semi-directifs a dévoilé que la simplicité visuelle des visualisations n° 5, 11, 13 et 7 a été très appréciée. Ces visualisations, même si l'ornementation se concentre sur le *data-ink*, sont épurées visuellement, ce qui expliquerait leur nombre bas de fixations. Quant aux visualisations à l'ornementation apposée, deux d'entre elles sont, au contraire, très chargées visuellement (visualisations n°1 et n°16). Deux autres d'entre elles, bien que simples, montrent deux diagrammes (n°9, n°17) tandis que la n°12 a été écrite dans les entretiens semi-directifs comme difficile à comprendre d'un point de vue textuel, ce qui explique le nombre plus élevé de points de fixations. En somme, même si nous avons invalidé l'hypothèse, elle est finalement porteuse de peu de sens étant donné la variété de notre corpus. Par contre, elle montre de nouveau que la simplicité paie : les visualisations n°6, 14 et 18 sont très chargées visuellement. On voit clairement sur la Figure 73 la grande différence de points de fixation entre ces visualisations et celles dont l'ornementation est épurée (n° 5, 11, 13, 7).

Puis, nous avons également rejeté l'hypothèse H5 : en comparaison à une situation dans laquelle le *lie factor* est fidèle, la durée moyenne des fixations n'est donc pas significativement plus élevée lorsque le *lie factor* est distorsif. Pour cette hypothèse, nous avons mené l'analyse sur les visualisations ornementées uniquement, contrairement à l'analyse des appréciations déclarées qui incluait également les visualisations de type standard. Ainsi, nous constatons ici très clairement que la mise en garde que nous formulions précédemment se vérifie encore : les visualisations au *lie factor* distorsif ne semblent pas plus complexes pour l'œil humain. De même, tout comme cela n'avait aucun

impact sur les appréciations affectives (« J'aime », « Belle »), nous constatons que le *lie factor* fidèle n'engendre pas forcément plus d'engagement d'un point de vue visuel que le *lie factor* distorsif. En somme, le *lie factor* distorsif présente un rapport aux données complètement faussé et le participant, à aucun niveau de l'expérimentation, ne se doute que le visuel ne reflète pas l'exactitude des données.

Pour terminer, nous avons partiellement validé l'hypothèse H6. Nous avons pu déduire que, au sein des visualisations ornementées

- un niveau de *data-ink* élevé induit un temps de vision significativement plus bas qu'un niveau de *data-ink* bas ;
- un niveau de *data-ink* moyen induit un temps de vision significativement plus bas qu'un niveau de *data-ink* bas ;
- le temps de vision ne varie pas forcément en fonction d'un niveau de *data-ink* moyen ou élevé.

Nous constatons ainsi qu'une visualisation dont le *data-ink* est bas (donc beaucoup d'ornementation et finalement peu de place aux données) augmente le temps de vision. Au-delà de la complexité ou de l'engagement, nous songeons à interpréter cela comme de la distraction. En revanche, entre une ornementation apposée, une ornementation complètement figurée sur le *data-ink* ou l'absence d'ornementation, le temps de vision n'est pas significativement différent. La remise en question du *data-ink ratio* se discute donc encore.

3.3.Exploration des zones d'intérêt

Les hypothèses suivantes se concentrent sur l'analyse des zones d'intérêt. Pour rappel, nous n'étudions que les icônes apposées. Cette analyse se concentre donc sur dix visualisations de données ornementées uniquement. L'objectif est d'évaluer si les icônes apposées causent de la distraction lors de la lecture des visualisations de données. Sur base des analyses quantitative et qualitative précédentes (Parties [III](#) et [IV](#)), nous avons dessiné les zones d'intérêt et formulé ces hypothèses. Étant donné leur nature et parce que cela ne se justifie pas, nous n'appliquerons pas de test statistique aux données fournies par Tobii Studio afin de répondre à ces hypothèses. Nous serons néanmoins en mesure d'observer et de commenter ces données.

Lors d'une première approche des données, nous avons supprimé trois zones d'intérêt appartenant à deux visualisations de données. La raison en est qu'une personne uniquement avait réalisé des fixations sur ces zones qui étaient donc des valeurs aberrantes.

Afin de répondre aux hypothèses, nous utiliserons trois métriques découvertes dans la littérature et trois nouvelles métriques proposées par Tobii Studio :

- Nombre de fixations ;
- Durée des fixations ;
- *Time-to-first fixation* : temps passé jusqu'à ce que l'individu réalise une première fixation sur la zone d'intérêt (secondes) ;
- *Fixations before* : nombre de fixations ayant été réalisées par l'individu avant qu'il ne fixe une zone d'intérêt ;
- *Visit duration* : laps de temps au cours duquel l'individu réalise des fixations, en secondes.

Nous utiliserons également le « Temps de vision » que nous avons calculé dans le cadre de la falsification des hypothèses précédentes.

L'hypothèse **H7** suppose que les ornements apposés captent rapidement l'attention : ils sont regardés avant la lecture du graphique. Dès la première approche des données, nous devons **rejeter l'hypothèse H7**. Effectivement, la moyenne des « *Time-to-first fixation* » s'élève à 11, 26 secondes pour une moyenne de temps de vision de 33,89 secondes. Ainsi, on peut dire que la première fixation sur une zone d'intérêt se réalise généralement à un tiers du temps de vision de la visualisation. L'ornement apposé n'est pas le premier élément observé par le participant. Par contre, le nombre moyen de fixations ayant lieu avant la première fixation dans la zone d'intérêt est de 39. Cela peut sembler assez peu lorsque l'on sait que le nombre moyen de fixations pour ces visualisations est de 2389. Le Tableau 62 résume ces différents nombres. Ainsi, la métrique « *Time-to-first fixation* » nous indique que les icônes apposées n'ont pas une grande propriété de captation de l'attention puisque ce n'est pas le premier élément observé sur l'image et les premières fixations apparaissent au tiers du temps de vision.

Tableau 62 Tableau synthétique des moyennes de métriques relatives à H7

<i>Time-to-first fixation</i> (zones d'intérêt)	11,26s	<i>Fixations before</i> (zones d'intérêt)	39
Temps de vision (Global)	33,89s	Nombre de fixations (Global)	2389

Ensuite, selon l'hypothèse **H8**, les ornements apposés ne sont pas des éléments importants : le nombre de fixations est faible dans ces zones d'intérêt par rapport au nombre moyen de fixations. Les données montrent qu'il y a en moyenne 10 fixations par zone d'intérêt. Effectivement, en sachant que le nombre moyen de fixation au sein de ces visualisations

est de 2389 pour un temps de vision moyen s'élevant à 33,89 secondes (c'est-à-dire environ 71 fixations par seconde), nous pouvons effectivement dire que ce nombre est bas et que les ornements apposés n'apparaissent donc pas comme des éléments importants.

L'hypothèse H8 est donc validée.

La dernière hypothèse concernant les zones d'intérêt est la suivante (**H9**). Les ornements apposés ne sont pas d'une grande complexité visuelle : la durée moyenne des fixations dans les zones d'intérêt est relativement basse par rapport à la durée moyenne des fixations pour les visualisations dans leur ensemble. La durée moyenne des fixations se déroulant dans les zones d'intérêt est de 0,16 secondes. En comparaison à la durée des fixations pour les visualisations dans leur ensemble, qui est de 0,19 secondes, la différence n'est que de 3 centièmes de seconde. Cet élément ne nous permet pas de valider ou d'invalidier l'hypothèse. Effectivement, cela pourrait également être interprété comme de l'engagement plutôt que de la complexité. En revanche, nous sommes en mesure de calculer le ratio de temps octroyé aux zones d'intérêt au cours de la lecture des visualisations. Pour se faire, nous utilisons la métrique « Temps de visite » (*Visit Duration*) qui indique le temps passé sur une zone d'intérêt et calculons le rapport entre cette métrique et le temps de vision global. Ainsi, il s'avère que sur le temps de vision total de ces visualisations de données, en moyenne, seul 0,82% du temps est accordé à l'observation des zones d'intérêt. Ce rapport est très bas : les participants voient simplement l'icône et n'y passent presque pas de temps. Cela nous permet de dire que les icônes ne sont pas des éléments complexes à comprendre. De même, les participants y prêtent en réalité très peu d'attention. Ce résultat est cohérent avec différents propos rapportés lors des entretiens semi-directifs par les participants. Ils prétendaient effectivement avoir vu ces ornements sans pour autant avoir passé du temps à les regarder. De même, ces icônes ont été décrites par les participants comme une mise en contexte. Nous pouvons ainsi conclure que les icônes sont certes engageantes et peu complexes, en tout cas lorsqu'elles sont apposées au *data-ink*.

Pour résumer, les zones d'intérêt sélectionnées sont loin d'être le premier élément observé par les participants. L'interaction des participants avec ces icônes apposées se produit à un tiers du temps de vision et dure très peu de temps : ce sont des éléments peu complexes, auxquels les participants accordent peu d'attention. En somme, nous pouvons affirmer que les icônes apposées ne causent aucune distraction quant à la concentration du lecteur sur le *data-ink*.

Chapitre 3 Détection des *patterns* de lecture : le *bundling*

Nous expliquons maintenant pourquoi le développement d'une nouvelle technique de visualisation a été nécessaire afin de détecter les différents chemins de lecture contenus dans les *gaze data*. Ensuite, nous présentons les observations les plus importantes que nous avons réalisées dans le cadre de cette analyse, afin d'aboutir à une série d'interprétations.

1. Pertinence d'une nouvelle technique de visualisation

Au-delà de la validation de simples hypothèses, l'exploitation des *gaze data* nous mène vers deux objectifs : identifier des chemins de lecture similaires entre tous les participants et vérifier la potentielle incidence de l'ornementation sur ces derniers.

En *eye tracking*, il est commun que l'analyse ne se fasse pas seulement d'un point de vue statistique mais également par l'interprétation de la visualisation des *gaze data* elle-même. Les données brutes fournies par l'*eye tracker* sont si complexes qu'il est parfois nécessaire de les simplifier pour aboutir à une analyse visuelle (Peysakhovich et Hurter, 2017). Nous avons précédemment indiqué que Tobii Studio fournissait ce type de visualisations (*gaze plots* et cartes de chaleur). Néanmoins, la difficulté dans notre cas est engendrée par le nombre de participants et le nombre de fixations qu'ils produisent sur une image. Grâce aux nœuds numérotés et aux liens, il est normalement possible sur un *gaze plot* de retracer le chemin parcouru par un ou deux utilisateurs, comme sur la Figure 74. Le *gaze plot* est une force parce qu'il permet d'identifier les différences dans les trajets oculaires réalisés. Mais une fois les trajets oculaires de tous les participants visualisés sur un même *gaze plot*, l'information est illisible. Le *gaze plot*, censé nous aider à détecter les différents chemins de lecture, est dans ce cas inexploitable, comme le montre la Figure 75. La version animée n'est pas une alternative convenable : au final, le résultat est identique et encombré.

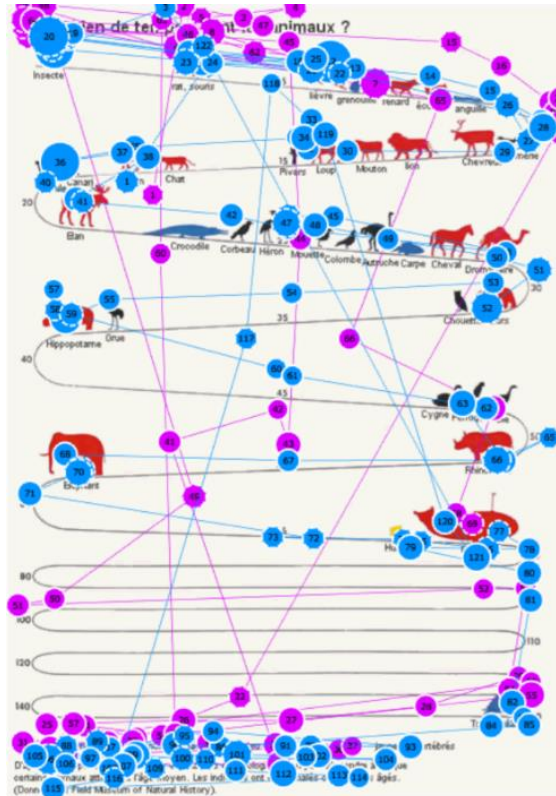


Figure 74 Trajets oculaires (gaze plots) de deux participants pour la visualisation n°3

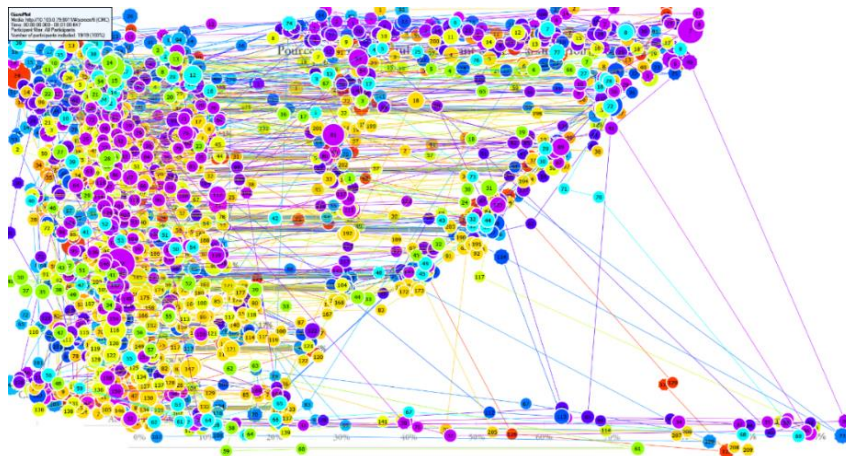


Figure 75 Trajets oculaires de la totalité des participants (20) pour la visualisation n°26

Quant à la carte de chaleur, elle montre très clairement soit les endroits où les points de fixation sont très nombreux, soit les durées de fixation les plus longues. Elles ne donnent en revanche aucune idée de flux : impossible de savoir où est le début et la fin d'un quelconque chemin.

Ainsi, pour analyser et tirer des conclusions de l'agrégation de 20 trajets oculaires différents et des milliers de points de fixation à afficher par visualisation, les visualisations fournies

par le logiciel Tobii Studio ne suffisent pas. Dans ce cadre, nous avons eu la chance de pouvoir bénéficier d'une technique d'agrégation des différents trajets oculaires : le *bundling*.

En juin 2019, nous avons réalisé un séjour de recherche au sein de l'équipe DEVI (Données, Économie et Visualisation Interactive) de l'École Nationale de l'Aviation Civile (ENAC) à Toulouse. L'objectif de ce séjour, encadré par le Professeur Hurter, était de vérifier si la technique de *bundling*, largement utilisée dans le domaine de l'aéronautique, pouvait s'appliquer dans notre cas en particulier.

Le principe du *bundling*⁴⁶ est celui-ci : les fixations et saccades sont extraites des *gaze data*. Les fixations sont groupées en *cluster* et les saccades sont regroupées ensemble. Le tout est visualisé sur une carte de flux (Peysakhovich et Hurter, 2017). Dans la version de *bundling* existant déjà auparavant dans le travail de Peysakhovich et Hurter (2017), les saccades sont représentées par des liens. Le *bundling* regroupe celles qui sont proches et qui vont dans la même direction. La couleur indique la direction des *bundles*⁴⁷ (grâce à une roue chromatique). La Figure 76 est un exemple de résultat issu de cette méthode qui, elle-même, est décrite dans (Peysakhovich et Hurter, 2017).

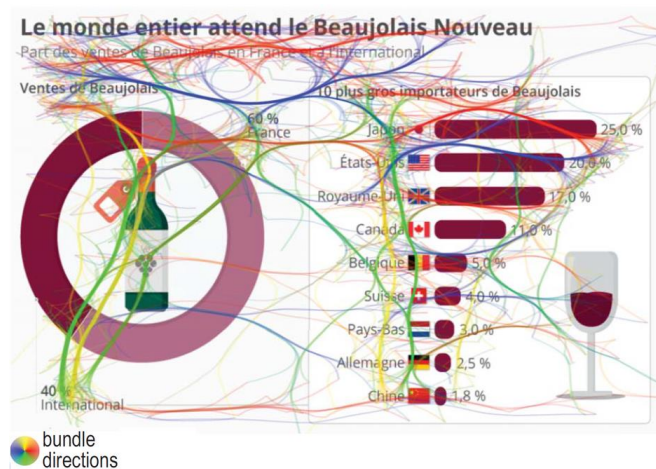


Figure 76 Exemple d'un premier résultat de bundling appliqué à la visualisation n°17

Néanmoins, afin de produire une visualisation plus adaptée à notre objectif, différentes améliorations ont été apportées à cette technique de visualisation par le Professeur Hurter. Une carte de densité des fixations a effectivement été ajoutée à l'arrière des *bundles*. Sur la Figure 77, on remarque que les principaux faisceaux commencent et se terminent dans les six principales zones de fixation (bleu foncé et rouge foncé dans la carte de densité).

⁴⁶ Méthodologie complète détaillée dans Andry T., Hurter C., Lambotte F., Fastrez P., Telea A. (2020), Interpreting the Effect of Embellishment on Chart Visualizations, soumis à CHI '21: ACM Symposium.

⁴⁷ Un *bundle* est un lien

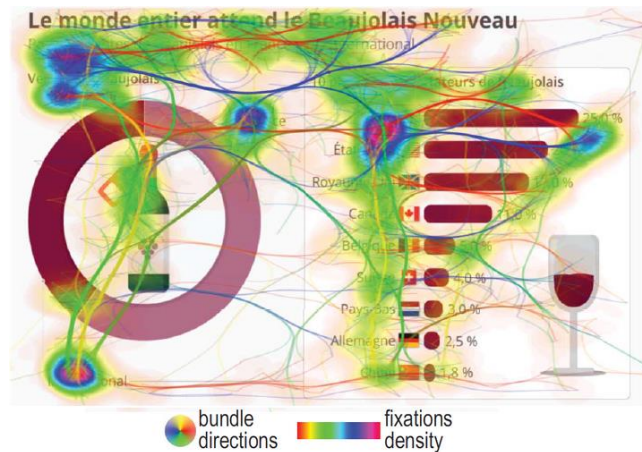


Figure 77 Exemple de carte de densité des fixations jointe aux bundles pour la visualisation n°17

La visualisation a ensuite été simplifiée : en l'état elle était encombrée d'ensembles de liens de faible densité qui bifurquent des *bundles* principaux, en raison de points de fixation répartis à peu près uniformément sur de grandes surfaces ainsi que de signaux faibles dans les zones où seules quelques saccades se fondent en faisceaux. Ces saccades correspondent à des trajets oculaires qui apparaissent à une faible fréquence. Il s'agit donc d'événements qui sont peu récurrents dans la lecture de l'image par les participants. Moins importants que le reste de l'information, ces liens ont été retirés de la visualisation.

La technique a encore été optimisée après cela afin d'aboutir à la version finale (Figure 78). L'épaisseur des faisceaux code la densité (densité des saccades allant dans le même sens et importance des fixations). Les faisceaux les plus importants apparaissent au premier plan de l'image ; ils sont plus saillants. La direction des faisceaux est codée grâce à une texture en forme de « V » pointant vers la direction suivie. Finalement, et c'est là l'originalité de cette technique, la luminance du faisceau (du bleu foncé au blanc) code la temporalité du trajet oculaire : le bleu foncé code ainsi le début du chemin de lecture et le blanc en code la fin (« *gaze time* » sur la Figure 78).

