

INFERÊNCIA DA MELHOR EXPLICAÇÃO*

Peter Lipton**

Tradução: Gabriel Chiarotti Sardi***

A ciência depende de julgamentos sobre a relação entre evidências e teorias. Os cientistas devem julgar se uma observação ou o resultado de um experimento apoia, refuta ou é simplesmente irrelevante para uma dada hipótese. De forma semelhante, os cientistas podem julgar que, dadas todas as evidências disponíveis, uma hipótese deve ser aceita como correta ou quase correta, rejeitada como falsa, ou nem uma coisa nem outra. Ocasionalmente, esses julgamentos evidenciais podem ser feitos com base em razões dedutivas. Se um resultado experimental estritamente contradiz uma hipótese, então a verdade da evidência implica dedutivamente a falsidade da hipótese. Na grande maioria dos casos, no entanto, a conexão entre evidência e hipótese é não demonstrativa ou indutiva. Em particular, isso ocorre sempre que uma hipótese geral é inferida como correta com base nos dados disponíveis, uma vez que a verdade dos dados não implica dedutivamente a verdade da hipótese. Sempre permanece a possibilidade de que a hipótese seja falsa, mesmo que os dados estejam corretos.

Um dos objetivos centrais da filosofia da ciência é fornecer uma explicação fundamentada desses julgamentos e inferências que conectam evidências à teoria. No caso dedutivo, esse projeto está bem avançado, graças a uma corrente produtiva de pesquisa sobre a estrutura do argumento dedutivo, que remonta à Antiguidade. O mesmo não pode ser dito para as inferências indutivas. Embora alguns dos problemas centrais tenham sido apresentados de forma incisiva por David Hume no século XVIII, nossa compreensão atual do raciocínio indutivo continua notavelmente precária, apesar dos intensos esforços de numerosos epistemólogos e filósofos da ciência.

* Texto originalmente publicado como um capítulo da obra *A companion to the philosophy of science* (W. H. Newton-Smith (ed). Oxford: Blackwell, 2000, p. 184-193). Tradução recebida em 15/03/2025 e aprovado para publicação em 10/09/2025.

** Peter Lipton foi professor *Hans Reusing* e chefe do Departamento de História e Filosofia da Ciência na Universidade de Cambridge, além de membro do King's College.

*** Doutor em Filosofia pela Universidade de São Paulo (USP). E-mail: gabrielchi@hotmail.com.

O modelo de Inferência da Melhor Explicação foi desenvolvido para oferecer uma explicação parcial de muitas inferências indutivas, tanto na ciência quanto na vida cotidiana. Uma versão do modelo foi desenvolvida sob o nome de “abdução” por Charles Sanders Peirce (1931) no início deste século, e o modelo foi consideravelmente desenvolvido e discutido nos últimos vinte e cinco anos. Sua ideia central é que considerações explicativas são um guia para a inferência e que os cientistas inferem a partir das evidências disponíveis a hipótese que, se correta, melhor explicaria essas evidências. Muitas inferências são naturalmente descritas dessa forma. Darwin inferiu a hipótese da seleção natural porque, embora não fosse implicada pelas suas evidências biológicas, a seleção natural forneceria a melhor explicação para essas evidências. Quando um astrônomo infere que uma estrela está se afastando da Terra com uma velocidade específica, ele faz isso porque o afastamento seria a melhor explicação para o desvio para o vermelho (*redshift*) observado no espectro característico da estrela. Quando um detetive infere que foi Moriarty quem cometeu o crime, ele o faz porque essa hipótese explicaria melhor as impressões digitais, manchas de sangue e outras evidências forenses. Ao contrário do que Sherlock Holmes possa sugerir, isso não é uma questão de dedução. As evidências não implicam necessariamente que Moriarty seja o culpado, já que sempre permanece a possibilidade de que outra pessoa seja a autora do crime. No entanto, Holmes está certo em fazer sua inferência, já que a culpa de Moriarty forneceria uma explicação melhor para as evidências do que qualquer outra pessoa.

