



**Attention, Semantics, and Transfer in Implicit Learning:
Insights for Theories of Consciousness
from an Adapted Contextual Cueing Paradigm**

Inauguraldissertation

zur

Erlangung des Doktorgrades
der Humanwissenschaftlichen Fakultät
der Universität zu Köln

nach der Promotionsordnung vom 18.12.2018

vorgelegt von

Felice Lluvia Tavera-Salyutov

aus Barcelona

Juni 2025

Diese Dissertation wurde von der Humanwissenschaftlichen Fakultät der
Universität zu Köln im September 2025 angenommen.

Abstract

Studying of human consciousness poses many challenges. Numerous approaches seek to define consciousness and how it may be scientifically studied. This dissertation advocates a theory-driven approach to the empirical study of consciousness. Specifically, it considers the Global Workspace Theory, the Higher-Order Thought Theory, and the Integrated Information Theory. It aims to identify the boundaries of unconscious processing proposed by these theories to clarify the functional role of consciousness. To this end, I developed a novel variant of the contextual cueing paradigm to examine implicit contingency learning as a proxy for unconscious processing in three studies. Study 1 tested the role of attention, specifically, whether implicit learning of contingencies between features depended on their task relevance, and whether cue competition effects occurred without awareness. Learning occurred independently of task relevance, with no evidence of cue competition. Study 2 examined the hypothesis of modularized unconscious processing by examining transfer effects of contingency knowledge from one feature to another without explicit instructions. Such transfer effects were observed, suggesting that contingency information can be exchanged between processing modules without being consciously accessible. Study 3 tested whether semantic information can be used in implicit learning. We found evidence of semantic category learning without awareness, but observed a reversed effect: not a learning benefit, but potential inhibitory effects based on the learned contingencies. Across all three studies, we implemented a refined test and analysis procedure for detecting explicit knowledge. This approach aimed to improve both the sensitivity and reliability of the awareness test assessment compared to conventional recognition-based measures. The findings are discussed in light of the three aforementioned theories of consciousness and their respective predictions. This dissertation contributes to refining theoretical models of consciousness, and opens new pathways for their empirical evaluation. It underscores the value of implicit learning paradigms alongside other methodological tools.

Content

Abstract	3
1 Introduction	6
2 Approaches to Consciousness Research	9
2.1 Atheoretical Approaches	10
2.2 Theoretical Approaches	16
3 Unconscious Processing.....	24
3.1 Experimental Paradigms of Subliminal Stimulus Presentation	27
3.2 Implicit Learning Paradigms	29
4 Measures of Conscious Awareness	33
4.1 Direct and Indirect, Objective and Subjective Measures.....	34
4.2 Confidence Measures.....	39
5 The Contextual Cueing Paradigm.....	42
5.1 Learning Mechanisms.....	43
5.2 Limitations	45
5.3 A New Variant of the Contextual Cueing Paradigm	49
6 Attention and Consciousness	52
6.1 Attention in Implicit Learning	54
6.2 Summary Study 1.....	57
6.3 Implications	61
7 Modular Processing of Unconscious Information	65
7.1 Transfer in Implicit Learning.....	66
7.2 Summary Study 2.....	70
7.3 Implications	71
8 Semantic Processing Without Awareness	74
8.1 Semantic Processing in Implicit Learning	78
8.2 Summary Study 3.....	79
8.3 Implications	81
9 Conclusion	87
9.1 Limitations and Future Directions	91
References.....	94

Publication List.....	137
Appendix A.....	140
Appendix B.....	184
Appendix C.....	206

1 Introduction

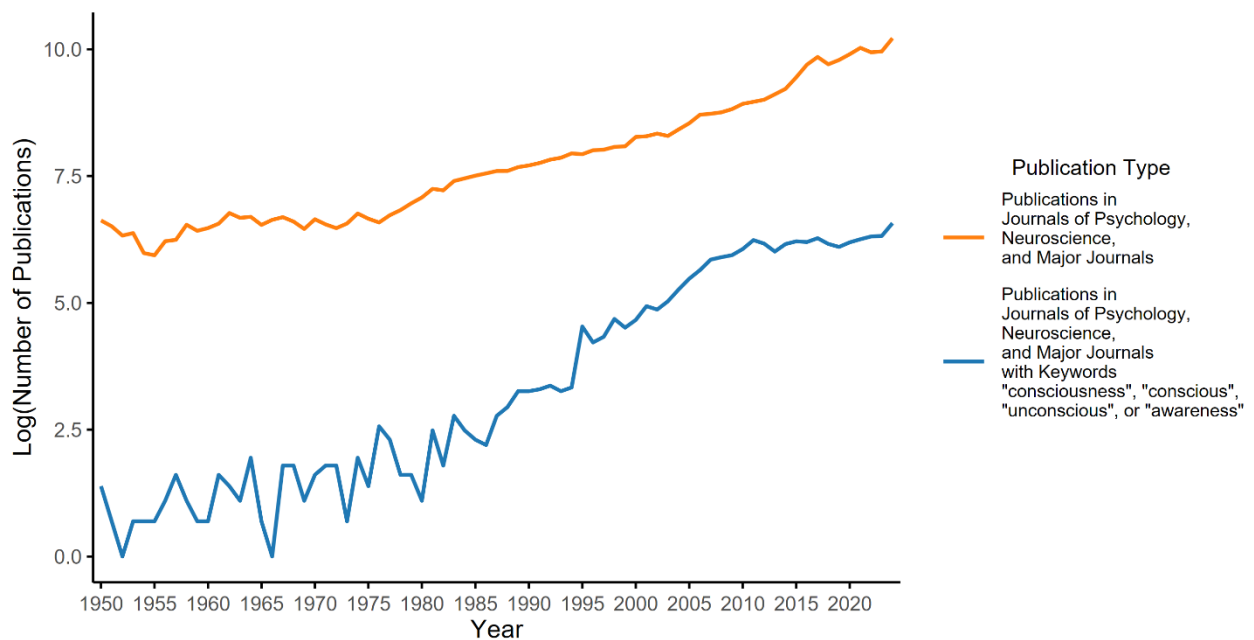
Consciousness might be the only remaining mystery to humankind, according to philosopher Daniel Dennett (1993). Thereby, he did not imply that we do not know anything about consciousness, or that we know less about consciousness than we know about, say, deep-sea creatures or the universe. Rather, he implied that we do not know how to approach the concept of consciousness. We have not agreed on a way to think about it, and about parameters that should be included in a theory of consciousness. And since Dennett wrote this statement in 1993, it has rather become more accurate than outdated. Over the past two decades in particular, research on consciousness has gained significant momentum, as reflected in the increasing number of respective publications (see Figure 1). As depicted in Figure 1, this growth outpaces the general increase in publications within the same journals, highlighting the rising scientific interest in the topic. Despite this intensified effort, spanning psychology, neuroscience, philosophy, and physics, there remains little agreement on a common conceptual ground, neither within disciplines nor in interdisciplinary work. This is signified by the myriads of theories on consciousness, and their vast differences. There have been tremendous endeavors to test hypotheses deducted from those theories, and again, there is a vast heterogeneity in doing so – a myriad of research questions, approaches, paradigms, mechanisms on quantum, molecular, neural, cognitive, and behavioral levels are investigated. In this work, I will argue for one approach that is promising for testing hypotheses derived from prominent (neuro-)psychological theories of consciousness, that is methodologically sound, and has relevant real-life implications.

In the following chapters, I will develop an argument for studying consciousness on the grounds of empirically testable theories, and for examining parameters that constitute their explanation of what consciousness is. I will also argue for an approach to develop experimental designs that allow for flexible manipulations of such parameters, and then to carefully measure

for conscious awareness¹. I will further assert that this is especially viable by implementing implicit learning paradigms.

Figure 1.

*Trends in Consciousness-Related vs. Other Publications Within the Same Journals
(Logarithmic Scale)*



Note. The graph was generated based on an advanced search in PubMed. Search keywords specified publications in the following journals: *Consciousness and Cognition*, *Neuroscience of Consciousness*, *Psychological Science*, *Psychological Review*, *Psychological Bulletin*, *Journal of Experimental Psychology: General*, *Journal of Experimental Psychology: Human Perception and Performance*, *Journal of Experimental Psychology: Learning, Memory, and Cognition*, *Cognitive Psychology*, *Cognition*, *Cognitive Science*, *Memory & Cognition*, *Journal of Memory and Language*, *Psychonomic Bulletin & Review*, *Attention, Perception, & Psychophysics*, *Behavior Research Methods*, *Trends in Cognitive Sciences*, *Wiley Interdisciplinary Reviews: Cognitive Science*, *Frontiers in Psychology*, *Philosophical Psychology*, *Journal of Cognitive Neuroscience*, *Neuropsychologia*, *Brain and Cognition*, *Brain Research*, *Nature Neuroscience*, *Frontiers in Human Neuroscience*, *Nature*, *Science*, *Nature Human Behaviour*, *Nature Communications*, *Scientific Reports*, *Cereb Cortex*, *J Neurosci*, *Behav Brain Sci*, *Proc Natl*

¹ Note that I use the terms *consciousness*, *awareness*, and *conscious awareness* interchangeably here.

Acad Sci U S A. Note that some of the keywords refer to PubMed-specific abbreviations rather than the full official journal names. The keywords *consciousness*, *conscious*, *unconscious*, and *awareness* were added to determine the consciousness-related publications.

2 Approaches to Consciousness Research

Different theories of consciousness set distinct foci with respect to the to-be-explained phenomena, and hence, the empirical work they inspire. They also differ in their level of explanation and approach to thinking about consciousness to begin with. There are different possibilities of approaching the goal to explain and test consciousness. One of the distinctions in approach lies in the question of what has to be explained: On the one hand, we need to answer the functional question of consciousness, which is why consciousness developed in the course of evolution, and what functions does it entail in our cognitive system (Dennett, 2014). This may be understood in the third-person perspective (Chalmers, 1997), such that we study consciousness in our fellow human beings. On the other hand, we should address the aspect of phenomenology of consciousness, which is also called “qualia” or “what-it-is-likeness” (Nagel, 1980). That is the first-person perspective (Chalmers, 1997), the subjective feeling of what it is like to be a conscious person.

Another central distinction between approaches to study consciousness is the explanatory basis or level of explanation: When formulating functions, mechanisms, and phenomena around consciousness, one can connect those to neurobiological mechanisms, hypothesizing localizable activation in the brain that is associated with conscious processing, and one can then test them empirically. Other frameworks are more abstract and philosophical, like the idea of a mind-body dualism (Descartes, 1901) that was soon widely viewed as reductionist and outdated or even unscientific (Fodor, 1981). Other theories are then even concerned with integrating the concept of consciousness into the cosmos (Hameroff & Penrose, 2014) or explaining it on the grounds of quantum mechanics, not on neurobiological grounds, aiming for a most fundamental theory of consciousness (Atmanspacher, 2004; Chalmers & McQueen, 2022; T. Li et al., 2019). Thus, different theories choose different aspects of consciousness they aim to explain, and different levels of explanation. In this dissertation, the focus will be frameworks that revolve

around functions and mechanisms of consciousness that are empirically testable in the scope of psychological methodology.

Within this, it is important to note that when taking a functionalist perspective on consciousness, there are still two approaches to distinguish. On the one hand, consciousness can mean a more global state of an organism, such as being awake versus being asleep or unconscious, or something in between which might be dreaming, mind-expanding drugs, or vegetative states (i.e., intransitive consciousness; Seth & Bayne, 2022; Weitzel & Bavishi, 2024). But consciousness can also refer to conscious awareness being directed at a specific object, asking whether a generally conscious person is consciously aware of a certain object or stimulus in their environment (i.e., transitive consciousness), which would then be considered content-based consciousness (Hohwy, 2009; Overgaard & Overgaard, 2010). In this dissertation, I will examine the latter, presupposing intransitive consciousness as a given. This means that I will investigate cases in which globally conscious individuals lack conscious awareness of specific stimuli or information.

2.1 Atheoretical Approaches

Before describing and discussing theories of consciousness, it is important to establish why we should formulate theories or models of consciousness to begin with. Could there be a simpler or more direct way to gain insight into the nature of consciousness?

Such an approach could be the investigation of neural correlates of consciousness (NCCs), which is the minimum neural activity sufficient for consciousness to occur (Baars, 1994; C. Koch et al., 2016). In empirical practice, that means to try and contrast one and the same cognitive process, once conscious, once unconscious, and propose that the difference seen in neural activation patterns is what constitutes consciousness (Klein et al., 2020). To give an example, Pins and Ffytche (2003) presented participants a hardly visible soft light-dark grating pattern with low contrast, and asked them in each trial to indicate whether they saw something

or not. They contrasted functional brain imaging in trials in which participants reported having seen or not seen the stimulus. Results suggested that activity in the occipital lobe (100ms after stimulus presentation) was a primary correlate of consciousness. While this seems a straightforward method to locate consciousness in the brain, there are several methodological problems with that approach. First, neural activity that is found to be exclusive to conscious trials could be contaminated by activity linked to response processes, thus overestimating the NCC (Marois & Ivanoff, 2005). This can be circumvented with no-report paradigms though (Tsuchiya et al., 2015). A second central problem remains, however, as, when looking for information that is perceived consciously versus not consciously (content-based consciousness), the organism as such is conscious in both conditions which makes it difficult to distinguish neural patterns of conscious contents from global state consciousness (Chalmers, 2000; Hohwy, 2009). And not only is this approach methodologically much more complex than it might appear at first glance, but it also inherently has limitations that are unlikely to be overcome by even an ideal method. Because the logic of looking for neural activation patterns that are exclusively linked to consciousness does not make the essential distinction between prerequisites and consequences of consciousness (Q. Li et al., 2014; Seth & Bayne, 2022; but see Sandberg et al., 2014, for a potential solution). Also, searching for NCCs is not “metaphysically neutral”, as often assumed, but it comes with theoretical presuppositions that might be viewed as reductionist (Klein et al., 2020). Because given the vast number of interactions between neuronal states and the possibilities of causal influence within them are no easy conditions for an experimental investigation (C. Koch et al., 2016). Thus, it is not trivial to gain any insight from mere brain activity patterns without specified hypotheses and the potential to falsify them empirically.

Within the debate around NCC research, one debate that has been specifically intriguing is the debate around the “unfolding argument” (Doerig et al., 2019; Usher et al., 2023). The unfolding argument revolves around the notion that consciousness does not only need to be

understood as a state, but as a process. That includes predicting not only one state of neuronal activity that produces consciousness, but observing neuronal activity over the course of time, during which consciousness “unfolds”. This is especially relevant for the integrative function of consciousness put forward by most theories of consciousness (Cleeremans & Frith, 2003). Information integration includes integration over time, and thus, consciousness has to be understood as integrating neuronal activity over time. Integration also entails a causal determination of neuronal activity. Because patterns of neuronal activity influence each other and create a causal chain of events that then characterizes consciousness. The current debate regarding the unfolding argument concerns the specific architecture of such causal processes (Doerig et al., 2019; Usher et al., 2023). There are models of consciousness that rely on recurrent feedback structures, such as recurrent processing or global workspace theories (Baars, 2005; Lamme & Roelfsema, 2000). Doerig et al. (2019) proposed that any recurrent neural network can in principle be replicated by a feed-forward network. They thus claim that theories that hypothesize recurrent feedback are overly complex and incompatible with the unfolding argument. Usher et al. (2023) on the other hand rejects this notion and claims that recurrent networks are, also functionally, much more complex, given their interactions, and so is consciousness itself. They further argue that the unfolding argument puts forward claims that are oversimplified, and draws conclusions that are broadly discarding many theories of consciousness and methodology that is used to study consciousness. This makes it a problematic starting point for any further research on the subject. This debate is not settled, but it remains an interesting question for NCC research and theories of consciousness, whether a feed-forward or recurrent feedback architecture should be adopted.

Similarly, it would be conceivable to approach consciousness research solely via computer simulations. This can be done with feed-forward or recurrent networks or large language models (LLMs). However unattainable that may sound, the goal would be to produce an

artificial conscious system by reaching the complexity of neuronal activation patterns and interactions. Yet, there has been a critique of this idea, stating that whichever systems built to mimic our neural networks lack decisive parts that are hypothesized to enable or contribute to human consciousness (Aru et al., 2023). The authors basically argue that there is a lack of “hardware” in current artificial systems. For example, that an algorithm does not have a thalamic structure, which many theories include in the mechanisms of consciousness (Dehaene & Naccache, 2001; Gennaro, 2004; Lamme, 2010; Tononi, 2004), nor dual-compartment pyramidal neurons, as put forward by the dendritic integration theory (Bachmann et al., 2020), or a global workspace structure (Baars, 2005), or an ascending arousal system (Solms, 2018; Solms & Friston, 2018). They claim that, “[t]opologically, present-day AI systems are extremely simple in comparison, which is among the reasons we are cautious in ascribing phenomenal consciousness to them.” (Aru et al., 2023). This argument is problematic on multiple levels. First, all theories on consciousness could be wrong, and thus, their reference to, or reliance on specific biological structures would become irrelevant. If the theories were right about, for example, the involvement or necessity of a thalamic structure for consciousness, why should an artificial system or other biological organisms not be capable of mimicking the functions of such a structure? The lack of “hardware” does not negate a potential functional equivalence (see also Rouleau & Levin, 2025). And even the mere lack of (functional) complexity of artificial systems as they are now, seems no valid argument against the possibility to create sufficiently complex systems in principle. What is most striking about that argument though, is the nonchalance with which the authors are claiming to be cautious to ascribe phenomenal consciousness to artificial systems, as if there was any basis on which one could do so. It is certainly a valid point to emphasize that current theories of consciousness aim to explain specifically *human* consciousness, and might explain other forms of consciousness, not attached to a human body, as a by-product. But to claim that because of the object of research (i.e., humans) that aims to explain a phenomenon such as consciousness, the phenomenon can in principle not occur in other

objects of research (other organisms or artificial systems), is circular reasoning. Even more so, as we would not even know what to look for as evidence of phenomenal consciousness (artificial or not) of a system to begin with (although, for recent attempts to adapt the Turing test see Gams & Kramar, 2024; C. Koch & Tononi, 2008). What this argument shows is how distant a functional computer model of human brain complexity still is. This is problematic, given that this would be the bare minimum to attempt to model consciousness. The challenges of the next steps would be to first determine *sine qua non* conditions for consciousness, and then being able to test a model or artificial system for consciousness.

An interesting argument that may put an atheoretical approach to consciousness in doubt is the thought experiment illustrated in Gidon et al. (2022). Assuming that one could account for the vast number of interactions and feedback circuits and complex connectivity in the brain – would that make a purely neuroscientific exploration of consciousness reasonable? Because hypothetically, it could be possible to model consciousness like that in the future. Specifically, by recording all action potentials in every individual neuron in a human brain during conscious processing (e.g., of a simplistic stimulus) and a response, and replaying them to the same brain, neuron by neuron with voltage clamps, overriding any other activation. The question that Gidon et al. (2022) then asks is: Does that human experience the stimulus consciously, and make the response consciously? As the neuronal activity is exactly the same in the replay as in the original situation, it is hard to argue that would not be, as one would have to find a factor beyond neuronal activity that produces or alters conscious experience – which bears the risk of Cartesian dualism (Dennett, 1993). And then, going one step further, Gidon et al. (2022) hypothesize to disconnect all synapses (e.g., chemically), and then replaying the neuronal activity. Lastly, in a third step, a region of the brain (e.g., in the case of a visual perception, the visual cortex) would be surgically removed from the rest of the brain. The question remains the same with these two further steps: Will there be conscious perception of the stimulus and response? Even if,

ultimately, the brain is no longer part of the body? This thought experiment is certainly intriguing, and is relevant for any neurological research approach to studying consciousness. It is crucial from an epistemological perspective to determine what neurological insights mean for consciousness research. There are certainly perspectives that criticize any purely neurological investigation of psychological phenomena as reductionist (Mausfeld, 2012), but Gidon et al. (2022) certainly raises an interesting point, challenging the stance to epistemologically reject neurological activity as a source of insight. But it also highlights the issues with an opposite position. If one were to reject the notion of consciousness being based on neuronal activity, one is at risk of explaining consciousness through parameters that go beyond empirically measurable phenomena, and scientific methodology. Should we be tempted to do that, we would not have come much further from the mind-body dualism of Descartes, and a homunculus-type pineal gland explanation of consciousness (but see Shapiro, 2011 for a critical appraisal). Interestingly, this dualism is something that researchers especially in the neurosciences seem to fall for too often, as illustrated by their designation of the brain and the entire person as two independent subjects (Mudrik & Maoz, 2015).

Taken together, the atheoretical approaches to consciousness are not reconcilable or viable with psychological methodology, but with computational, simulation or neurological methodology. In addition, they raise a myriad of issues concerning epistemology, including issues that apply more generally to the cognitive neurosciences. Further, although they might be independent from any specific theory or model of consciousness, they are not neutral in terms of a priori assumptions and epistemological convictions. Therefore, it does not seem expedient to solely aim for an atheoretical approach to study consciousness. In contrast, developing empirically testable theories of consciousness that are (relatively) transparent in terms of their a priori and theoretical assumptions, seems a viable way to move forward in consciousness research.

2.2 Theoretical Approaches

There are many theories of consciousness that take a functional, (neuro-)psychological approach to explaining consciousness. In a recent review, Seth and Bayne (2022) summarized and categorized prominent consciousness theories. Of those, I will go into more detail with the global workspace theory (GWT; Baars, 2005; Dehaene & Naccache, 2001), higher-order theories (HOT; Rosenthal, 2005), and integrated information theories (IIT; Tononi, 2004). For the sake of completeness, I will note that other prominent approaches are re-entry and predictive processing theories (RPT; Lamme & Roelfsema, 2000), and a theory that conceptualizes consciousness as a memory system (Budson et al., 2022).

For the research questions of the work at hand, GWT, HOT, and IIT accounts are most relevant. Therefore, there will be a short introduction into their key concepts, and a comparison of the three in terms of three key components that have been put forward to model consciousness – attention, modular processing, and semantic processing (see also Table 1). Those key components will be central for the empirical work of this dissertation.

GWT (Baars, 2005) has since its formation been transformed into a neurocognitive theory (Dehaene & Naccache, 2001), but had conceptually similar predecessors like the working memory model of Baddeley and Hitch (1974) or Dennett’s multiple draft theory of consciousness (Dennett, 1993). GWT can be categorized as a functionalist theory of consciousness. It does not centrally aim at explaining qualia, but instead, focuses on the question of when and how information processing becomes conscious. The theory’s approach is an architecture of modular processing circuits in which information remains unconscious, and it then assigns the role of the filter to conscious processing to attentional mechanisms. Those are not thought of as homunculi determining which information is processed consciously. They are rather conceptualized as neuronal activity, modulated by behavioral context, goals, and rewards, that determines the processing mode in a form of race, or, as Dehaene and Naccache (2001) put it,

‘neuronal Darwinism’. Once a stream of information processes wins this race, under certain circumstances, consciousness then emerges because information is widely spread across the brain, especially in cortical areas and the prefrontal cortex.

This framework can explain a myriad of empirical findings and also subjective phenomena (see e.g., Baars, 2017). First, it can explain why certain information can never be conscious – such as our neuronal or vegetative nervous system activity. Those processing networks are not long-distance connected via long-range axon neurons (Dehaene, Kerszberg, & Changeux, 1998) and thus have no access to the global workspace to begin with. Other information can, in principle, become conscious, due to the connectivity of its processing modules to the global workspace. But there are certain constraints for information to actually become conscious. First, the theory postulates that attention is a prerequisite for consciousness. This means, top-down attention that “mobilizes” or “broadcasts” information to the global workspace. There is certainly evidence for this claim (Alef Ophir et al., 2020; Hommel et al., 2006), but also doubts about its generalization (Tallon-Baudry et al., 2018; Tsuchiya & Koch, 2009). Secondly, processing of that information needs to be maintained over a certain amount of time to reach activity that is sufficient for conscious processing. This is an interesting point for research that aims to experimentally manipulate consciousness by strongly limiting presentation and processing time for stimuli (Kiesel et al., 2008). But not only are there prerequisites for conscious awareness to occur, but also consequences from conscious awareness that the theory predicts. The central claim here is that high-level, novel, and semantic information processing that is sustained over a longer time period can only occur with conscious awareness. High-level semantic integration of novel information across modalities, space, and time thus is the main function of consciousness, according to GWT (Mudrik et al., 2014). And although there is evidence that, for instance, semantic priming effects can occur in the absence of awareness, these effects only last a few hundred milliseconds (e.g., Naccache & Dehaene, 2001). In contrast,

highly integrated or semantic information that needs information retrieval from memory, and is sustained over a longer period of time, seems to require processing in the global workspace.

Secondly, there is the class of HOT (Rosenthal, 2005). One could say that they take the Cartesian *cogito ergo sum*, “I am, I exist, is necessarily true whenever it is put forward by me or conceived in my mind” (Descartes, 1641/1985), that is a self-assurance *requiring* consciousness, and turn it around to essentially explain consciousness through such metacognitive representations. That is, consciousness is the representation of the world’s representation of an organism. First-order states can be understood analogously to the encapsulated processing modules in GWT (Esler et al., 2022). There are some variants of HOT (for an overview, see Rosenthal, 2004), but all are based on the principle of postulating conscious awareness of a representation or mental state (a perception, a thought) as a consequence of a higher-order representation. That differentiates HOT from other theories of consciousness that propose a first-order perspective on consciousness (H. Lau & Rosenthal, 2011), suggesting that unconscious and conscious experience alike are based on first-order representations. But what determines whether a mental state is represented in a higher-order thought and whether it becomes conscious? Importantly, not all higher-order thoughts are necessarily conscious, but become increasingly conscious with increasing strength of representation, so with increasing metacognitive knowledge. Other than in GWT, attention does not play a central role in HOT. In HOT, attentional mechanisms can enhance first-order processing, and introspection is thought of as inner attention shifting between higher-order thoughts (van Gulick, 2004). However, attention is neither a prerequisite nor a sufficient factor to produce conscious awareness of a mental state. One could even claim that attentional mechanisms are just not well-defined in the scope of HOT (Hardcastle, 2004). What also differentiates HOT from GWT and other theories of consciousness is that HOT do not necessarily assign a function to consciousness. Conscious processes in HOT are not thought to be inherently different from unconscious processes, other than being

represented in a higher-order thought (Rosenthal, 2008, but see Hardcastle, 2004). That way, they do not predict influences of conscious awareness on performance (e.g., H. Lau & Passingham, 2006), and, in contrast to GWT and other theories, do not predict that high-level cognition and integration requires consciousness.

The third theory that will be roughly sketched here, is the IIT (Tononi, 2004). The theory is constantly under development (the latest update is IIT 4.0; Albantakis et al., 2023), but all versions of the theory share core principles. The theory provides a mathematical approach to consciousness, formalizing the approach that consciousness emerges from high levels of information integration. The aspect of information integration is certainly something that it shares with other theories of consciousness, but it is the only theory equating integration and consciousness. The level of integration can even be quantified by, broadly speaking, the difference between information processing in the individual parts of a system, and the processing in the system as a whole (the ϕ parameter). The idea is, somewhat similar to Gestalt theory (Miyahara & Witkowski, 2019), through information integration, the whole is more than the sum of its parts. And that the system cannot be reduced to (the sum of) its parts, because of the causal relations between them. In the case of consciousness, that means that the interplay of neurons produces consciousness. This also means that the system cannot be reduced to its independently working modules. IIT can thus explain why certain processes can never be conscious in principle. One example are processes in the cerebellum, a structure that is characterized by independent, modular, and mostly feedforward neuronal networks (Tononi, 2008). For such a modularized structure, IIT computes a low ϕ parameter, thus, no consciousness. Critically, such reasoning also means that any system with a high ϕ can potentially be conscious, and that the structure does not necessarily have to be based on neurons (Sheldrake, 2021). This aspect has recently been harshly criticized for its tendency to embrace a form of panpsychism (Merker et al., 2021). Many parameters thus distinguish IIT from GWT and HOT. For example, attention has a central

role in GWT, but not in IIT. Attention is independent from the level of integration, and is thus independent from consciousness. Also, IIT does not postulate modularized processing in unconscious states, on the contrary, as it equates the level of information integration through processing with consciousness, it is not compatible with a modularized processing view. In contrast to HOT, IIT does not include a metacognitive level. Instead, consciousness emerges from highly integrated information processing itself, not from a higher-level representation of it.

There are also attempts to unify or integrate theories along their shared features, claiming that current theories do not necessarily contradict each other (Storm et al., 2024). There is however a consensus that more empirical work, and a variety of methods and measurements of consciousness are vital to enrich the theoretical perspectives on consciousness (Melloni et al., 2023; Storm et al., 2024).

Table 1

Predictions from three theories of consciousness concerning the three aspects of attention, modularity and semantic processing.

Theory	Role of Attention	Modular Unconscious Processing	Semantic Information Processing
Global Workspace Theory (GWT)	One mechanism that produces consciousness. Information can be attended to be selected for global broadcasting.	Modules work independently. They do not communicate directly. Information integration can occur in the absence of consciousness, but is short-lived. For longer-term, novel integration and association, processing in the global workspace is required.	Can to some extent occur unconsciously, but full conceptual understanding requires consciousness. Complex semantic integration and novel association acquisition requires global broadcasting.
Higher-Order Thought Theory (HOT)	Is not a prerequisite for consciousness. Can amplify signals, but attended states are not necessarily conscious (e.g., blindsight patients).	Processing occurs in specialized modular systems that operate independently. Modules can exchange information before it becomes conscious.	Can remain unconscious if no higher-order representation is present. First-order thoughts are not qualitatively different than, and can be as complex as higher-order thoughts.
Integrated Information Theory (IIT)	Attention and consciousness are separate processes. Attention can modulate information flow but does not produce consciousness, so that even unattended stimuli can be conscious if they contribute to an integrated information structure.	Highly modular systems with weak integration lack consciousness. Information transfer between modules is only possible if the system as a whole reaches a high level of system integration.	Is possible unconsciously, but depends on the level of system integration. A system with low system integration can still process meaning, but cannot be conscious of it. Highly abstract concepts require high integration, making them likely to be conscious.

Despite having discussed the reasons to work theory-driven above, there certainly are pitfalls to working within the framework of few prominent theories, like it is the case in the field of consciousness. The risk of dominant theories in a field is the confirmation bias (Schnepf & Groeben, 2024). This means that research efforts are mainly directed at confirming hypotheses deducted from theories, instead of challenging and potentially falsifying them. And then to stick to the same empirical methods that keep confirming those hypotheses, instead of diversifying methods, and thus, empirical evidence. For the research on consciousness, this is illustrated by a recent review and the extensive ConTraSt database (<https://contrastdb.tau.ac.il/>; Yaron et al., 2022) of 412 empirical studies addressing four prominent theories of consciousness (GNW, HOT, IIT, and RPT). In the database and with its interactive tools, studies on consciousness can be searched and summarized according to categories like underlying theory, paradigm, task, stimulus modality, or measure of consciousness. In their review, the authors show that when hypotheses deducted from consciousness theories are tested empirically, the results are predictable merely on the basis of the methodological choices. For instance, it could be predicted from an experiment examining global state consciousness or content-based consciousness whether it supported GNW or IIT (Yaron et al., 2022). Further, the authors found that rather than rigorously testing predictions, many studies interpret the results with regards to implications on consciousness theories only post-hoc, instead of practicing a priori hypothesis testing. Only 15% of the studies challenged a theory instead of just confirming its predictions. One implication for the field is then, that a theory is hardly ever challenged or even discarded. As an example, one can see that only 27 studies challenge GWT by finding posterior NCC, whereas 77 studies find frontal and 83 studies find parietal activity and thus support GWT. This way, evidence on all theories are just piling up (Yaron et al., 2022). The same tendency can be seen in the developments of the NCC research. GWT and HOT, which both predict neural activity associated with the global workspace and higher-order thinking in frontal regions, have not lost popularity even as the search for NCC has narrowed to posterior cortical and sensory regions

rather than fronto-parietal regions (Koch et al., 2016) and studies using active stimulation were not able to manipulate consciousness via the prefrontal cortex (Raccah et al., 2021). This state of research field has triggered different solution approaches, such as a vast adversarial collaborations project that aims to test GWT and IIT against each other (Melloni et al., 2023).

To conclude, there are viable candidates for a broad theory of consciousness that view consciousness from a functional perspective, and are empirically testable. However, as though there are crucial advantages for studying consciousness on the basis of theories, to let those advantages count, experimental psychology as a discipline needs to be vigorous in challenging those theories.

3 Unconscious Processing

The review of different theoretical frameworks of consciousness has shown approaches to describing conscious processing, the emergence, and functions of consciousness. How can these approaches be evaluated empirically? I will argue that it is a fruitful approach to examine the scope and characteristics of *unconscious* processing to gain insight into the concept of consciousness. First, we can empirically identify functions that can occur *without* conscious awareness. We would thus posit constraints on theories that attribute such functions exclusively to consciousness. For example, high-level cognition is a function that GWT and IIT would ascribe to consciousness. When there is empirical evidence of some high-level processing in the absence of awareness (e.g., Mudrik et al., 2014; Naccache & Dehaene, 2001), this challenges those theories, or at least expands the scope of unconscious processing and thus confines the function of consciousness. The objective of examining unconscious processes is then to determine its scope and limits, which potentially, by exclusion, leaves the function and *raison d'être* of consciousness. This approach stems from a functional, but also evolutionary perspective on the issue – as we have conscious awareness, we need to explain why it is useful, and why it had an evolutionary advantage and developed the way it did (Velmans, 2014). Complementary to that, we can find functions that require awareness to delineate the role of consciousness in cognitive function. As an example, evidence supports the notion that conscious processing is needed for highly integrated, long-term sustained, and novel semantic information processing (e.g., Biderman & Mudrik, 2018; Moors et al., 2016; Treisman, 2003). This would then support theories of consciousness that claim such processing only in the scope of conscious processing, for example in the GWT, or defining it as integrated information processing that produces consciousness, as in IIT. In contrast, HOT would not require conscious processing for complex, high-level cognition, because processes can remain unconscious as long as there is no higher-

order thought representing it. This way, the investigation of unconscious processes can provide a valuable contribution to an evaluation of consciousness theories.

There is a vigorous debate about how to investigate unconscious processing, and especially, the differentiation of conscious and unconscious processes. There are some psychologists, neuroscientists, and philosophers that reject virtually any current evidence for unconscious processing or unconscious perception (Newell & Shanks, 2023; Peters et al., 2017; but see responses, e.g., Dijksterhuis et al., 2014). Aside from the question whether unconscious processing has been demonstrated convincingly in empirical work, there is also a debate concerning the quality of differentiation between conscious and unconscious processing. There are multiple-system views that differentiate the two qualitatively and claim two separate processing systems, as is suggested in GWT (Baars, 1997; Dehaene & Naccache, 2001). This view was also adopted by models that aim to differentiate conscious and unconscious learning processes (e.g., Keele et al., 2003; Sun et al., 2005). On the other hand, there are views that reject a multiple-system architecture and argue for a single-system approach in which conscious and unconscious processing only differs in characteristics of its representation, like in HOT (H. Lau & Rosenthal, 2011) or IIT (Tononi, 2004). This was also adopted for models of learning, concerning the question of characteristics of conscious and unconscious learning (Cleeremans & Jiménez, 2002).

In line with that, many theoretical frameworks would not subscribe to a dichotomous notion of unconscious versus conscious processing. Nonetheless, in empirical testing, methods regularly contrast conscious and unconscious processes, mostly in a dichotomous fashion (Ramsøy & Overgaard, 2004). Yet, to illustrate the alternatives to a strict dichotomy, one can shortly summarize the following approaches. Based on the GWT, a tripartite structure has been proposed, namely conscious, preconscious, and subliminal processes (Dehaene & Changeux, 2011; Dehaene et al., 2006). Other frameworks proposed a continuous scale for consciousness,

for example a continuum of information integration in the IIT that is scaled with the parameter Φ (Tononi, 2004), or a continuum along strength and duration of recurrent processing that is indicative of the level of awareness (Lamme, 2006). Finally, there are approaches that differentiate between the physical-neuronal and phenomenological aspect, and, concerning the latter aspect, emphasize the importance of qualitative assessment of awareness, instead of quantitative (dual-aspect approach; Chalmers, 1997). In the work at hand, the dichotomy will be accepted as a working definition for an empirical approach to contrast two different modes of processing, but it is not subscribed to it in a strictly theoretical sense.

Independent from the precise differentiation between conscious and unconscious processing, there are several functional arguments for a differentiation. First, there is the argument of efficiency. Conscious processing is capacity-limited (Marois & Ivanoff, 2005) and potentially needs more energy than unconscious processing (Schölvinck et al., 2008). In contrast, in unconscious processing, a myriad of processes operate in parallel across distributed networks (Kihlstrom, 2014). Also, building on the capacity argument, conscious processing is handling already filtered information, a fraction of what is processed in the unconscious system, often described as selected by attentional processes (e.g., Baars, 1997; J. Prinz, 2011; Tsuchiya & Koch, 2009). That is functionally crucial to avoid an overburdening of the system, given that visual information alone entails about 108 bits per second (Itti & Koch, 2000), and that it is essential for the organism to be capable of acting quickly and appropriately to its environment or according to its goals. Unconscious processes ensure quick processing, filtering action-relevant information enables action control, especially in situations that are novel or conflicted (D. A. Norman & Shallice, 1986). Taken together, these considerations demonstrate that unconscious processing is not negligible but a fundamental counterpart or complement to conscious processes.

However one views the differentiation between conscious and unconscious processes, there are methodological challenges to empirically test them. These lie specifically in the attempt to disentangle conscious and unconscious processing, so, to study the latter without contamination of the former. That is essential when aiming to determine the scope and limits of unconscious processing. In the last decades, methods to measure unconscious processes have been newly introduced, proven themselves useful, or been criticized and discarded. The following section illustrates the challenges the field is facing with current methods to study unconscious processing. From this, I will derive the argument for using implicit learning instead of other paradigms within this goal. It is worth noting that recently, a team of researchers has been working on a best practice manual for studying unconscious processing (Stockart et al., 2024). However, in their work, the authors do not weigh different arguments for different methods, but instead present the results of two surveys, having asked researchers in the field for their opinions. The recommendations that are deducted from these two surveys are thus not taken as a gold standard for this current work, as they are not deducted from empirical data, logical argument, or theoretical deliberations, but from alleged authority.

This is why, in the following, I will depict some of the most common methods examining unconscious information processing, and will discuss their issues. I will show that methods of implicit learning are less commonly considered for testing consciousness theories, and are, to a certain extent, a separated field from classical approaches. Thus, I aim to form an argument for the present work's approach to studying consciousness with implicit learning methodology.

3.1 Experimental Paradigms of Subliminal Stimulus Presentation

First, and maybe most prominent, is the notion to study unconscious processes by subliminal perception paradigms. Subliminal perception means that a stimulus is presented but rendered invisible. Its perception is demonstrated through any kind of influence imaginable on behavioral (e.g., Strahan et al., 2002), but also in neuroimaging measures (e.g., Eimer &

Schlaghecken, 1998, 2003). Typically, this is achieved by presenting the stimulus for a very brief time, often combined with visual masking techniques, or through manipulation of attentional resources or selection (for an overview, see Kouider & Dehaene, 2007). For instance, presentation times of 8, 16, or 33ms are commonly used in combination with visual masking. It could be questioned whether stimuli presented so briefly are even processed at all. But there is evidence that even sub-millisecond stimuli evoke brain responses, at least when they are not masked (Sperdin et al., 2015). There are several problems that have been raised over the last decades. First, there are technical issues. The hardware that was used in such experiments until the 2000s, cathode ray tube monitors, were criticized to produce artifacts depending on luminance and contrast of presented stimuli (García-Pérez & Peli, 2001). More so, the modern liquid crystal display monitors are not able to accurately time stimulus presentation especially if short presentation times are required (Ghodrati et al., 2015). There are modern technology solutions, but still, for much of the hitherto existing evidence it is questionable whether stimuli were in fact presented as shortly as claimed, if that was not checked by an external photodiode or other advanced technical checks (García-Pérez & Peli, 2001). Secondly, there is also a conceptual issue with presenting stimuli for a very brief time. Presentation time first and foremost manipulates representational strength. That is not the same as, and does not necessarily determine, consciousness. Being unaware of a visual stimulus could be different in terms of representation and processing than having poor representational strength (Kouider & Dupoux, 2004).

Methods that potentially do not conflate consciousness with representational strength, are manipulations of attention. For example, a more recently developed method is the continuous flash suppression technique (CFS; Tsuchiya & Koch, 2005). It is a combination of binocular rivalry and flash suppression. To one eye, a low-contrast target image is presented, while the dominant eye is presented with Mondrian patterns that are altered, thus “flashed” every 100ms. The result is, as has been repeatedly demonstrated, that the target image remains not consciously

perceived for several seconds. The popularity of this method since its introduction can be seen in increasingly more publications per year using it (Stein, 2019). Yet, it has extensively been under critique for several issues, first and foremost that it remains unclear what level of processing is happening during CFS (Pournaghdali & Schwartz, 2020).

The methods described in this chapter face challenges, both concerning validity and methodological details. Even though there have certainly been technological advances in rendering stimuli (allegedly) invisible, the epistemological gain is not compelling. Still, as can be seen from reviews and meta-level work on consciousness research (Melloni et al., 2023; Stockart et al., 2024; Yaron et al., 2022), these methods are the most prominent ones, especially when testing prominent theories of consciousness. Yet, there is a whole other perspective to approach the study of unconscious processing. The idea is to not render the stimulus invisible or not perceptible, and thus the stimulus potentially remains unconscious. Instead, the stimuli involved are perfectly visible or perceptible, but that knowledge about their statistical occurrence is not explicitly instructed and remains unconscious.

3.2 Implicit Learning Paradigms

For the acquisition of such knowledge, Reber (1967) coined the term *implicit learning*. Implicit learning is commonly defined as learning that occurs without any intention to learn and without conscious awareness of the learning process or its contents (Perruchet & Pacton, 2006; Reber, 1967). For example, Reber (1967) found that an artificially created grammar was learned by participants just by observing examples of grammatically correct letter strings. Participants were never told the rules of the grammar explicitly. That they learned them nevertheless was demonstrated by above chance performance when participants were asked to discriminate grammatical and non-grammatical letter strings. In an earlier experiment, Braine (1963) found that children (with literacy skills) could even produce grammatically correct material after a perceptual learning phase, but without ever having been told the rules. In these so-called

artificial grammar learning (AGL) paradigms, the grammar complexity and number of stimuli within the grammar is variable (Perruchet & Vinter, 1998; Schiff & Katan, 2014). AGL is a widely used paradigm to study implicit rule learning (for an overview, see Pothos, 2007). Yet, there is a lively debate about the question of the form of AGL knowledge representation of knowledge, and doubts about its implicit nature (Perruchet & Pacton, 2006; Pothos, 2007).

A second, closely related paradigm examining implicit learning, is the serial reaction time task (SRTT; Nissen & Bullemer, 1987). In its classical set-up, participants observe stimuli appearing at a fix number of positions on the screen, and are asked to respond to their appearance with response keys that are spatially mapped to the screen locations. What participants are not explicitly told is that the stimulus appearances follow a (deterministic or probabilistic) sequence. Nevertheless, when comparing blocks of sequentially presented stimuli, and randomly presented stimuli, response times are shorter in the former than in the latter blocks (for an overview, see Schwarb & Schumacher, 2012). This makes it an interesting extension of AGL because there is an online measure of learning in SRTT, which has proven to be a reliable measure of learning effects, at least on the group level (Oliveira et al., 2023). In the SRTT literature, there have also been extensive efforts to test for explicit sequence knowledge (e.g., Curran, 2001; for an overview see Schwarb & Schumacher, 2012). There are also further extensions of the SRTT, beyond its scope of examining spatial learning: It could be shown that not only motor sequences (key responses) could be learned, but also perceptual sequences (Eberhardt et al., 2017; Haider et al., 2014; Wilts & Haider, 2023).

Thirdly, there is the contextual cueing (CC) paradigm (Chun & Jiang, 1998). It is a visual search paradigm in which participants see displays with multiple distractor letters “L” and a target letter “T”. Unbeknownst to the participants, some letter displays are repeatedly presented throughout the experiment, so that a distractor configuration repeatedly appears with the same target location in multiple trials. This way, in such repeated trials, the target location would

be predictable on the basis of the spatial distractor configuration. What classical CC experiments have shown is that participants indeed learn these contingencies between distractor configuration and target location, as they show reduced response times for repeated, in contrast to novel displays. However, typically, this acquired knowledge remains implicit (Chun & Jiang, 1998; but see Vadillo et al., 2019).

It is worth noting that the term *implicit* is widely used in psychological research. For this work, it should not be confounded with the term as it is used in social cognition research, where it has extensively been shown that so-called implicit attitudes are not inaccessible to consciousness (Gawronski et al., 2006; Goedderz et al., 2024). In our use of the term implicit learning, it is implied that neither an explicit instruction regarding the learning content or process is given, nor that the learning contents are accessible to conscious awareness, and enable verbal report (Nissen & Bullemer, 1987). There are also several related terms that might seem synonymous to implicit learning, but need demarcation.

First, the term *incidental learning* is sometimes used in the context of implicit learning (J. R. Schmidt & De Houwer, 2019). Importantly, it does not signify the occurrence of implicit learning, but instead describes the task set-up as lacking explicit instruction to learn. Any learning that occurs incidentally, rather than intentionally, can be implicit, but it can also become explicit as the task progresses. Second, the notion of *statistical learning* is related to but different from implicit learning (Turk-Browne et al., 2005). Statistical learning will often be incidental and implicit, meaning that no instruction to learn the statistical regularities is provided, and regularities cannot be explicitly verbalized, but it is not necessarily so. Statistical learning can also occur intentionally and explicitly (Ren et al., 2024). Also, it is confined to the learning of statistical regularities, rather than, for example, implicit learning of social cues. The third term that might seem related is that of *procedural learning*. This is a common example of everyday relevance of implicit learning, like learning how to ride a bike or play the piano. But it is

confined to motor and sensory learning, while implicit learning also covers rule, sequence, and contingency learning. Lastly, *automatized processes* might seem to resemble the results of implicit learning in terms of effortless, fast task performance in the absence of conscious awareness (regardless of the criticism that this dichotomic conceptualization received; W. Schneider & Shiffrin, 1977; Shiffrin & Schneider, 1977; Willingham, 1998). Yet, the essential difference is that automatization most commonly develops from explicitly learned performance, that then becomes more automatized. Yet, it is not implicit in the sense that it does not involve or presuppose conscious knowledge (Shiffrin & Schneider, 1977).

Taken together, implicit learning is one approach to studying unconscious processing and then potentially deriving implications for theories of consciousness. One can deduce hypotheses about the conditions or limits of implicit learning from theories of consciousness. However, these deductions can be somewhat vague, given that the theories are rather broad and delineate the big picture of consciousness, not necessarily explicitly incorporating many of the parameters that determine and influence implicit learning as a phenomenon. Besides the more commonly used paradigms of subliminal stimulus presentation, paradigms of implicit learning constitute a fruitful approach to examine unconscious processes, and contrast them with conscious processes.

4 Measures of Conscious Awareness

Whichever paradigm we choose to study unconscious processing, we need a test for conscious awareness of stimuli or contents regardless. No paradigm so far can guarantee unconscious processing without the contamination by conscious processing. Now, it may be intuitive to just directly ask participants whether they have consciously perceived a stimulus, whether they noticed statistical regularities, or whether they otherwise consciously perceived what was supposed to be kept from conscious processing. However, already in the early research on consciousness, it was argued that introspection, that is, asking participants to verbally report their perceptions, may not be a valid measure of consciousness (C. W. Eriksen, 1960). The argument is that the verbalization of perceptual experience, its transfer into language, that is, abstract symbols, is probably going to be inadequate (C. W. Eriksen, 1960). The shortcomings of equating conscious awareness with verbal reportability can be illustrated in paradigms in which recall and recognition tests come to different results. For example, if asked what they perceived when presented with a masked stimulus, participants might be unable to freely report what they saw (recall), but might be able to choose from options and be above chance level (recognition; Micher et al., 2024). In such cases, it is difficult to conclude whether the stimulus has been processed consciously or not. However, one could also argue that in such cases, the recognition task performance is (partly) driven by implicit knowledge (Rünger & Frensch, 2010). For example, that participants were not explicitly aware of the stimulus, but its processing in the recognition task feels more fluent or there is a vague feeling of familiarity. These sentiments would then be driven by implicit processing, not by explicit.

Thus, obtaining a measure of conscious awareness is challenging, and it requires methodological rigor and effort. Many authors have concluded that obtaining multiple measures and combine or compare them might be advantageous, such as obtaining direct and indirect

measures of consciousness (Erdelyi, 1986; Holender, 1986; Reingold & Merikle, 1988). I will depict what those are in the following chapter.

4.1 Direct and Indirect, Objective and Subjective Measures

Thus, beyond simply asking for verbal reports, there are measures that have been proposed to test for conscious awareness. They are usually categorized into direct and indirect measures, and into objective and subjective measures. Table 2 provides an overview with examples of different measures in the scope of this categorization.

Direct, subjective measures are, as depicted above, reports requiring introspection. This can be done adopting a dichotomous view on consciousness, and subsequently, the response options would then be “seen” or “not seen” (with subliminal presentation paradigms), or “known” or “not known” (with implicit learning paradigms). For example, there is the perceptual awareness scale (PAS; Ramsøy & Overgaard, 2004). Such scales add labels that quantify perceptual awareness, but also qualitatively assess the experience in terms of “clearness” of the perception (e.g., the original PAS the labels are: 1=no experience, 2=brief glimpse, 3=almost clear experience, and 4=clear experience; Ramsøy & Overgaard, 2004). However, the issue with such measures is that participants could respond to any aspect of the stimulus (Michel, 2023a). For example, if they perceived the color, but not the shape of a briefly presented stimulus, they would respond that they perceived a ‘brief glimpse’ or ‘almost clear experience’. This will be independent from the task-relevance of that aspect of the stimulus. If the objective task is concerning shape, PAS scale responses ‘brief glimpse’ do not allow the conclusion that participants briefly perceived the shape, as it could also be that they briefly perceived the color of the stimulus. Thus, the measure would not necessarily inform the issue of whether the task-relevant aspect of the stimulus was consciously perceived or not, but only, if anything about the stimulus was perceived. This is why it is advised to adapt the PAS labels to the stimulus or feature of interest (Sandberg & Overgaard, 2015).

