

**Working Memory related to different subtasks can be
maintained in separate, non-interfering
stores.**

A Thesis Submitted to The Graduate School of Engineering and
Science of Bilkent University in partial fulfillment of the
requirements for
the degree of
Master of Science
in
Neuroscience

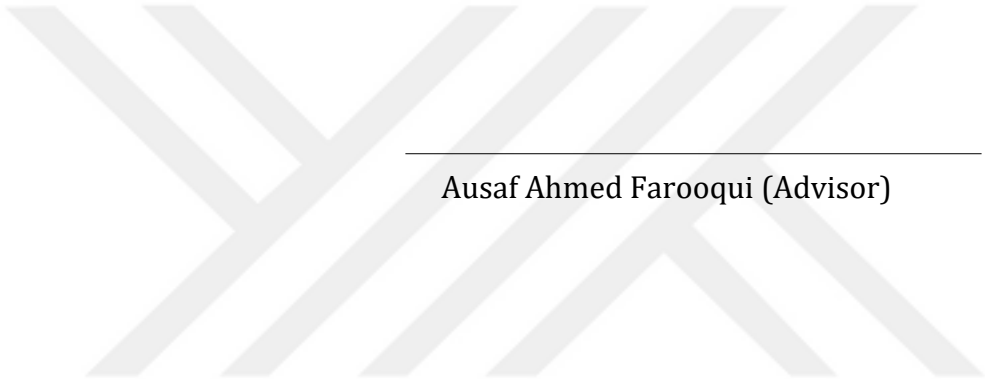
By
Aida Rezaei
July 2024

Working Memory related to different subtasks can be maintained in separate, non-interfering stores.

By Aida Rezaei

July 2024

We certify that we have read this thesis and that in our opinion it is fully adequate, in scope and in quality, as a thesis for the degree of Master of Science.



Ausaf Ahmed Farooqui (Advisor)

Hüseyin Boyacı

Buse Merve Ürgen

Approved for the Graduate School of Engineering and Science:

Orhan Arıkan
Director of the Graduate School

ABSTRACT

Working Memory related to different subtasks can be maintained in separate, non-interfering stores.

Aida Rezaei

M.S in Neuroscience

Supervisor: Ausaf Ahmed Farooqui

July 2024

How can WM maintain more information when its capacity is only around 4 items? Here, we explore the possibility that information related to separate subtasks do not count towards this limit, perhaps, because they are maintained in non-interfering stores. Across the two experiments, we investigated if increasing the WM load related to one subtask interfered with the execution of a concurrent but distinct second subtask.

In Experiment 1, participants first saw pictures that were to be kept in mind and used for a later subtask B. They then executed subtask A, while keeping in mind these subtask B pictures. Although subtasks A and B involved maintaining and updating identical sets of pictures, increasing the number of subtask B pictures did not interfere with subtask A execution, which was only affected when the number of pictures relevant to it was increased, suggesting

that subtask A and B pictures were maintained in separate non-interfering stores.

An objection might be that in Experiment 1, subtask B pictures were passively and not goal-directedly maintained. In Experiment 2, participants executed a more complex subtask that forced participants to maintain and update two separate sets of subtask B pictures while executing subtask A and subtask A also involved WM maintenance and updating of identical pictures. We found that even here increased load of subtask B pictures did not affect subtask A performance.

Thus, at least in some multitasking situations, information related to distinct subtasks can be maintained in separate non-interfering stores.

Keywords: *Working Memory, Working Memory Capacity, Multitasking, Subtask Independence, Task Interference*

ÖZET

Farklı alt görevlerle ilişkili çalışma belleği ayrı,
müdahale etmeyen depolarda korunabilir.

Aida Rezaei

Yüksek Lisans, Nörobilim

Tez Danışmanı: Ausaf Ahmed Farooqui

Temmuz 2024

Görsel Çalışma Belleği (ÇB) kapasitesi yalnızca 4 öge ile sınırlı iken bundan daha fazla görsel bilgiyi nasıl depolayabilir? Bu çalışmada, ÇB'nin farklı alt görevlere ilişkin bilgileri 4 öge sınırına takılmadan nasıl saklayabildiğini inceliyoruz. Kuvvetle muhtemel bu alt görevler birbirleriyle çakışmayan ayrı ÇB depolarında saklanmaktadır. Yaptığımız iki deneyde de, bir alt görevle ilgili ÇB yükünün artırılmasının, eş zamanlı olarak diğer alt görevin yerine getirilmesini etkileyip etkilemediğini araştırdık.

Deney 1'de, katılımcılar önce alt görev A'nın gerçekleştirilmesi için gösterilen resimleri akıllarında tutmaya çalıştılar. Alt görev A için akılda tutulan resimler, alt görev B'de de kullanıldı. Katılımcılar önce alt görev A'yı tamamladılar, bu görevi tamamlarken alt görev B ile ilgili resimleri de akıllarında tuttular. Alt görev A ve B'de akılda tutma ve güncelleme görevleri için kullanılan resimler aynı olmasına rağmen, alt görev B'de akılda tutulacak resimleri artırmak alt görev A performansını olumsuz etkilemedi. Alt görev A performansını sadece bu görevle ilgili resimleri artırmak etkiledi. Bu sonuç, alt görev A ve B ile ilgili

informasyonların birbirlerini etkilemeyen farklı ÇB depolarında tutulduğunu gösteriyor.

Alt görev B' deki resimlerin bir hedef doğrultusunda değil de pasif bir şekilde hafıza tutulduğu yönünde bir itiraz gelebilir. Deney 2'de katılımcılar yine alt görev B'de ortak kullanılan resimleri akılda tutmayı ve güncellemeyi içeren alt görev A' yı gerçekleştirirken, bu sefer alt görev B'nin iki farklı görev setinden oluştuğu daha karmaşık alt görev gerçekleştirdiler. Alt görev B'nin bilişsel yükünü artırmak bile alt görev A' daki performansı etkilmediğini bulduk.

Bu sonuçlar doğrultusunda, en azında çoklu görev gerektiren durumlarda, farklı alt görevlere ilişkin bilgilerin birbirine etkilemeyen ÇB depolarında tutabileceğini söylüyoruz.

Anahtar sözcükler: *Çalışma Belleği, Çalışma Belleği Kapasitesi, Çoklu Görev, Alt Görev Bağımsızlığı, Görev Karışması,*

ACKNOWLEDGEMENTS

I wish to extend my deepest appreciation to my advisor, Assistant Professor Ausaf Ahmed Farooqui, for his outstanding mentorship, assistance, and motivation during this academic inquiry. I am also thankful to the Bilkent University Neuroscience Graduate Program for providing an enriching educational environment and all the necessary resources for my studies. Special thanks to Asst. Prof. Huseyin Boyaci and Asst. Prof. Buse Merve Ürgen for joining my jury and offering warm and valuable feedback that significantly improved my work. I am grateful to Bilkent University, the National Magnetic Resonance Imaging Center (UMRAM), and the Aysel Sabuncu Brain Research Center, along with the entire Neuroscience Graduate Program family, for their unwavering support. My heartfelt appreciation goes to the Higher Cognitive Functions Lab members for their friendship, collaboration, and insightful discussions that greatly enriched my research experience. I am incredibly thankful to Gulsum Ozge Sengil, Tamer Gezici, and Ipek Ciftci for their constant support and encouragement throughout this journey. A special mention goes to SeyedFarid Sadati for his unwavering help and support from my first day at Bilkent University, especially for the times when he was the only supporter and believer. Ultimately, I am deeply thankful to my family for their support and unconditional love, which have served as the bedrock of my achievement.



To my mother Behnaz and my sister Sayna.

Contents

CHAPTER 1	1
INTRODUCTION.....	1
1.1 Understanding Working Memory	1
1.2 Exploring Working Memory Theories	3
1.2.1 Models of Memory Systems: The Multiple Component Model	8
1.2.2 Insights from State-based Models: Examining the Embedded Process Model.....	10
1.2.2.1 Concentric Model of Working Memory	11
1.3 Working Memory Capacity and Limitations.....	14
1.3.1.1 Degradation of Information	15
1.3.1.2 Investigating Interference	16
1.3.1.3 Constraints on Resources	16
1.3.1.3.1 Specific Allocation of Resources	17
1.3.1.3.2 Continuous Resource Allocation	19
1.4 Control of Working Memory Information: Attention-based Control of Working Memory Information.....	20
1.4.1.1 Processes of Consolidation.....	21
1.4.1.2 Refreshing and Removal	21
1.4.2 Hierarchical Control of Working Memory Information.....	23
1.5 Current Research	25
CHAPTER 2.....	28
Experiment 1	28
2.1 Methodology and Procedural Details	32

2.1.1	Participants.....	33
2.1.2	Setup Of Experiment	33
2.1.3	Data Analysis	38
2.2	Findings.....	39
2.2.1	Accuracy	39
2.2.1.1	Accuracy for three pictures at the beginning.....	39
2.2.1.2	Accuracy for 1 picture in the beginning.....	41
2.2.2	Presentation time	42
2.2.3	Reaction time.....	43
2.3	Discussion.....	47
CHAPTER 3		50
Experiment 2		50
3.1	Methodology and Procedural Details	52
3.1.1	Participants.....	53
3.1.2	Experimental Setup.....	53
3.1.3	Data Analysis	55
3.2	Findings.....	55
3.2.1	Accuracy	55
3.2.1.1	Accuracy for N-back	55
3.2.1.2	Accuracy for recall.....	57
3.2.2	Presentation time	58
3.2.3	Reaction time.....	60
3.2.3.1	Reaction time for N-back	60
3.2.3.2	Reaction time for recall.....	62
3.3	Discussion.....	63

CHAPTER 4	67
General Discussion	67
4.1 Implications	67
4.2 Elaboration On Discussion	70
4.3 Limitations and Future Directions	72
CHAPTER 5	74
BIBLIOGRAPHY	74



List of Figures and Tables

Figure 1.1: The illustration of WM theory by Baddeley & Hitch's:	6
Table 1.2: The illustration of WM theory by Cowan	7
Figure 1.3: The illustration of WM theory by Oberauer	7
Fig 2.1. Design of Experiment 1:	30
Fig 2.2. 1 back subtask example,	34
Fig 2.3. Two-back and three-back subtask example,	34
Figure 2.4. Subtask 1, session 1:	35
Figure 2.5. Subtask 2, session 2:	36
Figure 2.6. Subtask3, session 3:	37
Figure 2.7. Subtask 4, session 4:	38
Figure 2.8. Mean of accuracy for 3 pictures in beginning of experiment in 2 conditions with 95% CI error bars	40
Table 2.1 : ANOVA table of mean of accuracy for 3 pictures in beginning of experiment in 2 conditions.	40
Figure 2.9. Mean of accuracy for 1 picture in beginning of experiment in 2 conditions with 95% CI error bars	41
Table 2.2. ANOVA table of the mean of accuracy for 1 picture at the beginning of the experiment in 2 conditions.	41
Figure 2.10. average presentation time of all subtasks in experiment 1 with 95% CI error bars:	42
Table 2.3. ANOVA of experiment 1 recall presentation in all subtasks (Greenhouse-Geisser corrected)	43
Figure 2.11. Mean of reaction time for three-back subtask in 2 conditions with 95% CI error bars	44
Table 2.4. ANOVA table of mean of reaction time in three-back subtask for 1 picture in beginning and 3 pictures in beginning (Greenhouse-Geisser corrected)	45
Figure 2.12. Mean of reaction time for 1back subtask in 2 conditions with 95% CI error bars	46
Table 2.5. ANOVA table of mean of reaction time in 1back subtask for 1 picture in beginning and 3 pictures in beginning (Greenhouse-Geisser corrected)	47
Figure 3.1: The Design of Experiment 2.	54
Figure3.2: N-back subtask Accuracy for trials 5 to 9 with 95% CI error bars	56
Table3.1: ANOVA of N-back Accuracy for trials 5 to 9	56

Table 3.2: ANOVA of Recall accuracy for picture α and β between 3 episodes (Greenhouse–Geisser corrected)	58
Figure 3.4: mean of recall pictures Presentation time separately for picture α and β between episodes 1 ,2 and 3 with 95% CI error bars.	59
Table 3.4: ANOVA of presentation time for picture α and β between episode 1,2 and 3(Greenhouse–Geisser corrected)	59
Figure 3.5: Mean of RT for trials 1 to 9 in N-back subtask with 95% CI error bars	61
Table 3.5: ANOVA of RT mean in N-back subtask for trials 2 to 9 between 4 episodes, also trial 1 is eliminated because of updating info of working memory (Greenhouse–Geisser corrected)	61
Figure 3.6: Mean reaction time of recalled pictures in episodes 2,3 and 4 with 95% CI error bars	62
Table 3.6: ANOVA of mean of RT for recall picture α and β between episodes 2,3 and 4.	63

CHAPTER 1

INTRODUCTION

1.1 Understanding Working Memory

Working memory (abbreviated as WM) refers to our capacity to retain relevant information sourced from long-term memory or the external environment and maintained in a readily accessible state and, if need be, manipulated to support ongoing cognitive operations [1, 2]. To exemplify this idea through a real-world situation, Adams, Nguyen, and Cowan [3] offer the following. A teacher mentions to the class that Earth occupies the third position in the solar system and subsequently chooses a student to pinpoint Earth's location on a displayed solar system map. The student must keep in mind the teacher's statement regarding Earth's position in the solar system while processing the subtask. This entails initiating the counting process from the correct planet, rather than starting from the sun, and vigilantly monitoring the count to ensure it stops at three. Concurrently, other thoughts, such as anxiety about performing in front of the class, may also occupy WM. The authors posit that this intricate process can face various challenges resulting in WM errors, primarily because numerous cognitive processes vie for the limited capacity of WM.

Although the use of real-life examples can assist in conveying nature of WM, it's essential to acknowledge that a single definition cannot fully encapsulate the breadth of theories and applications linked to this concept [3]. Different researchers with varying research objectives often employ this term, resulting in multiple definitions. Nevertheless, WM is commonly described as having

certain characteristics, including its capacity to hold a 'limited' amount of information [3]. It plays a uniquely [3] vital role in cognitive processes and facilitates the execution of ecologically relevant subtasks in a hierarchically organized manner [4]. WM information is distinct from other forms of memory, and it possesses the capability to transition between a maintenance mode and a transformation mode. This transition allows it to quickly integrate representations from long-term memory (referred to as LTM) or current perception, creating new representations in conjunction with contextual information [2].

Any conceptualization of WM implies certain functions associated with it.

One framework posits that the primary function of WM is to hold information in a readily accessible and controllable manner for processing rather than mere storage for recollection [2]. As per this framework:

Firstly, our ability to execute different subtasks depending on our goals necessitates a temporary medium for binding existing representations from long-term memory (often referred to as "chunks") and reorganizing them into new representations. WM is the cognitive module where this is supposed to take place. These newly formed representations become the content of WM.

Secondly, the ability to manipulate the representations held in WM is essential. This manipulation should, at a minimum, involve mechanisms for selectively accessing these representations and maintaining the rules, goals, or plans required to manipulate them in a selective manner.

Another crucial aspect of WM is the selective maintenance of information that is most relevant to the current cognitive goal. This requirement introduces a potential conflict between two demands: safeguarding WM representations from interference or influence from perceptual input or long-term memory while also rapidly updating these representations with information from perception or long-term memory. Consequently, WM is proposed to operate in

a manner that allows it to transition between these seemingly conflicting modes to resolve this conflict effectively.

One might inquire about the rationale behind the prevalence of studying WM in the realm of human cognitive psychology [2]. Firstly, WM is believed to play a significant role in a broad spectrum of high-level cognitive functions, including cognitive control and planning. Furthermore, WM measures exhibit a stronger correlation with intellectual assessments, particularly fluid intelligence, in comparison to short-term memory (STM) [5] measures. Additionally, WM deficits are characteristic of numerous neurological and psychiatric conditions, including psychotic disorders such as schizophrenia, schizoaffective disorder, bipolar disorder with psychosis [6], attention deficit hyperactivity disorder (ADHD)[7], Alzheimer's Disease[8], Parkinson's Disease[9], and Major Depressive Disorder[10]. WM is extensively investigated and demonstrated across various sensory modalities, with visual and auditory modalities being the most extensively employed, alongside a linguistic/verbal approach[2]. However, it remains accurate to assert that both conceptual and empirical challenges persist within the WM concept, and many theoretical approaches coexist.

1.2 Exploring Working Memory Theories

The initial utilization of the WM concept originated in computing, stemming from the necessity to store pertinent information within problem-solving routines [3]. In the context of computers, WM was initially delineated as a distinct repository designed to temporarily hold multiple concurrent pieces of information, aiding ongoing subtasks [11]. Nonetheless, as the cognitive revolution took shape, suggesting that human cognition functions as an information-processing system, it became apparent that the human mind should harbor a mechanism resembling the structures and operations observed in computers [2]. During this influential era, in their 1960 publication authored by Miller, Galanter, and Pribram, the concept WM was introduced within the context of human cognitive psychology for the first time

[2-4]. Their conceptualization centred on the idea that WM played a pivotal role in executing 'plans' by acquiring 'data' from both external sources and within the organism itself [4]. In their framework, 'plans' functioned as the human equivalent of computer programs, comprising hierarchically organized processes that govern an organism's sequential operations. According to their model, when any plan segment was 'executed,' the relevant portions of that plan became consciously accessible and synchronized with other plans. To account for this, they proposed the existence of a specialized storage or state where such information could be retained during its execution without specifying the exact mechanism for this memory. This repository for rapidly accessible memory associated with plan execution was termed 'working memory' in their conceptualization. Importantly, they envisaged that several plans or sub-components of a single plan could be retained within WM simultaneously. Moreover, within a single 'plan,' there existed a hierarchical structure consisting of sub-plans and sub-sub-plans, with the lowest-level elements governing momentary actions[3].

Since the term WM found its way into human cognitive psychology, numerous theories, models, and frameworks have arisen to offer theoretical and empirical insights into these phenomena. Comparing these diverse theories of WM (or any other concept, for that matter) directly is challenging due to the variance in their empirical focus and often distinct definitions [3]. Nevertheless, an endeavor can be undertaken to categorize these models or frameworks into two different groups based on how scientists in 1960 [4] conceptualized the functioning of WM [12]. There are two potential perspectives regarding the status of WM information. It can either be construed as being stored in a 'dedicated mechanism,' implying the presence of a specialized WM mechanism. Conversely, it can be viewed as being placed within a 'special state,' representing an inherent aspect of the system that explicitly encodes the relevant information. This state has the capacity to undergo transient changes when the representational facets of the system supporting a cognitive process, be it language or sensory perception, alter

cognitively or neurally, as seen in synaptic changes in weights [2]. These state changes transpire as part of the system's regular operation, typically unrelated to immediate WM demands, such as when the system merely engages in perception[2]. It is feasible to position specific influential existing theories or frameworks within either of these categories. This categorization is fundamentally centred on the pivotal question of how WM operates, specifically whether it hinges on a distinct mechanism or emerges from the interactions of other cognitive and neural processes or the collaboration of various cognitive and neural systems [13]. Some early WM theories align with the former category as they emphasize the existence of discrete memory systems, and thus, they are denoted as 'Memory-systems Models.' In opposition, specific subsequent theoretical frameworks hint at the existence of a unified repository, suggesting that discrete memory representations, namely those about short-term or long-term memory, originate from varying states of activation present within this communal reservoir of information. [14]. For this reason, they are labeled as either 'Activation-based Models' or also named as 'State-based Models' [14].

