

Attention and evidence

David Thorstad

Forthcoming in *PPR*

Please cite published version

Abstract

In traditional models, agents have ample reason to gather cost-free evidence. Rational agents expect evidence gathering to weakly improve the quality of their decisions and the accuracy of their credences. These results are linked to more general principles of reflection and the value of knowledge taken to characterize genuine learning experiences. For attentionally-limited agents, I show that matters are more complex. Rational attentionally-limited agents still cannot expect to make worse decisions or hold less accurate credences as a result of gathering some rather than no evidence. But they can expect to make worse decisions and hold less accurate credences as a result of gathering more rather than less evidence about the world. This result highlights a neglected distinction between two versions of the requirement to gather cost-free evidence and a gap between evidence gathering and learning experiences.

1 Introduction

A.J. Ayer (1957) once asked why agents should gather evidence. I.J. Good (1966) responded by proving that under many conditions, rational agents cannot expect evidence gathering to decrease the quality of their decisions.¹ Patrick Maher (1990a; 1990b) proved that under many conditions, rational agents also cannot expect evidence gathering to decrease the accuracy of their beliefs. And in a more general setting, these results follow from principles of reflection (van Fraassen 1984) and the value of knowledge (Skyrms 1990) which are held to characterize genuine learning experiences (Huttegger 2014; Huttegger and Nielsen 2020; Skyrms 1990).

Here it is important to distinguish two questions. First, can some evidence be worse than no evidence? Second, can more evidence ever be worse than less evidence?² In the classical setting, we will see that both questions receive the same answer. Rational agents cannot expect to make worse decisions after gathering some rather than no evidence, or full rather than partial evidence about the world (Section 2). Likewise, rational agents cannot expect to hold less accurate beliefs after gathering some rather than no evidence, or full rather than partial evidence (Section 3). And in a more general setting, these results follow from principles of reflection (Section 4.1) and the value of knowledge (Section 4.2) taken to characterize genuine learning experiences.

¹See also Blackwell (1953).

²This paper focuses on a specific version of the second question: can complete evidence ever be worse than partial evidence? Blackwell (1953) proves the more general result that more informative experiments can never be expected to reduce decision quality relative to less informative experiments. I focus on the more specific version of the second question to avoid the complications introduced by a more general discussion of Blackwell informativeness.

However, matters are otherwise for agents with limited attention (Hertwig and Engel 2021; Kluger and DeNisi 1996; Lacko and Pappalardo 2004).³ For such agents, classic arguments may be halfway correct: some evidence cannot be worse than no evidence. Agents who wish to avoid receiving evidence of any kind may be able to simply ignore it. However, classic arguments are likely to be halfway incorrect: more evidence can be worse than less evidence. As any prosecutor will tell you before gleefully disclosing millions of pages of files to an overwhelmed defense team, attentionally-limited agents often learn better when they are presented with focused streams of highly relevant information rather than larger streams of less-relevant information (Lester et al. 2012).

The rest of this paper aims to make good on the previous paragraph. After introducing the nature and importance of attentional limits (Section 5), we will see that classic models of rational inattention (Section 6) are not up to the task (Section 7) because they do not allow us to distinguish the processes of gathering and attending to evidence. After developing a model which is expressive enough to distinguish questions about gathering and attending to evidence (Section 8), we will recover the desired results.

In particular, we will see that while attentionally-limited agents cannot expect to make worse decisions on the basis of some rather than no evidence, they can expect to make worse decisions on the basis of full rather than partial evidence (Section 9). In the same way, while attentionally-limited agents cannot expect to lose accuracy after gathering some rather than no evidence, they can expect to lose accuracy after gathering full rather than partial evidence (Section 10).

The reason for this will not be that principles of reflection and the value of knowledge fail, but rather that they no longer play their standard role in justifying claims about evidence gathering (Section 11). Reflection and the value of knowledge may continue to be correct characterizations of learning experiences, but for attentionally-limited agents, gathering evidence is not a complete learning experience. Attentionally-limited agents often have good reason to expect that they neither will nor should learn the complete contents of gathered evidence, so for such agents, questions about evidence-gathering and learning come apart.

2 Good’s Theorem: Evidence and decision quality

Suppose you are deciding between finitely many acts $\mathcal{A}$, such as competing courses of medical treatment.⁴ Suppose that a finite set Ω of world states are relevant to your decision. You have probabilistic credences c over all possible events, regarded as subsets of Ω . You assign utilities u to world states and seek to maximize the expected utility of the resulting world state.

In sum, you face the decision problem $\Gamma = (\Omega, \mathcal{A}, c, u)$. If forced to act now, you would take the treatment maximizing expected utility on your current credences, receiving an

³Other proposed reasons to avoid cost-free evidence include risk aversion (Buchak 2010; Campbell-Moore and Salow 2020), imprecise (Bradley and Steele 2016) or non-countably-additive priors (Kadane et al. 2008), imperfect control over future selves (Maher 1990a), act-dependent world states (Adams and Rosenkrantz 1980), rational modesty (Neth forthcoming), and externalist conceptions of evidence (Das 2023).

⁴I’ll regard acts as Savage acts, that is as functions $a : \Omega \rightarrow \mathcal{O}$ where $\mathcal{O}$ is an outcome space. Throughout this paper, I use examples with $\mathcal{O} = \Omega$.

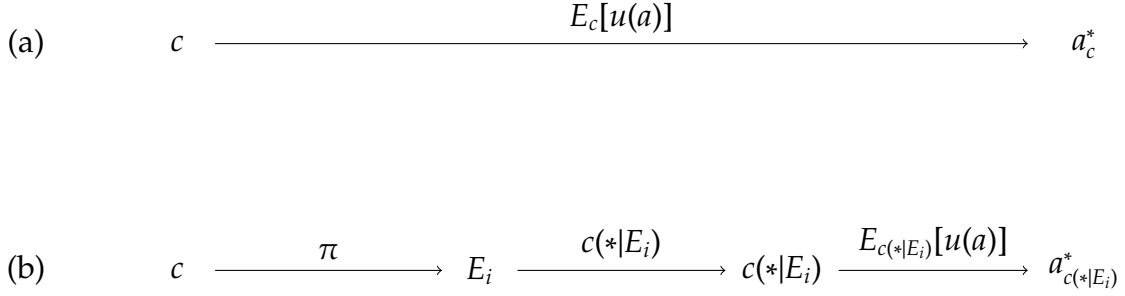


Figure 1: (a) Agents who forgo evidence gathering take an act maximizing expected utility on their current credences. (b) Agents who gather evidence conditionalize on the results of experimentation and take an act maximizing expected utility on their updated credences.

expected benefit equal to the expected utility of that treatment (Figure 1a). This values your current decision problem at $V(\Gamma) = \max_{a \in \mathcal{A}} E_c[u(a)]$.

However, suppose you are offered the chance to gather evidence before acting. For example, you might be given the opportunity to read an informative medical pamphlet π . Reading the pamphlet will provide an uncertain item of evidence E from a partition $\{E_i\}$ of Ω . For example, the pamphlet may tell you whether or not you have cancer, returning as evidence either the set of worlds in which you do have cancer or the set of worlds in which you do not have cancer. Because you have not yet read the pamphlet, you are not certain which evidence you will receive, but assign credence $c(E)$ to each item E of evidence you might receive.

Suppose you read the pamphlet and gain evidence E . You update your credences by conditionalization and face decision problem $\Gamma_\pi = (\Omega, \mathcal{A}, c(\cdot|E), u)$. You then take the act maximizing expected utility on your updated credences (Figure 1b). This values your updated decision problem at $V(\Gamma_\pi) = \max_{a \in \mathcal{A}} E_{c(*|E)}[u(a)]$.

If you were certain that you would receive some particular item of evidence E , the improvement in decision quality from gathering evidence E would be the value difference $V(\Gamma_\pi) - V(\Gamma)$ between your old and new decision problems.⁵ This quantity can be negative, since evidence can be misleading.

However, from an *ex ante* perspective you are uncertain what evidence you will receive. The *ex ante* improvement in decision quality from gathering evidence through process π is the expected value difference between your old and new decision problems, given uncertainty about the evidence that will result, $E_c[V(\Gamma_\pi)] - V(\Gamma)$.⁶ I.J. Good (1966) and David Blackwell (1953) independently proved that under a wide range of assumptions, gathering evidence cannot be expected to decrease the quality of the resulting decision.

$$\textbf{(Good's Theorem)} \quad E[V(\Gamma_\pi)] \geq V(\Gamma). \quad (1)$$

When evidence is costly, the expected improvement in decision quality need not outweigh the expected costs of gathering evidence. But when evidence is cost-free, the expected value of evidence just is the expected improvement in decision quality $E[V(\Gamma_\pi)] - V(\Gamma)$.

⁵In this expression, $V(\Gamma_\pi) = V(\Gamma_{c(*|E)})$ where E is the evidence you will actually receive from π .

⁶In this expression, $E_c[V(\Gamma_\pi)] = \sum_i c(E_i) V(\Gamma_{c(*|E_i)})$.

By Good's Theorem, this quantity must be non-negative. On this basis, it is argued that agents cannot rationally strictly prefer to avoid gathering cost-free evidence, since that evidence cannot yield an expected decrease in decision quality and has no other costs.

