

ESCOLA DE HUMANIDADES
PROGRAMA DE PÓS-GRADUAÇÃO EM FILOSOFIA
MESTRADO EM FILOSOFIA

NATÃ OLIVEIRA DA COSTA

**INFERÊNCIA À MELHOR EXPLICAÇÃO REINVENTADA:
O IMPACTO DO DESAFIO SCREEN-OFF**

Porto Alegre
2024

PÓS-GRADUAÇÃO - *STRICTO SENSU*



Pontifícia Universidade Católica
do Rio Grande do Sul

NATÃ OLIVEIRA DA COSTA

INFERÊNCIA À MELHOR EXPLICAÇÃO REINVENTADA:
O IMPACTO DO DESAFIO SCREEN-OFF

Dissertação apresentada como requisito para a
obtenção do grau de Mestre pelo Programa de
Pós-Graduação em Filosofia da Escola de
Humanidades da Pontifícia Universidade
Católica do Rio Grande do Sul.

Orientador: Prof. Dr. Claudio Gonçalves de Almeida

Porto Alegre

2024

Ficha Catalográfica

C837i Costa, Natã Oliveira da

Inferência à melhor explicação reinventada : O impacto do desafio screen-off / Natã Oliveira da Costa. – 2024.

77.

Dissertação (Mestrado) – Programa de Pós-Graduação em Filosofia, PUCRS.

Orientador: Prof. Dr. Cláudio Gonçalves de Almeida.

1. Inferência à melhor explicação. 2. Desafio screen-off. 3. Explicatividade. 4. Relevância epistêmica. I. Almeida, Cláudio Gonçalves de. II. Título.

Elaborada pelo Sistema de Geração Automática de Ficha Catalográfica da PUCRS
com os dados fornecidos pelo(a) autor(a).

Bibliotecária responsável: Clarissa Jesinska Selbach CRB-10/2051

NATÃ OLIVEIRA DA COSTA

INFERÊNCIA À MELHOR EXPLICAÇÃO REINVENTADA:
O IMPACTO DO DESAFIO SCREEN-OFF

Dissertação apresentada como requisito para a obtenção do grau de Mestre pelo Programa de Pós-Graduação em Filosofia da Escola de Humanidades da Pontifícia Universidade Católica do Rio Grande do Sul.

Aprovada em: ____ de _____ de 2024.

Banca Examinadora:

Prof. Dr. Claudio Gonçalves de Almeida – PUCRS

Prof. Dr. Luis Fernando Munaretti da Rosa – Washington University, St. Louis

Prof. Dr. André Luiz de Almeida Lisboa Neiva - UFAL

AGRADECIMENTOS

À Maiara, minha esposa, pelo suporte, encorajamento e companheirismo.

Aos meus filhos, Thomas e Melissa, pela alegria que me trazem todos os dias.

Aos meus pais, Alvacyr e Marlene, por todo amor, apoio e investimento.

Ao meu orientador, Claudio de Almeida, por acreditar na minha capacidade e, especialmente, por ser um exemplo de perspicácia e agudeza de espírito.

O presente trabalho foi realizado com apoio da Coordenação de Aperfeiçoamento de Pessoal
Nível Superior – Brasil (CAPES) – Código de Financiamento 001.

“Behind it all is surely an idea so simple, so beautiful, that when we grasp it – in a decade, a century, or a millennium – we will all say to each other, how could it have been otherwise?”

(Wheeler, 1986)

RESUMO

O presente trabalho tem como objetivo investigar o impacto do chamado desafio *Screen-off* (Sober & Roche, 2013; 2019) sobre as pesquisas em torno da Inferência à Melhor Explicação (IME). No segundo capítulo, a fim de melhor compreender o escopo da objeção, procuro esclarecer os contornos gerais da IME na literatura contemporânea: em especial, as diferentes atitudes que se pode manter em relação a seu status epistêmico, os desafios enfrentados em sua articulação e as principais propostas existentes de como contorná-los. No terceiro capítulo, passo à análise do argumento *Screen-off* propriamente dito e das respostas apresentadas até o momento. Em particular, examino a tese de que a explicatividade torna os graus de crença em uma determinada hipótese mais resilientes (McCain e Poston, 2014, 2017), a sugestão de que a falha reside, em verdade, na idealização de Onisciência Lógica inerente ao aparato Bayesiano (McCain e Poston, 2014, 2017; Lange, 2017), a tese de que a explicatividade é indispensável para a projeção indutiva de certos predicados (Climenhaga, 2017) e, por fim, a ideia de que considerações explanatórias possuem relevância epistêmica pelo fato de nos indicarem qual o tipo de explicação esperar para um dado *explanandum* (Lange, 2017, 2020, 2024). Concluo que nenhuma das respostas explanacionistas apresentadas consegue assegurar um papel epistêmico à explicatividade e que as concessões feitas no curso do debate modificam o discurso acerca da Inferência à Melhor Explicação substancialmente.

Palavras-chave: Inferência à melhor explicação. Desafio screen-off. Explicatividade. Relevância epistêmica.

ABSTRACT

The present work aims to investigate the impact of the Screen-off challenge (Sober & Roche, 2013; 2019) on the research on Inference to the Best Explanation (IBE). In the second chapter, in order to better understand the scope of the objection, I present the general outlines of IBE: in particular, the different attitudes one can have towards its epistemic status, the challenges faced in its articulation, and the main proposals on how to overcome them. In the third chapter, I analyze the Screen-off argument itself and the responses presented so far. In particular, I examine the thesis that explanatoriness makes the degrees of belief in a given hypothesis more resilient (McCain and Poston, 2014, 2017), the suggestion that the flaw actually lies in the idealization of Logical Omniscience inherent in the Bayesian apparatus (McCain and Poston, 2014, 2017; Lange, 2017), the thesis that explanatoriness is indispensable for the inductive projection of certain predicates (Climenhaga, 2017), and finally, the idea that explanatory considerations have epistemic relevance because they indicate what type of explanation one should expect for a given *explanandum* (Lange, 2017, 2020, 2024). I conclude that none of the explanationist responses presented is able to secure an epistemic role for explanatoriness and that the concessions made within them modify the discourse surrounding Inference to the Best Explanation substantially.

Keywords: Inference to the Best Explanation. Screen-off challenge. Explanatoriness. Epistemic relevance.

SUMÁRIO

1. INTRODUÇÃO	10
2. INFERÊNCIA À MELHOR EXPLICAÇÃO	13
2.1. DEFINIÇÃO	13
2.2. TIPOS DE EXPLANACIONISMO	16
2.3. ESTRUTURA, DESAFIOS E PROPOSTAS	18
2.3.1. A natureza do explanandum	19
2.3.2. As hipóteses elegíveis	20
2.3.3. A identificação da melhor explicação	22
2.3.3.1. <i>A delimitação do rol de explicações relevantes</i>	23
2.3.3.2. <i>A seleção e justificação das virtudes explicativas</i>	26
2.3.4. O caráter da conclusão	31
3. O DESAFIO <i>SCREEN-OFF</i>	33
3.1. BAYESIANISMO VIS A VIS INFERÊNCIA À MELHOR EXPLICAÇÃO	34
3.2. O ARGUMENTO DE SOBER & ROCHE	38
3.3. AS RESPOSTAS EXPLANACIONISTAS	42
3.3.1. Explicatividade, Peso da Evidência e Resiliência dos Graus de Crença .	43
3.3.2. O Problema da Onisciência Lógica	47
3.3.3. Considerações explanatórias e projeção indutiva	53
3.3.4. IME, Suposições de fundo e Tipos de Explicação	57
4. CONSIDERAÇÕES FINAIS	62
5. REFERÊNCIAS	64

1. INTRODUÇÃO

A capacidade de uma teoria ou hipótese de explicar determinada observação é comumente tida como sendo, em alguma medida, evidência de sua veracidade. É dizer, acredita-se que “obtemos evidência de que p ao percebermos que, se verdadeira, ela explicaria a evidência disponível” (McCain e Poston, 2024, p. 329). As Leis de Mendel, por exemplo, explicam a transmissão de caracteres hereditários; a Teoria do Big Bang explica a radiação cósmica de fundo (RCFM); e a Teoria Tectônica de Placas, dentre outras coisas, explica a ocorrência de terremotos, a formação de montanhas e o formato das costas continentais. E esses fatos seriam motivos – talvez *os* motivos – para que os aceitemos.

Essa posição, denominada Explanacionismo, é, aliás, mais do que uma reconstrução filosófica acurada. Sabemos que os próprios autores de algumas de nossas teorias mais bem-sucedidas concebiam a capacidade explicativa delas desse modo. Segundo Darwin (1962 [1872], p. 476), não era razoável “supor que uma teoria falsa pudesse explicar, de maneira tão satisfatória, como o faz a teoria da seleção natural, as diversas grandes séries de fatos” por ele listadas. De modo semelhante, Einstein (1915) afirmou que uma importante confirmação da Teoria da Relatividade Geral consistia no fato de que a precessão do periélio de mercúrio “é [por ela] explicada qualitativa e quantitativamente” (p. 831).

Importante destacar, esse apreço pela capacidade explicativa de uma hipótese não é exclusividade das ciências experimentais. Na Arqueologia, a Inferência à Melhor Explicação (IME) “é prática padrão já há quase um século” (Fogelin, 2007, p. 604). No Direito Penal, boa parte dos juristas também entende que a condenação do réu depende de a hipótese da culpa ser capaz de explicar a evidência disponível (Allen, 2019). Na área da Medicina, também há quem defenda “ser mais proveitoso para pesquisadores focar em inferências à melhor explicação do que continuar focando em indução” (Binney, 2023, p. 1). E de acordo com alguns pesquisadores, a negligência dessa dimensão do processo de avaliação de hipóteses seria um dos problemas a ser solucionado nos estudos conduzidos em Psicologia Cognitiva (Borsboom *et al*, 2021; Haig, 2009, 2018).

Na Filosofia, onde dificilmente uma controvérsia será resolvida com simples apelos à observação, considerações de cunho explanatório são ainda mais ubíquas. Na Epistemologia, uma das respostas mais proeminentes ao ceticismo cartesiano defende a resposta do senso comum com base em seus méritos explicativos (Vogel, 1990; 2008; 2013; McCain, 2016). Na Filosofia da Ciência, o Explanacionismo aparece mais claramente na defesa do Realismo Científico. Segundo seus proponentes, deveríamos aceitar a existência das entidades postuladas

por nossas teorias (elétrons, quarks, energia escura etc) por serem a melhor explicação disponível para sua adequação empírica (Musgrave, 1988; Psillos, 2005; Alai, 2023). E não seria difícil multiplicar os exemplos, com referências provenientes de áreas tão diversas quanto Ética (Bailey, 1997; Leiter, 2014), Filosofia da Mente (McLaughlin, 2010; Carruthers, 2015) e Filosofia da Religião (Dawes, 2012; Rasmussen & Leon, 2019).

Tal cenário pode sugerir que a relevância epistêmica da explicatividade seja tão incontroversa quanto uma tese filosófica pode ser. No entanto, para a surpresa de muitos, de acordo com Sober & Roche (2013), assim como o imperador, o explanacionista “has no clothes. Os autores sustentam que, uma vez que dados relativos à frequência¹ tenham sido tomados em conta, como nos paradigmáticos casos de diagnósticos médicos, é possível perceber que a explicatividade de uma hipótese – o fato de que, se verdadeira, ela explicaria as observações em questão – não possui qualquer valor evidencial².

O argumento é apresentado por meio de um exemplo envolvendo a relação do consumo de cigarros com o desenvolvimento de câncer de pulmão. Sober & Roche (2013) nos relembram que, previamente à descoberta de que o primeiro causa o último, cientistas já haviam notado uma correlação. Quanto maior o número de cigarros consumidos, maior era a chance de alguém contrair câncer de pulmão. Entretanto, como Ronald Fisher (1959) observou à época, a correlação poderia muito bem ser explicada de outras formas, como, por exemplo, apelando-se à existência de um fator comum, como uma predisposição genética, capaz de tornar o indivíduo suscetível a ambos os males. E aqui entra o ponto crucial para o argumento *screen-off*: isso implica que as chances de alguém vir a receber tal diagnóstico após consumir X maços de cigarro ao longo da vida era a mesma quer o consumo de cigarros fosse a causa – e, portanto, a explicação para o desenvolvimento da doença – ou apenas um efeito de uma causa comum. A explicatividade, então, nesse caso, e em inúmeros os outros semelhantes, “não possui qualquer relevância evidencial” (Sober & Roche, 2013, pág. 667).

Como se isso não fosse surpreendente o bastante, as respostas explanacionistas contêm algumas concessões curiosas. McCain e Poston (2014, p. 147), por exemplo, reconhecem que “acrescentar a alegação explicativa não altera o grau de confirmação”, enquanto Lange (2022, p. 87) afirma que “a ‘melhor explicação’ de um fato *não precisa* ser a hipótese mais plausível, considerando todos os aspectos” (grifei). Este último chega a declarar que “só pensamento mágico” sobre a natureza da IME “sugeriria que ela exige a existência de uma lista fixa de

¹ No original: *frequency data*

² Emprego “relevância evidencial”, “relevância epistêmica” e variações de forma intercambiável.

'virtudes explicativas', cuja posse por qualquer hipótese a tornaria mais plausível, independentemente do nosso conhecimento prévio” (2020, p. 40).

O objetivo geral da presente dissertação é investigar o impacto do desafio *Screen-off*³ (Sober & Roche, 2013; 2019) sobre o status da Inferência à Melhor Explicação (IME). Antes de adentrarmos à discussão propriamente dita, porém, algumas questões preliminares devem ser enfrentadas. Quarenta anos após sua primeira articulação por Harman (1965), Lipton (2004, p. 2) ainda pôde se queixar que a IME permanecia “mais um slogan do que uma teoria filosófica bem articulada”. Por essa razão, no segundo capítulo, buscar-se-á esclarecer em que, em primeiro lugar, consiste a Inferência à Melhor Explicação. Por meio de uma abordagem historicamente informada da literatura contemporânea, procuraremos identificar o que atualmente se pode afirmar de mais preciso sobre ela. Começando por sua definição, exploraremos as diferentes atitudes que se pode manter quanto ao status epistêmico da explicatividade, as dificuldades encontradas em cada uma das etapas em que (pelo menos para fins didáticos) se pode fracioná-la e as propostas existentes de como melhor contorná-las. O objetivo naturalmente é perquirir quais são seus elementos essenciais e, por consequência, elucidar se todos ou apenas alguns dos modelos disponíveis são vulneráveis ao argumento *Screen-off*.

De posse do arcabouço teórico necessário, no terceiro capítulo, passaremos à análise do escopo, da cogência e das implicações do desafio *Screen-off*. Após sintetizarmos a tese avançada por Sober e Roche (2013; 2019), avaliaremos a sugestão de que a explicatividade tem a propriedade de tornar graus de crença mais insuscetíveis de mudança em face a evidências futuras (McCain e Poston, 2014; 2017), de que ela é indispensável para a projeção indutiva de predicados (Climenhaga, 2017; McCain e Poston, 2014) e de que ela sinaliza quais tipos de explicações se pode esperar para um determinado *explanandum* (Lange, 2017; 2020; 2024). No percurso, também dedicaremos espaço ao Problema da Onisciência Lógica no Bayesianismo, que figura com bastante destaque na controvérsia.

Convém ressaltar, contudo, que o trabalho não é meramente exegetico. Ao longo do texto, quando oportuno, compartilho não só minha perspectiva sobre a qualidade de determinados argumentos como apresento algumas contribuições originais à discussão. As conclusões, assim como as linhas de pesquisa que se revelam mais promissoras daqui em diante, são sintetizadas nas considerações finais.

³ Assim como os autores que protagonizaram a disputa, utilizo “desafio *screen-off*”, “argumento *screen-off*” e “SOT” (*screen-off thesis*) como expressões sinônimas.

2. INFERÊNCIA À MELHOR EXPLICAÇÃO

2.1. DEFINIÇÃO

Os debates contemporâneos sobre a natureza da Inferência à Melhor Explicação remontam ao trabalho de Harman (1965), que também detém o mérito de ter cunhado a expressão (Cabrera, 2023). E embora a tese por ele defendida no referido artigo (vide próximo tópico) não tenha conquistado muitos adeptos, sua definição de IME ainda fornece um bom ponto de partida (p. 89):

Nesta espécie de inferência, do fato de que uma dada hipótese explicaria a evidência, o agente infere a veracidade da hipótese. Em geral, haverá várias hipóteses que poderiam explicar a evidência, de modo que se deve ser capaz de rejeitar todas as alternativas antes de se poder traçar essa conclusão. Assim, da premissa de que uma determinada hipótese forneceria uma explicação “melhor” do que qualquer outra hipótese, infere-se que a hipótese em questão é verdadeira.

Este capítulo como um todo, de certa forma, pode ser visto como uma explicação dos conceitos contidos nesse parágrafo. Alguns deles, consequentemente, serão objeto de tópicos posteriores. No entanto, dado o objetivo de identificar uma definição mais geral, por ora duas considerações se fazem oportunas. A primeira se refere à identificação da IME como sendo um tipo de inferência, e a segunda, à forma como ela se relaciona com a famosa Abdução, introduzida por Charles S. Peirce.

Dellsén (2024) recentemente propôs uma taxonomia que divide os diferentes modelos existentes na literatura em três grupos: inferenciais, probabilísticos e híbridos⁴. Como o nome sugere, os primeiros “envolvem comparar várias hipóteses explicativas concorrentes em termos de quão boa explicação cada uma forneceria, e então aceitar a hipótese que forneceria a melhor explicação” (p. 15). Por sua vez, os modelos probabilísticos não envolveriam inferir uma explicação, mas atribuir maior probabilidade àquela que melhor explica as evidências. O autor explica (p. 16):

[Modelos probabilísticos] sustentam que o raciocínio abduativo envolve atribuir probabilidades a hipóteses explicativas e atualizar⁵ essas probabilidades à medida que mais evidências são obtidas. Assim, se uma hipótese explica um dado recebido particularmente bem, enquanto outra explica o dado de forma ruim ou não explica, modelos probabilísticos dizem que você deve aumentar a probabilidade atribuída à primeira em detrimento da probabilidade atribuída à segunda.

⁴ No original, “inferential accounts”, “probabilistic accounts” e “hybrid accounts”.

⁵ No original, “update”, em referência à regra Bayesiana da condicionalização.

Embora Dellsén (2024) esteja tomando o raciocínio abdutivo como uma categoria mais ampla, da qual o modelo de Harman é apenas uma das espécies (a primeira provavelmente), a distinção entre modelos inferenciais e probabilísticos é problemática. E isso por duas razões.

Em primeiro lugar, ela passa a impressão de que, no caso dos ditos modelos probabilísticos, o agente nada estaria a inferir propriamente, mas “apenas atualizando seus graus de crença”. É verdade que vir a considerar uma hipótese mais provável do que inicialmente se supunha não implica inferir que ela é *a* verdadeira explicação para determinado fato ou observação. Para utilizar um exemplo mundano, avistar um colega, enquanto me desloco para uma aula, caminhando na direção oposta do local onde a disciplina é ministrada confere certo suporte à hipótese de que a aula foi cancelada, mas dificilmente seria o suficiente para me fazer abandonar a caminhada em direção à sala. Não obstante, seria equivocado dizer que nada foi inferido neste caso. Um agente Bayesiano ideal, no mínimo, inferiu que a hipótese é mais provável do que ele originalmente supunha⁶.

Além disso, designar parte dos modelos de *probabilísticos* também sugere que os pertencentes à primeira categoria (“inferenciais”) não envolveriam juízos de probabilidade, o que é enganoso. Inferências à Melhor Explicação, no sentido harmaniano, de fato, não *precisam* ser juízos probabilísticos. Por exemplo, se o conhecimento de fundo indica que as explicações sob consideração exaurem as possibilidades, e todas, à exceção de uma, são refutadas, a inferência de que a remanescente é a verdadeira decorre via dedução. Todavia, tais casos de inferência à única explicação (Bird, 2005; Woodward, forthcoming) são raríssimos, especialmente na Ciência e na Filosofia. Em razão da não-monotonicidade própria das inferências ampliativas, estas, como regra, são probabilísticas.

Por esses motivos, parece-me que, para evitar confusões desnecessárias, convém reconhecer que ambas as espécies de modelos são, no fim das contas, inferenciais e, na maior parte dos casos, probabilísticos. E se persiste a necessidade de sinalizar quais deles são mais afetos ao Bayesianismo, talvez seja melhor distingui-los simplesmente por modelos Bayesianos e Não-Bayesianos.

O caráter inferencial da IME levanta a questão de como ela se relaciona com a chamada Abdução, o segundo ponto a ser elucidado. Até o século XX, por influência de Aristóteles, nossas inferências eram tidas como sendo, basicamente, de dois tipos: dedução e indução

⁶ Agentes reais, por sua vez, podem inferir uma grande variedade de coisas em tais cenários, como “talvez a aula tenha sido cancelada”, “talvez o colega não esteja mais cursando a disciplina” e assim por diante.

(Cabrera, 2023). Entretanto, na segunda metade do século XIX, Charles S. Peirce afirmou ter identificado uma terceira espécie.

A Abdução, como a denominou, seria “o processo de formação de uma hipótese *explicatória*” (EP 2:216, 1992 [1903]). Referências como essa inicialmente levaram autores influentes a concebê-la como sendo praticamente idêntica à IME (Harman, 1965; Lipton, 2004, p. 56). Contudo, embora atualmente ainda existam autores que defendam a inconveniência ou mesmo impossibilidade de distingui-las de forma muito cartesiana (Schurz, 2017; Niiniluoto, 2018), o consenso se cristalizou no sentido de que se trata de operações bastante distintas (Minnameier, 2004; Mcauliffe, 2015; Campos, 2011).

Parte da confusão, é verdade, é atribuível ao próprio pragmatista inglês. A IME é bastante semelhante, senão idêntica, à chamada Hipótese ou Indução Hipotética, e Peirce admite ter confundido “Abdução com a segunda forma de indução, qual seja, a indução de qualidades” (PPM 277n3, 1997 [1903]). Ele teria se confundido em “quase tudo que escrevi antes desse século”⁷, ele diz, em uma carta a Paul Carus.

Atkins (2023, p. 69) esclarece:

[a indução qualitativa] é um raciocínio voltado a eliminar, corroborar ou confirmar hipóteses, colocando-as à prova. Isto é o que Peirce chama de hipótese em seus primeiros trabalhos. O segundo tipo é a abdução. A abdução é o gênero de inferência (a) que sugere hipóteses que podem ser dignas de serem perseguidas porque nos darão uma boa oportunidade para futuras investigações ou (b) que sugere a melhor explicação, dada a evidência que temos. Ela não elimina, corrobora ou confirma hipóteses colocando-as à prova de forma alguma.

A conotação (b) acima lembra muito a Inferência à Melhor Explicação. No entanto, percebe-se que a conclusão por ela autorizada é notadamente mais fraca. Isso se encontra ainda mais evidente na definição que se tornou canônica, apresentada em uma de suas palestras em Harvard de 1903 (EP 2:231, 1992):

- [i] O fato surpreendente C é observado.
- [ii] Mas se A fosse verdadeiro, C seria algo justamente o que nós observaríamos.
- [iii] Portanto, há razão para **suspeitar** que A seja verdadeiro. (grifei)

Peirce, portanto, não via a Abdução como apta a sozinha autorizar a conclusão mais forte de que a explicação entretida é, de fato, verdadeira. Na distinção sugerida por Hanson

⁷ É interessante que ainda hoje é possível encontrar contribuições em volumes renomados falhando em distinguir entre Abdução e Hipótese nos escritos de Peirce (Niiniluoto, 2024).