A

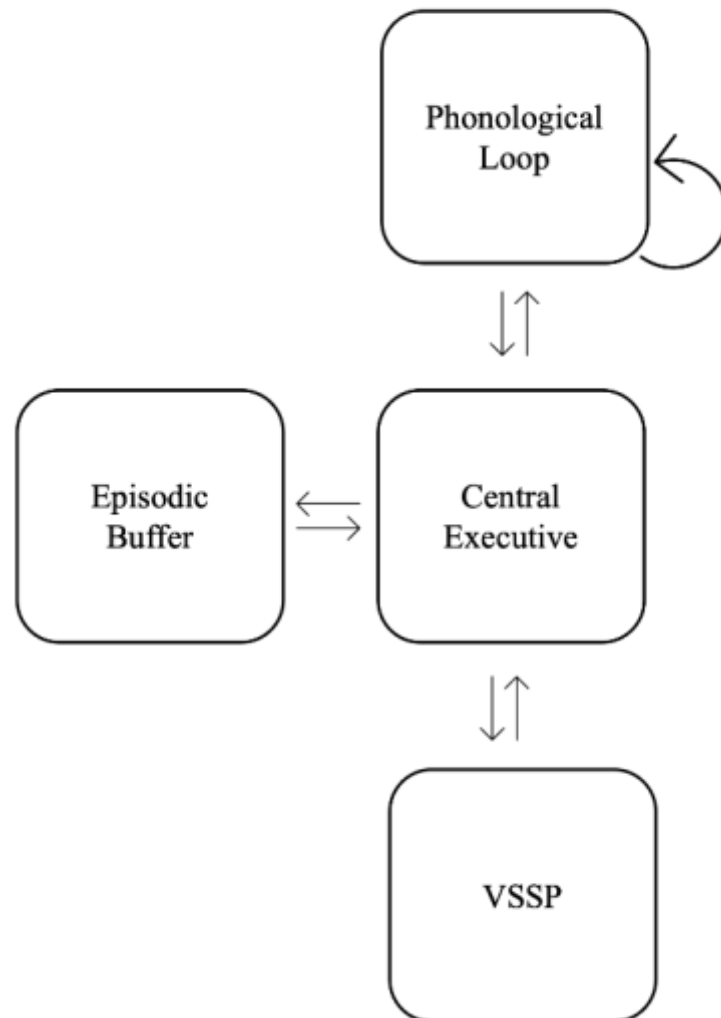


Figure 1.1: The illustration of WM theory by Baddeley & Hitch's: Baddeley and Hitch proposed the Multiple Component Model, which comprises a central executive, episodic buffer, visuospatial sketchpad (VSSP), and the phonological loop. Within the phonological loop, representations can be reactivated through articulatory rehearsal.

B

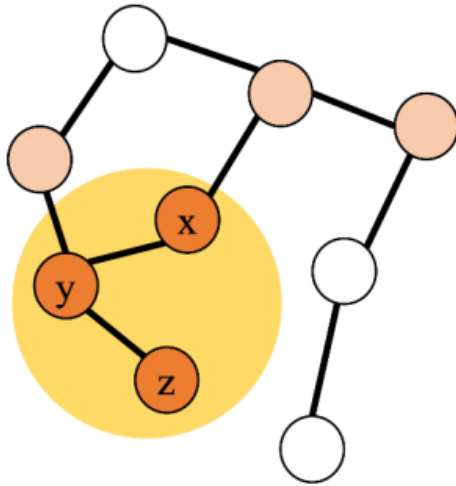


Table 1.2: The illustration of WM theory by Cowan

In Cowan's Embedded Process Model, all nodes reside within Long-Term Memory (LTM), while a specific subset of nodes, namely red and orange nodes, are activated in activated Long-Term Memory (aLTM). The representations of red nodes, falling under the capacity-limited Focus of Attention (FoA) denoted by the orange circle, form a subset of these activated nodes.

C

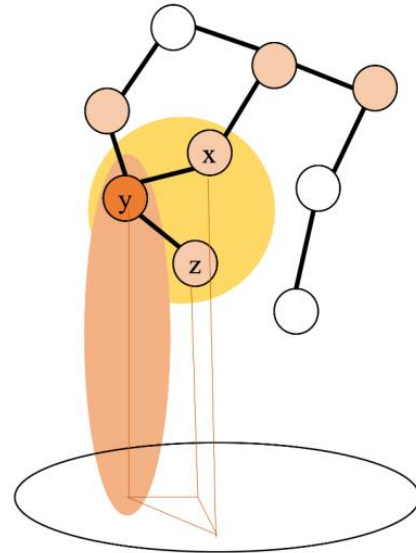


Figure 1.3: The illustration of WM theory by Oberauer

In Oberauer's Concentric Model, all nodes reside within long-term memory (LTM), with a portion of them existing in activated long-term memory (aLTM) depicted by orange nodes. Within this subset, certain nodes are situated in the immediate access area represented by an orange circle, linked by orange lines indicating a shared context or cognitive space within a black oval at the bottom. A specific element within the immediate access region is located in the Focus of Attention (FoA) depicted by an orange oval.

1.2.1 Models of Memory Systems: The Multiple Component Model

The multiple-component model stands as one of the most influential models in the history of WM [2]. Originating from the minds of Baddeley and Hitch [15], this model has undergone occasional revisions. Before its emergence, the prevailing STM model proposed by Atkinson and Shiffrin posited a hypothetical single store responsible for temporarily holding information. This single-store approach was considered the operative mechanism behind many STM paradigms of its time [16]. Furthermore, STM has attributed a central role in governing executive functions and systems relevant to learning and information retrieval [17]. However, Baddeley and Hitch [15] challenged this conceptualization by presenting empirical evidence suggesting that STM might play a different role in WM than previously thought. They argued for abandoning the single-module approach and advocated for a more comprehensive WM system capable of supporting reasoning, comprehension, and learning. In contrast to the prevailing unitary-store short-term memory models, the 'multiple component model' emphasized the contribution of these various systems to cognitive operations rather than focusing solely on memory [18]. This model defined three specialized components, each tailored to retain distinct forms of information: the phonological loop for verbal data, the visuospatial sketchpad for visual-spatial information, and the episodic buffer for episodic long-term memory content [18]. The fourth component, known as the central executive, assumed the role of controlling and supervising the 'slave systems' (the other components) while retrieving information from episodic and semantic long-term memory [18](Figure 1.1). It is proposed that the central executive functions more like a control system governing attentional processes.

In the context of the complex span subtask, a paradigm frequently employed to investigate WM in the framework proposed by Baddeley and Hitch, a list of items is initially held in memory [15]. Subsequently, another subtask, such as

mental arithmetic or reasoning, is executed, followed by the recall of the initial list of items. Baddeley and Hitch postulated that interference between the original items held in memory and the secondary distractor subtask would only occur when both subtasks depend on the same component. In other words, interference arises when both subtasks simultaneously use the same type of information codes stored within the same domain-specific store [3, 19]. For instance, interference is expected when two verbal or spatial subtasks are performed concurrently. Still, it is not anticipated when verbal and spatial subtasks are presented together within the same complex span subtask [3].

It's crucial to grasp how this model illustrates the relationship between WM and long-term memory. When the model was first introduced, it implied that the central executive had its own versatile memory system that worked separately from the specific functions of the other components [19, 20]. Subsequently, this aspect of the central executive's memory was initially removed from the model and later reintegrated under the name of the episodic buffer [21]. This component, or the previously excluded aspect of the central executive, played a critical role in explaining the handling of semantic information, which did not align with any domain-specific component. Significantly, it established a structure for comprehending how the proposed WM system interacts with long-term memory [20]. Although this model has seen WM and LTM as functionally distinct memory systems, it's essential to emphasize that WM's ability to operate across different modes relies heavily on long-term memory. These operations encompass activities such as retrieving the connections between two facets of an object [20], recombining information from long-term memory to interpret it in novel ways, discovering new information or solutions based on existing long-term memory data, and encoding the output of WM processes into long-term memory [22].

1.2.2 Insights from State-based Models: Examining the Embedded Process Model

For a considerable duration, human memory research has grappled with understanding the characteristics of information temporarily retained in memory and how it differentiates from long-term memory. Early frameworks had varying perspectives on how information transfers from long-term memory to temporary storage [23]. Some of these models, such as Hebb's memory framework, proposed that temporary memory essentially comprises information from activated cell assemblies, essentially being a heightened activation state within long-term memory. Norman's model of memory and attention also suggested that stimuli could automatically activate concepts from long-term memory [24].

Other frameworks delved into the theoretical question of whether memory could be activated outside of conscious awareness [23, 25]. Concepts like William James' [26] "primary memory" suggested that information could be temporarily held by bringing it into conscious awareness. One possible resolution to this theoretical issue was to consider short-term memory and awareness as synonymous, with the level of activation being the sole determinant of awareness.

While many different views emerged from early research, especially regarding the distinction between the focus of attention and activation aspects of short-term or WM information, Cowan's "embedded process model" emphasized the importance of this distinction. According to Cowan, WM encompasses all the information currently accessed or used through an embedded organization of the focus of attention (FoA), activated long-term memory (aLTM), and long-term memory (LTM). In this model, information in the focus of attention is a subset of aLTM, and aLTM is a subset of LTM [25] (Figure 1.2).

The focus of attention is seen as a limited-capacity state that holds conscious information for manipulation, while aLTM information exists outside of conscious awareness or attention. The focus of attention can be directed to a subset of activated memory either automatically (bottom-up) due to certain characteristics of the information or through central executive processes influenced by rules or instructions. One key aspect of this model is the differentiation between representations in aLTM and the focus of attention. The focus of attention has a discrete capacity limit of 3-4 chunks [27], whereas aLTM has no such item-related capacity limits; its functioning is primarily affected by interference and the passage of time [28]. Processing in the focus of attention remains unaffected by the information held in aLTM. Additionally, aLTM information is assumed to decay over time and be susceptible to interference, while the information in the focus of attention is safeguarded against forgetting [2]. Furthermore, representations in aLTM consist of pre-existing information that is activated, whereas individual representations in the focus of attention can be combined to form novel ones [25].

This model does not involve multiple modules, but its predictions regarding interference align with those of the multiple-component model. It suggests that interference is more likely to occur when the held items share common features, not necessarily because they are stored in the same module but because items in the activated portion of long-term memory with similar features will experience interference due to their reliance on the same neural apparatus for processing those features [3].

1.2.2.1 Concentric Model of Working Memory

For some time now, it has been obvious that WM is responsible for both storage and processing duties, and how these functions interact with one another remains a critical subject. While the simplest explanation for this interaction posits a single resource for both storage and processing, evidence from studies examining this interaction suggests scenarios where processing information remains unaffected by concurrently holding other information in

memory. Such evidence often arises from dual-task studies or complex-span WM subtasks that emphasize the dual nature of WM memorization and manipulation which was also a focus in Baddeley & Hitch's book, emphasizing the distinct operations of storage components and the central executive [12, 29-31]

Cowan's embedded WM model pays homage to these findings by proposing a distinction between the information available for immediate access and use (the focus of attention or FoA) and information held in the background for further processing (activated long-term memory or aLTM). This conceptualization suggests that information in aLTM can be maintained in the background without disruption from FoA. Consequently, selective attention or FoA emerges as a potential theoretical solution to the seemingly independent findings related to storage and processing in some complex span WM studies.

Although the idea of FoA or attentional focus in WM appears conceptually neat, Oberauer highlights different interpretations of FoA in various empirical studies. In Oberauer, Demmrich, Mayr, & Kliegl's study [32], participants performed a complex-span subtask involving holding three to six digits in mind while concurrently working on mental arithmetic problems. The digit list and mental arithmetic subtask were entirely independent in one condition. In another, participants had to replace some digits in the arithmetic subtask with the digits in their memory. The information held in memory affected performance only when the retained list and arithmetic subtask was not independent. This finding could be interpreted in two ways based on Cowan's embedded WM theory: either the entire list is brought into FoA during the arithmetic subtask, consuming some of FoA's capacity, or a relevant item is selectively brought into FoA from aLTM, with more extensive memory lists making it more challenging to find the appropriate item.

Another study by Garavan [33] provides evidence for the attentional focus in WM. Participants were tasked with independently counting squares and triangles presented serially on the screen. Results indicated that switching

between mental counters took significantly longer than updating the same counter. This suggests that one counter was held in FoA while the other was stored or memorized outside of FoA.

Although both studies offer evidence for FoA, they suggest two different functions attributed to it: Oberauer's interpretation posits that FoA can hold several items (around 4) simultaneously. In contrast, the second interpretation and Garavan's study imply a narrower attentional focus primarily directed at a single object [32, 33]

In response to these varying interpretations of FoA, Oberauer developed his concentric model of WM. This model conceptualizes WM as a concentric structure comprising three distinct functional representational regions. The first is aLTM, which temporarily memorizes information for later use. The second region is the region of direct access, which holds several chunks for current cognitive processing but has capacity limitations. The third region is the FoA, which can hold only one selected chunk at a time for the following cognitive operation. Oberauer asserts that Cowan's model's capacity limit of 4 ± 1 chunks aligns with the capacity limit of the region of direct access in this current model, corresponding to Cowan's FoA.

In a 2002 paper, Oberauer presents object-switching as further evidence for the necessity of a narrower attentional focus. When an element within a list designated to be held in the region of direct access (or Cowan's FoA) is reached, reaction times for accessing this object are higher if it was not the object used in the previous operation compared to the object accessed in the last operation [12].

In the concentric model of WM, capacity limits do not result from limited available resources but rather from the challenge of selectively accessing concurrently held WM representations. Oberauer introduces two mechanisms—overwriting and crosstalk—as factors influencing WM's capacity limits. Overwriting occurs when similar features in WM representations are overwritten, while crosstalk entails competition among

elements held in the region of direct access. It's important to note that crosstalk does not involve aspects in aLTM, as those are not considered for selection into FoA [12, 34].

To summarize, the concentric model of WM views WM as an organized set of representations with heightened access to current cognitive processes. Within these representations are hierarchically differential access levels among elements belonging to three different functional regions: aLTM, the regions (Figure 1.3).

1.3 Working Memory Capacity and Limitations

WM is often described as having a restricted capacity, and various theories have touched upon how they deal with issues related to this capacity. Nevertheless, it is worthwhile to allocate a dedicated section to delve into working memory capacity (referred to as WMC), given its intricate connection to numerous theoretical and empirical intricacies regarding the operation of WM. Assessments of WMC play a vital role in determining variations in cognitive and intellectual capabilities, as well as cognitive development across the lifespan, from childhood to old age. Consequently, gaining insight into why WM is seemingly limited in capacity is essential for comprehending the constraints on cognitive and intellectual abilities, their evolution over time, and the differences between individuals.

The breadth of research on WMC necessitates a selective focus on findings that illuminate the reasons behind the limitations in WM capacity. Oberauer [35] identifies three principal theoretical contenders that directly manifest the presence of WMC within this domain: Firstly, memory accuracy tends to decrease as the quantity of information to be retained in memory increases, commonly known as the set-size or memory load effect. As an illustration, individuals experience growing difficulty in recollecting visual stimuli as the number of stimuli in the set for memorization rises. This observation directly signifies the limited capacity of WM and WMC, indicating that amplifying the volume of information will have a negative impact on performance. Eventually,

the volume of information will exceed the capacity constraints. Secondly, the retention of information in WM is susceptible to being forgotten when a demanding distractor task, such as mental calculations or oral reading, is introduced after the presentation of items or sets to be remembered, prior to retrieval. In dual-task paradigms or complex-span settings, the adverse influence of the distractor task on memory is believed to stem from the additional cognitive load imposed on an already restricted capacity responsible for maintaining the stored information. Thirdly, measures of WMC highlight variations in capacity constraints among individuals, with some consistency in these limits observed across different testing methods and tasks in correlational investigations. These correlations indicate the potential existence of a unified construct of WMC.

The theories aimed at explaining the limitations of WMC strive to elucidate these fundamental findings, although it's evident that not every WMC hypothesis can equally explain the same interconnected phenomena with precision. Nevertheless, diverse perspectives can be synthesized into three overarching hypotheses concerning why WM is limited in capacity: decay, interference, and resource constraints.

1.3.1.1 Degradation of Information

The primary explanation for the limited capacity of WM revolves around the notion that WM representations degrade rapidly over time. Theorists proposing this concept, which underscores the significance of decay in WM, would also require mechanisms to safeguard WM representations against deterioration. Subvocal articulation (rehearsal) and attentional refreshing are believed to serve as such mechanisms. The central implication of this hypothesis regarding WM capacity would suggest that individuals can retain information only as long as they can rehearse or refresh it before it deteriorates irretrievably. When it comes to visual WM, as subvocal articulation is ineffective for rehearsing material, the debate intensifies concerning the role of refreshing. Refreshing, which is an attention-driven

mechanism for managing WM information, will be subjected to more extensive examination in the forthcoming sections.

1.3.1.2 Investigating Interference

The WMC interference hypothesis posits that limitations in capacity arise due to the interference among WM representations [35]. Two primary forms of interference are briefly discussed here, though other variations are explored in the literature. The first form, akin to Obeauer's 'crosstalk' in the concentric model section, is termed 'interference by confusion.' This interference results from the competition among items as multiple candidates vie for retrieval. The more activated representation is ultimately retrieved, with some theories suggesting that items remain active during the retention interval. Conversely, others argue that items are inactive after the presentation and are reactivated during retrieval, triggered by contextual cues provided by the subtask [35]. Confusion arises when contextual cues need more distinctiveness to activate the target representation selectively. Contextual cues encompass information regarding the context within which different sets of information should be maintained or cues distinguishing items, such as their order or spatial location. Consequently, the greater the distinguishability of contextual cues for individual elements or different sets of information, the lower the confusion. The second form of interference is termed 'interference by overwriting' [36]. This occurs when similar features of items in WM are overwritten, leading to interference. The severity of such interference increases when items exhibit more remarkable similarity than distinctiveness.

1.3.1.3 Constraints on Resources

The term 'resource' in cognitive psychology serves as an informal construct that denotes a set of phenomena indicating restricted accuracy and efficiency in cognitive operations. When the need arises to define 'resource' as an explanatory construct, it can be characterized as the quantity facilitating cognitive and psychological functions. This definition posits that the efficiency or success of a cognitive function or process is directly linked to the amount of

resources allocated to it [35]. In the WM domain, the utilization of the 'resource' concept is accompanied by specific assumptions. The maintenance of WM information and the execution of related cognitive operations necessitate a certain quantity of allocated resources. The success of maintenance and the speed and accuracy of associated operations are contingent upon the amount of resources allocated to them.

Within the visual WM literature, two distinct understandings emerge regarding the nature of resource allocation to WM representations: discrete allocation and continuous allocation of resources.

1.3.1.3.1 Specific Allocation of Resources

One of the initial inquiries into WMC revolved around the query of how many objects the mind can encompass simultaneously. Jevon's study, as referenced in [28], involved releasing black beans onto a surface, aiming to determine the precise number without explicit counting. The findings revealed accurate estimation for three to four items, yet enumeration (subitizing) exhibited increased errors with more significant item numbers. This marked one of the initial efforts to comprehend the constraints of ongoing cognitive processing, exploring these limits in the context of discrete units.

Similar endeavors to fathom the boundaries of temporal storage during the cognitive revolution are evident in Miller's renowned article, 'Magical number seven, plus and minus two'. Miller's investigations suggested that the number of items recallable in immediate subtasks and the amount of information apprehensible at a single time hovered around seven items [37]. However, he also emphasized that this limit might hinge on how information is grouped. Thus, the limits of recall in immediate subtasks might not necessarily pertain to individual items but to grouped information or 'chunks'. This secondary point introduces ambiguity regarding the theoretical underpinnings of Miller's magical number seven, as it does not strictly imply that short-term memory limitations revolve around seven chunks of information. The proposed common limit of seven items might represent a collective constraint

for individual items and chunked information in a specific immediate recall subtask, such as reciting a telephone number. In fact, the notion of a 'chunk' appears to be oversimplified, primarily because chunks may be hierarchically organized into super-chunks. Without precise knowledge of how information is chunked, using chunking measures as indicators of WMC could be problematic. Nevertheless, the idea of chunks provides the WMC literature with a quantifiable unit, and the seven-item capacity limit served as a practical guideline for many years before the cognitive revolution, suggesting the existence of seven 'slots' in short-term memory [28].

Considering substantial evidence, Cowan asserted that at least one mechanism related to WM should exhibit some capacity limits, possibly the focus of attention (FoA). He also delineated these capacity limits in terms of a specific number of items that can be simultaneously held. When chunking is permitted and regulated, Cowan argued for a limit of three to five chunks of information[28]. In the visual WM literature, the concept of discrete resource allocation is evident in various 'slot theories.' These theories propose that short-term or WM maintenance resources consist of a restricted number of discrete 'slots' allocated to temporarily held items or chunks [38]. If WM capacity is designed for N discrete items or chunks, information in WM can be exclusively from these N items at any given time. If the stimulus set surpasses N items or chunks, only information concerning a subset of N items can be maintained [35]. For example, in a change detection study by Luck and Vogel involving arrays of colored squares, participants' accuracy declined as the number of squares increased, leading to a set-size effect. From the distribution of accuracies, the authors inferred that participants could retain information about four items. Notably, this four-item limit persisted even when participants were tasked with detecting changes in two features from the same object (color and orientation) as opposed to one (color). Their interpretation suggested that this discrete capacity should be understood in terms of conjunctions of features or integrated objects rather than individual features.