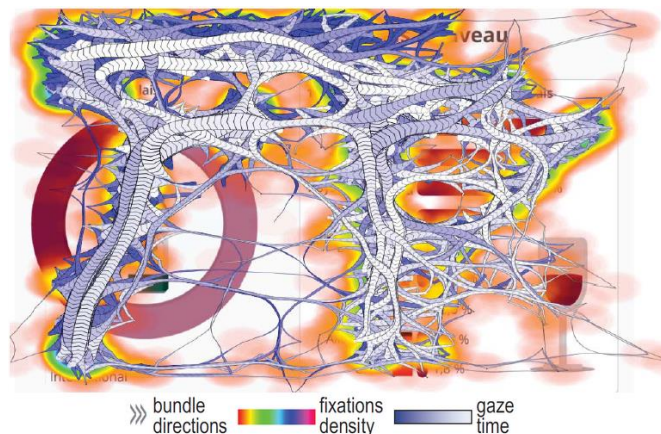


Figure 78 Version finale du bundling appliquée à la visualisation n°17

Ainsi, cette version finale permet de repérer, sans trop d'interférences, les faisceaux saillants (épais), leurs directions (grâce aux « V » en texture), l'ordre dans lequel les participants lisent l'image. Sur la Figure 78, on peut décrire succinctement l'ordre de lecture codé par la luminance comme ceci : titre d'abord, chiffres en bas à gauche ensuite, diagramme à barres à droite et enfin sous-titre du diagramme circulaire). On remarque également les zones de fixation (carte de couleurs du spectre des teintes). Cette technique de visualisation semble donc idéale pour analyser et comparer les trajets oculaires produits par les participants. Tous les résultats du *bundling* se trouvent en annexe 11 ou sur le [site web associé](#).

2. Analyse des chemins de lecture : observations et commentaire général sur le résultat du *bundling*

De manière générale, nous constatons bel et bien des *patterns* de lecture au sein des différentes visualisations faisant partie du corpus. Les différentes visualisations sont visibles en annexe 11. Le *bundling* agrège les fixations de tous les participants. On voit, sur la plupart des images, un résultat clair et des trajectoires qui se définissent et se répètent. La visualisation des *gaze data* grâce à la technique de *bundling* est donc un succès. Nous présentons ci-dessous les observations que cette technique nous a permis de réaliser.

2.1. Les cycles ou le « 3-stage model »

Pour commencer, nous avons observé qu'en général, la lecture d'une visualisation de données se fait de façon cyclique. L'œil produit un chemin de lecture qu'il reproduit d'une à deux fois en fonction de la complexité de l'image. Prenons pour exemple la Figure 79 qui montre la version ornementée et standard des mêmes données ainsi que le *bundling* qui y est appliqué (original à gauche, *bundling* à droite). Sur le *bundling* de ces deux images, on voit clairement en bleu foncé que la lecture commence par le titre avant de descendre et d'explorer les différentes étiquettes (données et titres des bâtonnets). En bleu clair, on constate la suite du trajet oculaire dans le temps. Elle est tout à fait similaire aux traces bleu foncé. Finalement, les traces blanches terminent le cycle en revenant au point de départ. Le *bundling* illustre bien les différentes étapes du *three-stage model* de Körner (2011). En effet, on distingue sur les images de droite sur la Figure 79 qu'une même étape est répétée de façon similaire une seconde fois et parfois une troisième fois. Selon Körner, l'œil effectue davantage de fixations qui suivent un autre schéma que celui réalisé au cours des deux premières étapes. Peut-être que la décision, prise au cours du développement de la technique de *bundling*, d'enlever les fixations et saccades les plus isolées de la

visualisation, nous empêche de voir ces fixations plus disparates. Néanmoins, le *bundling* permet de montrer qu'une dernière étape confirme en quelques sortes la lecture. De même, oralement, certains participants ont précisé qu'ils relisaient plusieurs fois la visualisation afin de confirmer qu'ils ont bien compris les données.

« Je relis toujours pour être sûre, même s'il n'y a pas d'efforts à faire »

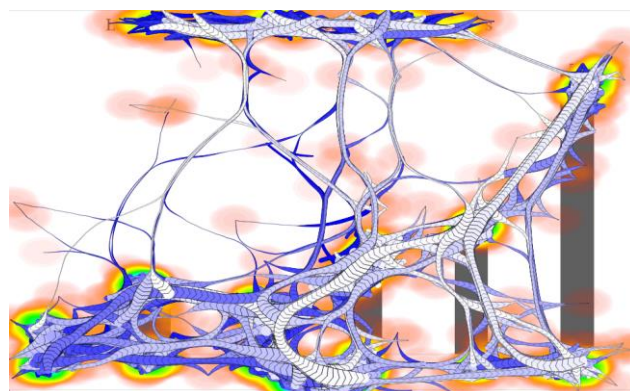
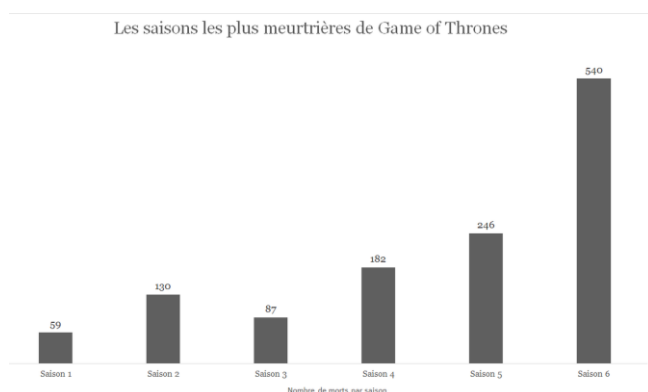
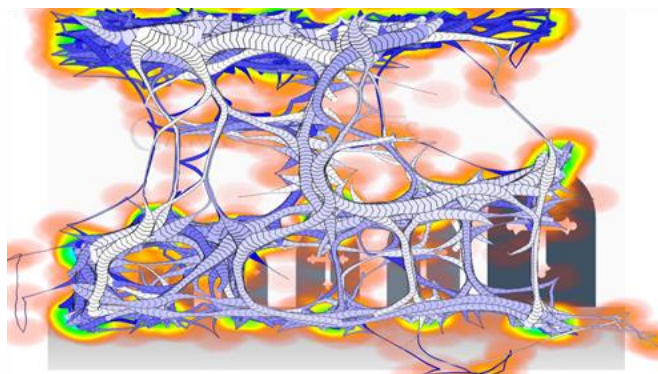
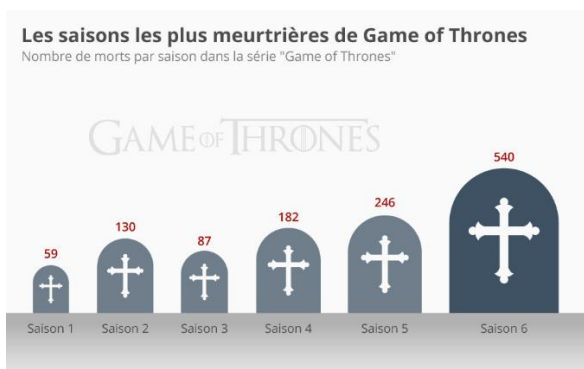


Figure 79 Visualisation n°11 (haut) et 31 (bas), version originale (gauche) et soumise au bundling (droite)

2.2. De la lecture

Ensuite, le *bundling* nous permet de prendre conscience de trois éléments très importants à propos de la lecture, si la visualisation de données en question comporte un titre⁴⁸ :

- Elle commence toujours par le titre ;
- Elle se termine en général par le titre ;
- S'il y a du texte en bas de l'image, la lecture se termine plutôt par le texte et non pas le titre.

Sur la Figure 80, qui présente le *bundling* de la visualisation n°21, nous constatons que les traces bleues les plus foncées se trouvent sur le titre. Les traces blanches, signifiant la fin

⁴⁸ Il n'y a pas de titre sur la [visualisation n°5](#) du corpus, par exemple

du trajet oculaire, reviennent sur le titre également. Il y a donc un retour au titre. Ce schéma se produit souvent lorsque la visualisation de données comporte un titre.

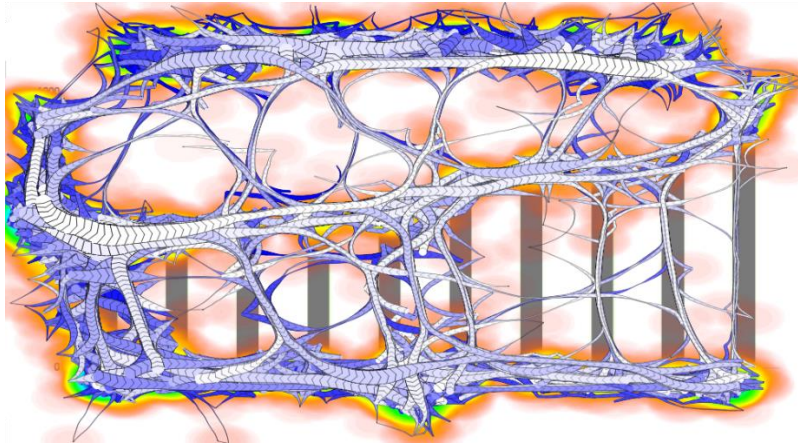


Figure 80 Bundling appliqué à la visualisation n°21

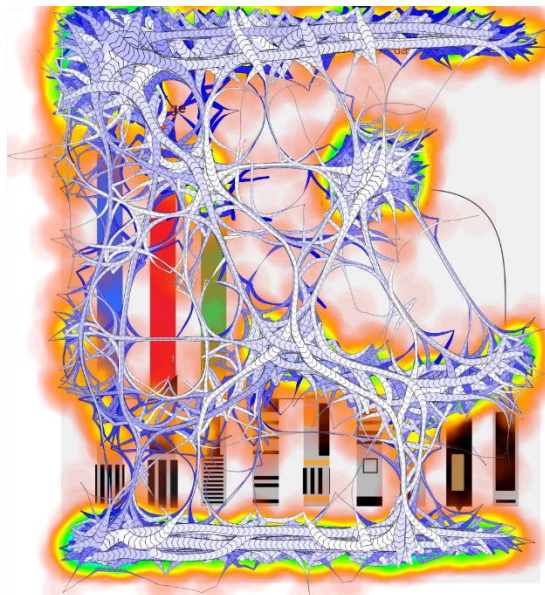


Figure 81 Bundling appliqué à la visualisation n°12

La Figure 81 représente le *bundling* appliqué à la visualisation n°12 du corpus. Une phrase explicative se trouve dans le bas de l'image. Sur cette image, nous constatons que la lecture commence à partir du titre et que les traces les plus claires se trouvent cette fois sur le texte au bas de l'image. La lecture ne se termine pas sur le titre cette fois-ci, mais sur le texte en pied de page. Ce phénomène est systématique : nous l'avons constaté de manière récurrente sur toutes visualisations soumises au *bundling*. Le fait que la visualisation soit standard ou ornementée n'y change strictement rien.

2.3. Les icônes apposées

Ensuite, nous avons observé que les icônes apposées à un graphique traditionnel dans une visualisation ornementée ne font pas systématiquement partie du chemin de lecture ou le sont, mais dans une moindre mesure : l'attention portée à ces éléments est très faible. Une fois le *bundling* appliqué, il arrive que les faisceaux passent sur une icône apposée parce qu'ils suivent une trajectoire spécifique (par exemple, la bouteille de la visualisation n°17 sur la Figure 82). Cela ne signifie pas forcément que l'œil s'est arrêté sur ces icônes apposées. Grâce à la carte de densité des fixations, il ressort de nouveau que les ornements apposés sont un élément sur lequel les participants ont passé peu de temps. Précédemment, l'exploration des zones d'intérêt s'était révélée utile en montrant également que le temps passé sur les icônes apposées était minime. On sait que l'individu voit un « flou » au cours de la saccade (Rayner, 1998) et on sait également que les participants ont qualifié ces icônes apposées d'éléments de contexte. Il arrive également que les faisceaux se dirigent vers ces ornements apposés. Sur la Figure 82 et la Figure 83, il s'agit du verre de vin, du passeport et de la serviette de plage. Ces faisceaux sont très étroits : leur densité est donc faible. Nous pouvons donc affirmer que les participants regardent ces éléments, mais en y accordant peu d'importance.

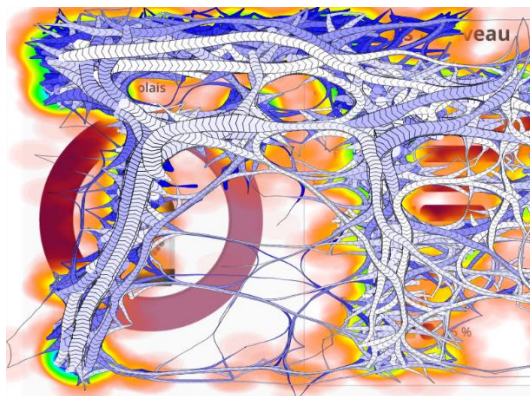


Figure 82 Bundling appliqué à la visualisation n°17

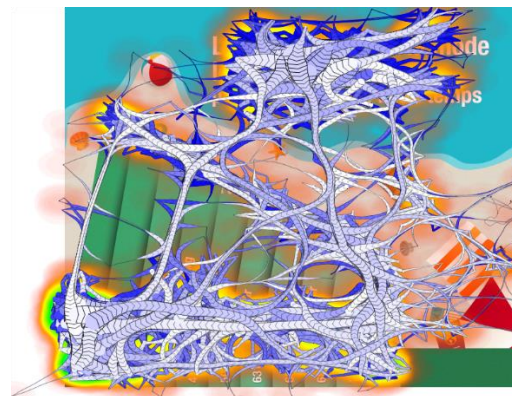


Figure 83 Bundling appliqué à la visualisation n°1

Le verbatim suivant résume bien l'idée de « voir sans regarder » les ornements qui englobent les visualisations de données. Les ornements n'ont donc pas d'importance visuelle mais sont décrits comme bénéfiques par les participants.

« La bouteille est tellement grande que je ne sais même pas si j'ai vraiment posé les yeux dessus parce qu'on la voit tout de suite, en englobant les images, donc on sait tout de suite de quoi on parle. Tant mieux, cela ça nous permet d'être dans le vif du sujet »

2.4.L'ornementation sur le *data-ink* et la place des étiquettes

Le *bundling* permet d'apporter une critique sur l'ornementation qui se constitue sur le *data-ink*. Les étiquettes de données sont très importantes : mal disposées, elles ne sont pas la priorité pour l'œil du participant qui se fixera plutôt sur l'ornementation constituée en *data-ink* et non sur les étiquettes de données. Les étiquettes sont petites et peu lisibles. Sur la visualisation n°18 (Figure 84, gauche), on voit que les étiquettes de données s'affichent sous les bâtonnets ornementés du diagramme. De plus, cette ornementation qui se trouve sur les traits relatifs aux données est assez complexe visuellement : les ornements utilisés ne sont pas du tout épurés. Les étiquettes sont petites et peu lisibles. Sur la version standard de cette visualisation (Figure 84, droite), les étiquettes de données sont directement disposées sur le *data-ink*. Leur position est importante et centrale. Appliqué à ces deux visualisations, le *bundling* révèle un fait important.

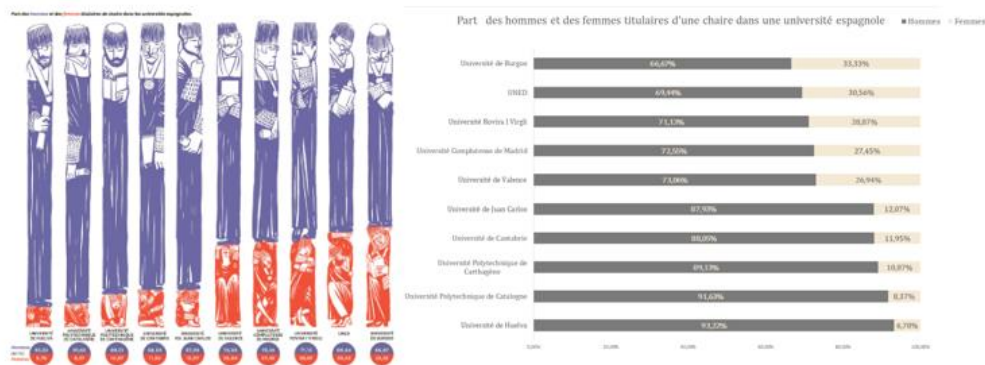


Figure 84 Visualisation n°18 du corpus, suivie de son homologue n°38

Sur la Figure 85 et la Figure 86, nous constatons ainsi que la lecture commence par le titre dans les deux cas. Sur la version standard, l'œil s'arrête sur la légende et, ensuite, passe sur les noms des universités et puis, successivement, sur les étiquettes de données de façon verticale. Toutes les informations importantes sont lues par les participants qui, ensuite, reproduisent le même comportement avant de revenir au titre. Sur la visualisation ornementée, nous constatons qu'après la lecture du titre, l'œil s'arrête sur les ornements (visages masculins) et puis sur la limite de la jauge montrant la différence entre masculin et féminin (mauve vs. rouge). La carte de densité de fixations montre d'ailleurs que les points de fixation sont nombreux en ces endroits. Ensuite, l'œil lit les différentes étiquettes relatives aux universités et les données qui sont situés au bas de l'image. Les participants reproduisent ce schéma de lecture une fois avant de revenir au titre. Le *bundling* montre clairement que dans le cas où les étiquettes de données sont présentées à la fin de la lecture (et donc au bas de l'image dans le cas de cette visualisation ornementée), le lecteur cherche un autre moyen d'évaluer les proportions au cours de sa lecture. Ainsi, de nombreux points

de fixation ont été réalisés sur la jauge dans le cas de la visualisation ornementée, mais pas dans le cas de la visualisation standard. De même, les propos récoltés dans les entretiens semi-directifs nous permettent de souligner que les participants avaient la sensation que les données importaient peu dans la visualisation ornementée.

T : Je suis en train de regarder sur le replay si tu regardes les chiffres.

A : Non je ne les regarde pas, j'ai juste regardé la proportion comme ça.

T : ça suffit alors ?

A : Ouais, là ça va mais à condition que le schéma ne soit pas erroné, ça suffit, il n'y a pas besoin de pourcentages !

T : Vous n'avez pas trop regardé les chiffres. Est-ce que c'est parce que ce n'est pas nécessaire ?

C : Heu, je ne connais pas les universités et... Là, dans un graphique comme ça, j'ai tout de suite compris et comme ça ne m'intéressait pas d'avoir les chiffres exacts...

