

Stand Back And Let It All
Be: Cognitive Scaffolds
To Help Academic
Writing In Psychology

STAND BACK AND LET IT ALL BE: COGNITIVE SCAFFOLDS TO HELP ACADEMIC WRITING IN PSYCHOLOGY

MICHAEL R.W. DAWSON



Stand Back And Let It All Be: Cognitive Scaffolds To Help Academic Writing In Psychology Copyright © by Michael R.W. Dawson is licensed under a Creative Commons Attribution-NonCommercial 4.0 International License, except where otherwise noted.

CONTENTS

Introduction	1
--------------	---

Part I. CHAPTER 1: EMBODIED COGNITION AND WRITING

1.1 Writing Is Hard	5
1.2 The Myth of Inspiration	8
1.3 Whole Inspiration	11
1.4 Fragmentary Inspiration	14
1.5 Cognitive Scaffolding	17
1.6 Extending the Mind	21
1.7 Writing With Your World	24
1.8 Index Card Affordances	27
1.9 A Writing Scaffold	33

Part II. CHAPTER 2:

SCAFFOLDING AN OUTLINE

2.1 Should You Outline?	39
2.2 Reverse Engineering An Outline	42
2.3 Meaningful Chains Of Paragraphs	51
2.4 A Desirable Outline	56
2.5 Moving Topics Into Your World	60
2.6 Step 1: Generate Your Broad Topics	63
2.7 Step 2: Organize Your Topics	67
2.8 Step 3: Enhance Your Existing Topics	70
2.9 Step 4: Develop Your Paragraph Topics	73
2.10 Step 5: Organize Your Paragraph Topics	77
2.11 Step 6: Write Your Topic Sentences	80
2.12 Step 7: Write Your Concluding Sentences	87
2.13 Step 8: Notate Your Scaffold Cards	92
2.14 Step 9: Move Your Outline Into A Computer	95
2.15 Index Cards – And Beyond	99

Part III. CHAPTER 3: SCAFFOLDING INITIAL TOPICS

3.1 Facing the Blank Set of Topics	107
3.2 Questions from Conventions	109
3.3 Follow A Model	121
3.4 Use Your Own Work To Scaffold	124
3.5 Let Your Findings Be Your Scaffold	130
3.6 Storyboard with Images	134
3.7 Brainstorming	137
3.8 Multiple Scaffolds For Writing	140

Part IV. CHAPTER 4: SCAFFOLDING DRAFTING AND REVISING

4.1 Your Outline Is A Scaffold	145
4.2 Use Standard Structures To Convert Your Outline To Your First Draft	148
4.3 Polishing Your First Draft	155
4.4 When Should You Stop Revising?	170

4.5 Scaffolding Responses To Reviewer Comments	175
--	-----

Part V. CHAPTER 5: YOUR WRITING WORLD

5.1 Your Writing World	183
5.2 Writing Equipment	185
5.3 Index Cards	188
5.4 Effective Displays Of Index Cards	193
5.5 Scaffolding What To Do Next	202
5.6 Scaffolds For Keeping Track Of Project Progress	206
5.7 Source Cards And Quote Cards	209
5.8 Zettelkasten	212
5.9 Notebooks	215
5.10 The Social World	218
5.11 Stand On The Shoulders Of Others	221
5.12 The Digital World And Its Scaffolds	225

Part VI. CITED LITERATURE

APPENDIX I: A SCAFFOLDING EXAMPLE	251
AI.1 Case Study: Large Language Models	252
AI.2 Can Large Language Models Inform Cognitive Science?	257
AI.3 Initial Topic Cards For The Essay	265
AI.4 Refining Topic Cards For The Essay	272
AI.5 Create Topic Sentences For Each Topic Card In The Essay	278
AI.6 Add Concluding Sentences To The Essay	284
AI.7 Create The First Draft Of The Essay	294
AI.8 Revise And Polish Drafts Of The Essay	307

INTRODUCTION

Can embodied cognitive science help you write?

In the introduction, I explore why you might find academic writing challenging. I argue mentors often hide – and fail to teach – academic writing processes, leaving students to learn how to write on their own. Often students try to write by following incorrect assumptions about academic writing processes. I discuss incorrect assumptions and consider how changing your assumptions might change how you approach writing. Different assumptions about how academics write lead me to consider how ideas from embodied cognition can help academic writing. You do not need to know much about embodied cognition to take advantage of its ideas about writing. The only idea to keep in mind is embodied cognition’s proposal that thinking need not exclusively be ‘inside your head’ but also can be manipulating objects in the world. The book explores how you can use objects in the world to help your writing.

PART I

CHAPTER 1: EMBODIED COGNITION AND WRITING

Use your world to debunk the myth of inspiration.

Why might you find writing challenging? If you believe the myth of inspiration, then you make writing difficult. The myth of inspiration makes writing difficult by presuming writing only goes on inside your mind. When you assume writing requires inspiration, you separate your mind from your world, and you dismiss your world's role in supporting your thinking. However, embodied cognitive scientists propose another view by arguing your mind extends into your world. Your world can replace your internal cognitive processes, and in so doing your world becomes a part of your mind. Chapter 1 considers writing from both perspectives and suggests you can use ideas from embodied cognition to make writing easier.

1.1 WRITING IS HARD

“Writing well is impossibly difficult” – Ernest Hemingway, Paris Review, 1958

If you or I asked successful writers to describe the process of writing, what would we hear? You or I would likely hear writing described as ‘hard’.

Consider some example quotes from famous writers (Jacobs & Hjalmarsson, 1999). Emile Zola complains “Giving birth to a book is always an abominable torture for me.” Thomas Mann declares “A writer is somebody for whom writing is more difficult than it is for other people.” James Joyce asserts “Writing in English is the most ingenious torture ever devised for sins committed in previous lives.” George Orwell warns “Writing a book is a horrible, exhausting struggle, like a long bout of some painful illness.” Writing seems hard – if not worse!

Because writing is challenging, you can find many books which provide advice to improve writing. Some books, called style guides, recommend rules for making writing clear and effective (Bierce & Freeman, 2009; Flaherty, 2009; Gordon, 1984; Messenger & Taylor, 1984; Pinker, 2014; Plotnik, 1982; Quiller-Couch, 1916; Shertzer, 1986; Strunk & White, 1959).

Some books provide conventions for citing references or structuring articles (American Psychological Association, 2020; Modern Language Association of America, 2021; University of Chicago Press., 2017). Different disciplines adopt different conventions. For example, psychologists follow the American Psychological Association's conventions.

Some books focus on practices to improve writing (Baker & Zinsser, 1987; Barzun, 1985; Bradbury, 1990; Cameron, 2022; Clark, 2008b, 2014; Elbow, 1981; Flower, 1989; Gardner & O'Nan, 1994; Germano, 2013, 2021; Hall, 2003; Hawker, 2015; Hoermann-Elliott, 2021; Kidder & Todd, 2013; Koch, 2003; Leith, 2018; Martin & Kroitor, 1979; McPhee, 2017; Pallant & Price, 2015; Rhodes, 1995; Saleses, 2021; Sargent & Paraskevas, 2005; Sawers, 2002; Wheelan, 2022; Williams & Bizup, 2017; Zinsser, 2006). I find such books particularly helpful; I constantly return to my favorites.

Some books appeal to cognitive science or to neuroscience for ideas about improving writing (Cron, 2012, 2016, 2021; Janzer, 2016; Prentiss & Walker, 2020; Storr, 2020). Such books claim better writing requires authors to understand the readers' mental operations.

In some books accomplished authors write about their craft to inspire struggling writers (Becker, 2020; Block, 2021; Brande, 1934; Dillard, 1989; Goldberg, 1986; Heighton, 2011; Hemingway, 1964; King, 2000; Lamott, 1995; Marche, 2023;

Orwell, 2005; Ueland, 1938; Wiebe, 2016). Such books often combine writing advice with memoirs about the writing life.

My book aims to improve academic writing. Academic writing must be particularly hard, because many books are written to guide specific writers: university researchers (Becker, 2020; Carpenter, 2020; Greene, 2013; Heard, 2022; Kail, 2019; Kumar, 2020; Rosnow & Rosnow, 1998; Sarnecka, 2019; Schimel, 2012; Silvia, 2007; Sword, 2012, 2016, 2017).

Why is academic writing hard? Academics rarely receive explicit training about how to write (Kail, 2019; Sarnecka, 2019; Silvia, 2007). As a result, academics rely upon conventional assumptions about the writing process. I believe conventional assumptions make academic writing *more* difficult. The current chapter explores some conventional ideas, describes why they make writing more difficult, and introduces alternative ideas which will help make academic writing easier.

My alternative ideas emerge from embodied cognition. A core theme unites these ideas: *cognitive scaffolding*. When you use a cognitive scaffold, you move your thinking from inside your mind to outside in your world. My book describes how cognitive scaffolds make writing easier; when I use writing scaffolds, I often feel I complete most of my work – by using my world – before I start ‘writing proper’. My book describes writing scaffolds to offer you a similar experience.

1.2 THE MYTH OF INSPIRATION

You make writing hard when you accept the myth of inspiration.

What makes writing hard? Most writers claim facing the blank page provides their greatest challenge. Neil Gaiman notes “Being a writer is a very peculiar sort of a job: it’s always you versus a blank sheet of paper (or a blank screen) and quite often the blank piece of paper wins.” Peter Elbow (1981, p. 14) wants to help “with the root psychological or existential difficulty in writing: finding words in your head and putting them down on a blank piece of paper.” Annie Dillard (1989, p. 59) claims the blank page teaches us to write, “the page in the purity of its possibilities; the page of your death, against which you pit such flawed excellences as you can muster with all your life’s strength.”

Many writers believe facing the blank page provides writing’s greatest challenge because they accept the *myth of inspiration* (Becker, 2020; Flower, 1989; Kidder & Todd, 2013). According to the myth of inspiration, writing proceeds as follows: first, inspiration provides ideas – complete sentences – to your mind; second, you move the inspired ideas from your mind to your world by writing them down. When

you accept the myth of inspiration, you believe inspiration must provide ideas *before* you write.

The myth of inspiration appears in many anecdotes about creative thinking or writing (Poincare, 1913). Poet A. E. Housman describes being inspired while walking: “As I went along, thinking of nothing in particular, only looking at things around me and following the progress of the season, there would flow into my mind, with sudden and unaccountable emotion, sometimes a line or two of verse, sometimes a whole stanza at once” (Housman, 1933, p. 46). When Housman returned home, he wrote his inspired verses down.

Housman’s anecdote reveals another property of the myth of inspiration: inspired ideas rise into consciousness unintentionally or inexplicably. “At the moment when I put my foot on the step the idea came to me, without anything in my former thoughts seeming to have paved the way for it” (Poincare, 1913, p. 388).

When you accept the myth of inspiration, you make the blank page challenging, because you *expect* inspiration to fill your mind with ideas to write about. Without inspiration, writing seems difficult or impossible. Perhaps inspiration distinguishes great writers from others: lesser writers lack inspiration and fail to meet the blank page’s challenge.

The myth of inspiration underlies much writing advice (Elbow, 1981; Goldberg, 1986, 2021; Janzer, 2016; Lamott, 1995). Some advice proposes you inhibit sentence creation – your inspiration — when you immediately evaluate

the sentences which come into your mind. Techniques like Elbow's freewriting encourage writers to increase creative flow – to enhance inspiration — by delaying criticism. Other writing aids provide prompts to spark your inspiration (Goldberg, 2021).

To me, a cognitive scientist, the myth of inspiration relates to a cognitive theory called the *disembodied mind* (Dawson, 2013). According to the disembodied mind, thought occurs within a mind separated from both your world and your body. I associate inspiration with the disembodied mind because inspiration occurs in the mind, not the world.

However, cognitive scientists explore other theories about the mind. Embodied cognitive scientists move thinking from inside the mind to outside in the world, a view called the *embodied mind* or *embodied cognition* (Clark, 1997; Dawson, 2013; Dawson et al., 2010; Shapiro, 2019). If you assume the embodied mind, then how might your assumptions about writing change? When your assumptions about writing change, how might you alter your methods for improving writing?

1.3 WHOLE INSPIRATION

Does inspiration provide complete sentences, paragraphs, or manuscripts?

Many writers believe writing requires us to compose complete sentences. Ernest Hemingway defeated writer's block by telling himself "All you have to do is write one true sentence. Write the truest sentence that you know" (Hemingway, 1964, p. 12). Natalie Goldberg notes "We think in sentences, and the way we think is the way we see" (Goldberg, 2016, p. 62). Many notetaking methods also encourage writing full sentences (Ahrens, 2022; Howard & Barton, 1986).

Writers who accept the myth of inspiration also think inspiration provides complete sentences – and possibly more. A famous example comes from Samuel Taylor Coleridge's account of writing his 1816 poem *Kubla Khan* (Lowes, 1927). Staying at a lonely farmhouse, ill, and after taking opium, Coleridge fell into a profound sleep in his chair while reading. While asleep, inspiration provided Coleridge with two or three hundred lines of poetry. "All the images rose up before [me] as things, with a parallel production of the correspondent expressions, without any sensation or consciousness of effort" (Lowes, 1927, p. 370). Awakening from his sleep, Coleridge experienced "a distinct recollection of the whole" and started to

transcribe his inspiration onto paper. The poem consists of only 54 lines, and not 300, because a visitor interrupted Coleridge during his transcribing. Coleridge's inspiration "had passed away like the images on the surface of a stream into which a stone has been cast" (Lowes, 1927, p. 370) when the poet returned to writing his ideas down.

You can easily find many other examples of writers thinking inspiration provides complete sentences. Stephen Koch's students "believed that *real* writers – the happy few – concocted stories through some magic process denied to lesser mortals like themselves. Stories came to 'real writers' – that mystery elite – complete, perfect, and intricate from the moment of inspiration onward" (Koch, 2003, p. 57).

Anne Lamott offers a similar account: "I sit down in the morning and reread the work I did the day before. And then I wool-gather, staring at the blank page or off into space. I imagine my characters and let myself daydream about them. A movie begins to play in my head, a motion pulsing underneath it, and I stare at it in a trancelike state, until words bounce around together form a sentence. Then I do the menial work of getting it down on paper, because I'm the designated typist" (Lamott, 1995, p. 54).

Dorothea Brande calls inspiration "releasing genius" (Brande, 1934). She describes an author who reports "after mulling an idea over till his head aches he comes to a kind of dead end; he can no longer think about his story or even understand why it once appealed to him" (Brande, 1934, p. 158).

But then inspiration strikes, releasing genius: “Much later, when he is least expecting it, the idea returns, mysteriously rounded and completed, ready for transcribing.” For Brande, releasing genius produces complete sentences.

Unfortunately, if writing requires inspiration, and inspiration delivers complete sentences, then you cannot learn how to write (Salesses, 2021). If writing requires inspiration, then writers are born, not made. Fortunately, inspiration may take different forms; some do not produce complete sentences. Alternative ideas about inspiration lead to writing practices which you *can* learn.

1.4 FRAGMENTARY INSPIRATION

Don't think in sentences.

According to the myth of inspiration, ideas inexplicably rise into consciousness (Poincare, 1913). I associate inspiration with the disembodied mind because inspiration exists in the mind before being brought into the world. For example, Mozart carried his pieces around in his head for days before setting them down on paper (Hildesheimer, 1983). A famous Victor Hugo quote captures the disembodied mind: “A writer is a world trapped inside a person.” The trapped world begins within the writer’s disembodied mind before being freed by being written into the world.

The disembodied mind belongs to Cartesian philosophy (Descartes, 1641/1996). Cartesian philosophy inspires traditional cognitive science, which argues cognition proceeds as a *sense-think-act cycle* (Dawson, 2013). In sense-think-act processing, you ‘seed’ cognition by sensing information in your world. Next, thinking manipulates information, building representations you use to understand your world, and to plan possible actions. Finally, you act on your world by following a plan provided by your thinking. The sense-think-act cycle occurs in a disembodied mind because thinking occurs sep-

arately from both sensing and acting. Thinking *must* occur before acting; thinking provides the only means for senses to (indirectly) inform actions.

The myth of inspiration agrees nicely with sense-think-act processing. When you accept the myth, you believe writing begins in your mind (via inspiration); only after ideas arise in your mind can you act on your world by writing your ideas down. Accordingly, “in writing classes, if nowhere else, it is entirely permissible to spend large chunks of your time off in your own little dreamworld” (King, 2000, p. 235). Anne Lamott’s writing process (Section 1.3) conforms to sense-think-act processing as well: she reads her previous work (sensing); inspiration creates a movie in her head (thinking); later she types her inspired ideas (acting).

However, a writer’s thinking may not provide complete sentences. “One does not usually think in full sentences or in single words, but in clumps of half-formed ideas that correspond very imperfectly to one’s intention” (Barzun, 1975, p. 57). Inspiration could be *fragmentary* (Koch, 2003).

Fragmentary inspiration offers a different account of the writing process. Ideas originate as small fragments which you write down. Later, after reading and thinking about the fragments, you develop larger ideas, full sentences, and completed manuscripts. The completed work comes last: “*you cannot know a story until it has been told*” (Koch, 2003, p. 6, his italics).

Ray Bradbury’s writing process embraced fragmentary

inspiration. He started his stories by free-associating fragmentary ideas to create lists. “I began to make lists of titles, to put down long lines of *nouns*. These lists were the provocations, finally, that caused my better stuff to surface. I was feeling my way toward something honest, hidden under the trapdoor on top of my skull (Bradbury, 1990, p. 17, his italics).

To produce a complete work, fragmentary inspiration requires you to conduct a conversation with your writing *as you write*. You write down your fragmentary ideas, read them, reflect upon them, and develop them into a larger fabric by writing more. ‘Writing’ becomes a dynamic back and forth between you and your written word. You let “the emerging story tell itself *through* you. As you tell it, you let the story give you your cues about where it is going to go next” (Koch, 2003, p. 6, his italics). John Gardner discovered “what every good writer knows, that getting down one’s exact meaning helps one to discover what one means” (Gardner, 1983, p. 19).

The dynamic back and forth between you and your written word radically changes conceptions of the writing process because the dynamic occurs between your mind and your world – the words you have already written down. The next section proposes the dynamic between your mind and your writing makes the writing process more embodied.

1.5 COGNITIVE SCAFFOLDING

Use your world to make writing easier.

A writer's back-and-forth conversation with what they write, required by fragmentary inspiration, illustrates what cyberneticists call a *feedback loop* (Ashby, 1956; Grey Walter, 1963; McCulloch, 1965; Wiener, 1948). Cyberneticists argued feedback loops provide a critical means for behavioral control. Cyberneticists defined a feedback loop as the constant back-and-forth between an agent's actions upon the world and how such changes guide the agent's future actions. Feedback tells an agent their distance from a goal. Feedback causes an agent to perform an action which brings the agent closer to their goal.

The feedback loop created by fragmentary inspiration involves three repeating events: 1) fragmentary ideas come into your mind; 2) you change your world by writing the fragments down; 3) your world (written words) provides you new ideas when you read and interpret what you have written. The cycle repeats when you write down your next ideas.

The feedback loop treats writing not as resulting from complete sentences arising from inspiration but rather as a writer's constant search for meaning which begins by exploring fragments which they have written down already. "The

Latin root of the word *invent* means ‘to find’. And since you cannot know what you have to say until you have said it, writers of both fiction and nonfiction ‘invent’ *through* finding” (Koch, 2003, p. 14, his italics).

The ‘finding’ Koch describes seems more embodied than the disembodied inspiration described by Housman. The feedback loop required by fragmentary inspiration moves some ‘writing’ from inside the mind to outside in the world. Writing now includes the world because the feedback loop requires what writers have already written in the world to affect their later ideas.

The feedback loop required by fragmentary inspiration relates writing to another idea from embodied cognitive science, cognitive scaffolding. Cognitive scaffolding occurs when objects in the world aid or replace mental processing (Clark, 1997). For instance, when a student takes lecture notes, they replace their internal memory with an environmental record. Lecture notes scaffold memory. Cognitive scaffolds belong to embodied cognitive science because scaffolds exist in the world, not in the disembodied mind.

Cognitive scaffolding plays an important role in writing. Many writers use, and many books about writing recommend using, notebooks or index cards to record ideas when they occur (Ahrens, 2022; Atchity, 1995; Brande, 1934; Butler & Burroway, 2005; Cameron, 2022; Goldberg, 1986; Heighton, 2011; Hemingway, 1964; Janzer, 2016; Koch, 2003; Lamott, 1995; Lowes, 1927; Luhmann, 1992; Raab, 2010;

Rhodes, 1995; Swain, 1974). Notebooks illustrate prototypical cognitive scaffolds.

How do scaffolds like notebooks aid writing? Scaffolds reduce demands on (internal) thinking. For example, writers use notebooks to scaffold memory by jotting down ideas as soon as they occur. Keeping the idea in memory, and recording it later, is too risky: “I used to think that if something was important enough, I’d remember it until I got home, where I could simply write it down in my notebook like some normal functioning member of society. But then I wouldn’t” (Lamott, 1995, p. 127). With a notebook, you don’t have to retrieve an idea from memory; instead, you simply read your notebook to remember the idea.

Furthermore, embodied cognitive scientists recognize when you move ideas from your mind to your world, you can scaffold more than your memory. When you attach ideas to objects in the world, you scaffold additional thinking processes, because embodied cognitive scientists view thinking as physically manipulating objects. Your world not only can your memory, but your actions on your world can generate new ideas for your writing. Cognitive scaffolds introduce a radical new idea: your world literally becomes part of your mind (Clark, 1997, 2008a; Clark & Chalmers, 1998; Shapiro, 2019).

Despite Chapter 1’s appeal to ideas like ‘feedback loop’ and ‘cognitive scaffolding’, you do not need to understand embodied cognition to take advantage of what it offers writing. For example, many students comfortably take notes to

scaffold their memory without needing to understand embodied cognition. You can feel reassured, though, by knowing embodied cognition helps explain why scaffolds for writing work.

1.6 EXTENDING THE MIND

Changing how you think about thinking changes how you think about writing.

Many cognitive scientists assume the disembodied mind and treat thinking as a sense-think-act cycle (Dawson, 2013, 2022). However, embodied cognitive scientists see problems with sense-think-act processing and propose an alternative: the *sense-act cycle*. Sense-act processing replaces internal thinking with acting on the world. If thinking involves a sense-act cycle, then you must reconsider the writing process.

A 20th century robot illustrates sense-act processing's power (Grey Walter, 1950a, 1950b, 1951, 1963). Grey Walter named his robot *Machina speculatrix*, but others called it the Tortoise because it looked like a toy tractor surrounded by a tortoise-like shell. The Tortoise generated very complex behavior. The *Daily Mail* reported “the toys possess the senses of sight, hunger, touch, and memory. They can walk about the room avoiding obstacles, stroll round the garden, climb stairs, and feed themselves by automatically recharging six-volt accumulators from the light in the room. And they can dance a jig, go to sleep when tired, and give an electric shock if disturbed when they are not playful” (Holland, 2003a, p. 2090).

However, the Tortoise's complex behavior did not arise from complex internal representations (a disembodied robot mind). It arose instead from direct links between two senses (detecting light and detecting touch) and two motors (one for steering, one for moving forward). Motor speeds changed depending upon whether the Tortoise sensed dim, moderate, or bright light. Motor speeds also changed when an obstacle bumped the Tortoise's shell. The Tortoise demonstrated sense-act processing's power, anticipating Herbert Simon's *parable of the ant* (Simon, 1969).

In the parable of the ant, Simon (1969) asks how you might explain an ant's complicated route while walking along a beach. Simon's answer uses sense-act processing, not complex thinking – the ant turns when it encounters obstacles. Like the Tortoise, the ant directly converts what it senses about the world into motor activity. “Viewed as a geometric figure, the ant's path is irregular, complex, hard to describe. But its complexity is really a complexity in the surface of the beach, not a complexity in the ant” (Simon, 1969, p. 24). Simon moves complexity from inside the ant's head to outside in the ant's world.

The Tortoise illustrates Simon's parable by generating interesting behavior without representing the world. Building, storing, and manipulating representations requires time-consuming processing, which slows behaving down. Systems which behave without representing – without ‘thinking’ – can act quickly by simply reacting to the world. “Models of the

world simply get in the way. It turns out to be better to use the world as its own model” (Brooks, 2002, p. 139).

Embodied cognitive science replaces sense-think-act processing with sense-act processing (Clark, 1997, 1999, 2003, 2008a; Shapiro, 2014, 2019). By adopting sense-act processing, embodied cognitive scientists generate theories which move the causes of complex behavior from internal thinking to the external world. As a result, embodied cognitive scientists propose the *extended mind hypothesis* (Clark & Chalmers, 1998). According to the extended mind hypothesis, no boundary exists between your mind and your world. “It is the human brain *plus* these chunks of external scaffolding that finally constitutes the smart, rational inference engine we call mind” (Clark, 1997, p. 180, his italics).

Changing how we think about thinking changes how we think about writing. The extended mind radically differs from the disembodied mind. As a result, ‘writing’ by an extended mind radically differs from ‘writing’ by a disembodied mind. How does the extended mind bring writing processes to life? What cognitive scaffolds can writers take advantage of?

1.7 WRITING WITH YOUR WORLD

You think by acting upon your world.

Embodied cognition makes cognitive scaffolding becomes central to cognitive theory. You can scaffold your writing when you treat writing as embodied thinking. I have already mentioned one example, the notebook. By moving your ideas into your notebook, you no longer need to worry about forgetting them. Importantly, written objects can scaffold more than memory by offering you other possible actions. For example, consider Peter Elbow's cut-and-paste revising method (Elbow, 1981).

Cut-and-paste revising assumes you have drafted paragraphs which you intend to improve. You do so, but not by writing. "You throw away your pen or pencil and revise with nothing but scissors and paste. You will be like a stone sculptor who never adds – only removes" (Elbow, 1981, p. 147).

Cut-and-paste revising requires you to read your draft, find good passages, and use scissors to cut them out. You spread out your clippings, rearrange them; you look for an emerging narrative. You arrange the clippings in their best order. You create a new draft by copying your rearranged clip-

pings, performing what little writing you require to connect fragments together.

Elbow (1981) describes a related method, collage writing, which begins with some completely unconnected writing fragments. In collage writing, you print each fragment out. You arrange the fragments on a table or floor and rearrange them to find the best order. Surprising meaning often emerges from such action. When you use collage writing, “you don’t worry about a thread at all, you just look for quality. You get an *implied thread* to assert itself by arranging the good bits in the right order” (Elbow, 1981, p. 150, his italics).

Elbow’s two methods presume you have written complete sentences, paragraphs, or passages. However, his methods also work with more fragmentary written objects. Imagine writing short ideas – like Ray Bradbury’s nouns – on their own individual objects, like index cards. Bradbury would search for a pattern in his list. But you could also discover patterns by using the collage method to rearrange index cards to find a surprising thread, while removing index cards which do not seem to fit.

Why might I say rearranging index cards scaffolds more than your memory? Consider the disembodied alternative: not only keeping all your fragmentary ideas in memory at once but also seeking meaningful relationships between them by rearranging them in your mind. When index cards provide scaffolds, you move ideas into your world. Having your ideas in physical format permits you to explore relationships by mov-

ing index cards around, by placing cards for related ideas closer to each other – in your world. You no longer remember your ideas, discovering relationships within ideas in your mind. Instead, you simply look at your cards; you see your ideas and their relationships without remembering them or thinking about them. Rearranging index cards in your world permits you to think about your ideas while reducing demands on cognition.

However, to serve as a useful scaffold, index cards – or any other scaffold – must be tailored to the actions an agent can perform. Scaffolds reduce demands on cognition by converting thinking into acting on the world. However, the actions on the world offered by a scaffold must be compatible with the agent's body. A scaffold offers certain potential actions to an agent whose body can manipulate the scaffold in particular ways. The same scaffold offers different potential actions to an agent with a different body. For a particular agent – a human writer – what potential actions do index cards offer to scaffold writing?

1.8 INDEX CARD AFFORDANCES

A writing scaffold must offer you appropriate affordances.

Why can index cards scaffold writing? To answer, consider another embodied concept: *affordance*. An affordance is a possible action the world offers an agent. An affordance depends upon an agent's, and the world's, physical properties (Gibson, 1979). Affordances “have to be measured *relative to the animal*. They are unique for that animal. They are not just abstract physical properties” (Gibson, 1979, p. 127, his italics). The same object offers different affordances to animals with different bodies because animal bodies dictate what actions an animal can perform. A doorknob offers ‘turnability’ to a human’s hand, but not to a cat’s paw.

A scaffold for writing *replaces* thinking with acting upon objects in the world (Shapiro, 2019). To aid writing, the scaffold must offer appropriate affordances. Consider the affordances offered to me by my preferred writing scaffold, 3” X 5” index cards.

Index cards offer ‘writability’. Most obviously, index cards offer ‘writability’: you can mark an index card’s surface with pen or pencil. Writability is crucial: to scaffold writing

you must move your thoughts from your mind to your world. You do so when you jot notes on index cards.

Index cards offer ‘readability’. Once you write on an index card, you can later read what you wrote on the card. Combined, writability and readability permit index cards to replace your memory. When you jot an idea on an index card, you no longer need to keep it in memory. You do not retrieve information from memory; instead, you inspect a card and read its contents. The index card reduces memory demands, freeing cognitive resources for other tasks, such as evaluating ideas.

Index cards enforce ‘writing short’. An index card’s small size restricts how much you can write: conciseness arises because index cards offer little space for wordiness or eloquence. When you write topics on index cards you must write short (Clark, 2014). I often express an idea by jotting down a short phrase; I rarely write a complete sentence.

Being forced to write short helps a writer, particularly during a project’s early stages. When a project begins, you need to generate ideas or topics. Evaluating ideas as you generate them inhibits creativity (Osborn, 1948, 1953). If you express topics in sentences, you invite unwanted evaluation. Writers automatically evaluate their sentences, reducing creative flow (Elbow, 1981). Index cards encourage you to jot down short phrases, allowing you to elude premature evaluation.

Index cards offer ‘arrangeability’. Index cards offer more than writability or readability. You can arrange index

cards on a two-dimensional surface. For instance, in my lab I use magnets to attach index cards to a large whiteboard. I can work with cards by moving them around on the surface: by placing cards in a particular order or by grouping related cards closer together.

Arrangeability replaces more than your memory. Thinking about ideas becomes rearranging your index cards in the world. By moving cards around, you physically explore answers to various questions: What topic needs do you need to present first? Do some topics belong in one section while others belong in a second? Are some topics redundant? You change card positions to represent answers to such questions. You literally *think* when you move your cards around; you can see your thinking's results when you examine index card positions in your display.

Index cards offer 'groupability'. Writing projects usually have a hierarchical structure, reflected in different sections or chapters. Notes related to a writing project also tend to belong to different types: ideas or topics, sources to cite, quotations to include, and so on. You can use index card 'arrangeability' to represent such structure, for example by placing cards related to the same section closer together.

However, index cards have other properties to help represent related ideas. For instance, you can find different colored index cards. You can use different colored cards to represent different ideas. I usually write topic notes on white cards,

but use pink cards for sources, green cards for quotes, blue cards for figures, and so on.

Using visual information like index card color makes relationships between different cards visually explicit. By looking at their colors, you can easily see different card types in your display. Visual ‘groupability’ also illustrates replacement, because you need not remember or think about relationships between ideas or between project components. Instead, you immediately see the relationships in your external display.

Index cards offer ‘portability’. Being small makes index cards extremely portable. I worked on the current chapter while on a family trip to Nova Scotia. I did not have room to pack my laptop in my carry-on baggage. Fortunately, I had room for index cards which I worked on during a long flight. Later, running low on blank cards, I easily found more to buy on my trip. Index card portability allows me to write anywhere.

Index cards offer ‘expendability’. Consider important writing advice originally offered by Sir Author Quiller-Couch: “Whenever you feel an impulse to perpetrate a piece of exceptionally fine writing, obey it—whole-heartedly—and delete it before sending your manuscript to press. *Murder your darlings*” (Quiller-Couch, 1916, pp. 234-235). You often must sacrifice ideas to improve your writing.

Index cards make discarding ideas easier. Index cards, being cheap and recyclable, are easy to throw away. When organizing ideas or topics into a narrative, you might find cards which don’t fit in. I can easily murder a card which does not fit

in if I haven't yet written sentences but instead have only jotted a short note on an index card. I simply toss the card into my recycle bin.

Index cards offer 'cooperativity'. Your index card display exists in public; more than one person can view or manipulate the index cards at the same time. The scaffold makes a writing project available to everyone on a writing team.

Index cards offer 'duality'. I sketch a scaffolding method for writing in the next section and describe it in detail in Chapter 2. The scaffold begins by writing broad topics on index cards and continues by converting broad topics into several, more specific, paragraph topics. With a satisfactory paragraph topic structure in place, you move away from jotting short notes (topics) to writing complete sentences. My method calls for writing the first and last sentence for each paragraph. Doing so takes advantage of the 'duality' offered by an index card: you can write on both of its sides. As you convert an organized set of topics into an organized chain of paragraphs, you use 'duality' by writing two sentences on the back of each paragraph's index card. The next section briefly describes my scaffolding method. Chapter 2 details the scaffolding method's logic.

Chapter 2 also explains differences between my methods and others which use index cards to outline (Atchity, 1995; Butler & Burroway, 2005). For instance, while the method described in Chapter 2 leans heavily on index cards, later chapters introduce different scaffolds, because if index cards can

scaffold writing, then other objects can too. Beginning in Chapter 3 the book introduces additional scaffolds for academic writing.

Implications of Affordances. Affordances are crucial for embodied thinking. The thinking you can do depends upon what affordances the objects being manipulated offer you as you think with the world.

On the one hand, the affordances offered by index cards are particularly powerful (Krajewski & Krapp, 2011). Krajewski and Krapp trace the history of index cards from the 16th century use of paper slips to catalog library holdings to the 20th century use of card catalog systems by businesses. Krajewski and Krapp argue index card affordances (in particular, the ones I have called writability, readability, arrangeability and portability) provide index card catalogs all the information processing abilities of digital computers. As a result, they call card catalogs ‘paper machines’ and justify the 1929 Fortschritt GmbH’s (a catalog producer) claim “card catalogs can do anything” (Krajewski & Krapp, 2011, p. 1).

On the other hand, index cards do not offer *every* affordance. Other objects offer different affordances which alter the (embodied) thinking you can do with them. For example, In section 5.4.2 I argue computerized index card systems offer different affordances – affordances which, for me, make them less useful for writing.

1.9 A WRITING SCAFFOLD

Topics before sentences.

How do I use index cards to scaffold my writing? Let me briefly sketch a method which I present in more detail in Chapter 2.

When I begin a writing project, I generate topics I need to communicate. I write each topic on its own index card; I move my topics from my memory to a display in my world. Then I manipulate my topics – I rearrange my cards – to seek a narrative structure. With the narrative structure taking shape, I insert cards representing sources to cite, quotes to use, figures to display, or tables to include into appropriate positions.

I enter into a feedback loop between my thinking and my index card positions. As I read and rearrange my topic cards, I ask what points I need to convey my topic to my reader. I use index cards to represent new subtopics which I add to my display. I continue to rearrange my index cards, seeking a satisfactory narrative structure, adding or removing index cards as required.

As my narrative takes shape, I begin to use index cards to represent more specific topics – each card represents a single

topic to communicate with a single paragraph. I still rearrange my cards seeking the best order for my paragraph topics.

When I feel happy with my paragraph topic order, I finally start writing complete sentences. I take an index card – upon which I have written a paragraph topic – and turn the card over. At the top of the index card’s reverse side, I write the paragraph’s first sentence: the topic sentence states the paragraph’s topic. I write a topic sentence for each paragraph topic card in my display.

After writing topic sentences, I next write each paragraph’s concluding sentence. I write the sentence at the bottom of each index card’s reverse side. I relate each concluding sentence to the topic sentence on the same card, as well as to the topic sentence on the next card.

Once I have written a topic sentence and a concluding sentence on each paragraph’s index card, I move the sentences from my index cards to a word processing document. I create a document by typing each pair of sentences into its own paragraph in our word processing file. I may also add place holders for sources, quotes, figures, or titles for sections and subsections. Afterwards, I have a remarkably complete, well-structured outline in electronic form.

