

Public Trust in Banks and Financial Institutions: A Knowledge Mining Approach

Soha Arian, Kasra Akbari, Eman Ajmal, & Joseph Martinez

EPPS 6323.001 - Knowledge Mining - Spring 2026

Professor Karl Ho

May 08, 2026

Executive Summary

Public confidence in banks and financial institutions shapes how people view economic security, trust in institutions, and the stability of the financial system. This project examines how confidence in banks differs by income, education, age, and survey year in the United States, using data from the General Social Survey. It also assesses whether an AI-generated summary can accurately report statistical findings without overstating results, overlooking uncertainty, or making unsupported causal statements. The main research questions are: How does confidence in banks differ by income, education, and age in the United States, and do these patterns remain stable over time? Can an AI-generated summary accurately reflect the statistical results?

This project uses a quantitative approach to analyze data from the General Social Survey cumulative file. The main outcome is confidence in banks and financial institutions. The analysis includes weighted descriptive statistics, ordered logistic regression, and predicted probabilities. The primary deliverable is an interactive Shiny application. Users can filter the data by year, age, sex, race, education, income, weighting, and period. The app also allows users to fit the ordered-logit model, view odds ratios, examine predicted probabilities, and export model results for AI summarization. The final Shiny app includes all these features and uses a filtered sample of 24,879 observations by default.

The main result is that most people report 'only some' confidence in banks, regardless of income group. While there are income differences, they are small. The highest income group is less likely to report high confidence in banks compared to the lowest income group. In the ordered-logit model, the highest income group has an odds ratio of 0.880 with a 95% confidence interval from 0.830 to 0.934, showing a statistically significant negative association. The middle income group, age, and education are not statistically significant because their confidence intervals include 1. Survey year shows a stronger effect than the demographic variables, especially after 2008. This suggests that public confidence in banks is influenced more by historical and institutional factors than by income alone.

This project adds to knowledge mining by transforming a long-term public opinion dataset into an interactive analytical tool. It also demonstrates the importance of carefully reviewing AI-generated summaries of statistical results. An effective AI summary should state that income differences exist but are small, that age and education are not significant in this

model, and that the analysis is observational, not causal. More broadly, the project provides a reproducible process for moving from raw survey data to cleaned variables, statistical modeling, interactive visualization, and controlled AI-assisted summarization.

Problem Statement

Public trust in banks matters because financial institutions rely on more than formal regulation, balance sheets, and market activity. They also depend on public confidence. When people believe banks are stable, fair, and useful, they may be more willing to use banking services, save money, borrow, invest, and participate in the formal financial system. When confidence is low, people may view banks with suspicion or avoid financial institutions altogether. This makes confidence in banks not only an economic issue, but also a question of institutional legitimacy and public trust.

The project begins from the idea that public confidence may not be evenly distributed across the population. People in different income groups may experience banks differently because they have different levels of financial security, different exposure to fees, different access to credit, and different relationships with financial institutions. Education may also matter because people with more schooling may have more financial knowledge, stronger ability to compare financial products, or greater skepticism toward institutions. Age may matter because older and younger respondents have lived through different financial events, including inflation, recessions, banking crises, and changes in digital banking. These differences make it important to examine confidence in banks by demographic group rather than only looking at one overall trend.

Historical context is also central to the problem. Confidence in banks may shift across time because banks are closely connected to economic crises, regulatory debates, and public perceptions of financial fairness. The 2008 financial crisis is especially relevant because it raised questions about risk-taking, regulation, public bailouts, and the relationship between ordinary households and large financial institutions. Prior research on institutional confidence shows that public trust is complex and can decline unevenly across different institutions and historical periods. More recent research on institutional trust also emphasizes that broad changes in public

confidence cannot be understood without looking at both demographic patterns and historical shocks.

This project addresses a gap between broad discussions of financial trust and the need for a reproducible, data-driven analysis of confidence in banks over time. Many public discussions describe confidence in banks as high, low, rising, or falling, but those descriptions often do not show whether the pattern differs by income, education, or age. The project therefore uses General Social Survey data to examine the relationship between public confidence and demographic variables across survey years. The project also uses a Shiny app to make the analysis more transparent and accessible to readers who may not use Stata or R. Together, these observations allow the project to study both public confidence in banks and the communication of statistical evidence.

Background & Literature

Public confidence in banks has changed over time because the banking system itself has changed. In the late 1970s and early 1980s, high inflation and rising interest rates put pressure on households, especially lower-income households. At the same time, the Reagan era marked a major shift toward financial deregulation. Restrictions on banks and savings-and-loan institutions were loosened, which helped create the conditions for the savings-and-loan crisis of the late 1980s and early 1990s. For many people, this crisis made bank failure feel real and showed that regulation, risk, and public trust were closely connected. Because of this, confidence in banks should not be understood only through individual characteristics like income, education, and age. It also needs to be understood through the historical period in which people were surveyed.

The 1990s and 2000s continued this pattern of financial change. The Clinton years included economic growth, but also further deregulation through the Gramm-Leach-Bliley Act of 1999, which allowed commercial banks, investment banks, and insurance companies to become more connected. The dot-com crash and the 2008 financial crisis later showed the risks of this more complex financial system. After 2008, public confidence in banks dropped sharply because many people saw the crisis as an institutional failure, not just an economic downturn. Dodd-Frank attempted to restore trust through stronger oversight and consumer protections, but confidence recovered slowly and remained below earlier levels. The Trump administration later

moved in the opposite direction by partially rolling back Dodd-Frank through the Economic Growth, Regulatory Relief, and Consumer Protection Act of 2018, which reduced oversight requirements for some smaller and mid-sized banks. This historical background helps explain why year effects matter in the analysis: people's trust in banks reflects not only who they are, but also the political, regulatory, and economic events they lived through.

The existing literature supports this connection between financial history and public trust. Sapienza and Zingales (2012) argued that the contraction in economic activity between late 2008 and early 2009 was driven partly by a collapse in trust in financial institutions, with trust reaching multi-decade lows. Guiso, Sapienza, and Zingales (2008) also showed that trust is important for financial participation, especially stock market participation. Cross-national studies find similar patterns. Hasan and Weill (2019) analyzed trust in banks across 52 countries and found that individual characteristics predict trust, but so does whether a country experienced the global financial crisis. Stevenson and Wolfers (2011) showed that countries with larger increases in unemployment experienced sharper declines in confidence. Knell and Stix (2015) found that personal experience with bank failures has a strong and lasting negative effect on trust. Together, these studies show that demographics matter, but their effects are limited unless they are understood alongside the larger historical and regulatory environment.

Methodology

This project applied a quantitative knowledge-mining approach to examine public confidence in banks and financial institutions in the United States. The analysis relied on data from the General Social Survey (GSS), a long-standing national survey that collects information on social attitudes, institutions, demographics, and public opinion. The primary focus was on confidence in banks and financial institutions. Specifically, the project investigated whether confidence levels vary by income, education, age, and survey year and whether an AI-generated summary can accurately represent the statistical results.

The analysis combined weighted descriptive statistics and ordered logistic regression, with support from an interactive Shiny application. Descriptive statistics were used initially to identify basic patterns in the data. Ordered logistic regression was then applied because the outcome variable, confidence in banks, is ordinal. Respondents indicated whether they had

hardly any, only some, or a great deal of confidence, which represents a progression from lower to higher confidence. This structure made an ordered model more suitable than a standard linear regression.

The project also incorporated an evaluation of AI-generated summaries. After developing the statistical model, the results were exported into a standardized prompt for the AI tool. This step assessed whether the AI could summarize the statistical findings accurately, without overstating results, making causal claims, or overlooking uncertainty. This process addressed the project's second research question regarding the accuracy of AI-generated summaries in reflecting model outcomes.