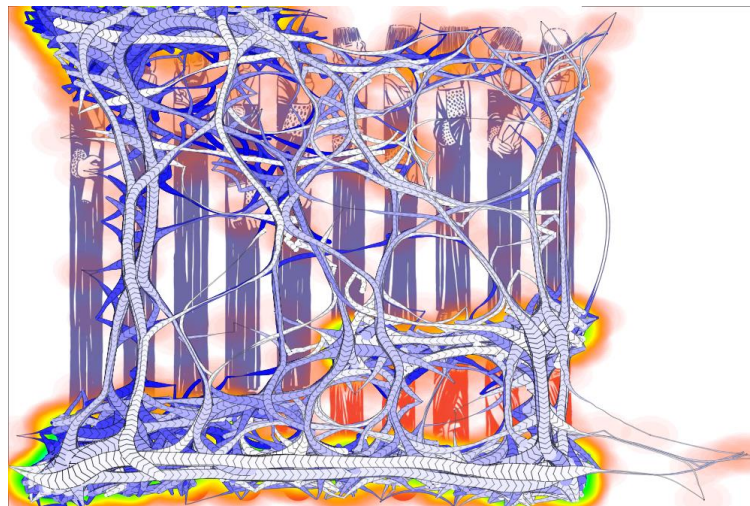


Figure 85 Bundling appliqué à la visualisation n°18 du corpus

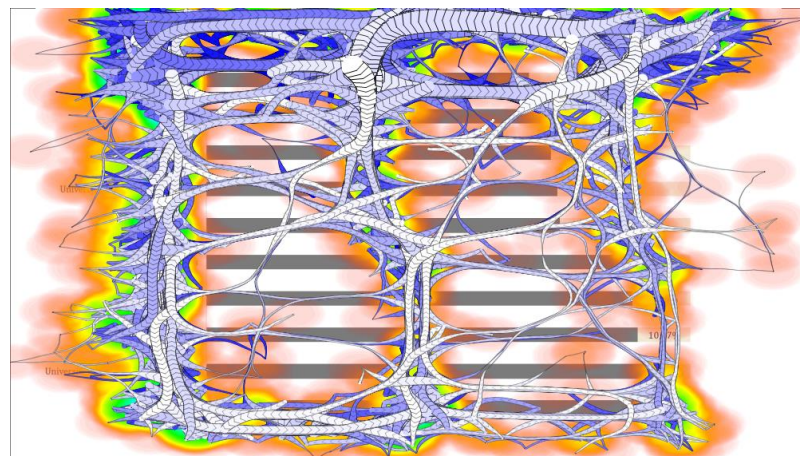


Figure 86 Bundling appliqué à la visualisation n°38 du corpus

Puisqu'il n'y a pas d'ornementation sur la visualisation standard, les participants regardent les données exactes avec plus d'attention. La visualisation n°18 a reçu des louanges de la part des participants qui voyaient en elle des propriétés attractives et un message émotionnel supplémentaire inscrit dans l'intention ornementale. Malgré tout, on voit que l'ornementation et surtout cette disposition des étiquettes n'incite pas le lecteur à lire les données exactes. En sachant que les participants ont passé davantage de temps à regarder l'ornementation que les chiffres exacts et en sachant que le nombre de fixations, la durée des fixations et le temps de vision sont plus élevés pour la visualisation ornementée que pour la version standard⁴⁹, nous déduisons que nous sommes face à une visualisation engageante avec une faible incitation à prendre connaissance des chiffres exacts.

2.5. La maximisation du data-ink ratio... ou de l'ornementation épurée

Nous pouvons ensuite observer la différence de *patterns* entre les visualisations qui sont visuellement simples (soit standard, soit ornementées mais épurées) et les visualisations dont l'apparence est chargée. Évidemment, cela est logique : moins il y a d'éléments à regarder, moins les *patterns* seront chargés. Néanmoins, cela indique l'importance de ne laisser au lecteur que peu de possibilités d'interprétation au lecteur. Ainsi, la Figure 87 et la Figure 88 montrent les mêmes données, ornementées ou non. Le visuel est simple.

Par contre, la Figure 89, qui montre le *bundling* appliqué à la visualisation n°9 du corpus, est très chargée : deux diagrammes, titres, sous-titres, textes explicatifs, pictogrammes, ... De nombreuses possibilités de chemins de lecture et d'interprétation sont laissées au lecteur. La simplification du graphique induirait sans aucun doute la simplification du *pattern* de lecture. Le respect du principe de maximisation *data-ink ratio* et la simplicité des visualisations ornementées ont un effet sur le *bundling* : plus l'image est simple et épurée, plus les *patterns* de lecture sont visibles.

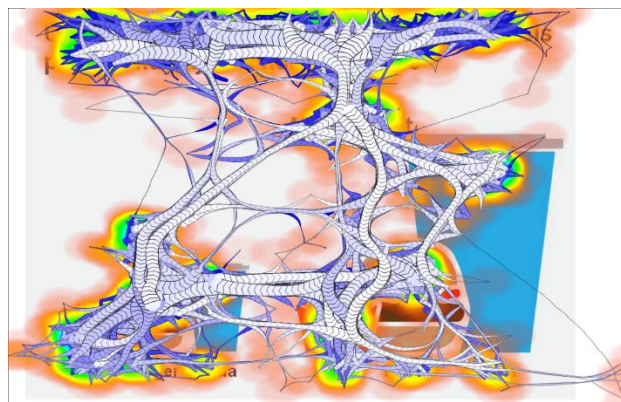


Figure 87 Bundling appliqué à la visualisation n°25 du corpus

⁴⁹ Standard vs. ornementé :

Nombre de fixations 2065 vs. 2540.

Durée des fixations

0,18 vs. 0,23s. Temps de vision 26 vs. 32s.

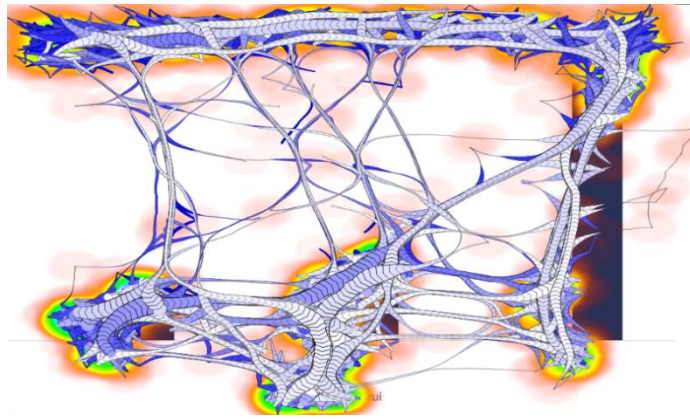


Figure 88 Bundling appliqué à la visualisation n°33 du corpus

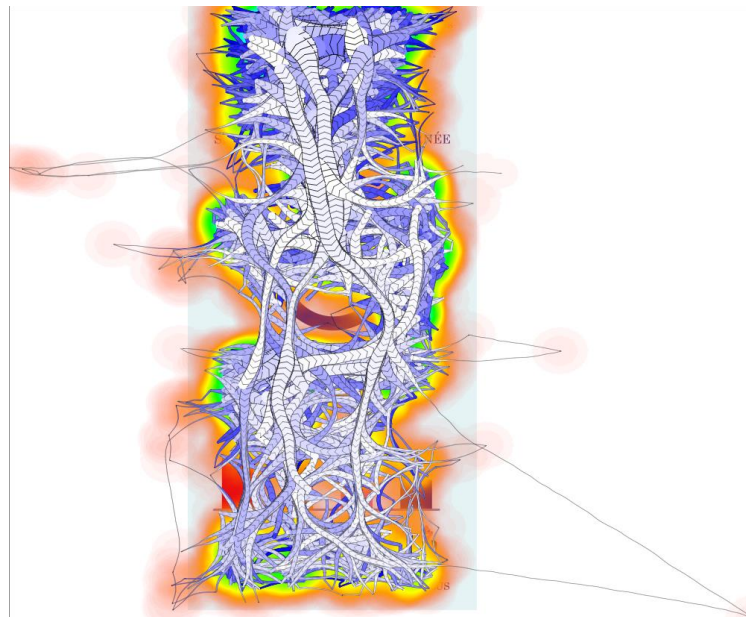


Figure 89 Bundling appliqué à la visualisation n°9 du corpus

En fait, j'étais perdu il y a beaucoup d'informations sur la même image (...). Ça me semble long pour le regard et là on voit bien que je suis perdu, je me demandais où est-ce que je devais regarder, je me recentre sur le titre, puis je redescends... elle m'a mis en "difficulté" cette image. Elle est grande, il y a beaucoup de texte au-dessus.

2.6. Les formes inhabituelles

Au contraire de l'observation précédente, nous observons sans difficulté que les visualisations de données à la forme inhabituelle (attribut : ornement) n'engendrent pas l'émergence d'un *pattern* de lecture clair. Il est difficile de dire dans quel sens se rendent les faisceaux : les arrêts, fixations et changements de direction sont très nombreux. En somme, nous avons ici visuellement la preuve que la complexité de ces visualisations, rapportée dans les appréciations déclarées et dans les entretiens semi-directifs, se vérifie

aussi physiologiquement. L'œil ne suit aucunement une structure claire pour identifier et comprendre l'information : le *bundling* illustre ce phénomène. Utiliser une forme graphique non habituelle n'encourage pas une lecture rapide et claire.

Celui où on voit tous des points de toutes les couleurs... C'est le pire parce qu'on ne comprend vraiment rien, les couleurs sont les unes à côté des autres... Mais celui-là comme toutes les couleurs sont à côté des autres et similaires, entre les légumes et les fruits, on ne sait même pas distinguer lequel est lequel. Je l'ai vraiment trouvé ignoble (Figure 90).

Il n'y a rien qui te dit dans l'image où regarder en premier donc fatalement tu vas un peu partout (Figure 91).

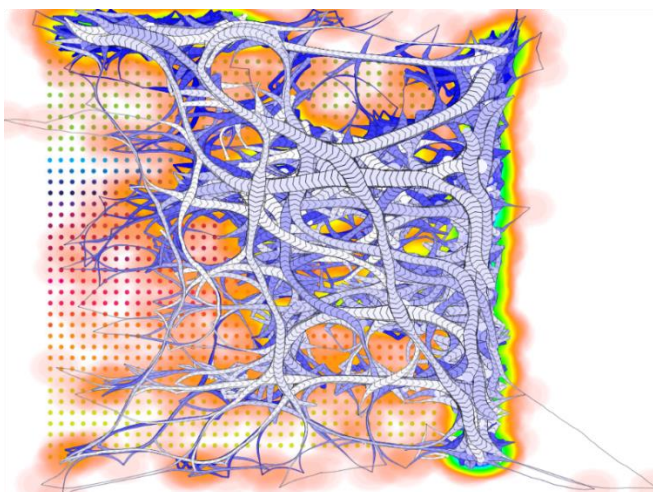


Figure 90 Bundling appliqué à la visualisation n°4

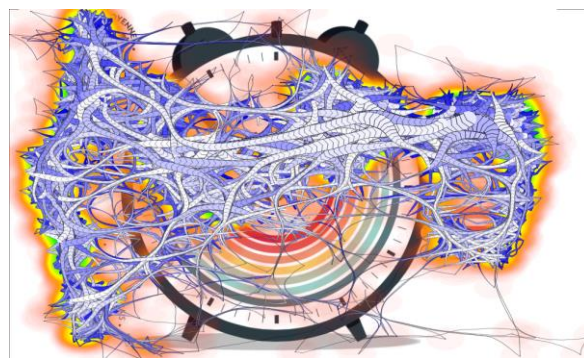


Figure 91 Bundling appliqué à la visualisation n°2

2.7.Ce qui influence le *pattern*

2.7.1. La légende et les axes

La disposition de la légende et l'affichage des axes dans le cas d'un diagramme à deux dimensions influence les différents *patterns* de lecture. Ces éléments ont d'ailleurs été mentionnés dans les entretiens semi-directifs : les participants réclamaient parfois un « ordre de lecture » logique en indiquant qu'une légende directement située sous le titre ou à côté de ce dernier simplifierait la tâche. Effectivement, on constate sur les chemins de lecture de visualisations affichant des légendes que le *pattern* est plus éparé et plus dense en aller-retours lorsque la légende n'est pas sous le titre. Sur la Figure 92, la légende est directement sous le titre. Les trajets visuels sont très clairs et peu complexes. Par contre sur la Figure 93, on constate davantage d'aller-retours qui alourdissent la lecture. La

disposition de la légende n'est normalement pas une question d'ornementation. Néanmoins, il arrive fréquemment que les concepteurs en choisissent la place à des fins esthétiques. Or, l'endroit fluidifiant le plus le comportement de lecture est soit sous le titre, soit à côté de celui-ci.

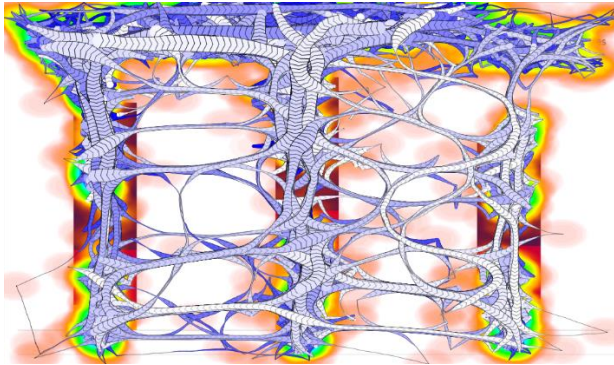


Figure 92 Bundling appliqué à la visualisation n°34 du corpus

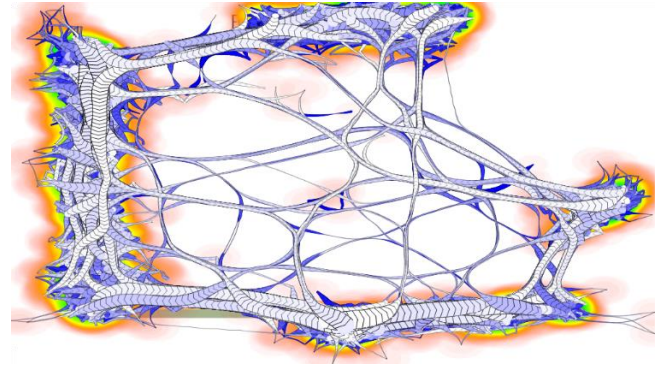


Figure 93 Bundling appliqué à la visualisation n°30 du corpus

Moi j'aime bien quand la légende vient d'abord, comme ça, après, je sais regarder le graphique et puis dire à quoi ça correspond. Je ne dois pas chercher l'information...Au début, ça ne m'a pas tapé aux yeux et puis c'est ce que je me suis dit au fur et à mesure de l'expérience. C'est mieux quand la légende est au début.

De plus, nous avons également constaté que la suppression des axes et l'affichage des étiquettes de données en bout du bâtonnet, qui est une application directe du principe de maximisation du *data-ink*, est complètement pertinente : les *pattern* de lecture sont plus clairs, plus précis et plus aérés, ce qui signifie une grande cohérence entre les comportements visuels des lecteurs. Lorsqu'un axe y est présent, l'œil y revient sans cesse. Lorsqu'il est absent et que les étiquettes de données sont utilisées en alternative, la lecture est plus rapide et plus fluide. Cette facilité a également été soulignée oralement :

C'est étonnant il y a d'habitude un axe sur le côté et là c'est en fait beaucoup plus clair d'avoir le nombre directement au-dessus de la barre.

2.7.2. L'orientation (visualisations standard, diagrammes en bâtonnets)

Lors de la conception de notre scénario d'expérimentation, nous avons choisi de réaliser une grande partie des visualisations standard en diagrammes en bâtonnets horizontaux, pour une facilité de lecture des étiquettes de données plus accrue. Au cours des entretiens semi-directifs, les participants ont admis qu'ils préféreraient les diagrammes verticaux plutôt qu'horizontaux, pour une question d'habitude. Cette préférence se vérifie dans les *patterns* de lecture : ils semblent plus clairs et épurés sur les visualisations standard dont le diagramme est vertical et non horizontal, ce qui signifierait que la lecture est plus fluide,

rapide et correspond davantage aux attentes et habitudes des participants (Figure 94). Même si, en théorie, dans certains cas, l'utilisation de diagrammes horizontaux se justifie, les habitudes et préférences des participants entrent en jeu et peuvent influencer leur comportement visuel, comme l'indiquent les facteurs humains et émotionnels de Kennedy et Hill (2017, 2018).

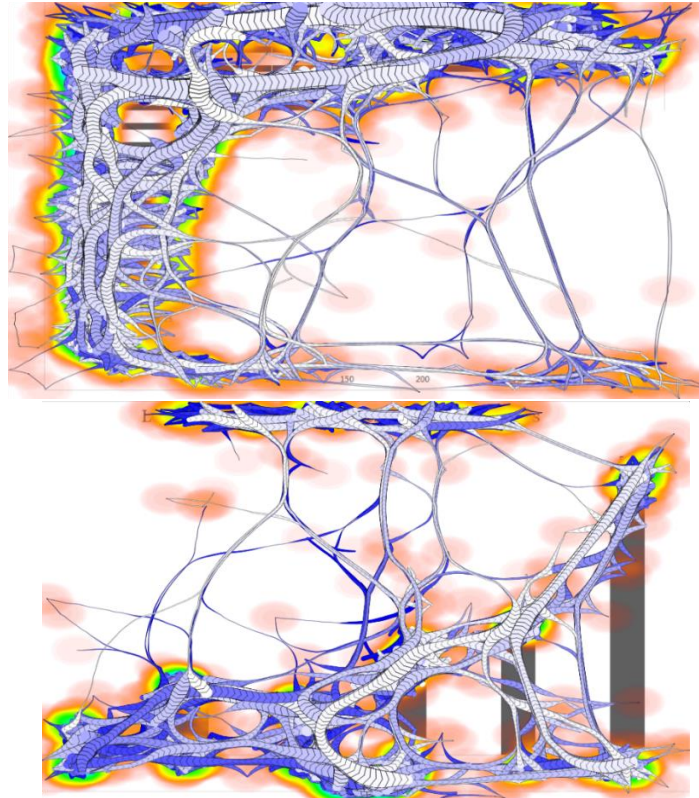


Figure 94 Bundling des visualisations n°24 (gauche) et n°31 (droite)

« J'avoue que je préfère les barres verticales. Parce qu'en général tu lis le titre, tu descends et tu as toutes les extrémités d'un coup et voir qui est le plus grand, le plus petit et analyser à ton aise quoi. Bizarrement, ça a toujours été utile pour mieux comprendre et te rendre compte du message à faire passer »

2.7.3. Le texte et les chiffres

Pour terminer, nous nous sommes focalisée sur le texte et les chiffres (ou étiquettes de données). En réalité, il s'agit des éléments principaux qui sont regardés sur une visualisation de données. La même chose se produit d'ailleurs lorsque l'ornementation se concentre sur le *data-ink*.

Plus particulièrement pour les visualisations dont l'ornementation se constitue sur le *data-ink*, c'est l'annotation et l'étiquetage de ces visualisations de données qui influence

l'émergence du *pattern* de lecture et non l'ornementation en soi. Cela peut être aisément constaté sur la Figure 95.

Dans l'ensemble, c'est ce qu'exprime le *bundling* : si les informations les plus importantes sont saillantes (contrairement à la visualisation n°18 du corpus), le trajet oculaire passe principalement par les éléments textuels et les chiffres (titres, légendes, nombres, etc.). Ce sont donc les éléments relatifs à la grammaire graphique qui guident le *pattern* de lecture, avec une rare influence de l'ornementation sur les points de fixation.

