

From Exclusion to Accountability

A Framework for Rebalancing Underwriting and Claims Incentives in the AI Era

Madhav Moganti

Independent Researcher, Moganti Labs LLC — mmoganti@outlook.com

Working paper · August 2026

Abstract

Insurance imposes two related costs on the people it is meant to protect: the underwriting cost of exclusion before a policy exists, and the claims cost of undervaluation after a loss occurs. Both stem from institutional information and effort asymmetry, but they differ in a way that matters for how artificial intelligence should be deployed against each. Part-One addresses life insurance underwriting, where hereditary or chronic conditions are priced from cohort-level mortality data that can be too coarse to reflect an individual's evidence-backed prognosis. We show the underwriting asymmetry is structural — a decline costs an insurer nothing today, while accepting a bad risk costs it immediately — and rejects an intuitive ex-post fix in favor of six mechanisms that correct the asymmetry directly. Part-Two addresses claims valuation, using total-loss auto claims as a case study: unlike underwriting, this asks a verifiable, backward-looking factual question, and documented litigation against a dominant valuation platform shows that consumer-side AI tools can close the cost-of-verification gap that lets systematically low offers go unchallenged. We conclude that AI's most defensible near-term role in insurance is collapsing the cost of independently checking an institution's number, not replacing its judgment — a role it fills cleanly in claims valuation today, and only partially, with guardrails, in underwriting. The underlying decision problem, the epistemic–aleatory distinction limiting individual-level verification, and the proposed mechanisms are formalized throughout and summarized in five tables and three figures.

1. Introduction: Two Asymmetries, One Root Cause

More than 100 million Americans currently lack the life insurance coverage they say they need (LIMRA & Life Happens, 2025). Some of this gap reflects genuine underuse — misunderstanding of cost, competing financial priorities, and a product category people buy for the benefit of others rather than themselves. But a meaningful share of the gap is not underuse at all. It is exclusion: people who sought coverage and were declined, rated at a price they could not sustain, or discouraged from applying because a hereditary condition made the process feel futile before it began.

This paper starts from one such case. An applicant with early-onset, hereditary diabetes — the same condition their parent has carried since a similarly young age, now sixty-plus years into the diagnosis and in good health at ninety-one — finds themselves priced out of a market that exists, in principle, to protect exactly this kind of family from financial hardship. That single data point is not an argument that underwriting is unnecessary; a parent whose actual outcome dramatically exceeds a standard mortality projection is one family's experience, not a repeal of population statistics. It is an argument that underwriting, as currently practiced, may be too coarse — and that the shift toward AI-driven underwriting is a genuine fork in the road for how coarse it stays.

This is not simply one family's good fortune. Real-world data tracking continuous glucose monitors and automated insulin delivery systems show measurable reductions in severe hypoglycemia, diabetic ketoacidosis, and related hospitalizations among users (Ambalavanan et al., 2024; Isaacs et al., 2021). Independent cost-effectiveness studies in France and Australia have each modeled a resulting gain in quality-adjusted life expectancy of roughly a full additional year or more compared to traditional fingerstick monitoring (Roze et al., 2020; Isitt et al., 2022). The exact size of that gain varies study to study and population to population; the direction does not. A parent's outcome that once looked like statistical noise is, increasingly, a documented and growing category — which raises its own question, taken up in Section 4.2, about who currently has access to the technology producing it.

Table 1. Modeled quality-adjusted life-year (QALY) gains from continuous glucose monitoring (CGM) versus self-monitored blood glucose (SMBG), by country.

Study	Country	Model used	Incremental QALYs
Roze et al. (2020)	France	IQVIA Core Diabetes Model	1.38
Isitt et al. (2022)	Australia	IQVIA Core Diabetes Model	1.199
Valentelyte et al. (2024)	Ireland	Sheffield-linked Markov model	0.76–2.993 (range)

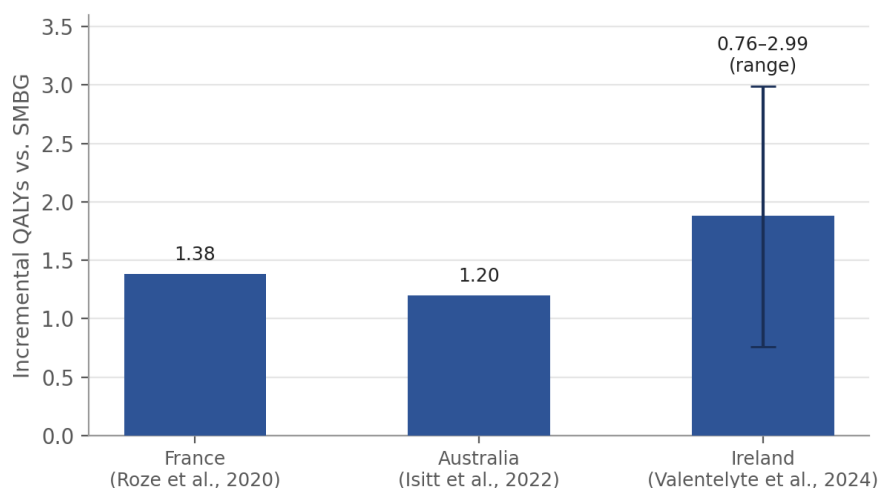


Figure 1. Incremental QALYs attributable to CGM use relative to SMBG, three independent health-system models. Error bar on the Ireland estimate reflects the published sensitivity range rather than a single point estimate.

The differences across rows in Table 1 are not noise to be averaged away; they reflect genuinely different populations, currencies, and cost-effectiveness models, which is itself the point — three independent health systems, using three different models, all find a positive, non-trivial QALY gain. A useful summary statistic from health economics is the incremental cost-effectiveness ratio, which expresses cost per QALY gained between two alternatives 1 and 2:

$$ICER = (C_2 - C_1) / (E_2 - E_1) \quad (1)$$

where C denotes cost and E denotes effectiveness in QALYs for each alternative. Every study in Table 1 reports an ICER below its jurisdiction's own willingness-to-pay threshold, which is the formal, standard basis on which health systems judge a technology worth covering — not an informal impression that the technology helps.

This paper treats the question in two parts. Part One examines the underwriting case just described, where the asymmetry runs between a costless decline and a costly bad acceptance. Part Two turns to a second, related case: how insurers value claims after a loss has already occurred, using a recent auto insurance example in which a policyholder's own AI-assisted analysis of the insurer's own ten selected comparable vehicles produced a valuation more than 30 percent higher than the adjuster's offer — and prompted the insurer to raise its settlement by a similar margin once challenged.

The two cases share a root cause — institutional information and effort asymmetry — but differ in a way that matters. Claims valuation asks a verifiable question about the past: what would this vehicle have sold for. Underwriting asks an unverifiable question about the future: how long will this person live. That difference determines how confidently AI can be expected to correct each one, and it is the organizing thread of what follows. Two futures are available for either domain. In one, AI simply automates the crude, effort-asymmetric version of today's practice, at a speed and scale no human department could match. In the other, AI is used to close the asymmetry itself — modeling individual factors precisely, monitoring outcomes against assumptions in real time, and putting a serious verification tool within reach of an

ordinary policyholder for the first time. Which future arrives depends less on the technology than on the business-model and regulatory incentives built around it. That is the subject of this paper.