Direct, objective measures on the other hand are what C. W. Eriksen (1960) proposed alternatively, given his critique of introspection as exclusive measure of awareness. Objective measures are commonly used in consciousness research still. Usually, they are two-alternative forced choice questions that concern the stimulus as a whole or specific features of the stimulus. For example, when square and diamond stimuli are used as masked primes in a classical priming task, the objective, direct measure would be participants' choice between squares and diamonds. This response can then be correct or incorrect, and one can compute whether a participant performed at chance level or above chance level. Or, following Signal Detection Theory (Green & Swets, 1966), participants demonstrate sensitivity or null sensitivity regarding the shown stimulus. In implicit learning paradigms, direct, objective measures would be recognition (e.g., Chun & Jiang, 1998) or generation tasks (e.g., Chun & Jiang, 2003). In recognition tasks, participants are asked whether they recognize a stimulus, sequence, or predictability. In generation tasks, they are asked to reproduce a sequence or regularity. For example, in CC, participants learn a contingency between contextual cues and the location of a target. Originally, and still in recent studies, participants are later only asked about recognition of the contextual cue (Bergmann & Schubö, 2021; Chun & Jiang, 1998). However, they are not asked about the learned contingencies. This is problematic because the recognition measure does not directly assess the specific knowledge of interest. Namely, awareness of the learned contingencies. This is why, for the studies included in this dissertation, we opted for the generation task instead of the recognition task. With a generation task, the reproduction of the regularity, in this case, the contingency between cue and target location, depends on the knowledge about this contingency. And not, like the cue recognition task, simply on a memory representation of the cue.

Then, there are indirect measures of conscious awareness. In current empirical research, those entail mostly objective measures, such as measures on the behavioral or neural level. For example, this could be a classical priming task. The measure would be a behavioral response to

a target stimulus response that was preceded by a masked prime stimulus. This way, the influence of the masked prime on the processing of the target stimulus can be determined. Such behavioral measures are commonly accuracy and response times, comparing trials in which prime and target are congruent and incongruent. Accordingly, in implicit learning, an indirect, objective measure would be response time or accuracy to stimuli that follow the learned rule (e.g., a sequence or grammar), contrasted with responses to stimuli that deviate from the rule.

Table 2.

Measures of Conscious Awareness Categorized by Direct/Indirect and Objective/Subjective Measures.

	Direct	Indirect
Objective	Two-alternative forced choice question on stimulus or on features of a stimulus; Recognition and generation tasks	Behavioral measures (response accuracy, response times); Neural measures (e.g., event-related potentials)
Subjective	Verbal report of perception of a stimulus; Visibility ratings (e.g., the Perceptual Awareness Scale, Ramsøy & Overgaard, 2004)	Indirect memory tests (e.g., familiarity, fluency, preference, liking)

There are several advantages and disadvantage of each type of measure of conscious awareness (Stockart et al., 2024). Notable limitations of direct measures are the following. First, it has been noted that it is essential that the objective task concerning the target stimulus should concern the feature of interest, and not another, potentially harder or easier aspect of the stimulus (T. Schmidt & Biafora, 2024). For example, if the congruity between prime and target lies

in the shape of the stimulus, the objective task should ask about that, and not about other features like color, identity or semantic meaning.

One issue with direct, objective measures is that it is unclear whether above chance level performance should be considered a result of unconscious processes (H. Lau & Passingham, 2006), or if above chance level performance is indicative of the involvement of conscious processes (Michel, 2023). In an attempt to settle this debate, Micher et al. (2024) found that the direct, objective task (forced-choice discrimination) was contaminated by unconscious processes, resulting in better performance. So, although participants reported not to have seen the stimulus in the PAS scale, they performed better than chance level in the forced-choice task. Although they felt like they were guessing, their responses were influenced by unconscious processing of the stimulus. This is why a direct, objective task should not be taken as a measure of conscious awareness on its own. Another limitation for direct, objective measures might arise from experimenters piloting their studies, and weakening stimulus presentation (e.g., reducing presentation time or contrast) until participants show below chance performance, but no processing of the stimulus is possible anymore, consciously or unconsciously (Michel, 2023b). This would result in an underestimation of unconscious processing, applying a criterion that is too strict and conservative to find unconscious processing. Lastly, when direct, objective tasks are interleaved with the actual task, this might lead to overspilling effects (Lin & Murray, 2014). So, for example, if participants are asked about features of the masked prime in each trial with response alternatives presented to them, they could potentially shift their attention to the inquired feature of the prime stimulus. This could then increase their ability to discriminate that feature in the masked prime.

Concerning indirect measures, it is noteworthy that the results of those often diverge from the results from the direct measure. That means, on an indirect measure, there are differences in response times or accuracy measures between, for example, congruent and incongruent

trials in a priming paradigm. But on the direct measure, for example, a two-forced choice task, participants do not perform above chance level. This has been commonly explained with a higher sensitivity of the indirect measure. However, Zerweck et al. (2021) found that direct and indirect measures are equally influenced by unconscious processing, and do not differ in sensitivity in a number priming paradigm. The divergent results from those two measures might thus stem from different factors. For example, there are two central caveats that need to be considered in obtaining indirect and direct measures.

First, Shanks and St. John (1994a) argue that the design of the awareness test should be constructed such that the retrieval context is kept similar to the former tasks' context. This is especially relevant for tasks in which participants learn contingencies without being instructed to do so. This learning is demonstrated through indirect measures such as response times and accuracy during the learning phase. Then, an awareness test determines whether the contingencies are explicitly learned or remain implicit. Keeping the task during the response time measure as similar to the task with the awareness measure would increase sensitivity of the awareness test. Because participants would be able to retrieve their contingency knowledge, should they have any, from the retrieval cues that the task context provides. The direct awareness test should thus be maximally similar in trial structure, and only differ in task instructions. In contrast, constructing a direct, rather abstract, verbal report task asking about any conscious knowledge, might decrease sensitivity of the awareness test, as participants are less likely to be able to retrieve the knowledge they might have by the lack of retrieval cues. Following this line of argument, in the experiments of this dissertation, the awareness tests are constructed such that they are more similar to the performance task.

The second caveat concerning awareness measures following performance measures is also relevant for sensitivity, but especially also for reliability of the awareness test. The issue with hitherto findings is that many studies were designed such that there was a large number of

trials for the indirect measure (hundreds of trials measuring response times), but only a small number for the direct measure. This might lead to an issue of sensitivity and reliability of the direct measure (Meyen et al., 2024; Vadillo et al., 2016). Especially in tasks with subliminal stimulus presentation, such a design can result in considerable noise in direct tests, given random attentional and cognitive parameters influencing to what extent a stimulus was perceived consciously. In other paradigms, this might be less of a problem, for example, when the direct test is a recognition task that is administered with optimal stimulus presentation, no time constraints and limited response options. Still, also in such cases, the number of trials remains an important factor for increased sensitivity and reliability (Meyen et al., 2024), and this is also considered in the experimental set-ups presented in this dissertation.

4.2 Confidence Measures

An additional measure on top of direct and indirect measures that can be taken into consideration when measuring conscious awareness, are confidence measures. This is an approach that relates to metacognitive aspects of consciousness, as depicted especially in HOT. After a direct, objective awareness test (e.g., recognition test), participants can be asked what their subjective conscious perception was (e.g., with the PAS; Ramsøy & Overgaard, 2004), or how confident they are regarding their response to the direct, objective awareness test (Michel, 2023a).

As we have seen from recent evidence, direct, objective awareness tests can also be informed by unconscious processing (Micher et al., 2024). Therefore, we should not conclude conscious perception or knowledge from above chance performance in such tasks. This is why a combined measure of a direct, objective task (i.e., correct or incorrect response) and a confidence measure (i.e., low or high confidence) has been proposed (Michel, 2023a). An additional advantage in applying confidence measures to inform a conscious awareness measure, is that they can be compared between studies, tasks, even modalities, and that they hold the potential

to learn more about metacognitive aspects in conscious and unconscious experience (Michel, 2023a).

There are variants of obtaining a confidence measure. There are approaches implementing a post-decision wagering task (Haider et al., 2011; Persaud et al., 2007). Typically, after participants have responded to the objective, direct task (e.g., two alternative forced-choice, recognition, or generation tasks), they are then asked to wager either a small or larger amount of money or points on their response, depending on the confidence they have in their response. When the high wager is placed with a correct response, they win the amount of money or points, but when they place it with an incorrect response, they lose this amount. The rationale of the wagering task is that participants would potentially be motivated to report their confidence correctly and even with little confidence wager the high amount. However, personality traits such as risk or loss aversion can influence the measure (Dienes & Seth, 2010; Fleming & Dolan, 2010). For example, participants with higher loss aversion might rarely wager higher amounts of money or points, not because they are less confident, but because they want to avoid losing higher amounts of money or points, should their response have been incorrect. Thus, Dienes and Seth (2010) empirically tested verbal confidence reports (binary judgement, ‘guess’ and ‘sure’) against wagering tasks, and found that the latter are not more sensitive to conscious knowledge, and are furthermore correlated with measures of risk aversion. Verbal confidence reports in contrast did not correlate with risk aversion measures. Thus, obtaining confidence in the form of wagering task could potentially underestimate confidence and thus, conscious knowledge (see also Konstantinidis & Shanks, 2014).

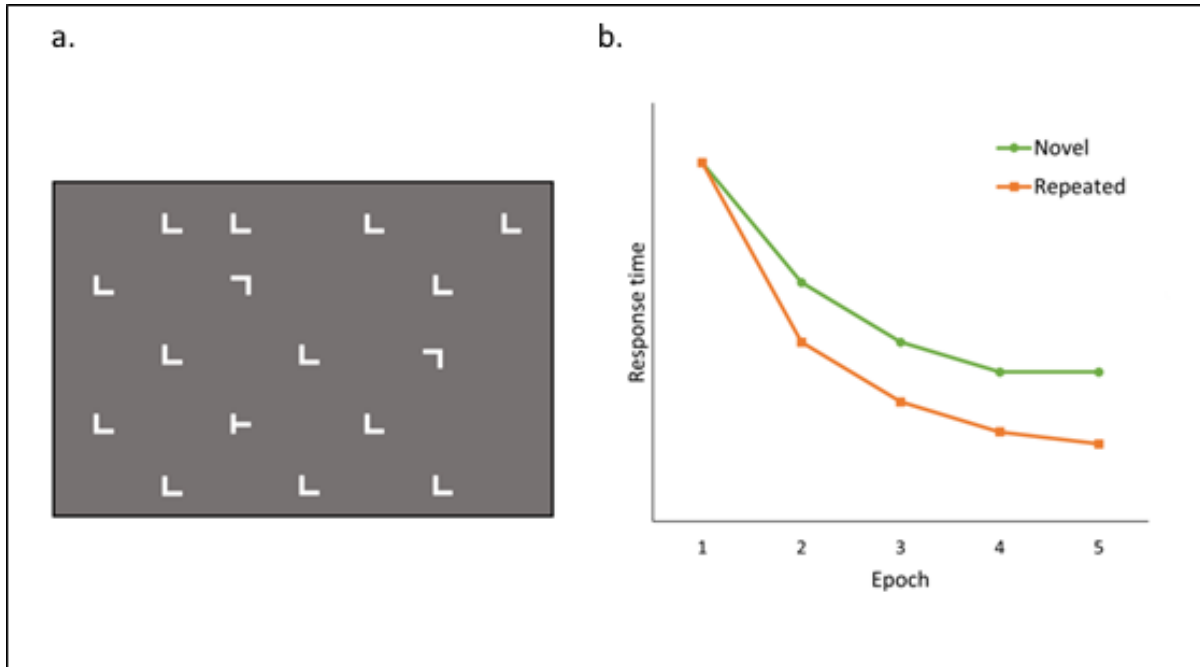
Following the debates on measures of consciousness, in the experiments of this thesis, multiple measures are used. Participants are asked to respond to a direct, objective task (generation task), which is, as discussed, preferable to a recognition task. Because the influence of implicit knowledge on this measure cannot be excluded (Micher et al., 2024), it is followed up

by a confidence measure. To avoid potential influences of inter-individual differences on the confidence measure, as has been discussed with wagering tasks, I opted for a relatively simple four-point Likert scale for confidence report (Konstantinidis & Shanks, 2014). With regard to research on the PAS, labelling the scale proved to be helpful to reduce ambiguity (Ramsøy & Overgaard, 2004; Sandberg & Overgaard, 2015). Therefore, we also labeled the Likert scale, with 1 = “complete guess” and 4 = “absolutely certain”. Concerning the analysis, we chose to analyze the objective and confidence measures following the rationale of wagering task analyses, combining the two measures (Persaud & McLeod, 2008; Persaud et al., 2007). Specifically, we compared the relative frequency of participants reporting high confidence, given their response was correct, with the relative frequency of participants reporting high confidence, given their response was incorrect (cf. Dienes & Seth, 2010). We also always conducted the analysis with the reverse base rates, so comparing correct responses given high confidence judgements with correct responses given low confidence judgements.

5 The Contextual Cueing Paradigm

As noted above, CC is a widely used and much studied paradigm to investigate implicit learning of contingencies that influence spatial attention (for an overview, see Jiang & Sisk, 2019). As a reminder, the typical CC paradigm is a visual search task, in which participants have to find a target letter among distractors (see Figure 2a). In this respect, it is similar to other visual search tasks (Wolfe, 2020). However, in classical visual search experiments, search displays are intentionally randomly generated to prevent any potential learning effects that would guide the search (e.g., Treisman & Gelade, 1980). In contrast, in the CC paradigm, visual search is used as a framework to investigate learning and its effect on attentional guidance. Therefore, some search displays are repeated throughout the learning phase.

The common finding is that there is a steeper response time decrease for repeated displays, in contrast to novel ones (see Figure 1b). This means that the contextual cues (spatial configurations of distractors) can be learned to predict target location and guide attention. As an effect, response times decrease in trials in which the target location is predictable. After learning, in a test phase, participants perform a recognition task in which they are asked to detect the repeated displays among novel ones (e.g., Bergmann & Schubö, 2021; Chun & Jiang, 1998). When participants are performing at chance level in the test phase, that has commonly been interpreted as evidence for implicit learning (Jiang & Sisk, 2019). As noted in the chapter above, the recognition task has been criticized for not targeting the contingency knowledge but only cue memory. Thus, alternatively, participants can be asked to do a generation task in which they are asked to indicate the target location, given a repeated or novel display (e.g., Chun & Jiang, 2003). Also in the generation task, it has commonly been found that performance is not significantly better than chance (e.g., Chua et al., 2003; Chun & Jiang, 2003).

Figure 2*Search Task and Performance in the Original Contextual Cueing Paradigm**(Adapted from Jiang & Sisk, 2019)*

Note. (a) A typical contextual cueing search display with “L” distractor letters and a “T” target letter. (b) A typical results graph from a contextual cueing experiment that shows decreasing response times over epochs for novel and repeated displays separately.

In the work at hand, I use the CC paradigm, but extend it to a variant in which it is more flexible with regard to stimuli, testable hypotheses, and condition comparisons. To explain the alterations that I undertook, I will first summarize what we know about mechanisms in CC learning, take up aspects that have been criticized about traditional CC, and then explain the new variant and highlight its advantages.

5.1 Learning Mechanisms

Although decades of research on the CC effect have passed, there is still no final answer to the question of what exactly produces the effect. The debate includes the question of whether CC is indeed an effect of attentional guidance. And if so, how exactly this guidance in the search

process is learned and executed. Further, the question is whether response-related processes play a role, and if so, how big this role is.

First, a central question is at which stage of processing the CC effect emerges (Sisk et al., 2019). The attentional guidance account posits that the effect arises at an early stage. In this account, a learned contingency between distractor configuration (context) and target location causes attentional guidance, thereby enhancing the efficiency of visual search. This would constitute an early locus of the effect, implying that CC directly modulates the search process. Empirical support for this view comes from studies demonstrating influence on early perceptual processing, indicated by event-related potentials (ERP) measures (Johnson et al., 2007), as well as influence on early eye-movements in eye-tracking procedures (Harris & Remington, 2017; Peterson & Kramer, 2001).

An alternative or complementary hypothesis, the response facilitation account, suggests that CC may instead, or additionally, be driven by response-related processes (Kunar et al., 2007). Here, the repeated contexts foster familiarity, which in turn facilitates response selection and motor execution, independent of the preceding search process. This would indicate a late locus, suggesting that the effect emerges after the target has already been found.

While some evidence supports a combined influence on both early and late processes (Schankin & Schubö, 2009a; Sewell et al., 2018), other findings point to a potential additional mechanism between post-attentional and pre-response stages (Schankin & Schubö, 2010). Also, more recently, interest in potential perceptual effects in CC have been debated (Sewell et al., 2018; Zhao & Ren, 2020).

The now prevailing perspective on the mechanism is an attentional guidance effect that is based on a rough attention allocation. This is supported by evidence from lateralized ERPs (e.g., the N2pc; Eimer, 1996) that indicated covert attentional shifts that were more efficient in

repeated search displays (Schankin & Schubö, 2009a). Also, using eye-movement measures, fewer saccades and fixations were observed in repeated configurations (Beesley et al., 2018; Goujon et al., 2015; Tseng & Li, 2004), suggesting a more efficiently guided search. At the same time, it seems that this guidance does not result in precise predictions and subsequently, precise eye movements to the target, but rather that the location of targets are only roughly estimated from the predictive context, but precise enough to benefit search and response times (Peterson & Kramer, 2001; Schankin & Schubö, 2009b). This raises the question whether these mechanisms are automatic or controlled. Beesley et al. (2018) argued that in CC, participants do not strategically search through the distractors, or process predictive distractors more than unpredictable ones. In that way, it is not a strategical and voluntary process. In line with that, it has been shown that an attentional focus on the predictive material (e.g., the predictive distractors) is not a necessary condition for the CC effect to occur (Conci & Mühlenen, 2011; Harris & Remington, 2017). Still, the process seems somewhat controllable as participants could suppress a once learned contingency when it was no longer advantageous (Luque et al., 2017; Manginelli & Pollmann, 2009).

5.2 Limitations

As much as the CC paradigm has contributed to the understanding of processes of attentional guidance and implicit learning, there are several limitations that I want to address, and draw conclusion for the methodology applied in the work at hand.

First, I would argue that the CC paradigm is limited in its scope of generalization. This limitation does not lie in an inherent lack of flexibility of the paradigm, but in the persistent use of a specific task set-up. There have been slight variations, for example, that the configuration of distractor stimuli in some studies follows a concentric circle set-up instead of a configuration in a grid structure. But over the first decades of CC research (Jiang & Sisk, 2019), and in recent research (Meyen et al., 2024), the set-up has always been roughly the same: The search displays

are comprised of letters on gray backgrounds, the contingencies are between distractor configuration and target location. The target is almost always the same, which is, randomly rotated “T”-shaped target letters, and participants’ task is then to report the orientation of the target letter. As we argue in Study 1 (Tavera & Haider, 2025), this set-up inherently emphasizes the spatial dimension above other dimensions (e.g., visual dimensions such as color and shape). The spatial configuration predicts a location in space where the target is to be found, and then the target is judged in terms of its orientation to the left or right, which again corresponds to the spatial dimension. The findings on the grounds of this paradigm are of course highly relevant for different processes factoring into the CC effect. However, they might be, given the ever-same task set-up, limited. And the limitation could be that findings from CC are only applicable to learning and attentional effects within one dimension (i.e., the spatial dimension), but not across dimensions (e.g., between color, shape, and space). Or, it could be even more limited in that the conclusions on cognitive processes drawn from the paradigm only apply to the spatial dimension in the first place. Whether the use of other, non-spatial visual cues, across-dimension contingencies, or variations in the response in the paradigm would still result in CC effects, remains unclear. Also, the spatial dimension cannot be viewed as just any visual feature in the cognitive system. We have argued (Tavera & Haider, 2025) that the spatial feature might have a special position and special processing pipelines in the cognitive system due to its inherent ties with the motor system (Paillard, 1991), but also, because its dominant role in learning has been shown empirically (I. Koch & Hoffmann, 2000; Kunar et al., 2013).

The second limitation of the CC paradigm is the debate over the implicit nature of the learning effect. In the chapter on awareness measures, I have already discussed the issue of their sensitivity and reliability. This issue does not remain a theoretical one, but is, given the hitherto studies of CC, an empirical one (Luque et al., 2017; Meyen et al., 2024). Because typically in CC studies, conscious awareness of the contingency between spatial configuration and target

location is tested with a recognition task. In such a task, participants are shown old and novel search displays, and asked to categorize them into old and novel (Bergmann & Schubö, 2021; Chun & Jiang, 1998). The important point here is, that it is different from awareness tests with subliminally presented stimuli, where one can, in principle, conduct many test trials as is determined necessary for sufficient sensitivity and reliability. In the case of recognition trials after CC, it is only possible to use a one-trial test for each search display, because repeating an old display in the test would make participants categorize it as old just by the familiarity from the test phase. Also, the validity of the measure is questionable. The recognition measure requires recognition of the configuration alone, but what really is learned in the learning phase is a contingency between configuration and target location. This violates the principle of having to test for what is actually learned in the task (Shanks & St. John, 1994). There have been attempts to tackle this problem by implementing not a recognition, but a generation task, modelled after procedures in SRTT experiments (Willingham et al., 1989). Here, participants are not asked whether the display is old or new, but where in this display they would predict the target to be (Chun & Jiang, 2003). But many studies using the CC paradigm, also recent ones, have not adopted this approach (e.g., Bergmann & Schubö, 2021; Bergmann et al., 2020). Also, many recent studies, have not administered any kind of recognition or knowledge test when the type of learning or knowledge is not their main research interest (Kobayashi & Ogawa, 2020; Kunar et al., 2006). Regardless of that, the generation task has been criticized, claiming that the way it is implemented lacks statistical power to find evidence of explicit knowledge (Smyth & Shanks, 2008). As discussed in the chapter on confidence measures, in the experimental procedures of the work at hand, we combine generation tasks with confidence measures to increase information density and implement a larger number of trials of the generation task to reach a certain statistical power and reliability.

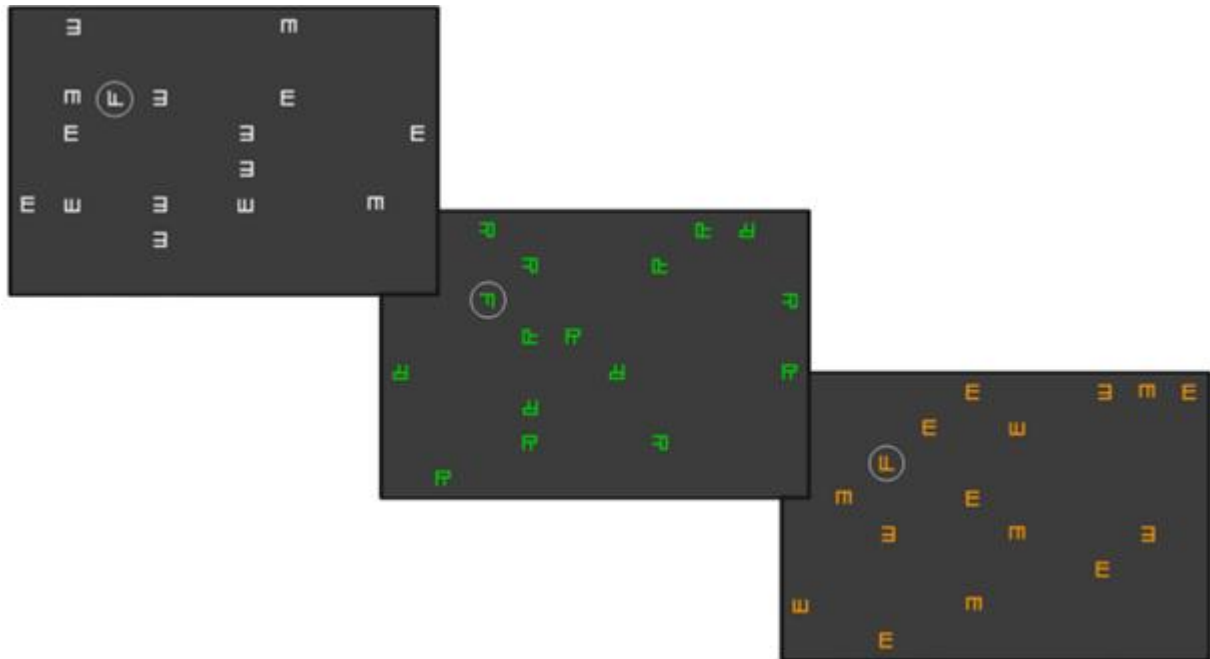
Given the methodological problem that the classical recognition tests with very few trials posit with regard to within-subject reliability and sensitivity, a further empirical issue with studies on CC are the rather small sample sizes (Vadillo et al., 2016). This again posits a power, sensitivity, and reliability problem regarding the overall CC effect and recognition measures that are computed across participants. Therefore, the argument of Vadillo et al. (2016) is that CC effects are contingent on explicit knowledge, despite of what the literature suggests. The authors argue that because of the flawed methodology of the recognition tests and underpowered samples, studies have just not detected the explicit nature of the acquired knowledge. However, there have been attempts to empirically investigate whether these empirical limitations in the CC literature are necessarily indicating an inevitable constraint of the paradigm, or are in fact just an empirical flaw. Colagiuri and Livesey (2016) argue that if the CC effect was dependent on explicit knowledge, one should find a positive correlation of explicit knowledge with the CC effect. However, in three experiments with samples of up to 600 participants, they did not find such a positive relationship. First, they tested the relationship on participant level, meaning that participants who recognized more repeated displays showed an increased CC effect. Second, they tested on the level of individual search displays, in the sense that the learning of only some configurations could drive the CC effect. This should reveal if a potential positive correlation could be concealed by the aggregation across all configurations. There was no compelling evidence for a positive correlation of knowledge and CC effect in either analysis. Some data rather suggested a negative relationship, such that more conscious knowledge might even attenuate a CC effect. Colagiuri and Livesey (2016) agree with Vadillo et al. (2016) that statistical power is considerably lacking in CC literature, and that participants demonstrate above chance level performance in the two-forced-choice recognition task (old-new) that is commonly used in the literature. Acknowledging these two conclusions from empirical data and simulations is however not equal to rejecting the notion of implicit learning in CC.

5.3 A New Variant of the Contextual Cueing Paradigm

To address the shortcomings of the original CC paradigm, as outlined above, we constructed a new variant. Our goal was to create a more versatile paradigm that can be used for more research questions, going beyond learning in the spatial dimension only, and increasing power and reliability to test for explicit knowledge of the learned contingencies.

Figure 3

Exemplary Search Displays in the Novel Variant of the Contextual Cueing Paradigm
(Adapted from Tavera & Haider, 2025)



Note. In the novel variant of the contextual cueing paradigm, search displays are comprised of stylized letter distractors in varying colors, and a target letter. This way, shape and color of the distractors can be manipulated, and used as predictors for target location respectively. Participants' task is to report whether the target letter has a short or long middle bar. Note that the circle around the target letter is for illustration only and was not shown in the experiment.

There are marginal differences in the design of stimulus material for CC, for instance, concerning display sizes, number of distractors and their distribution, as well as regarding number of possible target locations (Jiang & Sisk, 2019). For our material, we designed search displays on the basis of the material of Bergmann et al. (2019). As described in Tavera and Haider (2025), the search displays were built with letter arrays, as in the original variant of CC (see Figure 3). However, in our variant, not the spatial configuration of distractors was the predictive cue for target location. Instead, we rendered visual features of the distractors predictive. Therefore, the distractor and target shapes were varied (for similar material variations in CC, see Beesley & Shanks, 2012). Instead of always searching for a “T” shaped target letter in “L” shaped distractors, the distractors were now presented as stylized letters A, E, K, P, S, and W. The distractors could also be differently colored. The target was a stylized letter “F”, with a short or long second horizontal bar (see Figure 3). With these alterations, we attained two goals: First, the variation of distractor letters introduces the possibility to use letter shape and color as predictive cues for target location. Secondly, the target letter was designed such that participants’ task was no longer a spatial orientation task, but a judgement of target shape. Both these alterations reduced the emphasis on the spatial dimension of the task design.

In the new variant, contingencies can thus be manipulated more flexibly than in the original CC paradigm. One distractor shape, or multiple shapes, can predict the target location, either deterministically or probabilistically. Importantly, the shapes are not restricted to stylized letters, but could be geometrical shapes or even objects. The same way, colors can be predictive of target location, or color categories, or specific combinations of colors, or a combination of target shape and color serves as predictive cue. In principle, the spatial configuration could also be predictive, either additionally or independently from distractor shapes and colors. Also, unlike the original paradigm, the cues can all be predictive (e.g., every color predicts a different target location with some probability). But also, like in the original paradigm, only some cues

can be predictive (e.g., some colors predict a target location, while other colors are unpredic-
tive). As illustrated by these adaptations that can be implemented within the novel variant of
the CC paradigm, it makes a myriad of hypotheses testable, and allows for broader generaliza-
tion than the original CC paradigm.

6 Attention and Consciousness

In our first study, we aimed to examine the role of attention in implicit learning, using our new variant of the CC paradigm. As discussed in the chapter on theories of consciousness, there are different perspectives on the relationship between attention and consciousness. Thus, one can deduct different predictions for the possibility of implicit learning in the absence of attention.

The distinction between attention and consciousness is rather complex, and still topic of debate (for an overview, see Tsuchiya & Koch, 2009). In implicit learning literature, the role of attention is not well-defined. Generally, the question is whether implicit learning is dependent on attention, in that only content that is selectively attended, can be learned. The alternative view is that implicit learning is independent from attentional processes, thus a non-selective process that includes all contents in potential associative learning (for an overview, see Jiménez, 2003).

First of all, the definition of attention is crucial to the question of its role in implicit learning. As discussed in the first article (Tavera & Haider, 2025), attention is an often under-defined concept. This is potentially dangerous because it can lead to regressive reasoning, assigning all kinds of functions to a homunculus that one names attention, without having explained the mechanisms that make such functions possible or which parameters might influence them. Anderson (2011) even went as far as saying that there “is no such thing as attention”. Researchers should, in a way, try to explain attention away, instead of installing it as a homunculus. And they could do that best by not assigning the concept of attention a *causal* role in cognitive processing, but posit it as an *effect* of cognitive processes (Anderson, 2011).

Not only is it important to define attention conceptually, and its empirical operationalization, but also, to establish the distinction between consciousness and attention. In our case, in

the visual domain. Intuitively, one could picture both consciousness and visual attention as a kind of spotlight (Posner, 1980) or zoom lens (C. W. Eriksen & St. James, 1986): Whichever object, location or dimension is attended is also conscious. Or, the reverse, whichever is conscious, is attended. The causal relationship is equivocal. To test for dissociation or a kind of causal relationship between the two is challenging, given that research questions, proposed mechanisms, and methodology are overlapping in the literature on both concepts. It remains unclear whether they are comprised of independent or interdependent mechanisms (e.g., Lamme, 2003; Marchetti, 2012). Tsuchiya and Koch (2009) review cases in which something is attended, yet does not become conscious (e.g., priming with invisible stimuli, blindsight patients) and cases in which there is conscious awareness without attention (e.g., understanding the gist of a briefly presented visual scene, conscious awareness of objects in peripheral vision and secondary tasks). Still, there are arguments for the idea that attention is a necessary prerequisite for consciousness (Marchetti, 2012; J. Prinz, 2011). It is also difficult to distinguish the two concepts by means of neuropsychological empirical work, as hitherto research on the NCC might have been confounded with neural correlates of attention (Tsuchiya & Koch, 2009). Alternatively, there is the hypothesis that the two concepts can be dissociated by their neural mechanisms, but in the end feed the same process that then shapes experience and behavior (Tallon-Baudry, 2011), which then again makes it difficult to not conflate the two.

In our case, the differentiation between consciousness and awareness is fairly straightforward: We actively manipulate attention by rendering stimulus features task-relevant or task-irrelevant. Additionally, participants are not explicitly instructed about the contingencies in the learning paradigm, meaning that they can acquire them implicitly or explicitly during the learning phase. In a test phase, we will test for conscious awareness of the contingencies. In case we can establish that learning remains implicit, we can determine the role of attention in implicit

learning. More specifically, we test the hypothesis that contingencies involving task-irrelevant stimulus features are not learned.

6.1 Attention in Implicit Learning

In the work at hand, we define attention as the effect of the manipulation of task-relevance of a stimulus or stimulus feature (Tavera & Haider, 2025). That means, that when a feature, such as stimulus color, is task-irrelevant, it will not be attended, and, when another feature, such as stimulus shape, is task-relevant, it is consequentially attended. The consequence of the feature being task-relevant would then be an enhanced processing of the attended stimulus or feature (O'Craven et al., 1997), and its inclusion into a memory episode (Zivony & Eimer, 2022), which makes it available for learning, for example, of contingencies with other features (Logan & Etherton, 1994). This approach, operationalizing attention as task-relevance, has been applied before. In terms of different, and enhanced processing, it was found that ERP markers for early visual processing of stimuli differ significantly between task-relevant and task-irrelevant stimuli (Biehl et al., 2013). In behavioral studies, it could be shown that learning is contingent on task-relevance (e.g., Brosowsky & Crump, 2021), but there is also some evidence that learning is independent from task-relevance (e.g., Seitz & Watanabe, 2009). Lastly, in attempts to model human behavior, task-relevance is a central factor, for example in models of visual search (Navalpakkam & Itti, 2005). However, our approach is to specifically look at task-relevance in *implicit* learning, not just in behavior and learning generally.

In our paper (Tavera & Haider, 2025), we reviewed empirical work concerned with the role of task-relevance in implicit learning. There are studies that suggest that implicit learning processes can only include stimuli or stimulus features that are task-relevant. This was demonstrated in visual statistical learning of temporal sequences (Turk-Browne et al., 2005), in sequence learning tasks (Jiménez & Méndez, 1999, 2001), AGL tasks (Eitam et al., 2013; Eitam et al., 2009), and CC tasks (Jiang & Chun, 2001; Jiang & Leung, 2005; Vadillo, Giménez-

Fernández, et al., 2020). However, not all studies did unequivocally support the requirement of task-relevance for to-be-learned stimuli. There might have been power issues (Jiang & Chun, 2001) and issues of sensitivity of the chosen measure of learning, because there were effects of learning of task-irrelevant stimuli on other measures (Jiang & Leung, 2005). In three high-powered experiments, Vadillo, Giménez-Fernández, et al. (2020) replicated Jiang and Leung (2005), showing that when participants were instructed to ignore distractors of one color while attending only to distractors of another color, they only learned contingencies of the distractor configuration of the attended color, but not of the ignored color. Not only was there no learning effect visible in the response times for the ignored color, but also no evidence for a latent learning effect or a facilitative effect when the ignored color was then rendered the attended color. Evidence for the opposite hypothesis, claiming that also task-irrelevant features can be integrated in implicit learning processes, is sparse. There is compelling evidence from the CC paradigm (Endo & Takeda, 2004; Kunar et al., 2006; Kunar et al., 2013), and the flanker task (Miller, 1987) demonstrating learning effects also for ignored or task-irrelevant stimuli and features. However, those studies using the CC paradigm manipulated attention cross-dimensionally, comparing spatial configuration cues and identity cues (Endo & Takeda, 2004) or spatial configuration and color cues (Kunar et al., 2006; Kunar et al., 2013). It seems that only manipulating attention *within* a feature, for instance comparing a task-relevant color with a task-irrelevant color, leads to selective mechanisms that hinder learning of the irrelevant feature.

Another opportunity to potentially observe attentional mechanisms in implicit learning is the case of cue competition. Cue competition arises when there is not only one contingency between a cue and an outcome, but other cues are contingent as well. The most well-known cue competition effects include overshadowing and blocking (Mackintosh, 1971), as well as compound learning (Thein et al., 2008). Overshadowing occurs when two cues predict an outcome, but only one cue contingency is then learned. To be more exact, later research has found that

the other contingency is in fact learned and can be shown once the overshadowing contingency is extinguished, but is not initially reflected in behavior (Kaufman & Bolles, 1981; Matzel et al., 1985). Blocking is described as the observation that given one cue-outcome contingency has been learned, a second cue that is predicting the same outcome with the same contingency is not learned additionally. Compound learning occurs when two cues are presented together and both are learned as predictors of the outcome (Kehoe & Gormezano, 1980). This learning effect can be additive, meaning that the combined effect of the cues corresponds to the sum of their individual learning effects (Thein et al., 2008). Alternatively, it can be overadditive, meaning that the learning effect of the compound cue exceeds what would be expected from simply adding the effects of the individual cues, suggesting a potentiation of the effect by the compound (Durlach & Rescorla, 1980).

There has not been much research conducted to examine the effects of providing multiple cues in implicit learning (J. R. Schmidt & De Houwer, 2019). In the few studies that have been conducted, cue competition such as overshadowing or blocking effects have not been found (Beesley & Shanks, 2012; J. R. Schmidt & De Houwer, 2019). Only when one cue was more relevant to the task than the other, overshadowing effects were observed (Endo & Takeda, 2004). But to investigate these cue competition effects is an interesting approach to test two proposed learning mechanisms against each other for the case of implicit learning (Beesley & Shanks, 2012): Accounts of associative learning (Rescorla & Wagner, 1972), and propositional learning accounts (e.g., Mitchell et al., 2009). Propositional accounts posit that cue competition arises from forming conscious inferences about contingencies. And in the case of learning that remains implicit, cue competition would not be expected (Beesley & Shanks, 2012). In contrast, associative accounts would predict that cue competition arises from attentional mechanisms during learning. For example, in the case of overshadowing, there would be mechanisms that emphasize one cue, and its contingency would then be more strongly associated with the

outcome than the contingency of the other cue. This mechanism would apply to implicit and explicit learning alike. Accordingly, if cue competition is found in implicit learning, this would support an associative account, and the hypothesis that attentional mechanisms that solve the cue competition operate both in implicit and explicit learning. However, if there is no evidence for cue competition and selective mechanisms, this would support propositional accounts and the idea that for mechanisms of selective attention, conscious propositions are a prerequisite.

6.2 Summary Study 1

In Study 1 (see Appendix A), we tested both the effect of task-relevance, as a manipulation of attention, on implicit learning, as well as the case of cue competition in implicit learning. We did so, using our new variant of the CC paradigm. In this variant, the shape or the color of the distractors can be predictive of target location.

We conducted three experiments for Study 1. In a first experiment, we established the new variant of the paradigm, testing whether a task-relevant cue could be learned to predict the target location. In a second experiment, we tested whether also a task-irrelevant cue could be learned, and if so, to a lesser extent, or equally well. And in a third experiment, we examined the case of multiple redundant cues that predicted target location. In all three experiments, we tested for explicit knowledge by implementing a generation task (Chun & Jiang, 2003) coupled with a confidence measure.

In more detail, in the first experiment, we implemented a contingency between the shape of distractors and the target location. There were four possible target locations, and six different distractor shapes. Three of them were 100% contingent with one target location. The other three distractor shapes were equally often paired with all four possible target locations. This way, we had three 100% predictive and three unpredictable distractor shapes. Learning was defined as an effect of predictability (predictive vs. unpredictable distractor shapes), in interaction with our timing variable block (à 48 trials) on response times.

We analyzed the learning phase fitting mixed-effects models to the response time data. We found contingency learning, indicated by a significant interaction between predictability and block. In the analysis of the generation task, we did not find evidence for explicit knowledge of the contingencies between predictability and target location. Participants did not predict the target location associated with the respective predictive shapes above chance level, and were not more correct given a high confidence response, than given a low confidence response. This means that even when participants gave a correct response in the generation task, they randomly indicated high and low confidence, and thus, were not aware of their accuracy. Bayesian analyses revealed substantial evidence for a null difference. Also with the reverse base probabilities, their confidence was not higher given a correct response, than given an incorrect response. Additionally, we found a significant positive correlation between accuracy in the generation task and the CC effect in the learning phase. This could mean that implicit knowledge, that is shown in the learning phase, feeds into generation task performance. But we did not find any correlation between accuracy and confidence in the generation task. A positive correlation could potentially indicate explicit knowledge. Taken together, we have found no evidence of explicit knowledge of the contingencies.

In the second experiment, we kept the material and experimental procedure constant, while only changing the distractors in the search displays. This way, it was not the shapes but the colors of the distractors that were contingent with target location. Again, three distractor colors were predictive for target location, three other colors were unpredictable. The distractor shape was held constant. We obtained similar results as in Experiment 1. The interaction between predictability (predictive vs. unpredictable color) and block was significant in a mixed-effects model with the same fixed effects as in Experiment 1. The fixed effect estimate for the interaction was similar in size in Experiment 1 and 2. As in Experiment 1, we did not find evidence for explicit knowledge of the contingencies, indicated by strong evidence for the null

hypothesis when comparing relative frequencies of correct responses given high confidence, and correct responses given low confidence. The same result was found for the other base rate comparison.

In the third experiment, we investigated potential cue competition. We combined the material from Experiments 1 and 2, so that the shape and color both, and redundantly, predicted target location. Thus, were able to observe the interplay of multiple sources of predictive information, implemented here by two visual cues. The learning phase was structured like the ones in Experiments 1 and 2, but with compound cues (shape and color). In an additional test phase in Experiment 3, we implemented single cue blocks after the learning phase. In the color cue block, we introduced a novel distractor shape that had not been associated with any target location, but kept the distractor colors 100% contingent, as in the learning phase. And, vice versa, in the shape cue block, all distractors were presented in a novel color, but the distractor shapes still predicted target locations. Again, fitting mixed-effect models to the learning phase data with both cues predicting target location, we found a significant interaction of block and predictability, indicating learning of the cue – target location contingency. From that result alone, it remained ambiguous what was in fact learned.

There are three possible learning processes that can explain the data: First, it could be that one single visual feature cue was learned, either color or shape of the distractors, while the respective other one was overshadowed. Secondly, participants could have learned the compound cue, representing both visual cues in an integrated fashion. Thirdly, it is possible that participants learned both cues independently. In the first case, we should observe a predictability effect in only one of the single cue blocks. Following the second hypothesis, we should assume that there would be no predictability effect in the single cue blocks that would be comparable to the effect in the learning phase. From the third hypothesis, we should expect an effect of predictability, indicating learning in both single cue blocks.

What we found is in fact an overall significant predictability effect in the single cue blocks, in line with the third hypothesis. But from Bayesian analyses of the individual single cue blocks, there was no convincing evidence for a predictability effect. This may be due to a lack of statistical power that is due to the rather small number of 48 trials per single cue block. However, the number of trials is justified by the assumption that with more trials, participants would likely learn any single cue contingency anew. Thus, we conducted several explorative analyses. When comparing the single cue block response times with the learning phase that provided both cues, we do not find an increase in response times in predictive contexts. Such an increase could have been indicative of costs produced by the lack of one predictive cue. Also, the fixed effect estimates for predictability are roughly the same with shape as the predictive cue, as learned in Experiment 1 and the shape block in Experiment 3, and with color as the predictive cue, as learned in Experiment 2 and the color block in Experiment 3. This suggests that the effect of learning on response times is comparable for single cues or multiple redundant cues. And it suggests that learning might have occurred independently for both cues. As to the nature of this independent learning, we also compared the size of the CC effect (defined as subtracting response times in predictive trials from response times in unpredictable trials) in the learning and single cue blocks of Experiment 3, and found that it is almost double in size in the compound cue learning, when compared to the single cue blocks. This may indicate that the learning of the individual cues is additive in its effect on response times, an effect that has been shown before for animal (Thein et al., 2008) and human learning (Endo & Takeda, 2004). All our findings are compatible with the notion of independent, and even additive learning of both individual cues. However, there was still the possibility that our data looks compatible with that, but is really explained by the first hypothesis and individual differences. It could be that individual participants differ in their preference of one cue, and that, by individual learning effects of only one cue respectively, the overall result would look like both cues were learned. This is why we conducted a final explorative analysis, comparing the CC effect of each

individual participant for the color and shape single cue blocks separately. If individual learning effects of only one cue would explain our finding of learning of both cues on the group level, we should see a pattern that roughly half of the participants should show a learning effect in only one cue block, and the other half in the other cue block. However, this is not what we found. Instead, we found that the difference scores between the color and shape block CC effects revolved around zero. This indicated that both cues were learned to a similar extent. From the generation task and confidence measures, we again concluded that knowledge of the contingencies – involving color or shape – remained implicit.

6.3 Implications

Experiments 1 and 2 from Study 1 showed that contingencies with visual features that are relevant or irrelevant to the task are learned without explicit instruction, and that the knowledge of the contingencies cannot be reported explicitly.

This has implications for theories that propose attentional mechanisms as key components of their models of consciousness. First, the GWT assigns a central role to attention, claiming that once an information is attended, it is broadcasted into the global workspace and thus becomes conscious. In HOT, attention is thought to amplify processing, and making it more likely for information to become conscious. Lastly, in IIT, attention and consciousness are viewed more distinctly. Still, it is suggested that when attention contributes to integrative processes, it can contribute to the emergence of conscious processing.

In our studies, manipulating attention in terms of task-relevance, we do not find essential differences between task-relevant (attended), and task-irrelevant (not attended) features. Neither in the extent of learning a contingency between the feature and a goal-relevant target location, nor in the extent of conscious knowledge of those contingencies. Of course, one could argue that our definition and manipulation of attention is rather specific, and that given a different approach to attention, one could possibly find differences in attended and not attended

features. Also, one could suggest that our finding is limited to conclude the role of attention with regard to stimulus *features* but not stimuli as a whole, and argue that manipulating attention to whole sets of stimuli could yield different results (as in Jiang & Leung, 2005; Vadillo, Giménez-Fernández, et al., 2020). However, we can retain the following: When defining attention as a selection mechanism, not a finite resource (as in, e.g., Frensch et al., 1998), and as a mechanism resulting from task-relevance manipulation, not as a causal factor in cognitive processes (as in, e.g. J. H. Reynolds et al., 2000), we do not find evidence for a central role of attention in unconscious processing. This is generally broadly consistent with GWT, HOT and IIT. However, in GWT, attention has a central role, and one could deduct the hypothesis that contingencies between attended features should have a higher probability to become conscious, than contingencies between not attended features. This hypothesis could also hold for HOT and IIT, given that attention should play a role in the likelihood of information being processed consciously. And given the IIT account, the sheer process of visual search, scanning the distractors one by one while searching for the target, should increase feature integration (Treisman & Gelade, 1980), thus leading to more integrated information, which is then again more likely to become conscious according to IIT. However, one issue that remains with those implications is that the hypotheses deducted from the broad theories of consciousness remain vague with regard to their predictions in implicit learning processes.

Experiment 3 additionally provides insights into the mechanisms of cue competition in implicit learning. Our findings are compatible with other studies that did not find overshadowing and blocking effects in implicit learning (Beesley & Shanks, 2012; J. R. Schmidt & De Houwer, 2019). But our study goes beyond this by showing that overshadowing does not even occur when one of the provided cues is task-irrelevant and would thus be more likely to be overshadowed by a task-relevant cue, as it was demonstrated by Endo and Takeda (2004).

There are several conceivable practical implications of this study. Implicit learning in everyday life means that we learn statistical occurrences, rules, and contingencies without instruction or explicit awareness about them. But further, our results suggest that implicit learning is a rather automatic process that is not dependent on attention allocation. This means that potentially, we can learn a myriad of associations every day without being aware of it, and, importantly, without the stimuli being relevant to our task or goal. There has been a debate over the magnitude of the unconscious influence on behavior, such as consumer behavior or social judgement. In their review, Newell and Shanks (2014) criticize studies that claim a broad scope of unconscious processes influencing behavior for their poor methodology. Still, there are methodologically sound studies that show effects of implicit learning in the realm of learning body cues in social interactions (Heerey & Velani, 2010; E. Norman & Price, 2012). And some argue for a pronounced role of implicit learning in behavioral economics (Zizzo, 2000). Here, it would be interesting to consider the role of attention in those implicit learning situations, because it is a relevant factor in our overloaded everyday environments.

Further, there are efforts to investigate inter-individual differences in implicit learning to learn more about psychological and neurological conditions. For instance, there is the idea that schizophrenia could be associated with altered implicit learning processes. There is evidence of impaired (Horan et al., 2008), but also of intact (Danion et al., 2001) implicit learning in participants with schizophrenia. In addition, some research suggests that the ability to distinguish relevant from irrelevant stimuli plays a role in schizophrenia (Gray & Snowden, 2005). Because our novel variant of the CC paradigm provides the opportunity to test implicit learning and the role of task-relevance, testing participants with schizophrenia in this paradigm could contribute to this debate. This could also be true for other inter-individual differences, for example, when investigating potentially altered cognitive processing in individuals on the autism spectrum. There has been a line of research hypothesizing that there could be a lack of implicit

social learning in those individuals. However, evidence mostly shows no differences in implicit learning (for a meta-analysis, see Foti et al., 2015), also in children (Barnes et al., 2008). This evidence may be enriched and refined by further looking into implicit learning while manipulating attention.

7 Modular Processing of Unconscious Information

The holistic perception that we have of the world may be what most characterizes our conscious experience. We do not perceive visual features of a stimulus, like color and shape, or whole stimuli separately. In our conscious visual perception, we experience the world as integrated picture, almost like a movie played in front of our eyes (Damasio, 2000). But when we look at the neurophysiological basis of visual perception, information processing in our visual cortices is highly specialized and modular with regard to visual features (Ghose & Maunsell, 1999).

A central claim of many theories of consciousness is that processing of integrated information that produces such a holistic experience, requires consciousness, either in the form of a global workspace (Dehaene & Naccache, 2001), of a high level of integration (Tononi, 2004), or integration by higher-order thoughts (Rosenthal, 2005). In contrast, unconscious processing is thought to be modular. That means that visual, auditory, and other information are processed separately in independent and specialized modules (Abrahamse et al., 2010; Frost et al., 2015; Keele et al., 2003). Yet, it is still a matter of debate how these modules are defined, whether by modality, such as vision, hearing, tactility and olfaction (henceforth labelled modalities; e.g., Abrahamse et al., 2010; Keele et al., 2003), or even more refined, features within the modalities, such as color, shape, and location in vision (henceforth labelled features; Eberhardt et al., 2017; Moeller & Pfister, 2022; Wilts & Haider, 2023).

There are several ways to test whether modules are specialized with respect to modalities or features. A common approach is to test for concurrent learning of uncorrelated sequences. Modularized models of unconscious processing suggest that multiple sequences can be learned concurrently, as long as they are processed in encapsulated modules (Keele et al., 2003). Thus, the modality-based account predicts that sequences can be learned concurrently when they belong to distinct modalities, whereas the feature-based account predicts concurrent learning for

sequences that comprise of distinct features. It has been repeatedly demonstrated that independent sequences in different features (color, shape, location) can be learned concurrently, while two sequences instantiated within a feature cannot be learned (Eberhardt et al., 2017; Goschke & Bolte, 2012; U. Mayr, 1996; Wilts & Haider, 2023), which challenges a modality-based account, and supports a feature-based account. In Study 2, we investigated the hypothesis of independence of feature modules in our variant of CC paradigm.

7.1 Transfer in Implicit Learning

Another deduction from the hypothesis that the architecture of unconscious processing is comprised of encapsulated modules, is that these modules do not exchange information. This subsequently prevents information transfer across modules.

The modality-based account would thus predict that transfer can only occur for information within the same modality. For instance, that color information can be integrated with shape information, as both features are processed within the visual module. In contrast, visually perceived spatial information could not be transferred into motor, spatial information, because there is a distinction between the visual and the motor modules.

