



Apprentissage implicite de régularités contextuelles au cours de l'analyse de scènes visuelles

A. Goujon

► To cite this version:

A. Goujon. Apprentissage implicite de régularités contextuelles au cours de l'analyse de scènes visuelles. Psychologie. Université de Provence - Aix-Marseille I, 2007. Français. NNT : . tel-00543817

HAL Id: tel-00543817

<https://theses.hal.science/tel-00543817>

Submitted on 6 Dec 2010

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

UNIVERSITE DE PROVENCE – Aix-Marseille I
UFR Psychologie, Sciences de l'Education

THÈSE

En vue de l'obtention du grade de
Docteur de l'Université d'Aix-Marseille I
Formation doctorale : Psychologie

Présentée et soutenue publiquement par

Annabelle GOUJON

**APPRENTISSAGE IMPLICITE
DE RÉGULARITÉS CONTEXTUELLES
AU COURS DE L'ANALYSE DE SCÈNES VISUELLES**

Sous la direction de :

Evelyne MARMÈCHE,

André DIDIERJEAN et Pierre VAN ELSLANDE

Membres du Jury :

Muriel BOUCART, *Directeur de Recherche, CNRS & CHRU de Lille* (Rapporteur)

André DIDIERJEAN, *Professeur des Universités, Université de Franche-Comté* (Directeur)

Michèle FABRE-THORPE, *Directrice de Recherche, CNRS & Université Paul Sabatier* (Examinateur)

Patrick LEMAIRE, *Professeur des Universités, IUF, CNRS & Université de Provence* (Président du jury)

Evelyne MARMÈCHE, *Chargée de Recherche, CNRS & Université de Provence* (Directeur)

Pierre PERRUCHET, *Directeur de Recherche, CNRS & Université de Bourgogne* (Rapporteur)

Pierre VAN ELSLANDE, *Chargé de Recherche INRETS, Salon de Provence* (Directeur)

18 décembre 2007

REMERCIEMENTS

*Je tiens à remercier **Pierre Perruchet**, **Muriel Boucart**, **Michèle Fabre-Thorpe** et **Patrick Lemaire**, d'avoir accepté de faire partie du jury de thèse. C'est un honneur pour moi de vous faire part de mon travail.*

*Je voudrais ensuite dire à **Evelyne Marmèche** qu'elle a été une directrice formidable. Je lui suis tout d'abord d'une reconnaissance inestimable pour m'avoir orientée sur une thématique de recherche aussi passionnante. Ayant une formation initiale en biologie, elle m'a formée en psychologie cognitive, notamment en me recommandant une littérature riche et variée dans laquelle j'ai trouvé une nourriture intellectuelle extrêmement stimulante. Depuis mon DEA, elle a été très présente et disponible, parfois même les soirs et week-ends, tout en me laissant autonome dans la conduite de mes recherches et libre de mon temps. Elle a su me 'recadrer', lorsque parfois je m'égarais vers des recherches qui de toute évidence n'auraient abouti sur rien de concret. Jusqu'aux dernières lignes de ce mémoire, relecture après lecture, elle m'a fait part de ses compétences et de sa rigueur en matière de rédaction. Je voudrais aussi que de ces quelques lignes se dégagent ses qualités profondément humaines et généreuses. Ces dernières années à travailler sous son encadrement auront été extrêmement agréables. Merci pour tout.*

*Mes remerciements sont ensuite adressés à **André Didierjean**, qui avec Evelyne à co-dirigé ma thèse. Puisqu'il est aussi à l'origine de ma thématique de recherche, je lui en suis d'une grande reconnaissance. Même si la distance Besançon-Marseille n'est pas propice aux réunions hebdomadaires, il a toujours été présent, tout particulièrement ces dernières semaines, autrement dit dans les moments les plus importants de la rédaction de mon mémoire de thèse. Je le remercie pour le temps passé sur mon mémoire et pour les conseils avisés, les bonnes idées et les critiques pertinentes, dont il m'a fait part tout au long de ma thèse. Je lui exprime ma plus grande gratitude pour son encadrement mais aussi pour sa sympathie.*

*Je remercie **Pierre Van Elslande** de m'avoir dirigée dans le cadre de l'INRETS. Je lui suis reconnaissante de m'avoir soutenue même si mes recherches peuvent paraître, à première vue, à mille lieux des préoccupations de l'INRETS. S'il est vrai que la conduite automobile occupe peu de place dans ce travail, j'espère que nous continuerons ensemble l'étude en cours, dont la problématique est plus directement reliée aux préoccupations de l'INRETS.*

*Je remercie l'**Institut National de Recherche sur les Transports et leur Sécurité (INRETS)** de m'avoir financé ces trois dernières années.*

*Je remercie, encore une fois, **Patrick Lemaire**, directeur de l'équipe Plasticité Cognitive à laquelle je suis rattachée. Un grand merci pour les conseils et remarques constructives qu'il a apportés sur les travaux que j'ai présentés au cours de plusieurs réunions d'équipe.*

*J'exprime toute ma reconnaissance à **David Gimmig, Loïc Bonnier** et **Mathieu Gimmig**, qui m'ont amicalement fait part de leur expertise en matière informatique et donné de leur temps pour la programmation des stimuli. Je remercie tout particulièrement Loïc et David pour les nombreux services qu'ils ont pu me rendre tout au long de mon travail. Je remercie **Bastien Camus** et **Céline Parraud**, de m'avoir accompagnée pour photographier des scènes routières, mais aussi **Yousri Marsouki** pour ses multiples coups de main, **Guy Tiberghien** pour ses conseils statistiques et pour les références bibliographiques qu'il m'a communiquées, ainsi que **Fermin Moscoso del Prado** pour les relectures de manuscrits en anglais, et **Aline Pélissier** (responsable du service de documentation) pour l'aide documentaire qu'elle m'a fournie.*

*J'ai eu la chance au cours de mon doctorat d'être rattachée à deux laboratoires où règne une ambiance agréable. Ayant passé plus de temps à l'INRETS, je commencerai par remercier mes collègues de Salon de Provence, et en particulier **Sofia, Isabelle, Marilyne, Michèle, Anne-Laure, Christine, Katel, Christophe** et encore une fois **Bastien** et **Céline**. Et je remercie bien sûr chaleureusement tous les membres du LPC, et en particulier **Yousri, David, Magali, Laetitia, Emmanuelle, Caroline, Stéphanie, Delphine, Lènda, Jean, Xavier, Stéphanie**, ainsi que ceux qui sont partis, notamment, **Vincent** et **Laurence**.*

*Comme c'est souvent la seule chose que les gens lisent lorsqu'ils ont un mémoire de thèse entre les mains, je me garderai bien de faire des remerciements trop personnels. Je remercie ma famille, en particulier **mes parents, ma sœur** et **mes neveux** et bien sûr mes amis indispensables, **Fred, Julie, Flo, Karen, Laure, Fredo, Medhi, Laurence, Daphné, Ivan, Louisa, Ram, Yan, Sophie, Loule, Latif, Nico** et **Nidal**.*

A ceux qui apparaissent, disparaissent, réapparaissent ou ne font que passer. Et à tous ceux qui sans le savoir ont pu influencer de manière indirecte ce travail.

TABLE DES MATIÈRES

Avant-propos.....	5
--------------------------	----------

PREMIÈRE PARTIE

Chapitre I : La perception de scènes visuelles complexes.....	9
1. Les représentations visuelles.....	10
1.1. Des représentations visuelles éparées et volatiles vs. des représentations visuelles stables et détaillées.....	12
1.2. Arguments apportés par les données empiriques.....	17
1.3. La perception implicite : quel rôle dans les représentations visuelles ?.....	19
2. La sélection attentionnelle lors de l'exploration d'une scène visuelle.....	21
2.1. Le rôle de l'attention sélective sur le traitement de l'information visuelle.....	21
2.2. Guidage ascendant vs. guidage descendant de l'attention.....	22
2.3. Précocité des influences sémantiques sur la capture attentionnelle.....	27
2.4. Interaction entre traitements ascendants et traitements descendants au cours de l'analyse de scènes visuelles.....	30
3. Le rôle des régularités contextuelles dans la perception visuelle.....	31
3.1. Le contexte facilite la reconnaissance et l'identification des objets.....	32
3.2. La construction du contexte.....	36
3.3. Le rôle du contexte dans le guidage de l'attention.....	40
4. Conclusion et perspectives.....	42

Chapitre II : L'apprentissage implicite de régularités de l'environnement....	45
1. L'apprentissage implicite de régularités.....	46
1.1. L'apprentissage de grammaires artificielles (AGL).....	47
1.2. L'apprentissage séquentiel (SRT).....	48
2. L'apprentissage implicite : un débat terminologique.....	49
3. L'apprentissage implicite produit-il une connaissance inconsciente ?.....	51
3.1. Méthodes subjectives à la première personne.....	51
3.2. Méthodes objectives à la troisième personne.....	53
4. L'apprentissage implicite est-il un processus automatique, qui ne requiert pas d'attention ?.....	56
4.1. L'apprentissage implicite consomme-t-il des ressources attentionnelles ?.....	58
4.2. L'apprentissage implicite requiert-il une attention sélective ?.....	60
4.3. L'attention est-elle requise pour l'apprentissage implicite et/ou pour l'expression de la connaissance résultante ?.....	61
5. L'apprentissage implicite peut-il conduire à des connaissances abstraites ?.....	62
5.1. L'apprentissage implicite de règles abstraites.....	63
5.2. Les traitements implicites basés sur des régularités sémantiques.....	65
6. Conclusion et perspectives.....	68
 Chapitre III : L'indication contextuelle.....	 71
1. Arguments en faveur du caractère implicite de l'indication contextuelle.....	72
1.1. Les tâches directes d'accès aux connaissances.....	73
1.2. L'indication contextuelle persiste au cours du temps et résiste aux effets d'interférence.....	74
1.3. L'indication contextuelle spatial requiert-il une attention visuelle sélective ?.....	75
1.4. L'indication contextuelle et le syndrome amnésique.....	76
2. Comment les régularités contextuelles facilitent-elles la recherche visuelle ?.....	78
2.1. Le contexte guiderait l'attention de manière efficace.....	78

2.2. Cependant.....	79
3. Portée des effets d'indication contextuel dans des environnements arbitraires.....	82
3.1. Indication contextuel spatial.....	83
3.2. Indication contextuel basé sur l'identité spécifique des éléments contextuels.....	85
3.3. Indication contextuel temporel.....	87
4. L'indication contextuel dans des scènes du monde réel.....	88
4.1. Indication contextuel dans des scènes systématiquement répétées.....	88
4.2. Pourquoi les performances sont-elles différentes dans les scènes du monde réel ?.....	90
5. L'indication contextuel reflèterait une automatisation de la recherche visuelle.....	91
6. Perspectives.....	94

DEUXIÈME PARTIE

Problématique.....	97
---------------------------	-----------

Chapitre IV : Indication contextuel basé sur des régularités spécifiques et catégorielles d'environnements numériques.....	100
Partie 1 : Indication contextuel basé sur des régularités spécifiques.....	101
Expérience 1a.....	102
Expérience 1b.....	107
Partie 2 : Indication contextuel basé sur des régularités catégorielles.....	110
Expérience 2a.....	112
Expérience 2b.....	115
Expérience 3.....	118
Discussion générale.....	122

Chapitre V : Indicage contextuel sémantique et attention visuelle sélective dans des environnements lexicaux	126
Expérience 4.....	127
Expérience 5.....	135
Expérience 6.....	139
Expérience 7.....	147
Discussion générale.....	152
 Chapitre VI : Les schémas de scène : support de l'adaptabilité de la conduite et source de dysfonctionnements.....	 158
Expérience 8.....	162
Discussion.....	167
 Chapitre VII : Discussion générale.....	 171
1. Synthèse des résultats obtenus.....	172
2. Apprentissage implicite de régularités sémantiques.....	175
3. Apprentissage implicite et attention visuelle sélective.....	178
4. Le rôle de la conscience dans l'apprentissage des régularités.....	182
5. Le rôle des régularités dans l'analyse de scènes visuelles.....	185
6. Les représentations visuelles.....	187
7. Perspectives de recherche.....	190
 Références.....	 193
 Annexes.....	 212

AVANT-PROPOS

Au cours de l'analyse d'une scène visuelle, l'attention semble bien souvent guidée sans l'intervention d'une conscience réflexive de l'analyse de la scène. Cependant, si les mouvements des yeux sont en partie guidés par des traitements qui ne sont pas le fruit d'un contrôle intentionnel, peut-on pour autant dire que l'observateur n'est pas conscient des influences cognitives qui guident sa perception de la scène ? Quel rôle peut-on accorder aux traitements non conscients dans la perception visuelle ? Le travail exposé dans ce mémoire s'inscrit dans cette problématique. Les recherches menées visaient à explorer dans quelle mesure des mécanismes d'apprentissage implicite peuvent être déployés sur des régularités de l'environnement et si des connaissances inaccessibles à la conscience influencent les processus de sélection attentionnelle au cours de l'analyse de scènes visuelles.

Nos travaux trouvent leur origine dans deux domaines de recherche, celui de la perception visuelle et celui de l'apprentissage implicite. Les recherches dans le domaine de la perception visuelle amènent au constat suivant. L'homme évolue avec aisance dans un monde visuel riche, complexe et dynamique, alors que nous savons aujourd'hui que les représentations visuelles sont loin d'être une copie fidèle du monde qui se projette sur la rétine. Par exemple, lors d'une activité de conduite automobile, les processus attentionnels sélectionnent, dans un environnement visuel riche et en perpétuel remaniement, les informations permettant la mise en œuvre de comportements adaptés. Pourtant peu d'information semble accessible à la conscience à un instant donné. L'hypothèse la plus intuitive face à cette contradiction consiste à attribuer aux traitements non conscients un rôle central dans nos comportements et notamment dans la perception visuelle. En effet, l'homme est plus souvent à même d'interagir avec son environnement qu'il n'est capable d'expliquer comment. De surcroît, de nombreuses recherches montrent qu'il devient rapidement sensible aux régularités présentes dans son

environnement, même s'il est bien souvent incapable de rapporter explicitement ces régularités. Toutefois, les recherches dévolues à l'étude des traitements non conscients ont conduit certains auteurs à remettre en question un point de vue courant accordant aux influences inconscientes un rôle prépondérant sur nos actes. Certains vont même jusqu'à rejeter l'existence d'un inconscient cognitif et de connaissances inaccessibles à la conscience. La vie mentale consciente serait suffisante pour expliquer les phénomènes cognitifs. Sans attribuer aux traitements non conscients des capacités aussi sophistiquées qu'aux traitements conscients, les données empiriques offrent toutefois des arguments tangibles en faveur de l'existence de traitements non conscients susceptibles d'influencer nos comportements. Les travaux récents conduits avec le paradigme d'indiciage contextuel notamment, montrent de manière convaincante que des connaissances sur des régularités contextuelles peuvent être acquises de manière implicite au cours de l'analyse d'une scène visuelle. Mais dans quelle mesure des traitements échappant à toute expérience phénoménale peuvent-ils rendre compte des comportements ? Quel sont les niveaux de profondeur des traitements cognitifs non conscients ? De manière plus spécifique, quelle est la nature des connaissances inaccessibles à la conscience ? Des mécanismes d'apprentissage implicite peuvent-ils être déployés sur des régularités conceptuelles de l'environnement ou sont-ils limités à des aspects spécifiques et notamment perceptifs ?

La première partie de notre travail examinera la littérature autour de problèmes abordés dans le domaine de la perception visuelle et de l'apprentissage implicite. Le premier chapitre sera articulé autour du problème des représentations visuelles et de la sélection attentionnelle, en dégagant l'influence des connaissances relatives aux régularités contextuelles sur la perception visuelle. Le deuxième chapitre fera état de débats théoriques et empiriques inscrits dans le domaine de l'apprentissage implicite qui ont motivé nos recherches. Le troisième chapitre constitue une revue de la littérature sur le paradigme d'indiciage contextuel, paradigme que nous avons utilisé dans nos travaux expérimentaux. Après avoir présenté le cadre conceptuel, nous exposerons dans une deuxième partie les recherches expérimentales que nous avons conduites. Le quatrième et le cinquième chapitre présenteront deux études comportementales visant, d'une part, à tester si des connaissances relatives à des régularités spécifiques ou catégorielles et sémantiques peuvent être acquises implicitement, et d'autre part, à tester si ces mécanismes d'apprentissage implicite requièrent ou non une attention sélective. Dans le sixième chapitre sera décrite une étude inscrite dans un objectif plus

appliqué, visant, dans un premier temps, à explorer le développement d'un apprentissage de régularités contextuelles dans une tâche plus proche de celles impliquées dans une activité écologique, telle que la conduite automobile, et dans un deuxième temps, à tester l'hypothèse que les connaissances relatives aux régularités contextuelles peuvent parfois être source de défaillances fonctionnelles. Enfin, dans le dernier chapitre, nous concluons sur nos travaux de recherche en essayant d'en dégager leurs contributions dans les domaines de la perception visuelle et de l'apprentissage implicite vs. explicite.

PREMIÈRE PARTIE

CHAPITRE I

LA PERCEPTION DE SCÈNES VISUELLES COMPLEXES

Le monde visuel est composé d'une multitude de détails, de traits, de formes, d'objets et d'évènements en perpétuel remaniement. Lorsque nous observons une scène visuelle, notre expérience subjective nous renvoie à la complexité de notre environnement. Pourtant, la vision précise et détaillée est limitée à une partie très restreinte du champ visuel, la fovéa, et se dégrade rapidement lorsque l'on s'en éloigne. La perception détaillée de la scène et des objets qui la composent requiert un déploiement attentionnel. De cette réalité biologique découlent plusieurs questions. Comment le contenu des multiples fixations oculaires est-il intégré entre deux saccades pour former une représentation unifiée, stable et continue de l'environnement visuel ? Quelle information est mise en mémoire au cours des saccades oculaires ? Quels facteurs influencent le guidage de l'attention ?

Deux types de processus sont susceptibles d'orienter les processus attentionnels : d'une part, des processus ascendants, dirigés par les données, par les caractéristiques physiques d'une scène ou d'un objet, et d'autre part, des processus descendants, dirigés par les buts, les connaissances et les attentes de l'observateur. Si certaines propriétés physiques, certains traits saillants peuvent effectivement attirer l'attention, il est largement admis que l'attention est en grande partie guidée par des connaissances et des attentes relatives à la scène, et notamment, par des connaissances sur les régularités de l'environnement.

En effet, si le monde visuel est composé d'une multitude de détails, il n'en demeure pas moins fortement structuré. La présence de régularités contextuelles rend le monde visuel stable et prédictible. Notre compréhension immédiate du monde et l'orientation efficace des processus attentionnels pourraient ainsi tenir à l'existence de contextes signifiants structurés.

1. Les représentations visuelles

Notre impression phénoménologique lorsque nous observons une scène visuelle, est que le monde extérieur nous apparaît tel qu'il est, dans sa complexité et sa globalité. De surcroît nous naviguons avec aisance dans ce monde visuel riche et dynamique, en dépit d'une vision précise et détaillée limitée à la fovéa (Anstis, 1998). Compte tenu de cette adaptabilité, il était présumé dans les années 70 que notre perception stable et continue du monde au cours des saccades oculaires, tiendrait à l'existence en mémoire d'un tampon visuel intégrant l'ensemble des représentations sensorielles obtenues à partir des fixations multiples (pour une revue, Feldman, 1985). Chacune des représentations sensorielles serait ainsi retenue et organisée en accord avec la position dans laquelle elle a été encodée pour former une image composite, globale et détaillée de la scène (*hypothèse de fusion spatiotopique*, Irwin, 1992). L'attention serait alors guidée dans un environnement auquel l'observateur aurait accès dans son intégralité dès lors qu'il l'aurait abordé.

Aussi séduisante qu'elle puisse être, cette conception a néanmoins été confrontée à de nombreux arguments empiriques qui la rendent peu plausible. Depuis les années 80, de nombreux travaux ont révélé les limites drastiques de la mémoire visuelle entre deux fixations visuelles. Par exemple, les travaux d'Irwin, Yantis et Joindes (1983) ont montré que l'individu est incapable de combiner des patterns visuels simples au cours d'une saccade oculaire (voir aussi, Irwin, 1991, 1992 ; Irwin & Andrews, 1996). Leurs résultats expérimentaux sont incompatibles avec une conception selon laquelle la perception reposerait sur la superposition et la fusion du contenu visible retenu au cours des fixations oculaires successives. La mémoire transsaccadique ne permet pas d'intégrer les informations sensorielles successives et le système visuel ne construit pas une copie fidèle de la scène.

Cette conclusion est étayée par une abondante littérature révélant que des éléments particulièrement saillants d'un point de vue perceptif peuvent passer totalement inaperçus des participants. Par exemple, des participants ayant reçu pour instruction de compter le nombre de passes échangées par les membres d'une même équipe de basket, étaient incapables de percevoir un gorille traversant la scène en gesticulant (Simons & Chabris, 1999). Ce phénomène de « cécité inattentionnelle » (Mack & Rock, 1998) a été particulièrement exploré au moyen de tâches de détection de changements, très en vogue ces dernières années. Les tâches de détection de changements (e.g. le « flicker paradigm », Rensink, O'Regan, & Clark, 1997) mettent en exergue l'extrême difficulté que rencontrent parfois les sujets à détecter, localiser ou identifier un changement (e.g. addition, suppression, déplacement, rotation, changement de couleur) opéré sur un objet dans une scène (pour des revues, cf. Simons & Rensink, 2005 ; Rensink, 2002). Ce phénomène a été baptisé par Grimes (1996) « cécité au changement » suite à ses travaux montrant que 50% des participants ne remarquaient pas que les têtes des protagonistes d'une photographie présentée à plusieurs reprises, étaient interchangées, alors même que la consigne les invitaient à rechercher des changements. La cécité au changement est un phénomène robuste qui se manifeste aussi bien dans des scènes du monde réel que dans des affichages composés de stimuli sans signification, dès lors que le changement est concomitant avec une interruption empêchant la perception du mouvement qu'il occasionne. Ce phénomène a été observé quel que soit le type d'interruption, qu'elle soit le fait d'une saccade oculaire (Grimes, 1996 ; Hollingworth & Henderson, 2002), d'un clignement de paupières (Rensink, O'Regan, & Clark, 2000) ou qu'elle résulte de manipulations expérimentales comme l'administration d'un blanc (Rensink et al., 1997, 2000), d'une coupure dans un film (Levin & Simon, 1997) ou même lorsqu'il n'y a pas d'interruption réelle entre les différentes versions d'une même image si le changement se fait graduellement (Simons, Franconeri, & Reimer, 2000). Les effets de cécité au changement ne se limitent pas à des modifications subtiles de la scène. Bien au contraire, ils peuvent se manifester pour des changements importants, dès lors que ces derniers ne modifient pas l'un des aspects sémantiques essentiels de la scène. Par exemple, dans une étude de Levin et Simons (1997), la plupart des observateurs ne remarquaient pas qu'un acteur dans un film était remplacé par un autre lors d'un déplacement de la caméra.

Des résultats aussi spectaculaires laissent entrevoir les raisons d'un tel engouement pour le phénomène de cécité au changement, qui, outre son caractère populaire, permet *partiellement* d'aborder expérimentalement la question de la représentation des scènes

visuelles complexes. En effet, si l'homme disposait en mémoire d'une représentation intégrée de la scène, comme le présuppose l'hypothèse des représentations globales recomposées (i.e. *hypothèse de la fusion spatiotopique*), il n'éprouverait aucune difficulté à détecter des changements aussi considérables dans son champ visuel. Or, les travaux sur la détection de changements montrent que la mémoire transsaccadique n'assure pas l'intégration sensorielle et que les représentations visuelles ne survivent pas aux saccades oculaires. Les théories actuelles ont ainsi abandonné l'idée d'une image sensorielle détaillée recomposée. Mais si toutes s'accordent sur ce point, le détail de nos représentations visuelles est loin de faire l'objet d'un consensus. Le débat oppose ici les partisans d'une conception selon laquelle nos représentations seraient éparses et volatiles (e.g. Irwin & Andrews, 1996 ; Rensink, 2000a), et les partisans d'une conception selon laquelle nos représentations seraient extrêmement détaillées (e.g. Hollingworth, 2004 ; Hollingworth & Henderson, 2002). Le débat tourne autour de deux questions centrales. D'une part, quelle information est extraite et retenue au cours des saccades oculaires ? D'autre part, qu'est-ce qui est représenté à un instant donné ? Autrement dit, quels sont le contenu et la stabilité de nos représentations visuelles ?

1.1. Des représentations visuelles éparses et volatiles vs. des représentations visuelles stables et détaillées

Des représentations éparses et volatiles

Pour de nombreux auteurs, notre médiocrité à détecter des changements parfois considérables, atteste de l'existence de représentations sommaires et volatiles. Loin d'être globales et stables, nos représentations seraient davantage locales et transitoires. Pour Noë et O'Regan (2000), les effets de cécité au changement démontrent que notre perception d'un monde à la fois riche et détaillée est illusoire. Selon eux, la richesse et la complexité résident dans le monde lui-même, et non dans les représentations qu'on en a. Noë et O'Regan défendent même un point de vue selon lequel l'homme ne dispose pas de mémoire pour les informations visuelles. Le monde lui-même servirait de mémoire externe. Quel serait l'intérêt cognitif à disposer en mémoire d'une représentation hautement détaillée de la scène, puisque les objets sont rapidement accessibles par un simple mouvement de l'œil si nécessaire ?

En accord avec un point de vue selon lequel nos représentations seraient éparées et volatiles, la *théorie du fichier objet de la mémoire transsaccadique*¹ (Irwin, 1991 ; Irwin & Andrews, 1996) et la *théorie de la cohérence*² (Rensink, 2000a) constituent certainement les théories les plus détaillées et les plus influentes aujourd'hui. Dans ces théories, l'attention est nécessaire pour que les traits sensoriels soient intégrés en une représentation cohérente d'un objet et pour que cette représentation soit maintenue en mémoire visuelle à court terme (MVCT). Mais dès lors que l'objet n'est plus sous le focus attentionnel, la représentation de cet objet se dégraderait et seul resterait le code sémantique, « l'étiquette » de cet objet. La représentation d'une scène se limiterait en fait à une représentation détaillée mais temporaire du/des objet(s) présent(s) sous le focus attentionnel, à l'étiquette conceptuelle des objets précédemment attendus³, et à une représentation très schématique de la scène dérivée de son identification et de ses propriétés sémantiques et structurales spatiales globales. La théorie du fichier objet et la théorie de la cohérence sont toutes deux étroitement inspirées de la *théorie d'intégration des traits* proposée par Treisman et ses collaborateurs (Kahneman, Treisman, & Gibbs, 1992 ; Treisman & Gelade, 1980). Elles considèrent deux niveaux de représentation. Ces deux niveaux renvoient au contenu volatil de nos représentations, et se rapportent aux deux étapes de traitement, préattentive et attentive, évoquées par Treisman et Gelade en 1980 (pour une description détaillée de la théorie de la cohérence et de la théorie du fichier objet, cf. Annexe 1a).

A l'origine, la théorie du fichier objet a été développée en vue d'expliquer l'intégration visuelle au cours des saccades oculaires. Par la suite, Irwin et Andrews (1996) ont intégré la conception de la mémoire transsaccadique dans une théorie plus générale de la représentation de scène. Dans cette perspective, l'information retenue au cours d'un mouvement saccadique de l'oeil serait limitée à trois sources :

- les fichiers objets actifs maintenant les codes visuels (abstraits de la représentation sensorielle) à partir des objets attendus ou récemment attendus stockés en mémoire à court terme (MCT),
- l'activation indépendante des nœuds conceptuels en MLT, codant l'identité des objets individuels organisés selon leur localisation dans la scène,

¹ Traduit de l'anglais : object file theory of transsaccadic memory

² Traduit de l'anglais : coherence theory

³ Pour simplifier la lecture, le terme attendu est traduit de l'anglais « attended » et signifie donc « sur lequel l'attention est portée »

- des représentations schématiques de la scène dérivées de son identification globale, i.e. son « gist », sa signification générale ou catégorie de scène (e.g. scène de cuisine).

Seuls les fichiers objet maintiendraient une représentation visuelle des objets mais leur structure serait labile. Irwin et Andrews (1996, p. 130) résument ainsi ces propos : *“According to object file theory, relatively little information actually accumulates across saccades; rather, one’s mental representation of a scene consists of mental schemata and identity codes activated in long term memory and of a small number of detailed object files in short-term memory”*.

La théorie de la cohérence décrit quant à elle le contenu hypothétique de nos représentations visuelles. Parallèlement, Rensink (2000a) a intégré les principes fondamentaux de la théorie de la cohérence dans un modèle plus complet de la perception, *l’architecture triadique*, en faisant intervenir les connaissances préalables et les buts de l’observateur. L’architecture triadique de la vision vise essentiellement à expliquer l’impression de continuité entre les traits et le guidage efficace de l’attention au sein de l’environnement, en dépit d’une représentation cohérente limitée, selon Rensink, à l’unique objet susceptible d’être présent sous le focus attentionnel. L’architecture triadique de la vision postule l’existence de trois systèmes indépendants (pour une représentation schématique, cf. Figure 1).

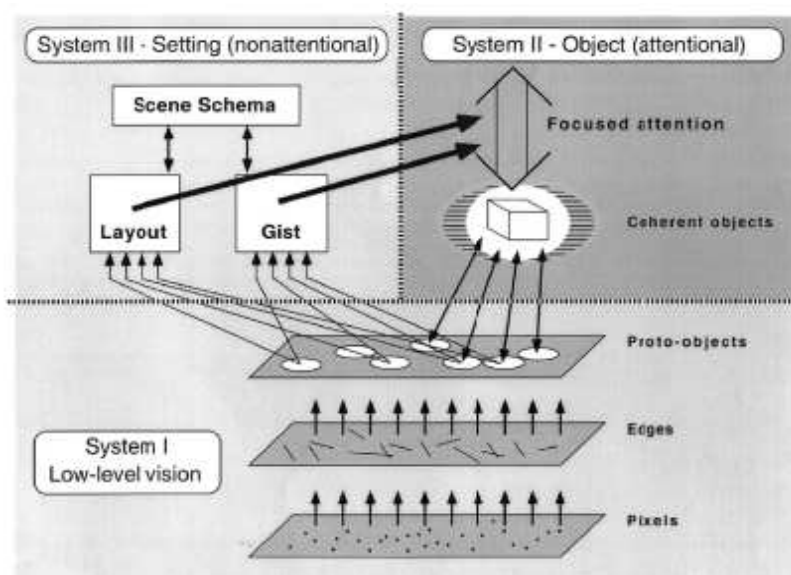


Figure 1. Architecture triadique. Adapté de Rensink (2000a)

Les deux premiers systèmes sont décrits en détail dans la théorie de la cohérence. Le premier est un système de bas niveau, conduisant à l'émergence de structures visuelles volatiles extrêmement détaillées mais peu cohérentes et disparaissant dès que la lumière cesse de pénétrer dans l'oeil. Le second système est un système attentionnel à capacité limitée qui transforme ces structures labiles en objets intégrés, i.e. en représentations stables et cohérentes. Le troisième système renvoie à un état cognitif plus stable fournissant un cadre de référence au déploiement attentionnel. Ce dernier système comprend certains aspects relatifs à la structure de la scène : sa signification générale (i.e. le « gist »), l'agencement spatial des objets qui la composent indépendamment de leurs propriétés visuelles et sémantiques, ainsi que les attentes dérivées des connaissances préalables inscrites en MLT sous forme de schémas de scène. Les schémas de scène décrivent l'ensemble des objets susceptibles d'apparaître dans une scène, ainsi que leur position relative les uns par rapport aux autres (Mandler & Parker, 1976 ; Mandler & Ritskey, 1977). Les informations visuelles véhiculées par les proto-objets⁴ activeraient rapidement une représentation schématique de la scène, résultant du traitement de la structure spatiale globale, de l'identification globale de la scène, et de la récupération d'un schéma de scène stocké en MLT. Au cours de l'analyse d'une scène, les trois systèmes interagiraient pour former une représentation émergente finalement beaucoup moins détaillée que ce que notre phénoménologie nous laisse penser. Pour Rensink, l'activation précoce d'un schéma de scène au cours du traitement donnerait cette impression de continuité entre les traits. Selon la théorie de la cohérence, la mémoire de la scène serait limitée au « gist », à l'agencement spatial et éventuellement aux identités abstraites des objets reconnus. La volatilité de nos représentations visuelles rendrait ainsi compte du phénomène de cécité au changement. Seul les changements effectués sur l'objet présent dans le faisceau attentionnel pourraient être détectés, puisque lui seul serait perçu de manière précise et cohérente.

Des représentations stables et détaillées

Les théories postulant que nos représentations visuelles sont éparses et labiles sont cependant loin de satisfaire l'ensemble de la communauté scientifique. Certains auteurs

⁴ Selon la théorie de la cohérence, les objets présents dans le champ visuel, même extérieur au faisceau attentionnel, subiraient des traitements préattentifs leur conférant un statut de proto-objets. Ces proto-objets pourraient être relativement complexes, mais présenteraient une cohérence spatio-temporelle extrêmement limitée.

pensent que loin d'être éparées, nos représentations sont au contraire hautement détaillées et qu'une grande quantité d'informations conceptuelles mais aussi visuelles est mise en MLT au cours des saccades oculaires. Pour Hollingworth et Henderson (2002 ; Hollingworth, 2004), il est invraisemblable d'envisager qu'aucune information visuelle spécifique ne soit mise en MLT. Selon eux une représentation relativement détaillée de la scène est construite à partir de l'extraction transsaccadique d'informations visuelles. Dans ce cadre, Hollingworth et Henderson ont proposé qu'à la fois la mémoire visuelle à court terme (MVCT) et la mémoire visuelle à long terme (MVLТ) jouent un rôle déterminant dans la construction d'une représentation visuelle. Comme dans la théorie du fichier objet, lorsque l'attention serait portée sur un objet de la scène, une représentation temporaire serait construite, un fichier objet, et serait maintenue en MCT. Mais de manière différente, les multiples fixations oculaires sur la scène permettraient une consolidation progressive des fichiers objets en MLT au sein d'une carte décrivant l'ensemble de la scène. Ainsi, l'accumulation des informations visuelles et conceptuelles relatives aux objets attendus conduirait à une représentation détaillée de la scène (pour une description plus détaillée de la *théorie de la perception de scène et de la mémoire* proposée par Hollingworth et Henderson, cf. Annexe 1b). Ici encore, l'attention jouerait un rôle fondamental à la fois dans l'encodage et dans la récupération des informations visuelles et conceptuelles : l'accès au contenu d'un fichier objet en MVCT dépendrait d'une attention spatiale allouée à l'objet indexé sous forme de fichier.

La *théorie de la perception de scène et de la mémoire* proposée par Hollingworth et Henderson, prédit donc que des changements concernant les propriétés visuelles (e.g. changement d'orientation, de couleur, de forme) d'un objet pourront être détectés et que les performances de détection seront d'autant meilleures que l'objet aura été traité de manière attentionnelle. La différence majeure avec les modèles de « représentations transitoires » tient à l'intervention de la MLT, permettant aux représentations visuelles de persister lorsque l'attention est détournée de l'objet. Le système visuo-cognitif tirerait ainsi avantage des capacités illimitées de la MLT pour accumuler de plus en plus d'informations pertinentes pour de futures explorations.

1.2. Arguments apportés par les données empiriques

Bien qu'a priori très divergentes, les théories postulant des représentations éparées et celles postulant des représentations détaillées partagent certains points communs. Dans les trois théories présentées précédemment, à la fois l'attention sélective et les connaissances activées précocement jouent un rôle majeur dans la construction des représentations des scènes visuelles. La différence fondamentale entre ces deux familles de théories (i.e. des représentations éparées et volatiles vs. stables et détaillées) réside dans la nature de l'information mise en MLT. Selon la théorie de la cohérence, aucune information perceptive n'est encodée en MLT. En revanche, selon la théorie de la perception de scène et de la mémoire, l'information visuelle disponible en MCT peut être consolidée en MVLT lorsque l'exposition de la scène est prolongée.

Même si la théorie de la cohérence présente certains attraits – notamment par l'idée que nos représentations sont en réalité évanescences – sa validité écologique reste, selon nous, assez peu plausible. Un certain nombre de données expérimentales sont probablement davantage compatibles avec les postulats décrits dans la théorie de la perception de scène et de la mémoire. En effet, la littérature sur la mémoire des images montre que l'homme est doté d'une prodigieuse capacité à se souvenir d'images présentées une seule fois (e.g. Standing, Conezio, & Haber, 1970 ; Shepard, 1967). Néanmoins, si la MLT pour les images de scènes visuelles perçues à la fréquence d'une par seconde est remarquablement bonne, lorsque ces dernières sont présentées pendant des durées très brèves (e.g. 125-333ms) la plupart d'entre elles ne sont plus accessibles au moment de la tâche de reconnaissance et semblent être oubliées (Potter, 1976 ; Potter & Levy, 1969 ; Potter, Staub, Rado, & O'Connor, 2002). En explorant le type d'information visuelle vs. conceptuelle retenu en MLT, Potter, Staub et O'Connor (2004) ont montré que lorsque les images sont présentées durant 173ms, si la signification générale semble persister quelques secondes en MCT, l'information « picturale » semble être rapidement oubliées. Cependant, lorsque ces dernières étaient présentées une seconde chacune, à la fois l'information conceptuelle et l'information perceptive étaient retenues en MVCT, puis consolidées en MLT. Il est certain qu'une exposition prolongée à une scène entraîne une accumulation et une consolidation des informations en mémoire. Les performances de reconnaissance des images (Loftus, 1972) ou des objets (Irwin & Zelinsky, 2002) s'améliorent en fonction du nombre de fixations réalisées sur l'image. De plus, Castelhamo et Henderson (2007) ont montré que la recherche d'un objet au sein d'une scène

visuelle est facilitée lorsque cette même scène est présentée brièvement (250ms) en amorçage. Par contre, lorsque la scène amorce n'était pas identique à la scène de recherche mais définie par la même catégorie sémantique, aucun bénéfice n'était observé, suggérant que les propriétés spécifiques d'une scène présentée brièvement facilitent davantage le guidage de l'attention au sein de la scène que la catégorie conceptuelle de cette scène, tout du moins dans ce type de tâche. Enfin, même si les effets de cécité au changement sont parfois spectaculaires, il arrive bien souvent que les sujets parviennent à détecter des changements sur des aspects visuels réalisés en dehors du focus attentionnel, et sans que les sujets n'encodent intentionnellement la propriété critique de l'objet avant que ce dernier ne change (e.g. Hollingworth & Henderson, 2002 ; Hollingworth, Williams, & Henderson, 2001 ; Tatler, Gilchrist, & Land, 2005 ; pour des arguments contraires, cf. Levin & Simons, 1997). De surcroît, les performances de détection de changements augmentent avec le nombre de fixations et avec la durée totale de fixations. Dans les travaux utilisant des paradigmes de détection de changements, les différences importantes entre les travaux de Rensink et collaborateurs et ceux de Hollingworth et collaborateurs pourraient résider dans les durées de présentation des scènes antérieures au changement (e.g. 240 ms dans Rensink et al., 2000, contre 20s dans Hollingworth & Henderson, 2002). Les effets d'accumulation et de consolidation pourraient expliquer les importantes différences dans les performances observées dans leurs travaux respectifs. Cette différence n'en reste pas moins congruente avec la théorie de la perception de scène et de la mémoire proposée par Hollingworth et Henderson, postulant que la mémoire ne se limite pas à la prise en compte des aspects conceptuels de la scène et que l'information peut être consolidée en MLT lorsque l'exposition de la scène est prolongée.

A noter que dans plusieurs revues de question récentes, Rensink et Simons (e.g. Simons & Ambinder, 2005 ; Simons & Rensink, 2005) ont fortement minimisé leurs conclusions initiales sur le phénomène de cécité au changement. Rensink et Simons (2005, p. 19) résument ainsi leur relecture des effets classiques de cécité au changement : *“In summary, the failure to detect change does not imply the absence of a representation unless all the requirements of scope are met. Given that no studies have done this, extend work on change blindness can say nothing about representations that subserve the static aspects of vision – it can only reveal limits on the representations that subserve the conscious perception of dynamic change. In fact, mounting evidence suggests that some information is preserved even when observers fail to detect changes”*. En effet, l'existence de cécité au changement

n'implique pas nécessairement l'existence de représentations éparses : des échecs à détecter des changements pourraient se manifester en dépit de représentations visuelles complètes et détaillées de la scène.

1.3. La perception implicite : quel rôle dans les représentations ?

Selon de nombreux auteurs, la cécité au changement ne signifie pas nécessairement que l'information critique soit absente de la représentation de la scène, mais que la détection explicite n'est pas toujours sensible à la présence de cette information. Lorsque l'on utilise des mesures alternatives (e.g. tests de choix forcés, durées de fixations sur la région subissant un changement), les participants semblent capables de détecter plus de changements qu'ils ne sont capables d'en rapporter verbalement (e.g. Fernandez-Duque & Thornton, 2000, 2003 ; Hollingworth & Henderson, 2002 ; Rensink, 2004 ; VanRullen & Koch, 2003 ; cependant, Mitroff, Simons, & Franconeri, 2002). Selon Simons et Ambinder (2005, p. 47), *“our conscious awareness of our visual environment is sparse even if our representations of it might not be”*. VanRullen et Koch (2003) ont proposé que les différents objets de la scène soient codés à des niveaux de représentation différents, accessibles par des mesures différentes et ayant des effets différents sur les comportements des participants. Leurs travaux montrent en effet que si les sujets ne peuvent rappeler que deux ou trois objets relatifs à une scène perçue très brièvement (250ms), ils seraient en mesure d'en reconnaître deux ou trois supplémentaires dans un test de choix forcé. Mais plus intéressant encore, en ayant recours à une tâche d'appariement, les auteurs ont mis en évidence des effets d'amorçage négatif par les objets non rappelés et non reconnus. Ce résultat suggère qu'une représentation de ces objets, inaccessible par des tests de rappel et de reconnaissance, peut néanmoins modifier le traitement de ce même stimulus lorsqu'il est de nouveau présenté aux participants. De nombreuses recherches suggèrent en effet que des informations inaccessibles à la conscience peuvent néanmoins être traitées par le système visuo-cognitif et influencer le comportement. Une littérature abondante rapportant des effets d'amorçage subliminaux offre des arguments en faveur de l'existence d'une perception implicite (pour une revue, Merikle & Daneman, 1998). A l'appui de ces résultats comportementaux, des études utilisant des méthodes d'imagerie cérébrale montrent que des entrées sensorielles traitées par le cerveau n'atteignent pas toujours un niveau de conscience permettant à l'individu un report explicite (pour une revue, Dehaene, Changeux, Naccache, Sakur, & Sergent, 2006). Ainsi, des stimuli

« invisibles » peuvent activer un ensemble de neurones sans pour autant que les traitements réalisés sur ces stimuli ne conduisent à l'émergence d'une expérience subjective. Pour Hollingworth et Henderson (2002) les effets de perception implicite démontrent que les performances de détections explicites sous-estiment le détail de nos représentations visuelles. L'idée que tous les aspects du traitement visuel ne conduisent pas nécessairement à une expérience visuelle n'est pas rejetée par Rensink. Bien au contraire, celui-ci défend l'existence d'une perception implicite, i.e. le traitement d'une information visuelle sans être accompagné d'une expérience visuelle ou d'une quelconque prise de conscience (Rensink, 2000b). Rensink (2000b) propose ainsi que les observateurs peuvent présenter une capacité à détecter des changements alors qu'ils ne manifestent aucune expérience consciente de ces changements (voir aussi Rensink, 2004).

Même si la littérature sur la détection de changements suggère que nos représentations visuelles conscientes sont plus pauvres que ce que notre phénoménologie nous laisse croire, certains aspects de la scène ne faisant pas l'objet d'une image mentale semblent néanmoins être traités par le système visuel et influencer les traitements ultérieurs. Peut-on pour autant dire que leur traitement donne lieu à une représentation ?

En outre, si le détail de nos représentations visuelles est loin de faire l'objet d'un consensus, la plupart des auteurs s'accordent sur le rôle déterminant de l'attention dans l'élaboration de nos représentations conscientes stables et cohérentes. En permettant aux traits visuels présents sous son joug d'être intégrés pour atteindre un état stable, l'attention enrichirait nos représentations d'une cohérence spatio-temporelle, leur assurant une sorte de persistance cognitive. Dans cette perspective, l'attention jouerait un rôle déterminant en modifiant la qualité et la nature des représentations visuelles sur lesquelles elle s'exerce. William James définissait l'attention en 1890 de la manière suivante: *"it is the taking possession by mind, in clear and vivid form of one out of what seem several simultaneously possible objects or train of thoughts. Focalisation, concentration, of consciousness are of its essence"*. Une attention focalisée sur les objets de la scène serait nécessaire à la formation d'une représentation cohérente de ces objets. Mais comment l'attention est-elle orientée au cours de l'exploration d'une scène visuelle ?

2. La sélection attentionnelle lors de l'exploration d'une scène visuelle

Les recherches actuelles sur la perception de scènes visuelles révèlent que nos représentations conscientes sont plus schématiques que ce que notre phénoménologie nous laisse penser. Afin de construire une représentation détaillée des éléments et événements présents dans notre champ visuel, l'attention sélective doit être déployée au sein de l'environnement. L'attention peut être abordée sous de multiples facettes, e.g. l'attention d'éveil, l'attention de contrôle, l'attention sélective. Nos travaux portant sur l'attention en tant que processus de sélection, nous limiterons notre approche à deux aspects : d'une part aux changements quantitatifs/qualitatifs que l'attention sélective applique à l'information ou aux traitements qui sont sous son contrôle, et d'autre part aux facteurs qui sont susceptibles d'orienter l'attention sélective dans l'espace.

2.1. Le rôle de l'attention sélective sur le traitement de l'information visuelle

Les recherches qui examinent l'attention visuelle sélective s'appuient couramment sur la métaphore du *spotlight* (Posner, 1980). La métaphore du spot de lumière rend compte de certains aspects phénoménologiques introspectifs de l'attention, comme par exemple la sensation que l'attention, à l'instar d'un faisceau de lumière mental, peut éclairer certaines régions de la scène. A la manière d'un zoom, la profondeur du focus attentionnel est ajustable par l'individu. L'attention peut ainsi être allouée à des régions de taille différente, mais avec la sensation que sa distribution spatiale diminue de manière graduelle lorsque augmente l'excentricité de son locus (qu'on pourrait assimiler à la fovéa). Cependant, étant donné que l'attention est dirigée sur les objets plutôt que sur les localisations, il paraît difficile de concevoir de manière littérale l'attention comme un spot de lumière. Si cette métaphore peut présenter une certaine valeur heuristique, elle doit être vue comme une figure de rhétorique et non comme un modèle de l'attention (Chun & Wolfe, 2001). Elle met néanmoins en exergue différentes propriétés centrales de l'attention sélective, et notamment celle de rehausser l'efficacité du traitement des événements présents au sein du faisceau de lumière.

Il est couramment envisagé que le rôle de l'attention sélective est de surmonter les limites des capacités de traitement. Dans cette perspective, puisque le champ visuel contient plus

d'information que le système visuel n'est capable de traiter à un instant donné, l'attention sélective serait nécessaire pour privilégier le traitement des informations pertinentes et pour inhiber les informations distractrices non pertinentes compte tenu des objectifs en cours (pour des revues, Chun & Wolfe, 2001 ; Pashler, 1998). Dans les modèles suggérant l'existence de limitations - exprimés sous forme de ressources (Kahneman, 1973), de filtre (Broadbent, 1958) ou de capacité limitée de mémoire de travail (Baddeley, 1993) - les processus attentionnels dicteraient les priorités de traitement. La priorité de traitement faciliterait le traitement d'une information et permettrait à l'information d'occuper à elle seule la totalité du champ de la conscience. L'attention « boosterait » en quelque sorte l'information attendue en lui permettant de 'gagner la course' au dépend des autres informations distractrices présentes, sur lesquelles elle exercerait un blocage, une atténuation, suppression ou inhibition.

Pour certains auteurs (e.g. Cowan, 1988 ; Posner & Dehaene, 1994 ; pour une revue, Camus, 2003), l'action de rehaussement de l'attention est considérée indépendamment des limites des capacités de traitement. Il s'agit d'une importante différence théorique dans le sens où l'attention aurait davantage un rôle d'amplification, de « magnification » sur le traitement de l'information attendue qu'un rôle d'atténuation sur le traitement des informations ignorées. L'information attendue est plus saillante perceptivement, les traitements sont plus profonds et les réponses pertinentes facilitées. Les modèles reposant sur l'existence de sources bruitées (internes ou externes) font intervenir des notions intéressantes, notamment par leur plausibilité neurophysiologique, comme celle d'activation, de sur-activation et de désactivation des traces mentales. Ces notions impliquent l'existence de seuils. En dessous d'un certain seuil d'activation, l'information n'est pas traitée, puis devient traitée de manière implicite, puis de manière explicite. Les processus attentionnels moduleraient les niveaux d'activation des informations et des représentations. Seule l'information attendue atteindrait ainsi un seuil explicite et conscient de traitement.

2.2. Guidage ascendant vs. guidage descendant de l'attention

Quels mécanismes sous-tendent la sélection attentionnelle ? Quels facteurs influencent le guidage de l'attention au sein de l'environnement ? On distingue classiquement des mécanismes ascendants, guidés par des facteurs exogènes, i.e. certaines caractéristiques saillantes de l'environnement visuel, et des mécanismes descendants, guidés par des facteurs

endogènes, i.e. les connaissances, les attentes et les objectifs de l'observateur. L'orientation exogène de l'attention serait automatique et involontaire. Par opposition, on considère souvent que l'orientation endogène de l'attention est intentionnelle. Il paraît cependant extrêmement réducteur de limiter l'orientation endogène de l'attention à des processus sous l'emprise d'un contrôle délibéré. De plus, l'orientation exogène de l'attention dépend toujours peu ou prou de l'état cognitif de l'observateur. Il est souvent difficile, même dans une situation peu complexe, de faire la part entre ce qui relève de mécanismes purement ascendants et ce qui relève de mécanismes purement descendants.

Guidage ascendant de l'attention

Il est cependant certain que l'attention peut être capturée de manière automatique par des stimuli extérieurs présents dans le champ visuel. Cette orientation exogène de l'attention (aussi appelée capture attentionnelle), déclenchée notamment par l'arrivée impromptue d'un stimulus dans le champ visuel et/ou par un stimulus saillant visuellement, présente un caractère irrépressible. Les lumières d'un véhicule dans la nuit attirent l'attention de manière automatique. Ou encore, l'apparition soudaine d'un événement dans le champ visuel constitue un facteur responsable de l'orientation exogène de l'attention (pour une revue, Egeth & Yantis, 1997). La survie de l'espèce nécessitant une aptitude à s'orienter rapidement vers des événements ou éléments saillants et/ou impromptus de la scène visuelle, on peut aisément comprendre pourquoi cette faculté est apparue et s'est maintenue au cours de l'évolution phylogénétique.

Si certains attributs visuels, comme la couleur rouge, peuvent constituer un trait saillant susceptible d'attirer de manière irrépressible l'attention, la saillance d'un objet ou d'un attribut est somme toute relative au contexte dans lequel il apparaît. Par exemple, un élément en mouvement dans un ensemble d'éléments statiques sera un facteur susceptible de capturer de manière automatique l'attention. Mais si cet élément en mouvement est présenté parmi un ensemble d'autres éléments en mouvement, il ne constituera plus un trait saillant du champ visuel. C'est donc plutôt le degré de ressemblance/dissemblance entre un élément cible et son contexte qui le rend susceptible d'attirer de manière exogène l'attention et donc le rend saillant perceptivement. Les singletons (e.g. une cible rouge parmi des distracteurs verts ou une cible verticale parmi des distracteurs horizontaux) constituent ainsi des déterminants de la

capture exogène de l'attention. Par exemple, Theeuwes (1992) demandait aux participants de rechercher un rond vert cible parmi des carrés verts distracteurs. Sur certains essais, l'un des carrés verts était remplacé par un distracteur rouge. Theeuwes observait alors qu'au détriment des performances, le singleton rouge non pertinent pour la tâche, capturerait de manière irrésistible l'attention et réduisait l'efficacité de la recherche. Des données neurophysiologiques étayaient les effets de capture exogène en montrant qu'un signal saillant visuellement provoque une plus forte réponse neuronale qu'un stimulus moins saillant (pour une revue, Desimone & Duncan, 1995).

Jonides et Yantis (1988) ont cependant montré que les singletons capturent plus volontiers l'attention s'ils partagent une propriété visuelle commune avec la cible. Si les événements extérieurs ne sont pas pertinents d'un point de vue comportemental, leur faculté à capturer l'attention sera minimisée. De plus, même si les singletons sont plus faciles à ignorer que les indices spatiaux ou les apparitions soudaines, il apparaît que le degré auquel les stimuli apparaissant brusquement dans le champ visuel capturent l'attention peut être considérablement modulable par un contrôle volontaire et semble dépendre de la stratégie de recherche adoptée par le sujet (pour une revue, Egeth & Yantis, 1997). Enfin, si des effets de guidage ascendant peuvent être observés dans des affichages relativement homogènes composés de stimuli simples, dès lors que les affichages sont plus complexes, les études échouent souvent à montrer une orientation purement exogène de l'attention (Theeuwes, 2004). Les facteurs susceptibles de constituer un fort signal ascendant dans des environnements homogènes ne semblent pas ou peu exister dans des environnements plus complexes. Les observateurs semblent adopter un mode de recherche dans laquelle la capture exogène potentielle des singletons est considérablement réduite. On peut donc s'interroger sur la contribution des processus purement ascendants dans l'orientation de l'attention au sein de scènes du monde réel. Par exemple, Henderson et al. (2007) ont montré que les données sur les mouvements des yeux des participants alors que ces derniers réalisaient une tâche de recherche visuelle au sein de photographies en couleur du monde réel, ne correspondaient pas aux prédictions faites par le « modèle de carte de saillance » développé par Itti et Koch (2000). Le modèle de carte de saillance est strictement ascendant et permet d'isoler la région présentant une saillance maximale et devenant une cible privilégiée pour une saccade ou un déplacement de l'attention, i.e. une capture oculomotrice et une capture attentionnelle. Les cartes de saillance sont calculées selon différentes dimensions incluant la couleur, l'intensité, le contraste, l'orientation, les jonctions de contours, les limites des bords, les ombres, ainsi

que les caractéristiques dynamiques susceptibles de constituer un facteur ascendant. Mais si les prédictions faites par ce modèle sur les régions susceptibles d'attirer l'attention, sont assez bien corrélées avec les effets classiquement observés lorsque les cibles sont définies de manière conjonctive ou disjonctive dans des affichages « simples », il échoue à rendre compte des performances dans la grande majorité des tâches de recherche utilisant des affichages comportant une plus grande hétérogénéité perceptive (Chen & Zelinsky, 2006 ; Lamy, Leber, & Egeth, 2004 ; Turano, Gerguschat, & Baker, 2003).

Guidage descendant de l'attention

Par opposition au guidage ascendant, l'attention peut être guidée de manière descendante par des facteurs endogènes et notamment via l'activation de connaissances en mémoire ou des attentes et des buts. Même si ce terme peut être discuté, on parle parfois de « contrôle cognitif » pour désigner la sélection de certains sites de fixation relativement aux besoins du système cognitif.

L'influence des connaissances et des attentes peut intervenir dans la sélection des traits visuels. Lors d'une recherche visuelle d'un objet connu, l'activation en mémoire d'une/de représentation(s) visuelle(s) associée(s) à cet objet influence le guidage de l'attention (Duncan & Humphreys, 1989). La recherche d'un objet connu, par exemple des clés, est facilitée par l'activation de leur description visuelle stockée en mémoire, avant même que ces dernières ne soient sous le regard. Des preuves empiriques relatives à ce phénomène sont apportées par des études d'enregistrement des mouvements des yeux montrant que les observateurs fixent préférentiellement les cibles ou distracteurs partageant des traits visuels communs avec la cible (e.g. Chelazzi, Duncan, Miller, & Desimone, 1998). De plus, des études d'enregistrement de neurones unitaires suggèrent que des feedback dérivés des représentations disponibles en MT modulent les réponses neuronales quand un stimulus cible a été sélectionné pour une réponse (Chelazzi, Duncan, Miller, & Desimone, 1998 ; Chelazzi, Miller, Duncan, & Desimone, 1993). Desimone, Duncan et collaborateurs interprètent ces effets relatifs à l'exposition antérieure de la cible, par un traitement descendant exerçant un feed back sur les régions situées le long de la voie visuelle. Ce traitement descendant biaiserait la compétition entre les objets présents dans l'affichage de recherche, en faveur de la cible. Selon le modèle de compétition biaisée de la sélection visuelle (Desimone & Duncan,

1995), les représentations neuronales des différents objets présents dans le champ visuel entrent en compétition et s'inhibent mutuellement pour accéder à un traitement de plus haut niveau. La sélection de l'objet serait contrôlée par la pré-activation des canaux neuronaux réceptifs à un objet pertinent particulier. Ainsi, sur la base des représentations disponibles en MT (via également l'activation de connaissances en MLT), des signaux descendants semblent biaiser la sélection en faveur de l'objet pour lequel les traits ont été pré-activés et donc « gagnant » la compétition sélective entre les objets dans le champ visuel (pour un modèle de recherche guidée par les traits visuels de la cible, cf. Wolfe, Cave, & Franzel, 1989).

Moore, Laiti et Chelazzi (2003) ont récemment proposé d'intégrer au modèle de compétition biaisée des composantes sémantiques dans la sélection attentionnelle. À l'aide d'un paradigme de recherche visuelle (pour un exemple, cf. Figure 2), les auteurs ont montré que des liens associatifs de type sémantique entre les représentations visuelles des objets peuvent affecter l'allocation attentionnelle, au même titre que des relations purement perceptives. Dans leurs expériences, les objets sémantiquement associés à la cible en MLT (e.g. trompette associée à la cible piano) étaient plus susceptibles d'être sélectionnés que des objets contrôles non reliés sémantiquement à la cible. Les objets associés à la cible semblaient également être traités en tant que cibles potentielles au niveau des mécanismes décisionnels. Les auteurs défendent l'idée selon laquelle la compétition entre les traits visuels pourrait reposer sur des relations conceptuelles en plus des relations visuelles.

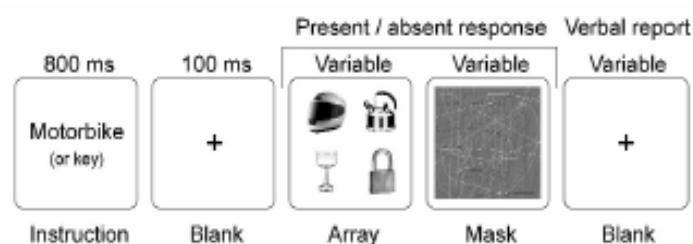


Figure 2. Procédure expérimentale utilisée par Moore, Laiti et Chelazzi (2003) dans l'Expérience 1. Les participants étaient exposés à une étiquette verbale (e.g. moto), puis ils réalisaient une tâche de recherche visuelle dans un affichage composé de 4 items objets parmi lesquels était ou non présent l'objet cible pré-spécifié par le nom amorce (ici moto). La réponse faisait apparaître un masque. Puis ils devaient rapporter le mieux possible les 4 objets présents dans l'affichage de recherche. Parmi les objets présents dans l'affichage, pouvaient être présents, l'objet cible, un objet associé à la cible (e.g. un casque de moto) et des objets contrôles (e.g. un verre). *Adapté de Moore et al. (2003)*

Si en effet l'attention peut être guidée de manière descendante par des connaissances sur les propriétés visuelles d'un objet pertinent de la scène, elle peut aussi être guidée par des connaissances sémantiques. Dans ce cadre, les travaux de Buswell (1935) montraient déjà que

les positions des fixations oculaires sont préférentiellement portées sur les éléments porteurs de sens dans la scène. Quelques décennies plus tard, Yarbus (1967) étendait ce résultat en montrant qu'une consigne invitant les sujets à estimer l'âge des personnages conduit à des fixations précocement et préférentiellement portées sur les visages, alors qu'une consigne les invitant à évaluer le statut social des personnages conduit à des fixations davantage distribuées sur leurs vêtements. Parallèlement, les travaux de Mackworth et Morandi (1967) indiquaient que non seulement les régions les plus informatives d'une image pour la compréhension de la scène reçoivent plus de fixations que les autres, mais qu'en plus, les régions les moins informatives ne reçoivent souvent aucune fixation. Ces résultats suggèrent que les objets dits « d'intérêt central » reçoivent un traitement attentionnel privilégié (Rensink et al., 1997).

2.3. Précocité des influences sémantiques sur la capture attentionnelle

L'influence de facteurs purement sémantiques sur le contrôle de l'attention ne fait a priori aucun doute. Si la démonstration expérimentale de ces effets ne présente pas grand intérêt, la précocité de cette influence est quant à elle moins évidente. La première fixation est-elle initialement affectée par une analyse sémantique parafovéale des régions de la scène ?

Loftus et Mackworth (1978) ont rapporté la première étude destinée à explorer directement l'influence de « l'informativité sémantique » sur la position des fixations. Les participants étaient exposés à des dessins de scènes dans lesquels un objet cible était, soit très informatif sémantiquement, soit peu informatif. L'informativité sémantique était ici définie par le degré selon lequel un objet donné était prédictible au sein d'une scène, les objets non prédictibles étant considérés comme les plus informatifs. Pour ce faire, deux objets dans deux scènes différentes étaient substitués. Par exemple, une scène de ferme et une scène de milieu sous-marin pouvaient l'une et l'autre contenir, soit un tracteur, soit une pieuvre. Les participants observaient ces scènes durant 4s et effectuaient ensuite une tâche de reconnaissance (pour une illustration de stimuli utilisés dans ce type d'étude, cf. Figure 3). Loftus et Mackworth ont rapporté trois résultats majeurs. Premièrement, le nombre de fixations étaient plus important pour les régions informatives sémantiquement que pour celles moins informatives. Deuxièmement, les sujets tendaient à fixer plus longtemps les objets inconsistants sémantiquement que les objets consistants. Et troisièmement, les sujets avaient

tendance à fixer les objets inconsistants sémantiquement immédiatement après la première saccade oculaire, suggérant qu'un traitement sémantique des régions extra-fovéales peut contrôler le déplacement des fixations. Ces résultats ont amené Loftus et Mackworth à conclure que les observateurs sont à même de déterminer en une seule fixation la consistance sémantique des objets individuels avec la scène en se basant sur la vision périphérique, et que l'information extraite pourrait exercer un effet immédiat sur le contrôle des mouvements des yeux (pour une hypothèse identique, Mackworth & Morandi, 1967). Cependant, quelques années plus tard, De Graef, Christiaens et d'Ydewalle (1990) n'ont pu apporter de preuve convaincante que, dans une tâche de recherche visuelle de cible, les objets sémantiquement incohérents sont plus précocement fixés que les objets moins informatifs. Les études de Mannan et al. (1995) et de Henderson, Weeks et Hollingworth (1999) corroborent cette conclusion en montrant que les fixations initiales sont contrôlées davantage par les traits visuels que par les traits sémantiques locaux. Le placement de la première fixation ne semble pas être contrôlé par une analyse sémantique périphérique des objets individuels de la scène. En revanche, si un objet informatif sémantiquement (i.e. ici inconsistant) est présent sous le regard, ce dernier tend à être fixé plus longtemps qu'un objet non informatif (consistant). De plus, le regard tend à retourner vers les objets inconsistants au cours de l'analyse de la scène visuelle.

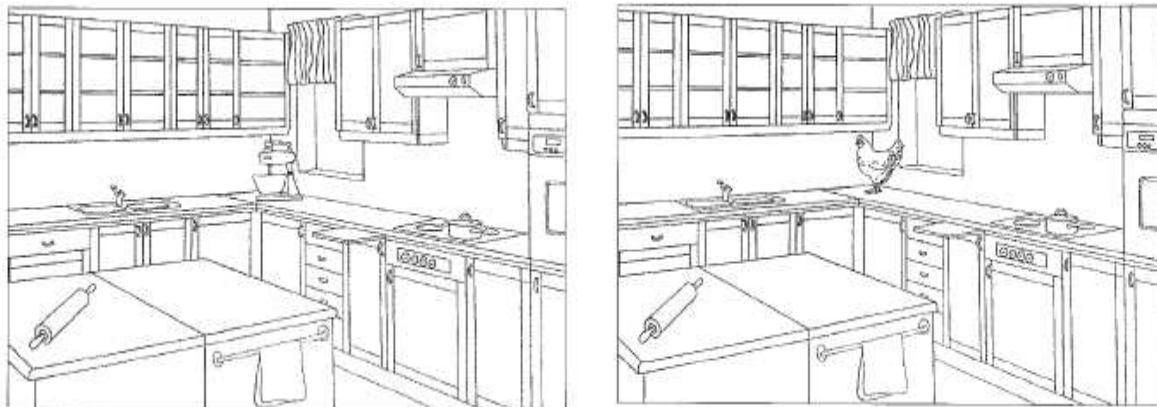


Figure 3. Exemples de scène consistante (à gauche) et inconsistante (à droite). L'objet cible consistant était le mixer dans la scène de gauche, et l'objet inconsistant était la poule dans la scène de droite. Adapté de Diepen & De Graef (1994)

Compte tenu de leurs résultats, Henderson et al. (1999) ont proposé que le placement des fixations dans une scène soit en premier lieu déterminé par une combinaison entre les caractéristiques visuelles des régions locales de la scène, la connaissance de la catégorie de la scène, et les propriétés visuelles globales de la scène. La position initiale des yeux ne

dépendrait pas selon eux de l'analyse sémantique des régions périphériques de la scène. Cependant, dès que l'analyse sémantique est rendue possible grâce à la vision fovéale, les fixations se porteraient ou retourneraient préférentiellement sur les régions les plus informatives d'un point de vue sémantique (résultat également observé par De Graef et al., 1990). Les yeux seraient donc initialement dirigés par des facteurs visuels (relatifs à la structure physique de l'image) et sémantiques globaux de la scène. Les aspects sémantiques locaux joueraient un rôle dans un deuxième temps, ce rôle étant de plus en plus important au fur et à mesure que la scène se déroule.

Des travaux récents suggèrent cependant qu'un traitement sémantique sur les objets individuels peut intervenir très précocement au cours du traitement. Dans une étude de Gordon (2004), les sujets étaient brièvement (53ms ou 147ms) exposés à une image, puis apparaissait un écran gris contenant un stimulus cible (& ou %) durant 107ms. Les participants réalisaient deux tâches. Dans un premier temps ils devaient identifier aussi rapidement que possible le stimulus cible. Dans un deuxième temps ils devaient rapporter, si possible, l'identité de l'objet sur lequel avait été superposé le stimulus cible. Le stimulus cible pouvait apparaître dans une localisation superposée, soit sur un objet consistant, soit sur un objet inconsistant, soit sur aucun objet. Les résultats montrent que pour des présentations de 147ms, les TR étaient plus courts lorsque le stimulus était superposé à un objet inconsistant qu'à un objet consistant avec la scène, suggérant que l'objet inconsistant attire davantage l'attention que les objets consistants. Par contre, aucun effet n'était observé lorsque la scène était présentée 53ms. Les données relatives à la tâche d'identification indiquent que ces détections étaient implicites. Ces résultats suggèrent que la perception précoce de la scène et des objets est non seulement déterminée par une catégorisation globale de la scène mais également par un traitement des objets individuels. De manière convergente, les travaux d'Underwood et Foushlam (2006) indiquent que les propriétés sémantiques d'un objet peuvent attirer l'attention lors de la première fixation, mais si et seulement si cet objet est saillant perceptivement. Auquel cas, cet objet sera plus enclin à attirer l'attention s'il est placé dans une scène dans laquelle il viole le « gist » que s'il est placé dans une scène dans laquelle il le respecte.

2.4. Interaction entre traitements ascendants et traitement descendants au cours de l'analyse de scènes visuelles

Les modèles visant à expliquer le guidage de l'attention dans une recherche visuelle par l'intervention de facteurs purement ascendants (e.g. le modèle de Itti & Koch, 2000) ou purement descendants (e.g. le modèle de Rao, Zelinsky, Hayhoe, & Ballard, 2002) rendent difficilement compte des données comportementales dans la plupart des tâches. Il est certain que le guidage de l'attention est déterminé par l'interaction entre des processus ascendants et descendants (pour des revues, Chun & Wolfe, 2001 ; Duncan & Humphreys, 1989 ; Egeth & Yantis, 1997). C'est pourquoi les modèles actuels se tournent de plus en plus vers une approche combinant des processus ascendants et descendants dans le guidage de l'attention. Une question centrale abordée dans ces travaux est de déterminer la contribution respective des processus ascendants et descendants dans le temps, à savoir si l'un de ces processus domine l'autre, et à déterminer comment évoluent ces influences au fur et à mesure que la scène se déroule. Comme dans l'architecture triadique proposée par Rensink (2000a), les modèles actuels sur la perception de scènes naturelles intègrent de plus en plus de composantes sémantiques comme facteur déterminant du déploiement attentionnel. De plus, ces modèles tiennent de plus en plus compte du fait que l'influence respective des facteurs ascendants et descendants évolue au fur et à mesure de l'exploration d'une scène.