Before continuing, it will be important to note a generalization of Good's Theorem. Let π_F be the act of gathering complete evidence about the world, in the sense that π_F always fully discloses the world state.⁷ Let π be any act of evidence gathering. Then the expected improvement in decision quality from gathering complete evidence through π_F is always at least as great as the expected improvement in decision quality from gathering evidence through the competing policy π .⁸

$$\textbf{(Good's Theorem Extended)} \quad E[V(\Gamma_{\pi_F})] \geq E[V(\Gamma_{\pi})]. \quad (2)$$

Again, when evidence is costly, the expected improvement in decision quality from gathering full information need not outweigh the expected costs of gathering it. But when evidence is cost-free, the expected value of gathering full information rather than gathering evidence through the competing policy π just is the expected improvement in decision quality $E[V(\Gamma_{\pi_F})] - E[V(\Gamma_{\pi})]$. By Good's Theorem Extended, this quantity must be non-negative. On this basis, the same argument suggests that agents cannot rationally strictly prefer any evidence gathering policy to the policy of gathering full information about the world.

Moreover, while agents can admit the possibility of making a worse decision after gathering partial evidence through some policy π than the decision they would have made before, they cannot admit the possibility of making a worse decision after gathering full evidence. Agents are certain that after gathering full evidence they will take the action maximizing utility in their actual world state. This means that even agents employing many popular alternative decision theories to expected value maximization will still be unable to strictly prefer partial over full cost-free information about the world. In this sense, Good's Theorem Extended is a more robust result than Good's Theorem.

3 Maher's Theorem: Evidence and credal accuracy

A natural objection to Good's Theorem is that it provides a practical rather than epistemic reason for agents to gather evidence. Good's Theorem says that agents cannot expect to make worse decisions on the basis of cost-free evidence. However, Good's Theorem does not tell agents concerned with epistemic goals such as accuracy how these goals will be affected by gathering cost-free evidence.

Let $I(c, p, w)$ measure the inaccuracy of credence function c on proposition p at world w . In this paper, I use a popular Brier score on which inaccuracy is squared distance from the truth, so that $I(c, p, w) = (c(p) - T(p, w))^2$, where $T(p, w)$ returns the truth value of p at w .⁹ However, the results in this paper hold for the broader class of strictly proper (or

⁷That is, for all $w \in \Omega$, $\pi_F(w) = \{w\}$.

⁸To see this, apply Jensen's inequality or note that $E[V(\Gamma_{\pi})] = \sum_i c(E_i) \max_{a \in \mathcal{A}} [\sum_w c(w|E_i) u(a, w)] \leq \sum_{i,w} c(E_i) c(w|E_i) \max_{a \in \mathcal{A}} [u(a, w)] = \sum_w c(w) \max_{a \in \mathcal{A}} [u(a, w)] = E[V(\Gamma_{\pi_F})]$.

⁹That is, $T(p, w) = 1$ if $w \models p$ and 0 if $w \not\models p$.

self-recommending) scoring rules.¹⁰

The inaccuracy of a credence function at world w will be the sum of the inaccuracy of its opinions about the elementary propositions $\{\{w_i\} : w_i \in \Omega\}$. That is, $\mathcal{I}(c, w) = \sum_{w'} \mathcal{I}(c, \{w'\}, w)$. And from the vantage point of any prior credence function c , we can calculate the expected inaccuracy of any credence function c' , including c , as $\mathcal{I}_c(c') =_{df} E_c[\mathcal{I}(c', *)] = \sum_w c(w) \mathcal{I}(c', w)$.

We know that agents cannot expect to make worse decisions as a result of gathering evidence. But can they expect to hold less accurate credences? Suppose that an agent with prior credences c gathers evidence through the experiment π . The expected inaccuracy of their posterior credences depends on the items of evidence that might result, as

$$\begin{aligned} \mathcal{I}_c(c_\pi) &= \sum_{w,i} c(E_i) c(w|E_i) \mathcal{I}(c(*|E_i), w) \\ &= \sum_i c(E_i) \mathcal{I}_{c(*|E_i)}(c(*|E_i)). \end{aligned} \quad (3)$$

By contrast, the expected inaccuracy of their prior credences can be expanded as

$$\begin{aligned} \mathcal{I}_c(c) &= \sum_{w,i} c(E_i) c(w|E_i) \mathcal{I}(c, w) \\ &= \sum_i c(E_i) \mathcal{I}_{c(*|E_i)}(c). \end{aligned} \quad (4)$$

A feature of the Brier score, and indeed of any strictly proper scoring rule, is that credences are self-recommending so that for any i we have $\mathcal{I}_{c(*|E_i)}(c(*|E_i)) \leq \mathcal{I}_{c(*|E_i)}(c)$. It follows immediately that:

$$\textbf{(Maher's Theorem)} \quad \mathcal{I}_c(c_\pi) \leq \mathcal{I}_c(c). \quad (5)$$

And indeed, Maher's Theorem can be shown in more general settings (Maher 1990a,b).¹¹

Maher's Theorem shows that from their own prior perspective c , agents can never expect to have more inaccurate credences as a result of gathering cost-free evidence. If they are unlucky, they may in fact come to have more inaccurate credences as a result of gathering misleading evidence. But while agents can regard this unfortunate scenario as possible, they cannot expect on balance to lose accuracy as a result of gathering cost-free evidence. As a result, it is argued, agents concerned with accuracy as well as decision quality will also have no grounds on which to rationally strictly prefer to avoid receiving cost-free evidence.

Before continuing, it will be important to note a generalization of Maher's Theorem. As before, let π_F be the act of gathering complete evidence about the world and let π

¹⁰Scoring rule $\mathcal{I}$ is *strictly proper* if for all probabilistic c we have $c = \underset{c'}{\operatorname{argmin}} [\mathcal{I}_c(c')]$ using the definition of $\mathcal{I}_c(c')$ introduced shortly.

¹¹Note that Maher (1990b) rejects the inference from Maher's Theorem to the value of gathering evidence because he thinks the assumption of a strictly proper scoring rule is unmotivated. On extensions beyond strictly proper scoring rules see Horwich (1982).

be any evidence gathering act. Then agents must expect gathering complete evidence to reduce inaccuracy by at least as much as any other evidence gathering act would, in the sense that:¹²

$$\textbf{(Maher's Theorem Extended)} \quad \mathcal{I}_c(c_{\pi_F}) \leq \mathcal{I}_c(c_{\pi}). \quad (6)$$

Maher's Theorem Extended shows that from their own prior perspective, agents can never expect to have more inaccurate credences as a result of gathering full rather than partial evidence about the world. Moreover, whereas agents can admit the possibility that misleading evidence can decrease the accuracy of their credences, agents cannot admit the possibility that gathering full evidence about the world will decrease the accuracy of their credences. The expected inaccuracy $\mathcal{I}_c(c_{\pi_F})$ of gathering complete evidence is always zero. In this sense, agents are certain that gathering full information about the world will make their credences as accurate as possible. This means that agents employing a variety of non-expectational epistemic decision theories also cannot strictly prefer partial to full information about the world. In this sense, Maher's Theorem Extended is a more robust result than Maher's Theorem.

4 Reflection and the value of knowledge

4.1 Reflection

Sections 2-3 discussed agents who update their beliefs in a particular way: by gathering and conditionalizing on evidence. What can be said in a more general setting where agents change their opinions in other ways?

Following the black-box learning model of Brian Skyrms (1990), suppose that an agent currently holds credences c and anticipates updating those credences to c_f at some fixed later time. She knows which credence functions c and c_f are and is certain that she will update from c to c_f at the given time. But she does not necessarily know anything about the process which will bring her to update from c to c_f .

One of the most-discussed principles of Bayesian epistemology is the principle of reflection. Although there are multiple inequivalent formulations of reflection (Huttegger 2013), a natural formulation says that agents should match their current credences to their anticipated future credences. That is, abusing notation slightly to let c_f stand for both a credence function and the proposition that the agent will adopt c_f at the given later time, reflection holds that for any credence functions c, c_f and any proposition p in their mutual domain:

$$\textbf{(Reflection)} \quad c(p|c_f) = c_f(p). \quad (7)$$

Reflection follows almost immediately for agents who update by conditionalization (Kadane et al. 1996; Briggs 2009), and it is widely agreed that in many cases, agents should obey Reflection (Green and Hitchcock 1994; Huttegger 2013).

However, suppose an agent knows or suspects that she will be brought from c to c_f by an epistemically dubious process such as memory loss (Talbot 1991; Arntzenius 2003)

¹²For a proof, note that $\mathcal{I}_c(c_{\pi_F}) = \sum_w c(w) \mathcal{I}_{c(*|w)}(c(*|w)) = 0$ which minimizes Brier inaccuracy.

or drug-induced hallucination (Christensen 1991). In these situations, many have held, it would be inappropriate for the agent to obey reflection. She should not be disposed to compound her anticipated future irrationality by believing irrationally now, even if this will restore a type of consistency between her present and anticipated future selves.