A Inferência da Melhor Explicação pode ser vista como uma extensão da ideia de explicações “autoevidentes”, nas quais o fenômeno a ser explicado, por sua vez, fornece uma parte essencial da razão para acreditar que a explicação está correta. Por exemplo, a velocidade de recessão de uma estrela explica por que seu espectro característico é desviado para o vermelho em uma quantidade específica, mas o desvio para o vermelho observado pode ser uma parte essencial do motivo pelo qual o astrônomo acredita que a estrela esteja se afastando a essa velocidade. Explicações autoevidentes exibem uma curiosa circularidade, mas essa circularidade é benigna. A recessão é usada para explicar o desvio para o vermelho, e o desvio para o vermelho é usado para confirmar a recessão, ainda assim, a hipótese de recessão pode ser tanto explicativa quanto bem fundamentada. De acordo com a Inferência da Melhor Explicação, esta é uma situação comum na ciência: as hipóteses são apoiadas pelas próprias observações que elas se propõem a explicar. Além disso, nesse modelo, as observações apoiam a hipótese precisamente porque ela as explicaria. A Inferência da Melhor Explicação, assim, inverte parcialmente uma visão comum da relação entre inferência e

explicação. De acordo com essa visão comum, a inferência é anterior à explicação. Primeiro, o cientista deve decidir quais hipóteses aceitar; então, quando solicitado a explicar alguma observação, ele recorrerá ao seu conjunto de hipóteses aceitas. Em contraste, de acordo com a Inferência da Melhor Explicação, é somente perguntando quão bem várias hipóteses explicariam as evidências disponíveis que o cientista pode determinar quais hipóteses merecem aceitação. Nesse sentido, a Inferência da Melhor Explicação defende que a explicação é anterior à inferência.

Existem dois problemas diferentes que uma teoria sobre indução na ciência pode tentar resolver. O problema da descrição consiste em fornecer uma explicação acerca dos princípios que regem a maneira como os cientistas avaliam as evidências e fazem inferências. O problema da justificação consiste em demonstrar que esses princípios são sólidos ou racionais, por exemplo, mostrando que eles tendem a levar os cientistas a aceitar hipóteses que são verdadeiras e a rejeitar aquelas que são falsas. A Inferência da Melhor Explicação tem sido aplicada a ambos os problemas.

As dificuldades do problema descritivo são às vezes subestimadas, porque se supõe que o raciocínio indutivo segue um padrão simples de extrapolação, com “Mais do Mesmo” como seu princípio fundamental. Assim, prevemos que o sol nascerá amanhã porque ele tem nascido todos os dias no passado, ou que todos os corvos são pretos porque todos os corvos observados são pretos. Esse modelo de “indução enumerativa”, no entanto, mostrou ser notavelmente inadequado como uma explicação da inferência na ciência. Por um lado, uma série de argumentos formais, mais notavelmente o chamado *paradoxo do corvo* e o *novo enigma da indução*, mostrou que o modelo enumerativo é excessivamente permissivo, tratando virtualmente qualquer observação como se fosse evidência para qualquer hipótese. Por outro lado, o modelo também é muito restritivo para explicar a maioria das inferências científicas. As hipóteses científicas normalmente apelam a entidades e processos não mencionados nas evidências que as sustentam, e frequentemente são elas mesmas inobserváveis e não apenas não observadas, de modo que o princípio de *Mais do Mesmo* não se aplica. Por exemplo, enquanto o modelo enumerativo poderia explicar a inferência que um cientista faz da observação de que a luz de uma estrela está desviada para o vermelho para a conclusão de que a luz de outra estrela também estará desviada para o vermelho, ele não explicará a inferência do desvio para o vermelho observado para a recessão não observada.

A tentativa mais conhecida de explicar essas inferências “verticais” que os cientistas fazem a partir de observações para hipóteses sobre a realidade, muitas vezes inobservável, que

estão por trás delas é o modelo Nomológico-Dedutivo. De acordo com esse modelo, os cientistas deduzem previsões a partir de uma hipótese (junto com várias outras “premissas auxiliares”) e depois determinam se essas previsões estão corretas. Se algumas delas não estiverem, a hipótese é desconfirmada; se todas estiverem, a hipótese é confirmada e pode eventualmente ser inferida. Infelizmente, embora esse modelo abra espaço para inferências verticais, ele permanece, como o modelo enumerativo, excessivamente permissivo, considerando dados como confirmadores de uma hipótese que, na verdade, são totalmente irrelevantes para ela. Por exemplo, uma vez que uma hipótese (H) implica a disjunção de si mesma e qualquer previsão (H ou P), e a verdade da previsão estabelece a verdade da disjunção (já que P também implica (H ou P)), qualquer previsão bem-sucedida será considerada como confirmadora de qualquer hipótese, mesmo que P seja a previsão de que o sol nascerá amanhã e H a hipótese de que todos os corvos são pretos.