(1958), para Peirce, a Abdução diz respeito a “razões para sugerir uma teoria”, enquanto a IME, a “razões para aceitar uma teoria”.

Mohammadian (2021) fornece o que provavelmente seja a reconstrução histórica da relação entre esses dois conceitos mais precisa disponível na literatura. Ele explica que, apesar das semelhanças (p. 4226):

Elas também têm duas diferenças importantes. Primeiro, apenas a abdução envolve a capacidade de gerar hipóteses. Segundo, na abdução, a classificação de hipóteses é feita antes de se realizar testes empíricos e, portanto, *hipóteses não testadas* são classificadas por esse processo. Mas, na IME, a classificação de hipóteses é feita *após* a realização de testes empíricos, e os objetos de classificação são *hipóteses empiricamente equivalentes* que sobreviveram aos aludidos testes (grifo no original).

As diferenças, como também elucida, decorrem de dois desenvolvimentos na história da Filosofia da Ciência: a superveniência da distinção feita por Reichenbach (1938) entre contexto da descoberta e contexto da justificação e, posteriormente, o surgimento da teoria da subdeterminação (Quine, 1975), que impulsionou a busca por critérios extra-empíricos de avaliação de hipóteses. Como o título do artigo, aliás, bem resume, *Abduction – the context of discovery + underdetermination = inference to the best explanation*.

À vista destes apontamentos, pode-se dizer que o rótulo Inferência à Melhor Explicação, quando utilizado de forma cuidadosa, designa todo processo inferencial (i) de natureza inerentemente comparativa (ii) no qual a seleção da hipótese a ser preferida se dá com base em critérios outros além da adequação empírica. A definição ora articulada possui algumas vantagens. Como a condição (i) é necessária, mas não suficiente, a definição nos permite distinguir a IME da chamada indução eliminativa. Ao mesmo tempo, não exclui o trabalho de autores proeminentes, como Douven (2022), que optam pelo termo Abdução, embora tratem, em verdade, da IME.

2.2. TIPOS DE EXPLANACIONISMO

Como qualquer corrente filosófica que se preze, o Explanacionismo também está disponível nas mais variadas versões. Há, assim, diferentes formas em que se pode conceber a explicatividade de uma hipótese como evidencialmente relevante⁸. Neste tópico, exploraremos seus diferentes matizes, procurando identificar quais deles o argumento *Screen-off*, *prima facie*, parece ameaçar.

⁸ Utilizo “*evidencialmente relevante*”, apesar de o primeiro vocábulo não ser um advérbio comumente utilizado em língua portuguesa, por ser a tradução mais fiel da expressão empregada (evidentially relevant) nos artigos analisados.

Lycan (2002) tem o que, para os nossos propósitos, se revela a taxonomia mais útil nesse ponto. Segundo o autor, pode-se adotar, pelo menos, quatro diferentes tipos de Explanacionismo, que ele denominou de fraco, robusto, feroz e global (p. 417)⁹.

O Explanacionismo Fraco corresponde à tese de que inferências à melhor explicação são capazes de justificar epistemicamente uma conclusão. O Robusto, por sua vez, acrescenta que tal justificação não é redutível a nenhuma espécie mais básica de inferência ampliativa. O Explanacionismo Feroz agrega a tese de que a Inferência à Melhor Explicação é a única ou a mais fundamental forma de inferência ampliativa, de modo que a indução enumerativa, por exemplo, pode ser a ela reduzida. Por fim, o Explanacionismo Global dá um passo além e advoga que todas as espécies de inferência, mesmo as dedutivas, são fundamentalmente casos de IME.

À exceção de Harman (1965), para quem todas as formas de indução enumerativa defensáveis¹⁰ poderiam ser reduzidas à Inferência à Melhor Explicação, a maior parte dos autores explanacionistas são do tipo robusto. Ou seja, concebem a IME como uma terceira espécie inferencial insuscetível de redução a outras formas de raciocínio.

Nada obstante, apesar de mais modesto, o Explanacionismo Robusto também é alvo de objeções. Fumerton (1980, 2017), por exemplo, sustenta que as ditas inferências explicativas é que, *contra* Harman, seriam redutíveis a combinações de espécies inferenciais mais fundamentais. De acordo com o primeiro, casos paradigmáticos de IME seriam entimemas cujas premissas, quando tornadas explícitas, revelam a existência de um argumento indutivo ordinário ou uma combinação de dedução e indução. Desse modo, a famigerada inferência de que alguém passou pela praia há pouco com base na observação das pegadas deixadas na areia não seria legitimada por ser esta a melhor explicação *sic et simpliciter*. Mas por ser aquela que dispõe de maior suporte indutivo (1980, p. 592)¹¹.

E apesar de ser a objeção mais popular, ela não é a única. Há também quem defenda que apelos à qualidade explicativa de uma hipótese “são, na verdade, apelos implícitos a suposições empíricas substantivas, e não a alguma forma privilegiada de inferência” (Day e Kincaid, 1994)¹². Essa, aliás, é a ideia fundamental subjacente à Teoria Material da Indução proposta por

⁹ No original, *weak, sturdy, ferocious* e *global*.

¹⁰ *Warranted* no original.

¹¹ Não sobeja registrar que, embora a posição de Fumerton seja minoritária, as respostas a seu argumento não são das mais persuasivas. Para as respostas existentes, vide Schurz (2017), Niiniluoto (2018) e Weintraub (2013; 2017).

Norton (2021). Infelizmente, as limitações de espaço e os objetivos estipulados na introdução não nos permitem perseguir esses pontos em maiores detalhes. De todo modo, tais exemplos ilustram que, caso os argumentos avançados por Sober e Roche (2013; 2014; 2017; 2019) não sejam convincentes, o explanacionista ainda tem muito trabalho a fazer.

O argumento *Screen-off* também tem por alvo o Explanacionismo Robusto. Como já mencionado (e conforme será analisado em detalhes no capítulo 3), S&R¹³ defendem que o contrafactual “H, se verdadeira, explicaria O” não confere qualquer credibilidade à hipótese (H). SOT, portanto, nega que a explicatividade seja evidencialmente relevante – e, por consequência, que ela seja hábil a justificar epistemicamente qualquer conclusão.

Sober e Roche não interagem diretamente com a taxonomia proposta por Lycan, mas é possível concluir que, para os autores, a opção remanescente, o Explanacionismo Fraco, ainda que esteja fora do escopo do desafio, não serve de reduto ao explanacionista. Em uma de suas respostas a McCain e Poston (2014), Sober e Roche (2014, p. 198) sugerem, inclusive, se tratar de uma tese filosoficamente desinteressante:

[...] a teoria da Inferência à Melhor Explicação supostamente fornece uma epistemologia fundamental. A ideia é que a capacidade explicativa é evidencialmente relevante *em si mesma*; a proposta não é que de que a capacidade explicativa às vezes está correlacionada com outras propriedades mais fundamentais de uma hipótese, que, no fim das contas, é que estão realizando todo o trabalho epistêmico. Se pessoas que vestem blusões vermelhos tendem a acertar mais respostas do que pessoas que não o fazem, então usar um blusão vermelho é evidencialmente relevante. Mas dificilmente alguém irá propor uma teoria da inferência que tenha como seu conceito fundamental usar blusões vermelhos.

Como se observa, o desafio *screen-off* é tão ambicioso quanto as demais objeções existentes mencionadas - e talvez tão ambiciosa quanto qualquer objeção anti-explanacionista em geral possa ser.

2.3. ESTRUTURA, PROPOSTAS E DESAFIOS

Josephson e Josephson (1994, p. 5) fornecem uma das formulações clássicas dos contornos gerais da IME - cuja utilidade, inclusive, é reconhecida por outros autores (Psillos, 2002):

1. D é um conjunto de dados¹⁴ (fatos, observações, pressupostos¹⁵).
2. H explica D (se verdadeira, explicaria D).

¹³ Assim como os protagonistas do debate, nos casos de coautoria, utilizo a abreviação dos nomes dos autores por suas iniciais.

¹⁴ No original, *data*.

¹⁵ No original, *givens*.

3. Nenhuma outra hipótese explica D tão bem quanto H o faz.
4. Portanto, H é provavelmente verdadeira.

Nos tópicos a seguir, a partir desse arcabouço, investigaremos como cada uma das etapas que integram a IME são compreendidas, os desafios encontrados na articulação de cada uma delas e as propostas existentes de como contorná-los.

2.3.1. A natureza do *explanandum*

Esta fase não é tão inocente quanto pode parecer à primeira vista. Implícita na ideia de que “D é um conjunto de dados” está a pressuposição de que eles demandam explicação ou, no mínimo, são passíveis de serem explicados. Afinal, se os fenômenos sobre os quais se debruça o filósofo ou o cientista não são passíveis de explicação, a inferência parece estar fadada ao erro.

Da Ética à Metafísica, inúmeras controvérsias parecem, em maior ou menor medida, pressupor que certos fatos demandam explicação. Debates importantes em Cosmologia e em Filosofia Religião, por exemplo, partem da premissa de que a sintonia existente entre os parâmetros que tornam a vida possível no universo (*fine-tuning*) é um exemplo. Para alguns autores, ele constitui evidência em favor da hipótese de que existem múltiplos universos (Isaacs & Hawthorne, 2022); para outros, evidência em favor de que existe um Deus (Swinburne, 2004, p. 177).

Curiosamente, não obstante a sua relevância, apenas recentemente a Epistemologia e a Filosofia da Ciência parecem ter dado atenção a esse ponto. Baras (2020a, 2020b, 2021, 2022), provavelmente o primeiro a se demorar sobre a questão, tem argumentado que essa propriedade (a que deu nome de *striking principle*) não é intrínseca a quaisquer fatos, mas é, antes, um produto de nossa psicologia, das expectativas que emergem de nossas experiências prévias. “Sem tal conhecimento empírico prévio, eu duvido que teríamos tal intuição”, ele pontua (2020b, p. 1505).

Como Bhogal (2022, p. 24) observou, *strikingness* permanece “um fenômeno insuficientemente explorado”. Muito trabalho há ainda a ser feito. Apesar disso, é interessante considerar quais seriam as implicações de o tempo vir a reivindicar a análise proposta por Baras. Nesse caso, parece que tanto o argumento do teísta quanto o do entusiasta dos muitos mundos perderiam parte de seu apelo. Afinal, a atratividade de ambas as propostas reside no suposto ganho explicativo obtido com os postulados que introduzem. Se, contudo, não há necessariamente algo a ser explicado, a capacidade de prover uma explicação parece conferir pouca credibilidade a elas.

Seja como for, qualquer que seja o veredito do tempo, ele não ameaça a continuidade de quaisquer investigações filosóficas. É questionável que nossos avanços na Física, por exemplo, dependam da convicção de que os fenômenos observados em si mesmos demandam explicação. A mera possibilidade de eles poderem vir a ser explicados é provavelmente combustível suficiente para a maior parte dos pesquisadores.

Por essa razão, talvez a melhor forma de conceber essa primeira etapa, até o momento, seja como o fez Peirce em seu esquema da Abdução. Como mencionado anteriormente (tópico 1.1.), o raciocínio abductivo começa com uma reação de surpresa (“O fato surpreendente C é observado”). E quer certos fatos demandem explicação ou não, certo é que muitos deles nos surpreendem - o que, no mais das vezes, será o bastante para nos motivar a tentar “fixar nossa crença” (Peirce, 1986 [1877])¹⁶.

Em razão disso, é aconselhável que, pelo menos por ora, se trabalhe com uma noção ampla daquilo que conta como *explanandum* de uma IME. Nessa leitura, parafraseando Josephson & Josephson (1994, p. 5), podemos dizer que “D é um conjunto de dados *para o qual gostaríamos de ter uma explicação*”.

2.3.2. As hipóteses elegíveis

A identificação daquilo que torna uma determinada proposição ou conjunto de proposições uma explicação é fundamental não só para o sucesso de qualquer modelo de IME, mas também para a própria pesquisa ora conduzida. Afinal, parece razoável supor que, para que se possa determinar a (ir)relevância evidencial da explicatividade, precisemos ser capazes de articular de forma minimamente precisa em que, em primeiro lugar, consiste *explicar* um fato ou observação.

Há dois principais obstáculos nessa fase. O primeiro concerne o fato amplamente reconhecido na literatura de que o conceito de explicação, assim como outros objeto de escrutínio filosófico, se mostrou deveras resistente à análise. É dizer, apesar dos contínuos refinamentos, não contamos com uma teoria unificada e imune a contraexemplos. E o segundo é que, não bastasse a ausência de uma teoria satisfatória, algumas das mais proeminentes de que dispomos parecem ser incompatíveis com os contornos gerais da IME. Em um trabalho recente sobre o tema, Prasetya (2024, p. 20) conclui:

[Neste artigo] explorei vários modelos de explicação científica. Dentre eles, um dá suporte à IME (a teoria unificacionista), um é compatível com a IME

¹⁶ Isso, a propósito, está em harmonia com a função desempenhada pela Abdução no pensamento mais maduro de Peirce, que, nos escritos de 1900 em diante, passa a concebê-la como sendo a apenas a primeira fase da investigação científica (Fann, 2012).

(modelo mecânico-causal proposto por Salmon), mas não lhe dá suporte, e os outros três são incompatíveis com a IME (o modelo Dedutivo-Nomológico-Probabilístico de Railton, o modelo da relevância estatística de Salmon e o modelo erotético de Van Fraassen).

O explanacionista parece então se ver entre Cila e Caríbdis. Se não endossar nenhuma teoria em particular, acaba com uma noção de *explicar* demasiadamente vaga, quiçá filosoficamente inócua. Se, por outro lado, for incapaz de fazer as pazes com a vagueza, as poucas opções disponíveis parecem forçá-lo a reconhecer que uma grande parcela do uso que fazemos da noção de *explicar* talvez fique de fora do alcance da Inferência à Melhor Explicação.

Talvez existam razões para o explanacionista preferir Caríbdis. Ele pode, por exemplo, (sem deixar de reconhecer que há explicações não causais¹⁷) restringir o âmbito de aplicação da IME a inferências de natureza causal. Dada a ubiquidade das inferências dessa natureza, a perda talvez não seja tão grande. Para Peirce, aliás, a Hipótese (ou Indução Hipotética) era a inferência “dos efeitos para sua [respectiva] *causa*” (grifei) (CP 2.536, 1974).

Porém, essa escolha talvez venha com um preço muito alto. Como veremos em detalhes no capítulo 3, o exemplo fornecido por Sober e Roche (2013) se propõe a demonstrar que, em certos casos, informações referentes à causa subjacente a uma determinada observação em nada alteram a probabilidade de a hipótese em questão ser verdadeira. Logo, se adotada essa concepção mais restrita do que significaria *explicar* algo, caso o argumento de S&R seja cogente, seria o fim do debate. A explicatividade, pelo menos nesses casos, seria demonstravelmente despida de valor epistêmico.

A maior parte dos explanacionistas contemporâneos, no entanto, explícita ou implicitamente, adota a postura sugerida por Lipton (2004). Em resposta a objeções levantadas por Salmon (2001) quanto à utilidade da IME na ausência de uma teoria satisfatória sobre a natureza da explicação, Lipton (2001, p. 100) propõe que o termo *explicação* seja tomado como primitivo:

[se] considerações de cunho explanatório são ou não um guia para a inferência não depende de termos ou não uma teoria adequada da explicação, assim como o nosso uso da gramática para compreender a nossa linguagem não depende da nossa capacidade de fornecer uma explicação explícita adequada da estrutura dessa gramática.

A sugestão não é tão insatisfatória quanto pode parecer. Em primeiro lugar, observe-se que a noção de explicação é um conceito dos mais familiares. Todos partilhamos de fortes intuições sobre a capacidade de uma hipótese de prover uma explicação para determinado

¹⁷ Vide Lange (2016) para uma análise pormenorizada do tema.

conjunto de observações. Ademais, é possível que a noção de *explicação* seja melhor concebida como um conceito prototípico, insuscetível de redução a um conjunto de condições suficientes e necessárias. Segundo Schurz (2014), a Filosofia da Ciência produziu três diferentes paradigmas sobre a natureza da explicação. De acordo com o que denominou de Paradigma da Expectabilidade, “a propriedade fundamental da explicação de um *explanandum* E consiste em tornar E mais esperado” do que inicialmente era (p. 67). No Paradigma da Causalidade, como o nome sugere, o traço distintivo seria a capacidade da explicação de elencar as causas do *explanandum*, sendo irrelevante se estas aumentam ou diminuem sua expectabilidade. E, por fim, no Paradigma da Unificação, “o principal objetivo de uma explicação consiste em fornecer uma compreensão mais profunda dos fenômenos, ao unificá-los”¹⁸ (idem).

Como Prasetya (2023) ressaltou, os dois programas de pesquisa até agora se desenvolveram de forma bastante independente. Em virtude disso, na ausência de uma síntese dos estudos até então desenvolvidos em cada um deles, presumiremos que as duas dificuldades mencionadas não são fatais, tomando o conceito de explicar como primitivo.

2.3.3. A identificação da melhor explicação

A noção de *explicar bem* não tem um sentido unívoco. Presume-se que uma boa explicação seja capaz de satisfazer um agente racional. No entanto, o contentamento ou descontentamento deste, que é uma das formas em que se pode tomar a ideia de uma proposição ser uma boa explicação, não é o que buscamos na Epistemologia e na Filosofia da Ciência. Ao contrário dos Psicólogos e Cientistas Cognitivos, que irão se ocupar precipuamente da descrição daquilo que consideramos uma boa explicação, estamos interessados nas características que *devem* ser assim tidas por agentes racionais.

Essa distinção se revela especialmente relevante quando consideradas as descobertas das últimas duas décadas na área de pesquisa das chamadas preferências explanatórias. Embora muitas das heurísticas que empregamos inconscientemente sejam bastante confiáveis (Gigerenzer e Selten, 2002), sabemos que nossas predileções inconscientes por certos atributos explicativos nem sempre são conducentes à verdade.

Por exemplo, quando confrontados com diferentes explicações para um determinado evento, costumamos preferir a que ostenta o menor número de efeitos não observados – mesmo não havendo indícios de que ela seja a mais provável de ser a correta. Isso é chamado de “viés

¹⁸ A ideia aqui é remanescente da famigerada consiliência de induções defendida por Whewell (1847).

de Escopo Latente” (Khemlani, Sussman, Oppenheimer, 2011; Johnston *et al*, 2017; Tsukamura *et al*, 2022)¹⁹. Outros estudos, agora já replicados em diferentes países (Mills e Frowley, 2014), indicam que temos uma inclinação a reconhecer agência, propósito e intenção por trás dos mais variados eventos e fenômenos no mundo natural. “O sol produz luz para que as plantas possam fazer a fotossíntese”, por exemplo. E para citar mais um exemplo, Lombrozo (2007) conduziu experimentos visando esclarecer como o número de causas postuladas em uma explicação afeta nossa predisposição a aceitá-la. Surpreendentemente, apesar de nossa tendência não ser insensível às informações disponíveis, observou-se que, “mesmo quando a explicação mais simples era 10 vezes menos provável que a alternativa que postulava duas causas, 41% dos participantes continuaram a preferi-la” (p. 246).

O objetivo filosófico, então, é identificar quais os atributos de uma explicação que, quer os consideremos satisfatórios ou não, são indicativos de que ela constitui uma representação verdadeira ou verossimilhante do *explanandum* sob consideração. Esta é, aliás, a etapa de que historicamente mais se ocuparam os estudiosos da IME.

Pode-se dizer que aqui são três os desafios: a delimitação do rol de explicações relevantes; a identificação das virtudes explicativas de importe epistêmico e a elucidação do porquê de elas serem dotadas de tal propriedade.

2.3.3.1. A delimitação do rol de explicações relevantes

A dificuldade de se delimitar o rol de explicações relevantes se encontra bem capturado na famosa objeção do Lote Ruim, formulada por Van Fraassen (1989). Em *Laws and Symmetry*, o autor argumenta (p. 142-143):

[A IME] seleciona apenas a melhor entre as hipóteses historicamente fornecidas. Não podemos observar uma competição entre as teorias que lutamos tão arduamente para formular e aquelas que ninguém propôs. Portanto, nossa seleção pode muito bem ser a melhor de um lote ruim. Acreditar é, pelo menos, considerar mais provável ser verdadeiro do que falso. Logo, acreditar na melhor explicação requer mais do que uma avaliação da hipótese dada. Requer um passo além do julgamento comparativo de que ela é melhor do que suas concorrentes atuais. Enquanto o julgamento comparativo é de fato um “pesar (à luz da) evidência”, o passo adicional - chamemos de passo ampliativo - não é. Acreditar que a melhor do conjunto X será a mais

¹⁹ Exemplo: Alfredo chega ao hospital queixando-se de forte dor na parte lateral da cabeça. Existem dois possíveis diagnósticos: doença X e doença Y. A primeira tem como sintomas apenas dor forte na parte lateral da cabeça. A segunda, além de dor no mesmo local, tem como sintoma a elevação da glicose no sangue. Os índices de glicose no sangue de Alfredo não foram testados. Nesse cenário, se ambos os diagnósticos antecipam o mesmo tipo de dor de cabeça, e ambas as doenças ocorrem com a mesma frequência, não há nenhuma a ser preferida. Ambos os diagnósticos são igualmente prováveis. Pelo menos até a checagem da glicose. Apesar disso, as pesquisas indicam que nós, como regra, julgamos a doença X como o diagnóstico mais provável.

provável de ser verdadeira requer uma crença prévia de que a verdade já é mais provável de ser encontrada em X do que não.

As explicações que se tem em vista na passagem são teorias científicas em sentido estrito. Todavia, embora a força do argumento fique mais visível nesse contexto, ele é igualmente aplicável à avaliação de explicações do dia a dia. Se você reside com apenas duas outras pessoas, não costumam receber visitas, e você é surpreendido com a observação de que sua garrafa de Coca Cola está pela metade, é razoável afirmar que a melhor explicação é que um dos dois outros moradores bebeu seu refrigerante. E para todos os nossos propósitos práticos, não há qualquer problema em inferir que isso é verdade. Pode, inclusive, ser um imperativo da racionalidade fazê-lo. O problema, no entanto, reside na justificação do salto de “*não tenho ciência* de nenhuma explicação melhor” para “*não há* nenhuma explicação melhor”.

A literatura conta com diversas propostas de resposta à objeção do Lote Ruim. Dentre as mais recentes, destacam-se Schupbach (2014), para quem “ela não é uma objeção à IME, mas sim ao conteúdo material envolvido em casos específicos de IME” (p. 63); Davey (2023), para quem não haveria lotes ruins, apenas formulações ruins de IME, aqui entendidas como modelos que autorizam a inferência à míngua de informações suficientes sobre o *set* de hipóteses candidatas; e Beni (2023), segundo o qual a capacidade de realizar inferências explanatórias maximiza as chances de sobrevivência de um organismo e, por isso, “vem com privilégios evolutivos” (p. 188).