The method described above begins from fragmentary inspiration, which views writing as a back-and-forth conversation between what writer’s think and what writer’s write. Such feedback requires you to move some writing processes from your mind to your world. The method then adopts an even

more embodied perspective. First, makes the outline process embodied by treating outlining as jotting short ideas onto index cards and conversing with the ideas by physically rearranging cards in a display. Second, the method reduces cognitive load by delaying composing – and evaluating – complete sentences until after you have discovered a strong narrative.

When I move my index cards from the physical world into a word processing document, I feel I haven't yet started 'writing'. My outline's organization and length, produced by my index card manipulations, usually amazes me. Furthermore, I feel less intimidated by my remaining writing because I need only add a few sentences to each paragraph, sentences dictated by my existing topic/concluding sentences pairs. My index card scaffold performed most of the hard work for me in advance.

Bruce Springsteen sings in *Jungleland* "And the poets down here don't write nothing at all/They just stand back and let it all be" (Springsteen, 1975). Springsteen's poets don't represent their world; they let their world represent itself. The scaffold which I detail in the next chapter makes me feel as if I have stood back and let my world write my manuscript for me. Chapter 2 describes how the scaffold makes academic writing easier. Later chapters introduce additional scaffolds to help academic writing.

PART II

CHAPTER 2: SCAFFOLDING AN OUTLINE

Use your world to outline a paper.

Chapter 1 introduced embodied cognition and argued you can take advantage of your world to help your writing. Chapter 2 provides a method for using your world to scaffold a detailed outline for a writing project. The method begins by generating potential topics for a paper, writing them on index cards, and putting the cards on display. By placing topics in your world you remove the need to keep the ideas inside your memory. You then manipulate the cards, seeking a sensible narrative. You also decompose broad topics into more specific topics, seeking ‘paragraph topics’ – topics specific enough to convey with a single paragraph. Chapter 2 describes how you convert paragraph topics into sentence pairs, producing a rich outline which makes completing your manuscript much easier. Chapter 2 also distinguishes my index card method from other systems which use index cards; the chapter ends by arguing the utility of index cards suggests different kinds of scaf-

fold for writing are worth exploring. Later chapters describe additional writing scaffolds.

2.1 SHOULD YOU OUTLINE?

Outline first, write last.

What is the best way to improve writing? Many books about writing answer by promoting *outlining* (Atchity, 1995; Butler & Burroway, 2005; Carpenter, 2020; Dillard, 1989; Hawker, 2015; Heard, 2022; Kumar, 2020; Rosnow & Rosnow, 1998; Sarnecka, 2019; Sawers, 2002; Schimel, 2012; Silvia, 2007; Swain, 1974; Wheelan, 2022). “On my list of maladaptive practices that make writing harder, Not Outlining is pretty high. ... Writers who complain about ‘writer’s block’ are writers who don’t outline” (Silvia, 2007, p. 79).

However, other books disagree with such advice and recommend avoiding outlining altogether (Greene, 2013; Kidder & Todd, 2013; King, 2000; Lamott, 1995). Some books present outlining’s benefits and costs and do not promote outlining as a method (Pinker, 2014; Sword, 2017). Some books fail to even mention outlining, (Brande, 1934; Flaherty, 2009; Kail, 2019; Strunk & White, 1959; Sword, 2012, 2016; Zinsser, 1988, 2006).

Why do books about writing not agree about outlining’s benefits? First, many successful writers do not outline. Stephen King offers advice for avoiding what he calls the

tyranny of the outline (King, 2000). E. L. Doctorow said “The most important lesson I’ve learned is that planning to write is not writing. Outlining a book is not writing. Researching is not writing. Talking to people about what you’re doing, none of that is writing. Writing is writing” (Weber, 1985).

Second, few agree about what ‘outlining’ means. To some, a discipline’s writing conventions provide an outline (Sarnecka, 2019). To others, an outline merely consists of a few phrases scribbled on paper (Clark, 2008b). Others use a catalog of written notes, with links from one note to another — a *zettelkasten* – to create a manuscript’s structure (Ahrens, 2022; Kadavy, 2021; Luhmann, 1992). Heard (2022) describes many different methods for outlining, including two sentence mini summaries, word stacks, concept maps, figure shuffling, and listing paragraph topic sentences. From her interviews with academic writers, Sword (2017) reports widely varying outlining practices, ranging from historian Kevin Kenny, who produced a fifty-page outline he then converted into a book in a single summer, to historian Russell Gray, who first plans a manuscript out in his head (only occasionally jotting the plan down) before sitting down and writing the full manuscript.

In short, writers who outline use widely varying methods, and many writers avoid outlining altogether. Despite such diversity, when I ask myself ‘Should I outline?’, I answer with a resounding ‘Yes!’. In the current chapter I explain my answer by introducing an outlining scaffold.

My method uses cognitive scaffolding to create an out-

line. My method moves ideas from my mind into my world where I can evaluate and organize them. My outline develops by external thinking: I read and manipulate index cards which hold my ideas. My method produces a paragraph-by-paragraph structure to insert into a word processing document. The structure provides an ordered set of paragraph topics; you draft two key sentences for each paragraph; you also can place other components (section headings, figures, tables, citations). When I move my outline into a word processor, I find I have already completed the hard work required to create my first draft.

2.2 REVERSE ENGINEERING AN OUTLINE

Within a good paper you can find a good outline.

What structure should your outline have? You need an outline which helps you to write. You also need an outline which you can easily and usefully scaffold. To propose such an outline, I follow some common writing advice (Sarnecka, 2019; Strunk & White, 1959).

First, writing guides advise you to treat the paragraph as writing's basic element because a good paragraph conveys a single topic (Strunk & White, 1959). Figure 2-1 illustrates 'one paragraph, one topic' by representing a paragraph with a cone. The cone's single point reminds you the paragraph conveys only one topic.

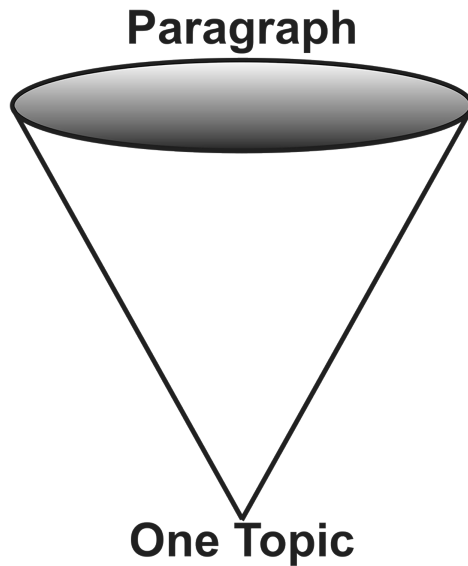


Figure 2-1. A paragraph is the basic element of writing; the cone shape illustrates a paragraph conveys only one topic.

Second, writing guides remind you about the importance of a paragraph's first and last sentences (Sarnecka, 2019; Strunk & White, 1959). A paragraph's first and last sentences convey a good paragraph's topic. We call a paragraph's first sentence the *topic sentence*. A topic sentence states the paragraph's topic to your reader. We call a paragraph's last sentence the *concluding sentence*. The concluding sentence brings the paragraph's topic home by summarizing or restating the topic for your reader.

A paragraph's topic sentence and concluding sentence relate meaningfully to one another because both convey the

same topic. “One easy way to check for coherence in a paragraph is to read just the topic sentence and the concluding sentence. If they aren’t on the same theme, the paragraph has wandered off track and needs some attention” (Sarnecka, 2019, p. 236). Figure 2-2 uses arrows to illustrate the meaningful relationship between a paragraph’s topic and concluding sentences.

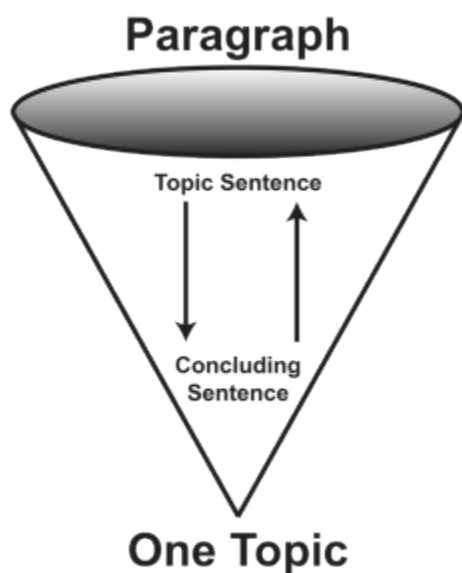


Figure 2-2. The topic sentence and the concluding sentence of a paragraph convey the same topic, which meaningfully links them to one another, as illustrated by the arrows.

Writing guides offer a third kind of advice by encouraging you to find examples of other writer’s work to admire, to learn from, or to try and emulate. I now use the three pieces of

writing advice to propose a desirable structure for an outline. I start with a writing sample I admire. I then use the other advice to ‘reverse engineer’ an outline which could create the writing sample. I do so by examining each paragraph’s topic and concluding sentences and by inferring a paragraph topic from the two sentences.

For my writing example, I chose ‘Connectionism and cognitive architecture: a critical analysis’, published in *Cognition*, and written by Jerry Fodor and Zenon Pylyshyn (Fodor & Pylyshyn, 1988). The paper criticizes cognitive scientists who use artificial neural networks. I chose the paper because it relates to my own neural network research, because of its influence (1792 citations as of February 2025), and because of its solid writing. You do not need any knowledge about cognition or connectionism to understand the reverse engineering I illustrate with the Fodor and Pylyshyn example.

How can I reverse engineer an outline from my writing sample? I pay heed to the paper’s paragraphs. I start by copying each paragraph’s first and last sentences. Table 2-1 provides the topic sentence and the concluding sentence for the first eight paragraphs of Fodor and Pylyshyn’s (1988) introduction.

After isolating the topic and concluding sentences, I next read each pair to figure out each paragraph’s topic. Remember, you state a paragraph’s topic with its topic sentence, and the paragraph’s concluding sentence should also refer to the paragraph’s topic. The final column in Table 2-1

provides my proposed topics for the introductory paragraphs in Fodor and Pylyshyn (1988).

Table 2-1 reverse engineers a paper which is important to my own writing and research experience. However, your interests likely differ from mine, making Fodor and Pylyshyn (1988) a poor example for you. Remember, though, you can perform the procedure I used to create Table 2-1 on a different paper, one related to your own interests, to produce an example more meaningful to you.

Table 2-1 illustrates a one-to-one mapping between each paragraph topic and each paragraph's topic and concluding sentences. In the current chapter I propose a scaffold for developing an outline which includes paragraph topics; the scaffold helps me convert each paragraph topic into the paragraph's first and last sentences. I begin describing my scaffold by considering the need to meaningfully organize paragraph topics.

Table 2-1. The first and last sentence of the paragraphs which make up the introduction to Fodor and Pylyshyn (1988), and the assumed topic of each paragraph.

Paragraph	Topic Sentence	Concluding Sentence	Topic
1	Connectionist or PDP models are catching on.	There are also, inevitably, descriptions of Connectionism as a Kuhnian “paradigm shift”.	Connectionism is interesting
2	The fan club includes the most unlikely collection of people.	Almost everyone who is discontent with contemporary cognitive psychology and current “information processing” models of the mind has rushed to embrace “the Connectionist alternative”.	Connectionism attracts scholars unhappy with the status quo

Table 2-1. The first and last sentence of the paragraphs which make up the introduction to Fodor and Pylyshyn (1988), and the assumed topic of each paragraph.

3	When taken as a way of modeling cognitive architecture, Connectionism really does represent an approach that is quite different from that of the Classical cognitive science that it seeks to replace.	The style of processing carried out in such models is thus strikingly unlike what goes on when conventional machines are computing some function.	Connectionism attracts researchers unhappy with the status quo because connectionism departs from standard approaches
4	Connectionist systems are networks consisting of very large numbers of simple but highly interconnected “units”.	The behavior of the network as a whole is a function of the initial state of activation of the units and of the weights on its connections, which serve as its only form of memory.	What are the properties of connectionism which make it different?

Table 2-1. The first and last sentence of the paragraphs which make up the introduction to Fodor and Pylyshyn (1988), and the assumed topic of each paragraph.

5	Numerous elaborations of this basic Connectionist architecture are possible.	The term ‘Connectionist model’ (like ‘Turing Machine’ or ‘Von Neumann machine’) is thus applied to a family of mechanisms that differ in details but share a galaxy of architectural commitments. We shall return to the characterization of these commitments below.	Many different versions of connectionist networks are possible
6	Connectionist networks have been analyzed extensively – in some cases using advanced mathematical techniques.	Of even greater interest is the fact that such networks can be made to learn; this is achieved by modifying the weights on the connections as a function of certain kinds of feedback.	Connectionist networks have been studied a lot, using mathematical techniques, because they learn

Table 2-1. The first and last sentence of the paragraphs which make up the introduction to Fodor and Pylyshyn (1988), and the assumed topic of each paragraph.

7	In short, the study of Connectionist machines has led to a number of striking and unanticipated findings' it's surprising how much computing can be done with a uniform network of simple interconnected elements.	Surely this is a proposal that ought to be taken seriously: if it is warranted, it implies a major redirection of research.	Studies of connectionism reveal surprising results: we should take connectionism seriously
8	Unfortunately, however, discussions of the relative merits of the two architectures have thus far been marked by a variety of confusions and irrelevances.	The arguments that then appeared to militate decisively in favor of the Classical view appear to us to do so still.	Twist: But, when you properly compare connectionist networks to classical models, classical models still win!

2.3 MEANINGFUL CHAINS OF PARAGRAPHS

Arrange your paragraphs in meaningful order.

Paragraphs do not work on their own; in good writing a link or bridge exists from one paragraph to the next. You convey your main topic by creating meaningful links between paragraphs. You create such links using your concluding sentences.

As noted in Section 2.2, a concluding sentence relates to the topic sentence in the same paragraph. However, the one paragraph's concluding sentence also provides a bridge to the next paragraph's topic sentence. Figure 2-3 illustrates one concluding sentence's links to two topic sentences. The links connect paragraphs together, creating a meaningful chain.

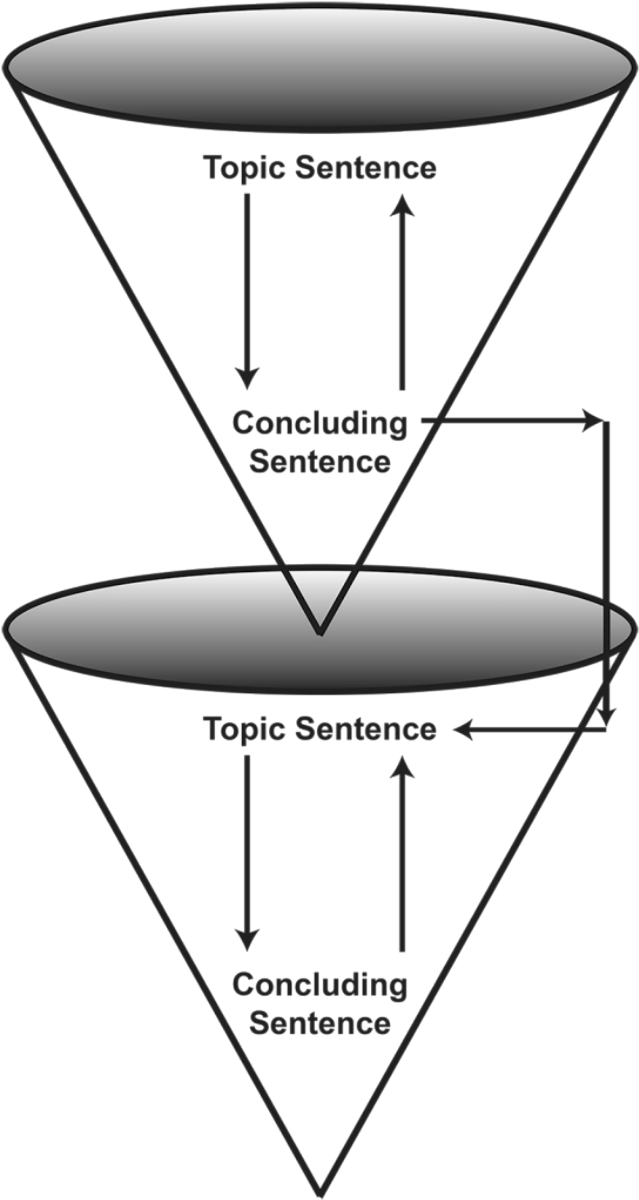


Figure 2-3. The concluding sentence of the first paragraph is meaningfully linked to two topic sentences: the one in the

concluding sentence's paragraph and the one in the next paragraph.

For example, you can create a meaningful chain of paragraphs by ordering your paragraph topics according to breadth. The first paragraph communicates the broadest topic, the next paragraph communicates a narrower topic, and so on. You smoothly move your reader from general points to more specific details, all the while keeping the paragraph chain on the larger topic. Figure 2-4 illustrates the approach by placing the paragraph chain inside a larger cone. The larger cone's narrowing represents the focusing of paragraph topics as one moves through the paragraph chain. The paragraph chain communicates the larger narrative's topic.

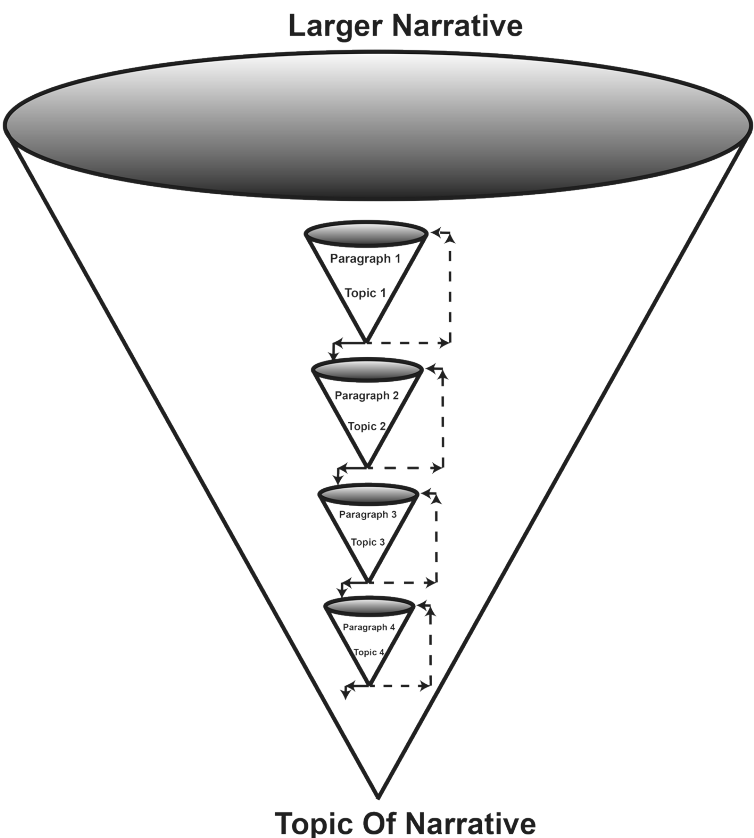


Figure 2-4. Ordering paragraph topics from general to specific communicates a larger narrative. In the figure, a larger inner cone communicates a broader topic than does a smaller inner cone.

Fodor and Pylyshyn (1988) organize their introduction according to Figure 2-4. The paragraph topics in Table 2-1 become more focused as you move down the table. For instance, Fodor and Pylyshyn begin by stating many find con-

nectionism interesting. They then focus by noting certain researchers find connectionism particularly interesting to certain researchers. They focus again by claiming connectionism attracts such researchers because connectionism differs from traditional approaches. They continue to focus throughout their introduction.

You can now see my Table 2-1 example reveals a detailed structure. The outline provides topics in order; the chain of topics moves smoothly from more general points to more detailed claims. The outline expresses each topic with a topic sentence and a concluding sentence. Links between concluding sentences and topic sentences establish the introduction's narrative structure. How can you create such a nicely structured outline – from scratch?

2.4 A DESIRABLE OUTLINE

You need a detailed outline to scaffold your first draft.

When I created Table 2-1, I performed the opposite of outlining, because I started with complete sentences and ended with paragraph topics. I presume Fodor and Pylyshyn did the reverse and began with topics which they later converted into paragraphs. Their final outline may have had a structure like the one presented in Figure 2-5.

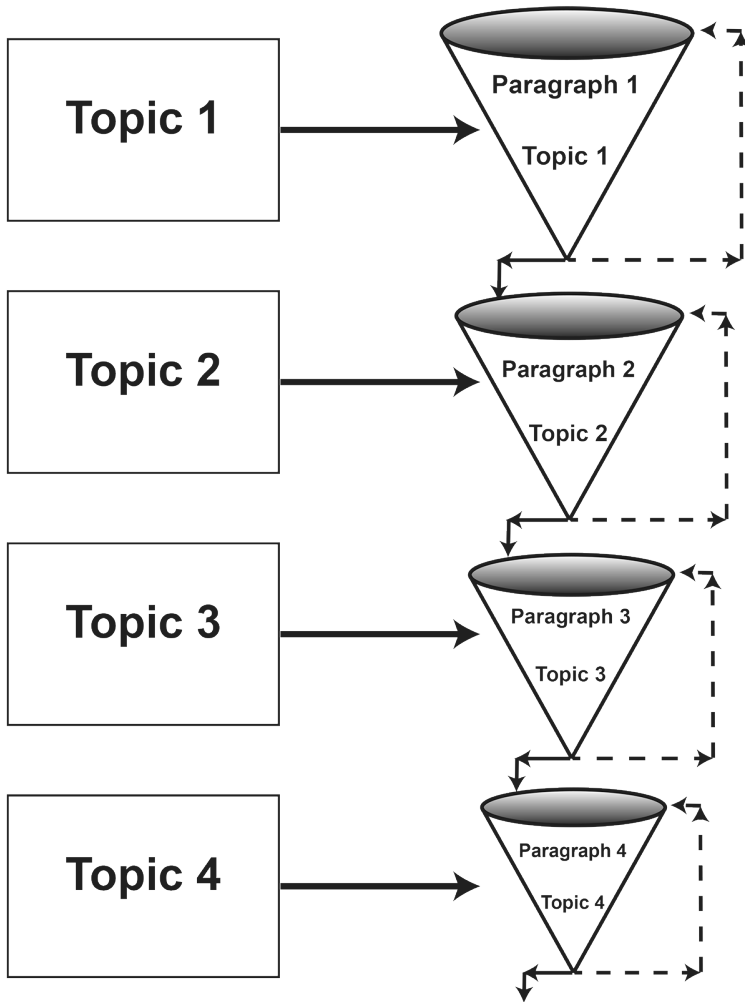


Figure 2-5. A general form of an outline in which topics are mapped into paragraphs.

Figure 2-5's most important property is the one-to-one correspondence between topics and paragraphs. Sarnecka (2019, p. 235) notes "organizing your writing into topic-sen-

tence paragraphs allows you to switch back and forth between outlines and drafts, which is magic when you are developing a complicated argument.”

Note the Figure 2-5 paragraphs are incomplete. Each cone in the figure represents a particular paragraph which will communicate a particular topic. When you begin to create an outline, your outline need only contain potential topics.

Figure 2-5 illustrates another important property: topic order. In creating your outline, you decide to present topics in a particular order to communicate your manuscript’s main message. For instance, you could organize your topics along the lines illustrated in Figure 2-4, moving from general topics to more specific topics.

Figure 2-5’s dual nature implies paragraph order on the right mirrors the topic order on the right. You reinforce paragraph ordering with links between paragraphs, illustrated by the arrows connecting paragraphs in Figure 2-5. Recall, from Figure 2-3, you link paragraphs by connecting each concluding sentence to two different topic sentences. Hence your most desirable outline would present topics in order, each communicated by a different paragraph; your outline would also provide each paragraph’s first and last sentence.

Imagine taking such an outline and typing each paragraph’s topic and concluding sentences into a word processing document. Your typed entries provide your manuscript’s skeleton: you have topics in a desired order, you have mapped each topic onto a single paragraph, and you have generated topic

and concluding sentences for each paragraph. You provide your manuscript's narrative with the meaningful links you have made between concluding sentences and topic sentences. To convert your document into a workable first draft, you only need to put flesh on your skeleton.

Adding flesh to the skeleton requires adding supporting sentences between each paragraph's topic and concluding sentences. But your manuscript's skeleton scaffolds what supporting sentences you add. Supporting sentences elaborate a paragraph's topic. Fortunately, you already have each paragraph's two most important sentences in your document, sentences which place strong constraints on what supporting sentences you can add.

You want to create a detailed outline for a manuscript, an outline with a structure like the one in Figure 2-5. However, such an outline takes effort to create. You rarely can generate the correct paragraph topics, in a desirable order, when you start a manuscript. Instead, you must work hard to produce the desired outline. Let me now show how you can use cognitive scaffolding to develop an outline whose structure is illustrated in Figure 2-5.

2.5 MOVING TOPICS INTO YOUR WORLD

Move your ideas from your mind to your world.

How do you scaffold the creation of an outline with a structure like the one in Figure 2-5? The structure in Figure 2-5 has one basic element: the topic. A single paragraph expresses each topic. You use the paragraph to inspire your cognitive scaffold by basing your cognitive scaffold on topics. To replace your need to represent topics in your mind, you represent each topic as an object in your world. You organize topics by rearranging objects in your world. For example, you can place objects representing related topics near one another in your world. Moving objects around in your world replaces thinking about topics in your mind.

I propose you represent topics using a particular object: the 3" x 5" index card, blank on one side and lined on the other. When a topic comes to mind, you make a note about the topic on an index card. Your note moves the topic from your mind to your world. In Chapter 1, I noted index cards offer many affordances which make them well-suited for scaffolding your outline.

For your scaffold, your index cards offer duality: you can write on both sides of an index card. You exploit duality

when an index card represents a paragraph topic: you write the paragraph topic on the index card's blank side; later, you write the paragraph's topic and concluding sentences on the card's lined side. As a result, a completed index card represents one horizontal slice of the Figure 2-5 outline structure, a slice illustrated in Figure 2-6.

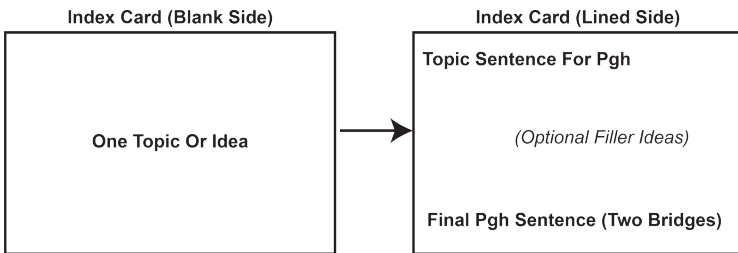


Figure 2-6. The information provided on an index card for a paragraph topic provides the topic on its blank side, and the topic and concluding sentences on the lined side. The card therefore represents one horizontal slice of the Figure 2-5 outline structure.

I cannot begin a writing project by creating index cards as detailed as Figures 2-5 or 2-6. Instead, when I begin a manuscript, I generate broad, unorganized topics. By thinking about – by physically rearranging — my broader topics, I organize my topics into a narrative structure for my manuscript. I can then consider what paragraphs I require to communicate each broader topic. As I consider, I develop topic cards for each potential paragraph (i.e., you create cards like the one on Figure 2-6's left). Later, I add sentences to the opposite side of each card (i.e., Figure 2-6's right).

When you use my method, you create index cards like Figure 2-6 later in your process. However, you can still use index cards earlier to scaffold how you generate and organize broader topics. You perform ‘bootstrapping’, where you begin with index cards for broad topics, you organize them, and then you convert them into larger sets of index cards representing finer details. Bootstrapping, in my method, uses index cards as a medium for fragmentary inspiration – for developing ideas by interacting with what you have already written down.

You eventually create a card set in which each card takes the Figure 2-6 form. When you have such a card set, your scaffold has succeeded; you have a complete outline. I now describe the steps I use to ‘bootstrap’ index cards into a desirable outline.

2.6 STEP 1: GENERATE YOUR BROAD TOPICS

What main topics must I communicate?

Table 2-2A. The scaffolding steps covered so far in Chapter 2.

Step 1	Generate broad topics: Write each topic on the blank side of an index card
--------	--

In Chapter 1, I claimed a writer’s greatest problem arises when they face the blank page. The outlining method I now describe aims to tame the blank page problem. However, isn’t building an outline from scratch just as challenging as facing a new manuscript’s blank page? I believe you will find building an outline from scratch less challenging if you realize you need not generate a complete outline right away. Instead, you simplify outline creation by *bootstrapping*.

In computer science, bootstrapping occurs when a computer first turns on; it loads and executes a small seed program. The seed program provides just enough instructions for the computer to load and start its operating system. Bootstrapping means you start with seed knowledge from which you create more detailed knowledge.

You bootstrap your outline by first generating general ideas about the main points to make in your manuscript (Table 2-2A). Later you refine your general ideas into more specific topics, creating an outline structured like Figure 2-5.

Before describing Step 1, please note Chapter 2 uses Table 2-2 to scaffold remembering my method's steps. The method has nine different steps. As I introduce a step in a new section, I revise Table 2-2 by adding a step. Table 2-2 grows to permit you to keep the whole method in mind with a single glance.

Generating ideas about your manuscript's main points becomes your first step in developing your outline. Where do ideas come from? You can use various techniques, techniques which I discuss in more detail in Chapter 3. In the current chapter, I use one example technique: brainstorming.

Brainstorming aims to generate ideas and foster creativity (Osborn, 1948, 1953; Young, 1940). Brainstorming agrees with one idea introduced in Chapter 1: when you evaluate, you inhibit your creativity. When you brainstorm, you generate as many ideas as possible without evaluating them at all.

Alex Osborn designed brainstorming for use by a working group which created ideas by following four simple ground rules (Osborn, 1948). First, group members withhold criticism until later. Second, group members welcome wild ideas: "The crazier the idea, the better; it's easier to tone down than to think up" (Osborn, 1948, p. 295). Third, a brain-

storming meeting aims to generate many ideas. Fourth, new ideas can result from combining previous ideas together, or from revising a previous idea. “Every idea, crackpot or crack-erjack, is written down. Even silly thoughts are helpful – they keep the group relaxed” (Osborn, 1948, p. 297).

You can take your first step, brainstorming general topics, either alone or with a writing team. When I brainstorm topics, I repeatedly ask and answer one question: “What idea could I write about in my manuscript?”. By answering this question, I generate topics – often in random order. I keep going until I run out of topics. I might take a break from brainstorming when I cannot generate new topics, but return to brainstorming later, which sometimes helps me generate new ideas (Young, 1940).

Importantly, I immediately jot down each idea I generate while I brainstorm – I write my topic down on the blank side of an index card. I write only one topic per index card, avoiding full sentences. I avoid full sentences because I do not want to inhibit generating ideas by criticizing how I express them (Elbow, 1981). I aim to generate as many potential topics as possible, immediately moving them from my mind to my world by using index cards.

When you carry out Step 1, you produce many potential topics for your manuscript, each displayed on its own index card. You can now proceed to Step 2: evaluating and organizing your topics. You evaluate and organize by acting upon your index cards, because Step 1 has created objects in your world

which you can manipulate. I discuss Step 2 in more detail in Section 2.7.

2.7 STEP 2: ORGANIZE YOUR TOPICS

Manipulating index cards is embodied thinking.

Table 2-2B. The scaffolding steps covered so far in Chapter 2.

Step 1	Generate broad topics: Write each topic on the blank side of an index card
Step 2	Organize and evaluate topics: Manipulate index cards into a plausible order

In Step 2 (Table 2-2B), you think about, evaluate, and organize the topics you created in Step 1. Importantly, in Step 2 you organize and evaluate *general* topics. You have not yet created paragraph-level topics, or written complete sentences, because you want to protect your creativity from premature evaluation (Elbow, 1981).

How do you think about, organize, and evaluate topics in Step 2? You exploit the index card affordances described in Chapter 1. You think about your topics by reading and rearranging the index cards you created in Step 1. Remember, you perform embodied thinking when you rearrange index cards.

The main activity you perform in Step 2 involves

grouping related topic cards together. You think about relationships between topics you have written on different cards. If you see a relationship between two cards, then you place them closer together in your display. By arranging cards according to their relatedness, you scaffold your thinking about topics, because you need not think about or remember relations between topics. Instead, you simply see topic relations revealed by card positions.

As you read and rearrange cards, you engage in another important activity: removing redundancies. You often discover two different cards express nearly identical topics. When you find such redundancy, you might throw one card away, removing it from your display. Or you might replace both cards with a new card whose topic encompasses the two redundant topics. Embodied thinking includes removing redundant index cards!

Another important activity in Step 2 involves arranging your topic cards in order. Remember your manuscript must present topics in a meaningful order, so, you begin to look for a narrative structure by rearranging your cards. Perhaps you will arrange topics in order from general to specific. Perhaps you will recognize your topics require a logical order, because your reader won't understand Topic Y unless you have already written about Topic X. The index card scaffold provides one advantage: you can rearrange cards as you consider different topic orders. By looking at one topic order on your

display surface you can see its advantages and disadvantages; you explore different orders by moving topic cards around.

Manipulating index cards while exploring topic orders typically requires arranging index cards on a two-dimensional surface. For instance, I often place similar cards beside each other horizontally on a display board. I represent topic order – moving from one general topic to the next – by placing adjacent topics near each other vertically on the display board.

As I develop my narrative structure by rearranging index cards on my display, I might realize I need another topic. I fix the problem by jotting down the missing topic on a new index card which I then add to my display. Also, I often become aware my topics belong to different sections or subsections of a manuscript. My display lets me arrange my cards so I can easily see my different sections. I often add new cards – usually cards which have a different color from my topic cards – upon which I jot section or subsection titles.

In summary, in Step 2 you think about the topics you generated in Step 1 by rearranging index cards on a display surface. By having the cards available in your world, you need not remember topics, because you can read them. By following a system when you rearrange your cards in two dimensions (e.g., by placing similar cards near one another, by developing topic narratives in columns, by using different colored cards to represent different things such as topics vs titles) you can see your manuscript's narrative structure take shape before your eyes.

2.8 STEP 3: ENHANCE YOUR EXISTING TOPICS

Can you express a topic as an organized set of subtopics?

Table 2-2C. The scaffolding steps covered so far in Chapter 2.

Step 1	Generate broad topics: Write each topic on the blank side of an index card
Step 2	Organize and evaluate topics: Manipulate index cards into a plausible order
Step 3	Enhance existing topic cards: Convert broad topics into finer detailed subtopics and write each subtopic on its own index card

In the 1982 film *Blade Runner*, protagonist Rick Deckard scans a grainy photograph into an ‘Esper’ supercomputer. The computer displays the photo on a screen while superimposing a grid. Using voice commands, Deckard explores what he sees: “Track 45 right. Stop. Center and stop.” The computer alters the display accordingly. Seeing something interesting, Deckard commands the machine to add details to a particular region: “Enhance 34 to 36.” The machine adds missing details while increasing the image’s resolution.

Analogous to Deckard’s approach, in Step 3 (Table

2-2C) you try to enhance each topic card you generated and organized in steps 1 and 2. You do so by replacing each topic card with a set of cards; each new card expresses finer details about the original topic. In Step 3, you analyze each topic card into several more specific subtopic cards.

How do I enhance my topics? I answer a basic question: what subtopics do I need to write about to communicate my broader topic? I write each answer down on a new index card. I then add my new cards to my display. I might replace the original topic card with the subtopics cards I generated to it. I might keep the original card in place and position the subtopics cards nearby to make my topics' hierarchical structure visible.

I created the current chapter by enhancing my Step 2 scaffold for it. In my Step 2 scaffold for Chapter 2, an index card read 'Describe my method'. Step 3 enhanced the topic card, replacing it with nine different cards, each naming one method step. Section titles in Chapter 2 (Step 1, Step 2 and so on) provide what I wrote on each new subtopic card when I enhanced the broad topic.

As you create new subtopic cards during Step 3, you also think about their organization. You must ask 'What subtopics order best communicates my broader topic?' as soon as you have some new cards. Remember you conduct embodied thinking when you rearrange cards; place your new cards in a plausible location when you add them to your scaffold.