1.1 Related Literature and Scope of Evidence

The underwriting asymmetry formalized in Section 2 is a claim about error costs, not about self-selection, and it is worth distinguishing from the classical adverse-selection literature it complements. Akerlof (1970) showed that when sellers know more about quality than buyers, a market can unravel toward only the worst-quality goods being traded; Rothschild and Stiglitz (1976) showed that in competitive insurance markets specifically, asymmetric information about risk type can leave no stable equilibrium at all, or force low-risk buyers into partial coverage to separate themselves from high-risk ones. Both results concern what happens — or fails to happen — when market prices risk under uncertainty about who is asking. This paper's asymmetry is a different, narrower claim that operates downstream of that pricing decision: given that an insurer must sort applicants under uncertainty, its own accounting already treats the two possible sorting errors unequally, independent of how well or poorly the sorting itself performs. The two literatures are complementary — Rothschild-Stiglitz explains why underwriting exists at all; this paper asks why the underwriting that exists is calibrated the way it is.

On the applied side, this paper's account of algorithmic bias in insurance underwriting draws on industry and regulatory analysis (KPMG, 2025; The Actuary, 2023; Grant Thornton, 2023; Sidley Austin, 2023) rather than a mature academic literature, because — as of this writing — peer-reviewed empirical work specifically on AI-driven life underwriting is thin; the same is true, more starkly, of claims-valuation automation, where the best available evidence is litigation and trade-press reporting rather than published research. This is a genuine limitation of the evidentiary base rather than a stylistic choice, and it bears on what kind of paper this is. This paper does not have access to proprietary carrier data — decline rates by condition, table-rating distributions, or cross-carrier experience studies — that would be needed to quantify how large the underwriting asymmetry actually is in dollars, as opposed to establishing that it exists and runs in a predictable direction. It should accordingly be read as a conceptual and incentive-design framework, grounded in documented regulatory, legal, and medical-outcomes evidence, rather than an empirical estimate of the asymmetry's magnitude. Closing that gap — through a carrier partnership, an SOA-sponsored experience study, or MIB-aggregated data on hereditary and chronic conditions specifically — is the most valuable next empirical step this paper does not itself take.

PART ONE — UNDERWRITING

2. The Hidden Asymmetry in Underwriting Economics

To design a better incentive structure, it helps to be precise about the one that exists today. When an insurer declines an application, the decision carries no cost to the insurer, however the applicant's life subsequently unfolds. If the applicant would, in fact, have died within a few years, the decline was “right,” though no one collects on having been right. If the applicant instead lives productively for decades past what the risk class implied, the decline was “wrong,” but again no one pays for it. The lost premium income is a real cost, but it is diffuse, invisible, and rarely tracked by the institution that incurred it.

Compare this to the opposite error: accepting an applicant who turns out to be a genuinely bad risk. This mistake is not diffuse. It shows up as a claim, on a specific policy, inside a specific underwriting cohort, and it is the error actuaries watch most closely — because enough of these errors, uncorrected, is what threatens an insurer's ability to pay every other policyholder's claim. That is the adverse-selection risk that gives medical underwriting its economic justification in the first place.

This can be stated formally. Let q_x denote an applicant's true probability of mortality within the relevant policy horizon, and let B denote the benefit amount at stake. Write L_A for the insurer's realized loss from a wrongful accept (a claim paid on a risk that should have been priced higher or excluded), and L_D for its realized loss from a wrongful decline (the foregone profit from a risk that would, in fact, have been profitable). Under current practice, the expected cost of each error type is:

$$E[L_D] \approx 0 \qquad E[L_A] = q_x \times B > 0 \qquad (2)$$

That single-line asymmetry — a decline's expected cost is approximately zero while an acceptance's expected cost scales directly with the benefit at risk — is the entire formal content of this paper's diagnosis. Table 2 makes the same point as a payoff matrix: the insurer's realized cost depends on the interaction between its decision and the applicant's true, unobserved type, and the current regime leaves one full quadrant costless.

Table 2. Realized cost to the insurer, by decision and by the applicant's (unobserved) true risk type.

	True type: Low risk	True type: High risk
Decision: Decline	0 (foregone profit, untracked)	0 (avoided claim)
Decision: Accept	Small positive (normal expenses)	$q_x \times B$ (claim paid; = L_A in Eq. 2)

Reading down the “Decline” row: both cells are approximately zero, regardless of the applicant's true type. Reading down the “Accept” row: one cell is small and the other is not. A cost-minimizing decision rule facing this matrix, and any genuine uncertainty about type, always prefers the top row when in doubt — not because sick applicants are disfavored as a matter of policy, but because the matrix itself only punishes one kind of mistake. Section 4 is, formally, a proposal to populate the empty top-row cells with a real number greater than zero, tied to pool-level performance rather than an individual prediction.

An underwriter facing genuine uncertainty is therefore facing a lopsided bet: err toward inclusion, and any mistake is expensive, immediate, and attributable; err toward exclusion, and any mistake is free, deferred indefinitely, and untraceable to the decision that caused it. Treated as a repeated game across millions of applications, this asymmetry does not require any individual underwriter, or any AI model, to be biased against sick people in order to produce a systematically exclusionary market. It only requires each decision-maker to respond rationally to the incentives already in place.

This asymmetry helps explain a pattern that recurs throughout underwriting guidance: the same medical file can be rated “Standard” at one carrier and several tables worse at another, because neither carrier bears a cost for guessing conservatively, while both bear a real cost for guessing generously and being wrong (New York Life, 2025; Memorial Merits, 2026). It also helps explain why the coverage gap

concentrates in exactly the population — hereditary and chronic conditions — whose individual variance around a cohort average is largest, and therefore hardest to price with confidence from cohort data alone.

2.1 The Legal Backdrop: Why Exclusion Is Permitted

This asymmetry persists in part because the law does not meaningfully constrain it. The Americans with Disabilities Act contains an explicit insurance carve-out, permitting insurers to structure coverage around underwriting and risk classification so long as that classification is grounded in actuarial data and consistent with state law (Federal Register, 2016). Courts have applied this carve-out to life insurance specifically, treating a denial as lawful once it rests on a legitimate risk classification rather than disability status alone. A second, narrower gap compounds the first: the Genetic Information Nondiscrimination Act bars health insurers and employers from using genetic test results, but explicitly does not extend to life insurance (Legal Information Institute, n.d.; U.S. Congress, 2008). Only a small number of states have closed this gap on their own, meaning a genetic test that could inform better medical care can, outside those states, simultaneously inform a worse underwriting outcome. Neither provision was written with malice; both reflect a decades-old judgment that risk classification is a legitimate, largely unregulated business practice. That judgment is precisely what the mechanisms in Section 4 are designed to work within, not around.