Data Source and Variables

The dataset used for this project was the General Social Survey cumulative file covering the available years in the project data. The Shiny application is titled “Public Confidence in Banks & Financial Institutions (GSS 1972–2024),” but the usable app output shown in the project file displays a working survey-year range from 1975 to 2018 after the relevant variables and missing values are handled. This means the broader GSS file may cover more years, but the final analysis depends on the years where the needed variables are available and usable.

The main outcome variable was confidence in banks and financial institutions. In the original GSS data, this variable was named "confinan." The project recoded it into a clearer, ordered variable called "conf_bank" or "conf_banks," depending on whether the work was being done in Stata or R. The response categories were recoded so that the lowest value represented “Hardly any,” the middle value represented “Only some,” and the highest value represented “A great deal.”

The main predictor variables were income, education, age, and survey year. Income was measured in two ways during the project. In the early Stata analysis, income was grouped into four within-year categories: lowest, low-middle, high-middle, and highest. In the final Shiny application, income was simplified into three groups: lowest, middle, and highest. This change made the interactive app easier to read and made the predicted probability output clearer for

users. The app also allowed users to choose between within-year income groups and pooled inflation-adjusted income groups.

The project also included demographic filters for sex, race, education, age, and survey year. These filters were included so users could narrow the sample and see whether the results changed for different groups. The app also allowed users to choose whether to apply GSS survey weights. Survey weights were important because they help make the survey results better reflect the population represented by the GSS.

Step-by-Step Procedure

1. Select the required variables from the GSS dataset.

These include survey year, respondent ID, age, education, degree, income, respondent income, confidence in banks, and the GSS survey weight. This step was used to confirm that the needed variables were present in the dataset before cleaning or modeling began. In Stata, this began with the command:

```
describe year id age educ degree income rincome confinan wtssall
```

2. Check the confidence-in-banks variable using a frequency table.

This shows how many respondents selected each answer category before recoding. This mattered because the team needed to understand the original coding before changing it into the final ordered outcome. In Stata, this was done using:

```
tab confinan
```

3. Remove observations that were missing the main variables needed for the analysis.

This includes missing values for confidence in banks, age, education, survey year, income, and the survey weight. This step made sure that the model used complete cases only. The analysis only included respondents who had usable information for the outcome, predictors, and weighting variable. In Stata, the command was:

```
drop if missing(confinan, age, educ, year, rincome, wtssall)
```

4. Create income groups within each survey year.

This was done because income values mean different things in different years. A person's income in the 1970s cannot be compared directly to a person's income in the 2010s without adjustment. To handle this, the Stata analysis sorted respondents by year and income, counted how many respondents were in each year, ranked respondents by income within that year, and then placed them into income groups. The early Stata version created four groups. The final Shiny version used three groups to make the app easier to interpret.

5. Recode the confidence-in-banks variable into an ordered outcome variable.

The original GSS coding was changed so that the values moved logically from low confidence to high confidence. The final categories were "Hardly any," "Only some," and "A great deal." This was necessary because the ordered logistic regression model assumes the outcome has a meaningful order.

6. Created age groups for descriptive tables.

Respondents were grouped into age categories: under 30, 30–49, 50–64, and 65 or older. This helped the team examine whether confidence patterns differed by age in a way that was easier to explain than using every single age value.

7. Apply the GSS survey weight, wtssall.

This was done so the descriptive tables and models would better reflect the survey design. In Stata, the command was:

```
svyset [pweight=wtssall]
```

After the weight was set, the team produced weighted descriptive tables showing confidence in banks by income group, degree, and age group.

8. Estimate an ordered logistic regression model.

The model tested whether income group, education, age, and survey year were associated with confidence in banks. In Stata, the model was:

svy: ologit conf_bank i.income_q i.degree c.age i.year, or

In the final Shiny application, the model formula was:

```
conf_banks ~ inc_q + age + educ + factor(year)
```

The Shiny model used MASS::polr() in R to estimate ordered logistic regression. The model output included coefficients, odds ratios, confidence intervals, and predicted probabilities. The uploaded app output shows this exact formula in the model tab.

This model generated predicted probabilities. Instead of saying only that one group has higher or lower odds, predicted probabilities estimate the chance that each income group falls into each confidence category. The Shiny app showed predicted probabilities for “Hardly any,” “Only some,” and “A great deal” across the lowest, middle, and highest income groups. In the uploaded output, “Only some” confidence is around 59.6% to 59.9% across income groups, while “A great deal” confidence decreases from 19.4% in the lowest income group to 17.5% in the highest income group.

Building the Shiny application helped the analysis to be explored interactively. The app allowed users to change the year range, age range, sex, race, degree, income measure, period, and weighting option. It also included separate tabs for overview charts, income patterns, model results, predicted probabilities, and AI export. The uploaded Shiny app output shows these tabs and filters, including the “Fit model” button and the “Export for AI” tab.

We created an AI export feature. After the model was fit, the app produced a standardized text export containing the filter settings, model formula, model output, odds ratios, predicted probabilities, and a fixed AI prompt. The purpose was to make sure that every AI summary was based on the same statistical information. This made the AI evaluation more consistent and reduced the risk that the AI tool would rely on outside assumptions.

Finally, the team evaluated the AI-generated summary. The AI prompt instructed the tool to identify significant and non-significant findings, describe the direction of relationships, explain predicted probability differences in plain language, avoid causal claims, and avoid exaggeration. The team then planned to compare the AI summary against the actual model

output to see whether it correctly reported direction, significance, size, uncertainty, omissions, and factual errors. This matches the project slideshow, which states that the AI summary would be compared directly to model results and scored for direction, size, significance, omissions, exaggerations, and factual errors .

Shiny Application Methodology

The Shiny application served as the project’s main interactive deliverable. It was not only a website; it was also a way to reproduce and inspect the analysis. The app used the cleaned GSS data files `gss_banks.rds` and `api_annual.rds`. These files had to be placed in the same folder as the Shiny app.R file before the app could run. If the files were missing, the app stopped and told the user to run the preparation script first.

Step-by-Step Procedure

1. Open RStudio and install the needed R packages: shiny, MASS, dplyr, tidyr, ggplot2, and scales.
2. Place app.R, gss_banks.rds, and api_annual.rds in the same project folder.
3. After opening app.R, click “Run App.”
4. Once the app opens, leave the default filters selected or adjust the sample by year, age, sex, race, or degree.
5. Choose the income measure, decide whether to apply survey weights, select the period mode, and click “Fit model.”

The app produces four main outputs. The first output is an overview of confidence in banks, including a bar chart and a trend chart over time. The second output shows confidence by income group. The third output shows the ordered-logit model results, including odds ratios and confidence intervals. The fourth output shows predicted probabilities by income group. The fifth tab exports the model results into a standardized AI prompt.

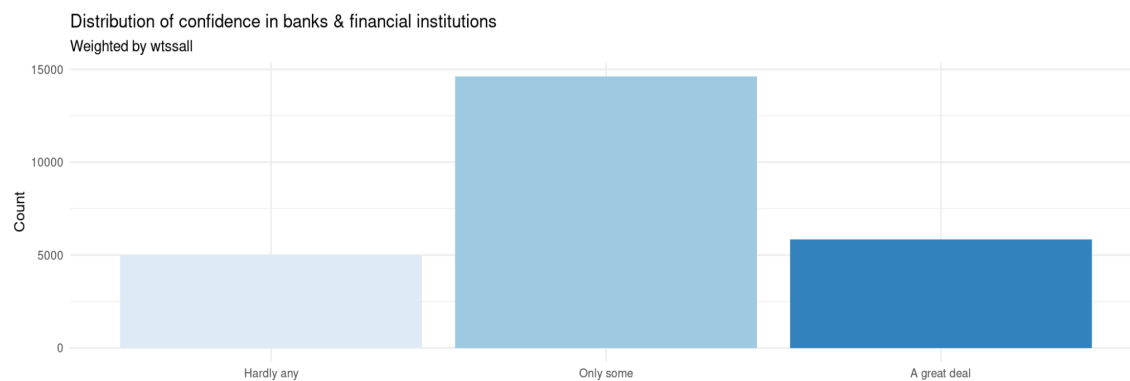
Findings

The analysis produced three main types of findings: descriptive patterns in confidence levels, changes in confidence over time, and model-based estimates of how income, age,

education, and survey year relates to confidence in banks. The findings are based on the cleaned General Social Survey confidence-in-banks dataset used in the Shiny application. Under the default filters, the app used a filtered sample of 24,879 observations and applied the ordered-logit model $\text{conf_banks} \sim \text{inc_q} + \text{age} + \text{educ} + \text{factor}(\text{year})$ with GSS survey weights (wtssall) . The results are organized around the research question rather than around the software steps, so the section focuses on what the evidence shows about public confidence in banks.