In *Blade Runner*, Deckard enhances the photograph

recursively. After finding something new, he commands Esper to further enhance an already enhanced image. When you evaluate subtopic card order, you also enhance recursively. As you consider each subtopic, you can ask whether you can also express the subtopic using ordered sub-subtopics. If so, then you replace the subtopic with new index cards, each expressing a sub-subtopic. Ideally such recursive enhancing continues until you feel each final topic you have can be expressed with a single paragraph (Figure 2-5).

As you enhance your original topic cards, you can also add cards to represent other manuscript components. For instance, during Step 3 you could add references you want to cite, figures or tables which you want to include, or quotes you want to use. You represent each additional item on its own index card; I use different colored cards to distinguish cards representing other components from cards representing topics or subtopics.

Your scaffold aids all your Step 3 activities. You don't have to remember your topics because you can read them. You remember your new subtopics using index cards; you think about subtopics by moving cards around in space. Step 3 elaborates the scaffold, but the scaffold supports its own elaboration.

Of course, Step 3 cannot go on indefinitely. At some point you must stop breaking topics into subtopics. When do you stop enhancing the scaffold?

2.9 STEP 4: DEVELOP YOUR PARAGRAPH TOPICS

Can a single paragraph communicate a subtopic?

Table 2-2D. The scaffolding steps covered so far in Chapter 2.

Step 1	Generate broad topics: Write each topic on the blank side of an index card
Step 2	Organize and evaluate topics: Manipulate index cards into a plausible order
Step 3	Enhance existing topic cards: Convert broad topics into finer detailed subtopics and write each subtopic on its own index card
Step 4	Convert subtopic cards into paragraph topic cards: Each index card represents a paragraph's topic

Step 3 can produce many subtopic cards as you enhance your broad topics. When do you stop generating subtopics? To answer, remember a paragraph serves as your manuscript's basic element because a paragraph expresses a single topic.

Figure 2-1 in Section 2.2 illustrated a *good* paragraph. I

can modify Figure 2-1 to define a *bad* paragraph (Figure 2-7): a bad paragraph attempts to communicate more than one topic.

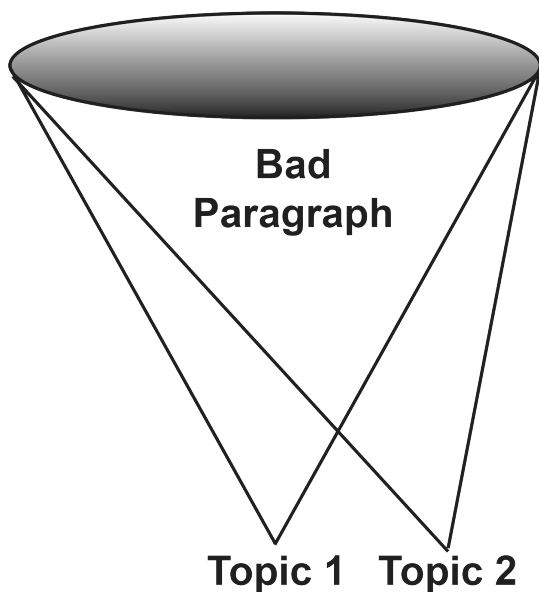


Figure 2-7. A bad paragraph attempts to communicate multiple topics.

In Step 4, you work to convert each subtopic into a paragraph topic. In Step 4, you aim to express each subtopic using a good paragraph – a paragraph which communicates one, and only one, topic. You stop enhancing topics (Step 3) when you can express each subtopic with a single paragraph.

You conduct Step 4 as follows. For each subtopic card in your display, you ask whether you can communicate the subtopic with one good paragraph. If you answer ‘Yes’, then you use the subtopic card as a paragraph topic. If you answer

‘No’, then then you return to Step 3 to break the subtopic down into sub-subtopics, intending to express each sub-subtopic using a single paragraph (Figure 2-8).

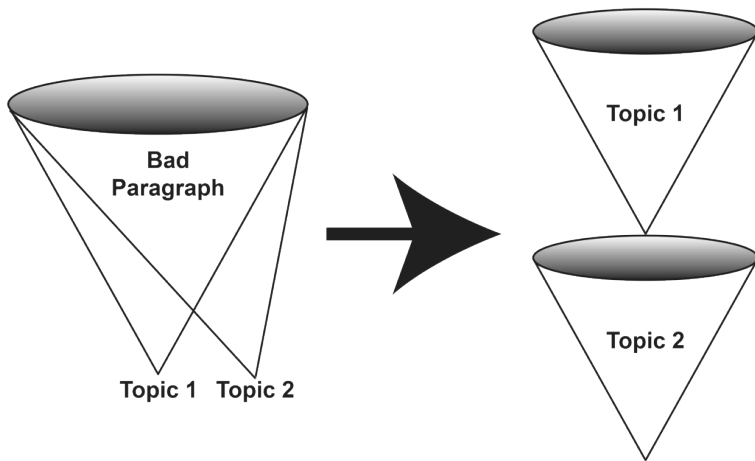


Figure 2-8. If you cannot express a subtopic from Step 3 using a single paragraph, break it into two expressible sub-subtopics.

To succeed, Step 4 requires you to have aimed with Step 3 to make each index card’s topic strongly related to a paragraph. If Step 3 succeeded, you usually find you expressed each subtopic with one paragraph. However, if in Step 4 you find one card requiring two or more paragraphs to express its topic, then you have started Step 4 prematurely. You should return to Step 3 and spend more time enhancing your topic cards.

Thus, a strong relationship exists between Steps 3 and 4 – you could consider Step 4 as Step 3’s final stage. You could

combine the two steps by stating “In Step 3, recursively enhance each topic card until each subtopic card can be expressed by one paragraph which communicates a single topic.”

Step 4 generates index cards which provide the paragraph topics ‘spine’ on the left of Figure 2-5. At the end of Step 4, you possess a complete topics chain for your manuscript, paragraph by paragraph. Each index card in your display holds one paragraph’s topic. However, you still have work to do. You need to confirm you organized your paragraph topics properly. To do so, you proceed to Step 5.

2.10 STEP 5: ORGANIZE YOUR PARAGRAPH TOPICS

Can I improve the order of my paragraph topics?

Table 2-2E. The scaffolding steps covered so far in Chapter 2.

Step 1	Generate broad topics: Write each topic on the blank side of an index card
Step 2	Organize and evaluate topics: Manipulate index cards into a plausible order
Step 3	Enhance existing topic cards: Convert broad topics into finer detailed subtopics and write each subtopic on its own index card
Step 4	Convert subtopic cards into paragraph topic cards: Each index card represents a paragraph's topic
Step 5	Organize your paragraph topics

When you complete Step 4, your scaffold displays paragraph topic cards which define your manuscript's detailed structure. But you cannot yet begin to write. Before you write, you should reconsider and confirm the order of your para-

graph topics. In Step 5 (Table 2-2E) you evaluate your paragraph topics to ensure your topics proceed in order.

How does Step 5 proceed? You can quickly get a manuscript's main drift by only reading each paragraph's first and last sentences (Sarnecka, 2019). The first and last sentences in a paragraph communicate the paragraph's topic provide links to create a narrative structure. Unfortunately, your scaffold does not – yet – contain any sentences. However, your scaffold does contain the next best thing, the paragraph topics – topics which topic sentences and concluding sentences will convey. So, in Step 5 you read and evaluate your paragraph topics.

You use your scaffold to support Step 5. You need not remember paragraph topics or their order. Instead, you read the paragraph topic you have written on each index card, moving in the order in which you arranged your cards on the scaffold. As you read the paragraph topics, you evaluate them. Does your paragraph topic order make sense? Does your paragraph topic order properly communicate your manuscript's main points?

In Step 5, you continue to manipulate the scaffold, particularly if you feel you can improve paragraph topic order. You might move index cards around to explore alternative topic orders. You can also evaluate positions of index cards representing section titles. In short, you use your scaffold to satisfy yourself with your manuscript's structure. You use your scaffold to evaluate your manuscript's structure even though you have not yet written any sentences.

During Step 5, you also continue to critically evaluate paragraph topics. Perhaps you feel your structure would improve by inserting another paragraph topic. To do so, you add a new paragraph topic card and place your new card in the desired spot in your display. If you detect redundancy in paragraph topics, you fix the problem by either removing a paragraph topic card or by replacing two cards with one which more clearly expresses the topic while removing redundancy.

When you finish Step 5, you have chosen your manuscript's paragraph topic order. You might now write a number on each card to indicate the order in which card topics will appear in your manuscript; you will soon remove them from your display and place them (in order) in a deck. Your ordered deck provides a 'spine' of detailed topics for your manuscript – the topic chain on the Figure 2-5's left.

Step 6, described in the next section, turns to building Figure 2-5's right side. You can finally start writing. You will see, though, all your scaffolding work makes writing your first sentences much easier.

2.11 STEP 6: WRITE YOUR TOPIC SENTENCES

Write a good sentence which states a paragraph’s topic.

Table 2-2F. The scaffolding steps covered so far in Chapter 2.

Step 1	Generate broad topics: Write each topic on the blank side of an index card
Step 2	Organize and evaluate topics: Manipulate index cards into a plausible order
Step 3	Enhance existing topic cards: Convert broad topics into finer detailed subtopics and write each subtopic on its own index card
Step 4	Convert subtopic cards into paragraph topic cards: Each index card represents a paragraph’s topic
Step 5	Organize your paragraph topics
Step 6	Write a topic sentence for each paragraph topic: Write it on the reverse side of a paragraph topic’s index card

Adler and van Doren (1972) describe writing as beginning with a skeleton which the author works to hide. The author’s aim “is to conceal the skeleton artistically or, in other

words, to put flesh on the bare bones” (Adler & van Doren, 1972, p. 90). After Step 5, your scaffold provides a highly structured skeleton. You can now put flesh on the skeleton by converting your topics into a partial manuscript: your outline.

Because each of your topic cards now corresponds to one paragraph’s topic, Step 6 (Table 2-2F) can begin putting flesh on your manuscript’s skeleton by adding topic sentences to each paragraph. Step 6 proceeds as follows: You process your index cards – your paragraph topics – in order. For each card, you read the topic you wrote on its blank side. You then turn the index card over and write the topic sentence for the paragraph at the top of the card’s reverse side (see Figure 2-6). Remember, a topic sentence states the paragraph’s topic.

Because you worked hard creating your scaffold, you can easily accomplish Step 6. You do not have to think a topic up from scratch, because you have already generated topics while building the scaffold. Instead, you face a more straightforward task: writing a complete sentence which communicates your paragraph’s topic.

To illustrate how you can easily convert topics into topic sentences, Table 2-3 provides some examples for the current section. The first column lists the paragraph topics for Section 2.11’s first eight paragraphs. The second column provides the topic sentence I wrote on each paragraph’s index card. Note the topic sentences in the table may not match the topic sentences in the book perfectly, because my topic sentences evolve when I revise my manuscript.

Importantly, Step 6 *only* requires you to write a topic sentence for each paragraph topic; you need not do anything else. You must fight the urge to write a concluding sentence underneath the topic sentence. (You write concluding sentences in Step 7). You postpone writing concluding sentences because each concluding sentence must relate meaningfully to the topic sentence on the current card as well as to the topic sentence on the next card (Figure 2-3). Therefore, you need all your topic sentences in place before you begin to write concluding sentences.

In Step 6 you finally write sentences. How much effort should you put into writing excellent topic sentences? The quick answer: you should write the best topic sentences possible in a fairly *short* time. You should *not* waste time or effort trying to write *perfect* sentences. Remember, Step 6 creates your *first drafts* of topic sentences. You do not need to write “one true sentence” on each index card in Step 6.

Admittedly, if you write better topic sentences now, you will have less work to do later when you revise your writing. So, if you *can* write a good topic sentence, then you *should*.

How can you improve your sentence writing? I shy away from providing advice for improving your sentences for several reasons. First, I recognize I can always improve my writing, and do not feel comfortable offering much advice about writing quality. Second, many better writers than I have written excellent books about how to improve writing. I advise you to find a book about improving writing which resonates with

you. I have my own personal favorites (Sword, 2012, 2016; Zinsser, 2006). I consult these books frequently, try to take their advice to heart, and try to use the advice to improve my writing.

My preferred books about improving writing agree on certain best practices. I should use simple sentences, I should use active verbs, and I should place related nouns and verbs close together in a sentence. I suggest you find your own favorite books about improving writing and treat them as I treat mine.

I know I could write better first drafts of my own topic sentences (see Table 2-3 for evidence). I write less than perfect initial topic sentences because I try to complete Step 6 quickly. I do not worry about writing poor topic sentences in Step 6. Anne Lamott provides excellent advice when she reports “the only way I can get anything written at all is to write really, really shitty first drafts” (Lamott, 1995, p. 21). She likes the freedom which shitty first drafts provide, because they make writing easier, because no one ever needs to see or judge them, and because she knows she will rework the draft later. Following Lamott’s advice, I’m content to jot down really, really shitty topic sentences for all the reasons she provides. I feel even more liberated, perhaps, because my shitty topic sentences appear long before my really, really shitty draft takes form!

Remember, you don’t use Step 6 to produce perfect sentences; instead, you use Step 6 to add plausible flesh to your manuscript’s skeleton. Step 6 produces more of your outline

and does not require you to *really* write yet! When you complete Step 6, each paragraph's index card will have a topic written on one side and a topic sentence written on its reverse side. With all your topic sentences created, you can move to Step 7 to add more flesh to your skeleton.

Table 2-3. Examples of converting paragraph topics into topic sentences for the first eight paragraphs of Section 2.11 of the current book.

Paragraph Topic	First Draft of Topic Sentence
One view of writing – put flesh on skeleton	Adler and van Doren (1972) describe writing as beginning with a skeleton which the author works to hide.
Work with scaffold has produced a skeleton – start to flesh it out	Enhance existing topic cards: Convert Our scaffold now has taken the form of a highly structured skeleton.
Step 6 begins this – write topic sentence for each paragraph topic card	Step 6 begins the process of adding flesh to the skeleton we have created by adding topic sentences to each paragraph.
Having a scaffold makes Step 6 easier	Step 6 is easily accomplished because of all our earlier work creating our scaffold.
Example topic sentences from paragraph topics	Write a topic sentence for each To illustrate the ease of converting topics into topic sentences, Table 2-2 provides some examples for the current section of the book.
Do this for all cards, because we need all topic sentences before doing Step 7	Write a concluding sentence for each Importantly, Step 6 requires us to write a topic sentence for each paragraph topic; we need not do anything else.
Writing sentences now – how good should they be?	In Step 6 we are finally writing sentences.

Table 2-3. Examples of converting paragraph topics into topic sentences for the first eight paragraphs of Section 2.11 of the current book.

Quick answer – write best sentence possible without taking too much time	The quick answer to this question is we should write the best topic sentences we can in a fairly short period of time.
--	--

2.12 STEP 7: WRITE YOUR CONCLUDING SENTENCES

Link your concluding sentences to pairs of topic sentences.

Table 2-2G. The scaffolding steps covered so far in Chapter 2.

Step 1	Generate broad topics: Write each topic on the blank side of an index card
Step 2	Organize and evaluate topics: Manipulate index cards into a plausible order
Step 3	Enhance existing topic cards: Convert broad topics into finer detailed subtopics and write each subtopic on its own index card
Step 4	Convert subtopic cards into paragraph topic cards: Each index card represents a paragraph's topic
Step 5	Organize your paragraph topics
Step 6	Write a topic sentence for each paragraph topic: Write it on the reverse side of a paragraph topic's index card
Step 7	Write a concluding sentence for each paragraph topic: Write it on the reverse side of a paragraph topic's index card

In Step 6, you wrote a topic sentence on each para-

graph topic's index card. In Step 7, you add concluding sentences (Table 2-2G). Again, you exploit your scaffold because the topic sentences you created in Step 6 guide the concluding sentences you write in Step 7. Make sure you write the concluding sentence well below an index card's topic sentence, leaving space between them on the index card (Figure 2-6). You will use the space to add additional notes later in the scaffolding process.

Recall a concluding sentence relates to two different topic sentences: the topic sentence which begins the paragraph containing the concluding sentence and the topic sentence which begins the next paragraph (Figure 2-3). You need to establish the two links for each concluding sentence you write. Accordingly, you need to complete Step 6 before you begin Step 7: you cannot write concluding sentences without already having both topic sentences available.

How do you write a concluding sentence for a paragraph? First, you read the topic sentence which starts the paragraph. Second, you read the topic sentence which starts the next paragraph. Third, you write a sentence which relates to both topic sentences.

To link a concluding sentence to the topic sentence in the same paragraph, you write a sentence which restates the topic in different, usually more specific, terms. To link a concluding sentence to the next paragraph's topic sentence, you use the concluding sentence to segue to the next topic sentence. You can build a bridge between a concluding sentence

and the next topic sentence by ensuring both sentences make similar points; often, both sentences will share words.

To illustrate, Table 2-4 provides the concluding sentences for each paragraph presented earlier in Table 2-3. Note how the concluding sentences try to connect to two different topic sentences. For instance, the first concluding sentence restates the first topic sentence's point; the same concluding sentence links to the second topic sentence because both sentences use the word 'skeleton', and both refer to 'writing puts flesh on a skeleton'.

Again, you shouldn't spend much time trying to make your concluding sentences perfect. Jot them down quickly; shitty concluding sentences are just as liberating as are shitty topic sentences. You will clean up your mess when you revise and polish your manuscript.

Step 7 ends when you have a concluding sentence on each paragraph card in your scaffold. When you complete Step 7, you have finished building your outline's core (Figure 2-5). You possess a set of index cards; you have a paragraph topic written on one side of each card; you have a topic sentence and a concluding sentence written on the other side. After one more evaluative step, you can move your scaffold into a word processing document.

Table 2-4. Examples of pairs of topic sentences and concluding sentences for the first eight paragraphs of the previous book section.

First Draft of Topic Sentence	First Draft of Concluding Sentence
Adler and van Doren (1972) describe writing as beginning with a skeleton which the author works to hide.	The author’s aim “is to conceal the skeleton artistically or, in other words, to put flesh on the bare bones”.
Enhance existing topic cards: Convert Our scaffold now has taken the form of a highly structured skeleton.	We can now begin to put flesh on the skeleton by converting it into a partial manuscript: an outline.
flesh to the skeleton we have created by adding topic sentences to each paragraph.	Remember, a topic sentence’s purpose is to state the topic of a paragraph.
Step 6 is easily accomplished because of all our earlier work creating our scaffold.	Step 6 is easily accomplished because of all our earlier work creating our scaffold.
Write a topic sentence for each To illustrate the ease of converting topics into topic sentences, Table 2-2 provides some examples for the current section of the book.	Note the topic sentences in the table may not match the topic sentences in the section perfectly, because my topic sentences will change during multiple revisions of my manuscript.
Write a concluding sentence for each Importantly, Step 6 requires us to write a topic sentence for each paragraph topic; we need not do anything else.	We need to have our topic sentences in place before we begin to write concluding sentences.
In Step 6 we are finally writing sentences.	How good should our topic sentences be?

Table 2-4. Examples of pairs of topic sentences and concluding sentences for the first eight paragraphs of the previous book section.

should write the best topic sentences we can in a fairly short period of time.	We should not waste time or effort trying to write perfect sentences.
--	---

2.13 STEP 8: NOTATE YOUR SCAFFOLD CARDS

Use your index cards to scaffold important details.

Table 2-2H. The scaffolding steps covered so far in Chapter 2.

Step 1	Generate broad topics: Write each topic on the blank side of an index card
Step 2	Organize and evaluate topics: Manipulate index cards into a plausible order
Step 3	Enhance existing topic cards: Convert broad topics into finer detailed subtopics and write each subtopic on its own index card
Step 4	Convert subtopic cards into paragraph topic cards: Each index card represents a paragraph's topic
Step 5	Organize your paragraph topics
Step 6	Write a topic sentence for each paragraph topic: Write it on the reverse side of a paragraph topic's index card
Step 7	Write a concluding sentence for each paragraph topic: Write it on the reverse side of a paragraph topic's index card
Step 8	Notate and deck your index cards

With Step 7 complete, you can now add important finishing touches to your scaffold before moving the scaffold into a word processing document. You have finished your most crucial outlining work, but some additional notes can make your future work proceed more smoothly.

Your earlier steps scaffolded your outline by laying index cards out on a two-dimensional surface. Arranging cards on a surface aided memory (you could see any card with a glance) and organizing (you could think about your project by repositioning your cards). However, any writing project must eventually become a paragraph chain. You need to convert your two-dimensional card display into a more conventional order.

In Step 8 (Table 2-2H) you remove index cards from your layout; you place them in order as a single deck. To preserve the narrative, you produced using earlier steps, you must ensure you keep your cards in order when you build your deck. As you construct your deck, I recommend writing a number on each card to indicate card order. (You might number the cards with pencil if you feel you might still revise card ordering!) Numbering your cards prevents problems when a cat knocks over your deck and you must put your index cards back in order.

After ordering your index cards, you perform one final notating step. You take each card in turn, flip the card over, and read the two sentences written on the card's back. You deliberately left space between sentences (see Figure 2-6) to provide

room for a few short notes. Soon you will write supporting sentences to complete each paragraph (Chapter 4). You can help create supporting sentences later by jotting down notes on each card to remind you about what your supporting sentences might say. You might list in point form what each supporting sentence will say. You might list some material you want to cite, or a figure or table you need to refer to. You don't need to add notes to every card, but your notes can provide important reminders when you convert your outline into a draft.

With your numbered cards in a deck, with some notes written between topic and concluding sentences, you can move to your final step: moving your scaffold into a word processing document for developing your first draft.

2.14 STEP 9: MOVE YOUR OUTLINE INTO A COMPUTER

An outline in a word processing document is another scaffold.

Table 2-2I. The scaffolding steps covered so far in Chapter 2.

Step 1	Generate broad topics: Write each topic on the blank side of an index card
Step 2	Organize and evaluate topics: Manipulate index cards into a plausible order
Step 3	Enhance existing topic cards: Convert broad topics into finer detailed subtopics and write each subtopic on its own index card
Step 4	Convert subtopic cards into paragraph topic cards: Each index card represents a paragraph's topic
Step 5	Organize your paragraph topics
Step 6	Write a topic sentence for each paragraph topic: Write it on the reverse side of a paragraph topic's index card
Step 7	Write a concluding sentence for each paragraph topic: Write it on the reverse side of a paragraph topic's index card
Step 8	Notate and deck your index cards
Step 9	Move your scaffold ? your outline ? into a word processor

After Step 8, you have created a rich structure in the form of Figure 2-5. You have created a set of paragraph topic cards and have created the topic sentence and concluding sentence for each paragraph. You now need to convert your index card structure into another form: an outline in a word processing document.

Step 9 (Table 2-2I) creates your outline. You open your word processor and type the information from your index

cards. You take the first index card from your deck and look at the side which holds the topic sentence and the concluding sentence. You type the two sentences into the same paragraph in your word processing document. If you have additional notes between the sentences, then you type them into the middle of the paragraph in the same format they have on the card. For instance, if you have a couple of point form reminders, then you type them in as incomplete sentences between the paragraph's topic sentence and concluding sentence.

You create an outline 'paragraph' for every index card in your deck. When you have finished typing, you possess a rich outline for your manuscript. Converting the outline into a draft requires you to add supporting sentences to each paragraph in your outline. But you add supporting sentences after you complete Step 9.

Importantly, Step 9 produces an outline while avoiding the blank page. I often find myself astounded at how many words I already have in my document when I finish Step 9. By putting effort into outlining a manuscript, while easing my effort by using a scaffold, my writing task now seems much less intimidating.

My outline's detail also makes writing seem easier after Step 9. I have created a narrative structure, I have represented the structure as a chain of paragraph topics, and I have created the two most important sentences for each paragraph. To convert my outline into a first draft, I only need to add supporting sentences to each paragraph.

My outline simplifies my adding supporting sentences because I have already created topic and concluding sentences. Supporting sentences provide data for, or elaborate, a paragraph's topic. The topic and concluding sentences already express each topic. I usually find having the two most important sentences in hand makes writing my supporting sentences straightforward.

When I add my supporting sentences, and therefore have converted my outline into a first draft, I still have hard work to do. Good writing requires me to repeatedly revise a document, simplifying and improving each sentence again and again. Ruby Wiebe believes “writing better is a never-ending possibility” (Wiebe, 2016, p. 23). A later chapter in the current book suggests additional scaffolds to aid revising a manuscript.

2.15 INDEX CARDS – AND BEYOND

Table 2-5. The scaffolding method discussed in Chapter 2.

Step 1	Generate broad topics: Write each topic on the blank side of an index card
Step 2	Organize and evaluate topics: Manipulate index cards into a plausible order
Step 3	Enhance existing topic cards: Convert broad topics into finer detailed subtopics and write each subtopic on its own index card
Step 4	Convert subtopic cards into paragraph topic cards: Each index card represents a paragraph's topic
Step 5	Organize your paragraph topics
Step 6	Write a topic sentence for each paragraph topic: Write it on the reverse side of a paragraph topic's index card
Step 7	Write a concluding sentence for each paragraph topic: Write it on the reverse side of a paragraph topic's index card
Step 8	Notate and deck your index cards
Step 9	Move your scaffold ? your outline ? into a word processor

Chapter 2 presented a method for scaffolding writing. The method uses index cards to build a scaffold which provides a manuscript's detailed outline. Table 2-5 lists the method's nine steps for your ease. Appendix I presents an example scaffold for a short essay.

The Table 2-5 method uses index cards to scaffold writing. Other writers have also proposed using index cards to help outline a writing project (Atchity, 1995; Butler & Burroway, 2005). How does the method described in Chapter 2 relate to other methods?

Let us focus on one method which has been described in detail (Atchity, 1995). Atchity's primary focus involves helping writers manage their time, and index cards play an important role. His 'card system' is used to scaffold a writer's initial research phase; index cards produced from such research are then used to organize the structure of a manuscript and scaffold its first draft.

Atchity's (1995) card system uses 5" X 7" index cards. The system begins with a large collection of blank cards; the number of cards in the collection is based on the desired length of the manuscript to be written. Atchity suggests each page of the manuscript will be scaffolded by four index cards.

The first phase of Atchity's (1995) card system involves library research, which begins with a planned usage of only half of the blank cards. What does the researcher record on their blank cards? The content can be quite varied. "Some cards will be quotations you read, copied out word for word;

others will be summaries of what you read; still others will be your own ideas that occur to you while browsing, ideas about your subject or your book's structure" (Atchity, 1995, p. 78).

When library research is complete research along different lines can proceed. Atchity (1995) suggests conducting interviews as one kind of research; if interviews are conducted, then blank cards will be filled with a participant's main points.

The next phase of Atchity's (1995) card system involves evaluating the index cards produced by the research. Evaluation begins by going through the collected cards one by one, deciding whether to keep the card or not. "Is this card good or not? Your aim is to throw away the ones that aren't" (Atchity, 1995, p. 82). The goal of evaluation is to remove cards.

The next phase of Atchity's (1995) card system is to organize the cards which remain after evaluation. You read the cards and sort them in a 'natural order' based on card content. You take all the remaining cards and read the first. It is placed in its own pile. You then read the next card to decide whether it belongs to an existing pile or should instead be placed in a new pile. You continue until all your research cards have been assigned to piles. Later, you repeat the process, moving pile to pile, considering whether a card should stay in a pile or be moved to another. In general, this phase groups content-related cards into the same pile. Atchity aims to use each pile to create one section or chapter of the to-be-written manuscript.

With research cards assigned to content-related piles,

the next phase is to process each pile. One activity in this phase is to reorder cards in a pile, sorting a pile's cards into a natural order based on their content. Another activity is to sort the order of the piles to plan the structure (e.g., chapter order) of the manuscript.

The final phase of Atchity's (1995) card system is to use the processed piles of cards as a scaffold for the first draft. Each pile has been sorted in a plausible order, and Atchity expects each manuscript page to be based on four research cards. So, writing the first draft means beginning with the first four cards from the first pile. You write what you need from Card 1 and then move on to the next card.

While the sorted research cards scaffold the first draft, Atchity (1995) is comfortable with abandoning them as writing proceeds. "Don't worry if you stop relying on the cards at any point; remember, they're only crutches: the real book is inside you" (Atchity, 1995, p. 86).

On the one hand, Atchity's (1995) card system takes advantage of many affordances offered by index cards and used in my method (Section 1.8). Index card 'writability' and 'portability' are important for conducting research. 'Readability' and 'expendability' are critical when you decide whether to keep a research card. 'Arrangeability', 'groupability' and 'readability' permit you to create piles of research cards organized into some natural order.

On the other hand, important differences exist between the two methods. In Atchity's (1995) card system

cards do not correspond to paragraphs but instead serve as research notes. All cards cannot be viewed simultaneously, but are instead examined in succession, because they are collected in stacks or piles. Cards are not used to directly scaffold paragraphs (because cards are not the source of topic or concluding sentences). Cards offer weak scaffolding, because they can be ignored when using card material to inspire writing of the first draft. Finally, cards help structure the order of contents, but do not alleviate the ‘myth of inspiration’ if creating the first draft requires the writer to produce new, full sentences for the first time.

From the perspective of embodied cognition, differences between my method and Atchity’s (1995) card system are not important. I believe my method is better suited to academic writing, but others might find Atchity’s more appropriate for their writing needs. The methods may differ, but both indicate index cards can scaffold writing.

However, other scaffolds are available. For instance, the Table 2-5 method requires you to begin by generating broad topics. To many writers, generating topics might offer as great a challenge as facing the blank page. Chapter 3 offers other scaffolds to help you generate topics when you begin your scaffold with Step 1 of the Table 2-5 method.

PART III

CHAPTER 3: SCAFFOLDING INITIAL TOPICS

One scaffold may bootstrap another.

Chapter 2 described how you can scaffold a rich outline for a writing project. The scaffold makes writing easier by rejecting the myth of inspiration. However, the Chapter 2 scaffold begins when you generate broad topics to launch the scaffold. Doesn't your first step require inspiration? Chapter 3 describes how you can decrease relying on inspiration when you generate your initial topics. Chapter 3 does so by providing additional scaffolds to support how you generate the broad topics from which your outline begins.

3.1 FACING THE BLANK SET OF TOPICS

Use one scaffold to create another.

Chapter 2 introduced a scaffold designed to avoid a key challenge: facing the blank page. The scaffold begins when you generate ‘seeds’ from which your outline grows (Section 2.6). Each seed provides a topic, written on its own index card to move topics from inside your mind to your world outside. Inspired by embodied cognition, writing becomes externalized thinking; you think by manipulating your index cards to develop a desired narrative.

The scaffold requires you to create seed topics from scratch in Step 1 (Section 2.6). Where do your seed topics come from? If creating seed topics forces you to face the blank page, then the scaffold’s goal is not achieved.

Chapter 3 provides different approaches to creating broad topics when you begin to create your scaffold. Each approach provides *another* scaffold for generating broad topics when your writing project begins. In short, Chapter 3 introduces additional scaffolds to help you create the scaffold described in Chapter 2.

All the scaffolds I describe in Chapter 3 treat creating topics as answering basic questions. I view the methods as

scaffolds because I provide you with the basic questions in advance. The scaffolds differ from one another by using different sources for the questions. Chapter 3 describes seven different scaffolds to help you generate broad topics to launch your outline. I begin by showing how a discipline's conventions often provide basic questions you can use to launch your outline's scaffold.

3.2 QUESTIONS FROM CONVENTIONS

Follow a no-fail recipe.

When scientists first started writing about their discoveries, no conventions existed to dictate what form their writing should take. In the 17th century, the first scientific journals published papers written as chronologically structured narratives, or even as descriptive letters (Sollaci & Pereira, 2004; Wu, 2011). Only in the 20th century did scientific journals begin to publish articles which followed a conventional structure.

Nowadays, we expect to read scientific articles organized according to a discipline's conventions. Conventions make articles easier to understand, because conventions satisfy readers' expectations about how researchers present information. Scientific writers follow accepted practices or conventions to increase the grasp – and impact – of their work. In Section 3.2 I describe how an important modern convention can scaffold creating topics.

By the 1940s, scientific writing in many disciplines began using a particular convention: the Introduction, Method, Results and Discussion (IMRaD) structure. By 1965, IMRaD had become the dominant convention, and now

serves, by far, as scientific writing's most common structure (Sollaci & Pereira, 2004).

IMRaD provides a logical narrative for a scientific paper. The introduction lays out a problem, states the problem's importance, relates what researchers know (or don't know) about the problem, and introduces research questions, hypotheses, and objectives. The method section describes how the writer studied their research questions. The results section conveys what the research methods found. The discussion interprets what the results mean, and states why readers should consider results as important.

Why has IMRaD become the dominant convention for scientific writing? First, IMRaD provides essential structure for successfully disseminating scientific research. "Most, if not all, editors and scientists agree that IMRaD provides a consistent framework that guides the author to address several questions essential to understanding a scientific study" (Wu, 2011, p. 1347). Second, you can apply IMRaD across diverse disciplines. Not surprisingly, many books about scientific writing discuss IMRaD (Heard, 2022; Sarnecka, 2019; Schimel, 2012; Sword, 2012).

One often finds IMRaD promoted within disciplines. For example, the 7th edition of the *Publication manual of the American Psychological Association* (APA) (American Psychological Association, 2020) provides conventions for psychologists to use to report research. Within the APA manual one finds an extremely detailed IMRaD structure. Table 3-1 of

the APA manual (pages 77-79) provides a long list of IMRaD components, including the structure for each IMRaD component and detailing the topics and subtopics to include in each section.

For example, the APA's Table 3-1 breaks the introduction into discussing three core subtopics: 'Problem', 'Review of Relevant Scholarship', and 'Hypothesis, Aims, and Objectives'. For the latter subtopic, the table (American Psychological Association, 2020, p. 78) instructs the writer to "state specific hypotheses aims and objectives including theories or other means to derive hypotheses, primary and secondary hypotheses, other planned analyses."

Some writing aids provide prompts to help overcome writer's block (Goldberg, 2021; Lamott, 1995). I follow suit by converting the statements of the APA manual's Table 3-1 into questions to use as prompts which appear (coincidentally) in my own Table 3-1 below. Answers to my IMRaD questions become your first topic cards (Step 1 in the Chapter 2 method). By giving you the IMRaD questions in advance, I provide you with a scaffold you can easily use to generate topic cards. The questions serve as a scaffold because you don't need to keep the questions in mind; you simply look at them and answer them, jotting the answers down on index cards.

I only provide questions for the APA's core IMRaD sections, ignoring other components (title page, author notes, abstract). Even so, the scaffold provides 104 different questions to use as prompts, because the APA manual's IMRaD struc-

ture is very detailed. I number each question to indicate its order relative to the table in the *Publication manual*. I also label each question to indicate the IMRaD topic and subtopic being addressed, using the terms provided in the *Publication manual's* table.

Why convert the *Publication manual's* statements about content into questions? To me, questions provide the most useful prompts for generating topics. First, most scientific articles aim to answer questions. Having questions in hand brings a paper's purpose to the foreground. Second, short answers to the questions almost always provide useful topics. Third, at times one question may generate more than one answer, and therefore efficiently generates more than one topic.

Table 3-1 below provides all my IMRaD questions. The table of questions *is* the scaffold, stepping away from exclusively using index cards to scaffold a paper's outline.

However, it can be convenient to convert the Table 3-1 scaffold into IMRad format. I print my own version of the table on labels, with one question per label. I then stick each label on its own index card, creating a deck of question cards to scaffold how I create topics. To use the deck, I take each prompt card in order. For each question, I generate one or more answers. I make each answer brief – a short note – and I write my answer on a new index card. If I feel a particular question is not relevant to my manuscript, I put the prompt card aside. When I finish answering the questions, I have a set of

index cards, each containing a short answer to a question, completing Step 1 of the Chapter 2 method.

When I choose to translate my Table 3-1 into index card form, I'm creating index cards which provide a different kind of scaffold than the index cards create in Chapter 2. First, I do not discard my IMRaD prompt cards because I can use them to prompt topics for many different papers. Second, I do not write on the prompt cards. I simply read them and write the answers on new index cards.

Table 3-1 which follows is a perfectly good scaffold on its own. Why might I convert it into a set of index cards? First, I can use a prompt card in public when I need a group of collaborators to collectively begin generating topics for a manuscript. Second, sometimes work proceeds more rapidly when Table 3-1 prompts are considered out of order (see Section 3-5 below). I find it easier to do so by breaking Table 3-1 questions into individual index cards.