What separates the situations in which Reflection is appropriate, such as gathering and conditionalizing on evidence, from the situations in which Reflection is inappropriate, such as hypnosis or hasty generalization? Many have suggested that the difference is whether the agent has undergone a genuine learning experience (Huttegger 2014; Huttegger and Nielsen 2020; Skyrms 1990). Reflection is appropriate for agents who undergo, or at least who are certain they will undergo, a change of belief based on genuine learning and nothing else. On this line, agents should think that learning will tend to improve their opinions and should therefore be disposed to defer to the beliefs formed as a result of genuine learning experiences. However, agents who know or suspect that they will change their beliefs at least partly through some process beyond a pure learning experience need not obey Reflection. For these agents, the bare fact that they may hold some opinions in the future does not commend those opinions to them now except insofar as they take those opinions to result from learning.

4.2 Value of knowledge

This connection between Reflection and learning is relevant because of connections between Reflection and Good’s Theorem. In the black-box learning setting, our agent thinks that belief change will shift her decision problem from $\Gamma = (\Omega, \mathcal{A}, c, u)$ to $\Gamma_{c_f} = (\Omega, \mathcal{A}, c_f, u)$. In this setting, the previous statement of Good’s Theorem generalizes to what is often called the Value of Knowledge relationship. This says that the expected value of the agent’s future decision upon shifting from prior c to posterior c_f cannot decrease from the expected value of the agent’s present decision. That is:

$$\textbf{(Value of Knowledge)} \quad E[V(\Gamma_{c_f})] \geq V(\Gamma). \quad (8)$$

Value of Knowledge is not a theorem. Indeed, Value of Knowledge often fails for agents who suspect their future selves of irrationality. When does Value of Knowledge hold?

Under quite general conditions including the present setup, we can show that Reflection and Value of Knowledge are equivalent (Huttegger 2014; Skyrms 1990).¹³ Agents obey Reflection just in case they also obey Value of Knowledge. What is going on here?

Many authors have suggested that this result shows that Value of Knowledge, like Reflection, should be regarded as a good formal characterization of learning experiences (Huttegger 2014; Huttegger and Nielsen 2020). Agents who think that they will change their opinions from c to c_f on the basis of a pure learning experience and nothing else should defer to their future credences, as Reflection asks, and should expect their future decisions to be weakly improved, as Value of Knowledge asks. But agents who allow that they may shift from c to c_f on the basis of other influences beyond learning need not be disposed to always defer to their future credences or expect their future decisions to be weakly improved. If this is right, then we should perhaps view Good’s Theorem as

¹³See Dorst (2020); Dorst et al. (2021) and Geanakoplos (1989) on connections between Value of Knowledge and deference under higher-order uncertainty.

a consequence of the more general Value of Knowledge relation on learning experiences, together with the fact that gathering and conditionalizing on evidence is a pure learning experience.

Before concluding, it is worth generalizing Value of Knowledge in the same way as before. Let c_F be the credences that result from learning the true world state, so that $c_F(w) = 1$ if w obtains and 0 otherwise. Then Value of Knowledge generalizes to the statement that:

$$\textbf{(Value of Knowledge Extended)} \quad E[V(\Gamma_{c_F})] \geq E[V(\Gamma_{c_f})]. \quad (9)$$

Value of Knowledge Extended, unlike Value of Knowledge, is a theorem so long as the agent is certain that c_F correctly identifies the true world state.¹⁴ This is true even if we do not assume that the agent knows which credence function c_F is. Agents need not know how they will change their opinions to expect that their decisions will be at least as good when based on complete information about the world as when formed in any other way. Indeed, as before the agent is certain that she will not make a worse decision after shifting to c_F than she would after shifting to c_f .

This rounds out the classic case for gathering cost-free evidence. Agents cannot expect to make worse decisions (Section 2) or hold less accurate credences (Section 3) as a result of gathering and conditionalizing on some rather than no cost-free evidence, or full rather than partial cost-free evidence. And on a natural view, this expectation is grounded in the fact that gathering and conditionalizing on cost-free evidence is a learning experience together with the idea that the generalized Value of Knowledge relation governs learning experiences.

5 Attention

The rest of this paper asks what happens to the classic case for evidence gathering when agents have limited capacities to attend to gathered evidence. The first order of business is to get clear on what attention is (Section 5.1) and why it matters (Section 5.2).

5.1 The nature of attention

Agents gain evidence about the world through perception. Some of the signals that we receive are passively acquired, such as the sounds entering my eardrums as I type these words. Others are actively acquired through processes of evidence gathering. We do not, and cannot learn the entire contents of incoming signals. Instead, we attend selectively to some aspects of incoming signals and ignore others.

On my preferred model (Section 8), attention maps incoming worldly signals into mental signals which are learned by the agent. This means that learning experiences are not fully characterized by incoming worldly signals, even when those signals are actively gathered. Rather, learning experiences are characterized by the application of attention to incoming signals. This paper is concerned with agents who learn through conditionalization. Therefore, my preferred model will treat agents as conditionalizing on the signals provided by attentional processes.

¹⁴That is, for all worlds w : $c(w|c_F(w) = 1) = 1$.

5.2 The importance of attention

It is often held that we are transitioning from an information economy into an attention economy (Castro and Pham 2020; Davenport and Beck 2001; Simon 1971). Whereas previous generations were limited in large part by the availability of information, we find ourselves overwhelmed by information and limited increasingly by our ability to attend to what is important. On this basis, many philosophers (Archer 2021; Clark ms; Siegel 2017; Irving and Glasser 2020), economists (Maćkowiak et al. 2023; Sims 2003) and cognitive scientists (Lieder and Griffiths 2020; van den Berg and Ma 2018) have held that rational cognition involves not only responding correctly to available information, but also making appropriate use of scarce attentional resources.

There is good evidence that agents avoid cost-free disclosures of excessive information, and that they often form more accurate beliefs and make better decisions as a result (Ben-Shahar and Schneider 2016; Golman et al. 2017; Hertwig and Engel 2021). For example, exclusion of evidence during legal fact-finding may increase the accuracy of fact-finding (Lester et al. 2012), and disclosure of complex financial information about loans may lead agents to pick inferior loans (Lacko and Pappalardo 2004). One natural explanation for these phenomena is attentional: agents avoid excessive intake of cost-free evidence as a means of allocating scarce attention towards the most important and decision-relevant evidence. If that is right, then we should expect theories of rational inattention to allow agents who gather less evidence to form more accurate beliefs and make better decisions.

It turns out that making good on this insight is not so easy. It is easy to see how cost-free evidence may fail to help agents without the attention to absorb it, but how can cost-free evidence harm agents? They could, it may seem, simply ignore evidence they do not wish to attend to. And in fact, the standard framework of rational inattention (Sims 2003) may imply exactly this result. Let us see how the standard framework goes wrong (Sections 6-7) and how to do better (Section 8).

6 Rational inattention: The standard framework

On the standard account of rational inattention, evidence gathering is mediated by attentionally constrained cognitive processes. When you read a medical pamphlet, you do not learn its contents E , but rather an attentionally-limited signal s extracted from the pamphlet.

More formally, assume as before that you face a medical decision problem $\Gamma = (\Omega, \mathcal{A}, c, u)$ with finite sets of states Ω and acts $\mathcal{A}$, probabilistic credences c and utilities u over states. Before acting, you may choose to read one of the available medical pamphlets. Because attention is limited, you must choose not only which pamphlet to read but also which elements of the pamphlet to attend to. The combined choice of a pamphlet to read and an allocation of attention to the pamphlet is modeled as the choice of an attention policy π from the set Π of possible attention policies.

Each attention policy $\pi \in \Pi$ is associated with a set $\mathcal{S}_\pi$ of signals that might result. If we model signals $\mathcal{S}_\pi$ as the events that those signals are sent, then $\mathcal{S}_\pi$ forms an event partition of Ω . The probabilities of receiving each signal depend on your attention policy as well as on the current state of the world. This means that each policy π is associated with a map $\pi : \Omega \rightarrow \Delta(\mathcal{S}_\pi)$ from world states to probability distributions over signals

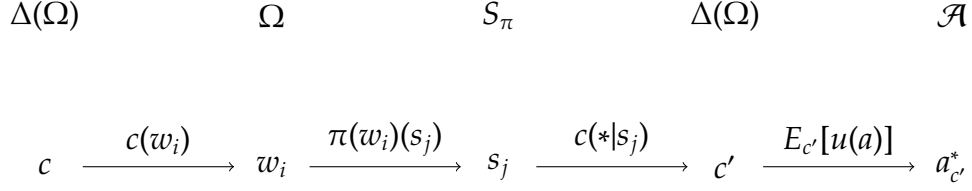


Figure 2: Standard rational inattention model: Agents conditionalize on the world-dependent signal s_j resulting from their attention policy π , then take the act $a_{c'}^*$ maximizing expected utility on their updated credences c' .

that might be sent in that world state. Once you receive some signal s_j you update by conditionalization to $c' = c(*|s_j)$ and take the act $a_{c'}^*$ maximizing expected utility on your updated credences c' (Figure 2).

Suppose for simplicity that there is only one medical pamphlet, so that attention policies π_i represent different strategies for allocating attention to the pamphlet. If attention were unlimited, it would always be optimal to pay full attention to the pamphlet, and it could never be worse in expectation to pay some attention to the pamphlet than to ignore the pamphlet, recovering Good's Theorem. To avoid this result, the standard rational inattention framework assumes that attention is costly. The cost of attention is driven by the expected reduction in uncertainty that results, cashing out the intuition that higher reductions in uncertainty require greater attention.