The clearest descriptive finding is that “Only some” confidence is the dominant response category. Figure 1 shows the overall distribution of confidence in banks and financial institutions from the Shiny app’s Overview tab. The tallest bar is the “Only some” category, while “Hardly any” and “A great deal” appear lower. This means respondents were not mostly choosing the strongest positive or strongest negative response. Instead, the most common position was moderate confidence. This matters because public confidence in banks is not best described as simply high or low. The main pattern is that many respondents occupy a middle position, suggesting limited or cautious confidence rather than full trust or complete distrust.

Figure 1. Distribution of confidence in banks and financial institutions

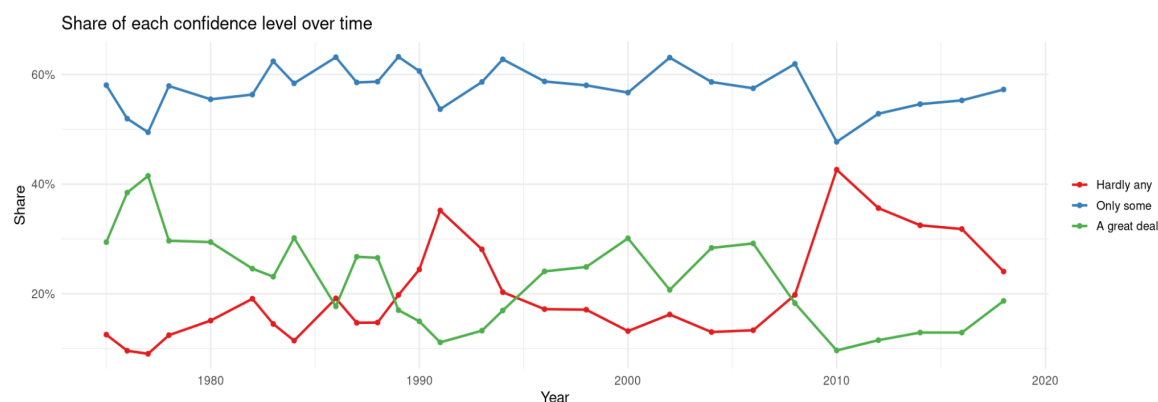


Source: Shiny app Overview tab. This figure shows the weighted distribution of responses to the GSS confidence-in-banks variable. The chart shows that “Only some” confidence is the most common response category.

The trend graph adds a historical pattern to the descriptive finding. Figure 2 shows the share of each confidence category over time. The “Only some” category remains high across much of the period, while “Hardly any” and “A great deal” shift more visibly across survey years. The graph shows that confidence in banks is not completely stable over time. There are periods where the

share of respondents reporting "Hardly any" confidence rises and periods where the share reporting "A great deal" falls. This supports the project's argument that confidence in banks should be studied historically, not only by income or demographic group. The Shiny app's overview output includes this trend chart alongside the distribution chart, making time one of the main visual findings of the project.

Figure 2. Share of each confidence level over time



Source: Shiny app Overview tab. This figure shows the weighted share of respondents reporting “Hardly any,” “Only some,” and “A great deal” of confidence across survey years. The graph shows that public confidence in banks changes over time rather than staying flat.

The income comparison shows that differences across income groups are present, but they are modest. Figure 3 shows the predicted probability of each confidence level by income group. Across the lowest, middle, and highest income groups, “Only some” confidence remains close to 60%. The lowest income group has a predicted probability of 59.9% for “Only some” confidence; the middle group also has 59.9%, and the highest group has 59.6%. These values are nearly identical. The graph therefore shows that income does not strongly change the most common response category. Respondents across all three income groups are still most likely to report “Only some” confidence.

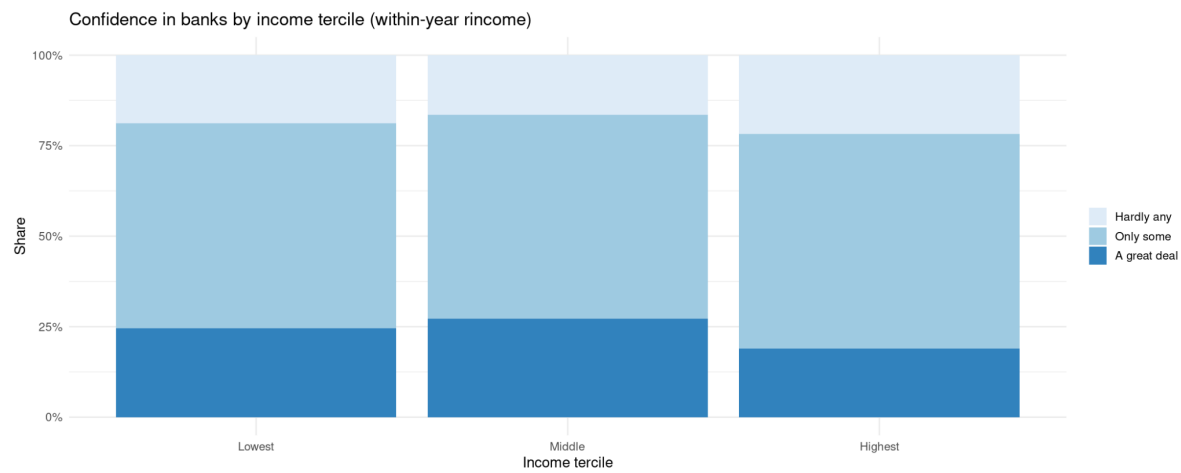


Figure 3. Predicted probability of each confidence level by income group

Source: Shiny app Predicted Probabilities tab. This figure shows predicted probabilities for “Hardly any,” “Only some,” and “A great deal” confidence across income terciles while holding age, education, and year at sample values.

Table 1: Predicted Probability of Confidence Category by Income Group

Income group	Hardly any	Only some	A great deal
Lowest	20.7%	59.9%	19.4%
Middle	21.1%	59.9%	19.0%
Highest	22.9%	59.6%	17.5%

Table 1 gives the same pattern in numeric form. The highest income group is slightly more likely to report “Hardly any” confidence and slightly less likely to report “A great deal” of confidence than the lowest income group. The predicted probability of “Hardly any” confidence increases from 20.7% in the lowest income group to 22.9% in the highest income group. The predicted probability of “A great deal” of confidence decreases from 19.4% in the lowest income group to 17.5% in the highest income group. These differences are small, so the correct finding is not that income sharply divides public confidence. A more accurate statement is that higher

income is associated with slightly lower predicted confidence in banks, while the overall income-group differences remain limited.

The ordered-logit model provides statistical support for part of the income pattern. The highest income group has a coefficient of -0.127550 and a t value of -4.2296. The odds ratio for the highest income group is 0.880, with a 95% confidence interval from 0.830 to 0.934 . Because the confidence interval does not include 1, this result is statistically significant. In plain language, respondents in the highest income group had lower odds of reporting a higher confidence category than respondents in the lowest income group, holding age, education, and survey year constant.