The process I used to convert the APA's IMRaD structure into prompts could be applied to other conventions. The example described in Section 3.2 comes from a table in the APA publication manual which provides the structure of a paper for describing a study which relies on quantitative analyses. The same manual provides another table which structures a paper which describes qualitative results, and still another table which describes a combination of quantitative and qualitative analyses. Both tables could be converted into topic scaffolds by having their contents converted into prompt

questions. Readers with interests in other disciplines can find guidelines for reporting their research and convert their guidelines into topic prompts as well.

Table 3-1. Questions for generating initial topics for a manuscript. Each question is labeled with a number indicating its relative position in Table 3-1 of the Publication Manual of the American Psychological Association.

APA IMRAD Questions to Scaffold The Introduction (Questions 1 to 18)

- | | |
|---|--|
| 1. How should my paper open to get my reader interested? | 2. What big question is my paper about? |
| 3. Why is it interesting or important to answer my big question? | 4. What have previous papers said about my big question? |
| 5. What previous answers exist for my big question? | 6. What problems exist with the existing responses to my big question? |
| 7. What gaps remain even with existing research on my big question? | 8. What new questions have arisen from the existing research on my big question? |
| 9. What more specific previous question is addressed in my current paper? | 10. What possible answers might there be to my more specific question? |
| 11. What are my hypotheses? | 12. What are my predictions? |
| 13. What is my theory related to my big question? | 14. What is the specific purpose of my research which I report in my paper? |
| 15. What is the main point I will make in my paper? | 16. How does my paper proceed? |
| 17. Do I need any figures in my introduction? | 18. Do I need any tables in my introduction? |

APA IMRAD Questions to Scaffold The Method (Questions 19 to 68)

Table 3-1. Questions for generating initial topics for a manuscript. Each question is labeled with a number indicating its relative position in Table 3-1 of the Publication Manual of the American Psychological Association.

19. Did I use any criteria to include or exclude participants from my study?	20. Who were my participants?
21. How many participants were in my study?	22. How did I choose my participants?
23. How did my participants give their informed consent?	24. How many participants agreed to participate?
25. Where did I collect my data from participants?	26. How were participants rewarded for participating?
27. How did I get ethics approval for studying my participants?	28. What was the intended sample size of my study?
29. How did I determine my sample size?	30. Did I achieve my intended sample size?
31. What is the power or precision of my results with my sample size?	32. What stimuli did I use in my study?
33. How or why did I choose the stimuli for my study?	34. Did I use any standardized measures of participants or their behavior?
35. What steps did a participant go through when in my study?	36. How was the order of participant steps determined?
37. What instructions were given to participants?	38. How are participants debriefed at the end of the study?

Table 3-1. Questions for generating initial topics for a manuscript. Each question is labeled with a number indicating its relative position in Table 3-1 of the Publication Manual of the American Psychological Association.

39. Where participants run individually or in groups?	40. Did I train people to improve the quality of my measurements?
41. What was the reliability of my measurements?	42. Did I take multiple measures to improve my measurements?
43. What equipment did I use to present stimuli to participants?	44. What equipment did I use to measure participant responses?
45. Were my participants aware of which condition they were participating in?	46. Were researchers aware of which condition a participant was in?
47. Why was masking important in my study?	48. How did I accomplish masking in my study?
49. How successful was the masking in my study?	50. What was the reliability of my measurements?
51. What was the convergent validity of my study?	52. What was the discriminant validity of my study?
53. What was the interrater reliability of my study?	54. What was the test-retest reliability of my study?
55. What was the design of my study?	56. What were my independent variables?
57. What were my dependent variables?	58. How did I assign my participants to conditions?
59. How many conditions did my participants participate in?	60. How did I decide to exclude any participants after collecting data?

Table 3-1. Questions for generating initial topics for a manuscript. Each question is labeled with a number indicating its relative position in Table 3-1 of the Publication Manual of the American Psychological Association.

61. How did I decide to deal with missing data?	62. How did I identify statistical outliers?
63. What were the properties of my data distributions?	64. How did I transform my measurements?
65. What was my approach to performing inferential statistics?	66. How did I protect against experiment-wise error for my hypotheses?
67. Do I need any figures in my method section?	68. Do I need any tables in my method section?
APA IMRAD Questions To Scaffold The Results (Questions 69 to 84)	
69. What steps did a participant go through when in my study?	70. How was the order of participant steps determined?
71. What instructions were given to participants?	72. How many participants were in each group at each stage of the study?
73. How are participants debriefed at the end of the study?	74. Were participants run individually or in groups?
75. What were the dates for my period of recruiting participants?	76. What were the dates for performing repeated measures?
77. What were the dates for performing follow-up measurements of participants?	78. How did I preprocess my data?

Table 3-1. Questions for generating initial topics for a manuscript. Each question is labeled with a number indicating its relative position in Table 3-1 of the Publication Manual of the American Psychological Association.

79. How did I summarize my data?	80. What are the general characteristics of my data (descriptive statistics)?
81. What statistical analyses did I perform on my data?	82. Why did I perform the statistical analyses I perform?
83. Do I need any figures in my results section?	84. Do I need any tables in my results section?
APA IMRAD Questions To Scaffold The Discussion (Questions 85 to 104)	
85. Should I remind my reader about my research question?	86. Should I summarize the method I used to study my question?
87. Should I remind my reader about my predictions?	88. Should I remind my reader about my theory?
89. What is my summary of my main findings?	90. How do my main findings relate back to my predictions?
91. How similar are my results to previous results?	92. What is the take home message of my results?
93. What is surprising about my findings?	94. How well will my results generalize to other settings or stimuli?
95. What are the implications of my main findings?	96. What theories do my findings support?
97. What theories do my findings refute?	98. What new theory is needed to explain my findings?

Table 3-1. Questions for generating initial topics for a manuscript. Each question is labeled with a number indicating its relative position in Table 3-1 of the Publication Manual of the American Psychological Association.

99. Why are my findings important?	100. What new questions do my findings raise?
101. Can my findings be applied to new areas?	102. What new benefits or treatments are suggested by my findings?
103. Do I need any figures in my discussion?	104. Do I need any tables in my discussion?

3.3 FOLLOW A MODEL

Don't reinvent the wheel.

Table 3-1 shows conventional structures, like IMRaD, provide a scaffold for creating initial topics for your outline. The IMRaD structure's scaffold offers the strength of generality: you can use the scaffold to seed topics for writing projects belonging to different disciplines, to seed topics for different projects belonging to the same discipline, or to seed topics for different projects from the same research lab.

However, the IMRaD scaffold's strength also creates a problem: projects from one lab, or from one researcher, might require more focus because many Table 3-1 questions do not relate to projects from a particular lab or to a researcher who requires a more customized scaffold.

Where might you find a more customized topics scaffold? A model manuscript provides one potential source. I define a model manuscript as a paper strongly related to your research, as well as a paper you consider well-crafted or well-written. You can use your model manuscript to create a customized topics scaffold.

How do you derive a topics scaffold from a model manuscript? You reverse engineer the model manuscript along the lines illustrated earlier in Table 2-1. You work through

the manuscript paragraph by paragraph, identifying each paragraph's topic sentence and concluding sentence. You then use the sentences to infer each paragraph's topic. Table 2-1 illustrated such reverse engineering for the introductory paragraphs of one paper. An analysis of a model paper requires reverse engineering all paragraphs in the model paper.

To create a custom scaffold, you then write a question for each paragraph topic you have reverse engineered. You design the question to elicit the paragraph's topic. Your questions provide a more customized scaffold; personal prompts for initial topics for a paper to emulate your model manuscript.

When I was a student, my model paper was written by two of my psychology professors, describing how mental imagery affected concept learning (Katz & Paivio, 1975). I turned to the Katz and Paivio paper as a model because I admired its writing — I wanted my own writing to exhibit similar craftsmanship. As a junior faculty member, I turned to a different model paper which described a new learning rule for artificial neural networks (Rumelhart et al., 1986). Once again, I chose the model paper because I admired its writing. I also chose it because my research interests had changed, and my new model was directly related to my new research.

If you decide to use a model paper to serve as a scaffold for generating topics, then your choice of model will be personal. You need to choose a well-written paper related to your

work, which you will reverse engineer to produce questions which prompt topics for a paper you want to write.

3.4 USE YOUR OWN WORK TO SCAFFOLD

If it ain't broke, then don't fix it.

Someone just starting in a research field might find the previous two sections' methods for scaffolding topics very useful. A new researcher benefits from paying attention to a discipline's conventions, or to a model manuscript. However, a more established researcher can rely on their own work to create a customized scaffold; one scaffold to use to launch several different writing projects. Established researchers often work on several related projects at the same time. By exploiting project similarity, you can develop a topics scaffold for repeated use.

A customized, reusable topics scaffold works well in a lab which explores one broad research topic, attempting to answer various questions related to the broad topic. Researchers in one lab might use different methods to explore different facets of the broader topic. They might also use similar methods to answer different questions inspired by the broad topic. Project similarities permit one topic scaffold to apply to several different projects carried out in one lab.

How might you create a tailored, reusable topics scaffold for yourself or for your lab members? You can use one or

more successful writing projects of your own to provide questions for scaffolding topics.

I will use my own lab's work to illustrate. My lab uses specific computer simulations, artificial neural networks, to provide a bridge between formal music theory and musical cognition. My lab members train networks to solve musical problems (e.g., to classify musical chords into particular chord types). They then examine the trained networks to determine how the networks solve musical problems, discovering new ideas for music theory, for musical cognition, or for finding relationships between these two research areas. When we write our research up, we try to describe the structure we discover in our networks, and to the importance of our discovery for music theory or musical cognition.

My lab has several publications related to our 'musical networks' project (Dawson, 2018; Dawson et al., 2020; Perez et al., 2023). Looking over the published papers, I see an emerging pattern. Our papers start by noting that many researchers use musical networks, but do not interpret network structure. The papers then argue many researchers mistakenly ignore network structure, because only by understanding network structure can a network contribute to theory. We then state our goal: to provide a case study to illustrate our approach, and to show the importance of understanding network structure. The papers then detail the case study, as well as the structure discovered in our trained net-

works. The papers end by noting the importance of the discovered structure.

I can convert the formula I see emerging from my lab's publications about musical networks into a series of questions; each question provides a prompt for one or more topics to launch a paper's outline. Table 3-2 (presented below) provides the questions. Table 3-2 provides a topics scaffold which my lab members can repeatedly use to begin outlining different papers related to our broader goals.

Table 3-2 represents a highly tailored version of the Table 3-1 IMRaD topics scaffold. The tailoring reflects the unique properties of the musical networks project. Table 3-2 eliminates many prompts from Table 3-1 because my lab's research uses computer simulations (not human participants), and because we rarely train networks in different experimental conditions. Table 3-2's early questions introduce my lab's general goal, using artificial neural networks to inform how we can study music. The later questions generate topics required to describe a particular project.

Of course, Table 3-2 provides a scaffold uniquely suited to my lab's goals. You will require a different scaffold, one customized to your own research goals. Fortunately, you can use the same approach to create your own custom topics scaffold: take one or more of your successful projects which relate to future work, reverse engineer the structure of your projects, and convert the structure into your own topic-generating questions to be used for your future projects.

Table 3-2. A topics scaffold for a musical networks project conducted by the Biological Computation Project (Dawson lab) in the Department of Psychology at the University of Alberta

Interpreting Musical Networks: Reusable Topic Scaffold:
Introduction (Questions 1-11)

- | | |
|---|--|
| 1. Why do cognitive scientists use computer simulation models? | 2. What are the properties of an artificial neural network (ANN)? |
| 3. Why are ANNs popular models for cognitive science? | 4. Is it easy to understand how an ANN solves a problem? |
| 5. If ANNs are hard to understand, can they inform cognitive science? | 6. Do methods exist for seeing how an ANN solves a problem? |
| 7. Can we illustrate a method for understanding a ANN's structure? | 8. What problem are we using as a case study for understanding ANNs? |
| 9. Why did we choose our problem to illustrate understanding an ANN? | 10. What main points will we make in our case study? |
| 11. How will the rest of the current paper proceed? | |

Interpreting Musical Networks: Reusable Topic Scaffold: Method
(Questions 12-24)

- | | |
|--|---|
| 12. What task did we train our ANNs to perform? | 13. Why did we choose this task to study with our ANNs? |
| 14. What kind of ANN did we train to perform our task of interest? | 15. Do we need a figure to illustrate the kind of ANN we trained? |
| 16. Why did we train this particular type of ANN on our problem of interest? | 17. How many patterns did we include in our training set? |

Table 3-2. A topics scaffold for a musical networks project conducted by the Biological Computation Project (Dawson lab) in the Department of Psychology at the University of Alberta

- | | |
|--|---|
| 18. Why does our training set include a particular number of patterns? | 19. How did we represent stimuli to be presented to our ANNs? |
| 20. How did we represent desired responses to be generated by our ANN? | 21. What do we mean when we say we train our ANNs? |
| 22. How do we initialize our ANNs before training begins? | 23. What specific learning rule do we use to train our ANNs? |
| 24. How do we train our ANNs on each pattern in a training set? | |

Interpreting Musical Networks: Reusable Topic Scaffold: Results (Questions 25-30)

- | | |
|---|---|
| 25. What criterion do we use to decide when training is to stop? | 26. Did our ANNs learn to solve the problem we trained them to perform? |
| 27. How long on average did it take our ANNs to learn to solve the problem? | 28. What method did we use to interpret our ANNs? |
| 29. Why did we choose our method for interpreting our ANNs? | 30. What did interpretation say about how our ANNs solved the problem? |

Interpreting Musical Networks: Reusable Topic Scaffold: Discussion (Questions 31-14)

- | | |
|---|---|
| 31. What is a good summary of our research goal? | 32. What is a good summary of our main results? |
| 33. How similar are my results to previous results? | 34. What is surprising about my findings? |

Table 3-2. A topics scaffold for a musical networks project conducted by the Biological Computation Project (Dawson lab) in the Department of Psychology at the University of Alberta

35. Will my results generalize to other tasks?	36. How do my results relate to music theory?
37. How do my results relate to musical cognition?	38. What new ideas are provided by my network interpretations?
39. What new questions are raised by my network interpretations?	40. Why are my findings important?
41. What is the take home message of my results?	

3.5 LET YOUR FINDINGS BE YOUR SCAFFOLD

What topics do your findings need you to communicate?

Tables 3-1 and 3-2 provide questions to help generate initial topics for the Chapter 2 scaffold. Each new scaffold cries out to ask its questions in order, starting from the beginning and proceeding linearly to the end. However, you should avoid generating topics by answering questions in a predetermined order. Starting a writing project at its beginning may not serve you well (Becker, 2020).

A better approach for using the prompts already presented in Chapter 3 involves asking questions in *any* desired order. Creating topics proceeds faster if you answer easier questions first, no matter where the questions appear in a scaffold. “Do whatever comes easiest first” (Becker, 2020, p. 60). You should only use scaffold ordering after you generate answers; you can place index cards holding your answers in a position signaled by a question’s number. But you don’t have to answer the questions in order. When you feel free to generate topics by answering questions from the scaffold’s middle first, then you also open the door to new approaches for scaffolding topics.

For example, I generally know my methods and results

before I know I want to write about in my introduction or discussion. So, I use a project's findings to scaffold the initial topics for my paper's scaffold. I generate my topics for the middle of my IMRaD structure first.

How do you start in the middle and use your findings to scaffold your initial topics? First, you generate the middle topics of your paper – the topics directly related to your findings. You *repeatedly* ask the following question: *What did I find?*

Each answer to the question generates a topic for you to write about in your paper's middle. Write the answer down on a new index card's blank side. Keep asking (and answering) the question. By the time you fail to generate new answers, you will have created several potential topics for your paper. You could take a minute to arrange their index cards, so the most important results come first. Arranging your topics helps you think about each finding's relative importance.

Next, for each topic generated by answering the question above create additional topics by answering other basic questions:

- How did I look for the result?
- What evidence reveals the result?
- Why is the result important?

Again, write each answer on a new index card, and place your new cards near the card upon which you wrote the

original topic which relates to your answers. Answering basic questions about each finding generates more topics for your paper's middle and raises ideas to address when you create your introduction or discussion. Generating topics for your paper's middle when you begin to scaffold will make it easier for you to later generate topics for your introduction or discussion.

After generating topics related to your findings, you can next proceed to generating topics for your paper's beginning and ending parts. Because you have already generated topics related to what you have found, and considered the importance of your findings, you can now create potential topics for your paper's other parts.

Another approach to generating more topics (having already generated topics about your findings) is to answer a small number 'stock questions'. For a paper's introduction, relevant questions include:

- What am I studying?
- Why is what I am studying important?
- What is known about what I am studying?
- What still needs to be learned about what I am studying?
- How do I plan to add to this knowledge with the project I am reporting?

Other stock questions can be used to generate topics for the discussion:

- What are my main findings?
- What theories do my findings support?
- What theories do my findings refute?
- What problems do my findings cause?
- What new questions do my findings raise?

Again, answering such stock questions generates topics for your paper's early and late sections. However, by starting in your paper's middle – by generating topics about your findings – you have already considered what you found as well as the importance of your results. Starting in the middle makes generating your early and late topics easier.

3.6 STORYBOARD WITH IMAGES

Let your images do your writing.

Many different outlining techniques exist. For example, screenwriters often use an outlining technique called *storyboarding*. A storyboard provides a linear sequence of images. Each image, called a panel, represents a single scene or event. The entire sequence delivers a visual outline of an entire screenplay (Pallant & Price, 2015). While we usually associate storyboarding with screenwriting, some use storyboarding to outline other kinds of writing (Bogard & McMackin, 2012; Lee et al., 2013; Seals, 2023).

Section 3.5 described how you can use your findings as a scaffold for generating initial topics for a paper. In many cases, researchers use figures, graphs, photographs or other images to present a project's main findings (Nersesian et al., 2022). When you have your images in advance, you can use storyboarding to generate initial topics. Your project's images can scaffold initial topics; storyboarding becomes a special case of using your results as a scaffold for generating seed topics (Section 3.5).

Storyboarding an IMRaD project requires you to create images for your findings before *any* writing begins.

Researchers often create images first because they use images to summarize discoveries while a project proceeds. Storyboarding works well with research which produces images before it generates writing. Once a project's images exist, you can create a storyboard for generating initial topics.

First, you create an index card for each image. Each index card merely serves as a placeholder; you only need a rough sketch or an image name on the index card. Second, you display your image cards and rearrange them to discover a plausible narrative. The arranged image cards provide your storyboard.

With the storyboard laid out, you next generate topics for each image. As discussed in Section 3.5, you can use stock questions to generate topics. You again write each answer to a question on a new index card and place the answer near the image card to which it relates. I provide some useful stock questions below:

- What does the image show?
- How does the current image relate to the previous image?
- How does the current image relate to the next image?
- Why is the image important?

Your images presumably illustrate your project's findings, so storyboarding generates topics for your paper's middle. So, Step 4 requires you to generate topics for your writing

project's start and end. To do so, you can turn to techniques discussed earlier in Chapter 3: using another topics scaffold (e.g., Table 3-1) or answering stock questions like those raised in Section 3.5. Remember, the stock questions you used to generate topics for each image made you reflect upon what you found and on the importance of your findings. Thus, storyboarding provides a basis to help you generate topics for your paper's other parts.

3.7 BRAINSTORMING

Free associate topics without evaluating them.

The scaffolding method described in Chapter 2 helps writing, in part, by separating creation from evaluation (Elbow, 1981). However, Chapter 3's scaffolds may bring creation and evaluation into conflict. The scaffolds provide questions to answer. Unfortunately, by using questions the scaffolds permit evaluation to prematurely creep back into the writing process. You might unconsciously strive to generate 'correct' or 'best' answers to a scaffold's prompts, evaluating your topics as you generate them.

How might you avoid premature evaluation when you first create topics? You could use a less structured method for topic creation, such as *brainstorming*, which I briefly introduced in Section 2.6. Advertising executive Alex Osborn developed brainstorming, a technique made famous by his book *Applied imagination* (Osborn, 1953). Osborn designed brainstorming to help groups generate as many ideas as possible by delaying evaluation. While immensely popular in the mid-20th century, brainstorming became less popular when experimental results raised questions about brainstorming's success (Bouchard, 1971; Taylor et al., 1958). Osborne promoted brainstorming with anecdotal data (Osborn, 1948, 1953); let

me follow suit by stating I find brainstorming very effective for creating initial topics for my Chapter 2 scaffold.

Groups use brainstorming to generate as many ideas as possible. To succeed, Osborn (1953) required a brainstorming group to follow four basic rules:

- Criticism is banned.
- Quantity is wanted.
- Freewheeling is welcomed (to generate wild ideas)
- Ideas already generated can be combined to create new ideas.

The rules for brainstorming – particularly banning criticism, and desiring many ideas – make it analogous to freewriting (Elbow, 1981). Brainstorming's aims to creatively generate many ideas by delaying criticism, an aim well suited to Step 1 of the Chapter 2 scaffold. To start a scaffold, you try to generate many potential broad topics, in any order, without evaluating them (e.g., without framing them as sentences). Brainstorming produces many broad topics if you repeatedly generate answers to one basic question:

- What could I tell my audience about my research project?

Repeatedly asking and answering the question above, while following brainstorming's basic rules, should produce

many broad topics to initiate a scaffold. Importantly, evaluating the generated topics occurs *later*— in Step 2 (Section 2.7).

Using brainstorming to generate topics helps inhibit other forms of criticism or evaluation. For instance, brainstorming need not pay attention to topic order, because you produce topics as they come to mind. If you focus on what topic importance, or on the best topic order, your evaluation occurs prematurely.

Brainstorming permits several individuals to work together to generate initial topics. A writing group can brainstorm topics together, writing each topic down on its own index card and immediately putting the new card on display for the whole group. Obviously, an individual can also use brainstorming to generate topics.

You should also consider brainstorming's emphasis on banning criticism when using other Chapter 3 scaffolds. For instance, the Table 3-1 IMRaD questions become more effective if you use each as a brainstorming cue; don't try to produce the best answer to each question.

3.8 MULTIPLE SCAFFOLDS FOR WRITING

How else can your world help your writing?

Writing is challenging. Chapter 1 argued we make writing challenging by believing the myth of inspiration: we wait for fully formed sentences to arise in our mind. If writing does not depend upon inspiration, believing the myth of inspiration makes writing more difficult.

Chapter 2 illustrated how embodied cognition can help us to abandon the myth of inspiration, making your writing easier. When using the Chapter 2 method, you use index cards to display your thoughts on a surface in your world. You need not remember your thoughts, because you can see them on display. Thinking about your ideas becomes manipulating the index cards in your display. However, to use the method you must start by generating broad topics which leads, again, to the myth of inspiration. Where do your topics come from?

Chapter 3 used embodied cognition to provide an answer: it describes several scaffolds to help generate your initial topics. The scaffolds provide questions whose answers become topics for your outline. The scaffolds offered in Chap-

ter 3 also move beyond this being a ‘book about using index cards to outline’. In Chapter 3 the prompt questions are the scaffold. While they could be placed on index cards, but they could just as easily be used by being read from the book. Furthermore, if you choose to move a set of prompt questions to index cards, then you will use the cards differently than is the case in the Chapter 2 method.

By providing additional scaffolds to launch topic creation for your outline’s scaffold, Chapter 3 shows you can use one scaffold to help create another, indicating *many* different scaffolds for writing exist. Chapter 4 explores additional scaffolds to support creating your first draft, revising and polishing your manuscript, and dealing with suggestions from editors and reviewers.

PART IV

CHAPTER 4: SCAFFOLDING DRAFTING AND REVISING

Your outline scaffolds your first draft; your first draft scaffolds your final paper.

The Chapter 2 scaffold provides you with a rich, well-structured outline which becomes your first draft. Chapter 4 describes how your outline scaffolds your first draft. Your outline already contains the topic sentence and the concluding sentence for each paragraph, which constrain what supporting sentences you can add to finish the paragraph. Once you finish your first draft, it then becomes a scaffold to help you complete your paper's final version. Chapter 4 also describes scaffolds to help you revise your finished product once an editor asks you to respond to reviewer suggestions.

4.1 YOUR OUTLINE IS A SCAFFOLD

Drafting is easy; writing is harder.

Chapter 2 described using index cards to scaffold a manuscript's outline. Chapter 3 proposed additional scaffolds to help you generate your initial topics. Chapter 4 now discusses steps you take after completing your outline. In the scaffold's final step, you copy your external scaffold – the sentences written on your index cards – into a word processor. You next convert the word processor version of your outline into a complete first draft.

I believe you can easily convert your outline into your first draft because your outline itself serves as a ready-made scaffold. Your outline establishes logical links between paragraphs, not to mention the two most important sentences for each paragraph. These two sentences place powerful constraints on what sentences you can add to create your first draft. We call the sentences you need to add the *supporting sentences* (Sarnecka, 2019) or the *development* (Greene, 2013). The development provides additional details which flesh out the relationship between the topic sentence and the concluding sentence. “One good way to write the development is to lay out well-defined steps that lead to some conclusion [from the

topic sentence]” (Green, 2013, p. 68). With topic and concluding sentences in hand, the development seems to write itself as you add flesh to your outline.

You should not expect your first draft to provide a polished, finished, or submittable manuscript (Becker, 2020). “The first draft is the down draft – you just get it down” (Lamott, 1995, p.24). When you fill in your outline, you should pay more attention to the points you make; you should pay less attention to your writing style.

You worry less about writing style when you create your first draft because you expect to revise and polish your draft later. “The second draft is the up draft – you fix it up. You try to say what you have to say more accurately” (Lamott, 1995, p. 24). However, Chapter 4 argues a scaffold – your first draft – helps you create your ‘up draft’.

When you write your first draft’s sentences, recognize you will later repeatedly work through your manuscript, revising and improving your sentences (Becker, 2020; Elbow, 1981; Germano, 2021; Lamott, 1995; McPhee, 2017; Zinsser, 2006). “The third draft is the dental draft, where you check every tooth, to see if it’s loose or cramped or decayed, or even, God help us, healthy” (Lamott, 1995, p 24). William Zinsser reminds you “very few sentences come out right the first time, or even the third time. Remember this in moments of despair” (Zinsser, 2006, p. 9). Zinsser’s reminder makes writing easier, because you realize you don’t need perfect sentences in your first draft; you repair broken sentences later. In Chapter 4, I

argue each (better) draft of your paper makes writing easier, too, because each draft scaffolds how you develop the next draft.

When revising ends, you possess a manuscript which satisfies you. For academic writing which reports research, you next submit your manuscript for review by a journal's editor. When reviews come back, you usually must return to your manuscript to for further revising and polishing — editors or reviewers usually ask for changes. Chapter 4 also shows how comments and suggestions from editors and reviewers provide another scaffold for additional revisions.

4.2 USE STANDARD STRUCTURES TO CONVERT YOUR OUTLINE TO YOUR FIRST DRAFT

Topic sentences and concluding sentences restrict potential development.

When you finish your Chapter 2 scaffold, you next convert the scaffold into a complete first draft by adding sentences to every paragraph. Remember, the Chapter 2 method assumes each paragraph communicates a single topic. A paragraph's topic sentence states the topic while the concluding sentence restates or summarizes the topic. You need to add supporting sentences to develop each paragraph's topic and to strengthen the link between the paragraph's topic and concluding sentences.

You can add appropriate supporting sentences by 1) thinking about your topic and concluding sentences and by 2) exploiting another scaffold, *standard paragraph structure* (Flower, 1989). A standard paragraph structure provides a

template which restricts the supporting sentences you add to your outline.

The *topic-restrict-illustration-topic* (TRIT) pattern (Flower, 1989) provides one standard paragraph structure which uses supporting sentences to link the topic and supporting sentences of a paragraph. With the TRIT pattern, the paragraph begins with the topic sentence. The second sentence – a supporting sentence – refines or restricts the topic. The remaining supporting sentences develop or illustrate the topic, and lead into the concluding sentence which restates the topic.

The paragraph below, taken from a recent book (Dawson, 2022), illustrates the TRIT pattern. The first sentence defines the topic as ‘how a researcher used a network to challenge an assumption’. The second sentence restricts the topic to particular evidence used to challenge the assumption. The next two supporting sentences develop the topic. The concluding sentence restates the topic sentence in more detail, supported by the paragraph’s middle sentences:

“Farah used PDP networks to challenge the locality assumption. Farah demonstrated that lesions to networks produce dissociations, arguing that this provides evidence against the locality assumption. Her argument was that networks are distributed systems, not local. Therefore, the locality assumption is not true of networks. That these networks still exhibit dissociations indicates that dissociations need not be caused by damage to localized brain functions” (Dawson, 2022, p 128).

For the TRIT pattern to work, logical links must exist between successive sentences in your paragraph (Flower, 1989). Sarnecka (2019) describes an effective technique to establish logical links: the *topic chain*. Sarnecka points out every sentence has a topic (its subject) and a comment (what the sentence says about the topic). She notes you create logical links when each supporting sentence shares the same topic. “A strong, clear link is formed when the topic of a sentence refers to something already mentioned, preferably in the previous sentence” (Sarnecka, 2019, p. 247).

The paragraph below illustrates a topic chain. The first sentence introduces the topic ‘feedback loop’. the remaining sentences in the paragraph have ‘feedback’ as their topic; two sentences include the word ‘feedback’, and the sentence starting with ‘The agent acts on the world’ provides a definition related to the topic:

“Cybernetics explained behaviour by appealing to the feedback loop. Feedback measures the distance between an agent’s current state and a goal state that the agent desires. The agent acts on the world to decrease the distance between the current state and the desired state. A feedback loop cycles back and forth between an agent’s actions and environmental changes, constantly measuring the distance from a desired goal to alter or guide the agent’s future actions” (Dawson, 2022, p. 52).

Flower (1989) identifies several other common paragraph structures, including the *problem-solution pattern*. When you use the problem-solution pattern, your paragraph

opens by stating a problem or question; your paragraph then communicates a solution or answer. For instance, the example below starts with a question; the remaining sentences provide a possible answer (humans differ from machines) and provide different evidence to support the answer:

“Do humans differ from animals and machines? Many scholars argue that humans are special because they use mental representations, a view rooted in 17th-century philosophy. Descartes) argued that only humans possess a soul or consciousness, and the soul’s essence is only to think. His notion of thinking resembles modern information processing. Modern echoes of Descartes are easily found. Bronowski writes that ‘man is distinguished from other animals by his imaginative gifts’. Bertalanffy argues that ‘symbolism, if you will, is the divine spark distinguishing the poorest specimen of true man from the most perfectly adapted animal’ (Dawson, 2022, p. 31).

Flower (1989) also describes another paragraph structure, the *cause-and-effect pattern*. In the cause-and-effect pattern, the topic sentence introduces a cause; the remaining sentences discuss the cause’s possible effects. In the example below, an assumption about human cognition provides the cause. The effects result from the cause; cognitive psychologists ‘must explain’ and ‘must adopt’ because of their assumption:

“Cognitive psychologists assume that cognition is computation: rule-governed symbol manipulation. As a result, they must explain cognition by appealing to processes that they

cannot observe directly. Cognitive psychologists must adopt a philosophy of science different from that of behaviorism. Behaviorists criticized the cognitive approach as being non-scientific because behaviorists believe that functional decompositions do not explain. In response, cognitive psychologists adopt a different approach to explanation, one no less scientific than the approach used by behaviorists” (Dawson, 2022, p. 69).

Flower (1989) describes another paragraph structure, the *chronological order pattern*. When you use the chronological order pattern, your topic sentence introduces a sequence, causing the reader to expect the writer to detail the sequence in the paragraph’s remaining sentences. The example paragraph below implies a sequence – ‘reversing the major method’ – and then provides the sequence (‘First’, ‘Second’, ‘And third’):

“We recall a digit sequence by reversing the major method. First, we recall a word from memory. Second, we extract the word’s consonant sounds. And third, we convert the consonant sounds into digits—our responses to the experimenter’s running a digit-span task. We repeat the process for each word held in primary memory” (Dawson, 2022, p. 41).

In Chapter 2, I described using index cards to scaffold your outline. The outline produced from the Chapter 2 scaffold becomes a new scaffold, one which helps you create your first draft. In all the paragraph patterns introduced above, topic and concluding sentences place powerful constraints on what supporting sentences you can add. The existing topic

and concluding sentences remove the myth of inspiration by restricting your freedom for finishing each paragraph's first draft.

Your outline can provide other constraints to guide the supporting information you add to create your draft. Notes on index cards remind you to cite particular sources in supporting sentences. Additional cards remind you to add or discuss specific quotes, images, or tables. Your work creating the outline means you require less work – or at least less inspiration – when you create your first draft.

Again, remember to treat your first draft as *only* a draft. You do not need perfection because you aim to 'get the first draft down' (Lamott, 1995). When creating your first draft, you do not need to create ideal supporting sentences. Expecting imperfections in your first draft provides you with another advantage. You need not write everything down when you create your draft.

For example, in my first draft of Chapter 4, I referred to other sections of the book, or I realized I did not have a date or a page number for a quote's citation. I did not immediately break away from drafting to look up the missing information immediately. Instead, I added question marks and highlighted them in yellow. The yellow reminds me to take a minute to find the missing information – later!

Highlighting issues to resolve later illustrates another advantage of your scaffold: the *principle of modular design*. Engineers adopt the principle of modular design when they

create parts of a project independently of other parts. You need not develop your whole project at once. By emphasizing paragraphs as your manuscript's core components, you compartmentalize your work, drafting one paragraph at a time (Lamott, 1995).

You can also exploit the principle of modular design by doing 'grunt work' on a draft when writing supporting sentences tires you out. For example, you can take time away from writing, but still develop the draft, by completing component tasks: dealing with highlighted missing information, drafting a figure or a table, adding citations, and so on.

As my draft develops – even in modules – I use my draft to scaffold itself. I like to take time to read my fully drafted paragraphs out loud before I proceed to add supporting sentences to flesh out other paragraphs in my outline (Becker, 2020). I find reading out loud helps me maintain my voice when I turn to adding new sentences later in the draft. Using my existing draft to scaffold its remaining paragraphs once again illustrates the principle of modular design.

4.3 POLISHING YOUR FIRST DRAFT

Craft your draft.

When you complete the process described in Section 4.2, your outline has become your first draft. With your draft down, you next polish up (Lamott, 1995) by revising your draft to improve your wording.

Polishing takes time and effort and requires multiple passes (Becker, 2020; Elbow, 1981; Germano, 2021; Lamott, 1995; McPhee, 2017; Zinsser, 2006). Improving writing requires multiple passes because each editing pass makes your manuscript better. William Zinsser (2006, p. 17) teaches an important lesson to his students: “If you give me an eight-page article and I tell you to cut it to four pages, you’ll howl and say it can’t be done. Then you’ll go home and do it, and it will be much better. After that comes the hard part: cutting it to three.” Editing your manuscript requires effort because you must pay attention to detail when you make choices about which words to cut or to rewrite.

Paying attention to detail while editing a manuscript means carefully considering a manuscript’s words while also keeping in mind rules for improving writing style. To keep such rules in mind, I often start a big editing project by revisit-

ing my favorite writing books (Strunk & White, 1959; Zinsser, 2006) which remind, inspire and motivate me. The style books I turn to refresh the core principles which guide my editing. I follow six crucial principles (Table 4-1).

Table 4-1. General principles to follow during polishing, with example sources.

Examples Of Principle 1: Choose words carefully

Have a point and make it by means of the best words (Barzun, 1985)

Choose your words with care (Greene, 2013)

Only use adjectives and adverbs when they add new information to a sentence (Sword, 2016)

Vary your prepositions (Sword, 2016)

Examples Of Principle 2: Think in paragraphs

Make the paragraph the unit of composition (Strunk & White, 1959)

Examples Of Principle 3: Use the active voice

Use the active voice (Strunk & White, 1959)

Favor the active voice (Greene, 2013)

Favor strong, specific, robust action verbs over weak, vague, lazy ones (Sword, 2016)

Limit your use of be-verbs (Sword, 2016)

Examples Of Principle 4: Be concrete

Use definite, specific, concrete language

Put statements in positive form (Strunk & White, 1959)

Do most of your descriptive work with concrete nouns and active verbs (Sword, 2016)

Anchor abstract ideas in concrete language and images (Sword, 2016)

Show don't tell (Sword, 2016)

Table 4-1. General principles to follow during polishing, with example sources.