O que se deseja, portanto, é uma explicação que permita inferências verticais sem permitir absolutamente tudo, e a Inferência da Melhor Explicação promete cumprir esse papel. A Inferência da Melhor Explicação legitima inferências verticais, porque uma explicação de algum fenômeno observado pode recorrer a entidades e processos que não são observados; mas ela não permite qualquer inferência vertical, uma vez que, obviamente, uma hipótese científica específica não explicaria qualquer observação, se verdadeira. Uma hipótese sobre a coloração dos corvos não explicará, por exemplo, por que o sol nascerá amanhã. Além disso, a Inferência da Melhor Explicação distingue entre diferentes hipóteses que explicariam a evidência, já que o modelo só legitima a inferência da hipótese que melhor a explicaria.

A Inferência da Melhor Explicação, assim, tem a vantagem de fornecer uma explicação natural para muitas inferências e de evitar algumas das limitações e excessos de outras abordagens familiares de inferência não demonstrativa. No entanto, se ela pretende ser um modelo sério de indução, a Inferência da Melhor Explicação precisa ser desenvolvida e articulada, e isso não tem se mostrado uma tarefa fácil. É preciso explorar mais, por exemplo, sobre as condições sob as quais uma hipótese explica uma observação. A explicação é, em si, um grande tema de pesquisa na filosofia da ciência, mas os modelos-padrão de explicação geram resultados decepcionantes quando são aplicados à Inferência da Melhor Explicação. Por exemplo, a abordagem mais conhecida da explicação científica é o modelo Nomológico-Dedutivo, segundo o qual um evento é explicado quando sua descrição pode ser deduzida a partir de um conjunto de premissas que inclui essencialmente ao menos uma lei. Este modelo tem muitas falhas. Além disso, ele é isomórfico ao modelo Hipotético-Dedutivo de

confirmação, o que reduziria, de forma decepcionante, a Inferência da Melhor Explicação a uma versão do hipotético-dedutivismo.

A dificuldade de articular a Inferência da Melhor Explicação se agrava quando abordamos a questão do que torna uma explicação melhor do que outra. Para começar, o modelo sugere que a inferência é uma questão de escolher a melhor entre aquelas hipóteses explicativas que foram propostas em um determinado momento, mas isso parece implicar que, a qualquer momento, os cientistas inferirão uma e apenas uma explicação para qualquer conjunto de dados. No entanto, os cientistas às vezes são agnósticos, relutantes em inferir qualquer uma das hipóteses disponíveis, e também estão às vezes dispostos a inferir mais de uma explicação, quando as explicações são compatíveis. “Inferência da Melhor Explicação” deve, portanto, ser entendida pela frase mais precisa, mas menos memorável: “inferência da melhor das explicações concorrentes disponíveis, quando a melhor é boa o suficiente”. Mas sob quais condições essa condição complexa é satisfeita? Quão boa é “boa o suficiente”? Ainda mais fundamentalmente, quais são os fatores que tornam uma explicação melhor do que outra? Os modelos-padrão de explicação praticamente não dizem nada sobre esse ponto. Isso não sugere que a Inferência da Melhor Explicação esteja incorreta, mas, a menos que possamos dizer mais sobre a [noção de] explicação, o modelo permanecerá relativamente pouco informativo.

Felizmente, algum progresso foi feito na análise da noção relevante de *melhor explicação*. Podemos começar considerando uma questão básica sobre o sentido de “melhor” que o modelo requer. Isso significa a explicação mais provável ou, em vez disso, a explicação que, se correta, proporcionaria o maior grau de compreensão? Em resumo, a Inferência da Melhor Explicação deve ser interpretada como uma inferência da explicação *mais provável* (*likeliest explanation*) ou da explicação *mais plausível* (*Loveliest explanation*)?¹ Uma explicação particular pode ser tanto provável quanto plausível, mas as noções são distintas. Por exemplo, se alguém diz que fumar ópio tende a fazer as pessoas dormirem porque o ópio tem um “poder dormitivo”, está dando uma explicação que tem muita probabilidade de estar correta, mas que não é nada plausível: ela fornece muito pouca compreensão. À primeira vista, pode parecer que a probabilidade seja a noção que a Inferência da Melhor Explicação

¹ Nota do tradutor: No original em inglês, Lipton utiliza as expressões: *likeliest* e *loveliest*. Optamos por traduzir como “mais provável” e “mais plausível” respectivamente. Contudo, uma outra possível tradução para *loveliest* é a expressão “mais proporcionadora de entendimento”, conforme pode ser encontrado na tradução do artigo *Inferência da Única Explicação* de Alexander Bird, realizada por Silva e Luz na revista *Cognitio*, São Paulo, v. 15, n. 2, 2014.