Il est par ailleurs certain que l'influence respective des facteurs ascendants et descendants sur l'allocation attentionnelle dépend largement des objectifs de l'observateur et des contraintes inhérentes à la tâche. Pour une même scène visuelle, les patterns des mouvements des yeux sont très différents si la tâche du sujet est d'observer une scène dans le seul but de mémoriser le plus de détails possibles (en vue d'une tâche de reconnaissance par exemple) que s'il est engagé dans une tâche de recherche visuelle. De manière assez intuitive au demeurant, les données empiriques suggèrent une contribution plus importante des effets de saillance visuelle lorsque les participants sont simplement invités à « regarder des images » que lorsqu'ils sont invités à rechercher un objet prédéfini peu saillant visuellement comme par exemple un fruit dans une cuisine (e.g. Underwood, Foulsham, van Loon, Humphreys, & Bloyce, 2006). Les poids des saillances peuvent donc être modulés par des influences cognitives telles que rechercher un objet spécifique. Par ailleurs, les influences purement ascendantes dans une tâche de recherche visuelle semblent plus importantes si l'identité de la cible n'est pas au préalable spécifiée au sujet (Chen & Zelinsky, 2006).

Il existe aujourd'hui de nombreux modèles visant à rendre compte du guidage de l'attention au cours de l'exploration d'une scène (e.g. Navalpakkam & Itti, 2005 ; Tatler Baddeley, & Gilchrist, 2005 ; Torralba, Oliva, Castelhano, & Henderson, 2006 ; Wolfe, Cave, & Franzel, 1989). Pour la plupart, ces modèles considèrent que le stimulus visuel (e.g. une scène) est en premier lieu filtré de manière ascendante pour former une carte de saillance visuelle. Dans un deuxième temps seulement, les influences cognitives relatives aux buts, connaissances et attentes interviendraient pour construire une « carte maître » des déplacements de l'attention. A chaque fixation, l'attention serait guidée vers la région de la scène la plus pondérée au sein de la carte de saillance. Au fur et mesure que l'observateur explore la scène, fixe et identifie chaque objet, le poids attribué aux objets de la scène se modifierait en fonction de leur pertinence sémantique et cognitive. Dans la plupart des modèles actuels, le poids des saillances perceptives tend à diminuer au cours de l'exploration au profit des saillances cognitives, incluant des relations sémantiques entre les objets et la scène.

3. Le rôle des régularités contextuelles dans la perception visuelle

Il est largement démontré que les connaissances en MLT et les attentes relatives à la scène en cours jouent un rôle déterminant dans le guidage de l'attention. Les attentes que l'observateur peut avoir sur une scène reposent sur l'existence de régularités dans l'environnement. En effet, si le monde visuel est composé d'une multitude de détails, il n'en demeure pas moins prédictible. Les objets de l'environnement ne sont pas disposés au hasard. Les événements ne se succèdent pas dans un ordre arbitraire. Le monde visuel est fortement structuré, de telle façon que les objets et les événements tendent à co-varier d'un point de vue spatial, temporel, spatio-temporel. Par exemple, un bureau comporte généralement des livres, des articles, des stylos, un ordinateur, plus rarement un panneau de signalisation ou encore moins un éléphant. Les objets dans une scène sont régis par les lois physiques, comme par exemple les lois de la gravitation. L'organisation d'une scène obéit aux contraintes inhérentes à l'univers lui-même. Un sofa ne flotte pas dans l'air. Une scène du monde réel est également organisée selon les aspects fonctionnels des objets qui la composent. Un réfrigérateur est

généralement rencontré dans une cuisine. Ces connaissances sur les régularités de l'environnement seraient inscrites en mémoire sous forme de schéma de scène. Comme nous l'avons déjà défini, les schémas de scène décrivent l'ensemble des objets et des événements susceptibles d'apparaître dans une scène ainsi que leur position relative les uns par rapport aux autres (Mandler & Parker, 1976; Mandler & Ritchey, 1977). Ces schémas de scène incluent des connaissances génériques sur le monde, ainsi que des connaissances fonctionnelles. On pourrait également étendre cette définition en proposant que les schémas de scène comprennent aussi des connaissances sur des co-variations purement visuelles.

3.1. Le contexte facilite la reconnaissance et l'identification des objets

La reconnaissance d'un objet est-elle facilitée par le contexte dans lequel il apparaît ? Dans une série de travaux, Biederman et collaborateurs (e.g. Biederman, 1972 ; Biederman, Glass, & Stacy, 1973) ont montré que la détection d'un objet particulier est plus rapide lorsque ce dernier est présenté dans un contexte cohérent que dans un contexte déstructuré. Le matériel utilisé par ces auteurs comportait une série de photographies d'environnements courants (e.g. une rue). A partir de ces photographies, une version déstructurée (« *jumbled* ») était créée en coupant la version originale en 6 pièces et en arrangeant ces 6 pièces de manière aléatoire (Cf. Figure 4). Les participants étaient exposés brièvement aux scènes cohérentes et déstructurées dans lesquelles un indice spatial indiquait la position d'un objet cible. La tâche des participants était d'identifier le plus rapidement possible l'objet présent à la position indiquée dans la scène. Les temps d'identification étaient alors bien plus courts dans les scènes normales que dans les scènes déstructurées. Des résultats similaires ont été observés à l'aide d'un paradigme de recherche de cible (Biederman et al., 1973) : les participants étaient plus rapides pour trouver la cible lorsqu'elle apparaissait dans des scènes normales que déstructurées. Néanmoins, la déstructuration des scènes originales introduisant de nouveaux contours, il est difficile de déterminer si les effets reposent effectivement sur une facilitation contextuelle et pas simplement sur une plus grande complexité des scènes déstructurées par rapport aux scènes normales (pour une critique de ces travaux, cf. Henderson & Hollinworth, 1999).

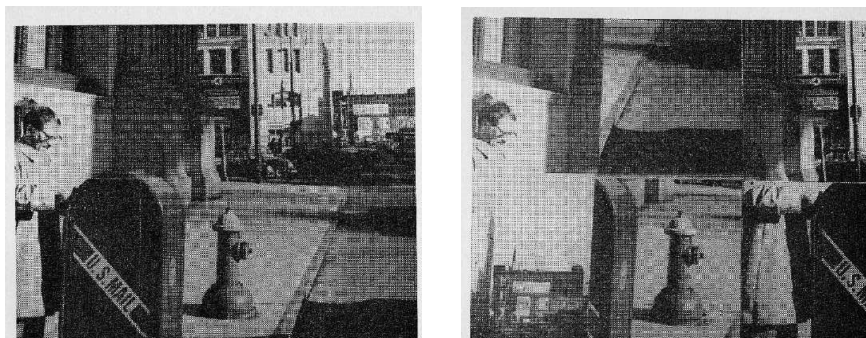


Figure 4. Exemples de scènes utilisées dans les travaux de Biederman et collaborateurs. A gauche : scène normale. A droite : scène déstructurée. *Adapté de Biederman, Rabinowitz, Glass et Stacy (1974)*

Les effets de consistance du contexte ont été par la suite étudiés en manipulant la consistance sémantique d'un objet dans une scène particulière. Par exemple, dans une étude de Palmer (1975), des dessins représentant des scènes du monde réel étaient présentés pendant 2 secondes (e.g. une cuisine), suivis d'une présentation rapide d'un objet cible isolé qui pouvait être sémantiquement consistant ou inconsistant avec la scène (i.e. susceptible ou non d'apparaître dans la scène). Les objets cibles inconsistants sémantiquement pouvaient être similaires ou différents d'un point de vue perceptif avec un objet consistant potentiel (pour un exemple d'objets similaires ou différents perceptivement, cf. Figure 5). Palmer a observé que les objets consistants étaient nommés plus précisément que les objets présentés sans contexte, qui eux étaient nommés plus précisément que les objets inconsistants. Et enfin, les objets inconsistants partageant de fortes similitudes visuelles avec un objet consistant potentiel conduisaient aux plus mauvaises performances. De tels effets facilitateurs d'un contexte consistant sur l'identification d'un objet cible (se traduisant par des temps de détection ou de dénomination plus courts pour des objets apparaissant dans un contexte congruent qu'incongru) ont été observés dans de nombreuses études (Biederman, Mezzanotte, & Rabinowitz, 1982 ; Biederman, Teitelbaum, & Mezzanotte, 1983 ; Boyce & Pollatsek, 1992 ; Boyce, Pollatsek, & Rayner, 1989 ; Davenport & Potter, 2004 ; cependant, voir Hollingworth & Henderson, 1998).

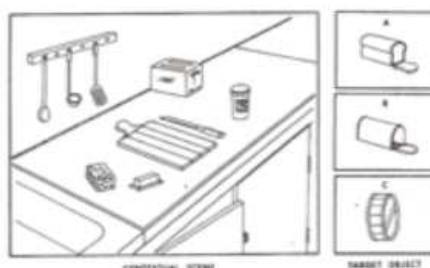


Figure 5. Exemple de stimulus utilisé par Palmer (1975). A droite : un objet consistant (une brioche), un objet inconsistant mais ressemblant visuellement à un objet consistant potentiel (une boîte aux lettres) et un objet inconsistant et différent visuellement (une roue). *Adapté de Palmer (1975)*

A quel niveau du traitement le contexte est-il susceptible d'exercer une influence sur l'identification d'un objet ? Différents modèles ont été proposés selon l'étape du traitement d'identification sur laquelle le contexte exercerait son influence. On présuppose classiquement que l'identification d'un objet implique trois niveaux de traitement. En premier lieu, l'image projetée sur la rétine est traduite en une série de primitives. En second lieu, ces primitives sont assemblées pour construire des descriptions structurales relatives aux objets présents dans la scène. Et enfin, ces descriptions sont mises en relation avec des descriptions élaborées, déjà stockées en MLT. Quand une correspondance est établie, l'objet est reconnu et son identification est rendue possible grâce à la disponibilité en mémoire de l'information sémantique sur ce type d'objet.

Le modèle du schéma perceptif propose que les attentes, dérivées des connaissances préalables sur la composition d'une scène, affectent de manière précoce le traitement perceptif des objets présents dans la scène (Biederman et al., 1982, 1983 ; Palmer, 1975). L'activation précoce d'un schéma de scène faciliterait la construction d'une description structurale des objets consistants avec la scène, et peut-être inhiberait la construction des descriptions relatives aux objets inconsistants avec le schéma de scène. *Le modèle de l'amorçage* propose quant à lui que le contexte influence l'identification des objets à un niveau plus tardif du traitement, à savoir lorsque la description structurale de l'objet est mise en relation avec une/des description(s) déjà en mémoire (Friedman, 1979 ; Palmer, 1975). Selon le modèle de l'amorçage, l'activation d'un schéma de scène amorcerait les représentations en mémoire des objets types consistants avec le schéma. Cet amorçage faciliterait l'identification d'un objet congruent en diminuant la quantité d'information nécessaire pour sélectionner en mémoire la représentation correspondant à l'objet visuel (Friedman, 1979). Comme le modèle du schéma perceptif, le modèle de l'amorçage prédit que l'identification des objets consistants avec la scène sera facilitée par rapport à l'identification des objets inconsistants. Ces deux modèles diffèrent en ce que le modèle de l'amorçage présuppose que les connaissances relatives à la scène influencent uniquement le critère d'analyse requis pour déterminer si un objet particulier est présent, sans influencer directement l'analyse perceptive de l'objet en question, alors que le modèle du schéma perceptif postule une influence précoce du contexte sur le traitement perceptif. Ces deux conceptions ont été unifiées au sein du *modèle d'activation interactive*, dans lequel les deux niveaux d'analyse pourraient être influencés par le contexte (Boyce et al., 1989 ; Metzger & Antes, 1983).

Le paradigme de détection d'objet introduit par Biederman et collaborateurs a certainement été l'outil le plus utilisé pour discriminer l'influence du contexte à un niveau précoce ou tardif du traitement d'identification (e.g. Biederman et al., 1982, 1983 ; Boyce et al., 1989 ; Hollingworth & Henderson, 1998). La procédure utilisée par Biederman et al. (1982) était la suivante. Au début de chaque essai, une étiquette désignant un objet cible était présentée aux participants. Puis apparaissait un dessin représentant une scène pendant 150ms, suivi d'un masque comportant un indice désignant une localisation. La tâche des participants était d'indiquer si l'objet cible était apparu dans la scène à la position indiquée. L'objet qui était apparu dans la localisation indiquée pouvait être consistant avec la scène ou violer les attentes établies par la scène selon une ou plusieurs dimensions. Ces violations pouvaient concerner, des probabilités épisodiques de co-occurrences (consistance sémantique), des dimensions relatives à la position ou à la taille de l'objet dans la scène, à la présence ou l'absence de support, ou encore des interpositions entre les objets (si l'objet chevauchait un des objets ou était transparent), (pour des exemples, cf. Figure 6). Biederman et al. ont observé que la sensibilité de détection (d') était meilleure quand l'objet indicé ne violait aucune contrainte imposée par la signification de la scène. Les performances (TR et réponses correctes) étaient dégradées quelle que soit la dimension de violation, mais elles l'étaient davantage lorsque la scène présentait des violations multiples (e.g. inconsistance sémantique et interposition). Puisque les violations structurales (i.e. support et interposition) n'avaient pas plus d'effet que les violations sémantiques (i.e. probabilité, position, taille), les auteurs en ont conclu que les influences contextuelles de nature sémantique devaient opérer durant l'analyse perceptive des objets. Ces résultats allaient donc en faveur du modèle du schéma perceptif de scène, suggérant une influence du contexte sur les étapes précoces du traitement perceptif d'identification (pour un résultat similaire, voir aussi Biederman et al., 1983).

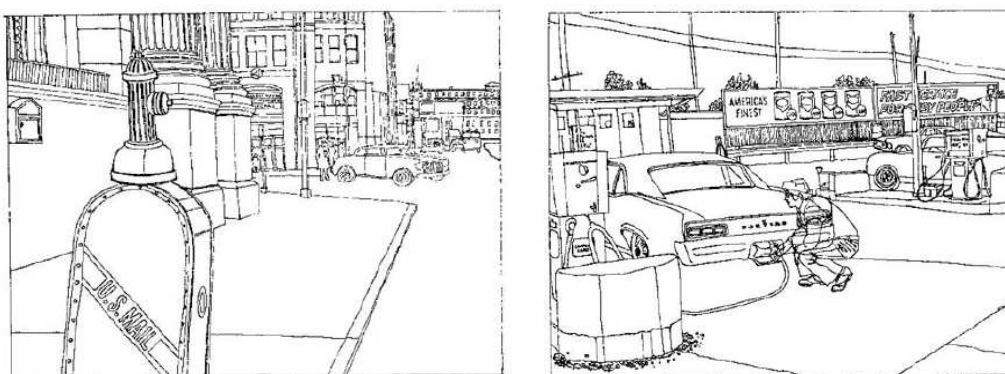


Figure 6. Exemples de violations de « support » et d' « interposition » dans l'étude Biederman et al., 1982. Adapté de Biederman et al. (1982)

Les travaux de Biederman et collaborateurs ne sont toutefois pas exempts de critiques. Hollingworth et Henderson (1998) ont ainsi rapporté un ensemble de biais dans le paradigme original de détection d'objet introduit par Biederman et al. (1982). En contrôlant alors les biais de réponses et en éliminant les autres sources de biais potentiels (contrôle du taux de fausses alarmes et de l'avantage « objet consistant-localisation probable »), les auteurs ne sont pas parvenus à observer un bénéfice dans la détection des objets consistants sémantiquement avec la scène par rapport aux objets inconsistants. Hollingworth et Henderson (1998) ont proposé une explication alternative aux effets de contexte : l'identification pourrait être fonctionnellement isolée de l'information contextuelle et des attentes relatives à la scène. Leur *modèle d'isolation fonctionnelle* stipule que le contexte n'influence pas la perception des objets, mais qu'il biaise le traitement postérieur à l'identification, tel que la génération de la réponse ou encore les stratégies de « guessing ». Par contre, une fois que la représentation d'un objet serait formée, le statut sémantique de l'objet pourrait influencer l'allocation attentionnelle visuelle et spatiale. Une étude récente de Davenport et Potter (2004) va cependant à l'encontre du modèle d'isolation fonctionnelle et montre, à partir de photographies cette fois-ci, que l'information relative à la consistance sémantique est non seulement disponible à partir d'images présentées brièvement (80ms), mais qu'en plus elle affecte la perception du contexte et l'identification des objets présents au premier plan.

Malgré les nombreux biais méthodologiques présents dans certains de ces travaux, les études portant sur l'influence du contexte dans l'identification des objets convergent vers un consensus accordant aux contextes congruents un rôle facilitateur dans l'identification des objets.

3.2. La construction du contexte

Les travaux portant sur les effets de contexte montrent que le contexte influence la perception des objets présents dans la scène. Comment le contexte exerce-t-il une influence sur la reconnaissance et l'identification des objets ? La facilitation contextuelle résulte-t-elle de l'information locale relative aux objets présents au sein de la scène ou résulte-t-elle d'un traitement global des « traits émergents de la scène » ? Quelle est la hiérarchie des traitements relatifs aux différentes entités d'une scène, à savoir le sens général de la scène (i.e. le « gist »)

et les objets qui la composent ? La reconnaissance de certains objets précède-t-elle la reconnaissance de la scène ou à l'inverse, l'identification d'une scène précède-t-elle la reconnaissance des objets ? En d'autres termes, comment le contexte est-il construit ? Est-ce l'identification du « gist » qui active un schéma de scène ou est-ce l'activation d'un schéma de scène qui permet l'identification du « gist » ?

La construction du contexte pourrait en effet reposer sur l'identification d'un ou plusieurs objets clés (cf. le modèle de l'amorçage, Friedman, 1979), ou encore sur leurs relations spatiales (e.g. De Graef et al., 1990). Mais le contexte pourrait aussi être inféré à partir de traitements de bas niveau uniquement, sans nécessiter l'identification des objets individuels (Biederman, 1988 ; Schyns & Oliva, 1994).

Selon Friedman (1979) les effets de contexte reposent sur les relations entre les objets présents dans la scène. L'identification d'un objet amorcerait la reconnaissance et la catégorisation des objets qui lui ont été associés dans le passé. Dans cette perspective, la facilitation contextuelle nécessiterait l'identification d'un ou plusieurs objets au sein de la scène (e.g. l'identification d'un lit amorcerait la reconnaissance d'un oreiller). Dans le même sens, les travaux d'Henderson, Pollatsek et Rayner (1987) ont montré que l'identification d'un objet cible est facilitée si les yeux se sont au préalable fixés sur un objet relié sémantiquement, et ce, même si les objets sont dépourvus de contexte (i.e. sans arrière-plan). Selon Henderson et al., la reconnaissance d'un objet amorcerait l'identification des objets. Ce résultat n'exclut néanmoins pas l'influence d'autres facteurs dans les effets contextuels.

Pour de nombreux auteurs, les effets de contexte reposent davantage sur l'extraction précoce d'une information globale dans l'image (e.g. Antes, Penland, & Metzger, 1981 ; Boyce et al., 1989). Boyce et al (1989) ont exploré au moyen d'un paradigme de détection d'objets indicés, si le bénéfice observé dans les contextes consistants tient à la signification globale de la scène ou à la présence des autres objets sémantiquement pertinents au sein de la scène. Les auteurs ont manipulé la consistance d'un objet indicé par rapport à la scène globale et/ou par rapport aux autres objets apparaissant dans la scène. Les temps de détection étaient alors plus courts lorsque l'objet indicé était sémantiquement consistant avec la scène dans laquelle il apparaissait, que lorsqu'il était inconsistant. Par contre, les auteurs n'ont pas observé d'effet de la consistance entre l'objet cible et les autres objets lorsque l'arrière-plan était absent. De plus, même lorsque les objets adjacents à l'objet cible ne lui étaient pas reliés

sémantiquement, l'arrière-plan à lui seul facilitait la détection de l'objet cible. Boyce et al. en ont conclu que la signification globale de la scène et les traits globaux, davantage que les objets spécifiques, ont un rôle fonctionnel dans l'identification des objets qui la composent (voir aussi Antes et al., 1981, pour un résultat concordant dans une tâche d'identification).

De manière convergente, la plupart des recherches actuelles soutiennent l'hypothèse selon laquelle la construction du contexte s'appuierait sur une information globale plutôt que sur des informations locales relatives aux objets individuels. D'une part, de nombreuses recherches ont montré que de très courtes durées de présentation sont suffisantes pour permettre à l'individu d'identifier et de dénommer une scène visuelle, c'est-à-dire d'en extraire le « gist ». Potter et collaborateurs (Potter, 1976 ; Potter & Levy, 1969), ont ainsi montré qu'une scène peut être identifiée en une centaine de millisecondes. Les travaux de Thorpe, Gegenfurtner, Fabre-Thorpe et Bülthoff (2001) ont révélé que les sujets sont capables d'indiquer si une scène contient ou non un animal dans des scènes présentées 28ms, et ce même lorsque la scène est excentrée par rapport au point de fixation. Mais surtout, il semblerait que les scènes visuelles soient reconnues aussi rapidement que les objets isolés (Biederman et al., 1982 ; Friedman, 1979 ; Intrub, 1997), parfois à partir d'informations extraites en une seule fixation (Henderson & Hollingworth, 2003). Enfin, les scènes visuelles peuvent être identifiées à partir d'images dans lesquelles sont préservées les relations spatiales entre les structures grossières de la scène mais dans lesquelles il manque le détail visuel nécessaire pour identifier les objets individuels (Schyns & Oliva, 1994).

Si la reconnaissance d'une scène ne nécessite pas la reconnaissance des objets, comment les représentations sur les scènes sont-elles construites ? De plus en plus d'auteurs s'accordent à penser que le contexte pourrait être construit de manière holistique, sans nécessiter la reconnaissance des objets individuels, et sans même nécessiter d'étapes de segmentation ou de groupement entre les traits. De plus, il semblerait que la mise en œuvre de ces traitements holistiques dans l'identification d'une scène (Oliva & Torralba, 2001) ou d'un objet (Thoma, Hummel, & Davidoff, 2004) ne requière pas l'intervention d'une attention visuelle sélective. La catégorisation de la plupart des scènes du monde réel pourrait être inférée à partir de l'agencement spatial uniquement ou à partir des traits globaux. Des études montrent que les traits basiques, comme la distribution spatiale des régions colorées (Oliva & Schyns, 2000 ; Rousselet, Joubert, & Fabre-Thorpe, 2005) ou encore la distribution des orientations (McCotter, Grosselin, Sowden, & Schyns, 2005) sont corrélés avec certaines catégories

sémantiques de scènes du monde réel. Schyns et Oliva (1994) ont émis l'hypothèse que les régularités relatives à l'organisation spatiale correspondant à une catégorie de scène particulière pourraient être suffisantes pour extraire l'essentiel de la scène et permettre son identification. Les variations chromatiques pourraient également participer à l'identification de la scène et à l'extraction rapide du « gist ». Le rôle des indices colorés dans la reconnaissance des scènes et des objets demeure cependant assez controversé (Oliva & Schyns, 2000 ; Rousselet, Joubert & Fabre-Thorpe, 2005). L'ensemble des recherches sur l'extraction rapide du « gist » suggère ainsi l'influence prédominante d'un traitement holistique sur un traitement analytique. Dans ce sens, Biederman (1988) a proposé que l'analyse des scènes puisse mettre en jeu le même mode représentationnel que celle relative aux objets individuels mais sur une plus large échelle spatiale. Cependant, certains travaux empiriques visant à explorer cette hypothèse laissent penser que les scènes ne sont pas représentées comme de « grands objets ». Des études d'imagerie cérébrale (Epstein & Kanwisher, 1998) suggèrent par exemple que des aires corticales distinctes peuvent être impliquées dans l'identification des objets et des scènes. Mais surtout, alors que les objets tendent à être fortement contraints par un ensemble de parties composantes et par les relations entre ces parties, une scène serait beaucoup moins contrainte par les objets et les relations spatiales entre ces objets (Hollingworth & Henderson, 1998). On peut néanmoins envisager qu'à une même étiquette de scène peuvent correspondre plusieurs prototypes très différents en termes de traits visuels.

Toutefois, les travaux de Davenport et Potter (2004) ont montré que les objets et la scène sont traités de manière interactive, et que les connaissances sur les co-occurrences objets/contexte influencent la perception. Les objets et l'arrière-plan pourraient mutuellement se contraindre. Par ailleurs, leurs travaux révèlent que l'objet en premier plan présente un statut privilégié dans le traitement. Aussi la reconnaissance d'une scène du monde réel pourrait-elle faire intervenir les deux processus. Si le traitement holistique n'est pas suffisant pour extraire le « gist », des traitements analytiques pourraient être parallèlement mis en œuvre pour accéder au sens de la scène. La facilitation contextuelle dans la reconnaissance d'un objet pourrait à la fois tenir à l'identification de la scène et à celle des objets individuels.

3.3. Le rôle du contexte dans le guidage de l'attention

En plus de faciliter l'identification des objets, les schémas de scène fournissent un cadre de référence qui permet de guider l'attention et les mouvements des yeux au sein de la scène (Rensink, 2000a). Par exemple, Shinoda, Hayhoe et Shrivastava (2000) ont montré que les sujets placés en situation de conduite détectent les panneaux « stop » principalement quand ceux-ci sont localisés aux intersections, et peu quand ils sont localisés au milieu de la chaussée, suggérant que la perception des conducteurs s'appuie sur les régularités de l'environnement (voir aussi, Land & Hayhoe, 2001). Non seulement les schémas de scène renseignent sur les objets susceptibles d'être présents dans la scène, mais en plus, ils renseignent sur les localisations potentielles de ces objets. L'activation précoce des connaissances sur les régularités de l'environnement optimise le déploiement attentionnel.

Le modèle développé par Oliva et Torralba (2001) et Torralba, Oliva, Castelhamo et Henderson (2006) vise à rendre compte des effets précoces du contexte sur le guidage de l'attention, en implémentant des mécanismes hypothétiques sous-jacents à l'extraction rapide du « gist » et à l'activation d'un schéma de scène. Le modèle *de guidage attentionnel par le contexte global* fait intervenir des traitements holistiques des traits globaux comme facteur précoce au déploiement attentionnel (pour une représentation schématique du modèle de guidage attentionnel par le contexte global proposé par Torralba et al. (2006), cf. Figure 7). Le traitement de la structure globale et les relations spatiales entre ses composants précéderait l'analyse des détails locaux (Navon, 1977). Ce modèle présuppose que la structure d'une scène peut être représentée par la moyenne des traits globaux de l'image pour une résolution spatiale minime. Selon Oliva et Torralba (2001) la représentation du contexte ne requerrait pas une analyse de la scène en objets individuels mais reposerait sur les propriétés statistiques globales de l'image. L'information contextuelle pourrait être intégrée avant la première saccade oculaire. Dans leur modèle, à la fois la saillance visuelle de l'image et les traits globaux du contexte sont enregistrés et calculés en parallèle puis intégrés précocement au cours du traitement visuel. La représentation schématique résultante de l'analyse holistique serait suffisante pour inférer la catégorie probable de la scène. Sur cette base, le rôle des traits globaux est d'activer les localisations susceptibles de contenir la cible, et donc de réduire les effets de saillance visuelle non pertinents pour la tâche de recherche. Par exemple, si la tâche est de déterminer si la scène contient un individu, la reconnaissance d'une scène de ville

activera certaines régions propices à la présence d'une personne et inhibera des régions peu probables, même si ces dernières s'avèrent être très saillantes perceptivement. Le modèle implémente ce facteur « connaissance » relatif à la scène, en pondérant les régions de la scène selon leur probabilité de contenir la cible par rapport à des scènes prototypiques. Par ailleurs, le modèle du guidage par le contexte global n'exclut pas un effet parallèle des effets de contexte objet-objet. S'il offre un chemin direct (« feed-forward ») au traitement visuel, en décrivant l'agencement spatial et l'information conceptuelle relative à la catégorie de la scène, sans nécessiter la segmentation des objets, il implémente également un traitement sur les traits locaux. La représentation globale du contexte visuel concourrait en parallèle avec le traitement des objets individuels. Ce modèle combine donc à la fois les sources locales et globales de l'information. Mais les traits globaux de la scène expliqueraient l'impact initial du contexte, contraignant rapidement l'analyse des objets locaux, alors que l'association objet-objet se développerait plus tardivement, de manière progressive, et dépendrait d'un traitement de segmentation de l'image en objets individualisés.

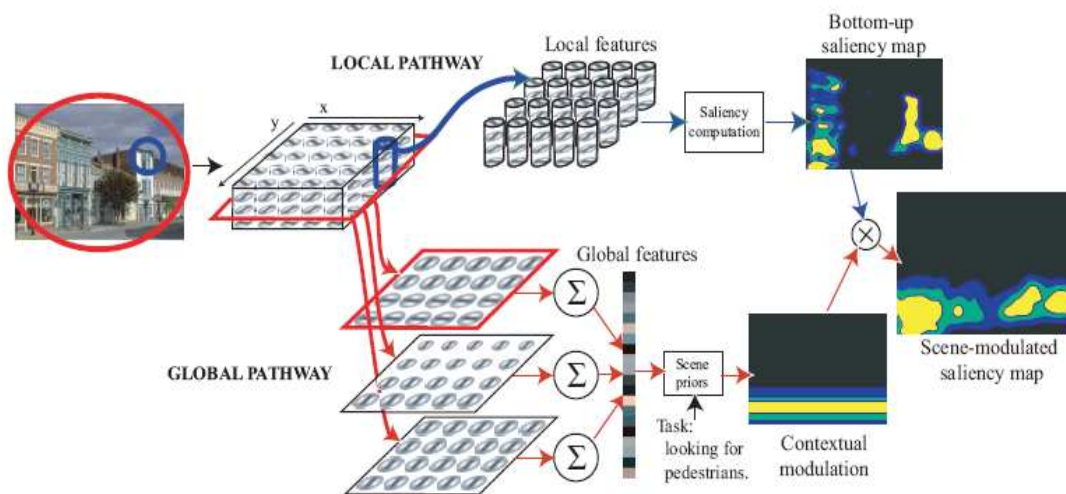


Figure 7. Modèle de guidage attentionnel, intégrant la saillance d'une image et les attentes relatives à la scène. L'image est analysée par deux « chemins ». Au cours de la première étape, commune aux deux chemins, l'image est filtrée par des filtres « multiscale-oriented ». Le chemin local représente chaque localisation spatiale indépendamment. Cette représentation locale est utilisée pour calculer les saillances de l'image et permettre la reconnaissance des objets sur la base de leur apparence locale. Le chemin global représente les entrées de l'image de manière holistique en extrayant les statistiques globales de l'image. Cette représentation globale peut être utilisée pour la reconnaissance de la scène. Dans ce modèle, le chemin global interagit avec les attentes relatives à la localisation de la cible dans l'image. *Adapté de Torralba et al. (2006)*

4. Conclusion et perspectives

Les travaux étudiant le contenu et la stabilité de nos représentations visuelles suggèrent que peu d'informations présentes dans notre champ visuel sont accessibles à la conscience à un instant donné. Dans les théories visant à expliquer la nature de nos représentations visuelles, l'attention sélective joue un rôle déterminant dans la construction d'une représentation stable et cohérente des objets de l'environnement. Ceci pose deux problèmes au cœur des recherches actuelles dans le domaine de la perception visuelle. Quels sont les facteurs responsables de notre impression de « continuité entre les traits visuels » d'une part, et comment l'attention est-elle orientée de manière efficace d'autre part ? Dans ce cadre, il est de plus en plus admis que les indices contextuels seraient activés très précocement au cours de la perception pour construire une représentation unifiée de l'ensemble de la scène. La présence de régularités contextuelles permettrait l'identification rapide de la scène (i.e. l'extraction du « gist »), qui à son tour activerait un schéma de scène, lequel faciliterait la reconnaissance et l'identification des objets qui la composent, ainsi que l'exploration de la scène. L'activation précoce de schémas de scène, via le contexte, rendrait ainsi compte de l'efficacité de nos comportements perceptifs. En permettant d'organiser la scène, les connaissances relatives aux régularités contextuelles facilitent les traitements visuels, notamment, l'interprétation du monde et l'orientation de l'attention sélective vers les aspects pertinents de l'environnement. L'efficacité de nos comportements perceptifs tiendrait à l'existence de contextes signifiants structurés, en dépit de la « pauvreté » de nos représentations conscientes.

Les travaux reposant sur les *modèles de facilitation contextuelle* soulèvent la question suivante : comment les connaissances sur les contextes sont-elles acquises et utilisées quand l'observateur est engagé dans une tâche particulière ? La sensibilité aux régularités contextuelles n'est pas innée, elle se développe avec l'expérience. On pourrait simplement envisager qu'à l'âge adulte, l'essentiel des schémas de scène est déjà en mémoire, et que l'individu ne fait que récupérer des connaissances acquises au cours des étapes précoces du développement. Pourtant, les travaux sur l'expertise montrent comment les traitements perceptifs, même de bas niveau, sont modulés par l'expérience. Des schémas de scènes se construisent tout au long de l'existence. Prenons le cas d'un expert et d'un novice en

mycologie à la recherche de girolles. L'expert sait où chercher, alors que le novice cherche au hasard. Un amateur chevronné dira qu'il faut chercher dans les endroits humides, qu'il faut tel ou tel type de végétation. Mais il dira aussi qu'il « sent » la girolle sans pouvoir expliquer pourquoi. Ce sont bien sûr ses connaissances sur les contextes propices qui orientent sa recherche. L'expert se différencie du novice par sa capacité à percevoir très rapidement les endroits fructueux pour sa recherche. Ses connaissances sur les contextes lui permettent à la fois d'encoder rapidement les éléments pertinents de la scène mais également d'orienter ses processus attentionnels. Ainsi, des connaissances inscrites en mémoire sous forme de schémas de scène s'acquièrent tout au long de l'existence.

Il reste néanmoins à définir les types de régularités contextuelles auxquels les observateurs deviennent sensibles au cours de l'exploration d'une scène et comment ces connaissances sont acquises. Les régularités contextuelles peuvent reposer sur l'identité propre des éléments constitutifs des contextes ou sur leur agencement spatial. Elles peuvent aussi concerner des régularités temporelles, correspondant à l'enchaînement plus ou moins déterministe d'une séquence d'événements. Mais elles peuvent également tenir à la signification globale du contexte, faisant intervenir des mécanismes préalables de catégorisation des éléments contextuels. Les *modèles de facilitation contextuelle* (Hollingworth & Henderson, 1998) soulignent ainsi le rôle déterminant de la signification globale du contexte dans les effets de contextes observés, plutôt que celui des propriétés perceptives spécifiques des éléments contextuels (Biederman et al., 1982 ; Friedman, 1979 ; Hollingworth & Henderson, 2000) ou celui des objets spécifiques présents dans la scène (Boyce et al., 1989).

Il reste également à déterminer comment ces connaissances sont acquises au cours de l'analyse d'une scène visuelle. Dans ce cadre, on peut se demander si des régularités sémantiques inhérentes à certains contextes catégoriels peuvent être apprises et utilisées de manière implicite, i.e. inconsciente et involontaire et donc, si l'apprentissage implicite participe à la construction de *schémas de scène*. L'homme n'est pas toujours à même d'expliquer les racines et les raisons de ses comportements néanmoins très adaptés. Réciproquement, on peut se demander dans quelle mesure la prise de conscience des régularités contextuelles facilite leur encodage, leur consolidation en MLT et leur récupération. Si les travaux inscrits dans le domaine de la perception visuelle évoquent souvent l'idée que les comportements perceptifs (e.g. les mouvements des yeux) sont

influencés par des traitements cognitifs non conscients, peu de ces travaux ont directement porté sur la nature implicite vs. explicite de l'apprentissage mis en évidence. Nos travaux s'inscrivent dans cette problématique. Quelle place peut-on accorder aux mécanismes d'apprentissage implicite dans l'acquisition de connaissances sur les régularités de l'environnement ?

CHAPITRE II

L'APPRENTISSAGE IMPLICITE DES RÉGULARITÉS DE L'ENVIRONNEMENT

Les travaux sur la perception de scènes visuelles montrent le rôle fondamental des connaissances relatives aux régularités de l'environnement dans les traitements perceptifs et notamment dans le guidage de l'attention. Comment ces connaissances sont-elles acquises au cours de l'analyse d'une scène visuelle ? L'hypothèse à l'origine de nos recherches est que des connaissances sur les régularités de l'environnement peuvent être acquises et récupérées de manière inconsciente et involontaire, i.e. de manière implicite. Dans cette perspective, trois questions faisant l'objet de débats fondamentaux autour du concept d'apprentissage implicite ont motivé nos travaux:

- Dans quelle mesure l'apprentissage implicite produit-il une connaissance inconsciente ?
- L'apprentissage implicite est-il un processus automatique, i.e. qui ne requiert pas d'attention ?
- L'apprentissage implicite peut-il conduire à des connaissances abstraites ?

Avant de présenter l'enjeu de ces questions pour la compréhension du fonctionnement cognitif, nous évoquerons très succinctement quelques situations expérimentales suggérant la mise en évidence d'apprentissages implicites.

1. L'apprentissage implicite de régularités

Le domaine de l'apprentissage implicite a été unifié autour de trois champs de recherche puisant leur origine dans une littérature hétérogène (pour des revues, Cleeremans, Destrebecqz, & Boyer, 1998 ; Goschke, 1997; Perruchet & Nicolas, 1998). Le premier champ de recherche est celui de l'apprentissage des grammaires artificielles à états finis, développé par Reber (1967) à partir des travaux conduits dans les années 50 par Chomsky et Miller en linguistique et psycholinguistique. La deuxième source d'investigation renvoie aux travaux sur les systèmes dynamiques de contrôle menés par Broadbent (1977) à des fins d'applications ergonomiques. Le troisième champ de recherche porte sur l'acquisition des habiletés motrices dans des tâches d'apprentissage de séquences répétées (Nissen & Bullemer, 1987). Mais si ces trois champs de recherche sont vus comme les fondements du domaine qui nous intéresse ici, le paradigme d'indigage contextuel⁵ (Chun & Jiang, 1998) et le paradigme d'apprentissage de régularités statistiques⁶ (Saffran, Aslin, & Newport, 1996) trouveront très certainement leur place dans les années à venir, au sein du domaine de l'apprentissage implicite.

De nombreuses autres situations expérimentales ont par ailleurs suggéré à leurs auteurs la mise en évidence d'un apprentissage implicite (pour une présentation de différentes situations d'apprentissage implicite, cf. Shanks, 2005). On pourra citer les travaux sur l'apprentissage de probabilité (Reber & Millward, 1968), la détection d'invariances statistiques (Lewicki, 1986), l'apprentissage perceptif (Cohen & Squire, 1980), le conditionnement classique (Weiskrantz & Warrington, 1979), l'apprentissage de signaux en mouvement cohérent (pour une revue, Seitz & Watanabe, 2005). Des effets d'apprentissage implicite à court terme ont également été mis en évidence dans le domaine de la perception visuelle, comme par exemple des effets d'inhibition de retour⁷ (Posner & Cohen, 1984) ou des effets d'amorçage pop-out⁸ (Maljkovic & Nakayama, 1994, 1996, 2000). L'ensemble de ces recherches mettent en exergue un phénomène que l'on peut désigner sous les termes d'« apprentissage implicite de régularités ». Ces recherches suggèrent que si le détail de l'information imprimée sur la rétine

⁵ Traduit de l'anglais, contextual cueing (ce paradigme sera décrit en détail dans le chapitre III).

⁶ Traduit de l'anglais, statistical learning (pour une description, cf. Annexe2).

⁷ IOR : inhibition of return: les sujets tendent à éviter les localisations déjà visitées, en faveur d'un déploiement attentionnel vers les objets et événements nouveaux.

⁸ Un effet d'amorçage pop-out se traduit par un effet facilitateur du guidage de l'attention par des indices amorçant l'identité ou la localisation d'une cible à détecter.

n'est pas consciemment retenu en mémoire, cela ne signifie pas une absence totale de persistance de cette information. Différents mécanismes peuvent permettre à l'information d'être préservée et de s'accumuler au cours des répétitions, biaisant notamment le déploiement de l'attention visuelle. Des traces en mémoire, non disponibles à la conscience, pourraient participer à ces mécanismes. Les importantes redondances visuelles qui existent dans l'environnement sont susceptibles de bénéficier d'un traitement mnésique à court terme et à long terme, même si elles dépassent largement les capacités de détection explicite.

Attardons nous à présent sur deux tâches classiques autour desquelles nos questions de recherche ont été les plus largement débattues, à savoir les tâches d'apprentissage de grammaires artificielles et les tâches d'apprentissage de séquences répétées.

1.1. L'apprentissage de grammaires artificielles (AGL)

Le terme « apprentissage implicite » a été utilisé pour la première fois par Reber en 1967 pour décrire certains processus d'acquisition en jeu dans ses travaux portant sur les grammaires artificielles à états finis (pour des revues, Nicolas, 1996 ; Pothos, 2007). Une tâche standard d'apprentissage de grammaire artificielle (AGL : « Artificial Grammar Learning ») met en œuvre des protocoles expérimentaux se déroulant en deux phases. Dans la première phase, il est présenté au sujet naïf plusieurs séries de lettres dont l'ordre de succession est régi selon des règles complexes arbitraires établies par l'expérimentateur (pour un exemple de grammaire artificielle, cf. Figure 8). Dans la recherche initiale de Reber (1967), il était demandé au participant d'apprendre des séries de consonnes. Dans d'autres versions de la tâche, il lui est simplement demandé d'observer attentivement les séries de lettres. Durant la deuxième phase, succédant immédiatement à la première, le sujet est informé que les séries de lettres qu'il vient d'étudier étaient construites à partir de règles formelles. Il est alors exposé à de nouvelles séries de lettres dont certaines respectent les règles de la grammaire artificielle et d'autres les violent. Le sujet a pour consigne de porter un jugement sur ces nouvelles séries en termes de respect ou de non respect des règles de construction. A l'issue de cette deuxième phase, les sujets parviennent généralement à classer les chaînes grammaticales et non grammaticales avec un taux de réussite supérieur au hasard, sans pour autant être capables de verbaliser les règles sur lesquelles leur jugement

semble être fondé ou encore de justifier leurs décisions. Sur la base de cette dialectique entre la performance de classification supérieure au hasard et l'incapacité des sujets à verbaliser les règles, Reber conclut que ces derniers ont acquis de manière incidente et automatique des connaissances inconscientes sur la structure abstraite de la grammaire artificielle. Selon Reber (1989, p. 226), la connaissance implicite acquise au cours d'une tâche d'AGL est *"deep, abstract, and representative of the structure inherent in the underlying invariance patterns of the stimulus environment"*. Cependant, cette lecture des résultats a fait l'objet d'une polémique empirique et théorique sur laquelle nous reviendrons à divers points de ce chapitre.

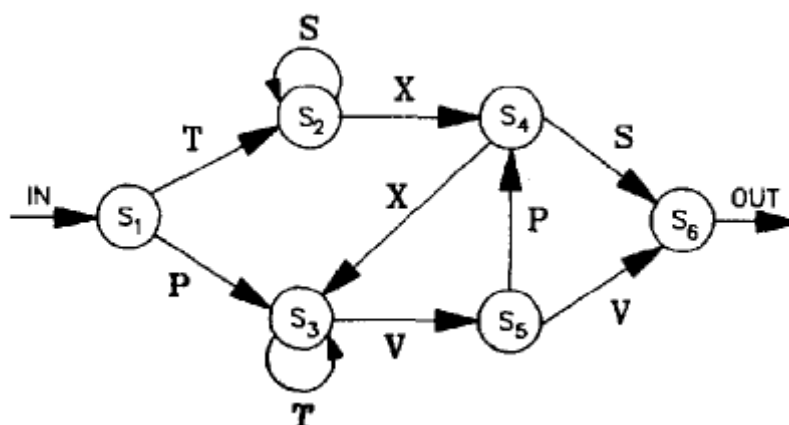


Figure 8. Exemple de grammaire artificielle à états finis. Les suites grammaticales sont générées en suivant l'ordre des flèches. Les items présentés durant la phase d'étude sont tous grammaticaux. Par exemple, les séries de lettres suivantes respectent les règles de la grammaire représentée ci-dessus : TXS, TSXS, TSSXS, TXXVV, PVV, PTVV, PTVPS. Les séries agrammaticales ne respectent pas les transitions de la grammaire. *Adapté de Reber, 1989*

1.2. L'apprentissage séquentiel (SRT)

Les situations d'apprentissage séquentiel (SRT : « Serial Reaction Time ») mettent en évidence des effets d'apprentissage sensori-moteur suite à la pratique d'une tâche qui consiste classiquement à pister une cible se déplaçant sur un écran d'ordinateur. Dans la tâche initiale développée par Nissen et Bullemer (1987) et dans de nombreuses études ultérieures, les participants sont exposés à de nombreux essais au cours desquels la cible apparaît dans l'une ou l'autre des quatre positions (1-4) alignées en bas de l'écran. La tâche du sujet est d'appuyer le plus rapidement et correctement possible sur la touche du clavier désignée au préalable comme correspondant à la position occupée par la cible. La réponse du sujet fait immédiatement apparaître l'essai suivant. Pour certains sujets, et ce à leur insu, les cibles apparaissent selon une même séquence répétée au sein d'un bloc (e.g., 3-4-2-1-2-4-3-1-2-3,

les chiffres désignant les positions de gauche à droite), alors que pour un groupe contrôle, les cibles apparaissent selon une séquence aléatoire. Les résultats révèlent que les TR des sujets entraînés sur une séquence répétée diminuent davantage et plus rapidement que pour ceux entraînés sur un matériel aléatoire. Et pourtant, cet apprentissage ne semble pas être associé à une prise de conscience de la répétition du matériel. L'amélioration des performances semble précéder la prise de conscience possible si une pratique suffisante leur est accordée. Ici encore, l'amélioration des performances dans le test indirect de mémoire (i.e. la tâche de pistage de cible), en dépit d'une incapacité des sujets à verbaliser qu'une même séquence est continuellement répétée, a été répliquée de nombreuses fois avec différentes versions de la tâche de SRT. Par exemple, si la séquence répétée est remplacée par un matériel aléatoire au cours de la tâche, la suppression du patron séquentiel provoque un allongement des TR (e.g. Shanks & Channon, 2002).

2. L'apprentissage implicite : un débat terminologique

De par son origine hétérogène, l'apprentissage implicite est loin de recouvrir une définition faisant l'objet d'un consensus. Dans une revue de question, Fresch (1998) énumère pas moins de onze définitions de l'apprentissage implicite. Pour bien des chercheurs, l'apprentissage implicite est synonyme d'apprentissage inconscient. Dans ce cadre, l'apprentissage implicite fait référence à l'acquisition de connaissances sur les relations structurelles entre les stimuli, sans que l'individu ne prenne conscience de la connaissance qui en résulte (Berry & Dienes, 1993 ; Dienes & Perner, 1999). Mais d'autres auteurs préfèrent associer le terme implicite à celui d'intentionnalité plutôt qu'à celui d'inconscient. Par exemple, Perruchet et Nicolas (1998, p. 15) définissent l'apprentissage implicite comme un « *mode d'adaptation dans lequel le comportement d'un sujet apparaît sensible à la structure d'une situation, sans que cette adaptation ne soit imputable à l'exploitation intentionnelle de la connaissance explicite de cette structure* ». Selon Whittlesea et Dorken (1993), plutôt que de se focaliser sur l'absence de « conscience », les recherches sur l'apprentissage implicite devraient davantage chercher à déterminer des critères mesurant le rôle de l'intention durant l'apprentissage ou la congruence entre les demandes de la tâche durant l'apprentissage et l'utilisation ultérieure de la connaissance. Frensch (1998) insiste quant à lui sur certains

aspects de l'acquisition de l'information. Il propose d'envisager l'apprentissage implicite comme une forme d'apprentissage dans laquelle les connaissances sur les relations structurales entre les objets ou les événements sont acquises de manière non intentionnelle, incidente ou automatique. La nature consciente ou inconsciente de la connaissance acquise ne doit, selon lui, faire l'objet du concept d'apprentissage implicite. Cependant, Destrebecqz (2000) souligne la nécessité de définir le caractère conscient ou non conscient des connaissances acquises. Meulemans (1998) met également l'accent à la fois sur l'aspect incident de l'apprentissage implicite et sur le degré d'accès à la conscience des connaissances acquises, qui restent, selon lui, à un niveau inconscient. En outre, le problème fondamental soulevé dans ce corps de recherche n'en reste pas moins de savoir si des changements dans les comportements peuvent émerger sans être corrélés à un changement de l'expérience subjective et de déterminer la nature des traitements « cognitifs » non conscients. Aussi, nous définirons l'apprentissage implicite comme un processus d'adaptation par lequel le comportement d'un individu devient, de manière incidente, sensible à une structure, sans que la connaissance qui en résulte ne soit rapportable verbalement, voire accessible à la conscience.

L'absence d'une définition précise du concept d'apprentissage implicite conduit inévitablement à rencontrer ce terme dans des situations où les phénomènes décrits sont bien différents. A la lecture de nombreux écrits portant sur la mémoire, l'apprentissage, la perception et les processus de décision, notre impression est que des termes sont utilisés pour d'autres, sans que leurs champs respectifs aient réellement été définis. Certaines recherches, notamment sur la perception visuelle, s'appuient sur le concept d'apprentissage implicite pour décrire, soit un processus d'acquisition, soit la connaissance qui en découle, alors que d'autres encore l'utilisent pour décrire un processus de récupération. Etant donné qu'un apprentissage délibéré peut conduire à une connaissance inaccessible à la conscience, et qu'inversement un apprentissage incident peut conduire à une connaissance accessible à la conscience ou encore qu'une connaissance verbalisable peut être récupérée de manière incidente, le terme implicite est dans bien des situations sujet à controverse. Cette ambiguïté tient certainement au fait qu'il est impossible d'étudier des phénomènes d'apprentissage sans passer outre le recours à un processus de récupération. Si l'on accepte que l'apprentissage implicite produit par définition une connaissance inaccessible à la conscience, la récupération de cette connaissance est inévitablement implicite. Par conséquent, les tâches mettant en évidence des effets d'apprentissage implicite mettent de manière inhérente en évidence des phénomènes de

mémoire implicite. La mémoire implicite transparait lorsque la performance dans une tâche est facilitée en l'absence d'une récollection intentionnelle et consciente des expériences passées (pour des revues, Nicolas, 1994 ; Schacter, 1987). La mémoire implicite fait ainsi référence à l'influence que peut avoir un événement passé sur le comportement d'un individu sans que ce dernier ne prenne conscience de l'influence de cet événement.

3. L'apprentissage implicite produit-il une connaissance inconsciente ?

Les débats terminologiques autour du concept d'apprentissage implicite trouvent en partie leur origine dans les nombreux désaccords sur les méthodes utilisées pour éprouver la nature (explicite vs. implicite) de l'apprentissage et des connaissances qui en résultent. On présuppose qu'un traitement est implicite lorsqu'un test indirect de mémoire (e.g. jugement de grammaticalité, tâche de SRT, tâche de recherche visuelle de cible) est sensible au traitement de l'information alors qu'un test direct de mémoire y est insensible. Mais le problème majeur est de s'accorder sur ce qui constitue une mesure directe adéquate au traitement conscient (pour des revues critiques séminales, Holender, 1986 ; Shanks & St. John, 1994). Deux approches expérimentales sont classiquement utilisées pour explorer la nature explicite/implicite de l'apprentissage, l'une basée sur des mesures introspectives de la conscience (méthodes subjectives à la première personne) et l'autre basée sur des mesures comportementales de la conscience (méthodes objectives à la troisième personne).

3.1. Méthodes subjectives à la première personne

Dans les premiers travaux sur l'apprentissage implicite, les arguments en faveur du caractère implicite de l'apprentissage reposaient essentiellement sur l'échec des participants à expliciter la connaissance résultant de cet apprentissage (e.g. Berry & Broadbent, 1984 ; Lewicki, 1986 ; Nissen & Bullermer, 1987 ; Reber, 1967). Cependant, sous l'impulsion de Dulany, Carlson et Dewey (1984), les reports verbaux libres comme indicateurs fiables ou suffisants des connaissances explicites ont été grandement remis en question (e.g. Perruchet & Amorim, 1992 ; Shanks & St John, 1994). Pour réaliser la tâche indirecte, les sujets pourraient

en effet s'appuyer sur une connaissance certes difficilement verbalisable, mais néanmoins accessible à la conscience. Il est possible que l'apprentissage puisse résulter d'hypothèses partielles, incomplètes, éventuellement provisoires mais néanmoins conscientes. Selon Shanks et St. John (1994), une tâche de rappel libre ne remplit pas les critères « d'information » et de « sensibilité » requis pour exclure la possibilité que l'apprentissage et la connaissance soient explicites. Le critère d'information est atteint lorsque la connaissance à la base de l'amélioration des performances correspond à celle que l'examineur cherche à mesurer dans la phase testant les connaissances explicites. Or, il est dans bien des situations difficile d'être certain que les sujets ont vraiment appris ce que l'expérimentateur teste par la méthode des protocoles verbaux. Rien n'assure que les réponses libres des sujets ne sont pas simplement dépendantes du type de réponse attendue par l'expérimentateur. Par exemple, lorsque dans une tâche d'AGL l'expérimentateur demande aux sujets s'ils ont détecté des règles sous-tendant l'ordre de succession des lettres, il n'est pas certain que les performances observées dans la tâche de classification résultent bien de l'apprentissage de règles abstraites (nous discuterons ce point dans le paragraphe 5.1). Le critère de sensibilité concerne quant à lui le niveau d'accès conscient aux connaissances acquises. Il implique que la phase de description explicite de la tâche permette de détecter l'ensemble des connaissances conscientes qui ont pu influencer la performance des sujets. Faute de quoi, l'amélioration de celle-ci pourrait être attribuée à l'acquisition de connaissances implicites simplement parce que la phase de description explicite de la tâche est moins sensible à l'expression de l'information consciente que la tâche indirecte de mémoire. Or, non seulement le rappel libre laisse au sujet l'option de ne pas exposer ses connaissances s'il ne souhaite pas en faire état (en raison de l'incertitude de sa réponse par exemple), mais il lui laisse aussi le temps d'oublier momentanément des connaissances qu'il pourrait récupérer explicitement au moment de la tâche indirecte (inconvenient des rapports verbaux rétrospectifs).

Cependant, si certains auteurs s'opposent à l'utilisation des rapports verbaux comme mesure sensible de la conscience (e.g. Shanks & St John, 1994), d'autres revendiquent la pertinence des méthodes introspectives, à condition que celles-ci soient combinées avec d'autres types de données et interprétées avec soin (e.g. Ericsson & Simon, 1984 ; Ziori & Dienes, 2006). Goschke (1997, p. 262) remarquait qu'« *il est difficile, sinon impossible, de ne pas se fier au rapport verbal des sujets ou au respect des instructions verbales par les sujets comme critère d'une connaissance inconsciente* ». En effet, l'expérimentateur n'aura jamais accès aux états mentaux du sujet. Aussi, d'autres méthodes basées sur un critère subjectif ont

été développées, comme par exemple l'utilisation d'une échelle de confiance sur laquelle le sujet doit estimer le degré de certitude du choix qu'il a fait de classer un item comme congruent ou non. Si les sujets réussissent mieux que du seul fait du hasard la tâche directe alors qu'ils croient avoir répondu au hasard (*critère guessing*) ou que leurs performances ne sont pas positivement corrélées avec leurs jugements de confiance (*critère zéro-corrélation*), l'apprentissage est supposé implicite (e.g. Dienes & Altman, 1997 ; Ziori & Dienes, 2006).

Face aux critiques relatives aux méthodes des protocoles verbaux, des tests alternatifs ont été élaborés, tels que les tests de choix forcé (e.g. reconnaissance, générations, prédictions). Contrairement au rappel libre, les tests de choix forcé présentent l'avantage de contraindre le sujet à donner une réponse et de ne pas assimiler l'incapacité de verbalisation à l'absence de conscience.

3.2. Méthodes objectives à la troisième personne

A l'inverse des rapports verbaux, souvent appelés mesures subjectives à la première personne, les tests de choix forcé sont considérés comme des mesures objectives à la troisième personne. Il existe différents types de tâches de choix forcé, comme par exemple, les tâches de reconnaissance, de génération libre ou indicée. Ces tâches sont présentées immédiatement après la tâche indirecte de mémoire. Dans les tâches de reconnaissance, les sujets sont généralement exposés à des items vus pendant la phase d'apprentissage et à des items nouveaux. Leur tâche est classiquement de porter sur chacun de ces essais un jugement de reconnaissance, en terme de « familier vs. non familier » par rapport à la phase précédente. Si les sujets ne trouvent pas plus familiers les essais vus pendant la phase d'apprentissage que les essais nouveaux, l'apprentissage est considéré comme étant implicite. Dans les tâches de génération libre, il est demandé aux participants de générer des réponses incluant le plus d'informations possibles dont ils se souviennent du matériel d'apprentissage. Ce type de tâche présente l'avantage d'engager des traitements plus similaires à ceux engagés dans la tâche indirecte de mémoire, et de pallier certaines critiques émises contre l'utilisation des jugements de familiarité comme mesure sensible de la mémoire explicite (Shanks & St. John, 1994). Dans les tâches de SRT, elles consistent par exemple à demander aux participants de générer des réponses incluant le plus possible d'information relative aux séquences d'entraînement. Il existe des variantes dans lesquelles la tâche est indicée par une bribe de la séquence répétée.

Les tâches à choix forcé comme mesures des influences implicites vs. explicites ont elles aussi été soumises à la critique (e.g. Perruchet & Amorim, 1992 ; Shanks & St. John, 1994). Selon Shanks (2005), les performances à ces tâches dépendraient fortement de la motivation des sujets et du type de stratégies utilisées, parfois différentes de celles utilisées dans les tests indirects de mémoire (e.g. encodage analytique vs. non analytique). De plus, les résultats obtenus dans ces tâches sont souvent contradictoires d'une étude à l'autre. Ces tâches apportent parfois des arguments favorables au caractère inconscient de la connaissance acquise, parfois des arguments contraires. Par exemple, dans les tâches d'AGL, Dulany et al. (1984) ont montré que les sujets jugent plus familiers les exemplaires anciens que nouveaux, suggérant que les sujets peuvent accéder consciemment à leur connaissance, au moins en partie. Des tâches de complétion de fragments (Dienes, Broadbent, & Berry, 1991) corroborent ces résultats en montrant que les performances des sujets dans des tests objectifs sont systématiquement corrélées aux jugements de grammaticalité. De même, à l'issue des tâches d'apprentissage séquentiel, les participants reconnaissent souvent mieux que du seul fait du hasard les fragments de séquences (e.g. Perruchet & Amorim, 1992 ; Shanks & Johnstone, 1999 ; Willingham, Nissen, & Bullemer, 1989) et semblent bien souvent capables de générer librement des séquences structurées (e.g. Perruchet & Amorim, 1992 ; Shanks & Johnstone, 1999). A noter cependant que des résultats différents ont été observés dans certaines situations (e.g. Jiménez, Méndez, & Cleeremans, 1996 ; Willingham et al., 1989) et que les performances observées dans ces tâches objectives ne sont pas toujours corrélées avec les jugements de confiance des participants (e.g. Dienes & Altman, 1997).

Si pour certains les tâches de choix forcé ne répondent pas de manière satisfaisante aux critères d'information et de sensibilité requis pour exclure la possibilité que la connaissance est explicite, pour d'autres, elles ne permettent en rien d'invalider l'influence de connaissances implicites dans les réponses. Pour les partisans de l'apprentissage inconscient, même si les participants utilisent des connaissances purement implicites dans le test indirect de mémoire, ils pourraient produire des réponses différentes du hasard dans les tâches de choix forcé sur la base d'autres considérations. Selon Goschke (1997, p. 257), *“Even if explicit knowledge as revealed by a direct task is in fact sufficient to account for performance in the indirect test, this does not necessarily imply that learning was not implicit”*. De même, Berry et Dienes (1993, p. 15) écrivaient *“performance cannot be at chance on all discriminative tests, otherwise no measure would remain to existence of implicit learning”*.

Par exemple, dans une tâche de génération libre conduite après l'apprentissage des séquences répétées, l'apprentissage sensori-moteur pourrait être reproduit de manière automatique sans pour autant que le sujet n'ait, au cours de l'apprentissage, conscience du développement de cette automaticité (Destrebecqz & Cleeremans, 2001). Ou encore, les jugements de reconnaissance pourraient refléter à la fois des récollections conscientes et des sentiments de familiarité qui eux ne sont pas nécessairement accompagnés de récollection consciente. Les performances dans les tests directs pourraient être « contaminées » par des connaissances implicites et explicites, sous-entendant que ces tâches ne constituent pas des mesures «*process-pure*» des connaissances implicites et explicites (pour une revue, Goschke, 1997).

Une critique néanmoins irréfutable à l'encontre des tâches à choix forcé comme preuve tangible des effets d'apprentissage implicite, est qu'ils présupposent l'acceptation d'une hypothèse nulle (i.e. l'hypothèse d'absence d'effet direct). De ce fait, il est impossible de démontrer statistiquement que l'information pertinente n'a pas été apprise consciemment. De la même manière que les « révisionnistes » de l'apprentissage implicite ne pourront jamais démontrer l'inexistence de connaissances inconscientes, les avocats de l'apprentissage implicite ne pourront jamais démontrer de manière satisfaisante leur existence par le recours à ces méthodes.

Aussi, la logique de dissociation, consistant à montrer que les conséquences d'un apprentissage implicite sont qualitativement différentes de celles d'un apprentissage explicite, est de plus en plus répandue. La procédure de dissociation de processus a été développée par Jacoby (1991) dans le domaine de la mémoire implicite pour évaluer les contributions respectives de la mémoire consciente et inconsciente dans des tâches de complèvement de fragments de mots⁹. Le présupposé à l'origine de cette méthode est que contrairement à la mémoire explicite, la mémoire implicite ne peut être affectée par une instruction explicite. Cependant, de manière générale, les procédures de dissociation appliquées aux paradigmes classiques d'apprentissage implicite permettent essentiellement d'estimer la contribution des processus automatiques et contrôlés dans la tâche directe ou indirecte de mémoire, plus que de vraiment démontrer la nature inconsciente ou consciente des connaissances acquises.

⁹ Dans une expérience de complèvement de fragments de mots, les participants étudient dans un premier temps une liste de mots. Durant une phase test, ils sont exposés à des fragments de mots qu'ils doivent compléter en incluant ou en évitant les mots de la liste d'étude. Si les participants ne sont pas capables d'exclure les mots de la liste d'étude dans la condition d'exclusion, cela présuppose l'influence de traitements automatiques et non conscients.

Compte tenu de résultats aussi discutables, certains auteurs (e.g. Perruchet & Vinter, 2002 ; Shanks & St John, 1994) rejettent l'intuition courante qui accorde une place importante aux traitements cognitifs non conscients dans nos comportements. Pour Shanks, les données de la littérature n'apportent aucune preuve tangible quant à l'existence de connaissances implicites, mais démontrent qu'au contraire « *Human learning is systematically accompanied by awareness* » (Shanks & St John, 1994). Selon lui, les performances dans les tâches d'apprentissage implicite sont systématiquement corrélées à la connaissance consciente. Perruchet, Vinter et Gallego (1997) proposent que si des processus d'apprentissage influencent le traitement de manière implicite, les connaissances acquises sont quant à elles systématiquement accessibles à la conscience. Si tous s'accordent à penser que l'individu peut devenir sensible à la structure d'une information sans pour autant être engagé dans une tâche d'apprentissage délibéré, la question majeure posée est de savoir si des changements comportementaux peuvent émerger sans être accompagnés de changements dans l'expérience subjective. Le caractère non intentionnel de l'apprentissage implicite amène à s'interroger sur son caractère automatique. L'apprentissage implicite peut-il être vu comme un apprentissage automatique, i.e. non intentionnel et non contraint par les limites attentionnelles ?

4. L'apprentissage implicite est-il un processus automatique, qui ne requiert pas d'attention ?

Dans la plupart des esprits, l'apprentissage implicite est un processus automatique. Ttzelgov (1997, p. 441) avait déjà fait le constat suivant : “*Psychologists frequently do not distinguish between automaticity and the absence of consciousness, even to the point of using the terms “automatic” and “unconscious” interchangeably*”. Le concept d'automaticité, au même titre que la plupart des concepts utilisés en psychologie, est désigné par un terme du langage courant ne possédant pas une rigueur formelle suffisante pour permettre d'en donner une définition cohérente. L'automaticité est couramment considérée comme un mode de fonctionnement rapide, « *effortless* », qui ne consomme pas de ressource attentionnelle, autonome, stéréotypé, inaccessible à la prise de conscience, dépendant de la tâche. Ces critères ne se recouvrant que rarement, l'absence de coût attentionnel et l'absence de contrôle

intentionnel ont été considérées comme les propriétés nécessaires et suffisantes pour qualifier un processus d'automatique. Dans la section suivante, nous nous placerons dans une conception de l'automatisme en tant que mode de fonctionnement autonome qui ne requiert pas d'attention. Mais même si l'on restreint le concept d'automatisme au problème de l'attention, nombreux sont ceux pour qui conscience et attention sont synonymes. Besner et Stolz (1999a, p. 453) remarque ainsi que : “ *In many account (...) consciousness is used as though it is synonymous with “attention” and unconsciousness with the absence of “attention”. Unconscious processes are therefore said to be automatic* ”.

De nombreuses théories considèrent en effet l'apprentissage implicite comme un traitement automatique, i.e. qui ne requiert pas d'attention (e.g. Berry & Dienes, 1993 ; Lewicki, 1986 ; Schneider & Shiffrin, 1977), comme “*an unselective and passive aggregation of information about the co-occurrence of environmental events and features*” (Hayes & Broadbent, 1988). Selon cette conception, l'apprentissage implicite serait un traitement absorbant les structures du monde (Lewicki & Hill, 1989) et permettrait une adaptation passive à ces structures (Reber, 1993). De nombreuses recherches suggèrent en effet que l'apprentissage et la mémoire implicite engagent des demandes attentionnelles minimales. Par exemple, plusieurs études inscrites dans le domaine de la perception visuelle semblent montrer que des objets non attendus laissent néanmoins des traces implicites pouvant être révélées par des méthodes indirectes (e.g. DeSchepper & Treisman, 1996 ; Moore & Egeth, 1997). Ces résultats ont amené un certain nombre d'auteurs à penser que l'attention n'est pas un pré-requis pour la perception et la mémoire, mais seulement pour la perception consciente et la mémoire explicite. Les traitements implicites seraient exempts des limitations attentionnelles (Mack & Rock, 1998). Cependant, les recherches inscrites dans le domaine de l'apprentissage implicite sont moins claires sur ce point. Certaines apportent des arguments en faveur d'un apprentissage automatique, alors que d'autres au contraire, suggèrent que les traitements implicites, au même titre que les traitements explicites, sont dépendants d'un traitement attentionnel (pour des revues en faveur de cette hypothèse, Chun & Turk-Browne, 2007 ; Perruchet & Pacton, 2006). Ces résultats très controversés semblent dépendre, entre autres, des méthodes employées pour tester le rôle de l'attention sur l'apprentissage implicite.

L'attention étant un concept lui-même difficile à définir, on peut comprendre les raisons de pareilles contradictions. Il est admis que l'attention n'est pas un processus unitaire correspondant à une opération mentale unique. L'attention peut être fragmentée en une variété

de composantes pouvant rendre compte de façon très différente de la diversité des changements qualitatifs qu'elle rend possible (pour une revue, Camus, 2003). L'attention fait à la fois référence à une « capacité mentale » ou « ressource limitée » utilisées dans des tâches qui requièrent un effort mental, et à un « processus de sélection », focalisant les ressources sur certains événements au dépend des autres. Le rôle de l'attention sur l'apprentissage en termes de ressources est classiquement étudié à l'aide d'une tâche concurrente à la tâche principale d'apprentissage. On présuppose que si cette condition d'exercice induit une baisse de la performance par rapport à une condition contrôle où chacune des deux tâches est effectuée isolement, la tâche principale consomme des ressources attentionnelles. En revanche, une absence de détérioration montre que le traitement ne consomme pas de ressources attentionnelles. Le rôle de l'attention en tant que processus de sélection est quant à lui souvent étudié en manipulant des conditions expérimentales dirigeant le focus attentionnel des participants sur certains aspects du matériel, de sorte que ces derniers soient attendus et que le reste du matériel soit ignoré.

4.1. L'apprentissage implicite consomme-t-il des ressources attentionnelles ?

Le paradigme de SRT a été l'un des outils les plus utilisés pour étudier si l'apprentissage implicite consomme des ressources attentionnelles. Dans leur recherche princeps, Nissen et Bullemer (1987) avaient déjà abordé cette question à l'aide d'une tâche concurrente à la tâche principale. Dans la condition de double tâche, un son aigu ou grave était présenté au cours de chaque essai. Les participants devaient compter le nombre de sons aigus. Cette tâche secondaire de comptage de sons (« *tone-counting* ») avait pour but de diminuer la quantité de ressources attentionnelles disponible pour réaliser la tâche de SRT. Nissen et Bullemer ont alors observé que l'apprentissage séquentiel explicite mais également implicite était éliminé lorsque les sujets réalisaient parallèlement à la tâche de SRT la tâche secondaire, suggérant que les mécanismes en œuvre dans l'apprentissage de séquences répétées requièrent des ressources attentionnelles. Ce résultat a par la suite été répliqué dans de nombreuses études (e.g. Shanks & Channon, 2002 ; Shanks, Rowland, & Ranger, 2005). Même lorsqu'un entraînement en condition de simple tâche précédait le test en condition de double tâche pour la même séquence répétée, l'apprentissage mis en évidence dans la condition de simple tâche était deux fois plus important que celui observé dans la condition de double tâche (Curran & Keele, 1993). Cependant, Cohen, Ivry et Keele (1990) ont montré que la tâche secondaire

affecte uniquement l'apprentissage de séquences ambiguës mais pas l'apprentissage de séquences uniques ou de séquences hybrides. Une étude de Jiménez et Méndez (1999) suggère également que l'apprentissage de séquences répétées n'est pas affecté par une tâche secondaire de comptage de symboles.