The middle-income group does not show a statistically clear difference from the lowest-income group. Its odds ratio is 0.976, with a 95% confidence interval from 0.908 to 1.049. Because that interval includes 1, the model does not provide enough evidence to say that the middle income group differs from the lowest income group. This is important because the income finding is not a simple pattern where every increase in income clearly lowers confidence. The statistically significant income result appears specifically for the highest income group compared with the lowest group.

Table 2: Selected Ordered-Logit Odds Ratios

Predictor	Odds ratio	95% CI lower	95% CI upper	Finding
Middle income	0.976	0.908	1.049	Not statistically significant
Highest income	0.880	0.830	0.934	Statistically significant negative association
Age	1.001	0.999	1.003	Not statistically significant
Education	0.993	0.984	1.002	Not statistically significant

Analyzing the Full Sample

Age and education are not statistically significant in the model shown. Age has an odds ratio of 1.001 with a 95% confidence interval from 0.999 to 1.003, while education has an odds ratio of 0.993 with a 95% confidence interval from 0.984 to 1.002. Both confidence intervals include 1. This means the model does not provide clear evidence that age or education is associated with confidence in banks after income and survey year are included. This does not mean age and education can never matter. It only means they are not statistically significant in this specific model and sample.

The survey year shows a stronger pattern than the demographic predictors. Several years' indicators have odds ratios below 1 and statistically significant confidence intervals, especially around and after the 2008 period. The model output shows odds ratios of 0.599 for 2008, 0.220 for 2010, 0.293 for 2012, 0.341 for 2014, 0.348 for 2016, and 0.535 for 2018. These values suggest that respondents in those survey years had lower odds of reporting higher confidence compared with the model's reference year. The size of these year effects is much larger than the income effect, which supports the finding that historical periods are a major part of the confidence pattern.

The graph-based and model-based findings point in the same direction. The descriptive graphs show that “Only some” confidence is the dominant response and that confidence changes across time. The predicted probability graph shows that income differences exist but remain small. The ordered-logit model confirms that the highest income group is statistically significant, while the middle income group, age, and education are not statistically significant in the shown model. Together, these results answer the research question by showing that confidence in banks varies somewhat by income and more strongly by survey year, but the demographic differences should not be overstated.

Across Decades

The education results are inconsistent across the decade-based samples. In 1975–1980, education has a negative coefficient and a t-value of -0.8753, which is not statistically significant using the $|t| > 1.96$ rule. In 1980–1990, the coefficient becomes positive, but the t-value remains

non-significant at 0.8505. In 1990–2000, education is still positive, with a larger but still non-significant t-value of 1.4257. The main change occurs in 2000–2010, when education becomes statistically significant and positive, with a t-value of 2.1854. However, in 2010–2018, the direction reverses: education has a negative coefficient and a statistically significant t-value of -4.4479. This pattern suggests that education does not have a stable relationship with confidence in banks across the samples. Its statistical significance appears only in the two most recent periods, but the direction differs between them, moving from a positive association in 2000–2010 to a negative association in 2010–2018.

The age results show a clearer pattern of change over time. In 1975–1980, age is positive and statistically significant, with a t-value of 9.6245. In 1980–1990, age remains positive and statistically significant, with a t-value of 4.9044. Beginning in 1990–2000, the direction changes: age becomes negative and statistically significant, with a t-value of -5.6107. This negative association continues in 2000–2010, where the t-value is -6.6129, and in 2010–2018, where the t-value is -2.5891. The major turning point is therefore between the 1980–1990 and 1990–2000 samples. Before that point, older respondents had higher estimated odds of reporting greater confidence in banks. After that point, older respondents had lower estimated odds of reporting greater confidence, holding the other model variables constant.

The income-group results are mixed and should not be overstated. For the middle-income group, the coefficient is negative in every period, but it is statistically significant only in 1980–1990, with a t-value of -2.2665, and in 2010–2018, with a t-value of -2.2667. In the other periods, the middle-income t-values do not exceed the |1.96| threshold. For the highest-income group, the coefficient is negative and statistically significant in 1975–1980, with a t-value of -2.0511. It remains negative but falls just below conventional significance in 1980–1990, with a t-value of -1.9063. The direction then changes in 1990–2000 and 2000–2010, where the highest-income coefficient is positive but not statistically significant. In 2010–2018, the highest-income coefficient is close to zero and clearly non-significant, with a t-value of 0.0640. Overall, the income results show some statistically significant differences in specific periods, but not a stable pattern across the full sequence of samples.

The AI-summary component produced an additional communication finding. The Shiny app includes an “Export for AI” tab, and the project plan states that the team would use a fixed prompt to generate an AI summary, compare the summary to the model results, and score whether it correctly reported direction, size, significance, omissions, exaggerations, and factual errors. The findings from the statistical analysis show exactly why this step matters. A faithful AI summary should say that “Only some” confidence is the most common response, income differences are modest, the highest income group is statistically significant, age and education are not statistically significant, and survey year appears important. A summary that simply says “income affects trust in banks” would be incomplete because it would hide the small size of the predicted probability differences and the non-significant results for several predictors.

Table 3: AI Summary Evaluation Criteria

Criterion	What the AI summary should do
Direction	Correctly state whether each association is positive, negative, or unclear
Significance	Distinguish statistically significant from non-significant findings
Magnitude	Use predicted probabilities to describe the size of differences
Causality	Avoid claiming that income, age, education, or year causes confidence
Uncertainty	Explain when the model does not provide clear evidence
Completeness	Include both significant and non-significant results
Tone	Avoid dramatic or opinion-based language

The findings support a careful answer to the research question. Public confidence in banks does vary by income, but the income differences are small in predicted probability terms. The highest income group is associated with lower odds of higher confidence, while the middle-income group is not statistically different from the lowest-income group. Age and education are not statistically significant in the model shown. Survey year appears more

important than the demographic predictors, especially around and after the 2008 period. The graphs make this pattern easier to see, while the model output confirms which differences are statistically supported.

Grader: Arian					Grader: Martinez				
ID	Response Name	ChatGP T Score	Claude Score	Copilot Score	ID	Response Name	ChatGP T Score	Claude Score	Copilot Score
1.1	Full Sample	4	5	4	1.1	Full Sample	3	4	2
1.2	Pre-2008	4	5	4	1.2	Pre-2008	4	4	4
1.3	Post-2008	4	5	4	1.3	Post-2008	4	2	4
1.4	Pre vs Post 2008	4	4	3	1.4	Pre vs Post 2008	5	5	4
2.1	1975-1980	4	5	4	2.1	1975-1980	4	5	4
2.2	1980-1990	4	5	4	2.2	1980-1990	5	5	3
2.3	1990-2000	4	5	4	2.3	1990-2000	4	5	4
2.4	2000-2010	4	4	4	2.4	2000-2010	4	5	4
2.5	2010-2018	4	5	4	2.5	2010-2018	4	5	3
2.Tot	2.1 to 2.5 exports	4	4	3	2.To t	2.1 to 2.5 exports	5	6	3

Grader: Arian					Grader: Martinez				
avg		4	4.7	3.8	avg		4.2	4.6	3.5

Source: *AI Prompt Response Scorecard. The scorecard shows the evaluations of two financial professionals.*

Three AI models were evaluated by Martinez and Arian on a scale from 1 to 5, with 1 meaning completely inaccurate and 5 meaning very accurate. The three models tested were ChatGPT, Claude, and Copilot. Out of the three, ChatGPT was the only model that did not hallucinate information. Claude performed well overall, but it made several mistakes when discussing post-2008 data, especially by giving inaccurate information about the time period being analyzed. Copilot also had accuracy issues, including one response where it incorrectly stated that a t-value of 3.28 was not significant.

This does not mean ChatGPT was perfect. Both ChatGPT and Copilot struggled at times to clearly explain whether the results showed significant evidence of a positive or negative correlation, especially when comparing differences between variables. Overall, Claude received the highest average scores, with a 4.6 from Martinez and a 4.7 from Arian. ChatGPT ranked second, with scores of 4.2 from Martinez and 4.0 from Arian. Copilot ranked last, with scores of 3.5 from Martinez and 3.8 from Arian.