In contrast, the feature-based account predicts that, within the visual modality, the module processing color does not have access to information represented in the shape module, and a spatial module could not enrich its information processing of stimuli with their color and shape in the absence of awareness. In the context of implicit learning, this would mean that learned contingencies within one feature, such as color, remain encapsulated information in the color module. Implicit knowledge of these contingencies should therefore not be accessible to the shape module. It is however important to note that in the feature-based account, the modules do not distinguish between perception and action (Hommel et al., 2001; W. Prinz, 1990). Thus, location information can be used for the guidance of visual attention (perception) and movements (action). This would apply to cue-target contingencies like in CC, but also sequential

contingencies like in the SRTT. However, Haider et al. (2020) found that a visually perceived stimulus location sequence could be transferred to a motor response location sequence, while knowledge of both remained implicit. However, location may play a special and pronounced role in our cognitive system and in learning (I. Koch & Hoffmann, 2000), Haider et al.'s finding concerning information transfer is not readily generalizable to other features. It remains ambiguous whether this is generalizable to visual modules, such as those for color and shape. In our study, we attempted to test whether such knowledge transfer is possible between two visual features, transferring color cue contingencies to shape contingencies, while knowledge of these contingencies remains implicit.

While the feature-based account is inconsistent with transfer between color and shape modules, would such a knowledge transfer be possible when certain prerequisites are met? In our paper (Tavera, Wilts, & Haider, unpublished), we proposed two theoretical candidates for mechanisms that would enable such information exchange or transfer between modules, and that do not necessarily require rejecting such an architecture. First, one could deduct such a mechanism from the theory of event coding and the concept of event files (Hommel, 1998; Hommel et al., 2001). This framework suggests that perception and action control occur via the construction of event files. These entities encode information from an experience, across all stimulus features and across perception and action (W. Prinz, 1997). This information can be used for a future encounter with the same or similar stimuli, by, for example binding together a stimulus and an appropriate response. This mechanism is not limited to task-relevant information, but is thought to be an automatic integration of all information available (Rothermund et al., 2005). This way, there would be a mechanism that binds together features from the processing of different modules into a structure, that is here called an event file. However, in this framework, it is often proposed that these bindings are transient, and there is hardly any empirical or theoretical connection to the research of learning, meaning a longer-lasting change in

behavior potential. Because there are recent advances to extend the framework in this direction (Arunkumar et al., 2024; Frings et al., 2020; J. R. Schmidt et al., 2020), we suggest that one deducted mechanism from the framework could be applied to explain knowledge transfer between feature-based modules. In the original work on event files (Hommel, 1998, 2004), the role of conscious awareness is not specified. But there is empirical work that suggests that the integration of information into an event file does not depend on top-down or attentional processes, but occurs automatically (Rothermund et al., 2005; Schmalbrock et al., 2023). This makes it a viable candidate for a mechanism in implicit learning, and an explanation for knowledge transfer in the absence of awareness. Therefore, in our study, we let participants form associations between shapes and colors, to enable a transfer from a learned contingency between shapes and target locations, to a contingency between colors and target locations.

A second framework that enables the deduction of a potential mechanism for knowledge transfer is learning in pre-conditioning situations (Holmes et al., 2022). In a typical preconditioning procedure, two stimuli (S) are associated with each other (S1-S2). Then, a response (R) is associated with one of the stimuli (S1-R). In a transfer phase, it is then shown that the second stimulus is also associated with the response (S2-R), although it has never been paired with it (e.g., Arunkumar et al., 2024). This is explained with an online integration account, suggesting that the S1-R association is bound into the S1-S2 association (Holmes et al., 2022). The role of conscious awareness in this mechanism has been examined recently, suggesting that the S1-S2 and S2-R associations need to be explicit for learning to occur (Arunkumar et al., 2024). The measure of awareness in this work is however questionable in terms of reliability, as it relies on one-trial assessments for each association. It is further noteworthy that the stimuli used in this experiment were visual (S1) and auditory (S2), thus testing a transfer of knowledge across modalities, not across features within modalities. Again, similarly to the event file framework, a

pre-conditioning account would also predict the transfer of a learned contingency from a shape cue to a color cue, given that shapes and colors were associated with each other.

The two frameworks differ in their prediction concerning the order of association formation. In the event file framework, the mechanism of knowledge transfer would be, that in every trial of the learning that retrieves event files with a shape and the respectively learned target location, the before matched color would be activated as well, being part of the event file. Then, in the transfer trials, when only color is activating the event file that was encoded before, the target location is activated along with it. One could argue whether this mechanism requires a S1-S2 matching before the S1-R learning, or whether the matching could be added after the learning. In the former order, it seems more probable that, once the one-to-one shape-color matching has been learned, shapes and colors are then over the course of the learning associated with target locations with increasing strength. In the latter order, the associative strength acquired over the course of the learning between shape and target location would have to be transferred to an association with color as well. This is not impossible against the backdrop of the event file framework, but there is also no mechanism in place to explain such a transfer of association. However, it remains an empirical question whether the mechanism does function in this way. The preconditioning account on the other hand is not confined to a fix order of association acquisition.

Additionally, since we cannot distinguish between the two proposed mechanisms of forming associations as in event files or as in pre-conditioning, we use the term matching to describe the learned relationship between shapes and colors.

There are several reasons to switch from the SRTT paradigm as used by Haider et al. (2020) to our adapted CC paradigm. In the SRTT, participants learn a visual sequence, whereas in our CC variant, they learn contingencies between features cues and target locations. This is potentially interesting to additionally test cross-dimensional learning. Commonly, in the

learning phase of the SRTT, the sequence is 100% predictable to ensure stable learning effects (e.g., Wilts & Haider, 2023). As a consequence, the learning effect cannot be assessed during the learning process. Only in a test phase, one can compare response times in sequential and random material to assess learning. In contrast, in our CC variant, there is an online measure for the learning effect during the learning process, as participants respond in every trial, and one can then compare between trials that are either predictive or unpredictable. However, it is yet uncertain whether these two paradigms test the same learning mechanism. Sequence learning is a kind of chaining across trials (Schuck et al., 2012), where one event is linked to a subsequent and/or a preceding one. In contrast, in CC, the contingency between the cue and the target location is within a trial. However, precisely for this reason, testing transfer between visual features within CC not only tests the generalizability to other features, but also to other learning paradigms.

7.2 Summary Study 2

In this study (see Appendix B), we tested whether transfer of contingency knowledge between two visual features, color and shape, is possible. We examined this in our novel variant of the CC paradigm. We designed the experiment such that we tested learning and transfer in three groups, a pre-matching group, a post-matching group, and a control group. All groups completed four experimental phases, the matching phase, learning phase, transfer phase, and generation task, but in different forms and orders.

The matching phase was the main experimental manipulation. The pre-matching group learned to associate shapes and colors prior to the learning phase. The post-matching group learned the shape-color associations only after the learning phase, but before the transfer phase. The control group did not learn any associations between shape and color, but did a filler task with the same set-up as in the other groups. The learning, transfer, and generation phases were the same for all three groups. In the learning phase, the distractor shapes predicted target

locations with a 70% contingency, but note that it did not predict correct responses. In the transfer phase, the distractors in the search displays were not characterized by shapes but by colors. The 70% contingencies were now transferred from the shapes to the colors, according to the shape-color matching from the matching phase. Lastly, in the generation phase, participants were asked to generate target positions from search displays, both with the distractor shapes from the learning phase, and the distractor colors from the transfer phase, and to indicate the confidence in their responses.

Fitting mixed-effects models to the data of the three groups, we found comparable learning effects in all three groups, indicated by significantly decreased response times for predicted trials (70%) when compared to unpredicted trials (30%). The learning effects did not differ significantly between groups. In the transfer phase, we observed a significant transfer effect only in the pre-matching group. But in an overall analysis, the transfer effects did also not differ significantly between groups. It thus remains unclear whether the observed effect in the pre-matching group can unequivocally be interpreted as evidence for knowledge transfer. For the shape and the color contingencies, we did not find above chance level accuracy in the generation task or evidence for explicit knowledge in the combined measures of accuracy and confidence judgement.

7.3 Implications

Our findings suggest that knowledge transfer is possible in the absence of awareness under certain conditions. In the pre-matching group, that first learned the shape-color matching, and then the contingencies between shapes and target locations, participants were able to demonstrate contingency knowledge in the transfer phase with the before matched colors. In the post-matching and control groups, we did not find evidence of knowledge transfer. However, we are careful with the interpretation of our findings. First, it is conceivable that what we observe with the response time difference between predictive and unpredictable trials is not in

fact knowledge transfer, but a facilitation effect for novel learning of the contingencies between colors and target locations. Explorative analyses show however that the response time difference between predictive and unpredictable distractor colors is present from the first trials in the transfer phase, and do not, as would be expected from novel learning, increase with trial number. The second caveat concerning our results is the lack of group differences in the transfer phase. From the models that we fitted, it can be seen that the groups do not significantly differ. Descriptively, it is still interesting to see that the transfer effect is largest in the pre-matching group, followed by the control group and then the post-matching group. It would be interesting to test other experimental conditions that manipulate potential covariate candidates. For example, the post-matching group could have had a disadvantage by having a task in between the learning and transfer phase. We do hypothesize that the kind of learning we are examining here is long-term learning, as studies show that CC effects prevail for multiple days (Bergmann et al., 2019; Chun & Jiang, 2003). But when taking into account the emphasis that the event file framework puts on the transience of bindings in event files, there might be mechanisms that dissolve such bindings or associations, for example whenever the task context drastically shifts (Gozli, 2019). It would be interesting for further research to look into such parameters and their influence on a transfer effect.

What do our results indicate for the scope of unconscious processing and theories of consciousness? The lack of knowledge transfer of contingencies that remain implicit in the control and post-matching groups is consistent with GWT, because more complex, integrative processes should require consciousness. Additionally, the information exchange that is required for transfer processes would not be in the scope of modularized processing as hypothesized by GWT (Baars, 1997).

Reconciling the findings with IIT accounts requires additional mechanism specifications. According to IIT, when information is highly integrated, and therefore highly flexible, it

should become conscious (Merker et al., 2021). Our finding, the acquisition of flexible, transferable knowledge that remains unconscious, does not correspond to mechanisms proposed in IIT. But as there is no possibility of quantifying integration processes in our study by computing ϕ , it is also not contradicting the theory clearly.

HOT entail that unconscious information processing can be highly complex, flexible and integrative, and does not become conscious as a consequence of its complexity, but as a consequence of forming a higher-order representation of the information (H. Lau & Rosenthal, 2011). In our study, we specifically ask participants not only for their contingency knowledge, but also for a metacognitive judgement, their confidence in their responses. We do not find that participants have metacognitive knowledge – they do not know that they (implicitly) know something about the contingencies. According to HOT, this means they are not conscious about it. But also for HOT, conscious awareness is not necessary for complex information processing such as transfer to occur.

In sum, our study provides valuable insight into unconscious processing and the characteristics of the modular processing architecture, while raising interesting questions for future research.

8 Semantic Processing Without Awareness

Semantic processing is commonly viewed as a main functional aspect of human consciousness. It is thought to be a consequence of integration processes (Ludwig, 2023). Accordingly, it is a central prediction of some theories of consciousness that consciousness emerges or is required whenever semantic information is processed (e.g., Dehaene & Naccache, 2001; Tononi, 2008). Given the centrality of high-level, semantic processing as a function of consciousness, it has been studied extensively. Empirically, one of the most used paradigms in psychological consciousness research is testing the influence of subliminally presented primes on the processing of semantically congruent or incongruent targets (Ludwig, 2023; Mudrik et al., 2014).

With this approach, a series of empirical studies has attempted to carve out the limits of unconscious processing with respect to semantics. There is evidence from ERP studies suggesting that semantic processing of words takes place even when they are not consciously perceived (Kiefer & Spitzer, 2000; Luck et al., 1996; Stenberg et al., 2000). This was shown by the finding that masked words produced the N400 ERP which is linked to semantic, integrative processing (E. F. Lau et al., 2008). Further, it has been shown that the amygdala specifically responds to masked images of fearful faces (Whalen et al., 1998). These studies are important, as they demonstrate that there is neurological evidence for semantic processing of invisible stimuli, and that this processing is done in the same neural pathways as conscious semantic processing. Yet, they do not allow conclusions from a phenomenological or behavioral standpoint. The question that remains is, how the activation of semantic processing networks, or the amygdala, translates into perception and behavior. There is evidence from behavioral studies that address this issue. For example, using subliminal priming paradigms, semantic processing of number words could be demonstrated (Dehaene, Naccache, et al., 1998; Naccache & Dehaene, 2001). The number words were used as masked primes, not consciously perceived by participants. Then,

participants had to respond to a two forced-choice response, indicating whether a target stimulus was a number above or below 5. They had to respond with the right or left hand. Dehaene, Naccache, et al. (1998) measured the lateralized readiness potential (LRP) that is associated with left or right hand response preparation. They could show covert motor priming, finding the LRP in accordance with the response primed by the masked number word. In a later study, they further showed the same priming effect of masked number words, but here, in response times for congruent or incongruent target stimuli (Naccache & Dehaene, 2001). Similar findings have been obtained with different stimuli and set-ups, finding response facilitation or inhibition by subliminal priming (for a review, see Eimer & Schlaghecken, 2003). There are even studies suggesting that there could be unconscious priming when prime and target are not of the same modality (Lamy, Mudrik, & Deouell, 2008). These findings were mostly interpreted as evidence for unconscious semantic processing. However, this implication was challenged by an alternative explanation of semantic priming effects. Kiesel et al. (2008) reviewed such findings and explain them on the basis of solely perceptual processing that can trigger actions. That means, for instance, that there is a learning of key responses triggered by the visual characteristics of numbers on the screen, not by their semantic meaning. With this action trigger account, subliminal priming effects would thus not require semantic processing. Therefore, it is still a matter of debate whether semantic processing is possible in the absence of awareness. Assuming that it is possible, it is still unclear what the limits of such processing are. For example, semantic priming effects might last a few milliseconds, but might not enable temporally stable representations without conscious awareness, and further, novel semantic integration might not be possible (for an overview, see Dehaene & Changeux, 2011; Dehaene & Naccache, 2001).

In addition to studies using priming paradigms, a substantial body of research has used visual scene processing to investigate semantic processing without awareness. To evaluate such studies, it is important to understand visual scene processing generally, involving conscious

processing. Visual scene processing is a highly complex process that requires processing of various features, integrating them into objects, and process the relationship between those objects on a semantic level (Biederman, 2017). There is, for example, a vast amount of research on eye movements in scene viewing, examining what is fixated and processed first and most when looking at a scene (Henderson, 2007; Itti et al., 1998), and the interaction of scene and object processing (Brandman & Peelen, 2017; Demiral et al., 2012). Meanwhile, it has often been investigated in combination with visual search tasks (for a review, see Wolfe, 2020). This allowed for conclusions about scene processing and its consequences for attentional guidance to a target (Eimer, 2014). For example, in a real-world scene, certain objects are generally more likely, and expected to be located in certain areas – in a kitchen, a knife is a likely object, and expected on the counter or table, not so much on top of a lamp or on the floor. It has therefore been shown that search processes in real-world scenes were based on such semantic guidance, but also on guidance by low-level target features (e.g., Bahle et al., 2018; Hayes & Henderson, 2019b).

Visual scene processing in the absence of awareness is one framework in which one can study unconscious semantic processing. Because of its complexity, visual scene processing is well suited to test the scope of complex, integrative, semantic unconscious processing. For instance, it can be tested by means of visual masking of scenes, and investigating consequential behavioral or neurological measures.

There are numerous studies that presented masked visual scenes with congruent or incongruent objects, for instance a man taking a baking sheet or a chess board out of the oven (e.g., Biderman & Mudrik, 2018; Faivre et al., 2019; Mudrik, Deouell, & Lamy, 2011; Mudrik & Koch, 2013; Mudrik et al., 2010). Many of them suggest that semantic scene processing, and, subsequently, detection of incongruity of an object within the scene, is possible in the absence of awareness. However, these findings have partly not been replicable, and thus called into

question (Biderman & Mudrik, 2018; Glanemann et al., 2016; Moors et al., 2016). Hence, the role of consciousness in semantic scene processing and integrative processes for the detection of incongruities remains unclear.

However, what we further know about semantic processing of visual scene is based on research on gist perception. The gist of a scene is a quick understanding of the meaning of the scene, possibly a semantic label or category, that is based on low-level feature, as well as on high-level semantic information (Oliva, 2005). Gist perception, operationalized as scene categorization has been found to occur within extremely short time frames. It was shown that presentation times as short as 26ms were enough to enable participants to categorize scenes into natural and human-made with more than 90% accuracy (Joubert et al., 2007; Rousselet et al., 2005). Importantly, rapid categorization of scenes was equally fast for novel and highly familiar scenes (Fabre-Thorpe et al., 2001). Also, participants observing flashing scenes for only 32ms performed at above 90% accuracy in the detection of a food item or an animal in a scene (De-lorme et al., 2000). It is therefore to be expected that participants upon presentation of a real-world scene, are able to extract its gist within a few milliseconds.

So, to recapitulate, several main findings of the hitherto reviewed research on semantic and scene processing are relevant to Study 3. First, from unconscious semantic processing research I deduct that there is the potential that unconscious semantic information influences behavior, as demonstrated in subliminal priming studies. Secondly, from general scene processing literature, we know that scene categorization can be performed in very short time frames and with high accuracy. From this literature, we also know that semantic guidance plays a role in visual search within scenes. From these main findings, we developed our design for the experiments of Study 3. It is important to note that Study 3 is not an examination of scene processing mechanisms themselves. Rather, I will make use of scene processing as a means to investigate implicit learning that involves semantic processing. In Study 1 of this dissertation, we showed

that low-level features such as color and shape could be learned as predictive cues in CC. In Study 3, we now extend this finding by first testing whether such visual cues can be implicitly learned with highly complex stimulus material, such as real-world scenes. Further, we investigate whether also semantic cues could be learned in the absence of awareness. Therefore, we take advantage of the ability of the cognitive system to quickly and accurately categorize visual scenes semantically. We will thus use scene categories as semantic cues in our variant of the CC paradigm. Additionally, we are taking into account that in our visual search task, semantic guidance would play a role if we instantiated real objects as targets. Because this guidance would potentially obscure any CC learning effect, we are using letters, such as in the original CC paradigm, that are meaningless in any visual scene.

8.1 Semantic Processing in Implicit Learning

Aside from general findings in the unconscious semantic and scene processing literature, there is also relevant evidence specifically from semantic scene processing in the framework of implicit learning. In our paper (Tavera, Abderahaman, & Haider, unpublished), we have reviewed the evidence for semantic processing in implicit learning. The idea in such studies is that not only stimulus-stimulus or stimulus-response associations between specific stimuli are learned, but rather between a category of stimuli and another category of stimuli or a response. This would then mean that the categorization of stimuli could take place in the absence of awareness. Or, if the categorization is performed consciously, that the learned contingencies remain implicit.

For instance, there is some evidence suggesting that visual sequences of objects of semantic categories can be learned implicitly (Brady & Oliva, 2008; Goschke & Bolte, 2007). Also within the CC paradigm, studies could demonstrate implicit learning of semantic word categories (Goujon et al., 2009) and semantic scene categories (Goujon, 2011). By having reviewed these studies in detail (Tavera, Abderahaman, & Haider, unpublished), we identified

some methodological issues that call into question whether the evidence is convincing. Specifically, one challenge in such studies is the isolation of semantic categorization from potential low-level categorization. The categories that served as predictive cues could potentially be built based on low-level similarities, thus not provide convincing evidence for semantic processing. Secondly, in all of the four studies reviewed, the awareness tests did not convincingly show that learning remained implicit, according to some of the standards set out in Chapter 4. Therefore, with Study 3, we aim to enhance the validity of an experimental approach to find implicit learning that involves semantic processing.

8.2 Summary Study 3

In Study 3 (see Appendix C), we first tested whether participants learn the contingency between the color of a complex, real-world scene, and target location within the CC paradigm. This is an extension of the finding of low-level visual feature learning in Study 1 to more complex stimulus material. Further, we tested whether a semantic scene category can be used as a cue when there is a contingency between the cue and target location. As in the studies before, we tested whether this contingency knowledge remained implicit when participants are not informed or instructed about any contingencies explicitly. It is important to note that this is a different approach from, for instance, unconscious priming literature. What those experimental approaches aim to show is semantic integration in the absence of awareness. In our Study 3, given that we find learning of semantic cues in our CC variant, it would not mean that semantic processing, in the sense of, for example, abstract reasoning, can take place in the absence of awareness. Because the semantic processing could involve conscious processing. Still, demonstrating learning would show that semantic cues can be bound into implicit learning. So, that semantic, integrated information can be used in unconscious processing to learn about contingencies, without the contingencies necessarily becoming conscious.

For Study 3, we conducted three experiments. For all three experiments, we used the same stimulus material, that was real-world, colored visual scenes. The scenes were comprised of four categories of functional rooms (bathroom, bedroom, kitchen, and living room) which constituted our abstract semantic categories. Additionally, the scenes were characterized by a color scheme (white, green, blue, and brown). As in Study 2, we implemented 70% contingencies between the cues and target locations.

In Experiment 1, we aimed to test whether our novel variant of the CC paradigm was suitable for examining visual search in complex, real-world scenes. We first tested whether a low-level feature such as color, operationalized as the general color scheme in a scene, can be learned as a cue to predict target location. In Experiment 2, we tested whether the contingencies between scene category and target location could be learned when participants were not instructed about them. In Experiment 3, we then tested whether performance would be similar or different when participants were explicitly instructed about the contingency between scene category and target location.

Interestingly, in Experiment 1, we found that participants did not learn to use the general color scheme of a scene as a cue for target location. Fitting a mixed-effects model to the data, we did not find a response time effect of predictability of the target location. Also, our awareness test, combining responses to the generation task and confidence, showed no evidence of explicit knowledge of the contingency between category and target location. In Experiment 2, we found an effect of predictability. However, the effect was not as hypothesized. In predictive trials, participants responded significantly *slower* than in unpredictable trials. In the generation task, they also showed above chance level performance, but no explicit knowledge as indicated by the combined accuracy and confidence measure, and no significant correlation between accuracy and confidence. In Experiment 3, given explicit instructions on cue contingencies, we found the same effect of predictability, even of roughly the same effect size, as in Experiment

2. Further, as expected from the explicit instructions, we found above chance level performance in the generation task and evidence for explicit awareness in combination with the confidence measure. We also found a moderate to large correlation between accuracy in the generation task, and confidence, as we would expect in the presence of explicit and metacognitive knowledge. In an additional, explorative analysis, we included only participants that performed above chance level in the generation task. Including only those participants, we found virtually the same effect of predictability as in the analysis with all participants, and, consequently, the same as in Experiment 2.

8.3 Implications

In Study 3, we aimed to show generalizability of an implicit learning effect in our variant of the CC paradigm, beyond low-level visual cues to semantic cues. Our results remain somewhat ambiguous.

In Experiment 1, we showed that color schemes in complex real-world scenes could not be learned as cues for target location. In contrast, we have shown in Study 1, that low-level cues, such as color, can be learned as a cue to predict target location, regardless of its response-relevance. We know that semantic scene processing is an automatic process (Joubert et al., 2007) that guides attention, also involuntarily (Hayes & Henderson, 2019b). Therefore, it could be claimed that color is always encoded as part of this process. But the role of color processing in complex real-world scene processing is highly complex and not well understood (Shevell & Kingdom, 2008). As summarized by Oliva and Schyns (2000, p. 179): “Existing data with real pictures [...] suggest that the color is never, always, and sometimes used to recognize a scene”. In the case of our study material, color cannot be used as diagnostic feature to recognize scene category, as is the case with natural scenes (Goffaux et al., 2005). Instead, here, color is a feature that is predictive of target location. However, it is not associated with scene category, as all four

scene categories were presented in all four colors. Therefore, it is unlikely that color processing supports understanding the scene category.

Nevertheless, in Study 1, we have shown that color could be learned as a cue, although it was not a task-relevant feature. Thus, we hypothesized that color cues could also be learned in real-world scenes, independent of their relevance to the task. That we did still not find learning could point to the complexity of the material in Study 3. Here, the cues were not defined as a consistent color, but as a broad color scheme including different hues and luminance, while there were other colors present within the scene. Thus, learning the color cue within complex, real-world scenes needed a certain amount of abstraction (van de Sande et al., 2010), generalizing over a spectrum of hues, defined as, for instance, “green” or “red”. Meanwhile, other colors that were marginally present in the scene had to be disregarded. This might be an integrative and statistical process that cannot be integrated into implicit learning episodes. In line with that, Delorme et al. (2000) have shown that in ultra-rapid scene processing with 32ms presentation times, the presence or absence of color in a scene did not affect scene categorization. They argue that in early stages of visual processing of a complex scene, color might not play a role in semantic categorization. Thus, in Experiment 1 of Study 3, color processing might not have played a prominent role because only very limited scene processing was necessary to do the visual search task.

The second main finding of Study 3 is the learning of scene category contingencies with target location in Experiments 2 and 3. Although this learning was reflected in significantly different response times for predictable versus unpredictable target locations, the direction of the effect was reversed. Similar reversed predictability effects have previously been observed, for instance, in SRTT experiments (I. Koch et al., 2020). However, the explanation they provided in that context is specific to their experimental setup and does not readily apply to our paradigm. As we discussed in our article (Tavera, Abderahaman, & Haider, unpublished), there

are some approaches to explaining this reverse CC effect, including effects of attentional inhibition (Tipper, 2001) or episodic retrieval (S. Mayr & Buchner, 2007). But in additional exploratory analyses, we did not find empirical support for such effects in our data. Since the underlying mechanisms proposed in these accounts remain underspecified, it is difficult to determine whether they can account for the effects observed in our data. Additionally, that we find this reverse CC effect in both implicit and explicit learning conditions is a central finding. It shows that the mechanism producing the effect may not be a top-down strategy, such as intentionally ignoring the predicted target location. Rather, it seems to be a process that stems from implicit processing, and is not overridden by explicit processes, such that the prediction is used intentionally to guide search.

To make these assumptions about the reverse CC effect in implicit and explicit learning requires a valid and reliable measure of explicit knowledge. Thus, the results from the generation task and confidence measure in Study 3 are another core finding. We found substantial evidence for explicit knowledge across participants when they were explicitly instructed with the contingencies, but not when they were not explicitly instructed. This is an important methodological validation of our explicit knowledge test. On the one hand, it served as a manipulation check regarding our instruction manipulation. On the other hand, it shows that our measure is, in principle, able and sensitive to detect explicit knowledge when it is prevalent. In Experiment 2, participants also indicated contingency knowledge by performing above chance level in the generation task, but there was no evidence for explicit knowledge. As discussed in the Introduction, we do not posit above chance level performance as evidence for explicit knowledge, as this performance might also be based on implicit knowledge (Jiménez et al., 1996; Reingold & Merikle, 1988). Therefore, we used the combined measure that reveals metacognitive knowledge about the knowledge (Michel, 2023a).

What we thus conclude from Study 3, is that our novel CC paradigm can only be partly generalized to other stimulus material. Implicit learning of low-level features seems disrupted when it requires a certain amount of abstraction of specific low-level features into broader feature categories. But we have shown that semantic cues can be learned both implicitly or explicitly. The reversal of the effect of learning is an interesting finding that remains to be studied further in future research.

Our finding of implicit learning of semantic category cues can be attributed back to the theories of consciousness that I have discussed in the beginning. However, that requires some rather speculative hypothesis deductions for implicit learning of semantic content, given the broadness of the theories. The question is whether semantic processing, specifically, semantically categorizing a real-world scene, and associating this category with a target location, is within the scope of unconscious processing. To examine this question against the backdrop of the theories, it might be important to note that this question is different from asking whether unconscious semantic processing, in the sense of learning something new, is within the scope of unconscious processing. Because other than in paradigms with subliminal stimulus presentation, in implicit learning, we generally consider participants aware of the presented stimuli. In the case of Study 3, we assume that participants consciously perceive the scenes, and are aware of the scene categories. Neither the semantic categorization itself nor the target locations are implicit, but the learned contingencies between the two are. This finding is difficult to reconcile with mechanisms put forward by theories of consciousness.

GWT posits that semantic processing requires information integration in the global workspace, but not, whether associations between semantic and low-level features can be learned without involving the global workspace. The semantic scene categories and the target locations are accessible in the global workspace. Assuming that the contingencies between them were accessible in the global workspace, it is unclear why the association between the two is

learned, but not accessible to verbal report. Thus, it is conceivable that top-down influence of conscious processes on unconscious processes, as proposed by GWT (Dehaene & Naccache, 2001), play a role here. In the case of Study 3, that could mean that the conscious perception of the scene and the recognition of its category feeds back into unconscious modules that guide attention, and possibly eye movements. However, this would also entail that this is a different mechanism from the top-down attentional amplification or mobilization (Dehaene & Naccache, 2001) of modularized information, because this would render that information accessible in the global workspace, and thus, conscious. Consequently, it seems that this mechanism is under-specified, and we need a clear set of conditions determining what information becomes conscious when involved in recurrent feedback loops between the unconscious module structure and the global workspace.

On the other hand, this finding seems more easily reconcilable with HOT. Although both scene and target location are consciously perceived, and there is a higher-order representation of the knowledge of the scene categories and the target locations, there is no higher-order representation of their contingency. Participants remain guessing when asked about the contingencies, indicating that it is merely a first-order representation that remains unconscious (Dienes & Scott, 2005). HOT postulates a dissociation between performance and higher-order representation which are conscious (e.g., H. Lau & Passingham, 2006). In our study, we also find a dissociation between response time differences in the learning phase (performance), and lack of metacognitive judgement in the confidence measure (higher-order representation). This dissociation can be explained by the lack of a higher-order representation of the association that was learned as a first-order representation, which influences performance but not metacognitive judgement.

Also, mechanisms from IIT can potentially be applied to understand our findings. As in GWT and HOT, IIT would posit that the semantic categories are processed consciously, as

determined by the high integration of information that they require. That the association between the categories and target locations remains implicit, could potentially be explained by hypothesizing that it is processed in specialized, lower-level networks involved in spatial attention and associative learning. These networks might not be accessible to conscious awareness, like aforementioned, whole areas like the cerebellum (Tononi, 2008). However, as it is still debated which brain structures are primarily involved in implicit learning, the validity of this line of argument remains unclear. A recent review has identified the basal ganglia, left inferior frontal gyrus, and hippocampus as associated with implicit learning outcomes, but has also emphasized the central role of the interconnectedness of brain regions for implicit learning (Williams, 2020). This would contradict the idea that implicit learning remains isolated in lower-level neural networks. However, one could propose that the representation of the learned association is just a change in neuronal activation pattern that bias attentional mechanisms. This alone lacks high levels of informational integration, and thus, remains unconscious. At the same time, it influences behavior. To hypothesize precise predictions here, more research would be needed examining the neurocognitive basis of implicit learning (Turk-Browne et al., 2009; Williams, 2020).

To summarize, there are approaches from the three theories of consciousness to explain our findings. Yet, they remain vague, and lack sufficient empirical evidence to refine the proposed mechanisms. Additionally, none of the theories, and no theories of attentional or episodic retrieval processes can, to my understanding, account for the effect that contingency learning can impair performance, both implicitly and explicitly. Future research should aim to specify the implicit mechanisms involved, especially given that similar performance patterns emerged under explicit learning conditions.

9 Conclusion

The main value of this work is to refine our understanding of human consciousness by examining the scope and limits of unconscious processes using an implicit learning paradigm. So far in the literature, there are different approaches to doing so. On the one hand, there is a strong atheoretical approach to search for NCC by examining conscious and unconscious perception (C. Koch et al., 2016). On the other hand, one of the most common empirical approach is to present stimuli subliminally to prevent them from being consciously perceived and processed (for a review, see Kouider & Dehaene, 2007). I argue that a promising way to obtain insight, and to test aspects of leading theories of consciousness lies within implicit learning paradigms. Thus, in this work, I have examined the scope and limits of unconscious processing by testing three major differential aspects of three leading theories of consciousness. I have examined the role of attention (Study 1), the potential for knowledge transfer (Study 2), and the potential of semantic processing (Study 3) in implicit learning. To do so, I developed a novel variant of the CC paradigm, and implemented advanced awareness measures as well as analysis approaches. I have discussed the findings in light of leading theories of consciousness.

A first question regarding the scope of unconscious processing is the role of attentional mechanisms within it. Attention is a key aspect of GWT, and there has been extensive research on the role of attention in implicit learning (e.g., Jiang & Leung, 2005; Jiménez & Méndez, 1999; Miller, 1987). However, attention is often under-defined, conflated with other mechanisms, or used as a homunculus explanation. Therefore, research and findings in this area have been quite heterogeneous. We thus have precisely defined attention as a consequence of task-relevance, and used it as a result of an experimental manipulation, not as causal factor. Thus, we can carefully interpret the conditions under which attention modulates implicit learning. We show that when attention is manipulated by means of task-relevance, it does not modulate implicit, cross-dimensional, within-trial learning. Instead, under these conditions, we find that the

implicit learning process seems to be indiscriminate, and automatically integrates features independent of their relevance. This finding is not in line with previous work proposing that implicit learning depends on selective attention (Abrahamse et al., 2010; Jiménez & Méndez, 1999), and thus raises the question of whether specific task conditions determine the role of attention. Such conditions may include how attention is defined and manipulated, the level at which selection occurs (e.g., object, stimulus, modality, or feature), the temporal structure of the contingencies (whether they occur within or across trials), the basis on which the contingencies are formed (object-, stimulus-, or feature-based), and the learning modality (whether the contingencies are cross-modal or confined to a specific modality). Irrespective of the potential influence of these parameters, we have shown that it can be, in principle, within the scope of unconscious processing to implicitly learn contingencies of task-irrelevant features. This broadens the potential influence of implicit learning on behavior.

But our finding on cue competition within the same study is consistent with hitherto evidence. We have not found cue competition, such as overshadowing or compound learning effects, which also fits our interpretation of Study 1 – that implicit learning processes are automatic and all-encompassing. That also entails that multiple feature contingencies, even if they are redundant with regard to their predictions, are learned independently from each other. There seems to be no attentional mechanism selecting one feature contingency or suppressing another. Just as we did not find attentional mechanisms (de-)selecting a feature based on its task-relevance when only one feature was provided.

While the findings of Study 1 point to a quite parsimonious model of automatic, non-selective implicit learning, the results from Study 2 rather add explanandum to a model of implicit learning – and unconscious processing more generally. At the same time, it is an important finding that expands supposed limits of unconscious processing. GWT and IIT suggest that unconscious processing occurs in a modularized architecture with little integration and

exchange of information in the absence of consciousness. From those theories, but also not specifically within HOT, we found no suitable mechanism that could explain our finding of implicit knowledge transfer. We thus proposed to take into account mechanisms from the frameworks of the event file (Hommel, 1998) and preconditioning (Holmes et al., 2022) literature. We thereby show that the three theories of consciousness reviewed here might fall short of explanations regarding transfer in implicit learning. Their focus lays more on unconscious perception and to some extent, unconscious information processing (such as information integration processes), but less so, on implicit learning processes and mechanisms. We have shown a lack of mechanisms explaining knowledge transfer in Study 2, and advocate for a specification of mechanisms within the unconscious processing system, that can explain our findings, but also findings of information integration in the absence of awareness more generally (for an overview, see Mudrik et al., 2014).

This is a point that is also made by our findings in Study 3. While our findings are compatible with HOT in principle, GWT and IIT can only insufficiently account for our findings that explicit semantic categories and target locations can be associated, but that this learning remains implicit. I have discussed potential future directions to reconcile the findings with the theories. For future research in GWT, the recurrent feedback into the independent modules could be better defined to explain implicit contingency learning of explicitly perceived content. In IIT, we should specify the neurological basis of implicit learning. This way, we would aim to distinguish implicit from explicit learning by comparing the interconnectedness of neuronal structures associated with them.

The three studies presented in this work show a rather broad scope of implicit learning. This calls into question the functional definition of consciousness as information integration process. Instead of a clear distinction of unconscious processing as highly modularized, inflexible system, and conscious processing as the integrative function, we have seen a flexible

unconscious processing system that is able to learn semantic cue contingencies, and that exchanged and integrated information across modules. With GWT and IIT assign functionally complex and temporally extended processes more exclusively to conscious processing (Dehaene & Naccache, 2001), it becomes complex to explain these findings. Conversely, it is easier to reconcile our findings with HOT. They claim that conscious processes are not, in principle, different from unconscious ones, just metacognitively represented (Rosenthal, 2008).

This work contributes also methodologically to the field of consciousness research. First, it shows that the common CC paradigm can be used more flexibly, and potentially to test more diverse research questions than traditionally done. This is especially interesting, because in the novel variant, we can test both the case of cross-dimensional and also cross-modal learning flexibly. As a new methodological tool, our CC variant could impact this separate, but currently strongly debated research question in the implicit learning literature (e.g., I. Koch et al., 2020), and also has interesting implications for theories of consciousness. Secondly, I have refined analysis pipelines for testing for implicit and explicit knowledge. For the learning phases, I have followed the approach of analyzing within-subject effects with mixed-effects models which increase statistical power and account for inter-individual differences in learning and performance (Weinfurt, 2000). This way of analysis should be considered for future research in implicit learning, given that effect sizes in response times are often rather small (e.g., Bergmann et al., 2019; Haider et al., 2012), and statistical power has been identified as an issue (Vadillo et al., 2016). For the awareness tests, I proposed an analysis of direct, objective task performance in combination with confidence measures (following Haider et al., 2011). This accounts for the metacognitive aspect of consciousness (H. Lau & Rosenthal, 2011), and provides a measure of awareness that is richer in information than a simple verbal report. Study 3 provided validation for this measure, as it successfully detected explicit knowledge in the explicitly instructed condition.

9.1 Limitations and Future Directions

Alongside the strengths and methodological advancements of this work, there are limitations to discuss. In all three studies, we applied the same procedure to measure awareness. Compared to prior work in implicit learning, and the CC paradigm in particular, the combination of a direct, objective awareness measure and a confidence measure, was a potentially more refined approach. Nevertheless, it is susceptible to critique. We ensured to keep the awareness test as similar in retrieval context as possible to the learning context (Shanks & St. John, 1994), we increased the number of trials to increase reliability (Vadillo, Linssen, et al., 2020), and provided Bayesian statistics to support our null findings in the explicit awareness tests (Vadillo et al., 2016). Still, one could argue that we did not implement enough trials in the generation task to obtain a reliable measure (Vadillo et al., 2016). However, our generation task trial numbers are comparable to the number of trials with which an independence of CC effect and explicit knowledge was reliably tested (Colagiuri & Livesey, 2016), which increases our confidence in the measure. Further, one could question why we did not exclude potentially consciously aware participants from our analyses. However, this procedure has been under critique, as it has been shown that participants cannot easily be classified as “aware” and “unaware” by a recognition test, due to measurement error (Vadillo et al., 2022), and one would potentially include participants with explicit knowledge, while also excluding participants without explicit knowledge. We thus opted for a group analysis, where performance per participant is linked to their confidence, therefore again accounting for the individual.

Further, each study in this work raises potential future directions. In Study 1, we recognize that our task-relevance manipulation is confounded with the low-level feature itself. That means, in our variant of the CC paradigm, shape is inherently task-relevant, and color task-irrelevant. It is not feasible to permute feature identity and task-relevance by rendering color task-relevant, because that would require a color discriminability of the target. That would then

lead to a color pop-out effect (e.g., Theeuwes & Lucassen, 1993) which would potentially conceal any other effects on response or search times. Nevertheless, it is conceivable that the novel CC paradigm can be extended to other designs and features. Basically, search difficulty can easily be varied (see also Jiang & Sisk, 2019), for example by set size (i.e., number of distractors), complexity of the cue (e.g., not one color, but color combination cues), and complexity of targets (e.g., easy or difficult to distinguish from distractors), or target responses (e.g., number of targets). Moreover, it could be varied as far as presenting cues that allow cross-modal learning, for instance, providing auditory cues that predict target location. Thus, there is potential for numerous research questions and experimental set-ups within our novel variant of the CC paradigm.

Regarding Study 2, we demonstrated one such research question that could be tested within the paradigm. To further support our finding, it could be shown that the knowledge transfer is not only possible from shape to color cues, but also vice versa, from color to shape cues. Further, one could extend the study by testing whether transfer between other cues is also possible. Again, these could be different visual cues, but one could also advance to the question of cross-modal learning by testing transfer between visual and auditory or motor cues.

Lastly, in Study 3, we have successfully extended our novel CC variant to more complex stimulus material. But we did not find learning of low-level cues within complex scenes, and have discussed potential reasons for that. In any case, it would be worthwhile to test the paradigm with different sets of stimulus material, to ensure that the lack of learning is not an artifact of the stimulus material, but generalizable. The noise that complex, real-world scenes bring to the data is a known challenge in scene processing literature, and this is why an increasing number of normed and validated scene data bases is published (Andrade et al., 2024; Greene, 2013; Mohr et al., 2016; Shir et al., 2021). We did not opt for one of them for our stimulus material, because they are, consistent with the most common research question in scene processing,

specifically curated to test the relationship between the processing of scenes and diagnostic objects. The material that we created was statistically analyzed in terms of hue variances across their semantic and color categorization. This is a good starting point to improve this scene material for future research. This would then be the first data-base, to my knowledge, that entails complex, real-world scenes in which low-level features and semantic meaning are not confounded. This way, one can investigate the influence of low-level and semantic characteristics separately.

Taken together, there are still challenges, especially regarding the methodological issue of testing for explicit knowledge, the generalizability of learning to other features and modalities, and the noise in data with complex stimulus material. More than anything, these challenges point to new research opportunities and endeavors to further explore consciousness by means of implicit learning research.

References

- Abrahamse, E. L., Jiménez, L., Verwey, W. B., & Clegg, B. A. (2010). Representing serial action and perception. *Psychonomic Bulletin & Review*, 17(5), 603–623.
<https://doi.org/10.3758/PBR.17.5.603>
- Akaike, H. (1998). Information theory and an extension of the maximum likelihood principle. In *Selected Papers of Hirotugu Akaike* (pp. 199–213). Springer, New York, NY.
https://doi.org/10.1007/978-1-4612-1694-0_15
- Albantakis, L., Barbosa, L., Findlay, G., Grasso, M., Haun, A. M., Marshall, W., Mayner, W. G. P., Zaeemzadeh, A., Boly, M., Juel, B. E., Sasai, S., Fujii, K., David, I., Hendren, J., Lang, J. P., & Tononi, G. (2023). Integrated information theory (IIT) 4.0: Formulating the properties of phenomenal existence in physical terms. *PLoS Computational Biology*, 19(10), e1011465. <https://doi.org/10.1371/journal.pcbi.1011465>
- Alef Ophir, E., Sherman, E., & Lamy, D. (2020). An attentional blink in the absence of spatial attention: A cost of awareness? *Psychological Research*, 84, 1039–1055.
<https://doi.org/10.1007/s00426-018-1100-x>
- Anderson, B. (2011). There is no such thing as attention. *Frontiers in Psychology*, 2, 246.
<https://doi.org/10.3389/fpsyg.2011.00246>
- Andrade, M. Â., Cipriano, M., & Raposo, A. (2024). Obscene database: Semantic congruency norms for 898 pairs of object-scene pictures. *Behavior Research Methods*, 56(4), 3058–3071. <https://doi.org/10.3758/s13428-023-02181-7>
- Aru, J., Larkum, M. E., & Shine, J. M. (2023). The feasibility of artificial consciousness through the lens of neuroscience. *Trends in Neurosciences*, 46(12), 1008–1017.
<https://doi.org/10.1016/j.tins.2023.09.009>
- Arunkumar, M., Rothermund, K., & Giesen, C. G. (2024). One link to link them all: Indirect response activation through stimulus–stimulus associations in contingency learning.

- Experimental Psychology*, 70(5), 259–270. <https://doi.org/10.1027/1618-3169/a000597>
- Atmanspacher, H. (2004). Quantum theory and consciousness: An overview with selected examples. *Discrete Dynamics in Nature and Society*, 2004(1), 51–73.
<https://doi.org/10.1155/S102602260440106X>
- Baars, B. J. (1994). A thoroughly empirical approach to consciousness. *PSYCHE*, 1(6), 1–18.
- Baars, B. J. (1997). In the theatre of consciousness: Global workspace theory, a rigorous scientific theory of consciousness. *Journal of Consciousness Studies*, 4(4), 292–309.
- Baars, B. J. (2005). Global workspace theory of consciousness: Toward a cognitive neuroscience of human experience. *Progress in Brain Research*, 150, 45–53.
[https://doi.org/10.1016/S0079-6123\(05\)50004-9](https://doi.org/10.1016/S0079-6123(05)50004-9)
- Baars, B. J. (2017). The Global Workspace Theory of Consciousness: Prediction and results. In S. Schneider & M. Velmans (Eds.), *The Blackwell companion to consciousness* (pp. 227–242). John Wiley & Sons Incorporated.
<https://doi.org/10.1002/9781119132363.ch16>
- Bachmann, T., Suzuki, M., & Aru, J. (2020). Dendritic integration theory: A thalamo-cortical theory of state and content of consciousness. *Philosophy and the Mind Sciences*, 1(II).
<https://doi.org/10.33735/phimisci.2020.II.52>
- Baddeley, A. D., & Hitch, G. (1974). Working Memory. In G. H. Bower (Ed.), *The psychology of learning and motivation: Advances in research and theory (Volume 8)* (pp. 47–89). Academic Press. [https://doi.org/10.1016/S0079-7421\(08\)60452-1](https://doi.org/10.1016/S0079-7421(08)60452-1)
- Bahle, B., Matsukura, M., & Hollingworth, A. (2018). Contrasting gist-based and template-based guidance during real-world visual search. *Journal of Experimental Psychology: Human Perception and Performance*, 44(3), 367–386.
<https://doi.org/10.1037/xhp0000468>

- Barnes, K. A., Howard, J. H., Howard, D. V., Gilotty, L., Kenworthy, L., Gaillard, W. D., & Vaidya, C. J. (2008). Intact implicit learning of spatial context and temporal sequences in childhood autism spectrum disorder. *Neuropsychology*, 22(5), 563–570.
<https://doi.org/10.1037/0894-4105.22.5.563>
- Barr, D., Levy, R. P., Scheepers, C., & Tily, H. (2018). *Random effects structure for confirmatory hypothesis testing: Keep it maximal*. Center for Open Science.
<https://doi.org/10.31234/osf.io/39mhs>
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2014, June 23). *Fitting Linear Mixed-Effects Models using lme4*. <https://arxiv.org/pdf/1406.5823>
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67(1).
<https://doi.org/10.18637/jss.v067.i01>
- Beesley, T., Hanafi, G., Vadillo, M. A., Shanks, D. R., & Livesey, E. J. (2018). Overt attention in contextual cuing of visual search is driven by the attentional set, but not by the predictiveness of distractors. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 44, 707–721.
- Beesley, T., & Shanks, D. R. (2012). Investigating cue competition in contextual cuing of visual search. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 38(3), 709–725. <https://doi.org/10.1037/a0024885>
- Bergmann, N., Koch, D., & Schubö, A. (2019). Reward expectation facilitates context learning and attentional guidance in visual search. *Journal of Vision*, 19(3), 10.
<https://doi.org/10.1167/19.3.10>
- Bergmann, N., & Schubö, A. (2021). Local and global context repetitions in contextual cueing. *Journal of Vision*, 21(10), 9. <https://doi.org/10.1167/jov.21.10.9>

- Bergmann, N., Tünnermann, J., & Schubö, A. (2020). Reward-predicting distractor orientations support contextual cueing: Persistent effects in homogeneous distractor contexts. *Vision Research*, 171, 53–63. <https://doi.org/10.1016/j.visres.2020.03.010>
- Biderman, N., & Mudrik, L. (2018). Evidence for Implicit-But Not Unconscious-Processing of Object-Scene Relations. *Psychological Science*, 29(2), 266–277. <https://doi.org/10.1177/0956797617735745>
- Biederman, I. (2017). On the semantics of a glance at a scene. In M. Kubovy & J. R. Pomerantz (Eds.), *Psychology library editions: Volume 16. Perceptual organization* (pp. 213–253). Routledge Taylor & Francis Group.
- Biederman, I., & Ju, G. (1988). Surface versus edge-based determinants of visual recognition. *Cognitive Psychology*, 20(1), 38–64. [https://doi.org/10.1016/0010-0285\(88\)90024-2](https://doi.org/10.1016/0010-0285(88)90024-2).
- Biehl, S. C., Ehrlis, A. C., Müller, L. D., Niklaus, A., Pauli, P., & Herrmann, M. J. (2013). The impact of task relevance and degree of distraction on stimulus processing. *BMC Neuroscience*, 14, 1–11. <https://doi.org/10.1186/1471-2202-14-107>
- Brady, T. F., & Oliva, A. (2008). Statistical learning using real-world scenes: Extracting categorical regularities without conscious intent. *Psychological Science*, 19(7), 678–685. <https://doi.org/10.1111/j.1467-9280.2008.02142.x>
- Braine, M. D. S. (1963). On learning the grammatical order of words. *Psychological Review*, 70(4), 323–348.
- Brandman, T., & Peelen, M. V. (2017). Interaction between Scene and Object Processing Revealed by Human fMRI and MEG Decoding. *The Journal of Neuroscience*, 37(32), 7700–7710. <https://doi.org/10.1523/JNEUROSCI.0582-17.2017>
- Brockmole, J. R., Castelhana, 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. <https://doi.org/10.1037/0278-7393.32.4.699>

- Brockmole, J. R., Hambrick, D. Z., Windisch, D. J., & Henderson, J. M. (2008). The role of meaning in contextual cueing: Evidence from chess expertise. *Quarterly Journal of Experimental Psychology* (2006), 61(12), 1886–1896.
<https://doi.org/10.1080/17470210701781155>
- Brockmole, J. R., & Henderson, J. M. (2006a). Recognition and attention guidance during contextual cueing in real-world scenes: Evidence from eye movements. *Quarterly Journal of Experimental Psychology* (2006), 59(7), 1177–1187.
<https://doi.org/10.1080/17470210600665996>
- Brockmole, J. R., & Henderson, J. M. (2006b). Using real-world scenes as contextual cues for search. *Visual Cognition*, 13(1), 99–108. <https://doi.org/10.1080/13506280500165188>
- Brosowsky, N. P., & Crump, M. J. C. (2021). Contextual recruitment of selective attention can be updated via changes in task relevance. *Canadian Journal of Experimental Psychology*, 75(1), 19–34. <https://doi.org/10.1037/cep0000221>
- Budson, A. E., Richman, K. A., & Kensinger, E. A. (2022). Consciousness as a memory system. *Cognitive Behavioural Neurol*, 35(4), 261–297.
- Chalmers, D. J. (1997). *The Conscious Mind: In Search of a Fundamental Theory*. Oxford University Press.
- Chalmers, D. J. (2000). What is a neural correlate of consciousness? In T. Metzinger (Ed.), *Neural Correlates of Consciousness: Empirical and Conceptual Questions* (pp. 17–40). The MIT Press. <https://doi.org/10.7551/mitpress/4928.003.0004>
- Chalmers, D. J., & McQueen, K. (2022). Consciousness and the collapse of the wave function. In S. Gao (Ed.), *Consciousness and Quantum Mechanics*. Oxford University Press.
- Chao, H.-F., Hsiao, F.-S., & Huang, S.-C. (2022). Binding of features and responses in inhibition of return: The effects of task demand. *Journal of Cognition*, 5(1), 49.
<https://doi.org/10.5334/joc.247>