More specifically, let $H(p) = -\sum_w p(w) \ln(p(w))$ be the Shannon entropy of a probability function p , a measure of its uncertainty (Shannon 1948). On the standard framework, attention policies are more costly if they are more informative, in the sense of offering a greater expected reduction in the uncertainty of the agent's credences. In particular, the cost $C(\pi)$ of policy π is proportional to the expected difference between the entropy of the agent's prior and posterior after applying π :

$$C(\pi) = \kappa(H(c) - E_c[H(c_\pi)]).^{15} \quad (10)$$

Here the non-negative scalar κ controls the cost of attention, with $\kappa = 0$ in the special case of unlimited attention. While the prior entropy is known, the posterior entropy is uncertain because it is uncertain which signal will be sent, and hence which posterior will result.

In the standard model of rational inattention, agents select an attention policy and resulting action to maximize the expected value of the resulting action less the attentional cost.¹⁶ Good's Theorem follows only under the further assumption that attention is unlimited, so that $\kappa = 0$ and there are no attentional costs.

¹⁵Here $E_c[H(c_\pi)] = \sum_{w,i} c(w) \pi(s_i|w) H(c(*|s_i))$.

¹⁶That is, they pick $\arg\max_{\pi} [E[V(\Gamma_\pi)] - C(\pi)]$ where $E[V(\Gamma_\pi)] = \sum_{w \in \Omega, s \in S_\pi} c(w) \pi(s|w) V(\Gamma_{c(*|s)})$.

7 Applying the standard framework

What does the standard rational inattention model say about the results of refusing cost-free evidence? It is hard to tell. The reason for this is that the attention policy π collapses two different processes: gathering evidence and attending to evidence. This means that we cannot, without further assumptions, conclude from the standard model that cost-free evidence can be harmful. This is because we will not be able to distinguish, within the standard model, a scenario in which cost-free evidence is harmful from a scenario in which it would be harmful to irrationally pay full attention to cost-free evidence, but not to gather and ignore it.

In fact, the standard rational inattention model is formally isomorphic to a natural model of costly evidence gathering. Suppose we take π to be an experiment, yielding evidence from the partition $\mathcal{S}_\pi$. Agents choose among various experiments π and select the experiment maximizing the expected value of the resulting action, less the cost of experimentation. Experiments, like attention, are assumed to be more costly if they lead to greater uncertainty reduction. On this reading, the standard rational inattention framework is isomorphic to a generalization of Good’s model in which evidence gathering is costly. We recover Good’s Theorem when evidence gathering costs κ are set to zero. But that was precisely what Good argued. So it is not clear how much progress has been made on the rationality of cost-free evidence refusal, beyond perhaps a formal reinterpretation of costly evidence gathering which combines the cost of attention with the cost of evidence gathering.

To do better, we will need a framework that can separate the processes of gathering and attending to evidence. I take up this task below.¹⁷

8 Gathering and attending to evidence

If we want to separate questions about gathering and attending to evidence, we need to modify the standard model by breaking policies into two components: *information policies* governing evidence gathering and *attention policies* governing attention to gathered evidence.

The set Π_1 of information policies is defined exactly as attention policies were defined in the standard model. That is, each information policy $\pi_1 \in \Pi_1$ is associated with an event partition $\mathcal{S}_{\pi_1} = \{s_i^1\}$ of possible signals, or *information outcomes*, and each π_1 maps states into probability functions over information outcomes.

If we assumed that agents update on the information outcome $s^1 \in \mathcal{S}_{\pi_1}$, we would recover standard evidence gathering models, with Good’s Theorem following again if evidence is assumed to be cost-free. However, attentionally-limited agents must select an attention policy π_2 specifying how attention is to be allocated to information outcomes. Attention policies are defined exactly as before, except that they take information outcomes rather than worlds as inputs.

More formally, attention policies $\pi_2 \in \Pi_2$ are associated with a set $\mathcal{S}_{\pi_2} = \{s_i^2\}$ of *attention outcomes* representing the possible outcomes of attending to gathered evidence. As before, $\mathcal{S}_{\pi_2}$ is an event partition of Ω , where signals are modeled as the events in which they

¹⁷This model builds on Lipnowski et al. (2020).

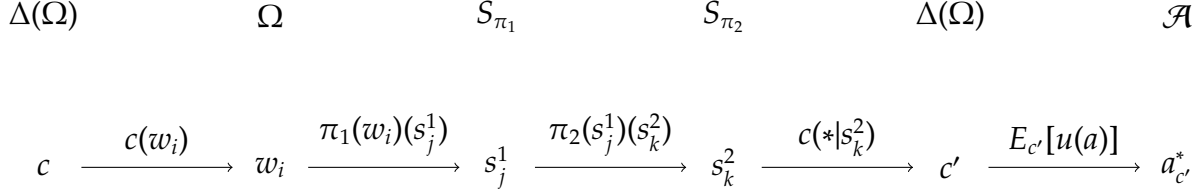


Figure 3: Revised rational inattention model: Agents gather evidence through cost-free experiment π_1 , then filter through attentionally-constrained process π_2 . They conditionalize on the resulting signal s_k^2 and choose an expected-utility-maximizing act given their updated credences.

are received. The probabilities of receiving each signal depend on the chosen attention policy as well as on the information outcome. This means that each attention policy π_2 maps information outcomes into the space $\Delta(\mathcal{S}_{\pi_2})$ of probability functions over resulting attention outcomes. π_2 thus induces a map $\pi_2 : \mathcal{S}_{\pi_1} \rightarrow \Delta(\mathcal{S}_{\pi_2})$, where we have restricted consideration to attention outcomes given the possible information outcomes from the previously chosen information policy.

Your decision problem is then given by Figure 3. You gather evidence according to your chosen information policy π_1 then attend to gathered evidence using your chosen attention policy π_2 , conditionalize on the resulting attention outcome and take an act maximizing expected utility given your updated credences. You must choose an information policy π_1 and a compatible attention policy π_2 to maximize the expected utility of your resulting choice.¹⁸

Because this is not a model of costly evidence gathering, information policies π_1 are assumed to be costless. Instead, the model captures attentional limitations by penalizing attention policies with costs proportional to expected uncertainty reduction. That is:

$$C(\pi_2) = \kappa(H(c) - E_c[H(c|s^2)]). \quad (11)$$

As before, setting κ to zero removes attentional limitations, recovering Good's Theorem.¹⁹

In this framework, there is a clear distinction between cost-free evidence gathering and costly attending. Information policies π_1 for evidence gathering are costless. Agents incur costs only for the information attended to through attention policies π_2 . Gathered

¹⁸That is, you pick π_1, π_2 to maximize $[\sum_{w,j,k} c(w)\pi_1(s_j^1|w)\pi_2(s_k^2|s_j^1)V(\Gamma_{c(*|s_k^2)})] - C(\pi_2)$ with the cost function $C(\pi_2)$ defined in the next paragraph.

¹⁹A referee asks whether this measure of attentional costs wrongly incorporates both the strength of informational signals and the costs of attentional processing, since final uncertainty reduction is a function of both quantities. Certainly there is a good case to be made for verifying the robustness of this result across alternative cost functions. I measure cost by uncertainty reduction for three reasons. First, this is the orthodox approach in rational inattention models, and I want to make minimal deviations from orthodoxy in order to clarify what drives my results. Second, because the model incorporates costless conditionalization on attentional signals s_2 , we might treat the standard cost function as a measure of the attentional cost of focusing on the evidence provided by s_2 , cashed out in information-theoretic terms through Shannon entropy reduction. Finally, varying the cost function makes solution methods for rational inattention problems considerably more complex, when they are known at all. However, I am not committed to monism about acceptable cost functions, and it would be a very important project to examine the results in this paper under alternative cost functions.

evidence can always be ignored at no cost, but it cannot be attended to costlessly. We can now ask what becomes of Good’s Theorem, Maher’s Theorem, Reflection and Value of Knowledge in our extended framework.

It will turn out that the original versions of Good’s Theorem and Maher’s Theorem continue to hold, but that Good’s Theorem Extended and Maher’s Theorem Extended, which were previously more robust than their original counterparts, now fail (Sections 9-10). Attentionally limited agents still cannot expect to make worse decisions or hold less accurate credences as a result of gathering some rather than no evidence. However, they can expect to make worse decisions and hold less accurate credences as a result of gathering full rather than partial evidence. This result reveals a disconnection between norms on learning and evidence gathering for attentionally-limited agents (Section 11).

9 Good’s Theorem revisited

9.1 Generalities

What happens to Good’s Theorem in this setting? The original theorem continues to hold. Let π_1 be an information policy and let π_2 be any compatible attention policy, optimal or not. Then $\pi_2 \cdot \pi_1$ is an experiment in the sense required by Good’s Theorem. Therefore, letting $\Gamma_{\pi_2 \cdot \pi_1} = (\Omega, \mathcal{A}, c_{\pi_2 \cdot \pi_1}, u)$ be the decision problem that results from applying π_1 and π_2 then conditioning on the results, we have:

$$\textbf{(Good’s Theorem for Inattentive Agents)} \quad E[V(\Gamma_{\pi_2 \cdot \pi_1})] \geq V(\Gamma). \quad (12)$$

Attentionally-constrained learning is still learning, so agents cannot expect even the most attentionally-constrained learning to decrease the expected quality of their decisions. They may prefer not to apply the attention policy π_2 on the grounds that it is costly, but they cannot prefer to avoid it on the grounds that π_2 decreases expected decision quality.