One of the most interesting results came from Claude's response to the five-decade export prompt. Instead of giving a shorter one- to two-paragraph answer, Claude produced a five-page Word document. Martinez considered this the strongest response in the entire experiment and even gave it a 6, despite the original grading scale only going up to 5. Arian also rated Claude highly overall, but the evaluation shows that the two graders did not always agree on which specific tasks each AI handled best.

Overall, Martinez and Arian generally agreed that Claude performed the best, ChatGPT was second, and Copilot was the weakest. However, their scores also show that there was not complete consistency in which tasks each grader believed the AI models handled most effectively. This suggests that while Claude had the strongest overall performance, the quality of

each model's response depended heavily on the type of prompt and the specific expectations of the grader.

Recommendations and Future Directions

Future research should examine why the relationship between education and confidence in banks changes across time rather than treating education as having one stable association with public trust. The decade-based results suggest that education is not consistently related to confidence in the same direction, which raises questions about whether broader historical, economic, or institutional conditions shape how education matters. We can explore how these patterns align with major financial events, changes in banking regulation, public exposure to financial news, or differences in financial knowledge across education groups; and test can be run on whether education interacts with survey year, income, age, or political and economic attitudes to better understand when education is associated with higher or lower confidence. The current analysis is observational, therefore, these findings should not be interpreted as showing that education causes changes in confidence. Instead, they point to education as a variable whose relationship with confidence in banks may depend on historical context and should be studied more directly in future work.

Conclusion

Taken together, the decade-based models show that age has the most consistent evidence of statistical significance, but its direction changes over time. Age is positively associated with confidence in banks in the earlier samples and negatively associated in the later samples. Education is less stable: it is not statistically significant in the first three periods, becomes positively significant in 2000–2010, and then becomes negatively significant in 2010–2018. Income-group differences appear in some periods, especially for the highest-income group in 1975–1980 and the middle-income group in 1980–1990 and 2010–2018, but these results do not form a consistent decade-by-decade trend. These findings show changes in statistical significance and direction, but they do not by themselves establish large substantive effects. The safest interpretation is that the relationships between education, age, income, and confidence in banks vary across historical periods, with age showing the clearest turning point and education showing a later shift in significance and direction.

References

- Algan, Y., & Cahuc, P. (2014). Trust, growth, and well-being: New evidence and policy implications. In P. Aghion & S. N. Durlauf (Eds.), *Handbook of economic growth* (Vol. 2A, pp. 49–120). Elsevier.
- Almond, G. A., & Verba, S. (1963). *The civic culture: Political attitudes and democracy in five nations*. Princeton University Press.
- Davis, J. A., Smith, T. W., & Marsden, P. V. (2024). *General Social Surveys, 1972–2022: Cumulative codebook*. NORC at the University of Chicago.
- Dincer, O. C., & Usaner, E. M. (2010). Trust and growth. *Public Choice*, 142(1), 59–67.
- Fine, B. (2001). *Social capital versus social theory: Political economy and social science at the turn of the millennium*. Routledge.
- Fukuyama, F. (1995). *Trust: The social virtues and the creation of prosperity*. Free Press.
- Guiso, L., Sapienza, P., & Zingales, L. (2004). The role of social capital in financial development. *American Economic Review*, 94(3), 526–556.
- Guiso, L., Sapienza, P., & Zingales, L. (2008). Trusting the stock market. *The Journal of Finance*, 63(6), 2557–2600.
- Hetherington, M. J. (2005). *Why trust matters: Declining political trust and the demise of American liberalism*. Princeton University Press.
- Knack, S., & Keefer, P. (1997). Does social capital have an economic payoff? A cross-country investigation. *The Quarterly Journal of Economics*, 112(4), 1251–1288.
- Lin, N. (2001). *Social capital: A theory of social structure and action*. Cambridge University Press.
- McCullagh, P. (1980). Regression models for ordinal data. *Journal of the Royal Statistical Society: Series B (Methodological)*, 42(2), 109–127.
- Mishler, W., & Rose, R. (2001). What are the origins of political trust? Testing institutional and cultural theories in post-communist societies. *Comparative Political Studies*, 34(1), 30–62.

- Newton, K., & Norris, P. (2000). Confidence in public institutions: Faith, culture, or performance? In S. J. Pharr & R. D. Putnam (Eds.), *Disaffected democracies: What's troubling the trilateral countries?* (pp. 52–73). Princeton University Press.
- Putnam, R. D., Leonardi, R., & Nanetti, R. Y. (1993). *Making democracy work: Civic traditions in modern Italy*. Princeton University Press.
- Putnam, R. D. (2000). *Bowling alone: The collapse and revival of American community*. Simon & Schuster.
- R Core Team. (2024). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing. <https://www.R-project.org/>
- Rothstein, B., & Stolle, D. (2008). The state and social capital: An institutional theory of generalized trust. *Comparative Politics*, 40(4), 441–459.
- Sapienza, P., & Zingales, L. (2012). A trust crisis. *International Review of Finance*, 12(2), 123–131.
- Stevenson, B., & Wolfers, J. (2011). Trust in public institutions over the business cycle. *American Economic Review*, 101(3), 281–287.
- Uslaner, E. M. (2002). *The moral foundations of trust*. Cambridge University Press.
- Venables, W. N., & Ripley, B. D. (2002). *Modern applied statistics with S* (4th ed.). Springer.
- Zak, P. J., & Knack, S. (2001). Trust and growth. *The Economic Journal*, 111(470), 295–321.

AI Usage Statement:

This project used generative AI tools in two distinct capacities: as an object of study and as a writing aid.

As an object of study. Three large language models; OpenAI's ChatGPT, Anthropic's Claude, and Microsoft Copilot were prompted to summarize the ordered logistic regression results exported from the Shiny application. These outputs were then evaluated by two members of the research team (Arian and Martinez) against fixed criteria for direction, statistical significance, magnitude, causal restraint, uncertainty, completeness, and tone. The AI-generated summaries themselves are presented as data in the Findings section and are not represented as the authors' own writing. Full prompts and unedited model responses are linked in the appendix.

As a writing aid. The authors used generative AI tools for limited drafting support, including help organizing the literature review, clarifying technical descriptions of the ordered-logit specification, and copy-editing for grammar and flow. All factual claims, statistical results, interpretation of model output, and substantive arguments were produced and verified by the authors. AI tools were not used to generate the data analysis, the statistical code, the Shiny application architecture, or the substantive conclusions of the paper. Any AI-assisted text was reviewed and revised by the authors, who take full responsibility for the final content.

The authors did not use AI tools to fabricate citations, generate references, or produce content that was presented as original analysis without disclosure.

Appendices

AI Prompt:

I am a graduate student working on a research project. Below are the statistical results from my ordered logistic regression analysis examining confidence in banks using General Social Survey data from 1972 to 2024. The output includes model coefficients, statistical significance levels, and predicted probabilities.

I need you to act as an academic research assistant and generate a concise, objective summary of the findings.

Identify statistically significant associations between income, education, age, and survey year and confidence in banks. Clearly state the direction of each significant relationship. Use predicted probabilities to explain the magnitude of differences where appropriate. Clearly distinguish between statistically significant and non-significant findings. Do not make causal claims, as the analysis is observational. Do not exaggerate effect sizes or use dramatic language. Do not include normative, evaluative, or opinion-based statements. Base the summary only on the statistical output provided. Do not introduce external information or assumptions. If the output does not support a claim, state this explicitly rather than inferring beyond the data.

The goal is to produce an accurate academic summary that faithfully represents the statistical findings without overstating or misinterpreting the results.