O que frequentemente passa despercebido é que, conquanto Van Fraassen a tenha intentado como um argumento contra a IME, a objeção do Lote Ruim também ameaça outras, senão todas, formas de inferência ampliativa. Tomemos um exemplo clássico de indução eliminativa, como o da descoberta da febre puerperal, no século XIX.

Isolando diversos fatores que poderiam ser os responsáveis pela notável diferença na predominância da doença nas duas salas obstétricas do Hospital Geral de Viena, Semmelweis (1861) veio a concluir que a diferença se devia ao fato de que, na sala com maior incidência da doença, os médicos frequentemente vinham de autópsias e, por isso, provavelmente traziam consigo “partículas cadavéricas” (1861, p. 59; Shorter, 1984, p. 54), que contaminavam as parturientes. Mesmo nessa hipótese, que não envolve IME²⁰, o êxito da inferência depende de

²⁰ O caso da febre puerperal foi oferecido por Lipton (1990) como um exemplo paradigmático de IME na história da ciência. No entanto, do que se pode depreender dos relatórios de Semmelweis, não é este o caso. Primeiramente, as possíveis explicações não foram postas lado a lado e então avaliadas com base em suas qualidades explicativas. Em vez disso, cada uma delas foi testada individualmente após ser conjecturada. O caso mais se adequa ao Método da Diferença proposto por Mill (1974 [1843]).

a hipótese verdadeira ter sido conjecturada. Em casos de indução eliminativa, isso pode não ser imediatamente aparente, pois, como regra, a identificação da verdadeira causa é sinalizada pela cessação do efeito. Contudo, tal observação é logicamente compatível com uma infinidade de outras explicações, igualmente hábeis de, em tese, produzir o mesmo efeito.

Mesmo o *framework* Bayesiano – endossado por Van Fraassen – é vulnerável a essa objeção. Como Hawthorne (1993) demonstrou de modo bastante persuasivo, “a inferência indutiva Bayesiana é essencialmente uma forma probabilística de indução por eliminação” (p. 99). Uma vez que as probabilidades antecedentes e condicionais pertinentes estejam minimamente especificadas, o aparato Bayesiano permite quantificar o impacto de cada observação na distribuição do conjunto de hipóteses sob análise. Todavia, nada nele nos informa quais hipóteses devem ou não ser consideradas. Para sermos justos, apesar de ser um problema para o Bayesiano em geral, a objeção do Lote Ruim não se volta contra Van Fraassen em particular, pelo menos não na discussão de teorias científicas. Isso devido à sua compreensão de que “a aceitação de uma teoria [científica] envolve apenas a crença de que ela é empiricamente adequada” (Van Fraassen, 1980, p. 12), sem qualquer comprometimento de ela ser verdadeira ou mais verossimilhante que suas rivais.

Por esses motivos, desafiadora que seja a objeção do Lote Ruim, ela é melhor compreendida como sendo, não um problema para a IME em particular, mas como um corolário da condição epistêmica humana²¹.

Dito isso, o cético talvez não disponha de tanta munção quanto supõe. Como é sabido, a relação de suporte evidencial não se dá em um vácuo. Primeiramente, observe que o esquema de IME desenvolvido até aqui é silente sobre a duração do processo avaliativo. O explanacionista não está obrigado a se comprometer com uma explicação *in actu oculi*. Esse aspecto foi recentemente bem explorado por Dellsén (2021), que introduziu a noção de consolidação explanatória da seguinte forma (p. 170):

Uma hipótese torna-se 'boa o suficiente' quando passa por uma consolidação explanatória, um processo em que evidências empíricas diretas a favor da hipótese, e repetidos fracassos em apresentar alternativas plausíveis, gradualmente se acumulam para apoiar a explicação otimista de que a razão

Para exames pormenorizados deste ponto, vide Schurz (2014, p. 185-187), Norton (2021, p. 262-267) e Pfister (2022).

²¹ O problema decorre do fato de que, quando inferimos uma causa a partir de um efeito (a forma de raciocínio que talvez melhor descreva os contornos gerais das ciências e da Filosofia), estamos, a rigor, cometendo a falácia da afirmação do consequente. A IME e o Bayesianismo, a propósito, podem ser concebidos como formas de lidar com essa verdade inconveniente.

pela qual não encontramos uma hipótese explicativa melhor é que não há tal hipótese a ser encontrada.

Além disso, como há muito sabemos (Good, 1967), a relação de suporte evidencial é uma “relação de três lugares” (Hájek & Joyce, 2013, p. 158). Ela se estabelece não apenas hipótese e evidência, mas hipótese, evidência e conhecimento de fundo.

Relembre o exemplo do misterioso desaparecimento da Coca Cola de sua geladeira. É possível, em princípio, que um terceiro indivíduo, além das duas pessoas que vivem com você, tenha entrado em seu apartamento, bebido seu refrigerante e ido embora deixando tudo o mais inalterado. Entretanto, à luz de todo o conhecimento de fundo de que você dispõe (de como o Universo se comportou até agora aos gostos das pessoas que vivem com você e à forma como assaltantes de residência costumam se comportar), não parece temerário eliminar a hipótese cética.

Importante salientar, a relevância desse ponto vai muito além de inferências mais cotidianas como a do exemplo acima. Sober (1991, p. 65-66) o coloca da seguinte forma:

O que encontramos em qualquer argumento indutivo articulado é um conjunto de suposições empíricas que permitem que as observações tenham impacto sobre hipóteses concorrentes. Essas suposições de fundo podem elas mesmas ser escrutinadas, e mais observações e teorias de fundo podem vir a ser oferecidas em seu apoio. Quando perguntados por que consideramos que observações passadas sustentam a crença de que o sol nascerá amanhã, respondemos citando nossa bem confirmada teoria sobre movimento planetário, e não o Princípio da Uniformidade da Natureza de Hume. Se desafiados a explicar por que levamos essa teoria científica a sério, responderíamos citando outras observações e outras teorias de fundo também.

Assim, muito embora sejamos incapazes de excluir todas as explicações logicamente consistentes com a evidência de que dispomos, à luz do conhecimento de fundo de que dispomos, parece que podemos, de fato, traçar uma “Inferência à Única Explicação” (Bird, 2005, 2010, 2022; Norton, 2021; Woodward, forthcoming).

2.3.3.2. A seleção e justificação das virtudes explicativas

É um eufemismo dizer que não há consenso quanto aos atributos explicativos a serem considerados na identificação da melhor explicação. Na verdade, é difícil encontrar dois autores que tenham elencado as mesmas virtudes explicativas.

Para citar alguns exemplos, Quine & Ullian (1978, p. 68-71) listaram conservadorismo, modéstia, simplicidade, generalidade e refutabilidade; Thagard (1978), consiliência, simplicidade e analogia; Psillos (2002), por seu turno, elegeu consiliência, completude, importância, parcimônia, unificação e precisão; Lipton (2004), mecanismo, precisão, escopo,

simplicidade, fertilidade e compatibilidade com conhecimento de fundo. Na mais recente tentativa de sistematização, Keas (2017) identificou não menos do que doze virtudes explicativas.

Esse cenário quase anárquico naturalmente inspira o temor de que diferentes modelos de IME sejam capazes de endossar diferentes conclusões como verdadeiras ou mais prováveis – o que, na ausência de um critério para preferir um modelo a outro (inferiríamos, via IME, qual seria o melhor modelo de IME?) – parece tornar qualquer opção arbitrária ou cometer a falácia da petição de princípio.

Uma proposta interessante e, em princípio, capaz de contornar essa dificuldade consiste em questionar a suposição amplamente aceita de que a melhor explicação é, necessariamente, aquela que melhor pontua de acordo com esses critérios. E a ideia de que se trata de coisas distintas possui um precedente proeminente na história da disciplina. No seminal *Inference to the Best Explanation*, Lipton (2004, p. 59) defende que “a melhor explicação [é] aquela que, se correta, seria a mais explicativa ou proporcionaria maior entendimento”.

Que essas duas concepções de qualidade explicativa não apontam necessariamente na mesma direção, foi bem ilustrado por Elliott (2021). A autora nos convida a considerar duas potenciais explicações para o fato (F) de que podemos observar a luz de estrelas tão distantes de nós. A primeira (H1) é a de que o universo tem cerca de 13,8 bilhões de anos e a velocidade da luz é constante no vácuo; e a segunda (H2), a hipótese de que o universo tem entre 6.000 e 10.000 anos, e a velocidade da luz vem diminuindo desde a criação do universo. Embora a primeira pontue melhor em termos de virtudes explicativas, “se eu soubesse que H2 é verdadeira, entenderia (F) tão bem quanto entenderia se eu soubesse, em vez disso, que H1 é verdadeira” (p. 177).

Esse exemplo pode sugerir que a concepção liptoniana de qualidade explicativa não é muito promissora. Mas talvez essa seja uma conclusão precipitada. Em *Inference to the Best Explanation, Cleaned Up and Made Respectable*, Schupbach (2017) avançou uma proposta bastante sofisticada de como articular a noção de melhor explicação. Lançando mão do *insight* de Peirce de que somos incitados a conjecturar uma explicação a partir de uma reação de surpresa, o autor propõe que a noção de poder explicativo seja concebida nos seguintes termos: “uma hipótese tem poder²² sobre uma proposição na medida em que torna essa proposição menos surpreendente - ou mais esperada - do que seria de outra forma” (p. 42). Schupbach

²² No original, “has power”. No sentido de “explanatory power”.

demonstra como ela satisfaz diversos desideratos desejáveis (p. 41-43) e, inclusive, testa sua performance em uma simulação computadorizada (p. 49-55).

Apesar de ser uma proposta interessante de como remediar o problema de identificação das virtudes explicativas, ela vem com alguns efeitos colaterais. O primeiro se relaciona à introdução da surpresa, um aspecto psicológico, como mediador da relação de suporte evidencial. Como destacado no tópico 2.3.3., a empreitada filosófica consiste no esclarecimento das características explicativas que *devem* ser tidas como epistemicamente desejáveis por agentes racionais, e não meramente aquelas que nós, muitas vezes equivocadamente, tendemos a identificar como tais. E, *prima facie*, parece plausível que diferentes pessoas possam esboçar surpresa em diferentes medidas, para não dizer em reação a fatos completamente distintos, sem a violação de qualquer dever epistêmico.

O segundo é que o quão esperada uma observação é sob determinada hipótese corresponde precisamente à chamada verossimilhança (likelihood), um dos termos que perfazem o Teorema de Bayes:

$$P(H|E^{23}) = \frac{P(H) \cdot P(E|H)}{P(E)}$$

Assim, o receio é que, apesar de coerente e precisa, a definição de poder explicativo oferecida (e o respectivo modelo de IME) seja parasitária no Bayesianismo. Em outras palavras, nessa perspectiva, a Inferência à Melhor Explicação talvez seja apenas um Bayesianismo sem *priors*²⁴.

O terceiro problema, diretamente relacionado a este último, se refere à própria idoneidade do conceito de poder explicativo proposto. Como as controvérsias na Estatística em torno dos chamados testes de significância testificam (Mayo, 2018), o fato de uma hipótese conferir uma probabilidade extremamente baixa a determinada observação, por si só, nada nos informa sobre a primeira ser verdadeira ou falsa. É possível que, sob as hipóteses rivais, a observação seja ainda mais improvável²⁵.

Mais uma vez, essa dificuldade também não precisa ser fatal para o Explanacionismo. Aqueles dentre nós de inclinação mais empirista podem, por exemplo, propor que as virtudes explicativas hábeis a integrar o rol são somente aquelas que contam com suporte empírico. É

²³ Simplificadamente, a $P(A|B)$ é a probabilidade A supondo que B é verdadeiro.

²⁴ Ou mais tecnicamente, seja equivalente ao chamado Verossimilhancismo.

²⁵ Para um excelente resumo dos problemas com o Frequentismo, ver (Titelbaum, 2022, p. 461-466)

possível que, ao se comparar as teorias que vieram a ser reivindicadas por experimentos com suas rivais hoje descartadas, venhamos a concluir que as primeiras ostentavam certas virtudes explicativas que as concorrentes não exibiam ou o faziam em grau notadamente menor. Contudo, como Cabrera (2023) comentou, “ainda não foi realizada uma investigação detalhada desse tipo” (p. 1886).

O desafio de prover uma justificativa para nossa preferência por certos atributos explanatórios pode ser separado da empreitada de identificá-los apenas para fins didáticos ou expositivos. Afinal, pode-se argumentar que é da dificuldade em justificá-los que nasce a dificuldade de identificar quais deles devem integrar o rol.

Para análise do argumento *Screen-off*, não é necessário considerar as tentativas de justificação de cada uma delas, e as limitações de espaço também sequer nos permitiriam fazê-lo. A título de ilustração, entretanto, consideremos o caso da simplicidade, que, como se verá, apesar de presente em praticamente todas as listas de virtudes explicativas, não é um atributo de fácil justificação.

As diferentes tentativas de justificar nossa preferência por explicações mais simples podem ser convenientemente divididas em três categorias: justificativas *a priori*, justificativas naturalistas e justificativas probabilísticas / estatísticas (Baker, 2022).

Para Descartes, por exemplo, todo movimento era retilíneo pela “imutabilidade e simplicidade da operação pela qual Deus preserva movimento na matéria” (1984 [1644], p. 39). Leibniz (1991 [1686], p. 40), no mesmo sentido, acreditava que “Deus escolheu o mundo mais perfeito, ou seja, aquele que é ao mesmo tempo o mais simples em hipóteses e o mais rico em fenômenos”. Mesmo Sir Isaac Newton veio a partilhar da mesma opinião. Embora em seus comentários metodológicos em *Mathematical Principles of Natural Philosophy* não tenha se proposto a justificá-la, em um comentário não publicado ao livro bíblico de Apocalipse, Newton fornece uma nova lista regras, a nona das quais prescreve:

Regra 9. Optar por aquelas construções que, sem esforço, reduzem as coisas à maior simplicidade. A razão para isso é evidente pela Regra precedente. A verdade sempre se encontra na simplicidade, e não na multiplicidade e confusão das coisas. Assim como o mundo, que a olho nu exhibe a maior variedade de objetos, parece muito simples em sua constituição interna quando observado por um entendimento filosófico, e tanto mais simples quanto melhor é compreendido, assim é com essas visões. É a perfeição das obras de Deus que todas são realizadas com a maior simplicidade. Ele é o Deus da ordem e não da confusão. E, portanto, assim como aqueles que desejam entender a estrutura do mundo devem se esforçar para reduzir seu conhecimento a toda simplicidade possível, assim deve ser ao buscar entender essas visões.

Desnecessário dizer, é uma explicação que atualmente satisfaz a poucos, talvez incluindo até mesmo teístas. E isso porque, mesmo pressupondo que Deus exista (o que é uma pressuposição e tanto), sabemos que os sistemas por ele colocados em movimento na natureza nem sempre trilham o caminho mais simples, como as famosas gambiarras evolutivas não nos permitem esquecer (Morris & Lundberg, 2011).

A segunda proposta *a priori* conta com proponentes ainda hoje. Segundo Swinburne (1997, p. 1), “é um princípio epistêmico *a priori* último que a simplicidade é evidência de verdade”. Sendo a parcimônia um critério pelo qual distinguimos o que é verdadeiro, não seria ela própria passível de justificação por critérios semelhantes. A preocupação aqui é óbvia. Como vimos anteriormente, há quem tome a noção de explicação como primitiva. Se a relevância dos atributos que tornam uma explicação melhor ou pior também for tomada como um fato bruto, há o claro receio de se acabar com nada além de um salto imaginativo insuscetível de desconfirmação - e, por extensão, também insuscetível de qualquer corroboração. É uma compreensão que, ironicamente, relembra as tentativas de salvar o modelo Ptolomaico adicionando mais e mais epiciclos.

Quine, o naturalista *par excellence*, também tentou prover uma justificação, em consonância com sua Filosofia, para o que via como uma das duas normas mais gerais da “arte de conjecturar” (1995, p. 49). Em *The Web of Belief*, ele avança a seguinte proposta (1978, p. 73):

A teoria da seleção natural de Darwin oferece uma conexão causal entre a simplicidade subjetiva e a verdade objetiva da seguinte maneira. Padrões subjetivos inatos de simplicidade que fazem as pessoas preferirem algumas hipóteses a outras terão valor de sobrevivência na medida em que favoreçam a previsão bem-sucedida. Aqueles que melhor predizem têm mais chances de sobreviver e reproduzir, pelo menos em um estado natural, e assim seus padrões inatos de simplicidade são transmitidos.

Apesar de mais consoante com nossas inclinações científicas modernas, pode-se temer que apenas trocamos uma “just so story” por outra. Conquanto Quine na passagem evidentemente não pretenda fornecer uma teoria detalhada desse processo, ela, não obstante, parece pressupor uma compreensão muito superficial de como o processo evolutivo se dá. Para citar apenas um dos problemas com essa sugestão, nem todos os traços e habilidades exibidos por organismos hoje foram selecionados (Gould & Lewontin, 1979) e, por isso, na ausência de um argumento estabelecendo essa premissa, ela passa longe de prover uma justificativa adequada.

Na segunda metade do Séc. XX, passou-se a buscá-la na Probabilidade e na Estatística. Jeffreys foi o primeiro a sugerir que “leis mais simples possuem maior probabilidade

antecedente” (1961, p. 47). A ideia, entretanto, revelou ter uma meia-vida bastante curta. Tomando o exemplo das equações lineares e quadráticas fornecido por Jeffreys, Popper (2005 [1959]) apontou que as equações lineares ($y = a + bx$) nada mais são do que equações quadráticas ($y = a + bx + cx^2$) em que a variável c é zero. E como uma proposição que implica outra será logicamente mais forte (e, portanto, menos provável), a proposta de Jeffreys vai de encontro aos axiomas da probabilidade.

As ferramentas estatísticas que se mostraram mais filosoficamente frutíferas vêm da área de seleção de modelos. Sabe-se que, conquanto um modelo mais complexo (com um maior número de parâmetros ajustáveis) tenha mais facilidade de prever o padrão em determinado conjunto de dados, a complexidade está inversamente relacionada com a acurácia. Modelos complexos encaixam-se não só ao padrão estatístico sendo investigado, mas também ao ruído existente nos dados. Na Estatística, para lidar com esse fenômeno, utilizam-se os chamados critérios de seleção de modelos, dos quais os mais populares são o Critério de Informação Bayesiana e o Critério de Informação Akaike. Desde a década de 90 (Forster & Sober, 1994), eles vêm sendo objeto de escrutínio filosófico e aplicados a uma variedade de contextos e questões filosóficas - do problema mente-corpo (Sober, 1996, 2009, 2022) ao ceticismo cartesiano (Shogenji, 2017). Contudo, não se trata de uma panaceia. Akaike, por exemplo, não é invariável à linguagem²⁶, de modo que o número de parâmetros de um modelo (e, consequentemente, o quão complexo ele é) depende da forma como ele é descrito (DeVito, 1997; Forster, 1999).

Este breve panorama obviamente não esgota todas as propostas existentes, e não se pode dizer que o projeto de tentar justificar nossa preferência por explicações mais simples foi abandonado pela comunidade (Lange, 2024). De todo modo, os exemplos supracitados ajudam a ilustrar o quão desafiadora se revelou a tarefa de justificar nossa preferência por certos atributos explicativos.

2.3. 4. O caráter da conclusão

A IME provavelmente não seria alvo de tantas objeções, senão fosse por sua grande ambição. Se a conclusão a ser inferida fosse apenas a de que a melhor explicação deve ser, por exemplo, a primeira a ser investigada, pouca resistência teria encontrado. Muito embora, é verdade, neste caso provavelmente boa parte de seu apelo também viria a se perder.

²⁶ No original, *language invariant*.

Um dos ambiciosos projetos para os quais ela foi recrutada é nada menos do que o de responder ao cético. Vogel (1990; 2008; 2013), por exemplo, avançou uma resposta ao ceticismo baseada nas “vantagens explicativas da visão do senso comum [de que não estamos em um típico cenário cético]” (1990, p. 658). Uma análise pormenorizada da cogência de seu argumento transcende tanto o escopo da presente dissertação como minha competência atual. Por ora, basta destacar que o oferecimento de uma resposta explanacionista ao desafio cético está longe de ser uma idiossincrasia sua. A proposta permanece sendo uma das, senão a mais, proeminente das respostas disponíveis (Beebe, 2009; Huemer, 2016; McCain, 2016; 2019; Warren, forthcoming).

Pouco tempo depois de Harman introduzi-la à literatura, a IME também foi recrutada para prover o “argumento último” (Musgrave, 1988, p. 229) em favor do Realismo Científico²⁷. De acordo com o famigerado Argumento do Milagre (Psillos, 2005, p. 76):

A melhor explicação para a confiabilidade instrumental da metodologia científica é que as declarações teoréticas que afirmam mecanismos ou conexões causais específicas por meio dos quais métodos científicos obtêm predições corretas são aproximadamente verdadeiras.

À luz da discussão dos tópicos anteriores acerca das dificuldades em se estabelecer a própria *relevância* epistêmica de aspectos explanatórios, talvez o advérbio *aproximadamente* possa parecer de pouca serventia para aplacar as inquietações explanacionistas. Entretanto, como Kuipers (2000) bem observou, o problema da formulação original da IME é que “ela conecta uma conclusão não comparativa (ser verdadeira) a uma premissa comparativa (ser a melhor teoria não refutada” (p. 171).

Por essa razão, alguns autores, mais proeminentemente Niiniluoto (1977; 1986; 1997; 2012; 2018; 2020), têm explorado a noção de verossimilitude ou proximidade da verdade. É plausível que a verossimilitude, diferentemente da veracidade, admita graus e, assim, seja uma noção inerentemente comparativa, assim como a Inferência à Melhor Explicação. É verdade que, para fazer sentido do conceito de verossimilitude para além do debate sobre realismo científico, há evidentemente bastante trabalho ainda a ser feito. Mas é sem dúvida um dos caminhos mais promissores.

Aliás, é possível que parte das dificuldades encontradas até aqui não decorra de um problema fundamental com a IME, mas sejam produto de uma expectativa irrazoável. Talvez a hercúlea tarefa de nos indicar a verdade – de que a IME foi historicamente incumbida – é que

²⁷ Como é sabido, o Realismo Científico é uma posição filosófica repleta de nuances, com aspectos ontológicos, semânticos, epistemológicos, teoréticos e metodológicos (Niiniluoto, 1999).

seja o principal entrave ao programa. Dada a condição epistêmica humana, talvez jamais sejamos capazes de dispensar a cláusula “as far as we know” de algumas de nossas inferências ampliativas.

3. O DESAFIO *SCREEN-OFF*

À semelhança do segundo capítulo, somos mais uma vez confrontados com uma encruzilhada metodológica. Assim como o panorama sobre a Inferência à Melhor Explicação, o debate em torno do desafio *Screen-off* pode, em princípio, ser apresentado de forma cronológica, biográfica ou temática.

Todavia, algumas peculiaridades da discussão sugerem a superioridade de uma abordagem mista. Embora a forma mais natural de se familiarizar com a respectiva literatura seja lendo os artigos em ordem cronológica, ela não se revela igualmente produtiva quando o objetivo é uma análise dos argumentos apresentados. Isso porque, além de Sober e Roche (2019) terem reformulado a tese *Screen-off*, alguns dos contra-argumentos reaparecem em respostas posteriores. A ideia de que a explicatividade é indispensável para a chamada projeção indutiva, por exemplo, aparece de forma bastante tímida em McCain e Poston (2014) e, mais tarde, Climenhaga (2017) a torna o objeto principal de sua contribuição. Do mesmo modo, o Problema da Onisciência Lógica figura tanto nos textos de McCain e Poston (2014; 2017) quanto em uma das objeções de Lange (2017).