However, Good’s Theorem Extended does not generalize to rationally inattentive agents. Let π_1^F be an information policy fully disclosing the world state.²⁰ Let π_1^P be an information policy partially disclosing the world state.²¹ Let π_2^{F*} and π_2^{P*} be optimal attention policies in response to π_1^F, π_1^P , respectively. Then we will see below that we can have:

$$E[V(\Gamma_{\pi_2^{F*} \cdot \pi_1^F})] < E[V(\Gamma_{\pi_2^{P*} \cdot \pi_1^P})]. \quad (13)$$

In this sense, rationally inattentive agents can expect to make better decisions on the basis of gathering and rationally responding to partial rather than full information.²²

²⁰That is, let π_1^F be any policy such that for each $s \in \mathcal{S}_{\pi_1^F}$, $c(*|s)$ assigns probability 1 to a single world.

²¹That is, let π_1^P be any policy such that for some $s \in \mathcal{S}_{\pi_1^P}$, $c(*|s)$ assigns nontrivial probability to at least two worlds.

²²A referee asks whether the position in this paper is unstable, because after gathering and optimally attending to partial information it could be rationally permissible to gather and optimally attend to the remaining information in one or more subsequent rounds. This instability is blocked by the following result. Suppose that π_1^* is an optimal information policy, in the sense that no competing information policy has higher expected value. Then after applying and optimally attending to π_1^* , it is always an optimal response to ignore all further gathered information. Thus, attentionally-limited agents would retain their

Of course, if we let π_2^F be the policy of fully attending to the signals received after full disclosure, matters are different. For any information policy π_1 and attention policy π_2 we will have:²³

$$E[V(\Gamma_{\pi_2^F, \pi_1^F})] \geq E[V(\Gamma_{\pi_2, \pi_1})]. \quad (14)$$

However, paying full attention is often suboptimal for rationally inattentive agents. As a result, while agents cannot expect to make better decisions than those they would after receiving and fully attending to complete evidence about the world, they can be certain that they neither will nor should fully attend to complete evidence if it is provided. Hence (14) provides few guarantees about rational or actual processes of evidence gathering.

9.2 Example: Full disclosure

To see how agents can expect to make better decisions after partial rather than full disclosure, consider an agent facing three world states w_1, w_2, w_3 and three acts f, g, h .²⁴ Each act performs optimally in one state and its performance decays quadratically as we move away from that state, so that the utilities of each act are given by Table 1. The agent currently applies a uniform prior $c = (1/3, 1/3, 1/3)$ over world states w_1, w_2, w_3 . Throughout this example and the rest of the paper, attention costs κ will be normalized to 1. Proofs are in the appendix.²⁵

	f	g	h
w_1	0	-1	-4
w_2	-1	0	-1
w_3	-4	-1	0

Table 1: Utilities of acts in world states.

Suppose first that we fully disclose the world state, letting π_1^F map w_1, w_2, w_3 with certainty into corresponding signals $s_{w_1}^1, s_{w_2}^1, s_{w_3}^1$, respectively. What shape does an optimal attention policy take? It turns out that optimal attention π_2^{F*} to full disclosure involves sending three signals s_f^2, s_g^2, s_h^2 with conditional probabilities given by Table 2.

updated credences when presented with information that they had rationally chosen not to gather or attend to.

²³This follows immediately from Good's Theorem Extended, putting $\pi_2^F \cdot \pi_1^F = \pi_F$ in our earlier notation and $\pi_2 \cdot \pi_1 = \pi$.

²⁴This example is adapted from Lipnowski and colleagues (2020).

²⁵A referee asks for an intuitive explanation of the suboptimality of full attention in a simple three-state example. The first thing to say here is that a three-state example is chosen for simplicity. Readers are welcome to imagine much larger examples to make the result more intuitive. Within the present example, one general way to understand the result would be as a vindication of the decision-theoretic wisdom that optimal behavior often varies with the choice of probability and utility functions, so that the present result is driven by the particular probabilities and utilities in this case. Another way to understand the result would be that fully attending to the world state can be very difficult, not only because of the number of world states, but also because of the complexity of evidence. A final way to understand the result would be to note that the agent in this case has much more to lose in confusing worlds w_1 and w_3 for one another than

	s_f^2	s_g^2	s_h^2
$s_{w_1}^1$	0.238	0.758	0.004
$s_{w_2}^1$	0.039	0.922	0.039
$s_{w_3}^1$	0.004	0.758	0.238

Table 2: Conditional probabilities $c(s_a^2 | s_{w_i}^1)$ of attention outcomes s_a^2 given information outcomes $s_{w_i}^1$.

As their name suggests, conditionalizing on the signals s_f^2, s_g^2, s_h^2 leads the agent to form posteriors c_f, c_g, c_h respectively, then to take the acts f, g , or h respectively. The posteriors c_a corresponding to each act as well as the prior probability of each posterior arising are given by Table 3.

Attention outcome	Associated act	Associated posterior	Prior	$E_c[V(\Gamma_{c_a})]$
s_f^2	f	$c_f = (0.845, 0.140, 0.015)$	0.094	-0.2
s_g^2	g	$c_g = (0.311, 0.378, 0.311)$	0.812	-0.622
s_h^2	h	$c_h = (0.015, 0.140, 0.845)$	0.094	-0.2

Table 3: Acts and posteriors corresponding to attention outcomes after rational attention to full disclosure, prior probabilities of each posterior, and prior expected utilities of resulting decision problems.

The expected value $E[V(\Gamma_{\pi_2^{F*}, \pi_1^F})]$ of the agent's resulting decision problem is found by multiplying the expected values of each posterior decision problem she might face together with the prior likelihood of facing each problem. That is, $E[V(\Gamma_{\pi_2^{F*}, \pi_1^F})] = (0.094)(-0.2) + (0.812)(-0.622) + (0.094)(-0.2) = -0.543$. This is not bad, and certainly an improvement over the expected value of the agent's prior decision problem, which sits at -0.666. But we can do better.

9.3 Example: Partial disclosure

Suppose that our partial disclosure policy π_1^P maps world states w_1 and w_3 to different signals but provides no information about w_2 , so that π_1^P is as in Figure 4.

Now that the agent has been deprived of the opportunity to learn about w_2 , optimal attention π_2^{P*} to partial disclosure shows greater sensitivity to the remaining states. Specifically, I show in the appendix that optimal attention involves sending two signals s_f^2, s_h^2 with conditional probabilities given by Table 4.

Each signal s_a^2 induces a corresponding posterior c_a and leads the agent to take the corresponding act a (Table 5). Note that because the induced posteriors have been deprived of information about w_2 , the agent has taken pains to become sharply opinionated about w_1 and w_3 , becoming over fourteen times as confident that one of these states occurs rather

she does in confusing either for w_2 . For this reason, optimal attention is directed primarily at discriminating w_1 from w_3 and not at discriminating w_2 from other worlds.

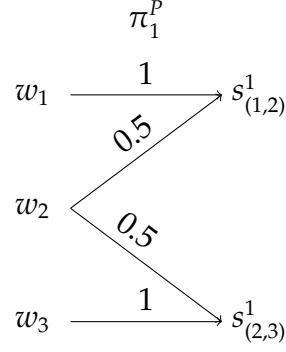


Figure 4: Partial disclosure π_1^P deterministically maps extreme states w_1, w_3 to associated signals but scrambles intermediate state w_2 between them.

	s_f^2	s_h^2
$s_{(1,2)}^1$	0.935	0.065
$s_{(2,3)}^1$	0.065	0.935

Table 4: Conditional probabilities $c(s_a^2 | s_b^1)$ of attention outcomes s_a^2 given information outcomes s_b^1 .

than the other. As a result, she is now certain to take one of the risky actions f or h and to reap expected utility -0.509 as a result of having learned enough about the extreme states w_1, w_3 to reward her risk. This is an improvement on the expected utility -0.543 of her decision after optimal attention to full disclosure.

Attention outcome	Associated act	Associated posterior	Prior	$E_c[V(\Gamma_{c_a})]$
s_f^2	f	$c_f = (0.623, 0.333, 0.044)$	0.5	-0.509
s_h^2	h	$c_h = (0.044, 0.333, 0.623)$	0.5	-0.509

Table 5: Acts and posteriors corresponding to attention outcomes after rational attention to partial disclosure, prior probabilities of each posterior, and prior expected utilities of resulting decision problems.

This completes our proof of (13), showing that Good's Theorem Extended no longer holds for rationally inattentive agents.

10 Maher's Theorem revisited

What happens to Maher's Theorem in this setting? As before, let π_1 be an information policy and let π_2 be any compatible attention policy, optimal or not. Then $\pi_2 \cdot \pi_1$ is an experiment in the sense required by Maher's Theorem. Therefore, letting $c_{\pi_2 \cdot \pi_1}$ be the

credence function that results from applying π_1 and π_2 then conditioning on the results, we have:

$$\textbf{(Maher's Theorem for Inattentive Agents)} \quad \mathcal{I}_c(c_{\pi_2, \pi_1}) \leq \mathcal{I}_c(c). \quad (15)$$

Attentionally-constrained learning is still learning, so agents cannot expect even the most attentionally-constrained learning to decrease the expected accuracy of their credences.