1.3.1.3.2 Continuous Resource Allocation

Concepts of capacity limits, such as those found in 'slot theories,' often stem from experimental paradigms utilizing categorical or discrete stimulus sets. These include subtasks like letter recall, digit span subtasks, and change detection subtasks, where the stimulus or changes in the stimulus are maintained discretely. In the realm of visual WM, a suggestion has been made that there are three or four object slots holding items in WM. The entry into these slots is proposed to occur in an all-or-none manner; an item is remembered only if it enters one of these slots, and in such cases, it is remembered with high precision. The item-limit perspective predicts a sharp decline in memory once a fixed number of items is held in memory [28, 39, 40].

However, recent research challenges this classical all-or-none view of WMC. Studies indicate that recall precision diminishes gradually as the number of items held in WM increases[41, 42]. Moreover, it is observed that more salient or goal-relevant items are recalled with higher precision, even at the cost of poorer memory for other items held [43]. A resource-oriented perspective, aligning better with these findings, conceives WM as still constrained in capacity, but the limited resources are flexibly distributed among WM items. This view underscores memory precision or quality as a means of understanding WM limits, focusing on the representational medium allocated to stimuli in a flexible manner[39]. The two fundamental premises of these resource models are that internal WM representations are inherently noisy, and the ratio of noise increases with the growing number of items in WM. The more resources allocated to an item, the more precise its representation.

Visual WM has been investigated using the delayed-estimation technique [41], where both stimulus and response spaces are continuous rather than discrete. This approach is applied to study the memory of visual features like motion direction, color, and orientation, and it measures recall variability as an indicator of WMC. Results show that recall variability gradually increases

as the number of stimuli in the set increases, supporting the idea of shared resources between items. Unlike item-limit models, which would anticipate fixed recall variability after a specific number of items in memory, proponents of continuous resource allocation argue for a variable distribution of resources among items and trials. Drawing on data from various studies, Van den Berg, Awh, and Ma [44] compared models and proposed that resources are variably allocated among items and trials rather than being distributed equally. The following part will go over attentional and non-attentional elements that affect how much WM is allocated to items and subtask trials, as well as how WM is controlled.

1.4 Control of Working Memory Information:

Attention-based Control of Working Memory Information

The limitations on WMC, particularly as proposed by discrete capacity perspectives, are considered indicative of the bottleneck effect in cognitive processing. Central selective attention is posited as the mechanism enabling WM operations to proceed despite inherent limitations, ensuring goal-directed processing remains focused on relevant aspects of the environment[45]. Kane propose a correlation between individuals' WM capacity and their attentional control in complex span subtasks [46]. Moreover, individuals with lower WM capacity are suggested to struggle with filtering out distractions during subtasks that demand WM storage [47]. The association between WM and measures of intellectual ability may stem from the shared requirement for attention control in both WM and intellectual subtasks [5]. However, further investigation is warranted to elucidate this relationship [20]. Performance and reasoning capabilities on a given subtask, as well as the tolerance for complexity, hinge on WM capacity and its allocation. Recent discussions on the role of attention in this allocation of WM resources have centred on three attentional processes: consolidation, refreshing, and removal [48].

1.4.1.1 Processes of Consolidation

Consolidation is proposed to be linked to directing attention towards a memory trace once it is no longer perceptually accessible, thereby temporarily stabilizing that memory trace. In essence, WM consolidation involves the transformation of temporary sensory representations into stable memory traces or representations that can be subsequently processed, specifically manipulated, and retrieved even after a delay period. While consolidation may seem akin to WM maintenance, a critical characteristic is its association with attentional processing immediately after the formation of a new transient sensory trace, aiming to stabilize its trace in WM maintenance [49, 50].

The initial evidence regarding the consolidation process is derived from the attentional blink phenomenon. Attentional blink occurs when there is a time gap of 100 to 400 milliseconds between the presentation of two items to be recalled, resulting in the failure to recall the second item. A pivotal finding supporting the necessity of consolidation processes to stabilize memory representations in WM is the requirement for a specific critical retention interval after the presentation of a to-be-memorized item and before its recall. This retention interval (the time immediately after the presentation of items until recall) is deemed crucial for consolidation in studies, particularly those involving simultaneous item presentations. Notably, a retention interval presented after some processing does not enhance memory performance as much as the initial retention interval, indicating its importance in reflecting the time course for consolidation [51].

1.4.1.2 Refreshing and Removal

According to one theory, refreshing entails the cyclical movement of attention to overall WM contents at any one time, regardless of whether they are brand-new or not. This cyclic process serves to keep the contents activated and prevent them from fading into oblivion. It functions as a central attention-based mechanism for WM maintenance. By continually refreshing, this process facilitates real-time thinking, ensuring that the selected encoded or

retrieved representations of WM remain active indefinitely. In scenarios like a dual-subtask paradigm, refreshing occurs when no attention-consuming processing is ongoing, thereby preventing WM representations from fading away. Attentional refreshing is also suggested to play a role in maintaining representations of various sensory modalities and multimodal information [48, 52, 53].

Theoretical and empirical investigations into refreshing propose the existence of two distinct types. The initial form of refreshing occurs outside of awareness and unfolds swiftly. In contrast, the second type of refreshing can be executed more gradually and intentionally. For example, it may involve consciously contemplating the content of WM or visualizing it in response to a cue, particularly in laboratory settings [48, 54].

Removal is posited as the active elimination of information no longer relevant to the current goal, aimed at creating space in WM, preventing interference with relevant information, and optimizing WM capacity utilization. This process constitutes an adaptive, goal-directed mechanism [55] crucial for keeping WM contents up to date. The concept of removing no longer relevant information from WM has recently gained attention, as the prevailing belief was that WM contents decay rapidly over time. However, studies on visual WM, for instance, challenge the notion of passive decay, suggesting a slow rate of decay (e.g., Ricker, Spiegel, and Cowan's [56] finding of a 10% loss in recognition ability of visual WM contents in 12 seconds). If decay were solely responsible, these findings would imply a visual WM capacity limit much higher than suggested by various views in the literature. Efficient removal of unnecessary information from WM is presumed to rely on faster processing [57].

The removal of no-longer-needed or temporarily goal-irrelevant information from WM, typically explored through the retro-cue paradigm, involves participants encoding a set of stimuli, followed by a cue indicating the relevance or irrelevance of specific stimuli for the next operation.

Subsequently, participants wait for a cue-target interval (a variable amount of time) before engaging in a recognition subtask. Oberauer observed a set-size effect of items cued as irrelevant on the WM operations of the remaining relevant items [58]. Significantly, this set-size effect diminished after 1 second, suggesting the complete removal of items cued as irrelevant from WM over time. Recent evidence indicates that the exclusion of no-longer-goal-relevant information from capacity-limited WM can occur either temporarily or permanently, depending on whether the information, now goal-irrelevant, will become goal-relevant in the future[55].

1.4.2 Hierarchical Control of Working Memory Information

Goal-directed behavior comprises temporally extended subtasks and actions characterized by an intrinsic yet variable hierarchical structure. These behaviors are composed of smaller subtask entities, such as individual events or acts, sequentially followed step by step. Prior accounts acknowledge the existence of entities like Miller, Galanter, and Pribram's 'plans,' which represent information on what to do, how, and when as a single entity. Once formed, plans govern the execution of lower-level steps, with each step potentially having its own lower-level sub-plans. Notably, while these sub-plans unfold sequentially, the overarching plan selects and executes them in sequence, exemplified in actions like 'adding sugar to coffee' involving subtasks such as 'getting a spoon,' 'scooping sugar,' and 'putting it in the cup.'

Computational approaches, especially in ecologically valid settings, often employ similar concepts when investigating hierarchical behaviors or routines. The interrelated components of an action are cast into a higher-level action termed 'temporal abstraction,' abstracting over temporally extended sequential lower-level acts. Once instantiated, a temporal abstraction operates until a termination point is reached. Such temporally extended goal-directed behaviors involving sequential execution of subtasks are referred to as 'subtask episodes [59, 60].'

Farooqui and Manly emphasize that task episodes, construed as a single cognitive entity, consist of smaller subtasks and processes. The interpretation of a task episode is not solely based on its physical structure but rather on the doer's construal. This construal can be manipulated with task-irrelevant cues, such as the order of trials, operations, or breaks, allowing investigation of task-episode-related phenomena without confounding factors related to WM, long-term memory recall, subtask rules, or experimental paradigms [61].

Hierarchical entities, known by different names such as plans or schemas, are studied in cognitive control literature. However, it's crucial to distinguish the control embodied by hierarchical entities (task episodes) from executive functions or attention. Hierarchical control organizes, controls, and runs hierarchical subtask entities, whereas executive and attentional control focus more on higher-level processes to monitor and implement lower-level processes strategically[61, 62].

Existing models propose that the hierarchical structure of sequential behavior should be directly reflected in the cognitive system as a hierarchy of plans or schemas. This hierarchical representation is suggested to influence WM processing, as evidenced by phenomena observed in sequential item or motor memorization and recall subtasks. The first item tends to have larger reaction times (RTs), particularly in complex or more extended sequences, reflecting the time needed for instantiation. Switch costs and benefits of repeating the same type of item are influenced by the hierarchical relationship between the subtask-related program and lower-level trial-related configurations[62].

Studies exploring hierarchical control of WM processes in unpredictable hierarchical behavior reveal that goal-related control processes still operate, even without precise anticipation of task episodes[62]. Farooqui and Manly suggest the presence of 'task episode-related programs,' conceptualized as cognitive entities controlling temporally extended sequences of behavior[61]. These programs are crucial for goal-directed cognitive control, guiding

cognitive processes until task completion. Farooqui and Manly's research indicates that longer reaction times and decreased accuracy at the beginning of a task episode are due to the time required for program instantiation, especially in complex tasks.

Neuroimaging studies further support the presence of task-episode-related programs, revealing distinctive activation patterns corresponding to hierarchical levels of task completion. The cognitive load of a program established at the beginning of a task episode is highest initially and decreases as the episode progresses or lower-level goals are accomplished. This result is consistent with findings in neuroimaging, indicating that brain regions linked to cognitive load exhibit peak activation at the start of an episode, gradually diminishing as the episode progresses. In summary, the synergy between behavioral and neuroimaging investigations enhances our comprehension of the hierarchical regulation of temporally extended goal-oriented behaviors [62].

1.5 Current Research

Visual WMC theories and research propose various conceptualizations of capacity, such as a 'magical' number of items or chunks (e.g., Cowan's four, Miller's seven) or a flexible allocation of resources to WM items. Despite these theoretical frameworks, everyday observations suggest that individuals can retain more information in WM than predicted by these models. Moreover, people adeptly handle multiple purposeful tasks simultaneously, involving the maintenance and processing of information, while flexibly managing seemingly hierarchical behavior associated with these tasks. These commonplace observations motivate the exploration in this thesis.

The primary objective of this thesis is to scrutinize the behavior of WMC concerning the hierarchical control imposed by goal-directed task entities. The execution of attentional and executive control within a subtask episode occurs within a context where a goal-directed entity oversees the behavior's execution [61]. This overarching control is believed to be facilitated by the

intermediation of subtask episode-related programs. If such is the case, then subtask-relevant information being maintained for subtask execution is likely to be maintained as part of such programs. This will mean that information related to different but concurrent subtasks will be held via different subtask related programs. These different programs may provide non-interfering stores for WM information related to different subtasks.

The secondary aim of this thesis is to provide an account or initiate a discourse on the interplay between attentional and hierarchical control of WMC. This involves delving into the intricate dynamics between attentional processes and the hierarchical structures inherent in goal-directed subtasks, shedding light on an area that remains unexplored in existing research.

How information related to different but simultaneous subtasks (e.g. during multi-tasking) are held in WM is not known. Most accounts, though, posit a unitary store for all WM items, regardless of whether they are related to the same or different concurrent subtasks. This is the case for accounts that see WM capacity in terms of a fixed number of items (e.g., ‘magical’ number seven or four) as well as for accounts that see this capacity as a limited resource (references). All of these accounts predict that behavior will show decrements as this unitary WM store is filled up with more and more subtask-relevant information.

In contrast to these accounts, Farooqui and Manly [61] posited that subtask-relevant information may be maintained as parts of the programs through which the subtask is being executed. Thus, in a multitasking situation where subtasks A and B are concurrently being executed, information related to subtask A will be maintained as part of the program underpinning its execution and information related to subtask B will be maintained separately as part of a distinct program underpinning subtask B execution. In such a case, increasing the information load related to subtask A may make the subtask A program more complex and decrease subtask A performance but not affect subtask B performance. In contrast, increasing the information load related to

subtask B may make the subtask B program more complex and decrease subtask B performance but not affect subtask A performance.

The experiments done as part of this study tested these two sets of predictions. We created situations where participants had to simultaneously keep in mind information related to different subtasks and tested if increasing the information load related to one subtask affected the performance on the other subtask.



CHAPTER 2

Experiment 1

The configuration of the continuous hierarchical behavior execution is manifested in the handling of WM data. Behavioral manifestations of WM operations indicated the establishment and limits of hierarchical subtask units, while neural indicators denoted the commencement and conclusion of these subtask sequences. However, before this thesis, the investigation of the relationship between WMC related to different subtasks and the execution of maintenance in separate non-interfering stores has been limited [61, 63, 64].

In the first experiment, we sought to uncover the interrelationships between the hierarchical implementation of WM operations and WMC. To accomplish this, we devised a novel temporally protracted visual WM experimental framework that incorporated two WM duties (subtask A and subtask B) possessing nearly identical requirements hierarchically arranged into a drawn-out subtask scenario. The newly introduced paradigm involves a subtask episode that comprises the sequential execution of subtasks A and B. This subtask episode requires participants to memorize and recall information from both subtasks, with the execution of subtask A being interleaved with that of subtask B. This subtask episode is organized hierarchically and aims to achieve higher goals related to goal-related behavior. Specifically, the completion of subtask A and subtask B is required to achieve these higher goals.

The most essential features of the elongation of this subtask are explained below:

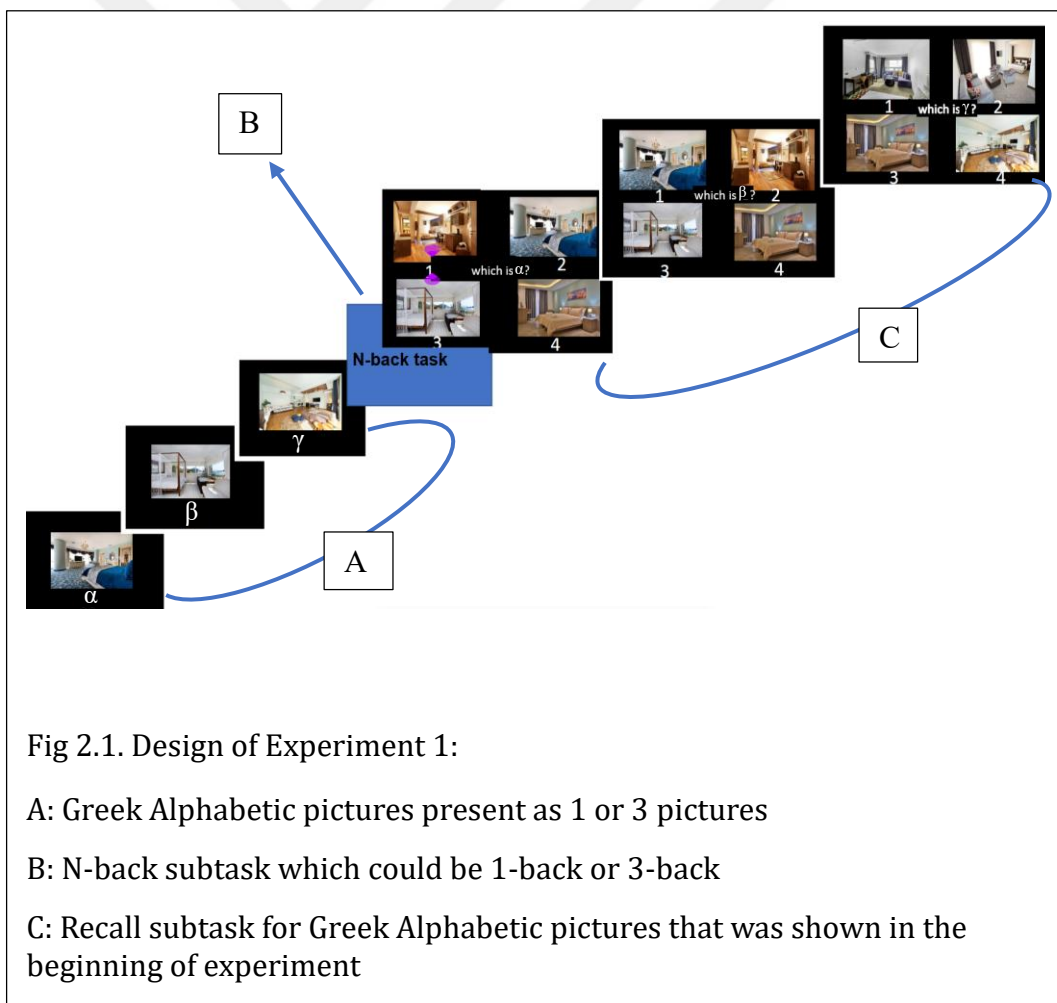
Firstly, the subtasks within an extended subtask episode are organized hierarchically and involve WM subtasks, which necessitate maintaining a relevant set of stimuli and manipulating them to align with the current subtask objectives.

Secondly, the interleaved subtasks share similar subtask demands regarding the stimuli to be stored, operations to be carried out, and overall goal of the subtask (memorizing visual information and recalling specific contents when prompted). Consequently, the representations of these subtasks in WM should be highly alike, and the subtask-related processes initiated by the subtasks should be nearly identical. The current experimental setup is distinctive as it comprises two subtasks that are identical in nature, possess similar subtask demands, and are structured hierarchically within an extended subtask episode.

Thirdly, the design of the extended subtask episode enables us to manipulate the loads of both subtasks independently. This facilitates the examination of WM load effects with this design. The memory-load or set-size effect, a widely observed phenomenon in research, indicates that increasing the amount of information leads to performance deterioration, highlighting the capacity limitations of WM. This effect suggests that the items in a set held in WM draw upon the same limited WM resources. Consequently, augmenting the information in a set result in either greater depletion of resources or reduced allocation of resources to individual items, leading to poorer recall performance when the WM load is higher. By utilizing the current experimental paradigm to independently adjust the loads of both subtasks, we can explore how the WM load of one subtask impacts its performance, influences the load of the other subtask, or affects the performance of the other subtask. These investigations can be analysed in the context of WMC concerns presented by the novel paradigm.

Fourthly, in the extended subtask paradigm being developed, visual stimuli are utilized, with a notable emphasis on the continuous nature of both the

stimulus and response domains, as opposed to discrete. This characteristic of the present experiment aligns with contemporary studies highlighting the continuous distribution of WM resources towards WM representations. Consequently, this approach enables the examination of set-size or load effects through the evaluation of memory precision using continuous measurements of accuracy and reaction time (RT) instead of categorical ones. Such a methodology proves to be more responsive to research on WMC, focusing on memory precision or quality as a means to comprehend WM limitations rather than the quantity of items retained [39]. The task design has been presented in Figure 2.1.



When considering the points mentioned above collectively, using the extended subtask paradigm in Experiment 1 facilitated an investigation into the correlation between WMC and control as instantiated by subtask entities organized hierarchically and executed concurrently.

The classical dual-task paradigm is a primary methodology extensively employed in WMC research. This is due to its emphasis on the dual nature of WM testing, specifically the maintenance and manipulation components. The extended subtask episode that was developed in Experiment 1 to elucidate the relationship between hierarchical control of WM and WMC exhibits a multitude of deviations from the classical dual-task paradigm:

1. While the dual-task paradigm is widely utilized in WMC studies, there is a requirement for additional inquiry into how the subtasks impact each other reciprocally. Typically, the focus has been on the impact of distractor subtasks on WM subtasks rather than the reverse. However, a newly introduced subtask paradigm presents an opportunity to explore the mutual effects of concurrent subtasks on one another. Notably, the similarity between the subtasks highlights the hierarchical organization of the function rather than a WMC issue arising from the unique demands of a given subtask.

2. In classical dual-task studies, the distractor subtask has primarily been an attentional subtask rather than a WM subtask, resulting in a focus on investigating attentional bottleneck or capacity limits. Consequently, in a classical dual-task context, it is rare to observe the impact of a second WM subtask on another WM subtask, let alone two subtasks of the exact nature.