3. Why an Ex-Post Penalty on Individual Declines Fails

A natural first response to this asymmetry is to attach a cost to the exclusion side of it directly: require insurers to compensate a declined applicant who outlives the mortality expectation implicit in their decline, much as a wrongful-denial doctrine already penalizes insurers for mishandling a claim. This proposal deserves to be taken seriously — and rejected in its literal form, for three specific reasons that any redesign has to solve for.

First, a decline does not come with a stated prediction. It is a judgment about a risk class — applicants who look like this file have a meaningfully worse aggregate mortality curve than standard — not a wager on any one person's death date. There is no “assumed timeline” on record to check an outcome against; one would have to be constructed after the fact from standard mortality tables for the relevant risk class, which simply re-derives the assumption the penalty is meant to punish, in a circle.

It is tempting to treat this as a data problem that better models will eventually close — give an AI enough biomarkers, genomic data, and lifetime history, and it should eventually be able to name a verifiable date. This conflates two different kinds of uncertainty, though. Let T denote an applicant's age at death, a random variable, and let I denote the information available about that applicant at underwriting. Any model's uncertainty about T decomposes as:

$$\text{Var}(T | I) = \sigma_{epi}(I)^2 + \sigma_{ale}^2 \quad (3)$$

More data and better models shrink the first term, $\sigma_{epi}(I)^2$, the epistemic component, the share of uncertainty that comes from imprecise information rather than the underlying process. They cannot shrink the second term, σ_{ale}^2 , the aleatory component: the genuine randomness in any single, unrepeated life course (an undiagnosed clot, an accident, a cell division that goes wrong) that persists even as I approaches complete information. Figure 2 illustrates the same decomposition visually (an idealized illustration, not

fitted to data): three distributions over age at death, built on progressively richer information, narrow toward a common floor width but never collapse to a single point.

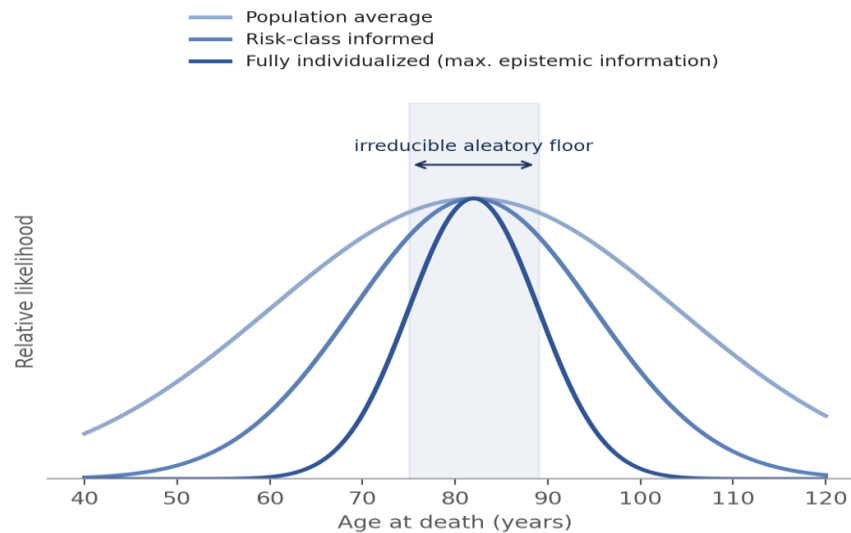


Figure 2 (illustrative). Conceptual illustration of Equation (3), not fitted to mortality data. Additional information narrows each distribution's spread, but all three retain a shared, non-vanishing minimum width — the aleatory floor.

A weather model with perfect atmospheric data still outputs a probability of rain, not a verifiable fact about tomorrow, because tomorrow happens exactly once; formally, a forecast is calibrated when, across every instance where it states probability q , the event occurs with empirical frequency q :

$$Pr(\text{event} \mid \hat{p} = q) = q \quad \text{for all } q \in [0, 1] \quad (4)$$

Calibration is a property of a large set of forecasts, not of any single one — there is no version of “was this instance right” that Equation (4) can answer for a single day, or a single applicant. A mortality model, however advanced, stands in exactly this relation to any one applicant. What improves with better AI is $\sigma_{epi}(I)^2$, the precision of the comparison group an applicant is judged against, the subject of Section 4.2, not the possibility of checking a single predicted outcome against a single life. Section 4's mechanisms are built around that distinction deliberately: they hold insurers accountable at the level Equation (4) actually operates on, a cohort, rather than an individual-level verification this paper does not expect any future model to deliver.

Second, and more importantly, the proposal corrects only one of the two errors Table 2 already lays out. A rule that attaches cost only to a wrongful decline — leaving a wrongful accept's cost exactly as it is today — gives insurers a clear rational response: become more conservative about acceptance generally, since acceptance now risks look-back liability on both ends, or price the expected cost of the new penalty into premiums across the whole book, spreading it across every policyholder rather than correcting the asymmetry Equation (2) identifies. Followed to its logical end, a fully functioning version of this proposal converges toward something resembling guaranteed-issue, community-rated life insurance — the reform path health insurance already took through the Affordable Care Act's guaranteed-issue mandate, and the same reform Section 5.3 shows can fail badly without a complementary mechanism keeping low-risk

applicants in the pool — but arrives there through an expensive, individually litigated, decades-long mechanism instead of a direct policy choice.

Third, the closest existing legal analogue confirms where the actual gap sits. Every insurance contract carries an implied covenant of good faith, and most U.S. states allow “bad faith” damages — beyond the policy's face value, occasionally punitive — when an insurer wrongfully denies a claim on a policy that already exists, including life insurance policies (Justia, 2026; McKennon Law Group, 2026). But this doctrine requires a contract to breach. A declined application never created one, so there is nothing for a court to call bad faith. The legal system already penalizes wrongful conduct toward people already inside the risk pool, and has almost nothing to say about the door itself — a gap Part Two of this paper returns to, from the other side.

The lesson is not that insurers should face no accountability for how they draw the boundary of coverage. It is that the accountability mechanism needs to operate on the asymmetry directly — correcting the fact that exclusion is cost-free — rather than attaching an ad hoc, individually litigated penalty to only one of its two symptoms.

4. A Framework: Six Mechanisms to Reward Inclusion and Accuracy

The mechanisms below share a design principle: none requires an insurer to accept unpriceable risk, and none asks a regulator or court to reconstruct an individual mortality prediction after the fact. Each instead makes accuracy and inclusion more profitable than blanket exclusion, using tools — real-time experience monitoring, individualized modeling, transparent scorecards — that are only computationally practical at the scale AI now makes possible. Table 3 summarizes all six against a common set of criteria before each is developed in full below.

Table 3. Summary of the six proposed mechanisms.