- Chua, K.-P., Chun, M. M., & 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, 28–71.
<https://doi.org/10.1006/cogp.1998.0681>
- Chun, M. M., & Jiang, Y. (2003). Implicit, long-term spatial contextual memory. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 29(2), 224–234.
<https://doi.org/10.1037/0278-7393.29.2.224>
- Clark, A. (2013). Whatever next? Predictive brains, situated agents, and the future of cognitive science. *Behavioral and Brain Sciences*, 36(3), 181–204.
- Cleeremans, A. (1997). Sequence learning in a dual-stimulus setting. *Psychological Research*, 60(1-2), 72–86. <https://doi.org/10.1007/BF00419681>
- Cleeremans, A., & Frith, C. (Eds.). (2003). *The unity of consciousness: Binding, integration, and dissociation*. Oxford Univ. Press. <https://doi.org/10.1093/acprof:oso/9780198508571.001.0001>
- Cleeremans, A., & Jiménez, L. (2002). Implicit learning and consciousness: A graded, dynamic perspective. In A. Cleeremans & R. M. French (Eds.), *Frontiers of cognitive science. Implicit learning and consciousness: An empirical, philosophical, and computational consensus in the making* (pp. 1–40). Psychology Press.
- Colagiuri, B., & Livesey, E. J. (2016). Contextual cuing as a form of nonconscious learning: Theoretical and empirical analysis in large and very large samples. *Psychonomic Bulletin & Review*, 23(6), 1996–2009. <https://doi.org/10.3758/s13423-016-1063-0>
- Conci, M., & Mühlénen, A. von (2011). Limitations of perceptual segmentation on contextual cueing in visual search. *Visual Cognition*, 19(2), 203–233.
<https://doi.org/10.1080/13506285.2010.518574>

- Conway, C. M., & Christiansen, M. H. (2006). Statistical learning within and between modalities: Pitting abstract against stimulus-specific representations. *Psychological Science*, 17(10), 905–912. <https://doi.org/10.1111/j.1467-9280.2006.01801.x>
- Curran, T. (2001). Implicit learning revealed by the method of opposition. *Trends in Cognitive Sciences*, 5(12), 503–504. [https://doi.org/10.1016/s1364-6613\(00\)01791-5](https://doi.org/10.1016/s1364-6613(00)01791-5)
- Damasio, A. R. (2000). *The feeling of what happens: Body and emotion in the making of consciousness*. Harcourt.
- Danion, J.-M., Meulemans, T., Kauffmann-Muller, F., & Vermaat, H. (2001). Intact implicit learning in schizophrenia. *American Journal of Psychiatry*, 158(6), 944–948. <https://doi.org/10.1176/appi.ajp.158.6.944>
- De Houwer, J., Beckers, T., & Vandorpe, S. (2005). Evidence for the role of higher order reasoning processes in cue competition and other learning phenomena. *Learning & Behavior*, 33(2), 239–249. <https://doi.org/10.3758/BF03196066>
- Dehaene, S., & Changeux, J.-P. (2011). Experimental and theoretical approaches to conscious processing. *Neuron*, 70(2), 200–227. <https://doi.org/10.1016/j.neuron.2011.03.018>
- Dehaene, S., Changeux, J.-P., Naccache, L., Sackur, J., & Sergent, C. (2006). Conscious, pre-conscious, and subliminal processing: A testable taxonomy. *Trends in Cognitive Sciences*, 10(5), 204–211. <https://doi.org/10.1016/j.tics.2006.03.007>
- Dehaene, S., Kerszberg, M., & Changeux, J. P. (1998). A neuronal model of a global workspace in effortful cognitive tasks. *Proceedings of the National Academy of Sciences*, 95(24), 14529–14534. <https://doi.org/10.1073/pnas.95.24.14529>
- Dehaene, S., & Naccache, L. (2001). Towards a cognitive neuroscience of consciousness: basic evidence and a workspace framework. *Cognition*, 79, 1–37.
- Dehaene, S., Naccache, L., Le Clec'H, G., Koechlin, E., Mueller, M., Dehaene-Lambertz, G., van de Moortele, P. F., & Le Bihan, D. (1998). Imaging unconscious semantic priming. *Nature*, 395(6702), 597–600. <https://doi.org/10.1038/26967>

- Del Viva, M. M., Punzi, G., & Shevell, S. K. (2016). Chromatic information and feature detection in fast visual analysis. *PloS One*, *11*(8). <https://doi.org/10.1371/journal.pone.0159898>
- Delorme, A., Richard, G., & Fabre-Thorpe, M. (2000). Ultra-rapid categorisation of natural scenes does not rely on colour cues: a study in monkeys and humans. *Vision Research*, *40*, 2187–2200. [https://doi.org/10.1016/S0042-6989\(00\)00083-3](https://doi.org/10.1016/S0042-6989(00)00083-3)
- Demiral, S. B., Malcolm, G. L., & Henderson, J. M. (2012). Erp correlates of spatially incongruent object identification during scene viewing: Contextual expectancy versus simultaneous processing. *Neuropsychologia*, *50*(7), 1271–1285. <https://doi.org/10.1016/j.neuropsychologia.2012.02.011>
- Dennett, D. C. (1993). *Consciousness explained*. Penguin.
- Dennett, D. C. (2014). The evolution of consciousness. In J. Toribio & A. Clark (Eds.), *Consciousness and emotion in cognitive science: Conceptual and empirical issues*. Routledge.
- Descartes, R. (1641/1985). *The philosophical writings of Descartes. Volume 2* (J. Cottingham, R. Stoothoff & D. Murdoch, Trans.) (Vol. 2). Cambridge Univ. Press.
- Descartes, R. (1901). *The method, meditations and philosophy of Descartes* (Vol. 7). Translated by John Vietch. M. Walter Dunne.
- Dienes, Z., & Scott, R. B. (2005). Measuring unconscious knowledge: Distinguishing structural knowledge and judgment knowledge. *Psychological Research*, *69*(5-6), 338–351. <https://doi.org/10.1007/s00426-004-0208-3>
- Dienes, Z., & Seth, A. K. (2010). Gambling on the unconscious: A comparison of wagering and confidence ratings as measures of awareness in an artificial grammar task. *Consciousness and Cognition*, *19*(2), 674–681. <https://doi.org/10.1016/j.concog.2009.09.009>

- Dijksterhuis, A., van Knippenberg, A., Holland, R. W., & Veling, H. (2014). Newell and Shanks' approach to psychology is a dead end. *The Behavioral and Brain Sciences*, 37(1), 25–26. <https://doi.org/10.1017/S0140525X1300068X>
- Doerig, A., Schurger, A., Hess, K., & Herzog, M. H. (2019). The unfolding argument: Why IIT and other causal structure theories cannot explain consciousness. *Consciousness and Cognition*, 72, 49–59. <https://doi.org/10.1016/j.concog.2019.04.002>
- Dreisbach, G., & Haider, H. (2008). That's what task sets are for: Shielding against irrelevant information. *Psychological Research*, 72(4), 355–361. <https://doi.org/10.1007/s00426-007-0131-5>
- Dreisbach, G., & Haider, H. (2009). How task representations guide attention: Further evidence for the shielding function of task sets. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 35(2), 477–486. <https://doi.org/10.1037/a0014647>
- Durlach, P. J., & Rescorla, R. A. (1980). Potentiation rather than overshadowing in flavor-aversion learning: An analysis in terms of within-compound associations. *Journal of Experimental Psychology: Animal Behavior Processes*, 6(2), 175–187. <https://doi.org/10.1037/0097-7403.6.2.175>
- Eberhardt, K., Esser, S., & Haider, H. (2017). Abstract feature codes: The building blocks of the implicit learning system. *Journal of Experimental Psychology: Human Perception and Performance*, 43(7), 1275–1290. <https://doi.org/10.1037/xhp0000380>
- Ehinger, K. A., & Brockmole, J. R. (2008). The role of color in visual search in real-world scenes: Evidence from contextual cuing. *Perception & Psychophysics*, 70(7), 1366–1378. <https://doi.org/10.3758/PP.70.7.1366>
- Eimer, M. (1996). The N2pc component as an indicator of attentional selectivity. *Electroencephalography and Clinical Neurophysiology*, 99(3), 225–234. [https://doi.org/10.1016/0013-4694\(96\)95711-9](https://doi.org/10.1016/0013-4694(96)95711-9)

- Eimer, M. (2014). The neural basis of attentional control in visual search. *Trends in Cognitive Sciences*, 18(10), 526–535. <https://doi.org/10.1016/j.tics.2014.05.005>
- Eimer, M., & Schlaghecken, F. (1998). Effects of masked stimuli on motor activation: Behavioral and electrophysiological evidence. *Journal of Experimental Psychology: Human Perception and Performance*, 24(6), 1737–1747. <https://doi.org/10.1037//0096-1523.24.6.1737>
- Eimer, M., & Schlaghecken, F. (2003). Response facilitation and inhibition in subliminal priming. *Biological Psychology*, 64(1-2), 7–26. [https://doi.org/10.1016/S0301-0511\(03\)00100-5](https://doi.org/10.1016/S0301-0511(03)00100-5)
- Eitam, B., Glicksohn, A., Shoval, R., Cohen, A., Schul, Y., & Hassin, R. R. (2013). Relevance-based selectivity: The case of implicit learning. *Journal of Experimental Psychology: Human Perception and Performance*, 39(6), 1508–1515. <https://doi.org/10.1037/a0033853>
- Eitam, B., Schul, Y., & Hassin, R. R. (2009). Goal relevance and artificial grammar learning. *Quarterly Journal of Experimental Psychology (2006)*, 62(2), 228–238. <https://doi.org/10.1080/17470210802479113>
- Endo, N., & Takeda, Y. (2004). Selective learning of spatial configuration and object identity in visual search. *Perception & Psychophysics*, 66(2), 293–302. <https://doi.org/10.3758/bf03194880>
- Erdelyi, M. H. (1986). Experimental indeterminacies in the dissociation paradigm of subliminal perception. *The Behavioral and Brain Sciences*, 9(1), 30–31. <https://doi.org/10.1017/S0140525X00021348>
- Eriksen, B. A., & Eriksen, C. W. (1974). Effects of noise letters upon the identification of a target letter in a nonsearch task. *Perception & Psychophysics*, 16(1), 143–149. <https://doi.org/10.3758/BF03203267>

- Eriksen, C. W. (1960). Discrimination and learning without awareness: A methodological survey and evaluation. *Psychological Review*, 67(5), 279–300.
<https://doi.org/10.1037/h0041622>
- Eriksen, C. W., & St. James, J. D. (1986). Visual attention within and around the field of focal attention: A zoom lens model. *Perception & Psychophysics*, 40(4), 225–240.
<https://doi.org/10.3758/BF03211502>
- Esser, S., Lustig, C., & Haider, H. (2022). What triggers explicit awareness in implicit sequence learning? Implications from theories of consciousness. *Psychological Research*, 86(5), 1442–1457. <https://doi.org/10.1007/s00426-021-01594-3>
- Fabre-Thorpe, M., Delorme, A., Marlot, C., & Thorpe, S. J. (2001). A limit to the speed of processing in ultra-rapid visual categorization of novel natural scenes. *Journal of Cognitive Neuroscience*, 13(2), 171–180. <https://doi.org/10.1162/089892901564234>
- Faivre, N., Dubois, J., Schwartz, N., & Mudrik, L. (2019). Imaging object-scene relations processing in visible and invisible natural scenes. *Scientific Reports*, 9(1), 4567.
<https://doi.org/10.1038/s41598-019-38654-z>
- Fiser, J., & Aslin, R. N. (2001). Unsupervised statistical learning of higher-order spatial structures from visual scenes. *Psychological Science*, 12(6), 499–504.
- Fiser, J., & Aslin, R. N. (2005). Encoding multielement scenes: Statistical learning of visual feature hierarchies. *Journal of Experimental Psychology: General*, 134(4), 521–537.
<https://doi.org/10.1037/0096-3445.134.4.521>
- Fleming, S. M., & Dolan, R. J. (2010). Effects of loss aversion on post-decision wagering: Implications for measures of awareness. *Consciousness and Cognition*, 19(1), 352–363. <https://doi.org/10.1016/j.concog.2009.11.002>
- Fodor, J. A. (1981). The Mind-Body Problem. *Scientific American*, 244(1), 114–123.
- Fodor, J. A. (1985). Précis of The Modularity of Mind. *Behavioral and Brain Sciences*, 8(01), 1. <https://doi.org/10.1017/S0140525X0001921X>

- Foti, F., Crescenzo, F. de, Vivanti, G., Menghini, D., & Vicari, S. (2015). Implicit learning in individuals with autism spectrum disorders: A meta-analysis. *Psychological Medicine*, 45(5), 897–910. <https://doi.org/10.1017/S0033291714001950>
- 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), 83–98. <https://doi.org/10.1007/s004260050015>
- Frings, C., Koch, I., Rothermund, K., Dignath, D., Giesen, C., Hommel, B., Kiesel, A., Kunde, W., Mayr, S., Moeller, B., Möller, M., Pfister, R., & Philipp, A. M. (2020). Merkmalsintegration und Abruf als wichtige Prozesse der Handlungssteuerung – eine Paradigmen-übergreifende Perspektive. *Psychologische Rundschau*, 71(1), 1–14. <https://doi.org/10.1026/0033-3042/a000423>
- Frings, C., Schneider, K. K., & Fox, E. (2015). The negative priming paradigm: An update and implications for selective attention. *Psychonomic Bulletin & Review*, 22(6), 1577–1597. <https://doi.org/10.3758/s13423-015-0841-4>
- Frost, R., Armstrong, B. C., Siegelman, N., & Christiansen, M. H. (2015). Domain generality versus modality specificity: The paradox of statistical learning. *Trends in Cognitive Sciences*, 19(3), 117–125. <https://doi.org/10.1016/j.tics.2014.12.010>
- Gams, M., & Kramar, S. (2024). Evaluating ChatGPT's consciousness and its capability to pass the Turing Test: A comprehensive analysis. *Journal of Computer and Communications*, 12(03), 219–237. <https://doi.org/10.4236/jcc.2024.123014>
- García-Pérez, M. A., & Peli, E. (2001). Luminance artifacts of cathode-ray tube displays for vision research. *Spatial Vision*, 14(2), 201–215.
- Gaschler, R., Frensch, P. A., Cohen, A., & Wenke, D. (2012). Implicit sequence learning based on instructed task set. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 38(5), 1389–1407. <https://doi.org/10.1037/a0028071>

- Gawronski, B., Hofmann, W., & Wilbur, C. J. (2006). Are "implicit" attitudes unconscious? *Consciousness and Cognition*, 15(3), 485–499. <https://doi.org/10.1016/j.con-cog.2005.11.007>
- Gennaro, R. J. (Ed.). (2004). *Higher-order theories of consciousness. An anthology*. John Benjamins Publishing Company. <https://doi.org/10.1075/aicr.56>
- Ghodrati, M., Morris, A. P., & Price, N. S. C. (2015). The (un)suitability of modern liquid crystal displays (LCDs) for vision research. *Frontiers in Psychology*, 6, 303. <https://doi.org/10.3389/fpsyg.2015.00303>
- Ghose, G. M., & Maunsell, J. (1999). Specialized representations in visual cortex: A role for binding? *Neuron*, 24, 79–85. [https://doi.org/10.1016/S0896-6273\(00\)80823-5](https://doi.org/10.1016/S0896-6273(00)80823-5)
- Gidon, A., Aru, J., & Larkum, M. E. (2022). Does brain activity cause consciousness? A thought experiment. *PLoS Biology*, 20(6), e3001651. <https://doi.org/10.1371/journal.pbio.3001651>
- Glanemann, R., Zwitserlood, P., Bölte, J., & Döbel, C. (2016). Rapid apprehension of the coherence of action scenes. *Psychonomic Bulletin & Review*, 23(5), 1566–1575. <https://doi.org/10.3758/s13423-016-1004-y>
- Goedderz, A., Rahmani Azad, Z., & Hahn, A. (2024). Awareness of implicit attitudes revisited: A meta-analysis on replications across samples and settings. *Collabra: Psychology*, 10(1), 126220. <https://doi.org/10.1525/collabra.126220>
- Goffaux, V., Jacques, C., Mouraux, A., Oliva, A., Schyns, P. G., & Rossion, B. (2005). Diagnostic colours contribute to the early stages of scene categorization: Behavioural and neurophysiological evidence. *Visual Cognition*, 12(6), 878–892. <https://doi.org/10.1080/13506280444000562>
- Golan, A., & Lamy, D. (2024). Attentional guidance by target-location probability cueing is largely inflexible, long-lasting, and distinct from inter-trial priming. *Journal of*

- Experimental Psychology: Learning, Memory, and Cognition*, 50(2), 244–265.
<https://doi.org/10.1037/xlm0001220>
- Goschke, T., & Bolte, A. (2007). Implicit learning of semantic category sequences: Response-independent acquisition of abstract sequential regularities. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 33(2), 394–406.
<https://doi.org/10.1037/0278-7393.33.2.394>
- Goschke, T., & Bolte, A. (2012). On the modularity of implicit sequence learning: Independent acquisition of spatial, symbolic, and manual sequences. *Cognitive Psychology*, 65(2), 284–320. <https://doi.org/10.1016/j.cogpsych.2012.04.002>
- Goujon, A. (2011). Categorical implicit learning in real-world scenes: Evidence from contextual cueing. *Quarterly Journal of Experimental Psychology (2006)*, 64(5), 920–941.
<https://doi.org/10.1080/17470218.2010.526231>
- Goujon, A., Brockmole, J. R., & Ehinger, K. A. (2012). How visual and semantic information influence learning in familiar contexts. *Journal of Experimental Psychology: Human Perception and Performance*, 38(5), 1315–1327. <https://doi.org/10.1037/a0028126>
- Goujon, A., Didierjean, A., & Marmèche, E. (2009). Semantic contextual cuing and visual attention. *Journal of Experimental Psychology: Human Perception and Performance*, 35(1), 50–71. <https://doi.org/10.1037/0096-1523.35.1.50>
- Goujon, A., Didierjean, A., & Thorpe, S. J. (2015). Investigating implicit statistical learning mechanisms through contextual cueing. *Trends in Cognitive Sciences*, 19(9), 524–533.
<https://doi.org/10.1016/j.tics.2015.07.009>
- Gozli, D. G. (2019). What is a task? In D. G. Gozli (Ed.), *Experimental Psychology and Human Agency* (pp. 83–111). Springer. https://doi.org/10.1007/978-3-030-20422-8_5
- Gray, N. S., & Snowden, R. J. (2005). The relevance of irrelevance to schizophrenia. *Neuroscience & Biobehavioral Reviews*, 29(6), 989–999. <https://doi.org/10.1016/j.neubio-rev.2005.01.006>

- Green, D. M., & Swets, J. A. (1966). *Signal detection theory and psychophysics*. Wiley.
- Greene, M. R. (2013). Statistics of high-level scene context. *Frontiers in Psychology*, 4, 777. <https://doi.org/10.3389/fpsyg.2013.00777>
- Haider, H., Eberhardt, K., Esser, S., & Rose, M. (2014). Implicit visual learning: How the task set modulates learning by determining the stimulus-response binding. *Consciousness and Cognition*, 26, 145–161. <https://doi.org/10.1016/j.concog.2014.03.005>
- Haider, H., Eberhardt, K., Kunde, A., & Rose, M. (2012). Implicit visual learning and the expression of learning. *Consciousness and Cognition*, 22(1), 82–98. <https://doi.org/10.1016/j.concog.2012.11.003>
- Haider, H., Eichler, A., & Lange, T. (2011). An old problem: How can we distinguish between conscious and unconscious knowledge acquired in an implicit learning task? *Consciousness and Cognition*, 20(3), 658–672. <https://doi.org/10.1016/j.concog.2010.10.021>
- Haider, H., Esser, S., & Eberhardt, K. (2020). Feature codes in implicit sequence learning: Perceived stimulus locations transfer to motor response locations. *Psychological Research*, 84(1), 192–203. <https://doi.org/10.1007/s00426-018-0980-0>
- Haider, H., & Frensch, P. A. (1996). The role of information reduction in skill acquisition. *Cognitive Psychology*, 30(3), 304–337.
- Hameroff, S., & Penrose, R. (2014). Consciousness in the universe: A review of the 'Orch OR' theory. *Physics of Life Reviews*, 11(1), 39–78. <https://doi.org/10.1016/j.plrev.2013.08.002>
- Hardcastle, V. G. (2004). HOT theories of consciousness: More sad tales of philosophical intuitions gone astray. In R. J. Gennaro (Ed.), *Higher-order theories of consciousness. An anthology* (pp. 277–294). John Benjamins Publishing Company.
- Harris, A. M., & Remington, R. W. (2017). Contextual cueing improves attentional guidance, even when guidance is supposedly optimal. *Journal of Experimental Psychology*:

- Human Perception and Performance*, 43(5), 926–940.
<https://doi.org/10.1037/xhp0000394>
- Hayes, T. R., & Henderson, J. M. (2019a). Center bias outperforms image salience but not semantics in accounting for attention during scene viewing. *Attention, Perception, & Psychophysics*. Advance online publication. <https://doi.org/10.3758/s13414-019-01849-7>
- Hayes, T. R., & Henderson, J. M. (2019b). Scene semantics involuntarily guide attention during visual search. *Psychonomic Bulletin & Review*, 26(5), 1683–1689.
<https://doi.org/10.3758/s13423-019-01642-5>
- Heerey, E. A., & Velani, H. (2010). Implicit learning of social predictions. *Journal of Experimental Social Psychology*, 46(3), 577–581. <https://doi.org/10.1016/j.jesp.2010.01.003>
- Henderson, J. M. (2007). Regarding scenes. *Current Directions in Psychological Science*, 16(4), 219–222. <https://doi.org/10.1111/j.1467-8721.2007.00507.x>
- Henderson, J. M., Brockmole, J. R., Castelhamo, M. S., & Mack, M. (2007). Visual saliency does not account for eye movements during visual search in real-world scenes. In R. P. G. van Gompel, M. H. Fischer, W. S. Murray, & R. L. Hill (Eds.), *Eye Movements: A Window on Mind and Brain*. Elsevier. <https://doi.org/10.1016/B978-008044980-7/50027-6>
- Henderson, J. M., & Hayes, T. R. (2017). Meaning-based guidance of attention in scenes as revealed by meaning maps. *Nature Human Behaviour*, 1(10), 743–747.
<https://doi.org/10.1038/s41562-017-0208-0>
- Hohwy, J. (2009). The neural correlates of consciousness: New experimental approaches needed? *Consciousness and Cognition*, 18(2), 428–438. <https://doi.org/10.1016/j.con-cog.2009.02.006>

- Holender, D. (1986). Semantic activation without conscious identification in dichotic listening, parafoveal vision, and visual masking: A survey and appraisal. *The Behavioral and Brain Sciences*, 9(1), 1–23. <https://doi.org/10.1017/S0140525X00021269>
- Holmes, N. M., Wong, F. S., Bouchekioua, Y., & Westbrook, R. F. (2022). Not "either-or" but "which-when": A review of the evidence for integration in sensory preconditioning. *Neuroscience & Biobehavioral Reviews*, 132, 1197–1204. <https://doi.org/10.1016/j.neubiorev.2021.10.032>
- Hommel, B. (1998). Event files: Evidence for automatic integration of stimulus-response episodes. *Visual Cognition*, 5(1-2), 183–216. <https://doi.org/10.1080/713756773>
- Hommel, B. (2004). Event files: Feature binding in and across perception and action. *Trends in Cognitive Sciences*, 8(11), 494–500. <https://doi.org/10.1016/j.tics.2004.08.007>
- Hommel, B. (2005). How much attention does an event file need? *Journal of Experimental Psychology: Human Perception and Performance*, 31(5), 1067–1082. <https://doi.org/10.1037/0096-1523.31.5.1067>
- Hommel, B. (2009). Action control according to TEC (theory of event coding). *Psychological Research*, 73(4), 512–526. <https://doi.org/10.1007/s00426-009-0234-2>
- Hommel, B., & Colzato, L. (2004). Visual attention and the temporal dynamics of feature integration. *Visual Cognition*, 11(4), 483–521. <https://doi.org/10.1080/13506280344000400>
- Hommel, B., Kessler, K., Schmitz, F., Gross, J., Akyürek, E. G., Shapiro, K., & Schnitzler, A. (2006). How the brain blinks: Towards a neurocognitive model of the attentional blink. *Psychological Research*, 70(6), 425–435. <https://doi.org/10.1007/s00426-005-0009-3>
- Hommel, B., Memelink, J., Zmigrod, S., & Colzato, L. S. (2014). Attentional control of the creation and retrieval of stimulus–response bindings. *Psychological Research*, 78(4), 520–538. <https://doi.org/10.1007/s00426-013-0503-y>

- Hommel, B., Müsseler, J., Aschersleben, G., & Prinz, W. (2001). The Theory of Event Coding (TEC): A framework for perception and action planning. *The Behavioral and Brain Sciences*, 24(5), 849-878. <https://doi.org/10.1017/s0140525x01000103>
- Horan, W. P., Green, M. F., Knowlton, B. J., Wynn, J. K., Mintz, J., & Nuechterlein, K. H. (2008). Impaired implicit learning in schizophrenia. *Neuropsychology*, 22(5), 606–617. <https://doi.org/10.1037/a0012602>
- Hox, J. J., Moerbeek, M., & van de Schoot, R. (2017). *Multilevel Analysis*. Routledge. <https://doi.org/10.4324/9781315650982>
- Huffman, G., Hilchey, M. D., & Pratt, J. (2018). Feature integration in basic detection and localization tasks: Insights from the attentional orienting literature. *Attention, Perception, & Psychophysics*, 80(6), 1333–1341. <https://doi.org/10.3758/s13414-018-1535-6>
- Huta, V. (2014). When to use hierarchical linear modeling. *The Quantitative Methods for Psychology*, 10(1), 13–28. <https://doi.org/10.20982/tqmp.10.1.p013>
- Itti, L., & Koch, C. (2000). A saliency-based search mechanism for overt and covert shifts of visual attention. *Vision Research*, 40(10-12), 1489–1506. [https://doi.org/10.1016/S0042-6989\(99\)00163-7](https://doi.org/10.1016/S0042-6989(99)00163-7)
- Itti, L., Koch, C., & Niebur, E. (1998). A model of saliency-based visual attention for rapid scene analysis. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 20(11), 1254–1259. <https://doi.org/10.1109/34.730558>
- Jarosz, A. F., & Wiley, J. (2014). What are the odds? A practical guide to computing and reporting bayes factors. *The Journal of Problem Solving*, 7(1). <https://doi.org/10.7771/1932-6246.1167>
- JASP Team. (2022). *JASP* (Version 0.16.3) [Computer software].
- Jiang, Y., & Chun, M. M. (2001). Selective attention modulates implicit learning. *The Quarterly Journal of Experimental Psychology: Section a*, 54(4), 1105–1124.

- Jiang, Y., & Chun, M. M. (2003). Contextual cueing: Reciprocal influences between attention and implicit learning. In L. Jiménez (Ed.), *Advances in consciousness research: Vol. 48. Attention and implicit learning* (pp. 277-196). Benjamins.
- Jiang, Y., & Leung, A. W. (2005). Implicit learning of ignored visual context. *Psychonomic Bulletin & Review*, 12(1), 100–106. <https://doi.org/10.3758/BF03196353>
- Jiang, Y., & Sisk, C. (2019). Contextual Cuing. In S. Pollmann (Ed.), *Spatial learning and attentional*. Springer Neuromethods.
- Jiménez, L. (Ed.). (2003). *Advances in consciousness research: Vol. 48. Attention and implicit learning*. Benjamins.
- Jiménez, L., & Méndez, C. (1999). Which attention is needed for implicit sequence learning? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 25(1), 236–259. <https://doi.org/10.1037/0278-7393.25.1.236>
- Jiménez, L., & Méndez, C. (2001). Implicit sequence learning with competing explicit cues. *The Quarterly Journal of Experimental Psychology: Section a*, 54(2), 345–369. <https://doi.org/10.1080/713755964>
- Jiménez, L., Méndez, C., & Cleeremans, A. (1996). Comparing direct and indirect measures of sequence learning. *Journal of Experimental Psychology. Learning, Memory, and Cognition*, 22(4), 948–969. <https://doi.org/10.1037//0278-7393.22.4.948>
- Johnson, J. S., Woodman, G. F., Braun, E., & Luck, S. J. (2007). Implicit memory influences the allocation of attention in visual cortex. *Psychonomic Bulletin & Review*, 14(5), 834–839. <https://doi.org/10.3758/bf03194108>
- Joubert, O. R., Rousselet, G., Fize, D., & Fabre-Thorpe, M. (2007). Processing scene context: Fast categorization and object interference. *Vision Research*, 47(26), 3286–3297. <https://doi.org/10.1016/j.visres.2007.09.013>
- Kabata, T., & Matsumoto, E. (2012). Cueing effects of target location probability and repetition. *Vision Research*, 73, 23–29. <https://doi.org/10.1016/j.visres.2012.09.014>

- Kaufman, M. A., & Bolles, R. C. (1981). A nonassociative aspect of overshadowing. *Bulletin of Psychonomic Society*, 18, 318–320. <https://doi.org/10.3758/BF03333639>
- Keele, S. W., Ivry, R., Mayr, U., Hazeltine, E., & Heuer, H. (2003). The cognitive and neural architecture of sequence representation. *Psychological Review*, 110(2), 316–339. <https://doi.org/10.1037/0033-295x.110.2.316>
- Kehoe, E. J., & Gormezano, I. (1980). Configuration and combination laws in conditioning with compound stimuli. *Psychological Bulletin*, 37(2), 351–378. <https://doi.org/10.1037/0033-2909.87.2.351>
- Kiefer, M., & Spitzer, M. (2000). Time course of conscious and unconscious semantic brain activation. *Cognitive Neuroscience*, 11(11), 2401–2407.
- Kiesel, A., Kunde, W., & Hoffmann, J. (2008). Mechanisms of subliminal response priming. *Advances in Cognitive Psychology*, 3(1-2), 307–315. <https://doi.org/10.2478/v10053-008-0032-1>
- Kihlstrom, J. F. (2014). Unconscious processes. In D. Reisberg (Ed.), *Oxford library of psychology. The Oxford handbook of cognitive psychology* (1. issued as an Oxford Univ. Press paperback, pp. 176–186). Oxford Univ. Press.
- Klein, C., Hohwy, J., & Bayne, T. (2020). Explanation in the science of consciousness: From the neural correlates of consciousness (NCCs) to the difference makers of consciousness (DMCs). *Philosophy and the Mind Sciences*, 1(II). <https://doi.org/10.33735/phil-misci.2020.II.60>
- Kobayashi, H., & Ogawa, H. (2020). Contextual cueing facilitation arises early in the time course of visual search: An investigation with the speed-accuracy tradeoff task. *Attention, Perception, & Psychophysics*, 82(6), 2851–2861. <https://doi.org/10.3758/s13414-020-02028-9>

- Koch, C., Massimini, M., Boly, M., & Tononi, G. (2016). Neural correlates of consciousness: Progress and problems. *Nature Reviews. Neuroscience*, 17(5), 307–321.
<https://doi.org/10.1038/nrn.2016.22>
- Koch, C., & Tononi, G. (2008). Can machines be conscious? *IEEE Spectrum*, 45(6), 55–59.
<https://doi.org/10.1109/MSPEC.2008.4531463>
- Koch, I., Blotenberg, I., Fedosejew, V., & Stephan, D. N. (2020). Implicit perceptual learning of visual-auditory modality sequences. *Acta Psychologica*, 202, 102979.
<https://doi.org/10.1016/j.actpsy.2019.102979>
- Koch, I., & Hoffmann, J. (2000). The role of stimulus-based and response-based spatial information in sequence learning. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 26(4), 863–882. <https://doi.org/10.1037/0278-7393.26.4.863>
- Konstantinidis, E., & Shanks, D. R. (2014). Don't bet on it! Wagering as a measure of awareness in decision making under uncertainty. *Journal of Experimental Psychology: General*, 143(6), 2111–2134. <https://doi.org/10.1037/a0037977>
- Kotabe, H. P., Kardan, O., & Berman, M. G. (2016). Can the High-Level Semantics of a Scene be Preserved in the Low-Level Visual Features of that Scene? A Study of Disorder and Naturalness. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 38. <https://escholarship.org/content/qt3p92837p/qt3p92837p.pdf>
- Kouider, S., & Dehaene, S. (2007). Levels of processing during non-conscious perception: A critical review of visual masking. *Philosophical Transactions of the Royal Society of London. Series B, Biological Sciences*, 362(1481), 857–875.
<https://doi.org/10.1098/rstb.2007.2093>
- Kouider, S., & Dupoux, E. (2004). Partial awareness creates the "illusion" of subliminal semantic priming. *Psychological Science*, 15(2), 75–81. <https://doi.org/10.1111/j.0963-7214.2004.01502001.x>

- 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. <https://doi.org/10.1037/0096-1523.33.4.816>
- Kunar, M. A., Flusberg, S., & Wolfe, J. M. (2006). Contextual cuing by global features. *Perception & Psychophysics*, 68(7), 1204–1216. <https://link.springer.com/content/pdf/10.3758/BF03193721.pdf>
- Kunar, M. A., Johnston, R., & Sweetman, H. (2013). A configural dominant account of contextual cueing: Configural cues are stronger than colour cues. *Quarterly Journal of Experimental Psychology*, 67(7), 1366–1382. <https://core.ac.uk/download/pdf/19210651.pdf>
- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. B. (2017). lmerTest Package: Tests in Linear Mixed Effects Models. *Journal of Statistical Software*, 82(13). <https://doi.org/10.18637/jss.v082.i13>
- Lamme, V. A. (2003). Why visual attention and awareness are different. *Trends in Cognitive Sciences*, 7(1), 12–18. [https://doi.org/10.1016/s1364-6613\(02\)00013-x](https://doi.org/10.1016/s1364-6613(02)00013-x)
- Lamme, V. A. (2006). Towards a true neural stance on consciousness. *Trends in Cognitive Sciences*, 10(11), 494–501. <https://doi.org/10.1016/j.tics.2006.09.001>
- Lamme, V. A. (2010). How neuroscience will change our view on consciousness. *Cognitive Neuroscience*, 1(3), 204–220. <https://doi.org/10.1080/17588921003731586>
- Lamme, V. A., & Roelfsema, P. R. (2000). The distinct modes of vision offered by feedforward and recurrent processing. *Trends in Neurosciences*, 23(11), 571–579. [https://doi.org/10.1016/s0166-2236\(00\)01657-x](https://doi.org/10.1016/s0166-2236(00)01657-x)
- Lamy, D., Antebi, C., Aviani, N., & Carmel, T. (2008). Priming of Pop-out provides reliable measures of target activation and distractor inhibition in selective attention. *Vision Research*, 48(1), 30–41. <https://doi.org/10.1016/j.visres.2007.10.009>

- Lamy, D., Mudrik, L., & Deouell, L. Y. (2008). Unconscious auditory information can prime visual word processing: A process-dissociation procedure study. *Consciousness and Cognition*, 17(3), 688–698. <https://doi.org/10.1016/j.concog.2007.11.001>
- Lau, E. F., Phillips, C., & Poeppel, D. (2008). A cortical network for semantics: (de)constructing the N400. *Nature Reviews. Neuroscience*, 9(12), 920–933. <https://doi.org/10.1038/nrn2532>
- Lau, H., & Passingham, R. E. (2006). Relative blindsight in normal observers and the neural correlate of visual consciousness. *Proceedings of the National Academy of Sciences*, 103(49), 18763–18768. <https://doi.org/10.1073/pnas.0607716103>
- Lau, H., & Rosenthal, D. M. (2011). Empirical support for higher-order theories of conscious awareness. *Trends in Cognitive Sciences*, 15(8), 365–373. <https://doi.org/10.1016/j.tics.2011.05.009>
- Leiner, D. J. (2024). *SoSci Survey* (Version 3.4.22) [Computer software]. <https://www.soscisurvey.de>
- Levin, Y., & Tzelgov, J. (2016). Contingency learning is not affected by conflict experience: Evidence from a task conflict-free, item-specific Stroop paradigm. *Acta Psychologica*, 164, 39–45. <https://doi.org/10.1016/j.actpsy.2015.12.009>
- Li, Q., Hill, Z., & He, B. J. (2014). Spatiotemporal dissociation of brain activity underlying subjective awareness, objective performance and confidence. *The Journal of Neuroscience*, 34(12), 4382–4395. <https://doi.org/10.1523/JNEUROSCI.1820-13.2014>
- Li, T., Tang, H., Zhu, J., & Zhang, J. H. (2019). The finer scale of consciousness: Quantum theory. *Annals of Translational Medicine*, 7(20), 585. <https://doi.org/10.21037/atm.2019.09.09>
- Logan, G. D., & Etherton, J. L. (1994). What is learned during automatization? The role of attention in constructing an instance. *Journal of Experimental Psychology: Learning*,

- Memory, and Cognition*, 20(5), 1022–1050. <https://doi.org/10.1037//0278-7393.20.5.1022>
- Luck, S. J., Vogel, E. K., & Shapiro, K. L. (1996). Word meanings can be accessed but not reported during the attentional blink. *Nature*, 383(6601), 616–618. <https://doi.org/10.1038/383616a0>
- Ludwig, D. (2023). The functions of consciousness in visual processing. *Neuroscience of Consciousness*, 2023(1), niac018. <https://doi.org/10.1093/nc/niac018>
- Luque, D., Vadillo, M. A., Lopez, F. J., Alonso, R., & Shanks, D. R. (2017). Testing the controllability of contextual cuing of visual search. *Scientific Reports*, 7, 39645. <https://doi.org/10.1038/srep39645>
- Mack, A., & Rock, I. (1998). Inattention blindness: Perception without attention. In R. D. Wright (Ed.), *Vancouver studies in cognitive science: Vol. 8. Visual attention* (pp. 55–76). Oxford Univ. Press.
- Mackintosh, N. J. (1971). An analysis of overshadowing and blocking. *Quarterly Journal of Experimental Psychology*, 23(1), 118–125. <https://doi.org/10.1080/00335557143000121>
- Mackintosh, N. J. (1976). Overshadowing and stimulus intensity. *Animal Learning & Behavior*, 4(2), 186–192. <https://doi.org/10.3758/BF03214033>
- Mackworth, N. H., & Morandi, A. J. (1967). The gaze selects informative details within pictures. *Perception & Psychophysics*, 2(11), 547–552. <https://doi.org/10.3758/BF03210264>
- Magen, H., & Cohen, A. (2007). Modularity beyond perception: Evidence from single task interference paradigms. *Cognitive Psychology*, 55(1), 1–36. <https://doi.org/10.1016/j.cogpsych.2006.09.003>

- Malcolm, G. L., Groen, I. I. A., & Baker, C. I. (2016). Making Sense of Real-World Scenes. *Trends in Cognitive Sciences*, 20(11), 843–856.
<https://doi.org/10.1016/j.tics.2016.09.003>
- Manginelli, A. A., & Pollmann, S. (2009). Misleading contextual cues: How do they affect visual search? *Psychological Research*, 73(2), 212–221.
<https://doi.org/10.1007/s00426-008-0211-1>
- Marchetti, G. (2012). Against the View that Consciousness and Attention are Fully Dissociable. *Frontiers in Psychology*, 3, 36. <https://doi.org/10.3389/fpsyg.2012.00036>
- Marois, R., & Ivanoff, J. (2005). Capacity limits of information processing in the brain. *Trends in Cognitive Sciences*, 9(6), 296–305.
<https://doi.org/10.1016/j.tics.2005.04.010>
- The MathWorks Inc. (2024). *MATLAB (R2024b)* (Version 24.2) [Computer software]. The MathWorks Inc. <https://www.mathworks.com>
- Matuschek, H., Kliegl, R., Vasishth, S., Baayen, H., & Bates, D. (2017). Balancing Type I error and power in linear mixed models. *Journal of Memory and Language*, 94, 305–315. <https://doi.org/10.1016/j.jml.2017.01.001>
- Matzel, L. D., Schachtman, T. R., & Miller, R. R. (1985). Recovery of an overshadowed association achieved by extinction of the overshadowing stimulus. *Learning and Motivation*, 16(4), 398–412. [https://doi.org/10.1016/0023-9690\(85\)90023-2](https://doi.org/10.1016/0023-9690(85)90023-2)
- Mausfeld, R. (2012). On some unwarranted tacit assumptions in cognitive neuroscience. *Frontiers in Psychology*, 3, 67. <https://doi.org/10.3389/fpsyg.2012.00067>
- Mayr, S., & Buchner, A. (2007). Negative Priming as a Memory Phenomenon. *Zeitschrift Für Psychologie / Journal of Psychology*, 215(1), 35–51. <https://doi.org/10.1027/0044-3409.215.1.35>
- 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(2), 350–364. <https://doi.org/10.1037/0278-7393.22.2.350>
- Melloni, L., Mudrik, L., Pitts, M., Bendtz, K., Ferrante, O., Gorska, U., Hirschhorn, R., Khalaf, A., Kozma, C., Lepauvre, A., Liu, L., Mazumder, D., Richter, D., Zhou, H., Blumenfeld, H., Boly, M., Chalmers, D. J., Devore, S., Fallon, F., . . . Tononi, G. (2023). An adversarial collaboration protocol for testing contrasting predictions of global neuronal workspace and integrated information theory. *PloS One*, 18(2), e0268577. <https://doi.org/10.1371/journal.pone.0268577>
- Memelink, J., & Hommel, B. (2013). Intentional weighting: A basic principle in cognitive control. *Psychological Research*, 77(3), 249–259. <https://doi.org/10.1007/s00426-012-0435-y>
- Merker, B., Williford, K., & Rudrauf, D. (2021). The integrated information theory of consciousness: A case of mistaken identity. *The Behavioral and Brain Sciences*, 45, e41. <https://doi.org/10.1017/S0140525X21000881>
- Meyen, S., Vadillo, M. A., Luxburg, U. von, & Franz, V. H. (2024). No evidence for contextual cueing beyond explicit recognition. *Psychonomic Bulletin & Review*, 31(3), 907–930. <https://doi.org/10.3758/s13423-023-02358-3>
- Michel, M. (2023a). Confidence in consciousness research. *Wiley Interdisciplinary Reviews: Cognitive Science*, 14(2), e1628. <https://doi.org/10.1002/wcs.1628>
- Michel, M. (2023b). How (not) to underestimate unconscious perception. *Mind & Language*, 38(2), 413–430. <https://doi.org/10.1111/mila.12406>
- Micher, N., Mazenko, D., & Lamy, D. (2024). Unconscious processing contaminates objective measures of conscious perception: Evidence from the liminal prime paradigm. *Journal of Cognition*, 7(1), 71. <https://doi.org/10.5334/joc.402>

- Miles, C. G., & Jenkins, H. M. (1973). Overshadowing in operant conditioning as a function of discriminability. *Learning and Motivation*, 4(1), 11–27.
[https://doi.org/10.1016/0023-9690\(73\)90036-2](https://doi.org/10.1016/0023-9690(73)90036-2)
- Miller, J. (1987). Priming is not necessary for selective-attention failures: Semantic effects of unattended, unprimed letters. *Perception & Psychophysics*, 41(5), 419–434.
<https://doi.org/10.3758/bf03203035>
- Miller, J. (2023). Outlier exclusion procedures for reaction time analysis: The cures are generally worse than the disease. *Journal of Experimental Psychology: General*, 152(11), 3189–3217. <https://doi.org/10.1037/xge0001450>
- Mitchell, C. J., Houwer, J. de, & Lovibond, P. F. (2009). The propositional nature of human associative learning. *The Behavioral and Brain Sciences*, 32(2), 183–98.
<https://doi.org/10.1017/S0140525X09000855>
- Miyahara, K., & Witkowski, O. (2019). The integrated structure of consciousness: phenomenal content, subjective attitude, and noetic complex. *Phenomenology and the Cognitive Sciences*, 18(4), 731–758. <https://doi.org/10.1007/s11097-018-9608-5>
- Moeller, B., & Pfister, R. (2022). Ideomotor learning: Time to generalize a longstanding principle. *Neuroscience & Biobehavioral Reviews*, 140, 104782.
<https://doi.org/10.1016/j.neubiorev.2022.104782>
- Mohr, J., Seyfarth, J., Lueschow, A., Weber, J. E., Wichmann, F. A., & Obermayer, K. (2016). Bois-Berlin object in scene database: Controlled photographic images for visual search experiments with quantified contextual priors. *Frontiers in Psychology*, 7, 749.
<https://doi.org/10.3389/fpsyg.2016.00749>
- Moors, P., Boelens, D., van Overwalle, J., & Wagemans, J. (2016). Scene integration without awareness: No conclusive evidence for processing scene congruency during continuous flash suppression. *Psychological Science*, 27(7), 945–956.
<https://doi.org/10.1177/0956797616642525>

- Morey, R. D., & Rouder, J. N. (2022). *BayesFactor: Computation of bayes factors for common designs. R package version 0.9.12-4.4*. <https://CRAN.R-project.org/package=BayesFactor>
- Mudrik, L., Breska, A., Lamy, D., & Deouell, L. Y. (2011). Integration without awareness: Expanding the limits of unconscious processing. *Psychological Science*, 22(6), 764–770. <https://doi.org/10.1177/0956797611408736>
- Mudrik, L., Deouell, L. Y., & Lamy, D. (2011). Scene congruency biases binocular rivalry. *Consciousness and Cognition*, 20(3), 756–767. <https://doi.org/10.1016/j.concog.2011.01.001>
- Mudrik, L., Faivre, N., & Koch, C. (2014). Information integration without awareness. *Trends in Cognitive Sciences*, 18(9), 488–496. <https://doi.org/10.1016/j.tics.2014.04.009>
- Mudrik, L., & Koch, C. (2013). Differential processing of invisible congruent and incongruent scenes: A case for unconscious integration. *Journal of Vision*, 13(13), 24. <https://doi.org/10.1167/13.13.24>
- Mudrik, L., Lamy, D., & Deouell, L. Y. (2010). Erp evidence for context congruity effects during simultaneous object-scene processing. *Neuropsychologia*, 48(2), 507–517. <https://doi.org/10.1016/j.neuropsychologia.2009.10.011>
- Mudrik, L., & Maoz, U. (2015). "Me & my brain": Exposing neuroscience's closet dualism. *Journal of Cognitive Neuroscience*, 27(2), 211–221. https://doi.org/10.1162/jocn_a_00723
- Naccache, L., & Dehaene, S. (2001). Unconscious semantic priming extends to novel unseen stimuli. *Cognition*, 80(3), 215–229. [https://doi.org/10.1016/s0010-0277\(00\)00139-6](https://doi.org/10.1016/s0010-0277(00)00139-6)
- Nagel, T. (1980). What is it like to be a bat? In N. Block (Ed.), *The language and thought series*. Harvard University Press. <https://doi.org/10.4159/harvard.9780674594623.c15>
- Navalpakkam, V., & Itti, L. (2005). Modeling the influence of task on attention. *Vision Research*, 45(2), 205–231. <https://doi.org/10.1016/j.visres.2004.07.042>

- Newell, B. R., & Shanks, D. R. (2014). Unconscious influences on decision making: A critical review. *Behavioral and Brain Sciences*, 37(1), 1–19.
- Newell, B. R., & Shanks, D. R. (2023). *Open minded: Searching for truth about the unconscious mind*. The MIT Press. <https://doi.org/146701>
- Nissen, M. J., & Bullemer, P. (1987). Attentional requirements of learning: Evidence from performance measures. *Cognitive Psychology*, 19(1), 1–32.
- Norman, D. A., & Shallice, T. (1986). Attention to action. In R. J. Davidson, G. E. Schwartz, & D. Shapiro (Eds.), *Consciousness and self-regulation: Advances in research and theory Volume 4* (pp. 1–18). Springer US. https://doi.org/10.1007/978-1-4757-0629-1_1
- Norman, E., & Price, M. C. (2012). Social intuition as a form of implicit learning: Sequences of body movements are learned less explicitly than letter sequences. *Advances in Cognitive Psychology*, 8(2), 121–131. <https://doi.org/10.2478/v10053-008-0109-x>
- Nuthmann, A., & Matthias, E. (2014). Time course of pseudoneglect in scene viewing. *Cortex; a Journal Devoted to the Study of the Nervous System and Behavior*, 52, 113–119. <https://doi.org/10.1016/j.cortex.2013.11.007>
- O'Craven, K. M., Rosen, B. R., Kwong, K. K., Treisman, A. M., & Savoy, R. L. (1997). Voluntary attention modulates fMRI activity in human MT-MST. *Neuron*, 18(4), 591–598. [https://doi.org/10.1016/S0896-6273\(00\)80300-1](https://doi.org/10.1016/S0896-6273(00)80300-1)
- Oliva, A. (2005). Gist of the Scene. In L. Itti, G. Rees, & J. K. Tsotsos (Eds.), *Neurobiology of attention* (pp. 251–256). Elsevier. <https://doi.org/10.1016/B978-012375731-9/50045-8>
- Oliva, A., & Schyns, P. G. (2000). Diagnostic colors mediate scene recognition. *Cognitive Psychology*, 41(2), 176–210. <https://doi.org/10.1006/cogp.1999.0728>
- Oliveira, C. M., Hayiou-Thomas, M. E., & Henderson, L. M. (2023). The reliability of the serial reaction time task: Meta-analysis of test-retest correlations. *Royal Society Open Science*, 10(7), 221542. <https://doi.org/10.1098/rsos.221542>

- 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.
<https://doi.org/10.1037//0278-7393.27.5.1299>
- Overgaard, M., & Overgaard, R. (2010). Neural correlates of contents and levels of consciousness. *Frontiers in Psychology*, 1, 164. <https://doi.org/10.3389/fpsyg.2010.00164>
- Paillard, J. (1991). Motor and representational framing of space. In J. Paillard (Ed.), *Brain and space* (pp. 163–182). Oxford University Press.
- Peirce, J., Gray, J. R., Simpson, S., MacAskill, M., Höchenberger, R., Sogo, H., Kastman, E., & Lindeløv, J. K. (2019). Psychopy2: Experiments in behavior made easy. *Behavior Research Methods*, 51(1), 195–203. <https://doi.org/10.3758/s13428-018-01193-y>
- Perruchet, P., & Pacton, S. (2006). Implicit learning and statistical learning: One phenomenon, two approaches. *Trends in Cognitive Sciences*, 10(5), 233–238.
<https://doi.org/10.1016/j.tics.2006.03.006>
- Perruchet, P., & Vinter, A. (1998). Learning and development: The implicit knowledge assumption reconsidered. In M. A. Stadler & P. A. Frensch (Eds.), *Handbook of implicit learning* (pp. 495–531). SAGE Publications.
- Persaud, N., & McLeod, P. (2008). Wagering demonstrates subconscious processing in a binary exclusion task. *Consciousness and Cognition*, 17(3), 565–575.
<https://doi.org/10.1016/j.concog.2007.05.003>
- Persaud, N., McLeod, P., & Cowey, A. (2007). Post-decision wagering objectively measures awareness. *Nature Neuroscience*, 10(2), 257–261. <https://doi.org/10.1038/nn1840>
- Peters, M. A. K., Kentridge, R. W., Phillips, I., & Block, N. (2017). Does unconscious perception really exist? Continuing the ASSC20 debate. *Neuroscience of Consciousness*, 2017(1). <https://doi.org/10.1093/nc/nix015>

- Peterson, M. S., & Kramer, A. F. (2001). Attentional guidance of the eyes by contextual information and abrupt onsets. *Perception & Psychophysics*, 63(7), 1239–1249.
<https://doi.org/10.3758/BF03194537>
- Pins, D., & Ffytche, D. (2003). The neural correlates of conscious vision. *Cerebral Cortex*, 13(5), 461–474. <https://doi.org/10.1093/cercor/13.5.461>
- Posner, M. I. (1980). Orienting of attention. *The Quarterly Journal of Experimental Psychology*, 32(1), 3–25. <https://doi.org/10.1080/00335558008248231>
- Pothos, E. M. (2007). Theories of artificial grammar learning. *Psychological Bulletin*, 133(2), 227–244. <https://doi.org/10.1037/0033-2909.133.2.227>
- Pournaghdali, A., & Schwartz, B. L. (2020). Continuous flash suppression: Known and unknowns. *Psychonomic Bulletin & Review*, 27(6), 1071–1103.
<https://doi.org/10.3758/s13423-020-01771-2>
- Poynton, C. A. (2012). *Digital video and HD: Algorithms and interfaces* (2nd ed.). *The Morgan Kaufmann Series in Computer Graphics*. Morgan Kaufmann.
<https://books.google.de/books?id=dSCEGFt47NkC>
- Prinz, J. (2011). Is attention necessary and sufficient for consciousness. In C. Mole, D. Smithies, & W. Wu (Eds.), *Attention: Philosophical and psychological essays* (pp. 174–203). Oxford University Press.
- Prinz, W. (1990). A common coding approach to perception and action. In O. Neumann & W. Prinz (Eds.), *Relationships between perception and action: Current approaches* (pp. 167–201). Springer Berlin Heidelberg. https://doi.org/10.1007/978-3-642-75348-0_7
- Prinz, W. (1997). Perception and action planning. *European Journal of Cognitive Psychology*, 9(2), 129–154. <https://doi.org/10.1080/713752551>
- Qualtrics. (2020). *Qualtrics* (Version 01/2024) [Computer software]. Qualtrics.
<https://www.qualtrics.com>

- R Core Team. (2021). *R: A language and environment for statistical computing* (Version 4.1.0) [Computer software]. R Foundation for Statistical Computing. <https://www.R-project.org/>
- Ramsøy, T. Z., & Overgaard, M. (2004). Introspection and subliminal perception. *Phenomenology and the Cognitive Sciences*, 3(1), 1–23.
<https://doi.org/10.1023/B:PHEN.0000041900.30172.e8>
- Reber, A. S. (1967). Implicit learning of artificial grammars. *Journal of Verbal Learning and Verbal Behavior*, 6(6), 855–863. [https://doi.org/10.1016/s0022-5371\(67\)80149-x](https://doi.org/10.1016/s0022-5371(67)80149-x)
- Reingold, E. M., & Merikle, P. M. (1988). Using direct and indirect measures to study perception without awareness. *Perception & Psychophysics*, 44(6), 563–575.
<https://doi.org/10.3758/bf03207490>
- Remillard, G. (2017). Independent learning of spatial and nonspatial sequences. *Canadian Journal of Experimental Psychology = Revue Canadienne De Psychologie Experimentale*, 71(4), 283–298. <https://doi.org/10.1037/cep0000133>
- Ren, J., Wang, M., & Conway, C. M. (2024). Can explicit instruction boost statistical learning? A meta-analytical review. *Journal of Educational Psychology*, 116(7), 1215–1237. <https://doi.org/10.1037/edu0000897>
- Rescorla, R. A., & Wagner, A. R. (1972). A theory of Pavlovian conditioning: Variations in the effectiveness of reinforcement and nonreinforcement. In *Classical conditioning: current research and theory* (Vol. II, pp. 64–99).
- Reynolds, G. S. (1961). Attention in the Pidgeon. *Journal of the Experimental Analysis of Behavior*, 4(3), 203. <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC1404062/pdf/jeabhav00196-0021.pdf>
- Reynolds, J. H., Pasternak, T., & Desimone, R. (2000). Attention increases sensitivity of V4 neurons. *Neuron*, 26, 703–714. [https://doi.org/10.1016/S0896-6273\(00\)81206-4](https://doi.org/10.1016/S0896-6273(00)81206-4)

- Rosenthal, D. M. (2004). Varieties of higher-order theory. In R. J. Gennaro (Ed.), *Advances in consciousness research. Higher-order theories of consciousness: An anthology* (Vol. 56, pp. 17–44). John Benjamins Publishing Company.
<https://doi.org/10.1075/aicr.56.04ros>
- Rosenthal, D. M. (2005). *Consciousness and Mind*. Clarendon Press.
- Rosenthal, D. M. (2008). Consciousness and its function. *Neuropsychologia*, 46, 829–840.
- Rothermund, K., Wentura, D., & Houwer, J. de (2005). Retrieval of incidental stimulus-response associations as a source of negative priming. *Journal of Experimental Psychology. Learning, Memory, and Cognition*, 31(3), 482–495. <https://doi.org/10.1037/0278-7393.31.3.482>
- Rouleau, N., & Levin, M. (2025). Brains and where else? Mapping theories of consciousness to unconventional embodiments. *PsyArXiv*. Advance online publication.
<https://doi.org/10.31234/osf.io/va5mk>
- Rousselet, G., Joubert, O. R., & Fabre-Thorpe, M. (2005). How long to get to the “gist” of real-world natural scenes? *Visual Cognition*, 12(6), 852–877.
<https://doi.org/10.1080/13506280444000553>
- Rünger, D., & Frensch, P. A. (2010). Defining consciousness in the context of incidental sequence learning: Theoretical considerations and empirical implications. *Psychological Research*, 74(2), 121–137. <https://doi.org/10.1007/s00426-008-0225-8>
- Sandberg, K., Andersen, L. M., & Overgaard, M. (2014). Using multivariate decoding to go beyond contrastive analyses in consciousness research. *Frontiers in Psychology*, 5, 1250. <https://doi.org/10.3389/fpsyg.2014.01250>
- Sandberg, K., & Overgaard, M. (2015). Using the perceptual awareness scale (PAS). In M. Overgaard (Ed.), *Behavioral methods in consciousness research* (pp. 181–195). Oxford Univ. Press.