D'autres explications ont par la suite été proposées pour rendre compte des effets délétères de la tâche secondaire sur les effets d'apprentissage implicite. Le déficit d'apprentissage dans la condition de double tâche pourrait tenir à des effets d'interférence spécifique plutôt qu'à un manque de capacité générale. Stadler (1995) a ainsi proposé que la tâche de comptage de sons puisse interférer avec l'apprentissage de la séquence en dégradant son organisation. Selon Frensch, Buchner et Lin (1994), la condition de double tâche pourrait diminuer le temps durant lequel deux stimuli consécutifs sont maintenus en MCT. De plus, certaines études ne montrent aucun effet de la difficulté de la sélection attentionnelle de la cible sur l'apprentissage des séquences (e.g. Rowland & Shanks, 2006). Même Shanks (2005, p. 206) reconnaît *“There is some evidence to support the claim that reductions in attentional capacity can be incurred without detectable cost on implicit learning and this plainly supports the idea that the latter is, to some extent, unintentional and automatic”*.

Dans le domaine de l'AGL, la plupart des travaux suggèrent quant à eux un effet délétère des doubles tâches sur les performances. Par exemple, Dienes, Broadbent et Berry (1991) ont montré qu'une tâche secondaire concurrente interfère avec les performances de classification, non seulement quand cette dernière est présentée pendant la phase d'étude mais également lorsqu'elle est présentée pendant la phase test (voir aussi, Dienes & Altman, 1997 ; cependant, Helman & Berry, 2003).

Les résultats concernant les effets délétères des tâches secondaires rapportés dans la littérature, sont somme toute assez équivoques, et ne permettent pas de déterminer de manière tranchée si l'apprentissage implicite consomme ou non des ressources attentionnelles. Le principe des doubles tâches reste par ailleurs discutable. Le rôle des tâches secondaires est de réduire la quantité de ressource attentionnelle disponible pour la tâche principale. Or la plupart des tâches secondaires utilisées dans les études rapportées ici mettaient en œuvre la modalité auditive, e.g. comptage de sons (pour d'autres types de tâches secondaires, cf. Heuer & Schmidtke, 1996). Etant donné que des ressources attentionnelles « distinctes » sont très certainement allouées aux traitements auditifs et visuels, il est fort probable que les tâches

auditives n'interfèrent pas complètement avec les tâches visuelles (Duncan, Martens, & Ward, 1997).

4.2. L'apprentissage implicite requiert-il une attention sélective ?

Les recherches testant le rôle de l'attention sélective dans l'apprentissage implicite sont tout autant litigieuses. Par exemple, dans une tâche de SRT, Mayr (1996) a observé des effets d'apprentissage concernant les séquences régulières des localisations de formes géométriques, alors que les participants étaient invités à répondre à l'identité de ces formes et non à leur localisation. Les localisations n'étaient bien sûr pas corrélées avec l'identité des formes. Un résultat contraire a cependant été rapporté par Willingham et al. (1989) lorsque la tâche des participants était de répondre à la couleur des cibles et non à leur identité. La distance entre les localisations spatiales, beaucoup plus importante dans l'étude de Mayr, pourrait expliquer ces différences. De manière concordante, Goschke (1996) a montré que des séquences de lettres et de localisations non corrélées peuvent être apprises simultanément, même lorsque les sujets devaient fixer un point de fixation au centre de l'écran. Pourtant, une étude de Jiménez et Méndez (1999) va à l'encontre de ces résultats. Dans leur étude, les participants devaient pister des symboles et répondre à la localisation qu'ils occupaient sur l'écran. Les régularités intégraient une relation prédictive entre les symboles et les localisations successives. Dans une condition de double tâche, les participants devaient compter certains des symboles utilisés. Dans la condition de simple tâche, les sujets n'avaient pas à compter les symboles, et donc ils ne leur prêtaient pas attention de manière sélective. Les résultats révèlent que l'apprentissage de la séquence simple n'était pas affecté par la tâche secondaire en situation de double tâche, suggérant que l'apprentissage séquentiel ne consomme pas de ressources attentionnelles. En revanche, seuls les participants assignés à la condition de double tâche apprenaient la relation symbole-localisation. Jiménez et Méndez en ont conclu qu'une attention sélective portée sur les items est requise pour qu'un apprentissage ait lieu.

De façon similaire, les travaux de Turk-Browne, Jungé et Scholl (2005) suggèrent que l'attention sélective module l'apprentissage des régularités visuelles temporelles dans une tâche d'apprentissage de régularités statistiques (Fiser & Aslin, 2002, pour une description du paradigme d'apprentissage de régularités statistiques, cf. Annexe 2). L'orientation de l'attention sélective était ici manipulée en faisant apparaître la moitié des formes en rouge et

l'autre moitié en vert. La consigne invitait les participants à ne prêter attention qu'aux formes de l'une ou l'autre de ces deux couleurs. Les résultats montrent que contrairement aux triplets attendus, les performances de discrimination sur les triplets ignorés n'étaient pas différentes du hasard. D'autres études corroborent l'hypothèse que l'attention sélective est cruciale à l'apprentissage implicite, notamment dans une tâche de détection de changements (e.g. Yi & Chun, 2005), alors que d'autres montrent des effets d'apprentissage perceptif sur des stimuli en mouvement en dehors du focus attentionnel (pour une revue, Seitz & Watanabe, 2005).

4.3. L'attention est-elle requise pour l'apprentissage implicite et/ou pour l'expression de la connaissance résultante ?

Des travaux récents révèlent un phénomène intéressant. Les processus attentionnels pourraient exercer une influence différente sur l'apprentissage implicite et sur l'expression de la connaissance résultante. Frensch, Lin et Buchner (1998) ont ainsi montré que l'attention serait nécessaire à l'expression de l'apprentissage séquentiel, mais pas à l'apprentissage lui-même. Les participants étaient exposés aux mêmes séquences répétées en conditions de simple tâche et de double tâche (la tâche secondaire était une tâche de comptage de sons). Ils réalisaient d'abord la tâche en condition de simple tâche puis en condition de double tâche ou l'inverse. Les sujets entraînés en condition de double tâche manifestaient un effet d'apprentissage immédiat lorsqu'ils étaient testés en condition de simple tâche. En revanche, l'effet d'apprentissage observé en condition de simple tâche disparaissait lorsqu'ils étaient exposés aux mêmes séquences en condition de double tâche. Ce résultat suggère que l'apprentissage des séquences répétées peut se développer parallèlement à la réalisation de la tâche secondaire, mais que l'expression de cet apprentissage est affectée par la tâche secondaire. Frensch, Wenke et Rüniger (1999) ont reproduit ces effets et confirmé que ces derniers ne tenaient pas à un manque d'interférence (voir aussi Curran & Keele, 1993). Cette dissociation entre le rôle de l'attention sur l'apprentissage et celui sur l'expression de la connaissance résultante a également été mise en évidence par Jiang et Leung (2005) avec le paradigme d'indication contextuel. Nous décrirons en détail ces résultats dans le Chapitre III dédié au paradigme d'indication contextuel.

Turk-Brown, Jungé et Scholl (2005) ont toutefois observé un résultat contraire avec le paradigme d'apprentissage de régularités statistiques. Lorsque les couleurs des formes étaient

inversées durant une phase de transfert, les formes ignorées durant la phase de familiarisation n'étaient pas mieux discriminées lorsqu'elles devenaient attendues durant la phase test. En revanche, les formes attendues pendant la phase de familiarisation persistaient à être mieux reconnues lorsqu'elles devenaient ignorées. L'apprentissage des co-occurrences entre les formes requerrait donc une attention sélective, mais pas la récupération. Néanmoins, les tâches de discrimination ne sont pas toujours sensibles aux connaissances implicites des sujets. L'inconvénient majeur de ce paradigme est de mesurer des effets d'apprentissage par le recours à une tâche directe d'accès aux connaissances (i.e. une tâche de reconnaissance).

En résumé, compte tenu de données empiriques aussi divergentes, il paraît nécessaire d'émettre des réserves sur le caractère automatique vs. non automatique de l'apprentissage implicite. Si un apprentissage de co-variations semble, dans certaines conditions, émerger en l'absence d'une attention focalisée ou semble ne pas consommer de ressources attentionnelles (tout du moins centrales), de nombreuses autres conditions expérimentales suggèrent le phénomène inverse. La procédure expérimentale, e.g. consigne, difficulté du matériel d'apprentissage, tests utilisés pour mettre en évidence des effets d'apprentissage, pourraient constituer des variables susceptibles de rendre compte de ces contradictions.

5. L'apprentissage implicite peut-il conduire à des connaissances abstraites?

La question sans doute la plus fondamentale posée dans les recherches sur l'apprentissage implicite concerne les niveaux de profondeur des connaissances inaccessibles à la conscience. Cette problématique peut être appréhendée sous deux aspects différents que nous traiterons séparément, car ils relèvent à notre avis de deux questions différentes : Peut-on apprendre implicitement des règles abstraites à partir d'un matériel complexe ? L'apprentissage implicite peut-il reposer sur des propriétés sémantiques déjà en mémoire ?

5.1. L'apprentissage implicite de règles abstraites

La question concernant le degré d'abstraction des connaissances résultant d'un apprentissage implicite a fait l'objet d'une littérature extrêmement controversée. C'est certainement à travers les travaux sur l'AGL que cette question a été la plus discutée. Selon Reber (1989), les résultats observés dans les tâches d'AGL signent l'existence d'un système d'abstraction non conscient pouvant enregistrer passivement une information. Par conséquent, ce qui est acquis serait une connaissance tacite relative à une représentation abstraite et inconsciente de la structure contenue dans l'information présentée. Mais cette lecture des résultats a rapidement fait l'objet d'une polémique. Pour de nombreux auteurs, les résultats obtenus dans les tâches d'AGL ne fournissent pas de preuve convaincante quant à l'existence d'une abstraction inconsciente de règles, ou d'une représentation inconsciente et abstraite des connaissances (pour des revues cf. Didierjean, 2001 ; Nicolas, 1996 ; Pothos, 2007). La position abstractionniste de Reber a été rapidement ébranlée par des explications alternatives basées, non sur l'abstraction implicite de règles, mais sur le stockage et sur la récupération d'exemplaires particuliers (e.g. Brooks & Vokey, 1991 ; Perruchet & Pacteau, 1990). Brooks (1978) a été le premier à souligner que la grammaticalité et la similarité entre les items de la phase d'étude et ceux de la phase test, sont systématiquement confondues dans les procédures classiques. Aussi les jugements de transfert pourraient-ils simplement tenir à la ressemblance analogique et à la similarité perceptive qui existe entre les exemplaires d'étude et de test. Pour tester cette hypothèse, Vokey et Brooks (1992) ont manipulé lors de la phase test, de manière orthogonale, la similarité et la grammaticalité des suites de lettres utilisées dans la phase d'étude. Leurs résultats ont montré une supériorité dans le classement des suites grammaticales sur les suites non grammaticales, mais également une supériorité dans le classement des suites similaires sur les suites dissimilaires. Vokey et Brooks ont dans un premier temps conclu à un modèle mixte, accordant encore un rôle aux processus d'abstraction en plus des mécanismes mnésiques. Mais parallèlement, Perruchet et Pacteau (Perruchet, 1994 ; Perruchet & Pacteau, 1990) ont avancé l'hypothèse selon laquelle les sujets pourraient mémoriser, non pas des exemplaires dans leur globalité (encodage holistique), mais des fragments d'exemplaires, c'est-à-dire des bigrammes ou trigrammes présentés avec une fréquence plus ou moins importante lors de la phase d'étude. Leurs travaux ont alors montré qu'une hypothèse en termes de fragments d'exemplaires comme unités perceptives de la connaissance, pouvait à elle seule rendre compte des performances observées dans les travaux de Vokey et Brooks (1992) et dans la plupart des recherches sur l'AGL (cf. Perruchet, 1994).

Mais plus encore, selon Perruchet et Pacteau, si l'on se place dans un cadre « mnémocentriste », le recours aux mécanismes d'apprentissage inconscient ne présente plus aucune utilité pour expliquer les données. Le taux de reconnaissance explicite des unités perceptives apparaît suffisant pour rendre compte des jugements de grammaticalité. De ce fait, selon Perruchet et Pacteau, la connaissance acquise durant la phase d'étude serait entièrement disponible à la conscience, rendant possible l'utilisation intentionnelle de cette connaissance lors de la phase test.

Le domaine de l'apprentissage implicite a alors été scindé en deux écoles, l'une prêchant une conception « abstractionniste » de l'apprentissage implicite, et l'autre prêchant une conception « mnémocentriste » ou « non abstractionniste ». Selon la conception abstractionniste, l'apprentissage implicite observé reposerait principalement sur l'extraction de propriétés abstraites générales (e.g. Reber, 1989), alors que selon la conception mnémocentriste, l'apprentissage reposerait exclusivement sur l'extraction de propriétés spécifiques du matériel d'étude. Pour les partisans de l'école « mnémocentriste », dans la phase d'étude, les participants ne feraient aucune induction mais « mémoriseraient » simplement des exemplaires d'items entiers (point de vue « exemplariste ») (Brooks, 1978 ; Brooks & Vokey, 1991) ou des fragments d'items (point de vue « fragmentariste », Perruchet, 1994 ; Perruchet & Pacteau, 1990). Lors de la deuxième phase expérimentale, les sujets porteraient un jugement de catégorisation en termes de ressemblances globales ou partielles entre les cas stockés en mémoire épisodique et les nouveaux cas à juger. Dans ce débat, l'une des principales difficultés pour dissocier l'utilisation de règles et l'utilisation d'exemplaires spécifiques tient au rôle des traits de surface dans la construction du matériel utilisé. Il est en effet difficile de savoir si la généralisation d'une solution provient de la projection des traits de surface ou si elle provient de la projection d'une structure abstraite (pour une revue, Didierjean, 2001).

Après des années de recherche et de débats houleux autour de l'AGL, et même si certains travaux suggèrent des effets de transfert positifs d'un groupe d'éléments à un autre (e.g. Altman, Dienes, & Goode, 1995¹⁰ ; Mathews, Buss, Stanley, & Blanchard-Fields, 1989¹¹ ;

¹⁰ Les auteurs ont montré un transfert intermodal et interdomaines dans des situations où les sujets ne semblaient pas avoir conscience des correspondances structurales.

¹¹ Les participants étaient exposés à des suites de lettres différentes au cours d'une session de transfert.

Reber, 1969¹²), on ne peut nier aujourd'hui que des mécanismes d'apprentissage associatif élémentaire ou de « *chunking* » sont suffisants pour rendre compte des performances observées dans la plupart des recherches d'AGL. Il est peu probable que les sujets aient acquis une connaissance qui s'apparente à la structure abstraite établie par l'expérimentateur.

Mais quand bien même les performances obtenues dans les tâches d'AGL reposeraient exclusivement sur l'apprentissage explicite de propriétés spécifiques, elles ne démontreraient en rien qu'aucun mécanisme d'apprentissage implicite sur des propriétés abstraites ne soit mis en œuvre dans d'autres types d'activités cognitives. Toutefois, à ce jour, peu de recherches ont rapporté des données mettant en évidence des effets d'apprentissage implicite de règles abstraites qui ne soient par la suite remis en question. Ainsi par exemple, les travaux conduits par Lewicki et collaborateurs (pour une revue, cf. Lewicki, Hill, & Czyzewska, 1992), suggérant des effets d'apprentissage de co-variations dans un contexte de jugements sociaux ont été soumis au même scepticisme (pour une revue, cf. Perruchet & Vinter, 2002).

5.2. Les traitements implicites basés sur des régularités sémantiques

Le débat théorique autour de l'AGL s'est par la suite généralisé à l'apprentissage implicite et à la cognition sans conscience. Nombreux sont ceux pour qui la littérature ne fournit pas de preuve convaincante en faveur de l'existence de mécanismes d'abstraction implicite de règles (e.g. Perruchet & Vinter, 2002 ; Shanks, 2005). Cependant, sans vraiment avoir recours à des processus d'abstraction aussi complexes que ceux nécessaires pour extraire les règles présentes dans les grammaires artificielles, on peut envisager que la cognition sans conscience puisse fonctionner par manipulation de connaissances conceptuelles déjà en mémoire. Des mécanismes d'apprentissage implicite peuvent-ils se développer à partir de connaissances déjà en mémoire ou sont-ils indépendants de l'état préalable du système qu'ils modifient ? L'apprentissage implicite peut-il reposer sur des propriétés sémantiques ? Si les travaux sur l'apprentissage de règles abstraites fait l'objet d'une littérature pour le moins étoffée, peu d'études se sont directement intéressées à l'apprentissage implicite de régularités sémantiques.

¹² Reber utilisa des symboles différents dans la phase d'étude et de test. L'information de surface changeait entre les deux phases.

Une étude de Lambert et Sumich (1996) suggère des effets d'apprentissage implicite sémantique. Leurs expériences se déroulaient de la manière suivante. Les participants devaient répondre le plus rapidement possible à la localisation d'une cible (un carré blanc) présentée dans le champ visuel droit ou dans le champ visuel gauche. Inspirés des travaux de Posner (1980), la cible pouvait être précédée d'un indice indiquant la localisation probable de la cible. L'indice était ici défini par un mot relatif à une catégorie sémantique (vivant vs. non vivant) présenté durant 67ms que le sujet devait ignorer. La catégorie sémantique était dans la majorité des essais associée à la même localisation de la cible (e.g. la catégorie vivant était dans la majorité des essais associée à la localisation gauche de la cible). Les résultats montrent qu'après de nombreux blocs d'essais, l'attention des sujets était orientée de manière plus efficace vers le champ visuel approprié. Les TR devenaient plus courts lorsque la cible apparaissait dans les localisations très probables que dans les localisations peu probables. Pourtant, les sujets ne rapportaient pas avoir détecté consciemment cette régularité. De plus, ces effets d'apprentissage étaient transférés à une condition dans laquelle les mots étaient nouveaux, suggérant que l'apprentissage repose sur la catégorie sémantique des mots et non sur la répétition de leurs traits visuels. Mais ici encore, l'apprentissage était défini comme implicite en raison de l'incapacité des participants à verbaliser les règles de co-occurrences. D'autres travaux ont exploré le rôle des connaissances préalables sur les effets d'apprentissage implicite et semblent indiquer des effets d'apprentissage implicite basé sur des propriétés sémantiques (e.g. Didierjean, 2007 ; Didierjean & Lemaire, 2005 ; Ziori & Dienes, 2006)

L'existence de traitements sémantiques non conscients a été plus largement explorée dans le domaine de la perception sans conscience (i.e. subliminale). Ce domaine de recherche révèle une histoire comparable à celle de l'apprentissage implicite, animée par des débats théoriques et méthodologiques similaires. Est-ce que la signification d'un mot peut être traitée et manipulée inconsciemment pour faciliter les réponses comportementales ? Suite aux travaux de Marcel (1983), suggérant des effets d'amorçage sémantique non conscients, de vives critiques ont été soulevées quant à l'interprétation des résultats observés (e.g. Holender, 1986). Ici encore les controverses tiennent à la pluralité des interprétations susceptibles de rendre compte des résultats. En effet, l'utilisation des mêmes items, à la fois comme amorces et comme cibles, permet d'interpréter les effets observés en termes d'amorçage basé sur les propriétés sémantiques des mots, mais également en termes d'amorçage de répétition, basé simplement sur leur forme spécifique.

Dans le prolongement d'une étude suggérant des effets d'amorçage subliminal sémantique (Draine & Greenwald, 1998), Abrams et Greenwald (2000) ont étudié le rôle joué par la répétition des items sur les effets d'amorçage observés. Au cours de cette étude, les sujets devaient évaluer le plus rapidement et correctement possible la valence (positive vs négative) de mots cibles perçus consciemment. La présentation d'un mot cible était précédée d'un mot amorce dont la valence pouvait être congruente ou non congruente avec celle de la cible. Comme dans l'expérience originale, une amorce congruente facilitait la réponse des sujets, alors qu'une amorce non congruente interférait avec leur réponse. Mais lorsque les auteurs ont examiné la généralisation des effets à des amorces nouvelles (i.e. jamais utilisées comme cibles), un effet d'amorçage n'était alors observé que pour les mots partageant des fragments de lettres communs avec les mots utilisés comme cible. Par exemple, après plusieurs classifications conscientes « négative » des mêmes mots cibles « smut » et « bile », l'amorce subliminale « smile » facilitait non pas une réponse positive, mais une réponse négative. Ce résultat montre d'une part le rôle considérable joué par les cibles traitées consciemment dans les effets d'amorçage, et d'autre part que les effets des amorces ne reposent pas sur un accès à leur sémantique, tout du moins dans cette situation. Bien au contraire, comme dans l'apprentissage séquentiel, les sujets auraient appris l'association entre des fragments de cible et la réponse motrice spécifique à donner. Ainsi, loin de mettre en évidence l'effet d'un traitement sémantique inconscient, ces résultats montrent comment l'apprentissage sensori-moteur peut être généralisé à des amorces composées de fragments de lettres perçus au préalable consciemment (pour un résultat concordant, Damian, 2001).

Néanmoins, plusieurs études récentes semblent démontrer de manière convaincante que les effets d'amorçage peuvent reposer sur un traitement sémantique de l'amorce et pas simplement sur un traitement de sa forme. Le biais de répétition évoqué précédemment a été neutralisé en utilisant des séries d'items amorces différents des séries d'items cibles. Des effets d'amorçage subliminal basés sur la signification d'images ont été observés alors que ces dernières n'avaient jamais été utilisées comme cibles (Dell'Acqua & Grainger, 1999). Un résultat convergent a été obtenu dans des études récentes mettant en œuvre une tâche de comparaison de magnitude numérique avec des amorces masquées (Greenwald, Abrams, Naccache, & Dehaene, 2003 ; Naccache & Dehaene, 2001), ou encore dans une tâche de classification lexicale en terme de valence positive/négative et dans une tâche de classification de genre (Klauer, Eder, Greenwald, & Abrams, 2007). Si ces résultats apportent des

arguments en faveur de l'existence de traitements sémantiques inconscients, les travaux de Dehaene et collaborateurs sont toutefois très controversés. Deux critiques sont principalement énoncées. La première relève de problèmes méthodologiques. Overgaard, Rote, Mouridsen et Ramsoy (2006) soulignent le manque de sensibilité des mesures de la perception consciente utilisées dans les travaux de Dehaene et collaborateurs. La deuxième fait référence au caractère indissociable de la conscience et de l'attention.

A défaut d'arguments irréfutables en faveur de traitements « abstraits » non conscients, certains auteurs ont proposé que les connaissances qui découlent d'un apprentissage implicite tendent à être relativement inflexibles, et portent davantage sur les traits de surface du matériel que sur sa structure abstraite (voir aussi Schacter, 1987 pour un point de vue plus nuancé mais néanmoins convergent dans le domaine de la mémoire implicite). Dans ce cadre, l'apprentissage implicite se caractériserait par une spécificité de transfert. Toutefois, la question de l'apprentissage implicite de régularités sémantiques reste ouverte.

6. Conclusion et perspectives

Au début des années 80, avant l'émergence des polémiques dans le domaine de l'AGL, Reber tirait comme conclusion de ses travaux que des mécanismes d'apprentissage inconscient auraient une généralité considérable (Reber, 1989). En permettant l'abstraction, les mécanismes d'apprentissage inconscients sous-tendraient selon Reber (voir aussi, Lewicki, et al., 1992) le développement du langage, de l'expertise et même plus généralement les règles de la socialisation. Pour de nombreux chercheurs, la dialectique qui peut exister entre le peu d'information qui émerge à la conscience et la complexité de nos comportements ne peut être expliquée que par l'intervention de processus d'abstraction non conscients qui structurent l'environnement. Mais ce point de vue est loin d'être admis de tous. Pour certains auteurs, l'apprentissage est toujours accompagné d'une prise de conscience (Perruchet & Vinter, 2002 ; Shanks & St John, 1994). Perruchet et Nicolas (1998, p. 23) rejettent le raisonnement selon lequel *« la simple observation de ce qui se passe en milieu naturel suffit à donner la certitude que les sujets ont abstrait inconsciemment les règles qui structurent leur environnement, et s'adaptent à des situations nouvelles en appliquant inconsciemment ces*

règles. Si l'on ne retrouve pas ce phénomène dans les situations de laboratoire, c'est que celles-ci ne constituent pas de reproductions valides des situations naturelles ». A travers le modèle SOC ("Self-Organizing Consciousness"), Perruchet et Vinter (2002, p. 298) défendent la thèse suivante: *"Our analysis leads to the surprising conclusion that there is no need for the concepts of unconscious representations and knowledge and, a fortiori, the notion of unconscious inferences: Conscious mental life, when considered within a dynamic perspective, could be sufficient to account for adapted behavior"*.

Cherchant à réconcilier ces deux approches du fonctionnement cognitif, Cleeremans et Jiménez (2002) ont revisité de nombreuses théories psychologiques sur la nature de la conscience, qu'ils ont regroupées de manière caricaturale en deux catégories, en fonction du rôle qu'elles lui attribuent : les théories « Commandant Data » et les théories « Zombie ». Selon Cleeremans et Jiménez, les théories « Commandant Data » impliquent la totale transparence de la cognition vis-à-vis de la conscience. Chaque connaissance exprimée à travers nos comportements est accessible à l'introspection. Dans les théories « Commandant Data », tout traitement cognitif s'accompagne de conscience. Sauf dans de rares circonstances, nous serions conscients de notre propre fonctionnement. Les seuls traitements qui sont, par essence, inconscients, sont neuronaux et non cognitifs (Perruchet & Vinter, 2002). Un tel système rejette toute possibilité d'un traitement cognitif inaccessible à la conscience. Selon ces théories, si la nature complètement consciente de la cognition n'est pas mise en évidence par les mesures expérimentales, c'est en raison de la faiblesse méthodologique des tâches utilisées (Shanks & St John, 1994). A l'inverse, les tenants des théories « Zombie » considèrent que le conscient et l'inconscient sont enracinés dans deux systèmes distincts mais de complexité équivalente (e.g. Reber, 1989). L'inconscient serait structuré de la même manière que le conscient. Un « dédoublement de traitement » serait présent en chacun de nous : une part de la cognition serait opaque et inconsciente. Nos actes seraient guidés par un zombie qui existe en nous et dont nous ignorons l'existence. Les connaissances à l'origine de nos comportements ne seraient ni explicites ni implicites, puisque par définition notre zombie est dépourvu d'expérience subjective. Notre zombie est inconscient, et la prise de conscience, lorsqu'elle est présente, offre une vision incomplète et imparfaite des états internes.

Sans accorder aux traitements non conscients les mêmes capacités qu'aux traitements conscients, nous défendons néanmoins l'existence d'un inconscient cognitif. Même si les

traitements non conscients sont moins sophistiqués que ce l'on pense communément, nous pensons que cela n'enlève rien à leur caractère adaptatif et de ce fait, à leur persistance au cours de l'évolution phylogénétique. C'est dans ce cadre conceptuel que nous nous inscrivons. Les travaux que nous avons conduits visaient à explorer le rôle de l'apprentissage implicite dans la perception visuelle, et à déterminer la nature des traitements cognitifs non conscients. Dans cette perspective, nous avons étudié si des connaissances spécifiques ou catégorielles et sémantiques peuvent être acquises de manière implicite et dans quelle mesure des connaissances inaccessibles à la conscience facilitent l'exploration d'une scène visuelle. D'autre part, nous avons étudié si l'apprentissage de régularités de l'environnement est automatique, i.e. ne requiert pas une attention sélective. Enfin, nous avons exploré le développement d'un apprentissage de régularités contextuelles dans une tâche plus proche de celles impliquées dans une activité écologique, telle que la conduite automobile.

Pour étudier le développement d'un apprentissage de régularités de l'environnement, nous avons eu recours au paradigme d'indilage contextuel, développé il y a quelques années par Chun et Jiang (1998). Le paradigme d'indilage contextuel permet d'étudier expérimentalement des mécanismes d'apprentissage de régularités contextuelles impliqués au cours de l'analyse d'une scène visuelle et de tester la nature intentionnelle vs. non intentionnelle de cet apprentissage ainsi que la nature consciente vs. inconsciente de la connaissance résultante. Ce paradigme constitue un outil potentiellement pertinent pour étudier des mécanismes d'apprentissage implicite vs. explicite intervenant dans la perception visuelle. En effet, comme nous allons en faire état dans le Chapitre III, le paradigme d'indilage contextuel offre des arguments convaincants en faveur d'un apprentissage involontaire conduisant à l'acquisition de connaissances inaccessibles à la conscience.

CHAPITRE III

L'INDIÇAGE CONTEXTUEL

Les recherches utilisant le paradigme d'indication contextuel, ou « *contextual cueing* », ont connu ces dernières années un essor considérable, avec plus de quarante publications dans ces cinq dernières années. Le paradigme d'indication contextuel, développé par Chun et Jiang (1998), permet d'explorer le développement d'une sensibilité à des régularités contextuelles lors de l'analyse de scènes visuelles. Ce paradigme met classiquement en œuvre une tâche de recherche visuelle de cible, au cours de laquelle les participants sont invités à détecter le plus rapidement et correctement possible la cible présente. L'originalité de ce paradigme par rapport aux tâches classiques de recherche visuelle (Schneider & Shiffrin, 1977 ; Shiffrin & Schneider, 1977) est de comporter une série de blocs d'essais « prédictifs » et « non prédictifs ». Dans un essai prédictif, une régularité du contexte est prédictive d'une caractéristique spécifique de la cible (e.g. sa localisation, son identité). Dans un essai non prédictif, le contexte ne comporte pas de régularités informatives sur la cible. Une sensibilité aux régularités contextuelles se traduit par une détection progressivement plus rapide de la cible lorsqu'elle est présentée dans un contexte prédictif que dans un contexte non prédictif. Chun et Jiang ont dénommé ce bénéfice « *contextual cueing* » ou indication contextuel.

Dans l'expérience principes d'indication contextuel (Chun & Jiang 1998, Expérience 1), les sujets étaient invités à détecter le plus rapidement possible la cible « T » orientée vers la droite ou vers la gauche, parmi onze « L » distracteurs pouvant être présentés dans 4 orientations différentes (pour un exemple, cf. Figure 9a). La tâche de recherche comportait 30 blocs de 24 essais dont 12 essais « anciens » prédictifs, et 12 essais « nouveaux » non prédictifs. Les 12 essais anciens étaient systématiquement répétés au cours de chaque bloc

alors que les essais nouveaux étaient générés aléatoirement à chaque nouvel essai. Dans cette expérience, la configuration spatiale des éléments distracteurs définissait le contexte visuel de la cible. Autrement dit, dans un essai ancien, l'agencement spatial des éléments contextuels était prédictif de la localisation de la cible. Comme en atteste la Figure 9b, l'analyse des TR au fur et à mesure des blocs d'essais révèle que les sujets deviennent rapidement sensibles aux contextes répétés. La cible était plus rapidement détectée lorsqu'elle apparaissait dans un contexte ancien que dans un contexte nouveau, traduisant un effet d'indication contextuel. Une expérience contrôle (Chun & Jiang, 1998, Expérience 2) montre qu'un changement dans l'identité des distracteurs au cours de la tâche de recherche n'affecte pas les performances d'indication contextuel, indiquant que l'apprentissage repose bel et bien sur l'apprentissage des configurations spatiales des items, et non sur le développement d'une expertise spécifique relative aux traits de surface des items présents dans l'affichage.

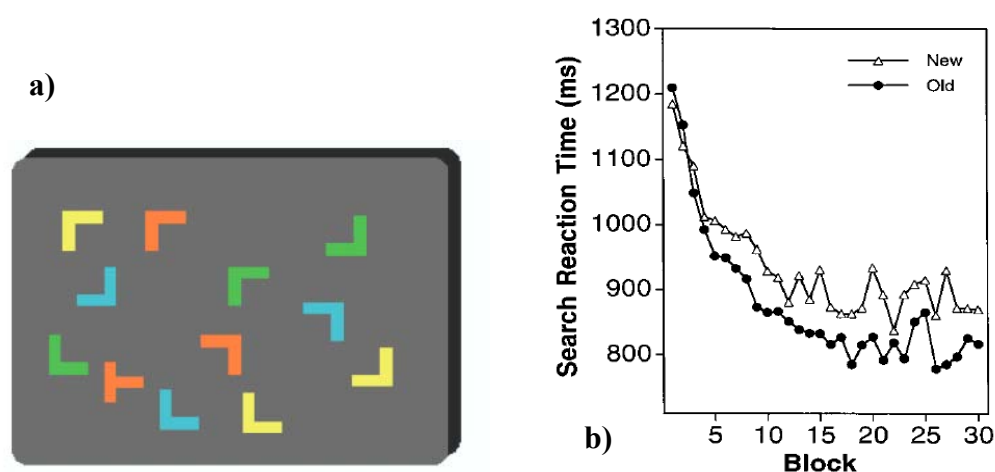


Figure 9. a) Exemple de configuration utilisée dans l'expérience principes de Chun et Jiang (1998). Les sujets devaient rechercher un T parmi des L. Dans un essai ancien (i.e. prédictif), l'agencement spatial des éléments de l'affichage était systématiquement répété au cours des blocs d'essais. b) Résultats observés dans la tâche de recherche visuelle. *Adapté de Chun et Jiang (1998)*

1. Arguments en faveur du caractère implicite de l'indication contextuel

Un aspect crucial des effets d'apprentissage mis en évidence avec la tâche classique d'indication contextuel est qu'ils sont le fruit d'un apprentissage implicite. Un certain nombre de résultats empiriques apportent en effet des arguments convaincants en faveur du caractère

implicite de l'apprentissage. Dans la section suivante, nous dresserons la liste de ces arguments.

1.1. Les tâches directes d'accès aux connaissances

A l'issue de la tâche princeps de recherche visuelle de cible (Chun & Jiang, 1998), les sujets étaient interrogés sur les régularités présentes dans le matériel. Aucun des participants n'a rapporté avoir remarqué que des contextes étaient systématiquement répétés au cours des blocs. Afin de tester le caractère implicite vs. explicite des connaissances acquises avec une tâche objective à la troisième personne, les participants étaient également soumis à une tâche de reconnaissance, au cours de laquelle ils étaient exposés à un nouveau bloc d'essais, construit comme ceux de la tâche précédente. Leur tâche était de juger la familiarité de chaque essai par rapport aux essais de la tâche de recherche. Les résultats de cette tâche révèlent que les participants n'étaient pas en mesure de discriminer les essais anciens (prédictifs) des essais nouveaux (non prédictifs), suggérant que l'apprentissage mis en évidence via l'effet d'indilage contextuel est de nature implicite. Ces résultats ont par la suite été reproduits dans de nombreuses études (e.g. Chun & Jiang, 1999, 2001, 2003 ; Howard et al., 2004 ; Olson & Chun, 2002). Donc contrairement aux paradigmes classiques d'apprentissage implicite (e.g. SRT, AGL), qui tendent souvent à montrer des effets différents du hasard dans les tâches de reconnaissance, les connaissances acquises sont ici inaccessibles par une tâche de reconnaissance.

D'autre part, afin de pallier certaines critiques concernant l'utilisation des jugements de familiarité comme mesure sensible de la mémoire explicite (Perruchet & Amorim, 1992 ; Shanks & St. John, 1994), Chun et Jiang (2003) ont développé une tâche directe d'accès aux connaissances explicites, plus proche de la tâche de recherche visuelle. L'objectif était d'engager des traitements les plus similaires possibles à ceux engagés dans la tâche de recherche. Dans ce test de mémoire explicite, les sujets devaient estimer la localisation de la cible dans des essais anciens où la cible avait été substituée par un distracteur. L'affichage était découpé en 4 quadrants et les sujets devaient indiquer lequel de ces quadrants était le plus susceptible de contenir la cible. La tâche comportait 12 configurations anciennes et 12 configurations nouvelles. Le pourcentage de réponses correctes observé par Chun et Jiang

(2003) était de 27% et ne différait pas significativement du hasard (25%). De manière plutôt probante, les participants dont les performances à la tâche directe de mémoire (i.e. génération indicée) était supérieures au hasard (donc susceptibles d'être conscients des régularités) manifestaient un effet d'indication contextuel inférieur aux participants non conscients des régularités. Ce test de « génération indicée » renforce la conclusion que la tâche classique d'indication contextuel met en œuvre des mécanismes d'apprentissage et de récupération implicites (cependant, Smyth & Shanks, sous presse).

Un troisième argument en faveur du caractère implicite de l'indication contextuel tient au caractère non intentionnel de l'apprentissage. En effet, les performances d'indication contextuel ne semblent pas être affectées par une instruction explicite incitant les sujets à rechercher activement les régularités (Chun & Jiang, 2003). Dans l'étude de Chun et Jiang (2003), lorsque l'expérimentateur informait les participants que certains essais étaient systématiquement répétés à chaque bloc, et qu'il les sollicitait à exploiter cette régularité pour faciliter leur recherche, aucun effet facilitateur n'était observé, ni sur les performances de recherche, ni sur les performances de « *guessing* ». Ces résultats montrent que les effets d'indication contextuel reposent sur des mécanismes d'apprentissage involontaire.

Les effets d'indication contextuel observés avec des configurations répétées semblent donc résulter d'un apprentissage implicite. Par ailleurs, dans la littérature sur l'apprentissage implicite, il est souvent stipulé que les connaissances acquises de manière implicite persistent davantage au cours du temps que celles acquises de manière explicite, qu'elles résistent aux effets d'interférence, à un manque de ressources attentionnelles, ainsi qu'à des désordres psychologiques tels que le syndrome amnésique. Qu'en est-il de l'indication contextuel ?

1.2. L'indication contextuel persiste au cours du temps et résiste aux effets d'interférence

A notre connaissance, deux études ont directement étudié la robustesse et la durabilité de l'apprentissage implicite mis en œuvre dans les tâches classiques d'indication contextuel. Dans leur étude portant principalement sur le caractère explicite vs. implicite de l'indication contextuel, Chun et Jiang (2003) ont également observé une persistance des effets d'indication après un délai d'une semaine entre deux sessions expérimentales, suggérant une rétention en

MLT. Jiang, Song et Rigas (2005) ont par la suite exploré plus largement la capacité de stockage et de rétention à long terme de la mémoire impliquée dans l'indication contextuel spatial, ainsi que les effets d'interférence. Les mêmes participants étaient exposés à cinq sessions d'entraînement sur cinq jours consécutifs et à une session test une semaine plus tard. Les cinq sessions d'entraînement étaient toutes similaires à la tâche princeps. Les résultats révèlent des effets d'indication contextuel comparables au cours de chacune des cinq sessions, suggérant que l'apprentissage du premier jour n'interfère pas, de manière proactive, sur les apprentissages ultérieurs. De plus, lorsqu'une semaine après la dernière session d'entraînement, les 60 configurations anciennes étaient présentées aux participants, non seulement les effets d'apprentissage observés le premier jour persistaient lors de la session test, mais en plus, ils étaient équivalents pour les configurations anciennes des cinq sessions d'entraînement. Cette étude montre que les contextes sont encodés en MLT et que l'indication contextuel est insensible aux interférences proactives et rétroactives. De surcroît, elle met en exergue les fortes capacités mnésiques à accumuler des configurations spatiales.

1.3. L'indication contextuel spatial requiert-il une attention visuelle sélective ?

A ce jour, le rôle de l'attention dans l'indication contextuel a exclusivement été exploré en considérant l'attention en termes de processus de sélection. Dans une étude de Jiang et Chun (2001), la moitié des participants devait chercher un T vert parmi des L distracteurs verts et des L distracteurs rouges. Les contextes verts définissaient les contextes attendus et les contextes rouges les contextes ignorés¹³. Pour l'autre moitié des participants, la cible à détecter était rouge et les contextes rouges définissaient les contextes attendus et les contextes verts les contextes ignorés. La tâche comportait 4 conditions expérimentales. Dans la condition « both-old », à la fois le contexte attendu et le contexte ignoré étaient répétés. Dans la condition « attended-old », le contexte attendu était répété mais le contexte ignoré était nouveau. Dans la condition « ignored-old », c'était l'inverse. Dans la condition « new », à la fois les contextes attendus et ignorés étaient nouveaux à chaque essai. Les résultats montrent que seuls les conditions « both-old » et « attended-old » conduisaient à des effets d'indication contextuel. Les performances observées dans la condition « ignored-old » n'étaient pas différentes de celles observées dans la condition « new ». Jiang et Chun (2001) en ont conclu

¹³ Afin de rendre la tâche plus difficile et d'induire une recherche visuelle sélective, les distracteurs étaient ici des L présentant de fortes ressemblances perceptives avec la cible T.

que l'apprentissage implicite des régularités contextuelles requiert une attention sélective (voir aussi, Kawahara, 2003).

Dans une étude postérieure, Jiang et Leung (2005) ont cependant montré que seule l'expression de l'apprentissage requiert une attention sélective. L'apprentissage en lui-même pourrait se développer de manière latente. Lorsque dans une phase de transfert les contextes au préalable ignorés devenaient subitement attendus (inversion dans les couleurs), ces derniers entraînaient immédiatement un effet d'indication. Inversement, lorsque les contextes au préalable attendus devenaient subitement ignorés, le bénéfice généré par ces contextes disparaissait immédiatement. Ces résultats montrent que si l'expression de la connaissance implicite requiert une attention sélective, l'apprentissage lui-même peut se développer sur des contextes extérieurs au focus attentionnel. Ces résultats sont concordants avec ceux rapportés par Frensch, Lin et Buchner dans une tâche de SRT (1998 ; voir aussi, Frensch, Wenke, & Rüniger, 1999). Cette étude constitue une source essentielle de nos travaux, puisque nous sommes inspirés de cette méthode pour étudier le rôle de l'attention sélective dans l'apprentissage de régularités sémantiques.

En utilisant le même plan expérimental que celui utilisé par Jiang et Leung (2005), Kim et Kim (2006) ont reproduit ce résultat et ont également montré que l'apprentissage des contextes ignorés concerne la relation spatiale entre la configuration des distracteurs et la cible attendue. La procédure était similaire à celle utilisée par Jiang et Leung, mais pendant la phase d'entraînement, l'affichage comportait un T dans la couleur attendue et un T dans la couleur ignorée. Durant un transfert, portant à la fois sur les distracteurs et sur les cibles (attendues et ignorées), aucun effet d'indication contextuel n'était observé. Un couplage temporel entre le contexte ignoré et la cible attendue serait donc nécessaire au développement d'un apprentissage latent.

1.4. L'indication contextuel et le syndrome amnésique

Les recherches dans le domaine de l'apprentissage implicite ont souvent montré que les capacités mnésiques des patients amnésiques étaient préservées dans des tâches d'apprentissage et de mémoire implicites. Pourtant, des patients amnésiques affectés par des lésions étendues du lobe temporal médian (i.e. hippocampe et régions voisines du lobe

temporal médian) manifestaient une incapacité à apprendre les configurations répétées (Chun & Phelps, 1999). Il semblerait toutefois que des lésions partielles de l'hippocampe ne soient pas suffisantes pour observer une dégradation de l'indiciage contextuel. Une lésion totale de l'hippocampe serait nécessaire (Manns & Squire, 2001). Des données en IRMf apportent des arguments supplémentaires en faveur de l'implication de l'hippocampe (Preston, Salidis, & Gabrieli, 2001). De surcroît, des participants auxquels avaient été administrés une dose relativement importante d'une benzodiazépine (i.e. le midazolam), connue pour provoquer un syndrome amnésique transitoire, ne manifestaient aucun effet d'indiciage contextuel (Park, Quinlan, Thornton, & Reder, 2004). Ces résultats s'opposent au point de vue selon lequel l'apprentissage implicite ne serait pas affecté par des états amnésiques. Mais ils ne remettent, à notre avis, pas en question le caractère implicite de l'apprentissage impliqué dans l'indiciage contextuel. Il semble évident que si les structures cérébrales impliquées dans un apprentissage particulier sont lésées ou anesthésiées, l'apprentissage sera dégradé, voire supprimé, que cet apprentissage soit explicite ou implicite.

Les recherches présentées dans les paragraphes précédents offrent des arguments en faveur de la nature implicite de l'apprentissage impliqué dans l'indiciage contextuel. Non seulement les tâches directes de mémoire apportent des arguments irréfutables, mais en plus, l'apprentissage impliqué persiste au cours du temps, résiste aux effets d'interférence proactives et rétroactives et semble émerger sans que l'attention sélective ne soit directement focalisée sur les régularités contextuelles. De surcroît, en plus d'être préservé du temps, l'apprentissage implicite des régularités contextuelles spatiales semble être préservé du vieillissement cognitif (Chun & Phelps, 1999 ; Howard et al., 2004). Contrairement aux performances observées avec une tâche de SRT, Howard et al. n'ont pas observé de détérioration avec l'âge. De plus, lorsque la difficulté de la tâche augmentait, les performances d'apprentissage restaient intactes chez les participants âgés, contrairement à celles observées avec une tâche de SRT. Pourtant, ces derniers trouvaient plus difficile la tâche d'indiciage contextuel que la tâche de SRT. Même si l'indiciage contextuel semble être fortement détérioré par des états amnésiques, le paradigme d'indiciage contextuel montre de manière incontestable que des contextes spatiaux répétés peuvent être appris implicitement au cours d'une recherche visuelle de cible et que ces connaissances implicites facilitent l'exploration de la scène. La question est alors de déterminer comment les régularités contextuelles facilitent la recherche visuelle.

2. Comment les régularités contextuelles facilitent-elles la recherche visuelle ?

2.1. Le contexte guiderait l'attention de manière efficace

L'amélioration des performances de recherche pourrait tenir à un apprentissage associatif entre le contexte et la localisation de la cible, mais elle pourrait aussi simplement tenir à une « facilitation du traitement perceptif », i.e. un meilleur encodage des items. Pour trancher entre ces deux hypothèses, Chun et Jiang (1998, Expérience 3) ont conduit une expérience au cours de laquelle la localisation de la cible variait au sein des contextes répétés. Un contexte répété n'était donc plus prédictif de la localisation de la cible. Aucun effet d'indication contextuel n'était alors observé, suggérant que l'indication contextuel repose sur un apprentissage associatif entre une configuration spécifique et une localisation de la cible.

En quoi l'apprentissage associatif contexte-localisation de la cible facilite-t-il la recherche ? L'hypothèse avancée par Chun et Jiang (1998) est qu'un contexte prédictif guide de manière efficace l'attention vers la localisation de la cible. Pour mettre à l'épreuve cette hypothèse, les auteurs ont manipulé le nombre de distracteurs définissant le contexte (8, 12 ou 16 distracteurs). La tâche classique d'indication contextuel mettant en œuvre une recherche sérielle, les temps de détection de la cible augmentaient proportionnellement avec le nombre de distracteurs présents dans l'aire de recherche. Cette expérience apportait trois résultats majeurs. En premier lieu, les intersections des droites relatives aux « TR x nombre de distracteurs » avec l'axe des ordonnées, n'étaient pas différentes et n'évoluaient pas différemment dans les conditions anciennes et nouvelles, suggérant que l'indication contextuel ne facilite pas la recherche en accélérant le traitement perceptif précoce, ni les traitements décisionnels impliqués dans la sélection de la réponse, ni encore les traitements impliqués dans l'exécution de cette réponse. En second lieu, l'analyse des pentes montre qu'au fur et à mesure de la tâche, la taille de la série (« *set size* », i.e. le nombre de distracteurs présents dans l'affichage) affectait davantage les TR dans la condition nouvelle que dans la condition ancienne. Et enfin, les effets d'indication contextuel augmentaient lorsque la taille de la série augmentait. Ces résultats offrent des arguments solides en faveur de l'hypothèse selon

laquelle l'indication contextuelle facilite la recherche en guidant le déploiement de l'attention vers la localisation de la cible.

Selon Chun et Jiang (1998), l'indication contextuelle résulterait d'une interaction entre mémoire et attention. La récupération d'une trace mnésique relative à l'association « contexte - localisation de la cible », acquise de manière implicite, orienterait l'attention sélective vers la localisation pertinente pour la tâche. D'autres données empiriques confortent cette hypothèse. Par exemple, les effets d'indication contextuelle sont d'autant plus importants que la tâche est rendue plus difficile, en augmentant la similarité perceptive entre la cible et les distracteurs par exemple (e.g. Chun & Phelps, 1999). Mais aussi, des études d'enregistrement des mouvements des yeux alors que les participants réalisaient la tâche classique corroborent cette hypothèse (Peterson & Kramer, 2001). Peterson et Kramer ont montré que les participants effectuent moins de saccades oculaires pour détecter la cible dans les essais anciens que nouveaux. De plus, lorsque des masques cachaient les éléments de l'affichage au début de l'essai (les localisations des items étaient donc indicées), les sujets tendaient davantage à fixer la cible dès la première saccade oculaire pour les configurations anciennes que sur les configurations nouvelles (pour un exemple, cf. Figure 10).



Figure 10. Exemple de configuration utilisée dans la recherche de Peterson et Kramer (2001). A gauche : affichage comportant des masques indiquant les localisations des items (cible et distracteurs). A droite, affichage de recherche. *Adapté de Peterson et Kramer (2001)*

2.2. Cependant...

Cependant, les travaux récents de Kunar, Flusberg, Horowitz et Wolfe (2007) appellent à une relecture des effets classiques d'indication contextuelle. Après de nombreuses tentatives, les

auteurs ne sont pas parvenus à reproduire le résultat rapporté par Chun et Jiang (1998) montrant une diminution plus marquée de la pente (TR x nombre de distracteurs) pour les essais anciens que pour des essais nouveaux. Or, cet argument était de loin le plus convaincant en faveur d'une interaction entre mémoire et attention. Si les contextes répétés ne facilitent pas le guidage de l'attention, à quoi pourraient alors tenir les effets d'indication contextuel et comment expliquer les résultats relatifs aux mouvements oculaires ? Nous avons évoqué quelques lignes auparavant l'hypothèse selon laquelle les effets d'indication contextuel tiennent au développement d'une facilitation lors de l'encodage perceptif des éléments contextuels. Les résultats sur les mouvements oculaires ne sont pas complètement incompatibles avec cette hypothèse. Les participants pourraient en effet encoder les éléments contextuels sous forme de chunks. Cet encodage « holistique » se traduirait lui aussi par une diminution du nombre de fixations (Didierjean, 2004 ; Reingold, Charness, Pomplun, & Stampe, 2001). De plus, les données empiriques tendent à montrer que les participants apprennent davantage les localisations individuelles que les configurations globales (Brady & Chun, 2007 ; Olson & Chun, 2002). Il semblerait que le contexte local de la cible soit suffisant pour observer des effets d'indication et que la force de ces effets ne soit guère plus faible que dans des affichages où l'ensemble des éléments contextuels sont répétés. Il reste donc possible que l'indication repose exclusivement sur la reconnaissance d'un chunk, formé par la cible et les distracteurs adjacents. Mais il n'est pas non plus exclu que la facilitation observée dans les essais répétés tienne à une réduction dans le temps de traitement de la cible elle-même. En faveur de cette hypothèse, Kunar, Flusberg et Wolfe (2006) ont montré que si la tâche du sujet est de décider si la cible est présente ou absente de l'affichage, aucun effet facilitateur n'est observé dans les essais répétés négatifs (i.e. lorsque la cible est absente). Mais surtout, lorsqu'une manipulation expérimentale interférait avec la sélection de la réponse, l'effet d'indication disparaissait (Kunar et al., 2007). Le deuxième argument avancé par Peterson et Kramer reste toutefois convaincant : lorsque des masques amorçaient la configuration globale, les participants tendaient à fixer la cible dès la première saccade oculaire pour les configurations anciennes. Mais surtout, si l'effet sur les pentes n'est pas significatif dans la tâche classique d'indication contextuel (i.e. 12 distracteurs), Kunar et al. (2006) ont rapporté des effets significatifs dans des versions différentes impliquant une recherche plus difficile et donc des TR plus longs (e.g. en augmentant la taille de la série ou la similarité entre la cible et les distracteurs). Dans l'ensemble, les données de la littérature permettent d'exclure la seule hypothèse d'une facilitation dans l'encodage perceptif ou dans la sélection de la réponse.

Un phénomène intéressant a par ailleurs été mis en évidence. Tous les individus ne manifestent pas des effets d'indication contextuel. Les recherches relatant ce phénomène indiquent que généralement une partie des sujets (environ 20%) ne manifestent pas de bénéfice dans la condition ancienne. Lleras et Von Muhlenen (2004) ont même rencontré des difficultés à mettre en évidence un effet d'indication contextuel spatial à partir de la tâche classique. Pour rendre compte de ce manque de significativité, les auteurs ont proposé que les stratégies de recherche mises en œuvre par les participants puissent influencer les performances d'indication contextuel. Pour tester cette hypothèse, les auteurs ont manipulé deux types de consignes, l'une encourageant les participants à élaborer une stratégie de recherche « active » de la cible, et l'autre les encourageant à rechercher de manière « passive » la cible. La stratégie de recherche active solliciterait davantage un déploiement rapide de l'attention item par item, alors que la stratégie de recherche passive favoriserait la mise en œuvre de traitement pop-out. Aucun effet d'indication contextuel n'était observé dans la condition active, alors que 80% des sujets manifestaient un bénéfice dans la condition passive. Ces résultats nous amènent à proposer que différents facteurs puissent être responsables des effets d'indication contextuel et que l'influence de ces facteurs puisse dépendre de la stratégie de recherche adoptée par le sujet.

Les influences multiples susceptibles de rendre compte de l'effet facilitateur des contextes spatiaux répétés pourraient expliquer pourquoi Jiang, Kim, Shim et Vickery (2006) n'ont pu dégager de corrélation entre les performances individuelles d'indication contextuel avec d'autres capacités cognitives. Des études de neuropsychologie (Chun & Phelps, 1999) et de neuroimagerie (Preston et al., 2001) suggérant une implication de l'hippocampe et des régions médio-temporales voisines dans l'indication contextuel spatial, les auteurs faisaient l'hypothèse que les tests impliquant la mémoire épisodique et spatiale constituaient des candidats susceptibles de rendre compte des capacités cognitives impliquées dans l'indication contextuel. Ils ont ainsi testé si les performances individuelles d'indication contextuel étaient corrélées avec celles obtenues dans des tests visant à mesurer la mémoire épisodique des sujets, leur capacité spatiale, leur mémoire visuelle de travail, leur capacité de navigation spatiale et la capacité de leur mémoire autobiographique. Aucune corrélation entre les performances d'indication contextuel et les performances obtenues dans ces différents tests n'était observée. De manière convergente, après avoir fait passer des tâches d'indication contextuel à plusieurs reprises aux mêmes participants, Jiang et al. (2005) n'ont pas observé de corrélation entre les performances obtenues au cours des différentes sessions. Les sujets manifestant un fort effet

d'indilage contextuel lors d'une session pouvaient manifester des performances médiocres dans la session suivante.

L'ensemble de ces données nous amènent à penser que les effets facilitateurs des contextes spatiaux répétés pourraient dériver d'influences multiples. Toutefois, l'hypothèse selon laquelle les contextes répétés facilitent la recherche en optimisant le déploiement de l'attention vers la localisation de la cible reste la plus admise, et certainement la plus probable. Si cette question nécessite encore d'être explorée de manière directe, des travaux ne testant pas directement cette question peuvent néanmoins apporter des éléments de réflexion. Dans la section suivante, nous rapporterons un ensemble de données démontrant que les effets d'indilage contextuel peuvent être généralisés à d'autres régularités que la répétition des configurations spatiales.

3. Portée des effets d'indilage contextuel dans des environnements arbitraires

Les travaux décrits dans les sections précédentes révèlent que le système visuo-cognitif est pourvu d'une étonnante capacité à mémoriser de manière implicite des contextes spatiaux répétés. Quelle est la portée et la généralisation des effets d'indilage contextuel ? Peut-on observer des effets de transfert à de nouveaux exemplaires ? Plus généralement, quelles sont les capacités du système visuo-cognitif à apprendre des régularités contextuelles et à quel type de régularités ce dernier devient-il sensible ? Si des effets d'indilage peuvent reposer sur d'autres régularités contextuelles que des régularités spatiales, l'apprentissage impliqué est-il toujours implicite ? Dans les travaux abordant ces questions, les contextes étaient soit définis par des plages de stimuli dépourvues de toute signification (e.g. des L), soit définis par des scènes du monde réel. Nous traiterons séparément ces deux types de recherches. Concernant, les travaux conduits avec des contextes « arbitraires », ces derniers ont essentiellement porté sur trois aspects du contexte : l'agencement spatial des éléments contextuels, l'identité des éléments contextuels, le déroulement temporel des indices contextuels.

3.1. L'indiciage contextuel spatial

Dans leur recherche princeps, Chun et Jiang (1998) ont montré que l'indiciage contextuel basé sur la configuration spatiale des contextes pouvait aussi bien émerger sur des items polychromes (Chun & Jiang, 1998, Expérience 1) que sur des items monochromatiques (Chun & Jiang, 1998 ; Expériences 2-6). De tels effets ont été reproduits dans un nombre considérable de recherches, attestant de leur robustesse. Des effets d'indiciage spatial ont été observés avec bien d'autres types de stimuli, tels que des formes « sans signification évidente » (Endo & Takeda, 2004), mais également dans des affichages en pseudo trois dimensions, en utilisant des indices picturaux donnant une impression de profondeur (Chua & Chun, 2003 ; Kawahara, 2003). L'indiciage contextuel spatial peut aussi reposer sur des configurations de couleurs, lorsque par exemple l'arrangement spatial de tâches colorées sur une matrice est prédictif de la localisation de la cible (Huang, 2006). Il semblerait d'ailleurs que des relations prédictives basées sur la teinte des éléments contextuels soient plus favorables au développement d'un effet d'indiciage que des relations basées sur la luminance des objets.

De nombreuses données empiriques suggèrent par ailleurs que les contextes réguliers induisant un effet d'indiciage contextuel sont relativement flexibles. Comme nous l'avons évoqué précédemment, un changement dans la nature des distracteurs en cours d'apprentissage n'empêche pas son développement et ne réduit pas l'indiciage contextuel, à condition que soit préservé l'arrangement spatial des éléments contextuels (Chun & Jiang, 1998). Ou encore, l'introduction de petites variations dans la position des éléments distracteurs dans les contextes anciens ne perturbe pas les effets d'indiciage contextuel (Chun & Jiang, 1998). D'autres expériences ont montré une tolérance au bruit, indiquant que l'intégralité de l'information contextuelle n'est pas nécessaire au développement d'un apprentissage (Olson & Chun, 2002). De plus, les régularités apprises par les participants dans l'indiciage contextuel spatial semblent pouvoir reposer sur les configurations globales mais également sur les localisations individuelles. En dégradant au cours d'une session de transfert, soit la configuration globale formée par l'ensemble des distracteurs, soit leur localisation individuelle, Jiang et Wagner (2004) ont montré que les participants peuvent extraire le pattern formé par l'ensemble des items et/ou les localisations individuelles de ces éléments¹⁴.

¹⁴ Dans l'Expérience 1, les essais anciens étaient recombinaisonnés de sorte que la moitié des distracteurs préservent leur emplacement et l'autre moitié occupent des emplacements nouveaux, procédure remodelant sensiblement la

D'autres études confirment que l'information globale du contexte n'est pas nécessaire au développement d'un effet d'indication contextuel. Par exemple, Song et Jiang (2005) ont montré qu'un contexte invariant composé de seulement trois des distracteurs adjacents à la cible serait suffisant, indiquant que le contexte spatial local de la cible peut à lui seul faciliter la recherche. Comme nous l'avons évoqué précédemment, une étude de Brady et Chun (2007) suggère même que le contexte local, davantage que le contexte global, facilite la recherche visuelle. Enfin, Ono, Jiang et Kawahara (2005) ont montré des effets d'indication contextuel inter-essais. Lorsque la configuration spatiale de l'essai N-1 (agencement spatial des distracteurs et de la cible) prédisait la localisation de la cible à l'essai N, un effet d'apprentissage était observé.

L'individu serait-il une « éponge », capable d'absorber et d'exploiter tout type de régularité spatiale ? Il semblerait que certaines régularités spatiales ne soient pas suffisantes pour produire des effets. Par exemple, Olson et Chun (2002) ont montré que si un contexte invariant éloigné de la cible peut à lui seul indiquer sa localisation, aucun effet d'indication n'émerge si un contexte variable interfère géométriquement entre le contexte invariant et la cible. Il semblerait de surcroît que la position relative des contextes invariants sur l'écran soit cruciale pour l'indication contextuel (Endo & Takeda, 2005). De plus, si les contextes répétés ne sont pas prédictifs de la localisation de la cible dès le début de l'entraînement, aucun bénéfice ne se développe lorsque ces derniers deviennent tardivement corrélés avec une localisation de la cible (Jungé, Scholl, & Chun, 2007). L'absence de régularité prédictive dans les étapes précoces de la tâche bloquerait l'apprentissage ultérieur. Des recherches ont également révélé une « hyper-spécificité » de l'indication contextuel. Une étude de Jiang et Song (2005a) montre ainsi que l'apprentissage de l'agencement spatial est « contingent » à l'identité des items. Conformément aux résultats obtenus par Chun et Jiang (1998), l'apprentissage des configurations spatiales définies par des distracteurs d'une forme particulière (identité 1) pouvait être transféré à des configurations identiques définies par des distracteurs de forme nouvelle (identité 2), même si la difficulté de la tâche était considérablement différente entre les conditions d'entraînement et de test. Mais lorsque la phase d'entraînement comportait à la fois des affichages composés de distracteurs d'identité 1

configuration globale du contexte (voir aussi Chun & Jiang, 1998 pour une procédure et un résultat similaires). Dans l'Expérience 2, la taille de l'affichage était modifiée et son emplacement sur l'écran était déplacé, de sorte que les items individuels n'occupent plus la même localisation sur l'écran, mais que les configurations globales demeurent inchangées. Les résultats révèlent des effets d'indication contextuel dans les deux conditions de transfert, suggérant que, non seulement la configuration spatiale est suffisante pour indiquer la localisation de la cible, mais également la localisation des items individuels.

et des affichages composés de distracteurs d'identité 2, l'apprentissage des configurations d'entraînement disparaissait lorsque l'identité de leurs distracteurs était inversée. Ce résultat est d'autant plus surprenant que les deux séries de distracteurs partageaient de fortes similitudes perceptives. Des résultats similaires ont été observés lorsque les deux séries de distracteurs différaient par leur couleur. En faveur de cette « spécificité de transfert », Chua et Chun (2003) ont montré que l'indication contextuel spatial est dépendant du point de vue selon lequel la scène est perçue par le sujet. Lorsque le point de vue selon lequel étaient présentées des scènes anciennes en trois dimensions était modifié par rapport aux essais d'apprentissage (i.e. rotation de 0°, 15°, 30° ou 45°), l'effet d'indication diminuait proportionnellement à la rotation effectuée sur la scène. De plus, les travaux de Jiang et Song (2005b) suggèrent que si l'apprentissage des contextes spatiaux se met en place à travers de multiples tâches (tâche de recherche, tâche de détection de changements), l'exploitation des contextes appris est en partie spécifique à la tâche. En combinant une tâche classique d'indication contextuel avec une tâche de détection de changements, les auteurs ont montré qu'une configuration spatiale apprise durant la recherche visuelle ne facilitait pas la détection de changements au sein d'une configuration identique, et inversement, un contexte appris durant une tâche de détection de changements ne facilitait que modérément la recherche visuelle. Les contextes appris durant une tâche semblent ainsi peu transférables à une autre tâche. Par ailleurs, tous les indices spatiaux ne conduisent pas à des effets d'indication contextuel. Par exemple, Endo et Takeda (2004) n'ont pas observé d'effet d'apprentissage lorsque l'agencement spatial prédisait l'identité de la cible et non sa localisation. Enfin, si la tâche est « trop facile », aucun bénéfice n'est observé. Les effets sont d'autant plus importants que la cible et les distracteurs sont difficiles à discriminer (e.g. Chun & Phelps, 1999).

3.2. Indication contextuel basé sur l'identité spécifique des éléments contextuels

Si la plupart des recherches réalisées ont porté sur des régularités spatiales, des effets d'indication contextuel ont également été mis en évidence lorsque les régularités contextuelles reposent sur l'identité des distracteurs, indépendamment de leur agencement spatial. Dans une expérience de Chun et Jiang (1999), les participants avaient pour consigne de rechercher parmi des formes sans signification une cible définie comme la seule forme symétrique par rapport à un axe vertical (pour un exemple, cf. Figure 11). Les régularités tenaient ici à des co-variations entre les formes présentes dans l'affichage et la forme de la cible, l'arrangement

spatial des éléments de l’affichage (i.e. cible et distracteurs) étant aléatoire à chaque essai. Un bénéfice était obtenu dans la condition constante (prédictive), suggérant que les formes des éléments contextuels peuvent indiquer la forme de la cible avec laquelle ces dernières co-varient. Il semblerait par ailleurs que l’identité d’un contexte local systématiquement associé à la cible puisse à lui seul en faciliter sa détection, indépendamment de la localisation des items dans l’affichage (Hoffman & Sebold, 2005).



Figure 11. Exemple de stimulus utilisé par Chun et Jiang (1999). Les participants devaient rechercher la seule cible symétrique par rapport à un axe vertical. Dans les essais prédictifs, la forme des distracteurs co-variait avec la forme de la cible.

Endo et Takeda (2004) ont par la suite reproduit et étendu ce résultat en montrant que « l’identité » des éléments contextuels pouvait également indiquer la localisation de la cible. La procédure utilisée était néanmoins différente. Plutôt que d’étudier de manière classique le développement progressif d’un apprentissage, Endo et Takeda ont testé si l’apprentissage pouvait être transféré à des conditions dans lesquelles seule l’identité des éléments ou leur localisation était préservée. Par exemple, au cours de l’apprentissage, à la fois la configuration spatiale et l’identité des objets étaient répétées dans les essais constants. La configuration et l’identité du contexte prédisaient donc de manière confondue, la localisation et l’identité de la cible. Au cours de la phase de transfert, seules les configurations invariantes ou bien seules les identités des distracteurs étaient préservées. Dans ces deux conditions, l’identité et la localisation de la cible étaient maintenues dans les essais constants. Les résultats indiquent que l’apprentissage était intégralement transféré lorsque seule la configuration spatiale du contexte était préservée (résultat classique), mais pas lorsque seule l’identité du contexte était préservée. Les sujets n’apprendraient donc pas l’identité des distracteurs lorsqu’à la fois l’identité et les configurations sont répétées au cours de l’apprentissage. Cependant, lorsque seule l’identité des distracteurs était prédictive de l’identité et de la localisation de la cible au cours de l’entraînement, des effets d’indigence basés sur l’identité du contexte étaient observés.

Les auteurs se sont alors demandés si les contextes invariants indiquaient l'identité et/ou la localisation de la cible. Un effet d'apprentissage était observé lorsque l'identité du contexte prédisait la localisation de la cible, mais pas lorsque la configuration du contexte prédisait l'identité de la cible.

L'indication contextuelle peut ainsi reposer sur des régularités relatives à l'identité ou la forme des éléments distracteurs, indépendamment de leur agencement spatial, que ces derniers indiquent la localisation de la cible ou son identité. Kunar, Flusberg et Wolfe (2006a) ont par ailleurs montré que la couleur du fond de l'affichage ou sa texture prédictive de la localisation de la cible pouvait faciliter sa détection, que celui-ci soit présenté simultanément avec l'affichage de recherche ou en amorçage.

3.3. Indication contextuelle temporelle

L'indication contextuelle peut également être basée sur des régularités temporelles, par exemple, à travers des environnements dynamiques composés d'items en mouvement selon des trajectoires prédictibles, i.e. une cible T parmi des distracteurs L (Chun & Jiang, 1999), ou encore, lorsque des événements visuels se déroulent dans un ordre temporel prédictible (Olson & Chun, 2001). Afin d'étudier si une structure contextuelle temporelle peut guider l'attention dans l'espace et dans le temps, Olson et Chun ont utilisé une procédure hybride dérivée de celles utilisées dans les tâches d'indication contextuelle et de SRT (Nissen & Bullemer, 1987). Le sujet devait rechercher une lettre cible potentielle (i.e. K vs. X) parmi une séquence de lettres présentées successivement de manière isolée (pour une illustration de la procédure, cf. Figure 12). La durée de présentation de ces lettres était variable, conférant un rythme visuel à la séquence. Ce rythme visuel, autrement dit, la durée de présentation des lettres, était répété de bloc en bloc, mais l'identité des lettres contextuelles était aléatoire. L'occurrence des deux cibles potentielles était quant à elle fixée dans la séquence. La durée d'apparition des lettres précédant la cible constituait donc ici la régularité prédictive du moment d'occurrence de la cible, indépendamment de son identité et de celle des lettres contextuelles. Durant une phase de transfert, lorsque cette séquence était dégradée, les TR augmentaient significativement, suggérant que la structure contextuelle temporelle invariante facilite la détection de la cible. Des effets d'indication contextuelle ont également été observés lorsque l'identité des lettres précédant la cible prédisait l'emplacement séquentiel de la cible, indépendamment de leur

durée d'apparition. Mais aussi, des séquences d'événements invariants définis à la fois dans l'espace et dans le temps pouvaient guider l'attention vers une localisation particulière de l'affichage.

