

Background

Blue Rose Research, in partnership with the Center for Shared AI Prosperity, conducts regular polls of policies that lawmakers, academics, and others are proposing to regulate and guide AI development. In this blog post, we'll explain how we measure how popular AI policies are, and why we think it's an effective way to understand what Americans think.

The Problem

Traditional issue and policy polls tend to encounter two specific problems, which inflate support numbers in unrealistic ways.

- First, there's a desire to incorporate the benefits of the policy into the description of the policy itself to help voters better understand. But this inherently bakes in one-sided messaging into the policy description, elevating support in an unrealistic way.
- Second, there's a phenomenon in survey research called "acquiescence bias," where respondents are more likely to choose the more positive option – "agree," "support," etc – than they are to disagree or oppose. This also has the effect of inflating baseline levels of support on issue and policy polling.

How We Solve It

Rather than trying to strip out acquiescence bias and fully remove all positive framing from our surveys, we choose an approach that provides voters with the policy details that they would see in a more neutral media framing, while simultaneously adding back the full persuasive lens from both sides by explaining the reasons supporters and opponents. This helps us determine which policies hold up better to public debate – especially on a complex topic where voters may not have existing perspectives or be familiar with the details of the pros and cons of any given policy.

Survey Language	Explanation
<i>Some policymakers are proposing to...</i>	Start with a neutral framing so the proposal isn't anchored to either side.
<i>make it illegal for AI companies to lie to the government about their safety practices. Under the plan, large firms that knowingly make false or misleading safety claims would face civil fines. The fines would grow with the size of the firm and with repeat offenses.</i>	Provide a description of the policy that is as neutral and non-partisan as possible.
<i>Some people support this policy [...because safety rules mean nothing if firms can fake their reports with no penalty. AI firms should face the same fines for lying that banks and drug firms do].</i>	Indicate that some people support this policy. Respondents are randomly chosen to receive an alternative that includes a policy-specific argument from supporters, to help assess levels of support in the face of arguments.
<i>Some people oppose this policy [...because experts often disagree on safety risks. Punishing firms over such claims would turn honest judgment calls into crimes].</i>	Indicate that some people oppose this policy. Respondents are randomly chosen to receive an alternative that includes a policy-specific argument from opponents, to help assess levels of support in the face of counterarguments.

Who do you agree with more? Supporters or opponents of this policy?

Then finally, ask which side they agree with. This helps to counter acquiescence bias. Support/agree/yes options tend to face acquiescence bias, but asking voters to choose between two options helps to avoid this dynamic.

How To Read Our Results

Once we've collected our survey data, we try to answer two questions with our results:

- 1) *How popular is this policy?*
 - a) We answer this question by estimating what percentage of the population supports this question
- 2) *How popular is this policy compared to other AI policies?*
 - a) After estimating how popular a policy is, we compare that policy to every other policy we've tested using this methodology. This lets us rank similar policies to identify the most popular policies to build a progressive vision of the United States as the world is transformed by AI.

Over the coming months, we'll be sharing our research about how popular AI policies are in our surveys. Keep an eye out, and come back to this post if you have any questions about how we learn about Americans' attitudes on AI.