AI Prompt Responses Link:

https://docs.google.com/document/d/1PBz1Lkd2qiJJsvGFDObcmdDkMkcVRuPh3WVaY_ywC8o/edit?usp=sharing

AI Prompt Response Scorecard:

<https://docs.google.com/spreadsheets/d/1YZJuyixuiNezRbifa2tVqiW8QzRJvIlcbIIanWmSna8/edit?usp=sharing>

Shiny Application: https://martinez2k18.shinyapps.io/epps_6323_final_project/

Shiny Application Code:

```
# -----  
  
# app.R  
# Shiny app: Public Confidence in Banks & Financial Institutions  
# General Social Survey, 1972-2024 (GSS cumulative file)  
#  
# REQUIRES gss_banks.rds and cpi_annual.rds -- run prep_data.R  
# once before launching. Click "Run App" in RStudio.  
#  
# Features:  
# - Income terciles (within-year, relative) OR pooled terciles  
#   of inflation-adjusted realrinc  
# - Year, age, sex, race, degree filters; weighted/unweighted  
# - Pre-2008 / 2008+ / full-sample / compare modes  
# - Ordered-logit model, odds ratios, predicted probabilities  
# - "Export for AI" button: downloadable standardized text  
#   blob for the team's AI-summary scoring pipeline  
# -----  
  
library(shiny)  
library(MASS)    # polr() -- load BEFORE dplyr so select() isn't masked  
library(dplyr)  
library(tidyr)  
library(ggplot2)  
library(scales)
```

```

# ---- Load prepared data -----
if (!file.exists("gss_banks.rds")) {
  stop("gss_banks.rds not found. Run prep_data.R first.")
}
if (!file.exists("cpi_annual.rds")) {
  stop("cpi_annual.rds not found. Run prep_data.R first.")
}
gss <- readRDS("gss_banks.rds")
cpi_annual <- readRDS("cpi_annual.rds")

year_min <- min(gss$year, na.rm = TRUE)
year_max <- max(gss$year, na.rm = TRUE)
age_min <- floor(min(gss$age, na.rm = TRUE))
age_max <- ceiling(max(gss$age, na.rm = TRUE))

sex_choices <- levels(gss$sex)
race_choices <- levels(gss$race)
degree_choices <- levels(gss$degree)

CUT_YEAR <- 2008 # pre/post threshold: 2008 and later = "Post"

# ---- UI -----
ui <- fluidPage(
  titlePanel("Public Confidence in Banks & Financial Institutions (GSS 1972-2024)",
  sidebarLayout(
    sidebarPanel(
      width = 3,
      h4("Filters"),
      sliderInput("years", "Survey year range:",

```

```

        min = year_min, max = year_max,
        value = c(year_min, year_max),
        step = 1, sep = ""),
sliderInput("agerange", "Age range:",
        min = age_min, max = age_max,
        value = c(age_min, age_max), step = 1),
checkboxGroupInput("sex", "Sex:",
        choices = sex_choices, selected = sex_choices),
checkboxGroupInput("race", "Race:",
        choices = race_choices, selected = race_choices),
checkboxGroupInput("degree", "Highest degree:",
        choices = degree_choices, selected = degree_choices),
checkboxInput("weighted",
        "Apply GSS survey weights (wtssall)", value = TRUE),
hr(),
h4("Income measure"),
radioButtons("income_measure", NULL,
        choices = c(
                "Within-year terciles of rincome (relative rank)" = "within",
                "Pooled terciles of realrinc (constant dollars)" = "pooled"
        ),
        selected = "within"),
hr(),
h4("Period"),
radioButtons("period_mode", NULL,
        choices = c(
                "Full sample" = "full",
                "Pre-2008 only" = "pre",
                "2008+ only" = "post",
                "Compare pre-2008 vs 2008+" = "compare"

```

```

    ),
    selected = "full"),
  hr(),
  h4("Ordered-logit model"),
  helpText("The model takes 10-30 seconds to fit. ",
    "Adjust filters first, then click."),
  actionButton("run_model", "Fit model", class = "btn-primary"),
  hr(),
  textOutput("n_text")
),
mainPanel(
  width = 9,
  tabsetPanel(
    id = "maintabs",
    tabPanel("Overview",
      br(),
      plotOutput("overview_bar", height = "320px"),
      plotOutput("trend_plot", height = "340px")),
    tabPanel("By income",
      br(),
      plotOutput("income_plot", height = "420px"),
      br(),
      h4("Row percentages: income tercile x confidence"),
      tableOutput("income_table")),
    tabPanel("Model",
      br(),
      helpText("Formula: conf_banks ~ inc_q + age + educ + factor(year)",
        uiOutput("model_output_ui"),
        br(),
        h4("Odds ratios with 95% CI"),

```

```

        uiOutput("or_output_ui")),
    tabPanel("Predicted probabilities",
        br(),
        helpText("Predicted probabilities at mean age, mean education, ",
            "and median year within each sample."),
        plotOutput("pred_plot", height = "440px"),
        br(),
        tableOutput("pred_table")),
    tabPanel("Export for AI",
        br(),
        helpText("Standardized export of model results + scoring prompt, ",
            "for the team's AI-summary evaluation pipeline. ",
            "Fit the model first on the 'Model' tab, then return here."),
        downloadButton("download_export", "Download export (.txt)",
            class = "btn-primary"),
        br(), br(),
        h4("Preview"),
        verbatimTextOutput("export_preview"))
    )
)
)
)

# ---- Server -----
server <- function(input, output, session) {

  # -- filtered reactive -----
  filtered <- reactive({
    req(input$sex, input$race, input$degree)
    df <- gss %>%

```

```

filter(year >= input$years[1], year <= input$years[2],
       age >= input$agerange[1], age <= input$agerange[2],
       sex %in% input$sex,
       race %in% input$race,
       degree %in% input$degree,
       !is.na(educ), !is.na(age))

if (input$income_measure == "within") {
  df <- df %>%
    filter(!is.na(rincome)) %>%
    group_by(year) %>%
    mutate(
      p50 = quantile(rincome, 0.50, na.rm = TRUE),
      p90 = quantile(rincome, 0.90, na.rm = TRUE),
      inc_q = case_when(
        rincome < p50 ~ 1,
        rincome >= p50 & rincome < p90 ~ 2,
        rincome >= p90 ~ 3,
        TRUE ~ NA_integer_
      )
    ) %>%
    ungroup()
} else {
  df <- df %>%
    filter(!is.na(realrinc)) %>%
    mutate(
      p50 = quantile(realrinc, 0.50, na.rm = TRUE),
      p90 = quantile(realrinc, 0.90, na.rm = TRUE),
      inc_q = case_when(
        realrinc < p50 ~ 1,

```

```

    realrinc >= p50 & realrinc < p90 ~ 2,
    realrinc >= p90 ~ 3,
    TRUE ~ NA_integer_
  )
)
}

df %>%
  mutate(inc_q = factor(
    inc_q,
    levels = 1:3,
    labels = c("Lowest", "Middle", "Highest")
  )) %>%
  filter(!is.na(inc_q))
})

# Apply the period_mode to produce the dataframe(s) the model uses.
# Returns either a single df or a named list of pre/post for compare mode.
period_data <- reactive({
  df <- filtered()
  switch(input$period_mode,
    "full" = df,
    "pre" = df %>% filter(year < CUT_YEAR),
    "post" = df %>% filter(year >= CUT_YEAR),
    "compare" = list(
      pre = df %>% filter(year < CUT_YEAR),
      post = df %>% filter(year >= CUT_YEAR)
    )
  )
})