Mechanism	Targets	Verified at	Existing precedent
4.1 Experience-rated inclusion pools	Cost asymmetry (Eq. 2)	Cohort	Mutual-insurer dividends
4.2 Protective-factor modeling	Temporal bias	Individual input, cohort-checked	NAIC bias-testing mandate
4.3 Graduated-accept-by-default	Blanket exclusion	Individual	Table ratings, guaranteed issue
4.4 Public accuracy scorecards	Opacity	Carrier	NAIC explainability principle
4.5 Right to reconsideration	Non-contestability	Individual	Claims appeal processes
4.6 Longevity risk transfer	Balance-sheet conservatism	Cohort	Enhanced annuities, longevity swaps

4.1 Experience-Rated Inclusion Pools

Mutual insurers already return “dividends” to policyholders when a book of business performs better than the mortality and expense assumptions priced into it, a mechanism refined over more than a century of

participating-policy structures. This paper proposes extending that logic to rated and substandard cohorts specifically: applicants accepted at a rated premium, rather than declined, would be pooled into a cohort tracked against its own priced-in mortality assumption using standard actuarial experience-study methodology (Society of Actuaries Research Institute, 2024). Formally, for a cohort with priced-in mortality assumption m_{priced} and observed experience m_{actual} , a dividend D is declared whenever experience beats assumption:

$$D = \alpha \times \max(0, m_{\text{priced}} - m_{\text{actual}}) \times \text{Premium} \quad (5)$$

for some payout fraction $\alpha \in (0, 1]$ set by the insurer's own solvency margin. If that cohort collectively outperforms the assumption — as a population of well-managed, hereditary-condition applicants plausibly would — a defined share of the resulting surplus returns to that cohort as reduced future premiums or a cash dividend. No one has to say what any single person was “supposed” to do; only whether the pool as a whole beat its own assumption, which is precisely what experience studies already measure, and precisely the quantity Equation (4) says can be checked.

Two operational questions determine whether this mechanism is real or merely notional: how α is set, and who verifies m_{actual} . On the first, mutual insurers already have a working answer — dividend scales are set annually through the same contribution-principle process used for ordinary participating policies, holding back a margin for required capital (state risk-based-capital rules) and for smoothing across volatile years before any surplus is released; α for a rated-cohort pool should sit below the ordinary dividend scale, not above it, precisely because a newly rated cohort has less claims history to smooth against. On the second, self-certified experience would defeat the mechanism's purpose; m_{actual} should be established the way statutory mortality experience already is, through an independent actuarial opinion filed with the domiciliary regulator, using the same NAIC financial-examination infrastructure that already audits reserve adequacy. Neither requirement is new institutional machinery — both already exist for other purposes — which is the point: the mechanism repurposes existing actuarial and regulatory audit trails rather than inventing new ones.

4.2 Mandated Individualized Protective-Factor Modeling

Regulators already require insurers to test AI-driven underwriting for bias (National Association of Insurance Commissioners [NAIC], 2023). This paper proposes extending that requirement in a constructive direction: any AI model used to underwrite a chronic or hereditary condition should be required to incorporate documented protective factors — duration since diagnosis, current control markers, treatment adherence, and parental or sibling longevity with the same diagnosis — as inputs that can move an applicant toward a better class, not only as inputs that penalize the diagnosis code itself. This directly counters what the actuarial literature calls “temporal bias”: the tendency of models trained on older cohorts to underweight outcomes, like modern diabetes management, that barely existed in the historical data (The Actuary, 2023).

This mechanism carries a real limitation worth stating plainly: control-marker quality is not evenly accessible. Continuous glucose monitor and automated insulin delivery adoption is strongly stratified by income and race, and the technology-use gap between the highest and lowest socioeconomic quintiles has

widened, not narrowed, over more than a decade of tracked data (Addala et al., 2021; Been et al., 2024). Table 4 and Figure 3 report one documented cross-section of the pattern.

Operationalizing this mechanism requires three concrete choices. First, validation standard: a factor should qualify as “protective” only once a peer-reviewed cost-effectiveness or outcomes study links it to reduced mortality or morbidity risk — the same evidentiary bar health systems already apply through the ICER threshold in Equation (1) — rather than an insurer's internal correlation mining, which is exactly the practice that produces the temporal- and proxy-bias failures described in Section 7. Second, data source and attestation: device-reported metrics (CGM time-in-range, A1D adherence logs) should be pulled directly from manufacturer or clinician data-sharing APIs under the applicant's authorization, not self-reported, to prevent the reconsideration right in Section 4.5 from becoming a venue for unverified claims. Third, and most important given Table 4: the access-adjusted baseline. Rather than scoring an applicant's control against the general population, the model should first predict expected control conditional on documented technology access (insurance-covered device status, prescribing specialty, zip-code-level access proxies), and then reward control that exceeds that conditional expectation. An applicant is judged against people with comparable access to the tools that produce good control, not against everyone — which is what keeps this mechanism from simply re-importing Table 4's disparity into the underwriting model under a new name.

Table 4. CGM uptake by patient race/ethnicity, single-center cohort of young adults with Type 1 diabetes (n = 300).

Group	Ever used CGM	Mean HbA1c
White (n ≈ 99)	71%	8.5% ± 1.8
Hispanic (n ≈ 102)	37%	n.r.
Black (n ≈ 96)	28%	10.7% ± 2.4

Source: Isaacs et al. (2021). n.r. = not separately reported in the source.

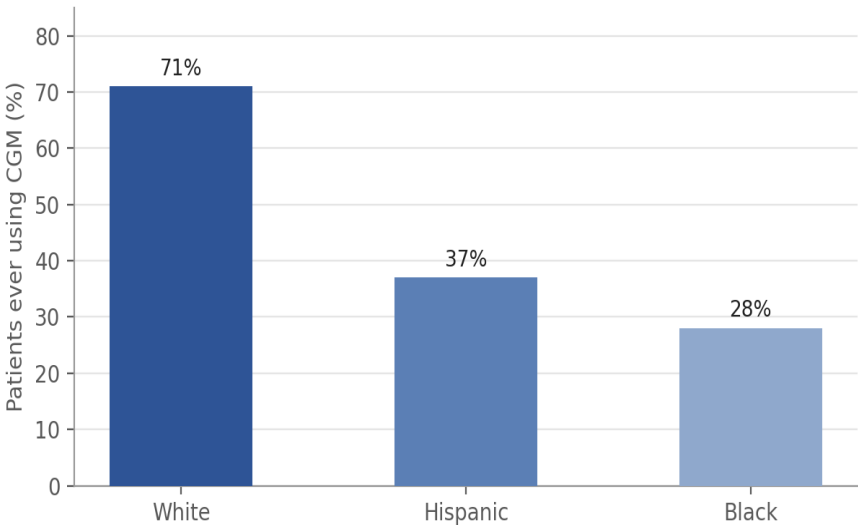


Figure 3. CGM uptake by patient race/ethnicity (data from Table 4).

The roughly two-to-one gap between the highest- and lowest-uptake groups in Table 4 is not a fringe finding; it recurs, in comparable magnitude, across the income-stratified literature as well. An underwriting model that rewards demonstrated control without adjusting for this risks quietly re-creating the socioeconomic and racial proxy-discrimination problem Section 7 already warns about, just relocated from a credit score to an HbA1c trend line. Any real implementation of this mechanism should weight control-marker evidence against a baseline calibrated to the applicant's documented access to the relevant technology, not against the general population.