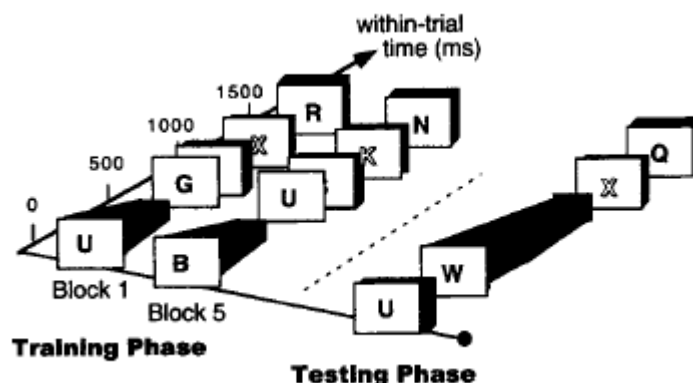


Figure 12. Illustration de la procédure utilisée par Olson et Chun (2001, Expérience 1). A noter que la durée de présentation des lettres contextuelles est constante au cours des répétitions, mais que les lettres sont variables. *Adapté de Olson et Chun (2001)*

Les travaux décrits montrent que les effets d'indication contextuel ne se limitent pas à des contextes spatiaux répétés dans leur intégralité, dénotant une certaine flexibilité de l'indication contextuel spatial. Ces derniers peuvent aussi être basés sur des régularités relatives à la forme des éléments contextuels ou à leur déroulement dynamique ou temporel. Un aspect essentiel des apprentissages mis en évidence dans des affichages composés de stimuli arbitraires est qu'ils sont systématiquement de nature implicite (cf. Chun & Jiang, 1999 ; Olson & Chun, 2001). Non seulement les participants sont incapables de verbaliser les régularités prédictives, mais surtout, leurs performances dans la tâche de reconnaissance ne sont jamais différentes de ce que le hasard prédit.

4. L'indication contextuel dans des scènes du monde réel

4.1. L'indication contextuel dans des scènes systématiquement répétées

L'indication contextuel a récemment été exploré à partir de scènes du monde réel. Dans ce cadre, les travaux de Brockmole et Henderson (2006a) montrent comment la recherche d'une

cible arbitraire peut être rapidement facilitée lorsqu'elle apparaît dans une scène du monde réel systématiquement répétée au cours de la tâche. Au cours de cette étude, les sujets étaient invités à rechercher une lettre cible arbitraire (un T ou un L) au sein de photographies de scènes du monde réel. La moitié des scènes était systématiquement répétée au cours de chaque bloc alors que l'autre moitié n'était présentée qu'une seule fois au cours de la tâche. Le principe était donc identique à celui de la tâche princeps (Chun & Jiang, 1998), mais le contexte de la cible était ici défini par une photographie du monde réel. Les résultats sont toutefois très différents de ceux observés avec des contextes arbitraires (pour les résultats relatifs aux études de Brockmole et Henderson (2006a) et de Chun et Jiang (1998), cf. Figure 13). D'une part, les effets d'indilage contextuel émergeaient très rapidement au cours de la tâche (5 fois plus rapidement que dans des contextes composés de « L ») et le bénéfice de recherche était 20 fois plus important que celui observé avec des affichages arbitraires. L'analyse des mouvements oculaires au cours de la recherche suggère que la répétition des essais facilite la recherche en guidant l'attention directement vers la cible : le nombre de fixations oculaires était parfaitement corrélé avec les données sur les TR. Mais surtout, les effets d'apprentissage étaient explicites. Non seulement les participants étaient à même de verbaliser l'existence de régularités relatives à la localisation de la cible dans des scènes répétées, mais en plus leurs performances dans la tâche de reconnaissance étaient très largement supérieures au hasard. Les participants manifestaient des performances de reconnaissance excellentes (97% de réponses « familier » pour des scènes anciennes contre 9% pour des scènes jamais vues). De surcroît, une tâche de génération indicée montre qu'ils étaient capables de positionner la cible de manière relativement précise au sein des scènes répétées (les erreurs de localisation étaient en moyenne de 1.7cm pour les scènes anciennes répétées contre 9.3cm pour les scènes vues une seule fois au cours de la tâche).

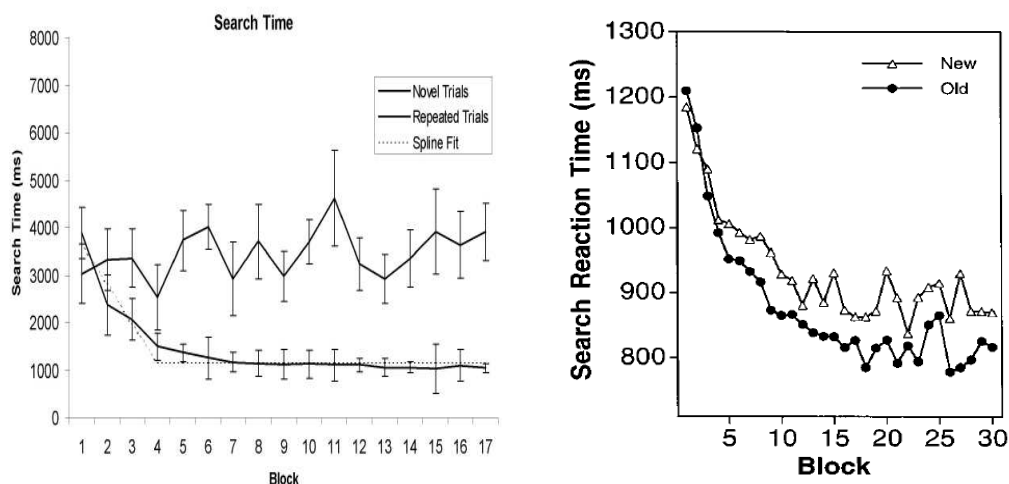


Figure 13. A gauche : Résultats observés par Brockmole et Henderson (2006a) à partir de scènes du monde réel. A droite : Résultats observés dans l'expérience princeps de Chun et Jiang (1998) à partir de configurations spatiales. *Adapté de Brockmole et Henderson (2006a) et de Chun et Jiang (1998)*

4.2. Pourquoi les performances sont-elles différentes dans les scènes du monde réel ?

Quelles sont les raisons de ces différences sensibles entre les tâches de recherche utilisant des contextes artificiels et des contextes « plus naturels » ? Selon Brockmole et Henderson (2006a), en favorisant la prise de conscience des régularités contextuelles, les indices sémantiques présents dans les photographies du monde réel faciliteraient leur encodage et leur récupération. En faveur de cette hypothèse, lorsque les scènes étaient inversées par rapport à un axe horizontal, le nombre moyen de répétitions nécessaires pour observer les mêmes performances d'indigage était environ deux fois plus élevé. L'inversion des scènes rendant plus difficile l'accès à leur sémantique, ce résultat offre un argument en faveur d'une hypothèse selon laquelle les effets précoces d'indigage contextuel observés avec des scènes du monde réel sont au moins en partie basés sur les caractéristiques sémantiques de la scène.

Brockmole et Henderson (2006b) ont testé autrement cette hypothèse en explorant la nature de l'information scénique utilisée pour guider l'attention vers les localisations connues de la cible. Pour ce faire, ils ont testé si les scènes étaient plus précocement représentées en termes de leur identité/signification générale ou en termes de l'arrangement spécifique de leurs propriétés visuelles. La tâche était la même que celle utilisée par Brockmole et Henderson (2006a), mais une phase de transfert était implémentée. Au cours du transfert, les scènes apprises étaient inversées selon un axe vertical. Ainsi, l'identité de la scène était préservée mais l'organisation spatiale des objets était changée, décalant la position de la cible

par rapport au cadre égocentrique de l'observateur. Lors de l'inversion en miroir d'une scène apprise, les yeux des observateurs se portaient initialement sur la position ancienne de la cible, comme si l'inversion en miroir de la scène n'avait pas été détectée par les participants. Cependant, les yeux se déplaçaient rapidement vers la nouvelle localisation de la cible. L'attention serait donc en premier lieu guidée par la reconnaissance de l'identité globale de la scène, sans référence à l'arrangement spatial spécifique des traits visuels. Mais rapidement, la reconnaissance d'autres informations visuelles locales telles que les traits visuels ou l'identité des objets individuels réorienterait l'attention au sein de l'image.

De manière peut-être plus convaincante, Brockmole, Castelhamo et Henderson (2006) ont montré que l'indication contextuelle serait davantage sous-tendue par le contexte global que par le contexte local de la cible. Lorsque le contexte local, défini par les objets adjacents de la cible était altéré durant une phase de transfert, mais que le contexte global de la scène était préservé, les effets d'indication contextuelle demeuraient intacts. En revanche, lorsque le contexte global était altéré mais que le contexte local était maintenu constant, le bénéfice d'indication contextuelle était totalement éliminé. Cependant, lorsque seul le contexte local ou seul le contexte global était répété au cours de l'apprentissage, des effets d'indication contextuelle étaient observés dans les deux conditions expérimentales. Le bénéfice observé était néanmoins beaucoup plus précoce et plus important lorsque le contexte global était répété. Ces résultats sont contraires aux résultats observés avec des plages de stimuli révélant que les cibles sont plus largement associées au contexte local qu'au contexte global (Brady & Chun, 2007 ; Jiang & Wagner, 2004 ; Olson & Chun, 2002).

5. L'indication contextuelle reflèterait une automatisation de la recherche visuelle

Chun et Jiang (1998) ont utilisé le terme de « carte contextuelle » pour définir les représentations en mémoire relatives au contexte visuel. Ces représentations comporteraient les informations visuelles invariantes dans une image ou une scène visuelle, auxquelles les sujets manifestent une sensibilité lors d'une recherche visuelle d'une cible. Si les observateurs ne sont pas sensibles à toute forme de régularités (e.g. Chua & Chun, 2003 ; Olson & Chun,

2002), la capacité à accumuler, et à utiliser des informations visuelles rencontrées semble cependant extrêmement puissante. Le système visuo-cognitif semble pouvoir encoder une grande quantité d'informations éphémères, les stocker, et les mettre en relation pour n'en récupérer que les aspects pertinents dans une situation donnée. Dans leur article princeps, Chun et Jiang (1998) ont formalisé le phénomène d'indilage contextuel dans le cadre des théories de l'automatisation basée sur la mémoire (Logan, 1988, 2002, 2004).

Dans *l'Instance Theory* (IT), Logan (1988) propose que l'automatisation des performances dans un domaine donné provienne de la récupération directe de solutions relatives aux problèmes ou situations traités antérieurement : une conduite est automatique lorsqu'elle est basée sur la récupération directe de solutions passées stockées en mémoire. Le cas présenté permettrait la récupération d'un exemplaire identique stocké en mémoire, et activerait la solution. La théorie présuppose que les novices commencent avec un algorithme général, suffisant pour la réalisation de la tâche. Après une exposition répétée à un même problème, ils apprendraient une solution spécifique à ce problème spécifique, qu'ils récupérerait de manière automatique lorsqu'ils sont de nouveau confrontés au même problème. Quand l'individu dispose en mémoire de la solution relative au problème, il aurait la possibilité de répondre en récupérant directement en mémoire la solution ou de répondre par un « calcul », à l'aide de l'algorithme adéquat. Ces deux types de processus entreraient en compétition au cours de la tâche, et le plus efficace « gagnerait la course » et « contrôlerait la réponse ». Dans la théorie de Logan, le poids de l'algorithme est inchangé au cours de l'expérience, alors que le poids du traitement mnésique augmente avec l'accumulation de cas identiques. Plus le nombre d'instances en mémoire augmente, plus la probabilité de récupérer en mémoire la solution d'un « exemplaire » augmente. Avec la pratique, l'individu abandonnerait complètement l'algorithme au profit de la récupération en mémoire. La mémoire finirait par dominer l'algorithme, jusqu'à le supplanter complètement. Selon Logan, la conduite est alors automatique. L'automatisation reflèterait un transfert d'une « conduite basée sur un algorithme » à une « conduite basée sur la mémoire ».

Dans une tâche de recherche visuelle de cible, les essais définissent des instances spécifiques et la localisation de la cible définit la solution au problème. Dans les étapes précoces de l'entraînement, la détection de la cible serait guidée par des *mécanismes attentionnels basés sur l'utilisation d'algorithmiques généraux* (Chun & Jiang, 1998). Ce traitement algorithmique correspondrait à la recherche de la cible, item par item. Mais au fur

et à mesure de la tâche et des répétitions, les traces mnésiques relatives à l'association « contexte-localisation de la cible » s'accumulent et la probabilité que l'une d'entre elles soit activée augmente. Dans les essais répétés, la recherche basée sur le recours à un algorithme serait ainsi progressivement relayée par une recherche basée sur la récupération de solutions en mémoire.

Dans sa théorie initiale, Logan (1988) décrit un modèle de « *pure instance* » où seuls les exemplaires strictement identiques au cas présenté peuvent être activés en mémoire et entrer en compétition avec l'algorithme. Il est vrai que la littérature suggère que l'indilage contextuel spatial est plutôt spécifique (e.g. Chua & Chun, 1998 ; Jiang & Song, 2005a). Cependant, de nombreuses recherches font état d'une certaine « flexibilité » dans le transfert de l'apprentissage. L'intégralité de l'information contextuelle n'est pas nécessaire au développement d'un apprentissage et à sa mise en évidence. Il semblerait que seules les informations pertinentes pour la réalisation de la tâche soient nécessaires à la récupération des représentations contextuelles. Cette récupération sélective des informations contextuelles pertinentes est notamment compatible avec le modèle EBRW (« Exemplar Based Random Walk »), développé par Palmeri et Nosofsky (1997 ; Palmeri, 1997). Ce modèle est une extension de l'IT, en intégrant à la fois ses éléments et ceux du modèle GCM (« Generalized Context Model») proposé par Nosofsky (1986) dans le champ de la catégorisation. Ce modèle étend l'automatisation basée sur la mémoire à la prise en compte des effets de similarités entre les exemplaires en mémoire et le nouveau cas à traiter, en s'appuyant sur les théories de la mémoire et de la catégorisation. Comme dans l'IT, le modèle EBRW considère l'automatisation comme le passage d'une conduite basée sur le recours à un algorithme à une conduite basée sur la récupération en mémoire d'exemplaires spécifiques. Mais la récupération est ici déterminée par le degré de similarité perceptive qu'entretient le cas avec les exemplaires en mémoire. Dans le modèle GCM de Nosofsky, les décisions de catégorisation sont basées sur la sommation des similarités relatives d'un item avec les exemplaires de chaque catégorie établie. Pour simplifier, plus un item partage de fortes similarités avec des items d'une catégorie A, plus il aura de chance d'être catégorisé comme membre de la catégorie A. Et plus il présente de fortes similarités avec les items d'une catégorie B, moins il aura de chance d'être catégorisé comme membre de la catégorie A, puisque ce dernier pourra aussi être classé dans la catégorie B. Dans le modèle EBRW, tous les exemplaires en mémoire concourent « contre l'algorithme », mais leur probabilité d'être récupérés est proportionnelle à la similarité relative qu'ils partagent avec l'item présenté. Les

exemplaires les plus susceptibles de gagner la compétition sont les plus similaires au cas présenté (voir aussi, Logan, 2002, 2004 pour une extension du modèle EBRW, inspirée des modèles inscrits dans le domaine de la catégorisation, de la mémoire et de l'attention visuelle).

Cette flexibilité des traitements automatiques basés sur la mémoire, amène à poser les questions suivantes : Sur quel(s) critère(s) de ressemblance la solution (e.g. localisation de la cible) d'un exemplaire stocké en mémoire est-elle activée ? Cette ressemblance tient-elle seulement aux traits perceptifs des exemplaires ou des effets de généralisation peuvent-ils être observés à partir de régularités sémantiques ?

6. Perspectives

Les recherches reposant sur l'utilisation du paradigme d'indication contextuel montrent que le système visuo-cognitif est doté d'une étonnante capacité à conserver de manière implicite des traces sur des régularités contextuelles spatiales. Ces recherches ont aussi permis de mettre en évidence des effets d'apprentissage implicite basés sur la forme des éléments contextuels, ainsi que sur le déroulement temporel des indices contextuels. Le paradigme d'indication contextuel constitue ainsi un outil expérimental adapté pour étudier des mécanismes d'apprentissage implicite intervenant dans la perception visuelle. L'avantage de ce paradigme expérimental est de mettre en évidence des effets d'apprentissage dans une tâche de recherche visuelle de cible, au cours de laquelle le sujet ne reçoit pas pour consigne d'étudier ou d'apprendre le matériel. Mais surtout, les effets d'apprentissage sont mis en évidence indirectement, sans avoir recours aux jugements des sujets. En plus de constituer une mesure objective des connaissances, cette tâche indirecte de mémoire (i.e. la tâche de recherche visuelle de cible) s'avère plus sensible aux connaissances des participants que des tâches directes, comme certaines tâches de jugements (e.g. de grammaticalité), définies par défaut comme mesure indirecte de la mémoire. D'autre part, les résultats aux tâches directes de mémoire, qu'il s'agisse des protocoles verbaux ou des tâches de choix forcé, démontrent de manière incontestable que les connaissances acquises ne sont pas accessibles à la conscience. Cependant, lorsque les contextes ne sont plus définis par des stimuli sans

signification mais par des scènes du monde réel, non seulement les participants reconnaissent mieux que du seul fait du hasard les scènes répétées, mais en plus, ils sont capables de rapporter la localisation de la cible au sein de ces scènes. Les associations scène-cible sembleraient être encodées explicitement en mémoire. La simple présence d'indices sémantiques dans les scènes du monde réel permettrait-elle une prise de conscience des régularités contextuelles ? Si des effets d'indication contextuel implicite ont été mis en évidence à partir de régularités perceptives, aucune étude n'a jusqu'alors rapporté des effets d'indication contextuel basés sur des propriétés catégorielles et sémantiques du contexte qui soient implicites. Cela signifie-t-il que l'apprentissage de régularités sémantiques est nécessairement associé à la prise de conscience de ces régularités ? L'apprentissage implicite est-il limité à la prise en compte de régularités perceptives ou peut-il être étendu à des aspects conceptuels ? Si des effets d'indication contextuel implicite sémantiques basés sur des régularités sémantiques peuvent être obtenus, il s'agira en outre d'étudier si ces effets requièrent ou non une attention visuelle sélective. Enfin, on peut se demander dans quelle mesure des effets d'indication contextuel peuvent être mis en évidence dans une activité plus écologique, par exemple au cours de la conduite automobile.

DEUXIÈME PARTIE

PROBLEMATIQUE

Les recherches sur la perception de scènes naturelles suggèrent que, contrairement à notre phénoménologie qui nous assure des représentations conscientes calquées du monde, ces dernières seraient davantage schématiques (e.g. Rensink, 2000a). Pourtant, l'attention est orientée de manière rapide et efficace vers les éléments pertinents de la scène, compte tenu des objectifs finalisés de l'observateur. Quels sont les facteurs responsables de cette adaptabilité, quand bien même peu d'information serait accessible à la conscience à un instant donné? Il est largement admis que des connaissances préalables, relatives à la scène et à la situation en cours, jouent un rôle déterminant sur le guidage de l'attention au sein de l'environnement. Au cours de l'analyse d'une scène visuelle, les informations ascendantes interagissent en permanence avec les connaissances sur les régularités de l'environnement pour faciliter le déploiement attentionnel. Ces connaissances sur les régularités de l'environnement seraient inscrites en mémoire à long terme sous forme de schémas de scène, lesquels décrivent les objets et les événements susceptibles d'apparaître dans une scène, ainsi que leur position relative les uns par rapport aux autres. L'activation précoce de ces schémas au cours des traitements visuels, contribuerait à expliquer leur efficacité, à la fois en facilitant l'interprétation de la scène et l'identification des objets qui la composent (Biederman, Mezzanotte, & Rabinowitz, 1982), mais également son exploration (Friedman, 1979 ; Henderson, Weeks, & Hollingworth, 1999). Néanmoins, les mécanismes d'apprentissage impliqués dans la construction des schémas de scène sont peu définis dans la littérature. Comment les connaissances sur les régularités de l'environnement sont-elles acquises et récupérées lors de l'exploration d'une scène visuelle, puisque peu d'information semble être mise en mémoire de manière explicite ?

L'hypothèse testée dans nos recherches est que des mécanismes d'apprentissage implicite participent à la construction des schémas de scène. Si le domaine de l'apprentissage implicite fourmille aujourd'hui encore de désaccords terminologiques et conceptuels, la plupart des auteurs s'accordent à définir l'apprentissage implicite comme un mode d'adaptation par lequel le comportement d'un individu devient, de manière incidente, sensible à une structure,

sans que la connaissance qui en résulte ne soit rapportable verbalement, voire accessible à toute conscience. Dans cette approche, deux questions fondamentales restent à élucidées. Quelle est la nature des connaissances inaccessibles à la conscience ? Sont-elles seulement de nature spécifique et notamment perceptive ou peuvent-elles être de nature conceptuelle en reposant sur une catégorisation sémantique des éléments présents dans la scène ? L'apprentissage des régularités sémantiques requiert-il nécessairement qu'une attention sélective soit portée sur les éléments contextuels ? La question des niveaux de profondeur des connaissances implicites fait l'objet d'une littérature pour le moins polémique. Mais si cette question a été largement débattue autour de l'apprentissage implicite de règles abstraites complexes, à ce jour, peu d'études ont rapporté des effets d'apprentissage implicite basés sur des propriétés sémantiques (cependant, Didierjean, 2007 ; Didierjean & Lemaire, 2005 ; Lambert & Sumich, 1996). Le rôle de l'attention dans l'apprentissage implicite est tout autant controversé. Dans le domaine de la perception visuelle par exemple, de nombreuses études rapportent des effets d'apprentissage sur des stimuli extérieurs au focus attentionnel. Pourtant, de nombreux auteurs défendent le point de vue selon lequel l'apprentissage implicite, au même titre que l'apprentissage explicite requiert une attention sélective.

Pour aborder expérimentalement ces questions, nous avons eu recours au paradigme d'indilage contextuel, développé par Chun et Jiang (1998). Dans la plupart des travaux conduits avec le paradigme d'indilage contextuel, des effets d'apprentissage implicite ont été observés à partir d'affichages arbitraires dans lesquels les régularités manipulées tenaient à des aspects spécifiques du contexte, tels que l'agencement spatial des éléments contextuels (Chun & Jiang, 1998), leurs formes spécifiques (Chun & Jiang, 1999 ; Endo & takeda, 2004) ou encore leur relation temporelle (Chun & Jiang, 1999 ; Olson & Chun, 2001). Des effets d'indilage contextuel ont également été obtenus en utilisant des scènes relatives au monde réel comportant des indices sémantiques (Brockmole & Henderson, 2006a). Mais lorsque les contextes ne sont plus définis par des stimuli arbitraires mais par des scènes du monde réel, les effets d'apprentissage sont explicites. Brockmole et Henderson (2006a) ont proposé que l'extraction rapide d'indices sémantiques présents dans les scènes du monde réel permette de prendre conscience des régularités et que cette prise de conscience des régularités contextuelles accélère l'apprentissage qui sous-tend l'indilage contextuel. L'apprentissage de régularités sémantiques est-il nécessairement associé à une prise conscience de ces régularités ?

Les travaux de recherche rapportés dans les chapitres suivants visaient à mieux comprendre les mécanismes impliqués dans l'apprentissage des régularités de l'environnement au cours de l'analyse de scènes visuelles. Dans cette perspective, trois études ont été conduites. Les procédures expérimentales utilisées au cours de ces trois études étaient

inspirées du paradigme d'indication contextuel. Dans l'Etude 1, nous avons testé si des régularités spécifiques du contexte d'une part, et des régularités catégorielles et sémantiques du contexte d'autre part, pouvaient être apprises de manière implicite et faciliter le guidage de l'attention. L'Etude 2 visait, en premier lieu à tester l'hypothèse d'un apprentissage implicite de régularités basées sur l'appartenance catégorielle et sémantique du contexte et en second lieu, à tester le rôle de l'attention sélective dans l'apprentissage implicite de régularités sémantiques. L'objectif de l'Etude 3 était d'explorer l'apprentissage de régularités contextuelles dans une tâche plus proche d'une activité écologique telle que la conduite automobile, et de tester l'hypothèse que les schémas de scène peuvent aussi être sources de défaillances fonctionnelles.

CHAPITRE IV

INDIÇAGE CONTEXTUEL BASÉ SUR DES RÉGULARITÉS SPÉCIFIQUES OU CATÉGORIELLES D'ENVIRONNEMENTS NUMÉRIQUES

Au cours de cette première étude, nous avons tout d'abord cherché à reproduire des effets d'indiçage contextuel tenant à l'identité spécifique des éléments contextuels, indépendamment de leur arrangement spatial. Nous avons d'autre part abordé un nouvel aspect des régularités contextuelles, à savoir l'existence de régularités reposant sur l'appartenance catégorielle et sémantique des éléments contextuels, et non sur leur identité spécifique. A supposer que des effets d'indiçage contextuel basés sur des propriétés catégorielles et sémantiques soient obtenus, il s'agissait, en outre, d'évaluer si l'apprentissage mis en évidence reste implicite ou si un apprentissage impliquant un processus de catégorisation basé sur des propriétés sémantiques des éléments contextuels est nécessairement associé à une prise de conscience de ces régularités contextuelles.

Cinq expériences réalisées au cours de cette étude sont rapportées ici. La partie expérimentale est découpée en deux parties. La première partie présente deux expériences visant à reproduire des effets d'indiçage contextuel basés sur des éléments contextuels spécifiques. La deuxième partie présente trois expériences visant à étudier si des effets d'indiçage contextuel peuvent être obtenus à partir de propriétés catégorielles sémantiques. Ces expériences mettaient toutes en oeuvre des tâches d'indiçage contextuel utilisant des affichages constitués d'items numériques à travers lesquels les participants devaient rechercher un nombre pré-défini. Dans les expériences décrites dans la Partie 1, l'identité

spécifique des nombres constitutifs du contexte était prédictive de la localisation de la cible. Par exemple, un contexte composé des nombres « 11 » et « 57 » était prédictif d'une localisation particulière de la cible. Dans les expériences décrites dans la Partie 2, ce n'était plus les éléments spécifiques du contexte qui était prédictif mais la propriété de parité/impairité. Par exemple, quand le contexte était composé de nombres pairs, la cible était localisée dans une zone particulière de l'affichage. Quand le contexte était composé de nombres impairs, la cible était localisée dans la zone opposée de l'affichage.

PARTIE 1

INDIÇAGE CONTEXTUEL BASÉ SUR DES RÉGULARITÉS SPÉCIFIQUES

Deux expériences ont été conduites afin d'étudier si des effets d'*indiçage contextuel* peuvent être obtenus à partir d'associations spécifiques entre un contexte numérique particulier (constitué d'un couple de nombres dupliqué 8 fois) et une localisation particulière de la cible à détecter (13 ou 28, par exemple). Le principe général était comparable à celui utilisé par Endo et Takeda (2004). Ces expériences comportaient deux phases : une tâche de recherche, immédiatement suivie d'une tâche de verbalisation puis de reconnaissance.

Dans la tâche de recherche, les participants étaient invités à détecter le plus rapidement et correctement possible laquelle des deux cibles prédéfinies (e.g. 13 ou 28) était présente parmi 16 items distracteurs (i.e. un couple de nombres dupliqué huit fois, distribués aléatoirement dans l'affichage). Les expériences étaient composées d'une série de blocs d'essais. Chaque bloc était composé d'essais « prédictifs », dans lesquels la spécificité du contexte était prédictive de la localisation de la cible, et d'essais « non prédictifs », dans lesquels le contexte n'était pas prédictif de la localisation de la cible. La Figure 14 illustre deux exemples d'essais prédictifs dans lesquels le contexte (le couple « 11-57 » dupliqué 8 fois) était prédictif de la localisation de la cible (13 ou 28). L'hypothèse testée était que si l'identité spécifique du contexte numérique peut indiquer la localisation de la cible, les temps de détection de la cible devraient être plus courts lorsqu'elle est présentée dans des contextes prédictifs que dans des contextes non prédictifs.

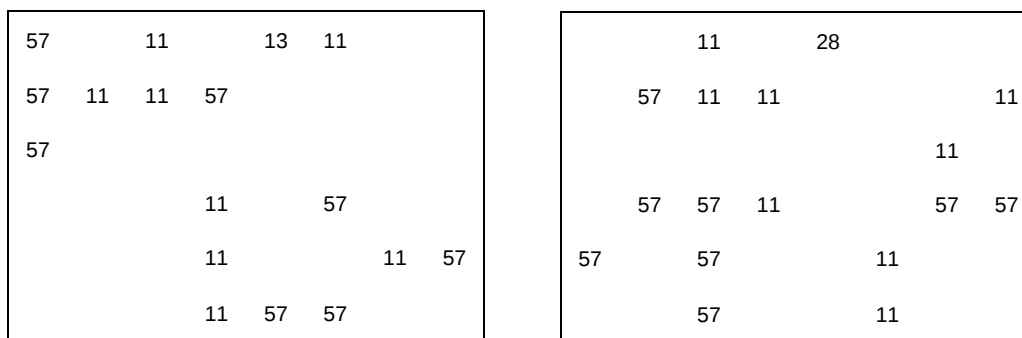


Figure 14 : Exemple de deux essais prédictifs. Dans un essai prédictif un couple spécifique de nombres contextuels (ici 11-57) était associé à une localisation particulière de la cible (ici la position occupée par la cible 13 ou 28). La configuration spatiale des éléments contextuels était aléatoire.

Par ailleurs, afin d'évaluer le caractère explicite vs. implicite des connaissances acquises lors de la tâche de recherche, les sujets étaient ensuite exposés à une tâche de verbalisation puis de reconnaissance (Chun & Jiang, 1998).

EXPÉRIENCE 1A

Méthode

Participants

Seize étudiants de l'Université d'Aix-en-Provence ont participé à l'expérimentation (11 femmes et 5 hommes, âgés en moyenne de 21 ans). Tous présentaient une acuité visuelle normale ou corrigée. Aucun n'avait connaissance du but de cette expérience.

Dispositif expérimental

L'expérience était pilotée par un ordinateur portable Macintosh dont l'écran mesurait 15 pouces. Les stimuli (type de police : Arial ; taille de la police : 24) étaient programmés par le logiciel Power Point et apparaissaient sur une grille invisible de 8 colonnes et 6 lignes. Il existait donc 48 localisations potentielles des items. L'expérience était générée à partir du logiciel Psyscope. Les sujets étaient installés à une distance approximative de 50 cm de l'écran.

Stimuli

Les cibles. Chaque essai comportait une seule cible pouvant prendre deux valeurs numériques. Deux couples de cibles ont été retenus: «13/18» d'une part, et « 23/28 » d'autre part, afin de neutraliser les effets potentiels liés à la parité/imparité du chiffre de la dizaine. En ce qui concerne les unités, les chiffres « 3 » et « 8 » ont été retenus, d'une part car l'un est pair et l'autre impair, et d'autre part car ils présentent a priori de fortes ressemblances perceptives. Les deux cibles pouvaient apparaître dans 12 localisations différentes sur la grille invisible 8x6 (les 12 localisations potentielles de la cible sont présentées en Annexe 3a).

Les distracteurs. Chaque aire de recherche comportait 16 distracteurs, i.e. une paire de nombres répétée huit fois. Ces distracteurs étaient des nombres à deux chiffres, différents des deux cibles potentielles. Leur dizaine et leur unité étaient toujours différentes de 3 et de 8. A chaque essai, les 16 distracteurs étaient distribués de façon aléatoire sur la grille invisible 8x6. Les 17 items (les 16 distracteurs et la cible) étaient tous présentés en noir sur un fond blanc.

L'expérience comportait deux types d'essais : des essais prédictifs et des essais non prédictifs. Dans un essai prédictif, le couple de distracteurs définissant le contexte était associé à une localisation particulière de la cible. Les 12 couples de distracteurs suivants ont été retenus : 11-57, 19-75, 55-71, 17-79, 51-77, 15-59, 20-44, 26-62, 42-60, 24-64, 40-66, 22-46. Dans un essai non prédictif, le couple de distracteurs était tiré aléatoirement sans remise parmi un ensemble de nombres de 10 à 99 (en excluant les couples choisis pour les essais prédictifs, ainsi que les nombres comportant les chiffres 3 et 8). Dans les essais non prédictifs, l'identité des distracteurs ne prédisait donc pas la localisation de la cible. Les cibles apparaissaient dans les mêmes localisations potentielles que dans les essais prédictifs (cf. Annexe 3a).

Procédure

L'expérience comportait deux phases expérimentales, une tâche de recherche, immédiatement suivie d'une tâche de verbalisation puis de reconnaissance.

Tâche de recherche. Les sujets avaient pour consigne de détecter aussi rapidement que possible la cible présente dans l'affichage, « 13 » ou « 18 » pour la moitié des sujets, et « 23 » ou « 28 » pour l'autre moitié. Les participants devaient répondre en appuyant sur le bouton du clavier correspondant à la cible présente. La touche assignée à chaque cible était contrebalancée entre les sujets.

La tâche de recherche comportait 24 blocs de 24 essais, soit au total 576 essais. Chaque bloc était composé de 12 essais prédictifs et de 12 essais non prédictifs, présentés dans un

ordre aléatoire. Pour chaque essai prédictif, l'association d'un contexte particulier (par exemple le contexte composé des nombres 11 et 57) avec une localisation particulière de la cible était systématiquement répétée à chaque bloc d'essais. Dans cette expérience, l'identité du couple de distracteurs était prédictive de la localisation de la cible, indépendamment de la configuration spatiale des distracteurs. Par contre, dans les essais non prédictifs, les couples de distracteurs étaient générés aléatoirement à chaque essai. Les effets potentiels de la nature de la cible (23/28 pour la moitié des sujets, 13/18 pour l'autre moitié), de sa localisation, et des associations « localisation de la cible – contextes pairs vs. impairs », ont été contrebalancés au moyen d'un carré latin.

L'expérience commençait par les instructions, suivies d'un bloc d'entraînement, composé de 12 essais non prédictifs pour familiariser les participants avec la procédure expérimentale. Immédiatement après, les participants réalisaient la tâche expérimentale de recherche composée des 24 blocs de 24 essais. Au sein d'un bloc, les 24 essais étaient présentés de façon aléatoire. Les participants appuyaient sur un bouton pour commencer le premier bloc d'essais. Après un délai de 500 ms, une plage de stimuli apparaissait sur l'écran. Les participants devaient rechercher la cible et, dès sa détection, ils devaient appuyer le plus rapidement possible sur la touche assignée à la cible présentée. La réponse du sujet faisait immédiatement apparaître, durant une seconde, un écran blanc comportant au centre un point de fixation noir. L'essai suivant était ensuite initié par l'ordinateur. Une pause était accordée à la fin de chaque bloc. Les sujets relançaient librement la suite de la tâche de recherche. Ils pouvaient passer immédiatement au bloc suivant ou prolonger la pause à leur gré.

Tâche de verbalisation puis de reconnaissance. A l'issue de la tâche de recherche, les deux questions suivantes étaient posées oralement aux sujets : « Avez-vous remarqué certaines régularités dans la construction du matériel ? », puis : « Avez-vous remarqué que certains contextes numériques de la cible étaient répétés au cours de l'expérience ? » Si les participants répondaient ne pas avoir remarqué que certains nombres étaient répétés au cours de l'expérience, ils étaient exposés à la tâche de reconnaissance. Les sujets n'étaient pas informés au préalable de cette tâche. Elle était constituée d'un nouveau bloc d'essais, généré de la même façon que ceux de la tâche de recherche. Les sujets avaient pour consigne d'observer successivement chacun des 24 essais et de juger du caractère familier ou non familier de chacun de ces essais par rapport à la tâche de recherche. Les sujets n'étaient soumis à aucune contrainte temporelle pour répondre.

La durée totale de l'expérimentation était approximativement de 35 minutes.

Résultats

Tâche de recherche

Les taux d'erreurs sont inférieurs à 2%, aussi bien dans la condition prédictive que dans la condition non prédictive. Ils ne seront pas discutés en détail. Nos analyses portent exclusivement sur les temps de réponse correspondant aux réponses correctes. Les TR sur les réponses correctes ont été regroupés en 6 époques, correspondant à 4 blocs d'essais successifs. Pour chacun des sujets, les moyennes des TR correspondant aux réponses correctes au sein d'une époque ont été calculées séparément pour chacune des conditions testées. Les TR supérieurs ou inférieurs à la « moyenne + 3 écart-types » ont été éliminés. Environ 1,6% des TR corrects ont été supprimés par cette procédure.

Une ANOVA à mesures répétées a été réalisée, avec les facteurs intra-sujets, condition (prédictive vs non prédictive) et époque (1-6). Les TR moyens pour chacune des deux conditions en fonction des époques sont présentés Figure 15. Cette analyse indique un effet facilitateur des contextes prédictifs sur les temps de détection de la cible. En effet, l'ANOVA montre un effet du facteur condition, $F(1,15)=10.491$, $p<.01$, un effet du facteur époque, $F(5,75)=29.682$, $p<.001$ et une interaction entre les facteurs « condition x époque », $F(5,75)=4.266$, $p<.01$. Des analyses partielles révèlent un effet significatif du facteur condition en faveur des essais prédictifs, à l'époque 2, $F(1,15)=5.40$, $p<.05$, à l'époque 3, $F(1,15)=8.91$, $p<.01$ et à l'époque 6, $F(1,15)=8.39$, $p<.05$, suggérant un effet d'indiciage contextuel, aux époques 2, 3 et 6.

Tâche de verbalisation puis de reconnaissance

Aucun sujet n'a répondu avoir remarqué que certains contextes ou certains nombres étaient répétés au cours de l'expérience. Tous les sujets ont donc été exposés à la tâche de reconnaissance. Les pourcentages de réponse « familier » ont été analysés. Une ANOVA à mesures répétées n'indique pas de différence significative entre les pourcentages de réponse « familier » obtenus dans la condition prédictive (55%) et ceux obtenus dans la condition non prédictive (50%), $F(1, 15)=1.178$, $p=.298$.

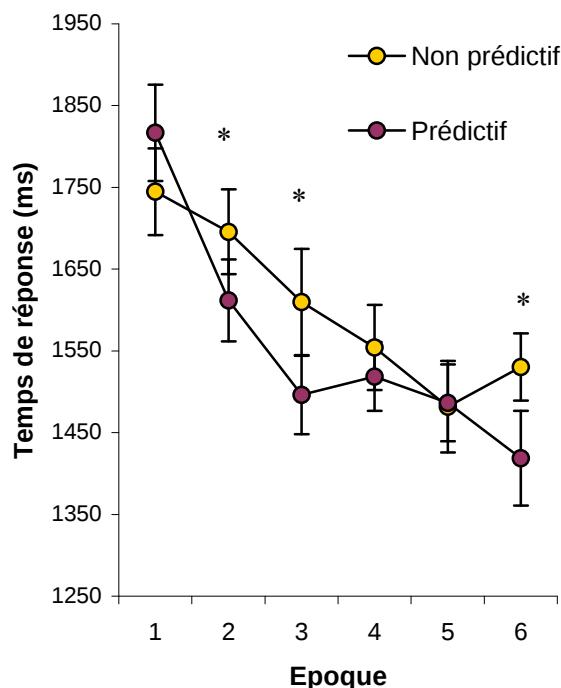


Figure 15. Temps de réponse dans les conditions prédictive et non prédictive.
Les barres d'erreurs correspondent aux SEM (N=16).

Discussion

Cette expérience suggère des effets d'indication contextuel lorsque des contextes numériques spécifiques sont prédictifs de la localisation de la cible. De plus, les résultats de la tâche de reconnaissance indiquent que les participants n'étaient pas à même de verbaliser ces régularités, ni même de les exploiter dans une tâche de reconnaissance. Cependant, cette expérience comporte un biais. En effet, l'identité des contextes prédictifs était systématiquement répétée, alors que l'identité des contextes non prédictifs était tirée aléatoirement à chaque essai. De ce fait, le bénéfice observé dans les essais prédictifs pourrait tenir, non pas à l'apprentissage associatif entre un contexte numérique spécifique et la localisation de la cible, mais à la répétition des contextes, indépendamment de leur caractère prédictif. Les participants pourraient avoir développé une stratégie dans des contextes systématiquement répétés leur permettant de discriminer plus rapidement les items distracteurs et par conséquent de traiter plus rapidement les contextes répétés. On notera entre parenthèses, qu'un tel biais existe aussi dans les expériences classiques d'indication contextuel, et notamment dans celles d'Endo et Takeda (2004).

EXPÉRIENCE 1B

Pour pallier cette critique, l'Expérience 1b a été conduite. Le principe était le même que celui de l'Expérience 1a, mais de manière différente, tous les contextes, prédictifs et non prédictifs, étaient répétés au cours de chaque bloc. Comme dans l'Expérience 1a, les contextes prédictifs étaient systématiquement associés à une localisation particulière de la cible. Dans les contextes non prédictifs, systématiquement répétés eux aussi, la localisation de la cible variait d'un bloc à l'autre. Ainsi, un contexte prédictif et un contexte non prédictif ne différaient que par le caractère prédictif vs. non prédictif de la localisation de la cible. Chaque bloc comportait donc des essais prédictifs, dans lesquels la spécificité du contexte prédisait la localisation de la cible, et des essais non prédictifs, dans lesquels la spécificité du contexte ne prédisait pas la localisation de la cible. Par ailleurs, afin de simplifier le plan expérimental, tous les sujets devaient rechercher la cible « 13 » ou « 28 ».

Méthode

Participants

Vingt-deux étudiants de l'Université d'Aix-en-Provence ont participé à l'expérimentation (12 femmes et 10 hommes, âgés en moyenne de 23 ans). Tous présentaient une acuité visuelle normale ou corrigée. Aucun n'avait connaissance du but de cette expérience.

Dispositif expérimental

Le dispositif expérimental était le même que dans l'Expérience 1a.

Matériel

Le matériel était le même que dans l'Expérience 1a, hormis les modifications suivantes. Les cibles potentielles étaient pour tous les sujets 13 et 28. Comme dans l'Expérience 1a, il y avait deux types d'essais, des essais prédictifs et des essais non prédictifs, définis par un couple de nombres dupliqués 8 fois dans l'affichage. Mais de manière différente, huit contextes prédictifs spécifiques et huit contextes non prédictifs spécifiques étaient établis. Comme dans l'Expérience 1a, la cible pouvait apparaître dans les mêmes localisations dans les essais prédictifs et dans les essais non prédictifs (pour les localisations potentielles des cibles, cf. Annexe 3b). Pour la moitié des participants, les huit contextes prédictifs étaient, 11-

57, 69-74, 55-70, 47-52, 22-26, 40-77, 45-54 et 14-66, et les huit contextes non prédictifs étaient 42-62, 56-65, 10-25, 21-50, 71-75, 29-49, 20-44 et 16-61. Pour l'autre moitié des participants, le plan expérimental était inversé.

Procédure

Tâche de recherche. La procédure était la même que dans l'Expérience 1a, hormis les modifications suivantes. Les participants avaient tous pour consigne de rechercher la cible 13 ou 28. Au cours de la phase de familiarisation, les participants étaient exposés à huit essais d'entraînement utilisant des couples contextuels différents de ceux utilisés au cours de la tâche de recherche effective. La tâche de recherche comportait 24 blocs de 16 essais (8 essais prédictifs et 8 essais non prédictifs), soit 384 essais au total.

Tâche de verbalisation puis de reconnaissance. Les participants étaient exposés aux mêmes questions que celles posées dans l'Expérience 1a, mais avec une question supplémentaire : « Avez-vous remarqué que certains contextes étaient associés à la localisation de la cible ? » La tâche de reconnaissance consistait en un nouveau bloc de 32 essais, dont 16 essais étaient générés de la même manière que dans la tâche de recherche, i.e. 8 essais prédictifs et 8 essais non prédictifs, ainsi que 8 essais « contre prédictifs » et 8 essais « non prédictifs de remplissage ». Dans un essai contre prédictif, la cible apparaissant dans un contexte prédictif était localisée dans une localisation très différente de celle normalement utilisée dans l'essai prédictif correspondant. Les 8 essais non prédictifs de remplissage ne servaient qu'à équilibrer le nombre de contextes prédictifs et non prédictifs au sein de la tâche de reconnaissance. Les sujets recevaient pour instruction pour chacun des 32 essais de juger si l'association entre les nombres contextuels et la localisation de la cible leur paraissait familière par rapport à la tâche de recherche.

Résultats

Tâche de recherche

Les taux d'erreurs étaient inférieurs à 1.6% dans les deux conditions expérimentales. Les TR supérieurs et inférieurs à la « moyenne + 3 écart-types » ont été éliminés des analyses, soit 1.2% des TR corrects. Les TR corrects ont été regroupés en six époques, recouvrant chacune quatre blocs successifs d'essais.

Une ANOVA à mesures répétées a été conduite avec les facteurs intra-sujet, condition (prédictive vs. non prédictive) et époque (1-6). Les TR moyens pour chaque condition à chaque époque sont présentés Figure 16. Ici encore un effet facilitateur des contextes prédictifs sur les temps de détection de la cible est observé. L'ANOVA révèle un effet du facteur condition, $F(1,21) = 5.986$, $p < .05$, un effet du facteur époque, $F(5,105) = 11.099$, $p < .001$, et une interaction entre les deux facteurs, $F(5,105) = 2.887$, $p < .05$. Des analyses partielles révèlent un effet du facteur conditions à l'époque 4, $F(1,21) = 5.26$, $p < .05$, et à l'époque 5, $F(1,21) = 19.06$, $p < .001$.

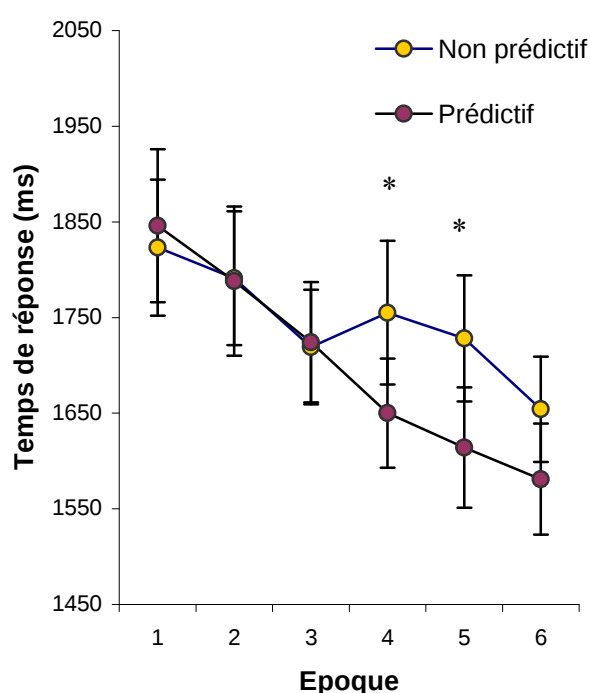


Figure 16. Temps de réponse dans les conditions prédictive et non prédictive.
Les barres d'erreurs correspondent aux SEM (N=22).

Tâche de verbalisation puis de reconnaissance

Aucun des participants n'a rapporté avoir remarqué que les contextes étaient répétés au cours de l'expérience. Tous ont été exposés à la tâche de reconnaissance. Les pourcentages de réponses « familier » ont été analysés. Une ANOVA à mesures répétées n'indique pas de différence dans les pourcentages de réponses « familier » observés dans les conditions prédictive (52%), non prédictive (52%) et contre prédictive (47%), $F(2, 42) < 1$.

Discussion

L'objectif de cette expérience était de généraliser les effets d'indication contextuel avec du matériel numérique, dans une situation où la régularité manipulée tenait à la spécificité du contexte, indépendamment de son agencement spatial. Les résultats indiquent que des effets d'indication contextuel peuvent être obtenus à partir d'associations prédictives entre certains contextes numériques spécifiques et une localisation particulière de la cible. Chun et Jiang (1998) ont montré que la configuration spatiale d'un contexte peut indiquer la localisation de la cible, et que l'identité des éléments composant l'environnement contextuel peut indiquer l'identité de la cible (Chun & Jiang, 1999). De manière similaire aux travaux d'Endo et Takeda (2004), nos résultats montrent que la spécificité des contextes peut indiquer la localisation de la cible. Cependant, dans cette expérience, les contextes prédictifs et non prédictifs étaient tous répétés au cours de la tâche de recherche. Par conséquent, les effets d'indication contextuel résultent, sans équivoque, d'un apprentissage associatif entre le contexte spécifique et la localisation de la cible. Par ailleurs, les résultats de la tâche de verbalisation puis de reconnaissance montrent le caractère implicite de l'apprentissage mis en évidence. En effet, au terme des 24 blocs d'essais, aucun participant n'a rapporté avoir remarqué une quelconque régularité quant à l'identité des nombres distracteurs. De surcroît, dans la tâche de reconnaissance, les participants n'étaient pas capables de différencier les essais prédictifs des essais non prédictifs ou contre prédictifs.

PARTIE 2

INDICATION CONTEXTUEL BASÉ SUR DES RÉGULARITÉS CATÉGORIELLES

Les Expériences 2a, 2b et 3, exposées dans cette deuxième partie visaient à déterminer si des effets d'indication contextuel peuvent également être obtenus lorsque les régularités contextuelles reposent sur l'appartenance catégorielle et sémantique des éléments contextuels plutôt que sur leur spécificité. Pour ce faire, nous avons rendu la propriété paire/impair du

contexte, prédictive de la localisation de la cible. Dans ces expériences, ce n'était donc plus des éléments spécifiques qui étaient prédictifs de la localisation de la cible, mais la propriété conceptuelle de parité/impairité. Le principe général était le même que dans les Expériences 1a et 1b. Chaque bloc d'essais était constitué d'essais prédictifs et d'essais non prédictifs. Dans la moitié des essais prédictifs, le contexte de la cible était composé exclusivement de nombres distracteurs pairs et dans l'autre moitié des essais, le contexte de la cible était composé exclusivement de nombres distracteurs impairs. Quand le contexte était pair, la cible apparaissait dans une zone déterminée de l'affichage (par exemple à gauche). Quand le contexte était impair, la cible apparaissait dans la zone opposée de l'affichage (par exemple à droite). Dans les essais non prédictifs, le contexte était composé d'autant de nombres pairs que de nombres impairs et la cible pouvait apparaître dans l'une ou l'autre de ces deux zones de l'affichage. Dans une expérience préliminaire, nous avons divisé l'affichage en deux parties égales par rapport à un axe vertical. Les localisations potentielles de la cible recouvraient l'ensemble de l'affichage. Cependant, nous ne sommes pas parvenus à montrer des effets d'apprentissage avec ce matériel. Il est fort probable que la régularité prédictive de la localisation de la cible était trop « discrète » pour conduire à un apprentissage associatif entre le contexte et la localisation de la cible et permettre un déploiement efficace de l'attention vers la zone où était localisée la cible, cette dernière étant trop étendue. C'est pourquoi dans les expériences suivantes, nous avons restreint les localisations potentielles de la cible aux extrémités gauche et droite de l'affichage, rendant ainsi la régularité plus « franche ». Les deux zones choisies pour les localisations potentielles des cibles sont présentées Figure 17.

C							C
	C					C	
	C					C	
C							C

Figure 17. Localisations des cibles dans les Expériences 2a, 2b et 3

EXPÉRIENCE 2A

Dans cette expérience, nous avons utilisé des chiffres comme cible et distracteurs. La tâche des sujets était de détecter la cible « 3 » ou « 8 » parmi un ensemble de 16 chiffres distracteurs. Dans la moitié des essais prédictifs, le contexte comportait exclusivement des chiffres distracteurs pairs et dans l'autre moitié des essais prédictifs, le contexte comportait exclusivement des chiffres distracteurs impairs. La propriété paire/impair du contexte prédisait la localisation droite/gauche de la cible (cf. Figure 17). Dans les essais non prédictifs, le contexte comportait autant de chiffres pairs que de chiffres impairs et la cible pouvait apparaître à gauche ou à droite de l'affichage.

Méthode

Participants

Trente étudiants de l'Université de Provence (15 femmes et 15 hommes, âgés en moyenne de 24 ans) ont participé à l'expérience.

Dispositif expérimental

Le dispositif était le même que celui des Expériences 1a et 1b.

Stimuli

Les cibles à détecter étaient 3 et 8. Un essai comportait toujours une et une seule cible. Chacune des cibles pouvait apparaître dans l'une des 8 localisations fixées sur la grille 8x6 (cf. Figure 17). Pour rendre la régularité saillante, ces 8 localisations appartenaient à deux zones restreintes de l'affichage. La cible apparaissait parmi un ensemble de seize distracteurs. Les 16 distracteurs étaient des chiffres de 0 à 9, hormis 3 et 8. Les 16 distracteurs étaient distribués de façon aléatoire sur la grille 8x6.

Procédure

La procédure expérimentale était la même que dans les expériences précédentes hormis les points suivants.

Tâche de recherche. Les sujets avaient pour consigne de détecter aussi rapidement que possible si la cible présente dans l'affichage était 3 ou 8. La tâche de recherche comportait 24

blocs de 16 essais. Chaque bloc comprenait 8 essais prédictifs et 8 non prédictifs, présentés dans un ordre aléatoire.

Parmi les essais prédictifs, 4 étaient des essais « prédictifs pairs », dans lesquels les distracteurs étaient exclusivement des chiffres pairs, et 4 étaient des essais « prédictifs impairs », dans lesquels les distracteurs étaient exclusivement des chiffres impairs. Pour la moitié des sujets, les contextes pairs étaient associés à une cible localisée dans la zone de gauche de l’affichage (i.e. dans l’une des quatre localisations de gauche présentées Figure 17), et les contextes impairs étaient associés à une cible localisée dans la zone droite de l’affichage (i.e. dans l’une des quatre localisations de droite présentées Figure 17). Pour l’autre moitié des participants, c’était l’inverse. Les 16 distracteurs définissant un contexte prédictif, pair ou impair, étaient tirés aléatoirement avec remise parmi la liste des quatre chiffres pairs, 0, 2, 4 et 6 ou parmi la liste des quatre chiffres impairs, 1, 5, 7, 9, respectivement.

Dans les essais non prédictifs, le contexte de la cible était composé de 8 chiffres pairs et de 8 chiffres impairs, en excluant les deux cibles potentielles (3 et 8). Pour s’assurer que les contextes des essais non prédictifs ne soient pas plus hétérogènes que ceux relatifs aux essais prédictifs, chaque contexte non prédictif ne pouvait comporter que quatre chiffres différents (deux pairs et deux impairs).

Le déroulement de l’expérience était le même que celui de l’Expérience 1.

Tâche de verbalisation puis de reconnaissance. Après la tâche de recherche, les participants étaient soumis aux questions suivantes : « Avez-vous remarqué des régularités dans le matériel ? », « Avez-vous remarqué que certains contextes étaient associés à la localisation de la cible ? » et « Avez-vous remarqué une règle qui prédisait la localisation de la cible ? ». La tâche de reconnaissance était constituée d’un nouveau bloc de 32 essais, dont 8 essais prédictifs et 8 non prédictifs, générés de la même façon que ceux de la tâche de recherche, plus 8 essais contre prédictifs et 8 essais non prédictifs de remplissage. Dans les essais contre prédictifs, la cible était localisée dans la zone opposée à celle relative aux contextes prédictifs. Si par exemple dans les contextes prédictifs pairs la cible était localisée à gauche de l’affichage, dans les essais contre prédictifs pairs, la cible était localisée à droite de l’affichage. La consigne donnée aux participants était la même que dans l’Expérience 1a.

Résultats et Discussion

Tâche de recherche

En condition prédictive comme en condition non prédictive, les taux d'erreurs étaient inférieurs à 2%. Comme dans les expériences précédentes, ces données ne seront pas analysées plus en détail. Les blocs d'essais ont été regroupés en 6 époques de 4 blocs successifs d'essais. Les TR correspondant aux erreurs ainsi que les TR supérieurs à la « moyenne + 3 écart-types » ont été éliminés des moyennes (2%).

Une ANOVA à mesures répétées avec les facteurs intra-sujets, condition (prédictive vs. non prédictive) et époque (1-6) a été réalisée. Les TR relatifs à ces deux conditions sont illustrés Figure 18. L'ANOVA montre un effet du facteur époque, $F(5,145)=26.966$, $p<.001$ et un effet du facteur condition, $F(1,29)=17.726$, $p<.001$, mais ne met pas en évidence d'interaction significative entre les facteurs « époque x condition », $F(5,145)= 1.377$, $p=.236$. Des analyses partielles révèlent un effet significatif du facteur condition uniquement à l'époque 3, $F(1,29)=5.06$, $p<.05$ et à l'époque 6, $F(1,29)=13.48$, $p<.001$.

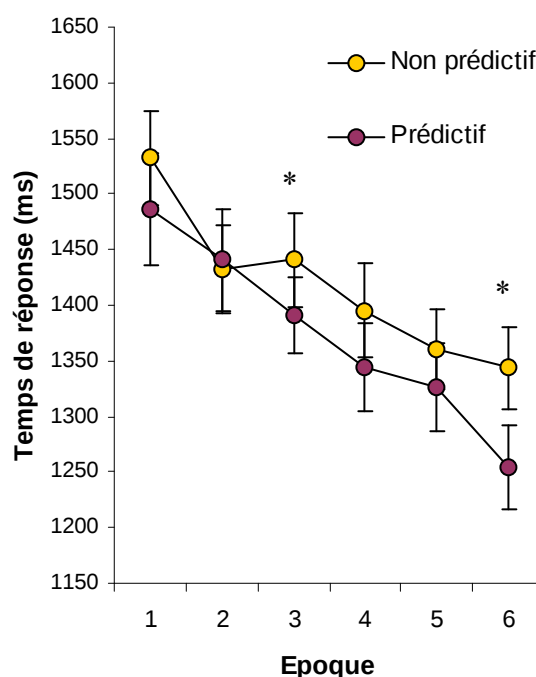


Figure 18. Temps de réponse dans les conditions prédictive et non prédictive. Les barres d'erreurs correspondent aux SEM (N=30).

Tâche de verbalisation puis de reconnaissance

A la fin de la tâche de recherche, aucun des 30 participants n'a évoqué la notion de parité/imparité. Tous ont donc été soumis à la tâche de reconnaissance. Une ANOVA à mesures répétées n'indique pas de différence significative entre les pourcentages de réponse « familier » obtenus dans les conditions prédictive (56%), non prédictive (55%), et contre prédictive (55%), $F(1,29) < 1$.

Ces résultats révèlent un effet d'indication contextuel lorsque la propriété de parité/imparité est prédictive de la localisation de la cible. Par ailleurs, le recueil des verbalisations et la tâche de reconnaissance indiquent que le bénéfice observé dans la condition prédictive résulte d'un apprentissage implicite.

EXPÉRIENCE 2B

L'objectif de l'Expérience 2b était de répliquer les données de l'Expérience 2a avec des nombres à deux chiffres. Le plan expérimental était identique à celui de l'expérience précédente, excepté que les sujets devaient détecter le nombre « 13 » ou « 28 » parmi un ensemble de nombres contextuels, i.e. 11, 15, 17, 19, 21, 25, 27 et 29 pour les nombres impairs, et 10, 12, 14, 16, 20, 22, 24, et 26 pour les nombres pairs.

Méthode

Participants

Vingt-six étudiants de l'Université de Provence (12 femmes et 14 hommes, âgés en moyenne de 24 ans) ont participé à l'expérience.

Dispositif expérimental

Le dispositif était le même que celui des expériences précédentes.

Stimuli

La cible à détecter était 13 ou 28. Les 16 nombres contextuels étaient des nombres à deux chiffres tirés dans la liste suivante pour les nombres impairs : 11, 15, 17, 19, 21, 25, 27, et 29, et dans la liste suivante pour les nombres pairs : 10, 12, 14, 16, 20, 22, 24, et 26.

Procédure

La procédure expérimentale était la même que dans l'Expérience 2a, hormis les modifications suivantes.

Tâche de recherche. Les participants recevaient pour instruction de rechercher le nombre 13 ou 28. Dans les essais prédictifs, les 16 nombres contextuels étaient tirés aléatoirement avec remise dans la liste des 8 nombres pairs pour les essais prédictifs pairs et dans la liste des 8 nombres impairs pour les essais prédictifs impairs. Dans les essais non prédictifs, le contexte était composé de 8 nombres pairs et de 8 nombres impairs, avec la contrainte que chaque contexte ne puisse contenir plus de 8 nombres différents (4 nombres pairs différents et 4 nombres impairs différents).

Tâche de verbalisation puis de reconnaissance. Les participants étaient exposés aux mêmes questions que dans l'Expérience 2a. Le principe de la tâche de reconnaissance était le même que dans l'Expérience 2a.

Résultats

Tâche de recherche

Dans les essais prédictifs et non prédictifs les taux d'erreurs étaient inférieurs à 1.5%. Les TR inférieurs et supérieurs à la « moyenne + 3 écart-types » ont été supprimés des analyses (1.6%). Comme dans les expériences précédentes, les blocs d'essais ont été regroupés en 6 époques, recouvrant chacune 4 blocs successifs d'essais.

Une ANOVA à mesures répétées a été conduite avec les facteurs intra-sujets, condition (prédictive vs. non prédictive) et époques (1-6). Les TR pour les deux conditions en fonction des époques sont illustrés Figure 19. Les résultats révèlent un effet du facteur condition, $F(1,25) = 5.049$, $p < .05$, un effet du facteur époque, $F(5,125) = 21.693$, $p < .001$, et une interaction entre ces deux facteurs, $F(5,125) = 2.404$, $p < .05$. Des analyses partielles indiquent un effet significatif du facteur condition en faveur de la condition prédictive, à l'époque 3,

$F(1,25) = 7.09, p < .05$, à l'époque 5, $F(1,25) = 8.66, p < .01$, et à l'époque 6, $F(1,25) = 4.47, p < .05$.

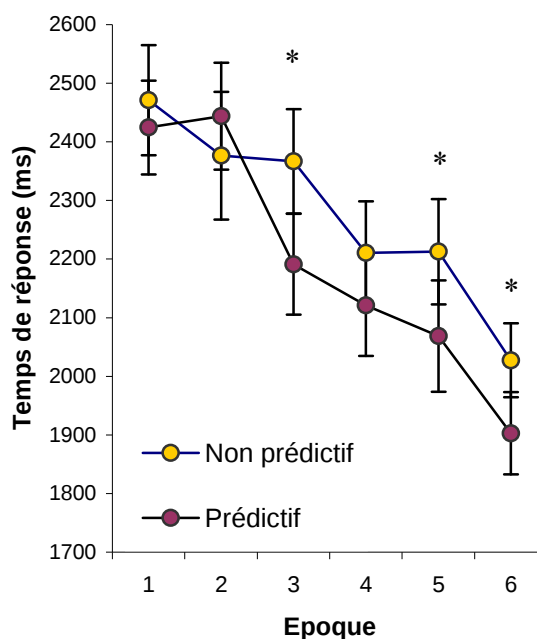


Figure 19. Temps de réponse dans les conditions prédictive et non prédictive.
Les barres d'erreurs correspondent aux SEM (N=26).

Tâche de verbalisation puis de reconnaissance

Aucun des participants n'a mentionné la propriété de parité/impairité. Tous ont été exposés à la tâche de reconnaissance. Une ANOVA à mesures répétées réalisées sur les pourcentages de réponses « familier » ne montre pas de différence significative entre les conditions prédictive, (55%), non prédictive (54%), et contre prédictive (56%), $F(2,50) < 1$.

Discussion des Expériences 2a et 2b

L'objectif des Expériences 2a et 2b était de tenter d'obtenir des effets d'indication contextuel reposant sur une régularité contextuelle de nature catégorielle et sémantique. Les résultats montrent un bénéfice sur les temps de détection lorsque la propriété paire/impair du contexte était prédictive de la localisation de la cible, que les affichages soient composés de chiffres ou de nombres. La réplication de ces effets atteste de leur robustesse. De surcroît, les données relatives à la tâche verbalisation puis de reconnaissance montrent que les sujets n'ont

pas détecté explicitement cette régularité, suggérant que ce bénéfice résulte d'un apprentissage implicite.

Cependant, ces résultats à eux seuls ne permettent pas de conclure à un indiçage contextuel basé sur une propriété catégorielle et sémantique. En effet, de même que dans les Expériences 1a et 1b, ce bénéfice pourrait tenir uniquement à l'extraction des traits spécifiques des items pairs ou impairs. Les chiffres ou nombres distracteurs étaient en effet tirés dans des listes spécifiques. Aussi les participants pourraient-ils simplement avoir appris des régularités basées non pas sur la propriété de parité/impairité, mais sur des régularités basées sur des traits spécifiques qui co-varient avec la propriété abstraite manipulée.

Afin de tester des effets d'indiçage contextuel basé sur la catégorie des contextes, nous avons conduit l'Expérience 3. Le but de l'Expérience 3 était de déterminer si l'indiçage contextuel peut être transféré à une nouvelle série de nombres, différents de ceux utilisés au début de la tâche de recherche.

EXPÉRIENCE 3

L'Expérience 3 visait à étudier si des effets d'indiçage contextuel peuvent être observés lorsque seule la propriété de parité/impairité est prédictive de la localisation de la cible. Pour ce faire, une phase de transfert a été implémentée durant la tâche de recherche. Puisque les Expériences 2a et 2b suggèrent que l'époque 3 constitue l'époque la plus favorable à la mise en évidence d'un effet d'indiçage contextuel, la phase de transfert a été implémentée à la troisième époque. Le principe était le même que celui de l'Expérience 2b : les participants devaient rechercher la cible 13 ou 28 parmi un ensemble de nombres contextuels à deux chiffres. Mais de manière différente, la tâche de recherche comportait seulement trois époques. A l'époque 3, les nombres contextuels étaient tirés dans des listes (paires et impaires) différentes de celles utilisées aux époques 1 et 2. Donc tout au long de la tâche de recherche, la propriété de parité/impairité du contexte prédisait la localisation de la cible (e.g. « contexte pair/cible à droite » vs. « contexte impair/cible à gauche »), mais à l'époque 3, les traits visuels des nombres contextuels étaient différents de ceux utilisés aux époques précédentes. Par conséquent, si un effet d'indiçage contextuel est observé à l'époque 3, i.e. au

moment du transfert, on pourra considérer que cet effet repose effectivement sur l'extraction de la propriété abstraite de parité/imparité.

Méthode

Participants

Vingt-huit étudiants de l'Université de Marseille ont participé à l'expérience (13 femmes et 15 hommes, âgés en moyenne de 21ans).

Dispositif expérimental

Le dispositif expérimental était le même que dans l'Expérience 1a.

Stimuli

La cible à détecter était 13 ou 28. Les 16 nombres contextuels étaient tirés aléatoirement avec remise dans les listes suivantes : 11, 15, 21 et 25 ou 17, 19, 27 et 29 pour les nombres impairs, et dans les listes suivantes : 10, 12, 20 et 22 ou 14, 16, 24 et 26 pour les nombres pairs.

Procédure

La procédure était la même que celle de l'Expérience 2b, avec les changements suivants.

Tâche de recherche. La tâche de recherche comportait 12 blocs de 16 essais, regroupés en 3 époques. La troisième époque constituait la phase de transfert. Une liste de nombres contextuels était utilisée aux époques 1 et 2, et une autre liste de nombres contextuels était utilisée à l'époque 3. Pour la moitié des participants, dans les essais prédictifs des époques 1 et 2 (8 blocs d'essais), les nombres contextuels étaient tirés aléatoirement avec remise dans la liste suivante pour les essais prédictifs impairs : 11, 15, 21, 25, et dans la liste suivante pour les essais prédictifs pairs : 14, 16, 24, 26. Les essais non prédictifs étaient générés en respectant les mêmes contraintes que dans l'Expérience 2b, mais avec les mêmes nombres que ceux utilisés dans les essais prédictifs (i.e. 11, 15, 21, 25, 14, 16, 24, 26). Dans les essais prédictifs de l'époque 3 (4 blocs de 16 essais), les nombres contextuels étaient tirés aléatoirement avec remise dans une nouvelle liste de nombres pour les essais prédictifs impairs, i.e. 17, 19, 27, 29, et dans une nouvelle liste de nombres pour les essais prédictifs pairs, i.e. 10, 12, 20, 22. Les contextes des essais non prédictifs étaient eux aussi générés à

partir de ces nouvelles listes de distracteurs. Pour l'autre moitié des participants, le plan expérimental était inversé.

Tâche de verbalisation puis de reconnaissance. La tâche de verbalisation puis de reconnaissance était la même que dans l'Expérience 2b mais avec les nombres contextuels utilisés aux époques 1 et 2 de la tâche de recherche.

Résultats

Tâche de recherche

Aussi bien dans la condition prédictive que dans la condition non prédictive, les taux d'erreurs étaient inférieurs à 1.5%. Les TR corrects supérieurs à la « moyenne + 3 écarts-types » ont été supprimés des analyses (1.6% des TR corrects). Les blocs d'essais ont été regroupés en 3 époques recouvrant chacune 4 blocs d'essais successifs.

Une ANOVA à mesures répétées a été conduite avec les facteurs, condition (prédictive vs. non prédictive) et époque (1-3). Les données sont présentées Figure 20. L'analyse montre un effet du facteur époque, $F(2,54) = 5.041$, $p < .01$, pas d'effet du facteur condition, $F(1,27) < 1$, mais une interaction entre les deux facteurs, $F(2,54) = 3.318$, $p < .05$. Des analyses partielles révèlent un effet significatif du facteur condition à l'époque 3 uniquement, $F(1,27) = 7.43$, $p < .05$.