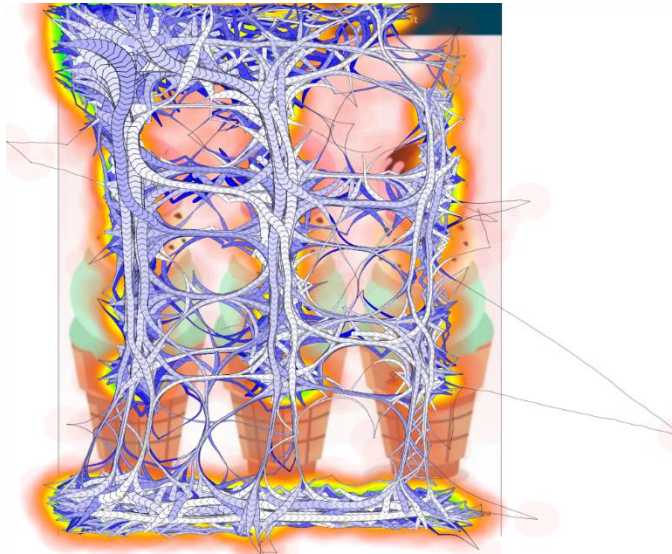


Figure 95 Bundling appliqué à la visualisation n°14 du corpus

Ce constat rejoint d'ailleurs ce qu'avaient déjà avancé Carpenter et Shah (1998) et Borkin et ses pairs (2013) : selon ces chercheurs, les individus passent plus de temps à lire les étiquettes et le texte plutôt qu'à regarder le graphique en soi et son ornementation. Nous pouvons ainsi ajouter que dans de bonnes conditions, l'ornementation ne perturbe pas (ou rarement) ce constat.

3. Interprétation des observations

Les différentes observations rendues possibles par le *bundling* permettent de tirer des conclusions importantes pour le domaine. Effectivement, à en croire Tufte (1983), l'ornementation qu'il appelle plutôt *chartjunk* est distrayante et négative. On pourrait dès lors penser que les trajets oculaires se focalisent énormément sur l'ornementation ou que les ornements sont les premiers éléments observés. Or, il n'en est rien. Évidemment, le propos est à nuancer en fonction des différents types d'ornementation présentés à l'utilisateur :

- Les formes inhabituelles sont, étant donné leur nouveauté et leur fantaisie, difficiles à saisir et d'une grande complexité visuelle : il est difficile d'identifier un chemin de lecture clair et précis. Ainsi, avec un objectif de communication rapide et exact, ces formes d'ornementation sont à éviter.
- Les icônes apposées : elles sont observées dans une faible mesure et font parfois rapidement partie du chemin de lecture en fonction de leur disposition. Elles ne sont pas réellement distrayantes et ne perturbent pas la lecture du *data-ink* qui, lorsqu'il constitue un diagramme traditionnel, suit le même *pattern* de lecture que les visualisations de type standard.
- Les icônes ou l'ornementation qui se constitue en *data-ink* : ce type d'ornementation est à prendre avec des pincettes. Le *bundling* montre clairement que les participants regardent principalement l'information chiffrée et textuelle, et finalement très peu l'ornementation. Néanmoins, présenter l'ornementation sur le *data-ink* est un double risque :
 - En transformant les formes, l'œil perçoit difficilement les aires représentées par les icônes et pictogrammes. Le risque d'atteindre un *lie factor* distorsif est plus grand ;
 - L'agencement des étiquettes de données et du texte peut être réalisé pour des raisons esthétiques et finalement être disposé de façon inadéquate, ce qui peut avoir une influence sur le *pattern* de lecture. Dans le cas où ce type d'ornementation est privilégié, l'importance de la disposition du texte, des étiquettes, en bref, du corps de l'information est de haute importance.

Ces commentaires nous poussent donc à confirmer la recommandation que nous formulions précédemment : le type d'ornementation le plus sûr quant à la bonne réception des données est l'ornementation apposée, qui est non distrayante et qui permet ainsi de bénéficier des forces de l'ornementation (contextualisation, engagement, mémorisation, plaisir, ...) sans perturber trop grandement la réception de l'information. La disposition de l'information sur le *data-ink*, lorsqu'elle se cantonne à un visuel très simple et épuré, ne montre pas de points négatifs. Elle reste néanmoins un pari risqué du côté de la conception, où l'intention se porte plutôt sur l'esthétique que sur la correcte représentation des données. Dans le cas du choix de ce type d'ornementation, il est primordial d'être très prudent quant à la disposition des étiquettes de données, du texte et de la présentation de l'information en général, au risque d'engendrer des chemins de lecture complexes.

Mais encore, le *bundling* montre tout le bénéfice des visuels simples et épurés, que la visualisation soit de type standard ou ornementé. Ainsi, il nous pousse à conseiller de

privilégier, toujours, les apparences les plus simples possibles, même lorsque l'on utilise l'ornementation. Les *patterns* de lecture sont plus clairs et distincts lorsque le visuel est simple. En connaissance de cause, il vaut mieux donc produire des visuels qui, eux aussi, seront d'une grande simplicité visuelle.

En ce qui concerne la simplicité, le *bundling* a également montré le bénéfice de la maximisation du *data-ink ratio* lorsqu'il s'agit de supprimer l'axe y et d'utiliser des étiquettes de données. Cela contribue à l'élaboration d'un *pattern* de lecture épuré. En cohérence avec nos propos précédents, la maximisation *data-ink ratio* est un grand bénéfice en ce qui concerne la présentation des données en soi. Cependant, ce principe pourrait tout de même tolérer l'ajout d'une icône simple en apposition afin de profiter des bénéfices de l'ornementation.

Finalement, le *bundling* montre visuellement les concepts de Körner (*3-stage model*, 2011), de Carpenter & Shah (1998) et de Borkin et ses pairs (2013) pour qui l'attention se porte de toute façon sur le texte et les nombres. Nous pouvons d'ailleurs y ajouter qu'une ornementation maîtrisée en fonction de nos commentaires ne perturbe pas l'application de ces concepts.

Conclusion L'analyse des *gaze data*: entre complexité, engagement et chemins de lecture

Les dernières analyses menées dans le cadre de cette thèse permettent de tirer des conclusions originales et inédites.

L'analyse statistique des fixations et l'exploration des zones d'intérêt ont montré que l'ornementation augmente systématiquement le nombre et la durée des fixations. Lorsque la visualisation est simple, cela peut être considéré comme de l'engagement, tandis que les formes inhabituelles mènent vers une difficulté de lecture et une complexité évidente. L'exploration des zones d'intérêt a, quant à elle, aisément montré que l'ornementation apposée ne distrait pas le lecteur.

Ensuite, le développement de la technique de visualisation appelée *bundling* pourrait, à l'avenir, se positionner comme un réel outil d'analyse de la réception de visualisations de données, et peut-être même d'autres dispositifs visuels. Au-delà des résultats apportés, notre recherche a permis de montrer que le *bundling* peut s'imposer comme un nouvel outil, pertinent dans le domaine *Infovis*. Quoi qu'il en soit, cette technique a rendu possible l'identification de *patterns* de lecture, de les comparer et d'en tirer des conclusions qui sont complètement en phase avec les analyses précédemment réalisées dans le cadre de cette thèse, tout en leur apportant de la consistance. En effet, tout comme les autres explorations de cette thèse, l'utilisation du *bundling* nous mène à nouveau à penser que la simplicité est la clé en visualisations de données et que même l'ornementation n'échappe pas à la règle. Ainsi, nous avançons que l'ornementation n'est pas à bannir, mais à utiliser sous certaines conditions, et surtout, de la manière la plus épurée possible, comme nous l'avons montré à différents instants de ce travail de recherche.

L'influence de l'ornementation des visualisations de données sur la construction de sens

Conclusion

1. Contexte de la recherche

Les recherches abordant l'ornementation évoluent beaucoup dans le domaine de l'*Infovis*, bien que la définition du terme « *embellishment* » soit floue dans la littérature anglo-saxonne. C'est pourquoi nous l'avons définie plus précisément lors de cette recherche, en décidant de nommer le phénomène « ornementation » en français. Au cours de cette thèse de doctorat, nous avons souhaité l'étudier à partir d'une perspective communicationnelle et du point de vue des sciences humaines. Si de telles démarches existent déjà, elles restent rares en regard des avancées sur le sujet dans la communauté *Infovis*.

Malgré les mises en garde de Tufte à propos des dangers du *chartjunk*, l'ornementation a montré son lot de bénéfices sur l'engagement et sur la mémorisation à long terme et à court terme. Aux prémices de notre travail de recherche, la question de l'ornementation gagnait un intérêt grandissant dans le domaine *Infovis*. Elle est aujourd'hui une thématique cruciale pour la communauté, comme en témoignent les travaux de Cui et ses collègues (2019) et Coelho et Mueller (2020) qui souhaitent proposer la création automatique de visualisations de données ornementées. Ainsi, avec les années, la façon d'aborder l'ornementation dans le domaine a changé : actuellement, elle n'est plus considérée comme un danger pour la correcte communication de l'information, bien que la plupart des chercheurs du domaine soient conscients de ce qu'elle peut générer une fois mal utilisée. Au contraire, les professionnels, eux, n'ont pas réellement de référentiel leur permettant d'éviter les erreurs en termes de conception de visualisation de données, du moins, dans le monde francophone. Pour les chercheurs, l'ornementation est devenue une pratique à étudier parmi tant d'autres. Il y a encore beaucoup de choses à apprendre sur le sujet, mais il est désormais possible de tenter de l'utiliser sans que ce ne soit un danger pour la compréhension et en mobilisant ses bénéfices. En revanche, dans le monde professionnel francophone, l'usage de l'ornementation des visualisations de données reste encore potentiellement un danger pour la compréhension des destinataires. Elle est largement utilisée. Nous nous sommes demandée ce qu'une telle chose induit sur le sens produit par les destinataires de ces visualisations de données ornementées. En tant que destinataires, nous avons choisi un

échantillon de personnes déjà sensibilisées à la façon dont on communique un message, notamment en ligne. Notre échantillon se compose pour rappel de 40 professionnels de la communication numérique.

Peu de chercheurs se sont attelés à comprendre comment l'ornementation peut influencer la construction de sens des destinataires des visualisations de données, en dépassant la simple étude de la compréhension de l'information. C'est pourtant ce que questionne notre problématique. Nous avons précédemment expliqué que nous concevions la construction de sens comme l'assemblage de différentes composantes qui ne cessent de s'influencer les unes les autres, à savoir la perception visuelle, la compréhension et la subjectivité menant aux comblements de discontinuité. Nous en étions donc venue à poser la question de recherche suivante :

Par rapport à une visualisation de données qui est correcte en regard des principes de conception (que nous appelons « visualisation standard »), qu'est-ce qu'une visualisation de données ornementée induit sur la construction de sens

(1) au niveau de la perception visuelle et de la compréhension ?

(2) au niveau du gap bridging ?

L'analyse et l'interprétation des résultats obtenus à l'issue de notre méthodologie expérimentale nous permettent désormais de répondre à cette question.

2. L'influence de l'ornementation sur la perception visuelle et l'impression de compréhension

L'analyse des appréciations déclarées ainsi que l'analyse des *gaze data* nous a permis de décrire très largement quelle est l'influence de l'ornementation sur la construction de sens des utilisateurs du point de vue de la perception visuelle et de la compréhension. Néanmoins, nous affirmons que notre recherche n'a pas réellement permis de décrire quelle est l'influence de l'ornementation sur la compréhension des données et du message *en soi*, mais plutôt sur la sensation ou sur l'impression de compréhension des participants au moment de l'expérience. Les informations qui sont à notre disposition, telles que les appréciations déclarées ou encore les transcriptions des entretiens semi-directifs, nous permettent de dire si les participants *pensent* avoir compris le message et non s'ils l'ont réellement compris⁵⁰. De même, l'ornementation mène parfois à des situations où les

⁵⁰ Pour pouvoir étudier la compréhension réelle, il aurait fallu que nous interrogeons les participants à propos des données contenues dans les visualisations, dans un questionnaire d'évaluation. Ce n'était néanmoins pas notre objectif.

utilisateurs *ne pensent pas* avoir compris le message, ce qui ne signifie pas pour autant qu'ils n'ont pas construit de sens.

En abordant la perception visuelle et la compréhension, nous nous sommes en réalité attelée à l'observation de phénomènes très spontanés, à savoir la vision et l'appréciation quasi-directe des visualisations, sans prise de recul. Les *gaze data* sont des données à caractère presque physiologique, tandis que les appréciations déclarées ont été complétées de manière très spontanée par les participants, nous donnant ainsi accès à l'impulsivité de la réception.

En ce qui concerne la **perception visuelle**, L'analyse des *gaze data* a montré que les visualisations de données ornementées récoltent un nombre de fixations et une durée moyenne de fixations plus élevées que les visualisations standard. Cela s'explique tantôt par une capacité à engager davantage les utilisateurs qui, par ailleurs, estiment ces visualisations plus belles que leurs homologues dans les appréciations déclarées. Les appréciations déclarées ont également montré que les participants jugeaient les visualisations ornementées moins claires que leurs homologues standard. Le nombre élevé de fixations peut également s'expliquer par une complexité visuelle plus élevée. Nous nous sommes ensuite penchée sur quelques caractéristiques de l'ornementation et nous avons découvert que le nombre de fixations est significativement plus élevé sur les visualisations proposant une ornementation dont la forme est inhabituelle (chargée, non familière, fantaisiste) en comparaison à l'ornementation qui consiste simplement en l'utilisation d'icônes (petites images ou pictogrammes familiers à l'œil humain). Ces formes inhabituelles sont donc bien plus complexes à décoder pour l'œil humain. La nouveauté de ce type de visualisation mais aussi les mauvais choix de conception graphique (type de graphique, disposition des informations, affichage du texte et des chiffres, ...) peuvent expliquer cette complexité évidente. Puis, nos analyses nous poussent à préconiser une ornementation constituée en icônes apposées à côté d'un graphique traditionnel : les appréciations déclarées montrent bien que les participants jugent les visualisations ornementées de cette façon plus claires que les visualisations dont l'ornementation se concentrait sur les traits relatifs aux données (*data-ink*). Néanmoins, l'analyse des *gaze data* a montré qu'il n'y a pas beaucoup plus de points de fixation sur l'un ou l'autre type d'ornementation. Ainsi, l'ornementation constituée en *data-ink* n'est pas forcément plus complexe visuellement que l'ornementation dont l'icône est apposée à un graphique traditionnel, même si les participants indiquent une sensation de clarté moins marquée. Rappelons tout de même la difficulté plus grande de comparer des aires entre elles (Cleveland, McGill, 1984), ce qui est souvent le cas pour l'ornementation en *data-ink*.

Mais encore, nous avons également pu montrer que visuellement, une visualisation au *lie factor* distorsif n'est pas plus complexe ou plus engageante qu'une visualisation dont le *lie factor* est fidèle. Cela montre ainsi que l'œil se laisse abuser par les effets de distorsion de l'information. Par contre, le temps de vision des visualisations dont le *data-ink* est bas est bien plus élevé que le temps de vision des visualisations dont le *data-ink* est moyen ou élevé. Il n'y a d'ailleurs pas de différence significative entre ces deux derniers, montrant ainsi qu'une ornementation constituée en *non data-ink* (donc apposée) n'apporte pas réellement de distraction.

Pour nous en assurer, nous avons étudié plus en détail la distraction que serait susceptible d'apporter l'ornementation apposée. La première fixation sur ces icônes apposées se réalise généralement à un tiers du temps de vision et le nombre de fixations est très faible sur ces éléments. Les icônes apposées ne captent pas rapidement l'attention et ne sont visuellement pas considérées comme importantes, ce qui correspond à la remarque souvent formulée au cours des entretiens semi-directifs par les participants qui estimaient avoir vu l'ornementation sans pour autant passer du temps à la regarder.

Pour terminer, nous avons pu commenter l'influence de l'ornementation sur les chemins de lecture réalisés par les participants et nous avons pu réaliser plusieurs observations à l'utilisation du *bundling*, méthode de visualisation des *gaze data* développée à l'ENAC. Le *bundling* a d'ailleurs pu confirmer certaines observations réalisées auparavant dans le cadre d'autres études (Carpenter et Shah, 1998 ; Körner, 2011 ; Borkin et al., 2016). Il a principalement montré que la lecture d'une visualisation se réalise de manière cyclique et que l'ornementation n'influence pas ou peu ces cycles, excepté lorsque la visualisation ornementée est une forme inhabituelle. Ainsi, des *patterns* de lecture très clairs émergent le plus souvent, sauf quand il s'agit de formes inhabituelles. De même, ce sont les informations écrites (texte, nombres, étiquettes, ...) et leur disposition dans l'espace qui structurent les *patterns* de lecture. Il s'agit d'ailleurs du principal élément observé sur une visualisation, comme déjà mis au jour par Carpenter et Shah (1998) ainsi que Borkin et ses collègues (2013). Nous avons également observé que les faisceaux passant sur les icônes apposées sont assez minces, ce qui converge avec notre affirmation précédente : les icônes apposées ne sont pas réellement une distraction. Ce que montre le *bundling* en réalité, c'est que ce n'est pas tant l'ornementation en tant que telle qui pousse à l'incompréhension de l'image. En fait, les visualisations ornementées sont parfois moins soignées en ce qui concerne la disposition du texte (légendes, étiquettes, nombres, ...). Nous avons ainsi montré que les participants préfèrent voir la légende apparaître sous le titre. Cela a été soulevé lors des entretiens semi-directifs et les *patterns* de lecture suivant ce modèle sont plus simples. De même, nous avons montré l'importance d'afficher les étiquettes de

données le plus tôt possible dans l'ordre de lecture, surtout lorsque la visualisation est ornementée. Le *bundling* montre également que les *patterns* de lecture qui émergent sur les visualisations dont l'ornementation est simple et épurée sont eux aussi simples et très clairs, tout comme ceux des visualisations standard. Cela pousse à penser qu'une ornementation épurée est à conseiller.

Finalement, nous pouvons apporter un commentaire clair et concis à propos de l'influence de l'ornementation sur la construction de sens des participants en ce qui concerne la perception visuelle. Notre recherche montre qu'il existe différents types d'ornementation qui ne sont pas forcément tous perçus visuellement de la même manière. En comparaison avec des visualisations de type standard, il existe des formes inhabituelles, qui ne permettent pas l'émergence d'un *pattern* de lecture clair et qui présentent un nombre beaucoup plus important de points de fixation, avec une durée moyenne de fixations importante également. L'étude visuelle de ce type d'ornementation montre qu'il complexifie l'information. Visuellement, l'œil s'y perd : il y a peu de structure dans les trajets oculaires réalisés au moment de la lecture des visualisations qui comportent ce type d'ornementation. Par ailleurs, nous avons constaté que lorsque l'ornementation est très simple et épurée, des *patterns* de lecture simples et similaires à ceux des visualisations standard émergent. Dans l'analyse du nombre de points de fixation, nous avons constaté qu'ils étaient plus nombreux pour les visualisations de données de type ornementé ce qui, dans ce cas en particulier, se traduit par de l'engagement. En effet, nous avons montré que l'ornementation, lorsqu'elle est une icône apposée, n'est pas une distraction. Lorsque l'ornementation se constitue en *data-ink*, le participant lit de toute façon les étiquettes de données sans observer les spécificités de l'ornementation. Notre constat est celui-là : sur les visualisations ornementées, les concepteurs prennent des libertés à propos de la disposition d'éléments essentiels comme la légende, les étiquettes de données ou encore des explications supplémentaires. Les éléments picturaux et ayant pour objectif d'ornementer sont finalement assez peu « regardés » par les participants. S'il ne s'agit pas de formes inhabituelles, et que les icônes choisies sont relativement simples, les ornements sont peu complexes, peu importants pour l'œil et ne peuvent pas être considérés comme une distraction. Ce sont les éléments textuels comme le titre, la légende, les nombres, les étiquettes de données ou le texte explicatif qui sont réellement regardés par les participants. Si l'ornementation peut avoir une influence sur la disposition de ces éléments, elle ne perturbe pas la réception de l'information lorsqu'on la considère pour ses propriétés picturales uniquement.