À vista dessas considerações, empregaremos uma abordagem predominantemente cronológica, porém com alguns agrupamentos temáticos. No tópico 3.1., prover-se-á um breve panorama das diferentes formas em que a relação entre a IME e o Bayesianismo foi concebida na literatura até o momento. No tópico 3.2., apresentaremos uma síntese do argumento *Screen-off* (S&R, 2013), com especial atenção às retificações e esclarecimentos feitos por Sober & Roche (2019). No tópico 3.3., passaremos à análise das réplicas explanacionistas. Em particular, consideraremos a proposta de que a explicatividade torna a probabilidade de uma dada hipótese mais resiliente (3.3.1); a objeção de que a falha reside, em verdade, no aparato Bayesiano (3.3.2); a tese de que a explicatividade é indispensável para a projeção indutiva de certos predicados (3.3.3); e, por fim, a proposta de que uma determinada hipótese pode obter “alguma credibilidade em virtude de sua capacidade de fornecer (se for verdadeira) a [um *explanandum*] E uma explicação do mesmo tipo que descobrimos que certos outros fatos possuem” (Lange, 2022, p. 86) (3.3.4).

3.1. BAYESIANISMO VIS A VIS INFERÊNCIA À MELHOR EXPLICAÇÃO

Sober e Roche (2013) não foram os primeiros a ponderar sobre como a Inferência à Melhor Explicação se relaciona com o Bayesianismo. Por isso, para que a discussão a seguir esteja devidamente contextualizada, convém explorar, ainda que brevemente, as diferentes formas em que essa interação foi, até o momento, concebida na literatura.

O relacionamento dos dois programas certamente não pode ser caracterizado como tendo sido de amor à primeira vista. Além da objeção do Lote Ruim, Van Fraassen (1989) pôs ainda outra pedra no sapato explanacionista. No capítulo 7 de *Laws and Symmetry*, o famoso empirista introduz um personagem, “o Bayesiano Peter, que se converte ao uso de algum tipo de Inferência à Melhor Explicação probabilística” (p. 161). O trecho é tanto esclarecedor quanto bem-humorado, merecendo sua reprodução *ipsis literis*.

Após descrever o processo de condicionalização, ele narra (p. 166):

Imagine agora que ele encontra um Pregador itinerante, que o convence de que as hipóteses de viés [do dado discutido no exemplo] são como leis [...] e não apenas se mostram “certas” [...]; mas fornecem também (depois do fato) uma boa explicação para a observação. Nosso [personagem] Pedro sente uma concordância intuitiva brotando em seu peito. Mas então o Pregador continua dizendo: em vista desse sucesso explicativo, você deve aumentar sua crença nas hipóteses mais explicativas.

Segundo Van Fraassen, ao aceitar o conselho do pregador, Peter logo descobrirá que isso o conduz a avaliações incoerentes.

A promessa do Bayesianismo é de que um agente que atualize seus graus de crença (à medida que as evidências se apresentam a ele) de acordo com suas diretrizes jamais exibirá graus de convicção incoerentes. Tal objetivo seria alcançado seguindo-se a chamada Regra da Condicionização, de acordo com a qual a confiança de um agente na veracidade de uma determinada hipótese (H), uma vez descoberta determinada evidência (E), deve ser equivalente à confiança que ele previamente mantinha em H condicionada a E.

A título de exemplo, um torcedor cujo grau de confiança na vitória de seu time é 90%, mas que, caso descobrisse que o principal atacante não irá jogar, restaria reduzido para 60%, ao ser informado de que o referido jogador ficará no banco de reservas, deve reduzir seu grau de convicção em 30%.

De acordo com Van Fraassen, ao autorizar um incremento no grau de convicção do agente com base na qualidade explicativa da hipótese, para além do que a Regra da Condicionização (RC) prevê, a Inferência à Melhor Explicação expõe o agente a uma aposta holandesa - uma série de apostas em que há garantia de perda. Simplificadamente, isso ocorre porque a observância da RC seria a única forma racional de os graus de crença somarem 1 (v.g.,

80% em p e 20% em não- p). Como Dellsén (2024) ressalta, a maior parte dos explanacionistas concorda com a objeção e procurar encontrar espaço para explicatividade de outras formas. Mas há uma exceção notável.

Douven (2013, 2017, 2021; 2022) já há algum tempo tem argumentado que conceder um “bônus probabilístico” à hipótese mais explicativa pode ser não só coerente como, inclusive, mais eficiente na busca da verdade. De forma resumida, isso seria possível – isto é, os graus de crença do agente poderiam somar 1 – ao também penalizar as hipóteses concorrentes menos explicativas.

Imagine, por exemplo, que há quatro hipóteses (H_1 , H_2 , H_3 e H_4) e que o grau de confiança do sujeito na primeira delas condicional em uma determinada evidência ($H_1 \mid E$) é 0.3, estando, por consequência, os 0.7 restantes distribuídos entre as três remanescentes ($H_2 \vee H_3 \vee H_4 \mid E = 0.7$). Conferir um *boost* à H_1 pelo fato de ela ser a melhor explicação, elevando o grau de confiança em sua veracidade para, digamos, 0.4, infringiria de fato os axiomas da probabilidade. No entanto, se esse crédito de 0.1 for descontado das explicações rivais, a coerência é preservada²⁸. Sua superioridade, em termos de eficiência, à Regra da Condicionização estaria evidenciada no fato de que, em simulações computadorizadas, agentes atualizando suas crenças da maneira indicada tendem a convergir mais rapidamente em relação à hipótese correta (Douven, 2022). Há, no entanto, dúvidas quanto à generalização de tais resultados, que envolvem a identificação do viés de uma moeda, para contextos epistêmicos mais complexos (Dellsén, 2024, p. 55).

Como mencionado, a maior parte das respostas à segunda objeção de Van Fraassen, no entanto, busca ser compatível com a Regra da Condicionização. As diferentes propostas podem ser divididas em três grupos: Compatibilismo de Constrição²⁹ (Dellsén, 2024), Compatibilismo Emergente (Henderson, 2014) e Compatibilismo Heurístico (McGrew, 2003).

De acordo com o chamado Compatibilismo de Constrição, considerações de natureza explicativa auxiliariam na fixação das probabilidades antecedentes (priors) das hipóteses e possivelmente no estabelecimento da variável da verossimilhança (likelihood). Nas palavras de Weisberg (2009), “idealmente, considerações explanacionistas complementariam os princípios objetivistas existentes, como o Princípio da Indiferença, resultando em uma forma de Bayesianismo mais aplicável e razoável” (p. 15). Esta é também a opinião de Huemer (2009, p. 354), para quem a “prioridade explicativa [de uma hipótese] pode vir a figurar em uma solução

$$\Pr'(H_j) = \frac{\Pr(H_j) \Pr(E \mid H_j) + f(H_j, E, \mathcal{H})}{\sum_{k=1}^n (\Pr(H_k) \Pr(E \mid H_k) + f(H_k, E, \mathcal{H}))},$$

²⁹ Estou seguindo aqui as nomenclaturas sugeridas por Dellsén (2024).

parcial do problema de como interpretar o Princípio da Indiferença”. A sugestão é atraente. Afinal, tudo o que o Teorema de Bayes nos diz é que, para uma dada tríade $P(H)$, $P(E|H)$ e $P(E)$, $P(H|E)$ deve ser equivalente ao produto das duas primeiras, dividido pela última. Nada nos informa, contudo, sobre o quão confiante um dado agente deve estar em cada uma delas.

Todavia, nossa discussão anterior sobre a dificuldade em se justificar a relevância epistêmica das virtudes explicativas subtrai um pouco do entusiasmo que, a princípio, essa proposta inspira. Se, por um lado, fixar as *priors* atribuindo maior probabilidade à hipótese tida por mais explicativa é compatível com o Bayesianismo, na ausência de uma justificativa para os critérios utilizados, há o receio de que fazê-lo seja tão vulnerável à crítica de subjetivismo quanto é a cepa *de finettiana*³⁰ de Bayesianismo.

O Compatibilismo Emergente, termo cunhado por Henderson (2014), consiste na ideia de que “um Bayesiano que adote restrições em suas [atribuições de] probabilidades, que sejam razoáveis em seus próprios termos, acabaria por favorecer teorias mais explanatórias”³¹ (p. 11). Esse é o sentido em que a IME emergiria do *framework* Bayesiano. A intenção é demonstrar que, “dadas certas *priors* ‘naturais’, o Bayesiano [já] leva em consideração aspectos explanatórios, mesmo que esses aspectos explanatórios não sejam responsáveis pela fixação das *priors*” (p. 2). No momento, é difícil dizer se esta constitui uma proposta de compatibilização superior à discutida nos parágrafos anteriores. Diante do insucesso de, até o momento, se prover uma solução ao Problema das *Priors* (Lin, 2022), pode-se argumentar que, neste caso, se estaria talvez trocando seis por meia dúzia.

A terceira espécie de compatibilismo, embora costume ser distinguida da anterior (Dellsén, 2024), poderia muito bem ser com ela agrupada. A concepção heurística da relação entre IME e Bayesianismo postula que a primeira, em verdade, é uma aproximação do raciocínio probabilístico ideal, adequada para agentes epistêmicos ordinários e falíveis como nós.

Okasha (2000) exemplifica bem a perspectiva. No cenário por ele delineado, uma médica se depara com uma criança de 5 anos se queixando de dores em uma das pernas e, a partir das informações fornecidas pela mãe, formula duas hipóteses: (H_1) a criança distendeu um músculo e (H_2) a criança rompeu um ligamento. “A médica examina a criança cuidadosamente

³⁰ O estatístico italiano Bruno De Finetti, tido como um dos pais do Bayesianismo moderno, entendia que a coerência era a única demanda de racionalidade a que os graus de crença de um agente estão sujeitos.

³¹ No original: “a Bayesian who adopts constraints on her probabilities which are reasonable on her own terms, would end up favouring more explanatory theories”.

e decide que H_2 oferece a melhor explicação para os sintomas observados; ela, portanto, aceita provisoriamente H_2 — embora nela não acredite completamente — e rejeita H_1 ” (p. 702).

Okasha, então, argumenta (p. 703):

Suponha que peçamos à médica que justifique seu raciocínio. Ela responde: “primeiramente, crianças pré-adolescentes raramente distendem músculos, mas frequentemente rompem ligamentos. Em segundo lugar, os sintomas, embora compatíveis com ambos os diagnósticos, são exatamente o que esperaríamos se a criança tivesse rompido um ligamento, mas não se ela tivesse distendido um músculo. Portanto, a segunda hipótese é preferível”. Esse raciocínio pode ser representado em termos probabilísticos da seguinte forma: “dada a informação de fundo, a probabilidade a priori de H_2 é maior do que a de H_1 ; a probabilidade da evidência condicional a H_2 é maior do que sua probabilidade condicional a H_1 , portanto, a probabilidade a posteriori de H_2 é maior do que a de H_1 ”. Assim representado, o uso da IME pela médica não é incoerente pelos padrões bayesianos; pelo contrário, ela usou considerações explanatórias como uma ajuda para calcular as *priors* e as verossimilhanças necessárias para aplicar o próprio teorema de Bayes.

Nas palavras de Lipton (2004, p. 107), nessa perspectiva considerações explanatórias “proveem uma heurística central que nós usamos para seguir o processo de condicionalização, uma heurística de que precisamos por não sermos muito bons em realizar cálculos de probabilidade”. Esta última parte conta com suporte empírico robusto (Batanero et al, 2009; Wilcox, 2024). Desde o trabalho seminal de Kahneman e Tversky (1973), sabemos que muitos de nossos julgamentos intuitivos vão de encontro ao veredito dado pela teoria da probabilidade. Todavia, não é igualmente claro que considerações de caráter explicativo efetivamente desempenhem tal papel. A menos que se disponha de uma teoria sobre como os dados probabilísticos se traduzem em considerações de cunho explanatório, a afirmação de que esse é o caso nada expressa, eu temo, senão um mero desejo³².

Em síntese, como se pode observar, muito embora o Bayesianismo não consista em uma negação da relevância evidencial de fatores explanatórios, conciliá-lo com a IME não é uma tarefa tão simples. As diferentes sugestões de como isso pode ser levado a efeito ainda carecem de maior refinamento e, por isso, se é de fato possível fazê-lo, permanece ainda, em certa medida, um debate em aberto³³.

³² Em uma nota de rodapé do artigo em que primeiro formulado o desafio Screen-off, Sober & Roche (2013, p. 665) comentam sobre a perspectiva de Okasha (2000). S&R concordam que ela é compatível com a Teoria Bayesiana de Confirmação, porém ressaltam que, se isso é tudo o que a IME envolveria, “protestaríamos, então, que o nome da teoria é enganoso; um rótulo melhor seria inferência para a melhor hipótese” (idem), pontuam.

³³ Repare que, subjacente a todas as propostas (inclusive a avançada por Douven, que abre mão da regra da condicionalização), está a premissa de que o Bayesianismo conta com fundamentos mais sólidos que a IME e, por isso, é dever desta se ajustar ao primeiro, e não o contrário. No entanto, é

3.2. O ARGUMENTO DE SOBER & ROCHE

Apesar de o desafio *Screen-off* ter sido primeiro articulado em Sober & Roche (2013), este não contém sua melhor formulação. Como os autores reconhecem, há nele “algumas passagens potencialmente enganadoras”³⁴ (2019, p. 130), sugerindo que seu escopo seria maior do que de fato pretendiam. Por essa razão, muito embora cobriremos o debate em ordem predominantemente cronológica, é importante, desde já, esclarecer como Sober e Roche (2019), em última análise, concebem as teses e seu escopo.

Impelidos pelas respostas de McCain e Poston (2014, 2017), Climenhaga (2017) e Lange (2017), Sober e Roche (2019) sintetizaram o desafio *Screen-off* como sendo composto de três teses. “Permita que BEST seja a proposição de que H é a melhor explicação rival disponível de O em termos das virtudes $v_1, v_2, \dots$ e v_n ” (p. 131) e “permita que HIGH seja a proposição de que a pontuação geral de H em termos de $v_1, v_2, \dots$, e v_n é alta” (p. 136). Nesse caso, são elas (p. 138):

- SOT: Existem muitos casos realistas em que a informação de fundo codificada em $\text{Pr}(-)$ inclui dados de frequência tais que O *screens-off* EXPL de H, de modo que $\text{Pr}(H \mid O \ \& \ \text{EXPL}) = \text{Pr}(H \mid O)$.
- SOT*: Existem muitos casos realistas em que a informação de fundo codificada em $\text{Pr}(-)$ inclui dados de frequência tais que O *screens-off* BEST de H, de modo que $\text{Pr}(H \mid O \ \& \ \text{BEST}) = \text{Pr}(H \mid O)$.
- SOT**: Existem muitos casos realistas em que a informação de fundo codificada em $\text{Pr}(-)$ inclui dados de frequência tais que O *screens-off* BEST & HIGH de H, de modo que $\text{Pr}(H \mid O \ \& \ \text{BEST} \ \& \ \text{HIGH}) = \text{Pr}(H \mid O)$.

Nas próximas seções, esclareceremos seu sentido e suas particularidades. Contudo, dois comentários não podem ser protraídos.

O verbo frasal “*to screen [something] off*”³⁵ é de difícil tradução. O sentido empregado na discussão vem da Estatística, onde se diz que uma variável *screens-off* outra, quando a primeira torna a última dela probabilisticamente independente.

A ideia talvez seja melhor aprendida por ilustração. Imagine que estamos tentando aferir (P) a probabilidade de um determinado sujeito S possuir determinado genótipo G. Descobrir (I) que seu irmão (tanto por parte de pai quanto por parte de mãe) possui o mesmo genótipo é

possível que os explanacionistas tenham se enamorado da teoria bayesiana rápido demais. Como é sabido, ela enfrenta críticas poderosas (Pollock, 2006; Harman, 1986; Holton, 2014). Vide tópico 2.3.2.

³⁴ No original, “there are some potentially misleading passages in Roche & Sober (2013)”.

³⁵ O Cambridge dictionary o traduz como “to separate one area from another using a wall or other vertical structure”. Disponível em: <https://dictionary.cambridge.org/dictionary/english/screen-off>.

probabilisticamente relevante, na medida em que aumenta nossa confiança na hipótese de que S possui G. Suponha, no entanto, que não sabemos (I) da existência do mencionado irmão. Nesse caso, descobrir (M) que a mãe de S possui G produz o mesmo efeito, ou seja, aumenta nossa confiança de que S também o possui. Ambas as descobertas (I e M), isoladamente, são relevantes. Nada obstante, perceba que, uma vez que descobrimos I, M se torna irrelevante à probabilidade P (e vice-versa). Nesse caso, dizemos que I *screens-off* M de P.

O argumento de Sober e Roche (2013) é uma espécie de *reductio ad absurdum*³⁶. Ele busca demonstrar que, ao atribuir relevância evidencial à explicatividade de uma hipótese, o explanacionista acaba obrigado a endossar uma conclusão absurda: o fato de o consumo de cigarros causar câncer de pulmão aumenta a probabilidade de um indivíduo qualquer acometido da doença ter fumado um elevado número de cigarros.

O motivo pelo qual talvez, à primeira vista, ela não soe absurda para muitos de nós é que hoje sabemos que o consumo de cigarros causa o desenvolvimento de câncer de pulmão. Por essa razão, os autores, inicialmente, nos convidam a voltar um pouco no tempo, a nos reportarmos à época em que esse fato não era conhecido. Décadas atrás, cientistas estudando a relação entre o consumo de cigarros e o desenvolvimento de câncer de pulmão notaram o seguinte: “ $\Pr(S \text{ terá câncer de pulmão} \mid S \text{ já fumou } x \text{ cigarros até agora}) > \Pr(S \text{ terá câncer de pulmão} \mid S \text{ fumou } y \text{ cigarros até agora})$, para todo $x > y$ ”³⁷ (p. 660). Os pesquisadores, portanto, primeiro observaram que havia uma correlação. Quanto maior o número de cigarros consumidos, maior a probabilidade de um indivíduo qualquer vir a desenvolver câncer de pulmão.

No entanto, como já vergastado na literatura, correlação não implica causalidade³⁸. Sober e Roche (2013, p. 661) nos lembram que à época R.A. Fisher (1959), por exemplo, avançou uma hipótese alternativa para a correlação. De acordo com o famoso estatístico, uma possível explicação era a de que o vício por cigarros e o desenvolvimento de câncer tivessem uma causa comum. Talvez um gene que predispucesse o indivíduo a ambos os males.

À luz dessas considerações, S&R, então, argumentam que, se a explicatividade da hipótese possui relevância evidencial, estaríamos comprometidos a endossar a seguinte proposição (p. 661):

Pr (S fumou ao menos 10,000 cigarros antes dos 50 anos de idade | S desenvolveu câncer de pulmão após essa idade e se S fumou essa quantidade de cigarros antes dessa idade e contraiu câncer, o consumo de cigarros

³⁶ Embora os autores não utilizem essa expressão.

³⁷ Troquei as variáveis i e j por x e y, respectivamente, para facilitar a leitura.

³⁸ Isso se relaciona também ao problema da subdeterminação, apontada por Duhem e Quine.

explicaria o câncer de pulmão) > Pr (S fumou ao menos 10,000 cigarros antes dos 50 anos de idade | S desenvolveu câncer de pulmão após essa idade) (grifei).

Dito de outra forma, sendo c uma medida estatística acurada da Pr (S fumou ao menos 10,000 cigarros antes dos 50 anos de idade | S desenvolveu câncer de pulmão após essa idade), descobrir que o consumo de cigarros é a real causa (chamemos de “E”) – e, portanto, a explicação do desenvolvimento da doença – em nada altera essa probabilidade. Ou de forma ainda mais simples, estabelecida a correlação entre consumo de cigarros e câncer de pulmão, a probabilidade p de um determinado indivíduo vir a desenvolver a doença é proporcional ao número n de cigarros consumidos *quer o câncer seja causado pelo cigarro, por uma causa comum ou por outro fator*. Como os autores resumem, sendo H a hipótese e O as informações sobre a observação de tais fatores na população relevante, a $\text{Pr}(H | O \& E) = \text{Pr}(H | O)$.

Outra evidência disso, sustentam, seria o fato de a relação de explicação ser assimétrica entre *explanans* e *explanandum*. Como o célebre exemplo do mastro da bandeira e da sua sombra denotam, o primeiro explica esta última, mas esta não explica aquele. Todavia, de acordo com o Teorema de Bayes, confirmação é uma relação simétrica: “ $\text{Pr}(Y | X) > \text{Pr}(Y)$ se e somente se $\text{Pr}(X | Y) > \text{Pr}(X)$ ” (p. 662).

O exemplo se torna mais vívido quando introduzida uma pequena modificação. Considere que a $\text{Pr}(\text{S fumou cigarros na juventude} | \text{S desenvolveu câncer de pulmão na velhice}) = \text{Pr}(\text{S desenvolveu câncer de pulmão na velhice} | \text{S fumou cigarros na juventude})$. Não há razão *a priori* para contar essa equação como verdadeira, mas tampouco há motivos para a descartá-la *a priori*. Suponha, então, que há suporte empírico para ela. Segundo S&R, nessa hipótese, se a explicatividade fosse evidencialmente relevante, estaríamos comprometidos com a seguinte inequação (p. 662):

$\text{Pr}(\text{S fumou cigarros na juventude} | \text{S desenvolveu câncer de pulmão na velhice e o consumo de cigarros na juventude explicaria o posterior desenvolvimento de câncer}) > \text{Pr}(\text{S desenvolveu câncer de pulmão na velhice} | \text{S fumou cigarros na juventude})$.

Da segunda metade do texto em diante, Sober & Roche antecipam e enfrentam quatro potenciais objeções (p. 663-667). Elas consistem (i) na observação de que dados sobre a frequência³⁹ nem sempre parecem ser a melhor base para inferência (p. 663); (ii) de que a descoberta de que a hipótese implica a evidência possui relevância confirmacional (p. 664); (iii) de que há casos em que uma variável neutralizada (*screened-off*) parece permanecer

³⁹ No original, *frequency data*.

evidencialmente relevante (idem); e (iv) de que a explicatividade, segundo alguns autores, pode ser um fator empregado na fixação da probabilidade inicial (prior) atribuída à hipótese (p. 666).

A primeira objeção é bem ilustrada pelos autores com um exemplo⁴⁰ (p. 663). A obtenção de 550 caras em 1.000 lances de uma moeda pode parecer que favorece a hipótese “a moeda é mais propensa a cair com a face cara para cima” sobre a hipótese “a moeda é justa”. Contudo, se soubermos que a moeda foi retirada de um recipiente onde a vasta maioria é justa, e apenas uma fração muito pequena delas possui tal propensão, a primeira hipótese se sagra a mais provável. Isso poderia ser visto como um caso em que a explicatividade possui valor evidencial. Sober e Roche concedem o cenário, mas rejeitam que a nova conclusão se deva à explicatividade. “Nossa resposta é concordar que informações de fundo influenciam a estimativa da probabilidade; frequências observadas não são a única fonte de informação”, comentam (p. 663).