deveria empregar, uma vez que os cientistas presumivelmente inferem apenas a hipótese mais provável entre as concorrentes que consideram. No entanto, provavelmente essa é a escolha errada, pois reduziria severamente o interesse do modelo, empurrando-o para a trivialidade. Os cientistas de fato inferem o que julgam ser a hipótese mais provável, mas o principal objetivo de um modelo de inferência é precisamente dizer como esses julgamentos são feitos, ou seja, fornecer o que os cientistas consideram serem os *sintomas* da probabilidade. Dizer que os cientistas inferem as explicações mais prováveis é perigosamente semelhante a dizer que grandes *chefs* preparam as refeições mais saborosas: verdadeiro, talvez, mas não muito informativo se alguém quiser saber os segredos do sucesso deles. Assim como a explicação do poder dormitivo dos efeitos do ópio, a “Inferência da Explicação Mais Provável” seria, ela mesma, uma explicação da prática científica que fornece apenas pouco entendimento.

O modelo deve, portanto, ser interpretado como “Inferência da Explicação Mais Plausível”. Sua afirmação central é que os cientistas tomam a plausibilidade como guia para a probabilidade, ou seja, a explicação que, se correta, proporcionaria o maior entendimento é a explicação que é considerada mais provável de estar correta. Isso, pelo menos, não é uma afirmação trivial, mas levanta pelo menos três desafios. O primeiro desafio é identificar as virtudes explicativas, as características das explicações que contribuem para o grau de compreensão/entendimento que elas fornecem. O segundo é demonstrar que esses aspectos da plausibilidade correspondem aos julgamentos de probabilidade, ou seja, que as explicações mais plausíveis tendem a ser também aquelas julgadas mais prováveis de estarem corretas. O terceiro desafio é mostrar que, considerando a correspondência entre plausibilidade e julgamentos de probabilidade, a primeira é, de fato, a guia dos cientistas para a segunda.

Para começar com o desafio da identificação, há várias candidatas plausíveis para as virtudes explicativas, tais como abrangência [alcance explicativo], precisão, mecanismo, unificação e simplicidade. Explicações melhores explicam mais tipos de fenômenos, explicam-nos com maior precisão, fornecem mais informações sobre mecanismos subjacentes, unificam fenômenos aparentemente distintos ou simplificam nossa visão geral do mundo. Todavia, algumas dessas características se mostraram surpreendentemente difíceis de analisar. Não há, por exemplo, um exame incontestável de unificação ou simplicidade, e alguns questionaram até mesmo se essas são características genuínas de hipóteses científicas, em vez de meras consequências aparentes resultantes da forma como elas acabam sendo formuladas, de modo que a mesma explicação será considerada simples se formulada de uma maneira, mas complexa se formulada de outra.

Uma abordagem diferente, mas complementar, para o problema de identificar algumas das virtudes explicativas foca na estrutura contrastiva de muitas perguntas do tipo “por quê”. Um pedido para a explicação de algum fenômeno frequentemente assume uma forma contrastiva: não se pergunta simplesmente “Por que P?”, mas “Por que P e não Q?”. O que conta como uma boa explicação depende não apenas do fato P, mas também do contraste com Q. Assim, o aumento de temperatura pode ser uma boa explicação de por que o mercúrio de um termômetro ter subido, em vez de ter descido, mas pode não ser uma boa explicação de por que ele subiu, em vez de quebrar o vidro. Dessa forma, é possível desenvolver uma explicação parcial do que torna uma explicação de um dado fenômeno melhor do que outra, especificando como a escolha da alternativa determina a adequação das explicações contrastivas. Embora muitas explicações, tanto na ciência quanto na vida cotidiana, especifiquem algumas das causas putativas do fenômeno em questão, a estrutura da explicação contrastiva mostra por que nem todas as causas servem. Grosso modo, uma boa explicação exige uma causa que “faz a diferença” entre o fato e [a alternativa de] contraste. Assim, o fato de Smith ter sífilis não tratada pode explicar por que ele, e não Jones, contraiu paralisia (uma forma de paralisia parcial), se Jones não tinha sífilis; mas não explicará por que Smith, e não Doe, contraiu paralisia, se Doe também tinha sífilis não tratada. Nem todas as causas fornecem explicações plausíveis, e uma explicação contrastiva ajuda a identificar quais fornecem explicações plausíveis e quais não.