- Schankin, A., & Schubö, A. (2009a). Cognitive processes facilitated by contextual cueing: Evidence from event-related brain potentials. *Psychophysiology*, 46(3), 668–679.
<https://doi.org/10.1111/j.1469-8986.2009.00807.x>
- Schankin, A., & Schubö, A. (2009b). The Time Course of Attentional Guidance in Contextual Cueing. In L. Paletta & J. K. Tsotsos (Eds.), *Lecture Notes in Computer Science: Vol. 5395. Attention in cognitive systems* (Vol. 5395, pp. 69–84). Springer.
https://doi.org/10.1007/978-3-642-00582-4_6
- Schankin, A., & Schubö, A. (2010). Contextual cueing effects despite spatially cued target locations. *Psychophysiology*, 47, 717–727. <https://online-library.wiley.com/doi/pdf/10.1111/j.1469-8986.2010.00979.x>
- Schiff, R., & Katan, P. (2014). Does complexity matter? Meta-analysis of learner performance in artificial grammar tasks. *Frontiers in Psychology*, 5, 1084.
<https://doi.org/10.3389/fpsyg.2014.01084>
- Schintu, S., Hadj-Bouziane, F., Dal Monte, O., Knutson, K. M., Pardini, M., Wassermann, E. M., Grafman, J., & Krueger, F. (2014). Object and space perception - is it a matter of hemisphere? *Cortex*, 57, 244–253.
<https://doi.org/10.1016/j.cortex.2014.04.009>
- Schmalbrock, P., Hommel, B., Münchau, A., Beste, C., & Frings, C. (2023). Predictability reduces event file retrieval. *Attention, Perception, & Psychophysics*, 85, 1073–1087.
<https://doi.org/10.3758/s13414-022-02637-6>
- Schmidt, J. R., & De Houwer, J. (2019). Cue competition and incidental learning: No blocking or overshadowing in the colour-word contingency learning procedure without instructions to learn. *Collabra: Psychology*, 5(1). <https://doi.org/10.1525/collabra.236>
- Schmidt, J. R., Giesen, C. G., & Rothermund, K. (2020). Contingency learning as binding? Testing an exemplar view of the colour-word contingency learning effect. *Quarterly*

- Journal of Experimental Psychology*, 73(5), 739–761.
<https://doi.org/10.1177/1747021820906397>
- Schmidt, T., & Biafora, M. (2024). A theory of visibility measures in the dissociation paradigm. *Psychonomic Bulletin & Review*, 31(1), 65–88. <https://doi.org/10.3758/s13423-023-02332-z>
- Schneider, W., & Shiffrin, R. M. (1977). Controlled and automatic human information processing: I. Detection, search, and attention. *Psychological Review*, 84(1), 1–66.
<https://doi.org/10.1037/0033-295X.84.1.1>
- Schnepf, J., & Groeben, N. (2024). The replication crisis as mere indicator of two fundamental misalignments: Methodological confirmation bias in hypothesis testing and anthropological oversimplification in theory-building. *New Ideas in Psychology*, 75, 101110.
<https://doi.org/10.1016/j.newideapsych.2024.101110>
- Schölvinck, M. L., Howarth, C., & Attwell, D. (2008). The cortical energy needed for conscious perception. *NeuroImage*, 40(4), 1460–1468. <https://doi.org/10.1016/j.neuroimage.2008.01.032>
- Schuck, N. W., Gaschler, R., & Frensch, P. A. (2012). Implicit learning of what comes when and where within a sequence: The time-course of acquiring serial position-item and item-item associations to represent serial order. *Advances in Cognitive Psychology*, 8(2), 83–97. <https://doi.org/10.2478/v10053-008-0106-0>
- Schwarb, H., & Schumacher, E. H. (2012). Generalized lessons about sequence learning from the study of the serial reaction time task. *Advances in Cognitive Psychology*, 8(2), 165–178. <https://doi.org/10.2478/v10053-008-0113-1>
- Schwarz, G. (1978). Estimating the dimension of a model. *The Annals of Statistics*, 6(2).
<https://doi.org/10.1214/aos/1176344136>
- Seitz, A. R., & Watanabe, T. (2009). The phenomenon of task-irrelevant perceptual learning. *Vision Research*, 49(21), 2604–2610. <https://doi.org/10.1016/j.visres.2009.08.003>

- Seth, A. K., & Bayne, T. (2022). Theories of consciousness. *Nature Reviews. Neuroscience*, 23(7), 439–452. <https://doi.org/10.1038/s41583-022-00587-4>
- Sewell, D. K., Colagiuri, B., & Livesey, E. J. (2018). Response time modeling reveals multiple contextual cuing mechanisms. *Psychonomic Bulletin & Review*, 25(5), 1644–1665. <https://doi.org/10.3758/s13423-017-1364-y>
- Shanks, D. R., & St. John, M. F. (1994). Characteristics of dissociable human learning systems. *Behavioral and Brain Sciences*, 17(3), 367–395. <https://doi.org/10.1017/S0140525X00035032>
- Shapiro, L. (2011). Descartes's pineal gland reconsidered. *Midwest Studies in Philosophy*, 35(1).
- Sheldrake, R. (2021). Is the sun conscious? *Journal of Consciousness Studies*, 28(3-4), 8–28.
- Shevell, S. K., & Kingdom, F. A. A. (2008). Color in complex scenes. *Annual Review of Psychology*, 59, 143–166. <https://doi.org/10.1146/annurev.psych.59.103006.093619>
- 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. <https://doi.org/10.1037/0033-295X.84.2.127>
- Shir, Y., Abudarham, N., & Mudrik, L. (2021). You won't believe what this guy is doing with the potato: The ObjAct stimulus-set depicting human actions on congruent and incongruent objects. *Behavior Research Methods*, 53(5), 1895–1909. <https://doi.org/10.3758/s13428-021-01540-6>
- Sisk, C. A., Remington, R. W., & Jiang, Y. V. (2019). Mechanisms of contextual cueing: A tutorial review. *Attention, Perception, & Psychophysics*, 81(8), 2571–2589. <https://doi.org/10.3758/s13414-019-01832-2>
- Smyth, A. C., & Shanks, D. R. (2008). Awareness in contextual cuing with extended and concurrent explicit tests. *Memory & Cognition*, 36(2), 403–415. <https://doi.org/10.3758/MC.36.2.403>

- Solms, M. (2018). The Hard Problem of Consciousness and the Free Energy Principle. *Frontiers in Psychology*, 9, 2714. <https://doi.org/10.3389/fpsyg.2018.02714>
- Solms, M., & Friston, K. (2018). How and why consciousness arises: some considerations from physics and physiology. *Journal of Consciousness Studies*, 25(5-6), 202–238.
- Sperdin, H. F., Spierer, L., Becker, R., Michel, C. M., & Landis, T. (2015). Submillisecond unmasked subliminal visual stimuli evoke electrical brain responses. *Human Brain Mapping*, 36(4), 1470–1483. <https://doi.org/10.1002/hbm.22716>
- Stein, T. (2019). The Breaking Continuous Flash Suppression Paradigm. In G. Hesselmann (Ed.), *Transitions between Consciousness and Unconsciousness* (pp. 1–38). Routledge. <https://doi.org/10.4324/9780429469688-1>
- Stenberg, G., Lindgren, M., Johansson, M., Olsson, A., & Rosén, I. (2000). Semantic processing without conscious identification: Evidence from event-related potentials. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 26(4), 973–1004. <https://doi.org/10.1037//0278-7393.26.4.973>
- Stockart, F., Schreiber, M., Amerio, P., Carmel, D., Cleeremans, A., Deouell, L., Dienes, Z., Elosegi, P., Gayet, S., Goldstein, A., Halchin, A.-M., Hesselmann, G., Kimchi, R., Lamy, D., Loued-Khenissi, L., Meyen, S., Micher, N., Pitts, M., Salomon, R., . . . Mudrik, L. (2024). Studying unconscious processing: Contention and consensus. Advance online publication. <https://doi.org/10.31234/osf.io/bkxzh>
- Storm, J. F., Klink, P. C., Aru, J., Senn, W., Goebel, R., Pigorini, A., Avanzini, P., Vanduffel, W., Roelfsema, P. R., Massimini, M., Larkum, M. E., & Pennartz, C. M. A. (2024). An integrative, multiscale view on neural theories of consciousness. *Neuron*, 112(10), 1531–1552. <https://doi.org/10.1016/j.neuron.2024.02.004>
- Strahan, E. J., Spencer, S. J., & Zanna, M. P. (2002). Subliminal priming and persuasion: Striking while the iron is hot. *Journal of Experimental Social Psychology*, 38(6), 556–568. [https://doi.org/10.1016/S0022-1031\(02\)00502-4](https://doi.org/10.1016/S0022-1031(02)00502-4)

- Stroup, W. W. (2013). *Generalized linear mixed models: Modern concepts, methods and applications. A Chapman & Hall book*. CRC Press.
- Su, W., Zhao, G., & Ma, J. (2024). Effect of cue validity on the contextual cueing effect. *Frontiers in Psychology, 15*, 1495780. <https://doi.org/10.3389/fpsyg.2024.1495780>
- Sun, R., Slusarz, P., & Terry, C. (2005). The interaction of the explicit and the implicit in skill learning: A dual-process approach. *Psychological Review, 112*(1), 159–192. <https://doi.org/10.1037/0033-295X.112.1.159>
- Tallon-Baudry, C. (2011). On the neural mechanisms subserving consciousness and attention. *Frontiers in Psychology, 2*, 397. <https://doi.org/10.3389/fpsyg.2011.00397>
- Tallon-Baudry, C., Campana, F., Park, H.-D., & Babo-Rebelo, M. (2018). The neural monitoring of visceral inputs, rather than attention, accounts for first-person perspective in conscious vision. *Cortex, 102*, 139–149. <https://doi.org/10.1016/j.cortex.2017.05.019>
- Tavera, F., Abderahaman, G., & Haider, H. (unpublished). Learning semantic and low-level cues in contextual cueing.
- Tavera, F., & Haider, H. (2025). The role of selective attention in implicit learning: Evidence for a contextual cueing effect of task-irrelevant features. *Psychological Research, 89*(1), 15. <https://doi.org/10.1007/s00426-024-02033-9>
- Tavera, F., Wilts, S., & Haider, H. (unpublished). Transferring predictions from one visual cue dimension to another in implicit learning.
- Theeuwes, J., & Lucassen, M. P. (1993). An adaptation-induced pop-out in visual search. *Vision Research, 33*(16), 2353–2357. [https://doi.org/10.1016/0042-6989\(93\)90113-b](https://doi.org/10.1016/0042-6989(93)90113-b)
- Thein, T., Westbrook, R. F., & Harris, J. A. (2008). How the associative strengths of stimuli combine in compound: Summation and overshadowing. *Journal of Experimental Psychology: Animal Behavior Processes, 34*(1), 155–166. <https://doi.org/10.1037/0097-7403.34.1.155>

- Tipper, S. P. (1985). The Negative Priming Effect: Inhibitory Priming by Ignored Objects. *The Quarterly Journal of Experimental Psychology: Section a*, 37(4), 571–590.
<https://doi.org/10.1080/14640748508400920>
- Tipper, S. P. (2001). Does negative priming reflect inhibitory mechanisms? A review and integration of conflicting views. *The Quarterly Journal of Experimental Psychology a*, 54(2), 321–343. <https://doi.org/10.1080/02724980042000183>
- Tononi, G. (2004). An information integration theory of consciousness. *BMC Neuroscience*, 5(42). <https://doi.org/10.1186/1471-2202-5-42>
- Tononi, G. (2008). Consciousness as integrated information: A provisional manifesto. *The Biological Bulletin*, 215(3), 216–242. <https://doi.org/10.2307/25470707>
- Treisman, A. (2003). Consciousness and perceptual binding. In A. Cleeremans & C. Frith (Eds.), *The unity of consciousness: Binding, integration, and dissociation* (pp. 95–113). Oxford Univ. Press. <https://doi.org/10.1093/acprof:oso/9780198508571.003.0005>
- Treisman, A., & Gelade, G. (1980). A feature-integration theory of attention. *Cognitive Psychology*, 12(1), 97–136. [https://doi.org/10.1016/0010-0285\(80\)90005-5](https://doi.org/10.1016/0010-0285(80)90005-5)
- Tseng, Y.-C., & Li, C.-S. R. (2004). Oculomotor correlates of context-guided learning in visual search. *Perception & Psychophysics*, 66(8), 1363–1378.
<https://doi.org/10.3758/BF03195004>
- Tsuchiya, N., & Koch, C. (2005). Continuous flash suppression reduces negative afterimages. *Nature Neuroscience*, 8(8), 1096–1101. <https://doi.org/10.1038/nn1500>
- Tsuchiya, N., & Koch, C. (2009). The relationship between consciousness and attention. In S. Laureys (Ed.), *The neurology of consciousness* (pp. 63–77). Academic Press.
<https://doi.org/10.1016/B978-0-12-374168-4.00006-X>

- Tsuchiya, N., Wilke, M., Frässle, S., & Lamme, V. A. (2015). No-report paradigms: Extracting the true neural correlates of consciousness. *Trends in Cognitive Sciences*, 19(12), 757–770. <https://doi.org/10.1016/j.tics.2015.10.002>
- 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. <https://doi.org/10.1037/0096-3445.134.4.552>
- Turk-Browne, N. B., Scholl, B. J., Chun, M. M., & Johnson, M. K. (2009). Neural evidence of statistical learning: Efficient detection of visual regularities without awareness. *Journal of Cognitive Neuroscience*, 21(10), 1934–1945. <https://doi.org/10.1162/jocn.2009.21131>
- Usher, M., Negro, N., Jacobson, H., & Tsuchiya, N. (2023). When philosophical nuance matters: Safeguarding consciousness research from restrictive assumptions. *Frontiers in Psychology*, 14, 1306023. <https://doi.org/10.3389/fpsyg.2023.1306023>
- Vadillo, M. A., Giménez-Fernández, T., Aivar, M. P., & Cubillas, C. P. (2020). Ignored visual context does not induce latent learning. *Psychonomic Bulletin & Review*, 27(3), 512–519. <https://doi.org/10.3758/s13423-020-01722-x>
- Vadillo, M. A., Konstantinidis, E., & Shanks, D. R. (2016). Underpowered samples, false negatives, and unconscious learning. *Psychonomic Bulletin & Review*, 23(1), 87–102. <https://doi.org/10.3758/s13423-015-0892-6>
- Vadillo, M. A., Linssen, D., Orgaz, C., Parsons, S., & Shanks, D. R. (2020). Unconscious or underpowered? Probabilistic cuing of visual attention. *Journal of Experimental Psychology: General*, 149(1), 160–181. <https://doi.org/10.1037/xge0000632>
- Vadillo, M. A., Malejka, S., Lee, D. Y. H., Dienes, Z., & Shanks, D. R. (2022). Raising awareness about measurement error in research on unconscious mental processes. *Psychonomic Bulletin & Review*, 29(1), 21–43. <https://doi.org/10.3758/s13423-021-01923-y>

- van de Sande, K. E. A., Gevers, T., & Snoek, C. G. M. (2010). Evaluating color descriptors for object and scene recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 32(9), 1582–1596. <https://doi.org/10.1109/TPAMI.2009.154>
- van Gulick, R. (2004). Higher-order global states (HOGS): An alternative higher-order model of consciousness. In R. J. Gennaro (Ed.), *Higher-order theories of consciousness. An anthology* (pp. 67–92). John Benjamins Publishing Company.
- Velmans, M. (2014). The evolution of consciousness. In D. V. Canter & D. A. Turner (Eds.), *Contemporary issues in social science. Biologising the social sciences: challenging Darwinian and neuroscience explanations*. Routledge.
<https://doi.org/10.4324/9781315685458>
- Wakefield Morys-Carter. (2021). *ScreenScale*. <https://doi.org/10.17605/OSF.IO/8FHQK>
- Weinberger, A. B., & Green, A. E. (2022). Dynamic development of intuitions and explicit knowledge during implicit learning. *Cognition*, 222, 105008.
<https://doi.org/10.1016/j.cognition.2021.105008>
- Weinfurt, K. P. (2000). Repeated measures analysis: ANOVA, MANOVA, and HLM. In L. G. Grimm & P. R. Yarnold (Eds.), *Reading and understanding MORE multivariate statistics* (pp. 317–361). American Psychological Association.
- Weitzel, L., & Bavishi, S. (2024). Disorders of consciousness. *Physical Medicine and Rehabilitation Clinics*, 35(3), 493–506. <https://doi.org/10.1016/j.pmr.2024.02.003>
- West, B. T., Welch, K. B., Galecki, A. T., & Gillespie, B. W. (2022). *Linear mixed models*. Chapman and Hall/CRC. <https://doi.org/10.1201/9781003181064>
- Whalen, P. J., Rauch, S. L., Etcoff, N. L., McInerney, S. C., Lee, M. B., & Jenike, M. A. (1998). Masked presentations of emotional facial expressions modulate amygdala activity without explicit knowledge. *The Journal of Neuroscience*, 18(1), 411–418.
<https://doi.org/10.1523/JNEUROSCI.18-01-00411.1998>

- Wheeler, M. E., & Treisman, A. M. (2002). Binding in short-term visual memory. *Journal of Experimental Psychology: General*, 131(1), 48–64. <https://doi.org/10.1037//0096-3445.131.1.48>
- Wickham, H. (2016). *ggplot2: Elegant graphics for data analysis* (Version 3.4.2) [Computer software]. Springer-Verlag New York. <https://ggplot2.tidyverse.org>
- Wickham, H., François, R., Henry, L., Müller, K., & Vaughan, D. (2023). *dplyr: A grammar of data manipulation* (Version 1.1.4) [Computer software]. <https://dplyr.tidyverse.org>
- Willenbockel, V., Sadr, J., Fiset, D., Horne, G. O., Gosselin, F., & Tanaka, J. W. (2010). Controlling low-level image properties: The SHINE toolbox. *Behavior Research Methods*, 42(3), 671–684. <https://doi.org/10.3758/BRM.42.3.671>
- Williams, J. N. (2020). The neuroscience of implicit learning. *Language Learning*, 70, 255–307. <https://doi.org/10.1111/lang.12405>
- Willingham, D. B. (1998). A neuropsychological theory of motor skill learning. *Psychological Review*, 105(3), 558–584. <https://doi.org/10.1037/0033-295x.105.3.558>
- 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. <https://doi.org/10.1037/0278-7393.15.6.1047>
- Wilts, S., & Haider, H. (2023). Concurrent visual sequence learning. *Psychological Research*, 1–15. <https://doi.org/10.1007/s00426-023-01810-2>
- Wolfe, J. M. (2020). Visual Search: How do we find what we are looking for? *Annual Review of Vision Science*, 6(1), 539–562. <https://doi.org/10.1146/annurev-vision-091718-015048>
- Yaron, I., Melloni, L., Pitts, M., & Mudrik, L. (2022). The ConTraSt database for analysing and comparing empirical studies of consciousness theories. *Nature Human Behaviour*, 6(4), 593–604. <https://doi.org/10.1038/s41562-021-01284-5>

- Zerweck, I. A., Kao, C.-S., Meyen, S., Amado, C., Eltz, M. von, Klimm, M., & Franz, V. H. (2021). Number processing outside awareness? Systematically testing sensitivities of direct and indirect measures of consciousness. *Attention, Perception, & Psychophysics*, 83(6), 2510–2529. <https://doi.org/10.3758/s13414-021-02312-2>
- Zhao, F., & Ren, Y. (2020). Revisiting contextual cueing effects: The role of perceptual processing. *Attention, Perception, & Psychophysics*, 82(4), 1695–1709. <https://doi.org/10.3758/s13414-019-01962-7>
- Zivony, A., & Eimer, M. (2022). The diachronic account of attentional selectivity. *Psychonomic Bulletin & Review*, 29, 1118–1142. <https://doi.org/10.3758/s13423-021-02023-7>
- Zizzo, D. J. (2000). Implicit learning of (boundedly) rational behaviour. *The Behavioral and Brain Sciences*, 23(5), 700–701. <https://doi.org/10.1017/S0140525X00613432>

This reference list includes not only all sources cited in the main body of the dissertation but also those cited exclusively within the three studies reproduced in Appendices A–C. To avoid redundancy and improve readability, separate reference lists are not included with the individual manuscripts in the appendix.

Publication List

The current thesis is based on three studies. One study has been published in a peer-reviewed journal. The other two studies are to be submitted into peer-review. The published version of this study, as well as the manuscripts for the other two studies, are provided below.

The following publications are included in this thesis:

Study 1

Tavera, F., & Haider, H. (2025). The role of selective attention in implicit learning: Evidence for a contextual cueing effect of task-irrelevant features. *Psychological Research*, 89(1), 15. <https://doi.org/10.1007/s00426-024-02033-9>

Authors' contribution. **Felice Tavera:** Conceptualization, methodology, formal analysis, visualization, writing (original draft, review, editing). **Hilde Haider:** Conceptualization, validation, resources, supervision, writing (review, editing).

Study 2

Tavera, F., Wilts, S., & Haider, H. (unpublished). Transferring Predictions from One Visual Cue Dimension to Another in Implicit Learning.

Authors' contributions. **Felice Tavera:** Conceptualization, Data curation, Formal analysis, Investigation, Methodology (lead), Project administration, Visualization, Writing – original draft, Writing – review and editing. **Sarah Wilts:** Conceptualization, Data curation, Investigation, Methodology (supporting), Validation, Writing – review and editing. **Hilde Haider:** Conceptualization, Supervision, Writing – review and editing

Study 3

Tavera, F., Abderahaman, G., & Haider, H. (unpublished). Implicit learning of low-level and semantic cues in an adapted contextual cueing paradigm.

Authors' contributions: **Felice Tavera:** Conceptualization, Data curation, Formal analysis, Investigation, Methodology (lead), Project administration, Visualization, Writing – original draft, Writing – review and editing. **Galla Abderahaman:** Conceptualization, Investigation, Methodology (supporting), Validation, Writing – review and editing. **Hilde Haider:** Conceptualization, Supervision, Writing – review and editing

Additional publications by the author that are not part of the present dissertation are the following:

Ruggeri, K., Stock, F., Haslam, S.A., ... **Tavera, F.**, ... Van Bavel, J. J., & Willer, R. (2024).

A synthesis of evidence for policy from behavioural science during COVID-19.

Nature, 625, 134–147. <https://doi.org/10.1038/s41586-023-06840-9>

Ruggeri, K., ..., **Tavera, F.**, ... García-Garzon, E. (2022). The globalizability of temporal

discounting. *Nature Human Behavior*. <https://doi.org/10.1038/s41562-022-01392-w>

Ruggeri, K., Berkessel, J., Bujold, P.M., Friedemann, M., Jarke, H., Steinnes, K. K.,

MacLennan, M., Quail, S. K., **Tavera, F.**, Verra, S., Naru, F., Cavassini, F., & Hardy,

E. (2021). A brief history of behavioral and decision sciences. In Ruggeri, K. (Ed.),

Psychology and Behavioral Economics: Applications for Public Policy (2nd ed.).

Routledge. <https://doi.org/10.4324/9781003181873>

Ruggeri, K., ..., **Tavera, F.**, ... Folke, T. (2020). Replicating patterns of Prospect Theory for

decision under risk. *Nature Human Behaviour*, 1-12. [https://doi.org/10.1038/s41562-](https://doi.org/10.1038/s41562-020-0886-x)

[020-0886-x](https://doi.org/10.1038/s41562-020-0886-x)

Mattes, A.*, **Tavera, F.***, Ophhey, A., Roheger, M., Gaschler, R., & Haider, H. (2020). Parallel and serial task processing in the PRP paradigm: a drift-diffusion model approach.

Psychological Research. <https://doi.org/10.1007/s00426-020-01337-w>

**shared first authorship*

Appendix A

Published Article

Title:

The role of selective attention in implicit learning: Evidence for a contextual cueing effect of task-irrelevant features

Authors:

Felice Tavera & Hilde Haider

Journal:

Psychological Research (2025)

DOI: <https://doi.org/10.1007/s00426-024-02033-9>

Copyright and Usage Notice:

This article is reproduced here with permission under the Creative Commons Attribution 4.0 International License (CC BY 4.0). This license permits use, duplication, adaptation, distribution, and reproduction in any medium or format, provided appropriate credit is given to the original authors and source, a link to the license is provided, and any changes are indicated. This dissertation includes the article in its peer-reviewed and accepted version, prior to copyediting and typesetting, in accordance with the publishing agreement.

Note:

References cited in this appendix are included in the general reference list. To enhance readability and avoid redundancy, separate reference lists are not provided for the individual appendices.

**The Role of Selective Attention in Implicit Learning:
Evidence for a Contextual Cueing Effect of Task-Irrelevant Features**

Felice Tavera¹ & Hilde Haider¹

¹University of Cologne

Author Note

Felice Tavera  <https://orcid.org/0000-0003-0994-6620>

Hilde Haider  <https://orcid.org/0000-0001-7293-3166>

We have no known conflicts of interest to disclose.

Felice Tavera was supported by the Klaus Murmann PhD scholarship awarded by the Stiftung der Deutschen Wirtschaft (German Economy Foundation).

Study materials can be obtained upon request, data and analysis scripts have been made available at the Open Science Framework (OSF) and can be accessed at

https://osf.io/3d57y/?view_only=1c5fe826a3704369b548b4c0fbd304f3

Corresponding author:

Felice Tavera

Department of Psychology

University of Cologne

Richard-Strauss-Str. 2

50931 Köln, Germany

Phone: +49-221 470-1471

Email: felice.tavera@uni-koeln.de

Abstract

With attentional mechanisms, humans select and de-select information from the environment. But does selective attention modulate implicit learning? We tested whether the implicit acquisition of contingencies between features are modulated by the task-relevance of those features. We implemented the contingencies in a novel variant of the contextual cueing paradigm. In such a visual search task, participants could use non-spatial cues to predict target location, and then had to discriminate target shapes. In Experiment 1, the predictive feature for target location was the shape of the distractors (task-relevant). In Experiment 2, the color feature of distractors (task-irrelevant) cued target location. Results showed that participants learned to predict the target location from both the task-relevant and the task-irrelevant feature. Subsequent testing did not suggest explicit knowledge of the contingencies. For the purpose of further testing the significance of task-relevance in a cue competition situation, in Experiment 3, we provided two redundantly predictive cues, shape (task-relevant) and color (task-irrelevant) simultaneously, and subsequently tested them separately. There were no observed costs of single predictive cues when compared to compound cues. The results were not indicative of overshadowing effects, on the group and individual level, or of reciprocal overshadowing. We conclude that the acquisition of contingencies occurs independently of task-relevance and discuss this finding in the framework of the event coding literature.

217 words

Keywords: Implicit learning, contextual cueing, cue competition, visual search, attention

In our daily environment, we process information about objects, their shapes, colors, locations, and so on. Thereby, we also register co-occurrences between such features. For instance, imagine your trips to the supermarket: If your favorite pasta comes in a blue package and is always in the same aisle at the supermarket, you will pick it up based on the location and color, without registering more details - you can act routinely in such environments. This kind of learning can occur without any intention to learn and usually, we are also not consciously aware about such learning processes or its contents. Therefore, it is termed implicit learning. It is an important feature of our cognitive system since it helps us to predict future events and thereby to act without effort (Clark, 2013). Another important characteristic of our system is that we learn to discriminate relevant from irrelevant information according to our action goals (Dreisbach & Haider, 2008, 2009; Haider & Frensch, 1996). If we want to buy our supermarket item, we will look for only blue packages, de-selecting other colors. This is a core ability of our attentional system and potentially shapes what we learn from our environment in such everyday actions. The goal of the current study is to ask for the role of selective attention of cues, here manipulated through their task-relevance, in implicit learning processes.

In the field of implicit learning, there has been a long-standing debate about the conditions that are required for such learning processes. When do we notice that certain features of stimuli are co-occurring in a systematic fashion? Do they need to be part of the current action goal or, more broadly, the task-set? Given a confined task context, do features need to be task- or response-relevant to be associatively learned? Or do we encode all the information about all the stimuli of the task at hand in a rather unselective manner and learning occurs automatically whenever the prediction error minimizes due to contingencies inherent in the environment?

Implicit Learning

In the lab, we can study implicit learning processes in several different paradigms, like the serial reaction time task (Nissen & Bullemer, 1987), statistical learning paradigms (Fiser &

Aslin, 2001; Reber, 1967), or in contextual cueing paradigms (Chun & Jiang, 1998), to only name a few. The research questions studied with these paradigms are rather similar, yet, research within the different paradigms is usually only loosely connected. Here, we focus mostly on the contextual cueing literature, but integrate also some findings from the other paradigms.

In the original contextual cueing paradigm, participants are instructed to do a visual search task and are asked to find a target letter “T” among a display of distractor letters “L”. For each block throughout an experiment, half of the displays are repeated distractor configurations that consistently predict a target location while the other half of displays are novel configurations. In each trial, participants are asked to report the orientation of the target letter. The contextual cueing effect (CC effect) is defined as a stronger decrease (steeper slope) in response time (RT) for the repeated configurations than for the novel configurations over the course of trials. Note that the configurations are not associated with the orientation of the target, and thus only the contingency between the distractor configuration and the target location can be learned, while the response remains unpredictable. The effect can be traced back to an enhanced efficiency in search, attentional guidance and selection, and, to a lesser extent, to response-related processes (Kobayashi & Ogawa, 2020; Kunar et al., 2007; Schankin & Schubö, 2009a, 2010; Sisk et al., 2019). It results in long-term implicit learning effects (Chun & Jiang, 2003). When asked to explicitly discriminate repeated spatial configurations from novel ones, participants are typically not able to do so, and they do not report having learned anything. Therefore, this learning process is assumed to be implicit (Colagiuri & Livesey, 2016; but see Vadillo, Linssen, et al., 2020).

The classical buildup of the contextual cueing task does not seem ideal to study our research question. Because originally, it emphasizes the spatial dimension above all else. When studying the question of the role of task-relevance in implicit learning, we want to compare different cues when they are task-relevant or irrelevant. In the classical contextual cueing, the

comparison between different cues would be inherently disbalanced: The predictive feature is the spatial configuration of the distractors, the visual search task is a spatial task, and the requested response is based on the spatial orientation judgement of the target.

Meanwhile, the contextual cueing paradigm has been used in different ways that suggest the possibility of reducing the dominance of the spatial dimension in the task. The empirical evidence supports that, within the paradigm, cues or contexts besides the spatial configuration of distractors are learned, and can guide attention. In the visual domain, multiple studies have shown a CC effect when repeating natural scenes or complex geometric patterns that predict target location, though, involving explicit learning (Brockmole et al., 2006; Brockmole & Henderson, 2006b; Ehinger & Brockmole, 2008; Goujon et al., 2012). With more simplistic stimulus material, it has been shown that background color and distractor identity can be implicitly learned to predict the target position. However, when color or shape cues in such form are predictive for target location on top of spatial cues (distractor configuration) being predictive, only spatial cue contingencies are learned, color and shape contingencies are overshadowed (Endo & Takeda, 2004; Kunar et al., 2006; Kunar et al., 2013). It has further been shown that spatiotemporal sequences can guide attention (Olson & Chun, 2001), illustrating the wide scope of environmental cues that the cognitive system uses for predictions. So, it seems that a number of features can be used as cues, and probably entirely task-irrelevant features like background color can be learned to predict the target position. Yet, the role of selective attention that might discriminate task-relevant from task-irrelevant stimuli or features, remains unclear in the field of implicit learning.

Attentional prerequisites for implicit learning

As a cautionary disclaimer: Attention is a widely used and too often under-defined term (Anderson, 2011). Here, we refer to attention as selective attention, not attention as a resource (as in, e.g., Frensch et al., 1998; Nissen & Bullemer, 1987). In the studies we will review here,

attention is also mostly operationalized as task-relevance. So, when a stimulus feature is task-relevant, it is considered to be attended, and is consequently integrated into the learning process. This is to be seen separately from the question if the feature is processed consciously or not. In many ways, consciousness and attention are closely related notions (Jiang & Chun, 2003; Mack & Rock, 1998; Tsuchiya & Koch, 2009). It is crucial in the definition of attention to avoid regressive reasoning in the form of invoking a homunculus that fulfils all assumed functions of attention, and is a causal, but unexplained factor in the cognitive system. Therefore, attention in our context is to be understood as the resulting effect when manipulating task-relevance, not as a causal factor on its own. A test for conscious knowledge of the learned contents must be an additional step and is not assumed to perfectly correlate with attending to the to-be-learned features (Tsuchiya & Koch, 2009).

There are two lines of argument with opposing predictions when it comes to attentional prerequisites of implicit learning. The first suggests that task-irrelevant features are not processed in a way that allows for integration into the learning process, either arguing that the features are not processed sufficiently, or that their representational strength is too weak to translate into behavior (Turk-Browne et al., 2005). The second argument suggests that task-irrelevant features are indeed processed to a degree that they can become part of contingencies which then form predictions (Kunar et al., 2013; Miller, 1987).

As to the first line of argument, there are studies that could demonstrate a learning effect only for relevant features. In visual search and also in statistical learning paradigms, participants were instructed to only pay attention to stimuli of one color, and to ignore stimuli of another color (Jiang & Chun, 2001; Jiang & Leung, 2005; Turk-Browne et al., 2005). Because learning of contingencies occurred for the attended color stimuli only, it was concluded that selective attention is a prerequisite for (implicit) learning. Similarly, participants were able to learn a spatial sequence of stimuli, but only additionally learned the contingencies with the identity of

these stimuli when they were instructed to count them (Jiménez & Méndez, 1999; Jiménez et al., 1993). Thus, only when the identity of the stimuli were made task- or response-relevant, they were learned (see also Dreisbach & Haider, 2008, 2009). Yet, Jiang and Leung (2005b) observed that contingencies in stimuli of an unattended color could be learned in some way, because even though learning did not manifest in behavior at first, it facilitated learning in a subsequent task. In a similar vein, the above mentioned results from Jiang and Chun (2001) cannot be interpreted unambiguously. In their third, higher-powered experiment, they found potential evidence for learning of contingencies also in a task-irrelevant color.

The second group of findings indicate that irrelevant information is also processed and respective learning contents used in future instances. For example, Miller (1987) used a variant of the Eriksen flanker task (B. A. Eriksen & Eriksen, 1974). In his experiments, the flankers were not, like originally done in this paradigm, of the same identity as the targets or were otherwise associated with a response. He observed that when these task-irrelevant flankers were associated consistently with a specific response, participants responded faster in these trials compared to when the flanker-response relation was changed. Hence, the irrelevant flankers were associated with the particular response. Similarly, Kunar et al. (2006, 2013) showed that in contextual cueing, task-irrelevant context features such as background color or texture were learned when they were predictive for target location.

An additional finding, however, is that context features like color, texture, or distractor identity are not learned when a spatial configuration is given as an additional cue (Endo & Takeda, 2004; Kunar et al., 2013). This suggests that the spatial configuration could overshadow the learning of other predictive features. This may not be surprising, because, as mentioned above, the task in contextual cueing paradigms inherently emphasizes the spatial dimension. In addition, in the literature on implicit sequence learning, for example, Koch and Hoffmann (2000) suggested that spatial relations of stimuli contributed significantly more to learning

effects than other stimulus features. But also generally, the spatial dimension might be distinctly represented in our cognitive system (U. Mayr, 1996; Paillard, 1991; Schintu et al., 2014). In fact, the spatial dimension might not even be a perceptual feature as such, as it is so tightly bound to the motor system (Gaschler et al., 2012; Goschke & Bolte, 2012; I. Koch & Hoffmann, 2000; Paillard, 1991).

With respect to findings on attentional mechanisms and learning specifically in the contextual cueing paradigm, these results suggest that their generalizability is strongly limited. The paradigm has, with very few exceptions (Endo & Takeda, 2004; Kunar et al., 2013), not been extended to test other, non-spatial stimulus features. This is particularly a problem when trying to draw conclusions about the learning of task-relevant and task-irrelevant features. Because either spatial features are overshadowing all other visual features (Kunar et al., 2013) because they are weighted more strongly according to the task requirements, or the spatial dimension is represented entirely differently, and thus shows different learning mechanisms than other visual features. Therefore, to conduct a more generalizable test on attentional mechanisms in implicit learning, we designed a novel variant of the task that de-emphasizes the spatial dimension. With this variant, we can contrast the learning of different visual features (color, shape) that are not problematic in terms of the task requirements, or, potentially, their general representation in the cognitive system.

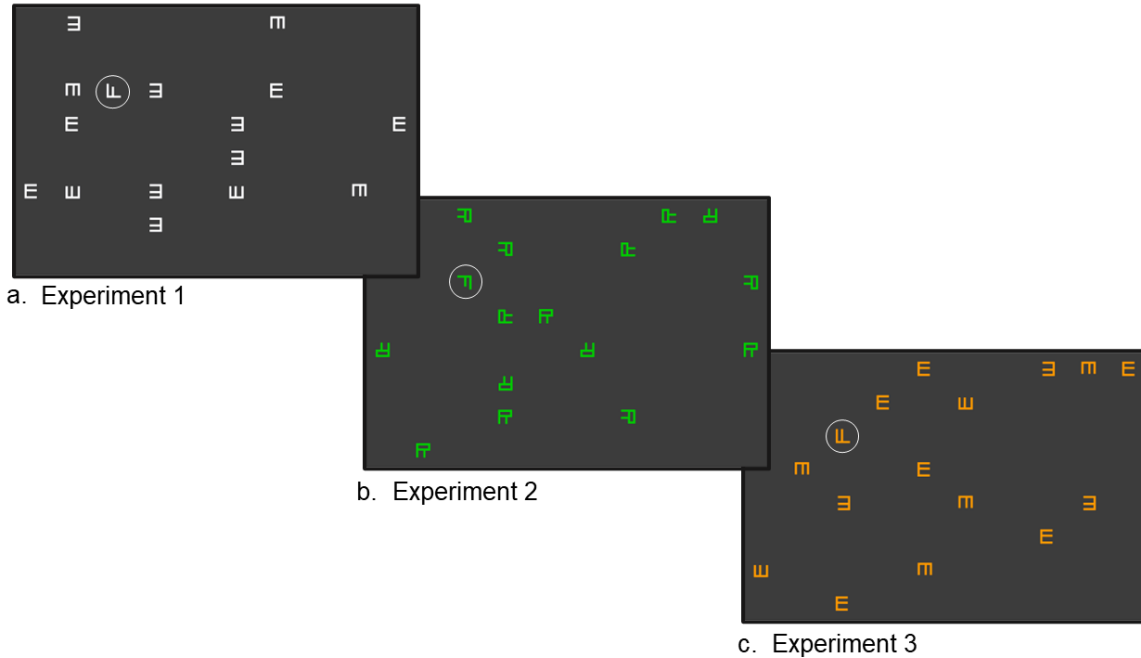
A second point noteworthy in the studies reviewed so far, is that the participants were not able to recognize predictive distractor configurations (Jiang & Chun, 2001; Kunar et al., 2006; Kunar et al., 2013) or recall the identity of flankers (Miller, 1987). This was respectively taken as evidence for incidental or implicit learning. However, Vadillo et al (2019) recently questioned the implicit nature of the CC effect, given non-sensitive awareness measures and issues with limited statistical power of many studies in the literature. We will address this with

a carefully designed test for conscious awareness, and discuss the issue in light of our results further in the General Discussion.

Overview of the study

The main goal of the current study was to examine whether selective attending, in terms of task-relevance, is needed to learn the contingencies between distractor features and target location within a contextual cueing paradigm. Importantly, whereas in the original contextual cueing paradigm, the spatial configuration of distractors is the cue for the target location, we implemented nonspatial features of the distractors as cues. We manipulated task-relevance of the predictive cue as following: The shape dimension is task-relevant because the task is to assess the target's shape (i.e., identity), and thus, the distractor shapes, which are the predictive cues, would be relevant and needed to be processed for the processing of the task. The color dimension, on the other hand, does not appear in any of the task's processes, neither the search process discriminating distractor from target shapes, nor for the response, that is referred to the distractor shape. The color dimension is thus considered task-irrelevant. A second question concerned cue competition. If more than one feature predicts the target location, will that lead to overshadowing of the task-irrelevant feature, as Kunar et al. (2013) have shown for spatial features? Or are such cue competition effects as overshadowing or blocking the result of explicit or deliberate processes, and thus do not occur in incidental learning paradigms such as our variant of contextual cueing (De Houwer et al., 2005; J. R. Schmidt & De Houwer, 2019)? A third question concerned the implicit nature of the acquired contingencies.

In all three experiments, the participants saw spatial configurations of distractors and had to find the target to answer whether a certain characteristic of the target was present. The spatial configurations of the distractors were novel in every trial and did not predict target location. Instead, either the shape (Experiment 1), the color (Experiment 2), or the color and shape (Experiment 3) of the distractors were predictive for the target location.

Figure 1*Search Displays in the Training Phase of the Experiments*

Note. Exemplary search displays for Experiments 1-3. The target letter “F” is shown in one of four potential target locations. It is circled only for illustration purposes, and was not highlighted in that way in the experiments. a. An exemplary search display with the E-shaped distractors for Experiment 1 with shape cues. The target letter F has a shorter second horizontal bar. b. A search display, exemplary green, for color as cue in Experiment 2. c. A search display with orange E-shapes with the compound cue of color and shape, exemplary for Experiment 3. The target letter F has equally long horizontal bars.

In Experiment 1, three of six distractor shapes each cued one of four potential target locations whereas for the other three shapes, targets were randomly assigned to the four potential target locations. Note that in this context, shape is a task-relevant cue in so far that it needs to be processed to discriminate the target from the distractor. In Experiment 2, we used the distractor color as a feature to cue the target location. Again, three colors each cued one particular target location, and the other three colors were randomly paired with the target locations.

The question here was if the predictive color would be learned as a cue for target location, even though it is neither task- nor response-relevant. To be more precise, color is neither relevant to the search task, as it does not distinguish target and distractors, nor relevant to the response as each color is equally likely to appear with each target identity and there is no color judgement required. In Experiment 3, we tested whether two distinct features of the distractors, shape (task-relevant) and color (task-irrelevant) would be learned to be associated with a target location as a compound, whether both features would be learned independently from one another, or if only one feature would be learned (overshadowing).

General Method

Stimuli. The search displays were 15x10cm in size, irrespective of the screen size of participant's computer monitors (see Procedure). The displays were constructed following the method of Bergmann et al. (2019), with minor variations, as described in the following. The displays consisted of 15 distractor letters on a dark grey background (RGB 60, 60, 60) rotated randomly by 0°, 90°, 180°, or 270°. They were organized in a 7x10 (invisible) grid, ensuring equal distance between adjacent stimuli, and distributed equally between the two horizontal halves of the display (see Figure 1). In Experiment 1, all 15 distractor letters of one search display were white (RGB 255, 255, 255) and shaped as one of the six stylized letters A, E, K, P, S, and W. In Experiment 2, all 15 distractor letters of one search display were R-shaped and were colored in the six colors green (RGB 1, 204, 0), orange (RGB 254, 153, 0), blue (RGB 0, 0, 254), red (RGB 254, 0, 0), pink (RGB 255, 0, 254), and cyan (RGB 1, 255, 255). In Experiment 3, the 15 distractor letters were distinctly shaped and colored, with each color matched to one shape (for example, S shaped distractors were always colored in pink, and so on). Colors were those of Experiment 2 and shapes those of Experiment 1. The color-shape matching was permuted across participants.

In one of four possible target locations (see Figure 2b), there was either an F with equally long horizontal bars, or an F with a shorter second horizontal bar as target letter (see Figure 1). All targets were randomly rotated to the right (90°) or the left (270°).

For half of the cues, contingencies between distractor characteristic and target location were fix. That is, three of the distractors' colors or shapes or shape/color combinations were always paired with one respective target locations: For example, for Experiment 1, three shapes were each 100% contingent with a target location. The shape-target location matching was permuted across participants. For the other three letters, the four target locations were equally likely. In Experiment 2, the same applies for color-target location matchings.

Procedure. All three experiments were conducted in accordance with the Declaration of Helsinki. Participants were recruited online via Prolific and were reimbursed according to Prolific's "ethical reward" standards. They were redirected to Pavlovia, where the experiments were uploaded from PsychoPy2 (version 2020.2.4; Peirce et al., 2019) and adapted to the JavaScript environment. Participants were informed about the procedure of the experiment and asked to give their informed consent.

Participants were first asked to follow a screen scaling procedure (Wakefield Morys-Carter, 2021). With the arrow keys on their keyboard, they were asked to adjust an image of a credit card on the screen to the size of an actual bank or credit card. This procedure ensured equal size of the search displays for every participant irrespective of the monitor size or aspect ratio.