```

```

output$n_text <- renderText({
  pd <- period_data()
  if (is.list(pd) && !is.data.frame(pd)) {
    paste0("Pre-2008: n = ", format(nrow(pd$pre), big.mark = ","),
          "\n2008+:  n = ", format(nrow(pd$post), big.mark = ","))
  } else {
    paste0("Filtered sample: n = ", format(nrow(pd), big.mark = ","))
  }
})

# -- Overview tab -----
output$overview_bar <- renderPlot({
  df <- filtered()
  validate(need(nrow(df) > 0, "No rows match your filters."))
  p <- ggplot(df, aes(x = conf_banks, fill = conf_banks))
  p <- if (input$weighted) {
    p + geom_bar(aes(weight = wtssall))
  } else {
    p + geom_bar()
  }
  p + scale_fill_brewer(palette = "Blues", direction = 1, guide = "none") +
    labs(title = "Distribution of confidence in banks & financial institutions",
         subtitle = if (input$weighted) "Weighted by wtssall" else "Unweighted",
         x = NULL, y = "Count") +
    theme_minimal(base_size = 13)
})

output$trend_plot <- renderPlot({
  df <- filtered()

```

```

validate(need(nrow(df) > 0, ""))
trend <- df %>%
  group_by(year, conf_banks) %>%
  summarise(w = if (input$weighted) sum(wtssall, na.rm = TRUE) else n(),
    .groups = "drop_last") %>%
  mutate(share = w / sum(w)) %>%
  ungroup()
p <- ggplot(trend, aes(x = year, y = share, color = conf_banks)) +
  geom_line(size = 1) + geom_point(size = 1.5) +
  scale_y_continuous(labels = percent) +
  scale_color_brewer(palette = "Set1") +
  labs(title = "Share of each confidence level over time",
    x = "Year", y = "Share", color = NULL) +
  theme_minimal(base_size = 13)
if (input$period_mode == "compare") {
  p <- p + geom_vline(xintercept = CUT_YEAR - 0.5,
    linetype = "dashed", color = "gray40") +
    annotate("text", x = CUT_YEAR, y = 0.02,
      label = "2008", hjust = -0.1, color = "gray40")
}
p
})

# -- By-income tab -----
output$income_plot <- renderPlot({
  df <- filtered()
  validate(need(nrow(df) > 0, "No rows match your filters."))
  p <- ggplot(df, aes(x = inc_q, fill = conf_banks))
  p <- if (input$weighted) {
    p + geom_bar(aes(weight = wtssall), position = "fill")
  }
})

```

```

    } else {
      p + geom_bar(position = "fill")
    }
  p + scale_y_continuous(labels = percent) +
    scale_fill_brewer(palette = "Blues", direction = 1) +
    labs(title = paste0("Confidence in banks by income tercile (",
      if (input$income_measure == "within")
        "within-year rincome" else "pooled realrinc",
      ")"),
      x = "Income tercile", y = "Share", fill = NULL) +
    theme_minimal(base_size = 13)
  })

output$income_table <- renderTable({
  df <- filtered()
  validate(need(nrow(df) > 0, ""))
  tab <- if (input$weighted) {
    df %>%
      group_by(inc_q, conf_banks) %>%
      summarise(w = sum(wtssall, na.rm = TRUE), .groups = "drop_last") %>%
      mutate(pct = w / sum(w)) %>%
      ungroup() %>%
      dplyr::select(-w)
  } else {
    df %>%
      count(inc_q, conf_banks) %>%
      group_by(inc_q) %>%
      mutate(pct = n / sum(n)) %>%
      ungroup() %>%
      dplyr::select(-n)
  }
})

```

```

    }
    tab %>%
      pivot_wider(names_from = conf_banks, values_from = pct) %>%
      mutate(across(where(is.numeric), ~ percent(.x, accuracy = 0.1)))
  })

# -- Model fit helper -----
fit_polr <- function(df, use_weights) {
  if (use_weights) {
    polr(conf_banks ~ inc_q + age + educ + factor(year),
          data = df, weights = wtssall, Hess = TRUE)
  } else {
    polr(conf_banks ~ inc_q + age + educ + factor(year),
          data = df, Hess = TRUE)
  }
}

# -- Model (event-driven) -----
mod_result <- eventReactive(input$run_model, {
  pd <- period_data()

  if (is.list(pd) && !is.data.frame(pd)) {
    validate(need(nrow(pd$pre) > 100 && nrow(pd$post) > 100,
                  "Need at least 100 rows in each period to compare. "))
    withProgress(message = "Fitting ordered-logit (both periods) ...", value = 0.2, {
      mod_pre <- fit_polr(pd$pre, input$weighted)
      incProgress(0.5)
      mod_post <- fit_polr(pd$post, input$weighted)
    })
    list(mode = "compare",

```

```

      models = list(pre = mod_pre, post = mod_post),
      data = pd)
} else {
  validate(need(nrow(pd) > 100,
    "Need at least 100 filtered rows to fit the model.))
  withProgress(message = "Fitting ordered-logit ...", value = 0.5, {
    mod <- fit_polr(pd, input$weighted)
  })
  list(mode = "single", model = mod, data = pd)
}
})

# -- Summary + OR rendering (handles single vs compare) -----
summarise_model <- function(mod) capture.output(print(summary(mod)))

or_table <- function(mod) {
  ct <- suppressMessages(confint(mod))
  data.frame(
    Term = names(coef(mod)),
    OR = exp(coef(mod)),
    Lower95 = exp(ct[, 1]),
    Upper95 = exp(ct[, 2]),
    row.names = NULL
  )
}

output$model_output_ui <- renderUI({
  req(input$run_model > 0)
  res <- mod_result()
  if (res$mode == "compare") {

```

```

fluidRow(
  column(6, h4("Pre-2008"),
    verbatimTextOutput("model_summary_pre")),
  column(6, h4("2008+"),
    verbatimTextOutput("model_summary_post"))
)
} else {
  verbatimTextOutput("model_summary_single")
}
})

output$model_summary_single <- renderPrint({
  summary(mod_result()$model)
})
output$model_summary_pre <- renderPrint({
  summary(mod_result()$models$pre)
})
output$model_summary_post <- renderPrint({
  summary(mod_result()$models$post)
})

output$or_output_ui <- renderUI({
  req(input$run_model > 0)
  res <- mod_result()
  if (res$mode == "compare") {
    fluidRow(
      column(6, h5("Pre-2008"), tableOutput("or_pre")),
      column(6, h5("2008+"), tableOutput("or_post"))
    )
  } else {

```

```

    tableOutput("or_single")
  }
})

output$or_single <- renderTable({ or_table(mod_result())$model }, digits = 3)
output$or_pre  <- renderTable({ or_table(mod_result())$models$pre }, digits = 3)
output$or_post <- renderTable({ or_table(mod_result())$models$post }, digits = 3)

# -- Predicted probabilities -----
make_preds <- function(mod, df) {
  newdat <- data.frame(
    inc_q = factor(levels(df$inc_q), levels = levels(df$inc_q)),
    age  = mean(df$age, na.rm = TRUE),
    educ = mean(df$educ, na.rm = TRUE),
    year = median(df$year, na.rm = TRUE)
  )
  pred <- predict(mod, newdata = newdat, type = "probs")
  cbind(newdat["inc_q"], as.data.frame(pred))
}

pred_df <- reactive({
  res <- mod_result()
  if (res$mode == "compare") {
    pre <- make_preds(res$models$pre, res$data$pre) %>% mutate(period =
"Pre-2008")
    post <- make_preds(res$models$post, res$data$post) %>% mutate(period =
"2008+")
    bind_rows(pre, post)
  } else {
    make_preds(res$model, res$data) %>% mutate(period = "Full")
  }
})

```

```

    }
  })

output$pred_plot <- renderPlot({
  validate(need(input$run_model > 0,
    "Click 'Fit model' first to see predicted probabilities."))
  long <- pred_df() %>%
    pivot_longer(cols = c("Hardly any", "Only some", "A great deal"),
      names_to = "outcome", values_to = "prob") %>%
    mutate(outcome = factor(outcome,
      levels = c("Hardly any", "Only some", "A great deal")))

  p <- ggplot(long, aes(x = inc_q, y = prob, group = interaction(outcome, period),
    color = outcome, linetype = period)) +
    geom_line(size = 1) + geom_point(size = 2.5) +
    scale_y_continuous(labels = percent) +
    scale_color_brewer(palette = "Set1") +
    labs(title = "Predicted probability of each confidence level",
      subtitle = "Holding age, education, and year at sample means",
      x = "Income tercile", y = "Predicted probability",
      color = NULL, linetype = NULL) +
    theme_minimal(base_size = 13)

  # When not in compare mode, hide the redundant linetype legend.
  if (mod_result()$mode != "compare") {
    p <- p + guides(linetype = "none")
  }
  p
})

