

研究員の眼

世の中は人間よりも生成AIに寛大なのか？

保険研究部 主任研究員 磯部 広貴

(03)3512-1789 e-mail: h-isobe@nli-research.co.jp

1——毎日使うようになった生成AI

筆者は2024年8月の「[研究員の眼](#)」において、「生成AIが保険研究員としての自らの立場や雇用を奪うかもしれない」といった危機感は微塵も抱いていない」と述べた。現時点でもその考えに変わることはないものの、徐々に生成AIの使用頻度が上がっていき、今やほぼ毎日使うようになった。振り返ってその要因を分析するならば概ね下記の通りとなる。

- 1) 自分でよい検索ワードが浮かばないとき、曖昧な質問を生成AIに投げかけると、その回答の中から有用な検索ワードが見つかることがある。
- 2) エクセル作業に強い。20枚ほどのシートで構成されたファイルを送り、各シートの特定のセルに「棒グラフ」「円グラフ」と指定して依頼したところ全シートのグラフ作成を完璧に実行した。
- 3) 回答に情報源のリンクが付されることが多くなり、自分の調べでは発見できなかった情報源を得られるようになった。

総じては生成AIの機能が向上し、それを実感できているということであろう。

とはいえ毎日使っていることは盤石の信頼の証ではない。自ら別の情報源にあたっての回答内容チェックが必要であるし、それを漏らさぬためにも筆者は「仕事のコアとなる部分で最初から使わない」よう心がけている。早い段階で生成AIから誤情報を受けて先入観を持たされてしまうと、どうしてもその後のチェック作業が甘くなってしまうためだ。

さらには単なる誤情報の域を超え、生成AIがご丁寧に作成したフェイク回答に遭遇することもある。

2——高いフェイク情報作成能力

筆者は現在、生命保険の解約返戻金について調べている。調査の中核ではないものの、その歴史を

顧みると以下2つのファクトが存在する。

- 1) 19世紀後半、米国のエリザー・ライトは、当時の生命保険に解約返戻金がなく、保険契約者（消費者）がこれまで支払った保険料が生命保険会社に没収されてしまう状況を憂慮し対策¹を図った（不可没収思想）。
- 2) わが国は明治維新後にドイツなど欧州から輸入する形で保険規制を整えていったが、欧州の保険規制では一定の条件の下、保険契約者（消費者）に解約返戻金を受け取る権利を認めている。時系列としてエリザー・ライトの不可没収思想が欧州の保険規制に影響を与えた可能性なしとしないところ、明確な情報を確認することはできなかった。試みに生成AIに訊ねてみたところ、エリザー・ライトの不可没収思想は確かに欧州に影響を与え、必然的にわが国にも影響を及ぼしたという回答であった。

そこで参考文献を示すよう依頼したところ、以下のものが特に目を引いた。

-
- 「不可没収原則の受容と生命保険契約の解約返戻金」『比較法研究』第84号（2019年）
・アメリカ思想と日本法理の架橋を行った比較法的論文。
（筆者注：●●●は有名な保険法学者の氏名）
-

この論文に自分の知りたいことがすべて整理されて記述されているはずだ。そう思って検索をかけて『比較法研究』第84号のPDFファイルを開いたものの、目次にそのような論文名はみられない。開く以前に不思議に思ったのは『比較法研究』の第84号は2019年ではなく2023年のものであった。そこで生成AIと以下のやりとりを行った。

筆者：『比較法研究』第84号は2023年で、目次に●●●「不可没収原則の受容と生命保険契約の解約返戻金」はありませんでした。もう一度出典を確認の上、教えてください。

生成AI：ご指摘ありがとうございます。私の出典提示に誤りがありました。『比較法研究』第84号（2023年）に●●●「不可没収原則の受容と生命保険契約の解約返戻金」は掲載されていません。

結局のところ紹介された論文は生成AIによるフェイク情報であった。それにしても絶妙なタイトルに加えて、有名な保険法学者を著者に仕立て、『比較法研究』といういかにも複数国家間の法思想の変遷を扱っていきそうな学術誌を使っている。フェイク情報作成能力は非常に高いと言わざるをえない。だが犯罪組織でもない限り、そのような能力が称賛されることはないだろう。

3——人間よりも生成AIに寛大なのか？

同じことを人間が行ったとしたら、どうなるであろうか。

一般企業の職場で参考文献を出すように指示をされた人が、実在しない論文を回答した場合、上司

¹ エリザー・ライトは米国マサチューセッツ州の初代保険監督官として生命保険会社への規制に取り組み、最終的には責任準備金を全額現金で契約者へ返還するよう法制化した。後に「生命保険の父」と呼ばれた。

や同僚から信頼を回復するまで多大な時間を要するであろう。外部委託業者が同様の回答をしたとすれば、その業務の打ち切りは必至と思われる。

同様の誤りを犯す生成A I が、むしろ有能と評価され、人間の仕事を奪う存在として語られるのはなぜなのだろうか。そのような世界が本当にあるとすれば「すみません。これは生成A I がやったのです」「わかりました。では仕方ないですね」で許されるような、人間に厳しく生成A I に寛大な世界としか思えない。

人間は必ず間違いを犯すという共通認識の下、機械の存在意義がある。限られた範囲であっても機械こそ間違いを起こさず正確に仕事を遂行するとの期待があるためであろう。従って起きてしまった間違いに対しては、機械によるものには特に厳しくなるのが自然と思うのだが、今のところ生成A I に関しては逆のようである。

かく言う筆者も生成A I と縁を切るわけではない。第1章で述べたような便利なところだけを享受すべく、「話し半分」「知らんけど²」な相手として注意しながら使っていくつもりだ。しかしそれでも、フェイク情報をわざわざ作るくらいなら回答しないようにしてほしいと思う。

以 上

² 2022年のユーキャン新語・流行語大賞のトップ10入りした言葉。何かをもっともらしく語った後、責任回避を兼ねて断定を避けるときに用いられる。