Tâche de verbalisation puis de reconnaissance

Aucun des participants n'a mentionné la propriété de parité/imparité. Tous ont donc été exposés à la tâche de reconnaissance. Une ANOVA à mesures répétées ne montre pas de différence entre les pourcentages de réponses « familier » observés dans les conditions prédictive (50%), non prédictive (55%), et contre prédictive (50%), $F(2,54) = 1.29$, $p = .284$.

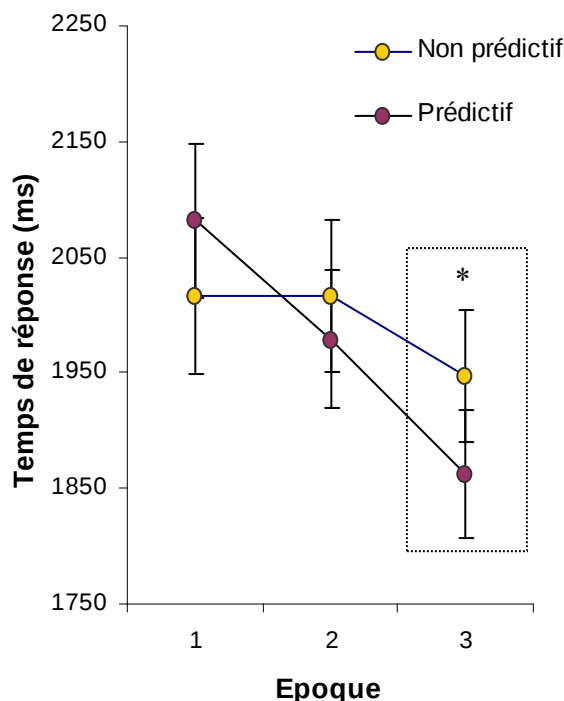


Figure 20. Temps de réponse dans les conditions prédictive et non prédictive.
 Les barres d'erreurs correspondent aux SEM (N=28). L'époque encadrée correspond au transfert.

Discussion

Cette expérience avait pour objectif de tester si l'indiciage contextuel peut reposer sur des régularités relatives à l'appartenance catégorielle et sémantique des nombres contextuels. Les résultats indiquent un effet d'indiciage contextuel, même lorsque de nouvelles séries de nombres sont utilisées au cours d'une phase de transfert, i.e. des nombres différents de ceux utilisés dans les essais précédents. Cet effet d'apprentissage ne peut donc être expliqué en terme d'apprentissage spécifique. L'indiciage contextuel observé à l'époque 3 repose effectivement sur la propriété catégorielle et sémantique de parité/impairité.

De plus, non seulement les participants étaient incapables de verbaliser les régularités catégorielles, mais en plus, ils étaient incapables de les utiliser pour discriminer les essais prédictifs des essais non prédictifs ou contre prédictifs dans une tâche de reconnaissance. Donc, même lorsque les régularités sont basées sur une propriété catégorielle – ici la parité/impairité – aucun des participants n'était à même d'exprimer explicitement cette régularité, ni même de l'exploiter dans une tâche de reconnaissance. Ces résultats montrent que les participants peuvent apprendre de façon implicite l'association entre une propriété

catégorielle, comme le concept de parité/imparité, et une position de la cible. Ils peuvent utiliser cette association dans une tâche de recherche afin d'accéder plus rapidement à la localisation de la cible, sans pour autant être capables d'exploiter ces régularités de manière explicite.

DISCUSSION GÉNÉRALE

L'objectif de cette étude était d'étudier si au cours de l'analyse visuelle, les observateurs sont capables d'apprendre, de manière implicite, tant des régularités contextuelles basées sur des caractéristiques spécifiques que sur des caractéristiques sémantiques et catégorielles. Les résultats mettent en évidence à la fois un apprentissage implicite basé sur des propriétés spécifiques des éléments contextuels (Expériences 1a et 1b), et un apprentissage implicite basé sur l'appartenance catégorielle sémantique des contextes (Expériences 2a, 2b et 3).

Les résultats des Expériences 1a et 1b généralisent les effets d'indilage contextuel obtenus par Chun et Jiang (1999) et Endo et Takeda (2004), où les régularités du contexte sont définies par la spécificité des éléments contextuels, indépendamment de leur agencement spatial. Cependant, la contribution de notre étude est d'avoir mis en évidence un apprentissage associatif entre l'identité spécifique des éléments constitutifs du contexte et la localisation de la cible, même lorsque les contextes non prédictifs sont systématiquement répétés au cours de la tâche de recherche. Les résultats des Expériences 2a et 2b montrent qu'un effet d'indilage contextuel peut également être obtenu lorsqu'une propriété catégorielle, en l'occurrence la propriété de parité/imparité des distracteurs, est prédictive de la localisation de la cible. Cependant, l'apprentissage observé pouvait aussi bien reposer sur la propriété catégorielle de parité/imparité que sur des propriétés spécifiques des éléments contextuels. Dans l'Expérience 3, après de nombreux essais d'apprentissage, une phase de transfert (utilisant des nombres contextuels différents de ceux utilisés au début la tâche) a été implémentée. Les résultats montrent qu'un effet d'indilage contextuel peut apparaître au moment du transfert, attestant d'un apprentissage basé sur la propriété de parité/imparité. Ces résultats montrent que les effets d'indilage contextuel ne sont pas limités à la prise en compte des traits spécifiques des éléments contextuels, notamment perceptifs, mais qu'ils peuvent être étendus à des aspects conceptuels, en reposant sur une catégorisation sémantique de ces éléments. Ainsi, selon leur

valeur prédictive, ce sont, soit des caractéristiques spécifiques des éléments contextuels, soit des propriétés catégorielles qui seront exploitées par le système visuo-cognitif, permettant la détection plus rapide d'une cible donnée.

Il est par ailleurs intéressant de noter que les effets d'indiciage contextuel ne sont pas constants au cours de la tâche. Il est fréquent que les effets d'indiciage contextuel apparaissent puis disparaissent. Dans l'Expérience 1a, un effet d'indiciage contextuel était obtenu aux époques 2, 3 et 6, dans l'Expérience 2a, aux époques 3 et 6 et dans l'Expérience 2b aux époques 3, 5 et 6. De manière intéressante, ces patterns de résultats partagent de fortes similitudes avec ceux obtenus par Chun et Jiang (1999, Expérience 1) et Olson et Chun (2002, Expérience 3). Une hypothèse en termes de compétition entre deux types de processus de recherche pourrait rendre compte de ce pattern : une recherche attentionnelle guidée par le recours à des algorithmes généraux, et une recherche attentionnelle guidée par la récupération en mémoire d'une solution spécifique. Chun et Jiang (1998, 2003) font en effet référence aux théories de l'automatisation basées sur la mémoire, et notamment à « l'Instance Theory », dans laquelle Logan (1988) propose que l'automatisation des performances dans un domaine donné provienne de la récupération directe de solutions relatives aux instances traitées antérieurement (voir aussi, Logan, 2002 ; Palmeri, 1997). Dans les tâches d'indiciage contextuel, les essais définissent en effet des instances spécifiques et les localisations de la cible constituent les solutions au test de recherche. Dans les essais non prédictifs et dans les premiers essais prédictifs, la recherche de la cible serait guidée par des *mécanismes attentionnels basés sur l'utilisation d'algorithmiques généraux*, correspondant à une recherche sérielle de la cible item par item. Après une exposition répétée aux essais prédictifs, des processus de récupération mnésique entreraient en compétition avec les traitements algorithmiques pour orienter l'attention de manière automatique vers la localisation de la cible (cf. Chun & Jiang, 1998 ; Peterson & Kramer, 2001). Nos résultats laissent supposer que vers la 3^{ème} époque, le système cognitif dispose de traces mnésiques relatives aux associations « contexte - localisation de la cible » suffisamment robustes pour optimiser le déploiement attentionnel. Au fur et à mesure de l'apprentissage et de sa consolidation, la solution serait de plus en plus accessible et sa récupération en mémoire de plus en plus rapide¹⁵. Cependant, les traitements de recherche basés sur l'utilisation d'algorithmes généraux deviennent eux aussi de plus en plus efficaces, comme en attestent l'amélioration des TR en condition non

¹⁵ A noter que cette conception diffère de celle de Logan, dans le sens où la probabilité qu'une trace mnésique soit récupérée dépend davantage de sa consolidation que de l'accumulation de traces relatives à un même cas.

prédictive. La stratégie de recherche basée sur l'algorithme pourrait elle aussi bénéficier d'une autre forme d'apprentissage (e.g. basé sur le traitement, où la quantité de ressources ou le nombre d'étapes nécessaires pour réaliser la recherche seraient réduit, Anderson, 1982)¹⁶. Il est possible qu'à un moment de la tâche (e.g. époque 4), la recherche basée sur l'algorithme devienne aussi, voire « plus efficace » qu'une recherche basée sur la récupération en mémoire de la « solution ». La cible serait alors détectée avant que la trace mnésique ne soit récupérée. A la fin de la tâche, les traces mnésiques seraient plus largement consolidées et leur activation serait de plus en plus systématique et/ou de plus en plus précoce, permettant de « gagner la course contre l'algorithme » et de « contrôler la réponse finale ». Ceci pourrait expliquer qu'un effet d'indilage soit souvent de nouveau significatif à la fin de la tâche.

Notre étude souligne en outre le caractère implicite des apprentissages mis en évidence. Dans les cinq expériences réalisées au cours de cette étude, aucun sujet n'a déclaré avoir remarqué des régularités prédictives de la localisation de la cible, ni même n'a été capable d'exploiter ces régularités pour différencier les essais prédictifs des essais non prédictifs ou contre prédictifs dans une tâche de reconnaissance. Le caractère implicite des apprentissages mis en évidence dans les tâches de recherche ne relève donc pas d'une simple difficulté des sujets à verbaliser leurs connaissances, mais bien d'une incapacité à accéder à ces connaissances de manière volontaire et consciente. Les effets d'indilage contextuel résultent donc bien d'un apprentissage implicite, même lorsque les régularités contextuelles reposent sur des propriétés catégorielles et sémantiques, et non sur des aspects spécifiques des éléments contextuels.

Dans la littérature sur l'analyse de scènes visuelles, de nombreux auteurs s'accordent à penser que les représentations visuelles sont sous-tendues par un *schéma de scène* stocké en MLT (Henderson & Hollingworth, 1999 ; Rensink, 2000a). Pour Rensink, cet aspect du contenu de nos représentations visuelles est fondamental car il fournit un cadre de référence qui permet de guider l'attention au sein de l'image vers les objets souhaités. Néanmoins, les mécanismes d'apprentissage impliqués dans la construction de *schémas de scène* sont peu définis. Au cours de notre étude, les observateurs semblent avoir appris de manière implicite où l'élément à détecter est localisé, dans un contexte spécifique, mais également dans un

¹⁶ Dans les théories de l'automatisation basée sur le traitement, l'apprentissage peut être transféré à des situations nouvelles familières par rapport aux items d'entraînement. La diminution des pentes relatives aux TR en fonction de la taille de la série dans la condition nouvelle (Chun & Jiang, 1998) étaye cette hypothèse.

contexte défini par l'appartenance catégorielle sémantique des éléments qui le composent. Ainsi, au cours de l'analyse d'une image ou d'une scène visuelle, le système visuo-cognitif semble pouvoir encoder implicitement, des relations spatiales entre les traits et les identités spécifiques des éléments, mais également des relations relatives à certaines propriétés catégorielles de ces éléments. Ces connaissances implicites, reposant sur l'extraction de co-occurrences basées sur des propriétés perceptives (e.g. Chun & Jiang, 1998 ; Fiser & Aslin, 2001), ou sur des propriétés plus conceptuelles des éléments de la scène, pourraient s'inscrire dans la construction de *schémas de scène* stockés en MLT. Ces schémas de scène orienteraient les traitements attentionnels de manière efficace et automatique. Les mécanismes d'apprentissage et de récupération implicites pourraient ainsi contribuer à expliquer la richesse de nos comportements et leur orientation adaptée, en dépit de tout contrôle et réalisation conscients.

CHAPITRE V

INDIÇAGE CONTEXTUEL SÉMANTIQUE ET ATTENTION VISUELLE SÉLECTIVE DANS DES ENVIRONNEMENTS LEXICAUX

Dans l'Etude 2, nous avons dans un premier temps cherché à généraliser les effets d'apprentissage implicite basés sur l'appartenance catégorielle et sémantique du contexte observés dans l'Etude 1 avec un nouveau matériel, i.e. lexical. Dans un deuxième temps, nous avons testé si l'apprentissage et la récupération implicites de régularités catégorielles et sémantiques requièrent une attention sélective. Quatre expériences reposant sur l'utilisation du paradigme d'indilage contextuel conduites avec des affichages composés de mots sont présentées dans le corps de texte. Trois expériences préliminaires sont présentées dans l'Annexe 4. De manière classique, chacune de ces expériences comportait une tâche de recherche, testant l'apprentissage de régularités contextuelles, suivie d'une tâche de verbalisation puis de reconnaissance, visant à éprouver le caractère implicite vs. explicite de l'apprentissage éventuellement mis en évidence dans la tâche de recherche.

Une série d'expériences préliminaires a été conduite avant de parvenir à mettre en évidence des effets d'indilage contextuel à partir d'un matériel lexical. Ces expériences sont détaillées en Annexe 4.

Les Expériences 4 et 5 testaient si des effets d'indilage contextuel basés sur l'appartenance catégorielle et sémantique des mots contextuels peuvent être obtenus. Pour ce faire, l'appartenance catégorielle des mots contextuels prédisait ou non la localisation de la

cible (i.e. le nom d'un vêtement ou d'un bâtiment). Par exemple, dans un contexte composé de noms de mammifère, la cible apparaissait toujours à gauche de l'affichage, alors que dans un contexte composé de noms de fruits/légumes, la cible pouvait aussi bien apparaître à gauche qu'à droite de l'écran.

Les Expériences 6 et 7 testaient quant à elles le rôle de l'attention sélective, sur l'apprentissage implicite sémantique, et sur son expression. Les procédures expérimentales étaient inspirées de celles utilisées par Jiang et Chun (2001) et Jiang et Leung (2005). Les participants devaient par exemple rechercher la cible écrite en vert parmi des mots verts et des mots rouges. Dans ces expériences, les régularités prédictives de la localisation de la cible étaient soit présentées dans la même couleur que la cible, soit dans une couleur différente.

EXPÉRIENCE 4

L'Expérience 4 testait si des effets d'indication contextuel peuvent être basés sur l'appartenance catégorielle et sémantique de contextes lexicaux. Le paradigme d'indication contextuel a ici été utilisé avec des affichages composés de mots. L'appartenance catégorielle des mots contextuels prédisait ou ne prédisait pas la localisation de la cible. Dans la tâche de recherche visuelle, les sujets avaient pour consigne de rechercher le nom d'un « vêtement » ou d'un « bâtiment » parmi un ensemble de mots contextuels. La cible n'était donc pas spécifique, mais était définie aux participants par deux catégories potentielles. Des catégories cibles ont été choisies plutôt que des mots spécifiques car des expériences préliminaires indiquent que l'utilisation de mots cibles spécifiques (e.g. « ourson » ou « lièvre ») rend difficile la mise en évidence d'effets d'indication contextuel basés sur des propriétés sémantiques des mots contextuels (Cf. Annexe 4). Il est probable que la recherche de mots cible spécifiques conduise à une recherche visuelle « superficielle », principalement basée sur des caractéristiques perceptives. L'utilisation de deux termes catégoriels pour désigner la cible contraignait ainsi les participants à effectuer un traitement sémantique des items de l'affichage pour accéder à la cible. Le mot cible pouvait apparaître dans deux zones restreintes de l'affichage (i.e. à droite vs. à gauche de l'affichage, Figure 21).

C							C
	C					C	
	C					C	
C							C

Figure 21. Localisations des cibles dans les Expériences 4 et 5.

A chaque essai, l'item cible était présenté parmi un ensemble de mots contextuels, appartenant tous à une même catégorie sémantique (i.e. mammifères, arbres/fleurs, oiseaux, ou fruits/légumes). Les mots contextuels étaient distribués de manière aléatoire dans l'affichage. La tâche de recherche visuelle comportait une série de blocs d'essais, chaque bloc comportant des essais prédictifs et des essais non prédictifs. Dans un essai prédictif, l'appartenance catégorielle et sémantique des mots contextuels était prédictive de la zone où était localisée la cible. Par exemple, quand le contexte était composé de noms de mammifères, la cible était toujours localisée à gauche de l'affichage (cf. Figure 22), quand le contexte était composé de noms d'oiseaux, la cible était toujours localisée à droite de l'affichage.



Figure 22. Exemples de deux essais prédictifs relatifs à l'association « mammifères – cible à gauche » utilisée dans les Expériences 4 et 5. Les mots cibles sont « BLOUSE » et « AUBERGE ».

Dans un essai non prédictif, la catégorie sémantique du contexte ne prédisait pas la zone où était localisée la cible. Par exemple, quand le contexte était composé de noms de fruits/légumes, la cible pouvait apparaître à gauche ou à droite de l'affichage, de même quand le contexte était composé de noms d'arbres/fleurs. Ainsi, la tâche de recherche visuelle

comportait deux catégories sémantiques prédictives (définissant les essais prédictifs) et deux catégories sémantiques non prédictives (définissant les essais non prédictifs).

La tâche de recherche était suivie d'une tâche de verbalisation puis de reconnaissance afin d'éprouver la nature explicite vs. implicite des connaissances acquises au cours de la tâche de recherche.

Méthode

Participants

Vingt-quatre étudiants français de l'Université d'Aix-en-Provence ont participé à l'expérimentation (19 femmes et 5 hommes) âgés en moyenne de 22,8 ans (étendue des âges : 18-26ans, $\sigma=3$). Tous présentaient une acuité visuelle normale ou corrigée. Tous étaient motivés par l'expérience et aucun n'avait connaissance du but et de la méthode de cette expérience.

Dispositif expérimental

L'expérience était pilotée par un ordinateur portable Macintosh dont l'écran mesurait 15 pouces. Les stimuli (Police : Courier New ; taille : 24) étaient programmés par le logiciel Power Point et apparaissaient sur une grille invisible de 8 colonnes et 6 lignes. Il existait donc 48 localisations potentielles des items. L'expérience était générée à partir du logiciel Psyscope. Les sujets étaient installés à une distance approximative de 50 cm de l'écran.

Stimuli

Les cibles. Chaque essai comportait une seule cible définie par deux catégories sémantiques potentielles. La cible pouvait être, soit un nom de « bâtiment », soit un nom de « vêtement ». Il existait 32 cibles potentielles, 16 pour chacune des deux catégories sémantiques établies (Cf. Annexe 5a). Toutes les cibles étaient des mots familiers de la langue française et l'appartenance catégorielle de ces mots ne suscitait, a priori, aucune équivoque. Chacune de ces 32 cibles apparaissait par permutations circulaires au cours de l'expérience, de façon à être présentée de manière équivalente dans chacune des conditions testées. Dans un essai, la cible pouvait apparaître dans l'une des 8 localisations présentées Figure 21.

Les contextes. Chaque aire de recherche comportait 14 mots contextuels. Ces 14 mots appartenaient tous à une même catégorie sémantique, clairement contrastée par rapport aux

deux catégories choisies pour les cibles. Il existait quatre catégories différentes de contexte : la catégorie « mammifères », la catégorie « oiseaux », la catégorie « arbres/fleurs » et la catégorie « fruits/légumes ». Chacune de ces catégories comprenait 13 mots. Les mots choisis, ainsi que leur fréquence d'usage sont présentés en Annexe 5b. La longueur moyenne des mots (de 5 à 7 lettres) était équilibrée dans chacune des 4 listes. Dans un essai, les 14 mots contextuels étaient tirés aléatoirement avec remise dans l'une des 4 listes catégorielles établies. Il y avait donc 4 types différents de contexte : contexte « mammifères », contexte « oiseaux », contexte « arbres/fleurs » et contexte « fruits/légumes ». Les 14 mots contextuels étaient distribués de façon aléatoire sur la grille invisible 8x6.

Les 15 items (les 14 mots contextuels et la cible) étaient tous présentés en noir sur fond gris.

Procédure

L'expérience durait environ 45 min et comportait deux phases expérimentales : une tâche de recherche visuelle et une tâche de verbalisation puis de reconnaissance.

Tâche de recherche. Les sujets avaient pour consigne de rechercher aussi rapidement que possible le nom d'un vêtement ou d'un bâtiment dans l'affichage. S'ils détectaient un nom de bâtiment, les participants devaient appuyer sur une touche désignée du clavier, s'ils détectaient un nom de vêtement, ils devaient appuyer sur une autre touche du clavier. La réponse à fournir en fonction de la cible présentée dans l'essai était contrebalancée entre les sujets. Dans la consigne, il était précisé aux sujets qu'un essai comportait toujours, soit un nom de vêtement, soit un nom de bâtiment, et jamais les deux.

La tâche de recherche comportait 24 blocs de 16 essais, soit au total 384 essais. Chaque bloc comprenait 8 essais prédictifs et 8 essais non prédictifs. Dans un essai prédictif, l'appartenance catégorielle des contextes prédisait la zone où était localisée la cible (e.g. à gauche de l'affichage dans un contexte « mammifères », et à droite de l'affichage dans un contexte « oiseaux »). Dans un essai non prédictif, l'appartenance catégorielle des contextes ne prédisait pas la zone où était localisée la cible. Les catégories prédictives vs. non prédictives étaient contrebalancées entre les participants. Ainsi, chaque catégorie était prédictive pour la moitié des participants et non prédictive pour l'autre moitié. La localisation de la cible (gauche vs. droite) pour chaque catégorie prédictive était aussi contrebalancée entre les participants. Chaque catégorie prédictive pouvait être associée aux localisations gauches (1/4 des participants) ou aux localisations droites (1/4 des participants). Les effets

potentiels des cibles ont été neutralisés en utilisant un plan expérimental par permutation circulaire établi en fonction des deux facteurs expérimentaux, bloc et condition.

L'expérience commençait par les instructions, suivies d'un bloc d'entraînement, composé de 8 essais non prédictifs visant à familiariser les participants avec la procédure expérimentale. Immédiatement après, les participants réalisaient la tâche expérimentale de recherche composée des 24 blocs de 16 essais. Au sein d'un bloc, les 16 essais étaient présentés de façon aléatoire. Les participants appuyaient sur un bouton pour commencer le premier bloc d'essais. Après un délai de 500 ms, une plage de stimuli apparaissait sur l'écran. Les participants devaient rechercher la cible (i.e. un nom de vêtement ou un nom de bâtiment) et, dès sa détection, ils devaient appuyer le plus rapidement possible sur la touche correspondant à la cible présente. La réponse du sujet faisait immédiatement apparaître, durant une seconde, un écran blanc comportant au centre un point de fixation noir. Dans la consigne, on demandait aux sujets de bien fixer ce point entre chaque essai. L'essai suivant était ensuite initié par l'ordinateur. Une pause était accordée à la fin de chaque bloc. Les sujets relançaient librement la suite de la phase test. Ils pouvaient passer immédiatement au bloc suivant ou prolonger la pause à leur gré.

Tâche de verbalisation puis de reconnaissance. A l'issue de la tâche de recherche, les trois questions suivantes étaient posées oralement aux participants : « Avez-vous remarqué certaines régularités dans la construction du matériel ? », « Avez-vous cherché des régularités ou cherché à mémoriser certains contextes ? » et « Avez-vous remarqué certaines associations entre le contexte et la cible, au cours de l'expérience ? » Si les participants répondaient ne pas avoir remarqué les régularités attendues, ils étaient exposés à la tâche de reconnaissance. Les participants n'étaient pas informés au préalable de cette tâche de reconnaissance. La tâche de reconnaissance consistait en un nouveau bloc de 32 essais, 16 étaient générés de la même manière que ceux de la tâche de recherche, i.e. 8 prédictifs et 8 non prédictifs, ainsi que 8 essais « contre prédictifs » et 8 « essais non prédictifs de remplissage ». Dans les essais contre prédictifs, la cible était localisée dans la zone opposée à celle utilisée dans les essais prédictifs (e.g. si dans un contexte prédictif mammifère la cible était localisée à droite de l'affichage, dans un essai contre prédictif mammifère, la cible était localisée à gauche de l'affichage). Les 8 essais non prédictifs de remplissage servaient juste à contrebalancer le nombre de contextes prédictifs et non prédictifs dans la tâche de reconnaissance. Les sujets recevaient pour instruction d'observer successivement chacun des 32 essais et de répondre à la question suivante : « Avez-vous un sentiment de déjà-vu concernant l'association entre le contexte et la localisation de la cible ? ». Si les sujets ne comprenaient pas la question, l'expérimentateur

leur reformulait la question : « Est-ce que dans ce contexte, il vous semble que la cible était localisée ici dans l’affichage ou est-ce qu’il vous semble qu’elle était localisée ailleurs, par exemple de l’autre côté de l’affichage ? ». Les sujets n’avaient pas de contrainte temporelle pour répondre.

Résultats

Tâche de recherche

Les taux d’erreurs étaient inférieurs à 1.61% pour les conditions prédictives et non prédictives. Ces résultats ne seront pas discutés plus en détail. Les TR corrects inférieurs ou supérieurs à « la moyenne + 3 écart-types » ont été supprimés des analyses (1.41% des réponses correctes). Les TR corrects ont été regroupés en 6 époques, recouvrant chacune quatre blocs successifs d’essais. Pour chaque participant, la moyenne des réponses correctes a été calculée à chaque époque et pour chaque condition.

Une ANOVA à mesures répétées a été effectuée avec les facteurs intra-sujets, condition (prédictive vs. non prédictive) et époque (1-6). Les TR moyens pour chaque condition en fonction des époques sont représentés Figure 23. L’ANOVA révèle un effet du facteur époque $F(5,115)=37.169$, $p<.001$, pas d’effet du facteur condition, $F(1,23)=1.885$, $p=.183$, mais une interaction entre les facteur « époque x condition », $F(5,115)=2.810$, $p<.05$. Les analyses partielles mettent en évidence un effet significatif du facteur condition à l’époque 3 uniquement, $F(1,23)=7.98$, $p<.01$, suggérant un effet d’indigence à cette époque. Un effet d’indigence contextuel était observé à l’époque 3 pour les contextes prédictifs mammifères (386 ms), $F(1,11) = 5.92$, $p < .05$, oiseaux (411 ms), $F(1,11) = 5.01$, $p < .05$, et fruit/légumes (241 ms), $F(1,11) = 5.24$, $p < .05$, mais pas pour les contextes prédictifs arbres/fleurs (67 ms), $F(1,11) < 1$ (pour le détail des résultats relatifs à chaque catégorie prédictives, cf. Annexe 6a).

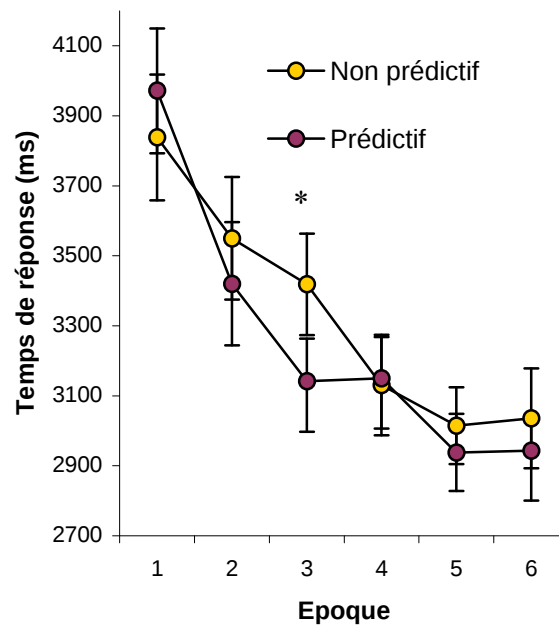


Figure 23. Temps de réponse dans les conditions prédictive et non prédictive.
Les barres d'erreurs correspondent aux SEM (N=24).

Afin de quantifier la vitesse d'apprentissage dans les contextes prédictifs et non prédictifs, nous avons estimé les paramètres des fonctions puissance ajustées aux données sur les TR moyens dans chaque condition (prédictive vs. non prédictive). Comme Chun et Jiang (2003), nous avons utilisé une fonction puissance car ce type de fonction est celui qui s'ajuste le mieux à nos données. L'équation de la fonction puissance est la suivante : $TR = a + bN^{-c}$, où TR correspond au temps de recherche, N est le nombre d'époques, la constante a représente l'asymptote d'apprentissage, la constante b est la différence entre la performance initiale et la performance asymptotique, et l'exposant c est le taux d'apprentissage. Nous avons imposé les mêmes contraintes que Chun et Jiang, à savoir, la performance asymptotique (a) doit être supérieure à 0, le niveau de la performance asymptotique finale doit être égal dans les conditions prédictive et non prédictive, et le taux d'apprentissage (c) doit être compris entre 0 et 1. Les fonctions obtenues pour chaque ensemble de données moyennes étaient les suivantes : dans la condition prédictive, $TR = 2535 + 1437^{-0.716}$ ($r^2=0.98$), et dans la condition non prédictive, $TR = 2535 + 1354^{-0.532}$ ($r^2=0.93$). Les taux d'apprentissage sont donc de 0.716 pour la condition prédictive et de 0.532 pour la condition non prédictive.

Tâche de verbalisation puis de reconnaissance

Aucun des participants n'a rapporté avoir cherché des régularités dans le matériel ou avoir cherché à le mémoriser. Aucun sujet n'a répondu avoir remarqué les régularités attendues. Tous les sujets ont donc été exposés à la tâche de reconnaissance. Une ANOVA à mesures répétées réalisée avec le facteur condition sur les pourcentages de réponse « déjà vu » n'indique pas de différence entre les conditions prédictive (46%), non prédictive (54%) et contre prédictive (55%), $F(2,27)=1.619, p=.209$.

Discussion

Les résultats montrent qu'un effet d'indication contextuel peut être obtenu avec un matériel lexical dans une situation où l'appartenance catégorielle et sémantique du contexte est prédictive de la localisation de la cible. En effet, un effet facilitateur sur la recherche visuelle est observé dans la condition prédictive par rapport à la condition non prédictive. Néanmoins, les analyses statistiques indiquent que l'indication contextuel n'est significatif qu'à l'époque 3. Ce bénéfice est toutefois significatif pour trois des quatre catégories utilisées dans cette expérience. De plus, l'estimation des pentes relatives aux courbes d'apprentissage montre que le taux d'apprentissage (le paramètre c) est plus important pour la condition prédictive (0.716) que pour la condition non prédictive (0.532). Pourquoi cet effet d'indication contextuel observé à l'époque 3, apparemment robuste, disparaît-il aux époques suivantes ? Ce résultat (que nous avons retrouvé dans une expérience préliminaire, cf. Annexe 4c), pourrait tenir à un changement de stratégie des participants au cours de la tâche de recherche. Au début de la tâche, les sujets ne connaissaient ni les cibles potentielles, ni les éléments contextuels potentiels. Ils étaient donc contraints à traiter les traits sémantiques des items pour trouver la cible. Mais, durant la tâche de recherche, les mêmes cibles étaient utilisées de manière cyclique. De même, les mots contextuels étaient tirés dans les mêmes listes spécifiques. Aussi, notre hypothèse est que les sujets pourraient rapidement disposer en mémoire de travail de connaissances sur les cibles et les éléments contextuels potentiels. Ils seraient dès lors capables de traiter les items en terme de « cible » vs « contexte », sans avoir besoin d'accéder à leur catégorisation sémantique pour effectuer leur recherche. De plus, les localisations potentielles de la cible étaient limitées à deux zones restreintes de l'affichage. Il est probable qu'au fur et à mesure que les sujets détectent cette régularité, ces derniers limitent leur recherche aux extrémités gauche et droite de l'écran. Le nombre d'éléments contextuels

traités diminuerait considérablement, et par conséquent les effets d'indiciage contextuel également (cf. Chun & Jiang, 1998, Expérience 4; Schneider & Shiffrin, 1977).

L'Expérience 4 montre donc qu'un effet d'indiciage contextuel peut être obtenu avec du matériel lexical, lorsque l'appartenance catégorielle et sémantique du contexte est prédictive de la localisation de la cible. De surcroît, la tâche de verbalisation puis de reconnaissance révèle le caractère implicite de l'apprentissage mis en évidence dans la tâche de recherche. En effet, à l'issue de la tâche de recherche, les participants n'étaient pas à même de rapporter les régularités contextuelles. Ils n'étaient pas non plus à même d'utiliser ces régularités pour différencier les essais prédictifs des essais non prédictifs ou contre prédictifs dans la tâche de reconnaissance.

Cependant, cette expérience ne permet pas de conclure à un apprentissage implicite qui serait vraiment basé sur l'appartenance catégorielle et sémantique du contexte, et non simplement sur la spécificité des items contextuels utilisés. En effet, comme dans les Expériences 2a et 2b, les éléments contextuels étaient tirés aléatoirement avec remise dans des listes spécifiques. L'apprentissage pourrait donc être basé sur la spécificité des éléments contextuels, et non sur leur appartenance catégorielle. Afin de tester si l'apprentissage implicite peut être basé sans équivoque sur les propriétés catégorielles du contexte, selon le même principe que dans l'Expérience 3, une phase de transfert a été implémentée dans l'Expérience 5. Tout au long de la tâche de recherche, les mêmes catégories prédisaient la localisation de la cible, mais au cours de la phase de transfert, les éléments contextuels étaient tirés dans des listes de mots différentes de celles utilisées au cours des essais précédents. L'Expérience 4 ayant montré un effet significatif d'indiciage contextuel à l'époque 3, la phase de transfert a été implémentée à cette époque.

EXPÉRIENCE 5

L'Expérience 5 visait à s'assurer qu'un effet d'indiciage contextuel peut être obtenu lorsque seule l'appartenance catégorielle et sémantique des mots contextuels est prédictive de la localisation de la cible. Pour ce faire, une phase de transfert a été implémentée à l'époque 3

de la tâche de recherche. Le principe méthodologique était le même que celui de l'Expérience 4, mais à l'époque 3, les mots contextuels étaient tirés dans des listes différentes de celles utilisées aux époques 1 et 2. Ainsi, si un effet d'indilage contextuel est obtenu à l'époque 3, cet effet sera basé, sans équivoque, sur l'appartenance catégorielle et sémantique des mots présents dans le contexte et non sur leurs caractéristiques spécifiques, notamment perceptives.

Méthode

Participants

Cinquante-six étudiants français de l'Université d'Aix-en-Provence ont participé à l'expérimentation (43 femmes et 13 hommes), âgés en moyenne de 23.4 ans (étendue des âges : 18-45, $\sigma=4$). Tous présentaient une acuité visuelle normale ou corrigée. Tous étaient volontaires et motivés par l'expérience. Aucun n'avait connaissance du but et méthode de cette expérience.

Dispositif

Le dispositif était le même que celui utilisé dans l'expérience précédente.

Stimuli

Les cibles étaient définies de la même manière que dans l'Expérience 4. Les éléments contextuels appartenaient aux mêmes catégories que dans l'Expérience 4, mais pour chacune des 4 catégories, une liste supplémentaire de 13 mots a été constituée. Ces listes sont présentées en Annexe 5b (listes 2).

Procédure

Tâche de recherche. La consigne et le principe général de la tâche étaient les mêmes que ceux de l'Expérience 4. L'expérience comportait 16 blocs de 16 essais, regroupés en 4 époques. Dans les 8 premiers blocs d'essais (époques 1 et 2), les éléments contextuels étaient tirés aléatoirement avec remise dans une série de listes (e.g. listes 1). Dans les 8 derniers blocs d'essais (époques 3 et 4), les éléments contextuels étaient tirés aléatoirement avec remise dans l'autre série de listes (e.g. listes 2). L'ordre de présentation de ces listes (listes 1 puis listes 2 vs. listes 2 puis listes 1) a été contrebalancé entre les participants au sein de chaque condition expérimentale. Le déroulement de l'expérience était le même que celui de l'Expérience 4.

Tâche de verbalisation puis de reconnaissance. La tâche était la même que celle de l'expérience précédente.

Résultats

Tâche de recherche

Les taux d'erreurs sont inférieurs à 1.53% aussi bien dans les essais prédictifs que dans les essais non prédictifs. Les TR corrects supérieurs ou inférieurs à la « moyenne + 3 écarts-types » ont été supprimés des analyses (1.65%). Une ANOVA à mesures répétées a été réalisée avec les facteurs intra-sujets, condition (prédictive vs. non prédictive) et époques (1-4). Les TR moyens relatifs à chaque condition pour les 4 époques sont présentés Figure 24. Les résultats révèlent un effet du facteur époque, $F(3,165)=80.451$, $p<.001$, pas d'effet du facteur condition, $F(1,55)=2.919$, $p=.093$, mais une interaction entre les deux facteurs, $F(3,165)=3.740$, $p<.05$. Des analyses partielles indiquent un effet du facteur condition à l'époque 3, $F(1,55)=9.05$, $p<.01$, ainsi qu'à époque 4, $F(1,55)=9.05$, $p<.01$, suggérant un effet d'indilage contextuel durant le transfert.

Pour s'assurer que l'effet observé à l'époque 3 ne tenait pas à un nouvel apprentissage, mais était bien la conséquence d'un apprentissage développé aux époques 1 et 2, nous avons analysé l'effet du facteur « condition » en se limitant à la première moitié des essais de l'époque 3. On observe une différence significative entre la condition prédictive (3314 ms) et la condition non prédictive (3609ms) dès la première moitié de l'époque 3 $F(1, 55) = 5.99$, $p<.05$.

Tâche de verbalisation puis de reconnaissance

Aucun sujet n'a répondu avoir cherché des régularités ou avoir cherché à mémoriser le matériel. Aucun sujet n'a répondu avoir remarqué les régularités contextuelles. Tous les sujets ont été exposés à la tâche de reconnaissance. Une ANOVA à mesures répétées réalisée sur les pourcentages de réponse « déjà vu » ne montre pas de différence entre les conditions « prédictive » (55%), « non prédictive » (58%) et « contre prédictive » (53%), $F(2,55)<1$.

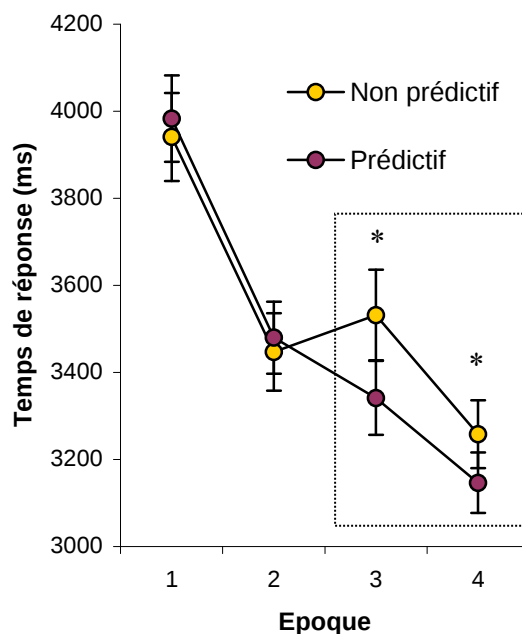


Figure 24. Temps de réponse dans les conditions prédictive et non prédictive.
 Les barres d'erreurs correspondent aux SEM (N=56). Les époques encadrées correspondent au transfert.

Discussion

Cette expérience visait à généraliser les effets d'indiciage contextuel basés sur l'appartenance catégorielle et sémantique des éléments contextuels. Les résultats de l'Expérience 5 montrent un effet d'indiciage contextuel à l'époque 3, alors que les mots contextuels étaient tirés dans des listes différentes de celles utilisées aux époques 1 et 2. L'analyse détaillée des résultats au cours de l'époque 3 montre que le bénéfice observé au cours de la phase de transfert est significatif dès les deux premiers blocs d'essais. Ce résultat suggère que l'apprentissage s'est effectivement développé au cours des époques 1 et 2, cet apprentissage se manifestant immédiatement après le transfert. L'apprentissage observé à l'époque 3 repose donc sans équivoque sur l'appartenance catégorielle et sémantique des mots contextuels et non sur leurs propriétés spécifiques.

De plus, les données relatives à chaque catégorie prédictive suggèrent que les quatre catégories tendent à conduire à un meilleur apprentissage que la condition non prédictive (pour les résultats détaillés relatifs à chaque catégorie prédictive, Cf. Annexe 6b). Des ANOVAs montrent un effet du facteur condition pour les contextes prédictifs oiseaux,

$F(1,27)=5.799$, $p<.02$, et fruits/légumes, $F(1,27)=4.607$, $p<.05$, et une interaction entre les facteurs époque x condition pour les contextes prédictifs mammifères, $F(3,81)=2.771$, $p<.05$, et arbres/fleurs, $F(3,81)=2.955$, $p<.05$.

Par ailleurs, au terme de la tâche de recherche, aucun des participants n'a rapporté avoir remarqué que certains contextes étaient associés à des localisations particulières de la cible. Dans la tâche de reconnaissance, les résultats ne montrent pas de différence entre les pourcentages de réponses « déjà vu » obtenus dans les conditions prédictive, non prédictive et contre prédictive. Les effets d'apprentissage observés dans la tâche de recherche sont donc de nature implicite.

Les résultats des Expériences 4 et 5 généralisent les effets d'indication contextuel sémantiques implicites obtenus dans les Expériences 2 et 3. Ils confirment que les effets d'apprentissage implicite observés à l'aide du paradigme d'indication contextuel peuvent être basés sur des propriétés catégorielles et sémantiques du contexte. Dans l'Expérience 6, nous avons testé dans quelle mesure une attention sélective portée sur les régularités prédictives est requise pour que se développe un effet d'indication contextuel.

EXPÉRIENCE 6

L'objectif était ici de tester si les effets d'indication contextuel basés sur des régularités relatives à l'appartenance catégorielle et sémantique du contexte requièrent une attention sélective. Pour ce faire, la procédure employée dans l'Expérience 4 a été couplée à celle utilisée par Jiang et Chun (2001). Le principe général était le même que dans l'Expérience 4, mais la cible (i.e. le nom d'un vêtement ou d'un bâtiment) était toujours écrite en vert pour la moitié des sujets et en rouge pour l'autre moitié des sujets. Dans un essai, la cible apparaissait parmi 9 mots contextuels affichés en vert et 9 mots contextuels affichés en rouge. Si la cible était rouge, le contexte rouge définissait le contexte attendu et le contexte vert définissait le contexte ignoré, et vice versa si la cible était verte. Les mots écrits en vert appartenaient tous à une même catégorie sémantique et les mots écrits en rouge appartenaient tous à une autre catégorie sémantique. La tâche de recherche comportait trois conditions expérimentales :

« prédictive attendue », « prédictive ignorée » et « non prédictive ». Dans un essai prédictif attendu, le contexte attendu (i.e. la catégorie sémantique des mots contextuels écrits dans la même couleur que la cible) était prédictif de la localisation de la cible, alors que le contexte ignoré (i.e. la catégorie sémantique des mots contextuels écrits dans la couleur différente de celle de la cible) n'était pas prédictif de la localisation de la cible. Dans un essai prédictif ignoré, le contexte ignoré prédisait la localisation de la cible, mais pas le contexte attendu. Dans la condition non prédictive, ni le contexte attendu, ni le contexte ignoré, ne prédisaient la localisation de la cible. La Figure 25 donne un exemple de deux essais prédictifs attendus et de deux essais prédictifs ignorés.



Figure 25. En haut: Exemple de deux essais prédictifs attendus pour l'association « contextes mammifères – cible à gauche ». Les mots cibles « BONNET » et « TAVERNE » apparaissent en vert parmi des noms de mammifères écrits en vert (catégorie prédictive attendue) et parmi des noms de minéraux/matériaux ou de fruits/légumes écrits en rouge (catégories non prédictives ignorées).

En bas : Exemple de deux essais prédictifs ignorés pour l'association « contextes oiseaux – cible à droite ». Les mots cibles « CHEMISE » et « FERME » apparaissent en vert parmi des noms d'oiseaux écrits en rouge (catégorie prédictive ignorée) et des noms de poissons/fruits de mer écrits en vert (catégories non prédictives attendues).

Afin de rendre les régularités plus saillantes, les localisations de la cible ont été restreintes aux extrémités gauche et droite de l’affichage. De plus, dans le but de minimiser l’emploi éventuel d’une stratégie consistant à limiter la recherche aux zones droite et gauche de l’affichage, des essais de remplissage ont été rajoutés. Dans les essais de remplissage, les localisations potentielles de la cible recouvraient l’ensemble de l’affichage (pour les localisations des cibles, cf. Figure 26).

		R					
C				R			C
C							C
		R					
						R	

Figure 26. Localisations potentielles de la cible utilisées dans les Expériences 6 et 7. C=localisations potentielles de la cible dans les essais tests. R= localisations potentielles de la cible dans les essais de remplissage.

Méthode

Participants

Trente-deux étudiants ou chercheurs français de l’Université de Marseille ont participé à l’expérimentation (17 femmes et 15 hommes), âgés en moyenne de 26ans (étendue des âges : 17-41 ans, $\sigma = 5$). Tous présentaient une acuité visuelle normale ou corrigée. Tous étaient volontaires et motivés par l’expérience. Aucun n’avait connaissance du but et méthode de cette expérience.

Dispositif

Le dispositif était le même que celui utilisé dans l’Expérience 4.

Stimuli

Les cibles. Les cibles étaient les mêmes que celles utilisées dans l’Expérience 4 (cf. Annexe 5a). De manière différente, la cible était écrite en vert (RGB : 0, 128, 36) pour la

moitié des sujets et en rouge (RGB :200, 50, 36) pour l'autre moitié des sujets. Les cibles pouvaient apparaître dans 8 localisations potentielles présentées Figure 26.

Les contextes. La cible apparaissait parmi 18 mots contextuels. La moitié des mots contextuels était écrite en vert et l'autre moitié en rouge. Les mots contextuels écrits en vert étaient tirés aléatoirement avec remise dans une liste catégorielle particulière et les mots écrits en rouge étaient tirés aléatoirement avec remise dans une liste d'une autre catégorie. Six catégories sémantiques ont été utilisées : mammifères, oiseaux, arbres/fleurs, fruits/légumes, poissons/fruits de mer et matériaux/minéraux. Les listes de mots utilisées au cours de cette expérience sont présentées en Annexe 5b. Parmi ces six catégories, deux constituaient les catégories prédictives, pour l'une la catégorie prédictive attendue et pour l'autre la catégorie prédictive ignorée. Les quatre autres catégories constituaient les catégories non prédictives, pour deux d'entre elles les catégories non prédictives attendues et pour les deux autres les catégories non prédictives ignorées. Une catégorie supplémentaire a été utilisée pour les essais de remplissage : fournitures de bureau.

Procédure

Tâche de recherche. Les participants avaient pour consigne de rechercher aussi rapidement et correctement que possible le nom d'un vêtement ou d'un bâtiment toujours écrit en vert pour la moitié des participants, et toujours écrit en rouge pour l'autre moitié des participants.

La tâche de recherche comportait 24 blocs de 16 essais, soit au total 384 essais. Chaque bloc comprenait 4 essais prédictifs attendus, 4 essais prédictifs ignorés, 4 essais non prédictifs et 4 essais de remplissage. Dans un essai prédictif attendu, les mots écrits dans la couleur attendue (e.g. vert si la cible était verte) appartenaient à la catégorie prédictive attendue (e.g. si le contexte attendu appartenait à la catégorie prédictive mammifère, la cible était toujours localisée à gauche de l'écran), alors que les mots écrits dans la couleur ignorée (rouge) appartenaient à une autre catégorie qui ne prédisait pas la localisation de la cible (i.e. une catégorie non prédictive ignorée). Dans un essai prédictif ignoré, les mots écrits dans la couleur ignorée (e.g. rouge si la cible était verte) appartenaient à la catégorie prédictive ignorée (e.g. si le contexte ignoré appartenait à la catégorie prédictive ignorée oiseaux, la cible était toujours localisée à droite de l'affichage), alors que les mots écrits dans la couleur attendue (vert) appartenaient à une catégorie qui ne prédisait pas la localisation de la cible (i.e. une catégorie non prédictive attendue). Dans les essais non prédictifs, les mots apparaissant dans la couleur attendue appartenaient à une catégorie non prédictive attendue et les mots

apparaissant dans la couleur ignorée appartenait à une catégorie non prédictive ignorée. Dans ces essais, la cible apparaissait dans les mêmes localisations potentielles que dans les essais prédictifs. Pour rendre la régularité plus saillante, les localisations ont été restreintes aux extrémités gauche et droite de l’affichage (Cf. Figure 26).

Les six catégories apparaissaient avec la même fréquence au sein d’un bloc. Les catégories attribuées à chaque condition expérimentale étaient contrebalancées entre les participants, ainsi que la localisation de la cible (gauche vs. droite) pour chaque catégorie prédictive attendue et ignorée. La couleur de la cible, vert ou rouge, était contrebalancée entre les participants. Dans les essais de remplissage, la cible apparaissait parmi des noms de fournitures de bureau et sa localisation potentielle recouvrait l’ensemble de l’affichage (Cf. Figure 26). La couleur vert vs. rouge de la cible était contrebalancée entre les participants.

Le déroulement de l’expérience était le même que celui de l’Expérience 4, hormis la consigne et la phase d’entraînement. La consigne informait les participants que les affichages de recherche étaient composés de mots verts et de mots rouges. Si la cible était verte, il leur était demandé de rechercher le nom d’un vêtement ou d’un bâtiment parmi les mots verts uniquement. Il leur était spécifié que la cible était toujours écrite en vert, jamais en rouge – et inversement si la cible était rouge. Dans les essais d’entraînement, les éléments contextuels appartenaient à la catégorie de remplissage.

Tâche de verbalisation puis de reconnaissance. Le principe était le même que celui des expériences précédentes. Mais de manière différente, la tâche de reconnaissance comportait 24 essais, dont 4 essais prédictifs attendus, 4 essais contre prédictifs attendus, 4 essais prédictifs ignorés, 4 essais contre prédictifs ignorés et 8 essais non prédictifs. Les essais contre prédictifs attendus et ignorés étaient construits selon le même principe que les essais contre prédictifs utilisés dans l’Expérience 4. La consigne était la même que dans l’Expérience 4.

Résultats

Tâche de recherche

Les TR obtenus sur les essais de remplissage ont été éliminés des analyses. Dans chacune des trois conditions expérimentales (prédictive attendue, prédictive ignorée, et non prédictive), les taux d’erreurs étaient inférieurs à 1.43%. Les TR sur les erreurs et les TR sur les réponses correctes supérieurs ou inférieurs à la « moyenne + 3 écart-types » (0.97% des

réponses correctes) ont été supprimés des analyses. Les blocs d'essais ont été regroupés en 6 époques. Ces résultats sont illustrés Figure 27. Une ANOVA à mesures répétées a été effectuée avec les facteurs, condition (prédictive attendue, prédictive ignorée et non prédictive) et époque (1-6). Elle montre un effet du facteur époque $F(5,175) = 69.822$, $p < .001$, mais pas d'effet du facteur condition, $F(2,70) = 1.765$, $p = .18$, et pas d'interaction entre les deux facteurs, $F(10,350) = 1.473$, $p = .15$. Cependant, une ANOVA à mesures répétées conduite avec le facteur condition (prédictive attendue et non prédictive) et époque (1-6) révèle un effet du facteur condition, $F(5,175) = 6.312$, $p < .05$, un effet du facteur époque, $F(1,35) = 56.216$, $p < .001$, et une interaction entre les deux facteurs, $F(5,175) = 2.332$, $p < .05$. Des comparaisons planifiées réalisées avec le facteur condition (prédictive attendue vs. non prédictive) à chaque époque indiquent un effet du facteur condition à l'époque 2, $F(1,35) = 5.93$, $p < .05$, à l'époque 3, $F(1,35) = 7.12$, $p < .05$, et à l'époque 5, $F(1,35) = 8.10$, $p < .01$. Les autres comparaisons (prédictive ignorée vs. non prédictive et prédictive attendue vs. prédictive ignorée) n'indiquent pas d'effet du facteur condition (tous les $F < 1$), et pas d'interaction entre les facteurs époque et condition, $F(5, 175) = 1.218$, $p = .301$ et $F(5, 175) < 1$, respectivement.

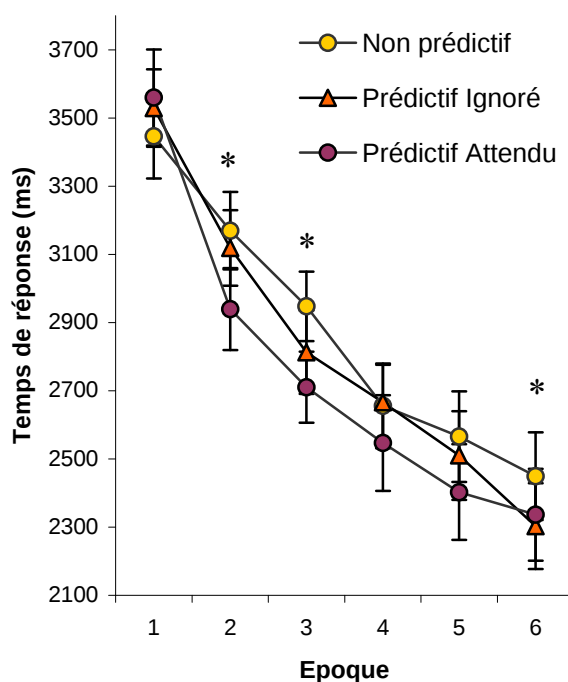


Figure 27. Temps de réponse dans les conditions prédictive attendue, prédictive ignorée et non prédictive. Les barres d'erreurs correspondent aux SEM (N=36).

Les fonctions puissances ajustées aux données sur les TR pour les trois conditions expérimentales étaient les suivantes : TR (prédictive attendue) = $1429 + 2123N^{-0.475}$ ($r^2=100$),

TR (prédictive ignorée) = $1429 + 2154^{-0.428}$ ($r^2=97$), et TR (non prédictive) = $1429 + 2089^{-0.359}$ ($r^2=95$).

Tâche de verbalisation puis de reconnaissance

Aucun sujet n'a répondu avoir remarqué les régularités attendues ou avoir cherché des régularités ou cherché à mémoriser le matériel. Tous les sujets ont été exposés à la tâche de reconnaissance. Une ANOVA à mesures répétées a été réalisée sur les pourcentages de réponses *déjà vu*. Elle ne montre pas de différence entre les pourcentages de réponse *déjà vu* obtenus dans les conditions prédictive attendue (53%), contre prédictive attendue (55%), prédictive ignorée (60%), contre prédictive ignorée (52%), et non prédictive (51%) $F(4,140) < 1$. Cependant, les comparaisons deux à deux montrent que les conditions prédictive ignorée et non prédictive étaient différentes $F(1,35) = 4.896$, $p < .05$ et que les conditions prédictive ignorée et contre prédictive ignorée étaient différentes de manière marginalement significative $F(1,35) = 3.500$, $p = .068$.

Discussion

Un effet d'indiciage contextuel est observé quand l'appartenance catégorielle et sémantique du contexte attendu est prédictive de la localisation de la cible, alors que le contexte ignoré est non prédictif. Néanmoins, les résultats ne montrent pas de différence significative entre les conditions prédictive ignorée et non prédictive. A noter cependant que les performances obtenues dans la condition prédictive ignorée sont intermédiaires entre celles obtenues dans la condition prédictive attendue et celle obtenues dans la condition non prédictive. Les taux d'apprentissage (0.475 pour la condition prédictive attendue, 0.428 pour la condition prédictive ignorée et 0.359 pour la condition non prédictive) supportent ce patron de résultat. Mais les analyses statistiques ne permettent pas de conclure à un effet d'indiciage contextuel dans la condition prédictive ignorée.

Les résultats relatifs à la condition prédictive attendue confirment les effets d'indiciage obtenus dans l'Expérience 4, mais le bénéfice obtenu dans la condition prédictive attendue semble plus robuste. Cette différence pourrait être due à certains changements dans la procédure expérimentale. D'une part, dans le but de rendre la régularité plus « franche » dans les essais effectifs, les localisations de la cible ont été réduites aux extrémités gauche et droite.

Ceci pourrait avoir optimisé le déploiement de l'attention vers la localisation de la cible. Mais surtout, des essais de remplissage ont été rajoutés afin que les localisations potentielles de la cible recouvrent l'ensemble de l'affichage. Ces essais de remplissage pourraient avoir minimisé l'emploi d'une stratégie consistant à limiter la recherche aux extrémités gauche et droite de l'écran. Les éléments contextuels effectivement traités demeureraient nombreux tout au long de la tâche, ce qui pourrait expliquer qu'un effet d'indication contextuel soit observé jusqu'à la dernière époque (cf. Chun & Jiang, 1998 ; Schneider & Shiffrin, 1977). De plus, même si certaines catégories prédictives attendues semblent plus favorables à la mise en évidence d'un effet d'indication contextuel (e.g. mammifères et oiseaux), les six catégories utilisées dans cette expérience produisaient un bénéfice sur les TR par rapport à la condition non prédictive (pour les résultats détaillées de chacune des catégories prédictives attendues et ignorées, cf. Annexe 6c).

Les résultats de la tâche de verbalisation puis de reconnaissance montrent une fois de plus le caractère implicite de l'apprentissage mis en évidence dans la condition prédictive attendue. D'une part, aucun des participants n'a rapporté verbalement les régularités manipulées, et d'autre part, la tâche de reconnaissance ne fait pas apparaître de différence entre les conditions prédictive attendue, non prédictive, et contre prédictive attendue. Cependant, de manière fort surprenante, les essais prédictifs ignorés étaient jugés plus familiers que les essais non prédictifs, et marginalement plus familiers que les essais contre prédictif ignorés, alors qu'ils ne conduisaient pas un effet d'indication contextuel significatif. Ce résultat laisse penser qu'il y a quand même eu un apprentissage sur les essais prédictifs ignorés. Il paraît néanmoins difficile de parler d'un apprentissage explicite compte tenu du fait que les participants étaient incapables de rapporter les catégories contextuelles apparaissant dans la couleur qu'ils devaient ignorer.

Les résultats obtenus suggèrent donc que l'attention sélective est nécessaire, ou tout du moins favorise l'indication contextuel basé sur les régularités manipulées dans cette expérience. En effet, les contextes prédictifs attendus entraînaient clairement des effets d'indication contextuel, alors que les contextes prédictifs ignorés n'en entraînent pas ou en entraînent moins – la condition prédictive ignorée semblant se situer entre la condition prédictive attendue et non prédictive. Les résultats de la tâche de reconnaissance suggèrent néanmoins qu'un apprentissage se serait développé sur les essais prédictifs ignorés. Les travaux de Jiang et Leung (2005) peuvent permettre d'expliquer ces résultats au premier abord contradictoires.

En effet, ces derniers ont montré (dans une tâche d'indication contextuel spatial) que si l'expression de l'apprentissage implicite est bien dépendante de l'attention sélective, ceci n'empêche pas qu'il puisse exister un apprentissage latent sur les essais prédictifs ignorés. En effet, les résultats de leur recherche ne mettent pas en évidence d'effets d'indication contextuel dans les contextes prédictifs ignorés. Mais en implémentant une phase de transfert dans laquelle les contextes prédictifs précédemment ignorés devenaient subitement des contextes prédictifs attendus, les auteurs ont montré que la performance était immédiatement facilitée. Jiang et Leung ont conclu à l'existence d'un apprentissage latent dans les essais prédictifs ignorés. Celui-ci serait indépendant de l'attention sélective allouée aux contextes, et ne se manifesterait pas directement dans les effets d'indication contextuel. En revanche, lorsque les contextes prédictifs attendus devenaient subitement ignorés, le bénéfice disparaissait aussitôt. Ce résultat suggère que l'attention est nécessaire à l'expression de l'apprentissage implicite. L'Expérience 7 a été conduite dans le but de tester ce phénomène dans le cadre d'un apprentissage implicite sémantique.

EXPÉRIENCE 7

Au cours de l'Expérience 7, nous avons testé le rôle de l'attention sélective sur l'apprentissage des régularités contextuelles catégorielles et sémantiques, et son rôle sur l'expression de la connaissance résultante. Comme dans les travaux de Jiang et Leung (2005), une phase de transfert a été implémentée. Les catégories contextuelles attendues au début de la tâche de recherche devenaient subitement ignorées au moment du transfert (la couleur des éléments contextuels étant inversée). Inversement, les catégories ignorées devenaient subitement attendues. A noter que la couleur de la cible ne changeait pas durant la phase de transfert. Etant donné que l'époque 3 semble être l'époque la plus favorable à la mise en évidence d'un effet d'apprentissage, la phase de transfert a été implémentée à l'époque 3. Par exemple, si la catégorie mammifère définissait des contextes attendus (e.g. vert si la cible était verte) aux époques 1 et 2, elle définissait les contextes ignorés (e.g. rouge) aux époques 3 et 4. Tout au long de la tâche de recherche visuelle, les mêmes catégories prédisaient la localisation de la cible (e.g. « mammifères-cible à droite » et « oiseaux-cible à gauche »), mais les couleurs dans lesquelles elles apparaissaient étaient inversées à partir de l'époque 3.

Sous l'hypothèse que l'attention sélective est nécessaire pour l'apprentissage des régularités, on ne devrait pas observer d'indication contextuel pour les contextes prédictifs subitement attendus à l'époque 3, succédant à deux époques où ces contextes étaient ignorés (condition prédictive ignorée puis attendue). Sous l'hypothèse que l'attention sélective est nécessaire à l'expression de la connaissance résultante, on ne devrait pas obtenir d'effet d'indication contextuel pour les contextes prédictifs subitement ignorés à l'époque 3, succédant à deux époques au cours desquelles ils étaient attendus (condition prédictive attendue puis ignorée).

Parallèlement, il s'agissait de s'assurer que les effets testés reposent bien sur des régularités relatives à l'appartenance catégorielle et sémantique du contexte. Ainsi, suivant le même principe que celui employé dans l'Expérience 5, les mots utilisés aux époques 3 et 4 étaient différents de ceux utilisés aux époques 1 et 2.

Méthode

Participants

Trente six étudiants ou chercheurs français de l'Université d'Aix-Marseille ont participé à cette expérience (19 femmes et 17 hommes), âgés en moyenne de 26 ans (étendue des âges : 17-42, $\sigma = 6$). Tous avaient une vue normale ou corrigée et aucun n'avait connaissance des raisons et méthodes de cette expérience.

Dispositif

Le dispositif était le même que celui de l'Expérience 4.

Stimuli

Le matériel et les catégories étaient les mêmes que ceux utilisés dans l'Expérience 6, excepté que deux listes de 13 mots ont été établies pour chacune des 7 catégories. Ces listes sont présentées en Annexe 5b.

Procédure

Tâche de recherche visuelle. La consigne et le principe général étaient les mêmes que ceux de l'Expérience 6, mais la tâche comportait 16 blocs de 16 essais et comprenait une

phase de transfert. Le transfert portait à la fois sur les mots (cf. Expérience 5) et sur la couleur des contextes. La tâche comprenait 4 époques, regroupant chacune 4 blocs successifs d'essais. Durant les 8 premiers blocs (époques 1 et 2) les mots contextuels étaient tirés dans une série de listes (e.g. listes 1). Durant les 8 derniers blocs d'essais (époques 3 et 4), les mots contextuels étaient tirés dans l'autre série de listes (e.g. listes 2). L'ordre de présentation de ces listes (e.g. listes 1 puis listes 2) a été contrebalancé entre les participants. Parallèlement, les couleurs des contextes étaient inversées par rapport aux époques 1 et 2, de telle manière que les contextes ignorés devenaient attendus et les contextes attendus devenaient ignorés. Ainsi, la catégorie définissant les contextes prédictifs attendus aux époques 1 et 2, définissait les contextes prédictifs ignorés aux époques 3 et 4, et, inversement, la catégorie définissant les contextes prédictifs ignorés devenait attendue. Les deux catégories non prédictives attendues, devenaient les catégories non prédictives ignorées, et les deux catégories non prédictives ignorées devenaient les deux catégories non prédictives attendues. Il y avait donc trois conditions expérimentales : « prédictive attendue puis ignorée », « prédictive ignorée puis attendue », et « non prédictive ». Comme dans l'expérience précédente, les contextes attribués à chaque condition expérimentale étaient contrebalancés entre les participants. Le déroulement de l'expérience était le même que celui de l'Expérience 6.

La tâche de verbalisation puis de reconnaissance. La tâche était la même que celle utilisée dans l'Expérience 6, mais avec les contextes utilisés avant le transfert (époques 1 et 2) et ceux utilisés après le transfert (époques 3 et 4). La tâche de reconnaissance comportait donc 48 essais.

Résultats

Tâche de recherche

Les réponses sur les essais de remplissage ont été éliminées des analyses. Les taux d'erreurs étaient inférieurs à 1.46%, et donc comme précédemment ils ne seront pas analysés plus en détail. Les TR sur les erreurs et les TR supérieurs ou inférieurs à la « moyenne + 3 écart-types » (1.28% des réponses correctes) ont été supprimés des analyses. Les TR sont présentés Figure 28. Une ANOVA à mesures répétées réalisée avec les facteurs intra-sujets, condition (prédictive attendue puis ignorée, prédictive ignorée puis attendue et non prédictive) et époque (1-4) montre un effet du facteur époque, $F(3, 105) = 32.971$, $p < .001$, mais pas d'effet du facteur condition, $F(2, 70) < 1$, et pas d'interaction entre les deux facteurs, $F(6,$

210) = 1.165, $p = .326$. Cependant, les analyses planifiées indiquent une différence significative entre les conditions prédictive ignorée puis attendue et non prédictive $F(1,35)=5.634$, $p < .05$ à l'époque 3, alors qu'aucune différence n'était observée entre les conditions prédictive attendue puis ignorée et non prédictive $F(1,35) < 1$. Comme dans l'Expérience 5, pour s'assurer que l'effet observé à l'époque 3 ne tenait pas à un nouvel apprentissage, mais était bien la conséquence d'un apprentissage développé aux époques 1 et 2, l'effet du facteur « condition » a été analysé en se limitant à la première moitié des essais de l'époque 3. Les résultats montrent que dès la première moitié de l'époque 3 la condition prédictive ignorée puis attendue (3217ms) différait significativement de la condition non prédictive (3606ms) $F(1,35)=4.34$, $p < .05$. Ils ne montrent aucune différence entre les conditions prédictive attendue puis ignorée (3572ms) et non prédictive (3606ms), $F(1,35) < 1$.

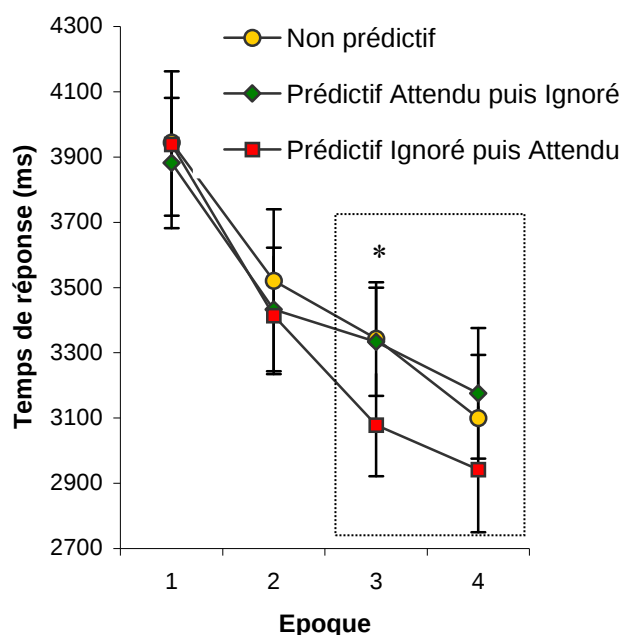


Figure 28. Temps de réponse dans les conditions prédictive attendue puis ignorée, prédictive ignorée puis attendue et non prédictive. Les barres d'erreurs correspondent aux SEM (N=36). Les époques encadrées correspondent au transfert.

Tâche de verbalisation puis de reconnaissance

Aucun sujet n'a répondu avoir remarqué les régularités attendues. Aucun n'a rapporté avoir cherché des régularités ou cherché à mémoriser le matériel. Tous les sujets ont été exposés à la tâche de reconnaissance. Les pourcentages de réponse *déjà vu* ont été analysés à l'aide de deux ANOVAs à mesures répétées. La première ANOVA compare les réponses sur

les contextes utilisés au cours des époques 1 et 2 de la tâche de recherche. Les résultats ne montrent pas de différence entre les pourcentages de réponses *déjà vu* obtenus dans les conditions prédictive attendue (44%), contre prédictive attendue (44%), prédictive ignorée (44%), contre prédictive ignoré (46%) et non prédictive (45%) utilisées avant le transfert $F(4, 140) < 1$. La deuxième ANOVA compare les réponses obtenues sur les contextes utilisés au cours des époques 3 et 4. Elle ne montre pas de différence globale entre les conditions prédictive attendue (57%), contre prédictive attendue (60%), prédictive ignorée (61%), contre prédictive ignoré (50%) et non prédictive (54%) utilisées après le transfert $F(4, 140) = 1.61$, $p = .174$. Néanmoins, les comparaisons planifiées montrent que la condition prédictive ignorée diffère significativement de la condition contre prédictive ignorée $F(1, 35) = 7.600$, $p < .01$ et marginalement de la condition non prédictive $F(1, 35) = 3.56$, $p = 0.068$. Elles ne mettent pas en évidence de différence entre les autres conditions expérimentales.

