

APPENDIX I

2025 NSCG Methodological Experiments: Minimum Detectable Differences

Minimum Detectable Differences for the 2025 NSCG Methodological Experiments

I. Background

This appendix provides minimum detectable differences for the proposed sample sizes in the 2025 NSCG methodological experiments.

Text Message Experiment

- Returning sample treatment group – 10,000 cases will receive text messages earlier in the day and 10,000 cases will receive text messages later in the day

Plain Language Letters Experiment:

- New sample treatment group – 10,000 cases will receive the experimental plain language letters

II. Minimum Detectable Differences Equation and Definitions

To calculate the minimum detectable difference between two response rates with fixed sample sizes, we used the formula from Snedecor and Cochran (1989) for determining the sample size when comparing two proportions.

$$\delta \geq \left((Z_{\alpha^*/2} + Z_{\beta})^2 \left(\frac{p_1(1-p_1)}{n_1} + \frac{p_2(1-p_2)}{n_2} \right) D \right)^{1/2}$$

where:

- δ = minimum detectible difference
- α^* = alpha level adjusted for multiple comparisons
- $Z_{\alpha^*/2}$ = critical value for set alpha level assuming a two-sided test
- Z_{β} = critical value for set beta level
- p_1 = proportion for group 1
- p_2 = proportion for group 2
- D = design effect due to unequal weighting
- n_1 = sample size for a single treatment group or control
- n_2 = sample size for a second treatment group or control

The alpha level of 0.10 was used in the calculations. The beta level was included in the formula to inflate the sample size to decrease the probability of committing a type II error. The beta level was set to 0.10.

The estimated proportion for the groups was set to 0.50 for the sample size calculations. This conservative approach minimizes the ability to detect statistically significant differences.

Design effects represent a variance inflation factor due to sample design when compared to a simple random sample. Because all experiment samples and the control will be representative, the weight distributions should be similar throughout all samples, negating the need to include a design effect. We do not expect to see a weight-based or sampling-based effect on response in any of the samples. However, for the sake of completeness, minimum detectable differences were calculated both ways, including and ignoring the design effect.¹

III. Pairwise Comparisons and the Bonferroni Adjustment

Because numerous treatment groups can be compared (as discussed in the research questions listed in Appendix I), α^* is adjusted to account for the multiple comparisons. The Bonferroni adjustment reduces the overall α by the number of pairwise comparisons so when multiple pairwise comparisons are conducted the overall α will not suffer. The formula is:

$$\alpha^* = \frac{\alpha}{n_c}$$

The adjusted alpha α^* is calculated by dividing the overall target α by the number of pairwise comparisons, n_c . It is worth noting that, despite being commonly used, the Bonferroni adjustment is very conservative, actually reducing the overall α below initial targets. An example showing how the overall α is calculated using an alpha level of 0.10, the Bonferroni adjustment, and 5 pairwise comparisons follows:

$$\alpha_{overall} = 1 - (1 - \alpha^*)^{n_c}$$

$$\alpha_{overall} = 1 - (1 - 0.02)^5 = 0.0961 < 0.100$$

$\alpha_{overall}$ is the resulting overall α after the Bonferroni correction is applied;
 $\alpha_{target} = 0.100$, and is the original target α level;
 $n_c = 5$, and is the number of comparisons (i.e., Appendix I research questions)
 $\alpha^* = \alpha_{overall} / n_c = 0.009$, and is the Bonferroni-adjusted α

In this example, the Bonferroni adjustment overcompensates for multiple comparisons, making it more likely that a truly significant effect will be overlooked.

Sample sizes were provided in section I of this appendix and are used in the formula. All minimum detectable differences using the Bonferroni adjustment were calculated and are summarized at the end of this appendix in table form.

¹ Design effects were calculated by examining the weight variation present in all cases in the 2023 NSCG new sample (5.22), and the returning sample (4.49).

IV. A Model-Based Alternative to Multiple Comparisons

Rather than relying on the Bonferroni adjustment for multiple comparisons, effects on response, cost per case or other outcome variables could be modeled simultaneously to determine which treatments have a significant effect on response.