Assumindo que uma abordagem razoável para as virtudes explicativas seja apresentada, o segundo desafio para a Inferência da Melhor Explicação diz respeito à extensão da correspondência entre a “plausibilidade” e os julgamentos de probabilidade. Se a Inferência da Melhor Explicação estiver na direção correta, então as explicações mais “plausíveis” devem também, em geral, ser consideradas mais prováveis. Aqui, a situação parece promissora, já que as características que identificamos provisoriamente como virtudes explicativas também parecem ser virtudes inferenciais, ou seja, características que dão suporte a uma hipótese. Hipóteses que explicam muitos fenômenos observados com alta precisão tendem a ser mais bem justificadas do que hipóteses que não o fazem. O mesmo parece se aplicar a hipóteses que especificam um mecanismo, que unificam, e que são simples. A sobreposição entre as virtudes explicativas e inferenciais certamente não é perfeita, mas pelo menos alguns casos de hipóteses que são prováveis, mas não “plausíveis”, ou vice-versa, não representam uma ameaça particular à Inferência da Melhor Explicação. Como vimos, a explicação do efeito soporífico do ópio baseada no “poder dormitivo” é muito provável, mas

nada “plausível”; mas isso não representa uma ameaça ao modelo, quando corretamente interpretado. Certamente existem explicações mais profundas para o efeito de fumar ópio, em termos de estrutura molecular e neurofisiologia, mas essas explicações não competem com a explicação banal, então o cientista pode inferir ambas sem violar os preceitos da Inferência da Melhor Explicação.

A estrutura da explicação contrastiva também ajuda a enfrentar o desafio de correspondência, porque os contrastes nas questões do tipo “por que” frequentemente correspondem aos contrastes nas evidências disponíveis. Uma boa ilustração disso é fornecida pela investigação de Ignaz Semmelweis, no século XIX, sobre as causas da febre puerperal, uma doença frequentemente fatal contraída por mulheres que davam à luz no hospital onde Semmelweis realizava suas pesquisas. Semmelweis considerou muitas explicações possíveis. Talvez a febre fosse causada por “influências epidêmicas” que afetavam os distritos ao redor do hospital, ou talvez fosse causada por alguma condição no próprio hospital, como superlotação, dieta inadequada ou tratamento inadequado, grosseiro. O que Semmelweis percebeu, no entanto, foi que quase todas as mulheres que contraíram a febre estavam em uma das duas alas de maternidade do hospital, o que o levou a fazer a óbvia reflexão contrastiva e, em seguida, a descartar as hipóteses que, embora logicamente compatíveis com sua evidência, não marcavam uma diferença entre as alas. Isso também o levou a inferir uma explicação que explicaria o contraste entre as alas, a saber, que as mulheres estavam sendo inadvertidamente infectadas por estudantes de medicina que iam diretamente de autópsias para exames obstétricos, mas apenas examinavam as mulheres da primeira ala. Essa hipótese foi confirmada por um procedimento contrastivo adicional, quando Semmelweis fez os médicos desinfetarem as mãos antes de entrar na enfermaria: a hipótese da infecção passou a explicar não apenas por que as mulheres na primeira, e não na segunda ala, contraíam febre puerperal, mas também por que as mulheres na primeira repartição contraíam a febre antes, mas não depois da introdução do regime de desinfecção. Esse padrão geral de argumentação, que busca explicações que não apenas explicam um determinado efeito, mas também contrastes específicos entre casos em que o efeito ocorre e casos em que ele está ausente, é muito comum na ciência, como, por exemplo, sempre que são realizados experimentos controlados.

Isso nos leva ao desafio da orientação. Mesmo que seja possível apresentar um esclarecimento sobre a “plausibilidade explicativa” (o desafio da identificação) e mostrar que as virtudes explicativas e inferenciais coincidem (o desafio da correspondência), ainda resta argumentar que os cientistas julgam que uma hipótese tem mais chances de estar correta

porque é plausível, como afirma o modelo de Inferência da Melhor Explicação. Assim, um crítico do modelo pode afirmar que explicações prováveis também tendem a ser plausíveis, mas argumentar que a inferência se baseia em outras considerações, sem relação com explicação. Por exemplo, pode-se argumentar que inferências a partir de dados contrastivos são, na verdade, aplicações do método da diferença de Mill, que não faz um apelo explícito à explicação, ou que a precisão é uma virtude porque previsões mais precisas têm uma probabilidade prévia menor e, portanto, fornecem um suporte mais forte como uma consequência elementar do cálculo de probabilidades.