Discussion

Les principaux résultats de cette expérience concernent la phase de transfert. Ils révèlent un effet d'indilage contextuel dans la condition prédictive auparavant ignorée lorsqu'elle devient subitement attendue et aucun effet d'indilage contextuel dans la condition auparavant attendue lorsqu'elle devient subitement ignorée. L'analyse détaillée des résultats obtenus à l'époque 3 montre que le bénéfice global observé dans la condition subitement attendue est significatif dès les premiers essais du transfert. Ceci montre que l'apprentissage mis en évidence au moment du transfert s'est développé au cours des époques antérieures, i.e. aux époques 1 et 2. Ce résultat suggère donc que les régularités ont été apprises avant le transfert, lorsque les contextes prédictifs étaient ignorés des participants. Un bénéfice par rapport à la condition non prédictive est observé à l'époque 3 pour toutes les catégories prédictives subitement attendues (574ms pour les mammifères, 338 ms pour les arbres/fleurs, 294 ms pour les oiseaux, 218ms pour les fruits/légumes et 199ms pour les minéraux/matériaux), sauf pour la catégorie poissons/fruits de mer (-45ms) (pour les résultats détaillés relatifs à chacune des catégories prédictives, cf. Annexe 6d). En revanche, la condition auparavant attendue n'entraîne aucun effet d'indilage contextuel lorsqu'elle devient subitement ignorée. Ce résultat suggère que l'apprentissage ne s'exprime pas sur les contextes ignorés. L'utilisation de mots nouveaux au cours de la phase de transfert atteste que ces résultats reposent sans équivoque sur les propriétés sémantiques des contextes. Cette expérience suggère donc qu'un

apprentissage latent s'est développé sur les contextes ignorés au cours des époques 1 et 2. Mais pour que cet apprentissage s'exprime, une attention sélective directement focalisée sur les régularités contextuelles serait requise.

Les résultats de la tâche de reconnaissance montrent que les essais prédictifs attendus dans la deuxième partie de la tâche de recherche n'étaient pas jugés plus familiers que les essais non prédictifs ou contre prédictifs attendus. Par contre, et de manière congruente avec les résultats de l'Expérience 6, les essais prédictifs ignorés utilisés dans la deuxième partie de la recherche étaient quant à eux jugés plus familiers que les essais contre prédictifs ignorés et jugés plus familiers, de manière marginalement significative, que les essais non prédictifs. Ce résultat suggère que les performances d'indilage contextuel obtenues dans la tâche de recherche et les performances observées dans la tâche de reconnaissance ne sont pas nécessairement corrélées.

Ainsi l'apprentissage implicite des régularités sémantiques ne nécessiterait pas que l'attention soit directement focalisée sur les régularités. Par contre, l'expression de cet apprentissage requerrait une attention sélective portée sur les régularités sémantiques. Ces résultats corroborent ceux observés par Jiang et Leung (2005) à partir de configurations spatiales et étendent cet aspect de l'apprentissage implicite aux régularités sémantiques.

DISCUSSION GÉNÉRALE

Au cours de l'analyse visuelle, le système cognitif a la capacité d'apprendre implicitement des régularités présentes dans le matériel (Chun & Jiang, 1998). L'étude présente apporte des précisions nouvelles sur les mécanismes d'apprentissage implicite impliqués dans une tâche de recherche visuelle. Les apports majeurs des expériences réalisées tiennent en trois points. D'une part, des effets d'indilage contextuel basés sur des régularités catégorielles et sémantiques du contexte ont été mis en évidence. D'autre part, ces effets résultent d'un apprentissage implicite. Et enfin, si une attention sélective portée sur les régularités sémantiques semble nécessaire à la mise en évidence d'un effet d'indilage contextuel, il s'avère qu'un apprentissage latent peut néanmoins se développer indépendamment d'une attention focalisée sur ces régularités sémantiques.

Indiçage contextuel basé sur des régularités sémantiques

Les quatre expériences réalisées au cours de cette étude ont permis de mettre en évidence des effets d'indiçage contextuel basés sur des propriétés sémantiques de contextes composés de mots. Dans les Expériences 4 et 6, la recherche visuelle était facilitée lorsque l'appartenance catégorielle et sémantique du contexte était prédictive de la localisation de la cible. Cependant, les mots contextuels étant tirés dans des listes spécifiques, ces résultats peuvent aussi bien être interprétés en termes d'apprentissage basé sur les propriétés catégorielles et sémantiques du contexte, qu'en termes d'apprentissage spécifique. Pour pallier cette critique, les Expériences 5 et 7 comportaient une phase de transfert, au cours de laquelle les mots étaient tirés dans des listes différentes de celles utilisées au début de l'expérience. Un effet d'indiçage contextuel était obtenu dès les premiers essais du transfert, indiquant que l'apprentissage repose, sans équivoque, sur l'appartenance catégorielle et sémantique du contexte.

Par ailleurs, lorsque nous analysons les effets d'indiçage relatifs à chaque catégorie prédictive, il s'avère que certaines catégories sont plus favorables que d'autres au développement d'un apprentissage. De manière globale, il semblerait que ce soit les catégories les plus homogènes (ici, mammifères et oiseaux), indépendamment de la fréquence d'usage des mots composant la catégorie ou de la difficulté préalable des contextes, qui seraient les plus à même de susciter un effet d'indiçage contextuel sémantique. Par exemple, il existe une forte relation sémantique entre brebis, béliet, agneau et chèvre ou entre bison, buffle et taureau. On peut faire l'hypothèse que l'un des facteurs qui détermine l'indiçage contextuel sémantique est la force des liens sémantiques entre les différents items d'une catégorie. Il n'existe cependant, à notre connaissance, aucune base de données en français permettant de dégager des indices psycholinguistiques susceptibles de valider cette hypothèse. A l'appui de tels résultats, les travaux de Moores, Laiti et Chelazzi (2003) suggèrent que les associations sémantiques en MLT, via l'amorçage conceptuel, peuvent influencer l'allocation attentionnelle lors d'une recherche visuelle.

Indiçage contextuel sémantique et indiçage contextuel spatial

Dans quelle mesure les taux d'apprentissage observés avec des contextes sémantiques sont-ils comparables aux taux d'apprentissage observés avec des contextes spatiaux ? Compte tenu des importantes différences entre la méthode utilisée dans les expériences classiques et celle utilisée dans les expériences présentées ici, il est difficile de comparer la force de nos effets avec ceux observés dans la littérature, à la fois en terme de vitesse d'apprentissage

(exprimé par la valeur de c) et en terme d'amplitude (bénéfice dans la condition prédictive par rapport à la condition non prédictive). Si l'on s'intéresse aux effets d'indication en termes d'amplitude, on peut constater que le bénéfice moyen maximal est supérieur dans nos expériences (environ 200ms) que dans les expériences classiques (environ 100ms). Cependant, le niveau initial des performances étant également beaucoup plus élevé, il paraît nécessaire de pondérer ces valeurs par rapport aux performances initiales. Dans les Expériences 4 et 6, l'amélioration était de 5-6% supérieure en condition prédictive par rapport à la condition non prédictive (e.g. pour l'Expérience 6 : $(2500-2300)/3500$), contre 7% environ dans les tâches d'indication spatiale (e.g. pour l'Expérience 1 de Chun & Jiang, 1998 : $(880-800)/1075$). Le bénéfice apporté est donc relativement comparable. Compte tenu des importantes différences entre le plan expérimental utilisé dans nos expériences et celui utilisé dans les expériences classiques d'indication contextuel, il nous paraît peu approprié de comparer la vitesse d'apprentissage en nous appuyant sur les taux d'apprentissage dégagés par l'estimation des fonctions ajustées aux données. En effet, dans les tâches classiques, une même configuration ancienne est présentée une seule fois dans un bloc, alors que dans notre étude, quatre contextes prédictifs définis par la même catégorie sémantique apparaissent au cours de chaque bloc. Le nombre d'expositions nécessaires aux essais prédictifs pour observer des effets d'indication contextuel basés sur des aspects sémantiques (une vingtaine en moyenne) semble être largement supérieur à celui observé dans la littérature (la plupart du temps autour de cinq répétitions d'une configuration ancienne). L'émergence de ces effets semble par ailleurs varier selon les catégories. Par exemple, la catégorie mammifères semble conduire à des effets d'indication plus précoces que la catégorie arbres/fleurs. Aussi, bien que les effets d'indication contextuel sémantique, impliquant un processus de catégorisation et de généralisation, semblent émerger plus tardivement que les effets d'indication spatiaux basés sur des propriétés spécifiques, il nous semble délicat d'affirmer de manière catégorique ce point de vue, compte tenu des différences importantes dans les méthodes utilisées.

L'indication contextuel sémantique résulte d'un apprentissage implicite

Un aspect crucial des effets d'indication contextuel observés est qu'ils sont le fruit d'un apprentissage implicite. L'apprentissage implicite renvoie à un processus d'adaptation par lequel le comportement d'un individu devient, de manière incidente, sensible à une structure, sans que la connaissance qui en résulte ne soit rapportable verbalement, voire accessible à la conscience (Cleeremans et al., 1998 ; Dienes & Berry, 1997 ; Reber, 1989). Dans les expériences conduites au cours de l'étude présente, à la fois le processus d'acquisition et les

connaissances qui en résultent étaient implicites. D'une part, à l'issue de la tâche de recherche, aucun des participants n'a rapporté avoir cherché des régularités ou avoir cherché à mémoriser le matériel expérimental. D'autre part, aucun des participants n'était à même de verbaliser les régularités contextuelles facilitant leur recherche, même lorsque l'expérimentateur sollicitait leur report. De plus, dans les tâches de reconnaissance, les participants ne jugeaient pas plus familiers les essais prédictifs conduisant à des effets d'indication contextuel, que les essais non prédictifs ou contre prédictifs. Ces données indiquent donc que les effets d'apprentissage observés dans les différentes tâches de recherche, via les effets d'indication contextuel sont de nature implicite. Non seulement les participants n'ont pas délibérément appris les régularités, mais en plus, ces connaissances ne sont pas accessibles consciemment.

Ces résultats prolongent ceux que nous avons obtenus dans l'Etude 1 en utilisant des affichages numériques dans lesquels les régularités contextuelles reposaient sur la propriété de parité/impairité des nombres utilisés (Goujon, Didierjean, & Marmèche, 2007). Ces résultats sont ici généralisés à d'autres catégories sémantiques à l'aide d'un matériel lexical. Ils montrent que les effets d'indication contextuel basés sur des régularités sémantiques ne sont pas nécessairement associés à une prise de conscience de ces régularités. La prise de conscience associée à l'apprentissage des régularités dans des scènes du monde réel (Brockmole & Henderson, 2006a, 2006b) ne tient donc pas à la simple présence d'indices sémantiques. Cette étude offre des arguments empiriques supplémentaires en faveur de l'existence d'un apprentissage implicite conceptuel. L'apprentissage implicite ne se limite donc pas à la prise en compte des traits spécifiques et notamment perceptifs du contexte, mais il peut aussi reposer sur des propriétés plus conceptuelles, en l'occurrence l'appartenance catégorielle et sémantique des éléments contextuels. A noter cependant que dans notre étude, il ne s'agit pas de l'apprentissage d'une connaissance abstraite nouvelle comme cela a pu être défendu dans la littérature sur l'apprentissage implicite (cf. Reber, 1989) mais de l'association entre une connaissance conceptuelle déjà en mémoire (e.g. le concept de « mammifère ») et une position de la cible.

Indication Contextuel et Attention Sélective

Les Expériences 6 et 7 montrent qu'une attention sélective portée sur les éléments contextuels est nécessaire à l'expression de l'apprentissage implicite, mais qu'un apprentissage latent peut néanmoins se développer sur les contextes ignorés. Dans l'Expérience 6 une attention sélective portée sur les régularités contextuelles était nécessaire à

la mise en évidence d'un effet d'indication contextuel. Lorsque la catégorie prédictive apparaissait dans une couleur ignorée par les participants (i.e. une couleur différente de celle de la cible), elle n'entraînait pas de bénéfice significatif sur les TR, alors qu'un effet d'indication contextuel était obtenu lorsque la catégorie prédictive apparaissait dans une couleur attendue par les participants (i.e. de la même couleur que celle de la cible). Dans l'Expérience 7, afin d'examiner de manière distincte le rôle de l'attention sélective, sur l'apprentissage, et sur son expression, une phase de transfert a été implémentée au cours de la tâche de recherche (cf. Jiang & Leung, 2005). Les catégories auparavant attendues devenaient ignorées durant le transfert et les catégories auparavant ignorées devenaient attendues. Les données montrent que les contextes prédictifs auparavant attendus n'entraînaient pas d'effet d'indication contextuel lorsqu'ils devenaient subitement ignorés. Ce résultat suggère qu'une attention sélective est nécessaire à l'expression de la connaissance. En revanche, les contextes prédictifs auparavant ignorés entraînaient un effet d'indication immédiat lorsqu'ils devenaient subitement attendus. Les participants auraient donc appris des régularités relatives aux contextes ignorés, mais pour qu'il y ait expression de la connaissance résultante, il semble nécessaire qu'une attention sélective soit portée sur les éléments contextuels. L'utilisation de mots nouveaux au cours de la phase de transfert atteste que ces effets reposent sans équivoque sur les régularités catégorielles et sémantiques du contexte.

Ainsi cette étude suggère que l'apprentissage implicite des régularités sémantiques peut se développer de manière latente, alors que l'attention n'est pas directement focalisée sur ces régularités. En outre, pour que la connaissance résultante s'exprime dans une tâche d'indication contextuel, il apparaît nécessaire que les régularités prédictives soient focalisées par l'attention visuelle. Si ces résultats sont consistants avec ceux obtenus par Jiang et Leung (2005) avec des contextes spatiaux, ils n'en sont pas moins inattendus compte tenu du niveau de profondeur des traitements à réaliser pour dégager les régularités sémantiques catégorielles.

L'étude présente montre comment le guidage de l'attention est facilité par des connaissances implicites relatives à des régularités sémantiques de l'environnement. De plus, elle montre que ces connaissances inaccessibles à la conscience peuvent être acquises sans que l'attention sélective ne soit focalisée sur les aspects sémantiques des éléments présents dans la scène. Le système visuo-cognitif serait ainsi capable d'encoder de manière latente des informations perceptives (Jiang & Leung, 2005), mais également des informations sémantiques de l'environnement. Des connaissances implicites sur les régularités sémantiques

de l'environnement pourraient ainsi s'inscrire dans la construction des schémas de scène et contribuer à rendre compte de l'adaptabilité des traitements perceptifs, en dépit de la pauvreté de nos représentations conscientes.

CHAPITRE VI

LES SCHÉMAS DE SCENE : SUPPORT DE L'ADAPTABILITE DE LA CONDUITE ET SOURCE DE DYSFONCTIONNEMENTS

Les Etudes 1 et 2 montrent comment l'acquisition de connaissances implicites relatives aux régularités contextuelles optimise le déploiement de l'attention au sein d'une image. L'objectif de l'Etude 3 était d'étudier dans quelle mesure les régularités de l'environnement influencent le comportement dans une activité écologique, telle que la conduite automobile.

Il est largement reconnu que les informations perceptives issues de l'environnement routier interagissent en permanence avec les connaissances préalables stockées en MLT du conducteur, pour orienter l'attention sélective vers les aspects pertinents de l'environnement. Ces connaissances sont supposées être organisées hiérarchiquement sous forme de schémas ou programmes qui guident les fonctions visuo-motrices (e.g., Michon, 1985 ; Summala & Räsänen, 2000). La perception, comme l'action, est organisée et guidée par des attentes qui peuvent être confirmées ou déçues selon leurs conséquences (Neisser, 1976). Dans une étude de Shinoda, Hayhoe et Shrivstava (2000), la perception semblait dépendre d'une recherche active en accord avec une planification guidée à la fois par les objectifs de l'observateur et par des probabilités apprises sur l'environnement. En manipulant la consigne et la localisation de panneaux stop, les auteurs ont montré que les sujets détectaient rarement les panneaux quand ils étaient soumis à une condition de « passager » et encore plus rarement quand ces panneaux

étaient situés ailleurs qu'au niveau d'intersections. De manière convergente, Theeuwes & Hagenzieker (1993) ont montré que l'attention des conducteurs était essentiellement portée sur les localisations du trafic où l'information pertinente est attendue. Ces schémas se développeraient en accord avec des feedbacks sur leur adaptabilité.

Si les schémas de scène sont certainement le fondement de l'efficacité des traitements perceptifs impliqués dans la conduite, ils sont également candidats pour rendre compte de certaines défaillances fonctionnelles. Theeuwes et Hagenzieker (1993) ont montré, chez des conducteurs expérimentés, une détérioration considérable des performances quand les stimuli apparaissent dans des localisations non prédictibles. Les conducteurs anticipent en permanence les événements à venir. Pourtant, tout ne peut être anticipé. Grâce à la recherche en mémoire de situations similaires rencontrées auparavant, le système cognitif simplifie la richesse et la complexité de la situation en cours (Reason, 1993). Mais cette simplification de l'environnement rend les réponses cognitives très dépendantes des régularités observées dans le passé. Les événements imprévus peuvent parfois conduire à des erreurs dans l'organisation des traitements perceptifs (Hayhoe, 2000). Les études détaillées d'accidents (Van Elslande et al., 1997) révèlent ainsi de nombreuses défaillances liées à ce fonctionnement intrinsèque du système cognitif. La récupération en mémoire de schémas, liée à des attentes sur la situation présente, constitue une classe majeure de dysfonctionnement. Un autre type de défaillance dû à la récupération de schémas de scène tient à leurs conséquences sur le comportement du conducteur. Par exemple, un sur-apprentissage peut conduire à une « minimisation des exigences de la tâche » (Van Elslande, Alberton, Nachtergaële, & Blanchet, 1997). Cette sur-expérience de la situation amène parfois le conducteur à adopter une conduite « en mode automatique » conduisant à une « absence de recherche active des informations pertinentes » ou à une « saisie d'information sommaire et/ou précipitée » (Van Elslande et al., 1997). Dès lors qu'un événement imprévu surgit, le conducteur se trouve dans l'incapacité d'identifier à temps la situation pour la mise en œuvre d'une stratégie de conduite adaptée à cette situation.

L'origine de ces défaillances a été essentiellement déterminée sur la base des recueils verbaux de conducteurs suite à un accident. De telles défaillances peuvent-elles être mises en évidence dans des situations expérimentales? L'expérience réalisée au cours de cette étude visait dans un premier temps à tester dans quelle mesure des régularités contextuelles systématiquement associées à une même action peuvent faciliter la prise de décision en

conduite. Dans un deuxième temps, l'objectif était de tester l'hypothèse que les schémas de scène peuvent être source de dysfonctionnements. Pour ce faire, le principe expérimental était de tester si une exposition répétée à des régularités contextuelles prédictives de la réponse à donner facilite la prise de décision, et de tester si l'altération de ces régularités entraîne une dégradation des performances. Afin de rendre la tâche plus proche d'une recherche visuelle mise en œuvre dans la conduite, les participants exposés à des photographies de situations routières, étaient invités à juger le plus rapidement possible la situation en vue d'une action : « continuer » vs « s'arrêter ». Les participants devaient ainsi s'appuyer sur les éléments routiers présents dans la scène pour réaliser la tâche. Donc, contrairement aux travaux classiques d'indication contextuel, cette expérience tenait compte des connaissances pragmatiques préalables des participants dans la construction de nouvelles attentes.

Comme dans les études précédentes, les participants étaient exposés à une série de blocs d'essais prédictifs et non prédictifs. Les essais étaient définis par des photographies de scènes routières décrivant une situation de conduite impliquant une prise de décision (i.e. feu, passage-piétons ou rond-point). Des carrefours différents ont été sélectionnés pour les essais prédictifs et non prédictifs. Les participants étaient exposés aux mêmes carrefours au cours de chaque bloc, aussi bien pour les essais prédictifs que non prédictifs, mais l'environnement dynamique était différent à chaque essai (pour deux exemples d'essais différents relatifs à un même carrefour, cf. Figure 29a). Contrairement aux expériences précédentes, dans les essais prédictifs, la scène était prédictive de la réponse à fournir. Par exemple, le passage-piétons illustré Figure 29a était systématiquement associé à la réponse « continuer ». Par contre, dans les essais non prédictifs, la scène n'était pas prédictive de la réponse à fournir (cf. Figure 29b). Nous faisons l'hypothèse que des carrefours systématiquement associés à la même réponse devraient conduire à des réponses plus rapides que des carrefours non associés à la même réponse.

Afin d'étudier si l'apprentissage des régularités contextuelles peut conduire à l'émergence de défaillances fonctionnelles, une phase de transfert a été implémentée dans le prolongement de la phase d'apprentissage. Au cours du transfert, les régularités contextuelles étaient altérées de sorte que les carrefours prédictifs devenaient contre prédictifs de la réponse. Si par exemple le passage-piétons illustré Figure 29a était prédictif de la réponse « continuer » au

cours de la phase d'apprentissage, la réponse correcte était « s'arrêter » au cours du transfert (cf. Figure 29c).

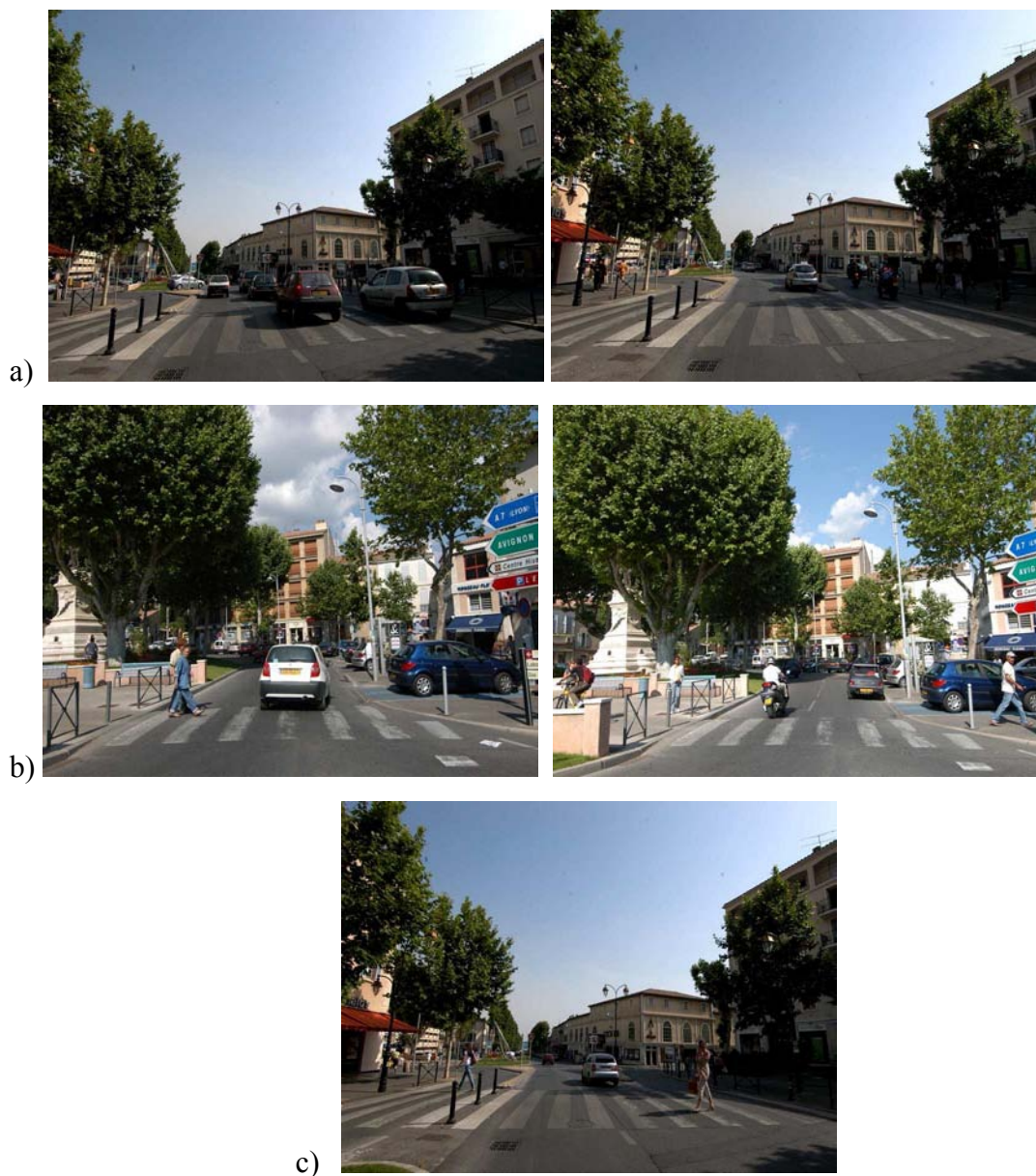


Figure 29. (a) Exemple de deux essais prédictifs de la réponse « continuer ». (b) Exemple de deux essais non prédictifs de la réponse. (c) Exemple d'essai contre prédictif relatif à l'exemple a).

EXPÉRIENCE 8

Méthode

Participants

Douze salariés de l'INRETS (6 femmes et 6 hommes) âgés en moyenne de 32ans ont participé à cette expérience. Tous présentaient une acuité visuelle normale ou corrigée. Hormis un participant, aucun n'avait connaissance du but de cette étude.

Dispositif

L'expérience était pilotée par un ordinateur PC dont l'écran mesurait 17 pouces. L'expérience était générée à partir du logiciel DMDX. Les sujets étaient installés à une distance approximative de 50 cm de l'écran.

Stimuli

Cinq cent quatre-vingt-huit photographies de scénarios routiers ont été utilisées au cours de cette expérience. Douze de ces photographies étaient utilisées dans la phase de familiarisation. Les 576 autres photographies correspondaient à 48 photographies de 12 carrefours différents, à savoir, 4 rond-points, 4 passage-piétons et 4 feux (pour deux exemples de chaque carrefour, cf. Annexe 7). Les 48 photographies relatives à un carrefour particulier étaient différentes les unes des autres. Pour chacun de ces carrefours, 24 étaient associés à la réponse « s'arrêter » et les 24 autres étaient associées à la réponse « continuer ». La réponse correcte était, *a priori*, celle attendue par rapport aux règles établies par le code de la route.

Procédure

Les participants étaient exposés à une tâche de « prise de décision » en vue d'une action et à une phase de verbalisation. L'expérience durait environ 25 minutes.

Tâche de prise de décision. Les participants étaient successivement exposés à des photographies de carrefours routiers. Ils avaient pour consigne de juger le plus rapidement et correctement possible s'ils devaient s'arrêter ou s'ils pouvaient continuer. S'ils jugeaient qu'ils devaient s'arrêter, ils devaient appuyer sur une touche prédéfinie à gauche du clavier, et s'ils jugeaient qu'ils pouvaient continuer, ils devaient appuyer sur une touche prédéfinie à

droite du clavier. Les touches gauche vs. droite associées aux réponses « s'arrêter » vs « continuer » étaient reliées la position des pédales de frein et d'accélération.

La tâche comportait 25 blocs de 12 essais, soit au total 300 essais. Les 24 premiers blocs constituaient la phase d'apprentissage et comportaient chacun 6 essais prédictifs et 6 essais non prédictifs, relatifs aux 12 carrefours différents. Parmi ces 12 carrefours, 2 feux, 2 rond-points et 2 passage-piétons définissaient les essais prédictifs et 2 feux, 2 rond-points et 2 passage-piétons définissaient les essais non prédictifs. Dans un essai prédictif, le carrefour était prédictif de la réponse correcte. Par exemple, le passage-piétons illustré Figure 29a était systématiquement associé à la réponse « continuer » au cours des 24 premiers blocs d'essais. Le participant était donc exposé aux 24 scènes différentes associées à la réponse « continuer » relatives à ce carrefour. Dans un essai non prédictif, le carrefour n'était pas prédictif de la réponse. Par exemple, le passage-piétons illustré Figure 29b était dans la moitié des blocs associé à la réponse « s'arrêter » et dans l'autre moitié des blocs associé à la réponse « continuer ». Pour chaque carrefour non prédictif, 12 scènes associées à la réponse « s'arrêter » et 12 scènes associées à la réponse « continuer » ont été sélectionnées aléatoirement. Pour l'autre moitié des participants exposés aux mêmes carrefours non prédictifs, les 12 autres scènes associées à la réponse « s'arrêter » et les 12 autres scènes associées à la réponse « continuer » étaient utilisées. Au sein d'un bloc, une scène non prédictive de la réponse « s'arrêter vs. continuer » était tirée aléatoirement avec remise afin de déséquilibrer le nombre de réponse « avancer » et « s'arrêter » au sein d'un bloc et de minimiser le développement d'une éventuelle stratégie de réponse. Les 12 carrefours, prédictifs vs. non prédictifs, étaient contrebalancés entre les participants.

Le 25^{ème} bloc d'essai constituait un bloc de transfert. Ce bloc de transfert comportait 6 essais non prédictifs et 6 essais contre prédictifs. Les 6 essais non prédictifs étaient générés de la même manière que ceux de la phase d'apprentissage. Pour les 6 essais contre prédictifs, 6 scènes relatives aux 6 carrefours prédictifs étaient sélectionnées aléatoirement parmi les scènes associées à la réponse contraire. Si par exemple, le passage-piétons illustré Figure 29a était prédictif de la réponse « continuer » au cours de la phase d'apprentissage, une scène contre prédictive relative à ce passage-piétons était tirée aléatoirement parmi les 24 scènes associées à la réponse « s'arrêter ».

L'expérience débutait par les instructions, suivies d'un bloc d'entraînement, composé de 12 essais pour familiariser les participants avec la procédure expérimentale. Ces 12 essais étaient définis par 12 photographies de carrefours différents de ceux utilisés dans la tâche effective. La tâche comportait 25 blocs de 12 essais. Les 24 premiers blocs constituaient la

phase d'apprentissage et le 25^{ème} bloc constituait la phase de transfert. Les participants étaient libre de lancer la tâche expérimentale dès qu'ils le souhaitaient en appuyant sur une touche pour commencer le premier bloc. Après un délai de 500ms, la photographie d'un carrefour apparaissait à l'écran. Les participants devaient juger le plus rapidement et correctement possible s'ils devaient s'arrêter ou s'ils pouvaient continuer. La réponse du sujet faisait immédiatement apparaître un écran blanc comportant au centre un point de fixation noir. L'essai suivant était ensuite initialisé. Au sein d'un bloc, les 12 essais étaient présentés de manière aléatoire. A la fin de chaque bloc, une pause était accordée aux participants. Ils étaient libres de relancer la suite de la tâche ou de prolonger la pause à leur gré. Les 24 blocs de la phase d'apprentissage étaient présentés de manière aléatoire, mais l'implémentation du bloc de transfert était fixée au 25^{ème} bloc.

Tâche de verbalisation. A l'issue de la tâche, les participants étaient soumis aux questions suivantes : « Avez-vous remarqué que les lieux étaient répétés au cours de la tâche ? » « Avez-vous remarqué que certains carrefours étaient systématiquement associés à la même réponse ? » et « Avez-vous remarqué qu'au cours du dernier bloc, la réponse « avancer ou s'arrêter » systématiquement associée à certains carrefours était différente par rapport aux blocs précédents ? »

Résultats

Tâche de prise de décision

Erreurs

Les pourcentages d'erreurs ont été regroupés en 7 époques. Les 6 premières correspondent à la phase d'apprentissage et la septième époque correspond à la phase de transfert. Les 6 premières époques regroupent 4 blocs successifs d'essais et la septième époque correspond au bloc 25. Pour chaque sujet, les pourcentages d'erreurs au sein d'une époque ont été calculés séparément pour les deux conditions prédictive et non prédictive. Les pourcentages d'erreurs sont présentés Figure 30a.

Une ANOVA à mesures répétées a été réalisée, avec les facteurs intra-sujets, condition (prédictive vs non prédictive) et époque (1-7) sur les pourcentages d'erreurs. L'ANOVA montre un effet du facteur époque, $F(6,66)=2.745$, $p<.05$ et une interaction entre les facteurs

condition x époque, $F(6,66)=2.620$, $p<.05$, mais pas d'effet simple du facteur condition, $F(1,11)<1$. Une ANOVA restreinte à la phase d'apprentissage (époque 1-6) n'indique pas d'effet du facteur condition, $F(1,11)=2.940$, $p=.11$, pas d'effet du facteur époque, $F(6,66)<1$ et pas d'interaction entre le facteur époque et condition, $F(6,66)=.873$, $p=.51$. Des analyses partielles ne mettent pas en évidence d'effet du facteur condition aux 6 époques d'apprentissage, mais les pourcentages d'erreurs étaient marginalement plus important dans la condition contre prédictive (13.9%) que dans la condition non prédictive (5.6%) à l'époque 7, i.e. pendant le transfert, $F(1,11)=3.67$, $p=.081$.

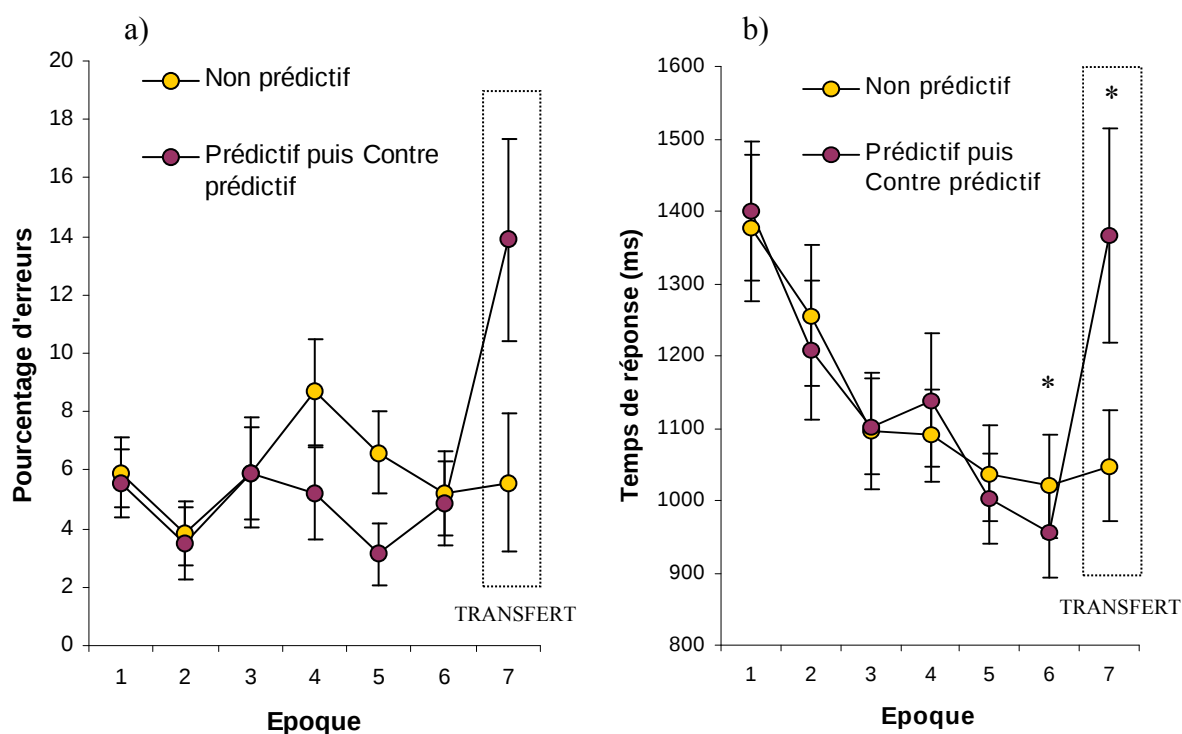


Figure 30. a) Pourcentage d'erreurs dans les conditions non prédictive et prédictive puis contre prédictive. b) Temps de réponse dans les conditions non prédictive et prédictive puis contre prédictive.

La condition contre prédictive était présentée à l'époque 7.

Les barres d'erreurs correspondent aux SEM (N=12).

Temps de réponse

Les TR sur les réponses correctes ont également été regroupés en 7 époques. Les 6 premières époques correspondent à la phase d'apprentissage et la septième époque correspond au bloc 25, i.e. au bloc de transfert. Pour chacun des participants, les moyennes des TR correspondant aux réponses correctes au sein d'une époque ont été calculées séparément pour chacune des conditions testées. Les TR supérieurs ou inférieurs à la « moyenne + 3 écarts-types » ont été éliminés.

Une ANOVA à mesures répétées a été réalisée, avec les facteurs intra-sujets, condition (prédictive vs non prédictive) et époque (1-7). Les TR moyens pour chacune des deux conditions en fonction des époques sont présentés Figure 30b. L'ANOVA montre un effet du facteur époque, $F(6,66)=12.376$, $p<.001$ et une interaction entre les facteurs « condition x époque », $F(6,66)=4.582$, $p<.001$, mais pas d'effet simple du facteur condition, $F(1,11)=1.625$, $p=.229$. Une ANOVA restreinte à la phase d'apprentissage indique seulement un effet du facteur époque, $F(5,55)=22.153$, $p<.001$, donc pas d'effet du facteur condition et pas d'interaction entre le facteur époque et condition, tous les $F_s<1$. Toutefois, des analyses partielles révèlent que les TR étaient plus courts dans la condition prédictive (956ms) que dans la condition non prédictive (1019ms) à l'époque 6, $F(1,11)=6.37$, $p<.05$, et inversement, les TR étaient plus courts dans la condition non prédictive (1047ms) que dans la condition contre prédictive (1366ms) à l'époque 7, i.e. au cours du transfert, $F(1,11)=6.75$, $p<.05$.

Phase de verbalisation

Tous les participants ont remarqué que les mêmes carrefours étaient répétés au cours de la tâche. Dix sur douze ont remarqué que certains carrefours étaient systématiquement associés à la même réponse. Cinq d'entre eux ont détecté le transfert au cours du dernier bloc. Seuls les participants rapportant avoir remarqué que certaines carrefours étaient toujours associés à la même réponse manifestaient un bénéfice sur les TR dans la condition prédictive par rapport à la condition non prédictive à l'époque 6 (bénéfice moyen : 85ms), ainsi qu'un coût sur les TR (coût moyen : 420ms) et sur les réponses correctes (10% d'erreurs supplémentaires par rapport à la condition non prédictive) au cours du transfert. Inversement, au vu des moyennes marginales, les participants rapportant ne pas avoir détecté que les régularités étaient altérées au cours du dernier bloc étaient les plus affectés par le transfert, aussi bien en terme de TR (coût moyen : 530ms) qu'en terme d'erreurs (16.66% d'erreurs supplémentaire dans la condition prédictive par rapport à la condition non prédictive). Les corrélations sont présentées Annexe 8. Cependant, compte tenu du faible effectif de participants, il est difficile d'attester de la validité de ces résultats.

Discussion

Cette étude visait dans un premier temps à tester si des régularités contextuelles systématiquement associées à une même action en conduite facilitent la prise de décision. Pour ce faire, les participants étaient exposés de manière répétée à des régularités prédictives de la décision correcte à prendre en vue d'une action (« s'arrêter » vs. « continuer »). Contrairement à nos attentes, l'apprentissage des régularités contextuelles est peu marqué durant la phase d'apprentissage. Un effet d'indication contextuel n'est observé que tardivement, i.e. à l'époque 6. Les résultats concernant les erreurs au cours de la phase d'apprentissage ne permettent pas de conclure sur un effet d'apprentissage. Deux interprétations peuvent être envisagées : soit les participants deviennent tardivement sensibles aux régularités prédictives, soit ils adoptent une stratégie de prudence tout au long de la tâche, rendant difficile la mise en évidence d'une automatisation des réponses.

Dans un deuxième temps, l'objectif de cette expérience était de tester si des régularités de l'environnement peuvent conduire à l'émergence de défaillances fonctionnelles dans une activité de conduite automobile. Pour ce faire, après une exposition répétée aux régularités prédictives, une phase de transfert était implémentée. Au cours de la phase de transfert, les régularités prédictives de la réponse pendant l'apprentissage étaient altérées. Le résultat principal de cette expérience montre une augmentation des TR dans la condition contre prédictive par rapport à la condition non prédictive. De surcroît, les taux d'erreurs tendaient à être plus importants dans la condition contre prédictive, suggérant que non seulement les participants répondent plus lentement lorsque les régularités sont altérées, mais qu'en plus, ils ont tendance à commettre plus d'erreurs. Si ce dernier résultat n'est pas significatif, on peut toutefois supposer, au vu des écart-types, que des effets délétères sur les réponses correctes soient validés statistiquement avec un effectif de sujets plus important.

Même si les résultats de la phase d'apprentissage restent fragiles, les performances de la phase de transfert attestent que les participants sont devenus sensibles aux régularités prédictives. Pourrait-on observer des effets délétères durant le transfert alors qu'aucun effet d'indication n'apparaît au cours de la phase d'apprentissage ? Pour tester cette hypothèse, il serait intéressant d'implémenter une phase de transfert à l'issue de l'époque 4, avant que l'apprentissage ne soit mis en évidence par un effet d'indication contextuel. Il est en effet

important de déterminer si, au cours de la phase d'apprentissage, les sujets développent une stratégie de « prudence », ne permettant pas la mise en évidence d'un apprentissage ou s'ils ne deviennent sensibles aux régularités que dans les derniers essais.

Cette expérience valide l'hypothèse selon laquelle les schémas de scène peuvent parfois conduire à l'émergence de défaillances fonctionnelles, i.e. à une augmentation des erreurs dans la prise de décision. Les données, même sur un faible effectif de participants, montrent que la prise de décision des participants dans une activité de conduite est plus lente lorsque les régularités prédictives de l'environnement sont altérées. Ce ralentissement ne constitue pas en soi une défaillance fonctionnelle. On peut même supposer qu'au contraire ce ralentissement témoigne d'une adaptation à des situations incongrues. Toutefois, un ralentissement dans la prise de décision peut indirectement conduire à une saturation des ressources attentionnelles nécessaires pour le traitement d'autres informations présentes dans l'environnement. Les résultats sur les erreurs (bien que non significatifs) suggèrent que les attentes peuvent conduire à des erreurs de décision. Ce résultat est congruent avec l'idée selon laquelle les schémas de scène en mémoire influencent le comportement des participants en « minimisant les exigences de la tâche ». Il est probable qu'en conduisant à des réponses « automatiques », la répétition des régularités amène parfois les participants à effectuer une « recherche visuelle sommaire et précipitée », pouvant conduire à des erreurs (Van Elslande et al., 1997). Ce résultat étaye l'hypothèse selon laquelle des attentes relatives à une situation particulière peuvent conduire à la récupération en mémoire d'un schéma de scène inapproprié et à l'émergence de défaillances fonctionnelles.

Ces résultats peuvent être mis en relation avec les théories de l'automatisation basée sur la mémoire. L'automatisme est un mode de fonctionnement essentiel dans la conduite. Il est plutôt admis que la conduite est composée d'activités routinières réalisées rapidement, sans effort, avec peu de conscience, soit, automatiquement. Nous avons déjà évoqué que si ce mode de fonctionnement semble intuitif, il n'en demeure pour autant pas simple à définir de manière consensuelle. Cette étude s'inscrit dans une approche du concept d'automatisme en tant que résultat d'un processus mnésique (Logan, 1988). La récupération en mémoire des schémas de scène permettrait une orientation automatique de l'attention vers les éléments pertinents de la scène et/ou une récupération automatique de la solution relative à la tâche. Le recours à ce type de processus permettrait de répondre rapidement, et économiquement en terme de ressources attentionnelles allouées au traitement. En contre partie, une recherche

active et sérielle d'informations pertinentes pour la situation en cours, basée sur le recours à des algorithmes généraux de recherche, offrirait une plus grande flexibilité des réponses, permettant notamment de répondre aux événements imprévus. Mais elle serait plus coûteuse en temps et en ressources allouées au traitement. Nos résultats sont concordants avec cette conception. Si les résultats de la phase d'apprentissage mettent en évidence un effet facilitateur tardif des régularités contextuelles sur les réponses des participants, les résultats relatifs au transfert montrent que la perturbation de « routines », même récentes, peut affecter les performances, suggérant une récupération automatique des connaissances en mémoire dans la prise de décision.

Par ailleurs, les données sur les verbalisations fournissent des prémisses intéressantes. De manière peu surprenante, tous les participants ont remarqué que les carrefours étaient répétés au cours de chaque bloc. Ce résultat est concordant avec les résultats rapportés par Brockmole et Henderson (2006a) et avec les données de la littérature suggérant que l'homme est doté d'une prodigieuse capacité à mémoriser des images visuelles (e.g. Shepard, 1967). De plus, les participants (2/12) n'ayant pas remarqué que certains carrefours étaient systématiquement associés à la même réponse ne manifestaient aucune tendance d'indigage contextuel durant la phase d'apprentissage et aucun effet sur les TR et sur les réponses correctes durant le transfert, laissant penser qu'ils ne sont pas devenus sensibles aux régularités. Le faible effectif de participants ne nous permet cependant pas de valider une hypothèse accordant à la conscience un rôle dans l'apprentissage des régularités. Cependant, si tous les participants ayant remarqué les régularités prédictives durant l'apprentissage manifestaient un coût sur les TR durant le transfert, seule la moitié d'entre eux rapportaient avoir remarqué que les régularités étaient altérées durant le dernier bloc. Si l'apprentissage est explicite, dans le sens où la connaissance sur les régularités est accessible à la conscience pour tous les participants présentant un effet d'indigage durant l'apprentissage, il n'est pas exclu que cette connaissance soit récupérée de manière implicite, i.e. non intentionnelle, voire inconsciente. En effet, les participants rapportant ne pas avoir remarqué le transfert étaient pourtant les plus affectés par celui-ci. Plusieurs interprétations sont toutefois envisageables. Il est peu étonnant que les participants ne détectent pas le transfert pour des essais sur lesquels ils ont répondu incorrectement. Ces derniers pourraient adopter une stratégie explicite basée sur la récupération directe de la solution, sans chercher les indices pertinents dans l'environnement pour répondre. Ce fait ne remettrait pas en question une récupération explicite, i.e. intentionnelle et consciente de la connaissance sur les régularités apprises. Néanmoins, les

résultats relatifs aux TR durant le transfert sont plus énigmatiques. Comment expliquer que les sujets n'ayant pas remarqué le transfert soient ceux dont les performances sur les TR étaient les plus affectées par celui-ci ? Cette étude mériterait des mesures complémentaires pour tester l'hypothèse selon laquelle les connaissances sur les régularités contextuelles sont récupérées de manière implicite.

Un grand nombre d'accidents de la route ont pour origine des défaillances fonctionnelles perceptives (Gross & Feldman, 1995 ; pour une revue, Chapman & Underwood, 1998). Les études détaillées d'accidents (EDA) indiquent que 33 % des accidents de la route peuvent être imputés à des défaillances attentionnelles liées à des problèmes « de détectabilité », « d'organisation défectueuse de la prise d'information » ou encore des problèmes liés à « une absence de recherche active d'information » (Van Elslande et al., 1997). La conduite automobile consiste en partie à rechercher des informations pertinentes dans un environnement complexe et dynamique. Quelles sont les stratégies de recherche mises en œuvre par les conducteurs lors de l'analyse d'une scène visuelle routière ? Quels sont les facteurs susceptibles d'être responsables d'une conduite parfois défectueuse ? Les résultats de cette expérience suggèrent que les attentes relatives à une situation spécifique peuvent être source de dysfonctionnements. Nous proposons ainsi que les mêmes facteurs responsables d'une conduite « efficace » peuvent aussi conduire à l'émergence de défaillances fonctionnelles.

On reconnaîtra toutefois que les données rapportées dans cette étude sont fragiles pour les inscrire dans un cadre théorique robuste. Non seulement cette expérience manque terriblement d'un effectif suffisant de participants pour valider ou invalider les effets tendanciellement observés, mais surtout, certaines conditions expérimentales seraient fortement utiles pour rendre compte des effets de répétition des carrefours sur le guidage de l'attention vs. la sélection de la réponse. Dans une expérience en cours, une condition « nouvelle » est rajoutée au matériel, dans laquelle des carrefours ne sont présentés qu'une seule fois au cours de la tâche. Si les prémisses de cette étude apportent des résultats prometteurs, un certain nombre de variables supplémentaires nous renseignera considérablement sur l'influence des régularités contextuelles sur les processus perceptifs, en distinguant leurs effets sur le guidage de l'attention de leurs effets sur la sélection de la réponse. Par ailleurs, tester l'effet de « l'expérience des conducteurs » pourra apporter des données pertinentes pour aborder le rôle des connaissances préalables dans la construction de nouvelles attentes.

CHAPITRE VII

DISCUSSION GÉNÉRALE

Les études menées ces dernières années autour des représentations visuelles montrent que contrairement à notre phénoménologie qui donne l'impression de percevoir une scène visuelle de manière stable et détaillée dans son ensemble, les représentations visuelles seraient davantage schématiques et transitoires. Certaines théories postulent même qu'aucune information visuelle n'est mise en mémoire au cours des saccades oculaires (Noë & O'Regan, 2001 ; Rensink, 2000a ; voir cependant, Hollingworth & Henderson, 2002 ; Simons & Rensink, 2005). Pourtant, afin de voir plus en détail les objets présents dans le champ visuel, l'attention sélective est orientée de manière rapide et efficace vers les éléments pertinents de la scène, compte tenu des objectifs finalisés de l'observateur. Nous avons vu dans le chapitre I, qu'en structurant la scène, les régularités contextuelles facilitent la reconnaissance des objets qui la composent, ainsi que son exploration. Les connaissances sur les régularités de l'environnement sont de bons candidats pour rendre compte de l'efficacité des traitements perceptifs, et notamment du guidage de l'attention au sein de l'environnement visuel. Ces connaissances seraient inscrites en mémoire sous forme de schémas de scène. L'hypothèse qui constitue le fil rouge de nos recherches est que les connaissances relatives aux régularités de l'environnement peuvent être de nature explicite mais également de nature implicite. Dans cette perspective, nous avons évoqué, dans le Chapitre II, le scepticisme qui règne autour de l'existence de connaissances inaccessibles à la conscience, en présentant les difficultés méthodologiques pour éprouver la nature consciente vs. inconsciente des connaissances mises en évidence (Perruchet & Vinter, 2002 ; Shanks & St. John, 1994). Puis nous avons énoncé deux questions faisant l'objet de débats empiriques et théoriques dans le domaine de

l'apprentissage implicite. Ces deux questions sont à l'origine de nos travaux. La première question est de savoir si l'apprentissage implicite peut reposer sur des régularités contextuelles de nature spécifique d'une part, et de nature catégorielle et sémantique d'autre part. La deuxième question porte sur le rôle de l'attention dans l'apprentissage implicite. Nous avons évoqué de manière succincte la littérature relative à ces deux questions, en soulignant quelques « contradictions » dans les données empiriques. Enfin, l'objectif de nos travaux en cours est d'explorer l'apprentissage des régularités contextuelles dans une tâche plus écologique, telle que la conduite automobile, en testant l'hypothèse que les régularités de l'environnement peuvent parfois être sources de défaillances fonctionnelles. Dans l'ensemble de nos études, nous avons eu recours au paradigme d'indication contextuel ou « *contextual cueing* », développé par Chun et Jiang en 1998. Les résultats empiriques obtenus à partir de ce paradigme montrent que des régularités contextuelles peuvent faciliter la recherche d'un élément particulier dans une scène visuelle, via des mécanismes d'apprentissage implicite ou explicite. Le Chapitre III constituait une revue de question sur les recherches réalisées ces dernières années avec ce paradigme. Nous avons présenté un certain nombre d'arguments montrant que le paradigme d'indication contextuel constitue un outil expérimental pertinent pour explorer des mécanismes d'apprentissage implicite vs. explicite de régularités contextuelles. Les chapitres IV, V et VI, rapportaient l'ensemble de nos travaux de recherche, dont nous allons faire une brève synthèse dans la section suivante.

1. Synthèse des résultats obtenus

Trois études sont rapportées dans ce travail. Dans l'Etude 1, cinq expériences utilisant le paradigme d'indication contextuel ont été conduites avec des affichages composés de nombres. Les Expériences 1a et 1b, montrent que des régularités reposant sur des caractéristiques spécifiques du contexte peuvent faciliter le guidage de l'attention dans une tâche de recherche visuelle de cible. En effet, des effets d'indication contextuel ont été mis en évidence lorsque des affichages numériques spécifiques étaient prédictifs de la localisation de la cible. Les Expériences 2a, 2b et 3, montrent que l'indication contextuel peut être étendu à la prise en compte de régularités contextuelles catégorielles et sémantiques. La recherche visuelle était ici facilitée lorsque la propriété de parité/impairité du contexte prédisait la localisation de la

cible. Un transfert sur des nombres distracteurs nouveaux atteste que ces effets reposent sans équivoque sur le concept de parité/imparité, et non sur des traits visuels ou caractéristiques spécifiques des nombres présentés initialement. Dans ces cinq expériences, des effets d'indication étaient observés sans que les participants ne soient conscients des régularités facilitant néanmoins leur recherche visuelle. Non seulement les participants étaient incapables de verbaliser les régularités contextuelles, mais surtout, ils étaient incapables d'exploiter ces régularités dans une tâche de reconnaissance. Cette étude montre qu'au cours de l'analyse d'une scène visuelle, des mécanismes d'apprentissage implicite peuvent être déployés sur des régularités contextuelles spécifiques mais également sur des régularités contextuelles catégorielles et sémantiques, selon leur valeur prédictive.

L'objectif de l'Etude 2 était, dans un premier temps, de généraliser les effets d'apprentissage implicite de régularités basées sur l'appartenance catégorielle et sémantique du contexte à d'autres catégories sémantiques, et dans un deuxième temps, de tester si cet apprentissage et l'expression de cet apprentissage requièrent une attention sélective. Le paradigme d'indication contextuel a été utilisé avec des affichages lexicaux, dans lesquels la catégorie sémantique du contexte prédisait ou non la localisation de la cible. Les Expériences 4 et 5 généralisent les effets d'indication implicite basés sur l'appartenance catégorielle et sémantique du contexte observés dans l'Etude 1. L'Expérience 6 indique un effet d'indication contextuel lorsque le contexte prédictif apparaît dans la couleur attendue (i.e. la même couleur que la cible), mais pas lorsqu'il apparaît dans une couleur ignorée (i.e. une couleur différente de celle la cible). Cependant, lorsque les contextes antérieurement ignorés devenaient subitement attendus, ces derniers facilitaient immédiatement la performance (Expérience 7). En revanche, lorsque les contextes antérieurement attendus devenaient subitement ignorés, aucun bénéfice n'était observé. Ici encore, les effets d'apprentissage observés étaient de nature implicite. Les résultats de cette étude révèlent ainsi que si l'expression de la connaissance sémantique implicite requiert une attention sélective, un apprentissage latent peut néanmoins se développer sur des régularités sémantiques extérieures au focus attentionnel.

L'Etude 3 visait, d'une part, à tester l'apprentissage de régularités dans une tâche plus proche de celle impliquée dans une activité de conduite automobile, et d'autre part, à tester si les connaissances relatives aux régularités de l'environnement sont parfois source de dysfonctionnements. Les participants avaient pour instruction de juger de situations routières

en vue d'une action, i.e. « avancer » vs. « s'arrêter ». La moitié de ces situations répétées était prédictive de la réponse alors que l'autre moitié était non prédictive. Contrairement à nos attentes, la répétition des situations prédictives conduisait à des effets d'indigence contextuel tardifs. Toutefois, lorsque les réponses correctes relatives aux situations spécifiques prédictives devaient incongrues par rapport aux essais précédents, les performances se dégradaient fortement. Par ailleurs, les rapports verbaux suggèrent que l'apprentissage des régularités était explicite. Les participants manifestant une sensibilité aux régularités rapportaient tous avoir détecté que certaines situations étaient systématiquement associées à la même réponse. Ces résultats suggèrent que des connaissances explicites sur des régularités contextuelles dans une situation de conduite peuvent faciliter la prise de décision en vue d'une action, mais également conduire à l'émergence de défaillances fonctionnelles. Toutefois, compte tenu du faible effectif des participants, les données rapportées dans cette étude sont encore trop fragiles pour en dégager de réelles conclusions. Des expériences complémentaires sont nécessaires pour valider ou invalider les premiers résultats observés.

L'ensemble des résultats rapportés dans ce mémoire de thèse peut être résumé en cinq points :

- L'apprentissage implicite peut reposer sur des régularités contextuelles de nature spécifique mais également catégorielle et sémantique de l'environnement.
- Des mécanismes d'apprentissage implicite peuvent être déployés sur des régularités sémantiques extérieures au focus attentionnel. Cependant, pour que la connaissance s'exprime, une attention portée sur les régularités apprises était ici requise.
- La prise de conscience des régularités contextuelles dans des scènes du monde réel semble favoriser leur encodage en mémoire et/ou leur récupération.
- Les connaissances sur les régularités de l'environnement peuvent être sources de défaillances fonctionnelles.
- Les connaissances implicites ou explicites relatives aux régularités de l'environnement guident les processus de sélection attentionnelle.

Compte tenu de la fragilité des résultats rapportés dans l'Etude 3, ces derniers seront peu rediscutés ici. Le présent chapitre vise essentiellement à discuter les données observées dans les Etudes 1 et 2.

2. Apprentissage implicite de régularités sémantiques

Le premier résultat majeur de nos recherches (Etudes 1 et 2) montre que des mécanismes d'apprentissage implicite peuvent être déployés sur des régularités spécifiques mais également sur des régularités catégorielles et sémantiques. Dans la littérature sur l'apprentissage implicite, il est parfois stipulé que l'apprentissage implicite repose exclusivement sur des propriétés spécifiques, et de ce fait, que le transfert de la connaissance acquise est limité à des exemplaires partageant des traits perceptifs similaires. En montrant que des connaissances inaccessibles à la conscience peuvent porter sur des propriétés catégorielles et sémantiques du matériel d'apprentissage, nos résultats sont contraires à ce point de vue.

Bien qu'à notre connaissance, aucune base de données ne nous permette de valider cette hypothèse, nos résultats suggèrent que la force des liens sémantiques qui existent entre les différents items des catégories utilisées (i.e. les mots dans l'Etude 2) conditionne la force des effets d'indiciage contextuel. Cette hypothèse peut être mise en relation avec les théories connexionnistes visant à modéliser l'organisation des concepts en mémoire. Ces théories considèrent que les concepts sont associés en MLT par des liens de force variable qui dépendraient notamment des relations sémantiques que partagent ces concepts (pour une revue, McClelland & Rogers, 2003). Dans ces modèles notamment, lorsque l'on apprend quelque chose sur « chat », on apprend indirectement sur « chien » et les autres animaux, car les connaissances correspondantes sont toutes reliées entre elles. Ces effets indirects n'ont pas besoin d'atteindre un état de conscience pour exister. L'activation transitoire d'un concept activerait de manière diffuse les concepts qui lui sont associés au sein d'un patron de connectivité. L'apprentissage pourrait donc simplement conduire à une modification des poids des connections neuronales sans pour autant impliquer de modification des représentations correspondantes et sans que ces effets indirects ne soit nécessairement accompagnés d'une prise de conscience. Les poids des connexions changeraient de manière auto-organisée, en fonction de l'état d'activation des unités et des feed-back de l'environnement. Les mécanismes d'apprentissage implicite, même déployés sur des connaissances conceptuelles, peuvent être entièrement expliqués par des mécanismes neurophysiologiques et peuvent être modélisés par des systèmes connexionnistes, sans avoir recours à la conscience pour rendre compte de leur existence.

Nos résultats sont compatibles avec les modèles connexionnistes mais semblent difficilement conciliables avec les présupposés définis par le modèle SOC (« Self-Organizing Consciousness »), développés par Perruchet et Vinter (2002). Le modèle SOC se situe dans une perspective « mentaliste » de la cognition en réponse à la théorie computationnelle de l'esprit qui présuppose l'existence d'un inconscient cognitif. Dans un cadre mentaliste, l'inconscient cognitif n'a aucune place. Le modèle SOC rejette la conception selon laquelle la manipulation et la computation d'informations impliquées durant un apprentissage puissent se dérouler à un niveau inconscient seulement. Les traitements associatifs n'opéreraient que sur des représentations conscientes. *“Processes and mechanisms responsible for the elaboration of knowledge are intrinsically unconscious, and the resulting mental representations and knowledge are intrinsically conscious. No other components are needed”* (Perruchet et al., 1997, p. 44). En écho aux théories connexionnistes, Perruchet et Vinter proposent que l'apprentissage repose sur la nature auto-organisée de la pensée consciente et non sur celle des réseaux connectés. Dans une perspective mentaliste, l'apprentissage implicite de régularités sémantiques n'existe pas.

Sans minimiser le niveau de profondeur de traitement requis pour extraire une catégorie sémantique derrière un mot, on peut néanmoins se demander si l'on a véritablement besoin de s'appuyer sur des processus d'abstraction pour expliquer les performances observées dans l'Etude 1 et l'Etude 2. Selon Reber (1993, p. 120-121), *“an abstract representation is assumed to be derived yet separate from the original instantiation. Abstract codes contain little, if any, information pertaining to the specific stimulus features from which they were derived; the emphasis is on structural relationships among stimuli”*. Si l'on s'en tient à cette définition, les apprentissages mis en exergue dans nos travaux relèvent de processus d'abstraction. Toutefois, nos effets reposent sur un apprentissage associatif à partir de catégories conceptuelles préexistantes en mémoire. Les « règles » dégagées par les participants associaient de manière directe une catégorie sémantique à la localisation de la cible. L'extraction de ces catégories impliquent-elles vraiment un processus d'abstraction ? A notre avis, nos travaux ne démontrent pas l'existence de processus d'abstraction qui ne soient pas accompagnés d'expérience subjective. Ils démontrent néanmoins que le système cognitif est capable d'acquérir de nouvelles connaissances à partir de connaissances sémantiques déjà en mémoire, et de généraliser ces connaissances à des exemplaires définis par la même catégorie sémantique. Il ne s'agit donc pas de dire que l'individu est capable d'abstraire ou de former inconsciemment un nouveau concept, mais qu'il est néanmoins capable d'extraire

implicitement des régularités sémantiques à partir d'exemplaires spécifiques, d'encoder en MLT des connaissances relatives à ces régularités sémantiques, et de généraliser ces connaissances implicites à des exemplaires partageant des traits sémantiques communs, mais des traits visuels différents. La thèse exposée dans ce manuscrit ne défend pas l'existence d'un système d'abstraction inconscient, pas plus que l'existence de représentations inconscientes, mais elle offre des arguments en faveur d'un inconscient cognitif, capable d'opérer sur des connaissances sémantiques préexistantes en mémoire.

Il faut en outre souligner que dans des expériences préliminaires de l'Etude 2, lorsque les cibles étaient définies par des mots spécifiques, nous ne sommes pas parvenus à mettre en évidence d'effet d'indiciage. En effet, nous n'observions pas d'effet d'apprentissage lorsque la tâche n'imposait pas aux participants une contrainte sémantique. Nous avons interprété cette absence d'effet en proposant que la recherche de mots spécifiques mette en œuvre une recherche visuelle superficielle, essentiellement basée sur les caractéristiques perceptives des stimuli. Comme dans le modèle « Guided Search » proposé par Wolfe, Cave et Franzel (1989), les traits visuels pourraient être comparés aux descriptions visuelles en mémoire sur la cible potentielle. Les stimuli partageant des similarités perceptives avec les cibles potentielles auraient une plus grande probabilité d'être visités. Pourtant, les résultats de l'Expérience 7 montrent que l'attention sélective n'est pas requise pour l'apprentissage. Comment expliquer alors qu'une contrainte sémantique soit nécessaire à la mise en évidence d'un apprentissage ? Ce fait peut renforcer l'idée qu'un traitement attentionnel profond et notamment sémantique est requis à l'expression des connaissances implicites sémantiques. Mais il n'est pas non plus exclu que lorsque les cibles sont définies de manière spécifique, les TR soient trop courts pour permettre la mise en évidence d'un effet d'indiciage contextuel sémantique. On peut aussi envisager qu'indépendamment de l'attention focalisée, la nature du traitement effectué sur les items, attendus et ignorés, dépende de l'objectif des sujets et des contraintes inhérentes à la tâche à accomplir induite par la consigne (exigeant ou non un traitement sémantique). Cette hypothèse est concordante avec des travaux indiquant que des effets d'amorçage sémantique ne peuvent être mis en évidence que dans des tâches exigeant un traitement sémantique. Becker, Moscovitch, Behrmann et Marlene (1997) ont même suggéré que le niveau de traitement opéré sur la cible doit être identique à celui requis sur l'amorce pour que des effets d'amorçage sémantique soient observés. Selon eux, cela expliquerait pourquoi il est difficile d'obtenir des effets d'amorçage sémantique à long terme dans une tâche de décision lexicale, alors que de tels effets peuvent être plus facilement observés à travers des tâches de

catégorisation lexicale. Une tâche de décision lexicale pourrait en effet être rapidement effectuée sans engager de traitements sémantiques, tout du moins suffisamment profonds pour établir un lien associatif entre l'amorce et la cible. En accord avec cette conception, il est fort probable que la tâche oriente la nature des traitements réalisés sur les stimuli attendus et ignorés. Ce point de vue rejoint partiellement celui défendu par Wright et Whittlesea (1998), dans le sens où ce qui est appris des stimuli dépendrait de l'activité cognitive qui est déployée, autrement dit de la tâche en cours. Cependant, pour Wright et Whittlesea, le sujet exerce toujours une stratégie d'apprentissage, ce qui n'est pas le cas dans nos travaux.

Pourtant, dans l'Etude 1 conduite avec des nombres, nous avons obtenu des effets d'indilage basés sur la propriété de parité/imparité sans pour autant imposer une contrainte sémantique aux participants. Les participants recevaient pour instruction de chercher un nombre spécifique (e.g. 13 ou 28). Pour expliquer cette différence, nous pouvons nous appuyer sur une hypothèse formulée par Greenwald, Abrams, Naccache et Dehaene (2003) selon laquelle l'accès à la sémantique des mots serait plus long que l'accès à la sémantique des nombres arabiques. En effet, de nombreux travaux ont montré une activation automatique du « sens » des nombres (e.g. magnitude, parité/imparité) dans des tâches où la signification des nombres était inutile (Brysbaert, 1995 ; Dehaene, Bossini, & Giraux, 1993 ; Fias, Brysbaert, Geypens, & d'Ydewalle, 1996). De manière différente, l'activation du sens des mots serait plus lente du fait des multiples aspects sémantiques qui s'y rapporte.

3. Apprentissage implicite et attention visuelle sélective

Le deuxième résultat majeur de nos travaux (Etude 2, Expérience 7) révèle que l'apprentissage implicite de régularités sémantiques peut émerger sans que l'attention sélective ne soit directement focalisée sur ces régularités. Au regard d'une littérature extrêmement controversée sur le rôle de l'attention dans l'apprentissage implicite (cf. Chapitre II.4), ce résultat est pour le moins surprenant, compte tenu du niveau de profondeur de traitement requis pour extraire des régularités sémantiques. En effet, de nombreuses recherches suggèrent que pour un matériel complexe, seuls les stimuli sélectivement attendus en liaison avec leur pertinence pour la tâche bénéficient d'un traitement suffisamment profond

pour être mémorisé de manière durable (e.g. Cohen et al., 1990). L'attention en tant que processus de sélection permettrait aux entrées visuelles attendues d'être traitées plus rapidement, plus profondément que les autres, de telle manière que ces dernières auront plus de chance d'influencer une réponse comportementale ou d'être mémorisées (Desimone & Duncan, 1995 ; Egeth & Yantis, 1997). Il est possible que les nombreuses expositions aux régularités sémantiques ignorées aient les mêmes conséquences que des expositions moins nombreuses sur des régularités sémantiques attendues. Cependant, de manière fort intrigante, le bénéfice observé dans la condition « prédictive ignorée puis attendue » de l'Expérience 7 était équivalent à celui observé dans la condition « prédictive attendue » de l'Expérience 6. Ce résultat est concordant avec le résultat observé par Jiang et Leung (2005) à partir de régularités spatiales (pour un résultat similaire, cf. Kim & Kim, 2006).

Si l'Expérience 7 suggère que des mécanismes d'apprentissage peuvent se mettre en place sur des stimuli ignorés, elle montre néanmoins que l'attention sélective est nécessaire à l'expression de la connaissance résultante. Il s'agit là du troisième résultat majeur de nos travaux. L'Expérience 7 montre clairement que l'indigage contextuel ne se manifeste que sur des régularités sémantiques prédictives focalisées par l'attention. Un phénomène identique a déjà été rapporté dans la littérature avec le paradigme d'indigage contextuel (Jiang & Leung, 2005 ; Kim & Kim, 2006) mais également avec le paradigme de SRT (e.g. Frensch, Lin, & Buchner, 1998). L'attention visuelle sélective pourrait ainsi être nécessaire, non pas pour l'encodage des régularités, mais pour la récupération ou pour l'expression de la connaissance résultante.