All sample cases, auxiliary sample data, and treatments are included in the model below, which predicts a given treatment's effect on response rate (or other outcome variable).

$$y = \beta_0 + \beta_1 I + \alpha X + \varepsilon$$

Assuming response rate is the outcome variable of interest:

$\mathcal{Y}$ is the average response propensity (response rate) for the entire sample;

β_0 is the intercept for the model;

β_1 is a vector of effects, one for each treatment;

I is a vector of indicators to identify a treatment in β_1 ;

α is a scalar vector;

X is a matrix of auxiliary frame or sample data; and

ε is an error term.

Once data collection is complete, the average response propensity is equal to the response rate. In the simplest case, no treatment has any effect (the 2nd term would drop out), and no auxiliary variables explain any of the variation in response propensities (the 3rd term would drop out). In that case, the average of the response propensities, and thus the response rate, would just equal:

$$y = \beta_0 + \varepsilon$$

However, a more complicated model gives information about each treatment's effect (2nd term) while taking into account sample characteristics (3rd term) that might augment or reduce the effect of a given treatment.

As a simple example, ignore the error term, and assume the overall mean response propensity was 72%. Also, assume the mean response propensity for a given treatment group was 83%. If only terms 1 and 2 were included in the model (no sample characteristics accounted for), the given treatment appears to have increased the response propensity by 11%. However, if the

sample was poorly designed, or if a variable not included in the sample design turned out to be a good predictor of response, there is value in adding the 3rd term. If auxiliary information added by the 3rd term shows that the cases in a particular sample group are 5% more likely to respond than the average sample case (because of income, internet penetration, age, etc.), this would suggest that while the treatment group had a response propensity 11% higher than the average, 5% came from sample person characteristics, and only 6% of that increase was really due to the treatment.

This method has several benefits over the multiple comparisons method. First, the number of degrees of freedom taken up by the model is the number of treatment groups plus one for the intercept, which is fewer than the number of pairwise comparisons that might be conducted.

Second, because confidence intervals are calculated around the β_1 values, it is easier to observe a treatment's effect on the outcome measures. Third, variables can be controlled for in the model, making significant results more meaningful. While we are striving to ensure the experimental samples are as representative (and as similar) as possible, the ability to add other variables to the model helps control for unintended effects.

The method uses response propensities, not the actual response rate. While the mean response propensity after the last day of data collection equals the overall response rate, it is important to note how the propensity models are built. If they are weighted models, weighted response propensities should be used in this model. The weights could be added as one of the auxiliary variables included in the X matrix.

V. Comments

It is worth noting from the calculations below that even using the Bonferroni adjustment, and conducting all pairwise comparisons, a difference of 6% - 8% in outcome measures should be large enough to appear significant, when the design effect is excluded from the calculations. Because the experimental samples are all systematic random samples, and should have similar sample characteristics and weight distributions, excluding the design effect seems appropriate.

Minimum Detectable Differences for the 2025 NSCG Methodological Experiments

Minimum Detectable Difference Equation for Response Rates

$$\delta \geq \left((Z_{\alpha^*/2} + Z_{\beta})^2 \left(\frac{p_1(1-p_1)}{n_1} + \frac{p_2(1-p_2)}{n_2} \right) \times deff \right)^{1/2}$$

δ	=	minimum detectible difference
$\star\delta$	=	minimum detectible difference without using design effect
α^*	=	alpha level adjusted for multiple comparisons (Bonferroni)
$Z_{\alpha^*/2}$	=	critical value for set alpha level assuming a two-sided test
Z_{β}	=	critical value for set beta level
p_1	=	proportion for group 1
p_2	=	proportion for group 2
$deff$	=	design effect due to unequal weighting
n_1	=	sample size for group 1
n_2	=	sample size for group 2

Text Message Experiment (returning sample)

10,000 cases in each Experimental Group

α^*	=	0.100	
$Z_{\alpha^*/2}$	=	1.645	
Z_{β}	=	1.282	$\delta = 0.0495$
p_1	=	0.672	$\star\delta = 0.0194$
p_2	=	0.672	
$deff$	=	6.5	
n_1	=	10,000	
n_2	=	10,000	

Plain Language Letters Experiment (new sample)

10,000 cases in Experimental Group

α^*	=	0.100	
$Z_{\alpha^*/2}$	=	1.645	
Z_{β}	=	1.282	$\delta = 0.0401$
p_1	=	0.598	$\star\delta = 0.0157$
p_2	=	0.598	
$deff$	=	6.5	
n_1	=	49,200	
n_2	=	10,000	