Poursuivons désormais avec nos conclusions sur l'influence de l'ornementation des visualisations de données sur **l'impression de compréhension** des utilisateurs. Comme

nous l'avons précédemment expliqué, nous ne prétendons pas avoir accès à la compréhension des utilisateurs mais plutôt à leur propre sensation ou impression de clarté ou de compréhension.

Les appréciations déclarées nous ont permis de tirer un constat clair et indiscutable : systématiquement, les visualisations de type standard semblent plus claires que les visualisations de type ornementé. Ainsi, on peut dire que l'ornementation diminue la sensation de clarté, même si elle est presque toujours préférée des participants (*item* « J'aime »), ce qui indique qu'elle est plus engageante que le minimalisme. Par contre, dans l'analyse des appréciations déclarées, nous avons noté que la présence d'ornementation ne diminue pas l'appréciation de compréhension. Cela signifie que les participants estiment presque toujours avoir compris l'information et cela semble logique en regard de notre design expérimental : nous n'avons imposé aucune limite de temps aux participants lorsqu'ils regardaient les visualisations de données. Ils pouvaient ainsi jouir du temps nécessaire pour lire l'entièreté de la visualisation et l'interpréter. Le fait que la sensation de clarté soit systématiquement plus basse pour les visualisations de type ornementé signifie qu'au premier coup d'œil ou, du moins, à la première lecture, l'information semble moins intelligible sur les visualisations ornementées. De nouveau, les participants ont indiqué que l'ornementation qui se constitue en icônes semble bien plus claire que celle qui consiste en des formes inhabituelles.

De même, l'utilisation du *bundling* nous a permis d'identifier que les *patterns* de lecture identifiés sur les visualisations ornementées ne sont pas forcément plus chargés ou compliqués (excepté en ce qui concerne l'ornementation qui consiste en des formes inhabituelles). De temps à autre, les ornements sont observés, dans une faible mesure. Mais en général, les *patterns* de lecture des visualisations ornementées ne sont pas plus complexes que ceux des visualisations de type standard. La compréhension est présente, même si le temps passé sur les visualisations ornementées est plus long, soit parce qu'elles sont plus complexes visuellement, soit parce qu'elles sont plus engageantes.

Dans les entretiens semi-directifs, les participants ont également noté plusieurs facteurs qui peuvent favoriser l'incompréhension. En réalité, peu de ces facteurs sont liés à l'ornementation, ce qui converge avec notre affirmation précédente : l'ornementation favorise la prise de liberté des concepteurs qui disposent mal l'information. Ainsi, par exemple, le manque d'informations, l'ordre de lecture peu intuitif, l'incompréhension du vocabulaire, un mauvais étiquetage ou encodage de données sont des éléments qui favorisent l'incompréhension. Les participants ont aussi noté qu'ils se rendaient compte que certaines visualisations avaient été pensées en fonction de leur apparence et non en

fonction de la communication des données. L'effort esthétique démesuré peut donc parfois devenir un facteur d'incompréhension.

La présence d'ornementation n'a pas influencé l'intérêt que portaient les participants aux visualisations de données. C'est bien le sujet qui les intéressait ou pas, et non l'apparence de la visualisation. Par contre, l'ornementation permet en général davantage d'engagement des utilisateurs envers la visualisation. L'analyse de l'*item* « J'aime », ainsi que des entretiens semi-directifs, montre que même si le sujet de la visualisation ne les intéresse pas, les participants auront davantage tendance à s'impliquer dans la lecture de celle-ci si elle est ornementée.

Certains points positifs de l'ornementation ont été soulignés, comme le fait que l'icône permette de se rappeler plus facilement du sujet de la visualisation, ce qui intervient en quelque sorte dans le processus de construction de sens. Les participants ont également noté que la surcharge visuelle peut être provoquée par l'ornementation et diminuer la sensation de clarté. De même, l'ornementation est inutile si les données sont mal présentées, puisque les participants estiment dès lors que la compréhension est moins aisée.

Pour conclure, nous pouvons dire qu'en fonction du type d'ornementation, la construction de sens, du point de vue de la compréhension, fluctue fortement. À condition de disposer d'un temps de lecture illimité et que l'ornementation soit simple, épurée, se constitue en icônes et que la plupart des règles de conception graphique soient bien respectées (correcte disposition du texte, nombres, légendes, ...), la compréhension, en tant que processus d'interprétation, ne semble pas impactée. Sans ces conditions, le processus de compréhension s'avère plus long et la clarté diminue. Quoi qu'il en soit, la sensation de clarté, elle, comprise en tant qu'intelligibilité au premier coup d'œil, sera toujours moindre pour les visualisations ornementées que pour les visualisations de type standard.

3. L'influence de l'ornementation sur la construction de sens individuelle

Au-delà de l'influence de l'ornementation des visualisations de données sur la perception visuelle et sur la sensation de compréhension, nous avons également souhaité étudier en quoi les caractéristiques personnelles mais aussi la subjectivité des participants entrait en compte dans la construction de sens.

Ainsi, nous avons constaté qu'à de nombreuses reprises, l'ornementation cause des discontinuités qui nécessitent d'être comblées, conformément à la *sense-making methodology* de Dervin (2003). Le plus souvent, les individus en arrivent à une situation

finale dans laquelle ils comprennent la visualisation, sans pour autant n'éprouver que des sensations positives. Les situations d'échec quant à la compréhension de la visualisation sont également nombreuses, menant les participants à ressentir des émotions négatives, ce qui ne signifie pas pour autant qu'aucun sens n'a été créé lors de la lecture de la visualisation. Au cours de l'analyse des données qualitatives, selon une adaptation du *verbing* de Dervin (2003), nous nous sommes aperçue qu'il existe autant de stratégies de comblement de discontinuité qu'il existe de discontinuités. Nous avons tenté de rassembler les stratégies qui se ressemblaient, à travers les différents participants, et nous nous sommes ensuite rendue compte que quelques caractéristiques des visualisations ornementées influençaient les stratégies de comblement de discontinuité. Malgré tout, là n'est pas l'essentiel : la variété des utilisateurs, de leurs expériences et de leurs personnalités, même en occupant une profession similaire, engendre une multitude de processus de construction de sens. Ainsi, il y a une part de construction de sens réalisée par l'individu en fonction de ce qu'il est uniquement, de son expérience, de son vécu et de ses caractéristiques, indépendamment de ce qu'est la visualisation. Cet aspect restera toujours incontrôlable pour le concepteur de visualisations de données qui ne connaît pas finement son audience.

Et pour cause, c'est à travers les actions des individus que le sens se crée, comme nous l'avons vu, et pas uniquement à travers la disposition de l'information sur la visualisation de données. L'ornementation cause des émotions qui poussent l'individu à agir et donc à faire sens. Ces émotions peuvent tant être positives que négatives. Alors qu'on pourrait interpréter la préconisation minimaliste de Bertin et Tufte comme une volonté d'éviter l'émergence d'émotions négatives, et donc d'actions de construction de sens influencées par cette négativité, il convient de rester alerte. L'absence d'ornementation cause elle aussi des émotions, des réflexions et des interrogations, comme l'a montré l'analyse thématique au cours de laquelle certains participants se sont demandé pourquoi les concepteurs des visualisations standard ne souhaitaient pas « prendre soin » du lecteur en rendant la visualisation « attrayante ».

La première impression laissée par la visualisation, inconscience et viscérale, est fondamentale : elle cause des émotions qui prédisposent l'individu à s'impliquer dans la lecture de la visualisation ou non, en fonction de sa personnalité. L'utilisabilité de la visualisation, à travers sa facilité de compréhension, est également un élément important, engendrant l'attention du lecteur ou son abandon. En raison de cela, nous avons montré que le design émotionnel selon Norman (2013) peut complètement s'appliquer à la visualisation de données, en combinant trois niveaux de design différents, à savoir le niveau viscéral, comportemental et réflexif. La visualisation qui fonctionnerait le mieux pour ces participants ne serait donc pas la plus fonctionnelle au sens strict du terme, mais celle qui

provoquerait une réaction positive au niveau viscéral, à première vue (« c'est beau », « je me sens attiré »), au niveau comportemental, en montrant une bonne utilisabilité, et au niveau réflexif, permettant ainsi à l'utilisateur de se positionner par rapport à ce qu'il a lu. Ainsi, on voit dans nos différents entretiens semi-directifs que, comme Norman l'énonce (2013), le plaisir de lecture des visualisations de données statiques fait en sorte qu'elles « fonctionnent » mieux. Dans les entretiens, les participants disent favoriser les visualisations ornementées parce qu'elles sont plus jolies, ont un caractère attractif, amusant, attrayant, ... En somme, elles génèrent une émotion positive qui prédispose à s'engager auprès de la visualisation de données. Mais d'autres éléments encore, comme les affinités du lecteur avec le sujet ou une grande aptitude critique, peuvent encore altérer cet engagement.

En effet, puisque les émotions et l'expérience des individus influencent les différents compléments de discontinuité lors de la lecture d'une visualisation, il semblait important de se pencher sur ces différentes caractéristiques. Justement, Kennedy et Hill (2017, 2018) parlent de facteurs humains et émotionnels qui favorisent ou défavorisent l'engagement envers la visualisation. Nous avons constaté que ces facteurs apparaissaient très concrètement dans le discours des participants. Ils expliquent souvent quelques difficultés quant à leurs compétences comme le vocabulaire ou encore les compétences mathématiques. Bien que les visualisations présentées dans le corpus étaient très accessibles, ce manque de compétences impacte pour certains participants leur confiance dans leur propre capacité à lire un graphique. De même, l'intérêt pour le sujet de la visualisation impacte souvent le temps que l'individu décide d'accorder à une visualisation. Les sensations de beau, de laid ou encore l'aspect scolaire, vieillot ou trop peu sérieux des visualisations prédisposent, dans différentes mesures, les participants à lire et à s'engager envers les visualisations. De nombreuses caractéristiques personnelles influencent donc la réception de l'information, en fonction de la présentation de cette dernière évidemment. Selon nous, la profession des utilisateurs est une caractéristique qui a influencé la réception des visualisations dans le cadre de notre recherche.

Additionnellement à tout cela, à travers l'analyse thématique, nous avons constaté que les participants expriment différentes sensations, notamment provoquées par l'ornementation. La plus commune de ces sensations est celle d'un effort qui semble diminué. Ainsi, alors que dans les faits, les participants produisent davantage de points de fixation sur les visualisations ornementées, ils ont en réalité l'impression que leur effort de compréhension est soulagé, notamment grâce à d'autres sensations conjointes comme l'expérience agréable ou encore l'amusement. Dans le même sens, les visualisations de type standard, pourtant estimées très claires, ont souvent provoqué une plus grande sensation d'effort à

fournir. Dès lors, nous constatons un décalage entre l'effort réel – notamment traduit par la perception visuelle ou par le fait que les visualisations ornementées soient estimées moins claires que leurs homologues standard – et la sensation d'effort. Selon nous, c'est la force des émotions qui produit ce décalage et qui, finalement, pousse l'individu à s'engager envers la visualisation.

Nous pouvons conclure que, bien qu'importantes, la qualité et la présentation des données ne sont pas les seuls éléments qui poussent l'individu à faire sens du message qui lui est présenté. De nombreux paramètres qui sont hors de contrôle de la conception entrent en compte dans ce processus : les émotions engendrées, les actions qui en découlent et ce, en fonction des différentes expériences et personnalités des utilisateurs. Néanmoins, à bien des égards, l'influence de l'ornementation a été très positive sur la construction de sens des participants qui, quasi-systématiquement, se sont dit plus motivés à lire et à s'engager auprès d'une visualisation ornementée plutôt qu'envers une visualisation standard. Pour autant, les participants ne se sont pas forcément montrés ouverts à l'ornementation dans chaque cas : en tant que pratique de visualisation, nous savons désormais qu'il convient de la maîtriser afin qu'elle convienne le plus finement possible aux attentes des utilisateurs, tout en respectant les principes de base de conception des visualisations de données. Nous savons également que certaines pratiques d'ornementation, comme l'utilisation de formes inhabituelles ou de formes causant des illusions d'optique désagréables, dérangent les utilisateurs.

4. Recommandations relatives à l'utilisation de l'ornementation

Conjointement à notre question de recherche, nous avons émis la volonté d'émettre, si possible, une ou plusieurs recommandations en termes de conception de visualisations de données ornementées et ce, en stimulant au mieux possible la perception visuelle, la compréhension et la subjectivité de l'utilisateur.

La principale recommandation que nous souhaitons effectuer dans le cadre de ce travail de recherche concerne la simplicité des visualisations de données et, principalement, de l'ornementation. Effectivement, l'idée de Tufte selon laquelle « *less is more* » s'appliquerait selon nous également à l'ornementation. Ainsi, plutôt que d'éviter d'employer l'ornementation, nous proposons d'utiliser une ornementation qui se veut elle aussi minimaliste. Dans une certaine mesure, la plus-value de l'ornementation est réelle lorsqu'elle consiste en l'utilisation d'icônes épurées. Néanmoins, de grandes questions persistent encore quant à la disposition de ces icônes. Il ne faut en revanche pas confondre

la position occupée par ces icônes dans l'espace et les différentes positions de toutes les composantes du graphique (titre, légende, etc.).

Il est possible de respecter les principes de conception des visualisations de données qui existent déjà (Bertin, 1967 ; Tufte, 1983, 2001) en se permettant d'inclure des icônes à des fins d'ornementation. Grâce aux résultats du *bundling*, nous pouvons par ailleurs affirmer que, s'il y a besoin d'une légende, sa position est appréciée sous le titre, afin de favoriser un ordre de lecture logique. De même dans le cadre de diagrammes en bâtonnets, les étiquettes de données placées au-dessus des bâtonnets sont appréciées et fluidifient le chemin de lecture sans provoquer sans cesse un retour du regard à l'axe y. De même, il va de soi que les étiquettes de données (nombres) doivent apparaître de manière claire et évidente (taille des nombres suffisamment grande et disposition centrale dans le chemin de lecture). En se cantonnant à l'utilisation d'icônes épurées, nous proposons d'éviter de générer de nouvelles formes fantaisistes de visualisations de données si leur but est esthétique uniquement. Notre recherche a montré que cette intention cosmétique est très mal reçue par les utilisateurs.

L'information nous manquant dans le cadre de cette recommandation concerne néanmoins l'espace occupé par l'icône apposée. En effet, où doit-on la disposer ? En réalité, notre travail n'a pas traité cette question-là en profondeur. Nous avons montré qu'il existe différents types d'ornementation (icône apposée, sur les traits relatifs aux données, en fond, ...) et nous avons étudié la différence de réception entre l'ornementation qui est apposée ou se constitue en *data-ink*. Nous pouvons donc seulement nous permettre d'expliquer quelle est la meilleure des deux méthodes selon nous. Les résultats de nos analyses n'ont pas montré de différences en termes de clarté et d'appréciation entre l'ornementation qui est apposée ou qui se constitue en *data-ink*. De même, le *bundling* n'a pas montré que l'ornementation apposée induisait une perturbation dans le chemin de lecture. Selon nos résultats, une méthode n'est pas forcément meilleure qu'une autre. Néanmoins, en regard des tâches perceptuelles et élémentaires de Cleveland et McGill (1984), nous recommandons de disposer l'icône à côté d'un graphique traditionnel. Effectivement, utiliser l'icône en la disposant sur le *data-ink* revient à comparer des aires. Or, l'œil humain évalue avec très peu de précision les différences entre les aires. De même, du point de vue de la conception, notre recommandation est prudente parce qu'elle permet de contrôler facilement le *lie factor*. Cela est moins aisé avec l'utilisations d'aires, d'autant plus dans le domaine professionnel où dans les faits, le *lie factor* n'est pas calculé. Nous le savons, cela ne dit pas pour autant où placer l'icône sur une visualisation dont l'ornementation est apposée à un graphique traditionnel.

L'application de cette recommandation ne semble pas difficile. Nous nous permettons donc de l'illustrer en fonction de ce qu'elle donne comme directives en l'état. La Figure 96 présente ainsi le nombre de voitures neuves immatriculées en Belgique en 2018 et 2019 par marque. Les données, issues de la Fédération Belge et Luxembourgeoise de l'Automobile et du Cycle, sont présentées dans un diagramme en bâtonnets. La légende est proche du titre conformément aux remarques des participants. Les étiquettes de données se trouvent au-dessus des bâtonnets afin de soulager le trajet visuel vers un éventuel axe y. L'icône représente une voiture et est directement liée au sujet de la visualisation. Elle est apposée au diagramme en bâtonnets, sous le titre, où l'espace le permet. Additionnellement à cela, nous nous sommes permise d'ajouter sous chaque marque de voiture le logo associé. Même si nous savons que de tels éléments ne sont pas forcément regardés, nous savons désormais qu'ils apportent de la contextualisation à la visualisation. Ce point était souligné par les participants qui appréciaient les icônes de marques et les drapeaux.

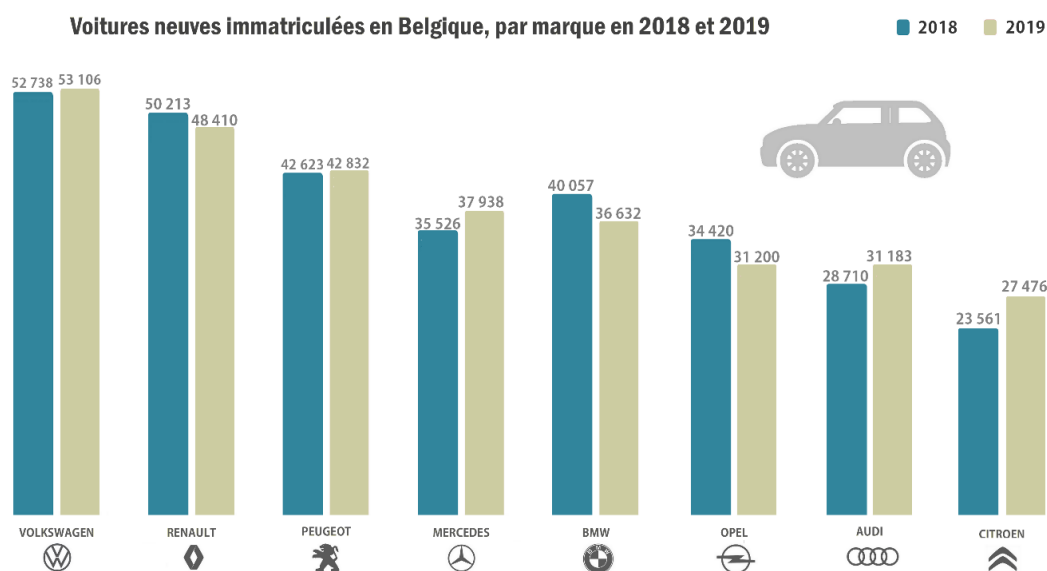


Figure 96 Prototype de visualisation suivant la recommandation principale de la recherche. Source des données : Datadigest (2020), "Immatriculations de voitures neuves par marque", repéré à <https://www.febiac.be/public/statistics.aspx?FID=23&lang=FR>

Du point de vue subjectif, une telle recommandation permet de tirer profit des différents bénéfices de l'ornementation, montrés par la littérature et soulignés dans cette recherche. Par exemple, utiliser une icône est un rappel au sujet et stimule l'aspect agréable d'une visualisation (nécessitant évidemment d'autres paramètres contrôlables comme l'utilisation de la couleur ainsi que des paramètres incontrôlables comme les affinités du lecteur avec le sujet de la visualisation). Il est également évident que le contexte d'utilisation de la visualisation de données influera sur les facteurs humains et émotionnels poussant les individus à s'engager envers la visualisation. Quoi qu'il en soit, proposer d'ornementer en utilisant une icône apposée pourrait être un pas vers le design émotionnel de Norman (2013), amenant un objet reconnaissable et agréable stimulant le jugement

esthétique au premier regard (niveau viscéral), garantissant clarté et utilisabilité (niveau comportemental). La connaissance de son audience est également importante afin d'atteindre ses intérêts (niveau réflexif).