De fato, em razão da reconhecida não monotonicidade dos raciocínios ampliativos, a maior ou menor quantidade de informação levada em conta influencia diretamente no resultado. E não se encontra evidente no exemplo que a diferença se resume a alguma consideração de caráter estritamente explanatório. A meu ver, teríamos uma resposta ao desafio *Screen-off* nesse caso se, a partir de um mesmo conjunto de informações, a explicatividade de uma das hipóteses é que fizesse a balança pesar em favor dessa.

A segunda objeção retira sua força do fato de que, quando a hipótese implica a observação, isto é, $\Pr(O | H) = 1$, a menos que esta já fosse tida como certa (caso em que integraria o conhecimento de fundo e teria a probabilidade 1), ela invariavelmente proverá algum grau de confirmação. A resposta oferecida no artigo, *prima facie*, também é bastante razoável. O aparato Bayesiano pressupõe que o agente detém onisciência sobre as relações lógicas mantidas entre as diferentes proposições que compõem sua distribuição de probabilidade. Por consequência, condicionarizar sobre relações lógicas ou puramente matemáticas não impacta as probabilidades posteriores, “haja vista que (por assim dizer) elas já são levadas em consideração desde o início” (p. 664). Isso certamente implica que, se a explicatividade se reduz à constatação da existência de tais relações, um agente Bayesiano ideal não traça Inferências à Melhor Explicação. A relevância desse ponto para agentes epistêmicos reais será discutida mais à frente, em conexão com o problema mais geral da idealização no Bayesianismo (vide 3.3.2).

⁴⁰ Modifiquei um pouco o exemplo para ficar mais claro.

Por sua vez, a terceira objeção se baseia no fato de que uma variável tornada irrelevante (*screened off*) pode, não obstante, interferir na probabilidade da teoria sob investigação. A ilustração empregada (p. 664) envolve a transmissão de caracteres hereditários entre três organismos relacionados por descendência/ascendência. Para simplificar, o exemplo assume que a reprodução é uniparental. Nesse caso, a probabilidade de o organismo neto compartilhar o traço hereditário t dado que o organismo pai o detém não é afetada (*screen-off*) pelo fato de que o organismo avô também compartilha a característica. Posto formalmente, significa dizer que (a) $\Pr(\text{o neto tem } t \mid \text{o pai tem } t) = \Pr(\text{o neto tem } t \mid \text{o pai tem } t \ \& \ \text{o avô tem } t)$. Todavia, “seria equivocado concluir disso que o fato de o avô possuir o traço T não fornece nenhuma evidência sobre se o neto também o possui ou não” (p. 664). Formalmente mais uma vez, apesar de (a), é possível que $\Pr(\text{o neto tem } t \mid \text{o avô tem } t) > \Pr(\text{o neto tem } t)$. A resposta de Sober e Roche a esse ponto é que, muito embora isso seja possível em certos casos, não há motivo para se crer que o é no da explicatividade (p. 665).

A quarta e última objeção nos remete a uma das propostas de conciliação da IME com o Bayesianismo previamente discutidas (3.1). Como mencionado, alguns autores reputam plausível que considerações de cunho explanatório possam ser utilizadas na fixação das probabilidades iniciais (*priors*) das hipóteses (Weisberg, 2009; Huemer, 2009). Para Sober e Roche (2013), o problema com a proposta é que, como se costuma dizer na literatura sobre Bayesianismo, “a *prior* de hoje é a (probabilidade) posterior de ontem” (Steel, 2001; Philips & Guttman, 1998; Czégel *et al*, 2022). Isso porque, como (no modelo altamente idealizado do Bayesianismo clássico) o agente atualiza seus graus de crença estritamente pela regra da condicionalização, a *prior* em questão resulta de um processo de condicionalização anterior. Por isso, dizer que a explicatividade ajudaria na fixação das probabilidades iniciais apenas empurraria o problema de justificar a relevância de considerações explanatórias uma etapa para atrás.

3.3. AS RESPOSTAS EXPLANACIONISTAS

Conforme indicado na introdução, nas seções que compõem esse tópico, analisaremos as respostas explanacionistas até aqui apresentadas. Mais especificamente, avaliaremos a sugestão de que a explicatividade tem a propriedade de tornar graus de crença mais insuscetíveis de mudança em face de evidências futuras (McCain e Poston, 2014; 2017); a tese de que a idealização de Onisciência Lógica do aparato Bayesiano nos impede de apreciar a relevância evidencial da explicatividade (McCain e Poston, 2014; Lange, 2017); a tese de que ela é indispensável para a projeção indutiva de predicados (Climenhaga, 2017; McCain e Poston,

2014) e de que ela sinaliza quais tipos de explicações se pode esperar para um determinado *explanandum* (Lange, 2017; 2020; 2024).

3.3.1. Explicatividade, Peso da Evidência e Resiliência dos Graus de Crença

O principal contra-argumento de Poston e McCain (2014) consiste na tese de que, ainda que se conceda que o argumento *screen-off* é cogente, a explicatividade seria ainda evidencialmente relevante em um sentido distinto. “As probabilidades de primeira ordem não são afetadas [pela explicatividade]; mas a resiliência dessas probabilidades, sim” (p. 148). Por se tratar de uma noção impopular mesmo nas obras sobre Bayesianismo, primeiramente exploraremos em que consiste a proposta, para então nos voltarmos à análise do argumento propriamente dita.

Embora o artigo não contenha nenhuma referência a Keynes (1921), a ideia de resiliência visa capturar uma dimensão da relação de suporte evidencial primeiro identificada pelo famoso economista. Em *A Treatise on Probability*, Keynes (1921, p. 77) comenta:

À medida que o conjunto de evidências relevantes à nossa disposição aumenta, a probabilidade [da hipótese] pode diminuir ou aumentar [...] mas algo parece ter aumentado em ambos os casos — temos uma base mais substancial sobre a qual fundamentar nossa conclusão [...] Novas evidências às vezes diminuirão a probabilidade [da hipótese], mas sempre aumentarão seu “peso”.

O problema, como Keynes pôde perceber, é que o aparato Bayesiano parece ser completamente oblíquo a essa importante dimensão. As probabilidades computáveis por meio do teorema revelam a maior ou menor probabilidade da hipótese, mas nada nos informam sobre o quão substanciais as evidências em questão são. Isso fica bastante claro no famoso exemplo do Ônibus Azul (Urbaniak & Di Bello, 2021; Steele & Colyvan, 2023), adaptado do precedente jurisprudencial norte-americano *Smith v. Rapid Transit, Inc*⁴¹.

O exemplo, tomado por muitos juristas como indicativo da inadequação, ou pelo menos insuficiência, do maquinário probabilístico em contextos jurídicos, envolve a seguinte situação: uma pessoa é atropelada por um ônibus. Não há testemunha. Porém, sabe-se que o veículo é necessariamente de uma das duas empresas de transporte operantes na região, a empresa Azul e a empresa Verde, e que a primeira detém 90% dos ônibus em circulação, e a última, os 10% remanescentes. Parece que, se tudo o que importa para o veredito é a probabilidade da hipótese dada a evidência disponível, o julgador estaria autorizado a condenar a empresa Azul a

⁴¹ Para um relato detalhado do caso, vide Risinger (2024).

indenizar a vítima (presumindo-se que 90% constitua prova além da dúvida razoável), o que não parece ser uma decisão bem fundamentada.

A resposta probabilista padrão a essa preocupação se vale do conceito de resiliência. No seminal *Resiliency, Propensities, And Causal Necessity*, Skymrs (1977, p. 705) introduz o conceito, descrevendo-o como sendo “uma propriedade de estabilidade, semelhante ao conceito de robustez na estatística. Uma probabilidade resiliente é aquela que é relativamente insensível a perturbações na nossa estrutura de crenças”. Adaptando um exemplo primeiro entretido por Popper, Skyrms (p. 705-706) o ilustra da seguinte forma:

Você recebe uma moeda e deve atribuir graus racionais de crença à probabilidade de que ela irá cair com a face cara para cima e à probabilidade de que ela cairá com a face coroa para cima. Você não tem certeza se a moeda é justa. Você acredita que há alguma chance de que ela seja tendenciosa, seja para cara ou para coroa, mas não tem mais motivos para pensar que ela seja tendenciosa mais para um lado do que para o outro. Pela "simetria da ignorância", por assim dizer, você chega à conclusão de que $\Pr(\text{coroa}) = 1/2$. Agora compare isso com a situação em que você joga a moeda um grande número de vezes e obtém cerca de 50% coroas; você examina a moeda e a encontra fisicamente simétrica, etc. Agora você tem uma grande quantidade de conhecimento disponível que não tinha no primeiro caso e, no entanto, quando perguntado sobre seu grau racional de crença de que a moeda cairá com a face coroa para cima na próxima jogada, você dará a mesma resposta: $\Pr(\text{coroa}) = 1/2$. A conclusão parece ser que o conhecimento adicional simplesmente não está refletido nos seus graus de crença nos resultados dos experimentos de lançamento de moeda.

De acordo com Poston e McCain, a explicatividade desempenharia um papel análogo aos vários lances da moeda. No exemplo do diagnóstico oncológico, apesar de a descoberta de que o cigarro causa câncer não aumentar a probabilidade de um determinado indivíduo vir a desenvolver a doença (para além daquela já inferida das frequências relativas observadas), a obtenção de uma explicação para a correlação tornaria nossos graus de confiança mais estáveis. A explicatividade não só seria compatível com o Bayesianismo como talvez seria necessária ao adequado manuseio do último.

Poston e McCain também proveem um exemplo próprio (p. 149). No cenário proposto, Sally e Tom são informados de que há 1.000 “x-esferas” no interior de uma urna. Ambos sabem que ela é composta de esferas vermelhas e azuis. A única diferença é que Sally também sabe que, em razão da estrutura atômica das x-esferas, se elas forem armazenadas em quantidades díspares, isto é, se houver mais esferas vermelhas do que azuis (e vice-versa), em poucos minutos os átomos de todas as x-esferas decaem espontaneamente, dando ocasião a uma grande explosão. Os autores, então, nos pedem para imaginar que, de um sorteio de 10 esferas sem

reposição resultam 5 esferas vermelhas e 5 esferas azuis. Após elas serem devolvidas à urna, a probabilidade que tanto Sally quanto Tom racionalmente devem atribuir ao lance aleatório de uma esfera azul é 0.5. No entanto, há uma diferença evidencialmente relevante no conhecimento detido por Sally não refletido em seu grau de crença. Segundo argumentam, isso se torna mais evidente ao considerarmos o que ocorreria na hipótese de um sorteio desigual entre esferas vermelhas e azuis. Se viesse a sair um número maior de bolas azuis, por exemplo, Tom ajustaria seu grau de confiança na retirada de uma nova bola azul para além de 0.5, enquanto Sally supostamente permaneceria no mesmo patamar.

Sober e Roche (2014) não acharam o argumento muito convincente. Para eles, o fator determinante tem a ver com diferença no conhecimento de fundo presumido, e não com a explicatividade em si (p. 196). A fim de tornar a diferença mais saliente, os autores apontam que há diferença entre (i) “Sally, mas não Tom, saber que, se as x-esferas azuis e vermelhas forem armazenadas em números desiguais, haverá uma enorme explosão” e (ii) “Sally, mas não Tom, ter uma explicação do porquê de a probabilidade de sair uma x-esfera azul em um sorteio aleatório ser 0,5”. Para Sober e Roche, no exemplo, (i) é que torna os graus de crença de Sally resilientes, não (ii).

Mas para S&R este não seria o único problema com o contraexemplo. Para os autores, o cenário introduzido por Poston e McCain (2014) modifica a disputa substancialmente em, pelo menos, três formas (p. 196). Primeiramente, porque o argumento *screen-off* tem por objeto um sentido bem específico de relevância evidencial, qual seja, o de incremento da probabilidade de a hipótese ser verdadeira, de acordo com a Teoria Bayesiana de Confirmação. É dizer, para a explicatividade (E) ser evidencialmente relevante, a probabilidade da hipótese à luz de determinada observação, mais a explicatividade, deve ser superior à probabilidade da hipótese dada a evidência sozinhas – a inequação $\Pr(H \mid O \ \& \ E) > \Pr(H \mid O)$ deve ser verdadeira.

Em segundo lugar, o cenário descrito por P&M modifica a disputa porque, enquanto o argumento *Screen-off* trabalha com o contrafactual “H, se verdadeira, explicaria O”, a resposta de Poston e McCain (2014) entretém o que ocorreria se *o agente* possuísse tal explicação. Ou seja, SOT está restrito à justificação proposicional, e não à justificação epistêmica em geral. Por fim, S&R salientam que a tese se refere à relevância evidencial da explicatividade em relação à hipótese considerada isoladamente, em si mesma (e um *background knowledge*). O caso de Tom e Sally, no entanto, ao depender da existência de uma alternativa H*, introduz um aspecto comparativo que transcende o âmbito original do debate.

Na tréplica, McCain e Poston (2017) responderam que a distinção entre (i) e (ii) (do penúltimo parágrafo) não refuta o contraexemplo sobre Sally e Tom, pois “a propriedade de

Sally possuir a explicação é a mesma propriedade de ela saber o fato relevante sobre as x-esferas” (p. 124). Distingui-las, como S&R sugerem, seria equivalente a negar que água tem a propriedade de apagar fogo, pois, uma vez que as propriedades do hidrogênio, por exemplo, são levadas em consideração, a presença de água se torna irrelevante uma vez introduzida H_2O . Afinal, água é H_2O (p. 124). Além disso, M&P comentam que, se fatos sobre a relação causal existente contam como explanatórios, como mantém a Teoria Causal da Explicação, não estaria ocorrendo *screen-off* propriamente, haja vista que, nesse caso, os fatos já foram levados em consideração. Como (i) e (ii) são equivalentes, condicionarizar sobre ambas seria contar a evidência duas vezes.

Como nos dois outros artigos que publicaram sobre o assunto (2017; 2019), Sober e Roche abordaram as respostas de Climenhaga (2017) e Lange (2017), respectivamente, M&P tiveram a última palavra nesse ponto. Não parecem, contudo, ter ficado também com a razão. Como S&R ressaltaram (penúltimo parágrafo), o contrafactual da explicatividade se refere à hipótese em si, e não à posse da explicação pelo agente. Importa se a hipótese, em caso de verdadeira, constituiria uma explicação, e não se, caso o agente soubesse determinado dado, este influenciaria a probabilidade da hipótese. O exemplo da água, apesar de incontroverso, não é análogo ao dos cigarros. “Água” é um dos rótulos que damos a essa substância líquida que resulta da combinação de duas moléculas de hidrogênio com uma de oxigênio. A explicatividade, por outro lado, *não é* a evidência, alguma suposição de fundo ou mesmo a hipótese; mas uma propriedade da hipótese, capturada pela proposição de que, se fosse verdadeira, a hipótese explicaria a observação em questão.

Sober e Roche (2014) não perseguem o ponto em maiores detalhes, provavelmente porque, como esclarecido, o argumento *Screen-off* trata apenas da relevância evidencial entendida como incremento na probabilidade da hipótese. Como a resiliência seria uma dimensão à parte, ela simplesmente estaria fora de seu alcance.

Permita-me, então, sintetizar o que compreendo como o resultado dessa discussão. Ao apelar para a noção de resiliência, M&P (2014; 2017) chamaram a atenção para o fato de que há diferentes sentidos em que se pode entender algo como evidencialmente relevante. Importante que seja, a concepção pressuposta por S&R não é a única e, portanto, é possível que, mesmo que o argumento *screen-off* seja sólido, a explicatividade seja epistemicamente relevante de outras formas. Nada obstante, isso não significa dizer que a explicatividade, de fato, torne a convicção em uma dada hipótese mais resiliente. Que ela possa e até mesmo seja frequentemente tomada dessa forma é menos controverso, porém nada nos revela, senão um aspecto da *psicologia* explanacionista. A observação de milhares de lances de uma mesma

moeda claramente fornece evidência da existência de algum viés ou não. Mas, para concluir que a explicatividade possui o mesmo efeito, precisamos de um argumento. Simplesmente assumir que este é o caso *begs the question*.

3.3.2. O Problema da Onisciência Lógica

Em sua resposta aos contra-argumentos de McCain e Poston (2014), Sober e Roche (2014) pontuaram que o argumento *Screen-off* “baseia-se na suposição de que agentes racionais são logicamente oniscientes, de modo que todos os fatos puramente lógicos e matemáticos estão codificados em Pr” (p. 195). As limitações que essa idealização impõe às conclusões a serem tiradas de SOT foram enfatizadas tanto por McCain e Poston (2017, p. 128) quanto por Lange (2017, p. 5). Nesse tópico, após uma breve introdução à origem dessa suposição, analisaremos os argumentos apresentados pelos autores.

Segundo Titelbaum (2022, p. 18), a Epistemologia Bayesiana envolve o endosso de, pelo menos, duas teses:

1. Agentes [epistêmicos] possuem atitudes doxásticas que podem ser representadas, de forma útil, atribuindo números reais às respectivas proposições.
2. As demandas que a racionalidade impõe sobre tais atitudes doxásticas podem ser representadas por meio de restrições matemáticas a essas atribuições de números reais, estreitamente relacionadas ao cálculo de probabilidades.

Bayesianismo, portanto, implica Probabilismo (mas não o contrário), “a perspectiva filosófica de que a racionalidade exige que as crenças de um agente formem uma distribuição de probabilidade, isto é, satisfaçam os axiomas de Kolmogorov (idem, p. 33)”. Como vimos no tópico 3.1., os axiomas são demandas de coerência e, por isso, *prima facie* bastante razoáveis. Entretanto, como frequentemente acontece na Filosofia, suposições aparentemente inocentes não raro escondem grandes perplexidades.

Um dos axiomas de Kolmogorov (1963) é o de que toda tautologia em uma dada linguagem L possui probabilidade 1. Como atribuir probabilidade 1 equivale a tomar um dado evento como certo, e uma tautologia não pode ser falsa, trata-se de um requisito à primeira vista também bastante plausível. Todavia, como é sabido, todo conjunto de proposições fechado sob os conectivos lógicos tradicionais (conjunção, disjunção, negação etc) dá origem a um número infinito de tautologias, e se é irracional não estar absolutamente certo de todas elas, isso significa que, de acordo com o Probabilismo, devemos ser oniscientes sobre todas as implicações lógicas das proposições que entretemos.

O problema foi bem capturado por Garber (1983, p. 105)⁴²:

O agente bayesiano que não é logicamente onisciente é incoerente e parece violar a única condição necessária para a racionalidade sincrônica com a qual todos bayesianos podem concordar. Esta é uma assimetria que cheira à temida distinção analítico-sintética. Mas, deixando de lado os escrúpulos sobre o status metafísico ou epistêmico dessa distinção, a assimetria no tratamento do conhecimento lógico e empírico é, à primeira vista, absurda. Não deveria ser mais irracional não saber o menor número primo maior que um milhão do que não saber o número de volumes na Biblioteca do Congresso.

Com essa discussão como pano de fundo⁴³, McCain e Poston (2017) oferecem duas respostas à afirmação de que “a explicatividade não tem relevância confirmacional, uma vez que fatos puramente lógicos e matemáticos tenham sido levados em consideração” (Sober & Roche, 2014).

Primeiramente, os autores salientam que ela pressupõe que considerações explanatórias seriam completamente independentes da relação de probabilidade existente entre a hipótese e a evidência em questão. Mas, “se conhecer os fatos explicativos equivale a conhecer certas relações probabilísticas entre a hipótese e a evidência, então essa resposta falha” (p. 125). Segundo argumentam, esse seria o caso no exemplo das x-esferas, em que “o conhecimento dos fatos causais *simplesmente é* o conhecimento de fatos explanatórios” (grifo no original) (idem).

Em segundo lugar, M&P defendem que há ainda uma questão mais ampla envolvendo o que conta ou não como um fato lógico ou matemático (p. 125):

É um fato lógico ou matemático que água extingue fogo? Se uma pessoa tivesse conhecimento completo da teoria química, poderia deduzir que a água extingue o fogo em um ambiente normal (ou que a probabilidade disso seria alta)? Presume-se que sim. Mas se esses tipos de fatos são fatos de que um agente logicamente onisciente estaria ciente, então claramente, uma vez que eles são levados em conta, a explicatividade se torna evidencialmente irrelevante. Novamente, isso ocorre porque os fatos lógicos/matemáticos abrangem os fatos explicativos. Assim, a explicatividade parece ser irrelevante em termos de evidência, mas isso é simplesmente porque esses fatos já foram considerados. Portanto, se é assim que se deve entender a tese SOT de R&S,

⁴² O Problema da Onisciência Lógica nos remete ao problema mais geral da idealização na Epistemologia Bayesiana. Como é sabido, mesmo a ideia de *credences* ou graus de crença - a matéria prima de toda a teoria - depende de certas idealizações (Pollock, 2006; Harman, 1986; Holton, 2014). Ninguém provavelmente colocou o problema de forma tão incisiva, para não dizer ácida, do que Horgan (2017, p. 250): “a epistemologia formal bayesiana é semelhante, no aspecto relevante, a disciplinas ultrapassadas como a Alquimia e a Teoria do Flogisto: não tem por objeto nenhum fenômeno real”.

⁴³ Naturalmente, o problema vem mantendo bayesianos ocupados desde sua primeira formulação por Savage (1967). A compreensão e a avaliação do mérito das considerações tecidas por McCain e Poston (2017) e Lange (2017), no entanto, não nos impõem que persigamos esse ponto em maiores detalhes.

isso equivale à proibição de contar os mesmos fatos duas vezes. Isso não deveria preocupar nenhum explanacionista.

Como não temos uma resposta de Sober e Roche a essas considerações, podemos apenas conjecturar como eles provavelmente responderiam. A afirmação de que (FLM) “a explicatividade não tem relevância confirmacional, uma vez que fatos puramente lógicos e matemáticos tenham sido levados em consideração”, de fato, parece pressupor que Sober e Roche os concebem como sendo de natureza explanatória. Nesse caso, ainda que a idealização referente à onisciência lógica os exclua, parece que deveriam, no mínimo, concordar que, como nenhum de nós é logicamente onisciente, para todos os efeitos práticos, a explicatividade é evidencialmente relevante. Ou de forma mais resumida, a explicatividade é evidencialmente relevante para agentes reais - que, muitos de nós diriam, é o que realmente importa.

Dito isso, não está claro que o argumento *Screen-off* dependa da aceitação de FLM. Repare na forma como McCain & Poston (2017) empregam a expressão “fatos explanatórios”. Num dos excertos supratranscritos, os autores afirmam que “o conhecimento dos fatos causais *simplesmente é* o conhecimento de fatos explanatórios”. Mesmo que se ignore as complicações de se introduzir a noção de conhecimento⁴⁴ na disputa, é importante relembrar que a posição explanacionista não tem a ver com relevância evidencial de “fatos explanatórios” (o que quer que eles signifiquem), mas com uma proposição em particular - a de que (E) “obtemos evidência de que p ao percebermos que, se verdadeira, ela explicaria a evidência disponível” (McCain, Poston, 2024, p. 329). $\Pr(H | O \& E) = \Pr(H | O)$ é a única equação defendida por SOT⁴⁵.

