

Por Samuel de Jesus Monteiro de Barros

Por muito tempo, a subscrição de risco foi uma arte artesanal disfarçada de ciência, onde se tinha um subscritor experiente, uma tabela de fatores e um certo grau de intuição respaldada por estatística. Era opaca, sim, mas era uma opacidade humana, sujeita a auditoria por entrevista, a "por que o senhor recusou este risco?". Em 2026, trocamos essa opacidade por outra, mais eficiente e infinitamente mais difícil de interrogar, afinal, agora estamos falando dos modelos. A ironia é que o setor vendeu essa troca como modernização, quando na prática substituiu uma caixa-preta que ao menos sabia explicar-se, ainda que mal, por outra que frequentemente não sabe, e que, pior, nem sempre é obrigada a tentar.

Essa troca deixou de ser um problema de modelagem estatística para se tornar, com uma velocidade que pegou parte do mercado desprevenida, um problema de governança corporativa. E é sobre essa migração, do departamento de dados para a sala do conselho, que vale a pena parar para pensar.

Da promessa de eficiência ao dever de fundamentação, de onde vem gatilho regulatório?

O ano de 2026 é, por definição legal, o primeiro ano de vigência plena da Lei nº 15.040/2024, o novo marco do contrato de seguro no Brasil. Entre suas exigências está o dever de fundamentação adequada para negativas de cobertura, o que, na prática, eleva o ônus probatório da seguradora exatamente nos casos em que a recusa se baseia em um modelo automatizado, especialmente quando a suspeita é de fraude identificada por IA. Essa elevação do ônus probatório, já apontada em análises jurídicas recentes sobre o tema, por si só já deveria bastar para tirar o tema do departamento de tecnologia e colocá-lo na pauta de risco legal.

A Susep fez sua parte para institucionalizar essa exigência, por meio da Resolução Susep nº 72/2025, formalizou seu Plano de Regulação para 2026, sinalizando que a fraude tecnológica e os controles associados devem figurar entre os temas prioritários do exercício regulatório. Em paralelo, e este é o ponto que mais deveria interessar a quem se senta em conselhos de administração, tramita na Câmara o Marco Legal da IA (PL nº 2.338/2023), já aprovado no Senado, que deve impor avaliação preliminar de risco obrigatória a sistemas de IA generativa e de propósito geral, sob coordenação de um futuro sistema nacional vinculado à ANPD, autoridade que já atua na prática como reguladora de fato do tema, mesmo antes de a lei ser sancionada, e que incluiu a IA entre seus eixos prioritários de fiscalização para o biênio 2026-2027.

Ou seja, o vácuo normativo que por anos serviu de justificativa confortável para adiar a discussão sobre governança algorítmica está se fechando por múltiplas frentes ao mesmo tempo, setorial (Susep), geral (Marco Legal da IA) e de proteção de dados (ANPD), antes mesmo que o setor tenha amadurecido um consenso sobre o que, exatamente, "explicar uma decisão algorítmica" deveria significar na prática.

Explicabilidade como obrigação, não como boa prática!

É tentador tratar explicabilidade (XAI, na sigla em inglês) como um daqueles conceitos que soam bem em relatório de sustentabilidade e morrem antes de chegar ao comitê de risco. O movimento regulatório em curso sugere o contrário. Explicabilidade tem definição operacional cada vez mais precisa como a capacidade de uma seguradora demonstrar quais critérios e variáveis influenciaram uma decisão automatizada (recusa de apólice, precificação diferenciada, negativa de sinistro). Não é uma aspiração ética, é, cada vez mais, um requisito documental que precisa resistir ao escrutínio de um perito, de um juiz ou de um órgão regulador.

O mercado, a seu crédito, não está inteiramente alheio a essa transição. Análises setoriais recentes descrevem 2026 como a "Era da Produção" da IA no seguro brasileiro, o momento em que modelos

generativos deixam de ser projeto-piloto e passam a operar precificação em tempo real e regulação automatizada de sinistros de baixo risco, com a explicabilidade citada como o "grande diferencial" competitivo do período. É uma leitura correta, mas otimista demais em um ponto: trata explicabilidade como vantagem a ser adotada por escolha estratégica, quando o pano de fundo regulatório aponta para uma obrigação que independe da vontade de diferenciação de cada seguradora.