Pour terminer, nous estimons que cette recommandation est une porte ouverte afin de proposer à l'avenir de nouvelles règles de conception graphique, incluant l'ornementation et en adéquation avec le nombre exponentiel de possibilités qui existe aujourd'hui en *Infovis*. Cela fait partie des quelques perspectives laissées par cette recherche.

5. Limites de la recherche et perspectives futures

Malgré ses apports, la présente recherche comporte beaucoup de limites dont la force est de pouvoir facilement se traduire en perspectives.

Pour commencer, comme nous l'avons précisé précédemment, les individus faisant partie de notre échantillon font tous partie d'un même monde professionnel. Ils sont tous à minima professionnels de la communication ou formés en communication et la pratique professionnelle de la plupart d'entre eux évolue dans l'écosystème numérique. Comme l'ont montré Quispel et Maes (2014), les professionnels du design graphique et les profanes (personnes de « la vie de tous les jours ») n'évaluent pas de la même façon l'ornementation des visualisations de données⁵¹. Il y a fort à parier que la profession influence la façon dont l'utilisateur considère la visualisation de données. Nous avons justement choisi les professionnels de la communication numérique car, habitués à la transmission de messages sur des dispositifs numériques, leur avis quant aux visualisations qui leur ont été présentées était marqué. Nos conclusions ne sont donc pas généralisables au citoyen lambda mais elles se traduisent en perspectives, en étudiant la question en fonction de différents profils d'utilisateurs.

De plus, nous n'avons pas traité l'aspect dynamique et interactif des visualisations de données. Nous nous sommes cantonnée aux visualisations statiques. Or, l'interaction est de plus en plus utilisée aujourd'hui, laissant davantage de libertés à l'utilisateur. Il serait sans doute judicieux de s'interroger sur l'ornementation de visualisations interactives dans des recherches futures. Les dynamiques d'ornementation y sont sans doute différentes.

Une autre limite de notre recherche, qui est la plus importante, concerne le corpus qui est très varié. Effectivement, lors des expérimentations, nous pensions que cette phase serait

⁵¹ Dans leur étude, l'ornementation concerne davantage des effets de style ajoutés à des graphiques traditionnels (ombrages, textures, etc.).

exploratoire et viserait ensuite à étudier des phénomènes plus restreints (par exemple, un type d'ornementation en particulier ou une situation particulière). La variété du corpus testé dans le cadre des expérimentations pose question puisque, parfois, il a été difficile d'affirmer un constat basé sur uniquement une seule visualisation ornementée du corpus. Néanmoins, c'est bien cette variété qui nous a permis de comparer et de mettre des différences en évidence. Si l'expérience était à refaire aujourd'hui, nous apporterions une attention plus grande aux visualisations ornementées sélectionnées afin de faire partie du corpus. De même, les sujets et sources des visualisations étaient très variés également, tandis que seul le diagramme en bâtonnets a été testé dans le cadre de cette étude. On pourrait aujourd'hui imaginer une étude se concentrant sur d'autres types de visualisations de données, d'autres ornementations, d'autres sujets. En cela, cette limite laisse place à d'importantes perspectives.

Nous avons conclu que la visualisation de données peut aujourd'hui être considérée comme une pratique qui se maîtrise. Une mauvaise maîtrise de la pratique engendrerait ainsi les dangers de l'ornementation que l'on connaît, tandis qu'une bonne maîtrise permettrait de profiter des avantages de l'ornementation. Même si nous avons amené des pistes vers des principes de conception de visualisations ornementées, tout l'enjeu reste aujourd'hui encore de pouvoir émettre des principes de conception beaucoup plus clairs et précis. Cela nous semble tout à fait possible, dans la mesure où les recherches futures s'attèleraient à qualifier plus clairement quels sont les différents types d'ornementation (au-delà des pictogrammes étudiés par Haroz et ses pairs (2015)). Une telle typologie permettrait ensuite d'étudier indépendamment les forces et faiblesses de différents types d'ornementation afin de recommander certaines règles à suivre. Selon nous, il est important de déterminer quelles sont les dispositions les plus adéquates de l'ornementation dans l'espace délimité des visualisations de données. Si la recommandation qui est émise dans le cadre de cette thèse émerge du croisement de plusieurs analyses, nous restons cependant convaincue que bien d'autres recommandations peuvent encore voir le jour, tant que la simplicité de l'ornement reste de mise.

Par exemple, il serait possible de vérifier avec un plus grand niveau de contrôle notre recommandation qui consiste à favoriser l'ornementation en icônes, apposée à un graphique traditionnel. Cette perspective se discute évidemment, mais en s'inspirant du protocole expérimental établi dans le cadre de cette recherche, nous pensons qu'il est possible d'étudier la validité d'une recommandation de ce type en induisant intentionnellement un effet de contexte qui, généralement, est évincé du laboratoire. Ainsi, c'est de façon idéaliste que nous nous permettons d'illustrer cette perspective à travers un cas fictif.

Imaginons donc un cas d'étude au cours duquel les personnes faisant partie de l'échantillon ont tous vécu une situation similaire. En offrant plus de contrôle que les expérimentations ayant fait l'objet de cette thèse de doctorat, ce cas fictif permettrait de confirmer et de faire émerger de nouvelles recommandations. L'échantillon peut également être défini en fonction de caractéristiques personnelles des utilisateurs, visant à découvrir davantage d'éléments relatifs à leur subjectivité. Imaginons également que les visualisations du corpus étudié aient toutes un élément de contexte en commun, qu'elles proviennent d'une source similaire et qu'elles soient également représentatives d'un type d'ornementation en particulier. Dans ce cas imaginaire, nous proposons l'ornementation apposée à un graphique traditionnel, qui se constitue en icônes simples. Le contexte proposé dans le cadre de ce cas imaginaire est celui de la pandémie de COVID-19 ayant frappé le monde entier au cours de l'année 2020. En annexe 12 se trouvent huit visualisations de données ornementées de la même façon, qui touchent de près ou de loin la thématique de la COVID-19⁵². La Figure 97 en montre deux exemples.

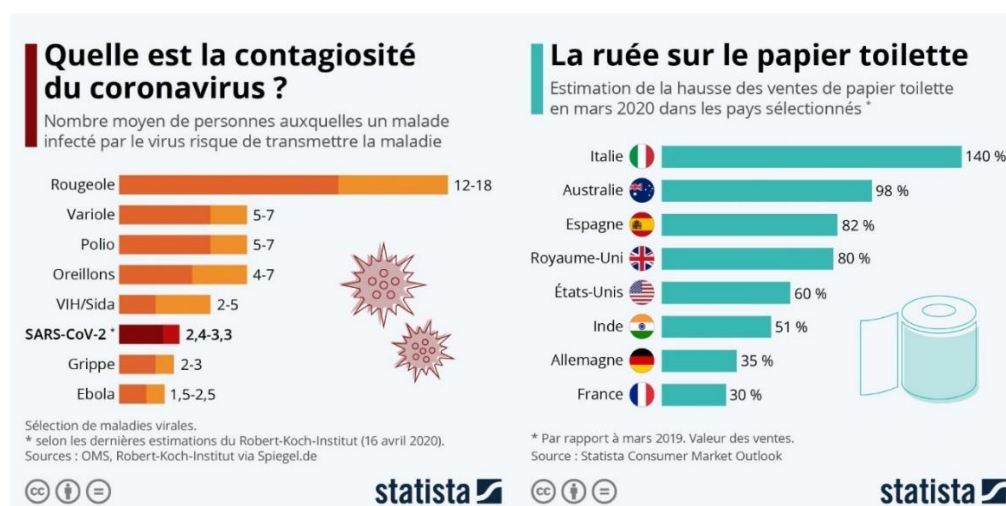


Figure 97 Gaudiaut T. (2020), "Quelle est la contagiosité du coronavirus ?", suivi de de "La ruée sur le papier toilette", repérés à <https://fr.statista.com/infographie/20653/comparaison-degre-de-contagion-coronavirus-avec-autres-virus/>, et <https://fr.statista.com/infographie/21368/hausse-des-ventes-papier-toilette-par-pays-pendant-epidemie-coronavirus-confinement/>

Dans le cadre de la pandémie, ces visualisations sont médiatiques et presque anecdotiques, à ne pas confondre avec les *dashboards* d'analyse d'évolution des cas recensés de coronavirus. L'expérimentation pourrait être présentée aux participants comme une étude portant sur la perception des informations relayées par rapport à la COVID-19. L'*eye tracker* relèverait de nouveau les données de regard au cours de l'expérimentation et des mesures telles que « aimer », « beau » et « claire » seraient de nouveau relevées sur des

⁵² Pour plus de simplicité dans l'illustration de ce cas fictif, nous n'avons de nouveau sélectionné que des diagrammes en bâtonnets. On pourrait également sélectionner un autre type de graphique ou, justement, varier les types de représentations graphiques présentes dans le corpus.

échelles de Likert. Les visualisations étant toutes similaires, l'objectif est de voir s'il y aurait une disparité entre les différentes données récoltées. De nouveau, la méthode de l'entretien semi-directif pourrait être employée en abordant la perception qu'ont les participants de la pandémie, en suivant la méthode de *sense-making methodology* de Dervin (1993) et en explorant plus en finesse les facteurs humains, émotionnels de Kennedy et Hill (2017, 2018). L'évitement des effets de séquence serait à réfléchir. Ce cas fictif pourrait montrer si les participants ont regardé l'ornementation, s'ils ont apprécié et trouvé les visualisations jolies et claires, en espérant que le score soit haut pour toutes ces mesures afin de valider notre recommandation première. En revanche, au vu de l'atmosphère mondiale négative et anxiogène résultant de la pandémie, il serait possible que les résultats soient différents et que la cause en soit expliquée dans les entretiens des participants qui auraient éventuellement mal vécu cette période. Ce cas pourrait être répliquable à d'autres types de visualisations et d'autres publics afin de mener à des comparaisons pertinentes.

Par ailleurs, les visualisations de données étudiées dans le cadre de cette recherche sont des objets médiatiques. Cette réflexion est assez tardive. Initialement, notre volonté était de partir d'objets simples et accessibles à tous avant de nous concentrer sur une utilisation professionnelle des visualisations de données. Au vu du déroulement de la recherche et du grand nombre de données récoltées, nous en sommes restée à la simple étude des visualisations de données comme objet médiatique. Selon nous, dans ce cadre, le contexte manque, à un double niveau : celui du dispositif dans lequel s'intègrent les visualisations (réseaux socio-numériques, sites web, applications numériques, ...) et celui du quotidien des utilisateurs. Choisir, c'est renoncer : en optant pour des conditions de laboratoire, nous n'avons pas pris en compte ces éléments de contexte, pourtant très importants dans l'appropriation des visualisations de données. Prendre attention à ces variations de contexte permettrait encore de développer davantage les conclusions que nous avons tirées dans le cadre de cette recherche. En contrepied de la proposition que nous venons de formuler, permettant un plus grand niveau de contrôle et donc la formulation de principes de design, il convient d'accorder une attention toute particulière aux différents effets de contexte à prendre en compte dans les recherches futures. Il serait ainsi imaginable de mener une ethnographie, par exemple en contexte professionnel (si c'est celui-là que l'on souhaite explorer), afin d'étudier comment sont utilisées les visualisations de données dans le monde professionnel. De même, à l'heure où les représentations graphiques foisonnent sur le web, mener une étude netnographique permettrait d'évaluer les intentions communicationnelles des concepteurs et des utilisateurs des visualisations de données ornementées. Cela serait d'ailleurs l'opportunité d'ouvrir le champ des questions de recherche sur cet objet qu'est la visualisation de données (cf. *infra*).

Pour terminer, l'interdisciplinarité de cette recherche ouvre différentes perspectives tant dans le domaine de *l'Infovis* que celui de la communication. Pour *l'Infovis*, force est de constater qu'il est aujourd'hui possible d'émettre des principes de conception permettant de contrôler au mieux la pratique d'ornementation des visualisations de données. Du côté des sciences de l'information et de la communication, ce long travail de recherche n'est encore qu'un préambule quant à l'étude d'un objet médiatique de plus en plus utilisé afin de communiquer l'information et ce, avec toutes les spécificités qu'il suscite, à l'heure d'une *datafication* exponentielle. L'originalité de ce travail tient en plusieurs caractéristiques : son sujet, qui est à la croisée des disciplines et surtout sa méthodologie complexe, qui pousse les sciences humaines en laboratoire afin d'augmenter mutuellement les apports des approches quantitative et qualitative.

6. Réflexion théorique et épistémologique : quels apports de la recherche aux SIC ?

Les visualisations de données sont incontestablement un objet de communication et pourtant, elles sont rarement abordées et étudiées en sciences de l'information et de la communication (SIC). Dans d'autres domaines, elles préoccupent les chercheurs depuis un bon moment. C'est ainsi que nous nous sommes tout naturellement tournée vers des théories venant d'autres disciplines (statistiques, informatique, design d'information, ...) pour nourrir notre propre recherche. D'une part, les principes de conception des visualisations d'information, dans lesquels nous avons constaté l'intention minimaliste de Bertin et Tufte, viennent du design d'information. D'autre part, c'est dans la communauté *Infovis* et parfois en informatique ou en *human computer interaction* (HCI) que nous nous sommes informée à propos de l'ornementation. Le concept de *sense-making* et de *gap bridging*, dont Dervin est la principale protagoniste (Dervin et al., 2003), vient quant à lui bel et bien du domaine de la communication. Par contre, ce cadrage théorique ne suffisait pas à lui seul pour commenter la construction de sens de manière complète : c'est pourquoi nous nous sommes tournée vers la perception visuelle et la compréhension. La richesse des sciences de l'information et de la communication tient dans son interdisciplinarité, que nous avons largement expérimentée dans le cadre de cette recherche.

« La communication est ainsi, au sens propre, une interdiscipline dont les théories et méthodes proviennent de champs disciplinaires voisins, historiquement antérieurs, que celle-ci adopte en les adaptant »

(Lits, 1999, p.11).

Cette recherche montre que les SIC ont encore à gagner de se nourrir de théories et méthodes venant d'ailleurs, surtout pour des objets de communication, comme la visualisation de données, qui ont été assez peu étudiés dans le domaine par le passé. Il est néanmoins important de ne pas considérer cette interdiscipline comme un fourre-tout dans lequel s'entassent des théories et méthodes qui, épistémologiquement, sont incohérentes entre elles. « La confrontation de modèles théoriques hétérogènes suppose en effet une mise en perspective de ceux-ci, pour éviter le rétrécissement de ces théories à des boîtes à outils, qui apparaîtraient comme dénuées de tout positionnement idéologique » (Lits, 1999, p.11). Justement, nous avons souhaité être précautionneuse à ce sujet : comment faire dialoguer entre elles des théories et méthodes aux différents positionnements épistémologiques ? Pour être plus concrète, comment concilier des approches d'horizons si différents ?

En effet, alors que nous avons expliqué notre positionnement constructiviste, les deux auteurs qui sont au cœur de notre questionnement, Bertin et Tufte, tiennent une position déterministe. Leur vision de la communication est à sens unique : le destinataire est supposé comprendre le message tel qu'il a été encodé. Pour Bertin, il est possible de limiter les interprétations possibles du graphique en utilisant un système de signe monosémique. Il fait fi du sujet interprétant, ne se préoccupant pas de l'utilisateur. Pour Tufte, qui milite contre le *chartjunk* et prône la maximisation du *data-ink ratio*, et donc du minimalisme, il n'y a qu'une réalité révélée par les données : la bonne visualisation montre la vérité si elle montre correctement les données. Cela suppose donc que les utilisateurs ne développent pas leur propre construction du monde social, ce qui entre en tension avec la vision de la réalité développée dans le cadre de cette recherche. Néanmoins, les principes de conception des visualisations de données sont incontournables. Élaborer un questionnement sur l'ornementation des visualisations de données sans les prendre en compte est impossible : il était impératif qu'ils fassent partie de l'aventure. Ils sont une référence qui permet tout de même de contrôler, autant que faire se peut, le message communiqué par les visualisations de données. Par la suite, la découverte des facteurs humains de Kennedy et Hill (2017, 2018) a permis d'établir de la nuance à propos de la façon de considérer les visualisations de données puisque ces chercheuses l'étudient d'un point de vue sociologique. Dès lors, leur discours sur l'utilisateur et sur ce qui influence son engagement envers la visualisation d'information s'est révélé être une clé nous ayant permis de se saisir pleinement de l'objet d'étude, en accord avec notre propre positionnement épistémologique. En connaissance de cause, rien n'indiquait qu'il était impossible de combiner ces différentes approches. Par ailleurs, d'autres approches étaient encore à combiner : comment étudier finement le sens produit à partir d'objets visuels ? Notre vision

de la construction de sens, dans le cadre de cette recherche, se veut complète. Elle aborde ainsi quelque chose qui est physique : la vision. De nouveau, l'étude de phénomènes physiques comme celui-là n'est pas fréquente en SIC. Pourtant, alliée au concept de *sense-making* selon Dervin, l'étude de la perception visuelle des visualisations de données grâce à l'*eye tracking* s'est révélée être une force pour notre analyse.

L'essence même des SIC, à savoir l'interdisciplinarité, est donc au cœur de ce travail de recherche. C'est bien ce qui l'a fondé, en apportant au domaine un objet d'étude neuf, accompagné de nouvelles perspectives de recherche. En tant qu'objet de communication évoluant constamment sur les supports numériques, la présentation des données, rendue possible par différentes pratiques de visualisation, se doit de tenir une place importante dans la recherche en communication. C'est d'ailleurs l'un des apports majeurs de la recherche. La visualisation de données n'est plus simplement un objet technique dont l'expérience d'utilisation est évaluable : c'est un media à l'implication sociologique indéniable. En étudier la construction de sens est un enjeu pour les SIC, dans une société où la communication par les données s'accroît.