However, the extension of Maher's Theorem fails for rationally inattentive agents in the same way that Good's Theorem Extended fails for them. Letting π_1^F, π_1^P be full and partial disclosure policies and π_2^{F*}, π_2^{P*} be optimal attentional responses to each, we can have:²⁶

$$\mathcal{I}_c(c_{\pi_2^{F*}, \pi_1^F}) > \mathcal{I}_c(c_{\pi_2^{P*}, \pi_1^P}). \quad (16)$$

That is to say, agents can expect to hold more accurate credences as a result of rational attention to partial information than they would as a result of rational attention to full information. As before, some information cannot be worse than no information, but it does not follow that full information is better than partial information.

In the present example, we can show (Appendix) that the prior expected inaccuracy $\mathcal{I}_c(c_{\pi_2, \pi_1})$ of the agent's posterior credence is the sum of the posterior expected inaccuracies $\mathcal{I}_{c_a}(c_a)$ weighted by the prior credence $c(c_{\pi_2, \pi_1} = c_a)$ that each posterior will come about. After rational attention to full information, the posterior expected inaccuracies and the prior probabilities of each posterior are given by Table 6. This puts the expected inaccuracy after full disclosure at $(0.266)(0.094) + (0.664)(0.812) + (0.266)(0.094) = 0.590$. This is better than the prior expected inaccuracy $\mathcal{I}_c(c)$ of 0.666, but again we can do better.

	$\mathcal{I}_{c_a}(c_a)$	$c(c_{\pi_2, \pi_1} = c_a)$
c_f	0.266	0.094
c_g	0.664	0.812
c_h	0.266	0.094

Table 6: Posterior expected inaccuracies and prior probabilities of posteriors after rational attention to full disclosure.

After rational attention to partial disclosure through π_1^P , the posterior expected inaccuracies and prior probabilities of each posterior are given by Table 7. This puts the expected inaccuracy after partial disclosure at 0.499, an improvement over full disclosure. Just as our agent expects partial rather than full disclosure to improve the final quality of her decisions, she also expects partial disclosure to improve the accuracy of her credences.

²⁶As a referee notes, there are two ways to interpret attentional costs κ in this section. The first would be as the epistemic cost of attention, determined perhaps by the opportunity cost of foregone accuracy improvements that compete for scarce attentional resources. The second would be as the practical cost of attention, understood as before. Both interpretations yield distinct and important interpretations of the results of this section.

	$\mathcal{I}_{c_a}(c_a)$	$c(c_{\pi_2, \pi_1} = c_a)$
c_f	0.499	0.5
c_h	0.499	0.5

Table 7: Posterior expected inaccuracies and prior probabilities of posteriors after rational attention to partial disclosure.

11 Reflection and the value of knowledge revisited

What happens to Reflection and Value of Knowledge in this framework? Our agent is certain to update by conditionalization, and under most conditions including the present setup, conditionalizers obey Reflection (Kadane et al. 1996; Briggs 2009). Likewise, our agent continues to satisfy both versions of the value of knowledge relation.

Suppose again that our agent holds priors c and knows that at some future time she will shift to posteriors c_f . We might, but need not also suppose she knows that this transition will occur through the combination of an unknown information and attention policy. We still have:

$$\textbf{(Value of Knowledge)} \quad E[V(\Gamma_{c_f})] \geq V(\Gamma). \quad (17)$$

Likewise, suppose an agent is certain that c_F correctly identifies the world state. Then as before, even if we no longer assume the agent is certain which credence function c_F is, we have:

$$\textbf{(Value of Knowledge Extended)} \quad E[V(\Gamma_{c_F})] \geq E[V(\Gamma_{c_f})]. \quad (18)$$

But we need to be careful in interpreting Value of Knowledge Extended.

Previously, we assumed that all gathered evidence is learned. This allowed us to gloss Value of Knowledge Extended as the claim that the agent expects gathering complete evidence about the world to bring about weakly better decisions than those brought about by any competing process of opinion change. When we remove the assumption that all gathered evidence is learned, no such gloss is true. For example, let π_1^F be a full disclosure information policy and π_2^{F*} be an optimal attentional response. It is not true that for any posterior c_f , we have:

$$E[V(\Gamma_{c_{\pi_2^{F*}, \pi_1^F}})] \geq E[V(\Gamma_{c_f})]. \quad (19)$$

The failure of (19) means that agents can expect some opinion changes to produce better decisions than those produced by gathering and rationally attending to complete evidence about the world, even when they know nothing about how the competing opinion change will be produced.

As before, if π_2^F is the policy of fully attending to π_1^F , we will have:

$$E[V(\Gamma_{c_{\pi_2^F, \pi_1^F}})] \geq E[V(\Gamma_{c_f})]. \quad (20)$$

But since agents can rationally expect, or even be certain that their opinions will not shift to $c_{\pi_2^F, \pi_1^F}$, (20) is not naturally glossed as a defense of gathering full evidence. It is most

naturally glossed as a statement about how agents must regard the omniscient credence function.

This result reinforces the lesson that Reflection and Value of Knowledge are characterizations of learning, not evidence gathering. When all evidence gathered is fully learned, as in the case of agents who conditionalize on all gathered evidence, then Reflection and the Value of Knowledge will continue to support reformulated versions of results such as Good's Theorem Extended and Maher's Theorem Extended. But when gathering evidence is driven apart from learning, we will not recover these reformulated results. We still recover the unexpanded results. But these results no longer tell us anything terribly direct about the value of gathering evidence. They tell us instead about the value of learning experiences, of which evidence gathering is only one part.

12 Taking stock

This paper began by presenting the classic case for gathering cost-free evidence. We saw that agents who are certain to conditionalize on all gathered evidence cannot expect to lose decision quality (Section 2) or credal accuracy (Section 3) as a result of gathering some evidence rather than no evidence, or complete evidence rather than partial evidence. We saw that both results are naturally interpreted in light of principles of Reflection (Section 4.1) and the Value of Knowledge (Section 4.2) which are often held to characterize genuine learning experiences.

We asked how these results fare for agents with limited attention (Section 5). Although standard models of rational inattention (Section 6) do not yield clear verdicts about the results of evidence gathering (Section 7), an extended model which separates the processes of gathering and attending to evidence (Section 8) does better. In this framework, we saw that agents still cannot expect to lose decision quality (Section 9) or credal accuracy (Section 10) as a result of gathering and rationally attending to some evidence rather than no evidence, but they can expect to lose decision quality and credal accuracy as a result of gathering and rationally attending to full rather than partial evidence. Since inattentive learning is still learning, these results are not driven by failures of Reflection or the Value of Knowledge (Section 11). Rather, we saw that they are driven by a reading of Reflection and Value of Knowledge as characterizations of learning experiences. When learning involves further processes beyond evidence gathering, Reflection and Value of Knowledge lose some of their original force in underwriting connections between evidence gathering, credal accuracy and decision quality.

References

- Adams, Ernest and Rosenkrantz, Roger. 1980. "Applying the Jeffrey decision model to rational betting and information acquisition." *Theory and Decision* 12:1–20.
- Archer, Sophie (ed.). 2021. *Salience: A philosophical inquiry*. Routledge.
- Arntzenius, Frank. 2003. "Some problems for conditionalization and reflection." *Journal of Philosophy* 100:356–70.

- Ayer, A.J. 1957. "The conception of probability as a logical relation." In S. Körner and M.H.L. Pryce (eds.), *Observation and interpretation: A symposium of philosophers and physicists*, 12–30. Butterworths.
- Ben-Shahar, Omri and Schneider, Carl. 2016. *More than you wanted to know: The failure of mandated disclosure*. Princeton University Press.
- Blackwell, David. 1953. "Equivalent comparisons of experiments." *The Annals of Mathematical Statistics* 24:265–72. doi:10.1214/aoms/1177729032.
- Bradley, Seamus and Steele, Katie. 2016. "Can free evidence be bad? Value of information for the imprecise probabilist." *Philosophy of Science* 83:1–28.
- Briggs, R.A. 2009. "Distorted reflection." *Philosophical Review* 118:59–85.
- Buchak, Lara. 2010. "Instrumental rationality, epistemic rationality, and evidence-gathering." *Philosophical Perspectives* 24:85–120.
- Campbell-Moore, Catrin and Salow, Bernhard. 2020. "Avoiding risk and avoiding evidence." *Australasian Journal of Philosophy* 98:495–515.
- Caplin, Andrew and Dean, Mark. 2013. "Behavioral implications of rational inattention with Shannon entropy." National Bureau of Economic Research Working Paper 19318, <http://www.nber.org/papers/w19318>.
- Castro, Clinton and Pham, Adam. 2020. "Is the attention economy noxious?" *Philosophers' Imprint* 20:1–13.
- Christensen, David. 1991. "Clever bookies and coherent beliefs." *Philosophical Review* 100:229–47.
- Clark, Tez. ms. "Ducks in a row: On order and inquiry." ms.
- Das, Nilanjan. 2023. "The value of biased information." *British Journal for the Philosophy of Science* 74:25–55.
- Davenport, Thomas and Beck, John. 2001. *The attention economy: Understanding the new economy of business*. Harvard Business School Press.
- Dorst, Kevin. 2020. "Evidence: A guide for the uncertain." *Philosophy and Phenomenological Research* 100:586–632.
- Dorst, Kevin, Levinstein, Ben, Salow, Bernhard, Husic, Brooke, and Fitelson, Branden. 2021. "Deference done better." *Philosophical Perspectives* 35:99–150.
- Geanakoplos, John. 1989. "Game theory without partitions, and applications to speculation and consensus." Cowles Foundation Discussion Paper 914.
- Golman, Russell, Hagmann, David, and Lowenstein, George. 2017. "Information avoidance." *Journal of Economic Literature* 55:96–135.