4.3 Graduated-Accept-by-Default Underwriting

Several tools already sit between standard approval and decline: table-rated policies, simplified issue, and guaranteed issue with a graded death benefit (NerdWallet, 2025; Western & Southern, 2025). This paper proposes making that ladder the computational default in AI-driven underwriting — a model should calculate the highest table rating at which a risk remains actuarially viable before returning a decline, rather than defaulting to decline whenever an applicant falls outside a narrow standard-issue band. Outright decline would be reserved for cases where no actuarially sound price exists at all, rather than used as the default response to ordinary chronic-condition uncertainty.

4.4 Public Inclusion and Accuracy Scorecards

Because the NAIC Model AI Bulletin already requires insurers to document bias-testing and maintain explainable systems (NAIC, 2023; Holland & Knight, 2025), this paper proposes a modest extension: publish, by carrier, the approval rate for major chronic and hereditary conditions at any tier — standard, rated, or guaranteed-issue — alongside how that carrier's accepted cohorts subsequently perform against their own priced-in assumptions. A carrier that is chronically over-conservative, whose rated cohorts persistently and substantially outperform the mortality assumptions it charged for, would be visible in the data, creating a reputational and competitive incentive to compete on accurate inclusion rather than on how narrowly it can screen applicants out.

The metric that matters is narrow and specific: a realized-to-priced mortality ratio, by condition category and rating tier, reported as an index rather than a raw rate. That specificity is also the mechanism's main operational risk. A raw approval rate by condition is competitively sensitive — it discloses underwriting appetite to rivals — and a sufficiently fine-grained cross-tabulation risks re-identifying individual applicants in a thin market segment. Both problems have a standard answer: a designated statistical agent, structured on the model the NAIC already uses for other cross-carrier statistics, collects raw carrier data under confidentiality and publishes only the aggregated index, suppressing any cell below a minimum cohort size (a k-anonymity threshold, in the terminology Section 7 uses for algorithmic fairness generally). This is closer to how the Medical Information Bureau or a state's all-payer claims database already operates than to a new disclosure regime.

4.5 A Structured Right to Algorithmic Reconsideration

The NAIC bulletin's transparency and explainability principle (NAIC, 2023) implies a right that is rarely operationalized in practice: if a system must be explainable, it should also be contestable. This paper proposes a formal, time-bound right for an applicant or their broker to submit specific individualized

evidence — family longevity despite the same diagnosis, updated clinical control data — after an initial AI-generated decline or rating, with a requirement that the carrier document, in plain language, whether and how that evidence changed the outcome. This converts an opaque, final decision into a documented, appealable one, directly addressing the “black box” concern that has driven much of the regulatory interest in explainable AI (Buchanan Ingersoll & Rooney, 2025). Part Two of this paper describes a close analogue already operating in claims handling.

Bounding the administrative burden this creates is straightforward if the right is scoped narrowly: limited to an enumerated evidence list (not an open-ended appeal), a fixed response window comparable to the claims-appeal timelines many states already mandate, and routed through a broker or agent of record where one exists, so the marginal cost falls on an intermediary already compensated for the case rather than on the carrier's underwriting staff directly. The realistic failure mode is not carrier burden but applicant gaming — cherry-picked or embellished supporting evidence — which is best addressed the way underwriting already handles material misrepresentation on an initial application: reconsideration evidence becomes part of the contestable record, subject to the same two-year scrutiny Section 3 describes, rather than a one-way ratchet that only ever helps the applicant.

4.6 Longevity Risk Transfer to Expand Underwriting Capacity

Enhanced, or impaired-life, annuities already price the mirror image of this paper's subject: a shorter-than-average life expectancy earns the buyer a higher guaranteed income, and the annuity provider bears the loss if the buyer outlives that shortened estimate (LV=, 2026). Reinsurers and capital-markets counterparties already trade this exact risk at scale, through longevity swaps and related instruments, because pension funds and annuity books want to hedge against paying out longer than expected (RGA, 2011; Artemis, 2018). This paper proposes that life insurers use the same market to hedge the opposite exposure — the risk that a rated or guaranteed-issue cohort, accepted under the mechanisms above, lives longer than priced for — freeing up underwriting capacity to accept more marginal risks without holding as large a conservatism margin on the insurer's own balance sheet.

This mechanism is the most speculative of the six, and honestly so. The longevity-swap market's current liquidity and pricing infrastructure is built around pension-scale block transactions — tens of thousands of lives, highly standardized — not the smaller, more idiosyncratic rated-cohort exposures this paper describes. Whether that market can price a hereditary-condition cohort of a few hundred lives at all, let alone efficiently, is genuinely unresolved, and the likely near-term path is aggregation: several carriers pooling rated-cohort exposure through a shared reinsurance treaty large enough to interest an existing longevity-swap counterparty, rather than any single carrier transacting directly. Readers should weigh 4.6 accordingly — as the mechanism with the clearest economic logic and the least-proven implementation path of the six.

5. Guardrails, Costs, and Counterarguments

None of the above is a case for abandoning risk-based pricing. Adverse selection is a real, measurable phenomenon, and a market that ignored it would eventually fail the people already inside it, not only the people trying to get in. Every mechanism proposed here operates on the accuracy and transparency of risk

classification — rewarding insurers for pricing correctly and inclusively — rather than forcing acceptance of risks that cannot be priced sustainably at any premium.

5.1 Adverse Selection Does Not Stop at the Pool Boundary

There is a second-order adverse-selection risk internal to Section 4.1 worth naming directly: a rated pool is not immune to the same self-selection dynamic that motivates underwriting in the first place. Applicants near the favorable edge of a rated tier — whose true risk is closer to standard than the tier assumes — have the strongest incentive to shop for a carrier willing to rate them better, or to wait until their next reconsideration window, leaving the pool that actually forms skewed toward its own unfavorable edge relative to what was priced. This is precisely the phenomenon Rothschild and Stiglitz (1976) describe, recurring one level down. It argues for conservative initial calibration of α and for experience studies segmented finely enough to detect within-tier drift before it erodes a cohort's solvency margin, not for abandoning the mechanism.

5.2 Costs, Spillovers, and Industry Resistance

A single carrier that unilaterally becomes more inclusive, without the rest of the market moving too, risks attracting a disproportionate share of the applicants competitors are declining — a textbook adverse-selection exposure for that one carrier, even if the broader shift would be sound at an industry level. This is why several mechanisms above — the AI Bulletin-linked scorecards, mandated protective-factor modeling — are proposed as industry-wide regulatory requirements, following the precedent the NAIC has already set, rather than as unilateral strategy any single insurer could safely adopt alone. Two further costs deserve equally direct treatment. First, premium spillover: Section 4.1's dividends are funded from within a rated cohort's own surplus by construction, not drawn from the standard pool, so that mechanism should not raise standard-tier premiums; 4.2 through 4.5, by contrast, carry real compliance and systems costs that carriers will price into their book generally, and some pass-through to the standard pool should be expected and disclosed rather than assumed away. Second, small-carrier burden: real-time experience monitoring and access-adjusted modeling favor insurers with existing data-science infrastructure, which is an argument for the shared statistical-agent model in Section 4.4 to also host smaller carriers' compliance function, the way rating bureaus already do for smaller property-casualty insurers, rather than treating each carrier's compliance as a bespoke build. Industry resistance to all of this should be expected as a rational response to compliance cost and first-mover disadvantage, not dismissed as bad faith — which is exactly why the mechanisms are designed as regulatory floors rather than appeals to voluntary adoption.