```

```

output$pred_table <- renderTable({
  req(input$run_model > 0)
  pred_df() %>%
    mutate(across(where(is.numeric), ~ percent(.x, accuracy = 0.1)))
})

# -- Export for AI -----
build_export <- function() {
  res <- mod_result()

  # Filter summary -----
  hdr <- c(
    "=====",
    "GSS Confidence in Banks -- Ordered Logit Export for AI",
    paste("Generated:", format(Sys.time(), "%Y-%m-%d %H:%M:%S %Z")),
    "=====",
    "",
    "FILTER SETTINGS",
    paste0(" Year range:   ", input$years[1], "-", input$years[2]),
    paste0(" Age range:    ", input$agerange[1], "-", input$agerange[2]),
    paste0(" Sex:         ", paste(input$sex, collapse = ", ")),
    paste0(" Race:        ", paste(input$race, collapse = ", ")),
    paste0(" Degree:       ", paste(input$degree, collapse = ", ")),
    paste0(" Survey weights: ",
      if (input$weighted) "Yes (wtssall)" else "No"),
    paste0(" Income measure: ",
      if (input$income_measure == "within")
        "Within-year terciles of rincome"

```

```

else
  "Pooled terciles of realrinc (constant dollars)",
paste0(" Period mode:   ", input$period_mode),
""
)

# Outcome distribution helper -----
outcome_dist <- function(df) {
  d <- if (input$weighted) {
    df %>% group_by(conf_banks) %>%
      summarise(w = sum(wtssall, na.rm = TRUE), .groups = "drop") %>%
      mutate(pct = w / sum(w))
  } else {
    df %>% count(conf_banks) %>% mutate(pct = n / sum(n))
  }
  paste0("   ", format(d$conf_banks, width = 14),
    " ", percent(d$pct, accuracy = 0.1))
}

model_block <- function(label, mod, df) {
  coef_mat <- summary(mod)$coefficients
  or      <- or_table(mod)
  preds   <- make_preds(mod, df)

  c(
    paste0("---- ", label, " ----"),
    paste0(" Observations:   ", format(nrow(df), big.mark = ",")),
    " Outcome distribution:",
    outcome_dist(df),
    ""
  )
}

```

```

" COEFFICIENTS (log-odds scale):",
capture.output(print(round(coef_mat, 4))),
"",
" ODDS RATIOS (95% CI):",
capture.output(print(
  data.frame(OR = round(or$OR, 3),
    Lower95 = round(or$Lower95, 3),
    Upper95 = round(or$Upper95, 3),
    row.names = or$Term))),
"",
" PREDICTED PROBABILITIES (at mean age/educ, median year):",
capture.output(print(
  preds %>%
    mutate(across(where(is.numeric), ~ round(.x, 3))),
  row.names = FALSE)),
""
)
}

body <- if (res$mode == "compare") {
  c(
    "MODEL",
    " Formula: conf_banks ~ inc_q + age + educ + factor(year)",
    " Method: Ordered logistic regression (MASS::polr)",
    " Two models fit separately on pre-2008 and 2008+ samples.",
    "",
    model_block("PRE-2008", res$models$pre, res$data$pre),
    model_block("2008+", res$models$post, res$data$post)
  )
} else {

```

```

c(
  "MODEL",
  " Formula: conf_banks ~ inc_q + age + educ + factor(year)",
  " Method: Ordered logistic regression (MASS::polr)",
  "",
  model_block("FULL MODEL", res$model, res$data)
)
}

```

```
prompt <- c(
```

```
"=====",
```

```
"STANDARDIZED AI PROMPT",
```

```
"=====",
```

```

"You are summarizing the results of an ordered logistic",
"regression that models confidence in banks and financial",
"institutions in the United States, using data from the",
"General Social Survey (1972-2024). Confidence is ordered:",
 "\"Hardly any\" < \"Only some\" < \"A great deal\".",
 "",
"Using ONLY the statistical output above, produce a concise",
"summary (3-5 sentences) that:",
" 1. Reports the direction of each predictor's estimated",
"    association (do not imply causation).",
" 2. Notes statistical significance (treat |t| > 1.96 as",
"    significant at the 0.05 level).",
" 3. Describes the magnitude of predicted-probability",
"    differences across income terciles in plain terms",
"    (e.g., \"about 3 percentage points\").",

```

```

" 4. Avoids: unsupported causal claims, exaggeration of",
"   effect size, omission of statistical uncertainty, and",
"   biased or normative language.",
"",
"Return only the summary text."
)

paste(c(hdr, body, prompt), collapse = "\n")
}

export_text <- reactive({
  req(input$run_model > 0)
  build_export()
})

output$export_preview <- renderPrint({
  validate(need(input$run_model > 0,
    "Fit the model first on the 'Model' tab."))
  cat(export_text())
})

output$download_export <- downloadHandler(
  filename = function() {
    paste0("gss_banks_ai_export_",
      format(Sys.time(), "%Y%m%d_%H%M%S"), ".txt")
  },
  content = function(file) {
    writeLines(export_text(), file)
  }
)

```

```
}
```

```
# ---- Run -----  
shinyApp(ui = ui, server = server)
```

Data Scrapping Stata Code

```
describe year id age educ degree income rincome confinan wtssall
```

```
tab confinan
```

```
drop if missing(confinan, age, educ, year, rincome, wtssall)
```

```
sort year rincome
```

```
by year: gen N_year = _N
```

```
by year: gen rank_income = _n
```

```
gen income_q = .
```

```
by year: replace income_q = 1 if rank_income <= N_year*0.25
```

```
by year: replace income_q = 2 if rank_income > N_year*0.25 & rank_income <= N_year*0.50
```

```
by year: replace income_q = 3 if rank_income > N_year*0.50 & rank_income <= N_year*0.75
```

```
by year: replace income_q = 4 if rank_income > N_year*0.75
```

```
label define income_q_lbl 1 "Lowest" 2 "Low-Middle" 3 "High-Middle" 4 "Highest"
```

```
label values income_q income_q_lbl
```

```
tab income_q
```

```
tab confinan
```

```

recode confinan (1=3) (2=2) (3=1), gen(conf_bank)
label define conf_lbl 1 "Hardly any" 2 "Only some" 3 "A great deal"
label values conf_bank conf_lbl

gen age_group = .
replace age_group = 1 if age < 30
replace age_group = 2 if age >= 30 & age < 50
replace age_group = 3 if age >= 50 & age < 65
replace age_group = 4 if age >= 65
label define age_lbl 1 "<30" 2 "30-49" 3 "50-64" 4 "65+"
label values age_group age_lbl

svyset [pweight=wtssall]

svy: tab conf_bank income_q, row
svy: tab conf_bank degree, row
svy: tab conf_bank age_group, row

svy: ologit conf_bank i.income_q i.degree c.age i.year, or

margins income_q, predict(outcome(1)) ///
    predict(outcome(2)) ///
    predict(outcome(3))

marginsplot, ///
    title("Predicted probability of confidence in banks by income group") ///
    xtitle("Income group") ///
    ytitle("Probability") ///
    legend(order(1 "Hardly any" 2 "Only some" 3 "A great deal"))

```

tab confinan
tab income_q
tab conf_bank

Shiny Application Exports:

<https://drive.google.com/drive/folders/17eV3IE5QzdaBamZ74jQ-HgegAbkpnpWz?usp=sharing>