Limit your use of abstract nouns and nominalizations (Sword, 2016)

Examples Of Principle 5: Pay attention to word order

Keep related words together (Strunk & White, 1959)

Place the emphatic words of a sentence at the end (Strunk & White, 1959)

Do not allow a noun and its accompanying verb to be separated by more than twelve words (Sword, 2016)

Avoid using more than three prepositional phrases in a row (Sword, 2016)

Examples Of Principle 6: Omit needless words

Omit needless words (Strunk & White, 1959)

Omit needless words (Greene, 2013)

Weed out the jargon (Barzun, 1985)

Look for all fancy wordings and get rid of them (Barzun, 1985)

Use it and this only when you can state exactly which noun each word refers to (Sword, 2016)

Avoid overuse of ‘academic ad-words’ (Sword, 2016)

Avoid using that more than once in a sentence or three times in a paragraph (Sword, 2016)

Beware of sweeping generalization which begin with ‘There’ (Sword, 2016)

‘*Choose words carefully*’ provides my first core principle for revising. Strunk and White (1959, p. 17) restate “choose

words carefully” as “make every word tell.” ‘Choose words carefully’ reminds me every word I use affects how easily a reader will understand my writing. “Prefer the short word to the long; the concrete to the abstract; and the familiar to the unfamiliar” (Barzun, 1985, p. 18).

On its own, ‘choose words carefully’ does not guide my specific word changes; I unpack such changes using other principles I discuss below. However, ‘choose words carefully’ guides me by reminding me I must concentrate; I must pay attention to every word. If I cannot do so, then I need to stop editing for a while.

‘Choose words carefully’ also reminds me to use the principle of modular design. I do not have to edit my whole manuscript at once. I can revise smaller parts of my manuscript – individual sections, or even single paragraphs – and then step away when my attention lapses.

Many books about writing suggest taking a break as a method for improving writing (Cameron, 2022; Heard, 2022; Janzer, 2016; King, 2000). “Most people get more done in three one-hour blocks of work separated by breaks than in one four-hour stint” (Janzer, 2016, p. 25). Indeed, the effort and time required to revise makes revising impossible to do without breaks. “I can never get it all down, and besides, there are times when I have to step away from the table, notebook, and turn to face my own life” (Goldberg, 2016, p. 115).

When I recognize I require a break, I must also pay attention to *how* I break away from my project. I prefer to

break away when I know what I will write next. “I always worked until I had something done and I always stopped when I knew what was going to happen next. That way I could be sure of going on the next day” (Hemingway, 1964, p. 12).

I find knowing how to get back to revising just as important as knowing when, and how, to step away from a manuscript. I prefer to re-read the paragraphs I’ve already edited, material which comes right before the next material I plan to revise. Re-reading refreshes my memory about what I’ve done already, gets my head back into my manuscript, and sets the standard for my next editing pass.

Finally, ‘choose words carefully’ reminds me to concentrate on every word, encouraging me to read my words out loud (Becker, 2020). I read out loud slower than I read to myself silently; reading out loud forces me to focus more attention on my words. As I hear what I read out loud, if I hear something strange or find myself being confused, then I know I must do more revising!

My second core principle for revising is ‘*think in paragraphs*’. For writing, to think in paragraphs is to “make the paragraph the unit of composition” (Strunk & White, 1959, p. 11). The scaffold detailed in Chapter 2 follows the principle by developing paragraphs which communicate one, and only one, topic.

While editing, ‘think in paragraphs’ leads me to check if each paragraph communicates a single topic. If a paragraph tries to express two topics (Figure 2-7), then I must revise. I

must split the bad paragraph into two good paragraphs (Figure 2-8), or I must keep the bad paragraph – after eliminating the paragraph’s second topic.

‘Think in paragraphs’ also leads me to check the logical flow between successive paragraphs in each paragraph. Flower (1989) recommends identifying paragraph types (TRIT etc., see Section 4.2) and then examine a paragraph’s sentences to ensure they adhere to the paragraph’s pattern. If the paragraph uses a topic chain structure (Sarnecka, 2019), then I check if each supporting sentence in a paragraph uses the same topic. Germano (2022) shows how reading paragraphs out of order helps evaluate paragraph logic.

‘Think in paragraphs’ also reminds me to reconsider paragraph order. Once I have used my outline to flesh out paragraphs, the sentences I have added might suggest moving the paragraph to a different position in my manuscript. For example, if a paragraph makes a more specific point than I expected, I consider moving it to later in the manuscript; if the paragraph communicates a more general topic, I consider moving it to earlier in the manuscript (assuming I’m organizing topics from general to specific as described in Section 2.3).

Finally, ‘think in paragraphs’ reminds me to check paragraph bridges – the link from a concluding sentence in one paragraph to the topic sentence of the next paragraph. I certainly must check the bridge if I move a paragraph to a new location. However, I check all paragraph bridges whether I have moved paragraphs or not.

Of course, if I find any issue when I use ‘think in paragraphs’ to guide my editing, then I must rewrite. I must also rewrite when I find problems from edits guided by the other principles described below. Because editing requires rewriting, I like to copy my current manuscript (the first draft or a later revision) before I do any new editing. I will only change the wording in the copy. If revision fails, then I can return to an older version. As a result, when editing I produce different versions of the same manuscript which I can use to track my editing.

‘*Omit needless words*’ provides the third general principle to guide my revising. During outlining (Chapter 2) and drafting (Section 4.2) I deliberately pay little attention to writing style (Elbow, 1981). I recognize I will improve my wording later. However, by not attending to style, I expect to see a wordy first draft.

I know my reader will have trouble understanding a wordy manuscript. I want to make my manuscript concise and plain. “The whole world will tell you, if you care to ask, that your words should be simple & direct. Everybody likes the other fellow’s prose to be *plain*” (Barzun, 1985, p. 17). As I revise to simplify my manuscript’s language, I try to shorten my paragraphs. An editor performed a favor to Stephen King by providing a formula for shortening a draft: “You need to revise for length. Formula: $2^{\text{nd}} \text{ draft} = 1^{\text{st}} \text{ draft} - 10\%$ ” (King, 2000, p. 222). King copied the formula onto cardboard and for display on the wall beside his typewriter.

With such advice in mind, I look for phrases to shorten or remove on every editing pass. For instance, I know when I use ‘of’, I usually have written a phrase which I can trim. Table 4-2’s first three rows illustrate phrase shortening made on Chapter 4’s first draft. Any good book about writing style will provide advice for shortening sentences.

Table 4-2. Examples of edits which removed unnecessary words from the current chapter.

Original	Revised	Words Removed
“The first step of the Chapter 2 method”	“Chapter 2’s first step”	4
“Parts of my favorite style books”	“My favorite writing books”	2
“reminds me to take advantage of the principle of modular design”	“reminds me to use the principle of modular design”	3
“Notes on index cards remind us to cite particular sources in a set of supporting sentences.”	“Notes on index cards remind us to cite particular sources”	7
“To guide the supporting information we add to create the draft”	“To guide the supporting information we add”	4
“In filling in the development, we pay more attention to the points being made; we are less concerned about writing style. We don’t worry about writing style because ...”	“In filling in the development, we don’t worry about writing style because ...”	17

On every editing pass through a manuscript, I also seek unnecessary words to delete. Unnecessary words may state something obvious given the sentence’s context. Unnecessary

words may restate an idea already provided by another sentence. Table 4-2's last three rows illustrate deleting unnecessary words.

Towards my final editing stages, I look for particular words which I can almost always delete. For example, I can almost always delete the word 'that' without changing a sentence's meaning.

Finding unnecessary words requires me to concentrate. As a result, every editing pass reveals more words to remove; I miss such words in earlier passes when my concentration lapses. I find myself frequently amazed at how much I can shorten a manuscript after already removing some unnecessary words. I remember Zinsser's advice – reduce the original by 50% with the first editing pass; reduce the edited version by 25% with the second editing pass, or King's cardboard formula.

The fourth general principle which guides my editing is '*be concrete*'. Many style books encourage using concrete language (Leith, 2018; Rhodes, 1995; Schimel, 2012; Strunk & White, 1959; Sword, 2012, 2016). "If those who have studied the art of writing are in accord on any one point, it is this: the surest way to arouse and hold the reader's attention is by being specific, definite, and concrete" (Strunk & White, 1959, p. 15).

Concrete language uses nouns which describe real world objects; concrete language uses active verbs. I can visualize concrete language more easily than I can visualize counterpart, abstract language. Consider an intentionally abstract

sentence: “The investigation was intended to provide a determination of whether embodied cognitivism was beneficial to academic writing.” Now consider its concrete version: “I investigate whether manipulating index cards helps academic writing.” Which sentence do you find easier to visualize, remember or understand?

Style books recommend using concrete language to make your writing easier to understand. “You should prefer concrete language – visual images and real-world situations – to abstract language, because these ask less of the reader’s brain” (Leith, 2018, p. 18). Cognitive psychologists have long known people find concrete ideas easier to visualize and to remember than abstract concepts (Paivio, 1971).

However, academic writing produces an uncomfortable conflict between using concrete and abstract language (Becker, 2020; Schimel, 2012). On the one hand, scientists collect concrete data. On the other hand, scientists typically interpret their data by proposing abstract concepts. I cannot imagine a scientific paper which does not use abstract language at all.

Academic papers usually contain much more abstract language than concrete language (Becker, 2020; Sarnecka, 2019; Schimel, 2012; Sword, 2012, 2016). For instance, academic papers contain many nominalizations – nouns created from verbs or adjectives. Common nominalizations in academic writing include ‘investigation’ (from the verb *investigate*), ‘discussion’ (from the verb *discuss*), *assumption* (from the

verb assume), and comparison (from the verb compare). Nominalizations hide the action from which we derive them. Academic papers usually use jargon which only a specialized audience understands. For instance, papers in my general field (cognitive science) often use jargon with specialized meaning, such as ‘cognitivism’, ‘connectionism’, ‘computation’, ‘intentionality’, and ‘representation’.

Because authors aim their academic papers towards a specialized audience, and typical academic articles emphasize abstractions, authors have trouble replacing abstract language with concrete language. We often require abstract language. For example, I had trouble reducing my use of the word ‘cognition’ in a book with the title ‘*What is cognitive psychology?*’ (Dawson, 2022). Nevertheless, readers find abstract language hard to understand. You need to reduce your use of abstract language; you need to replace abstract language with concrete language as often as possible.

Hence, guided by ‘be concrete’, as I edit, I pay careful attention to my words. Do I use concrete nouns? Are my verbs active? Do I use too many abstractions or nominalizations? By asking such questions as I read word by word, and by replacing abstract words with concrete ones, I make my writing more simple and direct (Barzun, 1985).

‘*Use the active voice*’ provides the fifth general principle to guide my editing. Many style books encourage using the active voice (Heard, 2022; Leith, 2018; Rhodes, 1995; Schimel, 2012; Strunk & White, 1959). ‘Use the active voice’

complements my fourth principle, ‘be concrete’. “Sentences in the active voice with concrete verbs make vigorous prose” (Rhodes, 1995, p. 102).

When a sentence uses the active voice, the sentence’s subject performs the action named by the sentence’s verb. For example, ‘we need not develop the whole project at once’ uses the active voice because the subject ‘we’ performs the action ‘develop’. In contrast, when a sentence uses the passive voice, the sentence’s subject does not perform the sentence’s action. For instance, ‘the whole project need not be developed at once’ uses the passive voice, because the subject ‘project’ does not perform the action ‘be developed’.

Leith (2018) provides an easily remembered test to determine whether a sentence uses the active or passive voice. Leith suggests adding the phrase ‘by zombies’ after the sentence’s main verb. If the modified sentence seems grammatical – like ‘the whole project need not be developed *by zombies* at once’ – then the sentence uses the passive voice. If the sentence does not seem grammatical – like ‘we need not develop *by zombies* the whole project at once’ – then the sentence uses the active voice.

I try to use the active voice as often as possible, making my sentences become concrete and easier to understand. Because the active voice demands present tense, I also wind up using fewer passive ‘*be*-verbs’ or past tense verbs like ‘was’ or ‘were’ (Sword, 2016). The active voice “is clear, concise and direct. It is also visual and evocative” (Schimel, 2012, p. 134).

However, controversy arises when writers use the active voice – or other clarifying techniques – in scientific writing (Becker, 2020). Most scientific articles use the passive voice, which suits research’s past tense (the reported research has already been conducted!) as well as the aloof, objective third person sentence framing frequently seen in a scientific report. Students often learn to avoid the active voice and to avoid using the first person ‘I’ or ‘we’. I do *not* steer my students towards the passive voice because writing in the active voice is better writing.

4.4 WHEN SHOULD YOU STOP REVISING?

Improve your writing, but don't strive for perfection.

I perform the revising process I described in Section 4.3 iteratively, because I conduct repeated editing passes on every manuscript I write. Every time I revise, my manuscript improves. But I never believe I can turn my writing into a perfect paper. At what point should I stop revising and finally send my manuscript out for review?

In academic writing's 'publish or perish' crucible, you need a plan to determine when to stop revising. Every editing pass will make your paper better, improving the odds a journal's editor will accept your paper for publication. But you also feel pressure to publish lots, and to publish quickly (Silvia, 2007). How do you find the 'sweet spot' where you revise enough to write excellent papers, but don't spend too much time editing?

Some successful writers perform a fixed number of editing passes. Stephen King has an advanced plan for revising: "Now let's talk about revising the work – how much and how many drafts? For me the answer has always been two drafts and a polish" (King, 2000, pp. 208-209). For his short stories, Ray Bradbury wrote his first draft on Monday, his second on Tues-

day, his third on Wednesday, his fourth on Thursday, and his fifth on Friday. On Saturday he wrote his sixth draft – and sent the manuscript out (Bradbury, 1990). You can quite plausibly plan how many revisions you will perform.

Alternatively, you might perform enough editing passes to make problems difficult to find. Realizing reviewers will request changes I cannot predict and recognizing my manuscripts will never achieve perfection, I feel content to submit my manuscript when I can't find any problems without expending enormous effort. For example, years ago I collaborated a great deal with Richard Wright. Richard and I had different views about comma usage. When we polished a manuscript to the point we only disagreed about commas, we knew we could submit the manuscript.

My own approach for working on a longer manuscript also ends editing when I feel happy about my writing. I perform multiple editing passes, but don't have a set number in mind before editing begins. I focus on different concerns with each editing pass. I begin by dealing with bigger issues; I focus on more particular problems during later editing passes.

Before any serious editing begins, I finish a complete draft of any manuscript (e.g., article or book). I delay editing until I have a complete draft to separate generating ideas from evaluating how I express the ideas (Elbow, 1981).

Creating a complete first draft also informs me about issues to deal with once editing (finally) begins. For instance, while writing the current book I felt unhappy about my use

of voice when I drafted my first four chapters. I only found a voice I liked when I drafted the fifth chapter. I knew my early editing passes would need to focus on voice to unify voice across my first draft.

For the current book, I conducted a very broad first editing pass. My first editing pass had the primary goal of rewriting anything which didn't make sense during my reading. My first editing pass had the secondary goal of shortening sentences when I saw obvious fixes. Fixing bigger problems in a first editing pass makes smaller problems easier to find in my later editing passes.

My second editing pass focused on voice, which needed attention as I mentioned above. I searched for words like 'one', 'us', or 'we' to find sentences where I could more effectively use 'I' or 'you'. I also looked for opportunities to convert 'the' into the possessive 'my'. My second pass was still broad – voice was a big issue – but because voice was the only issue it examined, Pass 2 was more focused than Pass 1.

My third editing pass resembled Pass 1 but reversed priorities. Shortening sentences whenever possible became my primary goal. Fixing writing which did not make sense became my secondary goal. I made correcting nonsense a secondary goal in Pass 3 because I expected I had already removed nonsense during Pass 1! However, I still looked for nonsense in case I added some when changing voice during Pass 2.

My first three editing passes demanded me to concentrate. My fourth pass improved manuscript quality but

required less effort than earlier passes. Pass 4 focused on dealing with word processor flags (e.g., spelling), fixing citations, checking figure and table numbers and captions. I needed to do such editing at some point. I chose to do my easier editing in Pass 4 to give me a break from the demands of the earlier editing passes.

My fifth editing pass returned to harder work, shortening more sentences. I did so by searching for words which signal when I use too many words: ‘of’ and ‘that’. I can usually remove these two words provided I fix (and shorten) the sentences from which I delete them.

I used software to scaffold my sixth editing pass. Helen Sword provides a free app to accompany her book *The writer’s diet* (Sword, 2016). The app runs within my word processor and reveals problems which Sword discusses in her book. Her app finds many problems which I missed during my earlier editing. I used her app in stages – dealing with ‘be-verbs’ first, then with ‘it, this, that, there’, then with ‘zombie nouns’, then with ‘ad-words’ (adverbs) and ended by dealing with ‘prepositions’. I worked through Sword’s categories in particular order because earlier passes required more revising on my part. Later passes required less writing – I often would simply delete an offending word. When I reached dealing with prepositions, I kept more than Sword would recommend – because I spent so much time talking about moving ideas ‘from’ one world ‘into’ another. I felt many of my prepositions worked because of their concrete nature.

During my seventh, and final, editing pass I found and fixed minor technical problems – spelling, punctuation, citations, and formatting. I think my final editing dots the i's and crosses the t's. Such editing requires attention to detail; I waited a week after Pass 6, and kept away from the manuscript, before I performed Pass 7. I also experimented with having my computer read my manuscript out loud to me so I could discover problems I missed earlier.

When I finished Pass 7, I felt happy about my manuscript. I knew some minor problems must still exist but did not feel compelled to hunt them down. After all, I expected reviewers to ask me to revise my manuscript later; hopefully I would find and fix some remaining issues while dealing with reviewer comments (Section 4.5). Being happy about my manuscript meant the time had come to write a cover letter and ship the product.

4.5 SCAFFOLDING RESPONSES TO REVIEWER COMMENTS

Every comment requires a response.

After receiving an editorial decision, you usually will return to revising your manuscript. A common request in a review is ‘revise and resubmit’, requiring you to make moderate changes. I suggest two steps to launch further revision. First, read the comments from the editor and the reviewers as dispassionately as possible. Second, read your manuscript’s current version while keeping reviewers’ comments in mind.

You can easily find negative reviewer comments upsetting, particularly after all your hard work creating and polishing your manuscript. I try to put my emotional responses to reviews on the back burner. I treat each comment as constructive criticism I can use to further improve my manuscript. Some reviewers make such a perspective difficult – but I try!

I adopt a constructive attitude towards reviewer comments because I realize I must respond to every comment to convince an editor to publish my revised manuscript. Being angry or defensive about comments does not help. I try to find

positive suggestions in every comment, even comments I disagree with.

Because every comment requires a response, and because I must communicate every response to an editor, I create scaffolds to help me revise my manuscript. My scaffolds ensure I respond to every comment and help me create my cover letter for resubmitting my revised manuscript.

I use two scaffolds to support revising from reviewer comments. First, I create a word processing document by pasting in all comments from the review letter. I organize the document to put each comment in its own paragraph. I add a label at each comment's beginning – 'Reviewer 1 comment:', 'Reviewer 2 comment:' and so on.

Next, I use my comments document to build my second scaffold: an index card for each comment. I copy each comment into a new word processing document, a label template; I put each comment onto its own label. I print the labels and paste each label onto an index card (i.e., on the blank side). I use my index cards to organize my revising.

After creating my index card labels, I return to my first scaffold and highlight every reviewer comment in yellow. When I address a comment, I will remove its yellow highlighting to signal I dealt with the comment.

My two scaffolds serve different purposes, as you will see. However, both scaffolds share one goal: tracking which comments I have addressed, and which comments I still need

to deal with. The two scaffolds help me remember to address every comment.

The index card scaffold helps me organize my revising steps. I follow the principle of modular design: I work on one comment at a time. Furthermore, I don't address comments in their order in the letter I received from the editor. Instead, I sort the index cards, putting the comments I feel easiest to address on the top, and the hardest comments to address on the bottom. I deal with the comments in their order in my sorted deck.

After sorting index cards, I begin revising. I work on one index card at a time. I take the top card – the easiest comment to fix – read the label and then address the comment by changing my manuscript. I then update my first scaffold. I add a new paragraph below the comment in the first scaffold, labeled 'Author's response:'. I use the index card I just dealt with to help me describe my editing in response to the comment. I also remove the yellow highlighting from the comment. I mark the index card as completed and put it in a new pile: my 'done cards. I'm rewarded by seeing my 'done cards' pile growing.

I usually find the upper cards in my scaffold easy to deal with. Many reviewers point out spelling mistakes, or sentences which don't make sense because of missing words or other simple problems. Early on, my 'done' pile grows quickly.

As I work through my index cards, I encounter more challenging comments. Some might require rewriting sections

or adding new paragraphs. I might aid such revisions by developing a miniature version of a Chapter 2 scaffold to organize writing my new paragraphs.

Some comments, when I reach them, ask for changes I do not want to make. I jot down my argument for *not* making the change on the lined side of the comment's index card. Later, I will add my argument beneath the reviewer's comment in the other scaffold. When I add my argument to the first scaffold, I feel I have addressed the comment.

I work through each index card, revising my manuscript and altering the first scaffold (the document filled with comments and my added responses) again and again. When I finish revising, all my index cards lie in the 'done' pile; my first scaffold will have no yellow highlighting and will have an author's response paragraph beneath each comment. Both scaffolds signal I have responded to every comment in the editor's letter.

Now I reap the benefits provided by updating the first scaffold every time I revised my manuscript. Each time I revised the scaffold filled with comments, I also described what I did to address a comment. I can now convert the scaffold into a cover letter to use when I resubmit my manuscript. I add some general material at the beginning of the scaffold – thanking the editor for the reviews, describing my main revising, expressing hope the editor will accept the manuscript because of my revisions – and then I write 'the specific responses to each com-

ment are provided below'; I then add the final version of my first scaffold, completing my cover letter.

With my edits finished, and with my cover letter completed, I can send my manuscript back to the editor. The scaffolds ensured I dealt with every comment, helped me deal with the easy problems first, and provided my cover letter. I believe I have increased the odds the editor will accept my revised manuscript because I responded to every comment in the review.

PART V

CHAPTER 5: YOUR WRITING WORLD

Your writing world is your writing mind.

Advocates for embodied cognition argue your mind extends beyond your brain; you think by acting upon your world. When you take the extended mind seriously, you recognize you can use your world to scaffold your mind. A deeper insight occurs if you recognize your world as part of your mind; you can change your mind by changing your world. Chapter 5 explores how you can change your writing by changing the world in which you write.

5.1 YOUR WRITING WORLD

Change your writing by changing your writing world.

A recent book challenges accepted practices for teaching creative writing by arguing separating separate meaning in fiction from meaning in the real world (Salesses, 2021). Salesses argues fiction writing must return to its cultural and historical context. “Reading and writing are not done in a vacuum. What people read and write affects how they act in the world” (Salesses, 2021, p. 6).

Embodied cognition takes Salesses’s (2021) position further by proposing writing does more than affect your actions upon your world. When you write – or when you do any thinking – you act on your world. Consequently, the embodied view of thinking – and writing – does not separate your mind from your world (Clark, 1997; Clark & Chalmers, 1998; Shapiro, 2019).

As a result, embodied cognition claims you can change your mind, and how you think, by changing your world (Dawson et al., 2010). Thus, you can change your writing by changing your writing world.

Books which offer writing advice tacitly endorse the embodied perspective (Butler & Burroway, 2005; Swain, 1974;

Sword, 2017). Often books advise writers to write in particular spaces at regular times using familiar equipment. Why? Your writing world offers affordances which affect what and how you write or think.

Julia Cameron describes how changing worlds changes her writing. She describes four different writing stations in her house, each supporting a different kind of writing (Cameron, 2022). For instance, one holds an uncomfortable couch upon which she writes short notes, finished before the couch causes her back to ache. During the day, as she writes, she moves from station to station to suit her mood. Each station's affordances draw forth a different writing style.

What properties characterize your writing world/mind? What worldly properties can you change to alter your writing/thinking? Chapter 5 describes different objects in your world you can use to scaffold academic writing.

5.2 WRITING EQUIPMENT

Dwell on your work, not on your tools.

In embodied cognitive science, *equipment* mediates our actions on the world (Heidegger, 1927/1962; Winograd & Flores, 1987). Equipment serves as an intermediary between an agent's body and an agent's actions on the world. For example, the Chapter 2 scaffold presumes a writer uses pen or pencil as equipment.

Heidegger argued you do *not* experience equipment as objects in your world. Instead, you only experience the affordances which your equipment offers. “That with which our everyday dealings proximally dwell is not the tools themselves. On the contrary, that with which we concern ourselves primarily is the work” (Heidegger, 1927/1962, p. 99). Heidegger calls equipment's invisibility *readiness-to-hand*.

Winograd and Flores (1987) argue readiness-to-hand shows direct engagement with your world; you only become aware of your equipment when the structural coupling between world, equipment, and agent breaks down. When your pen does not write smoothly, or when your computer keyboard has keys which stick, readiness-to-hand disappears – you suddenly become aware of your equipment.

You need to choose equipment – your pens, or pencils, or keyboards; your paper; your reference material – to maintain readiness-to-hand while you write. “A successful word processing device lets a person operate on the words and paragraphs displayed on the screen, without being aware of formulating and giving commands” (Winograd & Flores, 1987, p. 164). The invisibility of artifacts – readiness-to-hand – signals good design (Dourish, 2001; Norman, 1998, 2002, 2004). When you choose your writing equipment, you design your writing world; you strive for your design to achieve readiness-to-hand.

Writing equipment offers affordances which depend on both bodies and worlds. For example, the dimensions of objects you write upon constrain writing. Jack Kerouac’s novel *On the road* provides one famous example (Kerouac, 1957). Kerouac composed on his typewriter. Being forced to replace typewriter paper disrupted Kerouac’s readiness-to-hand, interrupting his creative flow. Kerouac scaffolded his creative flow by taping together paper sheets to create a 120-foot scroll on which he typed. The scroll eliminated paper-changing interruptions. Chapter 2 offered a more mundane example by advising you to use small index cards in order to restrict how much you can write when you create outline cards. The smaller your world, the shorter you write.

Equipment belongs to the world in which you write. However, embodied cognitive scientists view your world as part of your mind. Your chosen equipment determines not

only how you write, but also determines your writing mind. What choices can you make when designing your writing mind?

5.3 INDEX CARDS

Take advantage of index card affordances.

Chapter 2 described scaffolding your outline for a manuscript. The scaffold defeats the myth of inspiration by moving ideas from your mind to your world. Thinking becomes manipulating objects in the world. To succeed, the objects you manipulate must offer appropriate affordances to you. Section 5.3 describes the objects which bring Chapter 2's scaffold into being.

5.3.1 Index Cards

The index card provides the primary object for the Chapter 2 scaffold. As discussed in Section 1.8, index cards offer many affordances to aid writing: writability, readability, writing short, arrangeability, groupability, portability, expendability, cooperativity, and duality. Such affordances mean index cards play key roles in other scaffolds discussed in Chapter 5. Your index card choices determine the nature of your writing scaffold.

For instance, index cards come in different sizes and different colors. I prefer using 3" x 5" index cards which I can easily find in various colors. As discussed in Section 1.8, these

properties offer me different affordances which aid writing and outlining. My preferred, small, index cards force me to write short when I begin my scaffold. Card color permits me to make different topics or ideas immediately visible.

5.3.2 Printed Labels For Index Cards

Using index cards to build the Chapter 2 scaffold permits me to support group work on a project. Different people can view the index cards at the same time and can work together to manipulate the cards to refine the scaffold. Such group work requires legible information on displayed cards. Also, I may repeatedly use some information on cards because the same information applies to different projects. Therefore, I like to have stock index cards ready for repeated use because of their more permanent, easy-to-read format.

For example, I like to have stock questions for prompting topics (Chapter 3) in a more permanent form. Similarly, related projects conducted within my lab may lean on the same basic information. For instance, my lab's work on musical networks emphasizes musical triads. I have several cards — one for each triad — to remind me of each triad's basic properties, such as its type, its component notes, and its other characteristics.

I use computer-printed labels to create legible information for display on reusable index cards. I create labels in a

word processor template, print the labels out, and stick each label on its own index card. For instance, I print Table 3-1 on Avery 5066 labels. Each page contains 30 different labels, and each question fits on one label. Stock questions become transformed into legible, more permanent index cards for repeated use. Furthermore, having a label pasted on an index card distinguishes the card from handwritten cards added when I answer stock questions. Thus, labels can make index cards visibly distinct.

5.3.3 Visible Distinctiveness

In many research settings, many different activities occur at the same time. Different projects may occur simultaneously, may involve different people, and sometimes different people will work on the same project. Managing scaffolds to deal with multiple projects or multiple people requires visual distinctiveness. I use visual cues to make differences between projects or people immediately visible.

One method to make differences visible uses spatial location: different projects, or the work of different people, become visibly distinct when you display the appropriate index cards in different places. For example, my lab contains many different corkboards; different corkboards hold cards related to different projects or to the work of different people. Corkboard size provides another reason to favor small index

cards: you can display more small index cards on one corkboard.

When you cannot spatially segregate index cards for different projects or people, you can use other methods to make cards visibly distinct. For instance, I sometimes use colored dots to distinguish one set of cards from another.

5.3.4 Date Stamps

The Chapter 2 scaffold emphasizes using index cards to represent topics. Index cards can also support other activities related to a project, such as recording research questions you want to answer, sources you need to read, and so on (e.g., Section 5.5). I like to add dates to my index cards. When did I ask a question? When did I answer a question? When did I read a source? When did I add a topic card to my display? Because I try to date many different cards, I use an ink pad and date stamp for adding dates. I also use specialized pre-inked stamps (e.g., shapes like stars, or the word ‘rush’) to prioritize cards. Again, I try to make card properties immediately visible. Visibility makes scaffolds functional.

5.3.5 Storing Completed Index Cards

Chapter 3 provided a several supplementary scaffolds:

stock questions for helping generate seed topics. You can handily create more permanent cards for repeated use by printing stock questions on labels and by sticking the labels on index cards.

I like to group such cards in a deck, wrap the deck with an elastic band, and store the deck away for future use. Your writing world should include places to store your card decks. My lab has several, including plastic boxes for index cards, small metal cabinets with drawers for cards, and a large filing cabinet with each drawer designed for card storage.

5.4 EFFECTIVE DISPLAYS OF INDEX CARDS

To think – to write – is to act on your world.

5.4.1 Effective Displays

Some authors display their work's progress to aid writing. Tracey Kidder and his editor Richard Todd spread Kidder's book manuscripts across Todd's office floor. "Spreading the pages across the floor in itself lent the illusion of distance and control as we walked among the piles like a pair of Gullivers" (Kidder & Todd, 2013, p. 158). Annie Dillard uses a large conference table. "You lay your pages along the table's edge and pace out the work. You walk along the rows; you weed bits, move bits, and dig out bits, bent over the rows with full hands like a gardener" (Dillard, 1989, p. 46). To provide a useful scaffold, you must consider how you will display and act upon your index cards.

The medium for displaying index cards also offers you affordances. You need a two-dimensional display so you can organize your index cards spatially. You need a large enough display to hold the cards you need to use. For group projects,

you need a display visible to all group members who work together with the cards.

When I use a Chapter 2 scaffold, I employ various display surfaces. At home, depending on how many cards I need to work with, I will lay my cards out on a lap desk, on a coffee table, or on the kitchen table. All my tables permit me to inspect cards on display and to rearrange them. As I don't live alone, I must clean up my cards as soon as possible.

Corkboards provide another, more permanent, display. I use push pins to attach my cards to a corkboard; I can easily position and reposition my cards as an outline's scaffold develops. I tend to use 20" by 30" corkboards, big enough to lay out many cards, but small enough to hide (with card positions intact) if a workspace needs tidying. I also have a larger corkboard mounted on a cart which I can move around during group work or lab meetings.

My lab has an extremely large magnetic whiteboard. I use magnets to place index cards on the whiteboard. I recognize the whiteboard displays index cards publicly; I use the whiteboard to display index cards related to my lab's most important project. The whiteboard offers additional affordances not offered by coffee tables or corkboards. I can use dry erase markers to add lines or labels to help further organize my index cards.

The affordances offered by index cards and displays place requirements on the space in which you write. Your space must offer room for creating index cards, for their display and

manipulation, and for storing cards for later use. You must structure your writing space to provide all the affordances you need.

5.4.2 Why Not Use A Computer?

I began my graduate career in the early 1980s. As a master's student I did not have a personal computer. Instead, I used a mainframe computer which I never saw. I would go to a punch card room in the Social Sciences Centre at the University of Western Ontario to type out data and a program to analyze the data on blank punch cards. The punch card machine would print readable type at a card's top and punch holes in the card to make the information readable for the computer. When finished, I would take my cards down to the building's basement. I would hand my cards to an operator, who would feed them into a card reading machine, sending my information to the mainframe. Then the attendant handed my cards back to me. I would head back upstairs, returning hours later to pick up a printout of my results.

Later, a room containing terminals replaced the punch card room; the terminals connected me remotely to the mainframe. Once, half-jokingly, my Master's supervisor told me he missed the days of punch card data. He felt when he carried his punch cards his data felt real; his data seemed less real when submitted via the terminal.

Section 5.3 focused on physical scaffolds for writing,

emphasizing index cards and media for displaying them. Perhaps my interest in such material objects reflects my own old-fashioned longing for punch cards. Perhaps my preference for index cards reflects my inability to adapt to modern technology.

Remember, the Chapter 2 scaffold produces a rich outline which resides within a word processor. Why use index cards at all? Why not use a computer to develop the outline from the outset? After all, as Paul Simon wrote in *The boy in the bubble*, “These are the days of miracle and wonder/This is the long distance call” (Simon, 1986). Perhaps the days of miracle and wonder demand me to use modern technology to outline and develop my writing projects.

Certainly, many programs exist for developing outlines. Microsoft Word provides an outline view for generating a multi-level document outline. Similarly, many different apps can create and process virtual index cards on a tablet. Why not use a computer to scaffold your outline?

The physicality of index cards offers different affordances than those offered by programs for creating virtual cards. I have tried out various ‘index card programs’ and have always felt limited in how I can act upon virtual cards. For instance, index card software limits how I can arrange cards, because most apps assume users desire neat, nonoverlapping card displays. By offering different affordances, programs for virtual index cards alter my ability to think about my writing.

Always consider your affordances! The scaffold

described in Chapter 2 views thinking as physically manipulating index cards. The physical cards and their physical display permit me to think in a particular way. If I change the affordances — as moving from real cards to virtual cards demands — I change how I think.

For me the affordances offered by physical index cards beautifully support my thinking and writing. You should use other media, like virtual index cards, when different media offer useful affordances. However, virtual index cards and other electronic outliners do not work for me because I miss my preferred affordances.

Nowadays computers offer extremely tempting aids for writing. Large language models like ChatGPT will do the writing for you when you provide ChatGPT a prompt. Table 5-1 provides one example of ChatGPT's writing abilities – and lets ChatGPT suggest why you should avoid using it as a writing tool.

Table 5-1. Response of ChatGPT to a prompt about why it should not be used to create a paper’s first draft. The response was obtained February 10, 2025, from chatgpt.com

Prompt

Write me a three-paragraph essay explaining why I should not use ChatGPT to write the first draft of a paper.

Response

Table 5-1. Response of ChatGPT to a prompt about why it should not be used to create a paper's first draft. The response was obtained February 10, 2025, from chatgpt.com

“While ChatGPT can be an incredibly useful tool for brainstorming ideas and assisting with various writing tasks, it should not be relied upon to write the first draft of a paper. One major reason is that ChatGPT lacks the personal experience and specific knowledge that you, as the writer, bring to your topic. Academic papers often require a deep understanding of the subject matter and the ability to present your own analysis, which AI, despite its impressive capabilities, cannot fully replicate. Writing your own first draft allows you to engage critically with the material and develop your unique voice, something that AI-generated text may lack.

Another reason to avoid using ChatGPT for your first draft is that it can lead to a lack of originality. While ChatGPT can produce well-structured content, it draws from a vast amount of publicly available information and may not generate truly fresh insights. The result can be a paper that feels formulaic or even derivative, rather than one that demonstrates your independent thought and creativity. Originality is often a key aspect of academic writing, and by writing your own first draft, you can ensure that your work reflects your unique perspective on the topic.