All three experiments consisted of three main parts: A short practicing phase, a training phase, and lastly, a generation task. In the first 12 practice trials, they were shown displays with white, L-shaped distractor letters to get used to the task for the training. In each trial of the training (see Figure 2a), a fixation cross was presented for 500ms. Then, the search display

appeared for a maximum of 3,000ms or until the response. The response window started with the appearance of the search display and lasted 4,000ms. Participants were instructed to search for the target letter “F” among distractor letters, and to identify if the second bar was short or long, respectively pressing the “S” or “L” key on their keyboard with their index fingers as quickly and accurately as possible. The trial ended with a feedback text (“correct” or “incorrect”) that appeared on the screen for 600ms and was followed by a blank inter-trial-interval of 500ms. The training consisted of 15 blocks of 48 trials each. Participants were given the opportunity to take a short self-paced break after every block.

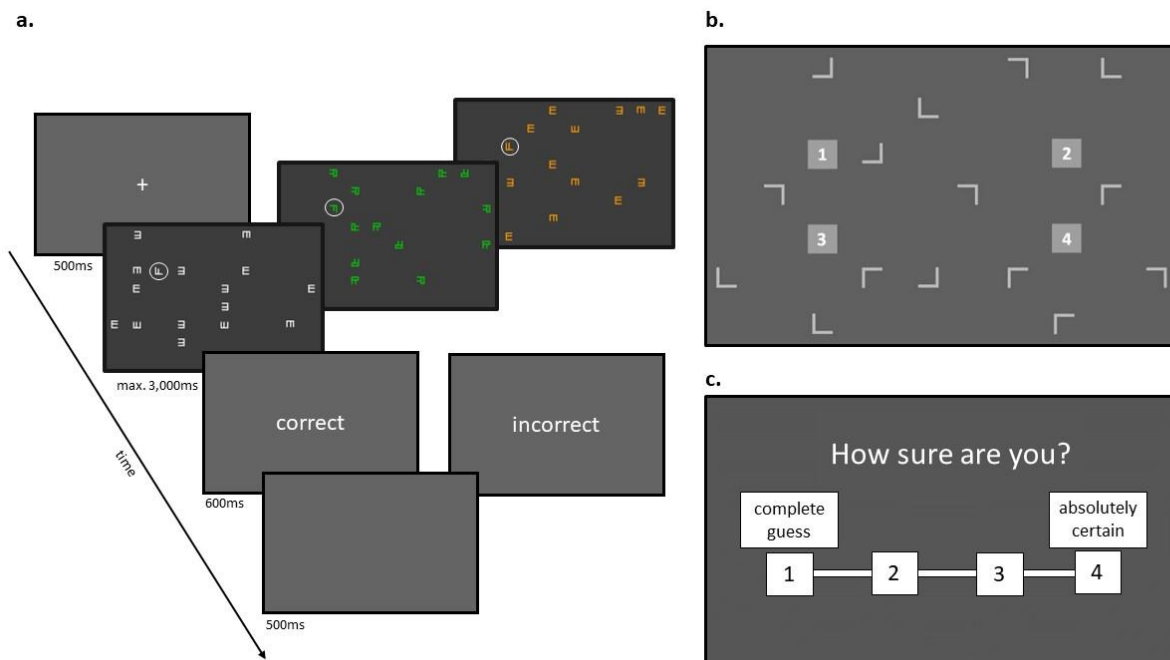
After training, a so-called generation task (Chun & Jiang, 2003) started to assess participants’ awareness about the cue-target location contingency. The generation task contained one block of 48 trials. It was designed such that participants were provided with a similar retrieval context as in the learning environment, as participants were shown the search displays of the training phase. This similarity of the environment and task during training and test phases provide similar sensitivities of both tests, thus increasing the chance to detect potential conscious knowledge (Shanks & St. John, 1994). In the test phase, participants were presented with search displays of the training phase, just that there was no target letter, but instead, the four potential target locations were marked with the numbers 1-4 (see Figure 2b). Participants were instructed to indicate in which target location they think the target letter was presented using the number keys on their keyboard. Afterwards, a visual scale from 1 (labeled “complete guess”) to 4 (labeled “absolutely certain”) appeared on the screen (see Figure 2c), and participants were asked to indicate their confidence with their generation response, again using the number keys on their keyboard.

To finish the study, participants lastly were redirected to Qualtrics (Qualtrics, 2020) or SoSci Survey (Leiner, 2024) to respond to some questions about the experiment. They were asked to report technical issues, their ideas on the purpose of the study or if they noticed

anything, if and why the task became more difficult or easier, and if they had noticed any regularities or contingencies.

Figure 2.

Procedure for Experiments 1-3



Note. a. The structure of a training trial. b. An example of a display in the generation task in which the four possible target locations are marked with the numbers 1–4. c. The 4-level confidence scale.

General data analysis. The statistical analysis was conducted in R Statistical Software (version 4.1.0; R Core Team, 2021). We used the dplyr package for most data manipulation (Wickham et al., 2023), the lme4 package for fitting models (REML; Bates et al., 2014) with restricted maximum likelihood (REML) model fit, and the lmerTest package (Kuznetsova et al., 2017) with the Satterthwaite's method for *t*-tests. Note that for the χ^2 tests for model

comparisons, the models are refitted using maximum likelihood (ML). Graphs were created with the ggplot2 package (Wickham, 2016). The study’s design and analysis were not pre-registered. Data and analysis scripts are available on OSF.

For the training, we excluded incorrect trials, and trials with the one target location out of the four that appeared less frequently. Because there were three cues, that is, three colors, shapes, or color/shape compounds, matched to one target location, and the other three colors or shapes were equally often paired with all four target locations, one target location is consequently never predicted by a cue, and is also less frequent than the other three target locations². In trials with the less frequent target location, we would obtain significantly longer RTs, because of the common probability cueing effect (Golan & Lamy, 2024). This is why, for the RT analysis, we excluded trials with the less frequent target location, and only compared predictive and unpredictable trials for the three equally frequent target locations. Further, to account for the intertrial priming effect (Golan & Lamy, 2024; Kabata & Matsumoto, 2012), that is, shorter RTs for trials in which the target location is repeated from trial $n - 1$, we excluded such target location repetition trials as well. Because of those necessities to exclude trials based on the design, we decided not to exclude any more trials based on outlier analysis. This is also in line with recent analyses that outlier exclusion procedures for RT analyses might add biases and power issues, and thus do more harm than good (Miller, 2023).

For the RT analyses, we fitted a mixed-effects model for two reasons. First, the method does not require aggregating data from multiple blocks into epochs, and thus less data is summarized (for a similar analysis, see e.g., Bergmann et al., 2019). Second, with a mixed-effects model, we are able to account for the repeated-measures design more efficiently by including

² We worked with four target locations following the task set-up and materials of Bergmann et al. (2019; 2020), and worked with three predictive and three unpredictable colors/shapes, as we speculated that only two predictive colors/shapes would be too easy to learn and potentially result in explicit knowledge, and four predictive colors/shapes might have been too difficult to learn, as we would have had to present eight colors/shapes in total. The scope of learning with respect to the number of predictive cues in this paradigm is something for future research to determine.

subject as random effect (Huta, 2014; Weinfurt, 2000). It should be noted that conducting power analyses with mixed-effects models is a challenge due to the complexity of parameter and variance estimation, particularly with unknown random effect structures and their interactions (West et al., 2022). Because we had no prior data from our paradigm to obtain such estimates, we chose to refrain from conducting a power analysis. Yet, the sample sizes and number of observations in all three experiments are larger than the recommended minimum for mixed-effects model analyses (Hox et al., 2017).

We selected a model with two fixed effects: context (predictive or unpredictable color/shape/color-shape compound) and block (as time variable) as factors. Context as a dichotomous factor was coded with contrasts -0.5 and 0.5, and the block factor was coded with block-1 for better interpretability. Our fixed effects were deduced from theoretical considerations, and on top of that, tested in model fits, but we decided on random effects solely based on the data. The argument here is that, on the one hand, it has been argued that maximal models are best for keeping the Type I error low while at the same time not significantly decreasing statistical power (Barr et al., 2018). On the other hand, however, simulations have shown that the statistical power to detect significant fixed effects can in fact be increased when opting for a random effect structure that fits the data better, as compared to implementing the full model (and more so for complex models; Matuschek et al., 2017; for a similar assessment of model selection in repeated-measures designs see also Stroup, 2013). To balance the Type I error rates and statistical power, Matuschek et al. (2017) suggest to select a model based on a selection criterion such as the Akaike information criterion (AIC; Akaike, 1998) or the Bayes information criterion (BIC; Schwarz, 1978) that assess goodness-of-fit. We decided to follow this line of argument to identify the most parsimonious model while balancing the Type I error and power. Thus, we compared various potential random effect structures based on the AIC, and additionally provide χ^2 significance tests of log-likelihood (-2LL) changes from nested models. The results from all

potential random effect structure models and model comparisons are accessible via the analysis script uploaded to OSF.

Following up on results from frequentist t -tests, we report Bayes factors that additionally indicate the strength of evidence for the null or alternative hypothesis. We computed Bayes factors (BF) using JASP (JASP Team, 2022) with the JASP default priors for t -tests (following Morey & Rouder, 2022; Cauchy distribution with a width of $r = .707$). The semantic labels for BF interpretation are taken from Jarosz and Wiley (2014).

For the generation trials, we tested the objective performance (target placement) against chance level (25% because of four response alternatives). Because one of the four target locations is far less frequent than the other three, one could even postulate, that chance level is rather 33%, as if it was choosing between the three more frequent target locations. Still, we opted for the more conservative approach to test against 25%. For the analysis, we selected only the predictive shapes and colors, because (implicit) knowledge about cue and target location associations could only be acquired for those. To assess conscious awareness of the learning contents, we followed the “consciousness-selectivity” argument of Michel (2023a), proposing that differences in consciousness lead to differences in metacognitive efficiency, distinguishing correct from incorrect responses (for a similar method see Persaud & McLeod, 2008). To assess the relationship between the objective performance measure and the subjective confidence measure, we computed variables for relative frequencies of correct response under the condition of high confidence (correct|high) and low confidence (correct|low). The logic here is that participants with explicit knowledge of a pairing of cue and target location should be able to make a metacognitive assessment of their knowledge (Haider et al., 2011). Thus, when knowing a pairing explicitly, their response should be correct, and their confidence should be high. Note that we summarized confidence ratings of 1 and 2 as low, and ratings of 3 and 4 as high confidence,

taking a more conservative approach that considers individual response biases and regression to the mean (both would potentially result in an avoidance of the extreme scale values).

Experiment 1

In the first experiment, as a way to introduce our variant of the contextual cueing paradigm, we were using distractor identity features (i.e., shape) as cues, instead of the original spatial configuration context cue. We tested whether different shapes can be learned to predict target location, and thus facilitate and speed visual search processes. First, we constructed the task so that three of the six possible shapes of distractors were 100% predictive of target location.

Method

Participants. 30 participants were recruited via *Prolific* (15 female, 1 diverse; $M_{\text{age}}=41.37$; $SD_{\text{age}}=13.50$). Participants were prescreened for living in the UK (participation took place during daytime for all participants), being fluent in English, have normal or corrected-to-normal vision, and had not taken part in a previous contextual cueing experiment of our lab.

Results

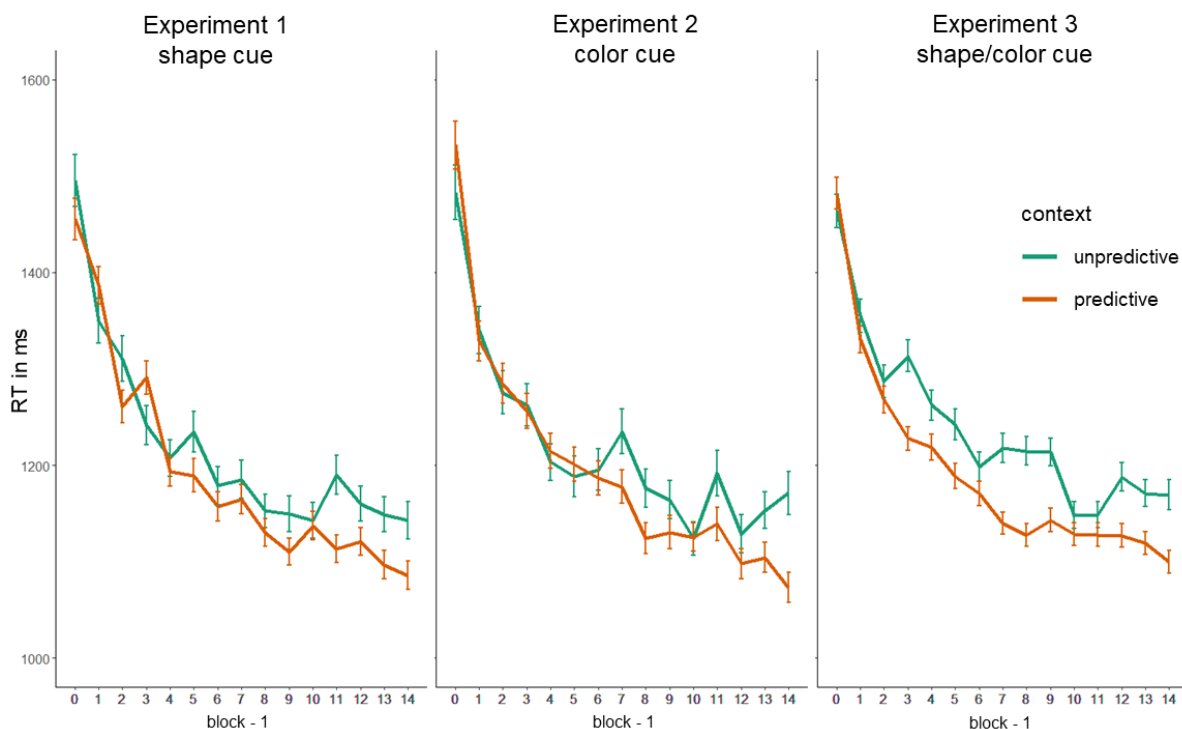
Training. The removal of incorrect responses, repeated target location trials, and trials with the less frequent target location resulted in a 21.59% trimming. Mean accuracy was 95.51% ($SD=0.21$), mean RT for the cleaned data set was 1193.73ms ($SD=436.37$). RTs over the course of the blocks, separated by predictive and unpredictive context, are shown in Figure 3.

For RT as dependent variable, we first tested whether the fixed effect structure hypothesizing an interaction of context and block was the best fit for the data. A comparison of AICs of models with no random effect structure and no fixed effect (AIC=254040.9), only block as

fixed effect (AIC=253372.1), an additive (AIC=253359.2), and an interaction effect (AIC=253354.7), confirmed that the interaction term model indeed yielded the best fit. In a next step, we compared random effect structures. Allowing for random slopes for context across participants (AIC=250835) yielded a better fit than only random intercept for participants (AIC=250864; $\chi^2(2)=32.287$, $p<.001$). Other and more complex random effect structure models did not converge or produced a singular fit. In the random slopes and random intercept model, the fixed effects context and block, as well as their interaction, were significant (see Table 1 and Figure 4).

Figure 3.

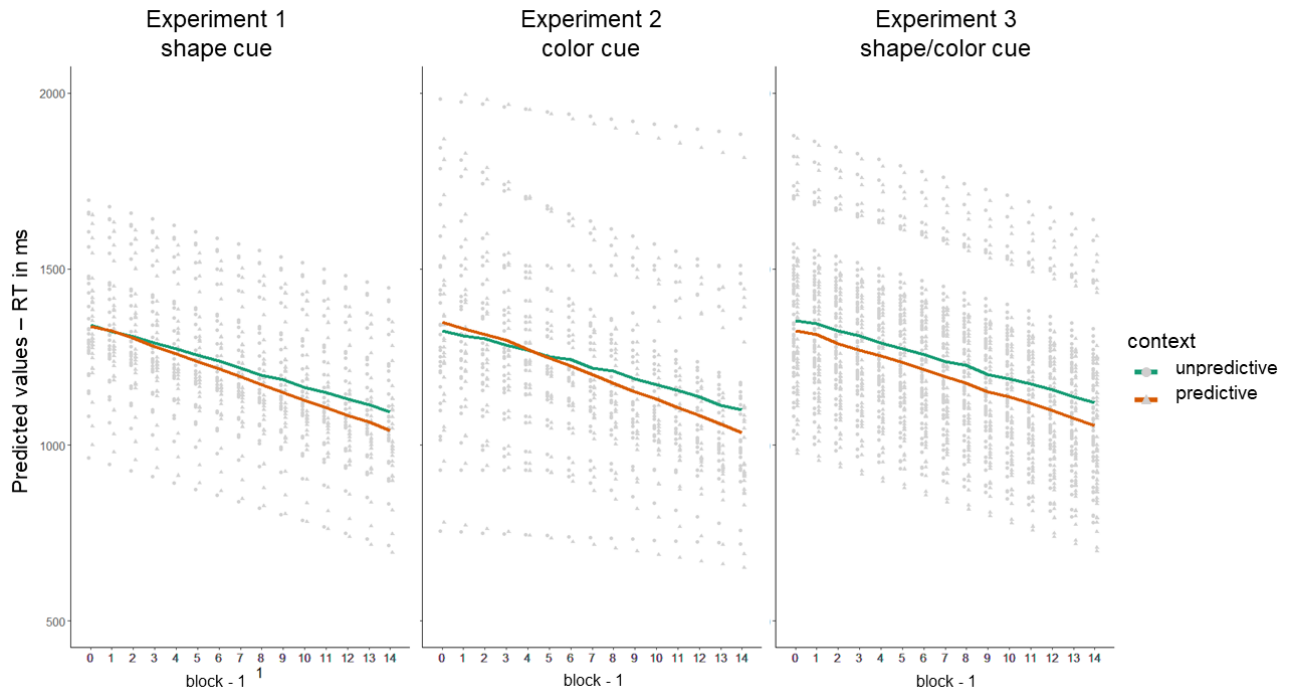
Response Times in Milliseconds by Block and Context (Predictive/Unpredictive)



Note. Error bars indicate standard errors.

Figure 4.

Predicted Response Times in Milliseconds per Participant, Block and Context (Predictive/Unpredictive) as Predicted by Mixed Effects Models for Experiments 1-3



Note. The light grey points indicate the individual intercepts/slopes for participants, the dot and triangle indicate the difference between predictive and unpredictable contexts. The thick lines indicate the overall effect of context.

Generation Task. Overall accuracy in the generation task for the predictive contexts was 27.22% ($SD=0.45$) which was not significantly above chance level, $t(29)=1.138$, $p=.132$. Mean confidence rating (1 – 4 scale) was 1.64 ($SD=0.50$) for the predictive shapes, and 1.59 ($SD=0.51$) for unpredictable shapes, which was not significantly different in a paired, one-tailed t -test, $t(29)=-1.244$, $p=.112$ ($BF_{01}=1.452$; interpreted as anecdotal evidence for the null hypothesis). We pooled confidence ratings for predictive shapes of 1 ($n=364$) and 2 ($n=267$) as low, and ratings of 3 ($n=70$) and 4 ($n=19$) as high confidence. The difference between

relative frequency of correct|high (23.66%) and correct|low (25.62%) responses was not significant in a paired, one-tailed t -test, $t(29)=-.317$, $p = .623$ ($BF_{01}=6.439$; interpreted as substantial evidence for the null hypothesis)³. A Pearson's product-moment correlation test showed a significant correlation within participants, between accuracy in the generation task and the CC effect, in the following simply defined as predictive – unpredictable RT, $r(28)=.367$, $p=.047$, but no correlation between accuracy and confidence in the generation task, $r(28)=.238$, $p=.206$.

Discussion

In Experiment 1, we could show that participants learn to predict the target location from the shapes of distractors. The interaction of context and block was significant, due to steeper RT slopes for predictive shapes than for unpredictable shapes. This search speed advantage for predictive shapes develops over the course of blocks, indicating learning to use the cues to find the target. The regression coefficient of the interaction of context and block indicates an RT difference increase between predictive and unpredictable contexts by -4.07ms with every block. The significant main effect for block indicates that there is also a general training effect. Experiment 1 thus showed that a task-relevant visual feature that characterizes distractor identity can be learned to cue the target location. We do not find a main effect of context, which is to be explained by the slowly emerging learning such that the RT difference between contexts is not consistent but only present in roughly the second half of the training phase.

As indicated by a lack of relation between objective measure (the generation task) and the confidence measure, we would argue that the knowledge of contingencies between cue and target location remained implicit.

³ As a sanity check, in all three Experiments, we also compared the relative values for correct|high and incorrect|high which are computed on the basis of all high certainty judgements instead of all correct responses. In all three experiments, the results resemble the results from the correct|high and correct|low comparison.

Table 1.*Response Time Analysis Results for Experiments 1-3*

	<i>Estimate</i>	<i>SE b</i>	95% CI <i>b</i>		df	<i>t</i>	<i>p</i>
			<i>LL</i>	<i>UL</i>			
Experiment 1 (shape)							
Intercept	1345.15	30.98	1284.42	1405.88	43	43.415	<.001***
Context (pred)	3.01	15.11	-26.60	32.63	94	0.199	.842
Block	-19.81	0.72	-21.21	-18.41	16878	-27.703	<.001***
Context x Block	-4.07	1.43	-6.87	-1.26	16886	-2.845	.004**
Experiment 2 (color)							
Intercept	1351.26	51.20	1250.92	1451.61	28	26.40	<.001***
Context (pred)	25.60	11.91	2.26	48.95	16436	2.150	.032*
Block	-19.91	2.56	-24.93	-14.88	28	-7.766	<.001***
Context x Block	-6.65	1.44	-9.47	-3.83	16436	-4.621	<.001***
Experiment 3 (compound)							
Intercept	1347.36	27.09	1294.27	1400.46	59	49.740	<.001***
Context (pred)	-26.87	12.52	-51.41	-2.34	141	-2.147	.032*
Block	-18.47	0.54	-19.52	-17.41	32194	-34.45	<.001***
Context x Block	-2.68	1.07	-4.78	-0.57	32196	-2.496	.013*
Experiment 3 (single cue blocks)							
Intercept	1118.51	26.82	1065.94	1171.08	57	41.710	<.001***
Context (pred)	-32.51	11.42	-54.90	-10.11	4491	-2.846	.004*

Note. SE = standard error; CI = confidence interval; LL = lower limit; UL = upper limit. Mixed-effects model computed coefficient, standard error and confidence interval for the coefficient, degrees of freedom, *t*-value, and *p*-value are displayed for each predictor and experiment. Degrees of freedom are rounded.

*** $p \leq .001$, ** $p \leq .01$, * $p \leq .05$

Experiment 2

In Experiment 1, we established a novel variant of the contextual cueing task with a non-spatial but task-relevant feature. Participants had to process the respective shape in order to find the target letter F. In the second experiment, we tested if participants even learned that different colors of the distractors predicted different target locations. As argued above, color is an entirely task-irrelevant feature.

Method

Participants. 30 participants were recruited via *Prolific* (19 female; $M_{age}=40.13$; $SD_{age}=12.20$). One participant was excluded from analysis due to poor performance in the training (48.88% accuracy). Participants were prescreened for living in the UK (to ensure that it was daytime), being fluent in English, having normal or corrected-to-normal vision, and not having taken part in a previous contextual cueing experiment of our lab.

Results

Training. Incorrect responses, trials with the less frequent target location or target location repetition were excluded from analysis (23.63% trimming). Mean accuracy was 94.49% ($SD=.23$), mean RT in the cleaned data set was 1205.86ms ($SD=470.25$). Mean RTs over the course of the blocks, separated by predictive and unpredictable context are displayed in Figure 3.

For RT as dependent variable, we first tested, whether the fixed effects structure hypothesizing an interaction of context and block, was the best fit for the data. Comparing AICs of models with no random effect structure and no fixed effect ($AIC=249809.7$), only block as fixed effect ($AIC=249243.1$), an additive ($AIC=249237.2$), and an interaction effect ($AIC=249223.6$), revealed that the interaction term model yielded the best fit. In a next step, we compared random effect structures. Allowing for random slopes for the factor context

(AIC=244395) was not a better fit than the random intercept model (AIC=244394; $\chi^2(2)=3.642$, $p=.162$), but random slopes for the factor block fit significantly better in comparison to the random intercept model (AIC=244144; $\chi^2(2)=243.1$, $p<.001$). More complex models did not converge or produced a singular fit. In the random slopes and random intercept model, the fixed effects context and block, as well as their interaction, were significant (see Table 1 and Figure 4).

Generation Task. Overall accuracy for the predictive colors was 30.03% ($SD=.46$) which was significantly above chance level, $t(28)=1.83$, $p=.039$, $d=-0.34$, $BF_{10}=1.639$ (interpreted as anecdotal evidence for the alternative hypothesis). Mean confidence rating (1 – 4 scale) was 1.52 ($SD=0.72$) for the predictive colors and 1.52 ($SD=0.71$) for unpredictable colors, which was not significantly different in a paired, one-tailed t-test, $t(28)=.041$, $p=.516$ ($BF_{01}=5.225$; interpreted as substantial evidence for the null hypothesis). We summarized confidence ratings of predictive colors of 1 ($n=414$) and 2 ($n=216$) as low, and ratings of 3 ($n=52$) and 4 ($n=14$) as high confidence. The difference between relative frequency of correct|high (19.78%) and correct|low (29.49%) responses was not significant in a paired, one-tailed t-test, $t(28)=-1.66$, $p=.946$ ($BF_{01}=12.268$; interpreted as strong evidence for the null hypothesis). A Pearson's product-moment correlation test showed no significant correlation between accuracy in the generation task and CC effect, $r(27)=-0.17$, $p=.385$, and no correlation between accuracy and confidence in the generation task, $r(27)=.22$, $p=.245$.

Discussion

As in Experiment 1, the significant interaction between context and block indicates learning of the cue and target location association. Also here, the RT slope for predictive colors is steeper than for unpredictable colors across training blocks. Other than in Experiment 1, we also obtain a significant main effect of context in the opposite direction as hypothesized

(predictive contexts are slower than unpredictable contexts overall). As can be seen from Figure 4, this finding is based on the reverse effect in the first block of the training phase. This can be demonstrated when, in a post-hoc analysis, we exclude the first block and compute the model (however, with only random intercepts as the random slope model did not converge). Then, the main effect of context is no longer significant, $b=17.60$, $p=.181$, but block, $b=-14.56$, $p<.001$, and the interaction effect is still significant, $b=-5.84$, $p<.001$. Testing the same model, but with context and block as additive factors, both main effects are significantly negative, context, $b=-26.47$, $p<.001$, and block, $b=-15.03$, $p<.001$. Both models fitted to the data with the first block removed, with additive and interaction effect respectively, are in line with a context learning effect. When not accounting for the interaction, both main effects are negative, which means that with block number, RTs decrease significantly, and for context, that RTs in predictive contexts are significantly shorter than in unpredictable contexts. In our original models, the main effect for context is overlaid by an interaction effect that stems from the first block. This interaction is however still significant when removing the first block. This significant interaction effect of context and block is consistent with the pattern of a context learning effect across blocks.

Building on previous research that has assigned task-relevance a major role in implicit learning (Jiang & Chun, 2001; Jiang & Leung, 2005; Jiménez & Méndez, 1999; Turk-Browne et al., 2005), color contingencies with target location should not have been learned. Nevertheless, we observe an RT advantage for predictive versus unpredictable colors over the course of the blocks that indicate learning. The model predicts a learning effect, which manifests in the RT difference between predictive and unpredictable contexts that increases by -6.65ms with every block.

In the generation task, we find a performance that is significantly above chance level. First of all, according to the Bayes Factor, there is no strong evidence for the hypothesis that

participants are indeed better than chance level. Secondly, we do not exclude the possibility of a higher-than-chance performance even under the premise of implicit learning. There is the chance that, given the training and the test phase are so similar, implicit knowledge spills over to the test phase, and causes higher performance (Michel, 2023a, but see Shanks & St. John, 1994). However, the awareness test lies in the association with a contingency measure, as this provides a metacognitive judgement. Here, we find no evidence of higher confidence in correct responses, which would be expected if participants had a metacognitive awareness of their knowledge of the contingencies, making it explicit knowledge (Dienes & Seth, 2010).

Experiment 3

In this experiment, we aimed to test the effects of cue competition. The task and set-up are the same as in Experiments 1 and 2, but here, we provided distractors that were characterized by a one-to-one mapped shape and color. Shape and color redundantly predicted target location. So, both, independently or integrated (as a compound or one overshadowing the other), could be learned to be used as a cue for target location. To test this question, additionally to the 15 training blocks, two blocks with 48 trials each were implemented. The distractors here contained either only color but not shape information (colored “R” distractors), or only shape but not color information (white distractors in six shapes) respectively. It was counterbalanced between participants if the first of the two blocks was the color or the shape block. With these blocks, we were able to test whether the learning effects for each individual cue would be additive, indicating independent learning effects, or underadditive, indicating compound learning or overshadowing effects.

Method

Participants. For Experiment 3, the number of participants was doubled relative to Experiments 1 and 2. This aimed to increase statistical power when testing effects in the single

cue blocks, as there were only 48 trials per participant and cue. An increase in the number of participants was preferred to an increase in the number of trials in the single cue blocks, because of the risk of participants learning the contingencies anew. Thus, 60 participants were recruited via *Prolific* (40 female; $M_{\text{age}}=39.29$; $SD_{\text{age}}=11.98$). They were prescreened for living in the UK, being fluent in English, have normal or corrected-to-normal vision, and had not taken part in a previous contextual cueing experiment of our lab. Two participants were excluded due to poor performance in the training phase (49.02% and 51.84% accuracy respectively).

Results

Training. Incorrect responses, target location repetitions and the infrequent target position trials were excluded (22.62% trimming). Mean accuracy was 95.06% ($SD=0.22$), mean RT in the cleaned data 1196.71ms ($SD=463.97$). Figure 3 displays mean RTs over the course of the blocks, separated by predictive and unpredictable context. For the fixed effect structure, comparing AICs of models with no random effect structure and no fixed effect (AIC=488753.4), only block as fixed effect (AIC=487808.3), an additive (AIC=487731.8), and an interaction effect (AIC=487729.5), revealed that, again, the interaction term model yielded the best fit. Adding random slopes for context yielded a better fit (AIC=480843) than only random intercepts (AIC=480963; $\chi^2(2)=124.09$, $p<.001$). More complex random effect structures did not converge or produced a singular fit. In the random intercept and random slopes model, the fixed effects of block and context as well as their interaction were significant (see Table 1 and Figure 4).

Single Cue Blocks. Mean RT in these blocks was 1113.89ms ($SD=432.23$), mean accuracy was 96.17% ($SD=0.19$). The trimming was conducted with the same procedure as in the training phase and affected 18.80% of the data. In the color block, mean RT was 1117.03ms ($SD=423.84$), mean accuracy was 96.30% ($SD=0.19$). In the shape block, mean RT was 1110.68ms ($SD=440.73$), mean accuracy was 96.05% ($SD=0.19$).

In a mixed model with no random effects, only context as fixed effect (AIC=68558.07) proved the best fit, contrasted with the null model (AIC=68562.11), only cue as factor (AIC=68563.86), both factors additive (AIC=68559.96), or in an interaction term of both factors (AIC=68561.57). A model with random intercepts and random slopes for the two cues (color and shape) per participant provided a better fit for the data (AIC=67671.53) than a model with only random intercepts for participants (AIC=67693.61; $\chi^2(2)=26.083$, $p<.001$). Post-hoc Bayesian t-tests revealed that the difference between predictive and unresponsive contexts in the shape block was not significant, $t(57)=-1.419$, $p=.081$, $BF_{01}=1.477$ (interpreted as anecdotal evidence for the null hypothesis), but in the color block, it was significant, $t(57)=2.22$, $p=.015$, $d=-0.17$, $BF_{10}=2.727$ (interpreted as anecdotal evidence for the alternative hypothesis).

Generation Task. For the compound stimuli, overall accuracy in the predictive contexts was 27.73% ($SD=.45$) which is not significantly above chance level, $t(57)=1.58$, $p=.060$. Mean confidence rating (1 – 4 scale) was 1.69 ($SD=0.82$) for predictive contexts and 1.69 ($SD=0.83$), which was not significantly different in a paired, one-tailed t-test, $t(57)=-0.159$, $p=.437$ ($BF_{01}=6.129$; interpreted as substantial evidence for the null hypothesis). We summarized confidence ratings for predictive contexts of 1 ($n=720$) and 2 ($n=410$) as low, and ratings of 3 ($n=234$) and 4 ($n=28$) as high confidence. The difference between relative frequency of correct|high (21.56%) and correct|low (25.54%) responses was not significant in a paired, one-tailed t-test, $t(57)=-1.10$, $p=.862$ ($BF_{01}=13.700$; interpreted as strong evidence for the null hypothesis). A Pearson's product-moment correlation test showed no significant correlation between accuracy in the generation task and the CC effect, $r(56)=0.04$, $p=.796$, and no correlation between accuracy and confidence in the generation task, $r(56)=.06$, $p=.630$.

Discussion

In Experiment 3, as a form of replication and extension of Experiments 1 and 2, the main effect of context and its interaction with block were significant, indicating learning of the predictivity of the compound cues in the search task (Table 1). In contrast to the results of Experiments 1 and 2, the difference between predictive and unpredictable contexts in RT emerges much faster, within the first block. This is also why the main effect of context is strongly negative, with an estimate of -26.87. The estimate for block is roughly the same as in Experiments 1 and 2, indicating that the general training effect is similar in all experiments. As a consequence, the interaction effect of context and block is less strong, estimating an increase in the RT difference between contexts of -2.68ms per block, given that the RT difference in contexts emerges faster across blocks as in Experiments 1 and 2.

In the single cue blocks, results remain somewhat more ambiguous. By presenting only one feature of the shape-color cue, either only shape or only color, we aimed to test the learning of single feature contingencies with target location. However, our analyses do not provide strong evidence for learning of either single feature contingency. Although the frequentist test for a difference between predictive and unpredictable colors is significant, Bayesian analysis suggests that the evidence for such an effect is only anecdotal. This might be a design-inherent limited power because we have only 48 trials per single feature cue per subject to be able to test the difference between predictive and unpredictable contexts. However, if we had presented more blocks for the single feature cue, there probably would have been a new learning of the single cue contingencies, which is not what we were aiming to test.

As can be seen in Figure 5, it is noteworthy that the RTs with predictive colors and shape contexts in the single cue blocks are similarly short as in the last block with compound cues, indicating that there are no costs of switching from compound to single cues. Still, the context difference becomes smaller, and that is because of shorter RTs in the unpredictable contexts. This

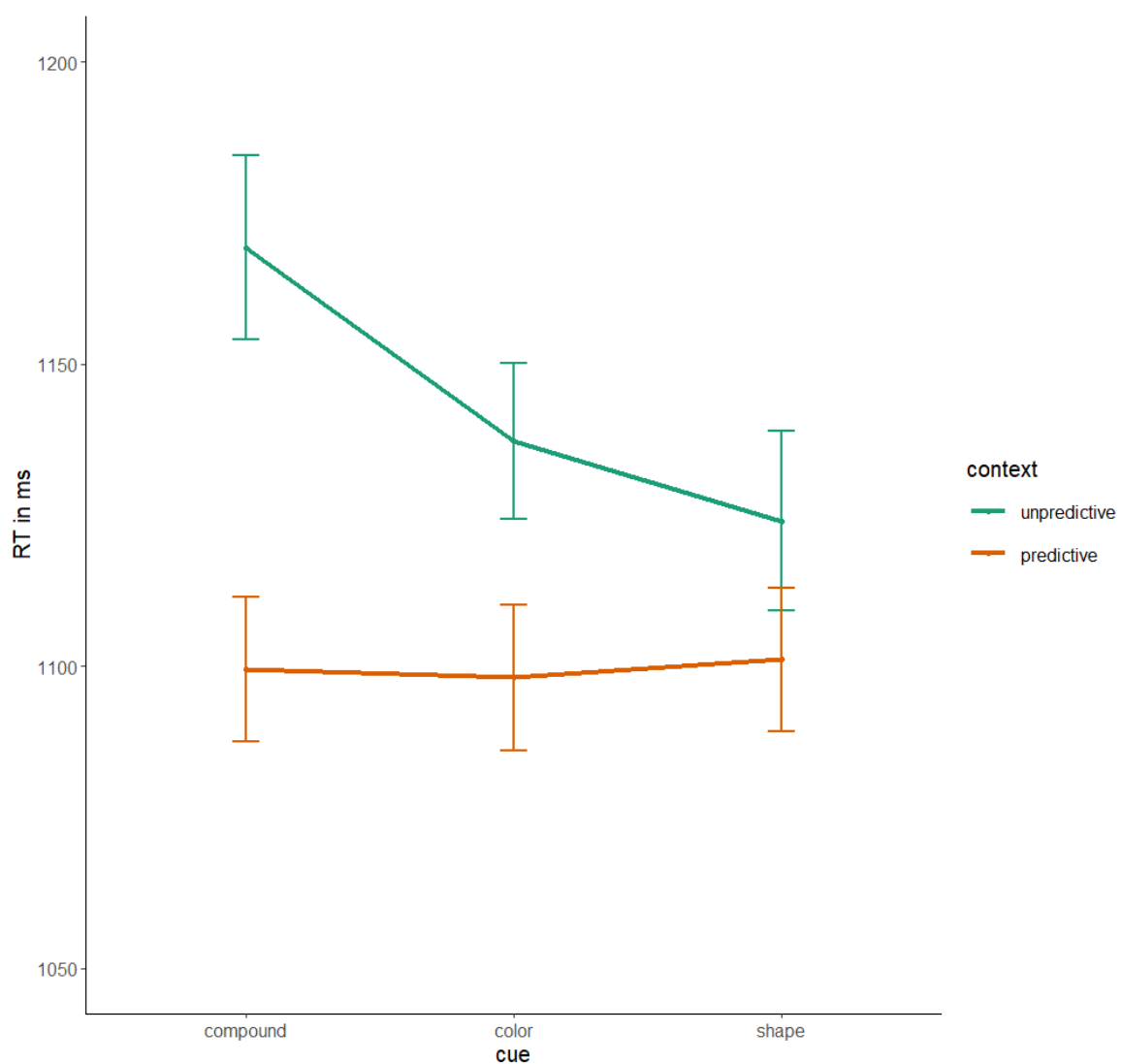
is not to be expected from the change of a compound cue to a single feature cue. One possible explanation would be that the single cues make it easier to detect the target in the unpredictable context. However, when comparing the RTs in unpredictable contexts in the single cue blocks with RTs in Experiments 1 and 2, where the same single feature cue displays are presented, mean RTs are virtually the same. Therefore, we suspect that the accelerated RTs in unpredictable contexts might rather be the result of unsystematic variation of the RTs in the unpredictable contexts that we similarly observed in Experiments 1 and 2. When looking at those RTs across blocks in Figure 3, there is some variation and overlapping standard error bars, however, still with RTs consistently larger in unpredictable than in predictive contexts. Given this variation in the unpredictable context RTs, and the argument of no observable costs from compound cue to single cues for predictive context RTs, one might argue that predictions from single cues could be used as well as from the compound cue.

To further explore the relationship between compound and single cue learning, we conducted an explorative analysis comparing all three experiments. We fitted a mixed-effects model for each training phase of each experiment and the single cue blocks in Experiment 3, including only context (predictive/unpredictive) as fixed effect and random intercepts for subjects. Then, we compared the fixed effect estimates for context over the three experiments (Figure 6). For Experiment 3, the estimate was roughly double ($b=-69.98$, 95% CI $[-102.80, -37.17]$) compared to the estimates in Experiment 1 ($b=-20.44$, 95% CI $[-33.02, -7.87]$) and Experiment 2 ($b=-23.88$, 95% CI $[-36.47, -11.29]$). The estimates for the single cue blocks (color: $b=-32.01$, 95% CI $[-63.42, -0.59]$; shape: $b=-23.56$; 95% CI $[-56.56, 9.44]$) in Experiment 3 are comparable to those of Experiments 1 and 2, except for their variance estimation, given that the estimation is based on a small number of trials in the single cue blocks of Experiment 3. And also descriptively, the CC effect in the compound blocks of Experiment 3 is almost double the size ($M_{CC}=71.41$, 95% CI $[53.08, 90.74]$) relative to the color ($M_{CC}=31.36$, 95% CI $[14.49, 48.22]$)

and shape ($M_{CC}=35.56$, 95% CI [17.24, 53.88]) single cue blocks respectively. Taken together, the learning effects of the individual cues seem to be additive with respect to their compound presentation.

Figure 5.

Single Cue Block Response Times in Experiment 3



Note. Response times in the last training block of Experiment 3 with compound cues, color and shape, and in single cue blocks (only color cue vs. only shape cue) per context (unpredictive vs. predictive). Error bars indicate standard error.

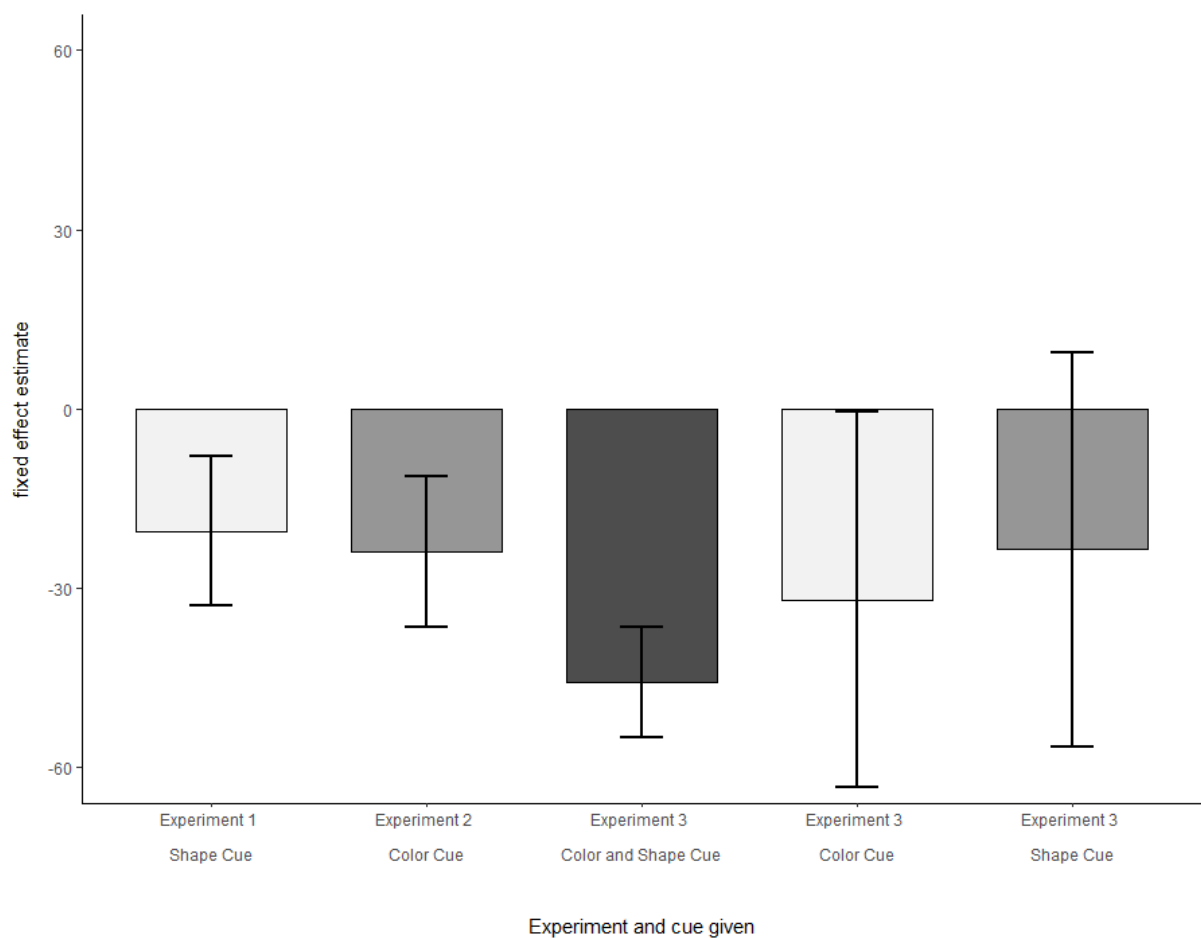
So, we obtain similar results from the two different approaches, a within-experiment and an across-experiment analysis. Results from both analyses are compatible with the notion of an additive learning effect. When looking at CC effects in the compound and single cue blocks of Experiment 3, we have a considerably larger effect of the compound cue. When comparing across experiments, we observe the same pattern, descriptively in the RT differences, and also in the fixed effects estimates for context in the three experiments. Only when taking into account the data pattern in Experiment 3, where we observe no RT costs in the predictive contexts, switching from a compound to a single cue, one might lean toward a different interpretation. It could mean that the cues are learnt independently, resulting in an underadditive effect. However, in terms of the CC literature, it is unconventional to interpret performance in the predictive contexts only, instead of the comparison between unpredictable and predictive contexts in the sense of the CC effect. We therefore interpret the results as supporting an additive CC effect.

The interpretation of the results at the group level remains somewhat ambiguous. On the one hand, it is conceivable that participants learned to predict the target location from both single cues, but benefited even more from a compound cue, producing additive learning effects. On the other hand, the smaller single cue learning effects at the group level could also be the result overshadowing effects at the individual level, meaning that one group of participants learned only the color cue, and the other group only learned the shape cue. This would potentially also result in the observed pattern of seemingly additive learning effects. From the planned analysis, we can see already that the model that fit the data best, was one that allowed random slopes for the two cues across participants. This might point to the possibility that participants differed in their color and shape CC effects. Therefore, we would argue that analyzing the cue competition effects on the group level is not sufficient for a nuanced interpretation. There could be overshadowing effects on an individual level, resulting in ambiguous group effects (and

weak evidence in terms of the Bayesian analysis; G. S. Reynolds, 1961). To explore these effects, we conducted post-hoc explorative analyses on the individual level.

Figure 6.

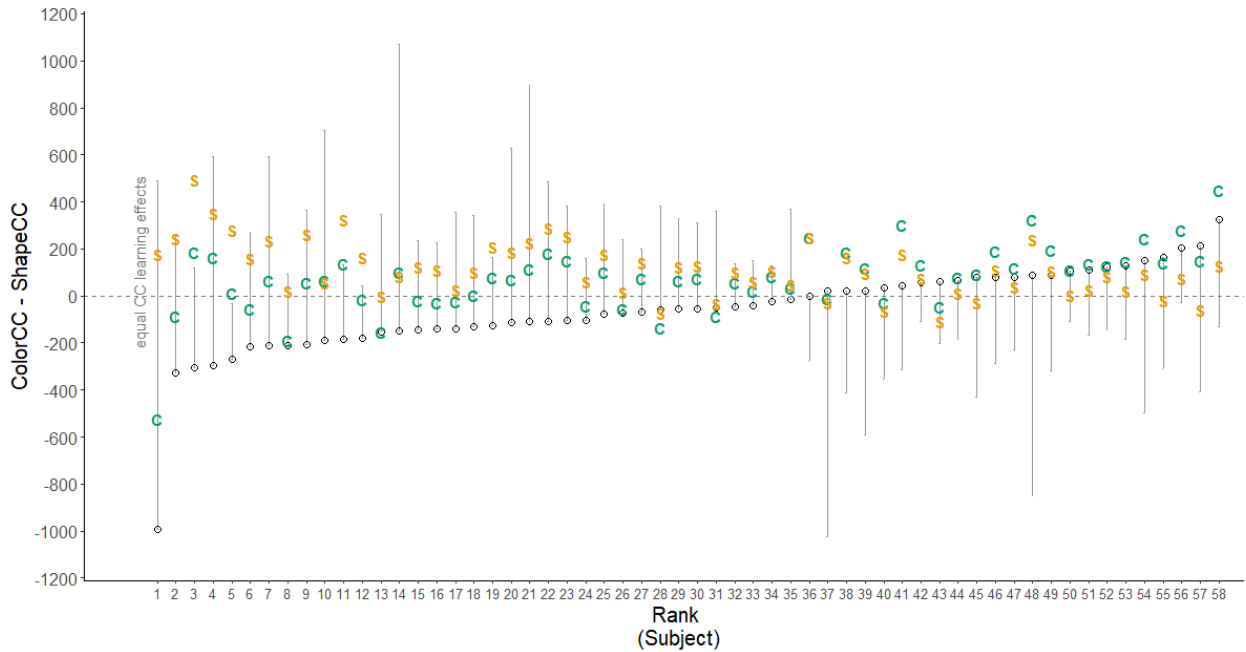
Comparison of the Fixed Effect Estimates for Context in the Training of Experiments 1-3 and the Single Cue Blocks of Experiment 3



Note. The fixed effect estimates refer to a mixed-effects model with context as fixed effect, and random intercepts for subject. Error bars indicate 95% confidence intervals.

Figure 7.

Comparison of the contextual cueing effects for color and shape cues in Experiment 3 by subject (ranked)



Note. The CC effect is computed by subject, deducting RTs in predictive contexts from RTs in unpredictable contexts, per single cue block, color (ColorCC; displayed as green “C”) and shape (ShapeCC; displayed as orange “S”). These two difference measures are then summarized into a difference measure contrasting the CC effect in the color versus in the shape block. Subjects are ranked based on the contrast measure, and it is then plotted as diamonds in the graph. The grey, dashed line indicates no difference between the color and shape block. Points below the zero-line indicate a stronger CC effect in the shape block, and points above the line indicate a stronger effect in the color block. Error bars indicate confidence intervals (95%), but are only shown one-sided towards zero for better visualization.

We computed estimates for CC effects per subject separately for the shape and the color cues. For each individual participant, we computed trial-wise RT differences for predictive and unpredictable contexts. Trial-wise in such a fashion that the trial pair from which the RT difference was computed had the same cue feature, the same target position, and the same target

identity, that is, the same response (long, short). From these differences between comparable trials, we then computed means and confidence intervals (95%) for a CC effect (unpredictive – predictive RT) per participant and per cue. For the final difference measure, we then deducted the CC effect of shape from the CC effect of color, reasoning that if participants learned both feature contingencies roughly equally well, it should result in equal CC effects, and thus in an around zero difference measure. If this difference measure is substantially above zero, it would indicate a more pronounced learning of the color contingencies (suggesting overshadowing of shape). If it is below zero, it suggests a stronger CC effect of the shape contingencies (suggesting overshadowing of color). To illustrate these effects on the individual level, Figure 7 displays the CC effects separately for color and shape cues as well as the difference measure for each participant rank-ordered according to the size of the difference between the CC effect for color minus the CC effect for shape. We observe that most participants have CC effect differences of around zero. Therefore, we would not argue that individual overshadowing effects are driving the weak CC effects on the group level in the single cue blocks.

General Discussion

The main goal of the current study was to investigate the role of selective attention in the implicit acquisition of contingencies between features. We implemented these contingencies in a novel variant of the contextual cueing paradigm using identity cueing instead of the classical spatial configuration cueing. For the purpose of testing the role of selective attention, we manipulated the task-relevance of distractor features that predicted the target location. In Experiment 1, the predictive feature was the task-relevant shape of the distractors. In Experiment 2, it was the task-irrelevant feature color. In Experiment 3, we aimed to test cue competition effects and therefore presented compound cues of color and shape.

The results of the first two experiments showed that participants learned to predict the target location from the shape (Experiment 1) and from the color as well (Experiment 2). The

RT differences between predictive and unpredictable search contexts emerged over the course of the training blocks in both experiments.