3. WM subtasks used in dual-task paradigms often share the characteristic of categorical or discrete stimulus and response spaces. This has led to an all-or-none understanding of WMC, with the number of items remembered being the main focus rather than considering the possibility of continuous allocation of WM resources. The extended subtask paradigm presented in this thesis involves two hierarchically organized visual WM subtasks, with information from one subtask being temporarily withheld while performing the other and

then retrieved. This experimental paradigm offers insight into how WM information from different hierarchically organized subtasks of the exact nature within the same extended subtask episode is handled in terms of WMC.

The main goal of Experiment 1 was to explore the connection between the hierarchical management of WM and WMC using the recently introduced expanded subtask framework. Mainly, the study aimed to examine the impact of the withheld subtask A WM load on the concurrently processed subtask B performance and the effect of subtask B WM load on the subtask A performance, which were pivotal to achieving this aim.

2.1 Methodology and Procedural Details

The experimental process received approval from Bilkent University's Research Ethics Committee. Participants were attracted using a convenience sample technique, and email lists belonging to the institution were used to send out recruitment emails. Individuals aged between 18 and 65, possessing regular or appropriately corrected vision, and without any neuropsychological disorder were prerequisites for inclusion in this research.

The Pavlovia platform (www.pavlovia.org) was used for our online tests. With the help of the PsychoPy constructor, experiments were coded. Using the PsychoJS package, Python code produced by the builder was transformed into JavaScript so that investigations could be carried out on a Pavlovia-hosted web server. Each participant was individually introduced to the experimenters during a scheduled Zoom conference, during which the participants were given instructions on the method and the experiment's guidelines. The experimenters individually met each participant in a Zoom meeting, giving them the information they needed to provide their informed consent to be signed electronically. They were informed that if they completed the whole work, they would either receive course credit or a payment of £50 for their time.

The experiment link was given to the participants by the experimenters; they would click on it and run the test locally on their PCs. They would be given the connection to the main experiment after completing enough practice under the instructor's supervision through a screen share, at which point the Zoom meeting would come to a conclusion. After each session, the participants conducted an experiment at their own pace with suggested breaks. Through the Pavlovia server, data was generated and gathered. Subsequently, participants were instructed to email the experimenter to signal the completion of the study and specify their chosen payment method. The charge was transferred to their bank accounts if they chose monetary recompense. Finally, they received an email debriefing regarding the trial.

2.1.1 Participants

Recruitment took place among 19 healthy volunteers (12 females) in the age range of 21-55 ($M = 27.31$, $SD = 7.01$). Before participating in the trial, participants received digital versions of the informed consent documents and electronically signed them. Participants received course credits or, if they attended both sessions, a pay-out of £50 as compensation. Participants were questioned about their biological sex and birthdate prior to taking part in order to compile descriptive data.

2.1.2 Setup Of Experiment

The experimental subtask was a variation of the n-back subtask. The subtask commonly referred to as the n-back subtask necessitates that the individuals participating in it decide whether or not each of the stimuli in a given sequence matches the one that had appeared precisely in a number of items ago. For example, look at Fig 2.2. In the 1back subtask, participants see the target, and they are going to see if the target is the same as the previous picture or not. As you see in Fig 2.2. the target is [t], and one picture before that is [T].

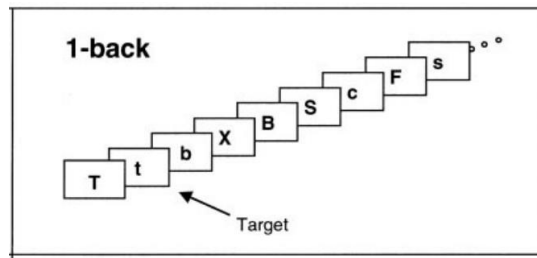


Fig 2.2. 1 back subtask example,

which the target is going to be compared with the 1 previous picture.

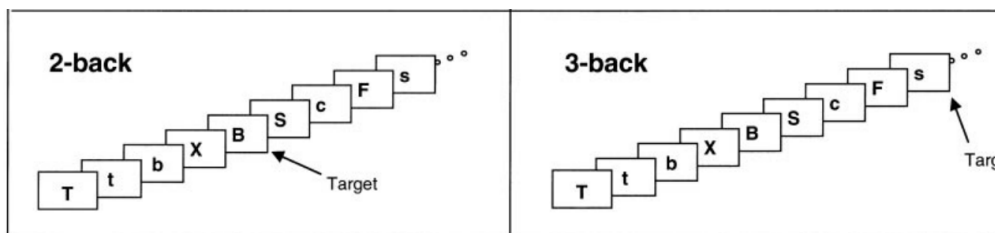


Fig 2.3. Two-back and three-back subtask example,

which the target is going to be compared with the two and three previous pictures respectively.

Also, the target should be compared with the two previous pictures in the two back subtasks. As you see in Fig 2.3, the target is [B], which will be compared with two previous pictures, which is [b]. The same for the three back subtask targets, which is [s], is going to be compared with three pictures before, which is [c].

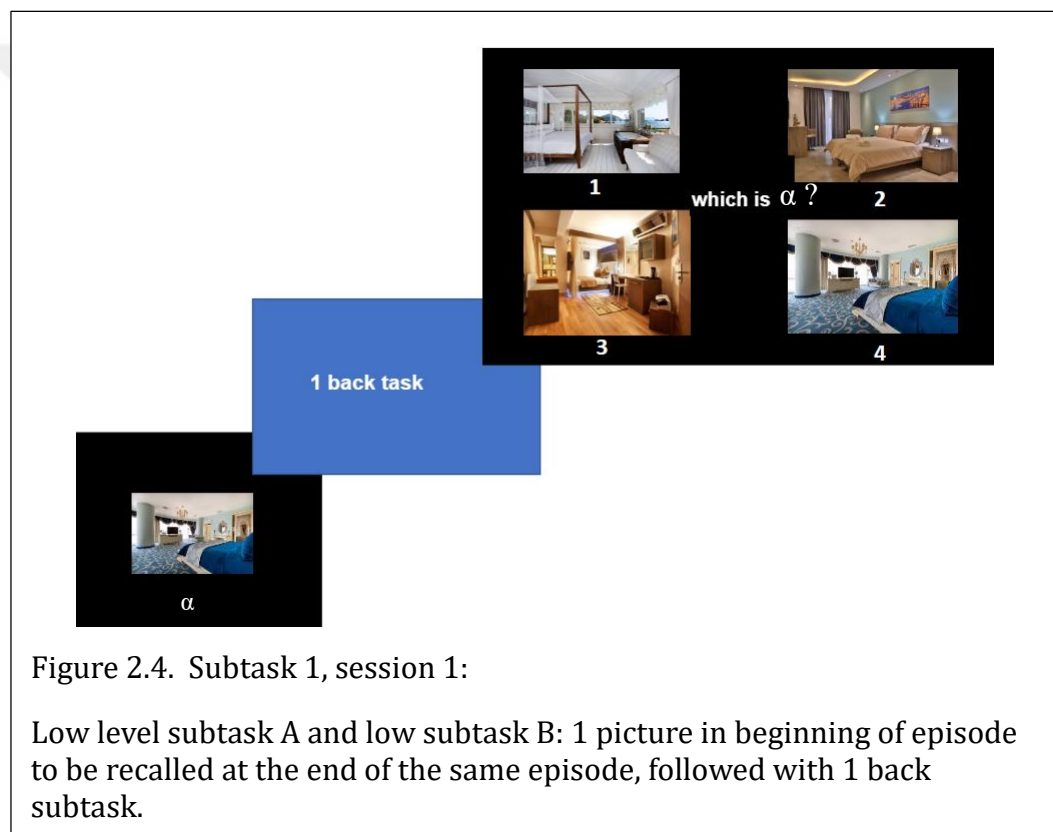
In our experiment, participants saw indoor scenes in the middle of the screen. They saw one picture at a time.

The pictures were randomly selected among 33 pictures for each experimental run. We use the N-back subtask as subtask B. It was a 1-back subtask for low-level and a two or three-back subtask for high-level subtask B. The subtask required comparing whether the current picture on the screen was the same as the 1,2, or 3 previous pictures or not.

Each session initiates with subtask A, during which participants observe either a single picture representing the low level or three pictures

representing the high level of subtask A. Following the completion of subtask A, participants then move on to engage in subtask B.

In total, the experiments have four sessions, and each participant has done all of the sessions. The first session is for the low level of subtask A and the low level of subtask B. As shown in Fig 2.4, the session starts with one labeled picture in the beginning, which is for the low level of subtask A, and then the low level of subtask B will be done, which is a one-back subtask.



In session two, subtask A is at a low level, which means one labeled picture at the beginning of the session will be shown, and then subtask B is at a high level, which means two or three-back subtasks will be done. (See fig 2.5.)

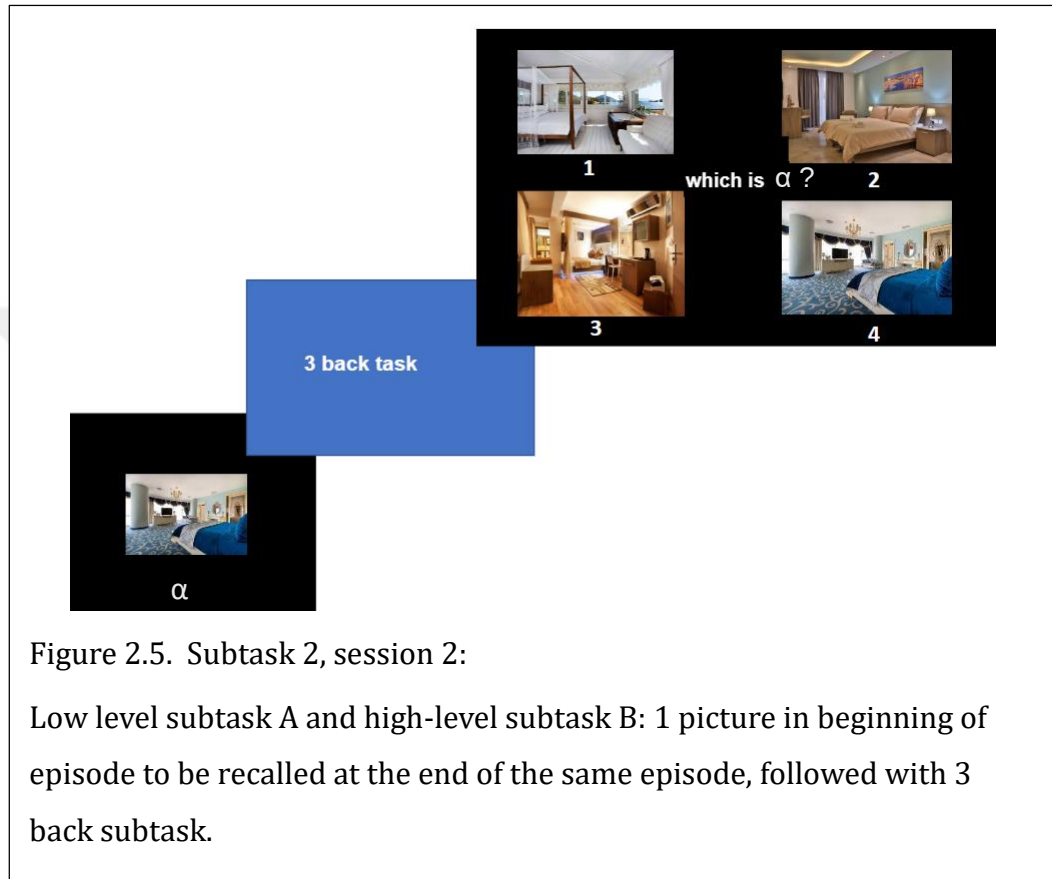
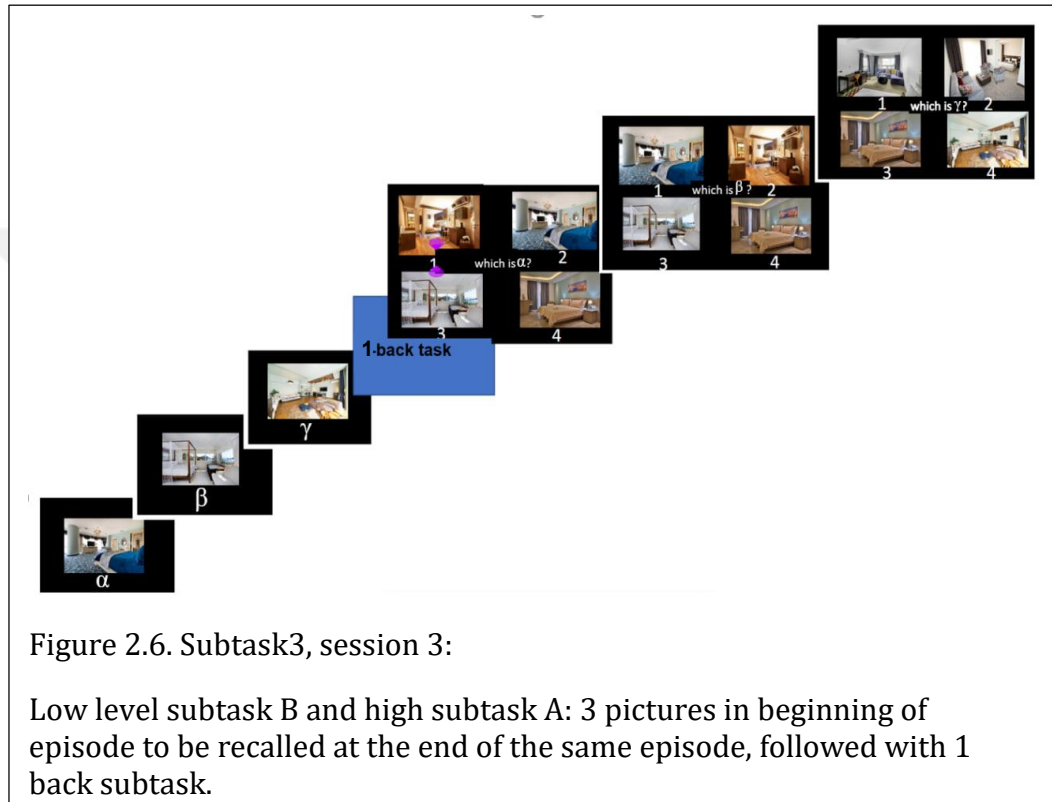


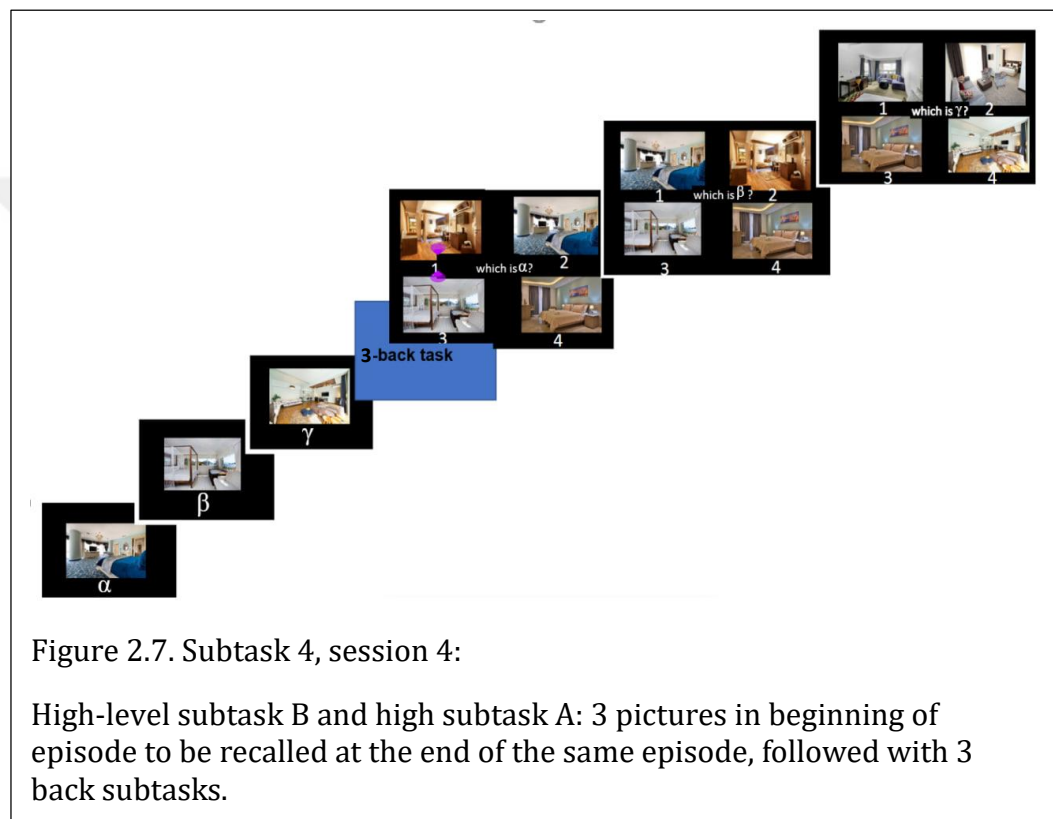
Figure 2.5. Subtask 2, session 2:

Low level subtask A and high-level subtask B: 1 picture in beginning of episode to be recalled at the end of the same episode, followed with 3 back subtask.

In Session three, subtask A is at a high level, which means three labeled pictures at the beginning of the session will be shown, and then subtask B is at a low level, which means one back subtask will be done. (See fig 2.6.)



In session four, subtask A is at a high level, which means three labeled pictures at the beginning of the session will be shown, and then subtask B is also at a high level, meaning two or three-back subtask will be done. (See Fig 2.7.)



2.1.3 Data Analysis

The unprocessed data is primed for a variety of analyses within R-studio [65]. The analytical procedures and visual representations of the analysis outcomes were generated utilizing R (R Core Team, 2023) and R-Studio (RStudio Team, 2023) by utilizing the mean and 95% Confidence Interval (CI) data acquired from RStudio analyses. Tables presenting descriptive statistics, ANOVA results, and pairwise comparisons are extracted from JASP [66, 67].

We carried out a 2x2 Repeated Measures Analysis of Variance (ANOVA) to examine the influence of subtask A load (defined by the initial presentation of one picture or three pictures, encompassing both the one-back and three-

back tasks) and subtask B load (initially presenting one picture or three pictures, including both the one-back and three-back tasks) on all output measures of interest such as accuracy, presentation time and reaction time for all conditions of experiments. We also performed planned pairwise comparisons with Bonferroni correction to get meaningful ANOVA results.

2.2 Findings

2.2.1 Accuracy

The mean of accuracy in different conditions was calculated as the sum of correct answers in a given condition divided by the total number of trials of that condition. Separately, three-back (Table 2.1) and one-back (Table 2.2) subtasks for both conditions.

