

Pre-Analysis Plan: Generative AI Use and Well-Being

Louis Fréget^{1,2}, Ariell Reshef³, and Claudia Senik^{4, 3}

¹Centre pour la recherche économique et ses applications (CEPREMAP)

²University Paris-Dauphine

³Paris School of Economics

⁴Sorbonne University

January 3, 2026

Abstract

Against the backdrop of increasing use of generative AI, we study the relationship between AI use and well-being using a large scale survey on a representative sample of the French population. First, we conduct descriptive analyses documenting how baseline AI use relates to both subjective and objective loneliness. Second, we run a randomized controlled trial in which participants are assigned either to engage in daily personal conversations with a generative AI for 28 days or to continue their usual behavior. This allows us to uncover the causal effect of daily conversations with generative AI on subjective loneliness and well-being as well as on objective loneliness (real-life interactions).

1 Hypotheses

We consider two competing ways in which guided daily AI conversations may affect our primary outcomes of subjective well-being and loneliness, and two competing implications for objective loneliness and real-life social connectedness. These yield opposing predictions for the primary and secondary outcomes.

1.1 Mechanisms related to subjective loneliness and subjective well-being

Support / augmentation mechanism. Daily guided personal conversations with AI may provide structure for emotional reflection, facilitate the articulation of concerns, and offer responsive feedback. Under this mechanism, treatment increases subjective well-being. Life satisfaction increases, loneliness decreases (both UCLA-3 and single-item loneliness), and affect improves, as reflected in higher reported happiness and lower reported depression. Participants may also report a higher valuation of AI use (higher willingness-to-accept forgoing AI) and rate the conversations as pleasant on average.

Rumination mechanism. Alternatively, daily AI conversations about personal matters may heighten focus on problems and encourage rumination rather than resolution. Under this mechanism, subjective well-being deteriorates or does not improve. Life satisfaction may not increase or may decline; subjective loneliness may increase (UCLA-3 and single-item loneliness); and affect worsens, with lower happiness and higher reported depression. Willingness-to-accept forgoing AI may increase even if well-being does not, and conversations may be experienced as less pleasant for some participants.

1.2 Mechanisms related to objective loneliness and real-life social connectedness

Complement mechanism. If AI use supports emotional regulation and coping, this may strengthen engagement with real-world social relationships. Under this mechanism, social connectedness improves. Relationship satisfaction, family satisfaction, and friendship satisfaction increase; perceived social support increases (higher probability the respondent has at least one person to rely on); and objective loneliness decreases, as reflected in fewer lunches and dinners eaten alone and more meetings with friends or relatives. Participants may place greater value on AI without substituting human interactions.

Substitution mechanism. Conversely, AI may partially replace human social interaction if it becomes a convenient or psychologically comfortable alternative. In this case, reliance on AI could crowd out efforts to engage socially. Under this mechanism, social connectedness deteriorates. Relationship satisfaction, family satisfaction, and friendship satisfaction decline; perceived social

support decreases; and objective loneliness increases, with more meals eaten alone and fewer meetings with friends or relatives. One additional possibility is that when participants discuss conflicts, AI responses may implicitly align with participants’ own perspectives in a way that reinforces grievances rather than facilitates repair. This mechanism is more likely to operate when participants report a high perceived bias of the AI toward their own viewpoint in conflict situations.

Given these competing mechanisms, the direction of effects on our primary and secondary outcomes is theoretically ambiguous.

2 Study Design and Procedures

2.1 Setting, Population, and Timeline

Adults residing in France are recruited from an online panel by a French company (Bilendi). The questionnaire will be administered in French. Wave 1 (baseline) collects demographics and all outcomes; randomization follows immediately. Participants then complete Wave 2 (endline) about 28 days later with identical outcome measures.

We conducted a pilot study with 1000 participants to test survey flow, outcome measurement, and treatment compliance. The pilot confirmed the feasibility of the design, the clarity of the instructions, and indicated no differential attrition by treatment status.

2.2 Intervention

Treatment participants are instructed to engage in at least one brief daily conversation for 28 consecutive days with a generative AI tool of their choice (e.g., ChatGPT, LeChat, Mistral AI...) on personal topics (experiences, feelings, goals, concerns). Control participants continue their usual behavior. The research team does not access the content of the conversations; participants privately record their days of compliance.

2.3 Randomization

Randomization unit: individual. Assignment ratio: 1:1 (treatment vs. control). Implementation: survey platform (computerized). No clustering.

2.4 Sample Size and Power

$N = 10,000$ at baseline. Assumed attrition 20%. Planned total $N = 8,000$ (approximately $n = 4,000$ per arm). Two-sided $\alpha = 0.05$. Under these assumptions, the minimum detectable effect (MDE) at 80% power is ≈ 0.06 SD. Accordingly, the power to detect an effect of 0.10 SD (a common benchmark in survey experiments) exceeds 99%. Estimates are conservative, as we do not account for baseline adjustments in power.

2.5 Sample Restrictions

We exclude respondents who fail to meet the baseline attention check.

3 Outcomes and Measurement

3.1 First Stage outcomes

1. **First stage (compliance).** Number of days conversing with an AI about personal matters over the 28-day intervention, self-reported at endline (this our key first stage measure).
2. **AI-use intensity.** Frequency of generative AI use across all domains in the survey (e.g., work, emotional support, companionship, health information...).