A generation task which also contained a confidence measure indicated that these learned associations were not explicitly represented. Participants were not able to report the correct target location according to a predictive feature above chance (only in Experiment 2, but Bayesian analysis provided no substantial evidence), and were not more confident in their respective response when they had responded with the correct target location. This indicated that participants did not have metacognitive access to the acquired information, enabling them to distinguish between their correct and incorrect responses (Michel, 2023a).

What do these findings offer in terms of understanding the role of selective attention in implicit learning? Attentional or selective mechanisms are essential to our cognitive system, in the visual system alone, we are bombarded with information of about 10^8 bits per second (Itti & Koch, 2000; Marois & Ivanoff, 2005). This requires mechanisms of selection, chunking, and binding (Fiser & Aslin, 2005; Wheeler & Treisman, 2002). As reviewed above, a number of studies suggested that task-relevance of a predictive feature, manipulated by instruction or by the nature of the task, is necessary for it to be learned implicitly (Jiang & Chun, 2001; Jiang & Leung, 2005; Jiménez & Méndez, 1999; Turk-Browne et al., 2005). What is implicitly assumed when arguing for a central role of selective attention in implicit learning is that implicit learning is subject to capacity limits. However, this seems to contradict the widely confirmed finding that people can learn more than one contingency in parallel (Conway & Christiansen, 2006; U. Mayr, 1996; Wilts & Haider, 2023). In addition, our current finding suggests that also contingencies involving task-irrelevant cues can be learned. Thus, there might not be such a compelling argument for a functionally imperative role of selective attention in implicit learning.

To solve this contradiction, it might be useful to refer to research in action control, because here, research has been going in a similar direction. In the framework of the Theory of

Event Coding (TEC; Hommel et al., 2001), an event file is thought to be formed when we integrate stimulus features and responses into an episode that can then be activated by the respective stimulus or response features it entails (Hommel, 1998). In multiple series of experiments, it has been tested what the attentional prerequisites for a stimulus or context feature are to be integrated into an event file. Conclusions from such experiments were that features are integrated into an event file when they are task-relevant (Chao et al., 2022; Hommel, 2005; Huffman et al., 2018), specifically, also if they can be used to discriminate targets from distractors (Hommel & Colzato, 2004). More recently, the modeling of the mechanism has been refined, as it has been proposed that the selectivity of integrating features does not lie in the encoding and building of an event file, which is now thought to be automatic, but rather at the retrieval stage of the event file (Hommel et al., 2014; Schmalbrock et al., 2023). Thus, it is not the question whether a feature is integrated a priori, but whether the weighting of a feature (Hommel et al., 2014; Memelink & Hommel, 2013) enables the retrieval of the episode (event file) in a future occurrence. The paradigms that are used in the context of action control, often rely on trial-by-trial observations, examining the effect of a trial n feature and response on a trial $n+1$. We believe that, with our longer-term learning context, we can extend the scope of studying the processing of features beyond this trial-by-trial frame (Moeller & Pfister, 2022). In our view, one can integrate our findings into the TEC framework, in a way that features are learned to predict events or actions when they activate respective event files that contain such information, irrespective of the features' task or response relevance. Applied to our current findings, a possible assumption concerning the underlying mechanism is that all features of a trial are integrated into an event file. Given that one feature is contingently paired with the target location (e.g. color), the retrieval of that episode containing the correct target location is strengthened over time (Hommel, 1998; Rescorla & Wagner, 1972). Consequently, it would not be task-relevance (or selective attention) that modulates implicit learning but rather the retrieval of episodes (event file), and the question of implicit learning is whether a particular feature is

capable of triggering the retrieval of a certain episode. If so, it leads to performance benefits, or, as we coined it here, implicit learning. This mechanism seems to be effective with task-relevant cues (distractor shape) and with task-irrelevant cues (distractor color).

We acknowledge that our manipulation of task-relevance differs from the studies presented in the introduction. We provided a context in which all distractors needed to be evaluated with respect to their shape matching the target shape or not. This way, color was not task- or response-relevant. However, it may have been processed stronger than in the case of irrelevant stimuli in previous studies, in which, for example, stimuli of a certain color did not have to be searched at all (Jiang & Chun, 2001; Jiang & Leung, 2005; Turk-Browne et al., 2005). Yet, what is unique to our design, is the distinction between features on the higher level, marking shape as task-relevant, and color as task-irrelevant, instead of marking one specific shape or one specific color as task-relevant or not. We argue that this is the more relevant question when it comes to specifying the building blocks of implicit learning. In that question, we test theoretical accounts that postulate processing in feature-specific modules that may not be able to integrate information from different features that are not attended (Baars, 2005; Eberhardt et al., 2017; Keele et al., 2003). In our experiments, we find such learning effects across features, not just within one feature. Although not compatible with feature-wise processing in independent modules, this finding is in line with the underlying learning mechanism we proposed above. Because when information in a trial is encoded into an event file, contingencies within or across features can, in principle, learned to be associated.

A notable limitation of our experiments is that task-relevance in our contextual cueing variant is confounded with the respective feature of the cue. That is, shape is task-relevant, and color is task-irrelevant. We cannot balance these two factors, because when target color were to be the relevant feature, we would have a pop-out effect that would hardly be affected by predictability of the distractors' shapes or colors. We had no reason to believe that the two visual

features (*ceteris paribus*) would differ in their potential to be associated with target position. Other researchers have found CC effects for (background) color (Kunar et al., 2006; Kunar et al., 2013), but also learning effects for irrelevant but predictive shapes (Levin & Tzelgov, 2016), and even letters (Miller, 1987). From visual recognition and visual scene processing literature, we would even hypothesize that there is a primacy of shape information over color information (e.g., Biederman & Ju, 1988; Del Viva et al., 2016). Extrapolating this to our experiments, the likelihood of color contingency learning would be further reduced. But note that this is an effect found with more complex stimuli, and might be traced back to complexity reduction, therefore not being transferrable to our simple stimulus set-up. Thus, although a limitation of our design is that these two conditions, task-relevance and feature dimension, cannot be disentangled, there is no compelling argument as to why the feature dimension should be the main contributor of the effect. More so, there would be an argument to hypothesize the opposite effect of task-relevance and feature dimension. Thus, from our experiments, we would deduce that, in principle, task-relevant and task-irrelevant features can be integrated and used for predictions (Experiments 1 and 2).

In Experiment 3, we used colors and shapes as compound cues and, after training, tested learning of both features in isolation. In the interpretation of the results from the single cue blocks, the picture is more nuanced. First, overall, we observe no costs in the RTs in predictive contexts from compound cue to single cues, that is, from training to the single cue blocks. Additionally, the descriptive differences between predictive and unpredictable contexts in the compound blocks (CC effects) are almost double the size than the differences in the shape and color single cue blocks. The same relationship is found when comparing the fixed effects estimates for context in the compound training of Experiment 3, which are almost double the estimates of the trainings in Experiment 1 (shape) and Experiment 2 (color). Thus, the learning effects in single cue experiments (Experiments 1 and 2) and the single cue block effects seem additive

with respect to the learning effect in the compound cue blocks of Experiment 3. Such summation effects have been shown in operant conditioning, when comparing compound cue and single cue learning in animals (Mackintosh, 1976; Miles & Jenkins, 1973; Thein et al., 2008).

Yet, our results from the single cue blocks remain somewhat ambiguous. It remains unclear if there are reliable context effects in the single cue blocks whatsoever. That the mixed model with random slopes for cue feature per participant fit the data best, is a first indicator for individual variance in the learning effects of the two cues. However, in an exploratory individual participant analysis, we do not see convincing evidence for overshadowing effects of any of the two features within participants. This is in itself interesting though, because an overshadowing effect would have been probable not only due to feature saliency (Mackintosh, 1976) or individual preferences (G. S. Reynolds, 1961), but also because the shape feature was task-relevant and was thus more probably going to overshadow the task-irrelevant color cue. We manipulated task-relevance to alter attentional processes, and overshadowing effects are also believed to build on attentional mechanisms (Mackintosh, 1971), in the sense that although more than one contingency can be learned, not all learned contingencies are translated into behavior (Kaufman & Bolles, 1981; Matzel et al., 1985). Thus, attentional processes would have been influenced by task-relevance and cue competition effects, and it is conceivable that cue competition effects would be influenced by the task-relevance of such cues. However, we observe no advantages for the task-relevant cue. This might point to a reciprocal overshadowing that has been observed in animals (Mackintosh, 1976; Miles & Jenkins, 1973), meaning that both features overshadow each other, resulting in a result pattern of a summation effect, as described above.

In a recent article, J. R. Schmidt and De Houwer (2019) noted that there is surprisingly little research on the issue of cue competition, especially in implicit learning (but see Beesley & Shanks, 2012; Endo & Takeda, 2004 for evidence from contextual cueing; Cleeremans, 1997;

Jiménez & Méndez, 2001, for evidence from implicit sequence learning). In multiple large studies, J. R. Schmidt and De Houwer (2019) found no evidence for blocking or overshadowing in an implicit learning paradigm. They also labeled the predictive features (shapes and words) in their experiments as task-irrelevant, because the response itself was only based on color. In that respect, their findings are consistent with ours: Task-irrelevant features that are predictive (although, in their case, for response), are still learned. In their case, they are even learned equally strongly, without overshadowing or blocking each other. That fits our take on interpretations of Experiments 1 and 2 – independently from task-relevance, cue contingencies can be learned. Our addition from Experiment 3 is that cue competition in our variant of an incidental learning paradigm does not result in overshadowing effects, even though one cue is task-relevant and the other is task-irrelevant. Rather, our results are compatible with the notion of independent learning of cues, resulting in additive learning effects in compound presentation.

One last issue concerning our findings might be doubts about the implicit nature of the knowledge in the contextual cueing paradigm. We claim that while participants' performance in the training phase reflected learning of the contingency between the respective feature and the target location, they were unable to express this knowledge explicitly. There is a long-lasting debate whether the common variant of contextual cueing is in fact based on non-conscious learning. It has been argued that studies have failed to correctly test for conscious knowledge (Luque et al., 2017; Vadillo et al., 2016; Vadillo, Linssen, et al., 2020), especially because they are underpowered, and measurement error leads to wrong conclusions regarding the implicit nature of the CC effect. In an attempt to empirically add to the debate, Colagiuri and Livesey (2016) tested samples of over 600 participants, and found no positive relationship between explicit knowledge and the cueing effect. Nevertheless, we take the criticism on the conventional testing for explicit knowledge seriously. Contextual cueing studies originally implement a recognition task: They show participants old and novel spatial configurations, and ask them to

categorize them into old and novel (Chun & Jiang, 1998). This means that they use a one-trial test for each configuration with often small sample sizes, and it does not seem surprising that there is a reliability and power issue here (Smyth & Shanks, 2008). This is why we did not implement a recognition task, as in the common variant, but a task that mirrors exactly the task that was provided in the training to enhance sensitivity of our test (Shanks & St. John, 1994). Participants thus had every chance to express any knowledge or intuition from the training in the generation task. Additionally, we were not restricted to a one-trial test, as one is in the recognition tasks. Rather, we presented participants with the same cue (color or shape, with random spatial configurations) multiple times, making the measure more reliable (Smyth & Shanks, 2008). Note that we can also expect that, given conscious awareness, the task to reproduce the contingencies between color or shape and target location is considerably easier than to recognize spatial configurations of distractors, and recall target location from that. So, we would expect less false-negatives (i.e., participants have explicit knowledge but cannot demonstrate that in the task) a priori. With our method of testing both an objective performance measure and a subjective confidence measure, and in addition testing for their interdependence (Michel, 2023a), we propose that what we observed here is indeed implicit knowledge.

Conclusion

Concluding, in our variant of the contextual cueing paradigm that utilizes identity cueing instead of the original, spatial cueing, we find compelling evidence for the learning of contingencies involving task-relevant and task-irrelevant cues. Further, when implementing compound cues in the learning phase, and testing the individual features of the cues in a subsequent test phase, we do not find evidence of overshadowing, neither on the group, nor on the individual level. There seem to be no costs of switching from the compound cues to the single cues with regard to RTs for predictive cues. Transferring current debates from the literature on event files and binding to our broader learning paradigm, we suggest that similar mechanisms can

account for our results. That would mean that in learning tasks, designed to observe binding, or, like ours, observe implicit learning, all available features of a trial are bound into an episode, and such features can be used to retrieve said episode, providing performance benefits such as shorter RTs. The retrieval of an episode from a given cue does not seem to depend on the task-relevance of said cue, but on the predictive value it provides. This would ultimately mean that attention, operationalized as a consequence of task-relevance, does not play a major role in the modulation of implicit learning. Taking into account previous empirical findings and theoretical accounts, our conclusion might be limited to situations in which task-relevance is manipulated across and not within features, and when contingencies occur within, but not across trials.

Appendix B

The following manuscript refers to Study 2. It is included as part of the cumulative dissertation.

Title:

Transferring Predictions from One Visual Cue Dimension to Another in Implicit Learning

Authors:

Felice Tavera, Sarah Wilts, & Hilde Haider

Note:

References cited in this appendix are included in the general reference list. To enhance readability and avoid redundancy, separate reference lists have not been provided for individual appendices.

Transferring Predictions from One Visual Cue Dimension to Another in Implicit Learning

Felice Tavera¹, Sarah Wilts¹, & Hilde Haider¹

¹University of Cologne

Author Note

Felice Tavera  <https://orcid.org/0000-0003-0994-6620>

Sarah Wilts  <https://orcid.org/0000-0003-2397-0969>

Hilde Haider  <https://orcid.org/0000-0001-7293-3166>

The authors have no known conflicts of interest to disclose.
Felice Tavera and Sarah Wilts were each supported by the Klaus Murmann PhD scholarship awarded by the Stiftung der Deutschen Wirtschaft (German Economy Foundation).
Study materials can be obtained upon request, data and analysis scripts have been made available at the Open Science Framework (OSF) and can be accessed at
https://osf.io/3z5qv/?view_only=0878eeddad464b88b8c0a03660bcb0b7

Corresponding author:

Felice Tavera

Department of Psychology

University of Cologne

Richard-Strauss-Str. 2

50931 Köln, Germany

Phone: +49-221 470-1471

Email: felice.tavera@uni-koeln.de

Abstract

The question whether learned contingencies can be transferred from one feature to another is essential for various theoretical frameworks concerning implicit learning. We tested contingency knowledge transfer in an adapted contextual cueing paradigm. In the training phase, distractor shapes predicted the location of a target, whereas, in the transfer phase, distractor colors predicted the target location. We tested transfer from the training phase in three groups: In a pre-matching group, shapes were associated with colors in a phase preceding training. In a post-matching group, they were associated after training. In a control group, participants did not learn to associate shapes and colors. Participants in all groups learned the contingencies. However, transfer was observed only in the pre-matching group. In all groups, the associations of shapes or colors with target locations remained implicit. We discuss implications for proposed mechanisms of implicit learning and preconditioning.

143 words

Keywords: Implicit learning, contextual cueing, visual search, attention

Introduction

In many models attempting to describe information processing and associative learning, especially implicit learning, information input is thought to be processed in more or less encapsulated modules, depending on their dimension (Baars, 2005; Dehaene & Naccache, 2001; Fodor, 1985; Keele et al., 2003; Magen & Cohen, 2007). Given such modular processing architectures, implicitly learned associations are thought to be represented in independent modules without any information exchange between them. This has specific empirical implications. For example, following Keele et al. (2003), in their “unidimensional” system, co-occurrences or sequences of events within one dimension would be processed and learned independently from co-occurrences or sequences within another dimension. Integrative, and then potentially explicit representations, only occur when this information is processed within the “multidimensional” system (Keele et al., 2003).

In support of this assumption, empirical evidence shows that two uncorrelated sequences instantiated by two distinct dimensions (modalities such as perceptual, motoric), can be learned concurrently (stimulus color and response location: Haider et al., 2014; U. Mayr, 1996; Remillard, 2017). Thus, implicit learning in one dimension does not interfere with implicit learning in a different dimension. As it has been shown that two uncorrelated sequences within one modality but instantiated by distinct features can be learned concurrently (color and shape within the visual modality: Conway & Christiansen, 2006; Wilts & Haider, 2023), an open question concerning these encapsulated modules is whether they are specialized on a dimension in the sense of modalities (e.g., perceptual, motoric) or of features within one modality (e.g., color, shape, location within the visual modality).

Irrespective of this, the assumption of encapsulated modules implies that implicitly acquired contingencies instantiated by one specific dimension should not be transferrable to contingencies instantiated by another dimension. Yet, Haider et al. (2020) showed that a visually

perceived stimulus location sequence could be expressed as a motor response key location sequence. However, both the to-be-learned and the to-be-transferred sequences relied on a spatial dimension. To better understand the building blocks of encapsulated modules, it is critical to test whether such knowledge transfer can be found outside the spatial dimension.

The present study aims to show knowledge transfer, specifically transfer of response contingency knowledge between two visual-perceptual features (color and shape), in an adapted contextual cueing paradigm (Tavera & Haider, 2025). Here, in the training phase, the target location is predictable from only the shape of the distractors. In the transfer phase, the distractors are presented in a novel shape, but in colors that now predict the target location. Importantly, participants learn the connection between shapes and colors in a matching phase either before (pre-matching group) or after (post-matching group) the training phase, while a third group (control group) does not learn this association. Our hypothesis is that, only when we induce such a shape-color matching, any learned contingency from shapes can be transferred to the respective colors.

Method

Participants. Via Prolific (www.prolific.com), we recruited 64 participants for the pre-matching group, 60 participants for the post-matching group, and 60 participants for the control group. In the pre-matching group, seven participants were excluded due to <60% accuracy in the training phase. Following the same criterion, one participant in the post-matching group and three participants in the control group were excluded. This way, 57 participants remained in the pre-matching group (36 women, 1 diverse; $M_{\text{age}}=34.51$; $SD_{\text{age}}=10.90$), 58 in the post-matching group (35 women, $M_{\text{age}}=36.40$; $SD_{\text{age}}=12.94$), and 57 in the control group (35 women; $M_{\text{age}}=32.84$; $SD_{\text{age}}=9.82$). Participants were prescreened for living in the UK, being fluent in English, having normal or corrected-to-normal vision, not being color-blind, and having not taken part in a previous contextual cueing experiment of our lab.

Stimuli. Search displays were constructed following the method Tavera and Haider (2025) with minor variations. The search displays were 15x10cm in size. The displays consisted of 16 letters on a dark grey background (RGB 60, 60, 60) rotated randomly by 0°, 90°, 180°, or 270°. They were organized in an 8x10 (invisible) grid, and distributed equally between the two horizontal halves of the display (Figure 1b).

In the training phase, all 15 distractor letters of one search display were white (RGB 255, 255, 255) and shaped as one of the four stylized letters A, E, P, S. In the transfer phase, all 15 distractor letters of one search display were R-shaped and were colored in one of the four colors green (RGB 1, 204, 0), orange (RGB 254, 153, 0), pink (RGB 255, 0, 254), and cyan (RGB 1, 255, 255). In one of four possible locations (Figure 1c), the target letter F appeared with either equally long horizontal bars or a shorter second horizontal bar. All targets were randomly rotated to the right (90°) or the left (270°) and were always the same color as the distractors.

Contingencies between distractor features and target locations were probabilistic. In 70% of the trials, the distractor shapes (training phase) or colors (transfer phase) were matched to one target location. The match between shape/color and target location was counterbalanced across participants. In 30% of the trials, the remaining three target locations were equally likely and randomly alternated for each shape or color. We refer to the 70% trials with predictable target locations as "regular trials" and the 30% trials with random target locations as "deviant trials".

Procedure. The experiment was conducted in accordance to the Declaration of Helsinki. Participants were recruited online via *Prolific* and were reimbursed according to *Prolific*'s ethical reward standards. They were redirected to *Pavlov*, where the experiments were uploaded from *PsychoPy2* (version 2020.2.4; Peirce et al., 2019) and adapted to the *Javascript* environment. For a short questionnaire, they were then redirected to *Sos*cisurvey (Leiner, 2024).

Participants were informed about the procedure of the experiment and asked to give their informed consent. To ensure equal size of the search displays for every participant irrespective of the monitor size or aspect ratio, all participants were first asked to follow a screen scaling procedure (Wakefield Morys-Carter, 2021).

As described in Figure 1, the experiments consisted of four main parts: A matching task, a training phase, a transfer phase, and a generation task. The matching phase (Figure 1a) was identical in the pre- and the post-matching group, but preceded the training phase in the pre-matching group and followed it in the post-matching group. Participants learned a one-to-one mapping between letters (A, E, P, S) and colors (blue, red, green, orange). In the control group, the matching task preceded the training phase but contained no one-to-one mapping between letters and colors. Yet, in all groups, participants observed that the letter was consequently colored in the respective color. The matching phase consisted of two blocks with 48 trials each.

The subsequent training phase was identical for all participants (Figure 1b). Before the training, they were given 16 practice trials in which the distractor letters were white L-shaped letters. The training consisted of 20 blocks of 40 trials each. In each block, each distractor shape was presented ten times, and the target letter was equally often the short or long F, and equally often rotated 90° or 270°. Participants were given the opportunity to take a short self-paced break between blocks.

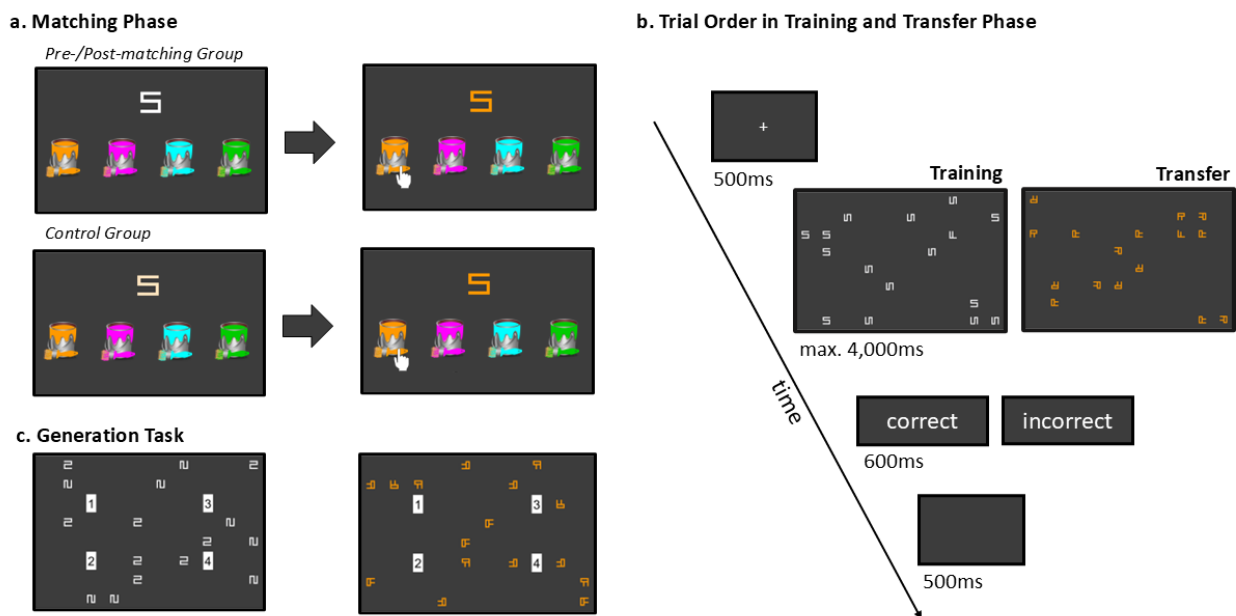
In the transfer phase (Figure 1b), participants were informed that the distractors were no longer letters but the same letter in distinct colors. Apart from that, the transfer phase followed the same procedure as the training phase and consisted of 4 blocks of 40 trials each.

In the generation task (Figure 1c), participants were presented with search displays similar to the ones of the training and transfer phases. It contained two blocks of 48 trials each, so that participants saw twelve search displays of each of the four shapes and colors from the

training and transfer phases. Shape and color displays were presented block-wise and the block sequence was alternated between participants. After each generation task trial, a visual scale from 1 (labeled “complete guess”) to 4 (labeled “absolutely certain”) appeared on the screen, and participants were asked to indicate their confidence concerning their target placement, again using the number keys on their keyboard.

Figure 1

Procedure of the Matching, Training, and Transfer Phase and Exemplary Displays of the Generation Task



Note. a. In the pre- and the post matching group, participants' task was to figure out which letter goes with which color by clicking the paint pots. The letter appeared in the respective color when they selected the correct one. Otherwise, nothing happened. In the control group, each letter appeared lightly tinted in one of the four colors. Participants had to click on the respective paint pot with the more saturated version of the same color. b. In the training/transfer phases, each trial started with the presentation of a fixation cross for 500ms, followed by the presentation of the search display for max. 4,000ms or until participants' response. The response window started with the appearance of the search display and lasted 5,000ms or until the response.

The trial ended with a feedback text (“correct” or “incorrect”) for 600ms. A blank inter-trial-interval of 500ms followed. Participants were instructed to search for the target letter “F” among distractor letters and to decide if the second bar of the F was short (S-key) or long (L-key). They were instructed to do this with their index fingers and as quickly and accurately as possible. c. In the Generation Task no target letter was presented. Instead, the four potential target locations were marked with the numbers 1-4. Participants were instructed to indicate the target location they think the target letter was presented before using the number keys on their keyboard.

Data analysis. The statistical analysis was conducted in *R Statistical Software* (version 4.1.0; R Core Team, 2021). We used the *dplyr* package for most data manipulation (Wickham et al., 2023), the *lme4* package for fitting models (REML; Bates et al., 2015) with restricted maximum likelihood (REML) model fit, and the *lmerTest* package (Kuznetsova et al., 2017) with the Satterthwaite's method for *t*-tests. Note that for the χ^2 tests for model comparisons, the models were refitted using maximum likelihood (ML). Graphs were created with the *ggplot2* package (Wickham, 2016). The study's design and analysis were not pre-registered. Data and analysis scripts are available on OSF.

For the training and transfer phases, we excluded incorrect trials from analysis. Additionally, in the training, we excluded trials with a target location repetition, as we expected shorter RTs for trials in which the target location is repeated from trial $n - 1$ according to the intertrial priming effect (Golan & Lamy, 2024; Kabata & Matsumoto, 2012). In the transfer phase, there were no target location repetitions. We did not conduct any outlier analyses excluding additional trials to avoid biases or power issues (Miller, 2023).

For the RT analyses, we fitted mixed-effects models with pre-defined fixed effects, and random effects selected based on the data (Huta, 2014; Weinfurt, 2000) to keep the Type I error probability minimal (Matuschek et al., 2017; Stroup, 2013). The two fixed effects were trial type (regular or deviant target location) and block (as time variable). We compared various

potential random effect structures based on the AIC (Akaike, 1998). The results from all potential random effect structure models and model comparisons are accessible via the analysis script uploaded to OSF.

Results

Matching phase. Descriptive statistics per group and block are shown in Table 1. For the pre-matching group, a paired, one-tailed t -test showed that RTs in block 2 of the task were significantly faster than in block 1, $t(56)=4.61$, $p<.001$, $d=.61$, as was the case in the post-matching group, $t(57)=4.99$, $p<.001$, $d=.66$. The same effect was shown in the control group, $t(56)=3.10$, $p=.001$, $d=.41$. Due to the nature of the task, the control group was overall faster in their responses than the post-matching and the pre-matching groups ($p<.001$ in a post-hoc, Bonferroni corrected t -test). The latter groups did not differ significantly ($p=1$ in a post-hoc, Bonferroni corrected t -test).

Training. Descriptive statistics per group are displayed in Table 1. The learning curves over blocks, separately displayed for regular and deviant trials, are displayed in Figure 2. Due to incorrect and target location repetition trial exclusion, we trimmed 7.48% of the data.

First, analyzing data across all groups, a model with group as factor was not a better fit ($AIC=1907509$) than the null model ($AIC=1907505$), $\chi^2(2)=0.227$, $p=.893$, but models with block ($AIC=1904433$) or with trial type ($AIC=1907486$) were ($p<.001$). Adding both factors to the model, the interaction of trial type and block yielded a significantly better fit ($AIC=1904400$) than the additive model ($AIC=1904415$), $\chi^2(1)=16.613$, $p<.001$. With the interaction model for the fixed effects, we tested random effect structures against each other, but a random slope (per trial type) and intercept (per subject) model produced a singular fit. The random intercept model showed a non-significant effect of trial type ($b=4.98$, $SE=4.84$, 95% CI [-4.52, 14.47], $p=.304$), but a significant effect of block ($b=-22.51$, $SE=0.45$, 95% CI [-23.40, -

21.62], $p < .001$), as well as a significant interaction between trial type and block ($b = -3.71$, $SE = 0.91$, 95% $CI [-5.49, -1.92]$, $p < .001$). The random intercept for subject was significant as well ($\tau_{00 \text{ subject}} = 36864.32$).

Table 1*Descriptive Statistics per Group and Phase*

	All groups	Pre-Matching group	Post-Matching group	Control group
Sample size	172	57	58	57
Matching Phase				
Block 1 response time (ms)	1791.23 (710.80)	1976.66 (700.35)	1912.66 (648.35)	1482.23 (690.62)
Block 2 response time (ms)	1496.22 (531.06)	1661.69 (537.93)	1581.44 (566.39)	1244.04 (381.01)
Training Phase				
Response time (ms)	1147.41 (480.71)	1145.85 (482.22)	1156.46 (489.52)	1145.98 (479.89)
Accuracy (%)	97.10 (16.77)	97.02 (17.00)	96.90 (17.32)	97.39 (15.96)
Transfer Phase				
Response time (ms)	1134.01 (498.83)	1155.51 (527.65)	1114.78 (483.27)	1132.09 (483.76)
Accuracy (%)	96.48 (18.42)	96.03 (19.52)	96.67 (17.94)	96.74 (17.75)

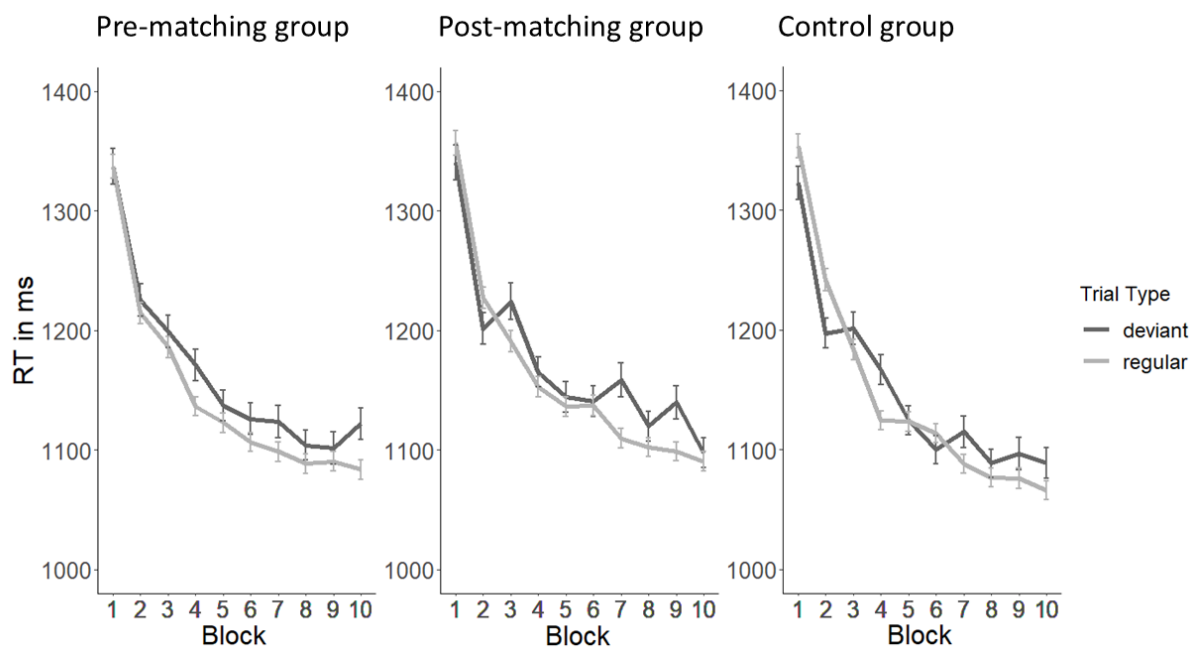
Note. Mean response times and accuracies per group. Numbers in brackets indicate standard deviations.

Second, we fitted models to the data of the individual groups. For the pre-matching group, an additive model of block and trial type with random intercepts for participants yielded the best fit. In the post-matching and control groups however, an interaction model with block and trial type and random intercepts yielded the best fit. Models with a more complex random effects structure yielded a singular fit. The parameters for the models for each of the three groups are reported in Table 2. The full breakdown of model comparisons is available in the

supplementary material on OSF. Concluding from this analysis, we observe learning effects in all three groups in the form of a significant effect of trial type. In the pre-matching group, the main effect of trial type is significant, in the post-matching and control group, trial type is significant in interaction with block. As can be seen from the response time trajectories, in the latter groups, the trial type effect in the first two blocks is almost reverse, resulting in the significant interaction effect.

Figure 2

Response times of the Training Phase per Block and Trial Type



Note. The graph depicts the response times for deviant and regular trials for the ten blocks of the training phase. Note that the block is coded as block -1 for better interpretation. Error bars indicate standard errors.

Table 2*Response Time Analysis Results for the Training Phase per Group*

	<i>Estimate</i>	<i>SE b</i>	95% CI <i>b</i>		df	<i>t</i>	<i>p</i>
			<i>LL</i>	<i>UL</i>			
Overall							
N=172							
Intercept	121257.86	14.84	1228.78	1286.95	178	84.77	<.001
Trial Type	4.98	4.84	-4.52	14.47	127138	1.03	.304
Block	-22.51	0.45	-23.40	-21.62	127138	-49.51	<.001
Trial Type * Block	-3.71	0.91	-5.49	-1.92	127138	-4.08	<.001
ICC=0.15							
Pre-matching Group							
N=57							
Intercept	1255.56	25.54	1205.50	1305.62	58	49.16	<.001
Trial Type	-18.19	4.54	-27.09	-9.28	42099	-4.00	<.001
Block	-22.19	0.72	-23.61	-20.77	42099	-30.61	<.001
ICC=0.17							
Post-matching Group							
N=58							
Intercept	1262.91	28.49	1207.07	1318.76	59	44.33	<.001
Trial Type	6.69	8.37	-9.71	23.10	42791	0.80	.424
Block	-21.56	0.78	-23.10	-20.02	42791	-27.45	<.001
Trial Type*Block	-4.06	1.57	-7.14	-0.98	42791	-2.58	.010
ICC=0.20							
Control Group							
N=57							
Intercept	1256.77	23.04	1211.62	1301.92	59	54.56	<.001
Trial Type	17.45	8.37	1.04	33.86	42243	2.08	.037
Block	-24.19	0.79	-25.73	-22.65	42243	-30.78	<.001
Trial Type*Block	-5.06	1.57	-8.14	-1.98	42243	-3.22	.001
ICC=0.14							

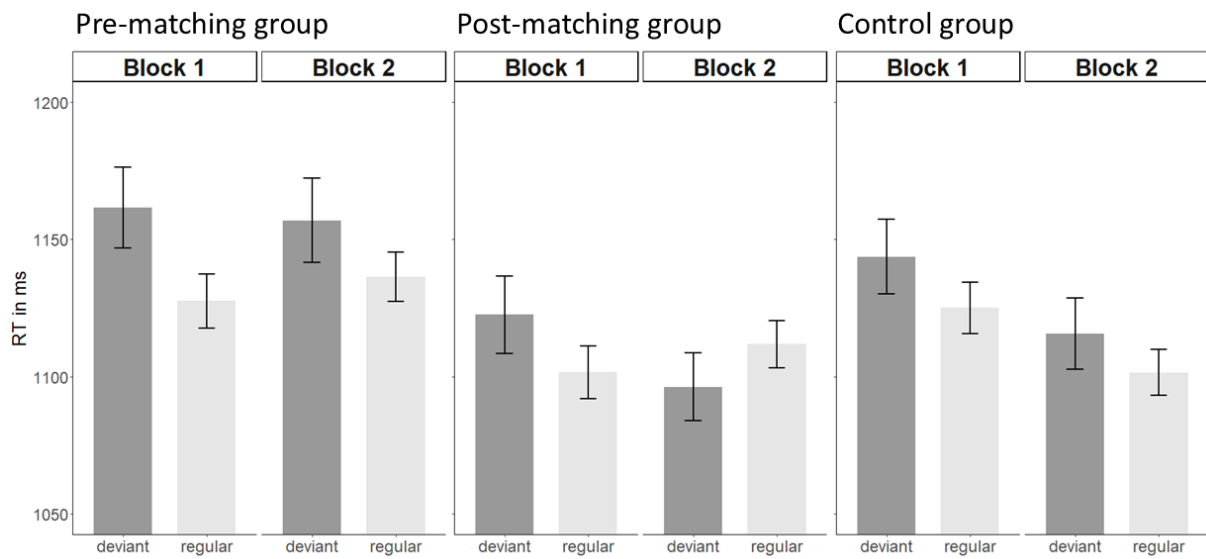
Note. *SE* = standard error; *CI* = confidence interval; *LL* = lower limit; *UL* = upper limit. Mixed-effects model computed coefficients, standard errors and confidence intervals for the coefficients, degrees of freedom, *t*-values, and *p*-values are displayed.

Transfer phase. The exclusion of incorrect responses and the first 15 trials of the transfer phase yielded a 12.58% trimming of the data.

First, analyzing data across all three groups, a model with group as factor was not a better fit ($AIC=359859$) than the null model ($AIC=359856$), $\chi^2(2)=0.846$, $p=.656$, but a model including trial type as factor ($AIC= 359851$) was, $\chi^2(1)=6.48$, $p=.011$. The interaction of trial type and block did not yield a better model fit, $\chi^2(1)=1.80$, $p=.179$.

Figure 3

Transfer Response Times per Group, Block, and Trial Type



Note. Error bars indicate standard errors.

The response times per group, block, and trial type are displayed in Figure 3. We analyzed the factors block and trial type in mixed effects models per group. For the pre-matching group, a random intercept model with only trial type yielded the best fit, revealing trial type as significant factor ($b=-25.04$, $SE=10.59$, 95% CI $[-45.80, -4.28]$, $p=.018$, $\tau_{00\ subject} =$

70714.19). For the post-matching group, neither the factor block, nor the factor trial type improved the model fit. To compare with the pre-matching group, in a random intercept model with only trial type, trial type was not a significant factor ($b=-1.43$, $SE=9.93$, 95% CI [-20.88, 18.03], $p=.0886$, $\tau_{00\ subject} = 54412.89$). Lastly, in the control group, a random intercept model with block yielded the best fit, with block as significant factor ($b=-26.86$, $SE=9.34$, 95% CI [-45.17, -8.54], $p=.004$, $\tau_{00\ subject} = 45087.63$). To compare with the other two groups, in a random intercept model, trial type was not a significant factor ($b=-18.79$, $SE=10.05$, 95% CI [-38.48, -0.90], $p=.061$, $\tau_{00\ subject} = 45080.06$). For a more direct comparison, we also computed the full, additive model with random intercepts for all three groups (Table 3).

Generation Task Analysis. For the generation trials, we tested mean accuracy against chance level (25% because of four response alternatives), computed mean confidence, and the relative frequencies of correct responses under the condition of high confidence (correct|high) and low confidence (correct|low). We summarized confidence ratings of 1 and 2 as low, and ratings of 3 and 4 as high confidence. Participants with explicit knowledge of a pairing between cue and target location should be able to make a metacognitive assessment of their knowledge (Haider et al., 2011; for a discussion of this method see Tavera & Haider, 2024). Therefore, we tested the relative frequencies of correct|high and correct|low against each other in a paired, one-tailed t -test.

We conducted the same analysis for each group, and we computed correlations between mean accuracy and mean confidence per participant, and observed no significant correlations, neither for the shape cue, nor for the color cue. The results are summarized in Table 4.

Table 3*Response Time Analysis Results for the Transfer Phase per Group*

	<i>Estimate</i>	<i>SE b</i>	<i>95% CI b</i>		<i>df</i>	<i>t</i>	<i>p</i>
			<i>LL</i>	<i>UL</i>			
Pre-matching Group							
N=57							
Intercept	1104.96	119.22	871.25	1338.67	4047	9.27	<.001
Trial Type	-25.15	10.60	-45.92	-4.38	7878	-2.374	<.018
Block	3.96	9.86	-15.36	23.28	7878	0.40	.688
ICC=0.27							
Post-matching Group							
N=58							
Intercept	1136.90	111.31	918.71	1355.10	4745	10.21	<.001
Trial Type	-1.35	9.93	-20.81	18.11	8065	-0.136	.892
Block	-2.66	9.26	-20.81	15.49	8065	-0.29	.774
ICC=0.24							
Control Group							
N=57							
Intercept	1426.38	111.60	1207.62	1645.14	5371	12.78	<.001
Trial Type	-18.05	10.05	-37.74	1.64	7938	-1.80	.072
Block	-26.42	9.35	-44.73	-8.10	7938	-2.83	.005
ICC=0.21							

Note. Additive Models with random intercepts for subjects. *SE* = standard error; *CI* = confidence interval; *LL* = lower limit; *UL* = upper limit. Mixed-effects model computed coefficients, standard errors, and confidence intervals for the coefficients, degrees of freedom, *t*-value, and *p*-value are displayed.

Table 4*Results for the Three Groups in the Generation Task*

	<i>Accuracy</i> <i>Mean % (SD)</i>	<i>t-test against</i> <i>chance (25%)</i>	<i>Confidence</i> <i>Mean (SD)</i>	Correlation ac- curacy and con- fidence	Relative frequencies correct high vs. correct low	t-test correct high > correct low
Pre-matching Group						
Shape	25.44 (0.06)	$p = .300$	1.63 (0.83)	0.13 ($p=.330$)	14.03 vs. 25.57	$p = 1$
Color	25.58 (0.08)	$p = .281$	1.62 (0.82)	-0.02 ($p=.908$)	12.62 vs. 26.27	$p = 1$
Post-matching Group						
Shape	24.35 (0.07)	$p = .765$	1.61 (0.91)	-0.18 ($p=.165$)	12.55 vs. 23.70	$p = 1$
Color	23.71 (0.07)	$p = .553$	1.59 (0.88)	-0.06 ($p=.677$)	17.52 vs. 23.21	$p = .950$
Control Group						
Shape	25.26 (0.08)	$p = .404$	1.60 (0.88)	0.17 ($p=.200$)	12.18 vs. 24.48	$p = 1$
Color	24.89 (0.06)	$p = .553$	1.58 (0.84)	-0.13 ($p=.342$)	14.84 vs. 25.18	$p = 1$

Note. Per group and visual cue, the table shows mean accuracy in percent, and the p -value when it is tested against chance performance (25%). Then, it shows mean confidence, and the correlation coefficient between accuracy and confidence, as well as the p -value from a Pearson correlation test. Lastly, it contrasts relative frequencies of correct responses given high confidence, and the relative frequencies of correct responses given low confidence, and the p -value of testing the values against each other in a paired, one-tailed t -test.

General Discussion

This study aimed to test the constraints of transfer of implicitly learned contingencies. In line with the assumption of modules processing information of one dimension (Keele et al., 2003), transfer should not be observed if the dimension changes from training to transfer. Here, we tested whether implicitly learned contingencies between a visual feature (shape) and a target location can be transferred to another visual feature (color) only when participants learn the association between these features. Accordingly, participants learned either a shape-color

association beforehand, after the training, or not at all. We implemented this in an adapted contextual cueing paradigm (Tavera & Haider, 2025), in which participants could learn to predict the location of a target by the distractor shapes with a 70% probability.

We showed that participants in all groups learned the contingencies between shapes and target locations in the training phase. In the transfer phase, participants in the pre-matching group showed contingency knowledge between colors and target locations whereas the post-matching group and the control group did not. In all groups, the awareness tests indicate that participants did not have explicit knowledge of either contingency between one of the visual features and target location.

These results have interesting implications for frameworks that formulate mechanisms of implicit learning. One concept that is found in the global workspace theory (Baars, 2005), the dual-system model (Keele et al., 2003), and other theories on learning and consciousness, is that implicit learning operates through encapsulated modules that are defined by processing one particular dimension (Eberhardt et al., 2017; Hommel, 2009). Accordingly, more integrated and distributed processing of information would, at some point, require conscious processing. Our findings, however, represent a case in which implicitly learned response contingencies are transferred from one cue to another under the condition that the processing of the two cues is linked before learning.

Similarly, Haider et al. (2020) showed transfer from a visual screen location sequence to a motor response key sequence if the screen locations were matched to the keys. However, the spatial dimension might be distinctly represented in our cognitive system (U. Mayr, 1996; Paillard, 1991) and might even be considered a non-perceptual dimension as it is so closely bound to our motor system (Gaschler et al., 2012; Goschke & Bolte, 2012; I. Koch & Hoffmann, 2000). Hence, participants in Haider et al. (2020) might have learned a sequence of eye movements from the visual sequence on the screen, and transferred it to a sequence of finger

movements. This learning might thus not have been a transfer between features, but a learning process that remains coded in the spatial dimension. What is new in our findings is that we show transfer between two perceptual features. In the training phase, participants learn a contingency between shape and target location, and then transfer this knowledge to a contingency between color and target location. So, the transfer is between two perceptual features, shape and color.

We used the contextual cueing paradigm as an interesting alternative candidate to sequence learning. There are some essential differences between the two paradigms, and although both paradigms are commonly used to study implicit learning, the literature seems to be quite separated. Contextual cueing is cross-dimensional learning to begin with, as participants learn contingencies between a visual feature and a spatial target location. It is then an interesting question whether such cross-dimensional contingencies can be transferred to other features. Sequence learning is a kind of chaining across trials (Schuck et al., 2012), where one event has to be linked to a subsequent (and preceding) one. For example, in a location sequence, locations are interlinked (U. Mayr, 1996). In the contextual cueing paradigm, the contingency across two features (e.g., spatial configuration and location) are learned within a trial. So, unlike in sequence learning, where for example one color should activate the subsequent color, in contextual cueing, one color activates a target location, and not another color. We do not have a reason to assume that these two paradigms test different learning mechanisms. But in terms of the description of features as the building blocks of implicit learning (Eberhardt et al., 2017), it is an important step that we cannot only show transfer between perceptual and motor information (Haider et al., 2020), but also transfer across perceptual features.

Our findings are compatible with two not mutually exclusive frameworks: the theory of event coding (Hommel et al., 2001), and sensory preconditioning (Holmes et al., 2022). Note that both frameworks are not tailored to modelling specifically implicit knowledge, and for

preconditioning effects in particular, it has been suggested that it requires conscious representation of all associations involved in the process (Arunkumar et al., 2024). Also, the methodology around both frameworks deviates from the paradigm used in this study, as it typically examines effects from trial n to trial $n+1$ (Hommel, 2004; but see J. R. Schmidt et al., 2020) or short time frames (Holmes et al., 2022). In contrast, in our paradigm, learning effects are observed over a rather lengthy period of time. Applying mechanisms of the event file framework to our findings, we could assume that event files are built in the matching phase. Only in the pre-matching group, these contain an association between the two features shape and color. In the training phase, when a shape predicts a target location, participants not only learn these shape-location associations, but the respective shape would also co-activate the associated color. In the transfer phase, the target-location would then be cued by color just as well as by shape. Note that this mechanism requires a preceding feature (shape-color) matching. This is in line with our findings regarding the post-matching group. Because there is no shape-color association, no associative strength between color and the target location can be developed through the course of the training phase. This refines the mechanism, showing that a post-training stimulus-stimulus association is not integrated into an event file in a way that the second feature is then also associated with the target location. This would explain why the post-matching group does not transfer the response association to the second feature.

The second framework that might apply here, especially because of the similarity in its methods and research questions, is preconditioning. In a typical preconditioning procedure, two stimuli, $S1$ and $S2$, are associated with each other ($S1-S2$). Then, one of them is associated with a response ($S1-R$). In a transfer phase, it is then shown that the second stimulus is also associated with the response ($S2-R$), although it has never been paired with it (e.g., Arunkumar et al., 2024). In our study, we would term the shape feature $S1$, and the color feature $S2$, and the target location would be R , in the sense of an eye-movement toward the target location. The

preconditioning phenomenon can be explained through an online integration account, suggesting that the *SI-R* association is bound into the *SI-S2* association (Holmes et al., 2022). This mechanism however does not predict an effect of the order of association learning, in that it requires an *SI-S2* association specifically before the training phase. It would thus not explain the difference between our pre-matching and the post-matching group.

One last important issue that is relevant for the interpretation of our results is the assumption that we test implicit learning effects here. Many studies using the contextual cueing paradigm have been criticized for not reliably testing for conscious knowledge (Luque et al., 2017; Vadillo et al., 2016; Vadillo et al., 2019). Therefore, we changed the common contextual cueing protocol (Chun & Jiang, 1998), having few trials of a recognition task, into a task that is maximally similar to the training. This should improve sensitivity of our knowledge test, and the increase in trials should improve its reliability (Smyth & Shanks, 2008). Also, we combine this objective measure with a subjective confidence measure, and test their interdependence (Michel, 2023). As we find no evidence for metacognitive sensitivity, that is, confidence in correct responses as indication for explicit knowledge, we propose that what we observed here is indeed implicit knowledge.

Conclusion

In this study, we have shown that implicitly learned contingencies are transferrable from one visual cue to another. This transfer seems to be dependent on an association between the two visual cues formed before training, as a control group and the post-matching group that did not acquire this association in advance to the training phase did not show the transfer effect. These are relevant findings for the study of the cognitive architecture that enables implicit learning. As discussed, it has implications for models proposing implicit learning mechanisms and also for models of consciousness. Not all models predict the amount of information integration in the absence of conscious awareness that we demonstrate with the here reported knowledge

transfer effect. Frameworks like the theory of event coding might entail such processes. In addition, the findings extend such frameworks of action control by showing that learned associations between features had comparably long-lasting effects on performance.

Appendix C

The following manuscript corresponds to Study 3. The manuscript is included as part of the cumulative dissertation.

Title:

Implicit learning of low-level and semantic cues in an adapted contextual cueing paradigm

Authors:

Felice Tavera, Galla Abdarahaman, & Hilde Haider

Note:

References cited in this appendix are included in the general reference list. To enhance readability and avoid redundancy, separate reference lists have not been provided for individual appendices.