2.2.1.1 Accuracy for three pictures at the beginning

Tables 2.1 show the effects of conditions and trial positions on RTs accuracy. Mean of RTs accuracy differed across the three-back and one-back conditions ($F_{1,18} = 10.364$, $p=0.005$, $\eta^2_p = 0.365$). However, RTs accuracy were not different across the three trial positions ($F_{2,36} = 0.291$, $p = 0.750$, $\eta^2_p = 0.016$). The effect of the condition was not different across different trial positions ($F_{2,36} = 0.252$, $p = 0.779$, $\eta^2_p = 0.014$). [see also Figure 2.8 and 2.9]

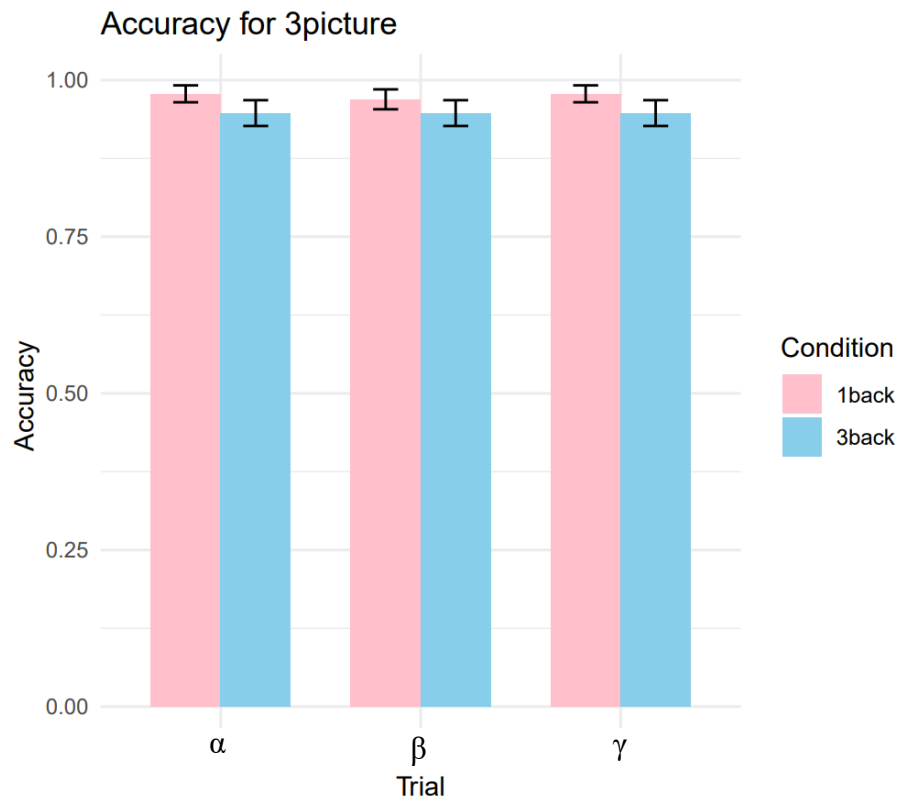


Figure 2.8. Mean of accuracy for 3 pictures in beginning of experiment in 2 conditions with 95% CI error bars

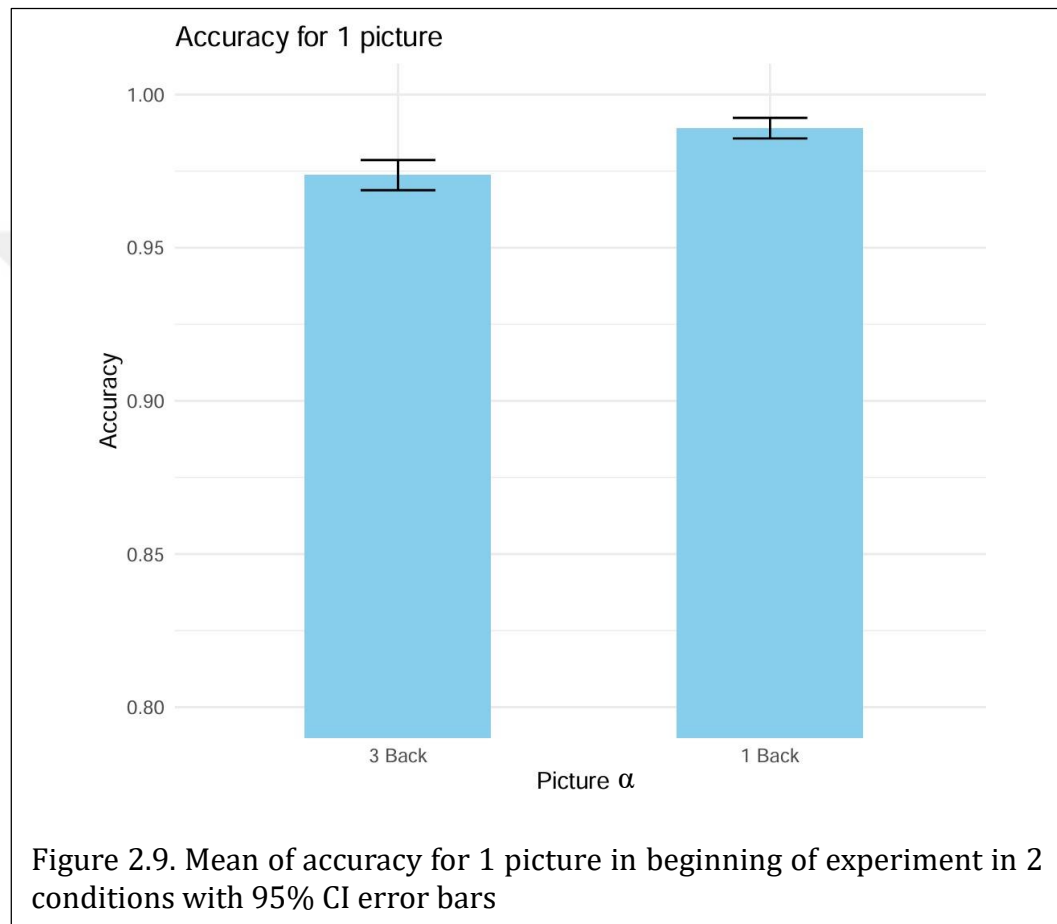
Within subject effects

Cases	Sum of squares	df	Mean square	F	p	η^2_p
Condition	0.022	1	0.022	10.364	0.005	0.365
Residuals	0.038	18	0.002			
Trial	4.873×10^{-4}	2	2.437×10^{-4}	0.291	0.750	0.016
Residuals	0.030	36	8.384×10^{-4}			
Condition* Trial	4.873×10^{-4}	2	2.437×10^{-4}	0.252	0.779	0.014
Residuals	0.035	36	9.670×10^{-4}			

Table 2.1 : ANOVA table of mean of accuracy for 3 pictures in beginning of experiment in 2 conditions.

2.2.1.2 Accuracy for 1 picture in the beginning

Table 2.2 shows the effects of conditions on RTs for the 1 picture in the beginning. RTs were not different across the three-back and one-back conditions mean ($F_{1,18} = 2.827$, $p = 0.110$, $\eta^2_p = 0.136$). [see also Figure 2.9]



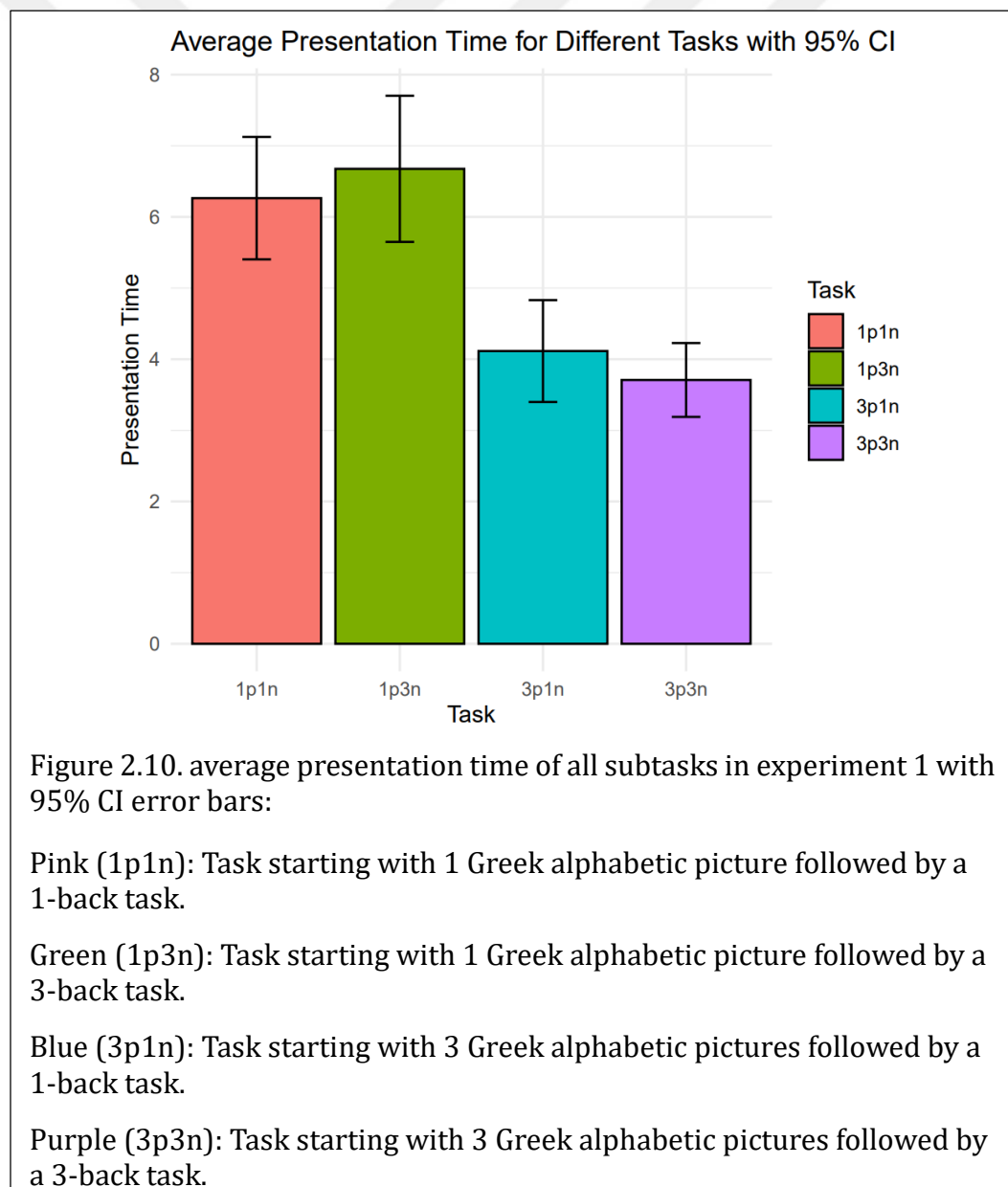
Within subject effects

Cases	Sum of squares	df	Mean square	F	p	η^2_p
Conditions	0.002	1	0.002	2.827	0.110	0.136
Residuals	0.014	18	7.919×10^{-4}			

Table 2.2. ANOVA table of the mean of accuracy for 1 picture at the beginning of the experiment in 2 conditions.

2.2.2 Presentation time

The presentation time was calculated as the mean amount of the time that participants took to memorize the pictures shown at the beginning of episodes as 1 picture and three pictures and followed by n-back subtasks, [the time between the episode started and the participant made the button response to move onto the following picture], in seconds (Figure 2.10). Table 2.3 shows the effects of recall presentation RT. 2x2 Repeated Measures ANOVA (Table 2.3) revealed that there was a difference across the four subtasks ($F_{1,2,23,1} = 6.071, p = 0.016, \eta^2_p = 0.252$)



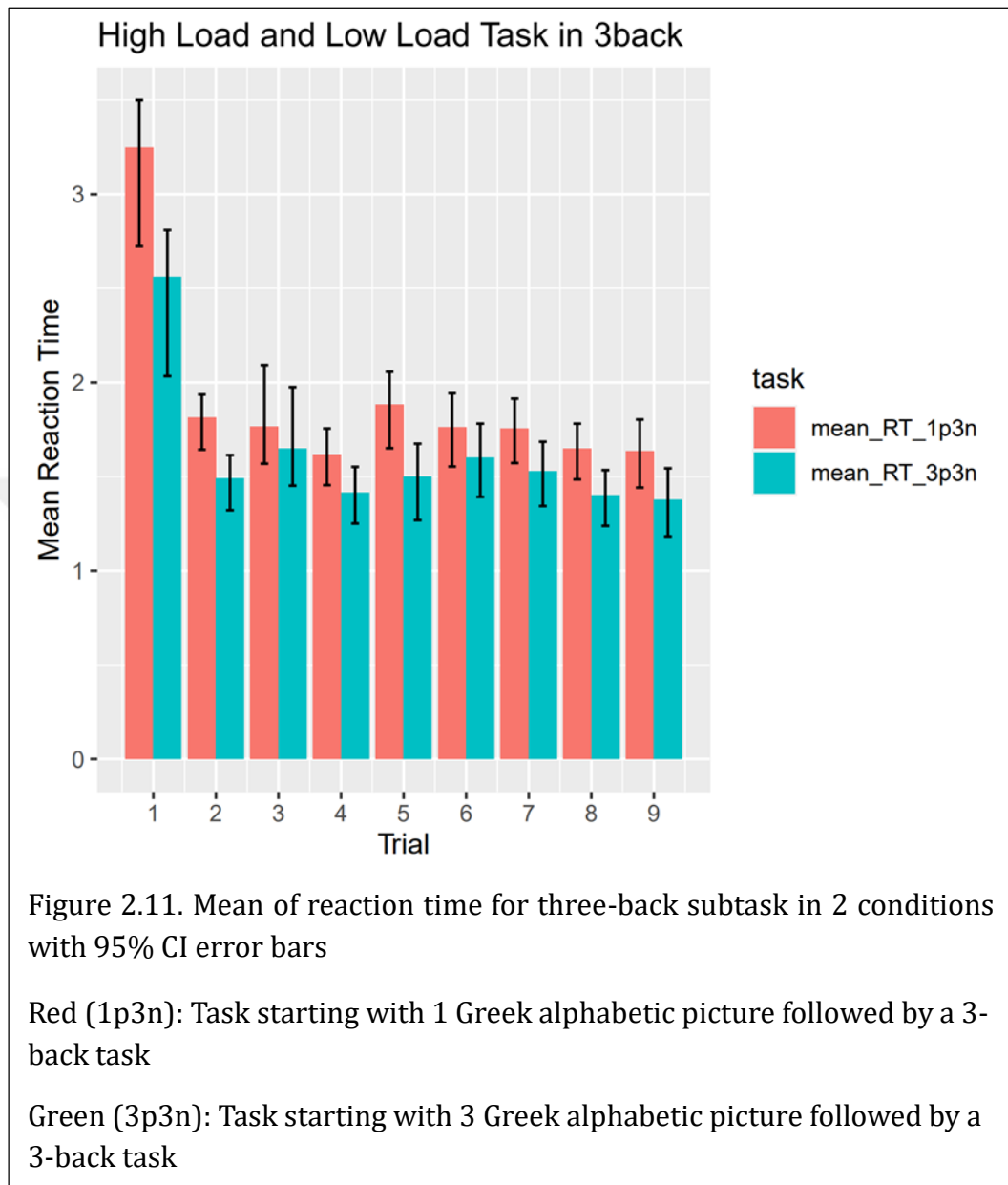
Within subject effects

Cases	Sum of squares	df	Mean square	F	p	η^2_p
Subtasks	127.563	1.286	99.194	6.071	0.016	0.252
Residuals	378.220	23.148	16.339			

Table 2.3. ANOVA of experiment 1 recall presentation in all subtasks (Greenhouse–Geisser corrected)

2.2.3 Reaction time

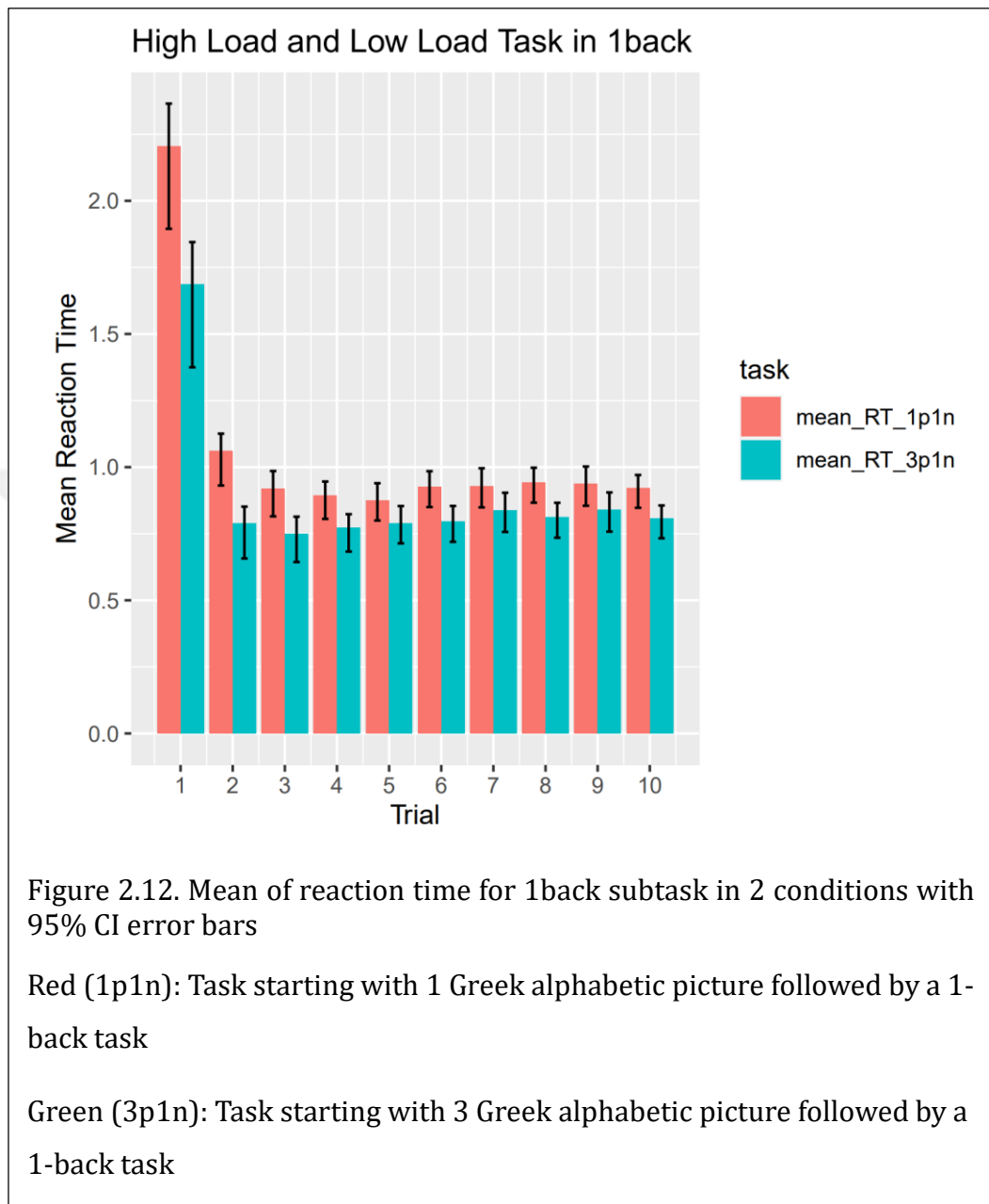
The mean of reaction time was calculated as the mean of the time for participants to respond to each picture in the N-back subtask [9 pictures for three-back and 10 pictures for one-back] in 2 different conditions [Three pictures in the beginning Figure 2.11], [One picture in beginning Figure 2.12] in seconds. Table 2.4 shows the effects of subtasks and trials on reaction time for three-back subtasks in 2 conditions. RTs differed across pictures in the beginning subtask and three pictures ($F_{1,18} = 4.737$, $p = 0.043$, $\eta^2_p = 0.208$). Also, RTs were different trials ($F_{3.3,60.3} = 17.24$, $p < 0.001$, $\eta^2_p = 0.489$). The effect of subtasks was different across different trial positions, as well ($F_{3.2,58.7} = 12.745$, $p < 0.001$, $\eta^2_p = 0.415$). [see also Figure 2.11]



Within subject effects						
Cases	Sum of squares	df	Mean square	F	p	η^2_p
Subtask	6.444	1	6.444	4.737	0.043	0.208
Residuals	24.490	18	1.361			
Trial	33.210	3.351	9.91	17.241	< 0.001	0.489
Residuals	34.672	60.320	0.575			
Subtask *	27.845	3.263	8.534	12.745	< 0.001	0.415
Trial						
Residuals	39.326	58.730	0.67			

Table 2.4. ANOVA table of mean of reaction time in three-back subtask for 1 picture in beginning and 3 pictures in beginning (Greenhouse–Geisser corrected)

Table 2.5 shows the effects of subtasks and trials on reaction time for one-back subtask in 2 conditions. RTs were different across one picture in the beginning subtask and three pictures in the beginning subtask reaction time mean ($F_{1,18} = 47.472$, $p < 0.001$, $\eta^2_p = 0.725$); also, RTs were different across trials ($F_{2.5,46} = 33$, $p < 0.001$, $\eta^2_p = 0.647$). The effect of subtasks was different across different trial positions, as well ($F_{2,36.7} = 31.453$, $p < 0.001$, $\eta^2_p = 0.636$). [see also Figure 2.12]



Within subject effects

Cases	Sum of squares	df	Mean square	F	p	η^2_p
Subtask	3.756	1	3.756	47.472	< 0.001	0.725
Residuals	1.424	18	0.079			
Trial	21.125	2.559	8.256	33.007	< 0.001	0.647
Residuals	11.52	46.057	0.25			
Task *	19.622	2.041	9.616	31.453	< 0.001	0.636
Trial						
Residuals	11.229	36.731	0.306			

Table 2.5. ANOVA table of mean of reaction time in 1back subtask for 1 picture in beginning and 3 pictures in beginning (Greenhouse–Geisser corrected)

2.3 Discussion

Experiment 1 aimed to investigate whether increasing the WM load related to subtask B (recalling pictures labelled with: **α , β and γ**) impacts the concurrent execution of subtask A (n-back subtask). This was done to explore the possibility that WM items related to different tasks or subtasks are maintained as part of different, non-interfering programs.

