

5Qs: Logue on New Research Suggesting Many Homeowners Don't Understand Their Insurance Policy

News August 14, 2026 Author(s)Bob Needham

[ADMINISTRATIVE LAW](#)[CORPORATE AND SECURITIES LAW](#)[LAW AND SOCIAL SCIENCE](#)[TAX LAW](#)



Homeowners who take the time to read their insurance policy often don't grasp what they're reading, according to a study co-authored by Professor Kyle Logue.

Past research has shown that homeowners are not great about reading their insurance policies in the first place. But Logue and his colleagues conclude that even when homeowners do read a policy, they are very likely to misunderstand it —even to the point of believing something is covered when it explicitly is not.

Logue—the Douglas A. Khan Collegiate Professor of Law—recently answered five questions about the research:

1. How did you become interested in this particular research question?

I've been teaching insurance law since I started at Michigan, so I have a lot of experience with insurance law and insurance contracts. I give copies of insurance policies to students, and a big part of the class is reading the policies carefully—and reading cases that interpret the language and helping the students understand it.

One thing that becomes clear is that the average person, if they read their policy, probably couldn't understand it. But there isn't much empirical research on that question, and a lot turns on the answer.

Our whole system of insurance regulation assumes that if we make companies disclose things clearly, consumers can protect themselves. If that assumption is wrong, the system needs rethinking. So that's what we really wanted to focus on.

2. How did you approach this project?

We surveyed 2,500 homeowners—representative of the country—giving them scenarios that any homeowner could recognize, where something bad happens to their house or to someone in their house. Then we asked whether the loss would be covered by a standard homeowner's policy.

We divided all those people into two groups. We asked one group if the loss was covered without giving them any insurance policy language—asking them to draw on their knowledge of their own policy, which would likely be a standard homeowner’s policy.

We asked the second group the same questions, but we gave them the relevant language—usually a paragraph—which answered the question of whether there was coverage. They had the policy language right in front of them.

The intuition, which is also the assumption behind disclosure regulation, was that people with the answer in front of them should do better. We wanted to see whether that’s true.

3. What did you find?

On average, having the policy language helped a little—much less than you’d hope. But the average hides the interesting part.

For some questions, the language didn’t help at all, and for some, it actually made things worse; on those questions, people would have answered more accurately without reading anything.

On top of that, reading the language made people more confident in their answers, including the wrong ones. So the policy sometimes delivered the worst combination: a wrong answer, confidently held.

4. How would you explain this?

Our theory—which is not proven, but is consistent with our evidence—is that the problem was the structure of the language, not the wording. All homeowner’s insurance policies are subject to state regulations about readability, so wording should not be a major issue.

The problem seems to have been that they would read two-thirds or three-fourths of the provision, and that portion might grant coverage. Then they either would stop reading or they wouldn’t read the last part as carefully, and that last part would be critical. It would say,

“You’re covered, you’re covered, you’re covered”—and then at the end, it would say, “... unless you’re not covered.”

You can imagine if someone’s reading something and they think, “Oh, yeah, I understand this. This looks like I’m covered.” They stop reading, and then they’re in trouble. It gives them the wrong answer, and it also gives them excessive confidence in their wrong answer.

5. What does your research suggest for solutions to this problem?

Where there are provisions that are likely to lead people astray, they could have a checkbox that says, “I will read this to the very end.” That sounds trivial, but the whole problem is that people stop reading too early, and small interventions that keep people reading are exactly what the evidence supports testing.

Another possibility is getting assistance from artificial intelligence. The same team is now designing a follow-up study with the same structure, the same scenarios and questions, but where some participants get access to a large language model and can ask it anything they want. We’ll be able to see whether the AI helps, whom it helps, and—because we’ll record what people ask and what the model tells them—why. It’s a genuinely open question. AI assistance has helped consumers in some settings and backfired in others, and nobody has tested it on insurance.

One thing I’m not sure about is whether there should be a doctrinal solution. You could imagine courts looking at provisions that are misleading, and that insurance companies have been told are misleading, and construing the ambiguity against the insurer, or even refusing to enforce an exclusion buried where readers predictably miss it. Insurance law already has doctrines pointed in this direction; the question is whether evidence like ours should push courts to use them more aggressively.

Or, rather than using courts, you could go to the state regulators. The regulator or commissioner of insurance in each state will decide whether policy language is acceptable. You could say to the regulators, “We’ve identified some policy language that’s problematic and misleading, and we need this to be changed.”