L'apport de cette recherche pour les SIC ne concerne pas seulement l'objet de recherche : la méthodologie mise en place en fait l'originalité. De nouveau, réaliser des expérimentations en laboratoire est une méthode assez peu commune en SIC. D'un point de vue post-positiviste, en général, les expérimentations servent à expliquer un phénomène grâce à, notamment, des mesures quantitatives. Or, c'est en partant d'une posture constructiviste et interprétative que nous avons élaboré cette méthode. Le test d'*eye tracking* réalisé dans les expérimentations a relevé des mesures claires, précises et factuelles. Néanmoins, il ne fallait pas céder à une tentation explicative : malgré l'utilisation de ce dispositif, c'est bien le paradigme constructiviste, avec sa démarche interprétative, qui nous a permis d'élaborer une compréhension profonde du phénomène, en lien avec les autres données récoltées au cours des autres étapes de l'expérience. Ainsi, nous montrons qu'il est possible de mener une recherche en communication avec un tel positionnement épistémologique, en réalisant des expérimentations en laboratoire. La transformation digitale touche aujourd'hui beaucoup de domaines et la technologie, impliquant de nombreux objets connectés, se crée une place réelle dans les foyers et les entreprises. Que ce soit à un niveau médiatique, professionnel ou social, la digitalisation amène de nombreux objets de recherche pour la recherche en communication. L'utilisation d'outils de mesures biométriques, comme l'*eye tracker*, a de beaux jours devant elle en recherche en communication, avec des méthodes particulières à développer spécifiquement pour cette interdiscipline. Mais encore, en ce qui concerne notre propre méthodologie, la combinaison de récolte d'appréciations, de la réalisation d'un test d'*eye tracking* et

d'entretiens semi-directifs s'apparente à la notion de bricolage, comme pratiquée par de nombreux chercheurs dans le domaine. Véritable exploitation créative des ressources, le bricolage s'apparente à un mélange de prise de risque et d'ingénierie méthodologique (Meunier, Lambotte et Choukah, 2013). Souvent, c'est sur le terrain que le bricolage est pratiqué par les chercheurs en SIC, tandis que ce qu'en ce qui nous concerne, c'est au niveau, d'une part, de la conception du protocole expérimental et d'autre part, des méthodes d'analyse des données récoltées, que le bricolage s'est produit. Collaborer avec des chercheurs directement impliqués dans le domaine de la visualisation d'information pour élaborer une technique de visualisation des données de regard, spécifiquement applicable à notre recherche, en est un exemple. La mise en place de notre méthodologie s'est déroulée dans un environnement contrôlé, montrant l'intérêt des outils de mesure biométriques pour l'étude de nouveaux objets de recherche en SIC.

Quoi qu'il en soit, la visualisation de données, principalement propulsée sur le web, se conçoit comme un protocole de communication. Elle n'échappe pas aux mêmes ambitions dans le champ des SIC que celles d'autres protocoles induits par l'émergence de la transformation digitale, qui « doivent être analysés, pensés dans leur intégration sociale, afin que soient perçus les nouveaux paradigmes qu'ils disposent, dans l'organisation de la société comme dans la saisie cognitive du savoir » (Lits, 1999, p.11). Pour y parvenir, la multiplicité des théories et des méthodes, maîtrisées consciemment en fonction des approches proposées, était une condition sine qua non. Cet alliage prudent a ainsi mené vers de consistants résultats, permettant de cadrer l'utilisation de l'ornementation de visualisations de données statiques. Il propose aussi une ouverture des perspectives, permettant d'élaborer de nouveaux principes de conception des visualisations de données ornementées.

Bibliographie

- Adam, F., Pomerol, J. C. (2008), Developing practical decision support tools using dashboards of information, In *Handbook on Decision Support Systems 2* (pp. 151-173), Springer, Berlin, Heidelberg.
- Andry T. (2020), Visualisation de données et design émotionnel peuvent-ils se conjuguer ? *Communiquer. Revue de communication sociale et publique*, 28, p.53-71.
- Albert, W., & Tullis, T. (2013), *Measuring the User Experience, Second Edition: Collecting, Analyzing, and Presenting Usability Metrics* (2 edition), Amsterdam ; Boston, Morgan Kaufmann.
- Badjio, E. P. F. (2005). *Evaluation qualitative et guidage des utilisateurs en fouille visuelle de données*, Dissertation doctorale, Université Lumière Lyon 2.
- Bateman, S., Mandryk, R. L., Gutwin, C., Genest, A., McDine, D., & Brooks, C. (2010), « Useful junk ? : the effects of visual embellishment on comprehension and memorability of charts », *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, ACM, p. 2573-2582.
- Brasseur, C. (2015), « Usages visuels des données et big data », *Informations, données et documents*, n°2, p.44-46.
- Bertacchini, Y. (2009). *Petit Guide à l'usage de l'Apprenti-Chercheur en Sciences Humaines & Sociales*. Coll les ETIC, Presses Technologiques, Toulon.
- Bertin, J. (1967), *Sémiologie graphique*, Paris, Mouton.
- Bertin, J. (1981), "Théorie matricielle de la graphique", *Communication et langages*, 48(1), 62-74. <https://doi.org/10.3406/colan.1981.1409>.
- Borgo, R., Abdul-Rahman, A., Mohamed, F., Grant, P. W., Reppa, I., Floridi, L., & Chen, M. (2012), An empirical study on using visual embellishments in visualization. *IEEE Transactions on Visualization and Computer Graphics*, 18(12), 2759–2768.
- Borkin, M. A., Vo, A. A., Bylinskii, Z., Isola, P., Sunkavalli, S., Oliva, A., & Pfister, H. (2013). What Makes a Visualization Memorable? *IEEE Transactions on Visualization and Computer Graphics*, 19(12), 2306-2315. <https://doi.org/10.1109/TVCG.2013.234>.
- Borkin, Michelle A., Bylinskii, Zoya, Kim, Nam Wook, et al. (2016), Beyond Memorability: Visualization Recognition and Recall, *IEEE Transactions on Visualization and Computer Graphics*, vol. 22, no 1, p. 519-528.
- Bylinskii, Z., Borkin, M. A., Kim, N. W., Pfister, H., & Oliva, A. (2017), Eye Fixation Metrics for Large Scale Evaluation and Comparison of Information Visualizations. In M. Burch, L. Chuang, B. Fisher, A. Schmidt, & D. Weiskopf (Éds.), *Eye Tracking and Visualization* (p. 235-255). Springer International Publishing. https://doi.org/10.1007/978-3-319-47024-5_14
- Cairo, A. (2012). *The Functional Art: An introduction to information graphics and visualization*. New Riders, San Francisco, US [Format EPUB].
- Cairo, A. (2016). *The truthful art: data, charts, and maps for communication*. New Rider, San Francisco, US [Format EPUB].

- Campbell, D. T., & Stanley, J. C. (1967). *Experimental and quasi-experimental designs for research* (2. print). Boston: Houghton Mifflin Comp.
- Card, S. (2012), Information Visualization, in *Human Computer Interaction Handbook : Fundamentals, Evolving Technologies, and Emerging Applications, Third Edition*, CRC Press, p. 516-548.
- Cardon, D. (2012), "Regarder les données", *Multitudes*, 49(2), 138. <https://doi.org/10.3917/mult.049.0138>
- Cardon, D. (2015), *A quoi rêvent les algorithmes. Nos vies à l'heure des big data*, Paris, Seuil
- Carpenter, P. A., & Shah, P. (1998), A model of the perceptual and conceptual processes in graph comprehension, *Journal of Experimental Psychology: Applied*, 4(2), 75.
- Catellin, S. (2004), "L'abduction: une pratique de la découverte scientifique et littéraire", *Hermès, La Revue*, (2), 179-185.
- Charters, E. (2003), The use of think-aloud methods in qualitative research an introduction to think-aloud methods, *Brock Education Journal*, 12(2). *Visualization. Proceedings*, IEEE, p. 74-81.
- Chauvin, S. (2005), *Visualisations heuristiques pour la recherche et l'exploration de données dynamiques: l'art informationnel en tant que révélateur de sens*, Dissertation doctorale, Université Paris VIII Vincennes-Saint Denis
- Chi E. H. H., & Card S. K. (1999), Sensemaking of evolving web sites using visualization spreadsheets , *Information Visualization, 1999.(Info Vis' 99) Proceedings. 1999 IEEE Symposium on*, IEEE, p. 18-25.
- Cleveland, W. S., & McGill, R. (1984), Graphical Perception: Theory, Experimentation, and Application to the Development of Graphical Methods, *Journal of the American Statistical Association*, 79(387), 531-554. <https://doi.org/10.1080/01621459.1984.10478080>
- Coelho, D., & Mueller, K. (2020), Infomages: Embedding Data into Thematic Images. In *Computer Graphics Forum*, Vol. 39, No. 3, pp. 593-606.
- Colin, C. (2017), "Design", in *Encyclopædia Universalis*, PDF en ligne, consulté à l'adresse <http://www.universalis-edu.com/encyclopedia/design/>
- Collins English Dictionary, (s. d.), Embellishment definition and meaning, en ligne, consulté 31 octobre 2018, à l'adresse <https://www.collinsdictionary.com/dictionary/english/embellishment>
- Cui, W., Zhang, X., Wang, Y., Huang, H., Chen, B., Fang, L., ... & Zhang, D. (2019), Text-to-viz: Automatic generation of infographics from proportion-related natural language statements, *IEEE transactions on visualization and computer graphics*, 26(1), 906-916.
- Dasgupta, A., & Kosara, R. (2012), The importance of tracing data through the visualization pipeline, *Proceedings of the 2012 BELIV Workshop: Beyond Time and Errors-Novel Evaluation Methods for Visualization*.
- Derèze G. (2009), *Méthodes empiriques de recherches en communication*, Bruxelles, De Boeck.

- Dervin, B. (1992), From the mind's eye of the user: The sense-making qualitative-quantitative methodology, *Qualitative research in information management*, 9, 61-84.
- Dervin, B. (1998), Sense-making theory and practice : An overview of user interests in knowledge seeking and use, *Journal of Knowledge Management*, 2(2), 36-46. <https://doi.org/10.1108/13673279810249369>
- Dervin, B., Foreman-Wernet, L., & Lauterbach, E. (2003), *Sense-Making Methodology Reader: Selected Writings of Brenda Dervin*, Cresskill, N.J: Hampton Press.
- Dufrenne, M., (2017), "Esthétique et philosophie" in *Encyclopædia Universalis*, en ligne, consulté à l'adresse <http://www.universalis-edu.com/encyclopedie/esthetique-esthetique-et-philosophie/>
- Elling, S., Lentz, L., & De Jong, M. (2011), Retrospective think-aloud method: using eye movements as an extra cue for participants' verbalizations, *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (pp. 1161–1170). ACM.
- Fallery, B., & Rodhain, F. (2007), Quatre approches pour l'analyse de données textuelles : lexicale, linguistique, cognitive, thématique, In *XVI ème Conférence de l'Association Internationale de Management Stratégique AIMS* (p. pp 1-16). Montréal, Canada: AIMS, <https://hal.archives-ouvertes.fr/hal-00821448>
- Few, S. (2006), « Information dashboard design. Chapter 1 : Clarifying the vision », O'Reilly, Newton.
- Few, S., Edge, P. (2011), The chartjunk debate. *Visual Business Intelligence Newsletter*, no. June, 1–11.
- Fredriksson S. (2015), Du design d'information à la visualisation de données : un enjeu de transmission de sens auprès de la société civile, *I2D – Information, données & documents* 2015/2 (Volume 52), p. 36-36.
- Friendly, M. (2008), A brief history of data visualization, *Handbook of data visualization*, Springer Berlin Heidelberg, pp. 15-56.
- Gignac, G. E. (2019), *How2statsbook (Online Edition 1)*, Perth, Australia, <http://www.how2statsbook.com/>
- Harsha P., Hines C. (2016), Radar Charts and the paradigm of cognitive fit; Implication for accounting research and practice, *Journal of Business and Accounting*, Vol 9 N°1, Missouri State University.
- Hearst, M. (2003), Information visualization: Principles, promise, and pragmatics, *In Handouts of the tutorial at CHI 2003 Conference on Human Factors in Computing Systems*.
- Heer, J., Bostock, M., & Ogievetsky, V. (2010), A tour through the visualization zoo, *Communications of the ACM*, 53(6), 59-67.
- Hill, S., Wray, B., & Sibona, C. (2017), Minimalism in Data Visualization: Perceptions of Beauty, Clarity, Effectiveness, and Simplicity, *In Information Systems & Computing Academic Professionals*. Austin, Texas. Consulté à l'adresse <http://proc.conisar.org/2017/pdf/4514.pdf>

Howell, D., Bestgen, Y., Yzerbyt, V., & Rogier, M. (2008). « *Méthodes statistiques en sciences humaines. Ouvertures Psychologiques* », De Boeck Supérieur, Louvain-La-Neuve.

Hullman, J., Adar, E., & Shah, P. (2011), Benefitting InfoVis with Visual Difficulties, *IEEE Transactions on Visualization and Computer Graphics*, 17(12), 2213-2222. <https://doi.org/10.1109/TVCG.2011.175>

Hurter, C. (2010), *Caractérisation de visualisations et exploration interactive de grandes quantités de données multidimensionnelles. Interface homme-machine*, Dissertation doctorale, Université Paul Sabatier – Toulouse, <https://tel.archives-ouvertes.fr/tel-00610623>

Inbar, O., Tractinsky, N., & Meyer, J. (2007), Minimalism in information visualization: attitudes towards maximizing the data-ink ratio, (185 p.). ACM Press. <https://doi.org/10.1145/1362550.1362587>

Jacques, J. (2016). *Définition des compétences propres à l'organisation des collections d'informations personnelles numériques*, Dissertation doctorale, Université catholique de Louvain <http://hdl.handle.net/2078.1/174446>

Just, M. A., & Carpenter, P. A. (1976), Eye fixations and cognitive processes, *Cognitive Psychology*, 8(4), 441-480. [https://doi.org/10.1016/0010-0285\(76\)90015-3](https://doi.org/10.1016/0010-0285(76)90015-3)

Kaufmann, J. C., (2011), *L'entretien compréhensif*, Armand Colin, Malakoff, France.

Kennedy, H., & Hill, R. L. (2017), The pleasure and pain of visualizing data in times of data power, *Television & New Media*, 18(8), 769-782.

Kennedy, H., & Hill, R. L. (2018), The feeling of numbers: Emotions in everyday engagements with data and their visualisation. *Sociology*, 52(4), 830-848, <https://doi.org/10.1177/0038038516674675>

Kennedy, H., Hill, R. L., Aiello, G., & Allen, W. (2016), The work that visualisation conventions do, *Information, Communication & Society*, 19(6), 715-735.

Kennedy, H., Hill, R., Allen, W. and Kirk, A. (2016), Engaging with (big) data visualisations: factors that effect engagement and resulting new definitions of effectiveness, *First Monday*, 21(11), <http://firstmonday.org/ojs/index.php/fm/article/view/6389>

Kim, S. D. (2008), La raison graphique de Saussure, *Cahiers Ferdinand de Saussure*, 23-42.

Kirk, A. (2016), *Data visualisation: a handbook for data driven design*, Sage.

Kirk A. (2015), Views from « Seeing Data » research (Part 1), en ligne, consulté le 7 novembre 2017, à l'adresse <http://www.visualisingdata.com/2015/10/views-from-seeing-data-research-part-1/>

Kirsh, D. (2010), Thinking with external representations, *Ai & Society*, 25(4), p. 441-454.

Körner, C. (2011), Eye movements reveal distinct search and reasoning processes in comprehension of complex graphs, *Applied Cognitive Psychology*, 25(6), 893-905. <https://doi.org/10.1002/acp.1766>

- Kosara, R. (2007), Visualization criticism-the missing link between information visualization and art, In *Information Visualization, 2007. IV'07. 11th International Conference* (p. 631–636). IEEE.
- Lau, A., & Moere, A. V. (2007), Towards a model of information aesthetics in information visualization. In *Information Visualization, 2007. IV'07. 11th International Conference* (p. 87–92). IEEE.
- Lezon Riviere A., Ihadjadene M. (2014), Construction de sens et pratiques informationnelles chez les chefs militaires, *Journal of Human Mediatized Interactions/Revue des Interactions Humaines Médiatisées*, 15(2).
- Lima, M. (2013), *Cartographie des réseaux: l'art de représenter la complexité*, Paris, Eyrolles.
- Lits, M. (1999), Présentation. Cinquante années de recherches en communication, *Recherches en communication*, 11(11), 9-19.
- Lloveria, V. (2014), Data design-moi un mouton: De la data visualisation au data storytelling chez Michael Paukner, *Communication et organisation*, (46), 99-112. <https://doi.org/10.4000/communicationorganisation.4747>
- Maxwell, S. E., Delaney, H. (1990), *Designing experiments and analyzing data*. Monterey, CA: Wadsworth Publishing Company.
- Maylor, H. (2001), Beyond the Gantt chart: Project management moving on, *European management journal*, 19(1), 92-100.
- Metz, C. (1971), Réflexions sur la « Sémiologie graphique » de Jacques Bertin, *Annales. Histoire, Sciences Sociales*, 26(3), 741-767. <https://doi.org/10.3406/ahess.1971.422441>
- Meunier, D., Lambotte, F., & Choukah, S. (2013), Du bricolage au rhizome: comment rendre compte de l'hétérogénéité de la pratique de recherche scientifique en sciences sociales?. *Questions de communication*, (23), 345-366.
- McCandless, D. (2010), La beauté de la visualisation des données | TED Talk., vidéo en ligne, consulté 31 octobre 2018, à l'adresse https://www.ted.com/talks/david_mccandless_the_beauty_of_data_visualization?language=fr
- McCandless, D. (2011), *Datavision*. Robert Laffont, Paris.
- McCandless, D. (2012), *Information is beautiful*, Collins, London.
- Michaud, Y. (2017), "Art (aspects esthétiques), Le beau" in *Encyclopædia Universalis*, consulté à l'adresse <http://www.universalis-edu.com/encyclopedie/art-aspects-esthetiques-le-beau/>
- Negura, L. (2006), L'analyse de contenu dans l'étude des représentations sociales, *SociologieS*. Consulté à l'adresse <http://journals.openedition.org/sociologies/993>
- Neilsen, J. (2006), F-Shaped Pattern of Reading on the Web: Misunderstood, But Still Relevant (Even on Mobile), en ligne, consulté 25 octobre 2018, à l'adresse <https://www.nngroup.com/articles/f-shaped-pattern-reading-web-content/>
- Norman, D. A. (2008), The way i see it. Simplicity is not the answer, *Interactions*, 15(5), 45. <https://doi.org/10.1145/1390085.1390094>

Norman, D. A. (2013), *The design of everyday things* (Revised and expanded edition). New York, New York: Basic Books.

Payne, J. W. (1976), Task complexity and contingent processing in decision making: An information search and protocol analysis, *Organizational behavior and human performance*, 16(2), p. 366-38.

Peirce, C. S. (1978), *Écrits sur le signe* (rassemblés, traduits et commentés par G. Deledalle), Paris, Seuil.

Penelaud, O. (2010), Le paradigme de l'énaction aujourd'hui. Apports et limites d'une théorie cognitive « révolutionnaire », *PLASTIR*, 18(1), 1-35

Pernice, K., & Nielsen, J. (2009), Eyetracking methodology: How to conduct and evaluate usability studies using eyetracking, *Nielsen Norman Group Technical Report*

Peysakhovich, V., & Hurter, C. (2017), Scanpath visualization and comparison using visual aggregation techniques, *Journal of eye movement research*, vol 10 (n°5), pp. 1-14, ISSN 1995-8692

Quispel, A., & Maes, A. (2014), Would you prefer pie or cupcakes? Preferences for data visualization designs of professionals and laypeople in graphic design, *Journal of Visual Languages & Computing*, 25(2), 107-116. <https://doi.org/10.1016/j.jvlc.2013.11.007>

Ramly, N. N., Nor, F. M., Ahmad, N. H., & Aziz, M. H. (2012), Comparative Analysis on Data Visualization for Operations Dashboard, *International Journal of Information and Education Technology*, 2(4).