Uma das três objeções ao argumento *Screen-off* avançadas por Lange (2017) também se relaciona com o Problema da Onisciência Lógica. O autor sustenta que “parece haver casos em que E é uma necessidade lógica” (p. 5), como, por exemplo, quando a observação (O) é uma lei da natureza. Nesses casos, “embora Roche e Sober estejam corretos ao afirmar que $\Pr(H|O\&E) = \Pr(H|O)$, essa igualdade é trivial, uma vez que $\Pr(E) = 1$ ” (idem). Ele sintetiza (p. 5):

Em outras palavras, se E é uma necessidade lógica, então o fato de que o teste *Screen-off* considera (sob onisciência lógica) E como evidencialmente irrelevante não reflete o fato de que a explicatividade é irrelevante em termos de evidência, mas apenas o fato de que, no maquinário bayesiano, qualquer necessidade lógica é irrelevante em termos de evidência.

⁴⁴ A relevância da explicatividade seria proposicional e não epistêmica.

⁴⁵ A igualmente problemática identificação das relações de causalidade e explicatividade será explorada no próximo tópico 3.3.3.

A resposta de Sober e Roche (2019) a Lange (2017) se dedica precipuamente às duas outras objeções levantadas no aludido artigo, relegando a esse ponto apenas um comentário em uma de suas notas de rodapé⁴⁶ (p. 139). S&R anotam que, ainda que Lange esteja certo, e E possa ser uma verdade logicamente necessária em alguns casos, há diversos outros, como o dos cigarros, em que isso não é verdade e nos quais, portanto, E seria evidencialmente relevante.

Perceba, no entanto, que está longe de ser incontroverso que, quando o *explanandum* é uma verdade necessária, a explicatividade é demonstravelmente dotada de valor epistêmico (apesar de a idealização da onisciência lógica impedir o aparato de Bayesiano de reconhecê-la como tal). Parece-me que a atribuição de maior credibilidade à hipótese nesses casos pressupõe a identificação de coisas notadamente distintas: (i) o fato de que X, se verdadeiro, explicaria Y e (ii) o fato de que, se X é verdadeiro, X é a explicação de Y. Noutras palavras, uma coisa é o contrafactual de que uma determinada hipótese, caso verdadeira, constituiria uma explicação para determinado evento ou fenômeno; outra, o fato de que a hipótese em questão efetivamente é a verdadeira explicação.

No texto, Lange (2017) discute dois exemplos. O primeiro envolvendo o fato de que a existência de uma lei da natureza sobre a conservação de energia, se verdadeira, explicaria o porquê desse fenômeno, e o segundo, de que as chances objetivas do decaimento de um determinado átomo dentro de um dado período de tempo explicariam o fenômeno (p. 4-5). E apesar de ser evidente que a existência da lei e da probabilidade objetiva seriam, respectivamente, (ii) as explicações para a conservação da energia e para o decaimento do átomo do modo observado, *i* deixa em aberto se este é, de fato, o caso.

O problema da Onisciência Lógica também aparece no texto de Poston e McCain (2014) em conexão com outra linha de resposta à SOT. Além do argumento relativo à resiliência supostamente conferida pela explicatividade aos graus de crença, P&M (2014) afirmam que “há um uso mais amplo da explicatividade que o argumento [*screen-off*] não impugna” (p. 146), que corresponde às situações em que não dispomos de dados acerca da probabilidade

⁴⁶ Literalmente, S&R fazem três comentários. Mas, como se pode ver, podem ser sintetizados em uma única observação: We have three comments. First, even if Lange is right that there are cases where $\Pr(H \mid O \& EXPL) = \Pr(H \mid O)$ because EXPL is a necessary truth, this leaves it open, as per SOT, that there are also many realistic cases in which the background information codified in $\Pr(-)$ includes frequency data such that O screens-off EXPL from H in that $\Pr(H \mid O \& EXPL) = \Pr(H \mid O)$. Second, EXPL isn't a necessary truth in cases like our smoking case. Third, as we explain in Sections III IV, and V, SOT and its cousins SOT* and SOT** are far from trivial, given their implications with respect to various versions of IBE.

antecedente (prior) da(s) hipótese(s) sob consideração⁴⁷. Como o argumento *Screen-off* trata apenas de casos em que as probabilidades antecedentes são determináveis (como ilustra o exemplo dos cigarros), seria temerário afirmar que ele generaliza para estes últimos casos. Seria perfeitamente possível que, nesses outros cenários, a explicatividade fosse evidencialmente relevante mesmo no sentido de incremento da probabilidade.

Embora Sober e Roche, como vimos, mais tarde tenham acabado restringindo o escopo da tese (2019), no artigo de 2014, em resposta a essa proposta, defenderam que a explicatividade seria irrelevante mesmo quando a natureza das hipóteses envolvidas não parece admitir a obtenção de dados sobre a frequência. O motivo seria o seguinte. Suponha que H implica O . Nesse caso, a $\Pr(O | H) = 1$ e, por consequência, se H e O não possuem probabilidade 0 ou 1, ao descobrir O , um agente racional deve aumentar sua confiança na veracidade de H . Se a explicatividade possui relevância evidencial nesses casos, S&R argumentam, ao descobrir posteriormente E (isto é, que H explica O), tal fato conferiria ainda mais confirmação à teoria. Todavia, nessa hipótese o *boost* probabilístico seria tão arbitrário quanto no caso do diagnóstico oncológico.

É interessante que, segundo Sober e Roche (2014), “comentários similares também se aplicam a casos envolvendo explicações de natureza probabilística” (p. 195). Para ilustrar esse ponto, eles nos propõem um novo cenário fictício: uma máquina produz tanto moedas justas quanto moedas com um viés de 95% para o resultado cara, porém em uma proporção desconhecida. Uma moeda qualquer por ela produzida é lançada, e o resultado é *cara*. “Suponha que você calcule $\Pr(H | O)$ em parte notando que $\Pr(O | H) = 0,95$. A $\Pr(H | O \ \& \ E) > \Pr(H | O)$? Claramente não. O que é verdadeiro é uma igualdade: $\Pr(H | O \ \& \ E) = \Pr(H | O)$ ” (p. 195). De acordo com os autores, mais uma vez, a aparente relevância evidencial da explicatividade se reduz à relevância dos fatos lógico-matemáticos presumidos na onisciência lógica.

Em Poston & McCain (2017), encontramos uma réplica a esses apontamentos semelhante à anteriormente considerada. Os autores disputam a premissa assumida por S&R de

⁴⁷ Vale ressaltar, isso pode ocorrer por duas razões distintas - e Poston e McCain (2014) tem em vista apenas uma delas. Em primeiro lugar, é possível que a informação simplesmente esteja indisponível, não obstante seja perfeitamente concebível obtê-la. É o caso, por exemplo, da descoberta de uma nova doença cuja incidência em determinada população ainda não foi estudada. Nós os desconhecemos, mas os dados existem. A limitação, pode-se dizer, é apenas epistemológica. No entanto, há casos – e esses é que dariam amparo à tese explanacionista – em que a indisponibilidade de tal informação parece ser inerente à natureza da própria hipótese ou explicação. Considere, por exemplo, a Teoria da Evolução ou a Teoria da Relatividade Geral. A probabilidade antecedente (prior) destas não só parece inescrutável como sequer parece existir ou fazer sentido. Teríamos, assim, aqui uma limitação ontológica.

que “descobrir que H implica O não é conhecimento explicativo” (p. 126). Segundo P&M, negar o caráter explanatório da relação de implicação seria um equívoco por três motivos.

Inicialmente, porque, embora não seja imune a contraexemplos, o modelo dedutivo-nomológico proposto por Hempel & Oppenheim se mostra adequado em alguns casos, e ele é um modelo de explicação. Além disso, descobertas matemáticas decorrem da manipulação de relações lógicas e probabilísticas *a priori* e, não obstante, parece ser inequívoco que dispomos de *explicações* matemáticas para as mesmas (p. 126-127). Por fim, P&M também advertem que não se deve identificar o caráter explicativo de uma hipótese ou proposição com o “estado psicológico resultante de compreender alguma coisa” (p. 127). O erro de S&R, conjecturam, possivelmente estaria em pressupor que o explanacionismo está comprometido com a relevância evidencial de “momentos aha!” (p. 127).

O embate ilustra as dificuldades em se tomar o conceito de explicação (e correlatos, como explicatividade, explanatório etc) como primitivo nessa exploração. Perceba que, muito embora dificilmente Lipton concordaria em caracterizar o que chamou de *loveliness* com “momentos aha!”, as ideias não são totalmente estranhas uma à outra. Afinal, para ele, a melhor explicação seria “aquela que, se correta, seria a mais explicativa ou *proporcionaria maior entendimento*” (Lipton, 2004, p. 59) (grifei).

Mas, ainda que rejeitemos essa caracterização, concebendo o caráter explicativo de uma hipótese unicamente em termos de virtudes explicativas, é questionável se a relação de implicação *per se* de fato pode ser concebida como explicativa. O que os sucessos e os insucessos do modelo dedutivo-nomológico (em particular o citado exemplo da sombra do mastro de bandeira) parecem ilustrar é justamente o contrário: relações explicativas não espelham perfeitamente relações de implicação.

Por essas razões, mesmo sem olvidar o problema da idealização para o Bayesianismo, é forçoso concluir que a explicatividade – entendida como a proposição (E) de que, se verdadeira, a hipótese explicaria a evidência disponível – permanece não reivindicada.

Vejamos, então, se melhor sorte assiste aos demais participantes do debate.

3.3.3. Considerações explanatórias e Projeção Indutiva

A terceira linha de resposta ao argumento *screen-off* proeminente nessa discussão é a de que a explicatividade seria necessária para a projeção indutiva de predicados. Poston e McCain (2014) afirmam que “não é possível obter projeção indutiva para casos não examinados sem considerações explicativas” (p. 150) e que, “da forma como entendemos, essa é a lição do problema do Grue. Regularidades do tipo Grue não são projetáveis porque não são explicativas”

(idem). O texto, contudo, não defende a tese. Em uma das notas de rodapé, recomendam ao leitor interessado “Huemer (2009) e Poston (forthcoming) para uma defesa dessa afirmação” (pág. 153).

Em *How Explanation Guides Confirmation*, malgrado não empregue o termo “projeção indutiva”, Climenhaga (2017) avança um argumento para o mesmo efeito, sustentando que determinada observação é capaz de confirmar uma teoria “*apenas* na medida em que temos evidências de que a conexão explicativa descrita existe” (p. 368) (grifei).

Importante ressaltar, a ideia de que a relação de suporte evidencial depende da existência de uma conexão explanatória entre hipótese e observação não é nova. Achinstein (1983; 2001; 2005; 2024) a defende já há quatro décadas. Segundo argumenta, a teoria de que suporte evidencial é melhor compreendido como relevância probabilística positiva (RPP), com sua noção de confirmação incremental, possui implicações pouco palatáveis. Em seu contraexemplo mais recente, Achinstein (2024, p. 105) aponta que, a menos que se introduza a exigência de uma conexão explanatória, não há como negar que o fato de entrarmos em um elevador, por exemplo, constituiria evidência de que sofreremos um acidente de elevador.

O argumento de Climenhaga (2017) é igualmente ambicioso. Pretende demonstrar que não só a igualdade defendida por S&R “não é em geral verdadeira, mas também é falsa nos próprios tipos de casos em que Roche e Sober se concentram, envolvendo dados de frequência” (p. 359).

Climenhaga (2017) inicia observando que não é imediatamente óbvio para muitos de nós que SOT é falsa, por envolver um contrafactual (H, se verdadeira, explicaria), e em circunstâncias normais, “só obteríamos evidências para um contrafactual sobre um caso específico como este obtendo evidências para afirmações explicativas mais amplas, como ‘em geral, fumar causa câncer’” (p. 363). Por essa razão, ele propõe o seguinte (idem):

Digamos que há uma conexão explicativa entre dois fenômenos X e Y se, e somente se, pelo menos às vezes, X causa Y, Y causa X, ou X e Y tiverem uma causa comum. Seja C1 a hipótese de que fumar causa câncer, C2 a hipótese de que câncer causa o consumo de cigarros, e C3 a hipótese de que eles têm uma causa comum. Então, $\sim [C1 \vee C2 \vee C3]$ é a afirmação de que não há conexão explicativa entre fumar e câncer.

À luz dessas observações, Climenhaga propõe uma revisão de SOT, que ele designa SOT*, nos seguintes termos: “sendo C1 a proposição de que às vezes fumar causa câncer, $P(S \text{ fuma} \mid S \text{ tem câncer} \ \& \ C1) = P(S \text{ fuma} \mid S \text{ tem câncer})$ ”⁴⁸ (p. 363).

⁴⁸ Omiti a referência ao conhecimento de fundo (K) para facilitar a leitura.

SOT*, no entanto, seria falsa, pois, “na hipótese (extremamente improvável) de que não há nenhuma conexão explicativa entre fumar e desenvolver câncer, os dados de frequência observados são uma enorme coincidência” (idem).

O texto prossegue com uma intimidadora série de derivações dessa equação até o que se pode considerar sua premissa fundamental (PF): “ $P(S \text{ fuma} \mid S \text{ tem câncer} [C1 \vee C2 \vee C3]) > P(S \text{ fuma} \mid S \text{ tem câncer})$ ” (p. 365). E esta implicaria que SOT* é falsa na medida em que “presumivelmente, qualquer um de C1, C2 e C3 também autoriza a extrapolação dos nossos dados sobre frequência” (idem).

Sober e Roche (2017) abordam esse argumento. Os autores destacam que, se PF implica $P(S \text{ fuma} \mid S \text{ tem câncer} \& C1) > P(S \text{ fuma} \mid S \text{ tem câncer})$, como Climenhaga pretende, também seria verdadeiro que $\Pr(S \text{ tem câncer} \mid S \text{ fuma} \& \neg C1) = \Pr(S \text{ tem câncer} \mid \neg C1)$, mas esse não seria o caso porquanto “a suposição de que fumar nunca causa câncer deixa em aberto a possibilidade de que muitas vezes eles possuem uma causa comum e, assim, não seja ‘coincidência’ ou ‘acaso’” (p. 587). E essa, acrescentam, não é a única forma de se perceber que PF é falsa. Utilizando um exemplo proposto por Sober (2001), S&R explicam (idem):

Seja F a classe de dias nos últimos 200 anos em que os preços do pão britânico estiveram acima da média, e G a classe de dias em que os níveis do mar em Veneza estiveram acima da média, e suponha que tanto os preços do pão britânico quanto os níveis do mar em Veneza aumentaram monotonicamente durante esses dois séculos. Então, embora não haja nenhuma conexão explicativa entre F e G, para um dia d selecionado aleatoriamente nesses 200 anos, $\text{Freq}_s(d \text{ está em } G \mid d \text{ está em } F) > \text{Freq}_s(d \text{ está em } G)$ e a associação não é uma coincidência; novas amostras retiradas desses 200 anos exibirão confiavelmente o mesmo padrão.

Climenhaga (2017) reconhece que, em casos como o do exemplo, PF pode ser falsa, mas somente se o conhecimento de fundo não contiver toda a informação temporal pertinente. Além disso, acrescenta que Sober deveria considerar PF plausível, pois, na obra de onde retirado o exemplo acima, este afirma que a inferência para uma causa comum “é racional porque é frequentemente o caso de que uma explicação de causa comum prevê uma correlação onde uma explicação de causas independentes não o faz” (2001, p. 342-343).

Sobre a primeira consideração, S&R pontuam que há também exemplos devido a informação relativa à posição das variáveis no espaço (p. 587). E em relação à segunda, Sober ressalva que, naquele contexto, o ponto sendo estabelecido se restringe à noção de favorecimento de acordo com a chamada Lei da Veromissilhança. É dizer, onde CC representa causa comum e SC causas independentes ou separadas, $\Pr(O \mid CC) > \Pr(O \mid SC)$ - do que não decorre, contudo, que $\Pr(CC \mid O) > \Pr(SC \mid O)$ (p. 588).

Note-se que, mesmo que superadas essas objeções, para que o contra-argumento de Climenhaga tenha alguma força contra SOT (e não apenas SOT*), ele deve estabelecer outra premissa não menos controversa. Climenhaga (2017) pede que imaginemos saber “que nada além de fumar causará câncer em S (e que S não terá câncer sem motivo)” (p. 368). Nesse caso, a probabilidade de que S não era fumante, dado que ele possui câncer, é, de fato, zero - ou seja, a $P(\neg C1 \mid S \text{ tem câncer}) = 0$. Entretanto, está longe de ser evidente porque deveríamos entretê-la. Nas palavras de Sober & Roche, tudo o que Climenhaga demonstra com tais estipulações é que “suposições de fundo podem ser elaboradas de modo a tornar E evidencialmente relevante” (2017, p. 586).

Mas não é este o elefante na sala. A meu ver, Climenhaga (2017) faz outra suposição ainda mais problemática, que é identificar explicatividade com causalidade. Relembre que, de acordo com sua definição, “há uma conexão explicativa entre dois fenômenos X e Y se, e somente se, pelo menos às vezes, X causa Y, Y causa X, ou X e Y tiverem uma causa comum” (p. 363). Entretanto, mesmo que todas as demais suposições impugnadas por S&R fossem defensáveis, parece-me que esta torna todas as conclusões dela decorrentes um tanto anticlimáticas, senão triviais. S&R, para não dizer qualquer filósofo, jamais disputariam que a existência de uma relação causal é evidencialmente relevante. Todavia, está longe de ser claro que $\sim [C1 \vee C2 \vee C3]$ corresponda à negação de que existe uma conexão explicativa. Corresponde, é verdade, à negação de que há uma relação causal entre uma variável X e uma variável Y qualquer, mas dizer que uma conexão causal é explicativa é diferente de afirmar que uma conexão causal é uma conexão explicativa.

Por esse motivo, pouco surpreende quando, após diversas derivações, o autor chega à conclusão de que “ $\Pr(S \text{ fuma} \mid \text{desenvolve câncer} \wedge E) > \Pr(S \text{ fuma} \mid S \text{ desenvolve câncer})$ ” (p. 366). O truque, afinal, está no nexo causal entre o consumo de cigarros e o desenvolvimento da doença.

O insucesso do argumento de Climenhaga (2017) evidentemente não significa que a projeção indutiva independa da existência de uma conexão explicativa. Seria uma conclusão, no mínimo, precipitada. Em princípio, é possível que as fontes indicadas por P&M (Huemer, 2009; Poston, 2014) sejam mais bem-sucedidas na defesa da tese. Por isso, antes de passarmos aos próximos argumentos em defesa da explicatividade, é aconselhável considerar o que estas têm a dizer.

O artigo de Huemer (2009), citado anteriormente (tópico 31.), apresenta a noção de Prioridade Explanatória⁴⁹ e defende a possibilidade de concebê-la como “uma solução parcial para o problema da interpretação do Princípio da Indiferença” (p. 354). Como se pode ver, o trabalho vai além da questão em apreço. Felizmente, a discussão em Poston (2014) dedicada especificamente ao tema (da necessidade de uma conexão explicativa para a projeção indutiva) “se baseia fortemente na discussão [encontrada na obra] de Huemer”⁵⁰ (p. 175).

No tópico intitulado *Bayesian Explanationism*, Poston (2014) advoga que considerações explanatórias são necessárias para que se possa evitar o ceticismo indutivo, lançando mão do seguinte exemplo (p. 174):

Suponha que eu tenha duas moedas em meu bolso. Uma delas tem um viés perfeito para caras (ou seja, uma moeda com duas caras) e a outra é justa. Eu escolho uma ao acaso e vou jogá-la dez vezes. Os primeiros sete lançamentos resultam todos em caras. Considere duas hipóteses. H_1 afirma que a moeda justa foi escolhida e, no entanto, todos os seus lançamentos resultaram em caras. H_1 é uma hipótese estranha, mas nada na metodologia bayesiana impede hipóteses esquisitas. H_2 é uma hipótese mais simples, que afirma que a moeda viciada foi a escolhida. Qual hipótese tem maior confirmação após observar sete caras? Claramente, H_2 . Por quê? Porque H_2 explica a evidência, enquanto H_1 apenas implica a evidência. H_1 deixa a sequência positiva de sete caras totalmente misteriosa.

O caro leitor deve ter percebido que a razão dada para a suposta superioridade de H_2 vem com certa ironia. Isso porque a explicatividade da relação de implicação, defendida com cunhas e dentes em Poston e McCain (2014) contra o argumento *Screen-off*, aqui simplesmente desaparece. Infelizmente, Poston (2014) não fornece um critério por meio do qual se possa determinar quando uma relação de implicação é dotada de explicatividade e quando ela não o é. Segundo defende, “a conclusão de que a virtude explicativa de H_2 é o fator crucial em sua confirmação é compatível com a atualização bayesiana” (p. 174). E o motivo seria o seguinte (p. 174-175):

Começamos atribuindo probabilidades iniciais às duas hipóteses. A moeda a ser selecionada é justa ou viciada. Assim, $\Pr(M_j) = \Pr(M_v) = 0,5$. H_2 , a hipótese de que a moeda viciada foi a selecionada, é apenas $\Pr(M_v)$. $H_1 = (M_j) \& \text{resultados caras em cada lançamento}$. Seja D_7 a observação de que os sete primeiros lançamentos resultaram todos em caras. $\Pr(D_7 | H_1) = \Pr(D_7 | H_2) = 1$. Usando a forma *odds* do Teorema de Bayes:

$$\frac{\Pr(H_1 | D_7)}{\Pr(H_2 | D_7)} = \frac{\Pr(H_1)}{\Pr(H_2)} \times \frac{\Pr(D_7 | H_1)}{\Pr(D_7 | H_2)}$$

⁴⁹ No original, *explanatory priority*.

⁵⁰ No original, “the following paragraphs rely heavily on Huemer’s discussion”.

Uma vez que a última razão é 1, a confirmação depende da primeira razão, a saber, $\Pr(H_1) / \Pr(H_2)$. A probabilidade de H_2 é 0,5. E a probabilidade de H_1 é $0,5 \times 1/2^7$. Claramente, D7 favorece H_2 sobre H_1 .

Dois comentários. Primeiro, observe que Poston (2014) pressupõe a viabilidade da questionável distinção feita por Bird (2017) entre virtudes explicativas externas e internas. Como consta de forma explícita do trecho citado, a diferença não residiria na forma como cada uma das hipóteses se relaciona com o *explanandum* - e, portanto, pode-se argumentar, não se encontra nos méritos explicativos de qualquer delas. Não bastasse, parece haver uma leitura mais plausível do porquê de estarmos justificados em preferir a segunda hipótese. E a diferença, creio, reside simplesmente na arbitrariedade da suposição de que todos os lances da moeda justa deram cara. O sorteio da moeda justa poderia, em princípio, ter gerado 127 outras combinações⁵¹, enquanto a moeda enviesada, apenas a observação que temos. Excluir essas possibilidades na formulação da hipótese sem motivo não parece denotar a alegada superioridade explicativa da hipótese rival, mas meramente o emprego de uma metodologia canhestra.