5.3 A Historical Warning: Risk Pooling Without Safeguards

The clearest cautionary precedent for anything resembling broader risk pooling comes from 1990s health insurance reform, and it argues for exactly the guardrails this section has just specified rather than against reform generally. In 1993, Washington State and New York each adopted guaranteed issue and community rating in their individual health insurance markets without an accompanying purchase mandate. In Washington, premiums rose roughly 78 percent and enrollment fell roughly 25 percent over the following three years; by 1999, 17 of the state's 19 private individual-market insurers had exited, and the remaining two had announced their intent to follow (Seattle Times, 2017). New York's individual market saw premium

increases on the order of 40 percent and the exit of its largest individual insurer by 1996 (Seattle Times, 2017). The mechanism was exactly the one Rothschild and Stiglitz's framework predicts: without a mandate keeping healthy people in the pool, community rating made coverage relatively cheaper for the sick and relatively more expensive for the healthy, the healthy left first, and the resulting pool's true cost pushed premiums up again, repeating until the market could no longer sustain itself.

The lesson for this paper is not that pooling fails — mechanism 4.1 pools only applicants already accepted at a rated premium, a materially narrower and already-underwritten population, not an open community-rated market — but that any mechanism resembling broader risk-sharing needs a corresponding mechanism keeping the pool from unraveling at its edges, which is exactly the function Section 5.1's within-tier monitoring and Section 4.4's public scorecards are meant to serve. A reform that reduces exclusion without an answer to “what keeps the good risks from leaving” is not a smaller, safer version of this paper's proposal; it is the 1993 experiment with different branding.

PART TWO — CLAIMS

6. A Second Case: Claims Valuation and the Cost of Verification

Underwriting is not the only point at which an insurer's information advantage translates into a lower payment. This section examines a claims-side case that is, in one important respect, more tractable than the underwriting problem examined so far — and, in another respect, a sharper illustration of what AI can and cannot be expected to fix.

6.1 How Total-Loss Valuation Actually Works

When an insurer deems a vehicle a total loss, it owes the policyholder the vehicle's actual cash value (ACV) immediately before the loss, ordinarily established by identifying comparable vehicles for sale in the local market and adjusting for differences in mileage, condition, and equipment. In practice this process is dominated by a single commercial platform: CCC Information Services, whose CCC ONE software is estimated by at least one certified appraiser organization to control more than 85 percent of the total-loss valuation market in the United States (My Fair Claim, 2026). CCC ONE and comparable platforms apply two adjustments to each comparable vehicle's listed price: a condition adjustment, typically negative and often applied uniformly across every comparable without any individual inspection, and a “projected sold adjustment,” a downward deduction intended to estimate a final negotiated sale price from an advertised asking price (Irvin Legal, 2024).

6.2 Two Explanations, Not One

There are two distinct, and mutually reinforcing, explanations for why an insurer's opening offer in this process tends to run low.

The first is that the valuation methodology itself is documented, in specific instances, to skew downward. State Farm faces active lawsuits alleging exactly this: *Brewer v. State Farm*, filed in the Western District of North Carolina in October 2025, alleges an automated CCC ONE function reduced comparable-vehicle retail costs by roughly 4 to 9 percent — about 5.5 percent on the plaintiff's own \$56,772 vehicle, a \$3,328 difference — in violation of North Carolina's total-loss valuation regulation (Repairer Driven News, 2025); a federal court denied State Farm's motion to dismiss the case in February 2026, finding the complaint stated a plausible claim for breach of contract and unfair trade practices (Law360, 2026). *Goode v. State Farm*, filed in the Northern District of Alabama in December 2025, alleges the same underlying methodology (My Fair Claim, 2026). Similar suits have been filed in at least six other states; two, in Kentucky and Mississippi, have already been dismissed, a reminder that this litigation is unsettled and not uniformly favoring plaintiffs (Lawyer Monthly, 2025).

The second explanation does not require the tool to be biased at all: it only requires that most claimants do not independently verify the insurer's number, and that an adjuster has limited authority to move off an initial offer without being presented with specific counter-evidence (Auto Claim Consultants, n.d.). Under these conditions, an opening offer calibrated to the low end of a plausible range is a rational starting point regardless of the underlying methodology's accuracy, because the expected cost of a lower offer — some claimants push back, most do not — is lower than the expected cost of a higher one, where every claimant benefits immediately. An independent analysis of closed total-loss claims where policyholders invoked their appraisal clause between 2023 and 2025 found the pattern directly: every insurer-hired appraiser in the dataset valued the claim below the eventual settlement, by margins ranging from 9 to 74 percent, with a mean gap of roughly 34 percent across all claims examined (AppraiseltNow, 2025) — a figure that lines up closely with the roughly 30 percent gap in the case that opens this Part, suggesting that example is representative of a broader pattern rather than an outlier. These two explanations compound rather than compete: a downward-biased valuation tool makes the low anchor look more defensible on paper, while the anchor itself does the majority of the practical work of reducing the average payout (SnapClaim, 2026).

The State Farm litigation is not an isolated data point. A related Massachusetts case against Commerce Insurance made similar allegations about the same underlying valuation methodology; it was dismissed in 2025, though on procedural standing grounds rather than a finding that the practice itself was lawful, leaving the substantive question unresolved rather than settled in the insurer's favor (Agency Checklists, 2026). Table 5 summarizes the documented cases as of this writing; readers should expect this table to be stale quickly, given how fast this docket is moving.

Table 5. Documented litigation concerning CCC ONE total-loss valuation methodology, as of early 2026.

Case	Jurisdiction	Core allegation	Status
<i>Brewer v. State Farm</i>	W.D. North Carolina	Automated deductions of ~4–9% regardless of local market	Motion to dismiss denied, Feb. 2026
<i>Goode v. State Farm</i>	N.D. Alabama	Same underlying CCC ONE methodology	Active
Suits in AK, IL, KY, MS, TN, WV	Various	Same underlying methodology	Mixed — KY, MS dismissed

Case	Jurisdiction	Core allegation	Status
Commerce Insurance	Massachusetts	Same underlying valuation methodology	Dismissed — standing, not merits

6.3 The Legal Mechanism That Already Fits

Section 3 of this paper noted that the bad-faith doctrine governing wrongful claim denials applies only once a contract already exists, and therefore has nothing to say about underwriting declines, which occur before any contract is formed. A total-loss claims dispute is the case that doctrine was built for: it concerns the handling of a claim on a policy the claimant already holds. An unreasonably low settlement offer, issued without a genuine investigation into the actual local market, is a far more plausible candidate for bad-faith claims handling than any underwriting decision could be (Justia, 2026).