Seguindo o espelho europeu da supervisão prudencial à supervisão de conduta

Se o Brasil ainda está consolidando esse arcabouço, vale olhar para onde a discussão já amadureceu mais. A experiência europeia é instrutiva porque documenta uma transição de ênfase regulatória que o Brasil parece destinado a repetir, com alguma defasagem: a supervisão do setor segurador, historicamente ancorada no eixo prudencial (solvência, reservas, adequação de capital), vem cedendo espaço crescente a uma supervisão comportamental, centrada nos resultados para o consumidor, na equidade dos processos e na adequação dos produtos oferecidos. A EIOPA, autoridade europeia de seguros, já consolidou princípios específicos para uso de IA no setor com proporcionalidade, não discriminação, transparência, explicabilidade, supervisão humana, governança de dados e robustez dos modelos.

O deslocamento conceitual mais interessante nessa literatura é a ideia de que regras sobre decisões automatizadas, como as do regime europeu de proteção de dados, não deveriam ser lidas como uma restrição à tecnologia, mas como um modelo de governança das decisões com impacto significativo sobre o cliente. Em outras palavras: o problema nunca foi a IA decidir; é a IA decidir sem que ninguém consiga, depois, reconstruir e justificar o raciocínio. Essa é precisamente a fronteira entre eficiência operacional e risco de conduta, e é precisamente onde a responsabilidade deixa de ser da equipe de dados e passa a ser do conselho.

Por que isso tem que ser pauta de conselho?

Aqui está o argumento central deste texto, e ele não é sutil: um conselho de administração que delega inteiramente a governança de modelos de IA ao time técnico está cometendo o mesmo erro que cometeria ao delegar integralmente a gestão de risco de crédito ao sistema de scoring, sem supervisão. A diferença é que, no caso da IA aplicada a subscrição e sinistros, o risco não é apenas prudencial (a seguradora pode precificar mal um risco), é também um risco de conduta, jurídico e reputacional, cujo custo se materializa em ação regressiva, em multa administrativa e, cada vez mais, em ônus probatório invertido contra a seguradora em juízo.

Três implicações práticas decorrem disso, e nenhuma delas é um exercício de futurologia:

1. Primeiro, o dever de documentação deixa de ser apenas técnico. Um "laudo técnico claro, fundamentado e compreensível", para usar a formulação que já aparece em análises jurídicas sobre fraude por IA no setor, precisa ser compreensível não só para o perito, mas para as próprias seguradoras, para autoridades policiais e para o Poder Judiciário. Isso é, por definição, um problema de comunicação e de governança documental, não de arquitetura de rede neural.
2. Segundo, a supervisão humana deixa de ser um checkbox de compliance para se tornar um requisito de desenho de processo: decisões automatizadas de alto impacto (negativa de cobertura, recusa por suspeita de fraude) precisam ter um ponto de revisão humana que seja genuíno, não decorativo, sob pena de a "supervisão humana" virar, ela mesma, mais um item de opacidade a ser questionado.
3. Terceiro, e este é o ponto mais desconfortável para conselhos acostumados a tratar tecnologia como pauta operacional, a pergunta que deveria estar na mesa não é "qual ferramenta de IA adotar", mas sim "quem, neste conselho, é capaz de exigir e avaliar criticamente a explicação de uma decisão algorítmica antes que um regulador ou um juiz o faça por nós". Sem essa capacidade instalada em nível de governança, a seguradora está apenas terceirizando seu risco de conduta para um fornecedor de tecnologia e descobrirá isso, tipicamente, no pior momento possível: durante uma fiscalização ou um litígio.

Uma pauta que ainda não amadureceu e que não vai esperar

O que há de comum entre a Lei 15.040/24, o Plano de Regulação da Susep para 2026, o Marco Legal da IA em tramitação e os princípios já consolidados pela EIOPA é a convergência, vinda de direções diferentes, para uma mesma exigência: a de que decisões algorítmicas com impacto sobre o segurado sejam demonstráveis, não apenas eficientes. O mercado brasileiro de seguros tem, nesse aspecto, uma vantagem pouco explorada, a de observar o amadurecimento regulatório europeu antes que a fiscalização doméstica atinja a mesma maturidade, e um risco simétrico: o de tratar essa janela como tempo de folga, em vez de tempo de preparação.

A explicabilidade algorítmica não vai continuar sendo, por muito mais tempo, um tópico de painel em conferência de inovação. Vai ser, cada vez mais literalmente, uma linha de pauta na ata de reunião do conselho. a pergunta que precisa ser feita antes que seja feita por outra pessoa, em outro tom, e com poder de multa.

(09.09.2026)