Rayner, K. (1998), Eye movements in reading and information processing: 20 years of research, *Psychological bulletin*, 124(3), 372.

Reinhard, C. D., & Dervin, B. (2012), Comparing situated sense-making processes in virtual worlds: Application of Dervin's Sense-Making Methodology to media reception situations, *Convergence*, 18(1), 27-48.

Reinsel D., Gantz J., Rydning J. (2020), *The Digitization of the World. From Edge to Core*, Data Age 2025. An IDC White Paper, PDF en ligne, repéré à <https://www.seagate.com/files/www-content/our-story/trends/files/dataage-idc-report-final.pdf> le 25 septembre 2020.

Romelaer (2005), Chapitre 4. L'entretien de recherche. Dans Roussel, P., & Wacheux, F. (dir), *Management des ressources humaines : Méthodes de recherche en sciences humaines et sociales*, p. 101-136, Bruxelles, De Boeck Supérieur.

Rouve, M. E., & Ria, L. (2008), Analyse de l'activité professionnelle d'enseignants néo-titulaires en réseau « ambition réussite »: études de cas. *Travail et formation en éducation*, (1).

Sack W. (2004), Aesthetics of information visualization, PDF en ligne, consulté 31 octobre 2018, à l'adresse https://www.researchgate.net/publication/228636081_Aesthetics_of_information_visualization

Saulnier, A., Thièvre, J., & Viaud, M.-L. (2006), La perception du mouvement dans la visualisation - Le cas des graphes, *Revue RIHM*, Vol. 7 (N°2).

- Saville, D. J., & Wood, G. R. (1991), Latin Square Design, In *Statistical Methods: The Geometric Approach* (p. 340-353). Springer, New York, NY. https://doi.org/10.1007/978-1-4612-0971-3_13
- Savolainen, R. (2006), Information use as gap-bridging : The viewpoint of sense-making methodology. *Journal of the American Society for Information Science and Technology*, 57(8), 1116-1125. <https://doi.org/10.1002/asi.20400>
- Scaife, M., & Rogers, Y. (1996), External cognition : How do graphical representations work? *International Journal of Human-Computer Studies*, 45(2), 185-213. <https://doi.org/10.1006/ijhc.1996.0048>
- Schmitt, D. (2012), *Expérience de visite et construction des connaissances : le cas des musées de sciences et des centres de culture scientifique* (Doctoral dissertation, Université de Strasbourg), consulté à l'adresse <https://tel.archives-ouvertes.fr/tel-00802163/document>
- Sicard, M. (1997). Les paradoxes de l'image. *Hermès*, 21, 45-54.
- Skau, D., Harrison, L., & Kosara, R. (2015), An evaluation of the impact of visual embellishments in bar charts, In *Computer Graphics Forum* (Vol. 34, p. 221–230), Wiley Online Library.
- Sow, A. B., & Chaudiron, S. (2017), Analyse des pratiques info-communicationnelles des chercheurs sénégalais selon l'approche sense-making, *Colloque international "Les sciences de l'information documentaire au service de la recherche, de la formation, de l'intégration et du développement durable"*, Dakar, Sénégal, <https://hal.archives-ouvertes.fr/hal-01832931>
- Spence, I. (2006), William Playfair and the psychology of graphs, *American Statistical Association, Alexandria*, p. 2426-2436.
- Spurgin, K. M. (2006). The Sense-Making Approach and the Study of Personal Information Management. *Personal Information Management*, 3.
- Telea A. (2014), *Data visualization: principles and practice*, CRC Press.
- Theureau, J. (1992), *Le cours d'action : Analyse sémiologique. Essai d'une anthropologie cognitive située*, Berne: Peter Lang
- Theureau, J. (2002), Cours d'expérience, cours d'action, cours d'interaction: essai de précision des objets théoriques d'étude de l'activité individuelle-sociale, *Objets théoriques, objets de conception, objets d'analyse & situations d'étude privilégiées*, Nouan-Le-Fuzelier.
- Theureau, J. (2006), *Le cours d'action: Méthode développée*, Toulouse, Octarès
- Thinès G. (2016), « Perception », In *Universalis éducation – Encyclopædia Universalis* [HTML en ligne] consulté le 22/02/17), sur <http://www.universalis-edu.com/encyclopedie/perception/>
- Tobii (2016), Tobii Pro X3-120 screen-based eye tracker, en ligne, consulté en mars 2018, à l'adresse <https://www.tobii.com/product-listing/tobii-pro-x3-120/>
- Tufte, E. R. (1983), *The visual display of quantitative information*, Cheshire, CT: Graphics.
- Tufte E. R. (1991), *Envisioning information*, Cheshire, CT: Graphics.
- Tufte, E. R. (1997), *Visual explanations: images and quantities, evidence and narrative*.Cheshire, CT: Graphics.

Vande Moere, A. V., Tomitsch, M., Wimmer, C., Christoph, B., & Grechenig, T. (2012), Evaluating the effect of style in information visualization, *IEEE transactions on visualization and computer graphics*, 18(12), 2739–2748 .

Vandeschrick, C., Wautelet (2003), J-M, *De la statistique descriptive aux mesures des inégalités.*, Academia/Bruylant/L'Harmattan, Louvain-la-Neuve/Paris (ISBN:2-87209-694-9).

Verhaegen, P. (2010), *Signe et communication*, De Boeck, Bruxelles.

von Glasersfeld, E. (1994), Pourquoi le constructivisme doit-il être radical ?, *Revue des sciences de l'éducation*, 20(1), 21. <https://doi.org/10.7202/031698ar>

Zen, M., & Vanderdonckt, J. (2014), Towards an evaluation of graphical user interfaces aesthetics based on metrics (p. 1-12), *IEEE*, <https://doi.org/10.1109/RCIS.2014.6861050>

Zhu, B., & Chen, H. (2008), Information visualization for decision support, *In Handbook on Decision Support Systems 2*, Springer Berlin Heidelberg, p. 699-722.

Table des figures et tableaux

Figures

Figure 1 : The purpose Map, tirée de Kirk A. (2016)	16
Figure 2 : Vente de magazines depuis 2005, graphique à titre d'exemple, par l'auteure ..	20
Figure 3 : Graphique minimaliste selon Tufte, utilisé dans l'expé. de (Inbar et al., 2007).	32
Figure 4 : Bateman et al. (2010, p.2576), " Example charts used in the study ».....	35
Figure 5 : Skau et al. (2015), "Embellished bar charts simplified for use in the study" ...	37
Figure 6 : Borgo et al. (2012), "Examples of stimuli in the study"	38
Figure 7 : Styles de visualisations repérés dans Vande Moere et al. (2012).	39
Figure 8 : Design expérimental exploité dans Borkin, M. A. et al (2016).	39
Figure 9 : Ornementation étudiée par Borkin et ses collègues.....	43
Figure 10 : Différencier chartjunk et ornementation, par l'auteure	53
Figure 11 : Entre le canard et le raw chart	54
Figure 12 : Les composantes du sense-making en visualisation de données par l'auteure	56
Figure 13 : The Visualization Pipeline - A technical View, tirée de A. Telea (2014)	68
Figure 14 : Influence des principes de conception et de l'ornementation sur (...)	71
Figure 15 : Data Match par Paris Match (2015).....	77
Figure 16 : « La nuit de sommeil du 31 décembre », raw visualization.....	77
Figure 17 : Statista (2017), « Le monde entier attend le beaujolais nouveau »,.....	78
Figure 18 : Statista (2017), « Le succès du champagne dans le monde »	78
Figure 19 : Usability Lab. Plan de C. Peeters annoté par nos soins	81
Figure 20 : Tobii Pro (2017), Tobii Pro X3-120	82
Figure 21 : Tobii Pro (2017), "What happens during the eye tracker calibration"	83
Figure 22 : Saville & Wood (1991, p.359), Randomization method for a carré latin.....	86
Figure 23 : Randomisation de l'ordre des visualisations de données	86
Figure 24 : Appréciations déclarées utilisées dans les expérimentations.....	89
Figure 25 : Exemples de résultats des appréciations déclarées	97
Figure 26 : Prise de vue de l'entretien semi-directif avec la caméra sur pied	98
Figure 27 : Appréciations déclarées utilisées dans les expérimentations.....	106
Figure 28 : Box plot : distribution des appréciations déclarées en fonction (...)	114
Figure 29 : Box plot : distribution des appréciations déclarées en fonction (...)	115
Figure 30 : Box Plots illustrant les distributions des groupes de l'attribut « Ornement »	123
Figure 31 : Box Plots représentant les distributions pour les variables de (...)	129
Figure 32 : Box Plots des variables relatives à l'attribut "Pictogramme"	135
Figure 33 : Box Plots présentant les distributions des données en fonction (...)	137
Figure 34 : Box Plots représentant les distributions des données en fonction (...)	140
Figure 35 : Box Plots représentant les distributions de l'attribut « Embellished data-ink »	145
Figure 36 : Visualisation n°7 du corpus	151
Figure 37 : Visualisation n°10 du corpus	152
Figure 38 : Visualisation n°17 du corpus	153
Figure 39 : Visualisation n°18 du corpus	155
Figure 40 : Visualisation n°14 du corpus	155
Figure 41 : Catégories de codes utilisées dans l'analyse thématique.....	164
Figure 42 : Codes apparaissant dans le contexte de la « sensation d'effort à fournir » ..	172
Figure 43 : Codes apparaissant dans le contexte de la « sensation d'effort amoindri » ..	173
Figure 44 : Nombres de codes menant à l'incompréhension / la confusion	180
Figure 45 : Exemple de schéma de gap bridging caractérisé par les verbing.	200
Figure 46 : Nuage de mot représentant les différentes discontinuités.....	201

Figure 47 : Nuage de mot représentant les différents pontages	202
Figure 48 : Nuage de mot représentant les différentes situations finales.....	203
Figure 49 : Nombre de schémas de gap bridging identifiés.....	205
Figure 50 : Miniatures des visualisations n°8, 11, 13, 15, 17	206
Figure 51 : Miniatures des visualisations n°5, 12, 18, 7 et 16.	207
Figure 52 : Miniatures des visualisations 18 et 19.....	208
Figure 53 : Miniatures des visualisations n°2, 4,10	209
Figure 54 : Miniatures des visualisations n°1, 3, 6, 9 et 14	209
Figure 55 : Schéma de gap bridging relatif à la visualisation n°4	211
Figure 56 : Schéma de gap bridging relatif aux visualisations chargées visuellement ...	211
Figure 57 : Schéma de gap bridging relatif à une discontinuité de type affectif.....	212
Figure 58 : Schéma modèle n°1 de gap bridging à situation finale positive	213
Figure 59 : Schéma modèle n°2 de gap bridging à situation finale positive	214
Figure 60 : Schéma modèle n°3 de gap bridging à situation finale positive	215
Figure 61 : Schéma modèle n°4 de gap bridging à situation finale positive	216
Figure 62 : Schéma modèle n°5 de gap bridging à situation finale négative	217
Figure 63 : Schéma modèle n°6 de gap bridging à situation finale négative	217
Figure 64 : Schéma modèle n°7 de gap bridging à situation finale négative	218
Figure 65 : Schéma modèle n°8 de gap bridging à situation finale négative	218
Figure 66 : Schéma modèle n°9 de gap bridging à situation finale négative	219
Figure 67 : Élément du corpus : visualisation ornementée et adaptée à l'expérimentatio	225
Figure 68 : Exemple d'un gaze plot pour un participant, visualisation n°5 du corpus	233
Figure 69 : Exemple de carte de chaleur pour la visualisation n°5 du corpus	233
Figure 70 : Visualisation des zones d'intérêt pour la visualisation n°17 du corpus.....	240
Figure 71 : Nombre moyen de fixation par visualisation.....	242
Figure 72 : Temps de vision moyen pour chaque visualisation	242
Figure 73 : Moyenne du nombre de fixations pour l'attribut « Figuration »	253
Figure 74 : Trajets oculaires (gaze plots) de deux participants pour la visualisation n°3	258
Figure 75 : Trajets oculaires de la totalité des participants pour la visualisation n°26 ...	258
Figure 76 : Exemple d'un premier résultat de bundling appliqué à la visualisation n°17	259
Figure 77 : Exemple de carte de densité des fixations jointe aux bundles.....	260
Figure 78 : Version finale du bundling appliquée à la visualisation n°17	260
Figure 79 : Visualisation n°11 (haut) et 31 (bas), version originale (gauche) (...).....	262
Figure 80 : Bundling appliqué à la visualisation n°21	263
Figure 81 : Bundling appliqué à la visualisation n°12	263
Figure 82 : Bundling appliqué à la visualisation n°17	264
Figure 83 : Bundling appliqué à la visualisation n°1	264
Figure 84 : Visualisation n°18 du corpus, suivie de son homologue n°38	265
Figure 85 : Bundling appliqué à la visualisation n°18 du corpus	266
Figure 86 : Bundling appliqué à la visualisation n°38 du corpus	266
Figure 87 : Bundling appliqué à la visualisation n°25 du corpus	267
Figure 88 : Bundling appliqué à la visualisation n°33 du corpus	268
Figure 89 : Bundling appliqué à la visualisation n°9 du corpus	268
Figure 90 : Bundling appliqué à la visualisation n°4.....	269
Figure 91 : Bundling appliqué à la visualisation n°2	269
Figure 92 : Bundling appliqué à la visualisation n°34 du corpus	270
Figure 93 : Bundling appliqué à la visualisation n°30 du corpus	270
Figure 94 : Bundling des visualisations n°24 (gauche) et n°31 (droite)	271
Figure 95 : Bundling appliqué à la visualisation n°14 du corpus	272
Figure 96 : Prototype de visualisation suivant la recommandation	285
Figure 97 : Gaudiaut T. (2020), "Quelle est la contagiosité du coronavirus ?", suivi de.	291

Tableaux

Tableau 1 : Classification des visualisations de données, par l'auteure	15
Tableau 2 : Exemples réels des différents types de chartjunk.....	27
Tableau 3 : Classification des pictogrammes dans (Haroz et al., 2015).....	44
Tableau 4 : Exemple de spécificités des visualisations ornementées.....	78
Tableau 5 : Facteurs humains et émotionnels selon Kennedy et Hill (2017, 2018).....	79
Tableau 6 : Randomisation de l'ordre de présentation des visualisations de données.....	87
Tableau 7 : Données disponibles pour l'analyse, par item	106
Tableau 8 : Nombre d'appréciations récoltées et analysables par item	107
Tableau 9 : Fréquences des données pour l'attribut "Ornement"	108
Tableau 10 : Fréquences des données pour l'attribut "Figuration"	108
Tableau 11 : Fréquences des données pour l'attribut "Pictogramme"	109
Tableau 12 : Fréquences des données pour l'attribut "Lie Factor"	110
Tableau 13 : Fréquences des données pour l'attribut "Data-Ink"	110
Tableau 14 : Fréquences des données pour l'attribut "Embellished Data-Ink"	110
Tableau 15 : Statistiques descriptives des appréciations par groupes de visualisations..	111
Tableau 16 : Résultat du test de Brown Forsythe pour l'analyse des données en (...)....	117
Tableau 17 : Moyenne des rangs des données classées selon l'attribut « Type ».....	118
Tableau 18 : Test de Brown Forsythe sur toutes les variables de l'attribut « Ornement »..	124
Tableau 19 : Test de Brown Forsythe pour les variables « Aucun » et « Icônes »	125
Tableau 20 : Test de Brown Forsythe pour les variables « Icônes » et « Forme (...).....	126
Tableau 21 : Test de Brown Forsythe pour les variables « Aucun » et « Forme (...).....	127
Tableau 22 : Test de Brown Forsythe concernant l'attribut "Figuration".....	130
Tableau 23 : Test de Brown Forsythe sur les variables "Absence" et "Apposé »	131
Tableau 24 : Test de Brown Forsythe sur les variables « Apposé » et « Sur le Data-Ink	132
Tableau 25 : Test de Brown Forsythe sur les variables « Absence » et « Sur le Data-In	133
Tableau 26 : Test de Brown Forsythe sur l'attribut « Pictogramme »	136
Tableau 27 : Test de Brown Forsythe sur l'attribut "Lie Factor"	138
Tableau 28 : Test de Brown Forsythe sur l'attribut « Data-Ink »	141
Tableau 29 : Test de Brown Forsythe sur les variables « Bas » et « Moyen » (...).....	142
Tableau 30 : Test de Brown Forsythe sur les variables « Moyen » et « Élevé » (...).....	143
Tableau 31 : Test de Brown Forsythe sur les variables « Bas » et « Élevé » de (...)	143
Tableau 32 : Test de Brown Forsythe sur l'attribut "Embellished Data-Ink"	146
Tableau 33 : Test de Brown Forsythe pour les variables « Absence » et « bas » (...)....	147
Tableau 34 : Test de Brown Forsythe pour les variables « Absence » et « moyen (...) .	147
Tableau 35 : Test de Brown Forsythe pour les variables « Absence » et « Élevé » (...) .	148
Tableau 36 : Test de Brown Forsythe pour les variables « Bas » et « Moyen » (...).....	148
Tableau 37 : Test de Brown Forsythe pour les variables « bas » et « élevé » de (...)	149
Tableau 38 : Test de Brown Forsythe pour les variables « Moyen » et « Élevé » (...) ..	149
Tableau 39 : Brefs résultats de l'étude de Quispel et Maes (2014)	161
Tableau 40 : Tableau de synthèse sur la sensation d'effort	172
Tableau 41 : Exemple d'analyse d'un verbatim propre au gap bridging (...).....	199
Tableau 42 : Verbins des discontinuités.....	202
Tableau 43 : Verbins des pontages.....	203
Tableau 44 : Verbins des situations finales	204
Tableau 45 : Parallélisme entre le design émotionnel et une situation de gap bridging..	224
Tableau 46 : Expression orale des niveaux de design émotionnel	226
Tableau 47 : Sélection de métriques oculométriques.....	237
Tableau 48 : Présentation des métriques oculométriques dans le tableur pour analyse ..	240
Tableau 49 : Test de normalité de la durée des fixations, par type (Raw, Junk).....	244

Tableau 50 : Test d'homogénéité des variances pour la durée moyenne des fixations ...	244
Tableau 51 : Test de Mann Whitney U pour la durée des fixations en fonction du type	244
Tableau 52 : Test de normalité du nombre de fixations par type (Raw, Junk)	245
Tableau 53 : Test de Levene pour le nombre de fixations	245
Tableau 54 : Test de Welch sur les moyennes des rangs du nombre de fixations (...)	246
Tableau 55 : Test de Welch sur le nombre de fixations, en fonction de l'attribut (...)	248
Tableau 56 : Test de Welch sur les rangs du nombre de fixations, en fonction de l'attribut « Figuration » (Apposé, Sur le Data-Ink)	248
Tableau 57 : Test de Welch sur le nombre de fixations, en fonction du lie factor.....	249
Tableau 58 : Test de Brown Forsythe sur les moyennes des rangs du Temps de vision	250
Tableau 59 : Test de Brown Forsythe sur le temps de vision en fonction du (...)	251
Tableau 60 : Test de Brown Forsythe sur le temps de vision en fonction du (...)	251
Tableau 61 : Test de Brown Forsythe sur le temps de vision en fonction du (...)	251
Tableau 62 : Tableau synthétique des moyennes de métriques relatives à H7	255