More immediately useful than litigation, however, is a contractual mechanism already present in most auto and many homeowners policies: the appraisal clause. This provision allows either party to demand an independent appraisal when they disagree on a covered loss's value; if the policyholder's and insurer's appraisers cannot agree, a neutral umpire is jointly selected, and the resulting value becomes binding on both parties (Auto Claim Consultants, n.d.). The clause exists precisely because the ordinary claims process gives an adjuster little room to move without compelling counter-evidence, and because, absent that evidence, the low end of a valuation range costs the insurer nothing to offer — which is exactly what the AppraiselNow data above shows happening at scale, not just in one case.

6.4 Why the Claims Case Is Cleaner Than the Underwriting Case

Claims valuation asks a verifiable, backward-looking factual question: what would this specific vehicle have sold for, in this specific market, at this specific time. That question has a real answer, checkable against actual listings and sales records, independent of anyone's judgment. Underwriting, examined in Part One, asks an unverifiable, forward-looking probabilistic question: how long will this specific applicant live. No dataset can resolve that question for an individual at the time a policy is underwritten — only for a population, in aggregate, and only after the fact.

This distinction matters for what AI can be expected to do in each domain. In claims valuation, an AI tool applying the same comparable-sales methodology an insurer's own platform uses can, in principle, be checked against public data and found more accurate — precisely what occurred in the case that opens this Part, where a policyholder's own analysis of the insurer's ten selected comparable vehicles produced a valuation more than 30 percent higher, and the insurer raised its offer by a similar margin once challenged. In underwriting, an AI model can only ever produce a different probability estimate, never a verified correction, because the ground truth will not be known for decades, if ever, for any single applicant. Claims-side AI tools therefore stand on considerably firmer epistemic footing than underwriting-side ones, and the two should not be argued for with the same degree of confidence.

What the two cases share is not certainty but cost structure. In both, an incumbent institution benefits from most counterparties lacking the expertise or resources to independently verify its number, and in both, the historical remedy has been a paid specialist — an independent broker in underwriting, an independent

appraiser in claims — available mainly to those who know the specialist exists and can afford one. What AI plausibly changes is not the existence of that remedy but its cost, moving it from a several-hundred-dollar professional service toward something closer to free and immediate. That shift, not superior intelligence, is the more defensible core of the argument this paper is making.

7. Why AI Changes the Calculus Now

The same automation trend runs through both parts of this paper, though it cuts differently in each. In claims valuation, AI's role is comparatively narrow and well defined: replicate a checkable, comparable-sales methodology at a cost approaching zero, so that the verification step once reserved for policyholders who could afford an independent appraiser becomes available to any policyholder before they sign a release.

In underwriting, the same automation could go two ways. It could simply scale today's exclusionary defaults — sorting applicants into wide risk buckets defined by decades-old cohort data, at a speed no manual underwriting department could match. Or it could run the mechanisms in Section 4: experience-rated inclusion pools in real time, individualized protective-factor modeling across an entire book of business, per-applicant reconsideration on demand — all tasks that would have been prohibitively labor-intensive to run manually. The documented risk is that AI-driven underwriting instead substitutes broad, loosely validated, non-medical data — credit scores, purchasing habits, social media activity — for the medical judgment it replaces, in pursuit of speed rather than accuracy (Grant Thornton, 2023; Sidley Austin, 2023). That substitution, not an excess of health data, is the actual mechanism by which AI could make today's exclusionary pattern worse rather than better. Actuarial and consumer-advocacy research converges on this warning from different directions: industry analyses note that models trained on historical claims and underwriting data can encode and scale whatever bias already shaped who was approved, declined, or charged more in the past (KPMG, 2025), while consumer advocates describe the same dynamic, when it runs along racial lines, as a technological continuation of older discriminatory practices rather than a break from them (Public Citizen, 2022).

The business case for the more accountable path, in both domains, is not purely a matter of fairness. LIMRA's own research attributes a large share of the underwriting-side coverage gap to cost misperception and product complexity rather than declines alone (LIMRA & Life Happens, 2025), meaning a meaningful, underserved market already exists for a carrier that can profitably underwrite it. On the claims side, an insurer whose valuations can already withstand independent, publicly checkable scrutiny needs no lawsuit, no regulatory order, and no new technology to defend itself — only a methodology it is willing to show its work on. Carriers that build toward both are positioned to compete on accuracy while a regulatory environment — more than twenty states and counting on AI governance alone — continues to tighten around those that do not (NAIC, 2023; Quarles Law Firm, 2025).

8. Conclusion: The Cost of Verification

The applicant whose family history motivated Part One of this paper is not asking insurers to ignore risk, and the policyholder whose claim motivated Part Two is not asking insurers to overpay for a car. In both

cases, the request, translated into economic terms, is that an institution should bear some cost for how confidently it estimates — whether that estimate is a mortality assumption or a comparable vehicle's true sale price — in the same way it already bears cost for the opposite kind of error. An individual, after-the-fact penalty tied to an invented mortality timeline cannot deliver that without undermining the risk pooling that makes insurance affordable. A policyholder's independently produced valuation, checked against the same public data the insurer's own tool used, can deliver exactly that, immediately — because the underlying question has a real answer to check against.

That contrast is this paper's thesis in miniature. Removing the invisible subsidy that makes underwriting exclusion cost-free is a genuine, buildable project, but a probabilistic one that will always involve judgment calls, guardrails, and industry-wide coordination to avoid adverse selection. Closing the cost-of-verification gap in claims valuation is a narrower, more immediate one: the tools to check a specific factual number against public data already exist, are getting cheaper every year, and require no new law, regulation, or actuarial framework to start using today. Insurers that treat both shifts as inevitable, and compete on accuracy and transparency rather than on how effectively they can rely on asymmetric effort, are better positioned for the AI era than those that do not. That is the version of AI-powered insurance worth building — and worth insisting on, before the alternative version finishes scaling first.

References

- Addala, A., Auzanneau, M., Miller, K., et al. (2021). A decade of disparities in diabetes technology use and HbA1c in pediatric Type 1 diabetes: A transatlantic comparison. *Diabetes Care*, 44(1), 133–140. <https://pmc.ncbi.nlm.nih.gov/articles/PMC8162452/>.
- Agency Checklists. (2026). Class action filed against Commerce total loss payments rejected. <https://agencychecklists.com/2025/06/23/class-action-filed-against-commerce-total-loss-payments-rejected-76154/>.
- Akerlof, G. A. (1970). The market for “lemons”: Quality uncertainty and the market mechanism. *Quarterly Journal of Economics*, 84(3), 488–500.
- Ambalavanan, J., Isaacs, D., & Lansang, M. C. (2024). Diabetes technology: A primer for clinicians. *Cleveland Clinic Journal of Medicine*, 91(6), 353–360. <https://www.ccjm.org/content/91/6/353>.
- AppraiselNow. (2025). Independent appraisals increase insurance settlements by 34% on average [Data analysis]. Reported in Ritz Herald. <https://ritzherald.com/appraiselnow-data-reveals-independent-appraisals-increase-insurance-settlements-by-34-on-average/>.
- Artemis. (2018). What is longevity risk transfer? <https://www.artemis.bm/library/what-is-longevity-risk-transfer/>.
- Auto Claim Consultants. (n.d.). The appraisal clause: A comprehensive guide. <https://autoclaimconsultants.com/right-to-appraisal>.
- Been, R. A., Lameijer, A., Gans, R. O. B., van Beek, A. P., Kingsnorth, A. P., Choudhary, P., & van Dijk, P. R. (2024). The impact of socioeconomic factors, social determinants, and ethnicity on the utilization of glucose sensor technology among persons with diabetes mellitus: A narrative review. *Therapeutic Advances in Endocrinology and Metabolism*. <https://journals.sagepub.com/doi/10.1177/20420188241236289>.
- Buchanan Ingersoll & Rooney PC. (2025). When algorithms underwrite: Insurance regulators demanding explainable AI systems. <https://www.bipc.com/when-algorithms-underwrite-insurance-regulators-demanding-explainable-ai-systems>.