- Good, I.J. 1966. "On the principle of total evidence." *British Journal for the Philosophy of Science* 17:319–321. doi:10.1093/bjps/17.4.319.
- Green, Mitchell and Hitchcock, Christopher. 1994. "Reflections on reflection: van Fraassen on belief." *Synthese* 98:297–324.
- Hertwig, Ralph and Engel, Christoph (eds.). 2021. *Deliberate ignorance*. MIT Press.
- Horwich, Paul. 1982. *Probability and evidence*. Cambridge University Press.
- Huttegger, Simon. 2013. "In defense of reflection." *Philosophy of Science* 80:413–33.
- . 2014. "Learning experiences and the value of knowledge." *Philosophical Studies* 171:279–88.
- Huttegger, Simon and Nielsen, Michael. 2020. "Generalized learning and conditional expectation." *Philosophy of Science* 87:868–83.
- Irving, Zachary and Glasser, Aaron. 2020. "Mind wandering: A philosophical guide." *Philosophy Compass* 15:e12644.
- Kadane, Joseph, Schervish, Mark, and Seidenfeld, Teddy. 1996. "Reasoning to a foregone conclusion." *Journal of the American Statistical Association* 91:1228–35.
- . 2008. "Is ignorance bliss?" *Journal of Philosophy* 105:5–36.
- Kluger, Avraham and DeNisi, Angelo. 1996. "The effects of feedback interventions on performance: A historical review, a meta-analysis, and a preliminary feedback intervention theory." *Psychological Bulletin* 119:254–84.
- Lacko, James and Pappalardo, Janis. 2004. "The effect of mortgage broker compensation disclosures on consumers and competition: A controlled experiment." Federal Trade Commission Bureau of Economics Staff Report, <https://www.ftc.gov/reports/effect-mortgage-broker-compensation-disclosures-consumers-competition-controlled-experiment>.
- Lester, Benjamin, Persico, Nicola, and Visschers, Ludo. 2012. "Information acquisition and the exclusion of evidence in trials." *Journal of Law, Economics and Organization* 28:163–82.
- Lieder, Falk and Griffiths, Thomas. 2020. "Resource-rational analysis: Understanding human cognition as the optimal use of limited computational resources." *Behavioral and Brain Sciences* 43:E1.
- Lipnowski, Elliott, Mathevet, Lauren, and Wei, Dong. 2020. "Attention management." *American Economic Review: Insights* 2:17–32.
- Maćkowiak, Bartosz, Matějka, Filip, and Wiederholt, Mirko. 2023. "Rational inattention: A review." *Journal of Economic Literature* 61:226–73.
- Maher, Patrick. 1990a. "Symptomatic acts and the value of evidence in causal decision theory." *Philosophy of Science* 57:479–98.

- . 1990b. “Why scientists gather evidence.” *British Journal for the Philosophy of Science* 41:103–119. doi:10.1093/bjps/41.1.103.
- Neth, Sven. forthcoming. “Rational aversion to information.” *British Journal for the Philosophy of Science* forthcoming.
- Shannon, Claude. 1948. “A mathematical theory of communication.” *The Bell System Technical Journal* 27:379–423.
- Siegel, Susanna. 2017. *The rationality of perception*. Oxford University Press.
- Simon, Herbert. 1971. “Designing organizations for an information-rich world.” In Martin Greenberger (ed.), *Computers, communications, and the public interest*, 37–72. Johns Hopkins Press.
- Sims, Christopher. 2003. “Implications of rational inattention.” *Journal of Monetary Economics* 50:665–90.
- Skyrms, Brian. 1990. *The dynamics of rational deliberation*. Harvard University Press.
- Talbott, William. 1991. “Two principles of Bayesian epistemology.” *Philosophical Studies* 62:135–50.
- van den Berg, Ronald and Ma, Wei Ji. 2018. “A resource-rational theory of set size effects in human visual working memory.” *eLife* 7:e34963.
- van Fraassen, Bas. 1984. “Belief and the will.” *Journal of Philosophy* 81:235–56.

A. Appendix

A.1 Full disclosure: Posteriors

Since c_f, c_g, c_h are probability functions, we have:

$$c_f(w_1) + c_f(w_2) + c_f(w_3) = 1 \quad (21)$$

$$c_g(w_1) + c_g(w_2) + c_g(w_3) = 1 \quad (22)$$

$$c_h(w_1) + c_h(w_2) + c_h(w_3) = 1 \quad (23)$$

By Caplin and Dean (2013) Theorem 1, for each pair a, b of acts from f, g, h and each state w_i we have:

$$c_a(w_i) = c_b(w_i)e^{u(a, w_i) - u(b, w_i)} \quad (24)$$

Solving (24) for c_g, c_h in terms of c_f and substituting into (22)-(23) gives:

$$e^{-1}c_f(w_1) + ec_f(w_2) + e^3c_f(w_3) = 1 \quad (25)$$

$$e^{-4}c_f(w_1) + c_f(w_2) + e^4c_f(w_3) = 1 \quad (26)$$

Solving the system of equations (22), (25), (26) for c_g and calculating c_g, c_h in terms of c_f using (24) gives c_f, c_g, c_h as in Table 8.

	Exact			Approximate		
	w_1	w_2	w_3	w_1	w_2	w_3
c_f	$\frac{e^4}{e(e-1)(1+e)^2}$	$\frac{e^3-(1+e+e^2)}{e(e-1)(1+e)^2}$	$\frac{1}{e(e-1)(1+e)^2}$	0.845	0.140	0.015
c_g	$\frac{e^3}{e(e-1)(1+e)^2}$	$\frac{e^4-e(1+e+e^2)}{e(e-1)(1+e)^2}$	$\frac{e^3}{e(e-1)(1+e)^2}$	0.311	0.378	0.311
c_h	$\frac{1}{e(e-1)(1+e)^2}$	$\frac{e^3-(1+e+e^2)}{e(e-1)(1+e)^2}$	$\frac{e^4}{e(e-1)(1+e)^2}$	0.015	0.140	0.845

Table 8: Exact and approximate probabilities $c_a(w_i)$. Subsequent calculations use approximate probabilities.

A.2 Full disclosure: Prior probabilities of posteriors

As before, let us abuse notation to let c_f, c_g, c_h stand for both credence functions and the propositions that these credence functions are adopted as the agent's posterior. By the law of total probability and Reflection, we have for each world w_i :

$$c(w_i) = \sum_{a \in \{f, g, h\}} c(c_a) c(w_i | c_a) \quad (27)$$

$$= \sum_{a \in \{f, g, h\}} c(c_a) c_a(w_i) \quad (28)$$

Plugging in the previously calculated values for $c_a(w_i)$ and the prior $(1/3, 1/3, 1/3)$ to (28) gives a system of three equations:

$$1/3 = 0.845c(c_f) + 0.311c(c_g) + 0.015c(c_h) \quad (29)$$

$$1/3 = 0.140c(c_f) + 0.378c(c_g) + 0.140c(c_h) \quad (30)$$

$$1/3 = 0.015c(c_f) + 0.311c(c_g) + 0.845c(c_h) \quad (31)$$

Which is solved at $c(c_f), c(c_g), c(c_h) = 0.094, 0.812, 0.094$.

A.3 Full disclosure: Signal inducing posteriors

Posteriors c_f, c_g, c_h are induced by signals s_f^2, s_g^2, s_h^2 . Since π_1^F is a full disclosure policy, it produces signals $s_{w_1}^1, s_{w_2}^1, s_{w_3}^1$ with $c(s_{w_i}^1 | w_j) = 1$ if $i = j$ and 0 if $i \neq j$. This means that the signals $s_{w_i}^1$ have priors $c(s_{w_i}^1) = c(w_i) = 1/3$.

To find the optimal attention policy π_2 inducing posteriors c_f, c_g, c_h , we need to find $c(s_a^2|s_{w_i}^1)$ for each $a \in \{f, g, h\}$ and $i \in \{1, 2, 3\}$. Note that for a, i as before, we have:

$$c(s_a^2|s_{w_i}^1) = \frac{c(s_{w_i}^1|s_a^2)c(s_a^2)}{c(s_{w_i}^1)} \quad (32)$$

$$= \frac{c(w_i|s_a^2)c(s_a^2)}{c(s_{w_i}^1)} \quad (33)$$

$$= \frac{c_a(w_i)c(c_a)}{c(s_{w_i}^1)} \quad (34)$$

We saw above that $c(s_{w_i}^1) = 1/3$ for each i . Posteriors $c_a(w_i)$ were found in Section A.1. Priors $c(c_a)$ were found in Section A.2. Substituting into (34) and solving gives:

	s_f^2	s_g^2	s_h^2
$s_{w_1}^1$	0.238	0.758	0.004
$s_{w_2}^1$	0.039	0.922	0.039
$s_{w_3}^1$	0.004	0.758	0.238

Table 9: Conditional probabilities $c(s_a^2|s_{w_i}^1)$ of attention outcomes s_a^2 given information outcomes $s_{w_i}^1$.