Finally, using ChatGPT for the first draft of a paper can hinder the development of essential writing skills. Writing a first draft is an important process that involves refining your ideas, organizing your thoughts, and improving your ability to express complex concepts. By relying on AI to do the initial work, you miss out on practicing these vital skills, which are crucial for academic success and personal growth as a writer. Over time, this dependency could limit your ability to communicate effectively without the aid of external tools, ultimately impeding your development as a scholar.”

ChatGPT's response in Table 5-1 illustrates one seductive reason for using it: ChatGPT can create a sensible series of paragraphs. With the right prompt, ChatGPT could create

a solid first draft for you, sidestepping all the scaffolding discussed in earlier chapters. Why should you avoid such temptation?

You should avoid modern scaffolds like ChatGPT because using large language models to write confuses why academics write in the first place. You might turn to ChatGPT if you believe writing's purpose is to create a manuscript. However, writing has a very different purpose – you write to understand (Howard & Barton, 1986; Zinsser, 1988). Writing is a form of thinking. Furthermore, your ability to write about a topic reflects how much you understand about it. If I have trouble writing about something, I realize I need to step away from writing and return to thinking about the topic to understand more about it.

The scaffolds provided in earlier chapters reflect the need to think or understand before you write. By being inspired by embodied cognition, the scaffolds bring 'thinking' to life as the manipulation of physical objects. The scaffolds are designed to support your thinking about your paper – before you start to write it.

ChatGPT offers writing without requiring you to perform much thinking or understanding; I find this very concerning. One recent study highlights the danger offered by large language models (Gerlich, 2025). Gerlich, studying nearly 700 participants, found greater amounts of AI tool usage was associated with poorer critical thinking abilities. If a

major goal of writing is to understand, then you should avoid turning to tools like ChatGPT to create your first draft.

5.5 SCAFFOLDING WHAT TO DO NEXT

You are only as good as your next publication.

Academic writing occurs as a late step in a long process. You must choose a project, decide upon research questions, collect data, and analyze results long before writing begins.

In my lab, joy comes from having three different research projects in three different states at the same time: we have one project accepted for publication; we have submitted a second project to an editor; we have just started a third project. Maintaining lab bliss requires solving the problem of what to do next: choosing the next project to pursue after sending a completed manuscript off for review. Index cards scaffold project selection and preparation.

5.5.1 'What To Do Next?' Cards

I frequently use index cards to scaffold how I answer the question 'What project do I do next?'. I use several different kinds of index cards to decide what project I should pursue next.

First, I use *'What If?' cards*. Each 'What If?' card holds a single idea which I believe I should explore. Generally, I generate 'What If?' cards by brainstorming ideas; I usually inform my brainstorming by considering questions which come to mind from work which I have already completed. Given the context of existing work, I repeatedly answer questions like 'Wouldn't it be cool if ...?', or 'What if ...?'.

Because I use my previous work to provide context for 'What If?' cards, I scaffold my previous work with other index cards. For instance, I often use *'What Do We Know?' cards*. Each 'What Do We Know?' card expresses a fact or result which I confidently believe. A 'What Do We Know?' card also includes a brief note providing a source for the fact – a journal article, a book, one of my own papers, etc. Creating 'What Do We Know?' cards moves knowledge from my head into my world, helping me remember all the results to use to choose my next project.

I supplement my 'What Do We Know?' cards with *'What Don't We Know?' cards*, which state results which I predict, but which I have not yet established. Each card serves as a seed for a future project designed to convert an unknown into a known.

After generating the different 'What To Do Next?' cards, I think about future projects by displaying and manipulating my cards. I group related cards together. I place cards in different positions to separate highly interesting ideas from less

interesting ideas; to separate easier to pursue ideas from harder to pursue ideas; and so on.

Organizing ‘What To Do Next?’ cards helps me choose a new project. For instance, a grouping of related ‘What If?’, ‘What Do We Know?’ and ‘What Don’t We Know?’ cards can inspire, and can provide a rationale for pursuing, a new project. For instance, my cards might lead to an arrangement I interpret as: ‘I know w and x , but I don’t know y . Could I discover y by doing z ?’

I find my ‘What To Do Next?’ cards useful for keeping me on track, reminding me about the questions I want to answer and how I came up with them. I therefore display these cards in a relatively permanent layout for regular inspection. While on display, I can easily perform other actions on my ‘What To Do Next?’ cards. I can date stamp a card to indicate when I started trying to answer a question, or when I found an answer. I can use colored dots, or I can physically group cards, to indicate different lab members working on different tasks.

The final manipulation of my ‘What To Do Next?’ cards occurs when the research on the new project finishes. I organize the cards, gather them into a deck, and file them away for later use. For instance, I might find them useful to pull out when I outline the paper which describes the project’s results.

5.5.2 To Do (Right Now) Cards

After deciding what project to do next, I need to

decide what steps I need to take, decide the order in which to take them, and decide who will do them. To me, such ‘To Do’ tasks seem more urgent; I call them ‘To Do Right Now’ tasks.

Index cards offer a very useful scaffold for a project’s early stages. I write each task on its own index card, creating several ‘To Do’ cards. I usually find such cards easy to generate; I grab some blank cards and repeatedly answer the question ‘What do I need to do to start the project?’. After generating the cards, I arrange them in order – which tasks need to be done first, which tasks can be done quickly, and so on. When a task is initiated, I use a date stamp to record when the task began; the card can also be used to indicate who is completing the task. When the task is over, I date stamp the card to record when the task was completed.

‘To Do’ cards provide a useful scaffold, because they display what I need to do, my active tasks, and my completed tasks. I don’t have to keep all tasks and their states in memory. For group projects, I can put ‘To Do’ cards on public display to serve as a collective memory to guide multiple researchers working on the same project. On public display, ‘To Do’ cards motivate team members by showing progress on a project when we stamp a card ‘completed’.

5.6 SCAFFOLDS FOR KEEPING TRACK OF PROJECT PROGRESS

Step back and reflect upon what you discover.

The index cards described in Section 5.5 scaffold selecting and launching a new research project. Once a project begins, I like to track its progress using index cards. I do so by updating various ‘to do’ cards (marking them as completed, date stamping them, etc.). I also create new cards to remind me what new knowledge I have acquired or what new questions have arisen. I create my new cards in a format which I use to communicate project progress to others. Each card holds a short note about the project. The short format forces me to ‘think in paragraphs’: each summary card could be used later as a paragraph topic card in a Chapter 2 scaffold.

I found creating project summary cards particularly useful during the covid pandemic, when I had to communicate with my collaborators remotely. I created my summary cards by creating labels with a word processor. I could print the labels and stick them on index cards for my own use and could also email the word processor file to my collaborators for easy reading.

I find creating summary cards with labels handy because summarizing project progress often requires including preliminary graphs, figures, tables, or statistics. I prefer 2" by 4" labels (e.g., Avery 5163); Avery fits 10 such labels on a single sheet. I find such labels large enough for different kinds of information, but small enough to force me to think in paragraphs.

In general, when I create labels to summarize a project's current state, the labels present different information. Most labels present facts as short statements: statements about what we did, or statements about what we found. Some labels present results, such as summary statistics or statistical tests. Some labels provide small tables or figures which summarize results.

Creating summary labels forces me to reflect on a project and to carefully examine current results. I focus on what I discovered, on what I still need to discover, and on what findings seem particularly important, surprising, or novel. Thus, creating summary index cards helps me begin to think about elements I use to create a Chapter 2 scaffold: what main points should I write about, and what audience would show interest in such research? In short, each summary card can serve as a potential topic card, or as evidence required to illustrate a topic when I build my Chapter 2 scaffold. Therefore, creating summary index cards helps me bootstrap the scaffold for my outline – particularly if I stick my labels on the blank sides of index

cards, leaving their lined sides free for writing topic and concluding sentences!

5.7 SOURCE CARDS AND QUOTE CARDS

Read before you write.

5.7.1 Source Cards

Section 5.5 and 5.6 described how to use index cards to scaffold planning tasks and tracking their progress. When I begin a new project, I usually start by collecting relevant information. In particular, I find source materials – books and articles – relevant to my project, and then I read them. I use index cards to scaffold my information gathering.

When I compile a list of sources to read, I create an index card for each source. I could write out the full citation, but often I find a short note works. Why create an index card for each source? First, each card reminds me of a source I need to retrieve and read. Second, a source's index card provides a place to jot some short notes about key points I read and wish to refer to later. By adding such notes, and by stamping a finished date on the card after reading the source, index cards distinguish what I have read already from what I still need to read. Third, I can add a completed source card to my outline later. I

put a source card in the location in my scaffold where I want to refer to the source in my paper. The source card reminds me to cite the source, reminds me where in the writing project to cite the source, and reminds me what information in the source I need to mention.

After creating new source cards for a project, I suggest filing them away for use in future projects. I keep my source cards organized alphabetically in a file cabinet to consult for future projects.

5.7.2 Quote Cards

I create different kinds of source cards. Quote cards provide one useful example. To create a quote card, I take an index card upon which I write a quote which I want to include in a manuscript. Along with the quote, I write the citation to use when I include the quote (i.e., Author, Year, Page).

I usually create my quote cards while reading books. Following excellent advice (Adler & Van Doren, 1972), I make a book my own by marking the book up. Marking the book up includes underlining sentences I find interesting, important, or quotable. To create quote cards for a book, I go back to the book when I have finished reading and scan its pages, looking for underlining. I copy each underlined quote and add its citation to a new quote card.

Like source cards, quote cards scaffold memory. Placing a quote card amongst a set of topic cards reminds me to

include the quote; the quote card's position reminds me where I would like to place the quote in my manuscript. However, quote cards provide additional affordances.

First, quote cards offer portability. After creating a quote card, I find the card easier to carry around than the book from which I found the quote.

Second, quote cards offer non-repeatability. I often discover I use the same quote more than once in a book-length manuscript. I use quote cards as a remedy. If I only include a quote written on a quote card in my outline cards, and if I only have one such card for each quote, then I will only use the quote once.

Third, quote cards remind me about what I have read. As mentioned, I usually compile quote cards while reading a book. When I finish, I file the book's source card away, and I file the book's quote cards away with the source card. When I later pull the source card out, I remind myself about the book by reading my quote cards.

5.8 ZETTELKASTEN

Luhmann wrote by communicating with his slip-box.

The scaffold described in Chapter 2 creates a topics-based outline from scratch, assuming I already know quite a bit about a project, and I want to start writing. I use the Chapter 2 method to bootstrap my manuscript's outline; I do not use the method to collect or organize information long before writing can begin. However, other scaffolds exist for organizing knowledge for later use.

Richard Rhodes describes how he collects information and ideas for potential writing projects (Rhodes, 1995). Rhodes keeps pencils and blank index cards throughout his house. Whenever an idea occurs, he fills out an index card. He stores his index cards in a file he calls 'Futures' which provides materials for potential projects. "I have ten years of notes on three-by-fives toward a work of fiction I've been planning" (Rhodes, 1995, p. 32). Rhodes's system requires him to return regularly to the Futures file to read his filed cards.

Rhodes (1995) also developed a physical system for classifying cards by topic, and for removing topic-related cards from his file. He used index cards with numbered holes punched around the edge. He associated each hole with a particular topic. He coded topics covered by a particular card by

using a notching punch to remove the top of the hole which corresponded to a card's topic. The notches enabled Rhodes to remove all the cards related to a topic from his system: he lined up all his cards, threaded a knitting needle through a topic hole, and shook the stack. Those cards which had notched holes (because they were related to the topic) fell out.

Rhodes' (1995) method resembles another system for storing and organizing ideas, *zettelkasten*, which is German for 'slip-box' (Ahrens, 2022; Helbig, 2019; Kadavy, 2021; Krajewski & Krapp, 2011; Luhmann, 1992). We most closely associate *Zettelkasten* with German sociologist Niklas Luhmann (Luhmann, 1992). Beginning in the early 1950s, he began a filing system called a slip-box, which contained notes (e.g., index cards). Luhmann's slip-box would grow to hold over 90,000 notes.

Each note in a slip-box holds a single idea expressed in complete sentences. Each note has a unique identifying number which only identifies a note; the number provides no information about a note's content. Later, Luhmann browsed through his slip-box, searching for related ideas. When he discovered a relationship, he added the links to the notes. For instance, if Luhmann found a relationship between note 65 and note 97, he added the number 97 to note 65 and he added the number 65 to note 97.

While discovering links between notes in the slip-box, Luhmann also developed a parallel filing system for topics. Luhmann created a note for each topic which interested him,

labelled the note with the topic, and added the numbers of topic-related index cards in the slip-box to the topic's index card. Luhmann could then use the topic card to retrieve a set of notes, as well as linked notes, from the slip-box.

Luhmann used his *zettelkasten* to discover new ideas. He did so by pulling a note from the slip-box and by then following the note's links to other notes. He also searched for new links between existing notes by browsing through filed cards – sometimes randomly – searching for new relationships to record.

Luhmann used his *zettelkasten* to write. When he retrieved related notes, he had links between ideas laid out in a sequence, and each note in the sequence contained full sentences. “Every question that emerges out of our slip-box will naturally and handily come with material to work with” (Ahrens, 2022, p. 46).

Rhodes's Futures file and Luhmann's *zettelkasten* relate to the Chapter 2 method because both scaffold memory for ideas. Because both systems succeed by searching through filed information, both can benefit by being converted into digital file systems. Rhodes (1995) has transcribed his notes into a computer; many digital *zettelkasten* also exist (Ahrens, 2022; Kadavy, 2021).

5.9 NOTEBOOKS

A writer always carries pen and paper.

According to many books about writing, notebooks provide the most common scaffold (Brande, 1934; Cameron, 2022; Clark, 2008b; Goldberg, 1986; Janzer, 2016; Lamott, 1995; Raab, 2010). A notebook primarily scaffolds a writer's memory. Writers often carry a notebook which they use to immediately record observations and ideas before forgetting them. Hemingway constantly wrote in his notebooks while seated in Parisian cafes: "The blue-backed notebooks, the two pencils and the pencil sharpener (a pocketknife was too wasteful), the marble-topped tables, the smell of early morning, sweeping out and mopping, and luck were all you needed" (Hemingway, 1964, p. 91).

Notebooks offer additional affordances. Regularly writing in notebooks makes writing habitual. Some writers advise writing in a notebook just after awakening as a method to become more receptive to inspiration (Brande, 1934; Cameron, 2022).

Notebooks serve as essential tools for creative writing. For instance, Charles Darwin kept field notebooks during his voyage on the *Beagle* and used later notebooks to explore his ideas about evolution (Darwin et al., 1987). Darwin's note-

books, available at the website Darwin Online, provided the core material for his published works (Darwin, 1988; Darwin et al., 1962).

Notebooks for academic writing can take other useful forms. I encourage my students to use an idea log while reading shorter sources (i.e., articles or book chapters). When you finish reading a source, you enter a date into an idea log, and then add two paragraphs. The first describes the just-read content in your own words. The second describes what you find interesting in the source, or how an idea in the source relates to your other ideas.

Idea logs, as sequences of two-paragraph entries, seem like short-form notebooks. A book summary illustrates a longer form notebook. When I read a book, I use pen or pencil to make the book my own (Adler & Van Doren, 1972). I summarize key points in the margin, where I also make notes about other ideas my reading suggests. I also underline important passages for later quoting. When I finish the book, I create a summary by typing my marginalia and underlined quotes into a word processing file. I organize my summary by using the book's chapter and section titles. I have posted some of my book summaries on my website: http://www.bcp.psych.ualberta.ca/~mike/Pearl_Street/Margin/index.html.

I find any useful notebook only becomes so through regular reading and re-reading. Writing a book summary helps me consolidate a book's material right after I finish reading. However, book summaries become far more useful later, when

I reread them, because they scaffold my memory for previously consumed material.

5.10 THE SOCIAL WORLD

Your writing world contains people too.

When embodied cognitive scientists extend the mind into the world, they naturally focus on the physical objects used to scaffold cognition. Chapter 5 has almost exclusively discussed physical objects like index cards and their displays. However, you also have a social world which can also scaffold social cognition (Vygotsky, 1986).

Vygotsky (1986), for example, emphasized assistance to foster cognitive development. He defined the difference between a child's ability to solve problems without aid and their ability to solve problems when assisted as the *zone of proximal development*. Children with larger zones of proximal development did better in school. Vygotsky criticized instructional methods which required children to solve problems without help. "The true direction of the development of thinking is not from the individual to the social, but from the social to the individual" (Vygotsky, 1986, p. 36). The zone of proximal development emerges when you scaffold cognition with your social world.

Vygotsky also emphasized other social scaffolds, like language: "Real concepts are impossible without words, and

thinking in concepts does not exist beyond verbal thinking. That is why the central moment in concept formation, and its generative cause, is a specific use of words as functional ‘tools’” (Vygotsky, 1986, p. 107). Clark (1997, p. 180) agrees cognition is not only scaffolded by objects, but also by social and cultural contexts: “Advanced cognition depends crucially on our abilities to dissipate reasoning: to diffuse knowledge and practical wisdom through complex social structures, and to reduce the loads on individual brains by locating those brains in complex webs of linguistic, social, political, and institutional constraints.”

The outlining method described in Chapter 2 permits social scaffolding by moving ideas from your mind to your world. The method need not depend on a single mind. Several collaborators can generate ideas and put them on display; the index cards on display provide a collective memory. Collaborators can work together to rearrange and to elaborate index cards as well. When a social group develops an outline, the zone of proximal development of group members can expand.

Sarnecka’s (2019) writing workshops also illustrate social scaffolding for academic writing. Sarnecka’s formal writing workshop takes the form of a graduate course which incorporates reading, instruction, and writing. However, her workshop’s writing component can take many different forms. I find her write-on-site group particularly interesting. A write-on-site group has people gathering at a prearranged location to sit quietly together and do their own writing. “These groups

can be very helpful for people who feel isolated or stuck in their writing practice – there is something both comforting and energizing about writing in the quiet company of others” (Sarnecka, 2019, p. 15). In other words, the social act of ‘writing alone together’ may aid the writing process.

Another social scaffold for writing involves meeting others to talk about writing. Many books about writing describe writers meeting to discuss their writing lives with each other (Cameron, 2022; Goldberg, 1986; Hemingway, 1964; Janzer, 2016; Lamott, 1995). Such discussions include seeking encouragement, finding support for writing setbacks, or getting honest feedback about writing. “Ask a knowledgeable friend to be brutally honest. Tell him or her that you’re looking for your unique quirks” (Janzer, 2016, p. 129).

Sarnecka (2019) notes we rarely find social scaffolds, like seeking brutally honest feedback from colleagues, in academic writing. “Academics don’t talk nearly as much about writing as they talk about other research skills, such as experimental design or statistical analysis” (Sarnecka, 2019, p. 2). I hope when you realize the social world scaffolds writing you will seek others out to talk about your writing. I routinely share my own writing with students and colleagues. Academic writing improves when we use social initiatives to support academic writing.

5.11 STAND ON THE SHOULDERS OF OTHERS

Keep good advice close to hand.

Academic writers should strive to improve their craft (Becker, 2020; Germano, 2021; Kail, 2019; Sarnecka, 2019; Sword, 2012). How can we improve academic writing? I suggest standing on the shoulders of others. I try to improve my writing by reading excellent writing, including excellent writing about writing better. Fortunately, many excellent books about writing better exist (Section 1.2). I find two types of books particularly useful for improving my writing. I have my favorites in both types, and repeatedly turn to them for advice. Thus, I keep my favorites close to hand in my writing space.

The first type of book I rely upon offers specific advice about style: sentence structure, common faults, and important rules to follow (use no unnecessary words!). Such books typically include examples or exercises to help me improve my writing.

I have four favorites of the first type of book. Helen Sword wrote two; I find both extremely useful because they focus on academic writing (Sword, 2012, 2016). I repeatedly

use a third, William Zinsser's *On writing well* (Zinsser, 2006). I frequently return to the first few chapters of the 30th anniversary edition of Zinsser's book to find writing advice. *The elements of style* by William Strunk Jr., with revisions by E. B. White, (Strunk & White, 1959) serves as my fourth go to style book. I recently found a hard cover copy at a used bookstore; I often turn to its second chapter for writing advice.

Many other books of the first type exist (Baker, 1973; Barzun, 1985; Bierce & Freeman, 2009; Carpenter, 2020; Clark, 2008b, 2014; Germano, 2021; Greene, 2013; Heard, 2022; Kail, 2019; McPhee, 2017; Messenger & Taylor, 1984; Pinker, 2014; Quiller-Couch, 1916; Rosnow & Rosnow, 1998; Sarnecka, 2019; Schimel, 2012; Shertzer, 1986; Silvia, 2007; Wheelan, 2022; Williams & Bizup, 2017). You might find other books more useful than those which count as my favorites. Do not feel obligated to read what I read. Instead, find your own favorites, and add them to your writing space.

The second type of book I rely upon offers insight into the writing life or about the act of writing. Such books provide glimpses into how accomplished writers approach their craft. I find myself turning to such books when my urge to write lags. Books about the writing life reignite my desire to write.

I also keep four favorite books of the second type close at hand. I love reading Annie Dillard's beautifully written *The writing life* (Dillard, 1989). Most books which recommend books about improving and inspiring writing include Stephen King's *On writing* on their list (King, 2000). I always enjoyed

Tracy Kidder's Pulitzer prize winning book *The soul of a new machine* (Kidder, 1981). I find the account of creating such a book, described in Tracy Kidder and Richard Todd's *Good prose*, a useful addition to my writing space (Kidder & Todd, 2013). I recently discovered Richard Rhodes's *How to write* (Rhodes, 1995), which now belongs with my favorites. Rhodes's book contains many ideas about dealing with problems faced by writers.

Again, you can find many other books which belong to the second type (Brande, 1934; Cameron, 2022; Goldberg, 1986, 2021; Heighton, 2011; Hemingway, 1964; Janzer, 2016; Koch, 2003; Kumar, 2020; Marche, 2023; Orwell, 2005; Salesses, 2021; Sargent & Paraskevas, 2005; Sword, 2017; Ueland, 1938; Zinsser, 1988). I advise you, again, to find your own favorites to keep close at hand in your writing space.

Professional writers advise keeping a third type of book close at hand for repeated use. The third type of book provides reference material for both writing and content. Rhodes (1995), for example, keeps many reference books at hand, including several dictionaries (including foreign language dictionaries, a medical dictionary, a science dictionary), an anatomy text, and books of quotations.

You should keep books which improve or inform the act of writing nearby for repeated consultation. You should have other material at hand as well. Chapter 3 suggested a scaffold for generating topics using model writing which you strive

to emulate. Choose such models; add them to your writing space.

Importantly, keeping models of writing close at hand only works if you use them. You should analyze your models and mark them up as you determine *why* they work. The same point applies to the first two types of books described in the current subsection. Why do you find the books useful? Why do they inspire? Why do your favorite books demand repeated reading? Answers to such questions point you towards improved writing.

5.12 THE DIGITAL WORLD AND ITS SCAFFOLDS

Imaginary worlds offer scaffolds too.

Chapter 5 has presented various physical objects to use as scaffolds. My emphasis on physical objects comes from embodied cognition. Embodied cognition argues we can replace cognition by using physical objects and their affordances. Thinking becomes manipulating physical objects. “The great benefit of experience is that your senses gather information directly and you feel it. No collection of documents is ever as rich in felt detail as experience itself” (Rhodes, 1995, p. 60).

However, modern writing inevitably involves using a different physical object, the digital computer. Even the final step of the Chapter 2 scaffold requires you to transfer information from index cards to a word processing document.

The digital computer scaffolds writing but does so differently than other physical objects discussed in the current book (see Section 5.4.2). We do not physically manipulate a computer as a worldly object on its own. Instead, a computer

provides a new, digital, world which offers its own affordances to scaffold writing.

The most obvious digital scaffold: the word processor. Word processors offer new affordances by permitting editing of a virtual document. Books about writing which appeared during the personal computer revolution rave about the new affordances offered by word processors. “The computer is God’s gift, or technology’s gift, to rewriting and reorganizing. It puts your words right in front of your eyes for your instant consideration – and reconsideration; you can play with your sentences until you get them right” (Zinsser, 2006, p. 87).

Word processors offer other affordances by flagging spelling mistakes or poor grammar. I encourage students to pay attention to potential writing problems identified by their word processor. I myself find a word processor’s suggestions useful, even when I disagree with them.

Other computer software can help improve writing by inspecting a digital document. Many of my students recommend using Grammarly to check spelling and grammar. Helen Sword offers a free app to accompany her book *The writer’s diet* (Sword, 2016). (Sword, 2016). You can use the app directly with Microsoft Word to receive visual feedback about how closely your writing follows Sword’s general suggestions about excellent writing style.

Bibliographic software like EndNote provides another digital scaffold for writing. Such software permits you to cite sources by inserting tags into a word processing document.

You can then convert tags into proper in-text citations which populate a complete reference section. Such software removes the drudgery from checking what sources you cite or from adding a reference section by hand. Bibliographic software will also automatically reformat citations to many different standards, which I find handy when I send interdisciplinary research to different journals which use different citation formats.

Bibliographic software can also create a searchable digital library for use during a project's early stages. For instance, I have an EndNote library which holds over 5,200 different citations. I typically create an EndNote entry by importing information directly from a library database. The import usually includes an abstract along with authors, title, publication, and so on. I also frequently download a source's PDF which I save on a lab disk drive; I link the PDF to the reference stored in EndNote. As a result, I have a large electronic library on my desktop which I use to search for material when I research a project.

Other software packages scaffold academic writing. Statistical software performs data analysis and provides results to describe. Software generates figures to include in a manuscript. I sometimes use voice recognition software, particularly when dictating index cards into a word processor.

Using various software packages to scaffold writing leads, in turn, to needing an additional scaffold: an archive of files containing documents, data analyses, images, references,

and so on. I find such an archive useful because new material often begins by modifying old material. Therefore, you must take pains to transparently organize an archive, so you can quickly find old files when required. Need I remind you (or myself) to back files up?

While speaking about the digital world and its scaffolds, let me return to an issue raised earlier in Section 5.4.2: Given the modern act of writing – including the Chapter 2 scaffold – ultimately aims to move a document into the digital world, why even discuss non-digital scaffolds? I believe physical worlds and digital worlds offer different affordances. According to embodied cognition, different affordances lead to different thinking. Some affordances offered by the physical world aid writing, but software packages may not offer similar affordances.

For example, consider writing by applying pen to paper versus writing by typing on a keyboard. Each offers different affordances which impact writing. For instance, Brenda Ueland preferred typing because she could typewrite nearly as fast as she could think. “This is a help. It makes the involuntary spilling of one’s thinking more possible” (Ueland, 1938, p. 140). In contrast, Julia Cameron recognizes her writing benefits from patience, which she achieves by writing by hand. “Handwriting keeps pace with our thoughts” (Cameron, 2023, p. 95).

Perhaps, if your writing world works, you have no need to defend or change what you do: “What equipment you use

for writing doesn't matter so long as it helps you write" (Rhodes, 1995, p. 33). But a two-dimensional index card display offers vastly different affordances than those offered by a digital text on a screen. Different worlds offer different affordances; you must realize different affordances support different writing and thinking. Such awareness helps you design your writing world, a world which offers you affordances which you need.

PART VI

CITED LITERATURE

Adler, M. J., & Van Doren, C. (1972). *How to read a book* (Rev. and updated ed.). Simon and Schuster.

Ahrens, S. (2022). *How to take smart notes : one simple technique to boost writing, learning and thinking— for students, academics and nonfiction book writers* (Second ed.). Independently published by Sonke Ahrens.

American Psychological Association. (2020). *Publication manual of the American Psychological Association* (Seventh edition. ed.).

Ashby, W. R. (1956). *An Introduction To Cybernetics*. Chapman & Hall.

Atchity, K. J. (1995). *A writer's time : making the time to write* (Rev. ed.). W.W. Norton.

Baker, R., & Zinsser, W. (1987). *Inventing the truth : the art and craft of memoir*. Houghton Mifflin.

Baker, S. W. (1973). *The practical stylist* (3d ed.). Crowell.

Barzun, J. (1985). *Simple & direct : a rhetoric for writers* (Rev. ed.). Harper & Row.

Becker, H. S. (2020). *Writing for social scien-*

tists : how to start and finish your thesis, book, or article (3rd ed.). University of Chicago Press.

Berkeley, I. S. N., Dawson, M. R. W., Medler, D. A., Schopflocher, D. P., & Hornsby, L. (1995). Density plots of hidden value unit activations reveal interpretable bands. *Connection Science*, 7, 167-186.

Bierce, A., & Freeman, J. (2009). *Ambrose Bierce's Write it right : the celebrated cynic's language peeves deciphered, appraised, and annotated for 21st-century readers* (1st U.S. ed.). Walker & Co.

Block, L. (2021). *A writer prepares*. LB Productions.

Bloom, L. Z. (1981). Why graduate students can't write: Implications of research on writing anxiety for graduate education. *Journal of Advanced Composition*, 2(1/2), 103-117. <http://www.jstor.org/stable/20865491>

Bogard, J. M., & McMackin, M. C. (2012). Combining Traditional And New Literacies In A 21st-Century Writing Workshop [Article]. *Reading Teacher*, 65(5), 313-323. <https://doi.org/10.1002/trtr.01048>

Bouchard, T.J. (1971). Whatever happened to brainstorming? [Article]. *Journal of Creative Behavior*, 5(3), 182-189. <https://doi.org/10.1002/j.2162-6057.1971.tb00889.x>

Bradbury, R. (1990). *Zen in the art of writing*. Joshua Odell Editions.

Brande, D. (1934). *Becoming a writer*. J. P. Tarcher

Brooks, R. A. (2002). *Flesh And Machines: How Robots Will Change Us*. Pantheon Books.

Butler, R. O., & Burroway, J. (2005). *From where you dream : the process of writing fiction / Robert Olen Butler ; edited, with an introduction by Janet Burroway*. Grove Press.

Cameron, J. (2022). *Write for life : creative tools for every writer : a six-week artist's way program* (First edition. ed.). St. Martin's Essentials.

Carpenter, S. (2020). *The craft of science writing : selections from The Open Notebook*.

Chomsky, N. (1965). *Aspects Of The Theory Of Syntax*. MIT Press.

Chomsky, N. (1966). *Cartesian Linguistics: A Chapter In The History Of Rationalist Thought* ([1st ed.]. Harper & Row.

Chomsky, N. (1995). *The Minimalist Program*. MIT Press.

Chomsky, N., Roberts, I., & Watamull, J. (2023, March 8). The false promise of ChatGPT. *New York Times*.

Clark, A. (1997). *Being There: Putting Brain, Body, and World Together Again*. MIT Press.

Clark, A. (1999). An embodied cognitive science? *Trends in Cognitive Sciences*, 3(9), 345-351. <Go to ISI>://000082310000006

Clark, A. (2003). *Natural-born Cyborgs*. Oxford University Press. <http://www.loc.gov/catdir/toc/fy041/2002042521.html>

Clark, A. (2008a). *Supersizing The Mind: Embodiment, Action, And Cognitive Extension*. Oxford University Press.

Clark, A., & Chalmers, D. (1998). The extended mind (Active externalism). *Analysis*, 58(1), 7-19. <Go to ISI>://000073222300002

Clark, R. P. (2008b). *Writing tools : 50 essential strategies for every writer* (College ed.). CQ Press. Table of contents only <http://www.loc.gov/catdir/toc/ecip0811/2008008290.html>

Clark, R. P. (2014). *How to write short : word craft for fast times* (First edition. ed.).

Contreras Kallens, P., Kristensen-McLachlan, R. D., & Christiansen, M. H. (2023). Large Language Models Demonstrate the Potential of Statistical Learning in Language [Article]. *Cognitive Science*, 47(3), Article e13256. <https://doi.org/10.1111/cogs.13256>

Cron, L. (2012). *Wired for story : the writer's guide to using brain science to hook readers from the very first sentence* (1st ed.). Ten Speed Press.

Cron, L. (2016). *Story genius : how to use brain*

science to go beyond outlining and write a riveting novel (before you waste three years writing 327 pages that go nowhere) (First edition. ed.). Ten Speed Press.

Cron, L. (2021). *Story or die : how to use brain science to engage, persuade, and change minds in business and in life* (First edition. ed.). Ten Speed Press.

Darwin, C. (1988). *On the origin of species, 1859*. New York University Press. Contributor biographical information <http://www.loc.gov/catdir/enhancements/fy0808/88022464-b.html>

Publisher description <http://www.loc.gov/catdir/enhancements/fy0808/88022464-d.html>

Darwin, C., Barrett, P. H., Gautrey, P. J., Herbert, S., Kohn, D., & Smith, S. (1987). *Charles Darwin's notebooks, 1836-1844 : geology, transmutation of species, metaphysical enquiries*. British Museum (Natural History) ;

Cornell University Press.

Darwin, C., Engel, L., & American Museum of Natural History. (1962). *The voyage of the Beagle*. Doubleday.

Dawson, M. R. W. (1998). *Understanding Cognitive Science*. Blackwell.

Dawson, M. R. W. (2004). *Minds And Machines: Connectionism And Psychological Modeling*. Blackwell Pub. <http://www.loc.gov/catdir/toc/ecip041/2003005918.html>

Dawson, M. R. W. (2009). Computation, cognition – and connectionism. In D. Dedrick & L. Trick (Eds.), *Cognition, Computation, and Pyllyshyn* (pp. 175-199). MIT Press.

Dawson, M. R. W. (2013). *Mind, Body, World: Foundations Of Cognitive Science*. Athabasca University Press.

Dawson, M. R. W. (2014). Embedded and situated cognition. In L. A. Shapiro (Ed.), *The Routledge Handbook Of Embodied Cognition* (1 edition . ed., pp. 59-67). Routledge.

Dawson, M. R. W. (2018). *Connectionist Representations of Tonal Music: Discovering Musical Patterns By Interpreting Artificial Neural Networks*. Athabasca University Press.

Dawson, M. R. W. (2022). *What is cognitive psychology?* Athabasca University Press.

Dawson, M. R. W., Dupuis, B., & Wilson, M. (2010). *From Bricks To Brains: The Embodied Cognitive Science Of LEGO Robots*. Athabasca University Press.

Dawson, M. R. W., Medler, D. A., & Berkeley, I. S. N. (1997). PDP networks can provide models that are not mere implementations of classical theories. *Philosophical Psychology*, 10, 25-40.

Dawson, M. R. W., Perez, A., & Sylvestre, S. (2020). Artificial neural networks solve musical prob-

lems with Fourier phase spaces. *Scientific Reports*, 10(1), 7151. <https://doi.org/10.1038/s41598-020-64229-4>

Descartes, R. (1641/1996). *Meditations On First Philosophy* (Rev. ed.). Cambridge University Press. <http://www.loc.gov/catdir/description/cam027/95010664.html>

Dillard, A. (1989). *The writing life* (1st ed.). Harper & Row.

Dong, C., Li, Y., Gong, H., Chen, M., Li, J., Shen, Y., & Yang, M. (2023). A Survey of Natural Language Generation [Article]. *Acm Computing Surveys*, 55(8), Article 173. <https://doi.org/10.1145/3554727>

Dourish, P. (2001). *Where The Action Is: The Foundations Of Embodied Interaction*. MIT Press. <http://www.loc.gov/catdir/toc/fy036/2001030443.html>

Elbow, P. (1981). *Writing with power : techniques for mastering the writing process*. Oxford University Press.

Flaherty, F. (2009). *The elements of story : field notes on nonfiction writing* (1st ed.). Harper.

Flower, L. (1989). *Problem-solving strategies for writing* (3rd ed.). Harcourt Brace Jovanovich.

Fodor, J. A., & Pylyshyn, Z. W. (1988). Connectionism and cognitive architecture. *Cognition*, 28(1-2), 3-71.

Gardner, J., & O’Nan, S. (1994). *On writers and writing*. Addison-Wesley Pub. Co.

Gerlich, M. (2025). AI tools in society: Impacts on cognitive offloading and the future of critical thinking. *Societies*, 15(1), 6. <https://doi.org/https://doi.org/10.3390/soc15010006>

Germano, W. P. (2013). *From dissertation to book* (Second edition. ed.). The University of Chicago Press.