Federal Register. (2016). Regulations under the Americans with Disabilities Act. <https://www.federalregister.gov/documents/2016/05/17/2016-11558/regulations-under-the-americans-with-disabilities-act>.

Grant Thornton. (2023). Model bias rules target insurance practices. <https://www.grantthornton.com/insights/articles/insurance/2023/model-bias-rules-target-insurance-practices>.

Holland & Knight. (2025). The implications and scope of the NAIC model bulletin on the use of AI by insurers. <https://www.hklaw.com/en/insights/publications/2025/05/the-implications-and-scope-of-the-naic-model-bulletin>.

Irvin Legal. (2024). How to get the most money from insurance for your totaled car. <https://irvinlegal.com/more-money-total-loss-claim/>.

Isaacs, D., Bellini, N. J., Biba, U., Cai, A., & Close, K. L. (2021). Health care disparities in use of continuous glucose monitoring. <https://journals.sagepub.com/doi/10.1089/dia.2021.0268>.

Isitt, J. J., et al. (2022). Long-term cost-effectiveness of Dexcom G6 real-time continuous glucose monitoring system in people with type 1 diabetes in Australia. Diabetic Medicine. <https://onlinelibrary.wiley.com/doi/10.1111/dme.14831>.

Justia. (2026). Insurance bad faith law. <https://www.justia.com/injury/insurance-bad-faith/>.

KPMG. (2025). Unequal odds: Addressing and mitigating bias in life insurance. <https://kpmg.com/us/en/articles/2025/unequal-odds-mitigating-bias-life-insurance.html>.

Law360. (2026). State Farm unit can't escape undervalued vehicle dispute. <https://www.law360.com/commercialcontracts/articles/2446114>.

Lawyer Monthly. (2025). State Farm sued in NC over undervalued car payouts. <https://www.lawyer-monthly.com/2025/10/state-farm-north-carolina-class-action-2025/>.

Legal Information Institute, Cornell Law School. (n.d.). Genetic Information Nondiscrimination Act (GINA). https://www.law.cornell.edu/wex/genetic_information_nondiscrimination_act_gina.

LIMRA & Life Happens. (2025). 2025 Insurance Barometer Study. <https://www.limra.com/en/newsroom/news-releases/2025/adults-age-30-and-younger-overestimate-life-insurance-cost-by-1012-times/>.

LV=. (2026). Enhanced & impaired life annuities. <https://www.lv.com/pensions-retirement/enhanced-annuities>.

McKennon Law Group. (2026). Insurance bad faith FAQs. <https://mslawllp.com/faqs/insurance-bad-faith-faqs/>.

Memorial Merits. (2026). Life insurance with pre-existing conditions: What your agent won't explain about how underwriting actually works. <https://memorialmerits.com/life-insurance-pre-existing-conditions-underwriting-guide/>.

My Fair Claim. (2026). What every vehicle owner should know about the 2025 State Farm lawsuits. <https://myfairclaim.com/ccc-one-total-loss-valuations-state-farm-lawsuits/>.

National Association of Insurance Commissioners. (2023). Model bulletin: Use of Artificial Intelligence Systems by insurers. <https://regulations.ai/regulations/RAI-US-NA-NAICAIB-2023>.

NerdWallet. (2025). Guaranteed issue life insurance: Is it worth it? <https://www.nerdwallet.com/insurance/life/learn/guaranteed-issue-life-insurance>.

New York Life. (2025). Life insurance with pre-existing conditions. <https://www.newyorklife.com/articles/life-insurance-with-pre-existing-conditions>.

Public Citizen. (2022). Racism in, racism out. <https://www.citizen.org/article/algorithmic-racism/>.

Quarles Law Firm. (2025). Nearly half of states have now adopted NAIC model bulletin on insurers' use of AI. <https://www.quarles.com/newsroom/publications/nearly-half-of-states-have-now-adopted-naic-model-bulletin-on-insurers-use-of-ai>.

- Repairer Driven News. (2025). Total loss undervaluation lawsuit filed against State Farm in North Carolina. <https://www.repairerdrivennews.com/2025/10/20/total-loss-undervaluation-lawsuit-filed-against-state-farm-in-north-carolina/>.
- RGA. (2011). Underwriting beyond +400. On the Risk, 27(2). <https://www.rgare.com/docs/default-source/knowledge-center-articles/substandard-lives.pdf>.
- Robert L. McDorman. (2024). Total loss claims. <https://robertlmcorman.typepad.com/my-blog/total-loss/>.
- Rothschild, M., & Stiglitz, J. (1976). Equilibrium in competitive insurance markets: An essay on the economics of imperfect information. Quarterly Journal of Economics, 90(4), 629–649.
- Roze, S., Isitt, J. J., Smith-Palmer, J., et al. (2020). Long-term cost-effectiveness of the Dexcom G6 real-time continuous glucose monitoring system compared with self-monitoring of blood glucose in people with type 1 diabetes in France. Diabetes Therapy. <https://pmc.ncbi.nlm.nih.gov/articles/PMC7651823>.
- Seattle Times. (2017). Dismantling of state's health reforms in 1993 may offer lessons for Obamacare repeal. <https://www.seattletimes.com/seattle-news/politics/dismantling-of-states-health-reforms-in-1993-may-offer-lesson-for-obamacare-repeal/>.
- Sidley Austin LLP. (2023). Scrutiny of the use of AI by life insurers is increasing. <https://www.sidley.com/en/insights/publications/2023/05/scrutiny-of-the-use-of-ai-by-life-insurers-is-increasing>.
- SnapClaim. (2026). Dispute your CCC total loss valuation for a fair payout. <https://snapclaim.com/ccc-total-loss-valuation/>.
- Society of Actuaries Research Institute. (2024). AI risk essays: Mitigating biased decisions. <https://www.soa.org/4a3f62/globalassets/assets/files/resources/research-report/2024/ai-risk-essays/zuo-mitigate-biased-decision.pdf>.
- The Actuary. (2023). Unnatural bias: When AI is over-assumptive. <https://www.theactuary.com/2023/12/07/unnatural-bias-when-ai-over-assumptive>.
- U.S. Congress. (2008). Genetic Information Nondiscrimination Act of 2008, Pub. L. No. 110-233.
- Western & Southern. (2025). Understanding guaranteed issue whole life insurance. <https://www.westernsouthern.com/life-insurance/guaranteed-issue-whole-life-insurance>.