Our experiment involved a task episode comprising of two interleaved WM subtasks, A and B. This hierarchical subtask structure allowed us to examine whether the WM load of one subtask impacted the concurrent execution of the other subtask.

Specifically, we structured the experiment to monitor how variations in the WM load of subtask B (from 1 to 3 pictures) influenced the performance metrics (accuracy and reaction time) of subtask A, and vice versa.

We found that increasing the WM load of one subtask did not affect the performance of another subtask, suggesting that WM related different subtasks may be stored in separate, non-interfering stores. Our findings support the multi-component model of WM proposed by Baddeley and Hitch [15], which posits that WM consists of separate components that can operate independently. The lack of interference between subtasks A and B suggests that different types of WM content may indeed be processed and stored in distinct subcomponents. This is consistent with prior research demonstrating that tasks involving different modalities can be performed concurrently with minimal interference [18].

As evident in Figure 2.12, reaction time in the three back subtasks for one picture, in the beginning, was greater than that for the three pictures in the beginning for the same subtask. This means that when WM related to Subtask-B was increased from 1 picture to 3 pictures, the reaction time on Subtask-A didn't increase. In contrast, increasing the WM load from 1 to 3 pictures for Subtask-B did increase the time participants took to register the pictures (Figure 2.11).

Likewise, participants were more accurate in the one-back condition of subtask A compared to the three-back condition of the same subtask, which is as expected (Figure 2.8 and 2.9).

As showed in Table 2.2 for the 1 picture presented at the beginning in both the one-back and three-back conditions, our ANOVA results show that the accuracy in the one-back subtask is statistically similar to the accuracy in the three-back subtask in the same condition. In alignment with Table 2.2, Table 2.1 for three pictures in the beginning of the experiment for both low and high load of subtask-A, the ANOVA results indicate that the accuracy is statistically similar between conditions.

Specifically, we found that increasing the WM load of subtask B does not impair performance on subtask A and vice versa. Evidentially Figure 2.10 shows that the time that participants need to memorize the pictures of

Subtask-B in high load for both conditions of Subtask-A was less than the time for memorizing the low load of Subtask-B in the same conditions.

Our study also aligns with theories emphasizing the role of executive control in managing multiple tasks. According to Engle's framework[46], the central executive component of WM is responsible for coordinating multiple tasks and inhibiting interference. Our results suggest that the central executive effectively manages the increased load in one subtask without detrimentally affecting the other, highlighting its capacity for efficient resource allocation and task management.

The hierarchical subtask structure used in our experiment allowed us to search into how WM handles complex, interleaved tasks. The findings are in line with theories suggesting that WM can compartmentalize and prioritize tasks based on their hierarchical importance [25]. This compartmentalization prevents interference and ensures that increasing the load in one task does not overflow and disrupt another, which is a novel contribution to understanding hierarchical task management in WM.

Our experiment uniquely contributes to the hypothesis that WM items related to different tasks or subtasks are maintained in separate, non-interfering stores. This extends the understanding of WM capacity limits, suggesting that these limits may be task or subtask-specific rather than global. Previous studies have often focused on interference within a single modality or task type [12], but our findings indicate that WM's ability to manage multiple tasks may involve separate, isolated stores for different tasks.

The results have practical implications for cognitive load theory[68], particularly in educational and multitasking environments. Understanding that WM can handle multiple concurrent tasks without mutual interference can inform strategies for designing learning and working environments that optimize cognitive load distribution. For instance, tasks can be structured to capitalize on WM's compartmentalization ability, thereby reducing overall cognitive load and enhancing performance.

CHAPTER 3

Experiment 2

In experiment 1, performance on n-back, i.e., subtask A, did not decrease when more subtask B pictures were to be kept in mind. We found that increasing the number of subtask B pictures did not affect subtask A performance, and vice-versa. Thus, an increased WM load related to one subtask did not affect the concurrent WM subtask, suggesting that the WM items were maintained in separate, non-interfering stores.

An objection may be that in experiment 1, subtask B pictures were not maintained in a goal-directed WM store but were passively maintained as part of (e.g.) activated long-term memory. To rule this out, we did experiment 2.

In experiment 2, participants executed even longer trial episodes made of four sequential subtask episodes. Each episode was again made of subtasks A and B and involved the same maintenance of subtask B pictures while executing subtask A. But crucially, the B pictures kept in WM in one subtask-episode were to be recalled not in that episode but in the next, e.g., B pictures from episode 1 (or mB1) were recalled not after subtask A of the episode (i.e. A1) but at the end of next episode (A2).

Trial episodes began with the presentation of subtask B and its memorization. Firstly, a picture which was labeled with the letter 'A' was shown. As in Experiment 1, participants were to keep it in mind. This was followed by presenting the second picture of subtask B, again labeled as 'B,' let's call them [mB1]. Participants were then probed to do a N-back subtask as subtask A, which we call [A1]. Unlike experiment 1, this was not repeated in all episodes. Instead, participants recalled subtask B pictures of the first episode at the end of the second episode[rB1].

That means the second episode starts the same as episode 1, with memorizing two pictures for subtask B labeled with α and β named as [mB2]; then they will do subtask A which is an N-back subtask[A2]; then unlike episode 1, they will recall mB1 which we say[rB1].

Episode 3 was the same as episode 2; first, they do [mB3], then [A3], and then [rB2].

The last episode starts with subtask A, which means they start the episode with N-back[A4], and then they recall subtask B of the previous episode[rB3].

The overall structure thus was: mB1-A1 – mB2-A2-rB1 – mB3-A3-rB2 – A4-rB3.

Note that in episodes 2 and 3, participants keep four pictures in their mind, while in episodes 1 and 4, this number decreases to 2. This makes participants' long-term WM active in removing and replacing the information.

If the number of pictures affected the fidelity of recall of pictures, then performance should be poorer in episodes 2 and 3 compared to episodes 1 and 4. However, if increasing the number of pictures for recalling did not interfere with each other, then performance should be the same across episodes.

The current design of Experiment 2 allowed us to see if the finding of lack of across-subtasks load-effect in Experiment 1 is replicable in Experiment 2, and more crucially, to test the specific predictions of two alternatives for the findings of Experiment 1, namely the 'passive storage idea' and the separate maintenance of information that belongs to different subtask-related programs (programs-related view):

a. If subtask A and B pictures did not interfere with each other, then performance should be the same across episodes.

b. If the number of α pictures affected the recall of β pictures, then performance should be poorer on rB1 and rB2 (i.e., four pictures are held) compared to rB3 (i.e., two pictures are held).

c. If β pictures were held in a passive store, then performance should decline across β 2 and β 3 (i.e., across episode 2 and episode 3) due to interference between β pictures that will start to pile up.

d. If pictures β 1, β 2, and β 3 were organized as parts of different programs, this would predict that recollection of β 1 and β 2 pictures will not show decreased fidelity compared to the recall of β 3 pictures even though during the former four β pictures are being maintained. In comparison, only two pictures are being maintained during the latter.

e. Every WM account other than β pictures being organized as parts of different programs (including the 'passive store' views) would predict that recall fidelity should be poorer for β 1 and β 2 compared to β 3.

The aim of Experiment 2 was to see the replicability of the findings of Experiment 1 and to disambiguate between two alternatives for the core findings in Experiment 1 by simply testing their specific predictions mentioned above.

3.1 Methodology and Procedural Details

The experimental procedure was the same as in Experiment 1. In the second experiment, participants engaged in an N-back task, which varied between 2-back, 3-back, and 4-back conditions. Specifically, out of the 40 participants, one participant performed the 2-back task, while the remaining 39 participants completed either the 3-back or 4-back tasks.

The experiment also included high load and low load conditions between episodes, as mentioned previously at the beginning of the chapter. The high load conditions were Episode 2 and Episode 3, which involved more challenging tasks, such as the 3-back and 4-back tasks. The low load conditions were Episode 1 and Episode 4, represented by the 2-back task.

For the analysis, the mean reaction times (RT) reported represent the average RT across all 40 participants, regardless of the specific N-back condition they were assigned. This overall mean RT provides a comprehensive measure of participants' performance across the varying levels of task difficulty (2-back, 3-back, and 4-back) and load conditions.

3.1.1 Participants

Recruitment took place among 40 healthy volunteers (26 females) in the age range of 22-29 ($M = 24.5$, $SD = 1.78$). Before taking part in the trial, participants received digital versions of the informed consent documents and electronically signed them. Participants received course credits or, if they attended both sessions, a payout of £50 as compensation. Participants were questioned about their biological sex and birthdate prior to taking part in order to compile descriptive data.

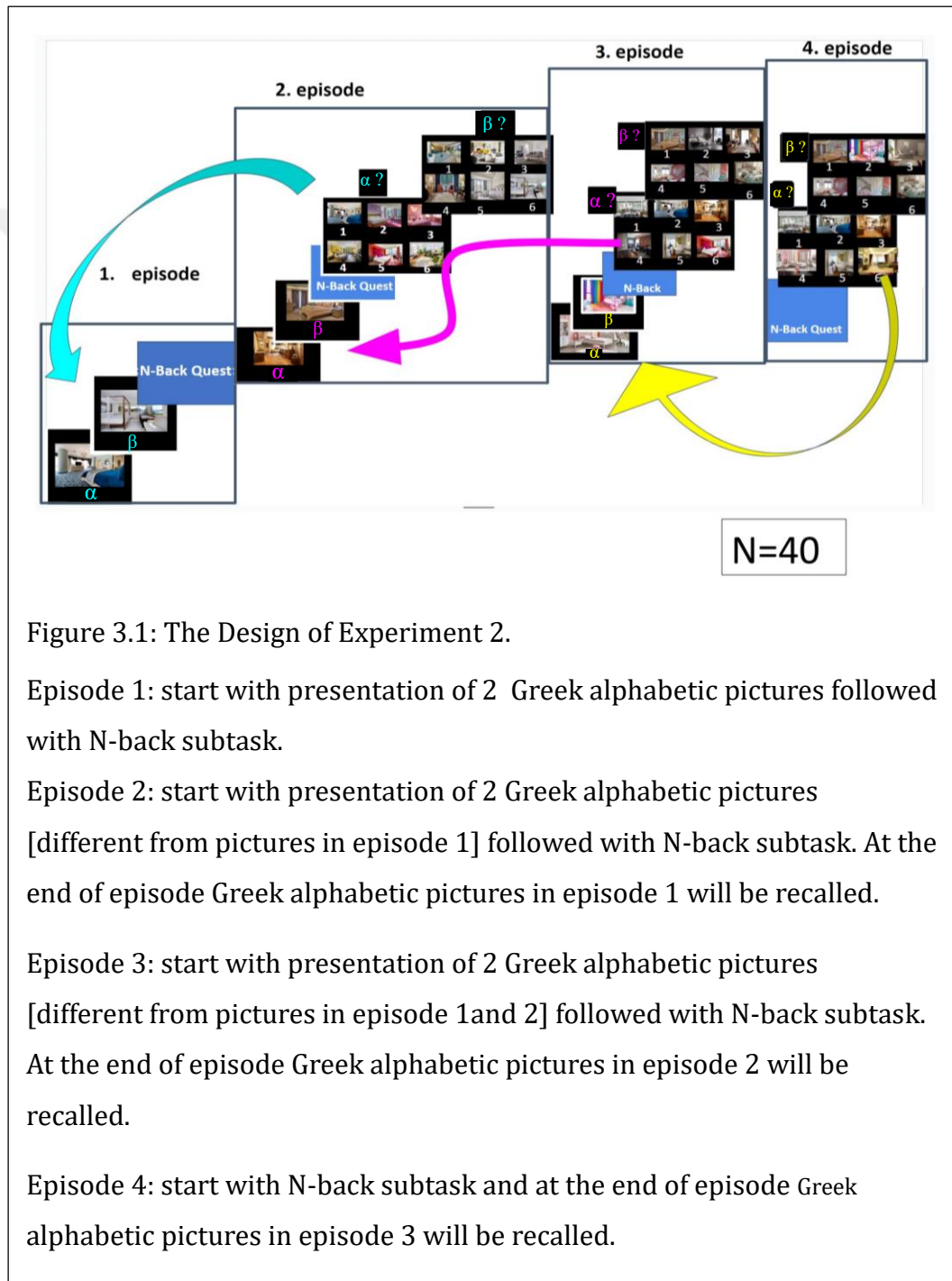
3.1.2 Experimental Setup

In experiment 2, participants executed even longer trial episodes of four sequential subtask episodes. The pictures were randomly selected among 70 pictures for each experimental run. Each episode was again made of subtasks A and B and involved the same maintenance of subtask B pictures while executing subtask A. But crucially, the β pictures kept in WM in one subtask-episode were to be recalled not in that episode but in the next, e.g., β pictures from episode 1 (or mB1) were recalled not after subtask A of the episode (i.e., A1) but after A2. The overall structure thus was:

mB1-A1 – mB2-A2-rB1 – mB3-A3-rB2 – A4-rB3

This forced the participant to keep two sets of β pictures in mind and recall not the immediately preceding β pictures but the one before. This can only be done if β pictures were goal-directed rather than passively maintained. This also meant that during subtasks A2 and A3, participants were maintaining two sets of β pictures ($\beta 1$ & $\beta 2$ and $\beta 2$ & $\beta 3$, respectively). However, during $\alpha 1$

and α 4, they only maintained one set of β pictures (β 1 and β 3, respectively) (see Figure 3.1).



3.1.3 Data Analysis

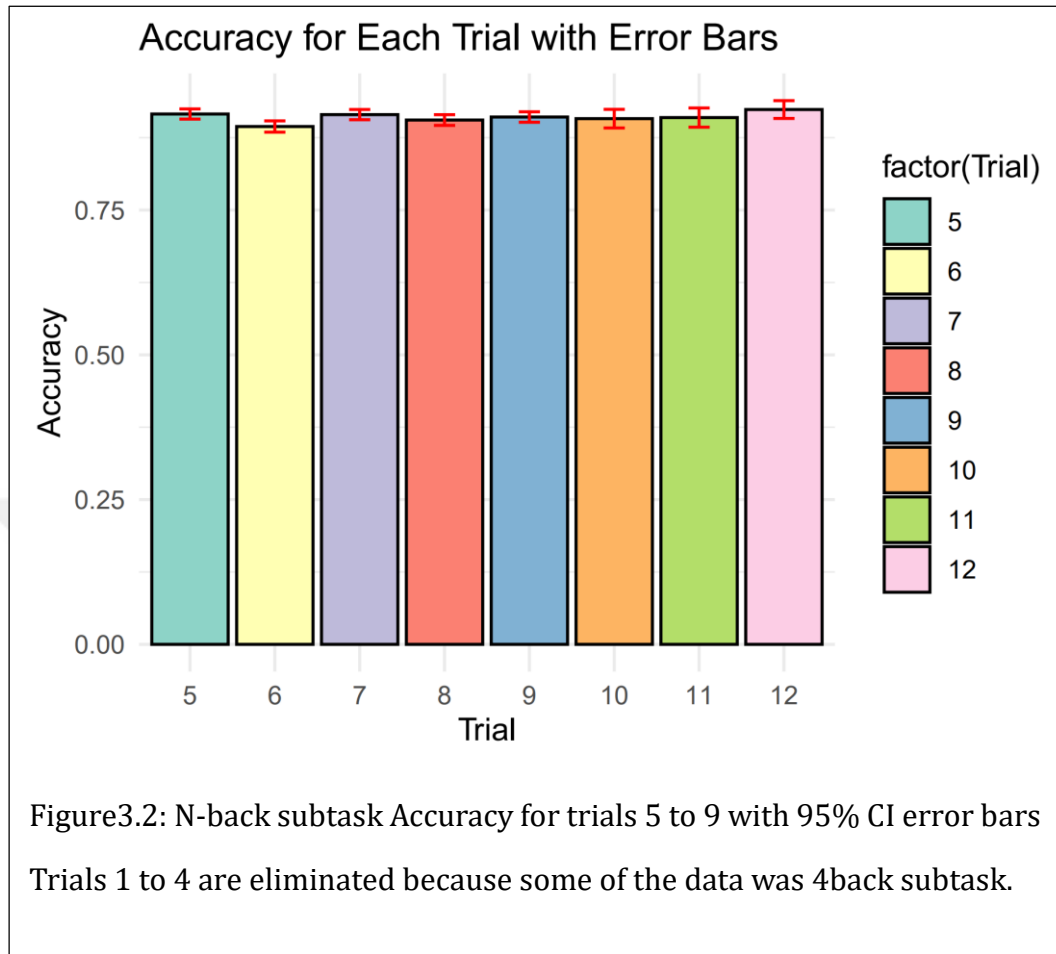
Same as experiment 1, all analyses were conducted on RStudio.

3.2 Findings

3.2.1 Accuracy

3.2.1.1 Accuracy for N-back

The mean of accuracy in different trials was calculated as the sum of correct answers in different episodes divided by the total number of trials of that condition. Separately, for the N-back subtask (Table 3.1) and recall subtask (Table 3.2). Tables 3.1 show the effects of episodes and trials positions on RTs accuracy. RTs accuracy were not different across the three-episode positions ($F_{2,2,85.9} = 0.708$, $p = 0.508$, $\eta^2_p = 0.018$). However, RTs accuracy differed across the trials mean ($F_{3,1,122.7} = 4.728$, $p = 0.003$, $\eta^2_p = 0.108$). The effect of trials was not different across different episode positions ($F_{3,4,133.2} = 0.753$, $p = 0.539$, $\eta^2_p = 0.019$). [see also Figure 3.2 and 3.3]



Within subject effects

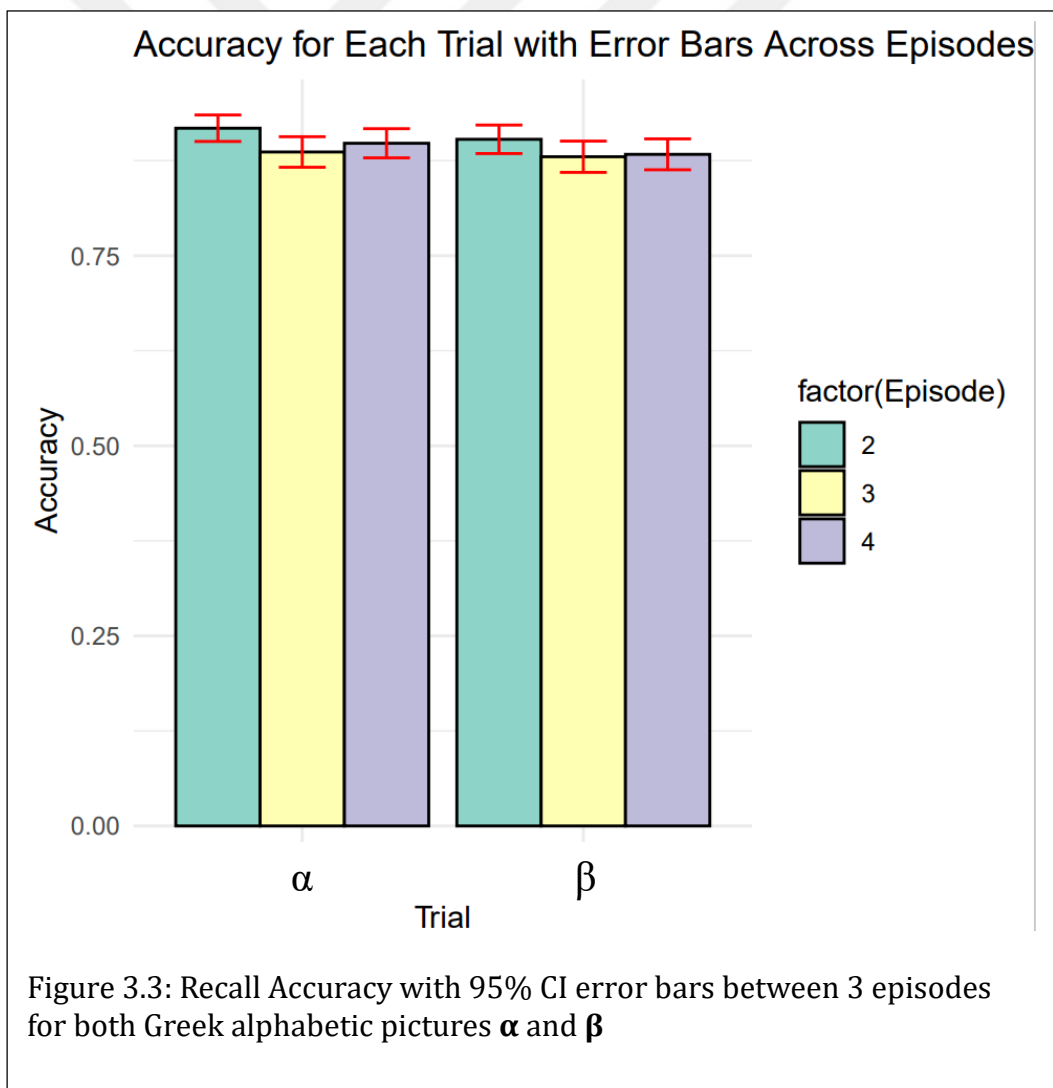
Cases	Sum of squares	df	Mean square	F	p	η^2
Episode	0.056	2.204	0.025	0.708	0.508	0.018
Residuals	3.091	85.974	0.036			
Trial	0.821	3.146	0.261	4.728	0.003	0.108
Residuals	6.768	122.711	0.055			
Episode * Trial	0.221	3.416	0.065	0.753	0.539	0.019
Residuals	11.451	133.222	0.086			