Germano, W. P. (2021). *On revision : the only writing that counts*. University of Chicago Press.

Gibson, J. J. (1979). *The Ecological Approach To Visual Perception*. Houghton Mifflin.

Goldberg, N. (1986). *Writing down the bones : freeing the writer within* (30th anniversary edition. ed.). Shambala.

Goldberg, N. (2021). *Writing Down the Bones Deck: 60 cards to free the writer within*. In. Boulder, CO: Shambala Publications.

Gordon, K. E. (1984). *The transitive vampire : a handbook of grammar for the innocent, the eager, and the doomed*. Times Books.

Greene, A. E. (2013). *Writing science in plain English*. University of Chicago Press.

Grey Walter, W. (1950a). An electro-mechanical animal. *Dialectica*, 4(3), 206-213.

Grey Walter, W. (1950b). An imitation of life. *Scientific American*, 182(5), 42-45.

Grey Walter, W. (1951). A machine that learns. *Scientific American*, 184(8), 60-63.

Grey Walter, W. (1963). *The Living Brain*. W.W. Norton & Co.

Hall, G. M. (2003). *How to write a paper* (3rd ed.). BMJ.

Hawker, L. (2015). *Take off your pants! Outline your books for faster, better, writing*. Running Rabbit Press.

Heard, S. B. (2022). *The scientist's guide to writing : how to write more easily and effectively throughout your scientific career* (2nd edition. ed.). Princeton University Press,.

Heidegger, M. (1927/1962). *Being And Time*. Harper.

Heighton, S. (2011). *Work book : memos & dispatches on writing*. ECW Press.

Helbig, D. K. (2019). Life without Toothache: Hans Blumenberg's Zettelkasten and History of Science as Theoretical Attitude. *Journal of the History of Ideas*, 80(1), 91-112. <https://doi.org/10.1353/jhi.2019.0005>

Hemingway, E. (1964). *A moveable feast*. C. Scribner's Sons.

Hildesheimer, W. (1983). *Mozart* (1st Vintage Books ed.). Vintage Books.

Hoermann-Elliott, J. (2021). *Running, thinking, writing : embodied cognition in composition*. Parlor Press.

Holmes, B., Waterbury, T., Baltrinic, E., & Davis, A. (2018). Angst about academic writing: Graduate students at the brink. *Contemporary Issues in Education Research*, 11(2), 67-72.

Howard, V. A., & Barton, J. H. (1986). *Thinking on paper* (1st ed.). W. Morrow.

Huerta, M., Goodson, P., Beigi, M., & Chlup, D. (2017). Graduate students as academic writers: writing anxiety, self-efficacy and emotional intelligence [Article]. *Higher Education Research & Development*, 36(4), 716-729. <https://doi.org/10.1080/07294360.2016.1238881>

Jacobs, B., & Hjalmarsson, H. (1999). *The quotable book lover*. Lyons Press.

Janzer, A. (2016). *The writer's process: Getting your brain in gear*. Cuesta Park Consulting.

Kadavy, D. (2021). *Digital zettelkasten: Principles, methods and examples*. Kadavy Inc.

Kail, R. V. (2019). *Scientific writing for psychology : lessons in clarity and style* (Second edition. ed.). Sage.

Katz, A. N., & Paivio, A. (1975). Imagery vari-

ables in concept identification. *Journal of Verbal Learning and Verbal Behavior*, 14(3), 284-293.
[https://doi.org/10.1016/s0022-5371\(75\)80072-7](https://doi.org/10.1016/s0022-5371(75)80072-7)

Kerouac, J. (1957). *On the road*. Viking Press.

Kidder, T. (1981). *The soul of a new machine*. Avon Books.

Kidder, T., & Todd, R. (2013). *Good prose : the art of nonfiction* (1st ed.). Random House.

King, S. (2000). *On writing : A memoir of the craft*. Scribner.

Koch, S. (2003). *The modern library writer's workshop : a guide to the craft of fiction* (Modern Library pbk. ed.). Modern Library.

Krajewski, M., & Krapp, P. (2011). *Paper machines : about cards & catalogs, 1548-1929*. MIT Press.

Kumar, A. (2020). *Every day I write the book : notes on style*. Duke University Press.

Lamott, A. (1995). *Bird by bird : some instructions on writing and life* (1st Anchor Books ed.). Anchor Books. <http://www.loc.gov/catdir/bios/random057/95010225.html>

LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436-444.
<https://doi.org/10.1038/nature14539>

Lee, S., Alfano, C., Carpenter, R., & Carpenter, R. G. (2013). *Invention in Two Parts: Multimodal*

Communication and Space Design in the Writing Center.
<https://doi.org/10.4018/978-1-4666-2673-7.ch003>

Leith, S. (2018). *Write to the point : a master class on the fundamentals of writing for any purpose.* The Experiment,.

Lippmann, R. P. (1989). Pattern classification using neural networks. *IEEE Communications magazine*, November, 47-64.

Lowes, J. L. (1927). *The road to Xanadu; a study in the ways of the imagination.* and N.Y.

Luhmann, N. (1992). Communicating with slip boxes (M. Kuehn, Trans.). In N. Luhmann & A. Kieserling (Eds.), *Universität als Milieu kleine Schriften hrsg. von André Kieserling* (pp. 53-61). Haux. <http://luhmann.surge.sh/communicating-with-slip-boxes>

Ma, H. L., Dawson, M. R. W., Prinsen, R. S., & Hayward, D. A. (2023). Embodying cognitive ethology [Article]. *Theory & Psychology*, 33(1), 42-58. <https://doi.org/10.1177/09593543221126165>

Manning, C. D. (2022). Human Language Understanding & Reasoning [Article]. *Daedalus*, 151(2), 127-138. https://doi.org/10.1162/daed_a_01905

Manning, C. D., Clark, K., Hewitt, J., Khandelwal, U., & Levy, O. (2020). Emergent linguistic

structure in artificial neural networks trained by self-supervision [Article]. *Proceedings of the National Academy of Sciences of the United States of America*, 117(48), 30046-30054. <https://doi.org/10.1073/pnas.1907367117>

Marche, S. (2023). *On writing and failure, or, On the peculiar perseverance required to endure the life of a writer* (First edition. ed.). Biblioasis.

Martin, L. J., & Kroitor, H. P. (1979). *The 500-word theme : discovery, organization, expression* (3rd ed.). Prentice Hall.

McCulloch, W. S. (1965). *Embodiments Of Mind*. MIT Press.

McPhee, J. (2017). *Draft no. 4 : on the writing process* (First edition. ed.). Farrar, Straus and Giroux.

Messenger, W. E., & Taylor, P. A. (1984). *Elements of writing : a process rhetoric for Canadian students*. Prentice-Hall Canada.

Mitchell, M., & Krakauer, D. C. (2023). The debate over understanding in AI's large language models [; Research Support, U.S. Gov't, Non-P.H.S.; Research Support, Non-U.S. Gov't]. *Proceedings of the National Academy of Sciences of the United States of America*, 120(13), e2215907120-e2215907120. <https://doi.org/10.1073/pnas.2215907120>

Modern Language Association of America.

(2021). *MLA handbook* (Ninth edition. ed.). The Modern Language Association of America,.

Nersesian, S., Sholzberg, M., Cushman, M., & Wolberg, A. S. (2022). The Journey to a Successful Illustrated Review [Review]. *Research and Practice in Thrombosis and Haemostasis*, 6(4). <https://doi.org/10.1002/rth2.12721>

Norman, D. A. (1998). *The Invisible Computer*. MIT Press.

Norman, D. A. (2002). *The Design Of Everyday Things* (1st Basic paperback. ed.). Basic Books.

Norman, D. A. (2004). *Emotional Design: Why We Love (Or Hate) Everyday Things*. Basic Books.

Orwell, G. (2005). *Why I write*. Penguin Books.

Osborn, A. F. (1948). *Your creative power: how to use imagination*. C. Scribner's Sons.

Osborn, A. F. (1953). *Applied imagination; principles and procedures of creative thinking*. Scribner.

Paivio, A. (1971). *Imagery And Verbal Processes*. Holt, Rinehart & Winston.

Pallant, C., & Price, S. (2015). *Storyboarding : a critical history*. Palgrave Macmillan.

Perez, A., Ma, H. L., Zawaduk, S., & Dawson, M. R. W. (2023). How Do Artificial Neural Networks Classify Musical Triads? A Case Study in Eluding

Bonini's Paradox. *Cognitive Science*, 47(1), e13233.
<https://doi.org/https://doi.org/10.1111/cogs.13233>

Piantadosi, S. T. (2023). *Modern language models refute Chomsky's approach to language*. *lingbuzz/007180*

Pinker, S. (2014). *The sense of style : the thinking person's guide to writing in the 21st century!* Viking.

Plotnik, A. (1982). *The elements of editing : a modern guide for editors and journalists*. Collier Macmillan.

Poincare, H. (1913). *The foundations of science* (G. B. Halsted, Trans.). The Science Press.

Prentiss, S., & Walker, N. (2020). *The science of story : the brain behind creative nonfiction*. Bloomsbury Academic/Bloomsbury Publishing, Plc.

Quiller-Couch, A. (1916). *On the art of writing; lectures delivered in the University of Cambridge, 1913-1914*. The University press.

Raab, D. (2010). *Writers and their notebooks*. University of South Carolina Press.

Rhodes, R. (1995). *How to write : advice and reflections*. William Morrow and Co.

Rosnow, R. L., & Rosnow, M. (1998). *Writing papers in psychology : a student guide* (4th ed.). Brooks/Cole.

Rumelhart, D. E., Hinton, G. E., & Williams,

R. J. (1986). Learning representations by back-propagating errors. *Nature*, 323(6088), 533-536.

Rumelhart, D. E., & McClelland, J. (1986). On learning the past tenses of English verbs. In J. McClelland & D. E. Rumelhart (Eds.), *Parallel Distributed Processing. Volume 2: Psychological And Biological Models* (pp. 216-271). MIT Press.

Salesses, M. (2021). *Craft in the real world : rethinking fiction writing and workshopping*. Catapult.

Sargent, M. E., & Paraskevas, C. C. (2005). *Conversations about writing : eavesdropping, inkshedding and joining in*. Nelson.

Sarnecka, B. (2019). *The Writing Workshop: Write More, Write Better, Be Happier in Academia*. Barbara Sarnecka.

Sawers, N. (2002). *Ten steps to help you write better essays & term papers : I wish I'd had this when I was in school!* (3rd ed.). NS Group.

Schimmel, J. (2012). *Writing science : How to write papers that get cited and proposals that get funded*. Oxford University Press.

Seals, D. R. (2023). Publishing particulars: Part 2. Tips for effective manuscript development [Review; Editorial; Research Support, N.I.H., Extramural]. *American journal of physiology. Regulatory, integrative and comparative physiology*, 324(3),

R393-R408. <https://doi.org/10.1152/ajpregu.00267.2022>

Shapiro, L. A. (2014). *The Routledge Handbook Of Embodied Cognition* (1 edition . ed.) [text]. Routledge.

Shapiro, L. A. (2019). *Embodied cognition* (Second Edition. ed.). Routledge.

Shertzer, M. D. (1986). *The elements of grammar*. Collier Macmillan.

Silvia, P. J. (2007). *How To Write A Lot* (1st ed.). American Psychological Association. Table of contents only <http://www.loc.gov/catdir/toc/ecip0618/2006023493.html>

Simon, H. A. (1969). *The Sciences of the Artificial*. MIT Press.

Simon, P. (1986). The Boy In The Bubble. On *Graceland*. Warner Bros. Records.

Sollaci, L. B., & Pereira, M. G. (2004). The introduction, methods, results, and discussion (IMRAD) structure: a fifty-year survey [Article]. *Journal of the Medical Library Association*, 92(3), 364-367. <Go to ISI>://WOS:000222581200012

Springsteen, B. (1975). Jungleland. On *Born To Run*. Columbia Records.

Stokel-Walker, C., & Van Noorden, R. (2023). The promise and peril of generative AI [Editorial

Material]. *Nature*, 614(7947), 214-216. <Go to ISI>://WOS:000928286200001

Storr, W. (2020). *The science of storytelling : why stories make us human and how to tell them better*. Abrams Press.

Strunk, W., & White, E. B. (1959). *The Elements Of Style*. MacMillan Publishing.

Swain, D. V. (1974). *Techniques of the selling writer* (New ed.). University of Oklahoma Press.

Sword, H. (2012). *Stylish academic writing*. Harvard University Press.

Sword, H. (2016). *The writer's diet : a guide to fit prose*. The University of Chicago Press.

Sword, H. (2017). *Air & light & time & space : how successful academics write*. Harvard University Press.

Taylor, D. W., Berry, P. C., & Block, C. H. (1958). Does group participation when using brainstorming facilitate or inhibit creative thinking? [Article]. *Administrative Science Quarterly*, 3(1), 23-47. <https://doi.org/10.2307/2390603>

Ueland, B. (1938). *If you want to write*. G.P. Putnam's Sons.

University of Chicago Press. (2017). *The Chicago manual of style* (17th edition. ed.). The University of Chicago Press. <http://www.chicagomanualofstyle.org/16/contents.html>

Veres, C. (2022). Large Language Models are Not Models of Natural Language: They are Corpus Models [Article]. *Ieee Access*, 10, 61970-61979. <https://doi.org/10.1109/access.2022.3182505>

Vygotsky, L. S. (1986). *Thought And Language* (Translation newly rev. and edited / ed.). MIT Press.

Walter, L., & Stouck, J. (2020). Writing the literature review: Graduate student experiences. *The Canadian Journal for the Scholarship of Teaching and Learning*, 11(1). <https://doi.org/10.5206/cjsotl-rcacea.2020.1.8295>

Weber, B. (1985, October 20 1985). The myth maker. *New York Times*, 25.

Weizenbaum, J. (1966). Eliza – a computer program for the study of natural language communication between man and machine. *Communications of the Acm*, 9(1), 36-45.

Weizenbaum, J. (1976). *Computer power and human reason*. W.H. Freeman.

Wheelan, C. J. (2022). *Write for your life : a guide to clear and purposeful writing (and presentations)* (First edition. ed.). W.W. Norton & Company.

Wiebe, R. (2016). *Where the truth lies : selected essays*. NeWest Press.

Wiener, N. (1948). *Cybernetics: Or Control And Communciation In The Animal And The Machine*. MIT Press.

Williams, J. M., & Bizup, J. (2017). *Style : lessons in clarity and grace* (Twelfth edition. ed.).

Winograd, T., & Flores, F. (1987). *Understanding Computers And Cognition*. Addison-Wesley.

Wu, J. (2011). Improving the writing of research papers: IMRAD and beyond [Editorial Material]. *Landscape Ecology*, 26(10), 1345-1349. <https://doi.org/10.1007/s10980-011-9674-3>

Young, J. W. (1940). *A technique for producing ideas*. Advertising Publications.

Zinsser, W. (1988). *Writing to learn* (1st ed.). Harper & Row.

Zinsser, W. (2006). *On Writing Well* (30th anniversary ed.). HarperCollins.

APPENDIX I: A SCAFFOLDING EXAMPLE

Does ChatGPT tell us anything about language or writing?

Large language models, like ChatGPT, provide a different scaffold for writing. Given a short prompt, ChatGPT can do the writing for you! Appendix I provides a brief essay which introduces large language models. When I created the essay, I used scaffolds described in earlier chapters. Appendix I provides a concrete example of my scaffolds.

AI.1 CASE STUDY: LARGE LANGUAGE MODELS

Is ChatGPT a writing scaffold?

Chapter 5 described various writing scaffolds but recommended avoiding one modern scaffold: large language models, such as ChatGPT. Large language models learn about a huge corpus of text (billions of words). Given a sentence, a large language model predicts what words should occur next. Given a short prompt, a large language model uses its predictive power to generate long, meaningful responses. Some scientists report using tools like ChatGPT to help edit scientific manuscripts (Mitchell & Krakauer, 2023).

However, I worry about calling a large language model a ‘scaffold’ because the model does most of the writing for you. Professors fret their students turn in exams or essays written not by themselves, but by a large language model. Scientific journals fear researchers use large language models to create research papers, or to write reviews of papers. Danger exists in such practice. (Stokel-Walker & Van Noorden, 2023) The journal *Science* refuses to publish articles created with the help of a large language model.

Table 5-1 provided one example of ChatGPT’s writing

to a prompt. Consider another example of what large language models can do. My ‘History of Modern Psychology’ students recently wrote a take-home, open-book exam in which they answered a general question with an essay. What if some students yielded to the temptation to use ChatGPT? I wondered what ChatGPT would do when prompted by my exam questions. The following paragraphs provide ChatGPT’s response to one question:

Exam Question Prompt: *“To some, the cognitive revolution against behaviorism returned psychology to its proper domain: the scientific study of consciousness. Others disagree, believing consciousness has little role in cognitive psychology. Adopt one of these two views, and argue in favor of it, using appropriate examples from the text and the lectures.”*

ChatGPT Response: *“I will argue in favor of the view that the cognitive revolution against behaviorism returned psychology to its proper domain: the scientific study of consciousness.*

Behaviorism, which dominated psychology in the early 20th century, rejected the study of mental processes and focused exclusively on observable behavior. This approach largely ignored the role of consciousness in shaping behavior and mental processes, leading to a gap in our understanding of human cognition.

The cognitive revolution, which emerged in the 1950s and 1960s, challenged the behaviorist approach and emphasized the importance of studying mental processes such as attention,

perception, memory, and reasoning. This shift marked a return to the scientific study of consciousness, as cognitive psychologists sought to understand the underlying processes that give rise to our subjective experiences.

One key example of the role of consciousness in cognitive psychology is the study of attention. Attention is a fundamental cognitive process that allows us to selectively attend to certain stimuli in our environment while ignoring others. However, attention is also a conscious experience that we are aware of and can reflect on. Cognitive psychologists have studied attention extensively, examining its neural basis, its role in perception and memory, and its relation to other cognitive processes such as emotion and decision making. These studies have shed light on both the objective and subjective aspects of attention, providing a more complete understanding of this complex cognitive process.

Another example is the study of consciousness in cognitive neuroscience. Advances in brain imaging technologies have allowed researchers to investigate the neural basis of consciousness and the mechanisms that give rise to our subjective experiences. These studies have revealed the complex interplay between brain regions involved in perception, attention, and memory, and have provided insights into the neural correlates of conscious awareness.

In conclusion, the cognitive revolution against behaviorism marked a return to the scientific study of consciousness in psychology. By emphasizing the importance of studying mental processes and their relation to conscious experience, cognitive psy-

chology has provided a more complete understanding of human cognition. The study of attention and consciousness in cognitive psychology and cognitive neuroscience exemplify the role of consciousness in cognitive psychology, supporting the view that consciousness has a central role in the scientific study of human cognition.”

Clearly, ChatGPT generated an interpretable answer to the exam question. It produced an impressive first paragraph which suggested ChatGPT understood the question raised two opposite positions and asked a student to defend one. ChatGPT chose one and then wrote an understandable essay to defend the choice. Given such ability, why did I recommend not using large language models in Chapter 5?

Appendix I addresses my understanding of ChatGPT and its cousins. To prepare, I read several articles about large language models. As I read, I began to think about how large language models relate to cognitive science. Appendix I provides an essay which communicates my thoughts to you.

While Appendix I deals with a topic you might find dry or uninteresting, I try to wrap the topic in a more enticing package. I created my essay using the scaffolds described in earlier chapters. In addition to my essay’s final version, Appendix I provides my intermediate work – the index cards I created – while outlining and drafting my essay.

Appendix I proceeds as follows. Section AI.2 presents my essay’s final version. Section AI.3 presents my initial topic

cards. Section AI.4 presents new topic cards which emerged when I refined and organized the Section AI.3 cards. Section AI.5 provides the topic sentences I created for each paragraph; Section AI.6 provides the concluding sentences. Section AI.7 provides the essay's first draft, which I created by adding supporting sentences between each topic sentence/concluding sentence pair. Section AI.8 describes how I worked to revise and polish the essay and compares the Section AI.7 first draft to the final version from Section AI.2.

AI.2 CAN LARGE LANGUAGE MODELS INFORM COGNITIVE SCIENCE?

We live during an artificial intelligence (AI) revolution driven by the *large language model* (LLM). LLMs, a kind of deep belief network, perform complicated tasks because they contain many intermediate processing layers (LeCun et al., 2015). LLMs also include features for processing language (Dong et al., 2023). LLMs learn to predict which words should come next. They achieve incredible accomplishments with such learning: “Give them a human language description or several examples of what one wants them to do, and they can perform tasks for which they were never trained” (Manning, 2022, p. 132).

LLMs can generate long, detailed, meaningful responses to short prompts. LLMs can accomplish many complex tasks, including editing scientific manuscripts, writing or checking programming code, or brainstorming ideas (Mitchell & Krakauer, 2023; Stokel-Walker & Van Noorden, 2023). OpenAI reports its most recent LLM, GPT-4, took a simu-

lated bar exam; GPT-4 placed in the top 10% of test takers. “What is clear is that these models *use language in a way that is remarkably human*” (Piantadosi, 2023, p. 4, his italics).

LLMs raise many questions in the popular press and in scholarly journals (Mitchell & Krakauer, 2023). Do LLMs understand language? Are LLMs intelligent? Are LLMs sentient or conscious? A recent New York Times headline exclaims “Microsoft Says New A.I. Shows Signs of Human Reasoning.” Scientific journals fear scientists ask LLMs to write research papers, and worry their reviewers ask LLMs to evaluate submitted papers.

A different question interests me: ‘Can LLMs inform cognitive science?’ Below, I argue LLMs can inform cognitive science – but only if researchers study an LLM’s inner workings to discover *how* the LLM generates responses.

LLMs excite interest because they generate detailed responses to short prompts. LLMs seem to handle language with human-like skill. However, researchers disagree about the relationship between LLM language processing and human language processing.

Cognitive scientists have studied human language for decades. Cognitive science’s most famous theory about language, generative grammar (Chomsky, 1965, 1966, 1995), proposes human cognition uses specialized rules to manipulate complex sentence representations. Generative grammar represents sentences as tree-like forms called phrase markers. A phrase marker encodes a sentence’s word order, the parts of

speech to which its words belong, and the sentence's hierarchical structure. Rules, called transformations, convert one phrase marker into a different phrase marker. For instance, a transformation can convert a phrase marker representing a statement into a phrase marker representing a question. For Chomsky, language – and all cognition – requires the rule-governed manipulation of symbols.

Cognitive scientists who believe human language is the rule-governed manipulation of symbols do not believe LLMs can contribute to cognitive science (Chomsky et al., 2023; Veres, 2022). Chomsky et al. point out “We know from the science of linguistics and the philosophy of knowledge that [LLMs] differ profoundly from how humans’ reason and use language. These differences place significant limitations on what these programs can do, encoding them with ineradicable defects.”

Other researchers argue LLMs provide alternatives to generative grammar (Contreras Kallens et al., 2023). UC Berkeley psychologist Steven Piantadosi agrees with Chomsky: LLMs do not use grammatical rules (Piantadosi, 2023). However, Piantadosi argues LLMs *refute* Chomskyan linguistics because of what LLMs accomplish without rules. “The success of large language models is a failure for generative theories because it goes against virtually all of the principles these theories have espoused. In fact, *none* of the principles and innate biases that Chomsky and those who work in his tradition have

long claimed necessary needed to be built into these models” (Piantadosi, 2023, pp. 14-15, *his italics*).

My own research compares theories based on rules and symbols to theories based on artificial neural networks. I recognize a historical precedent for Piantadosi’s position, a precedent relevant to determining whether LLMs can inform cognitive science. In the mid-1980s, cognitive science experienced its connectionist revolution. Cognitive scientists discovered new artificial neural networks, called multilayer perceptrons, powerful enough to model cognitive phenomena. Multilayer perceptrons contained intermediate processors called hidden units which gave them power: with enough hidden units a multilayer perceptron can learn any mapping between stimuli and responses (Lippmann, 1989).

Multilayer perceptrons caused the connectionist revolution because network proponents attacked theories which appealed to rule-governed symbol manipulation. For example, one network converted present-tense verbs into their past-tense form (Rumelhart & McClelland, 1986). Rumelhart and McClelland argued the network works without using grammatical rules: “We suggest that lawful behavior and judgments may be produced by a mechanism in which there is no explicit representation of the rule” (p. 217).

My own research focuses on problems with arguments like Rumelhart and McClelland’s (1986). Connectionist revolutionaries *assumed* networks abandoned symbols and rules, *but never provided evidence* to support their assumption, or to

show *how* their networks *replaced* symbols and rules. Instead, by assuming networks differed from traditional theories, when they trained networks to perform ‘symbolic tasks’ (like converting verb tenses) they claimed, ‘gee whiz we have a non-symbolic model of the task’. I call such work *gee whiz connectionism* (Dawson, 2009). I distance myself from gee whiz connectionism when I train networks but then analyze them to discover how my networks actually work.

When I look inside my trained networks, I discover symbol-like properties. For example, I trained one network to solve different logic problems. Inside the network I found logical rules like those taught to philosophy students (Berkeley et al., 1995). In another study, I trained networks to classify mushrooms as being edible or poisonous. Inside the network I discovered a traditional symbol/rule system called a production system (Dawson et al., 1997). My results reveal surprising similarities between network models and symbolic models, blurring the distinctions between the two (Dawson, 1998, 2004, 2013, 2018).

My students and I do not always find structures which replicate current symbolic theories. Instead, we often find new structures, symbolic in nature, but which differ from current proposals about symbols and rules.

For example, I train artificial neural networks to make musical judgements. Inside my networks I find structures strongly related to traditional music theory (Dawson, 2009, 2018; Dawson et al., 2020; Perez et al., 2023). However, my

network structures depart from traditional music theory in surprising ways.

Traditional music theory treats Western music as consisting of twelve different pitch-classes (C, C#, B, and so on). In contrast, my musical networks treat Western music as consisting of only six different pitch-classes. My networks treat pitch-classes which measure six semitones apart in traditional theory (such as C and F#) as being identical. In short, when I look inside my networks, I find alien – but formal – music theory.

My research makes me believe LLMs will only inform cognitive science when researchers stop simply assuming LLMs differ from rule and symbol models, and instead start studying the similarities and differences between LLMs and cognitive theories.

Why must we look inside LLMs to inform cognitive science? Cognitive scientists know methods completely different from human cognition can produce human-level performance. Consider Joseph Weizenbaum's program ELIZA (Weizenbaum, 1966). ELIZA conversed with humans, but Weizenbaum did not build language understanding into his program. "ELIZA shows, if nothing else, how easy it is to create and maintain the illusion of understanding, hence perhaps of judgment deserving credibility. A certain danger exists there" (Weizenbaum, 1966, pp.42-43). Examples like ELIZA show why cognitive scientists should compare *processes*, not *performance*.

I suspect LLMs processes differ dramatically from human cognition. LLMs use representations unrelated to any proposed by cognitive scientists. LLMs do not use complete sentences or individual words. Instead, they break sentences into tokens, components smaller than individual words. A token is represented as a hundreds-dimensional number vector; LLMs, assign similar vectors to related tokens. If LLM representations differ from human cognition, then LLMs do not refute Chomsky's approach. Instead, they refute using Chomsky's approach to explain LLMs!

LLM advocates recognize their models possess hidden procedures which manipulate language, but also realize the difficulties faced when searching for such procedures. "It becomes clear that it could be hard to determine what is going on, even though *the theory is definitely in there*" (Piantidosi, 2023, p. 8, his italics). To inform cognitive science, to defend claims like 'LLMs refute Chomsky', researchers must do the hard work to discover what methods LLMs use.

LLMs' size and complexity make the work hard. OpenAI's ChatGPT learns by adjusting roughly 175 billion different parameters. Another LLM, BERT, contains twelve different large layers of intermediate processors. Mitchell and Krakauer (2023, p. 1) note "the inner workings of these networks are largely opaque; even the researchers building them have limited intuitions about systems of such scale." Piantidosi (2023, p. 8, his italics) concurs: "In fact, we don't deeply understand *how* the representations these models create work."

Fortunately, researchers can create new techniques to understand a LLM's internal structure. Consider Christopher Manning's work (Manning, 2022; Manning et al., 2020). Manning probes an LLM's structure to determine whether the network represents structures found in generative grammar.

For example, Manning et al. (2020) examined LLM components called attention heads. An attention head determines the importance of one word in a sentence to other words in the sentence, or to words in the output being generated by the LLM. Related words receive higher attention. Manning et al. found attention strength captured linguistic properties. Higher attention linked objects to appropriate verbs, linked prepositions to appropriate objects, linked noun premodifiers to appropriate nouns, and so on.

Manning et al. (2020) also used a structural probe method to detect phrase marker trees represented by an LLM's processors. They measured the distance between different vectors represented in the network. Items whose vectors are close together in the LLM's space are also close together in a sentence's phrase marker. Manning et al. reconstructed phrase markers from network properties. In short, "these models learn and represent the syntactic structure of a sentence" (Manning, 2022, p. 131).

I hope similar research is on the horizon. As researchers explore LLM representations, and study how LLMs use representations to generate responses, we move closer to comparing LLMs to human cognition.

AI.3 INITIAL TOPIC CARDS FOR THE ESSAY

A paper begins as a collection of ideas you could write about.

I created the Section AI.2 essay by using the scaffold detailed in Chapter 2. I began by preparing – I needed material to give me ideas before starting the essay. I knew the main point I wanted to explore: the relationship between large language models (like ChatGPT) and cognitive science. However, I needed to learn a bit more about large language models; I also needed to determine what others might have written about my intended topic.

I prepared for a week, using Clarivate's Web of Science database to collect relevant articles (which I then read). I also searched through articles published in the New York Times. I remembered the Times had published a series which introduced large language models to the general public, and I knew the Times frequently published articles with alarming headlines concerning new developments in AI.

As I read, I related the new material to my own ideas about cognitive science. Potential topics which I could write about were popping into my head. With my reading finished, I felt prepared to start writing.

However, for me, starting to write is actually starting

to outline. My outlining process follows the method detailed in Chapter 2. My first step was generating broad topics (Section 2.6).

I performed the first step by taking blank index cards to use to jot down ideas as they came to mind. I began by writing down the basic thread which I planned to communicate (the first four entries in the left column of Table AI-1). For my essay about large language models and cognitive science I found my topics came to mind in related groups. For instance, I first generated some general topics related to how modern AI is being received, and then found myself generating ideas concerning particular properties of an example of modern AI, ChatGPT. As topics arose in related groups, I wrote a title to classify a group of topics on a separate index card.

Table AI-1 provides the topics which I jotted down on my index cards. Creating the topics required about a half hour on May 26, 2023. I laid out related topic cards in a column on a coffee table. I needed to see the topics I had already generated. The first card in each column was my title for related topics. I sat on my sofa generating topics for a while, until topic ideas dried up. I then spent some time looking at the cards I had created. On occasion looking at the cards reminded me of topics which I could add. In particular, I realized – after looking at my topics – I could talk about the evidence cognitive science used, evidence motivated by theory in cognitive science. So, the last topics I generated were the ‘Theory’ topics at Table AI-1’s end.

Table AI-1 lists the topics I generated in the order I generated them. The table can be read by reading down its columns, first the left column then the right one.

After creating my topic cards, two observations became apparent to me. First, I felt I had generated topics by ‘thinking in paragraphs’, because I felt capable of using a paragraph to express the idea jotted down on many different topic cards. Second, I realized I needed to prune topics. My plan was to write a fairly short essay. However, Table AI-1 shows I produced nearly 60 topics. If I was indeed thinking in paragraphs, and needed a paragraph to express each topic, then my planned essay would be far too long. As I moved to my scaffold’s next stages, I realized my primary task was to prune topics which I didn’t have space to include.

Table AI-1 The initial set of topics generated to outline an essay which explored how large language models might inform cognitive science.

Overall Thread Cards (May 26 2023)

Title: From ChatGPT to Cognitive Science

Point 1: ChatGPT leads to gee whiz connectionism

Point 2: For ChatGPT to inform cognitive science, we need detailed interpretation of its internal structure

Introductory Topics (May 26 2023)

Modern AIs perform amazing stuff

Modern AIs are generating lots of excitement and press

Lots of fear about modern AIs

Fear of AI expressed as scary questions

Fear of AI due to its performance

AI performance leads to less scary (more boring) question: what does it say about cognitive science?

Radical answer to this question: Piantadosi – modern AI refutes all of Chomsky!

But radical cognitive science answers are just new gee whiz connectionism

Westbury question at candidacy exam – don't need to look inside ChatGPT to inform cognitive science

My point: modern AI can inform cognitive science, but only if we look past performance and interpret internal structure

Overall Thread Cards (May 26 2023)

Title: From ChatGPT to Cognitive Science

Table AI-1 The initial set of topics generated to outline an essay which explored how large language models might inform cognitive science.

Point 1: ChatGPT leads to gee whiz connectionism

ChatGPT Topics (May 26 2023)

What is NLP?

What is ChatGPT

What does ChatGPT do?

What new ideas give ChatGPT power?

What is ChatGPT trained on?

What is ChatGPT's architecture?

How big is ChatGPT?

Classical Linguistics Topics (May 26 2023)

What is classical cognitive science?

Chomsky as classical prototypes

Syntax vs semantics; grammar; universal ; learning

Gee Whiz Connectionism Topics (May 26 2023)

Connectionism is old

1980s connectionist revolution due to network power (multilayer perceptron)

Connectionist revolution attacked classical rules

Problem: revolution assumed nets don't have rules

Musical networks capture informal (i.e., no rules)

Table AI-1 The initial set of topics generated to outline an essay which explored how large language models might inform cognitive science.

But if you look inside nets you find rules (logic network, mushroom network)

But if you look inside musical networks you find formal music theory

Moral: look inside networks to confirm the revolution

Kicker Topics (May 26 2023)

P. claims ChatGPT refutes Chomsky

Refutation: ChatGPT is better than any linguistics theory, but ChatGPT does not use rules

Issue: P. is doing gee whiz connectionism

Look inside ChatGPT: find hierarchies (Manning) – but P. misinterprets this result

What else will be found inside? ChatGPT is big!

ChatGPT may or may not refute Chomsky – to know, have to look inside

Theory Topics (May 26 2023)

Levels of analysis

Computation vs algorithm

Bonini's paradox

Strong vs weak equivalence

Theory vs technology

Humans: first and second order effects

Networks: look inside

Table AI-1 The initial set of topics generated to outline an essay which explored how large language models might inform cognitive science.

Theory Topics (May 26 2023)

Levels of analysis

Bonini's paradox

Strong vs weak equivalence

Theory vs technology

Computation vs algorithm

Humans: first and second order effects

Networks: look inside

Computation vs algorithm

AI.4 REFINING TOPIC CARDS FOR THE ESSAY

Kill your topics before they become needless paragraphs.

The next three steps in developing the scaffold are organizing and evaluating topics (Section 2.7); enhancing existing topic cards (Section 2.8); and converting subtopics into paragraph topics (Section 2.9). My next step was to carry these operations out. However, I felt when I generated topics I was already ‘thinking in paragraphs’, so I emphasized organizing and evaluating my topics. I believed most of my topic cards already expressed paragraph topics, so I had little need to enhance topics (i.e., dividing topics into subtopics) or to convert subtopics into paragraph topics. Nevertheless, the Table AI-1 topics were not organized in proper order; I needed to evaluate my topics and I had to remove many.

I conducted topic organization and evaluation in three phases. The first occurred on May 28, 2023. I took the Table AI-1 topic cards and arranged them in an order which made sense to me. I then placed the cards in a deck and worked through the deck from top to bottom. I looked at the top card on the deck and decided whether I needed to include the topic in the essay. (Remember, Table AI-1 sent me a clear signal: prune!) If I felt the topic card was necessary, I wrote the topic

(with some possible editing) on a blank index card; I was creating a new stack of index cards. If I felt the topic card was unnecessary, then I moved to the next card in my Table AI-1 deck.

At the end of the first pass of processing, I had generated a smaller deck of topic cards – my new set included only 22 topics. I laid the new cards out on my coffee table and evaluated their order. Happy with the narrative, I placed them into a deck and put the cards away.