A.4 Full disclosure: Expected utility

Note that:

$$E_c[V(\Gamma_{\pi_2, \pi_1})] = \sum_{a \in \{f, g, h\}} c(c_a)V(\Gamma_{c_a}). \quad (35)$$

Since:

$$V(\Gamma_{c_f}) = (0.845)(0) + (0.140)(-1) + (0.015)(-4) = -0.200 \quad (36)$$

$$V(\Gamma_{c_g}) = (0.311)(-1) + (0.378)(0) + (0.311)(-1) = -0.622 \quad (37)$$

$$V(\Gamma_{c_h}) = (0.015)(-4) + (0.140)(-1) + (0.845)(0) = -0.200 \quad (38)$$

Substituting (36)-(38) into (35) and solving gives $E_c[V(\Gamma_{\pi_2, \pi_1})] = -0.543$.

A.5 Full disclosure: Expected inaccuracy

Note that:

$$\mathcal{I}_c(c_{\pi_2, \pi_1}) = \sum_{w \in \Omega, a \in \{f, g, h\}} c(c_a) c(w|c_a) \mathcal{I}(c_a, w) \quad (39)$$

$$= \sum_{w \in \Omega, a \in \{f, g, h\}} c(c_a) c_a(w) \mathcal{I}(c_a, w) \quad (40)$$

$$= \sum_{a \in \{f, g, h\}} c(c_a) \mathcal{I}_{c_a}(c_a). \quad (41)$$

Calculating the values of $\mathcal{I}_{c_a}(c_a)$ and recalling the priors $c(c_a)$ found in Section A.2 gives:

	$\mathcal{I}_{c_a}(c_a)$	$c(c_{\pi_2, \pi_1} = c_a)$
c_f	0.266	0.094
c_g	0.664	0.812
c_h	0.266	0.094

So that $\mathcal{I}_c(c_{\pi_2, \pi_1}) = (0.266)(0.094) + (0.664)(0.812) + (0.266)(0.094) = 0.590$.

A.6 Partial disclosure: Posteriors

It is never optimal to have more attention outcomes than information outcomes, so there must be at most two attention outcomes. It is also never optimal to have two attention outcomes favoring the same act. This leaves four possible combinations of attention outcomes, labeled in terms of the acts they favor:

Case 1: c_f, c_g

Case 2: c_g, c_h

Case 3: c_f, c_h

Case 4: c_a for some a .

The argument against Case 1 will show by symmetry that we are not in Case 2. Ruling out Case 4 will then leave Case 3. Along the way, we will also identify the posteriors in Case 3, which are the posteriors that result from rational inattention.

A.6.1 Case 1: c_f, c_g

Note that $u(f, s_{(12)}^1) = 2/3(0) + 1/3(-1) = -1/3$ and calculating in the same way, we have:

Since c_f, c_g are probability functions we have:

$$c_f(s_{(12)}^1) + c_f(s_{(23)}^1) = 1. \quad (42)$$

$$c_g(s_{(12)}^1) + c_g(s_{(23)}^1) = 1. \quad (43)$$

	f	g
$s_{(12)}^1$	-1/3	-2/3
$s_{(23)}^1$	-3	-2/3

Table 10: Utilities of acts f, g given partial disclosure signals $s_{(12)}^1, s_{(23)}^1$.

And solving (24) for c_g in terms of c_f as before, then substituting into (43) gives:

$$e^{-1/3}c_f(s_{(12)}^1) + e^{7/3}c_f(s_{(23)}^1) = 1. \quad (44)$$

Solving (42) and (44) gives:

	c_f	c_g
$s_{(12)}^1$	.970	.695
$s_{(23)}^1$	.030	.305

Table 11: Posterior credences $c_a(s_k^1)$ in information outcomes after partial disclosure, Case 1.

Note that for $a \in \{f, g\}$ and $i \in \{1, 2, 3\}$ we have:

$$c_a(w_i) = \sum_{b \in \{(12), (23)\}} c_a(s_b^1) c_a(w_i | s_b^1). \quad (45)$$

(45) allows us to express c_f, c_g in terms of w_1, w_2, w_3 as in Table 12. But this is impossible, since:

$$1/3 = c(w_1) = c(c_f)c_f(w_1) + c(c_g)c_g(w_1) \quad (46)$$

$$= c(c_f)(0.647) + (1 - c(c_f))(0.462) \geq 0.462 \quad (47)$$

	c_f	c_g
w_1	.647	.462
w_2	.333	.333
w_3	.020	.205

Table 12: Posterior credences $c_a(w_i)$ in worlds after partial disclosure, Case 1.

A.6.2 Case 3: c_f, c_h

We can calculate the utilities $u(a, s_b^1)$ of acts given information outcomes as before:

Since c_f, c_h are probability functions we have:

	f	h
$s_{(12)}^1$	-1/3	-3
$s_{(23)}^1$	-3	-1/3

Table 13: Utilities of acts f, h given partial disclosure signals $s_{(12)}^1, s_{(23)}^1$.

$$c_f(s_{(12)}^1) + c_f(s_{(23)}^1) = 1 \quad (48)$$

$$c_h(s_{(12)}^1) + c_h(s_{(23)}^1) = 1 \quad (49)$$

As before, solving (24) for c_h in terms of c_f and substituting into (49) gives:

$$e^{-8/3}c_f(s_{(12)}^1) + e^{8/3}c_f(s_{(23)}^1) = 1 \quad (50)$$

Solving (48) and (50) for c_f and using (24) to find c_h in terms of c_f gives:

	c_f	c_h
$s_{(1,2)}^1$	0.935	0.065
$s_{(2,3)}^1$	0.065	0.935

Table 14: Posterior probabilities $c_a(s_b^1)$ of information outcomes s_b^1 .

As before, we can use (45) to express c_f and c_h in terms of w_1, w_2, w_3 as:

	c_f	c_h
w_1	0.623	0.044
w_2	0.333	0.333
w_3	0.044	0.623

Table 15: Posterior probabilities $c_a(w_i)$ of worlds w_i .

A.6.3 Case 4: c_a

Given a single attention outcome c_a , Reflection forces $c_a = c$. This yields no improvement in decision quality and incurs no attention cost, so has value zero, less than the value of Case 3. So Case 4 is suboptimal.

A.7 Partial disclosure: Prior probabilities of posteriors

As before, note that:

$$c(w_i) = \sum_{a \in \{f, h\}} c(c_a) c_a(w_i) \quad (51)$$

For worlds w_1, w_3 , this generates a system of two equations:

$$1/3 = 0.623c(c_f) + 0.044c(c_h) \quad (52)$$

$$1/3 = 0.044c(c_f) + 0.623c(c_h) \quad (53)$$

These equations are solved at $c(c_f) = c(c_h) = 0.5$.

A.8 Partial disclosure: Signal inducing posteriors

The signaling structure involves signals s_f^2, s_h^2 inducing c_f, c_h , respectively. Note that for all $a \in \{f, h\}$ and $b \in \{(1, 2), (2, 3)\}$ we have:

$$c(s_a^2 | s_b^1) = c(c_a | s_b^1) \quad (54)$$

$$= c(s_b^1 | c_a) c(c_a) / c(s_b^1) \quad (55)$$

$$= c_a(s_b^1) (0.5) / (0.5) \quad (56)$$

$$= c_a(s_b^1) \quad (57)$$

So we can read the signal structure directly from Table 14, giving Table 16.

	s_f^2	s_h^2
$s_{(1,2)}^1$	0.935	0.065
$s_{(2,3)}^1$	0.065	0.935

Table 16: Conditional probabilities $c(s_a^2 | s_b^1)$ of attention outcomes s_a^2 given information outcomes s_b^1 .

A.9 Partial disclosure: Expected utility

We can calculate $E[V(\Gamma_{c_f})] = E[V(\Gamma_{c_h})] = -0.509$, forcing $E[V(\Gamma_{\pi_2, \pi_1})] = -0.509$.

A.10 Partial disclosure: Expected inaccuracy

As before, note that:

$$\mathcal{I}_c(c_{\pi_2 \cdot \pi_1}) = \sum_{w \in \Omega, a \in \{f, h\}} c(c_a) c(w|c_a) \mathcal{I}(c_a, w) \quad (58)$$

$$= \sum_{w \in \Omega, a \in \{f, h\}} c(c_a) c_a(w) \mathcal{I}(c_a, w) \quad (59)$$

$$= \sum_{a \in \{f, h\}} c(c_a) \mathcal{I}_{c_a}(c_a). \quad (60)$$

Calculating yields $\mathcal{I}_{c_f}(c_f) = \mathcal{I}_{c_h}(c_h) = 0.499$ forcing $\mathcal{I}_c(c_{\pi_2 \cdot \pi_1}) = 0.499$.