3.3.4. IME, Suposições de fundo e Tipos de Explicação

Lange (2017) apresenta três objeções ao argumento screen-off (p. 1): (1) considerações explanatórias são evidencialmente relevantes porque, pelo menos em alguns casos, aumentam a probabilidade de uma hipótese ao restringir o espaço de possibilidades pela exclusão de algumas das hipóteses rivais; (2) “o teste scree-off (sob a suposição de onisciência lógica) não é capaz de capturar qualquer significância confirmatória que a explicatividade de H possa ter quando E é uma necessidade lógica” (p. 5); (3) nos casos em que E implica que H é uma explicação comum de um segundo *explanandum*, pace Sober & Roche, a explicatividade é evidencialmente relevante. A segunda foi explorada no tópico 3.2., e a primeira, como veremos a seguir, é melhor concebida como estando compreendida na terceira.

Sua principal objeção pode ser dividida em duas teses, uma negativa e outra positiva. Enquanto a primeira questiona a adequação de SOT, a segunda defende um sentido específico em que a explicatividade seria evidencialmente relevante. Sua conjunção dá origem a uma nova concepção de Inferência à Melhor Explicação, supostamente imune ao argumento de Sober e Roche (2013; 2019), porém, como veremos, profundamente distinta daquela defendida historicamente.

⁵¹ $(1/2)^7 = 1/128$

De acordo com a tese negativa, “a suposta distinção de Roche & Sober entre considerações explicativas serem confirmatórias em si mesmas ou apenas pela graça de opiniões de fundo⁵² é uma falsa dicotomia” (p. 2024, p. 527). Afinal, para o autor, esta é a forma por meio da qual *todos* os demais fatores que influenciam na credibilidade de uma hipótese adquirem sua relevância. Portanto, seria “injusto atribuir à IME uma conexão mágica entre considerações explicativas e confirmação” (2020, p. 39).

Apesar de apenas progressivamente ter se tornado o centro de sua resposta ao desafio (2020; 2022; 2024), a ideia já aparece em sua primeira colaboração à conversa. Lange (2017) inicia com a seguinte ilustração (p. 2):

Suponha que H seja [a hipótese de] que Jones é a pessoa que roubou a joia do cofre, O é [a observação de] que o único fio de cabelo encontrado dentro do cofre era loiro, e a informação de fundo nos diz que havia exatamente um ladrão e um fio de cabelo encontrado dentro do cofre. O conhecimento de fundo também nos diz que Jones tem cabelo loiro e que o fio de cabelo encontrado tem uma probabilidade significativa (embora não conclusiva) de ter sido deixado pelo ladrão durante o roubo (embora existam outras maneiras pelas quais o cabelo poderia ter entrado no cofre). A informação de fundo também nos diz que Jones é um suspeito, ao contrário de muitas outras pessoas com cabelo loiro - embora Jones seja um entre vários suspeitos com cabelo loiro e também haja uma probabilidade razoável de que o ladrão não esteja listado entre nossos suspeitos. O conhecimento de fundo também nos diz que, se o cabelo fosse de Jones, então Jones provavelmente seria o ladrão (já que ele o teria deixado durante o roubo); Jones não teria tido ocasião de acessar o cofre exceto para roubá-lo.

Como o fio de cabelo encontrado é da mesma cor do cabelo de Jones, o fato de ter sido encontrado no interior do cofre confere certo suporte a H, isto é, $\Pr(H|O) > \Pr(H)$. Porém, segundo Lange, o suporte conferido pela observação não é máximo, considerando que o “cabelo pode não pertencer ao ladrão e que, mesmo que pertença, o ladrão não precisa ser Jones, já que muitas outras pessoas (incluindo alguns outros suspeitos) têm cabelo loiro” (p. 2). A esse cenário, então, somos provocados a considerar o acréscimo de E: “Se Jones fosse o ladrão e o único fio de cabelo encontrado dentro do cofre fosse loiro, então o fato de Jones ser o ladrão explicaria por que o fio de cabelo encontrado no cofre é loiro” (idem). De acordo com Lange, embora as evidências permaneçam inconclusivas, E é evidencialmente relevante nesta hipótese, pois remove “a possibilidade de que, mesmo se Jones fosse o ladrão, tal cabelo não pertenceria a ele pois não foi deixado pelo [verdadeiro] bandido” (p. 3). Logo, $\Pr(H | O \& E) > \Pr(H | O)$.

O exemplo não é muito feliz. Como Sober & Roche (2019) bem ressaltaram, o caráter explicativo da hipótese estaria excluindo justamente uma hipótese na qual H é verdadeira.

⁵² No original, *background opinions*.

Afinal, H diz apenas que Jones é quem cometeu o crime, e não “Jones cometeu o crime e o cabelo encontrado é do autor do delito”. É possível, é verdade, conceber cenários em que a eliminação de uma possibilidade em que a hipótese seria verdadeira, ainda assim, aumenta sua probabilidade. Isso ocorre quando um número ainda maior de possibilidades concorrentes é eliminado. Sober e Roche (2019, p. 124-125) fornecem o seguinte exemplo: Uma carta é sorteada aleatoriamente de um baralho padrão bem embaralhado. Seja A, D e S compreendidos da seguinte maneira: A: A carta sorteada é um Ás. D: A carta sorteada é um Ouros. S: A carta sorteada é uma Espadas. Dado que $\sim(A \& D) \& \sim S$ implica $\sim(A \& D)$, segue-se que $\sim(A \& D) \& \sim S$ elimina A&D e, portanto, elimina uma possibilidade na qual D é verdadeiro. No entanto, $\sim(A \& D) \& \sim S$ aumenta a probabilidade de D de 1/4 para 12/38.

Contudo, “duvidamos que esse seja o caso no exemplo do roubo” (Sober e Roche, 2019, p. 124). Repare que é um tanto difícil conceber a ideia de um contrafactual sozinho ser capaz de eliminar alguma possibilidade. Parece que a alegada diminuição do espaço de possibilidades só seria possível com a obtenção da informação de que o fio de cabelo, de fato, pertence a Jones ou não, e não pela mera suposição.

Não obstante, em antecipação a uma potencial objeção ao exemplo⁵³, Lange (2017) ressalta que o explanacionista não está comprometido com a tese de que “a explicatividade de H *per se*, independente de quaisquer opiniões de fundo, possui relevância evidencial” (p. 4). “A descoberta do potencial explicativo de H faz (ou não) uma diferença evidencial exatamente da mesma maneira que qualquer outra descoberta: à luz das opiniões de fundo do agente” (idem), comenta.

O reconhecimento da sensibilidade da relação de suporte evidencial às suposições de fundo é provavelmente um dos poucos achados filosóficos incontroversos em Teoria da Confirmação. Este *insight*, aliás, remonta a uma das célebres controvérsias da história da Filosofia da Ciência. Em resposta a Hempel, Good (1967) argumentou que o Critério de Nicod é falso, pois a observação de um caso de uma hipótese mais geral é capaz de, pelo contrário, desconfirmá-la. O contraexemplo, apesar de bastante imaginativo, ilustra muito bem esse ponto. O famoso estatístico nos exorta a supor “que nós sabemos estar em um de dois mundos” (p. 322). Num deles, há 100 corvos pretos, nenhum corvo que não seja preto e 1 milhão de outros pássaros; no segundo, 1.000 corvos pretos, 1 corvo branco e também 1 milhão de outros pássaros. “Um pássaro é selecionado de forma equiprovável do mundo em que nos

⁵³ De que o crédito seria da informação provida, e não da explicatividade (p. 3).

encontramos, e se trata de um corvo preto” (idem). Tal observação, ironicamente, provê evidência robusta de que estamos no segundo mundo, onde nem todos corvos são pretos⁵⁴.

Perceba, no entanto, que tal entendimento não torna a tese defendida por Lange automaticamente plausível. Primeiramente, porque ele não está defendendo apenas que o “conhecimento de fundo pode fazer uma diferença crucial para a confirmação” (Earman, 1992, p. 67). Ela é mais forte: o fato de a explicatividade adquirir sua relevância somente por intermédio das suposições de fundo não seria um problema porque este é supostamente o meio pelo qual todas as demais “descobertas” que influenciam a relação de suporte evidencial obtêm sua relevância. E em segundo lugar, porque, ainda que este seja o caso, é possível, *prima facie*, que sejam sempre as suposições de fundo em si, independentemente de seu eventual caráter explanatório, que possuem relevância evidencial. Esta é, aliás, a conclusão a que chegou Norton (2021), o principal defensor de que a relação de suporte seria material *through and through*: “quanto mais de perto examinamos um exemplo [dado em defesa da IME], menos importante se torna o papel da explicação como uma noção distinta” (p. 267-268).

A resposta a estas preocupações é fornecida naquilo a que nos referimos anteriormente como sendo sua tese positiva. “Por que dar qualquer destaque especial, em uma teoria da confirmação, ao fato de H ser a ‘melhor (potencial) explicação’ de O?”, Lange (2020, p. 39) pergunta de forma retórica. A resposta, segundo o autor, é “que opiniões prévias sobre o tipo de explicação que algum fato conhecido provavelmente possui são um tipo importante de opinião prévia” (idem). Lange elabora (2020, p. 40):

Essa abordagem à Inferência à Melhor Explicação (IME) não associa [a qualidade de uma explicação] a nenhuma lista de “virtudes explicativas” específica. Nossa experiência prévia com as explicações de vários fenômenos que consideramos provavelmente semelhantes a O em termos explicativos pode nos levar a esperar que O tenha uma explicação desunificada e complicada, envolvendo várias coincidências — se esses forem os tipos de explicações que esses outros fenômenos tendem a ter. Nesse caso, a IME tenderia a promover a atribuição de maior credibilidade a hipóteses que, se verdadeiras, forneceriam a O explicações com essas características “não parsimônicas”.

O artigo veicula alguns exemplos, mas Dellsén (2024) contém um ainda mais claro. Quando um de nossos dispositivos eletrônicos domésticos para de funcionar, geralmente é

⁵⁴ A insuficiente apreciação desse ponto, hoje sabemos, é o problema fundamental com os critérios de adequação propostos por Hempel (1945a; 1945b). Como Eagle (2023, p. 19) bem coloca, “Esses contraexemplos às condições de Hempel revelam um problema mais fundamental. A resposta mais óbvia a eles não é rejeitar, por exemplo, a confirmação de instância ou a Condição de Consequência, etc., mas notar que elas se aplicam em alguns casos e não em outros”.

porque um único de seus componentes está com defeito. Em razão disso, se o controle remoto da televisão hoje parar de funcionar, haverá razões para crer que um único de seus componentes (o esgotamento das pilhas, o rompimento de um fio etc) é que lhe deu causa.

Uma das vantagens da proposta é que ela parece acomodar muito bem o fato amplamente conhecido (Salmon, 2001, p. 87) de que as características que tendem a ser mais conducentes à verdade não são insensíveis ao domínio de investigação. Nas ciências históricas, por exemplo, explicações mais simples e unificadoras correm um maior risco de serem uma generalização superficial, em vez de uma explicação adequada, na medida em que acontecimentos históricos, como regra, costumam ser o produto de um conjunto de diversos fatores. Lange, portanto, parece ter encontrado um papel evidencialmente relevante para considerações de cunho explanatório.

Apesar disso, talvez seja uma vitória com a qual poucos explanacionistas consigam vibrar. A forma como o modelo de Inferência à Melhor Explicação resultante – se ainda faz sentido utilizar o mesmo rótulo – se distancia dos contornos gerais da IME identificados no segundo capítulo é notável. À parte de não haver uma lista de virtudes explicativas veroconducentes, nesta concepção “a ‘melhor explicação’ de um fato *não precisa* ser a hipótese mais plausível, considerando todos os aspectos” (Lange, 2022, p. 87) (grifei). O autor continua (idem):

Mesmo reconhecendo o suporte conferido a uma dada hipótese por sua capacidade (se verdadeira) de fornecer uma boa explicação para algum fato, podemos justamente considerar uma hipótese concorrente que não seria a ‘melhor explicação’ desse fato como ainda assim melhor suportada pelo corpo completo de evidências disponíveis.

Lange (p. 88) destaca que, de acordo com Lipton, a qualidade de uma explicação seria um guia para inferência, mas não o único. Conquanto isso seja verdade, é mais controverso que Lipton concordaria que a melhor explicação pode não ser a mais plausível e, mais importante, se isso não põe em xeque todo o valor da atribuído a essa forma de inferência – sem falar no perigo de um deslizamento para o infalibilismo e, por consequência, para uma forma desinteressante de ceticismo⁵⁵.

Ainda mais fundamental talvez seja o problema de que, apesar de a semelhança entre os *explananda* poder ser identificada como uma consideração de natureza explicativa (por sugerir uma semelhança quanto ao tipo de explicação), a proposta não salva a explicatividade do argumento *Screen-off*. Conforme frisado ao longo deste texto, na equação $\Pr(H \mid O \ \& \ E) = \Pr$

⁵⁵ Agradeço ao meu orientador por me apontar esse ponto.

($H \mid O$), E corresponde à proposição “ H , se verdadeira, explicaria O ”. Para Lange, contudo, E parece ser equivalente a algo como “ H , se verdadeira, forneceria uma explicação do mesmo tipo da que aceitamos para outros *explananda* semelhantes a O ”.

Observe-se também que a proposta acaba por mitigar pelo menos parte das grandes aspirações que antes sopravam as velas do Explanacionismo. Ao se restringir a relevância epistêmica de considerações explanatórias ao papel indicado por Lange, parece ser necessário, por exemplo, abandonar a esperança de elas virem a prover uma resolução satisfatória do problema da indução. Mais do que isso. Nesta nova perspectiva, este último se torna ainda mais evidente. Afinal, por que o fato de O_1 ter uma explicação do tipo X torna mais provável que O_2 também possua? – poderá insistir o cético.

Dessa forma, conclui-se que, apesar de Lange (2020; 2024) ter identificado um papel evidencialmente relevante para considerações de natureza explicativa em sentido amplo, a explicatividade - da forma como definida por Sober & Roche (2013; 2019) - permanece vulnerável ao argumento *Screen-off*.

4. CONSIDERAÇÕES FINAIS

Para melhor compreensão do impacto do argumento *Screen-off*, iniciamos nossa investigação buscando identificar os contornos gerais da Inferência à Melhor Explicação. Nesse processo, verificamos que, em cada uma das etapas em que (para fins didáticos) é possível fracioná-la, a IME enfrenta desafios notáveis, muito embora também não faltem propostas de como contorná-los.

Ao nos debruçarmos sobre o debate provocado por Sober & Roche (2013), observamos que a explicatividade é um alvo em movimento. Apesar de cuidadosamente definida ainda no início da discussão, notamos que as respostas existentes ao argumento *Screen-off* frequentemente escondem outras noções disfarçadas. Quando pensávamos enfim ter encontrado um exemplo incontroverso de sua relevância evidencial, ao olhar mais de perto, logo percebíamos que o truque era feito, na verdade, por algo escondido na manga do explanacionista, seja a relação de implicação (McCain e Poston, 2014), a relação de causalidade (Climenhaga, 2017) ou alguma suposição de fundo mais específica (Lange, 2017). Nenhuma delas facilmente capturada pela ideia de que “obtemos evidência de que p ao percebermos que, se verdadeira, ela explicaria a evidência disponível” (McCain e Poston, 2024, p. 329).

Dito isso, todas as respostas – e inclusive o próprio argumento *Screen-off* – testificam de algo mais fundamental: a importância das suposições subjacentes a nossos argumentos. E tal constatação sugere uma linha de pesquisa aparentemente promissora. Trata-se, a propósito, de

um tema muito caro a Sober & Roche⁵⁶. Em *Evidence and Evolution* (2008), por exemplo, o primeiro defende que proponentes do Design Inteligente (DI) precisam de fundamentos independentes dos fatos biológicos sendo explicados para as suposições de que o designer desejaria criar os organismos com as exatas características que observamos. Contudo, razoável que tal demanda nos soe, traçar uma linha entre onde a hipótese termina e onde começam as suposições auxiliares (que dependem de justificação independente) não é uma tarefa fácil. Como Samson (2011) bem notou, nada impede o entusiasta do Design Inteligente de avançar uma teoria da qual as suposições necessárias já façam parte. Nesse caso, “eles evitam a principal objeção de Sober de que dependem de hipóteses auxiliares não defendidas” (p. 675).

A identificação de critérios por meio dos quais possamos distinguir uma suposição aceitável de uma suposição inaceitável, portanto, “é uma questão interessante e que não foi suficientemente investigada” (Samson, 2011, p. 677)⁵⁷. E sua relevância, é importante destacar, vai muito além de lançar alguma luz sobre como melhor responder a objeções criacionistas. Da razoabilidade de se postular a existência de múltiplos universos (Friederich, 2021, p. 75-78) ao teste da hipótese de sobrevivência *post-mortem* (Sudduth, 2016, p. 214-244), a determinação das suposições auxiliares adequadas em cada caso é crucial. Como Jantzen (2014, p. 310) coloca, “se a Filosofia é sobre alguma coisa, ela é sobre obter clareza sobre nossas suposições”⁵⁸.

O debate em torno do desafio *Screen-off* provavelmente ainda não chegou ao fim, como sugere o fato de as últimas publicações serem bastante recentes. Talvez ainda vejamos alguma teoria específica sobre como a explicatividade tornaria a probabilidade de uma hipótese mais resiliente. Ou talvez a explicatividade passe a ser concebida na literatura apenas da forma como conceptualizada por Lange. Uma coisa, porém, é certa: à vista de todas as concessões feitas pelo partido explanacionista, como Dorothy Gale comenta em *The Wizard of Oz*, “I’ve a feeling we’re not in Kansas anymore”.

⁵⁶ Para alguns exemplos, vide Sober (1996; 1999; 2000; 2018).

⁵⁷ No original, “an interesting and generally under investigated issue”.

⁵⁸ If philosophy is about anything, it is about getting clear on one’s assumptions.

5. REFERÊNCIAS

ACHINSTEIN, Peter. Positive Relevance. In: The Routledge Handbook of the Philosophy of Evidence. Routledge, 2024. p. 104-112.

ACHINSTEIN, Peter. Positive Relevance. In: **The Routledge Handbook of the Philosophy of Evidence**. Routledge, p. 329-341, 2024

ACHINSTEIN, Peter. The Concept of Evidence. Oxford: Oxford University Press, 1983.

ACHINSTEIN, Peter. **The book of evidence**. Oxford University Press, 2001.

ALAI, Mario. Scientific realism, metaphysical antirealism and the no miracle arguments. **Foundations of Science**, v. 28, n. 1, p. 377-400, 2023.

ALLEN, Ronald J. Legal probabilism—a qualified rejection: A response to Hedden and Colyvan. Available at SSRN 3355969, 2019.

ATKINS, Richard Kenneth. Peirce on inference: Validity, strength, and the community of inquirers. Oxford University Press, 2023.

BAILEY, Storm McClintock. **Inference to the best explanation and justification in ethics**. The University of Wisconsin-Madison, 1997.

BAKER, Alan. Simplicity. 'Stanford Encyclopedia of Philosophy. 2022.

BARAS, Dan. A strike against a striking principle. *Philosophical Studies*, v. 177, n. 6, p. 1501-1514, 2020b.

BARAS, Dan. Calling for explanation. Oxford University Press, 2022.

BARAS, Dan. Calling for explanation: An extraordinary account. 2020a.

BARAS, Dan. How can necessary facts call for explanation?. *Synthese*, v. 198, n. 12, p. 11607-11624, 2021.

BEEBE, James R. The abductivist reply to skepticism. **Philosophy and Phenomenological Research**, v. 79, n. 3, p. 605-636, 2009.

BENI, Majid D. Casting inference to the best explanation's lot with active inference. **Theoria**, v. 89, n. 2, p. 188-203, 2023.

BERNABEU, Carmen Batanero; FERNANDES, José António; GARCÍA, José Miguel Contreras. Un análisis semiótico del problema de Monty Hall e implicaciones didácticas. **Suma: Revista sobre Enseñanza y Aprendizaje de las Matemáticas**, n. 62, p. 11-18, 2009.

BHOGAL, Harjit. Strikingness: When Facts Call Out for Explanation. 2022.

BINNEY, Nicholas. Methods of Inference and Shaken Baby Syndrome. *Philosophy of Medicine*, v. 4, n. 1, p. 1-33, 2023.

BIRD, Alexander. Abductive knowledge and Holmesian inference. **Oxford studies in epistemology**, v. 1, p. 1-31, 2005.

BIRD, Alexander. Eliminative abduction: Examples from medicine. **Studies in History and Philosophy of Science Part A**, v. 41, n. 4, p. 345-352, 2010.

BIRD, Alexander. Inference to the best explanation, Bayesianism, and knowledge. **Best explanations: New essays on inference to the best explanation**, p. 97-120, 2017.

BIRD, Alexander. **Knowing science**. Oxford University Press, 2022.

BORSBOOM, Denny et al. Theory construction methodology: A practical framework for building theories in psychology. **Perspectives on Psychological Science**, v. 16, n. 4, p. 756-766, 2021.

CABRERA, Frank. Inference to the Best Explanation: An Overview. In: MAGNANI, Lorenzo (Ed.). **Handbook of Abductive Cognition**. Springer, p. 1863-1896, 2023.

CAMPOS, Daniel G. On the distinction between Peirce's abduction and Lipton's inference to the best explanation. **Synthese**, v. 180, p. 419-442, 2011.

CARRUTHERS, Peter. **The centered mind: What the science of working memory shows us about the nature of human thought**. OUP Oxford, 2015.

CLIMENHAGA, Nevin. How explanation guides confirmation. **Philosophy of Science**, v. 84, n. 2, p. 359-368, 2017.

CLIMENHAGA, Nevin. The structure of epistemic probabilities. *Philosophical Studies*, v. 177, n. 11, p. 3213-3242, 2020.

CZÉGEL, Dániel et al. Bayes and Darwin: How replicator populations implement Bayesian computations. **Bioessays**, v. 44, n. 4, p. 2100255, 2022.

DANIEL, GARBER. Old Evidence and Logical Omniscience in Bayesian Confirmation Theory. **Testing Scientific Theories**, 1983.****

DARWIN, Charles. *The Origin of Species*, 6th Edition. New York: Collier, 1962 [1872].

DAVEY, Kevin. There Are No Bad Lots, Only Bad Formulations of Inference to the Best Explanation. 2023.****

DAWES, Gregory W. **Theism and explanation**. Routledge, 2012.

DAY, Timothy; KINCAID, Harold. Putting inference to the best explanation in its place. *Synthese*, v. 98, p. 271-295, 1994.

DELLSÉN, Finnur. Explanatory consolidation: From ‘best’ to ‘good enough’. **Philosophy and Phenomenological Research**, v. 103, n. 1, p. 157-177, 2021.

DELLSÉN, Finnur. **Abductive Reasoning in Science**. Cambridge University Press, 2024.

DESCARTES, René. **Principles of Philosophy: Translated, with Explanatory Notes**. Springer Science & Business Media, 1984 [1644].

DEVITO, Scott. A gruesome problem for the curve-fitting solution. **The British Journal for the Philosophy of Science**, v. 48, n. 3, p. 391-396, 1997.

DOUVEN, Igor. Abduction. In Zalta, E. N., editor, *Stanford Encyclopedia of Philosophy*, 2021: <https://plato.stanford.edu/archives/sum2021/entries/abduction/>.

DOUVEN, Igor. Inference to the best explanation, Dutch books, and inaccuracy minimisation. **The Philosophical Quarterly**, v. 63, n. 252, p. 428-444, 2013.