Peut-on pour autant parler d'apprentissage passif ? Les travaux de Kim et Kim (2006) suggèrent que l'indigage contextuel sur les contextes antérieurement ignorés ne se manifestent que si la cible associée aux contextes répétés apparaît dans la couleur attendue au cours de l'apprentissage. De manière congruente, les travaux de Seitz et Watanabe (2003), conduits avec un paradigme de détection de mouvement cohérent, montrent que seules les directions de mouvement couplées temporellement avec les cibles conduisent à un apprentissage (pour une description de la tâche, cf. Annexe 9). Selon Seitz et Watanabe, la relation de co-occurrence entre les traits non pertinents pour la tâche et la présentation de la cible est essentielle à l'apprentissage des traits ignorés. Le développement d'une sensibilité aux mouvements cohérents serait le résultat d'un couplage temporel entre ces mouvements ignorés et une cible pertinente pour la tâche. Sur la base de leurs résultats, Seitz et Watanabe (pour une revue,

Seitz & Watanabe, 2005), ont proposé une théorie selon laquelle un apprentissage perceptif peut se mettre en place sur un trait perceptif ignoré, si ce dernier coïncide temporellement avec un signal de renforcement relatif au traitement de la cible par exemple. Selon eux, les mêmes mécanismes d'apprentissage seraient déployés sur des traits perceptifs attendus et ignorés, si tant est que ces derniers coïncident temporellement avec les signaux de renforcement relatifs à la tâche en cours. Le *modèle d'apprentissage unifié* développé par Seitz et Watanabe (2005) vise à expliquer l'apprentissage perceptif. En intégrant les principes computationnels développés dans les théories connexionnistes, nos résultats offrent des perspectives pour étendre le modèle d'apprentissage unifié à la prise en compte de traits sémantiques. L'apprentissage des associations « catégorie sémantique – localisation de la cible » serait la conséquence d'un couplage temporel entre la présence de régularités sémantiques et la localisation de la cible. La facilitation de la recherche visuelle constituerait ici un renforcement positif à l'apprentissage. Des feedback sur le caractère adaptatif ou non adaptatif de la performance influenceraient les mécanismes de sélection attentionnelle. Dans ce cadre, des recherches montrent qu'une motivation, comme l'obtention d'une récompense constitue un renforcement à l'apprentissage, même dans des tâches plutôt « triviales ». Par exemple, les travaux de Libera et Chelazzi (2006) montrent qu'un renforcement positif, comme une récompense monétaire pour la simple identification d'une cible, module l'apprentissage perceptif. Libera et Chelazzi proposent que si aucun feedback exogène n'est fourni à l'individu, la performance elle-même pourrait constituer une sorte de « récompense endogène ». Les réponses correctes ou rapides par exemple, pourraient constituer un renforcement positif et favoriser l'apprentissage.

Nos résultats signifient-ils qu'aucune forme d'attention n'est requise à un apprentissage conceptuel ? Selon Posner et Petersen (1990), le rôle de l'attention dans l'apprentissage peut être évalué de manière différente selon le type de sous-système attentionnel considéré. L'alerte, l'orientation et la fonction exécutive sont des systèmes attentionnels dissociables ayant des effets différents sur le traitement d'un stimulus. Dans le modèle d'apprentissage unifié, l'activation du sous-système d'alerte par un indice temporel, tel l'occurrence de la cible pour la tâche, serait nécessaire pour rehausser le traitement d'une partie étendue de la scène, incluant les traits attendus et ignorés présents dans le champ visuel. Mais l'activation du sous-système d'orientation, orientant les ressources attentionnelles vers les cibles potentielles (dans notre tâche, les mots écrits dans la couleur attendue), ne serait pas nécessaire pour qu'un apprentissage émerge sur les items ignorés. En accord avec ce point de

vue, l'apprentissage catégoriel mis en évidence sur les mots ignorés ne signifie pas qu'aucune forme d'attention n'est engagée dans cet apprentissage, mais qu'un apprentissage latent peut néanmoins se mettre en place sur des régularités conceptuelles présentes en dehors du focus attentionnel.

Nos résultats suggèrent-ils que l'apprentissage de régularités sémantiques était ici automatique ? Si l'on s'en tient à la définition la plus admise de l'automatisme, on peut supposer que les participants ont encodé de manière automatique les régularités sémantiques manipulées dans l'Etude 2. Néanmoins, la récupération de la connaissance ou l'expression de cette connaissance sur le comportement semble être le fruit d'un processus qui requiert une attention sélective. En outre, on peut envisager que l'expression de la connaissance sur le comportement implique de nombreux processus. De manière non exhaustive, elle met en œuvre un processus de catégorisation sémantique, un processus de couplage « catégorie-localisation de la cible » et un processus qui oriente les ressources attentionnelles vers la localisation de la cible. Compte tenu des résultats observés dans la condition « prédictive ignorée puis attendue » (Etude 2, Expérience 7), on peut spéculer que l'attention sélective est requise, soit pour le processus de couplage « catégorie-localisation de la cible », soit au niveau des conséquences de ce couplage sur les processus de sélection attentionnelle.

Pour conclure sur ces deux sections, nos résultats amènent à défendre l'existence d'un inconscient cognitif capable d'encoder des régularités sémantiques et d'agir de manière non-intentionnelle sur les comportements. Sans tomber dans une conception assimilant l'individu à un « zombie » (Cf. Cleeremans & Jiménez, 2002), on peut reconnaître que de nombreuses données rapportées dans la littérature, offrent des arguments convaincants en faveur de l'existence de traitements cognitifs non conscients. Nos résultats ne peuvent être interprétés en termes d'apprentissage intentionnel, conduisant à l'émergence d'une connaissance accessible à la conscience. Il ne s'agit pas d'affirmer que nos conduites sont guidées par des traitements échappant à toute expérience sensible, mais de reconnaître que certains traitements élémentaires inscrits dans des comportements « plus finalisés » (e.g. un mouvement de l'œil dans une tâche de recherche visuelle) peuvent être guidés par des connaissances échappant à toute expérience subjective, aussi bien lors de leur encodage que lors de leur récupération. Nous insistons néanmoins sur le fait que nos résultats ne mettent pas en évidence de processus d'abstraction de concepts inexistants en mémoire. Ces derniers

montrent des effets d'apprentissage associatif relativement simples et non des effets d'apprentissage implicite de règles complexes, comme cela a pu être suggéré dans les travaux sur les grammaires artificielles. Aussi nos résultats ne sont-ils pas incompatibles avec un point de vue selon lequel la déduction, le raisonnement et l'abstraction véritables relèvent de traitements impliquant et conduisant à une expérience phénoménale.

4. Le rôle de la conscience dans l'apprentissage des régularités

Si la thèse ici défendue accorde aux traitements cognitifs non conscients une influence sur les comportements, il est bon de souligner qu'elle ne minimise en rien le caractère adaptatif de la conscience. Compte tenu de la fragilité des résultats rapportés dans l'Etude 3, il nous paraît prématuré d'en tirer de réelles conclusions. Toutefois, les premiers résultats obtenus suggèrent que dans cette étude, la prise de conscience des régularités contextuelles favorise l'exploitation de ces régularités dans la sélection de la réponse. De plus, les participants rapportant verbalement ne pas avoir détecté le transfert semblaient être les plus affectés par celui-ci. Si ces résultats sont confirmés sur un échantillon de sujets plus important et par des mesures objectives des connaissances explicites, ils suggéreront que la prise de conscience de « contre-exemples » minimise l'émergence de défaillances fonctionnelles relatives à l'existence de routines susceptibles de causer une perte progressive de contrôle.

Dans une étude réalisée au sein de notre équipe¹⁷, visant à tester l'hypothèse d'un gradient dans la prise de conscience de connaissances sur des régularités présentes dans des scènes du monde réel, nous avons pu observer le rôle adaptatif de la conscience dans l'exploitation des régularités contextuelles. Le point de départ de cette étude provenait du caractère explicite des connaissances mises en évidence dans les tâches d'indication contextuel où les contextes sont définis par des scènes du monde réel. En effet, les travaux de Brockmole et Henderson (2006a) montrent que lorsque les contextes ne sont plus définis par des affichages arbitraires mais par des scènes du monde réel, l'apprentissage des régularités est explicite. Non seulement les participants reconnaissent mieux que du seul fait du hasard les scènes répétées,

¹⁷ Etude conduite par Veena Murdymootoo dans le cadre de son Master II de Psychologie

mais en plus, ils sont capables de rapporter avec précision la localisation de la cible au sein des scènes anciennes. L'hypothèse testée au cours de cette étude était que les connaissances sur les régularités contextuelles sont au préalable implicites puis deviennent rapidement accessibles à la conscience. Pour tester cette hypothèse, deux expériences ont été réalisées. La première expérience, identique à celle de Brockmole et Henderson (2006a), répliquait les données obtenues par ces auteurs. Dans une deuxième expérience, des tests directs de mémoire (i.e. verbalisation, reconnaissance et localisation) ont été introduits dans les étapes précoces de l'apprentissage. Les résultats suggèrent que les connaissances relatives aux régularités contextuelles sont en partie accessibles à la conscience dès le début de l'apprentissage. Toutefois, les résultats individuels montrent que si les taux d'apprentissage observés dans la tâche indirecte (i.e. recherche visuelle de cible) sont bien corrélés aux performances observées dans les tâches directes de mémoire (i.e. reconnaissance et estimation de la localisation de la cible) dans l'Expérience 1 (i.e. après 17 blocs d'essais), aucune corrélation n'apparaît dans l'Expérience 2 (i.e. après 5 blocs d'essais). Les participants manifestant un fort effet d'indice contextuel après 5 répétitions ne présentaient pas de meilleures performances aux tâches de mémoire explicite que les participants manifestant peu, voire pas d'indice contextuel. L'apprentissage des régularités précéderait-il la prise de conscience des régularités ? En revanche, après 17 répétitions, une forte corrélation était observée, suggérant que la prise de conscience des régularités favorise l'exploitation de ces régularités.

Par ailleurs, on peut émettre l'hypothèse que si les participants avaient détecté explicitement les régularités présentes dans les Etudes 1 et 2, l'apprentissage aurait été plus efficace ou plus stable au cours de la tâche. La prise de conscience des régularités contextuelles pourrait permettre au système cognitif de libérer l'organisme de traitements hasardeux de certains stimuli et de « synchroniser » les processus vers une direction plus cohérente, en adéquation avec son environnement (Dehaene & Naccache, 2001). Dehaene et Naccache (2001) proposent que la conscience permet à différents systèmes autonomes de communiquer. Il est certain qu'en donnant du sens à l'environnement visuel, la conscience oriente les processus qui en sont à l'origine même. En accord avec certaines théories, la conscience phénoménale serait nécessaire à la véritable flexibilité. Selon Perruchet et Vinter (2002), la conscience phénoménale offre au système cognitif une représentation globale adaptée à la situation en cours, permettant aux mécanismes d'apprentissage d'être déployés sur les meilleures représentations possibles. Elle serait un pré-requis indispensable pour le

raisonnement analytique et l'abstraction véritable notamment. Dans cette perspective, la fonction centrale de la conscience est d'offrir un contrôle adaptatif flexible sur les comportements (Cleeremans & Jiménez, 2002). L'évolution aurait favorisé l'émergence de systèmes de contrôle unifiés caractérisés par des états de conscience. Quel serait en effet le rôle adaptatif de la conscience dans l'évolution phylogénétique si, comme le propose Reber (1976) ou Lewicki et al. (1992), l'intervention de la conscience ne favorise pas l'apprentissage, voire lui fait obstacle?

Selon O'Reilly et Munakata (2000), l'apprentissage tendrait toujours à être associé à la conscience – même s'il arrive souvent que le contenu de la conscience échoue à être associé à une méta-connaissance, et donc demeure implicite. Ce point de vue amène à se questionner sur un aspect essentiel de nos travaux. Les mots à ignorer dans l'Etude 2 étaient-ils traités à un niveau supraliminaire ou subliminal ? Autrement dit, les participants ont-ils eu une expérience consciente des mots qu'ils devaient ignorer ? A en croire les principes énoncés dans la théorie d'intégration des traits, et si les participants se sont conformés à la consigne, nous sommes tentés de penser que les sujets n'ont pas accédé consciemment à la sémantique des mots ignorés. Néanmoins, rien ne garantit qu'au cours des saccades oculaires, les mots apparaissant dans la couleur à ignorer n'aient pas été traités de manière attentionnelle. De plus, il est possible que l'attention focalisée sur un item attendu permette de percevoir consciemment les mots adjacents. Il s'agit là d'une critique indéniable à la méthodologie utilisée dans les Expériences 6 et 7 (Etude 2). Rien n'assure qu'aucune attention focalisée n'ait été portée sur les mots à ignorer. Toutefois, les performances d'indilage à l'époque 3 étaient similaires dans la condition « prédictive ignorée puis attendue » de l'Expérience 7 et dans la condition « prédictive attendue » de l'Expérience 6. Ce résultat suggère qu'à défaut de pouvoir affirmer de manière certaine qu'aucune attention focalisée n'est requise à l'apprentissage implicite des régularités sémantiques, l'attention visuelle sélective est peu, sinon pas requise pour cet apprentissage.

Par extension, on peut se demander si l'attention sélective est vraiment nécessaire à la conscience. Un grand nombre d'auteurs pensent que nous sommes seulement conscients de ce qui est présent sous le focus attentionnel (pour des revues, cf. Dehaene et al., 2006 ; Merikle & Jordens, 1997). Il est vrai que les recherches sur la cécité au changement montrent qu'un changement sur un objet sera d'autant mieux détecté que l'attention a été récemment et/ou intensément allouée sur cet objet (e.g. Simons & Levin, 1997 ; Mack & Rock, 1998).

L'attention assurerait le maintien des informations en MT, nécessaire à la détection explicite des changements (Luck & Vogel, 1997). Pourtant, focaliser son attention sur un aspect particulier de la scène, n'empêche pas d'avoir une expérience visuelle de l'ensemble de la scène présente dans le champ visuel, certes schématique, mais d'avoir néanmoins conscience de l'ensemble de la scène. L'attention sélective est-elle vraiment nécessaire à la conscience ? Pour Koch et Tsuchiya (2007), non seulement l'attention et la conscience sont distinctes, mais l'attention n'est pas nécessaire pour la conscience. Ils suggèrent qu'on peut faire attention à des objets sans pour autant les percevoir consciemment. A l'inverse, un événement ou un objet peut être perçu consciemment en l'absence de traitement attentionnel. Lamme (2003) fait une distinction entre les « entrées » conscientes vs. non conscientes et les traitements de sélection attentionnelle à une étape indépendante : l'attention sélective ne déterminerait pas si des stimuli atteignent un état de conscience mais déterminerait si les items sont encodés de manière suffisamment stable en MT pour permettre un report ultérieur ou une comparaison avec des scènes subséquentes. Dans ce cadre, la cécité au changement ou la cécité inattentionnelle ne seraient pas des « échecs » de la conscience mais des échecs de la mémoire consciente. En d'autres termes, nous serions conscients de nombreuses entrées visuelles, mais sans attention, cette expérience consciente serait rapidement oubliée et ne pourrait donc être rapportée. En revanche, en assurant une permanence cognitive aux entités visuelles qui sont sous son joug, l'attention sélective permettrait à nos représentations conscientes d'atteindre un état stable et plus durable.

5. Le rôle des régularités contextuelles dans l'analyse de scènes visuelles

Le dernier point important de nos travaux concerne les conséquences des connaissances implicites ou explicites sur les traitements perceptifs et sur le comportement. L'Etude 1 et l'Etude 2 suggèrent que des connaissances implicites relatives à des régularités de l'environnement peuvent faciliter l'exploration d'une image. Plus spécifiquement, nos travaux illustrent comment des connaissances implicites sur des régularités spécifiques ou catégorielles et sémantiques peuvent influencer le guidage de l'attention. Toutefois, nous avons évoqué dans le Chapitre III que les effets classiques d'indigage contextuel pouvaient être interprétés autrement qu'en termes d'effet facilitateur sur le guidage attentionnel. Puisque

l'agencement des distracteurs était aléatoire aussi bien dans les essais prédictifs que non prédictifs, une hypothèse en terme de chunk peut difficilement rendre compte des effets d'indilage observés. D'autre part, les contextes prédictifs étant prédictifs de la localisation de la cible, indépendamment de son identité, et l'identité des distracteurs étant nouvelle au cours de la phase de transfert, un effet facilitateur dans la sélection de la réponse s'avère peu envisageable. Etant donné que les contextes prédictifs vs. non prédictifs ne différaient que par le caractère prédictif vs. non prédictif de la localisation de la cible, les effets d'indilage observés dans les Etudes 1 et 2 semblent bel et bien reposer sur une facilitation dans le guidage de l'attention. Nos résultats sont congruents avec ceux observés par Lambert et Sumich (1996) montrant qu'un indice relatif à une catégorie sémantique peut de manière implicite orienter l'attention sélective vers le champ visuel pertinent pour la tâche. Avec une procédure expérimentale très différente, de nombreuses études ont également montré que les traits sémantiques d'indices verbaux ou picturaux peuvent diriger l'attention vers des objets pertinents dans l'environnement. Par exemple, Moores, Laiti et Chelazzi (2003) ont montré que les liens associatifs sémantiques entre les objets affectent l'orientation de l'attention sélective lors d'une recherche visuelle de cible (voir aussi, Cooper, 1974 ; Huettig & Altmann, 2005 ; Kamide, Altmann, & Haywood, 2003). Comme dans nos travaux, ces études suggèrent que les effets sur les mouvements des yeux dépendent des effets de similarité conceptuelle.

Les études rapportées dans ce travail étayaient les modèles de facilitation contextuelle dans les traitements perceptifs. Nous proposons qu'en participant à la construction des schémas de scène, des connaissances explicites, mais également implicites sur les régularités de l'environnement guident les processus de sélection attentionnelle. Dans cette conception, les connaissances constitutives des schémas de scène ne seraient pas nécessairement accessibles à la conscience. Des expositions répétées aux régularités de l'environnement conduiraient à la mise en mémoire de ces régularités. Par ailleurs, on lit souvent que le guidage descendant de l'attention, via des facteurs cognitifs (connaissances, objectifs) est le fruit d'un contrôle intentionnel. Pourtant, il y a tout lieu de penser qu'une orientation non intentionnelle de l'attention puisse résulter de l'activation de connaissances en mémoire (e.g. Soto, Blanco, Heike, & Humphreys, 2005) Dans nos travaux par exemple, les connaissances implicites relatives aux régularités de l'environnement guidaient l'attention sans que ce guidage soit le fruit d'un contrôle délibéré.

6. Les représentations visuelles

Une vaste littérature sur les phénomènes de cécité au changement et de cécité inattentionnelle suggère que nos représentations visuelles sont loin d'être un calque de la réalité. Certaines théories (cf. Chapitre 1.1) postulent même que nos représentations seraient davantage locales et transitoires. Selon la théorie de la cohérence, la perception visuelle ne repose pas sur une représentation interne stable et détaillée du monde, mais plutôt sur une représentation virtuelle dynamique comportant une structure cohérente et détaillée labile et transitoire. De manière inhérente, la vision est un système dynamique, un système « de l'instant » dont les représentations détaillées sont en constante régénération. En accord avec cette conception, il est probable qu'il n'existe pas de représentations visuelles inactives en MLT. Les représentations visuelles ne sont ni stables, ni immuables. Selon moi, une représentation disparaît à l'instant même où elle émerge, aussitôt remplacée par une nouvelle. Les représentations visuelles ne sont pas figées, elles évoluent à chaque instant du traitement. Les contaminations ou révisions de la mémoire suffisent à donner la certitude que l'activation de connaissances en mémoire se traduit plus par la reconstruction d'un souvenir que par la récupération d'une représentation mentale en mémoire. Néanmoins, chaque représentation visuelle modifie de manière inéluctable la mémoire en laissant une trace, une signature neuronale, stigmate d'un traitement psychologique passé. Aussi, quand bien même nos représentations visuelles seraient-elles éphémères et disparaîtraient à l'instant même où elles émergent, elles ne sont pas sans conséquences sur la mémoire et sur les opérations visuelles ou visuo-motrices ultérieures. Dans quelle mesure ? Quelle information et comment cette information dérivée de nos représentations est-elle retenue en mémoire au cours de l'analyse d'une scène visuelle ? Comment affecte-t-elle les traitements visuels subséquents ? La thèse ici défendue est qu'au cours de l'analyse d'une scène visuelle, le système visuo-cognitif encode explicitement, mais également implicitement, non seulement des régularités perceptives mais aussi des régularités sémantiques.

Par ailleurs, dans les théories visant à expliquer la nature des représentations visuelles, l'attention focalisée est nécessaire à l'encodage et la consolidation des informations visuelles et sémantiques (e.g. Hollingworth & Henderson, 1998 ; Rensink, 2000). Nos travaux montrent pourtant que l'attention sélective n'est pas un pré-requis au développement des mécanismes

d'apprentissage (voir aussi Mack & Rock, 1998). Nous minimisons cependant ces résultats en précisant que nos travaux ne montrent pas qu'aucune forme d'attention ne soit nécessaire à l'apprentissage des régularités sémantiques, mais que de tels mécanismes peuvent se mettre en place sur des régularités qui ne sont pas directement focalisées par l'attention visuelle sélective. Nous continuons néanmoins à penser que l'attention en tant que processus de maintien et de vigilance est nécessaire à la mise en œuvre des mécanismes d'apprentissage.

Nous avons évoqué dans le chapitre I.1. des recherches révélant des effets de détection implicite de changements. Selon certaines conceptions (Cf. Figure 31), la perception explicite recouvre la *prise de conscience visuelle* (au cours de laquelle aucune impression phénoménologique, ou *qualia*, n'émergerait de ce traitement néanmoins conscient) et *l'expérience visuelle* (qui consiste en l'émergence de la phénoménologie). Les deux accompagnent l'expérience consciente à *propos de quelque chose*. Selon Rensink (2004), la perception explicite et la perception implicite mettent en œuvre des processus indépendants, la première impliquant l'intervention d'une attention focalisée, et la deuxième non (Merikle & Joordens, 1997). Cette dissociation fait référence aux opérations de « voir » et de « percevoir ». Ainsi, selon cette conception, à l'inverse de la vision explicite, la vision implicite apparaît quand les stimuli visuels ont été identifiés, mais pas à un niveau qui peut être explicitement rapportable par le sujet. En effet, dans bien des esprits, une représentation est explicite lorsqu'elle est rapportable verbalement. La « sensation de perception » correspondrait à une représentation implicite et par extension à une représentation inconsciente. On peut néanmoins s'interroger sur la valeur de cette dichotomie. Ce que l'on appelle couramment « sensing » correspond-il vraiment à une perception inconsciente ? Si les termes explicite/implicite renvoient de manière « plus consensuelle » à l'aptitude ou l'incapacité à rapporter de manière explicite (i.e. verbale) le contenu des états mentaux, peut-on vraiment dire qu'un percept non verbalisable soit vraiment perçu de manière inconsciente ? La sensation de voir n'est-elle pas une forme de perception consciente ?

Même si l'on admet que la simple sensation de perception est une forme de perception consciente, de nombreuses recherches démontrent l'influence de traitements perceptifs non conscients sur le comportement. Il est admis par le plus grand nombre que certains aspects de la scène ne faisant pas l'objet d'une image mentale sont néanmoins traités par le système visuel et influencent les traitements ultérieurs. Peut-on néanmoins dire que leur traitement donne lieu à une représentation visuelle ? Comment qualifier le contenu psychologique se

rapportant à un objet perçu de manière subliminale ? Peut-on parler de représentation visuelle non consciente ? Le terme représentation visuelle désigne-t-il uniquement le contenu de notre expérience visuelle subjective ? Auquel cas, seuls les objets et les événements de la scène conduisant à l'émergence d'une image mentale seraient représentés ? Quel rôle accorder à ces traitements non conscients dans les représentations visuelles ? Peut-on parler de représentations visuelles non conscientes ?

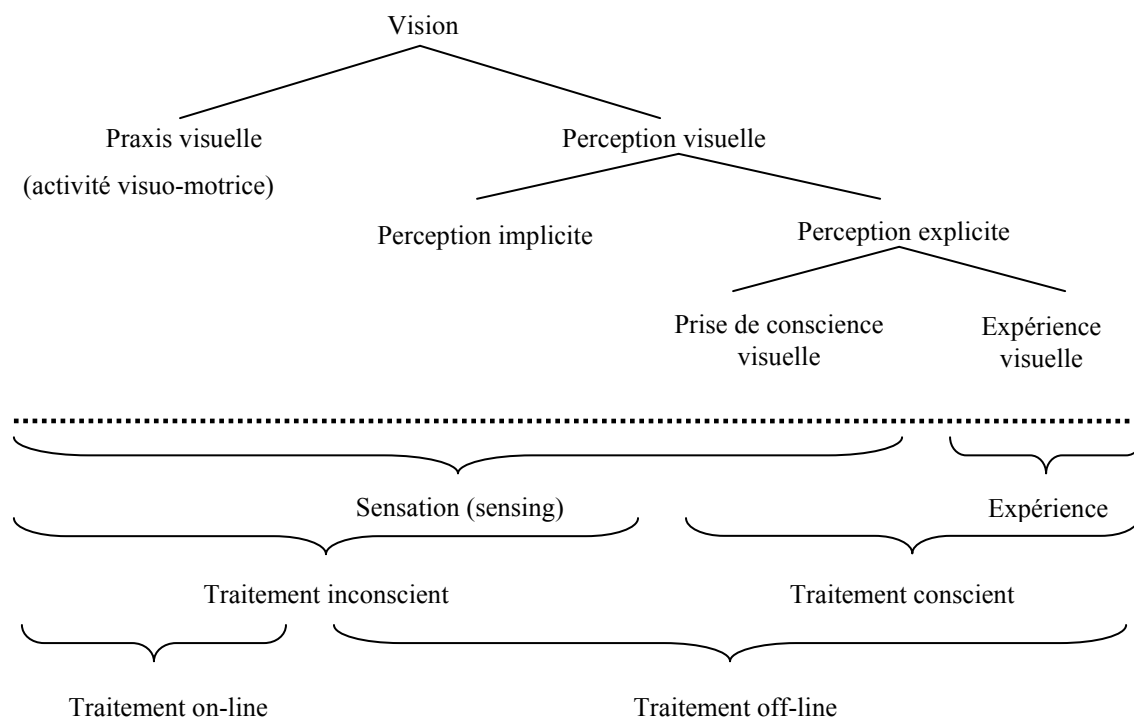


Figure 31. Taxonomie des opérations visuelles. Adapté de Rensink (2000b)

Pour conclure, les travaux de recherche présentés dans ce mémoire de thèse montrent que des connaissances relatives à des régularités spécifiques ou catégorielles et sémantiques peuvent être acquises sans être accompagnées d'une quelconque expérience subjective. Ces résultats offrent des arguments en faveur d'une cognition sans conscience, capable de mettre en mémoire des informations inaccessibles par des tests directs d'accès aux connaissances. De surcroît, ces connaissances inaccessibles à la conscience peuvent être acquises sans que l'attention sélective ne soit focalisée sur les aspects sémantiques des éléments présents dans la scène. Le système visuo-cognitif serait ainsi capable d'encoder des informations perceptives (Jiang & Leung, 2005), mais également des informations sémantiques de l'environnement

extérieures au focus attentionnel. Ces connaissances implicites peuvent influencer les traitements perceptifs en optimisant le guidage de l'attention au sein de l'environnement. Des connaissances implicites sur les régularités sémantiques de l'environnement pourraient ainsi s'inscrire dans la construction des schémas de scène et contribuer à rendre compte de l'adaptabilité des traitements perceptifs, en dépit de la « pauvreté » de nos représentations conscientes. Nos travaux soulignent par ailleurs le rôle adaptatif de la conscience dans l'apprentissage des régularités contextuelles, en montrant que la prise de conscience des régularités de l'environnement favorise l'exploitation de ces régularités.

7. Perspectives de recherches

Les recherches rapportées dans ce mémoire de thèse montrent que des connaissances sémantiques peuvent être acquises de manière implicite au cours de l'analyse d'une scène visuelle. Cependant, les travaux conduits avec le paradigme d'indilage contextuel montrent que lorsque les contextes ne sont plus définis par des affichages arbitraires mais par des scènes du monde réel, les connaissances acquises tendent à être associées à une représentation consciente des régularités de l'environnement. Quelle validité écologique peut-on accorder à l'ensemble des effets d'apprentissage implicite mis en évidence dans les tâches classiques d'indilage contextuel (e.g. Chun & Jiang, 1998) si aucun mécanisme d'apprentissage implicite n'est mis en œuvre dans des environnements naturels ? Peut-on néanmoins observer des effets d'apprentissage implicite dans des scènes du monde réel ? Quels sont les facteurs susceptibles de favoriser la prise de conscience des régularités contextuelles ? Quelle est la valeur adaptative de la conscience dans l'apprentissage de ces régularités ?

Dans la perspective d'aborder ces questions, il s'agit dans un premier temps de dégager ce qui différencie les aires composées de stimuli arbitraires des scènes du monde réel. Faute de mieux, les scènes visuelles peuvent être définies comme des vues sémantiquement cohérentes du monde réel, qui peuvent être nommées et qui présentent un arrière-plan duquel se détachent des objets organisés régulièrement dans l'espace (Henderson, 2003). Il est certes difficile d'assigner cette définition aux plages de stimuli composés d'éléments tel que des L, des formes non dénommables ou même des nombres et des mots. Dans les scènes du monde

réel, les connaissances préalables relatives aux objets présents pourraient fournir à l'observateur une représentation signifiante lui permettant de prendre conscience des régularités contextuelles.

Brockmole et Henderson (2006a) ont proposé que l'extraction rapide d'indices sémantiques présents dans les scènes du monde réel permet de prendre conscience des régularités et que cette prise de conscience des régularités contextuelles accélère l'apprentissage qui sous-tend l'indication contextuelle. Néanmoins, les travaux rapportés dans ce manuscrit montrent que la présence et l'apprentissage de régularités sémantiques ne sont pas nécessairement associés à une prise de conscience de ces régularités. De plus, contrairement aux travaux réalisés avec des scènes du monde réel, les effets d'indication contextuelle reposaient exclusivement sur les propriétés sémantiques des éléments contextuels. La simple présence d'indices sémantiques ne semble donc pas être le seul facteur responsable de la prise de conscience des régularités contextuelles. Un autre facteur susceptible de faciliter la prise de conscience des régularités contextuelles tient à la signification globale de la scène, i.e. le « gist ». L'identification et la dénomination immédiate de la scène visuelle pourraient en effet fournir à l'observateur un cadre de référence signifiant lui permettant d'associer de manière explicite la signification générale de la scène avec la localisation de la cible. Mais la prise de conscience des régularités contextuelles pourrait aussi tenir à l'extraction de facteurs purement visuels. En effet, dans les travaux conduits avec le paradigme d'indication contextuelle à partir de scènes du monde réel, les scènes étaient très hétérogènes en termes de leurs traits visuels. L'organisation des formes et des couleurs était très différente d'une scène à l'autre. En revanche, dans les travaux classiques d'indication contextuelle, les affichages de recherche sont composés d'éléments très similaires visuellement et ces derniers sont utilisés de manière systématique à chaque essai. Par exemple, dans la tâche classique d'indication contextuelle, les distracteurs ne sont que des L. Il n'est donc pas exclu que des facteurs purement visuels puissent à eux seuls favoriser la prise de conscience des régularités. La singularité perceptive des scènes du monde réel utilisées, dérivée de l'extrême hétérogénéité des traits visuels constitutifs, pourrait à elle seule permettre à l'observateur de rapidement encoder en mémoire explicite les scènes répétées.

Dans ce cadre, on peut s'interroger sur le rôle joué par les régularités spécifiques dans les effets d'indication. Dans quelle mesure des effets d'apprentissage explicite vs. implicite peuvent-ils être observés dans des scènes du monde réel lorsque les régularités manipulées

reposent exclusivement sur les propriétés perceptives spécifiques de scènes systématiquement répétées ? Réciproquement, quel rôle joue la catégorie de scène, lorsque les traits perceptifs relatifs à l'agencement des formes et des couleurs ne sont pas prédictifs ? Outre les indicateurs comportementaux, des indicateurs de neuroimagerie pourraient aider à valider l'existence d'apprentissages de nature différente.

Après avoir été longtemps censurée comme objet scientifique, l'étude de la conscience et du non conscient connaît un regain d'intérêt depuis plusieurs années. L'émergence relativement récente en psychologie cognitive du problème de l'inconscient et de la conscience s'accompagne d'une littérature foisonnante et de modèles théoriques de plus en plus circonstanciés. En dépit des débats terminologiques et des désaccords méthodologiques rencontrés dans ces domaines de recherche, la nécessité fondamentale d'aborder le problème de la conscience et du non conscient par une approche scientifique ne fait plus aucun doute. Par ailleurs, les techniques d'imagerie cérébrale ont connu des avancées considérables ces dernières années, offrant des indicateurs nouveaux et pertinents pour aborder les processus cognitifs et dégager une compréhension unifiée des phénomènes psychiques et cérébraux. Ces techniques sont-elles pour autant pertinentes dans l'étude de la conscience ? L'idée portée par les tenants des neurosciences est que les études en neuroimagerie constituent la voie royale vers une meilleure compréhension de la conscience (e.g. Koch & Crick, 2004). On ne peut cependant passer outre les limites méthodologiques propres aux études de neuroimagerie et du « fossé explicatif » qui sépare les processus physiques et les aspects qualitatifs de nos expériences. Pour de nombreux philosophes, la conscience phénoménologique n'est pas de l'ordre des phénomènes physiques, ce qui limite considérablement l'importance des corrélats neuronaux de la conscience. Dans ce sens, l'approche neuroscientifique de la conscience sera toujours insatisfaisante. Si la compréhension de la conscience ne se limite évidemment pas à savoir si le lieu d'interaction cerveau-esprit est la glande pinéale, le cortex frontal, le thalamus ou encore l'amygdale (il est d'ailleurs peu probable qu'il existe une structure cérébrale qui soit à elle seule responsable de la conscience), les techniques de neuro-imagerie offrent néanmoins des indicateurs du comportement et de la cognition, au même titre que les indicateurs classiques de la psychologie cognitive. Une approche pluricom pétente des sciences cognitives nous semble pertinente pour dégager une compréhension unifiée de la cognition.

RÉFÉRENCES

- Abrams, R. L., & Greenwald, A. G. (2000). Parts outweigh the whole (word) in unconscious analysis of meaning. *Psychological Science*, 11(2), 118-124.
- Altmann, G. T. M., Dienes, Z. n., & Goode, A. (1995). Modality independence of implicitly learned grammatical knowledge. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21(4), 899-912.
- Anstis, S. (1998). Picturing peripheral acuity. *Perception*, 27(7), 817-825.
- Antes, J. R., Penland, J. G., & Metzger, R. L. (1981). Processing global information in briefly presented pictures. *Psychological Research*, 43(3), 277-292.
- Baddeley, A. D. (1993). *La mémoire humaine*, Grenoble, PUG.
- Becker, S., Moscovitch, M., Behrmann, M., & Joordens, S. (1997). Long-term semantic priming: A computational account and empirical evidence. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 23(5), 1059-1082.
- Berry, D. C., & Broadbent, D. E. (1984). On the relationship between task performance and associated verbalizable knowledge. *The Quarterly Journal of Experimental Psychology A: Human Experimental Psychology*, 36(2), 209-231.
- Berry, D. C., & Dienes, Z. n. (1993). *Implicit learning: Theoretical and empirical issues*. Hillsdale, NJ, England: Lawrence Erlbaum Associates, Inc.
- Besner, D., & Stolz, J. A. (1999). Unconsciously controlled processing: The Stroop effect reconsidered. *Psychonomic Bulletin & Review*, 6(3), 449-455.
- Biederman, I. (1972). Perceiving real-world scenes. *Science*, 177(4043), 77-80.
- Biederman, I., Glass, A. L., & Stacy, E. W. (1973). Searching for objects in real-world scenes. *Journal of Experimental Psychology*, 97(1), 22-27.
- Biederman, I., (1988). Aspects and extensions of a theory of human image understanding. In Z. W. Pylyshyn (Ed.), *Computational processes in human vision: An interdisciplinary perspective* (pp. 370-428). Norwood, NJ: Ablex.

- Biederman, I., Mezzanotte, R. J., & Rabinowitz, J. C. (1982). Scene perception: Detecting and judging objects undergoing relational violations. *Cognitive Psychology*, 14(2), 143-177.
- Biederman, I., Rabinowitz, J. C., Glass, A. L., & Stacy, E. W. (1974). On the information extracted from a glance at a scene. *Journal of Experimental Psychology*, 103(3), 597-600.
- Biederman, I., Teitelbaum, R. C., & Mezzanotte, R. J. (1983). Scene perception: A failure to find a benefit from prior expectancy or familiarity. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 9(3), 411-429.
- Boyce, S. J., & Pollatsek, A. (1992). Identification of objects in scenes: The role of scene background in object naming. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 18(3), 531-543.
- Boyce, S. J., Pollatsek, A., & Rayner, K. (1989). Effect of background information on object identification. *Journal of Experimental Psychology: Human Perception and Performance*, 15(3), 556-566.
- Brady, T. F., & Chun, M. M. (2007). Spatial constraints on learning in visual search: Modeling contextual cuing. *Journal of Experimental Psychology: Human Perception and Performance*, 33(4), 798-815.
- Broadbent, D. E. (1958). *Perception and communication*. Elmsford, NY, US: Pergamon Press.
- Broadbent, D. E. (1977). Levels, hierarchies, and the locus of control. *Quarterly Journal of Experimental Psychology*, 29, 181-201.
- Brockmole, J. R., & Henderson, J. M. (2006a). Using real-world scenes as contextual cues for search. *Visual Cognition*, 13(1), 99-108.
- Brockmole, J., & Henderson, J. (2006b). Recognition and attention guidance during contextual cueing in real-world scenes: Evidence from eye movements. *Quarterly Journal of Experimental Psychology*, 59(7), 1177-1187.
- Brockmole, J. R., Castelano, M. S., & Henderson, J. M. (2006). Contextual Cueing in Naturalistic Scenes: Global and Local Contexts. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 32(4), 699-706.
- Brooks, L. R. (1978). Nonanalytic concept formation and memory for instances. In E. Rosch & B. B. Lloyd (Eds.), *Cognition and Categorization* (pp. 169-215). Hillsdale, NJ: Erlbaum.
- Brooks, L. R., & Vokey, J. R. (1991). Abstract analogies and abstracted grammars: Comments on Reber (1989) and Mathews et al. (1989). *Journal of Experimental Psychology: General*, 120(3), 316-323.
- Brysbaert, M. (1995). Arabic number reading: On the nature of the numerical scale and the origin of phonological recoding. *Journal of Experimental Psychology: General*, 124(4), 434-452.

- Buswell, G. T. (1935). *How people look at pictures: a study of the psychology and perception in art*. Oxford, England: Univ. Chicago Press.
- Camus, J.-F. (2003). L'attention et ses modèles. *Psychologie Française*, 48(1), 5-18.
- Castelhano, M. S., & Henderson, J. M. (2007). Initial scene representations facilitate eye movement guidance in visual search. *Journal of Experimental Psychology: Human Perception and Performance*, 33(4), 753-763.
- Chapman, P. R., & Underwood, G. (1998). *Visual search of dynamic scenes: events types and the role of experience in viewing driving situations*. In: G. Underwood, Editor, *Eye guidance in reading and scene perception*, Elsevier, Oxford (1988), pp. 369-394.
- Chelazzi, L., Duncan, J., Miller, E. K., & Desimone, R. (1998). Responses of neurons in inferior temporal cortex during memory-guided visual search. *Journal of Neurophysiology*, 80(6), 2918-2940.
- Chelazzi, L., Miller, E. K., Duncan, J., & Desimone, R. (1993). A neural basis for visual search in inferior temporal cortex. *Nature*, 363(6427), 345-347.
- Chen, X., & Zelinsky, G. J. (2006). Real-world visual search is dominated by top-down guidance. *Vision Research*, 46(24), 4118-4133.
- Chua, K.-P., & Chun, M. M. (2003). Implicit scene learning is viewpoint dependent. *Perception & Psychophysics*, 65(1), 72-80.
- Chun, M. M., & Jiang, Y. (1998). Contextual cueing: Implicit learning and memory of visual context guides spatial attention. *Cognitive Psychology*, 36(1), 28-71.
- Chun, M. M., & Jiang, Y. (1999). Top-down attentional guidance based on implicit learning of visual covariation. *Psychological Science*, 10(4), 360-365.
- Chun, M. M., & Jiang, Y. (2003). Implicit, long-term spatial contextual memory. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 29(2), 224-234.
- Chun, M. M., & Phelps, E. A. (1999). Memory deficits for implicit contextual information in amnesic subjects with hippocampal damage. *Nature Neuroscience*, 2, 844-847.
- Chun, M. M., & Turk-Browne, N. B. (2007). Interactions between attention and memory. *Current Opinion in Neurobiology*, 17(2), 177-184.
- Chun, M. M., Wolfe, J. M. (2001). *Visual attention*. In Goldstein, B. (Ed), *Blackwell Handbook of Perception* (pp. 272-310). Oxford, UK: Blackwell Publishers Ltd.
- Cleeremans, A., Destrebecqz, A., & Boyer, M. (1998). Implicit learning: News from the front. *Trends in Cognitive Sciences*, 2(10), 406-416.
- Cleeremans, A., & Jiménez, L. (2002). *Implicit learning and consciousness: A graded, dynamical perspective*. In R. M. French & A. Cleeremans (Eds.), *Implicit learning and consciousness* (pp. 1-40): Psychology Press.
- Cohen, A., Ivry, R. I., & Keele, S. W. (1990). Attention and structure in sequence learning.

- Journal of Experimental Psychology: Learning, Memory, and Cognition*, 16(1), 17-30.
- Cohen, N. J., & Squire, L. R. (1980). Preserved learning and retention of pattern-analyzing skill in amnesia: Dissociation of knowing how and knowing that. *Science*, 210(4466), 207-210.
- Conway, C. M., & Christiansen, M. H. (2005). Modality-Constrained Statistical Learning of Tactile, Visual, and Auditory Sequences. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 31(1), 24-39.
- Cowan, N. (1988). Evolving conceptions of memory storage, selective attention, and their mutual constraints within the human information-processing system. *Psychological Bulletin*, 104(2), 163-191.
- Curran, T., & Keele, S. W. (1993). Attentional and nonattentional forms of sequence learning. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 19(1), 189-202.
- Damian, M. F. (2001). Congruity effects evoked by subliminally presented primes: Automaticity rather than semantic processing. *Journal of Experimental Psychology: Human Perception and Performance*, 27, 154-165.
- Davenport, J. L., & Potter, M. C. (2004). Scene Consistency in Object and Background Perception. *Psychological Science*, 15(8), 559-564.
- De Gref, P., Christiaens, D., & d'Ydewalle, G. (1990). Perceptual effects of scene context on object identification. *Psychological Research*, 52, 317-329.
- Dehaene, S., & Naccache, L. (2001). Towards a cognitive neuroscience of consciousness: Basic evidence and a workspace framework. *Cognition*, 79(1), 1-37.
- Dehaene, S., Bossini, S., & Giraux, P. (1993). The mental representation of parity and number magnitude. *Journal of Experimental Psychology: General*, 122(3), 371-396.
- Dehaene, S., Changeux, J.-P., Naccache, L., Sackur, J. r. m., & Sergent, C. (2006). Conscious, preconscious, and subliminal processing: A testable taxonomy. *Trends in Cognitive Sciences*, 10(5), 204-211.
- Dell'Acqua, R., & Grainger, J. (1999). Unconscious semantic priming from pictures. *Cognition*, 73(1), 1-15.
- DeSchepper, B., & Treisman, A. (1996). Visual memory for novel shapes: Implicit coding without attention. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 22(1), 27-47.
- Desimone, R., & Duncan, J. (1995). Neural mechanisms of selective visual attention. *Annual Review of Neuroscience*, 18, 193-222.
- Destrebecqz, A. (2000). *Mesures directes et indirectes de l'apprentissage implicite. Etude expérimentale et modélisation*. Université Libre de Bruxelles, Bruxelles.
- Destrebecqz, A., & Cleeremans, A. (2001). Can sequence learning be implicit? New evidence

- with the process dissociation procedure. *Psychonomic Bulletin & Review*, 8(2), 343-350.
- Didierjean, A. (2004). Quelques facettes de l'expertise. Habilitation à Diriger des Recherches.
- Didierjean, A. (2001). Apprendre à partir d'exemples: Abstraction de règles et/ou mémoire d'exemplaires? *L'Année psychologique*, 101, 325-348.
- Didierjean, A. (2007). Parity Effects in an Implicit Learning Task. *British Journal of Psychology*.
- Didierjean, A., & Lemaire, P. (2005). Learning rules unconsciously while playing loto: A new task and its application to assess younger and older adults' implicit learning. *Current Psychology of Cognition*, 23, 249-270.
- Dienes, Z., & Altmann, G. T. M. (1997). Transfer of implicit knowledge across domains? How implicit and how abstract? In D. Berry (Ed), *How implicit is implicit learning?* (pp. 107-123). Oxford: Oxford University Press.
- Dienes, Z., & Perner, J. (1999). A theory of implicit and explicit knowledge. *Behavioral and Brain Sciences*, 22(5), 735-808.
- Dienes, Z., Broadbent, D., & Berry, D. C. (1991). Implicit and explicit knowledge bases in artificial grammar learning. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 17(5), 875-887.
- Draine, S. C., & Greenwald, A. G. (1998). Replicable unconscious semantic priming. *Journal of Experimental Psychology: General*, 127(3), 286-303.
- Dulany, D. E., Carlson, R. A., & Dewey, G. I. (1984). A case of syntactical learning and judgment: How conscious and how abstract? *Journal of Experimental Psychology: General*, 113(4), 541-555.
- Duncan, J., & Humphreys, G. W. (1989). Visual search and stimulus similarity. *Psychological Review*, 96(3), 433-458.
- Duncan, J., Martens, S., & Ward, R. (1997). Restricted attentional capacity within but not between sensory modalities. *Nature*, 379(6567), 808-810.
- Egeth, H. E., & Yantis, S. (1997). Visual attention: Control, representation, and time course. *Annual Review of Psychology*, 48, 269-297.
- Endo, N., & Takeda, Y. (2004). Selective learning of spatial configuration and object identity in visual search. *Perception & Psychophysics*, 66(2), 293-302.
- Endo, N., & Takeda, Y. (2005). Use of spatial context is restricted by relative position in implicit learning. *Psychonomic Bulletin & Review*, 12(5), 880-885.
- Epstein, R., & Kanwisher, N. (1998). A cortical representation of the local visual environment. *Nature*, 392(6676), 598-601.
- Ericsson, K. A., & Simon, H. A. (1984). *Protocol analysis: Verbal reports as data*.

Cambridge, MA, US: The MIT Press.

- Feldman, J. A. (1985). Four frames suffice: A provisional model of vision and space. *Behavioral & Brain Sciences*, 8, 265-289.
- Fernandez-Duque, D., & Thornton, I. M. (2000). Change detection without awareness: Do explicit reports underestimate the representation of change in the visual system? *Visual Cognition*, 7(1), 323-344.
- Fernandez-Duque, D., & Thornton, I. M. (2003). Explicit mechanisms do not account for implicit localization and identification of change: An empirical reply to Mitroff et al. (2002). *Journal of Experimental Psychology: Human Perception and Performance*, 29(5), 846-858.
- Fias, W., Brysbaert, M., Geypens, F., & D'Ydewalle, G. (1996). The Importance of Magnitude Information in Numerical Processing: Evidence from the SNARC Effect. *Mathematical Cognition*, 2(1), 95-110.
- Fiser, J. z., & Aslin, R. N. (2001). Unsupervised statistical learning of higher-order spatial structures from visual scenes. *Psychological Science*, 12(6), 499-504.
- Fiser, J. z., & Aslin, R. N. (2002). Statistical learning of higher-order temporal structure from visual shape sequences. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 28(3), 458-467.
- Frensch, P. A. (1998). One concept, multiple meanings: How to define the concept of "implicit learning". In M. S. P. A. F. (Hg.) (Ed), *Handbook of implicit learning* (pp. 47-104). Thousand Oaks, CA: Sage Publications.
- Frensch, P. A., Buchner, A., & Lin, J. (1994). Implicit learning of unique and ambiguous serial transitions in the presence and absence of a distractor task. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 20(3), 567-584.
- Frensch, P. A., Lin, J., & Buchner, A. (1998). Learning versus behavioral expression of the learned: The effects of a secondary tone-counting task on implicit learning in the serial reaction task. *Psychological Research*, 61(2), 83.
- Frensch, P. A., Wenke, D., & R nger, D. (1999). A secondary tone-counting task suppresses expression of knowledge in the serial reaction task. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 25(1), 260-274.
- Friedman, A. (1979). Framing pictures: The role of knowledge in automatized encoding and memory for gist. *Journal of Experimental Psychology: General*, 108(3), 316-355.
- Gordon, R. D. (2004). Attentional Allocation During the Perception of Scenes. *Journal of Experimental Psychology: Human Perception and Performance*, 30(4), 760-777.
- Goschke, T. (1996). *Implicit learning of stimulus and response sequences: evidence for independent learning systems*. Paper presented at Max-Planck Institute for Cognitive Neuroscience, Leipzig, Germany.
- Goschke, T. (1997). Implicit learning and unconscious knowledge: mental representation,

- computational mechanisms, and brain structures. In K. Lamberts & D. Shanks (Eds.), *Knowledge, concepts and categories* (pp. 247-33). Hove, UK: Psychology Press.
- Goujon, A., Didierjean, A., & Marmèche, E. (2007). Contextual cueing based on specific and categorical properties of the environment. *Visual Cognition*, 15(3), 257-275.
- Greenwald, A. G., Abrams, R. L., Naccache, L., & Dehaene, S. (2003). Long-term semantic memory versus contextual memory in unconscious number processing. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 29(2), 235-247.
- Gross, M., & Feldman, R. (1995). National Transportation Statistics: 1995 (DOT-VNTSC-BTS-94-3). Washington, DC: U.S. DOT/Bureau of Transportation Statistics.
- Gimes, J. (1996). On the failure to detect changes in scenes across saccades. In K. Akins (Ed), *Perception: Vancouver studies in cognitive science* (Vol. 5, pp. 89-110). Oxford, England: Oxford University Press.
- Hayes, N. A., & Broadbent, D. (1988). Two modes of learning for interactive tasks. *Cognition*, 28(3), 249-276.
- Hayhoe, M. (2000). Vision using routines: A functional account of vision. *Visual Cognition*, 7(1-3), 43-64.
- Helman, S., & Berry, D. C. (2003). Effects of divided attention and speeded responding on implicit and explicit retrieval of artificial grammar knowledge. *Memory & Cognition*, 31(5), 703-714.
- Henderson, J. M. (2003). Human gaze control during real-world scene perception. *Trends in Cognitive Sciences*, 7, 498-504.
- Henderson, J. M., & Hollingworth, A. (1999). High-level scene perception. *Annual Review of Psychology*, 50(1), 243-271.
- Henderson, J. M., & Hollingworth, A. (2003). Eye movements and visual memory: Detecting changes to saccade targets in scenes. *Perception & Psychophysics*, 65(1), 58-71.
- Henderson, J. M., Brockmole, J. R., Castelhamo, M. S., Mack, M., van Gompel, R. P. G., Fischer, M. H., et al. (2007). *Visual saliency does not account for eye movements during visual search in real-world scenes*. Amsterdam, Netherlands: Elsevier.
- Henderson, J. M., Pollatsek, A., & Rayner, K. (1987). Effects of foveal priming and extrafoveal preview on object identification. *Journal of Experimental Psychology: Human Perception and Performance*, 13(3), 449-463.
- Henderson, J. M., Weeks, P. A., Jr., & Hollingworth, A. (1999). The effects of semantic consistency on eye movements during complex scene viewing. *Journal of Experimental Psychology: Human Perception and Performance*, 25(1), 210-228.
- Heuer, H., & Schmidtke, V. (1996). Secondary-task effects on sequence learning. *Psychological Research/Psychologische Forschung*, 59(2), 119-133.
- Hoffman, J., & Seibald, A. (2005). When obvious covariations are not even learned implicitly.

European Journal of Cognitive Psychology, 17(4), 449-480.

- Holender, D. (1986). Semantic activation without conscious identification in dichotic listening, parafoveal vision, and visual masking: A survey and appraisal. *Behavioral and Brain Sciences*, 9(1), 1-66.
- Hollingworth, A. (2004). Constructing Visual Representations of Natural Scenes: The Roles of Short- and Long-Term Visual Memory. *Journal of Experimental Psychology: Human Perception and Performance*, 30(3), 519-537.
- Hollingworth, A., & Henderson, J. M. (1998). Does consistent scene context facilitate object perception? *Journal of Experimental Psychology: General*, 127(4), 398-415.
- Hollingworth, A., & Henderson, J. M. (2000). Semantic Informativeness Mediates the Detection of Changes in Natural Scenes. *Visual Cognition*, 7(1/2/3), 213-235.
- Hollingworth, A., & Henderson, J. M. (2002). Accurate visual memory for previously attended objects in natural scenes. *Journal of Experimental Psychology: Human Perception & Performance*, 28(1), 113-136.
- Hollingworth, A., Williams, C. C., & Henderson, J. M. (2001). To see and remember: Visually specific information is retained in memory from previously attended objects in natural scenes. *Psychonomic Bulletin & Review*, 8(4), 761-768.
- Howard, J. H., Jr., Howard, D. V., Dennis, N. A., Yankovich, H., & Vaidya, C. J. (2004). Implicit Spatial Contextual Learning in Healthy Aging. *Neuropsychology*, 18(1), 124-134.
- Huang, L. (2006). Contextual cuing based on spatial arrangement of color. *Perception & Psychophysics*, 68(5), 792-799.
- Huetting, F., & Altmann, G. T. M. (2005). Word meaning and the control of eye fixation: Semantic competitor effects and the visual world paradigm. *Cognition*, 96(1), 23-32.
- Intraub, H. (1997). The representation of visual scenes. *Trends in Cognitive Sciences*, 1(6), 217-222.
- Irwin, D. E. (1991). Information integration across saccadic eye movements. *Cognitive Psychology*, 23(3), 420-456.
- Irwin, D. E. (1992). Memory for position and identity across eye movements. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 18(2), 307-317.
- Irwin, D. E., & Zelinsky, G. J. (2002). Eye movements and scene perception: Memory for things observed. *Perception & Psychophysics*, 64(6), 882-895.
- Irwin, D. E., Andrews, R. V. (1996). *Integration and accumulation of information across saccadic eye movements*. Cambridge, MA, US: The MIT Press.
- Irwin, D. E., Yantis, S., & Joindes, J. (1983). Evidence against visual integration across saccadic eye movements. *Perception & Psychophysics*, 34(1), 49-57.

- Itti, L., & Koch, C. (2000). A saliency-based search mechanism for overt and covert shifts of visual attention. *Vision Research*, 40(10), 1489-1506.
- Jacoby, L. L. (1991). A process dissociation framework: Separating automatic from intentional uses of memory. *Journal of Memory and Language*, 30(5), 513-541.
- James, W. (1890). *The principles of Psychology*. New York: Holt.
- Jiang, Y., & Chun, M. M. (2001). Selective attention modulates implicit learning. *Quarterly Journal of Experimental Psychology: Section A*, 54(4), 1105-1124.
- Jiang, Y., & Leung, A. W. (2005). Implicit learning of ignored visual context. *Psychonomic Bulletin & Review*, 12(1), 100-106.
- Jiang, Y., & Song, J.-H. (2005a). Hyperspecificity in Visual Implicit Learning: Learning of Spatial Layout Is Contingent on Item Identity. *Journal of Experimental Psychology: Human Perception and Performance*, 31(6), 1439-1448.
- Jiang, Y., & Song, J.-H. (2005b). Spatial context learning in visual search and change detection. *Perception & Psychophysics*, 67(7), 1128-1139.
- Jiang, Y., & Wagner, L. C. (2004). What is learned in spatial contextual cuing-configuration or individual locations? *Perception & Psychophysics*, 66(3), 454-463.
- Jiang, Y., Song, J.-H., & Rigas, A. (2005). High-capacity spatial contextual memory. *Psychonomic Bulletin & Review*, 12(3), 524-529.
- Jiménez, L., & Méndez, C. s. (1999). Which attention is needed for implicit sequence learning? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 25(1), 236-259.
- Jiménez, L., Méndez, C. s., & Cleeremans, A. (1996). Comparing direct and indirect measures of sequence learning. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 22(4), 948-969.
- Jonides, J., & Yantis, S. (1988). Uniqueness of abrupt visual onset in capturing attention. *Perception & Psychophysics*, 43(4), 346-354.
- Jungé, J. A., Scholl, B. J., & Chun, M. M. (2007). How is spatial context learning integrated over signal versus noise? A primacy effect in contextual cuing. *Visual Cognition*, 15(1), 1-11.
- Kahneman, D. (1973). *Attention and Effort*, Londres, Prentice Hal.
- Kahneman, D., Treisman, A., & Gibbs, B. J. (1992). The reviewing of object files: Object-specific integration of information. *Cognitive Psychology*, 24(2), 175-219.
- Kamide, Y., Altmann, G. T. M., & Haywood, S. L. (2003). The time-course of prediction in incremental sentence processing: Evidence from anticipatory eye movements. *Journal of Memory and Language*, 49(1), 133-156.
- Kawahara, J.-i. (2003). Contextual cuing in 3D layouts defined by binocular disparity. *Visual*

Cognition, 10(7), 837-852.

- Klauer, K. C., Eder, A. B., Greenwald, A. G., & Abrams, R. L. (2007). Priming of semantic classifications by novel subliminal prime words. *Consciousness and Cognition: An International Journal*, 16(1), 63-83.
- Koch, C., & Crick, F. (2004). The neuronal basis of visual consciousness. In L. M. Chalupa & J. S. Werner (Eds), *The visual neuroscience*: MIT Press.
- Koch, C., & Tsuchiya, N. (2007). Attention and consciousness: Two distinct brain processes. *Trends in Cognitive Sciences*, 11(1), 16-22.
- Kunar, M. A., Flusberg, S. J., & Wolfe, J. M. (2006). Contextual Cuing by Global Features. *Perception & Psychophysics*, 68(7), 1204-1216.
- Kunar, M. A., Flusberg, S., Horowitz, T. S., & Wolfe, J. M. (2007). Does contextual cuing guide the deployment of attention? *Journal of Experimental Psychology: Human Perception and Performance*, 33(4), 816-828.
- Lambert, A. J., & Sumich, A. L. (1996). Spatial orienting controlled without awareness: A semantically based implicit learning effect. *The Quarterly Journal of Experimental Psychology A: Human Experimental Psychology*, 49(2), 490-518.
- Lamme, V. A. F. (2003). Why visual attention and awareness are different. *Trends in Cognitive Sciences*, 7(1), 12-18.
- Lamy, D., Leber, A., & Egeth, H. E. (2004). Effects of Task Relevance and Stimulus-Driven Saliency in Feature-Search Mode. *Journal of Experimental Psychology: Human Perception and Performance*, 30(6), 1019-1031.
- Land, M. F., & Hayhoe, M. (2001). In what ways do eye movements contribute to everyday activities? *Vision Research*, 41(25), 3559-3565.
- Levin, D. T., & Simons, D. J. (1997). Failure to detect changes to attended objects in motion pictures. *Psychonomic Bulletin & Review*, 4(4), 501-506.
- Lewicki, P. (1986). *Nonconscious social information processing*. San Diego, CA, US: Academic Press.
- Lewicki, P., Hill, T., & Czyzewska, M. (1992). Nonconscious acquisition of information. *American Psychologist*, 47(6), 796-801.
- Libera, C. D., & Chelazzi, L. (2006). Visual Selective Attention and the Effects of Monetary Rewards. *Psychological Science*, 17(3), 222-227.
- Lleras, A., & Von Mühlenen, A. (2004). Spatial context and top-down strategies in visual search. *Spatial Vision*, 17(4/5), 465-482.
- Loftus, G. R. (1972). Eye fixations and recognition memory for pictures. *Cognitive Psychology*, 3(4), 525-551.
- Loftus, G. R., & Mackworth, N. H. (1978). Cognitive determinants of fixation location during

- picture viewing. *Journal of Experimental Psychology: Human Perception and Performance*, 4(4), 565-572.
- Logan, G. D. (1988). Toward an instance theory of automatization. *Psychological Review*, 95(4), 492-527.
- Logan, G. D. (2002). An instance theory of attention and memory. *Psychological Review*, 109(2), 376-400.
- Logan, G. D. (2004). Cumulative progress in formal theories of attention. *Annual Review of Psychology*, 55(1), 207-234.
- Luck, S. J., & Vogel, E. K. (1997). The capacity of visual working memory for features and conjunctions. *Nature*, 390(6657), 279-281.
- Mack, A., & Rock, I. (1998). *Inattention blindness*. Cambridge, MA, US: The MIT Press.
- Mackworth, N. H., & Morandi, A. J. (1967). The gaze selects informative details within pictures. *Perception & Psychophysics*, 2(11), 547-552.
- Maljkovic, V., & Nakayama, K. (1994). Priming of pop-out: I. Role of features. *Memory & Cognition*, 22(6), 657-672.
- Maljkovic, V., & Nakayama, K. (1996). Priming of pop-out: II. The role of position. *Perception & Psychophysics*, 58(7), 977-991.
- Maljkovic, V., & Nakayama, K. (2000). Priming of popout: III. A short-term implicit memory system beneficial for rapid target selection. *Visual Cognition*, 7(5), 571-595.
- Mandler, J. M., & Parker, R. E. (1976). Memory for descriptive and spatial information in complex pictures. *Journal of Experimental Psychology: Human Learning & Memory*, 2(1), 38-48.
- Mandler, J. M., & Ritchey, G. H. (1977). Long-term memory for pictures. *Journal of Experimental Psychology: Human Learning & Memory*, 3(4), 386-396.
- Mannan, S., Ruddock, K. H., Wright, J. R., Findlay, J. M., Walker, R., & Kentridge, R. W. (1995). *Eye movements and response times for the detection of line orientation during visual search*. New York, NY, US: Elsevier Science.
- Manns, J. & Squire, L. R. (2001). Perceptual learning, awareness, and the hippocampus. *Hippocampus*, 11, 776-782.
- Marcel, A. J. (1983). Conscious and unconscious perception : experiments on visual masking and word recognition. *Cognitive Psychology*, 15 (2), 197-237.
- Mathews, R. C., Buss, R. R., Stanley, W. B., & Blanchard-Fields, F. (1989). Role of implicit and explicit processes in learning from examples: A synergistic effect. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 15(6), 1083-1100.
- Mayr, U. (1996). Spatial attention and implicit sequence learning: evidence for independent learning of spatial and nonspatial sequences. *Journal of Experimental Psychology:*

Learning, Memory, and Cognition, 22, 350-364.

- McClelland, J. L., & Rogers, T. T. (2003). The Parallel Distributed Processing Approach to Semantic Cognition. *Nature Reviews Neuroscience*, 4(4), 310-322.
- McCotter, M., Gosselin, F., Sowden, P., & Schyns, P. (2005). The use of visual information in natural scenes. *Visual Cognition*, 12(6), 938-953.
- Merikle, P. M., & Daneman, M. (1998). Psychological investigations of unconscious perception. *Journal of Consciousness Studies*, 5(1), 5-18.
- Merikle, P. M., & Joordens, S. (1997). Parallels between perception without attention and perception without awareness. *Consciousness and Cognition: An International Journal*, 6(2), 219-236.
- Metzger, R. L., & Antes, J. R. (1983). The nature of processing early in picture perception. *Psychological Research*, 45(3), 267-274.
- Meulemans, T. (1998). Apprentissage implicite, mémoire implicite et développement. *Psychologie Française*, 43(1), 27-37.
- Michon J. (1985). *A critical view of driver behavior models: what do we know, what should we do?* In Evans L., Schwing R. Eds: Human Behavior and Traffic Safety. New York: Plenum Press
- Mitroff, S. R., Simons, D. J., & Franconeri, S. L. (2002). The siren song of implicit change detection. *Journal of Experimental Psychology: Human Perception and Performance*, 28(4), 798-815.
- Moore, C. M., & Egeth, H. (1997). Perception without attention: Evidence of grouping under conditions of inattention. *Journal of Experimental Psychology: Human Perception and Performance*, 23(2), 339-352.
- Moores, E., Laiti, L., & Chelazzi, L. (2003). Associative knowledge controls deployment of visual selective attention. *Nature Neuroscience*, 6(2), 182.
- Naccache, L., & Dehaene, S. (2001). Unconscious semantic priming extends to novel unseen stimuli. *Cognition*, 80(3), 215-229.
- Navalpakkam, V., & Itti, L. (2005). Modeling the influence of task on attention. *Vision Research*, 45(2), 205-231.
- Navon, D. (1977). Forest before trees: The precedence of global features in visual perception. *Cognitive Psychology*, 9(3), 353-383.
- Nicolas, S. (1994). Reflexions autour du concept de mémoire implicite, *L'Année psychologique*, 94, 67-85.
- Nicolas, S. (1996). L'apprentissage implicite: le cas des grammaires artificielles. *L'Année psychologique*, 96, 459-493.
- Nissen, M. J., & Bullemer, P. (1987). Attentional requirements of learning: Evidence from

- performance measures. *Cognitive Psychology*, 19(1), 1-32.
- Noë, A., & O'Regan, J. K. (2000). Perception, attention and the grand illusion. *Psyche: An Interdisciplinary Journal of Research on Consciousness*, 6(15), NP.
- Noles, N. S., Scholl, B. J., & Mitroff, S. R. (2005). The persistence of objects files representations. *Perception and Psychophysics*, 67, 324-334.
- Nosofsky, R. M. (1986). Attention, similarity, and the identification-categorization relationship. *Journal of Experimental Psychology: General*, 115(1), 39-57.
- Nosofsky, R. M., & Palmeri, T. J. (1997). An exemplar-based random walk model of speeded classification. *Psychological Review*, 104(2), 266-300.
- Oliva, A., & Schyns, P. G. (2000). Diagnostic colors mediate scene recognition. *Cognitive Psychology*, 41(2), 176-210.
- Oliva, A., & Torralba, A. (2001). Modeling the shape of the scene: A holistic representation of the spatial envelope. *International Journal of Computer Vision*, 42, 145-175.
- Olson, I. R., & Chun, M. M. (2001). Temporal contextual cuing of visual attention. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 27(5), 1299-1313.
- Olson, I. R., & Chun, M. M. (2002). Perceptual constraints on implicit learning of spatial context. *Visual Cognition*, 9(3), 273-302.
- Ono, F., Jiang, Y., & Kawahara, J.-i. (2005). Intertrial Temporal Contextual Cuing: Association Across Successive Visual Search Trials Guides Spatial Attention. *Journal of Experimental Psychology: Human Perception and Performance*, 31(4), 703-712.
- O'Regan, J. K., & Noë, A. (2001). A sensorimotor account of vision and visual consciousness. *Behavioral and Brain Sciences*, 24(5), 939-1031.
- O'Reilly, R. C., & Munakata, Y. (2000). *Computational explorations in cognitive neuroscience: Understanding the mind by simulating the brain*. Cambridge, MA, US: The MIT Press.
- Overgaard, M., Rote, J., Mouridsen, K., & Ramsøy, T. Z. g. (2006). Is conscious perception gradual or dichotomous? A comparison of report methodologies during a visual task. *Consciousness and Cognition: An International Journal*, 15(4), 700-708.
- Palmer, S. E. (1975). The effects of contextual scenes on the identification of objects. *Memory & Cognition*, 3(5), 519-526.
- Palmeri, T. J. (1997). Exemplar similarity and the development of automaticity. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 23(2), 324-354.
- Park, H., Quinlan, H., Thornton, E., & Reder, L. M. (2004). The effect of midazolam on visual search: Implications for understanding amnesia. *Proc. Natl. Acad. Sci. U. S. A.* 101, 17879-17883.
- Pashler, H. E. (1998). *The psychology of attention*. Cambridge, MA, US: The MIT Press.

- Perruchet, P. (1994). Defining the knowledge units of a synthetic language: Comment on Vokey and Brooks (1992). *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 20(1), 223-228.
- Perruchet, P., & Amorim, M.-A. (1992). Conscious knowledge and changes in performance in sequence learning: Evidence against dissociation. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 18(4), 785-800.
- Perruchet, P., & Nicolas, S. (1998). L'apprentissage implicite: Un débat théorique. *Psychologie Française*, 43(1), 13-25.
- Perruchet, P., & Pacteau, C. (1990). Synthetic grammar learning: Implicit rule abstraction or explicit fragmentary knowledge? *Journal of Experimental Psychology: General*, 119(3), 264-275.
- Perruchet, P., & Pacton, S. (2006). Implicit learning and statistical learning: One phenomenon, two approaches. *Trends in Cognitive Sciences*, 10(5), 233-238.
- Perruchet, P., & Vinter, A. (2002). The self-organizing consciousness. *Behavioral and Brain Sciences*, 25(3), 297-388.
- Perruchet, P., Vinter, A., & Gallego, J. (1997). Implicit learning shapes new conscious percepts and representations. *Psychonomic Bulletin & Review*, 4(1), 43-48.
- Peterson, M. S., & Kramer, A. F. (2001). Contextual cueing reduces interference from task-irrelevant onset distractors. *Visual Cognition*, 8(6), 843-859.
- Posner, M. I. (1980). Orienting of attention. *The Quarterly Journal of Experimental Psychology*, 32(1), 3-25.
- Posner, M. I., & Cohen, Y. (1984). Components of visual orienting. In H. Bouma & D. G. Bouwhuis (Eds), *Attention and performance X* (pp. 55-66). Hove, UK: Lawrence Erlbaum Associates Ltd.
- Posner, M. I., & Dehaene, S. (1994). Attentional networks. *Trends in Neurosciences*, 17(2), 75-79.
- Posner, M. I., & Petersen, S. E. (1990). The attention system of the human brain. *Annual Review of Neuroscience*, 13, 25-42.
- Pothos, E. M. (2007). Theories of Artificial Grammar Learning. *Psychological Bulletin*, 133(2), 227-244.
- Potter, M. C. (1976). Short-term conceptual memory for pictures. *Journal of Experimental Psychology: Human Learning & Memory*, 2(5), 509-522.
- Potter, M. C., & Levy, E. I. (1969). Recognition memory for a rapid sequence of pictures. *Journal of Experimental Psychology*, 81(1), 10-15.
- Potter, M. C., Staub, A., & O'Connor, D. H. (2004). Pictorial and Conceptual Representation of Glimpsed Pictures. *Journal of Experimental Psychology: Human Perception and Performance*, 30(3), 478-489.