**Implicit learning of low-level and semantic cues
in an adapted contextual cueing paradigm**

Felice Tavera¹, Galla Abderahaman¹, & Hilde Haider¹

¹University of Cologne

Author Note

Felice Tavera  <https://orcid.org/0000-0003-0994-6620>

Galla Abderahaman

Hilde Haider  <https://orcid.org/0000-0001-7293-3166>

The authors have no known conflicts of interest to disclose.
Ethics approval has been obtained from the Ethics Commission of the Faculty of Human Sciences of the University of Cologne (reference number FTHF0209).
Study materials can be obtained upon request, primary data have been made available at the Open Science Framework (OSF) and can be accessed at https://osf.io/36ypv/?view_only=c8c2e09ad16741879b2c51b2f6b37c48

Corresponding author:

Felice Tavera

Department of Psychology

University of Cologne

Richard-Strauss-Str. 2

50931 Köln, Germany

Phone: +49-221 470-1471

Email: felice.tavera@uni-koeln.de

Abstract

When we search for something in our everyday environment, our attention is often guided by previously learned contingencies between objects and their typical locations. Such learning processes help us navigate complex scenes more efficiently. In this study, we investigated implicit and explicit learning of low-level (color) and semantic (scene category) cues in an adapted contextual cueing paradigm using complex real-world scenes. In Experiment 1, we found that contingencies between dominant colors of a real-world scene and target locations were not learned. In Experiment 2, we found that semantic scene categories were implicitly learned as predictive cues, however, the contextual cueing effect was reversed. Response times were longer for predictable than for unpredictable target locations. Interestingly, this reversed effect also emerged in Experiment 3, where participants were explicitly informed about the contingencies between scene category and target location. To assess explicit knowledge in all experiments, we implemented an objective generation task combined with a confidence measure. Our results support the validity of our test, as we found indicators of explicit knowledge only in the explicit learning condition (Experiment 3). These findings demonstrate the potential of the adapted contextual paradigm to investigate both low-level and semantic cue learning under implicit and explicit learning conditions. To account for the reversed contextual cueing effects, potential mechanisms such as attentional inhibition or episodic retrieval are discussed.

231 words

Keywords: Implicit learning, contextual cueing, visual search, attention

When we are looking for our favorite coffee mug in the morning, there are different strategies that we can use. We can broadly look around for anything of the mug's color. At the same time, we will most likely look for it in the kitchen counter, not so much in the bathroom. This simple example shows that our attention in visual search is guided by different mechanisms (Itti & Koch, 2000). In the lab, there has been extensive research on complex real-world scene processing (for a review, see Malcolm et al., 2016) and visual search within them (for a review, see Wolfe, 2020) to investigate such mechanisms of attentional guidance in complex environments. On the one hand, attentional guidance by saliency maps (Itti & Koch, 2000; Itti et al., 1998) have been the most prominent approach to explain shifts of attention on the basis of low-level features (e.g., color, intensity, contrast). Yet, there is evidence that attention is immediately guided to the most informative locations within a scene, regardless of low-level saliency (e.g., Mackworth & Morandi, 1967). Meaning maps represent such semantic guidance by mapping out the most semantically diagnostic locations within a scene (Henderson & Hayes, 2017). They have been found to explain variance in eye movement during scene viewing beyond saliency maps, as attention allocation to salient image features can be suppressed in favor of semantically relevant locations (Hayes & Henderson, 2019b). Still, it is difficult to distinguish the influence of low-level features (saliency) from the influence of semantic content. Because semantically relevant objects in a scene tend to be different with respect to low-level features as well, findings that suggest an influence of saliency could also be interpreted in favor of semantic guidance (Henderson, 2007; Henderson et al., 2007). However, meaning maps are based on human ratings of the “meaningfulness” of scene parts for the whole scene (e.g., Henderson & Hayes, 2017), and do not specify underlying cognitive processes. Thus, the mechanisms of how meaning maps are internally constructed and activated when confronted with a novel scene remains unclear. In the present study, we aim to investigate whether and how semantic knowledge about our visual environment is learned to guide attention in visual search.

What we know is that semantic scene processing is possible in very short time frames (Oliva, 2005). For example, with presentation times as short as 26ms, participants were able to categorize scenes into natural and human-made with more than 90% accuracy (Joubert et al., 2007; Rousselet et al., 2005).

Such findings however are ambiguous to whether such categorizations are driven solely by basic low-level feature analyses, or whether they require the integration of these features into a higher-level semantic understanding (Kotabe et al., 2016). If such integration and high-level processing are indeed necessary, this raises an important question: Can such processes occur in the absence of conscious awareness? Some theories of consciousness suggest that integrating information is a core function of consciousness (e.g., Dehaene & Naccache, 2001; Tononi, 2004). To test this hypothesis, there has been extensive research on whether semantic scene processing can occur without awareness. For instance, this has been tested rendering scenes invisible by only briefly presenting and masking them, or implementing continuous flash suppression techniques (e.g., Mudrik, Breska, et al., 2011; Mudrik & Koch, 2013). However, evidence for semantic processing of scenes in the absence of awareness has been called into question and was partly not replicable (e.g., Biderman & Mudrik, 2018; Moors et al., 2016). As a result, the issue is still debated. In the present study, we aim to determine whether low-level feature cues or semantic information can be learned *implicitly* to guide attention in visual search.

To address this question, we adapted the contextual cueing (CC) paradigm (Chun & Jiang, 1998), which combines visual search with contingency learning. In the CC paradigm, participants perform a visual search task. Commonly, they are asked to find a target letter among distractor letters on a display, and respond to a feature of the target. In some trials, unbeknownst to participants, a configuration of distractor letters is repeated, and therefore, the target position can be predicted from the configuration. In the same way, contingencies between the color or

shape of the distractors, and the target position, can also be learned (Tavera & Haider, 2025). Going beyond simplistic stimulus material, it has also been shown that real-world scenes can be learned to cue the target location (Brockmole & Henderson, 2006b). Thus, this approach allowed us to examine the effect of contingency learning on visual search performance with complex, real-world scenes. In the present study, the contingencies were implemented between either a low-level (color) or a semantic feature (scene category) of a real-world scene, and target location. Importantly, we also manipulated the type of learning, distinguishing explicit and implicit learning conditions. This allows us to test whether semantic processing can be involved in implicit learning processes. That would not be expected if awareness is required when learning processes include higher-level integrated information such as semantic scene categories.

There are already attempts to investigate semantic processing within implicit learning. As noted, the CC paradigm has been used with real-world scenes. Instead of repeating distractor configurations, scenes were repeated and thus predictive of target locations (Brockmole et al., 2006; Brockmole et al., 2008; Brockmole & Henderson, 2006a, 2006b; Henderson et al., 2007). Interestingly, these studies found that participants acquired explicit contingency knowledge right from the first repetition of a scene. So, although these findings were a proof of concept that the CC paradigm also produces reliable effects with real-world scenes, it did not contribute to the question of semantic processing in the absence of awareness. But with a slightly different approach, other literature has offered evidence for that.

For instance, there is a line of research with implicit visual sequence learning (Nissen & Bullemer, 1987). It has been shown that a visual sequence of abstract semantic categories (e.g., furniture, clothing, and animals) can be learned (Goschke & Bolte, 2007; Brady & Oliva, 2008). Although these are important findings, both studies may not have unequivocally demonstrate *semantic* and *implicit* processing. First, there are low-level features that are categorically different between the categories of the sequences. To give a few examples, furniture has more

straight lines and right angles than animals, which in turn almost all have four legs, whereas many clothing items have two sleeves. Goschke and Bolte (2007) did not address this issue, but Brady and Oliva (2008) implemented an additional test. To exclude the possibility that participants learned the sequence according to the predictability of low-level features, they conducted a test phase in which they transferred the visual sequence to a sequence of words naming the categories. Again, participants recognized sequential material with above-chance level accuracy, however, the learning effect on performance was no longer significant ($p=.05$, no effect size reported; Brady & Oliva, 2008).

In addition, there is an additional methodological issue. To assess explicit knowledge, participants were asked two questions: Whether they could identify “any patterns” (Brady & Oliva, 2008, p. 680), and whether they would be able to report what image category would follow the image of a mountain. In all four experiments, all participants were classified as unaware. This awareness test is problematic, because the first question is vague and ambiguous in terms of interpretations of “patterns”, and the second question asks for only one out of twelve contingencies. The probability to find explicit learning that might have occurred, is thus low (Shanks & St. John, 1994). In Goschke and Bolte (2007), for an awareness test, participants were asked to freely reproduce the category sequence, and to then categorize four sequences into old and new sequences. Participants were then excluded when they had higher hit rates than false alarm rates. Goschke and Bolte (2007) still found sequence learning for participants that did not show conscious awareness. This measure of awareness is certainly elaborate in terms of design and analysis. However, its reliability is questionable due to the small number of trials per participant (Vadillo et al., 2016).

Two studies that were conducted shortly after, tested whether abstract semantic knowledge could be used in the CC paradigm (Goujon, 2011; Goujon et al., 2009). Goujon et al. (2009) showed CC effects for word displays, in which the word categories (e.g., mammals,

birds, fruits/vegetables) predicted target locations. They also claim that this knowledge remained implicit, on the basis of their two-step awareness test. First, they asked whether participants noticed an association between context and target location. In the second step, they then did not test for contingency knowledge but solely asked participants for “*déjà-vu*” experiences when showing them predictive and counter-predictive word displays. This second test is very far from the tasks’ demand in training, and from a valid measure of contingency knowledge, since they do not ask about the contingencies to begin with. Also, performance in the “*déjà-vu*” task is difficult to interpret, given that by then, all participants already had knowledge about the prevalence of contingencies because of the question in the first step.

In her subsequent study, Goujon (2011) went one step further and did not test word categories, but scene categories as predictive cues in a CC task. There were eight categories of rooms that were either predictive or not predictive of the eight possible target locations. She did not find a CC effect unless participants either had done a scene categorization task beforehand, or the scene had first been presented for 1,500ms without the target first in each trial. She argues that this led to an enhancement of semantic processing of the scene, and thus to a semantic contextual cueing effect. Yet, this explanation is hardly consistent with the scene processing literature that shows how quick, automatic, and involuntary scene categorization occurs (Hayes & Henderson, 2019b; Joubert et al., 2007; Rousselet et al., 2005). Goujon (2011) then argues that the contingency knowledge remained implicit in her experiments, following tests similar to Goujon et al. (2009), which are problematic for the reasons discussed above. So, the results allow an alternative explanation: Participants who were simply made aware of the contingencies by emphasizing the semantic categories, could then learn the contingencies explicitly. However, the awareness tests did not detect this explicit knowledge because of their methodological shortcomings. Another explanatory factor concerning the first experiment in which Goujon (2011) did not find a CC effect might be the complexity of the design. Participants had to

learn that there were eight possible target locations paired with predictive and unpredictable scene categories. Especially problematic might be the second step, to differentiate predictive and unpredictable categories. Given the eight target locations, the probability for each target location is $1/8$. The four predictive scenes were 100% contingent with one target location respectively. In contrast, the unpredictable categories were not entirely unpredictable, because by design, they were shown with only four of the eight target locations. Thus, the probability of target locations was $1/4$, instead of $1/8$. So, there is some predictive value also to the allegedly unpredictable categories. This makes it a priori less likely to obtain a sufficiently large response time difference between these two conditions, given that in both, there should be a response time advantage.

Building on the literature reviewed above, we used material of abstract semantic categories in the present study, but were careful to keep low-level features as similar as possible. We did so by instantiating the semantic scene categories as categories of rooms in a house. That enabled us to make the rooms similar in terms of low-level features across categories. Further, we meticulously measured contingency awareness using not only a well-established, objective task (Chun & Jiang, 2003) but also a confidence measure to additionally assess metacognitive knowledge (Michel, 2023a; Tavera & Haider, 2025). Also, we aimed to reduce complexity of the learning design to examine whether this might have been a relevant factor in not finding a CC effect in Goujon (2011).

Overview over the experiments. In all experiments, we used the same complex, real-world scenes of room categories. To control for the influence of low-level features, each particular scene predicted the target location with a probability of 70%. In the remaining 30% of the trials, the target could occur at the other target locations with the same probability. In Experiment 1, color cues within the scenes predicted target location with a probability of 70%. Participants were not instructed about the contingencies. In Experiment 2, the semantic scene

category predicted target location, also with a probability of 70%, and participants were also not instructed about the contingencies. Experiment 3 was a replication of Experiment 2, except for explicit instructions of the contingencies between scene categories and target locations. In all three experiments, we then tested whether participants learned the contingencies, and whether this was reflected in their response times to predicted versus unpredicted target locations. We then additionally tested whether this knowledge was explicit or implicit.

Method

Participants. In Experiment 1, we obtained data from 63 participants ($M_{\text{age}}=42.14$, $SD_{\text{age}}=12.79$; 35 female). We excluded one participant due to chance-level performance in training, and excluded incorrect trials and response times >5 seconds. This resulted in a 15.21% trimming of the data. In Experiment 2, we collected data from 60 participants ($M_{\text{age}}=42.7$, $SD_{\text{age}}=11.52$; 43 female, 1 diverse). We excluded incorrect trials and response times >5 seconds which resulted in a 10.02% trimming of the data. In Experiment 3, we had 40 participants ($M_{\text{age}}=42.1$, $SD_{\text{age}}=10.47$; 22 female). We excluded incorrect trials and response times >5 seconds, resulting in a 9.78% trimming of the data.

Stimuli. The search displays were complex, real-world scene photographs. The scenes were presented in the size of 15×10 cm search displays, irrespective of the screen size of participant's computer monitors (see Procedure). Aspect ratio was normalized to 3 x 4. The material was comprised of four categories (bathroom, living room, bedroom, and kitchen). Additionally, the scenes were distinctively colored in one of four colors (blue, white, green, brown) in Experiment 1, and one of six colors (blue, white, green, pink, red, brown) in Experiments 2 and 3. Importantly, all semantic categories were shown in all colors with equal frequency, so that there were no contingencies between a category and a color. Thus, a total of 240 unique scenes were selected for Experiments 2 and 3, and, excluding all pink and red images, 180 scenes were selected for Experiment 1.

Search targets within the scenes consisted of the letters T or L overlaying the scene. They were placed on a 40×40 pixels area, in one of four predetermined locations on an invisible 8×6 grid. To ensure consistent visibility across different scenes and positions, target color was adjusted to the luminance of the target location area. The 50×50 pix area marking the target position was analyzed for its average luminance and RGB composition. Luminance was calculated using a standard formula for perceived brightness that accounts for human photosensitivity to different wavelengths: $E'Y = 0.299 \times E'R + 0.587 \times E'G + 0.114 \times E'B$ (where $E'Y$ represents the weighted luminance, and $E'R$, $E'G$, and $E'B$ denote the red, green, and blue channel values, respectively; Poynton, 2012). The resulting luminance was then compared to a threshold of 128, which is the midpoint of the 8-bit color range (0–255). Values above this threshold were classified as “bright”, and values below or equal to it as “dark”. Based on this classification, the target color was computed as the negative RGB contrast of the background (i.e., each RGB component was inverted), and then adjusted in brightness (i.e., brightened on dark backgrounds and darkened on bright backgrounds) to maximize contrast. This method ensured high and uniform visibility of the targets regardless of the background color and luminance. For a more detailed analysis of the scene statistics, see Appendix A.

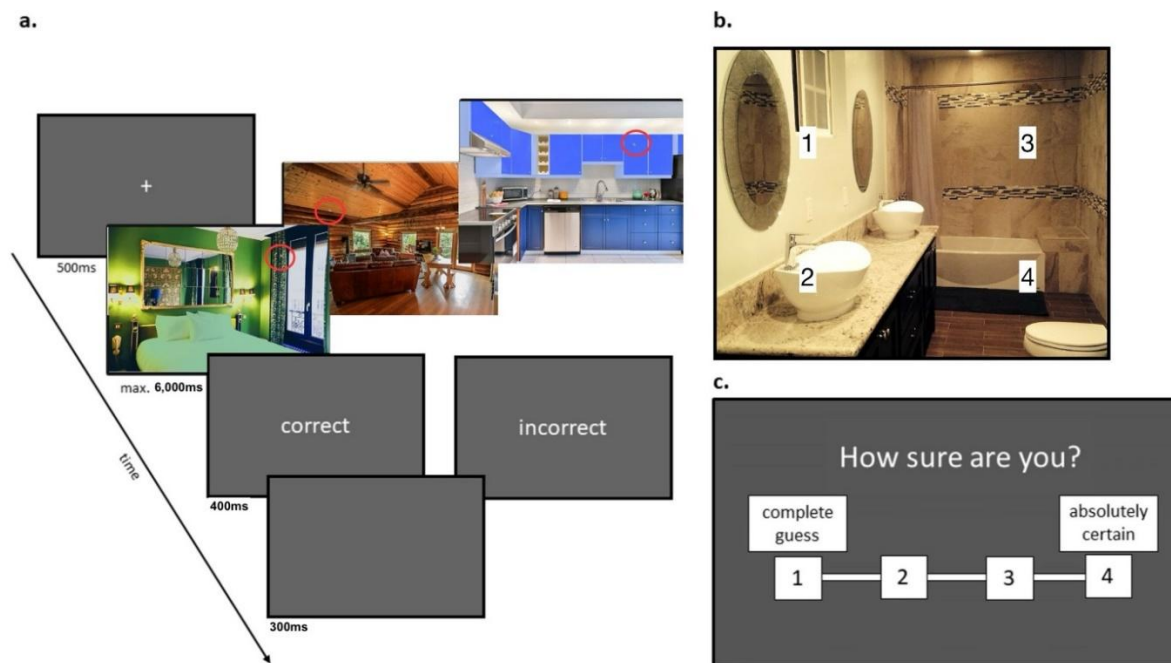
Procedure. All three Experiments were conducted in accordance with the Declaration of Helsinki, and approved by the Ethics Commission of the Faculty of Human Sciences at the University of Cologne. Participants were recruited online via Prolific and were reimbursed according to Prolific’s “ethical reward” standards. The experiments were built in PsychoPy2 (version 2022.5.2; Peirce et al., 2019) and run on Pavlovia (www.pavlovia.org) and SoSci survey (Leiner, 2024).

Participants were first asked to follow a screen scaling procedure (Wakefield Morys-Carter, 2021). With the arrow keys on their keyboard, they were asked to adjust an image of a credit card on the screen to the size of an actual bank or credit card. The size of all stimuli in

the experiment was subsequently presented relative to the scaling provided by participants. This procedure ensured equal size of the search displays for every participant irrespective of the monitor size or aspect ratio.

Figure 1

Overview of the Learning Phase, Generation Task, and Confidence Measure with Example Displays from Experiments 1, 2, and 3



Note. a. Schematic illustration of the trial structure during the learning phase with example search displays. For illustrative purposes, target letters are highlighted in red circles; they were not marked in the actual experiment. b. Example display from the generation task. Participants were asked to indicate the expected target location using the number keys on their keyboard. c. Confidence scale shown after each generation task trial. Participants rated their confidence in their response to the generation task using the number keys on their keyboard.

All three experiments consisted of three parts: A short practice phase, a training phase, and the awareness test. In the end, they were redirected to a short questionnaire.

First, participants were instructed to search for target letters in scenes, and were shown exemplary search displays. In Experiments 1 and 2, instructions were confined to informing participants about the procedure of the experiment, their task, and that target letters would come in different colors, depending on the scene. In contrast, in Experiment 3, participants were also informed about the underlying contingency between scene category and matching target location. It was made explicit which category was associated with which target location, and that the association applied to most scenes.

In the first 4 practice trials, they were shown a scene from each color (Experiment 1) or each category (Experiment 2 and 3). The training phase consisted of four blocks of 40 trials each (Experiment 1), or 60 trials each (Experiments 2 and 3). Trial procedure is displayed in Figure 1a. In each trial, first, a fixation cross was presented for 500ms. Then, the search display appeared for a maximum of 6000ms or until the response. The response window started with the appearance of the search display and lasted 6000ms. Participants were instructed to search for the target letter, pressing the “T” or “L” key on their keyboard with their index fingers as quickly and accurately as possible. All four colors (Experiment 1) or four scene categories (Experiments 2 and 3) respectively predicted one of four target locations with a probability of 70% (predicted trials). In the other 30% of trials (unpredicted trials), target locations were randomly distributed across the remaining three target positions. Importantly, in Experiment 1, scene category had no predictive value, and in Experiments 2 and 3, color was not associated with any target location. The matching between color (Experiment 1) or category, and target location (Experiments 2 and 3) was permuted across participants. The trial ended with a feedback text (“correct”, “incorrect” or “no response”) that appeared on the screen for 400ms and was followed by a blank inter-trial-interval of 300ms (see Figure 1a). Participants were given the opportunity to take a short self-paced break after every block.

After training, we implemented an awareness test using the so-called generation task (Chun & Jiang, 2003) and a confidence measure (Michel, 2023a). Specifically, the test aimed to assess participants' metacognitive awareness about contingencies between color or category, and target location. In 16 trials, novel scenes with no target letter were presented. The four potential target locations were marked with squares with the numbers 1–4 (see Fig. 1b). Participants were instructed to indicate in which location they think the target letter was presented using the number keys on their keyboard. Afterwards, a visual scale from 1 (labeled “complete guess”) to 4 (labeled “absolutely certain”) appeared on the screen (see Fig. 1c), and participants were asked to indicate their confidence in their generation response, again using the number keys on their keyboard.

In the final, short questionnaire, participants were offered to report technical issues, their ideas on the purpose of the study, if and why the task became more difficult or easier, if they had noticed any regularities, and their estimate of the percentage of predicted trials.

Data analysis. Statistical analyses were conducted using the R Statistical Software (R Core Team, 2021), the dplyr package for data manipulation (Wickham et al., 2023), the lme4 (Bates et al., 2015) and lmerTest (Kuznetsova et al., 2017) packages for fitting mixed-effects models to the data. Graphs were designed with the ggplot2 package (Wickham, 2016). The data of all three experiments, and the respective R analysis scripts are accessible on OSF.

We fit mixed-effect models to the data to account for the within-subjects design and the complexity of the stimulus material. In all three experiments, we tested different fixed effects structures against each other by refitting the models as maximum likelihood models, and comparing their AIC (Akaike, 1998) with χ^2 tests. In all three experiments, models with fixed, additive effects of prediction, block, scene category, scene color, and target location, and random intercepts for participants, yielded the best fit. Models including interactions of the fixed effects did

not yield a better fit, and models including random slopes did not converge or produced a singular fit.

Results

Training. In all three experiments, we fitted mixed-effects models to the data, and determined that the best fit was a model with predictability, block, scene category, scene color, and target location as factors, and a random intercept for participants. Any model including interactions, or models with random slopes for participants, did not yield a better fit than the model with additive factors. The model parameters for all three experiments are displayed in the Appendix B in Table A1. The descriptive data for RTs by block and predictability are shown in Figure 2.

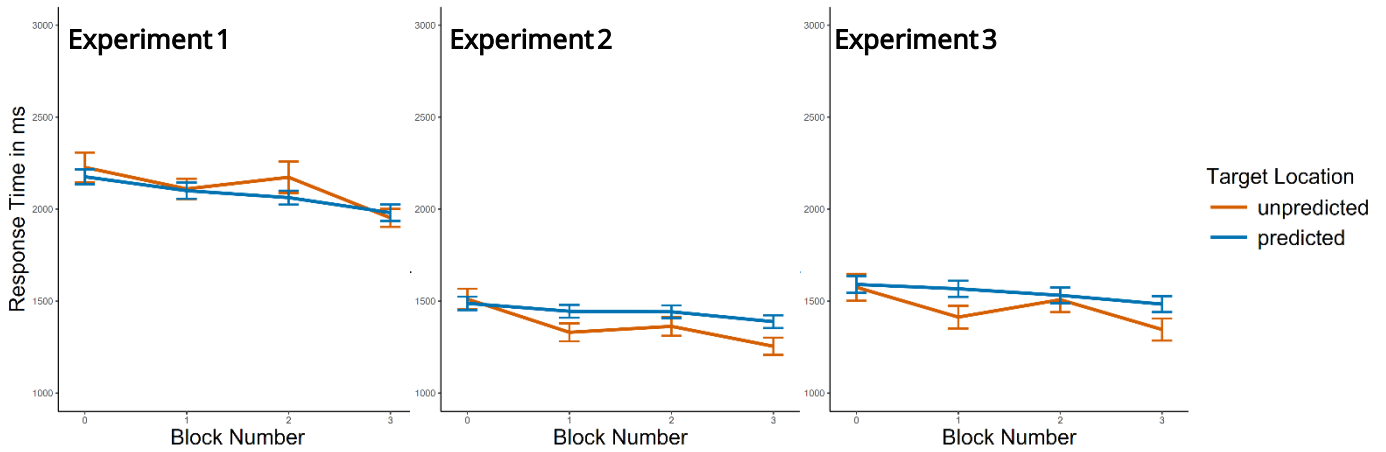
In Experiment 1, in which the scene color predicted target location, mean accuracy in the learning phase was 94.91% ($SD=.22$), and mean RT was 3178.78ms ($SD=4552.14$). The mixed-effects model revealed no significant effect of predictability of the target location on RT, but a significant decrease of RT over the course of the training blocks. It also showed that there are significant effects of the stimulus material characteristics, such as differences between RTs in the different scene categories and colors, and further, a significant effect of target position.

In Experiment 2, in which participants were not explicitly instructed about the contingencies between scene category and target location, mean accuracy in the learning phase was 97.33% ($SD=.16$), and mean RT was 2383.09ms, ($SD=4079.08$). In the mixed-effects model, there was a small but significant effect of predictability. In trials with predictable target location, RT was significantly longer than in trials with unpredictable target locations. Further, RT decreased with the number of blocks, and there were significant differences in RT for the different scene categories and scene colors, as well as for different target positions.

Figure 2

Mean Response Times in the Learning Phase by Block Number and Target Location

Predictability for Experiments 1, 2, and 3



Note. Error bars indicate standard errors.

In Experiment 3, participants were explicitly instructed about the contingencies between scene category and target location. Mean accuracy was 96.60% ($SD=.18$), and mean RT was 2238.90ms, ($SD=3316.21$). The results from the mixed-effects model were similar to those of Experiment 2: Predictability was also a significant factor, in that RTs were slower when the target locations were predictable, and faster when they were unpredictable. The effect of block was significant with the same magnitude as in Experiment 2, and characteristics of the scene, its category and color, as well as target location, also had significant effects on RT. In an exploratory analysis, we excluded participants who performed below chance performance in the generation task. The re-analysis with the 33 remaining participants revealed no significant differences for the effects of predictability or block, in contrast to the analysis with all participants of Experiment 3 (for a detailed analysis, see analysis script on OSF).

As can be seen in the descriptive data from all three experiments, it is noteworthy that RTs are slower in Experiment 1, than in Experiment 2 and 3. However, a comparison between the experiments would not be meaningful because the stimulus material is different in Experiment 1, including only scenes of four color categories, whereas Experiments 2 and 3 used six color categories. As participants' task and the experimental set-up are exactly the same in all experiments, the increased RT in Experiment 1 can only be attributed to the stimulus material. This should not be confounded with the factor of predictability. Rather, longer RTs would potentially increase the likelihood of a predictability effect emerging, as predictable target positions should have a more pronounced benefit. In contrast, generally faster RTs might increase the risk of masking a predictability effect due to a potential floor effect that limits the observable behavioral benefit of prediction.

Generation Task Analysis. In Experiment 1, performance in the generation task was not significantly above chance-level (24.09%), $t(61) = -.796$, $p = .786$. Mean confidence rating was 2.09 (SD=.86). The combined measure of generation task accuracy and confidence showed (see, e.g., Dienes & Seth, 2010, for a similar approach) no evidence of explicit knowledge, as the relative frequency of high confidence given correct responses (15.12%) was not significantly higher than low confidence given correct responses (25.00%), $t(61) = -3.315$, $p = .999$. Interestingly, there was still a significant correlation between generation task accuracy and confidence ($r = .27$), $t(60) = 2.1327$, $p = 0.037$.

In Experiment 2, performance in the generation task was significantly above chance-level (30.63%), $t(59) = 3.52$, $p < .001$. Mean confidence rating was 1.84 (SD=.93). The combined measure of generation task accuracy and confidence showed no evidence of explicit knowledge, as the relative frequency of high confidence given correct responses (21.19%) was not significantly higher than low confidence given correct responses (29.69%), $t(59) = -2.12$, p

= .981. There was no significant correlation between generation task accuracy and confidence ($r=.19$), $t(58) = 1.45$, $p = .152$.

In Experiment 3, performance in the generation task was significantly above chance-level (57.37%), $t(38) = 6.91$, $p < .001$. Mean confidence rating was 2.79 (SD=1.11). The combined measure of generation task accuracy and confidence showed evidence of explicit knowledge. The relative frequency of high confidence given correct responses (54.59%) was significantly higher than low confidence given correct responses (27.50%), $t(38) = 3.47$, $p < .001$. There also was a significant correlation between generation task accuracy and confidence ($r=.53$), $t(37) = 3.77$, $p < .001$.⁴

As an exploratory analysis, we further examined whether performance in the learning phase correlated with performance in the generation and confidence task. Therefore, we computed mean accuracy and confidence per participant, as well as their individual CC effect, defined as $RT_{\text{unpredictable}} - RT_{\text{predictable}}$. In Experiment 1, there was neither a significant correlation between accuracy and CC effect ($r=.12$), $t(60) = -0.92$, $p = .364$, nor between confidence and CC effect ($r=.04$), $t(60) = .29$, $p = .772$. In Experiment 2, there was neither a significant correlation between accuracy and CC effect ($r=-.01$), $t(58) = -0.09$, $p = .927$, nor between confidence and CC effect ($r=-.16$), $t(58) = -1.25$, $p = .216$. In Experiment 3, there was a significant correlation between accuracy and CC effect ($r=-.33$), $t(37) = -2.16$, $p = .037$, but not between confidence and CC effect ($r=-.31$), $t(37) = -2.01$, $p = .052$.

⁴ When excluding participants who were at or below chance-level performance from the data set, we find similar results. Performance in the generation task was significantly above chance-level (64.77%), $t(32) = 9.00$, $p < .001$. Mean confidence rating was 2.85 (SD=1.11). The combined measure of generation task accuracy and confidence showed evidence of explicit knowledge, as the relative frequency of high confidence given correct responses (62.83%) was significantly higher than low confidence given correct responses (30.41%), $t(32) = 3.66$, $p < .001$. There was a significant correlation between generation task accuracy and confidence ($r=.60$), $t(31) = 4.13$, $p < .001$.

General Discussion

In this study, we tested implicit and explicit learning of low-level and semantic cues in an adapted CC paradigm (Tavera & Haider, 2025) with complex real-world scenes. With this, we aimed to test the hypothesis that semantic processing is exclusively. Three important results were obtained. First, the results of Experiment 1 revealed that color was not learned to be a predictable cue. Second, Experiments 2 and 3 showed that the complex real-world scenes as predictive cues of target locations were learned without or with explicit instructions about the contingencies. However, in both cases the learning effect was reversed. Performance was worse in trials with predictable target location compared to trials with unpredictable target location. Third, the assessment of explicit knowledge revealed that only those participants who were informed about the contingency between scene category and target location had access to explicit knowledge. This validates our approach to test for implicit and explicit knowledge. Below, we will discuss these findings in further detail.

The first finding, that color was not used as a predictive cue, was relatively unexpected, because recently, Tavera and Haider (2025) have shown that color cues can be learned implicitly in the same adapted CC paradigm. In our past study, the stimulus material was simplistic letter displays, and the colors were well-defined and one-to-one contingent with target location. In contrast, Experiment 1 of the present study deviates from that in two ways (see Figure 1). First, the stimulus material here was more complex, as the real-world scenes were characterized not only by colors and shapes, but additionally, by a myriad of features, such as texture, intensity, and edge orientation. Second, the color cue was not one clearly defined hue, but rather a mixture of hues from one color family, while also other colors were present in the scene. Thus, learning to use the color cue to predict target location required a certain amount of generalization or categorization into color schemes.

So, there are two potential explanations. On the one hand, the complexity of the stimulus material could have eliminated learning, potentially because color was just one of many low-level features present in the scene. Multiple research groups have in fact investigated the role of color in scene processing but have come to diverse conclusions. For instance, Gegenfurtner, Wichmann, and Sharpe (1996) find that colored scenes are recognized better than black-and-white scenes in a memory task (see also Gegenfurtner & Rieger, 2000), whereas Nijboer, Kanai, Haan, and van der Smagt (2008) come to the opposite conclusion. They suggest that an advantage for colored scenes is restricted to scenes that imply a nameable gist. Analogously, assessing verbal labelling of scenes, Oliva and Schyns (2000) find that color information enhances categorization performance significantly. However, Delorme, Richard, and Fabre-Thorpe (2000) claim that monkey's and human's scene categorization performance does not strongly depend on color. Oliva and Schyns (2000) quite tellingly end their summary of the literature with the conclusion that "existing data with real pictures (...) suggest that the color is never, always, and sometimes used to recognize a scene!" (Oliva & Schyns, 2000, p. 179). It thus remains unclear whether color processing is a main or a negligible part of visual scene processing. For our stimuli in particular, color is not a diagnostic feature of the scene semantics. The different rooms are human-made environments and can thus come in any color. It would be different for stimulus material with, for example, natural scenes, in which green would be diagnostic for plant-related, and blue for water-related scenes. We can thus only speculate that with our stimulus material, color was not processed in such a way to be associated with target location.

On the other hand, the finding that color is not learned as a cue could also be interpreted as a generalization failure. As not one specific color was predictive for target location, identifying a cue required a generalization across different hues, and categorization into color schemes. This process may not be possible without explicit instruction.

The second main result of the current study is that participants learned the relation between scene categories as predictive cues. However, it is a bit more complex to discuss because we found a reversed CC effect, both with implicit and with explicit instructions about the contingencies between scene categories and target locations.

According to the classical CC effect, RTs should decrease for trials with predictable target locations when compared to trials with unpredictable target locations. Consistently in Experiments 2 and 3, we instead see increased RTs for predictable trials. There are some approaches to explaining this unexpected finding.

First, when conducting experiments with real-world scenes, an increased variance due to the stimulus material is to be expected. For example, the scenes are not perfectly balanced in terms of saliency, and objects within the scenes that might have produced pop-out effects, therefore slowing RTs to the actual task. This is why we included variance that can be explained by the category or color of the scene in the mixed-effects model. Therefore, we can, for instance, observe that RTs tended to be faster in scenes in white, and slower in kitchens (see Table A1 in Appendix B). Additionally, we constructed the experiment such that there were four target locations. In visual scene search, it is known that there is a general center bias (Hayes & Henderson, 2019a), which is insignificant in our experiments, because the four target locations are equally distant from the center of the scene. The center was additionally emphasized by a fixation cross to ensure equal first fixation across participants and trials. Additionally, research has shown a general leftward bias (Nuthmann & Matthias, 2014) in visual search in real-world scenes. This is also not relevant for our experiments, because all four target locations serve as both predictable and unpredictable target locations within the experiment. We do find significant differences in RTs between the target locations, most strikingly, a consistent RT advantage for the upper left target location (see Table A1 in Appendix B). However, we also included the

factor of target location into the mixed-effects model to account for these differences. The effect of predictability remains significant beyond these stimulus material effects.

Secondly, it is conceivable that the experimental design of 70% predictability between scene category and target location is not appropriate for our paradigm. However, if 70% was not a high enough cue validity, we should expect a null-effect for predictability. What we find instead is a predictability effect, only reverse. Additionally, this ratio of predictable to unpredictable target location trials has been well-proven in previous CC studies in our lab (Tavera, Wilts, & Haider, unpublished), and a recent systematic test of different cue validities in CC has shown that 75% predictability produced a reliable CC effect (Su et al., 2024).

Because the factors of stimulus material and experimental design cannot account for our reverse RT performance effect, we will discuss candidates for a theoretical explanation of this reverse predictability effect. An explanation of the effect would require a mechanism that produces slower RTs for an expected, frequent event, in contrast to an unexpected, infrequent event.

One such explanation could potentially be a mechanism similar to the one producing negative priming effects (Frings et al., 2015; Tipper, 1985, 2001). The exact mechanisms that may underlie such effects are still debated (for an overview, see Frings et al., 2015; S. Mayr & Buchner, 2007). Generally, the effect arises when a stimulus in trial t_{n-1} is irrelevant, and then hampers a second stimulus at trial t_n . A similar mechanism has been proposed specifically for visual search, suggesting that after trial t_{n-1} , features of distractors are inhibited, which influences visual search in trial t_n (Lamy, Antebi, et al., 2008). One can potentially transfer these effects and mechanisms to the present experiments. This would mean that in a trial t_{n-1} , when a scene was presented with an unpredicted target location, the originally predicted target location needed to be inhibited. As a consequence, it is conceivable that in trial t_{n-1} , when then the same scene category is presented with the predicted target location, attention allocation to this target location is still inhibited from trial t_{n-1} . However, in an exploratory analysis, we did not

find a more significant slowing of such trials, in comparison with trials with predicted target locations that were not preceded by a trial with the same scene category. Thus, we do not find evidence for a mechanism similar to a negative priming or distractor inhibition effect.

Regardless of the mechanism that can explain our finding of a reverse CC effect, there is another essential result of our Experiments 2 and 3. The comparison between the two offers insight into the potential difference between conscious and unconscious processing. In Experiment 2, we implemented the adapted CC paradigm in which participants could learn contingencies between semantic scene category and target location. They were not explicitly instructed to do so. In Experiment 3, while holding all other factors constant, we explicitly informed participants about the contingencies. Interestingly, we find no difference in behavioral patterns between the two experiments. In both experiments, there is a significant learning effect, but in a direction opposite from the hypothesized. This finding is essential because it demonstrates that the reverse CC effect is not a shortcoming of implicit processing, such that predictions for target locations are somehow inhibited. We find the same reverse CC effect also in participants who were explicitly instructed about the contingencies. This finding alone would have suggested that there is a top-down mechanism that inhibits target location predictions. A potential explanation could then have been that the 70% predictability resulted in a top-down strategy of inhibiting the predictions because of the 30% error probability. Because we find the reverse CC effect in both implicit and explicit learning conditions, we would rather suggest that it is a fundamental effect of the learning process in this paradigm. The mechanism that remains unclear from our experiments seems not to be influenced by explicit knowledge and potential subsequent top-down processes. This is an important finding, and further requires empirical investigation.

The argument that the CC effect is the same for explicit and implicit learning conditions is based on the premise that we were able to validly assess explicit knowledge in our

experiments. This is our third main finding: A validation of our explicit knowledge test which is a combination of an objective, direct task (generation task; Chun & Jiang, 2003) and a meta-cognitive measure (confidence measure; Michel, 2023a). We could show that this combination of two measures has the capacity to indicate explicit knowledge in principle. One could, in this context, view our Experiment 3 as a manipulation check. By explicitly instructing participants, we expected to find explicit knowledge in our measure. As expected, we found a relationship between accuracy and confidence in the explicit learning condition. This indicated meta-cognitive knowledge of participants, meaning that they had a conception of their own contingency knowledge or the lack thereof. In contrast, we found no such evidence for meta-cognitive, explicit contingency knowledge. Still, we found above chance level performance in the generation task in Experiment 2. This, however, is not indicative of explicit knowledge, as performance in the generation task might as well be driven by implicit processes (Jiménez et al., 1996; Michel, 2023a; Reingold & Merikle, 1988) which could then potentially result in intuitions that produce above chance accuracy (Weinberger & Green, 2022). Furthermore, our exploratory analyses of the correlations between learning and test phase performance yielded interesting results. This analysis was done with the notion that participants with explicit knowledge should show a more pronounced CC effect. At the same time, those participants should be able to exhibit this knowledge with high accuracy and high confidence, and, most importantly, a positive correlation between the two measures (Dienes & Seth, 2010). In Experiments 1 and 2, there were no significant correlations between accuracy or confidence in the generation task, and learning phase performance. In contrast, in Experiment 3, we found significant correlations between accuracy and learning phase performance, while simultaneously finding a positive correlation between accuracy and confidence. We would have expected a significant correlation between confidence and learning phase performance in Experiment 3, but this correlation did not reach significance. Yet, this might be due to relatively low confidence overall ($M=2.79$), which is surprising, given that participants were explicitly instructed about the contingencies.

Nevertheless, also with this exploratory analysis, we find indicators of explicit knowledge only in Experiment 3, in which participants were explicitly instructed about the contingencies.

Concluding, with these three experiments, we have shown further potential for the adapted CC paradigm (Tavera & Haider, 2025) by extending it to investigate CC effects with more complex stimulus material, in this case, real-world scenes. By doing so, we were first able to show that color cues, as a dominant low-level feature within scenes, could not be implicitly learned to predict target locations. Secondly, we found that semantic scene category cues, could be learned both in implicit and explicit learning conditions. However, we found a reverse CC effect in both such conditions. The mechanism behind this finding remains unclear, but should be further investigated, potentially in the realm of attentional inhibition or episodic retrieval mechanisms. Thirdly, we have shown that our test for explicit knowledge is in fact able to detect explicit knowledge in conditions of explicitly instructed contingency learning. It did not indicate explicit knowledge in our Experiments 1 and 2 where participants were not explicitly instructed about any contingencies.

Appendix A

The following analyses were conducted to ensure that the scene material that we created fit the requirements of our experiments. Therefore, we analyzed the scene images (without the target letters) using the SHINE toolbox (Willenbockel et al., 2010) for MATLAB (The MathWorks Inc, 2024). We analyzed RGB channel intensities per semantic scene category, and per scene color category. Our hypothesis for the RGB analysis was that the channel intensities would not be significantly different across the semantic scene categories, but significantly different across the scene color categories. That would mean that the color categories are distinguishable on the basis of low-level features, but the semantic categories are not. The descriptive data are displayed in Figures A1 and A2.

Figure A1

RGB Intensity Values per Semantic Scene Category

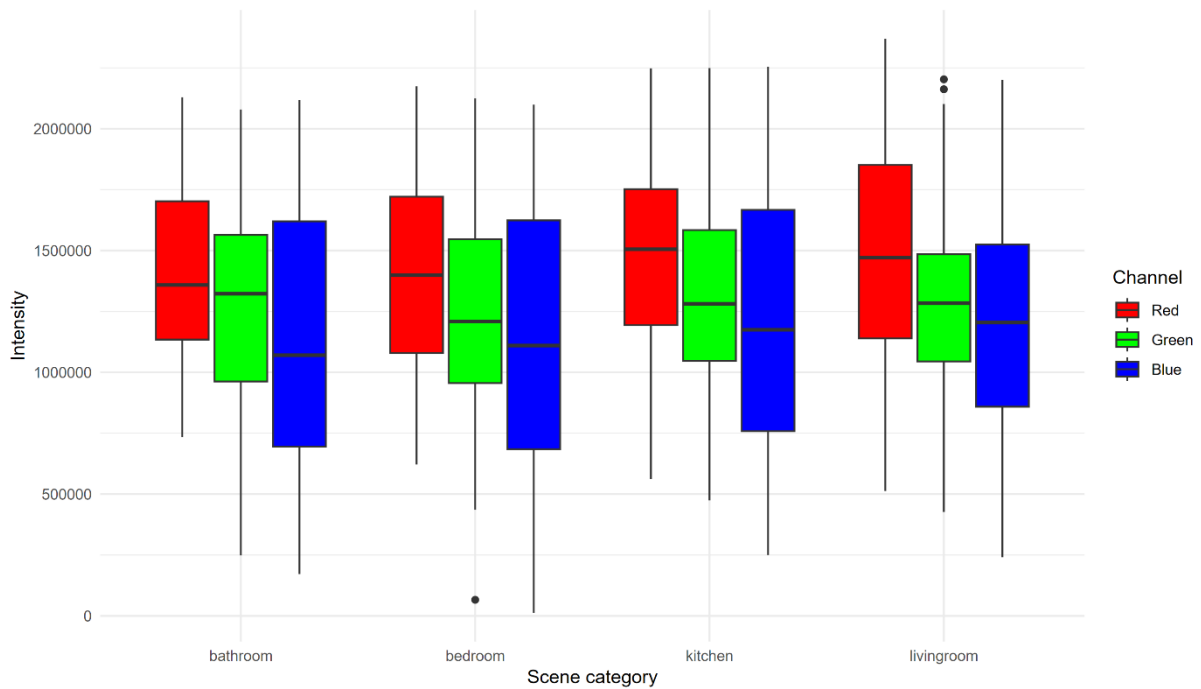
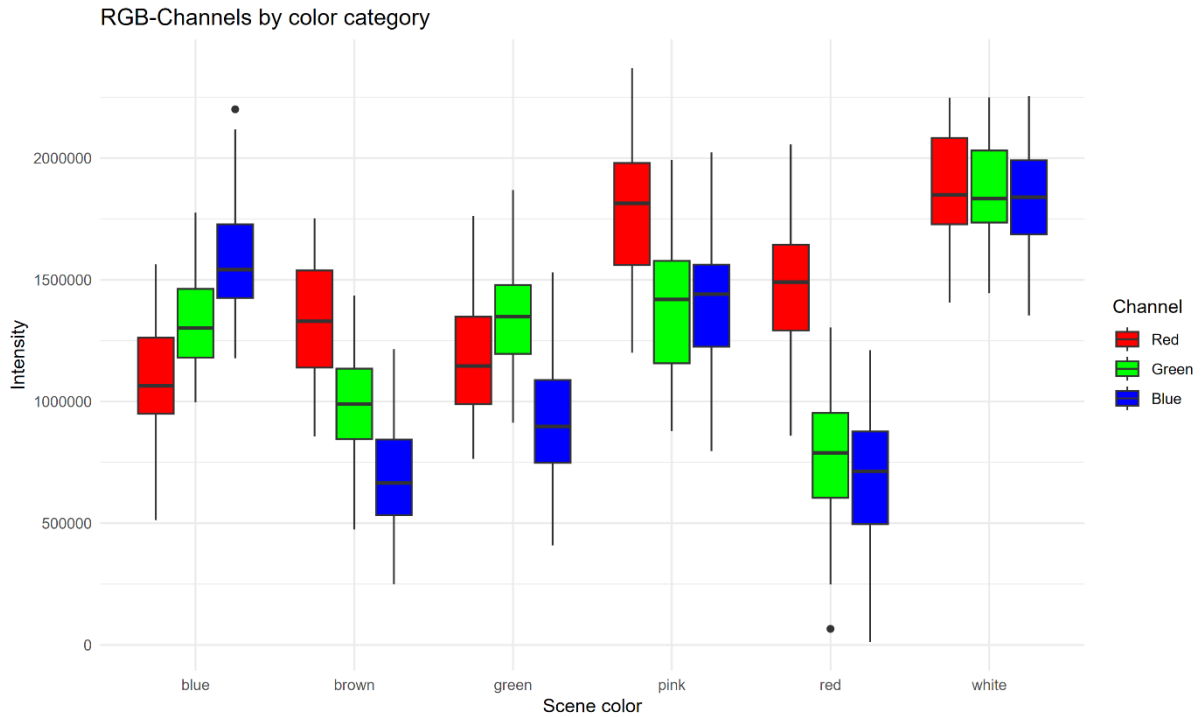


Figure A2*RGB Intensity Values per Scene Color Category*

We conducted ANOVAs to test whether RGB channel intensity values differed across semantic scene categories and across scene color categories. The analyses revealed that there were no significant differences across semantic scene categories in any RGB channel (red: $F(3, 236)=0.722, p=.54$; green: $F(3, 236)=.345, p=.793$; blue: $F(3, 236)=.25, p=.861$). In contrast, all differences across scene color categories in all RGB channels were significant (red: $F(5, 234)=59.05, p=.001$; green: $F(5, 234)=101.9, p=.001$; blue: $F(5, 234)=151., p=.001$).

Appendix B

Table A1

Fixed and Random Effects of the Linear Mixed-Effects Model Predicting Reaction Times for Experiments 1, 2, and 3

Experiment	Predictors	Estimate	95% CI		p
			LL	UL	
1	Intercept	2168.81	2081.93	2255.70	<.001
	Predictability [predictable]	-10.66	-45.41	24.09	.548
	Block	-62.43	-76.79	-48.07	<.001
	Scene category [bedroom]	-90.52	-134.67	-46.36	<.001
	Scene category [kitchen]	122.06	76.79	167.33	<.001
	Scene category [living room]	-66.53	-110.97	-22.09	.003
	Scene color [brown]	30.72	-14.06	75.51	.179
	Scene color [green]	76.28	30.91	121.66	.001
	Scene color [white]	-316.23	360.11	-272.35	<.001
	Target Location [2]	193.32	148.73	237.91	<.001
	Target Location [3]	29.66	-14.41	73.74	.187
	Target Location [4]	153.84	109.44	198.25	<.001
	Random Effects				
	σ^2	561002.24			
	$\tau_{00\text{subject}}$	74958.38			
	ICC	0.12			
	N_{subject}	62			
	Observations	8683			
	Marginal R^2 / Conditional R^2	0.064 / 0.175			
2	Intercept	1462.33	1379.58	1545.07	<.001
	Predictability [predictable]	59.56	27.57	91.55	<.001
	Block	-50.31	-62.52	-38.10	<.001
	Scene category [bedroom]	-58.53	-96.40	-20.67	.002
	Scene category [kitchen]	160.81	122.23	199.39	<.001
	Scene category [living room]	2.15	-35.99	40.29	.912
	Scene color [brown]	-77.83	-125.95	-29.71	.002
	Scene color [green]	53.61	4.12	103.11	.034
	Scene color [pink]	-149.08	-195.99	-102.17	<.001
	Scene color [red]	-131.23	-178.63	-83.83	<.001
	Scene color [white]	-324.03	-371.01	-277.05	<.001
	Target Location [2]	196.66	158.32	235.00	<.001
	Target Location [3]	43.72	5.48	81.96	.025
	Target Location [4]	233.05	194.56	271.54	<.001
	Random Effects				
	σ^2	617926.32			
	$\tau_{00\text{subject}}$	67021.74			
	ICC	0.10			
	N_{subject}	60			
	Observations	12957			
	Marginal R^2 / Conditional R^2	0.048/0.141			

Experiment	Predictors	Estimate	95% CI		p
			LL	UL	
3	Intercept	1509.00	1408.97	1609.03	<.001
	Predictability [predictable]	57.28	16.86	97.71	.005
	Block	-49.00	-64.40	-33.61	<.001
	Scene category [bedroom]	-51.60	-99.50	-3.70	.035
	Scene category [kitchen]	181.55	132.97	230.12	<.001
	Scene category [living room]	22.94	-25.17	28.81	.350
	Scene color [brown]	-31.83	-92.46	28.81	.304
	Scene color [green]	68.01	5.99	130.03	.032
	Scene color [pink]	-174.53	-233.85	-115.21	<.001
	Scene color [red]	-110.97	-170.76	-51.18	<.001
	Scene color [white]	-390.47	-449.66	-331.27	<.001
	Target Location [2]	242.16	193.87	290.45	<.001
	Target Location [3]	99.17	50.98	147.36	<.001
	Target Location [4]	269.40	221.00	317.81	<.001
Random Effects					
	σ^2	654977.32			
	$\tau_{00\text{subject}}$	61523.60			
	ICC	0.09			
	N_{subject}	40			
	Observations	8661			
	Marginal R^2 / Conditional R^2	0.059/0.140			

Note. Results from a linear mixed-effects model: response times ~ predictability + block + scene category + scene color + target location + (1|subject). Response time is predicted by target location predictability, block number, scene category (bathroom [reference], bedroom, kitchen, living room), scene color (blue [reference], brown, green, white), and target location (position 1 [reference], 2, 3, 4). The model includes random intercepts for subjects. Fixed effects are presented with unstandardized estimates, 95% confidence intervals (LL = lower limit, UL = upper limit), and *p*-values. Random effects include the residual variance (σ^2), the variance of the random intercept for subjects (τ_{00}), and the intraclass correlation coefficient (ICC). R^2 values represent marginal (fixed effects only) and conditional (fixed + random effects) model fit. *p* values < .05 are considered statistically significant and are shown in bold.