I conducted my second pass the next morning. I took my 22 topic cards and read them in order. As I read the cards, I treated each as expressing a paragraph topic, and I began to think about what the paragraph might say. I found myself satisfied with the narrative until I reached Card 18. At the time, Card 18 was followed by the card numbered 23 in Table AI-2. I had difficulty linking the two paragraph topics. So, I added a new card – the version of Card 19 which is crossed out in Table AI-2. I felt the new topic fixed the narrative. I then numbered all the topic cards because I felt they were ready to be used as cues for writing topic sentences.

When I moved to writing topic sentences later the same afternoon, I discovered narrative problems when I hit the newly added Card 19. I realized I needed to provide more information to make my narrative clearer. Such a discovery indicates I should have expended more effort on the two steps I mostly ignored, refining topics and expanding topics into para-

graph topics. Clearly, I had not thought in paragraphs while generating all of my topics.

Discovering my problem, I retreated from writing topic sentences, and returned to developing topic cards. I replaced Card 19 with a new topic, and I then added three additional topic cards to fix my narrative. These topic cards were created during the afternoon of May 29. With the new topic cards, I was able to finish developing topic sentences for my essay's paragraphs.

The second phase of scaffolding accomplished two goals. First, I reduced the number of topics to permit me to write a shorter essay. By evaluating my topics, I eliminated about half of the topics which appear in Table AI-1. Second, I organized the remaining topics into a solid narrative order. With topics (i.e., paragraph topics) in the correct order, the topics can be converted into sentences.

Table AI-2. Topic cards after organizing and evaluating the Table AI-1 topics. The first pass created most of the cards provided in the table. The second pass added one new card (the crossed-out version of Card 19) but did not include the four cards (19-21) written in bold font. The cards were numbered during the second pass. The third pass removed the old Card 19 and added the four topic cards in highlighted rows.

1. What is a large language model (LLM)?
2. LLMs can do amazing things – code, chat, see examples from Mitchell & K., Stokel-Walker
3. LLM abilities lead to big questions: Do they understand language? Are they intelligent? Hotly debated – Mitchell & K
- 3B. NY Times headline May 16 2023: ‘Microsoft Says New A.I. Shows Signs of Human Reasoning’
4. For me, more productive question: How might LLMs inform cognitive science?
5. My answer: Can only inform cognitive science if we look past performance and examine internal structure
6. Roadmap here? Decide later
7. LLMs are exciting because of their facility with language
8. In cognitive science, language is explained by appealing to symbols and rules – Chomsky NY Times essay?
9. Chomsky argues LLMs can’t inform cognitive science because LLMs don’t use rules. Veres agrees?
10. Piantadosi agrees LLMs don’t use rules like Chomsky – but then claims this means LLMs refute Chomsky
11. Piantadosi position has historical precedent
12. 1980s PDP nets cause connectionist revolution due to new powerful ANNs

Table AI-2. Topic cards after organizing and evaluating the Table AI-1 topics. The first pass created most of the cards provided in the table. The second pass added one new card (the crossed-out version of Card 19) but did not include the four cards (19-21) written in bold font. The cards were numbered during the second pass. The third pass removed the old Card 19 and added the four topic cards in highlighted rows.

13. Connectionist revolution attacked theories based on symbols and rules – see ‘What is cognitive psychology?’
14. My concern: connectionist revolutionaries assumed no rules, but didn’t look inside to support the claim – gee whiz connectionism
15. When I looked inside, I saw lots of formal structure (logic network, mushroom network examples)
16. But – network formal structure can be novel and informative (musical network examples)
17. Moral: to inform cognitive science, don’t gee whiz, look inside
18. Consequence: LLMs may indeed inform cognitive science, but only after you figure out how they work
19. Need to look inside because LLMs use very different representations (May 28 added, May 29 deleted)
19. Why do we need to look inside? Different methods produce same behavior; we need to look at methods
20. We know methods are different, because LLMs use novel representations.
21. We need to focus on methods, too, because performance is too tempting – mistake to speculate on child language learning
22. LLM proponents know theory is there – need to figure it out and see if the theory applies to humans

Table AI-2. Topic cards after organizing and evaluating the Table AI-1 topics. The first pass created most of the cards provided in the table. The second pass added one new card (the crossed-out version of Card 19) but did not include the four cards (19-21) written in bold font. The cards were numbered during the second pass. The third pass removed the old Card 19 and added the four topic cards in highlighted rows.

23. Problem: LLMs are big – go to ChatGPT manual, Mitchell & Krakauer quote

24. Researchers are developing interesting new techniques to explore innards of LLMs – manning emerging hierarchies paper

25. More work like Manning's required

26. Do LLMs refute Chomsky? Don't know now – look inside to find out!

AI.5 CREATE TOPIC SENTENCES FOR EACH TOPIC CARD IN THE ESSAY

Start your paragraph with a sentence which states the paragraph's topic.

The next step in developing my scaffold was to write a topic sentence for each paragraph topic (Section 2.11). When creating topic sentences, I perform the first 'proper writing' in my manuscript, because so far I have avoided complete sentences. However, I also realize my first sentences are first drafts. I try not to waste time writing perfect sentences. I try to generate topic sentences as quickly as possible because I will improve my wording later.

I wrote topic sentences for my essay as follows: I took each topic card in order. I read the topic on the blank side of the card. I turned the card over, and on the top of its (lined) side I wrote a topic sentence. I did so while keeping the card's topic in mind; my goal was to write a complete sentence which expressed the topic. Thus, the topic on one side of the index

card scaffolded creating the topic sentence on the card's other side.

As noted earlier, a paragraph's topic sentence supports a manuscript's narrative structure. After creating topic sentences, I felt I could obtain a stronger sense of my paper's narrative by reading the topic sentences in order. When I read my first set of topic sentences for my essay, I found my narrative disappeared toward the end. As described in Section AI.4, my problem narrative made me return to working with topic cards; I removed one topic card and added four new topic cards in the same location, before I was satisfied with my narrative.

Table AI-3 provides, in order, each topic sentence for each paragraph in my essay – at least for the version of the essay represented in my topic cards:

Table AI-3. The topic sentence generated for each topic card (each paragraph) in the essay.

Paragraph	Topic Sentence
1	We live in an artificial intelligence (AI) revolution fueled by a new invention called a large language model (LLM).
2	LLMs are revolutionary because they can generate long, detailed, meaningful responses to short text prompts.
3	LLMs’ performance has generated many questions in both the popular press and scholarly journals.
4	Speaking as a cognitive scientist, I feel such questions miss the key point. I am interested in the question ‘Can LLMs inform cognitive science?’.
5	Below, I argue LLMs may indeed be able to inform cognitive science – but only if researchers expend considerable effort to study the internal structure of LLMs in order to discover how LLMs produce their amazing behavior.
6	Modern AI’s excitement and controversy comes from an LLM’s ability to generate paragraphs of meaningful sentences in response to short prompts or questions.
7	Cognitive science has studied human language for decades. Cognitive science’s most influential account proposes human language involves specialized rules or processes manipulating complex mental representations of sentences.
8	Cognitive scientists who believe human language is the rule-governed manipulation of symbols do not believe LLMs inform cognitive science.

Table AI-3. The topic sentence generated for each topic card (each paragraph) in the essay.

9	UC Berkeley psychologist Steven Piantadosi agrees LLMs do not use grammatical rules. However, he then proceeds to argue an LLM's high level performance without using rules refutes Chomskyan linguistics.
10	My own research concerns cognitive science's foundations, with particular interest in the relation between theories based on rules and symbols and theories based on artificial neural networks. I therefore recognize as historical precedent for Piantadosi's position.
11	In the mid-1980s, cognitive science found itself in the midst of what is now called its connectionist revolution.
12	The rise of multilayer perceptrons in cognitive science caused a revolution because proponents of artificial neural networks attacked traditional theories which appealed to the rule-governed manipulation of symbols.
13	My interest in the connectionist revolution focused on a curious aspect of the revolutionaries' argument: they assumed networks abandoned symbols and rules, but never provided evidence to support their assumption, or to show what their networks used to replace symbols and rules.
14	Surprisingly, when I looked inside my trained networks, I discovered lots of structures which resembled theories based on symbols and rules.
15	Importantly, I did not usually find network structure which replicated existing formal theories. Instead, I found new formal structures which could inform a cognitive science based on symbols and rules.

Table AI-3. The topic sentence generated for each topic card (each paragraph) in the essay.

16	I feel my research on interpreting the structure of artificial neural networks demonstrates the peril of gee whiz connectionism.
17	The moral of my story about my own work is my suspicion LLMs will only inform cognitive science when researchers abandon mere assumptions about what makes LLMs different from rule and symbol models, and instead see evidence about both the similarities and differences between both types of models.
18	Why must we look inside LLMs to inform cognitive science? Cognitive scientists have long known psychologically plausible performance can be produced by methods completely unrelated to the processes of human cognition.
19	Indeed, I strongly suspect LLMs use methods radically different from those used by humans because they represent stimuli and responses with encodings unrelated to any proposed by cognitive scientists.
20	Because I suspect LLMs use methods unrelated to human cognition, I also believe Piantadosi’s (2023) claim ‘LLMs refute Chomskyan linguistics’ illustrates the danger of being seduced by network performance. Again, cognitive scientists recognize human-like performance is not sufficient to establish human-like processing.
21	LLM proponents recognize LLMs use some method to produce a remarkable facility with language.
22	However, understanding how LLMs convert stimuli into responses is extremely challenging, because LLMs are intimidatingly large and complex systems.

Table AI-3. The topic sentence generated for each topic card (each paragraph) in the essay.

23	Fortunately, researchers recognize the need to extract potentially novel theories or representations from LLMs and are developing new techniques to understand a LLM's internal structure.
24	My hope is more work of this sort is on the horizon.
25	Piantadosi (2023, p. 30) claims “large language models rewrite the philosophy of approaches to language. Do LLMs refute Chomsky’s approach? Do LLMs represent a new connectionist revolution for cognitive science? I believe we can’t answer such questions – yet.

AI.6 ADD CONCLUDING SENTENCES TO THE ESSAY

Concluding sentences link the current paragraph's topic to the topic of the next paragraph.

My next step in developing the scaffold was to add concluding sentences to each paragraph's index card (Section 2.12). Concluding sentences solidify a manuscript's narrative structure. A concluding sentence relates to the topic sentence of its own paragraph, but also relates to the topic sentence of the next paragraph.

I created concluding sentences for my essay by paying attention to the topic sentence on the index card I was about to modify while also paying attention to the topic sentence written on the next card. The two topic sentences provide powerful constraints on what concluding sentence can be written, and I tried to pay attention to the constraints as much as possible. However, I still remembered I was writing a first draft. So, rather than writing the best concluding sentence possible, I tried to speed things up by writing a reasonable sentence. I wrote the concluding sentence on the lower part of the lined

side of the current index card, then reached for the next card and repeated the process.

After writing a concluding sentence on an index card, I could add additional material – very short notes – in the space between the card’s topic sentence and index sentence. The notes I added were reminders to include particular material when I added supporting sentences later. During my reading and my outlining, I created additional cards. For example, when I read papers on large language models, I created quote cards by writing quotes down on pink index cards. I placed a quote card after the paragraph card which cited the quote; the note on the paragraph card reminded me to look ahead in my cards to retrieve the quote.

The paragraphs which follow provide the topic sentence and the concluding sentence for each paragraph card. Any material in bold font provides additional notes which I added to the index card to help write supporting sentences.

After I created my concluding sentences, I went through the cards one-by-one to read the manuscript in outline form. The manuscript’s meaning should be communicated by having a topic sentence and a concluding sentence for each paragraph. When reading the sentences, I identified additional material which did not seem to work. I have left the material in the following paragraphs but have crossed the words out to indicate I discarded the index card before fleshing out the first draft.

A scaffold created with the Chapter 2 method pro-

duces most of a paper before any real writing begins. The paragraphs which follow represent the outline which I moved from my index cards into my word processor (Section 2.14). The paragraphs below contain 1268 words. Given my goal was to write a short essay, my sense was I already had most of my essay – before I had even created a first draft. Table AI-4 provides my paragraph by paragraph outline.

Table AI-4. The outline (topic sentence, notes, concluding sentence) for each paragraph in the essay. Notes are indicated by bold text between topic and concluding sentences.

Paragraph	Outline
1	We live in an artificial intelligence (AI) revolution fueled by a new invention called a <i>large language model</i> (LLM). Give basic properties. LLMs are trained on a huge amount of text taken from the internet and learn to predict which new words should follow those presented to an LLM as a stimulus.
2	LLMs are revolutionary because they can generate long, detailed, meaningful responses to short text prompts. Give examples. “What is clear is that these models <i>use language in a way that is remarkably human</i>” (Piantadosi, 2023, p. 4, his italics).
3	LLMs’ performance has generated many questions in both the popular press and scholarly journals (M & K ref). Example questions, NYT headline. Such questions appear to be very polarizing; Mitchell and Krakauer report 51% of scholars believe LLMs understand language.
4	Speaking as a cognitive scientist, I feel such questions miss the key point. I am interested in a different question: ‘Can LLMs inform cognitive science?’
5	Below, I argue LLMs may indeed be able to inform cognitive science – but only if researchers expend considerable effort to study the internal structure of LLMs in order to discover how LLMs produce their amazing behavior. LLMs may provide new theories to cognitive science, but only if we look for a new theory inside an LLM.

Table AI-4. The outline (topic sentence, notes, concluding sentence) for each paragraph in the essay. Notes are indicated by bold text between topic and concluding sentences.

6	Modern AI’s excitement and controversy comes from an LLM’s ability to generate paragraphs of meaningful sentences in response to short prompts or questions. 304 eg? LLMs consistently generate long, well-written, interpretable and surprising responses to short, vague prompts.
7	Cognitive science has studied human language for decades. Cognitive science’s most influential account proposes human language involves specialized rules or processes manipulating complex mental representations of sentences. Such a theory is called a generative grammar. Give eg of components – phrase marker, transformation. Chomsky’s generative grammar, which explained language by appealing to special rules and symbolic structures, not only transformed linguistics but also shaped the theories in a broader discipline, cognitive science, for many decades.
8	Cognitive scientists who believe human language is the rule-governed manipulation of symbols do not believe LLMs inform cognitive science. Quote from Chomsky essay. Chomsky’s position mirrors a motto often stated by my own PhD supervisor, Zenon Pylyshyn: no cognition without computation. The motto claims we can only explain cognition by appealing to symbols and rules.

Table AI-4. The outline (topic sentence, notes, concluding sentence) for each paragraph in the essay. Notes are indicated by bold text between topic and concluding sentences.

9	UC Berkeley psychologist Steven Piantadosi agrees LLMs do not use grammatical rules. However, he then proceeds to argue an LLM's high level performance without using rules refutes Chomskyan linguistics. "The success of large language models is a failure for generative theories because it goes against virtually all of the principles these theories have espoused. In fact, none of the principles and innate biases that Chomsky and those who work in his tradition have long claimed necessary needed to be built into these models" (Piantadosi, 2023, pp. 14-15, <i>his italics</i>).
10	My own research concerns cognitive science's foundations, with particular interest in the relation between theories based on rules and symbols and theories based on artificial neural networks. I therefore recognize as historical precedent for Piantadosi's position on Chomskyan theory. I argue below the historical precedent is relevant to answering the question of whether LLMs can inform cognitive science.
11	In the mid-1980s, cognitive science found itself in the midst of what is now called its connectionist revolution. The new networks, called multilayer perceptrons, were powerful enough to serve as theories about human cognitive phenomena. R and M quote about hidden units.
12	The rise of multilayer perceptrons in cognitive science caused a revolution because proponents of artificial neural networks attacked traditional theories which appealed to the rule-governed manipulation of symbols. Grab stuff from What is Cognitive Psychology.

Table AI-4. The outline (topic sentence, notes, concluding sentence) for each paragraph in the essay. Notes are indicated by bold text between topic and concluding sentences.

13	My interest in the connectionist revolution focused on a curious aspect of the revolutionaries' argument: they assumed networks abandoned symbols and rules, but never provided evidence to support their assumption, or to show what their networks used to replace symbols and rules. I call their approach gee whiz connectionism (ref). I tried to distance myself from gee whiz connectionism by training multilayer perceptrons on various tasks, and by conducting detailed analyses of the internal structure of my trained networks.
14	When I looked inside my trained networks, I discovered lots of structures which resembled theories based on symbols and rules. Logic network. Mushroom network. I believed my results revealed surprising similarities between network models and symbolic models, blurring the distinctions between the two approaches.
15	Importantly, I did not usually find network structure which replicated existing formal theories. Instead, I found new formal structures which could inform a cognitive science based on symbols and rules. Music network example. In short, when I looked inside my networks, I found new kinds of formal structures for cognitive science to explore.
16	I feel my research on interpreting the structure of artificial neural networks demonstrates the peril of gee whiz connectionism.

Table AI-4. The outline (topic sentence, notes, concluding sentence) for each paragraph in the essay. Notes are indicated by bold text between topic and concluding sentences.

17	The moral of my story about my own work is my suspicion LLMs will only inform cognitive science when researchers abandon mere assumptions about what makes LLMs different from rule and symbol models, and instead seek evidence about both the similarities and differences between both types of models.
18	Why must we look inside LLMs to inform cognitive science? Cognitive scientists have long known psychologically plausible performance can be produced by methods completely unrelated to the processes of human cognition. ELIZA example. Examples like ELIZA show why cognitive scientists are more concerned about comparing processes than comparing performance.
19	Indeed, I strongly suspect LLMs use methods radically different from those used by humans because they represent stimuli and responses with encodings unrelated to any proposed by cognitive scientists. Representation eg? If LLM representations are unrelated to human cognition, then LLMs do not refute Chomsky's approach. Instead, they refute the applicability of Chomsky's approach to the explanation of LLMs!

Table AI-4. The outline (topic sentence, notes, concluding sentence) for each paragraph in the essay. Notes are indicated by bold text between topic and concluding sentences.

20	Because I suspect LLMs use methods unrelated to human cognition, I also believe Piantadosi’s (2023) claim ‘LLMs refute Chomskyan linguistics’ illustrates the danger of being seduced by network performance. Again, cognitive scientists recognize human-like performance is not sufficient to establish human-like processing. Language development eg – NETTALK? An LLM’s high performance signals potential relevance to human cognition. However, the signal means ‘look inside the LLM to see how it functions’.
21	LLM proponents recognize LLMs use some method to produce a remarkable facility with language. Piantadosi quote. To inform cognitive science, to defend claims like ‘LLMs refute Chomsky’, researchers must do the hard work to discover what methods LLMs use, and to compare the discovered methods to those discovered by research on human cognition.
22	However, understanding how LLMs convert stimuli into responses is extremely challenging, because LLMs are intimidatingly large and complex systems. How big is ChatGPT; how big is its training set? Mitchell and Krakauer (2023, p. 1) note “the inner workings of these networks are largely opaque; even the researchers building them have limited intuitions about systems of such scale.” Piantadosi (2023, p. 8, his italics) concurs: “In fact, we don’t deeply understand how the representations these models create work.”

Table AI-4. The outline (topic sentence, notes, concluding sentence) for each paragraph in the essay. Notes are indicated by bold text between topic and concluding sentences.

23	Fortunately, researchers recognize the need to extract potentially novel theories or representations from LLMs and are developing new techniques to understand a LLM’s internal structure. Manning et al. examples. End with Manning 2022 p. 131 quote.
24	My hope is more work of this sort is on the horizon. As researchers explore LLM representations, as well as how the representations are used to generate responses, we move closer to relating LLM to human cognition.
25	Piantadosi (2023, p. 30) claims “large language models rewrite the philosophy of approaches to language. Do LLMs refute Chomsky’s approach? Do LLMs represent a new connectionist revolution for cognitive science? I believe we can’t answer such questions – yet.

AI.7 CREATE THE FIRST DRAFT OF THE ESSAY

Add supporting sentences to each paragraph to convert the outline into a draft.

The paragraphs at the end of Section AI.6 represent my essay's outline. The outline contains topic sentence/concluding sentence pairs, strung together in a narrative order. The outline includes several notes reminding me of additional points or material to insert.

The Section AI.6 outline began as a set of index cards; each card has a paragraph topic on one side, and a pair of sentences on the other. The completed set of index cards represents the end of using the Chapter 2 scaffold. I next entered the index card information into a word processor (creating, for instance, the paragraph sequence which ended Section AI.6).

To create my essay, my next step was to convert the outline into a first draft. I used the word processor to add supporting sentences between each topic sentence/concluding sentence pair (Section 4.2). When supporting sentences were added to every paragraph, the first draft was complete.

The outline provided a scaffold which constrained what supporting sentences were added to each paragraph. As noted in Chapter 4, every paragraph has a logical structure;

topic sentences, supporting sentences and concluding sentences are all related. I created my first draft by reading a paragraph's existing material in the outline, and then by immediately adding supporting sentences to flesh out the paragraph. I did not expect the supporting sentences to be perfect; after all, I didn't even expect the existing topic and concluding sentences to be perfect either! I planned to take several editing passes to revise and polish the draft (Sections 4.3 and 4.4).

Nevertheless, adding supporting sentences struck me as more serious writing. At a local level – within each paragraph – I tried to create sentences which worked together to communicate a topic. I did not worry about the whole manuscript's narrative structure when adding supporting sentences. I didn't worry about narrative structure because I was confident I paid sufficient attention to its design during my earlier outlining stages.

When all supporting sentences were added, I had created my first draft. Writers expect first drafts to be so terrible they should never be shared (Becker, 2020; Koch, 2003; Lamott, 1995). However, to illustrate my process for creating the Section AI.2 essay, I provide my (terrible) first draft below:

We live in an artificial intelligence (AI) revolution fueled by a new invention called a *large language model* (LLM). LLMs are built from deep belief networks, which are

artificial neural networks capable of learning to perform complicated tasks because they contain many layers of intermediate processors called hidden units (LeCun et al., 2015). LLMs differ from traditional deep belief networks by including additional architectural properties which aid their ability to learn and to process language (Dong et al., 2023). LLMs trained on a huge amount of text taken from the internet, learn to predict which words should follow from a stimulus sentence. “Give them a human language description or several examples of what one wants them to do, and they can perform tasks for which they were never trained” (Manning, 2022, p. 132).

LLMs are revolutionary because they can generate long, detailed, meaningful responses to short text prompts. LLMs are now commonly used to accomplish a variety of complex tasks, including editing scientific manuscripts, writing or checking programming code, and brainstorming ideas (Mitchell & Krakauer, 2023; Stokel-Walker & Van Noorden, 2023). OpenAI reports its most recent LLM, GPT-4, can pass a number of professional and academic benchmarks. For example, GPT-4’s score on a simulated bar exam placed it in the top 10% of test takers. “What is clear is that these models *use language in a way that is remarkably human*” (Piantadosi, 2023, p. 4, his italics).

LLMs’ performance has generated many questions in both the popular press and scholarly journals (Mitchell & Krakauer, 2023). A recent headline in the New York Times read “Microsoft Says New A.I. Shows Signs of Human Rea-

soning.” Do LLMs understand language? Are LLMs intelligent? Are LLMs sentient or conscious? Such questions are very polarizing; Mitchell and Krakauer report 51% of scholars believe LLMs understand language.

Speaking as a cognitive scientist, I feel such questions miss the key point. I am interested in a different question: ‘Can LLMs inform cognitive science?’ Below, I argue LLMs may indeed be able to inform cognitive science – but only if researchers expend considerable effort to study the internal structure of LLMs in order to discover *how* LLMs produce their amazing behavior. LLMs may provide new theories to cognitive science, but only if researchers look inside them to pull theories out.

Modern AI’s excitement and controversy comes from an LLM’s ability to generate paragraphs of meaningful sentences in response to short prompts or questions. For example, my third-year students in my ‘History of Modern Psychology’ class recently wrote a two-page essay in response to a broad final exam. I explored how OpenAI’s ChatGPT would respond if I only used an exam question as a prompt. I tested ChatGPT with six different possible questions. For each question, ChatGPT generated six paragraphs of well-written prose whose sentences were definitely related to a question’s theme. LLMs consistently generate well-written, interpretable and surprising responses to short, vague prompts.

Cognitive science has studied human language for decades. Cognitive science’s most influential account, genera-

tive grammar (Chomsky, 1965, 1966, 1995), proposes human language involves specialized rules or processes manipulating complex mental representations of sentences. In general, generative grammar represents sentences as a tree-like structure called a phrase marker which encodes the order of words in a sentence, the parts of speech to which words belong, and the hierarchical structure which organizes the sentence. Rules, called transformations, convert one phrase marker into a different phrase marker – for instance, to convert a statement into a question. By focusing on symbols and rules (i.e., phrase markers and transformations), Chomsky’s generative grammar not only transformed linguistics but also inspired theories in cognitive science for many decades.

Cognitive scientists who believe human language is the rule-governed manipulation of symbols do not believe LLMs inform cognitive science (Chomsky et al., 2023; Veres, 2022). For example, in a recent New York Times opinion piece, Chomsky et al. point out “We know from the science of linguistics and the philosophy of knowledge that [LLMs] differ profoundly from how humans reason and use language. These differences place significant limitations on what these programs can do, encoding them with ineradicable defects.” Cognitive scientists have, for decades, followed the motto ‘no cognition without computation’. The motto claims we can only explain cognition by appealing to symbols and rules, which cognitive scientists assume are core properties of computation (Dawson, 2013, 2022).

Others believe the success of LLMs suggest alternatives to generative grammar, like statistical language learners, are worthy of cognitive science's interest (Contreras Kallens et al., 2023). UC Berkeley psychologist Steven Piantadosi agrees with Chomsky et al. (2023) that LLMs do not use grammatical rules (Piantadosi, 2023). However, he then argues an LLM's high level performance without using rules *refutes* Chomskyan linguistics. "The success of large language models is a failure for generative theories because it goes against virtually all of the principles these theories have espoused. In fact, *none* of the principles and innate biases that Chomsky and those who work in his tradition have long claimed necessary needed to be built into these models" (Piantadosi, 2023, pp. 14-15, his italics).

My own research examines cognitive science's foundations, focusing on relations between theories based on rules and symbols and theories based on artificial neural networks. I therefore recognize a historical precedent for Piantadosi's position on Chomskyan theory, a precedent relevant to answering the question about whether LLMs can inform cognitive science.

In the mid-1980s, cognitive science found itself in the midst of what is now called its connectionist revolution. The new networks, called multilayer perceptrons, were powerful enough to serve as theories about human cognitive phenomena. The power of the new networks arose from their containing a layer of hidden units; with enough hidden units a

multilayer perceptron could in principle learn any mapping between stimuli and responses (Lippmann, 1989).

The rise of multilayer perceptrons caused a revolution in cognitive science because proponents of artificial neural networks attacked traditional theories which appealed to the rule-governed manipulation of symbols. For example, one network was trained to convert present-tense verbs into their past-tense form (Rumelhart & McClelland, 1986). Rumelhart and McClelland proposed their network indicated the past-tense network performed linguistics without using grammatical rules like those proposed by Chomsky: “We suggest that lawful behavior and judgements may be produced by a mechanism in which there is no explicit representation of the rule” (p. 217).

My interest in the connectionist revolution focused on a curious aspect of the revolutionaries’ argument: they *assumed* networks abandoned symbols and rules, *but never provided evidence* to support their assumption, or to show what their networks used to replace symbols and rules. I call their approach gee whiz connectionism (Dawson, 2009). I tried to distance myself from gee whiz connectionism by training multilayer perceptrons on various tasks, and by conducting detailed analyses of the internal structure of my trained networks.

When I looked inside my trained networks, I discovered structures which resembled theories based on symbols and rules. For example, my students and I trained one network to solve a number of different logic problems. When we looked

inside the network, we discovered formal rules of logic of the sort philosophy students would learn in an introductory logic course (Berkeley et al., 1995). In another study, my students and I trained networks to classify mushrooms as being edible or poisonous. When we looked inside, we found we could translate network states into a traditional symbol/rule system called a production system (Dawson et al., 1997). Such results reveal surprising similarities between network models and symbolic models, blurring the distinctions between the two approaches (Dawson, 1998, 2004, 2013, 2018).

Importantly, my students and I did not usually find network structure which *replicated* existing formal theories. Instead, we usually found new structures which could inform a cognitive science based on symbols and rules. For instance, my recent work on interpreting artificial neural networks trained to make musical judgements finds structures strongly related to traditional music theory (e.g., preference for particular musical intervals) or to the formal set theory of music (e.g., Fourier representations of musical sets) (Dawson, 2009, 2018; Dawson et al., 2020; Perez et al., 2023). However, I often discover the formal properties of networks depart in surprising ways from traditional music theory. For example, music theory usually represents Western music as consisting of twelve different pitch-classes (C, C#, B, and so on). In contrast, my musical networks generate a formal theory which consists of only six different pitch-classes, and which treats pitch-classes which are six semitones apart in traditional theory (such as C and F#) as

being identical. In short, when I looked inside my networks, I found new kinds of formal structures for cognitive science to explore.

My own research makes me suspect LLMs will only inform cognitive science when researchers abandon mere assumptions about what makes LLMs different from rule and symbol models, and instead seek evidence about both the similarities and differences between both types of models.

Why must we look inside LLMs to inform cognitive science? Cognitive scientists have long known psychologically plausible performance can be produced by methods completely unrelated to the processes of human cognition. One famous example was the conversational program ELIZA which carried out convincing conversations with human participants (Weizenbaum, 1966). ELIZA's performance deliberately did not require the program to understand language. "ELIZA shows, if nothing else, how easy it is to create and maintain the illusion of understanding, hence perhaps of judgment deserving credibility. A certain danger exists there" (Weizenbaum, 1966, pp.42-43). (When ELIZA's danger, and Weizenbaum's intent in creating ELIZA, were ignored Weizenbaum abandoned artificial intelligence research altogether (Weizenbaum, 1976)). Examples like ELIZA show why cognitive scientists are more concerned about comparing processes than comparing performance.

I strongly suspect LLMs use methods radically different from those used by humans because they represent stimuli

and responses with encodings unrelated to any proposed by cognitive scientists. For instance, while LLMs process sentences, they do not represent sentences as sentences, or even as a collection of words. First, they break words into smaller components, called tokens. Then they encode a token as a long vector of numbers in a scheme which assigns similar vectors to similar tokens. In one LLM, BERT, each token is represented by a 768-dimensional vector (Manning et al., 2020). To my knowledge, no cognitive scientist has proposed representing text using such high-dimensional codes or by decomposing words into smaller components. If LLM representations are unrelated to human cognition, then LLMs do not refute Chomsky’s approach. Instead, they refute the applicability of Chomsky’s approach to the explanation of LLMs!

LLM proponents recognize LLMs use some – potentially novel — method to produce a remarkable facility with language. “*The theory is definitely in there*” (Piantidosi, 2023, p. 8, his italics). To inform cognitive science, to defend claims like ‘LLMs refute Chomsky’, researchers must do the hard work to discover what methods LLMs use, and to compare the discovered methods to those discovered by research on human cognition.

However, understanding how LLMs convert stimuli into responses is extremely challenging, because LLMs are intimidatingly large and complex systems. For example, OpenAI’s ChatGPT is reported to have approximately 175 billion parameters which can be adjusted by learning and has been

trained on text consisting of approximately 300 billion words. Another LLM, BERT, consists of twelve different layers of intermediate processors. Mitchell and Krakauer (2023, p. 1) note “the inner workings of these networks are largely opaque; even the researchers building them have limited intuitions about systems of such scale.” Piantadosi (2023, p. 8, *his italics*) concurs: “In fact, we don’t deeply understand *how* the representations these models create work.”

Fortunately, researchers recognize the need to extract potentially novel theories or representations from LLMs and are developing new techniques to understand a LLM’s internal structure. For example, consider the work of Stanford linguist and computer scientist Christopher Manning (Manning, 2022; Manning et al., 2020). Manning and his colleagues have developed methods which probe the internal structure of an LLM in an attempt to determine whether the network represents structures found in generative grammar.

One of Manning’s studies (Manning et al., 2020) examines a language-specific component of an LLM, components called attention heads. An attention head determines the relevance of one word in a sentence presented to an LLM to other words in a sentence, or to different words in the output being generated by the LLM. The more related two words are, the greater the amount of attention is assigned to them. Manning et al. discovered attention being assigned to word pairs captured, in part, linguistic properties of word relationships. For instance, the amount of attention assigned linked

objects to appropriate verbs, linked prepositions to appropriate objects, linked noun premodifiers to appropriate nouns, and so on. Importantly such relationships are linguistic – represented in the hierarchical structure of a phrase marker – because two words may be far apart in a sentence but may still be linguistically related.

Manning et al. (2020) also describe a structural probe method which they use to detect phrase marker trees represented in processors in an LLM’s layers of hidden units. The method involves measuring the distance between different vectors (representing tokens) in the network. With an appropriate distance metric, items whose vectors are close together in the LLM’s space are also close together in the phrase marker structure representing words in a complete sentence. Manning et al. report they can use their distance metric to reconstruct a phrase marker from network properties. In short, “these models learn and represent the syntactic structure of a sentence” (Manning, 2022, p. 131).

My hope is more work of this sort is on the horizon. As researchers explore LLM representations, as well as how the representations are used to generate responses, we move closer to relating LLM to human cognition.

Piantadosi (2023, p. 30) claims “large language models rewrite the philosophy of approaches to language. Do LLMs refute Chomsky’s approach? Do LLMs represent a new connectionist revolution for cognitive science? I believe we can’t answer such questions – yet. Answering such questions

requires researchers to discover the nature of an LLM's representations, as well as how its representations are used to generate responses.

AI.8 REVISE AND POLISH DRAFTS OF THE ESSAY

Drafting is easy; revising is where the real work begins.

The paragraphs which end Section AI.7 provide my essay's first draft. A first draft is merely a manuscript's beginning. Completing the manuscript requires multiple editing passes as I described in Chapter 4.

I worked through my essay several times in order to convert it from a first draft into a completed essay (Section AI.2). My initial editing passes had three basic goals: to correct any flags generated by my word processor; to clean up any writing I didn't like or understand; and to make the manuscript as concise as possible. I typically only conduct one revising pass through a short manuscript per day; stepping away from the manuscript permits fresh eyes to detect new problems when revising continues.

Shortening the manuscript involved two general processes. First, I looked for phrases which I knew I could either shorten or remove. Second, I killed my darlings – I removed material I enjoyed adding earlier in the writing process, but which I recognized didn't really work (Quiller-Couch, 1916).

I believe I successfully converted my first draft into a

clearer, more concise manuscript. The first draft given in Section AI.7 contains 2331 words. In contrast, the final essay provided in Section AI.2 contains only 1676 words: a reduction of nearly 30%.

I used the Editor function in Microsoft word to further compare the essay's first and final versions in order to assess my editing. The Editor function delivers several basic counts of objects (e.g., words, sentences), and averages of such counts (e.g., words per sentence).

The Editor function also delivers a few readability measures. Flesch Reading Ease measures readability using the average number of syllables per word and the average number of words per sentence. It uses a 100-point scale; as the score increases, more people can readily understand the manuscript. Flesch-Kincaid Grade Level uses syllables per word and words per sentence to create a score indicating a grade level. For instance, a score of 4 means a fourth grader can understand the manuscript. The Editor also determines the percentage of passive sentences in the document. Its algorithm defines a passive sentence as one in which the subject does not perform the action of the sentence's verb; instead, the action of the verb is performed on the sentence's subject.

Table AI-5 provides the statistics the Editor computed for the essay's first and final versions. My editing did little to make the essay easier to understand for general readers; both versions have similar reading ease scores and grade level scores. However, the editing did make the final version shorter than

the first draft, and probably punchier: the percentage of passive sentences in the final version is a third of the percentage found in first draft. I removed about five words from the average sentence. Given the essay's topic is technical, I'm satisfied with my editing efforts.

Table AI-5. Readability statistics for the first and final drafts of the essay.

Statistic	Version	
	First Draft	Final Draft
Number of Words	2331	1676
Number of Characters	32515	24609
Number of Paragraphs	24	24
Number of Sentences	110	101
Average Number of Sentences Per Paragraph	4.5	4.2
Average Number of Words Per Sentence	21.1	16.5
Average Number of Characters Per Word	5.5	5.6
Readability: Flesch Reading Ease	30.7	31.2
Readability: Flesch-Kincaid Grade Level	13.8	12.7
Readability: Percent Passive Sentences	9.0	2.9