O defensor da Inferência da Melhor Explicação se encontra aqui em uma posição delicada. Ao tentar mostrar que as virtudes explicativas e inferenciais coincidem, ele inevitavelmente mostrará também que as virtudes explicativas correspondem a algumas das características citadas por outras abordagens concorrentes de inferência como as verdadeiras guias inferenciais. O defensor, assim, expõe-se à acusação de que são essas outras características, e não as virtudes explicativas, que desempenham o verdadeiro papel inferencial. Superar o desafio da correspondência, portanto, exacerbaria o desafio da orientação. No entanto, a situação não é desesperadora, pois existem pelo menos duas maneiras de argumentar que a plausibilidade é uma guia para os julgamentos de probabilidade. Como vimos, pelo menos algumas das abordagens concorrentes de inferência estão repletas de dificuldades, são inaplicáveis a muitas inferências científicas e incorretas em relação a outras. Se for mostrado que a Inferência da Melhor Explicação fornece uma explicação melhor para mais inferências do que qualquer outra abordagem disponível, isso é um forte motivo para supor que a plausibilidade é, de fato, uma guia para a probabilidade. Em segundo lugar, se houver uma boa correspondência entre plausibilidade e probabilidade, como o desafio da orientação assume, isso presumivelmente não é uma coincidência e, portanto, exige uma explicação. Por que as hipóteses que os cientistas julgam mais prováveis de serem corretas também são aquelas que forneceriam mais compreensão, se estivessem corretas? A Inferência da Melhor Explicação dá uma resposta muito natural a essa pergunta, semelhante à explicação darwiniana para o fato de os organismos tenderem a se ajustar bem aos seus ambientes. Se os cientistas selecionam hipóteses com base nas suas virtudes explicativas, a correspondência entre plausibilidade e julgamentos de probabilidade se segue naturalmente. A menos que os opositores do modelo possam apresentar uma explicação melhor para essa correspondência, o desafio terá sido superado.

Estivemos considerando as perspectivas da Inferência da Melhor Explicação como uma solução parcial para o problema de descrever a estrutura das inferências científicas, mas o modelo também foi aplicado a problemas de justificação. O problema mais fundamental da justificação indutiva se deve a David Hume (1975), que argumentou que não há boa razão para acreditar que nossas práticas indutivas sejam minimamente confiáveis, ou seja, que tendam a nos levar de observações verdadeiras a hipóteses ou previsões verdadeiras. Segundo Hume, para justificar a indução, teríamos que produzir um argumento convincente cuja conclusão fosse que a indução é geralmente confiável e cujas premissas não fossem elas mesmas baseadas em indução. As únicas premissas assim seriam os relatos de observações passadas e as verdades demonstrativas da lógica e da matemática. Todos os argumentos convincentes são ou dedutivos ou indutivos. Agora enfrentamos um dilema. Não pode haver um argumento dedutivo convincente para a confiabilidade da indução, uma vez que nenhum número de observações passadas (juntamente com as verdades demonstrativas) garante dedutivamente que a indução seja geralmente confiável. Em suma, observações passadas nunca implicarão que a indução será confiável no futuro. Tampouco há um argumento indutivo convincente para a indução, uma vez que qualquer tal argumento pressupõe justamente a prática que se supõe justificar. Por exemplo, argumentar que a indução provavelmente será confiável no futuro com base no fato de que tem sido confiável no passado cairia na falácia do círculo vicioso, mesmo que fosse aceito que a confiabilidade passada da indução pudesse ser conhecida a partir da observação. Assim, nossas práticas indutivas são injustificáveis.

Se o argumento de Hume for válido, não há razão alguma para acreditar em qualquer afirmação científica que vá além do que foi diretamente observado, o que, pelo menos, significa que não há razão para acreditar em nenhuma previsão, hipótese ou teoria científica. Isso é inacreditável, mas o argumento cético tem se mostrado extraordinariamente resistente e ainda não há uma resposta amplamente aceita para ele. No entanto, apesar de toda a sofisticação de Hume ao apresentar o problema da justificação, sua solução para o problema da descrição é bastante rudimentar. Ele parece ter aceitado uma versão do modelo de indução simples e enumerativo “Mais do Mesmo”, discutido anteriormente. Consequentemente, pode-se esperar que uma explicação mais sofisticada e precisa da prática indutiva tornaria possível evitar ou refutar o argumento cético de Hume. Em particular, às vezes se supõe que a Inferência da Melhor Explicação forneça tal explicação.