DOUVEN, Igor. Inference to the best explanation: What is it? And why should we care. **Best explanations: New essays on inference to the best explanation**, p. 4-22, 2017.

DOUVEN, Igor. **The art of abduction**. MIT press, 2022.

EAGLE, Antony. *Probability and Inductive Logic*. 2023.

EARMAN, John. Bayes or bust?: A critical examination of Bayesian confirmation theory. 1992.

EINSTEIN, Albert. Explanation of the Perihelion Motion of Mercury from the General Theory of Relativity. *The Collected Papers of Albert Einstein*. Princeton University, v. 6, p. 112-116, 1997.

ELLIOTT, Katrina. Inference to the best explanation and the new size elitism¹. 2021.

FANN, Kuang Tih. **Peirce’s theory of abduction**. Springer Science & Business Media, 2012.

FISHER, Sir Ronald Aylmer. **Smoking: the cancer controversy: some attempts to assess the evidence**. Edinburgh: Oliver and Boyd, 1959.

FOGELIN, Lars. Inference to the best explanation: A common and effective form of archaeological reasoning. **American antiquity**, v. 72, n. 4, p. 603-626, 2007.

FORSTER, Malcolm R. Model selection in science: The problem of language variance. **The British journal for the philosophy of science**, v. 50, n. 1, p. 83-102, 1999.

FORSTER, Malcolm; SOBER, Elliott. How to tell when simpler, more unified, or less ad hoc theories will provide more accurate predictions. **The British Journal for the Philosophy of Science**, v. 45, n. 1, p. 1-35, 1994.

FRIEDERICH, Simon. **Multiverse theories: A philosophical perspective**. Cambridge University Press, 2021.

FUMERTON, Richard A. Induction and reasoning to the best explanation. *Philosophy of Science*, v. 47, n. 4, p. 589-600, 1980.

FUMERTON, Richard. Reasoning to the best explanation. *Best explanations: New essays on inference to the best explanation*, p. 65-69, 2017.

GIGERENZER, Gerd; SELTEN, Reinhard (Ed.). **Bounded rationality: The adaptive toolbox**. MIT press, 2002.

GOOD, Irving John. The white shoe is a red herring. **The British Journal for the Philosophy of Science**, v. 17, n. 4, p. 322-322, 1967.

GOULD, Stephen Jay; LEWONTIN, Richard C. The spandrels of san marco and the panglossian paradigm: a critique of the adaptationist programme. **Conceptual Issues in Evolutionary Biology**, v. 205, p. 79, 1979.

HAIG, Brian D. Inference to the best explanation: A neglected approach to theory appraisal in psychology. **The American journal of psychology**, v. 122, n. 2, p. 219-234, 2009.

HAIG, Brian D.; HAIG, Brian D. An abductive theory of scientific method. **Method matters in psychology: Essays in applied philosophy of science**, p. 35-64, 2018.

HÁJEK, Alan; JOYCE, James M. Confirmation. In: **The Routledge Companion to Philosophy of Science**. Routledge, 2013. p. 146-159.

HANSON, Norwood Russell. The logic of discovery. **The Journal of Philosophy**, v. 55, n. 25, p. 1073-1089, 1958.

HARMAN, Gilbert. *Change in View*. Cambridge, MA: MIT Press.

HARMAN, Gilbert. The inference to the best explanation. *The philosophical review*, v. 74, n. 1, p. 88-95, 1965.

HAWTHORNE, James. Bayesian induction is eliminative induction. **Philosophical Topics**, v. 21, n. 1, p. 99-138, 1993.

HEMPEL, Carl G. Studies in the Logic of Confirmation (I.). *Mind*, v. 54, n. 213, p. 1-26, 1945.

HEMPEL, Carl G. Studies in the Logic of Confirmation (II.). *Mind*, v. 54, n. 214, p. 97-121, 1945.

HENDERSON, Leah. Bayesianism and Inference to the Best Explanation. *The British Journal for the Philosophy of Science*, p. 687-715, 2014.

HOLTON, Richard. "Intention as a Model for Belief." In *Rational and Social Agency: Essays on the Philosophy of Michael Bratman*, edited by Manuel Vargas and Gideon Yaffe. Oxford: Oxford University Press, 2014.

HORGAN, Terry. Troubles for Bayesian formal epistemology. **Res Philosophica**, v. 94, n. 2, p. 233-255, 2017.

HUEMER, Michael. Explanationist Aid for the Theory of Inductive Logic. **The British journal for the philosophy of science**, v. 60, n. 2, p. 345-375, 2009.

HUEMER, Michael. Serious theories and skeptical theories: Why you are probably not a brain in a vat. **Philosophical Studies**, v. 173, p. 1031-1052, 2016.

ISAACS, Yoaav; HAWTHORNE, John; SANFORD RUSSELL, Jeffrey. Multiple universes and self-locating evidence. *Philosophical Review*, v. 131, n. 3, p. 241-294, 2022.

JANTZEN, Benjamin C. **An introduction to design arguments**. Cambridge University Press, 2014.

JEFFREYS, Harold. *Theory of Probability*: Oxford University Press, 1961.

JOHNSTON, Angie M. et al. Little Bayesians or little Einsteins? Probability and explanatory virtue in children's inferences. **Developmental science**, v. 20, n. 6, p. e12483, 2017.

JOSEPHSON, John R.; JOSEPHSON, Susan G. (Ed.). **Abductive inference: Computation, philosophy, technology**. Cambridge University Press, 1996.

KEAS, Michael N. Systematizing the theoretical virtues. *Synthese*, v. 195, n. 6, p. 2761-2793, 2018.

KEYNES, John Maynard. **A treatise on probability**. Courier Corporation, 2013.

KHEMLANI, Sangeet S.; SUSSMAN, Abigail B.; OPPENHEIMER, Daniel M. Harry Potter and the sorcerer's scope: latent scope biases in explanatory reasoning. **Memory & cognition**, v. 39, p. 527-535, 2011.

KOLMOGOROV, Andrei Nikolaevich. The theory of probability. **Mathematics, Its Content, Methods, and Meaning**, v. 2, p. 110-118, 1963.

KUIPERS, Theo AF. **From instrumentalism to constructive realism: On some relations between confirmation, empirical progress, and truth approximation**. Springer Science & Business Media, 2013.

LANGE, Marc. A false dichotomy in denying explanatoriness any role in confirmation. **Noûs**, v. 58, n. 2, p. 520-535, 2024.

LANGE, Marc. *Because without cause: Non-casual explanations in science and mathematics*. Oxford University Press, 2016.

LANGE, Marc. How Simplicity Can be a Virtue in Philosophical Theory-Choice. **Erkenntnis**, v. 89, n. 3, p. 1217-1234, 2024.

LANGE, Marc. Putting explanation back into “inference to the best explanation”. **Noûs**, v. 56, n. 1, p. 84-109, 2022.

LANGE, Marc. The evidential relevance of explanatoriness: A reply to Roche and Sober. **Analysis**, v. 77, n. 2, p. 303-312, 2017.

LANGE, Marc. What Inference to the Best Explanation Is Not. **Teorema: Revista Internacional de Filosofía**, v. 39, n. 2, p. 27-42, 2020.

LAUDAN, Larry. **Science and values: The aims of science and their role in scientific debate**. Univ of California Press, 1984.

LEIBNIZ, Gottfried Wilhelm. Discourse on Metaphysics. In D. Garber and R. Ariew (eds. and trans.), *Discourse on Metaphysics and Other Essays*. Indianapolis: Hackett, p. 1–40, 1991.

LEITER, Brian. Moral skepticism and moral disagreement in Nietzsche. **Oxford studies in metaethics**, v. 9, p. 126-151, 2014.

LIN, Hanti. Bayesian epistemology. *Stanford Encyclopedia of Philosophy*, 2022.

LIPTON, Peter. *Inference to the Best Explanation*. 2 ed. London: Routledge, 2004.

LIPTON, Peter. Is explanation a guide to inference? A reply to Wesley C. Salmon. In: *Explanation: Theoretical approaches and applications*. Dordrecht: Springer Netherlands, p. 93-120, 2001.

LOMBROZO, Tania. Simplicity and probability in causal explanation. **Cognitive psychology**, v. 55, n. 3, p. 232-257, 2007.

LYCAN, William. Explanation and epistemology. In: Moser, P. (Ed.), *The Oxford Handbook of Epistemology*. Oxford: Oxford University Press, p. 408-433, 2002.

MAYO, Deborah G. **Statistical inference as severe testing: How to get beyond the statistics wars**. Cambridge University Press, 2018.

MCAULIFFE, William HB. How did abduction get confused with inference to the best explanation?. *Transactions of the Charles S. Peirce Society: A Quarterly Journal in American Philosophy*, v. 51, n. 3, p. 300-319, 2015.

MCCAIN, Kevin. In defense of the explanationist response to skepticism. **International Journal for the Study of Skepticism**, v. 9, n. 1, p. 38-50, 2019.

MCCAIN, Kevin. Skepticism and Elegance: An Explanationist Rejoinder. *International Journal for the Study of Skepticism*, v. 6, n. 1, p. 30-43, 2016.

MCCAIN, Kevin; POSTON, Ted. Explanation and Evidence. In: **The Routledge Handbook of the Philosophy of Evidence**. Routledge, 2024. p. 329-341.

MCCAIN, Kevin; POSTON, Ted. The evidential impact of explanatory considerations. In: **Best explanations: New essays on inference to the best explanation**, v. 121, p. 129, 2017.

MCCAIN, Kevin; POSTON, Ted. Why explanatoriness is evidentially relevant. **Thought: A Journal of Philosophy**, v. 3, n. 2, p. 145-153, 2014.

MCGREW, Timothy. Confirmation, heuristics, and explanatory reasoning. *The British Journal for the Philosophy of Science*, v. 54, n. 4, p. 553-567, 2003.

MCLAUGHLIN, Brian P. Consciousness, type physicalism, and inference to the best explanation. **Philosophical Issues**, v. 20, p. 266-304, 2010.

MILL, John Stuart. A system of logic ratiocinative and inductive. Ed. JM Robson. **Collected Works**, v. 7, 1974.

MILL, John Stuart. A system of logic, ratiocinative and inductive, being a connected view of the principles of evidence and the methods of scientific investigation, 2 vols. London: John W. Parker, 1843.

MILLS, Rónán; FROWLEY, Jason. Promiscuous teleology and the effect of locus of control. **The Irish Journal of Psychology**, v. 35, n. 2-3, p. 121-132, 2014.

MINNAMEIER, Gerhard. Peirce-suit of truth—Why inference to the best explanation and abduction ought not to be confused. *Erkenntnis*, v. 60, n. 1, p. 75-105, 2004.

MOHAMMADIAN, Mousa. Abduction— the context of discovery+ underdetermination= inference to the best explanation. **Synthese**, v. 198, p. 4205-4228, 2021.

MORRIS, Douglas W.; LUNDBERG, Per. Pillars of evolution: fundamental principles of the eco-evolutionary process. OUP Oxford, 2011.

MUSGRAVE, Alan. The ultimate argument for scientific realism. In: **Relativism and realism in science**. Dordrecht: Springer Netherlands, 1988. p. 229-252.

NIINILUOTO, Ilkka. Truthlikeness: Old and new debates. **Synthese**, v. 197, n. 4, p. 1581-1599, 2020.

NIINILUOTO, Ilkka. On the truthlikeness of generalizations. In: **Basic Problems in Methodology and Linguistics: Part Three of the Proceedings of the Fifth International Congress of Logic, Methodology and Philosophy of Science, London, Ontario, Canada-1975**. Dordrecht: Springer Netherlands, 1977. p. 121-147.

NIINILUOTO, Ilkka. Reference invariance and truthlikeness. **Philosophy of Science**, v. 64, n. 4, p. 546-554, 1997.

NIINILUOTO, Ilkka. Truthlikeness and Bayesian estimation. **Synthese**, p. 321-346, 1986.

NIINILUOTO, Ilkka. Truthlikeness: Old and new debates. **Synthese**, v. 197, n. 4, p. 1581-1599, 2020.

NIINILUOTO, Ilkka. Truth-seeking by abduction. Springer, 2018.

NIINILUOTO, Ilkka. **Critical scientific realism**. OUP Oxford, 1999.

NIINILUOTO, Ilkka. **Truthlikeness**. Springer Science & Business Media, 2012.

NORTON, John D. The material theory of induction. University of Calgary Press, 2021.

OKASHA, Samir. Van Fraassen's critique of inference to the best explanation. *Studies in History and Philosophy of Science Part A*, v. 31, n. 4, p. 691-710, 2000.

OKASHA, Samir. Van Fraassen's critique of inference to the best explanation. **Studies in History and Philosophy of Science Part A**, v. 31, n. 4, p. 691-710, 2000.

PARK, Seungbae. The Supremacy of IBE over Bayesian Conditionalization. *Problemos*, n. 103, p. 66-76, 2023.

PEIRCE, Charles Sanders. The Essential Peirce, Volume 2: Selected Philosophical Writings (1893-1913). Indiana University Press, 1992.

PEIRCE, Charles Sanders. The fixation of belief. *Writings of Charles S. Peirce: A chronological edition*, 3, 242-257, 1986.

PEIRCE, Charles Sanders. *Writings of Charles Sanders Peirce. A Chronological Edition*, v. 1, 1982.

PEIRCE, Charles Sanders. **Collected papers of charles sanders peirce**. Harvard University Press, 1974.

PEIRCE, Charles Sanders. **Pragmatism as a principle and method of right thinking: The 1903 Harvard lectures on pragmatism**. Suny Press, 1997.

PFISTER, Rolf. Towards a theory of abduction based on conditionals. *Synthese*, v. 200, n. 3, p. 206, 2022.

PHILIPS, R.; GUTTMAN, I. A new criterion for variable selection. **Statistics & probability letters**, v. 38, n. 1, p. 11-19, 1998.

POLLOCK, John L. **Thinking about acting: Logical foundations for rational decision making**. Oxford University Press, 2006.

POPPER, K. (2005). The logic of scientific discovery. Routledge.

POPPER, Karl. **The logic of scientific discovery**. Routledge, 2005.

PRASETYA, Yunus. Towards a Synthesis of Two Research Programmes: Inference to the Best Explanation and Models of Scientific Explanation. **Australasian Journal of Philosophy**, v. 101, n. 3, p. 750-764, 2023.

PRASETYA, Yunus. Which Models of Scientific Explanation Are (In) Compatible with Inference to the Best Explanation?. *The British Journal for the Philosophy of Science*, v. 75, n. 1, p. 209-232, 2024.

PSILLOS, Stathis. Simply the best: A case for abduction. In: *Computational logic: Logic programming and beyond: Essays in honour of Robert A. Kowalski part II*. Berlin, Heidelberg: Springer Berlin Heidelberg, 2002. p. 605-625.

PSILLOS, Stathis. **Scientific realism: How science tracks truth**. Routledge, 2005.

QUINE, Willard V. On empirically equivalent systems of the world. In: **The Nature of Scientific Theory**. Routledge, p. 353-368, 2014.

QUINE, Willard Van Orman. *From Stimulus to Science*. Cambridge: Harvard University Press, 1995.

QUINE, Willard Van Orman; ULLIAN, Joseph Silbert. **The web of belief**. New York: Random House, 1978.

RASMUSSEN, Joshua; LEON, Felipe. **Is god the best explanation of things?**. Springer International Publishing, 2019.

REICHENBACH, Hans. *Experience and prediction: An analysis of the foundations and the structure of knowledge*. 1938.

RISINGER, D. Michael. Off the Rails: The Surprising Story of Smith v. Rapid Transit, Inc. **Seton Hall Journal of Legislation and Public Policy**, v. 48, n. 3, p. 6, 2024.

ROCHE, William; SOBER, Elliott. Explanatoriness and evidence: A reply to McCain and Poston. **Thought: A Journal of Philosophy**, v. 3, n. 3, p. 193-199, 2014.

ROCHE, William; SOBER, Elliott. Explanatoriness is evidentially irrelevant, or inference to the best explanation meets Bayesian confirmation theory. *Analysis*, v. 73, n. 4, p. 659-668, 2013.

ROCHE, William; SOBER, Elliott. Inference to the best explanation and the screening-off challenge. **Teorema: Revista Internacional de Filosofía**, v. 38, n. 3, p. 121-142, 2019.

ROCHE, William; SOBER, Elliott. Is explanatoriness a guide to confirmation? A reply to Climenhaga. **Journal for General Philosophy of Science**, v. 48, p. 581-590, 2017.

SALMON, Wesley C. Explanation and confirmation: A Bayesian critique of inference to the best explanation. In: *Explanation: Theoretical approaches and applications*. Dordrecht: Springer Netherlands, 2001. p. 61-91.

SANSOM, Roger. Auxiliary Hypotheses in Evidence and Evolution. **Philosophy and Phenomenological Research**, v. 83, n. 3, 2011.

SAVAGE, Leonard J. Difficulties in the theory of personal probability. **Philosophy of Science**, v. 34, n. 4, p. 305-310, 1967.

SCHUPBACH, Jonah N. Inference to the best explanation, cleaned up and made respectable. **Best explanations: New essays on inference to the best explanation**, p. 39-61, 2017.

SCHUPBACH, Jonah N. Is the bad lot objection just misguided?. *Erkenntnis*, v. 79, p. 55-64, 2014.

SCHURZ, Gerhard. Patterns of abductive inference. **Springer handbook of model-based science**, p. 151-173, 2017.

SCHURZ, Gerhard. Unification and explanation: Explanation as a prototype concept. A reply to Weber and van Dyck, Gijsbers, and de Regt. *THEORIA. Revista de Teoría, Historia y Fundamentos de la Ciencia*, v. 29, n. 1, p. 57-70, 2014.

SEMMELEWEIS, Ignaz Philipp. **Die aetiologie, der begriff und die prophylaxis des kindbettfiebers**. Hartleben, 1861.

SHOGENJI, Tomoji. **Formal epistemology and Cartesian skepticism: In defense of belief in the natural world**. Routledge, 2017.

SHORTER, Edward. K. Codell Carter (translator and editor), Ignaz Semmelweis: The etiology, concept, and prophylaxis of childbed fever, Madison and London, University of Wisconsin Press, **Medical history**, v. 28, n. 3, p. 334-334, 1984.

SKYRMS, Brian. Resiliency, propensities, and causal necessity. **The Journal of Philosophy**, v. 74, n. 11, p. 704-713, 1977.

SOBER, Elliott. Evolution and the problem of other minds. **The Journal of Philosophy**, v. 97, n. 7, p. 365-386, 2000.

SOBER, Elliott. *Ockham's razors*. Cambridge University Press, 2015.

SOBER, Elliott. Parsimony and predictive equivalence. **Erkenntnis**, v. 44, n. 2, p. 167-197, 1996.

SOBER, Elliott. Parsimony and predictive equivalence. **Erkenntnis**, v. 44, n. 2, p. 167-197, 1996.

SOBER, Elliott. Parsimony arguments in science and metaphysics, and their connection with unification, fundamentality, and epistemological holism. In: **Levels of Reality in Science and Philosophy: Re-examining the Multi-level Structure of Reality**. Cham: Springer International Publishing, 2022. p. 229-260.

SOBER, Elliott. Parsimony arguments in science and philosophy—A test case for naturalism p. In: **Proceedings and addresses of the American Philosophical Association**. American Philosophical Association, 2009. p. 117-155.

SOBER, Elliott. Testability. In: **Proceedings and addresses of the American Philosophical Association**. American Philosophical Association, 1999. p. 47-76.

SOBER, Elliott. Venetian sea levels, British bread prices, and the principle of the common cause. **The British Journal for the Philosophy of Science**, v. 52, n. 2, p. 331-346, 2001.

SOBER, Elliott. **Reconstructing the past: Parsimony, evolution, and inference**. MIT press, 1991.

SOBER, Elliott. **The design argument**. Cambridge University Press, 2018.

STEELE, Katie; COLYVAN, Mark. Meta-uncertainty and the proof paradoxes. **Philosophical Studies**, v. 180, n. 7, p. 1927-1950, 2023.

SUDDUTH, Michael. **A philosophical critique of empirical arguments for postmortem survival**. Springer, 2016.

SWINBURNE, Richard. Simplicity as evidence of truth. 1997.

SWINBURNE, Richard. The Existence of God. 2 ed. Oxford: Oxford University Press, 2004.

THAGARD, Paul R. The best explanation: Criteria for theory choice. *The journal of philosophy*, v. 75, n. 2, p. 76-92, 1978.

TITELBAUM, Michael G. **Fundamentals of Bayesian epistemology 1: Introducing credences**. Oxford University Press, 2022.

TSUKAMURA, Yuki et al. How does the latent scope bias occur?: Cognitive modeling for the probabilistic reasoning process of causal explanations under uncertainty. In: **Proceedings of the Annual Meeting of the Cognitive Science Society**. 2022.

TVERSKY, Amos; KAHNEMAN, Daniel. Availability: A heuristic for judging frequency and probability. **Cognitive psychology**, v. 5, n. 2, p. 207-232, 1973.

URBANIAK, Rafał; DI BELLO, Marcello. Legal Probabilism. **Stanford Encyclopedia of Philosophy**, 2021.

VAN FRAASSEN, Bas C. Laws and symmetry. Oxford: Clarendon Press, 1989.

VAN FRAASSEN, Bas C. **The scientific image**. Oxford University Press, 1980.

VOGEL, Jonathan. Cartesian skepticism and inference to the best explanation. *Journal of Philosophy*, 87, 658–666, 1990.

VOGEL, Jonathan. Internalist Responses to Skepticism. In Greco, J. (ed.), *The Oxford Handbook of Skepticism*, 533–556, 2008.

VOGEL, Jonathan. The refutation of skepticism. In Matthias Steup & John Turri (eds.), *Contemporary Debates in Epistemology*. Chichester, West Sussex, UK: Blackwell. p. 72—84, 2013.

WARREN, Jared. The Sense-Data Language and External World Skepticism. forthcoming.

WEINTRAUB, Ruth. Induction and inference to the best explanation. **Philosophical Studies**, v. 166, p. 203-216, 2013.

WEINTRAUB, Ruth. Scepticism about inference to the best explanation. **Best Explanations: New Essays on Inference to the Best Explanation**, v. 188, 2017.

WEISBERG, Jonathan. Locating IBE in the Bayesian framework. **Synthese**, v. 167, n. 1, p. 125-143, 2009.

WHEWELL, W. The Philosophical Foundations of the Inductive Sciences, Founded on Their History, Vol. I-2. **Cambridge University Press, London**, v. 43, p. 0-8493, 1847.

WHEELER, John Archibald. How Come the Quantum? a. **Annals of the New York Academy of Sciences**, v. 480, n. 1, p. 304-316, 1986.

WILCOX, John E. Likelihood neglect bias and the mental simulations approach: An illustration using the old and new Monty Hall problems. **Judgment and Decision Making**, v. 19, p. e14, 2024.

WOODWARD, James. The Place of Explanation in Scientific Inquiry: Inference to the Best Explanation vs Inference to the Only Explanation, forthcoming.



Pontifícia Universidade Católica do Rio Grande do Sul
Pró-Reitoria de Pesquisa e Pós-Graduação
Av. Ipiranga, 6681 – Prédio 1 – Térreo
Porto Alegre – RS – Brasil
Fone: (51) 3320-3513
E-mail: propesq@pucrs.br
Site: www.pucrs.br