Table3.1: ANOVA of N-back Accuracy for trials 5 to 9

due to 4back subtask ANOVA analysis was done for trials 5 to 9(Greenhouse–Geisser corrected)

3.2.1.2 Accuracy for recall

Tables 3.2 show the effects of trials and episode positions on RTs accuracy for recall pictures. RTs accuracy for recall were not different across the three episode positions ($F_{1,7,66.9} = 0.256$, $p = 0.741$, $\eta^2_p = 0.007$). However, RTs accuracy were different across the trials mean ($F_{1,39} = 31.489$, $p < 0.001$, $\eta^2_p = 0.447$). The effect of trials was not different across different episode positions ($F_{1,7,68.4} = 2.032$, $p = 0.144$, $\eta^2_p = 0.05$). [see also Figure 3.3]

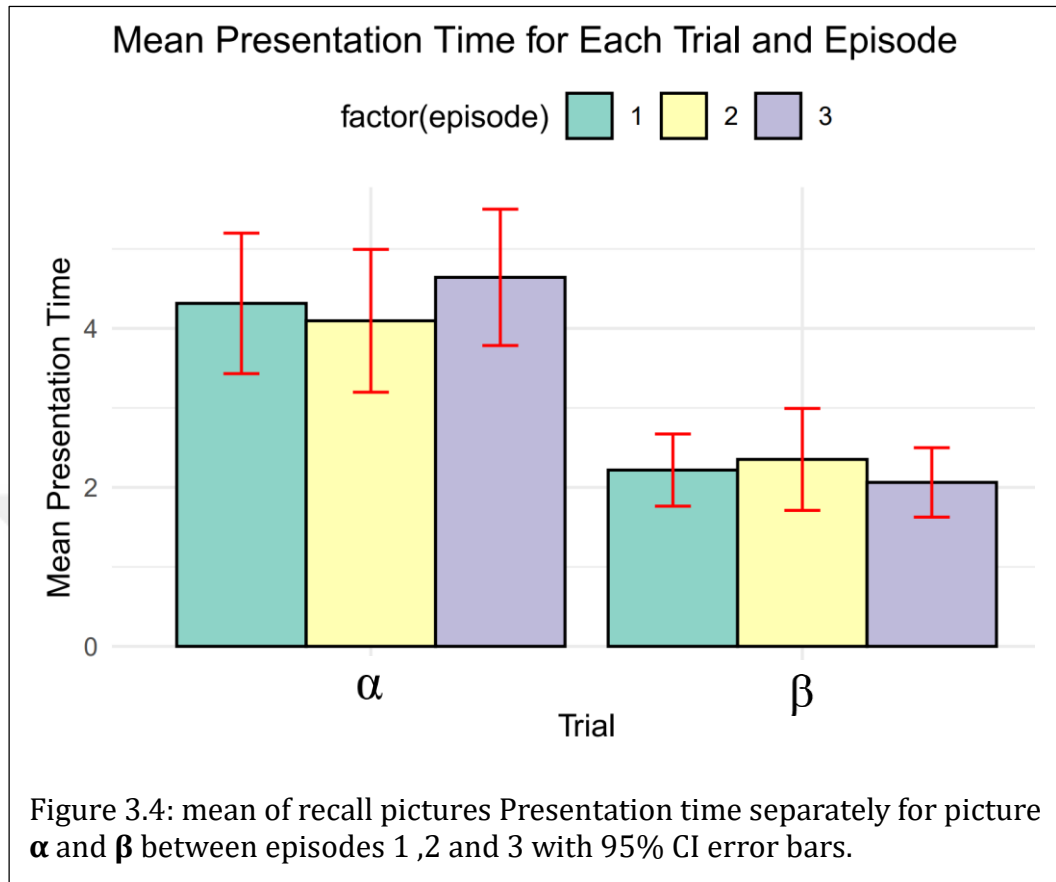


Within subject effects						
Cases	Sum of squares	df	Mean square	F	p	η^2_p
Episode	0.687	1.716	0.4	0.256	0.741	0.007
Residuals	104.6	66.939	1.563			
Trial	274.831	1	274.831	31.489	< 0.001	0.447
Residuals	340.382	39	8.728			
Episode *	7.046	1.754	4.017	2.032	0.144	0.05
Trial						
Residuals	135.231	68.410	1.977			

Table 3.2: ANOVA of Recall accuracy for picture α and β between 3 episodes (Greenhouse–Geisser corrected)

3.2.2 Presentation time

The presentation time was calculated as the mean amount of the time that participants took to memorize the pictures shown at the beginning of episodes 1,2,3 and followed by n-back subtasks and were asked to be recalled at the end of episode 2,3,4, [the time between the episodes started and participant made the button response to move onto the following picture], in seconds (Figure 3.4). Table 3.4 shows the effects of episodes and trials on presentation RTs. Presentation RTs were different across the episode mean ($F_{1.2,49.2} = 19.34$, $p < 0.001$, $\eta^2_p = 0.332$). Moreover, presentation RTs differed across the trial positions ($F_{1,39} = 25.896$, $p < 0.001$, $\eta^2_p = 0.399$). The effect of trials was different across different episode positions ($F_{1.7,69.2} = 16.158$, $p < 0.001$, $\eta^2_p = 0.293$). [see also Figure 3.4]



Within subject effects

Cases	Sum of squares	df	Mean square	F	p	η^2_p
Episode	148.130	1.263	117.25	19.34	< 0.001	0.332
Residuals	298.718	49.271	6.063			
Trial	47.35	1	47.35	25.896	< 0.001	0.399
Residuals	71.309	39	1.828			
Episode * Trial	87.084	1.774	49.076	16.158	< 0.001	0.293
Residuals	210.186	69.204	3.037			

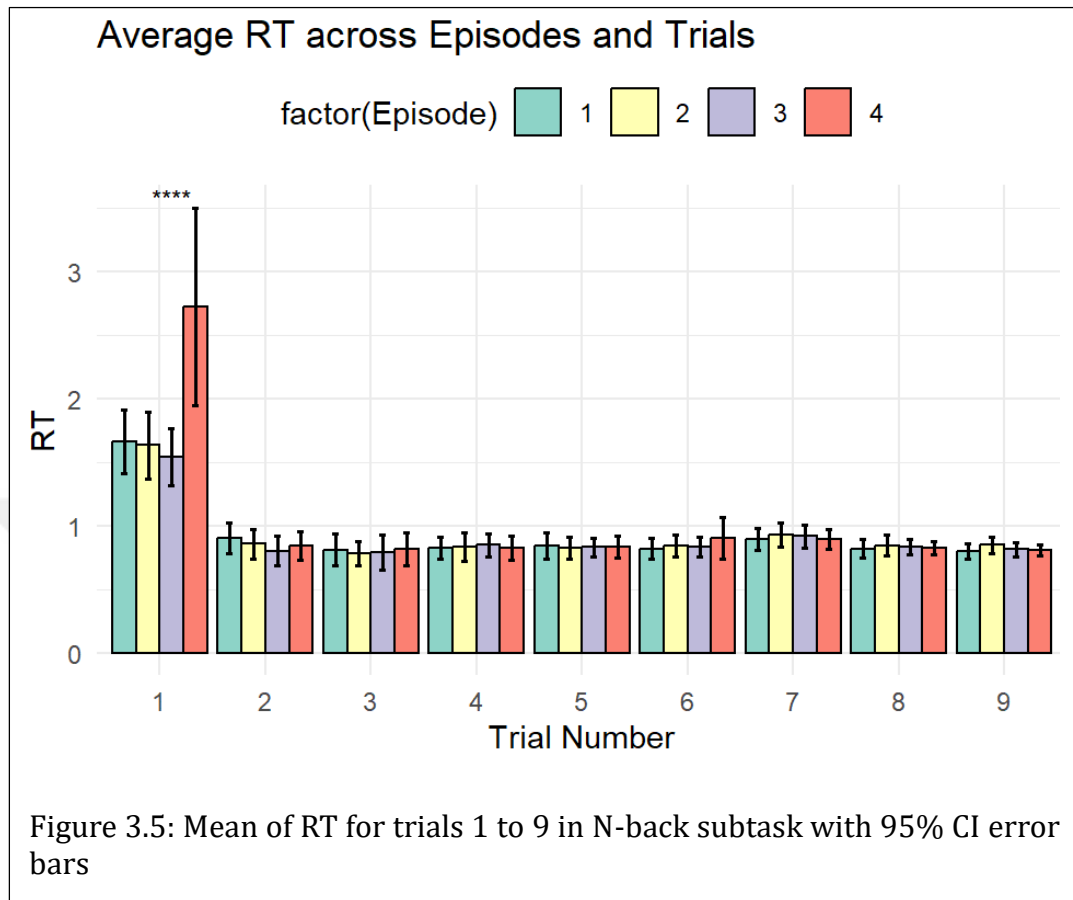
Table 3.4: ANOVA of presentation time for picture α and β between episode 1,2 and 3(Greenhouse–Geisser corrected)

3.2.3 Reaction time

3.2.3.1 Reaction time for N-back

The mean of reaction time was calculated as the mean of the time for participants to respond to each picture in the N-back subtask [9 trials] in seconds. Table 3.5 shows the effects of trials and episodes on reaction time for the N-back subtask. RTs were not different across episodes ($F_{2,6,104} = 0.35$, $p = 0.759$, $\eta^2_p = 0.009$). Also, RTs were not different across trials ($F_{2,7,108.3} = 2.004$, $p = 0.122$, $\eta^2_p = 0.049$). Moreover, the effect of episodes was not different across different trials positions, as well ($F_{6,236.1} = 1.015$, $p = 0.417$, $\eta^2_p = 0.025$). [see also Figure 3.5]

Due to the subtask changing from recall to N-back and updating information of WM, the first trial reaction time is much bigger in all episodes compared to other trials in the same episodes [it is explained in detail in “3.3 Discussion”]. For the ANOVA part, 1st trial is also eliminated, and ANOVA analysis is done among trials 2 to 9.



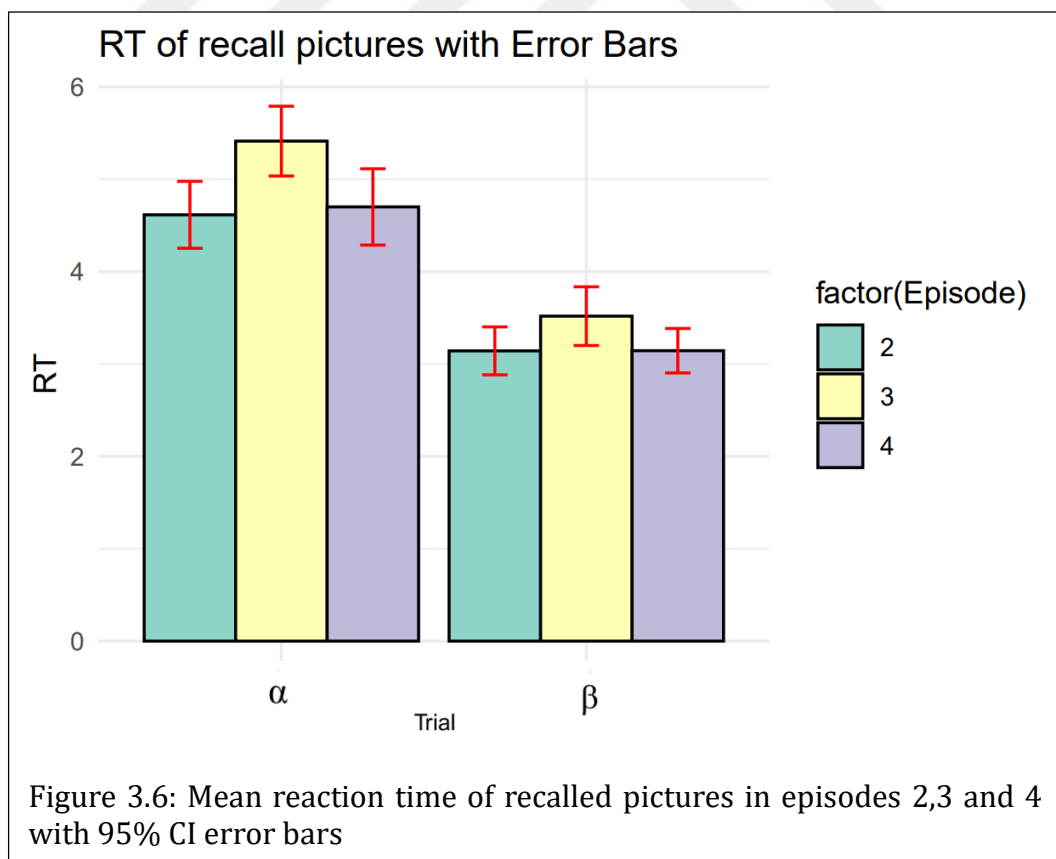
Within subject effects

Cases	Sum of squares	df	Mean square	F	p	η^2_p
Episode	0.025	2.675	0.009	0.359	0.759	0.009
Residuals	2.66	104.333	0.025			
Trial	1.151	2.777	0.414	2.004	0.122	0.049
Residuals	22.395	108.307	0.207			
Episode * Trial	0.496	6.055	0.082	1.015	0.417	0.025
Residuals	19.055	236.159	0.081			

Table 3.5: ANOVA of RT mean in N-back subtask for trials 2 to 9 between 4 episodes, also trial 1 is eliminated because of updating info of working memory (Greenhouse–Geisser corrected)

3.2.3.2 Reaction time for recall

The mean of reaction time was calculated as the mean of the time for participants to recall the memorized pictures at the beginning of episodes at the end of episodes 2,3,4 in seconds. Table 3.6 shows the effects of trials and episodes on reaction time for recall pictures. RTs were different across episodes ($F_{2,78} = 19.994$, $p < 0.001$, $\eta^2_p = 0.339$). Also, RTs were different across trials ($F_{1,39} = 156.203$, $p < 0.001$, $\eta^2_p = 0.8$). The effect of trials was different across different episode positions, as well ($F_{2,78} = 3.591$, $p = 0.032$, $\eta^2_p = 0.084$). [see also Figure 3.6]



Within subject effects

Cases	Sum of squares	df	Mean square	F	p	η^2_p
Episode	17.097	2	8.549	19.994	< 0.001	0.339
Residuals	33.349	78	0.428			
Trial	161.612	1	161.612	156.203	< 0.001	0.8
Residuals	40.350	39	1.035			
Episode *	1.996	2	0.998	3.591	0.032	0.084
Trial						
Residuals	21.674	78	0.278			

Table 3.6: ANOVA of mean of RT for recall picture α and β between episodes 2,3 and 4.

3.3 Discussion

In Experiment 2, we aimed to further explore the dynamics of WMC and task-specific storage by addressing potential criticisms from Experiment 1. The primary goal was to determine whether the observed independence between subtasks A and B could be attributed to passive maintenance of information.

The results from Experiment 2 largely align with those from Experiment 1, confirming that working memory can maintain more information than traditionally predicted by capacity limits. This finding suggests a task-specific or subtask-specific storage mechanism in WM. Specifically, participants demonstrated a significant ability to recall items from subtask B despite

intervening tasks, indicating robust separation between the memory representations of the two subtasks.

Reaction times, as depicted in the graphs(Figure3.5), showed a slight increase during more demanding tasks (Episodes 2 and3), reflecting the cognitive load required to manage multiple subtasks simultaneously, which is as we expect.

A notable observation in the results is the peak highlighted in Figure 3.5. This peak likely reflects the updating processes within WM. When participants were required to recall items from subtask B after an intervening task, the peak performance indicates effective updating and maintenance of these items within WM. This updating peak aligns with the hypothesis that WM operates with task-specific stores, allowing efficient retrieval and maintenance of information pertinent to ongoing tasks.

It's important to note that Trial 1 was excluded from the ANOVA analyses. The rationale behind this exclusion is based on the recognition that the first trial may serve as a baseline or practice trial, which might not accurately reflect participants' typical performance levels. By excluding Trial 1, we aimed to ensure a more accurate and reliable analysis of the subsequent trials, which better represent the participants' performance under the experimental conditions.

The findings from Experiment 2 have significant implications for our understanding of WM capacity. They suggest that WM is not a single, homogenous store but rather a flexible system capable of maintaining and updating multiple streams of information independently. This could mean that capacity limits observed in traditional WM tasks may be more reflective of task-specific constraints rather than inherent limitations of WM itself.

The findings from Experiment 2 offer significant theoretical implications for our understanding of working memory (WM) and task interference. Traditional models of WM often emphasize its limited capacity, suggesting that WM can only hold a small amount of information at any given time. However,

the results from this experiment challenge this view by demonstrating that WM can dynamically update and maintain separate streams of information with minimal interference.

The clear separation between the memory representations of subtasks A and B, even in the presence of intervening tasks, suggests that WM operates with task-specific storage mechanisms. This finding aligns with theories proposing that WM consists of multiple, independent stores that can manage different types of information simultaneously.

The peak observed in Figure 3.5 reflects the dynamic updating capability of WM. This ability to refresh and prioritize relevant information supports models that view WM as a flexible system capable of adapting to changing task demands. This flexibility is crucial for cognitive functions that require rapid switching between tasks, such as in multitasking environments.

The minimal interference observed between subtasks indicates that WM can effectively isolate and manage different tasks. This supports the idea that interference in WM is not a result of inherent capacity limitations but rather a function of task-specific demands and the efficiency of cognitive control mechanisms.

The findings from Experiment 2 can be compared with several key studies in the WM literature to highlight their novel contributions.

The multi-component model proposed in 1974 [15] suggests that WM consists of separate components (e.g., phonological loop, visuospatial sketchpad) that manage different types of information. Our findings extend this model by providing empirical evidence for task-specific stores within these components, demonstrating that WM can handle multiple streams of information related to different tasks without significant interference.

Cowan's model which presented in 1999 [25] posits that WM capacity is limited by the focus of attention, which can hold only a few items at a time. The results from Experiment 2 suggest that while the focus of attention may be

limited, WM can dynamically update and maintain task-specific information outside this narrow focus, supporting the idea of flexible allocation of cognitive resources.

Studies using dual-task paradigms often report significant interference effects, supporting the notion of limited WM capacity [69]. However, our findings suggest that when tasks are designed to minimize overlap in cognitive demands, WM can manage them with minimal interference. This highlights the importance of task design in understanding WM capacity and interference.

In conclusion, the findings from Experiment 2 significantly advance our understanding of WM by demonstrating its capacity for task-specific storage and dynamic updating. These results challenge traditional views of WM as a single, limited-capacity system and provide new insights into the mechanisms underlying task interference. This study contributes to the broader literature by highlighting the flexibility and efficiency of WM in managing multiple cognitive demands.

CHAPTER 4

General Discussion

Our findings in both experiments have important implications for our understanding of WM capacity and control. The classical dual-task paradigm, commonly used in WM capacity research, typically focuses on the impact of distractor subtasks on WM subtasks. However, our experimental paradigm allowed us to explore the mutual effects of concurrent subtasks on one another, highlighting the hierarchical organization of WM function.

Furthermore, our experimental paradigm deviates from the classical dual-task paradigm in several vital aspects. Firstly, unlike traditional dual-task studies where the distractor subtask is often an attentional subtask, our paradigm involved two WM subtasks of the exact nature. Secondly, our paradigm used visual stimuli with a continuous stimulus and response space, allowing for a more nuanced understanding of WM capacity beyond the number of items remembered.