Infelizmente, a Inferência da Melhor Explicação não resolve o problema de Hume. A descrição que ele deu acerca da indução estava incorreta, mas seu argumento cético não depende disso. Na verdade, o argumento parece depender de pouco mais do que o fato inegável de que os argumentos indutivos não são validamente dedutivos. Relatos de observações passadas nunca implicarão que futuras inferências de melhores explicações irão de fato selecionar hipóteses verdadeiras; e qualquer argumento de que a confiabilidade da inferência da melhor explicação seria a própria melhor explicação do que observamos incorreria em uma petição de princípio. Pode-se até mesmo argumentar que a Inferência da Melhor Explicação agrava o problema da justificação, já que é bastante incerto por que a hipótese que, se correta, proporcionaria a mais profunda compreensão é também, de fato, a mais provável de ser correta. Por que deveríamos supor que o nosso é o mais belo de todos os mundos possíveis? No entanto, essa preocupação adicional pode ser uma reação exagerada, pois o argumento cético de Hume sugere é que o sucesso de qualquer outro método de indução seria igualmente misterioso.

A Inferência da Melhor Explicação também foi invocada para resolver problemas mais modestos de justificação indutiva. Mesmo que o modelo não seja útil contra um cético indutivo completo, ele pode ter um papel na defesa do realismo científico, segundo o qual há boas razões para acreditar que teorias bem fundamentadas provavelmente são pelo menos aproximadamente verdadeiras, em oposição a posturas como o empirismo construtivo, segundo o qual só temos razões para acreditar que nossas melhores teorias são empiricamente adequadas, ou seja, que suas consequências observáveis são verdadeiras. (O empirismo construtivo foi desenvolvido em detalhes por Bas Van Fraassen (1980), que também é um crítico vigoroso da Inferência da Melhor Explicação). O empirista construtivo não é um cético indutivo, já que afirmar que todas as consequências observáveis de uma teoria são verdadeiras é uma afirmação muito mais forte do que meramente dizer que suas consequências observadas são verdadeiras; mas o realista vai além, sancionando também inferências verticais para a verdade das alegações de uma teoria sobre entidades e processos inobserváveis.

Talvez o exemplo mais conhecido dessa aplicação da Inferência da Melhor Explicação na defesa do realismo científico seja o chamado “argumento do milagre”, discutido por Hilary Putnam (1978). Ele considera que o modelo oferece uma boa solução para o problema descritivo e propõe que os filósofos possam fazer, por si mesmos, uma inferência da melhor explicação na defesa do realismo científico. Suponha que todas as muitas e variadas previsões derivadas de uma teoria científica particular sejam consideradas corretas: qual seria a melhor

explicação para esse sucesso preditivo? De acordo com Putnam, a melhor explicação é que a própria teoria é verdadeira. Se a teoria for verdadeira, então a verdade de suas consequências dedutivas se segue naturalmente; mas se a teoria for falsa, seria um “milagre” que todas as suas consequências observadas sejam consideradas corretas. Assim, por uma aplicação filosófica da Inferência da Melhor Explicação, estamos autorizados a inferir que a teoria é verdadeira, já que a “explicação da verdade” é a melhor explicação para o sucesso preditivo da teoria. Essa inferência de nível superior seria distinta das inferências de primeira ordem feitas pelos cientistas, mas segue a mesma forma.

Essa aplicação justificatória da Inferência da Melhor Explicação tem considerável apelo intuitivo, mas enfrenta três objeções. A primeira é que a “explicação da verdade” para o sucesso preditivo de uma teoria não é realmente distinta das explicações científicas substantivas que a teoria fornece e com base nas quais ela foi inferida pelos cientistas em primeiro lugar. Se isso for verdade, então o argumento do milagre não fornece nenhuma razão adicional para acreditar que a hipótese está correta: ele é meramente uma repetição da inferência científica que deveria justificar. No entanto, essa objeção pode ser respondida observando que os dois tipos de explicação possuem estruturas diferentes. As explicações científicas fornecidas por uma teoria são tipicamente causais, enquanto a “explicação da verdade” é lógica. A verdade de uma teoria não causa fisicamente que suas consequências sejam verdadeiras; a conexão explicativa é, ao invés disso, que um argumento válido com premissas verdadeiras deve também ter uma conclusão verdadeira.

