



SUPERHUMAN

Why Synthetic Respondents Can Never ‘Outperform’ Us

By Susan Schwartz McDonald, Ph.D.

(a 4-minute read written by a certifiably human author)

A question some people are asking

A recent post by my colleague, Michael Polster, speculated wryly about whether synthetic respondents might actually be better than human respondents at predicting human behavior, and it struck a chord that’s still reverberating. At a time when we’re all trying to figure out whether synthetic respondents can approach or approximate humans as survey-takers, this provocative idea really ups the ante.

“Can synthetics ever out-predict humans?” is an intriguing question with important technical and philosophical ramifications. *But (spoiler alert), the answer to this question will always be emphatically no, even though we’ve yet to see the full extent of what increasingly better-bred synthetics are capable of achieving.* The idea of a point beyond asymptote for synthetics *emulating* humans (as opposed to complex AI models *out-forecasting* humans) upends the very premise upon which synthetic respondents are based. ‘Humans doing human’ is a kind of absolute, never to be exceeded, like the speed of light. Here’s why.

The good-enough synthetic: utterly ‘sincere’ but no smarter than the humans they mimic

Let’s start with one of the first arguments made in favor of synthetic respondents, before we knew as much as we do now about what we can expect from synthetics. In those early days, the claim was made that even though synthetics will be prone to modeling error, they are incapable of insincerity. Unlike human respondents, they are never motivated to prank or posture when answering questions.

That argument fails to recognize that since synthetics are trained to perform like humans, their answers will reflect the data distribution patterns used by the training models that gave them “life.” Synthetic respondents cannot be earnest because they have no self-awareness. They are simply built to mimic humans, not improve on them, which means that if humans dissemble in surveys, then synthetics will too, even if unintentionally. What animates synthetics is imitation rather than motivation.

To err by over-prediction is human—and thus, *synthetically human* too

For much the same reason, it would be unreasonable to expect synthetics to out-perform humans in their predictions. But let’s first parse several different ways of thinking about prediction in order to fully support that case. One source of survey prediction error with which we all contend is the tendency of human respondents to over-predict new product adoption based on factors like survey acquiescence, ebullience, concept overpromise and limited information about the options they are evaluating. Most of us are accustomed to correcting for over-statement, although we also recognize the paradox that when products or ideas are truly revolutionary, we are likely to encounter the opposite sort of error: significant under-prediction. Most of us could not have fully anticipated even a decade ago what kind of lives we would live online in 2026, and more recently, how much human activity we would willingly outsource to AI. Revolutionary change occurs at a pace that defies imagination and thus offsets survey overstatement.

If we keep things simple, though, we’d have to assume that synthetic respondents – unless specifically post-trained to dampen their enthusiasm about “product concepts” or “target profiles” – will similarly overstate. The better their training, the more likely synthetics are to display the failures of humans when using scales to predict their future behavior. And if we attempt to correct for that problem in our synthetic respondent models, we will inadvertently, but inevitably, create new sources of synthetic divergence from humans that serve us poorly in the end. Differences in the way humans and synthetics use rating scales need to be programmed *out* of these models, not strategically built in, if we are to make synthetics more effective human proxies. We don’t want to benefit from lucky breaks produced by random modeling error even if it makes our predictions more accurate, since it means we can expect unlucky ones too. We want neither.

Synthetic respondents – as opposed to prediction models – should be judged by the stated mission we’ve established for them: to mimic humans rather than improve upon them. We are not looking for synthetics to be over-achievers. If the goal is *replicability* of human responses using synthetics, an occasional “improvement” over human powers of prediction should be considered neither a boon nor an excuse to rely on synthetics. It’s a form of misbehavior known as modeling error. Synthetics that could manage to come consistently closer to future reality in their behavioral or societal predictions are no longer human surrogates. They are in the evolutionary process of becoming a *forecasting model*, more powerful than even the most prescient of beings, human or synthetic. At which point, they are very impressive butterflies, no longer caterpillars. But we need caterpillars too.

The value of human predictions, even if they’re wrong

We ask humans to make more than one kind of prediction, and we tend to use different types of human predictions for different purposes. When we ask humans to predict what they are *personally* likely to do in the future, our goal is to gauge the appeal of new products or identify new opportunities, and we will take those predictions seriously, even if not literally. On the other hand, when we ask humans to make *general predictions*, like whether AI will take human jobs or how they expect climate change to affect them, we are merely gauging their mood in the moment, and we cannot expect to be guided by their view of the future.

So why do we need survey predictions in a world where turbo-charged forecast models might do a fine job of reading our future based on vast and variegated data inputs? Consumer surveys have long been a critical ingredient in market forecasting, and when effectively interpreted in context, they ought to make forecast models better. AI modeling does not obviate the value of asking humans to predict *themselves* (even if imperfectly) in the service of product development or market-shaping strategies. There is always great value in understanding what consumers think they will do and how they believe changing circumstances might affect that. The highest calling of a synthetic respondent, and the test of success for the concept, is to respond like a human, not outthink or out-predict it.

About the Author



Susan Schwartz McDonald, Ph.D.
Chief Executive Officer
215.496.6850
smcdonald@naxionthinking.com

Susan’s career focus has been on the development and protection of robust brands, and the research methodologies needed to support them. Much of her career has been spent consulting to clients on the development of life science and technology commercialization strategy but she also has extensive experience with consumer products and financial services. She has contributed to the evolution of many standard research techniques, and she writes frequently on industry topics and issues of broader interest. In deference to changing times, the snappier-paced McDonald Minute replaces Susan’s long-form Smartmouth blog, but with the same goal of connecting cultural themes to business challenges. She holds MA and PhD degrees from UPenn’s Annenberg School of Communication.

About NAXION

NAXION is a nimble, broadly resourced, employee-owned boutique that relies on advanced research methods, data integration, and sector-focused experience to guide strategic business decisions that shape the destiny of brands. Our century-long history of innovation has helped to propel the insights discipline and continues to inspire contributions to the development and effective application of emerging data science techniques. For information on what’s new at NAXION and how we might help you with your marketing challenges, please visit <https://www.naxionthinking.com/>