In conclusion, our study provides novel insights into the hierarchical implementation of WM operations and WMC. By demonstrating that increasing the WM load of one subtask does not impact the performance of the other subtask, we contribute to a better understanding of how WM maintenance operates in complex subtask scenarios.

4.1 Implications

1- *The WM information related to hierarchically organized tasks may be stored separately within their respective task-related programs.* In Experiment 1 and Experiment 2, we manipulated the WM loads of identical WM subtasks, interleaving them. Interestingly, the performance in each subtask was only affected by the WM load of its corresponding subtask, not by the WM load of concurrently executed subtasks. This indicates that within-subtask load

effects were present, but across-subtask load effects were not. These findings suggest that WM information for a specific subtask can be stored separately within its own subtask-related program, independent of other concurrently held subtask-related programs.

2- Storing WM information of different task-related programs separately in their respective WM systems could offer a novel approach to bypassing WM capacity limits. If WM information for distinct subtasks is stored in their corresponding subtask-related programs, this could present an alternative method for bypassing WM limits, which are commonly cited as around four items for visual WM [28]. By maintaining WM information as components of separate programs, similar to chunking a set of information, it may be possible to exceed these item limits.

The main inspiration behind this thesis stems from the observation that we retain more information in WM in everyday situations than any traditional views of WM capacity would suggest. We can simultaneously handle multiple intentional subtasks, manage, and process significant amounts of information for them, and effectively control seemingly hierarchical behaviors associated with these subtasks. The findings of this thesis suggest that WM capacity, as portrayed, might demonstrate how this control over WM information, which should exceed the limits proposed by existing literature, is achieved. Information about different intentional subtask-related programs is stored separately as part of their respective programs, surpassing WM capacity limits beyond what is typically proposed by existing literature. This thesis proposes a relationship between WM capacity and intentional subtask entities' hierarchical control of WM.

3- Attentional control may be influenced by hierarchical control and goal-directed task-episode-related programs. Subtasks and goals are not merely representations of declarative information about behavior or cognition; they are cognitive entities that guide cognitive processes until completion. We refer

to these goal-directed cognitive entities as "subtask episode-related programs" [61] or "subtask-related programs," and any extended sequence of behavior is encompassed by these programs. In the context of the current experiments, an extended subtask episode consists of subtasks A and B with their unique goal states. Therefore, their respective subtask-related programs encompass the sequential behavior, and even attentional control processes are likely mediated by these goal-directed entities within a given subtask episode.

Viewed through this lens, attentional control-related functions, such as selecting WM information from all available perceptual and long-term memory (LTM) information (e.g., recalling pictures labeled: α and β in the correct order at the end of an episode), along with the consolidation and retrieval of relevant information from LTM into WM, as well as the removal of information that is no longer goal-relevant, are governed by these subtask-related programs. Once established and initiated, these programs regulate attentional processes and hierarchical behavior systematically and sequentially.

4- Heightened difficulty levels among subtasks could result in more significant interference between the WMC of information linked to distinct yet simultaneous subtask-related programs. Introducing increased difficulty between the demands of concurrent subtasks resulted in eliminating interference-free execution of concurrent but separate subtasks in experiments, particularly when both subtasks remained crucial to achieving higher episode-related goals. Achieving interference-free maintenance of concurrently organized, hierarchical subtask-related programs might necessitate some degree of dissimilarity either in the subtask entities that govern episode execution or in the subtask information itself.

Several findings in the current thesis suggest that the WM items related to different subtasks were maintained as parts of different non-interfering programs, and increasing the WM load of one did not impair performance on the other subtasks.

The first finding is that in Experiment 1, the mean of RT was lower for the high-load subtask (3 pictures in the beginning) in both one-back and three-back conditions compared to the low-load subtask (1 picture in the beginning in the same N-back condition). Also, RTs accuracy were different across both conditions. One implication of this discovery is that there was a level of attentional selection within the subtask-program, leading to competition between actively maintained subtasks for retrieval to the attentional focus. However, once a subtask has been selectively accessed, its content is retrieved without interference from concurrent loads, as attentional processes are now focused solely on the WM operations within the accessed WM store.

In Experiment 2, an interesting observation was made: even after completing one of the subtasks, which rendered it no longer necessary for the larger subtask-episode, there was no evidence of interference in WMC across subtasks. This suggests that completing the larger subtask-episode did not impact the independence between the subtasks in WM.

These findings collectively suggest that the subtask-episode's hierarchical structure might influence WM information management by attentional control. This management appears to occur at different levels of the subtask-episode-related program. When multiple WMs are actively maintained, the process involves accessing the relevant subtask-related program and its corresponding WM, followed by the retrieval of the target within the contents of the accessed WM. Therefore, the selection of WM for further processing through attentional control could be guided by the subtask entities and the levels or subtasks they comprise.

4.2 Elaboration On Discussion

To the best of our knowledge, this thesis represents the initial exploration of the connection between WMC, and the control exerted by hierarchical subtask entities. Moreover, the discovery of the interference-free execution of simultaneously performed subtasks constitutes a unique contribution to the literature on WMC.

The novel extended subtask paradigm introduced in this thesis offers a fresh approach to investigating WMC related issues. Unlike the widely used dual-task or complex-span WM measures, this paradigm features two identical interleaved subtasks, both of which are WM subtasks. Additionally, WMC investigations are conducted using a continuous stimulus and response space. This subtask presents a promising alternative to dual-task paradigms for studying WMC and its control by hierarchically organized subtask entities.

In dual-task studies, when the distractor subtask negatively impacts the performance of the withheld subtask, it is typically interpreted as evidence that processing and storage share a common resource [69]. However, some previous studies have reported instances where such a distractor or concurrent WM demand did not affect the performance of the withheld subtask. It's important to note that these studies focused solely on a unidirectional effect, specifically the impact of the withheld subtask on the performance of the distractor subtask. Additionally, these studies utilized only one WM subtask, with the second subtask often being a non-WM subtask, such as mental arithmetic, in one study. Furthermore, the WM subtasks in these studies primarily involved simple maintenance and recall of information without the need for manipulation based on subtask demands. Therefore, the findings of the current thesis represent distinct and valuable contributions to this body of research[12, 29, 32].

The sustained neural activity observed in the prefrontal cortex (PFC) during delay periods in WM subtasks has traditionally been seen as indicative of WM and the maintenance of WM representations. However, recent suggestions propose that the ability to decode WM information from neuroimaging data may depend on whether the information is considered high-priority or functionally goal-relevant at a particular moment[70]. It is hypothesized that WM representations that are currently and temporarily irrelevant to the goal may enter a latent state characterized by baseline activity, indicating that attention has been withdrawn from them. This storage mechanism, involving

functionally latent states, is termed "activity silent WM" [71], and short-term synaptic plasticity is proposed as a mechanism for such storage. Studies indicate that WM items that are low-priority or temporarily goal-irrelevant may become behaviorally silent, having no impact on attention-based ongoing processing [72]. For instance, in some studies, high-priority WM items affected visual search, whereas low-priority or currently goal-irrelevant items (potentially activity silent) did not [72, 73]. In the current thesis, WM information encoded about a subtask-related program may adopt such a latent state. During periods when it remains goal-irrelevant, it might not affect WMC through a mechanism similar to activity-silent WM. However, how activity-silent WM information is encoded to achieve this remains to be seen, and whether the current paradigm facilitates such encoding requires further investigation.

4.3 Limitations and Future Directions

These findings open new avenues for future research to explore the boundaries of non-interfering stores in WM. Subsequent studies could investigate whether this compartmentalization extends to more complex and varied task combinations or if there are limits to the independence of WM stores. Additionally, neuroimaging studies could examine the neural correlates of this phenomenon, providing further evidence for the distinct storage and management of different WM tasks.

By situating our findings within the broader context of WM and task interference research, we demonstrate their relevance and contribution to existing theories. This comparative analysis not only validates our results but also highlights their novel implications, advancing the understanding of how WM operates under varying load conditions.

Future studies should explore the boundaries of WM's task-specific storage capabilities. Varying the types of intervening tasks or the nature of the

information to be recalled could provide deeper insights into the mechanisms underlying WM's flexibility. Additionally, employing neuroimaging techniques could help identify the neural substrates associated with these processes, offering a more comprehensive understanding of how WM supports complex cognitive functions.



CHAPTER 5

BIBLIOGRAPHY

1. Velichkovsky, B.B., *Consciousness and working memory: Current trends and research perspectives*. Consciousness and cognition, 2017. **55**: p. 35-45.
2. Postle, B.R. and K. Oberauer, *Working memory*. 2022, Oxford Univ. Press. In press Oxford, UK.
3. Adams, E.J., A.T. Nguyen, and N. Cowan, *Theories of working memory: Differences in definition, degree of modularity, role of attention, and purpose*. Language, speech, and hearing services in schools, 2018. **49**(3): p. 340-355.
4. Miller, G.A., E. Galanter, and K.H. Pribram, *Plans and the Structure of Behavior*. New York: Henry Holt and Company. 1960, Inc.
5. Engle, R.W., et al., *Working memory, short-term memory, and general fluid intelligence: a latent-variable approach*. Journal of experimental psychology: General, 1999. **128**(3): p. 309.
6. Lee, J. and S. Park, *Working memory impairments in schizophrenia: a meta-analysis*. Journal of abnormal psychology, 2005. **114**(4): p. 599.
7. Lenartowicz, A., et al., *Alpha modulation during working memory encoding predicts neurocognitive impairment in ADHD*. Journal of Child Psychology and Psychiatry, 2019. **60**(8): p. 917-926.
8. Huntley, J. and R. Howard, *Working memory in early Alzheimer's disease: a neuropsychological review*. International Journal of Geriatric Psychiatry: A journal of the psychiatry of late life and allied sciences, 2010. **25**(2): p. 121-132.
9. Costa, A., C. Caltagirone, and G.A. Carlesimo, *Prospective memory functioning in individuals with Parkinson's disease: a*

- systematic review. Prospective Memory in Clinical Populations, 2020: p. 197-219.
10. Channon, S., J.E. Baker, and M.M. Robertson, *Working memory in clinical depression: an experimental study*. Psychological medicine, 1993. **23**(1): p. 87-91.
 11. Baddeley, A.D., *Short-term and working memory*. The Oxford handbook of memory, 2000. **4**: p. 77-92.
 12. Oberauer, K., *Access to information in working memory: exploring the focus of attention*. Journal of Experimental Psychology: Learning, Memory, and Cognition, 2002. **28**(3): p. 411.
 13. Morrison, A.B., A.R. Conway, and J.M. Chein, *Primacy and recency effects as indices of the focus of attention*. Frontiers in human neuroscience, 2014. **8**: p. 6.
 14. Baddeley, A., *Working memory: Theories, models, and controversies*. Annual review of psychology, 2012. **63**: p. 1-29.
 15. Baddeley, A.D. and G.J. Hitch, *Development of working memory: Should the Pascual-Leone and the Baddeley and Hitch models be merged?* Journal of experimental child psychology, 2000. **77**(2): p. 128-137.
 16. Atkinson, R.C. and R.M. Shiffrin, *Human memory: A proposed system and its control processes*, in *Psychology of learning and motivation*. 1968, Elsevier. p. 89-195.
 17. Atkinson, R.C. and R.M. Shiffrin, *The control of short-term memory*. Scientific american, 1971. **225**(2): p. 82-91.
 18. Baddeley, A., *The episodic buffer: a new component of working memory?* Trends in cognitive sciences, 2000. **4**(11): p. 417-423.
 19. Baddeley, A., *Psychology of learning and motivation*. (No Title), 1974. **8**: p. 47.
 20. Cowan, N., *What are the differences between long-term, short-term, and working memory?* Progress in brain research, 2008. **169**: p. 323-338.
 21. Baddeley, A., *Fractionating the central executive*. Principles of frontal lobe function, 2002: p. 246-260.
 22. Baddeley, A.D. and R.H. Logie, *Working memory: The multiple-component model*. 1999.
 23. Cowan, N., *Evolving conceptions of memory storage, selective attention, and their mutual constraints within the human*

- information-processing system*. Psychological bulletin, 1988. **104**(2): p. 163.
24. Norman, D.A., *Toward a theory of memory and attention*. Psychological review, 1968. **75**(6): p. 522.
 25. Cowan, N., *An embedded-processes model of working memory*. Models of working memory: Mechanisms of active maintenance and executive control, 1999. **20**(506): p. 1013-1019.
 26. James, W., *The principles of psychology*. Vol. 1. 2007: Cosimo, Inc.
 27. Cowan, N., *The magical number 4 in short-term memory: A reconsideration of mental storage capacity*. Behavioral and brain sciences, 2001. **24**(1): p. 87-114.
 28. Cowan, N., *Working memory capacity*. 2012: Psychology press.
 29. Klapp, S.T., E.A. Marshburn, and P.T. Lester, *Short-term memory does not involve the "working memory" of information processing: The demise of a common assumption*. Journal of Experimental Psychology: General, 1983. **112**(2): p. 240.
 30. Fürst, A.J. and G.J. Hitch, *Separate roles for executive and phonological components of working memory in mental arithmetic*. Memory & cognition, 2000. **28**: p. 774-782.
 31. DeStefano, D. and J.A. LeFevre, *The role of working memory in mental arithmetic*. European Journal of Cognitive Psychology, 2004. **16**(3): p. 353-386.
 32. Oberauer, K., et al., *Dissociating retention and access in working memory: An age-comparative study of mental arithmetic*. Memory & Cognition, 2001. **29**: p. 18-33.
 33. Garavan, H., *Serial attention within working memory*. Memory & cognition, 1998. **26**: p. 263-276.
 34. Oberauer, K., *The focus of attention in working memory—from metaphors to mechanisms*. Frontiers in human neuroscience, 2013. **7**: p. 673.
 35. Oberauer, K., et al., *What limits working memory capacity?* Psychological bulletin, 2016. **142**(7): p. 758.
 36. Oberauer, K. and R. Kliegl, *Beyond resources: Formal models of complexity effects and age differences in working memory*. European Journal of Cognitive Psychology, 2001. **13**(1-2): p. 187-215.

37. Miller, G.A., *The magical number seven, plus or minus two: Some limits on our capacity for processing information*. Psychological review, 1956. **63**(2): p. 81.
38. Fukuda, K., E. Awh, and E.K. Vogel, *Discrete capacity limits in visual working memory*. Current opinion in neurobiology, 2010. **20**(2): p. 177-182.
39. Ma, W.J., M. Husain, and P.M. Bays, *Changing concepts of working memory*. Nature neuroscience, 2014. **17**(3): p. 347-356.
40. Luck, S.J. and E.K. Vogel, *The capacity of visual working memory for features and conjunctions*. Nature, 1997. **390**(6657): p. 279-281.
41. Wilken, P. and W.J. Ma, *A detection theory account of change detection*. Journal of vision, 2004. **4**(12): p. 11-11.
42. Bays, P.M., R.F. Catalao, and M. Husain, *The precision of visual working memory is set by allocation of a shared resource*. Journal of vision, 2009. **9**(10): p. 7-7.
43. Gorgoraptis, N., et al., *Dynamic updating of working memory resources for visual objects*. Journal of Neuroscience, 2011. **31**(23): p. 8502-8511.
44. Van den Berg, R., E. Awh, and W.J. Ma, *Factorial comparison of working memory models*. Psychological review, 2014. **121**(1): p. 124.
45. Fougner, D., *The relationship between attention and working memory*. New research on short-term memory, 2008. **1**: p. 45.
46. Kane, M.J., et al., *A controlled-attention view of working-memory capacity*. Journal of experimental psychology: General, 2001. **130**(2): p. 169.
47. Vogel, E.K., A.W. McCollough, and M.G. Machizawa, *Neural measures reveal individual differences in controlling access to working memory*. Nature, 2005. **438**(7067): p. 500-503.
48. Camos, V., et al., *What is attentional refreshing in working memory?* Annals of the New York Academy of Sciences, 2018. **1424**(1): p. 19-32.
49. Souza, A.S. and E. Vergauwe, *Unravelling the intersections between consolidation, refreshing, and removal*. Annals of the New York Academy of Sciences, 2018. **1424**(1): p. 5-7.
50. Bayliss, D.M., J. Bogdanovs, and C. Jarrold, *Consolidating working memory: Distinguishing the effects of consolidation, rehearsal and attentional refreshing in a working memory*

- span task*. Journal of Memory and Language, 2015. **81**: p. 34-50.
51. Dudai, Y., A. Karni, and J. Born, *The consolidation and transformation of memory*. Neuron, 2015. **88**(1): p. 20-32.
 52. Lemaire, B., et al., *What is the time course of working memory attentional refreshing?* Psychonomic Bulletin & Review, 2018. **25**: p. 370-385.
 53. Oberauer, K., *Working memory and attention—A conceptual analysis and review*. Journal of cognition, 2019. **2**(1).
 54. Johnson, M.R. and M.K. Johnson, *Toward characterizing the neural correlates of component processes of cognition*. Neuroimaging of human memory: Linking cognitive processes to neural systems, 2009: p. 169-194.
 55. Lewis-Peacock, J.A., Y. Kessler, and K. Oberauer, *The removal of information from working memory*. Annals of the New York Academy of Sciences, 2018. **1424**(1): p. 33-44.
 56. Ricker, T.J., L.R. Spiegel, and N. Cowan, *Time-based loss in visual short-term memory is from trace decay, not temporal distinctiveness*. Journal of Experimental Psychology: Learning, Memory, and Cognition, 2014. **40**(6): p. 1510.
 57. Oberauer, K. and S. Lewandowsky, *Modeling working memory: A computational implementation of the Time-Based Resource-Sharing theory*. Psychonomic bulletin & review, 2011. **18**: p. 10-45.
 58. Oberauer, K., *Removing irrelevant information from working memory: a cognitive aging study with the modified Sternberg task*. Journal of experimental Psychology: learning, Memory, and cognition, 2001. **27**(4): p. 948.
 59. Botvinick, M.M., Y. Niv, and A.G. Barto, *Hierarchically organized behavior and its neural foundations: A reinforcement learning perspective*. cognition, 2009. **113**(3): p. 262-280.
 60. Barto, A.G. and S. Mahadevan, *Recent advances in hierarchical reinforcement learning*. Discrete event dynamic systems, 2003. **13**: p. 341-379.
 61. Farooqui, A.A. and T. Manly, *We do as we construe: extended behavior construed as one task is executed as one cognitive entity*. Psychological Research, 2019. **83**(1): p. 84-103.
 62. Meyer, D.E. and D.E. Kieras, *A computational theory of executive cognitive processes and multiple-task performance:*

- Part I. Basic mechanisms. Psychological review, 1997. 104(1): p. 3.*
63. Schneider, D.W. and G.D. Logan, *Hierarchical control of cognitive processes: switching tasks in sequences*. Journal of Experimental Psychology: General, 2006. **135**(4): p. 623.
 64. Farooqui, A.A. and T. Manly, *When attended and conscious perception deactivates fronto-parietal regions*. Cortex, 2018. **107**: p. 166-179.
 65. Peirce, J., et al., *PsychoPy2: Experiments in behavior made easy*. Behavior research methods, 2019. **51**: p. 195-203.
 66. Singmann, H., et al., *afex: Analysis of Factorial Experiments.[R package]*. Retrieved from <https://cran.r-project.org/package=afex>, 2018.
 67. Alova, C.A.R., *Performance of College of Education Graduates in the Licensure Examination for Teachers: A Descriptive Study*. Online Submission, 2021.
 68. Sweller, J., *Cognitive load during problem solving: Effects on learning*. Cognitive science, 1988. **12**(2): p. 257-285.
 69. Mathy, F., M. Chekaf, and N. Cowan, *Simple and complex working memory tasks allow similar benefits of information compression*. Journal of Cognition, 2018. **1**(1).
 70. Mongillo, G., O. Barak, and M. Tsodyks, *Synaptic theory of working memory*. Science, 2008. **319**(5869): p. 1543-1546.
 71. Lewis-Peacock, J.A., et al., *Neural evidence for a distinction between short-term memory and the focus of attention*. Journal of cognitive neuroscience, 2012. **24**(1): p. 61-79.
 72. Mallett, R. and J.A. Lewis-Peacock, *Behavioral decoding of working memory items inside and outside the focus of attention*. Annals of the New York Academy of Sciences, 2018. **1424**(1): p. 256-267.
 73. van Moorselaar, D., J. Theeuwes, and C.N. Olivers, *In competition for the attentional template: Can multiple items within visual working memory guide attention?* Journal of Experimental Psychology: Human Perception and Performance, 2014. **40**(4): p. 1450.