A segunda objeção ao argumento do milagre é que, mesmo que a “explicação da verdade” seja distinta das explicações científicas, a inferência da verdade da teoria é viciada pelo mesmo tipo de circularidade que Hume invocou em seu argumento cético. Na prática, o argumento do milagre é uma tentativa de usar uma inferência da melhor explicação para justificar as inferências científicas da melhor explicação, então, como o objetor dirá, tal argumento incorre em petição de princípio quanto à confiabilidade dessa forma de inferência. Em particular, o empirista construtivo pode insistir que, embora ele permita a legitimidade de algumas formas de indução, as inferências da verdade de teorias que lidam com entidades inobserváveis são precisamente aquelas que estão em questão. Uma possível resposta à objeção da circularidade é argumentar que o círculo é quebrado em virtude da diferença entre inferências de explicações causais e de explicações lógicas, mas a objeção tem uma força considerável.

A terceira objeção ao argumento do milagre é que a verdade simplesmente não é a melhor explicação para o sucesso preditivo, de modo que o argumento falha em seus próprios termos. A maneira óbvia de desenvolver essa objeção é apresentar outra explicação que seja pelo menos tão boa. Por exemplo, o empirista construtivo pode afirmar que podemos explicar o sucesso preditivo de uma teoria supondo que ela seja empiricamente adequada, ou seja, que todas as suas consequências observáveis são verdadeiras, independentemente de a teoria ser verdadeira como um todo. No entanto, nesse caso, o defensor do argumento do milagre tem duas respostas prontas. Primeiramente, não está claro que a explicação em termos de adequação empírica seja tão “plausível” quanto a explicação da verdade, já que ela está perigosamente próxima de afirmar que as consequências da teoria são verdadeiras porque são verdadeiras, uma explicação extremamente pouco plausível, remanescente do apelo ao “poder dormitivo” do ópio. Além disso, mesmo que, como no caso do ópio, inferimos essa explicação, isso não impede uma inferência da explicação da verdade, pois as duas explicações são compatíveis: uma teoria pode ser tanto empiricamente adequada quanto verdadeira. A terceira objeção ao argumento do milagre pode, no entanto, se tornar mais convincente por meio de uma escolha melhor de explicações alternativas. Pois, dado qualquer conjunto de previsões bem-sucedidas, sempre existem, em princípio, muitas teorias incompatíveis com a teoria original, mas que compartilham essas consequências. A verdade de qualquer uma das teorias concorrentes também explicaria o sucesso preditivo que elas compartilham com a teoria original, e não está claro que essas explicações alternativas da verdade seriam menos “plausíveis” do que a original. Assim, a inferência da verdade da teoria original pode ser bloqueada.

Nenhuma das aplicações justificatórias da Inferência da Melhor Explicação que consideramos parece promissora. Se o modelo pode ajudar a resolver problemas de justificação indutiva, é provável que esses problemas se refiram a aspectos mais específicos das práticas indutivas dos cientistas. Por exemplo, o modelo foi aplicado de forma plausível em um argumento para mostrar por que é racional que os cientistas deem maior peso aos dados que uma hipótese prevê corretamente do que aos dados que estavam disponíveis quando a hipótese foi formulada e que ela foi construída para acomodar. Seja qual for o potencial justificatório da Inferência da Melhor Explicação, no entanto, o modelo pode ser considerado um sucesso filosófico se puder ser demonstrado que ele oferece uma descrição esclarecedora de alguns dos princípios inferenciais gerais que orientam a prática científica.

REFERÊNCIAS

FRAASSEN, B. van. **The scientific image**. Oxford: Oxford University Press, 1980.

GARFINKEL, A. **Forms of explanation**. New Haven: Yale University Press, 1981.

HARMAN, G. The Inference to the Best Explanation. **The Philosophical Review**, Ithaca, v. 74, n. 1, p. 88-95, 1965.

HEMPEL, C. **Aspects of scientific explanation**. New York: Free Press, 1965.

HUME, D. **An enquiry concerning human understanding**. Oxford: Oxford University Press, 1975.

LIPTON, P. **Inference to the Best Explanation**. London: Routledge, 1991.

PEIRCE, C. S. **Collected papers**. Cambridge: Harvard University Press, 1931.

PUTNAM, H. **Meaning and the moral sciences**. London: Hutchinson, 1978.

THAGARD, P. The best explanation: criteria for theory choice. **Journal of Philosophy**, New York, v. 75, n. 2, p. 76-92, 1978.