- Potter, M. C., Staub, A., Rado, J., & O'Connor, D. H. (2002). Recognition memory for briefly presented pictures: The time course of rapid forgetting. *Journal of Experimental Psychology: Human Perception and Performance*, 28(5), 1163-1175.
- Preston, A.R., Salidis, J., & Gabrieli, J.D.E. (2001). *Medial temporal lobe activation during implicit contextual learning*. Abstract Viewer and Itinerary Planner, 347.6. Washington, DC: Society for Neuroscience.
- Rao, R. P. N., Zelinsky, G. J., Hayhoe, M. M., & Ballard, D. H. (2002). Eye movements in iconic visual search. *Vision Research*, 42(11), 1447-1463.
- Reber, A. S. (1967). Implicit learning of artificial grammars. *Journal of Verbal Learning & Verbal Behavior*, 6(6), 855-863.
- Reber, A. S. (1969). Transfer of syntactic structure in synthetic languages. *Journal of Experimental Psychology*, 81(1), 115-119.
- Reber, A. S. (1989). Implicit learning and tacit knowledge. *Journal of Experimental Psychology: General*, 118(3), 219-235.
- Reber, A. S. (1993). *Implicit learning and tacit knowledge: An essay on the cognitive unconscious*. New York, NY, US: Oxford University Press.
- Reber, A. S., & Millward, R. B. (1968). Event observation in probability learning. *Journal of Experimental Psychology*, 77(2), 317-327.
- Reingold, E. M., Charness, N., Pomplun, M., & Stampe, D. M. (2001). Visual Span in Expert Chess Players: Evidence From Eye Movements. *Psychological Science*, 12(1).
- Rensink, R. A. (2000a). The dynamic representation of scenes. *Visual Cognition*, 7(1-3), 17-42.
- Rensink, R. A. (2000b). Seeing, sensing, and scrutinizing. *Vision Research*, 40(10-12), 1469-1487.
- Rensink, R. A. (2002). Change detection. *Annual Review of Psychology*, 53(1), 245-376.
- Rensink, R. A. (2004). Visual Sensing Without Seeing. *Psychological Science*, 15(1), 27-32.
- Rensink, R. A., O'Regan, J. K., & Clark, J. J. (1997). To see or not to see: The need for attention to perceive changes in scenes. *Psychological Science*, 8(5), 368-373.
- Rensink, R. A., O'Regan, J. K., & Clark, J. J. (2000). On the failure to detect changes in scenes across brief interruptions. *Visual Cognition*, 7(1), 127-145.
- Rousselet, G. A., Joubert, O. R., & Fabre-Thorpe, M. I. (2005). How long to get to the 'gist' of real-world natural scenes? *Visual Cognition*, 12(6), 852-877.
- Rowland, L. A., & Shanks, D. R. (2006). Attention modulates the learning of multiple contingencies. *Psychonomic Bulletin & Review*, 13(4), 643-648.

- Saffran, J. R., Aslin, R. N., & Newport, E. L. (1996). Statistical learning by 8-month old infants. *Science*, 274, 1926-1928.
- Saffran, J. R., Johnson, E. K., Aslin, R. N., & Newport, E. L. (1999). Statistical learning of tone sequences by human infants and adults. *Cognition*, 70, 27-52.
- Schacter, D. L. (1987). Implicit memory: History and current status. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 13(3), 501-518.
- Schneider, W., & Shiffrin, R. M. (1977). Controlled and automatic human information processing: I. Detection, search, and attention. *Psychological Review*, 84(1), 1-66.
- Schyns, P. G., & Oliva, A. (1994). From blobs to boundary edges: Evidence for time- and spatial-scale-dependent scene recognition. *Psychological Science*, 5(4), 195-200.
- Seitz, A. R., & Watanabe, T. (2003). Is subliminal learning really passive? *Nature*, 422(6927), 36-36.
- Seitz, A., & Watanabe, T. (2005). A unified model for perceptual learning. *Trends in Cognitive Sciences*, 9(7), 329-334.
- Shanks, D. R. (2005). Implicit learning. In K. Lamberts and R. Goldstone, *Handbook of Cognition* (pp. 202-220) . London: Sage.
- Shanks, D. R., & Channon, S. (2002). Effects of a secondary task on 'implicit' sequence learning: Learning or performance? *Psychological Research/Psychologische Forschung*, 66(2), 99-109.
- Shanks, D. R., & Johnstone, T. (1999). Evaluating the relationship between explicit and implicit knowledge in a sequential reaction time task. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 25(6), 1435-1451.
- Shanks, D. R., & St. John, M. F. (1994). Characteristics of dissociable human learning systems. *Behavioral and Brain Sciences*, 17(3), 367-447.
- Shanks, D. R., Rowland, L. A., & Ranger, M. S. (2005). Attentional load and implicit sequence learning. *Psychological Research/Psychologische Forschung*, 69(5), 369-382.
- Shepard, R. N. (1967). Recognition memory for words, sentences, and pictures. *Journal of Verbal Learning & Verbal Behavior*, 6(1), 156-163.
- Shiffrin, R. M., & Schneider, W. (1977). Controlled and automatic human information processing: II. Perceptual learning, automatic attending and a general theory. *Psychological Review*, 84(2), 127-190.
- Shinoda, H., Hayhoe, M. M., & Shrivastava, A. (2001). What controls attention in natural environments? *Vision Research*, 41(25), 3535-3545.
- Simons, D. J., & Ambinder, M. S. (2005). Change Blindness: Theory and Consequences. *Current Directions in Psychological Science*, 14(1), 44-48.

- Simons, D. J., & Chabris, C. F. (1999). Gorillas in our midst: Sustained inattention blindness for dynamic events. *Perception*, 28(9), 1059-1074.
- Simons, D. J., & Rensink, R. A. (2005). Change blindness: Past, present, and future. *Trends in Cognitive Sciences*, 9(1), 16-20.
- Simons, D. J., Franconeri, S. L., & Reimer, R. L. (2000). Change blindness in the absence of a visual disruption. *Perception*, 29(10), 1143-1154.
- Song, J.-H., & Jiang, Y. (2005). Connecting the past with the present: How do humans match an incoming visual display with visual memory? *Journal of Vision*, 5(4), 322-330.
- Smyth, A. & Shanks, D. R. (in press). Awareness in contextual cuing with extended and concurrent explicit tests. *Memory & Cognition*.
- Stadler, M. A. (1995). Role of attention in implicit learning. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21(3), 674-685.
- Standing, L., Conezio, J., & Haber, R. N. (1970). Perception and memory for pictures: Single-trial learning of 2500 visual stimuli. *Psychonomic Science*, 19(2), 73-74.
- Summala, H. & Räsänen, M. (2000). Top-down and bottom-up process in driver behaviour at roundabouts and crossroads. *Transportation Human Factor*, 2, 29-387.
- Tatler, B. W., Baddeley, R. J., & Gilchrist, I. D. (2005). Visual correlates of fixation selection: Effects of scale and time. *Vision Research*, 45(5), 643-659.
- Tatler, B. W., Gilchrist, I. D., & Land, M. F. (2005). Visual memory for objects in natural scenes: From fixations to object files. *The Quarterly Journal of Experimental Psychology A: Human Experimental Psychology*, 58(5), 931-960.
- Theeuwes, J. (1992). Perceptual selectivity for color and form. *Perception & Psychophysics*, 51(6), 599-606.
- Theeuwes, J. (2004). Top-down search strategies cannot override attentional capture. *Psychonomic Bulletin & Review*, 11(1), 65-70.
- Theeuwes, J., & Hagenzieker, M. P. (1993). *Visual search of traffic scenes: On the effect of location expectations*. In A. G. Gale et al. (Eds), *Vision in Vehicles 4*. Amsterdam: Elsevier, 149-158.
- Thoma, V., Hummel, J. E., & Davidoff, J. (2004). Evidence for Holistic Representations of Ignored Images and Analytic Representations of Attended Images. *Journal of Experimental Psychology: Human Perception and Performance*, 30(2), 257-267.
- Thornton, I. M., & Fernandez-Duque, D. (2000). An implicit measure of undetected change. *Spatial Vision*, 14(1), 21-44.
- Thorpe, S. J., Gegenfurtner, K. R., Fabre-Thorpe, M. I., & Bülthoff, H. H. (2001). Detection of animals in natural images using far peripheral vision. *European Journal of Neuroscience*, 14(5), 869-876.

- Torralba, A., Oliva, A., Castelhamo, M. S., & Henderson, J. M. (2006). Contextual Guidance of Eye Movements and Attention in Real-World Scenes: The Role of Global Features in Object Search. *Psychological Review*, 113(4), 766-786.
- Treisman, A. M., & Gelade, G. (1980). A feature-integration theory of attention. *Cognitive Psychology*, 12(1), 97-136.
- Turano, K. A., Geruschat, D. R., & Baker, F. H. (2003). Oculomotor strategies for the direction of gaze tested with a real-world activity. *Vision Research*, 43(3), 333-346.
- Turk-Browne, N. B., Jungé, J., & Scholl, B. J. (2005). The Automaticity of Visual Statistical Learning. *Journal of Experimental Psychology: General*, 134(4), 552-564.
- Tzelgov, J. (1997). Specifying the relations between automaticity and consciousness: A theoretical note. *Consciousness and cognition*, 6, 441-451.
- Underwood, G., & Foulsham, T. (2006). Visual saliency and semantic incongruity influence eye movements when inspecting pictures. *The Quarterly Journal of Experimental Psychology*, 59(11), 1931-1949.
- Underwood, G., Foulsham, T., van Loon, E., Humphreys, L., & Bloyce, J. (2006). Eye movements during scene inspection: A test of the saliency map hypothesis. *European Journal of Cognitive Psychology*, 18(3), 321-342.
- Van Elslande, P., Alberton, L., Nachtergaële, C., Blanchet G. (1997). *Scénarios types de production de "l'erreur humaine" dans l'accident de la route*. Rapport de recherche INRTES n°218 (180p).
- VanRullen, R., & Koch, C. (2003). Competition and selection during visual processing of natural scenes and objects. *Journal of Vision*, 3(1), 75-85.
- Vokey, J. R., & Brooks, L. R. (1992). Saliency of item knowledge in learning artificial grammars. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 18(2), 328-344.
- Watanabe, T., Nanez, J. E., & Sasaki, Y. (2001). Perceptual learning without perception. *Nature*, 413, 844-848.
- Weiskrantz, L., & Warrington, E. K. (1979). Conditioning in amnesic patients. *Neuropsychologia*, 17(2), 187-194.
- Whittlesea, B. W., & Dorken, M. D. (1993). Incidentally, things in general are particularly determined: An episodic-processing account of implicit learning. *Journal of Experimental Psychology: General*, 122(2), 227-248.
- Willingham, D. B., Nissen, M. J., & Bullemer, P. (1989). On the development of procedural knowledge. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 15(6), 1047-1060.
- Wolfe, J. M., Cave, K. R., & Franzel, S. L. (1989). Guided search: An alternative to the feature integration model for visual search. *Journal of Experimental Psychology: Human Perception and Performance*, 15(3), 419-433.

- Wolfe, J. M., Horowitz, T. S., & Michod, K. O. (2007). Is visual attention required for robust picture memory? *Vision Research*, 47(7), 955-964.
- Wright, R. L., & Whittlesea, B. W. A. (1998). Implicit learning of complex structures: Active adaptation and selective processing in acquisition and application. *Memory & Cognition*, 26(2), 402-420.
- Yantis, S., & Jonides, J. (1990). Abrupt visual onsets and selective attention: Voluntary versus automatic allocation. *Journal of Experimental Psychology: Human Perception and Performance*, 16(1), 121-134.
- Yarbus, A. L. (1967). *Eye movements and vision*. New York: Plenum.
- Yi, D.-J., & Chun, M. M. (2005). Attentional Modulation of Learning-Related Repetition Attenuation Effects in Human Parahippocampal Cortex. *Journal of Neuroscience*, 25(14), 3593-3600.
- Ziori, E., & Dienes, Z. (2006). *Subjective measures of unconscious knowledge of concepts*. *Mind and Society*, 5(1), 105-122.

ANNEXES

ANNEXE 1 : Description de la théorie de la cohérence (Rensink, 2000a), de la théorie du fichier objet (Irwin & Andrews, 1996) et de la théorie de la perception de scène et de la mémoire (Hollingworth & Henderson, 2002)

ANNEXE 2 : L'apprentissage de régularités statistiques

ANNEXE 3 : Etude 1. Localisation potentielle de la cible

ANNEXE 4 : Etude 2. Expériences préliminaires

ANNEXE 5 : Etude 2. Mots utilisés

ANNEXE 6 : Etude 2. Résultats relatifs à chaque catégorie prédictive

ANNEXE 7 : Etude 3. Matériel utilisé

ANNEXE 8 : Etude 3. Corrélations entre les performances durant la tâche de prise de décision et les verbalisations

ANNEXE 9 : L'apprentissage de signaux en mouvement cohérent

ANNEXE 1

Description de la théorie de la cohérence (Rensink, 2000a), de la théorie du fichier objet (Irwin & Andrews, 1996) et de la théorie de la perception de scène et de la mémoire (Hollingworth & Henderson, 2002)

Annexe 1a : Théorie de la cohérence et Théorie du fichier objet

La théorie du fichier objet et la théorie de la cohérence considèrent deux niveaux de représentation. Le premier niveau, plus largement détaillé dans la théorie de la cohérence, correspondrait à une représentation relativement détaillée de la scène mais peu cohérente et extrêmement instable dans l'espace et dans le temps. Cette représentation serait sous-tendue par un système visuel de bas niveau opérant à un niveau préattentif, dès les premières étapes du traitement visuel. Les traits et les dimensions basiques (couleur, orientation, fréquences spatiales...) seraient précocement et massivement traités en parallèle. Selon la théorie de la cohérence, les objets présents dans le champ visuel, même extérieurs au faisceau attentionnel, subiraient ces traitements préattentifs leur conférant un statut de proto-objets. Ces proto-objets seraient rétinotopiques, formés rapidement, et de manière continue par l'assemblage en parallèle de fragments relativement complexes conformément à leur localisation sur la rétine, sans requérir d'attention. La formation de ces structures visibles (i.e. accessibles à la phénoménologie) impliquerait le traitement des blobs et des bars, le traitement primaire de certaines propriétés de la scène telles que les courbes de surface, l'inclinaison, les ombres, mais également un traitement de groupement rapide et « peu cohérent » des traits. Ces proto-objets pourraient être relativement complexes, mais présenteraient une cohérence spatio-temporelle extrêmement limitée. En effet, étant continuellement remplacées par une nouvelle image se superposant à l'ancienne sur la rétine (i.e. un nouvel encodage visuel), ces structures de bas niveau seraient de manière inhérente volatiles. Ainsi, un croquis relativement détaillé de scène serait précocement construit au cours du traitement, mais ce dernier serait peu cohérent. Donc contrairement à notre phénoménologie qui nous donne l'illusion de percevoir la scène de manière cohérente dans son ensemble, la théorie de la cohérence stipule que notre représentation globale est extrêmement schématique.

Le second niveau de représentation fait appel à un système attentionnel qui permettrait de transformer ce que Rensink appelle proto-objet en objet unifié. Dans la théorie du fichier objet, Irwin (Irwin, 1992 ; Andrews & Irwin, 1996) s'appuie explicitement sur le concept de « fichier objet » développé par Treisman et collaborateurs (Kahneman, Treisman, & Gibbs, 1992) dans le cadre de la *théorie d'intégration des traits* (Treisman & Gelade, 1980). La théorie d'intégration des traits postule que l'attention focalisée est nécessaire pour que les traits visuels soient organisés et intégrés en un objet unifié. Lorsque l'attention est portée sur un objet, elle agirait comme un ciment intégrant toutes les dimensions basiques présentes dans le faisceau attentionnel en un objet unique. Une fois les dimensions intégrées, une représentation temporaire serait formée : un « fichier objet ». Ce « fichier objet » coderait les descriptions visuelles de l'objet ainsi que sa position spatiale au sein d'une carte mère directrice des localisations (Kahneman & Treisman, 1984). Selon Irwin, contrairement au concept de fichier objet défini dans la théorie d'intégration des traits, il ne s'agit pas d'une

représentation sensorielle, mais de codes visuels abstraits à partir d'une représentation sensorielle. Compte tenu des capacités de stockage extrêmement limitées de la MVCT, un nombre très limité de fichiers objets, 3 ou 4 fichiers objet, seraient maintenus en MVCT au cours d'une saccade oculaire (pour des arguments en faveur d'une MCT comportant 3-4objets (cf. Irwin, 1992 ; Kahneman, Treisman, & Gibbs, 1992). Ces fichiers objets seraient maintenus pendant un délai relativement long, conformément aux capacités de rétention de la MVCT, i.e. plusieurs secondes (Noles, Scholl, & Mitroff, 2005). Ces fichiers objets constitueraient le contenu primaire de la mémoire au cours des saccades, fournissant une impression de continuité entre les traits entre deux fixations.

De manière similaire, la théorie de la cohérence stipule que l'attention joue un rôle essentiel en fournissant à la scène et aux objets qui la composent une cohérence spatiale et temporelle. L'attention agirait de manière métaphorique comme « une main à 4, 5 ou 6 doigts qui cueillerait un petit nombre de proto-objets à partir du flux d'information en constante régénération » (Rensink, 2000a, p. 24). Comme dans la théorie du fichier objet, l'attention permet de construire une description cohérente de cet objet attendu en intégrant plusieurs de ses propriétés comme sa taille, sa forme, sa couleur. Mais de manière différente, la théorie de la cohérence postule que les structures sélectionnées par l'attention ne peuvent former qu'un seul et unique objet susceptible d'atteindre un état stable à un instant donné. De plus, l'attention focalisée serait ici nécessaire pour que cette représentation cohérente soit maintenue en MVCT. Sitôt que l'attention se détournerait de l'objet, ce dernier perdrait sa cohérence et retournerait à l'état de proto-objet. La représentation de l'objet se dégraderait, étant immédiatement remplacée par le nouvel encodage visuel. Le traitement attentionnel effectué sur un objet n'aurait pas ou peu de conséquence sur les représentations ultérieures. Néanmoins, Rensink (2000a) stipule que deux structures adjacentes sont spatialement et temporellement cohérentes si elles se rapportent à un même objet.

Ainsi, la théorie du fichier objet et la théorie de la cohérence diffèrent sur trois points. Rensink propose que seul un objet peut être maintenu en MVCT, alors que Irwin (1992) apporte des arguments en faveur de 3 ou 4 objets. La théorie de la cohérence stipule que les représentations sensorielles peuvent être retenues en MVCT, alors que la théorie du fichier objet stipule que la MVCT assure le maintien de représentations visuelles abstraites à partir des informations sensorielles (i.e. indépendamment de leur position dans l'espace). Enfin, selon la théorie de la cohérence, dès lors que l'attention est détournée de l'objet, les représentations visuelles seraient désintégrées, alors que les fichiers objets resteraient actifs en MVCT même lorsque l'attention est détournée de l'objet.

Annexe 1b : Théorie de la perception de scène et de la mémoire

Selon la théorie de la perception de scène et de la mémoire proposée par Hollingworth et Henderson (2002), lorsque l'attention serait tour à tour portée sur des objets qui composent la scène, des représentations visuelles et conceptuelles de haut niveau, abstraites à partir des propriétés sensorielles de bas niveau, seraient construites. Ces représentations coderaient une description relativement détaillée sur les caractéristiques perceptives et conceptuelles des objets et seraient dépendantes du point de vue selon lequel chacun de ces objets a été observé. Ces représentations abstraites seraient indexées sous forme de « fichiers objets » au sein d'une carte codant l'agencement spatial de la scène. Comme dans la théorie du fichier objet de la mémoire transsaccadique (Irwin, 1991), la conception de « fichier objet » diffère de celle

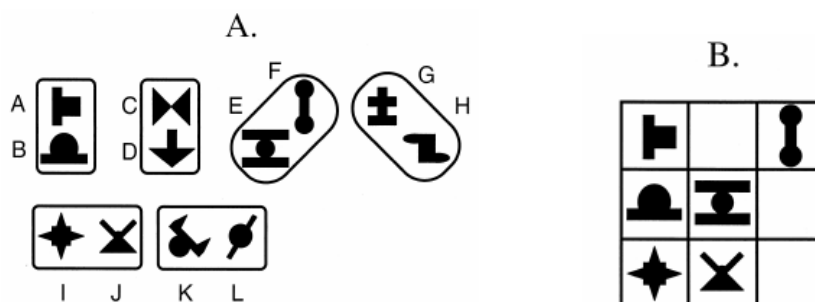
proposée initialement par Kahneman et collaborateurs (1992) dans le sens où serait préservées des représentations visuelles abstraites plutôt que des informations sensorielles, et donc ces codes abstraits supporteraient la rétention en MCT. Les fichiers objets seraient ainsi encodés non seulement en MVCT, mais également en MCT. Sur la base de ces représentations actives en mémoire à CT, un traitement approfondi des fichiers objets au cours des multiples fixations oculaires sur la scène permettrait une consolidation progressive en MLT au sein d'une carte décrivant l'ensemble de la scène.

Quand l'attention se détournerait de l'objet, les représentations sensorielles se dégraderaient immédiatement, mais des fichiers objets très détaillés seraient néanmoins retenus quelques instants en MVCT et beaucoup plus longtemps en MVLT. Ainsi, cette représentation visuelle relativement détaillée résisterait aux saccades oculaires. Au cours des multiples fixations sur une scène, l'accumulation des informations visuelles et conceptuelles relatives aux objets attendus conduirait à une représentation détaillée de l'ensemble de la scène. Ici encore, l'attention joue un rôle fondamental à la fois dans l'encodage et dans la récupération des informations visuelles et conceptuelles : l'accès au contenu d'un fichier objet en MVCT est supposé dépendre d'une attention spatiale allouée à l'objet indexé sous forme de fichier.

ANNEXE 2

L'apprentissage de régularités statistiques

Le paradigme d'apprentissage de régularités statistiques (SL : Statistical Learning) permet d'étudier la capacité d'un individu à extraire des co-occurrences à partir d'un matériel simple et fortement structuré sans instruction ni feedback particuliers. A l'origine, ce paradigme développé par Saffran, Aslin et Newport (1996) visait à mieux comprendre comment les enfants apprennent à segmenter le flux continu de la parole en mots en étudiant leurs capacités à découvrir des mots cachés dans un langage artificiel continu. Les enfants étaient exposés de manière passive à une séquence de syllabes comportant des relations de co-occurrences (e.g. « pi-go-la-bi-ku-ti-... »). Les douze syllabes probables utilisées dans l'expérience princeps étaient en fait structurées en triplets de syllabes, de telle manière que les trois syllabes apparaissaient toujours dans le même ordre. Au sein de la séquence, les triplets étaient présentés de manière aléatoire (e.g. ABC, GHI, DEF, ABC, JKL.....). Deux minutes d'exposition à cette séquence structurée suffisaient pour que des enfants de 8 mois discriminent les triplets de haute fréquence (e.g. GHI) par rapport à des triplets jamais entendus (e.g. AEI) ou de basse fréquence (e.g. BCG). Par la suite, le phénomène d'apprentissage de régularités statistiques a été étendu chez l'adulte avec des stimuli auditifs non linguistiques (e.g. Saffran, Johnson, Aslin, & Newport, 1999) et à d'autres modalités sensorielles, i.e. visuelle (Fiser & Aslin, 2001, 2002) et tactile (Conway & Christiansen, 2005). Par exemple, l'apprentissage implicite de régularités visuelles a été mis en évidence lorsque des objets co-variaient dans l'espace (Fiser & Aslin, 2001) et dans le temps (Fiser & Aslin, 2002).



Matériel utilisé dans l'étude de Fiser et Aslin (2001). Chaque exemplaire comportait 6 formes visuelles sélectionnées parmi un ensemble de 12 formes probables et organisées sur une grille visible 3x3 (Figure B). Ces formes sont suffisamment complexes pour être facilement discriminées et non familières. Le paradigme VSL (visual statistical learning) comporte deux phases expérimentales, une phase de familiarisation et une phase test. Dans cette étude, durant la phase de familiarisation, les participants étaient exposés à 144 exemplaires dont la durée d'apparition était de 1s et recevaient pour instruction de prêter attention à cette séquence continue. Mais à l'insu des sujets, les 12 formes utilisées étaient en fait organisées en 6 paires de base, consistant en 2 formes spécifiques agencées selon une relation spatiale particulière (Figure A). Durant la phase test, les sujets étaient exposés à des paires de base dont la moitié respectait les règles de co-occurrences et l'autre moitié les violait. Ils réalisaient une tâche de choix forcé à deux alternatives temporelles dans laquelle ils devaient juger du caractère familier vs. non familier de chacune de ces paires de base. Les résultats montrent que les sujets jugeaient plus familières les paires de base qui respectaient les règles de co-occurrences que celles qui les violaient.

ANNEXE 3

Etude 1 : Localisations potentielles de la cible

				C			C
C		C					
				C		C	
C			C				
					C		C
C			C				

Annexe 3a: Localisations potentielles de la cible dans l'Expérience 1a

				C			C
	C						
				C			
			C				
						C	
C			C				

Annexe 3b : Localisations potentielles de la cible dans l'Expérience 1b

ANNEXE 4

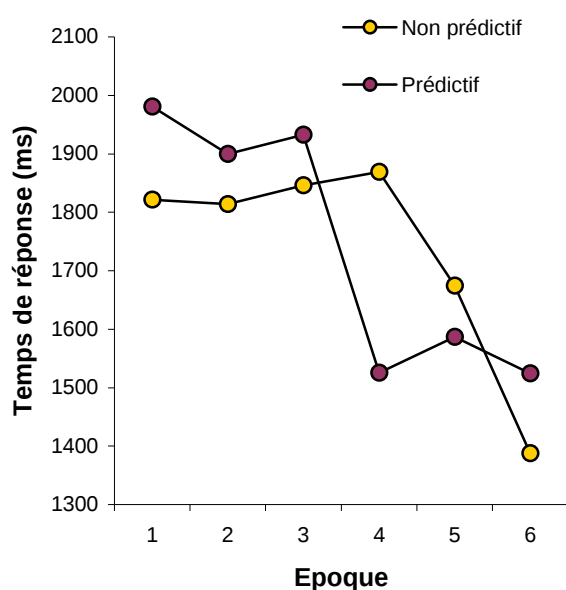
Etude 2. Expériences préliminaires

Annexe 4a : Expérience préliminaire 1

Méthode

Le principe de cette expérience était identique à celui des Expériences 2a et 2b réalisées avec des nombres, mais le paradigme d'indication contextuel a ici été utilisé avec des affichages composés de mots. Dans la tâche de recherche visuelle, trois participants avaient pour consigne de rechercher le nom « MULOT » ou « MULET » parmi un ensemble de mots contextuels. Le mot cible pouvait apparaître dans deux zones restreintes de l'affichage (i.e. à droite vs. à gauche de l'affichage, cf. Figure 21). A chaque essai, l'item cible était présenté parmi un ensemble de mots contextuels, appartenant, soit à la catégorie sémantique « poissons » soit à la catégorie sémantique « oiseaux ». La tâche de recherche visuelle comportait une série de blocs d'essais, chaque bloc comportant des essais prédictifs et des essais non prédictifs. Dans un essai prédictif, l'appartenance catégorielle et sémantique des mots contextuels était prédictive de la zone où était localisée la cible. Par exemple, quand le contexte était composé exclusivement de noms d'oiseaux, la cible était toujours localisée à gauche de l'affichage, quand le contexte était composé exclusivement de noms de poissons, la cible était toujours localisée à droite de l'affichage. Dans un essai non prédictif, la cible apparaissait parmi autant de noms de poissons que de noms d'oiseaux et la cible pouvait apparaître à gauche ou à droite de l'affichage.

Résultats et discussion



Les premiers résultats de cette expérience ne présentant aucun effet d'indication contextuel, le matériel a été reconsidéré.

Les cibles utilisées au cours de cette expérience (mulot vs. mulet) étaient très ressemblantes visuellement. Aussi est-il probable que les TR étaient trop courts pour permettre la mise en évidence d'un effet d'indication contextuel. De ce fait, après avoir fait passer trois participants seulement, d'autres cibles ont été utilisées dans l'Expérience préliminaire 2, à savoir « ourson » et « lièvre ».

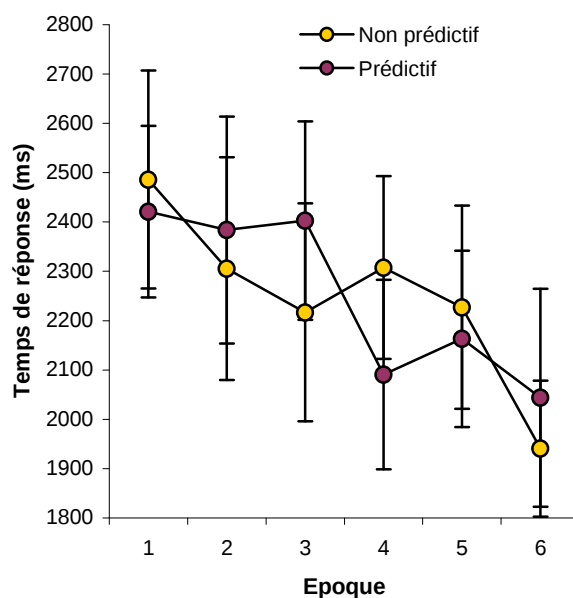
Temps de réponse dans les conditions prédictive et non prédictive observés dans l'Expérience préliminaire 1 (N=3)

Annexe 4b : Expérience préliminaire 2

Méthode

La tâche et le matériel étaient les mêmes que dans l'Expérience préliminaire 1, mais de manière différente, neuf participants avaient pour consigne de rechercher le nom « OURSON » ou « LIEVRE » parmi un ensemble de mots contextuels.

Résultats et discussion



Une ANOVA à mesures répétées a été réalisée, avec les facteurs intra-sujets, condition (prédictive vs non prédictive) et époque (1-6). Cette analyse montre seulement un effet du facteur époque, $F(5,40)=10.104$, $p<.001$, soit, pas d'effet du facteur condition, $F(1,8)<1$ et pas d'interaction entre les facteurs « condition x époque », $F(5,40)=1.895$, $p=.168$.

Temps de réponse dans les conditions prédictive et non prédictive observés dans l'Expérience préliminaire 2. Les barres d'erreurs correspondent aux SEM (N=9)

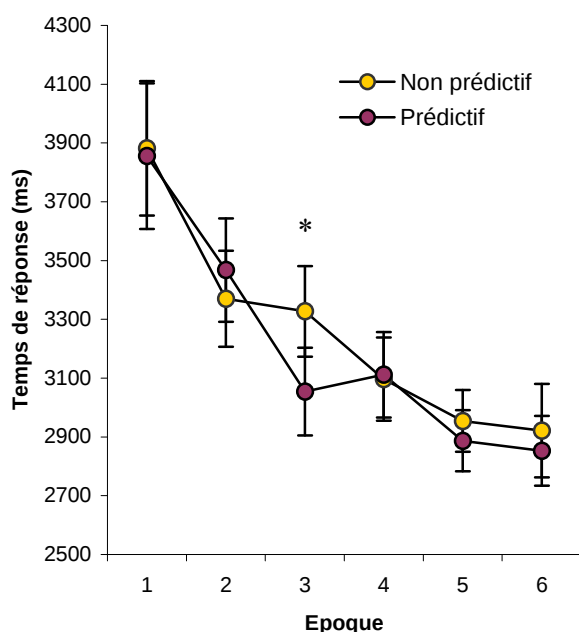
Cette expérience ne met pas en évidence d'effet facilitateur de la condition prédictive sur la condition non prédictive. Il est probable que la recherche de mots cibles spécifiques conduise à une recherche visuelle « superficielle », principalement basée sur des caractéristiques perceptives des items. Aussi, dans l'Expérience préliminaire 3, les cibles n'étaient plus spécifiques mais elles étaient définies aux participants par deux catégories sémantiques potentielles. Les participants avaient pour consigne de rechercher le nom d'un vêtement ou d'un bâtiment.

Annexe 4c : Expérience préliminaire 3

Méthode

Le principe expérimental était le même que celui utilisé dans l'Expérience préliminaire 1. Mais afin de contraindre les participants à réaliser un traitement sémantique sur les items de l'affichage, ils étaient invités à rechercher le nom d'un vêtement ou d'un bâtiment le plus rapidement et correctement possible ($N=16$). Donc, dans cette expérience, la cible n'était pas spécifique mais elle était définie aux participants par deux catégories potentielles. Comme dans l'Etude 1, les contextes prédictifs étaient composés exclusivement de noms de poissons ou exclusivement de noms d'oiseaux. Dans les essais non prédictifs, le contexte comportait autant de noms d'oiseaux que de noms de poissons et la cible pouvait apparaître à gauche ou à droite de l'affichage.

Résultats et discussion



Une ANOVA à mesures répétées a été réalisée, avec les facteurs intra-sujets, condition (prédictive vs non prédictive) et époque (1-6). Cette analyse montre seulement un effet du facteur époque, $F(5,75)=22.206$, $p<.001$, soit, pas d'effet du facteur condition, $F(1,15)=1.161$, $p=.298$, et pas d'interaction entre les facteurs « condition x époque », $F(5,75)=1.297$, $p=.274$. Des analyses partielles révèlent toutefois un effet significatif du facteur condition à l'époque 3, $F(1,15)=8.339$, $p<.02$.

Temps de réponse dans les conditions prédictive et non prédictive observés dans l'Expérience préliminaire 3. Les barres d'erreurs correspondent aux SEM ($N=16$).

Les résultats mettent en évidence un bénéfice dans la condition prédictive à l'époque 3 uniquement. Cependant, l'ANOVA globale ne permet pas de conclure à un effet d'apprentissage plus prononcé dans la condition prédictive que dans la condition non prédictive. Il est possible qu'en constituant pour le sujet des contre-exemples, les contextes non prédictifs mixtes (i.e. comportant des noms d'oiseaux et de poissons) rendent difficile l'apprentissage des régularités contextuelles. Pour simplifier l'apprentissage des régularités, des catégories prédictives et des catégories non prédictives distinctes ont été utilisées dans les Expériences présentées dans le corps de texte.

ANNEXE 5

Etude 2 : Mots utilisés

Annexe 5a : Cibles utilisées dans l'Etude 2, Expériences 4, 5, 6 et 7

Cibles utilisées dans les Expériences 4, 5, 6 et 7	
VETMENTS	BATIMENTS
GILET	USINE
BOTTE	FERME
TRICOT	MAIRIE
VESTON	PALAIS
ANORAK	MOULIN
BLOUSE	CHALET
BONNET	GARAGE
MOUFLE	MANOIR
CORSET	ÉGLISE
PYJAMA	CABANE
MAILLOT	GYMNASE
CHEMISE	COLLÈGE
CULOTTE	AUBERGE
FOULARD	TAVERNE
COSTUME	MAGASIN
BERMUDA	CHÂTEAU

**Annexe 5b : Listes des mots contextuels utilisés dans l'Etude 2, Expériences 4, 5, 6 et 7
(les fréquences d'usage sont issues de la base de données BRULEX)**

Premières listes de mots contextuels utilisées dans les Expériences 4, 5, 6, et 7											
FRUITS/LEGUMES Mean Freqfilms: 2,00 Mean Freqbook: 2,15			ARBRES/FLEURS Mean Freqfilms: 1,19 Mean Freqbook: 3,88			MAMMIFERES Mean Freqfilms:2,64 Mean Freqbook: 2,36			OISEAUX Mean Freqfilms:3,19 Mean Freqbook: 3,80		
	Ffilms	Fbook		Ffilms	Fbook		Ffilms	Fbook		Ffilms	Fbook
MELON	4,22	5,27	CHÊNE	2,89	16,49	BISON	2,83	1,28	AIGLE	4,28	7,91
RADIS	1,45	3,11	FRÊNE	0	1,62	ZÈBRE	1,33	3,04	HÉRON	0,18	1,08
COING	0,06	0,41	SAPIN	4,1	9,86	RENNE	0,84	0,47	CYGNE	2,47	4,66
RAISIN	4,76	4,86	ACACIA	0,06	3,24	AGNEAU	8,73	5,95	CANARD	17,11	16,15
ANANAS	3,37	3,51	CACTUS	2,65	2,3	CHÈVRE	6,45	10,14	PIVERT	0,48	0,81
MANGUE	0,3	0,74	ÉRABLE	1,99	1,15	BUFFLE	3,13	1,96	FAISAN	1,14	1,69
PAPAYE	0,42	0,14	JASMIN	1,2	4,19	JAGUAR	1,02	0,41	FAUCON	9,04	2,36
FRAISE	5,3	3,99	MUGUET	0,12	3,85	COBAYE	2,89	0,74	DINDON	1,27	0,74
POIREAU	0,6	0,88	FOUGÈRE	0,66	0,74	BELETTE	1,33	0,68	CIGOGNE	1,02	1,35
POTIRON	0,96	0,61	PIVOINE	0,24	0,74	CHAMEAU	3,31	5,41	PÉLICAN	0,3	0,81
ÉPINARD	0,12	0,41	ROMARIN	1,02	1,01	GUÉPARD	0,42	0,34	CORBEAU	3,86	4,19
CAROTTE	4,1	2,97	TILLEUL	0,3	5	MACAQUE	1,57	0,14	GOÉLAND	0,3	0,81
ABRICOT	0,36	1,15	CAMÉLIA	0,18	0,27	CARIBOU	0,48	0,14	MÉSANGE	0,06	6,89

Deuxièmes listes de mots contextuels utilisées dans les Expériences 5 et 7											
FRUITS/LEGUMES Mean Freqfilms:2,15 Mean Freqbook: 2,52			ARBRES/FLEURS Mean Freqfilms:1,04 Mean Freqbook: 3,62			MAMMIFERES Mean Freqfilms: 3,63 Mean Freqbook: 2,67			OISEAUX Mean Freqfilms: 3,94 Mean Freqbook: 3,87		
	Ffilms	Fbook		Ffilms	Fbook		Ffilms	Fbook		Ffilms	Fbook
FIGUE	1,87	1,35	CÈDRE	0,54	2,16	PANDA	1,69	0,14	MERLE	1,39	2,64
PRUNE	1,51	1,55	HÊTRE	0,18	3,38	FURET	0,66	0,14	POULE	26,63	16,82
NAVET	1,39	0,88	SAULE	0,96	1,96	TIGRE	10,3	4,86	HIBOU	2,65	2,36
CASSIS	0,3	3,45	BAMBOU	3,8	3,78	BREBIS	5,66	7,09	CAILLE	1,81	2,97
GOYAVE	0,06	0,2	CYPRÈS	0,54	8,51	BÉLIER	1,87	3,11	PIGEON	9,58	7,97
CERISE	2,77	3,31	MIMOSA	0,36	1,35	COYOTE	1,33	0,34	PINSON	0,54	1,08
BANANE	8,98	4,19	ROSEAU	0,54	2,97	GIRAFE	2,11	1,89	TOUCAN	0,12	0,14
CITRON	8,01	9,05	TULIPE	1,63	0,54	CASTOR	1,87	1,08	RAPACE	0,66	1,49
PRUNEAU	0,96	1,35	PALMIER	1,45	3,04	GAZELLE	1,02	2,77	FLAMANT	0,12	0,07
AIRELLE	0,12	0,07	LAURIER	0,36	3,58	HAMSTER	2,17	0,68	PERDRIX	1,14	1,49
HARICOT	1,2	1,22	LAVANDE	1,02	6,82	TAUREAU	9,88	10,07	MOINEAU	1,93	4,32
POIVRON	0,3	0,27	PLATANE	0,12	4,26	BABOUIN	1,39	0,54	MOUETTE	1,2	5,47
ASPERGE	0,48	5,88	SÉQUOIA	0,18	1,22	GORILLE	7,29	2,03	COLOMBE	3,49	3,51

Premières listes de mots contextuels utilisées dans les Expériences 6 et 7						Deuxièmes listes de mots contextuels utilisées dans l'Expérience 7					
POISSONS/FRUITS DE MER Mean Freqfilms:1,97 Mean Freqbook: 1,85			MINÉRAUX/MATÉRIAUX Mean Freqfilms: 5,77 Mean Freqbook: 13,19			POISSONS/FRUITS DE MER Mean Freqfilms:1,71 Mean Freqbook: 2,17			MINÉRAUX/MATÉRIAUX Mean Freqfilms: 4,26 Mean Freqbook: 14,08		
	Ffilms	Fbook		Ffilms	Fbook		Ffilms	Fbook		Ffilms	Fbook
BULOT	0,06	0	ACIER	19,1	33,38	MOULE	4,04	3,99	MÉTAL	14,58	41,82
MORUE	1,99	4,86	PLOMB	10,12	20,34	CRABE	5,6	7,3	RUBIS	2,71	3,11
CARPE	2,59	2,77	AMBRE	1,51	4,39	COLIN	0	0,54	BÉTON	6,08	15,2
HOMARD	5,12	3,51	BRONZE	6,93	18,58	CALMAR	1,2	0,27	CUIVRE	3,61	30,68
REQUIN	10,3	1,62	TITANE	2,05	0,2	SAUMON	5,18	3,65	LAITON	0,12	1,55
HARENG	1,69	2,91	SAPHIR	1,08	1,22	OURSIN	0,06	0,54	NICKEL	1,63	1,22
TRUITE	1,39	3,38	MARBRE	4,82	41,49	MERLAN	0,24	3,45	ARGILE	4,34	9,32
ROUGET	0,3	0,54	QUARTZ	0,42	1,82	GARDON	0,3	0,41	TOPAZE	0,18	0,81
SARDINE	0,54	1,28	BASALTE	1,99	0,47	ANCHOIS	1,57	2,57	GRANITE	0,24	0,47
DAURADE	0,18	0,74	URANIUM	0,06	0,95	ESPADON	1,08	0,47	BAUXITE	0,12	0,41
LIMANDE	0	0,61	CRISTAL	13,67	13,72	BROCHET	0,24	3,18	DIAMANT	14,34	14,12
SEICHES	0,24	0,68	GOUDRON	2,05	7,5	PIRANHA	0,72	0	MERCURE	0,96	1,76
POULPES	1,2	1,22	PÉTROLE	11,27	14,73	PIEUVRE	1,99	1,82	CHARBON	6,51	22,03

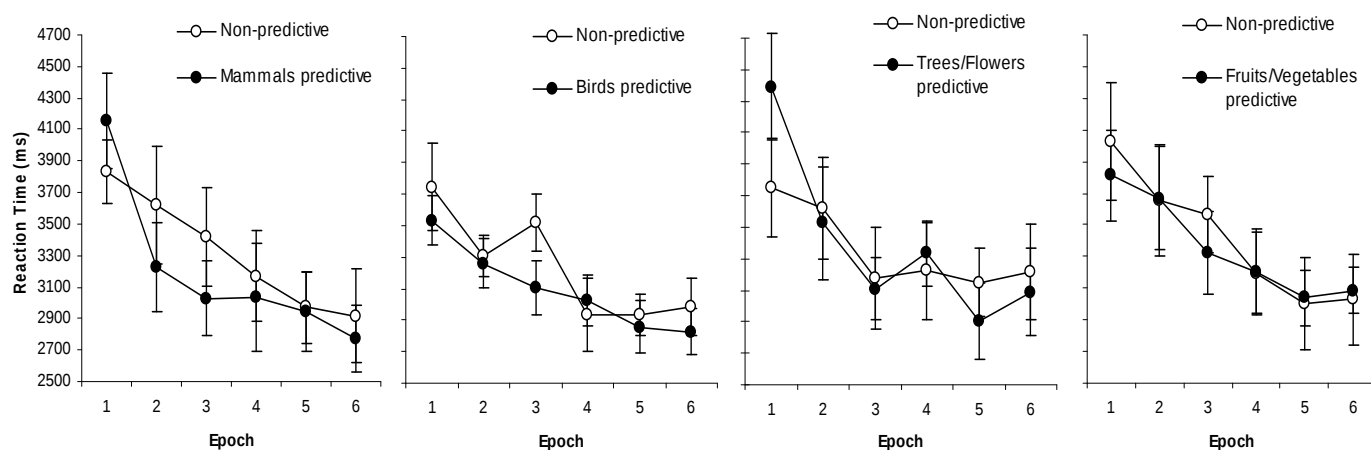
Listes des mots de remplissage utilisées
dans les Expérience 6 et 7

FOURNITURES DE BUREAU

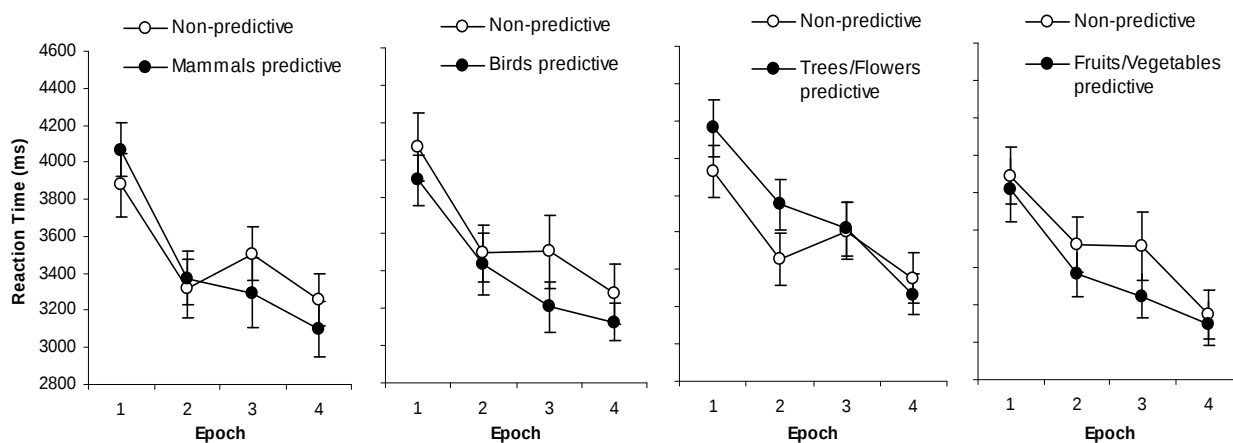
CRAIE	PLUME
ATLAS	GOMME
STYLO	COLLE
CARNET	AGENDA
CAHIER	MANUEL
FEUTRE	DESSIN
TIMBRE	CRAYON
CUTTER	TAMPON
CISEAUX	PALETTE
TROUSSE	CALEPIN
CROQUIS	COLLAGE
DOSSIER	TABLEAU
PINCEAU	ARCHIVE

ANNEXE 6

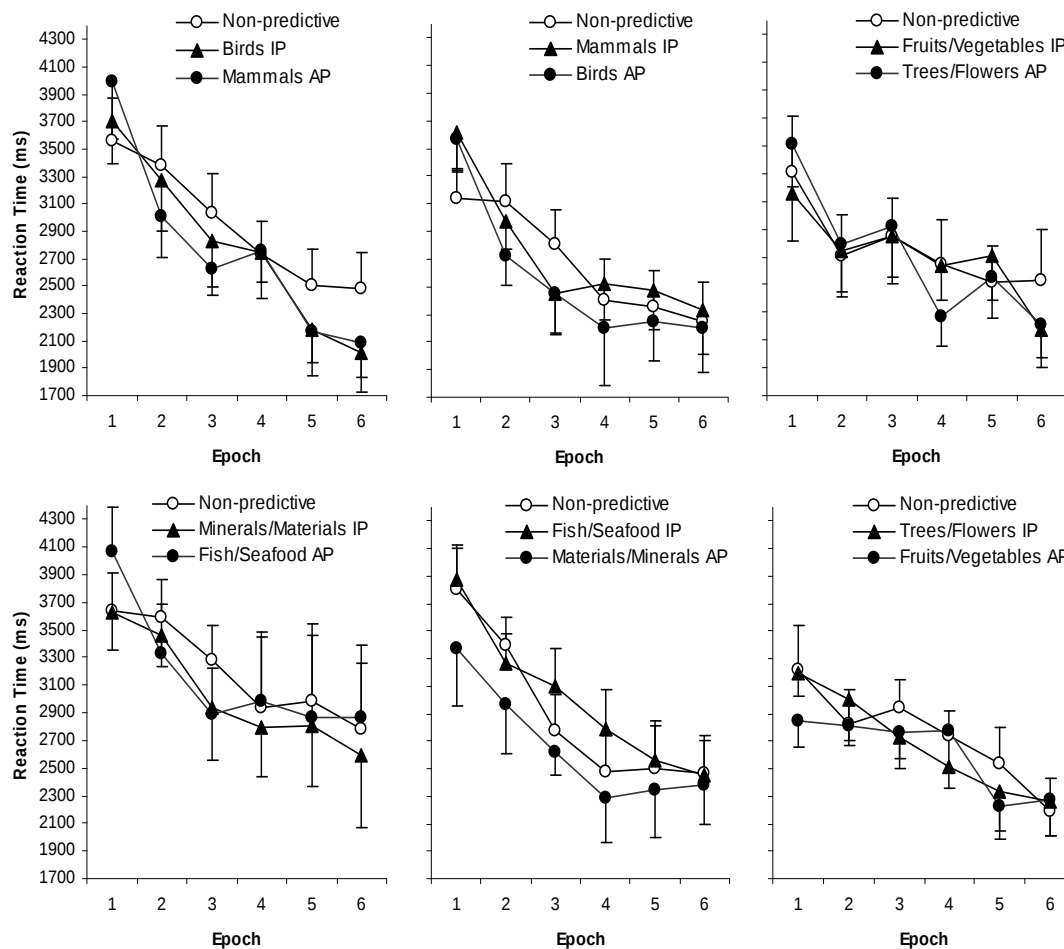
Etude 2 : Résultats relatifs à chaque catégorie prédictive



Annexe 6a : Résultats détaillés de l'Expérience 4. Temps de réponse moyens dans les conditions prédictives, mammifères, oiseaux, arbres/fleurs et fruits/légumes, ainsi que dans les conditions non prédictives correspondantes. Les barres d'erreur correspondent aux SEM (N=12).

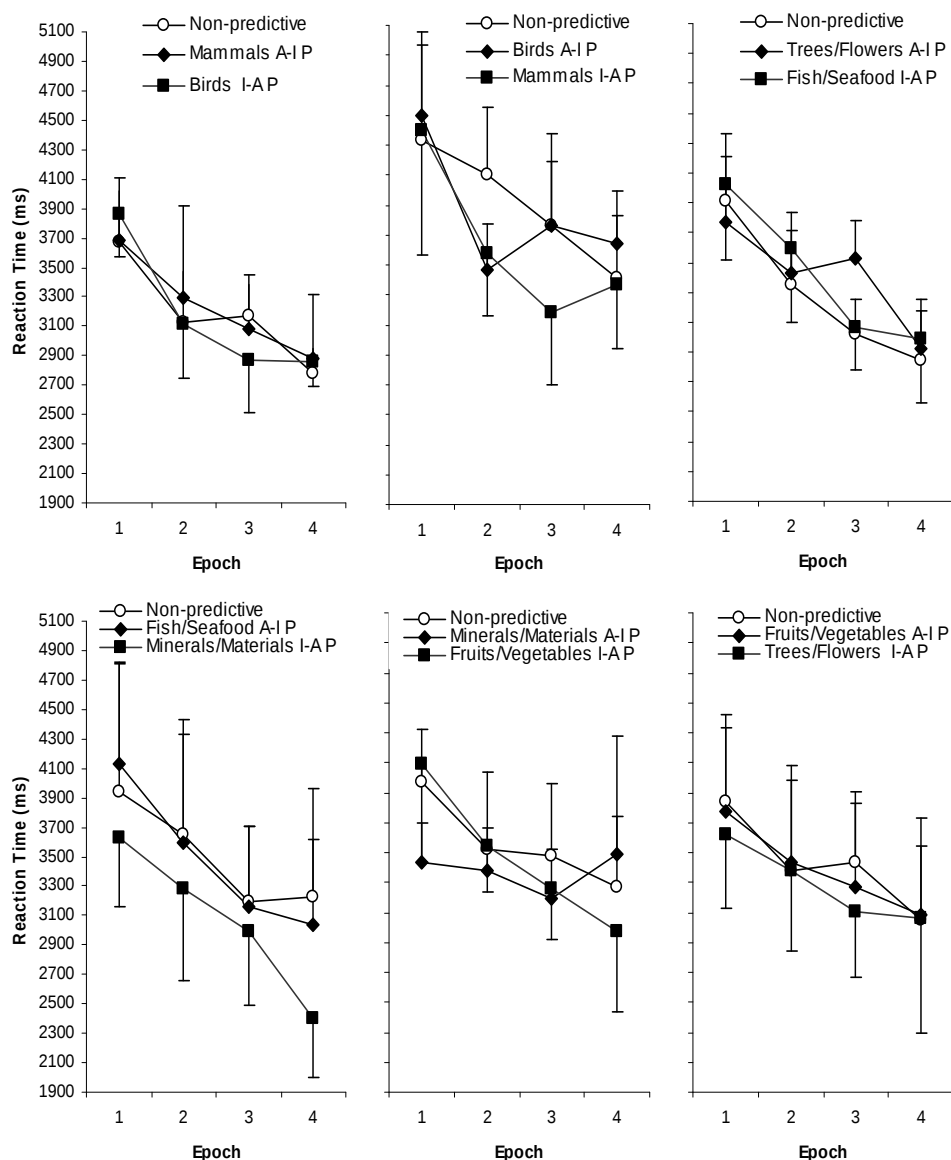


Annexe 6b : Résultats détaillés de l'Expérience 5. Temps de réponse moyens dans les conditions prédictives, mammifères, oiseaux, arbres/fleurs et fruits/légumes, ainsi que dans les conditions non prédictives correspondantes. Les barres d'erreur correspondent aux SEM (N=28).



Annexe 6c : Résultats détaillés de l'Expérience 6. Temps de réponse moyens dans les conditions prédictives attendues (mammifères, oiseaux, arbres/fleurs, fruits/légumes, matériaux/minéraux et poissons/ruits de mer), dans les conditions prédictives ignorées (mammifères, oiseaux, arbres/fleurs, fruits/légumes, matériaux/minéraux et poissons/ruits de mer), ainsi que dans les conditions non prédictives correspondantes. Les barres d'erreur correspondent aux SEM (N=6).

IP : Prédictive Ignorée, AP : Prédictive attendue

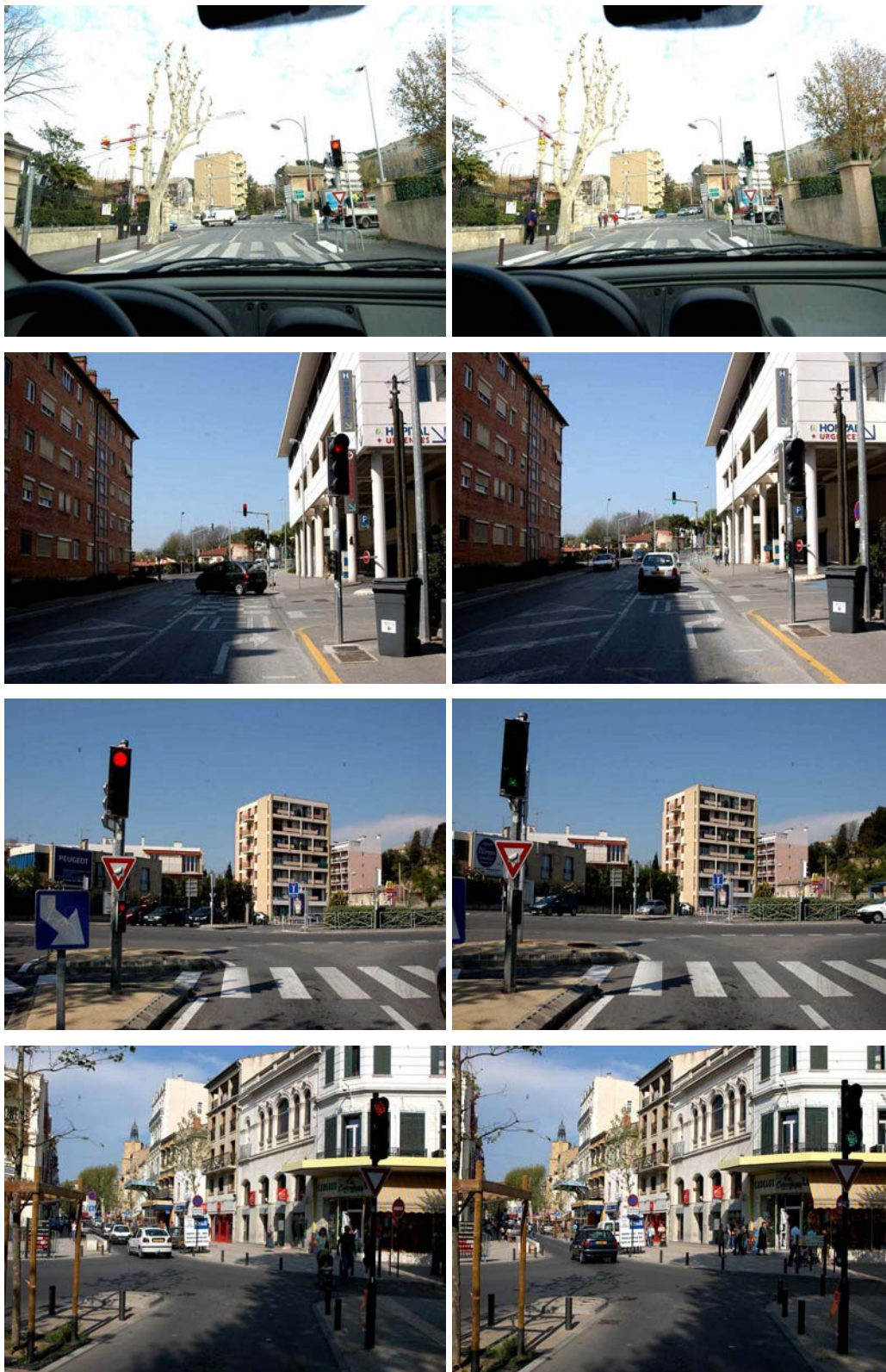


Annexe 6d : Résultats détaillés de l'Expérience 7. Temps de réponse moyens dans les conditions prédictives attendues puis ignorées (mammifères, oiseaux, arbres/leurs, fruits/légumes, matériaux/minéraux et poissons/fruits de mer), dans les conditions prédictives ignorées puis attendues, (mammifères, oiseaux, arbres/leurs, fruits/légumes, matériaux/minéraux et poissons/fruits de mer), ainsi que dans les conditions non prédictives correspondantes. Les barres d'erreur correspondent aux SEM (N=6).

A-I P : Prédictive attendue puis ignorée, I-A P : Prédictive ignorée puis attendue.

ANNEXE 7

Etude 3 : Matériel utilisé



Feux utilisés



Passages-piétons utilisés



Rond-points utilisés

ANNEXE 8

Etude 3. Corrélations entre les performances durant la tâche de prise de décision et les verbalisations

	TR (non prédictif) - TR (prédictif) à l'époque 6
Participants (n')ayant :	Bénéfice dans la condition prédictive
remarqué les régularités prédictives (N=10)	85ms
pas remarqué les régularités prédictives (N=2)	-45 ms
Participants ayant remarqué les régularités prédictives et (n')ayant :	
remarqué le transfert (N=5)	54ms
pas remarqué le transfert (N=5)	116ms

Annexe 8a : Différence entre les performances sur les TR observées dans la condition prédictive et celles observées dans la condition non prédictive à l'époque 6

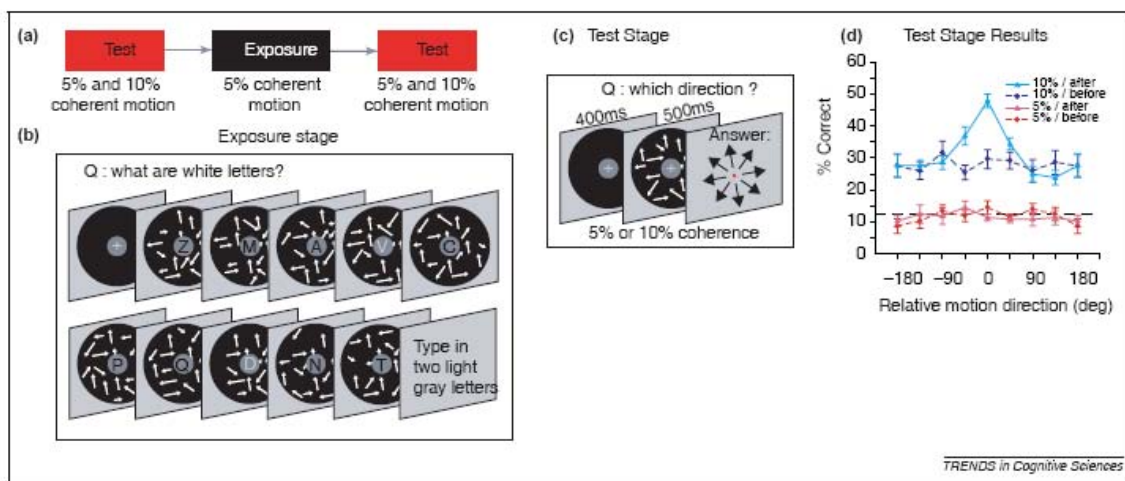
	Performances (contre prédictif) - Performances (non prédictif)	
Participants (n')ayant :	Coût sur les TR dans la condition contre prédictive	Pourcentage d'erreurs supplémentaires
remarqué les régularités prédictives (N=10)	420ms	10.00%
pas remarqué les régularités prédictives (N=2)	-190ms	0%
Participants ayant remarqué les régularités prédictives et (n')ayant :		
remarqué le transfert (N=5)	309ms	3.33%
pas remarqué le transfert (N=5)	530ms	16,66%

Annexe 8b : Différence entre les performances observées dans la condition contre prédictive et celles observées dans la condition non prédictive au cours du transfert

ANNEXE 9

L'apprentissage de signaux en mouvement cohérent

Un paradigme récemment développé par Watanabe et collaborateurs (pour une revue, Seitz & Watanabe, 2005) met en évidence une forme d'apprentissage perceptif lors d'expositions répétées à un signal en mouvement. Les participants sont exposés à des affichages de points dynamiques, dont un petit nombre se déplacent selon une trajectoire cohérente (signal), alors qu'un plus grand nombre se déplacent selon une trajectoire aléatoire (bruit) (Cf. Figure b). Des expériences préliminaires révèlent que lorsque moins de 8% des points se déplacent selon une trajectoire cohérente, la force du signal est inférieure au seuil de visibilité. Les participants ne parviennent ni à identifier, ni à discriminer la direction du mouvement, ni même à détecter un mouvement cohérent. Pourtant, les participants devenaient sensibles à la répétition de trajectoires cohérentes imperceptibles.



Les expériences comportent une phase d'exposition, précédée et suivie d'une phase test (Cf. Figure a). Durant la phase d'exposition, les participants réalisent une tâche d'identification de lettres présentées au centre de l'affichage, de telle manière que le mouvement des points constitue un trait non pertinent pour la tâche (Cf. Figure b). Durant cette phase, le ratio signal/bruit est toujours de 5% (i.e. 5% des points se déplacent selon une trajectoire cohérente et 95% selon des trajectoires aléatoires), de sorte que les sujets n'accèdent pas consciemment au signal en mouvement cohérent durant cette tâche. Durant les phases tests, précédant et succédant la phase d'exposition, les affichages peuvent comporter 5% ou 10% de cohérence. Les tâches à réaliser durant les phases tests étaient différentes d'une expérience à l'autre. Dans l'expérience princeps (Watanabe, Nanez, & Sasaki, 2001), les sujets devaient juger la direction du mouvement des points en indiquant parmi 8 directions, celle correspondant au mouvement qu'ils venaient de percevoir (Cf. Figure c). Les résultats, présentés Figure d, révèlent que pour des affichages présentant 5% de cohérence, les performances n'étaient pas différentes du hasard aussi bien lors de la phase test précédant la phase d'exposition que lors de la phase test lui succédant (environ 12,5% de réponses correctes). Ces résultats confirment que les participants ne perçoivent pas consciemment la direction d'un ensemble de points lorsque seulement 5% de ces points se déplacent selon une direction cohérente. En revanche, lorsque les affichages comportaient 10% de cohérence, les performances étaient largement supérieures au hasard (environ 30% de réponses correctes),

mais surtout, les performances étaient considérablement améliorées pour la direction à laquelle ils avaient été exposés durant la phase d'exposition. Les performances atteignaient 50% de réponses correctes dans la deuxième phase test pour les directions présentées durant la phase d'exposition, alors que le mouvement cohérent était imperceptible (le ratio dans la phase d'exposition était de 5% de cohérence). Aucun effet de la répétition du signal en mouvement cohérent n'était observé lorsque les deux phases tests se succédaient l'une à l'autre, attestant que cette amélioration ne tenait pas à la première phase test mais bien à la phase d'exposition. Des effets identiques étaient obtenus lorsque la tâche « test » des sujets était de juger la cohérence du mouvement perçu, ou encore de juger si deux patterns de points se déplaçaient selon la même trajectoire. Ainsi, une exposition à un mouvement cohérent invisible répété, peut rehausser l'efficacité d'indication de direction, de détection de mouvements cohérents, de discrimination de direction, mais seulement si ces signaux cohérents deviennent par la suite visibles, i.e. présentant une cohérence de 10%.

Ces études relèvent plusieurs phénomènes intéressants. D'une part, l'apprentissage repose sur des répétitions de signaux non pertinents pour la tâche dans laquelle l'attention du sujet est engagée. D'autre part, bien que la cohérence des signaux soit en dessous du seuil de visibilité, et que cette cohérence ne soit pas pertinente pour la tâche (tâche d'identification pendant la phase d'exposition), l'exposition répétée améliore les performances dans une tâche ultérieure. Ces résultats suggèrent que l'attention focalisée n'est pas nécessaire à l'apprentissage perceptif et donc n'est pas un critère décisif à la plasticité. Néanmoins, Seitz et Watanabe (2003) ont montré que le rehaussement de cette sensibilité nécessite un couplage temporel entre la présentation subliminale du stimulus en mouvement, non pertinent pour la tâche, et une tâche cible. Dans leur expérience, quatre directions de mouvements différentes étaient présentées de manière équivalente durant la phase d'exposition, mais une seule direction était couplée avec les stimuli comportant la cible (i.e. une lettre grise). Or, l'apprentissage n'émergeait que pour les directions qui étaient temporellement associées, ou couplées, avec les cibles de la tâche, et pas pour les autres directions de mouvements.

Résumé

En structurant le monde visuel, les connaissances relatives aux régularités de l'environnement facilitent l'orientation des processus de sélection attentionnelle vers les aspects pertinents. Dans ce cadre, quel rôle peut-on accorder aux traitements non conscients dans l'apprentissage de régularités contextuelles ? Dans quelle mesure des connaissances inaccessibles à la conscience influencent-elles la perception ? Le travail de recherche rapporté dans ce mémoire de thèse visait à mieux comprendre les mécanismes d'apprentissage, implicite ou explicite, impliqués lors de l'exploration d'une scène visuelle. Nos travaux expérimentaux montrent que des connaissances relatives aux régularités contextuelles peuvent être acquises de manière implicite et faciliter le guidage de l'attention au sein d'une image. Ils révèlent que des mécanismes d'apprentissage implicite peuvent reposer sur des régularités contextuelles spécifiques, mais également sur des régularités catégorielles et sémantiques. En outre, ces mécanismes d'apprentissage implicite peuvent être déployés sur des régularités sémantiques hors du focus attentionnel, mais pour que la connaissance s'exprime une attention sélective était ici requise. Par ailleurs, nos travaux montrent que si l'apprentissage de régularités contextuelles peut faciliter la prise de décision dans une tâche écologique telle que la conduite automobile, il peut également conduire à l'émergence de défaillances fonctionnelles. Sans minimiser le caractère adaptatif de la conscience dans la perception de scènes visuelles complexes, nos recherches amènent à défendre la thèse selon laquelle « l'inconscient cognitif » est capable d'encoder des régularités perceptives ou sémantiques présentes dans l'environnement, et de guider l'attention dans l'analyse d'une scène visuelle.

Mots clés

Perception visuelle, apprentissage implicite, indiçage contextuel, régularités sémantiques, schémas de scène

Abstract

When structuring the visual world, information on the contextual regularities of the environment facilitates the processes of attentional selection toward its relevant aspects. Within this framework, which part can be attributed to the non conscious processes involved in learning contextual regularities? To which extent can unconsciously available information influence perception? This thesis reports research aimed at understanding the learning mechanisms – implicit or explicit – that are involved in exploring a visual scene. In our experimental studies, we have observed that information relating to contextual regularities can be acquired implicitly and contribute to the guidance of attention within an image. The experiments revealed that implicit learning mechanisms can rely on specific contextual regularities but also on categorical and semantic regularities. In addition, these learning mechanisms can be used on semantic regularities out of the attentional focus. However, effective use of this information requires selective attention. Furthermore, in ecologically valid tasks such as car driving, our results suggest that if the learning can facilitate “decision making”, it can also lead to the emergence of dysfunctions. Keeping in mind the adaptive role of consciousness on the perception of complex visual scenes, our research supports the idea that the “cognitive unconscious” can encode perceptual or semantic regularities from the environment, and use these to guide attention in the analysis of a visual scene.

Key words

Visual perception, implicit learning, contextual cueing, semantic regularities, scene schemas.