3.2 Primary Outcomes

1. **Life satisfaction.** Single-item measure on a 0–10 scale. For analysis, standardized (mean 0, SD 1).
2. **Subjective loneliness.** (i) the UCLA-3 scale, consisting of three items (lack of companionship, feeling left out, feeling isolated from others), with responses coded 1 = hardly ever/never, 2 = some of the time, 3 = often ; (ii) a single-item self-reported loneliness question on a 0–10 scale.
3. **Self-reported affect.** Measured with two items: happiness yesterday (0–10) and depression yesterday (0–10).

3.3 Secondary Outcomes

- **Relationship satisfaction.** Satisfaction with partner (0–10; not applicable if there is no partner)¹, satisfaction with close others (family and close friends) (0–10), and general friendship satisfaction (0–10).
- **Objective loneliness.** (i) Number of lunches alone in the last 7 days (0–7), (ii) number of dinners alone (0–7), (iii) number of meetings with friends/relatives outside the household (0–7).
- **Perceived social support.** whether the respondent has someone on whom they can rely.
- **AI-use valuation.** Willingness-to-accept (WTA) to forgo generative AI for one week (numeric response in euros).
- **Pleasant or unpleasant personal conversations with AI.** 0 – 10 scale.

3.4 Index Construction and Standardization

Outcomes not originally measured on a 0-10 scale or on a 0-7 scale will be standardized (z-scores) for regression analysis. The UCLA-3 loneliness measure is constructed as a summed index (range 3-9) and subsequently standardized.

4 Empirical Strategy

4.1 Descriptive Analysis of the relationship between AI Use, Loneliness, and Life Satisfaction

Before turning to causal analysis, we will document descriptive associations between generative AI use, and our primary and secondary outcomes.

4.1.1 Measuring AI use in the descriptive analyses

We construct an index of socio-emotional AI use by summing three items (*AI use frequency* — *emotional support*, *friendly conversation/companionship*, *relationship advice*), all measured on a

¹Partner satisfaction is measured only among respondents who report having a partner at the relevant survey wave; analyses of this outcome use the corresponding subsample.

Likert-type frequency scale, given the excellent internal consistency observed in our pilot (Cronbach’s $\alpha = 0.93$). In complementary analyses, we will also consider a broader AI-usage index constructed by summing all AI-use frequency items across domains, justified by similarly strong internal consistency in the pilot.

4.1.2 Univariate and bivariate statistics

We will produce:

1. Summary statistics (means, standard deviations, and distributions) of all variables in the survey.
2. Binned scatter plots: Mean loneliness and life satisfaction by AI-use quantile, with 95% confidence intervals

4.1.3 Baseline Regressions (Correlational)

We will estimate the following descriptive regressions:

$$y_{i1} = \alpha + \delta \text{AIUse}_{i1} + \mathbf{X}'_{i1} \boldsymbol{\theta} + \varepsilon_i, \quad (1)$$

where y_{i1} is the primary or secondary outcome measured at baseline; the covariate vector $\mathbf{X}_{i1}$ includes age, gender, education, relationship status, number of friends at age fifteen, and household size. AIUse_{i1} is the (socio-emotional) AI usage at baseline defined in 4.1. We will enter it both as a continuous variable and as a discrete variable (quantile of AI-usage) in the regressions. Standard errors are robust.

We will also explore alternative specifications using broader metrics of AI use, for instance by summing all questions on AI usage frequency across domains to construct an index. In addition, we will estimate a specification in which socio-emotional AI use is the key predictor while controlling for work-related AI use; however, we treat this as a secondary robustness exercise, since work-related AI use may constitute a bad control.

4.1.4 Heterogeneity

We will document how baseline correlations between AI use, loneliness, and life satisfaction vary across:

- Age groups (18–24, 25–34, 35–44, 45–54, 55+)
- Gender
- Education levels
- Remote work frequency quantiles
- Relationship status.

4.2 Analyzing Experimental Results: Primary Estimand and Main Specification

The primary estimand is the ITT effect of assignment to treatment on endline outcomes. Our main specification is ANCOVA:

$$Y_{i2} = \alpha + \beta T_i + \gamma Y_{i1} + \mathbf{X}_{i1}' \boldsymbol{\theta} + \varepsilon_{i2}, \quad (2)$$

where Y_{i2} is the outcome at four weeks, T_i is the treatment assignment, Y_{i1} is the baseline measure of the same outcome included when available, and $\mathbf{X}_{i1}$ includes pre-specified covariates: age, gender, education (highest diploma), socio-economic status, telework frequency at baseline, relationship status, and household composition (number of members, presence of children). Standard errors are robust.

4.3 Alternative Specification

As a robustness check, we estimate difference-in-means regressions:

$$Y_{i2} = \alpha + \beta T_i + \varepsilon_i \quad (3)$$

without baseline adjustment, on Y_{i2} .

4.4 Multiple Hypothesis Testing

Our main correction is the Romano–Wolf stepdown procedure, implemented with resampling-based adjusted p-values. We will report both raw and Romano–Wolf adjusted p-values. We pre-specify three outcome families and apply Romano–Wolf stepdown adjustments within each: (i) primary subjective well-being outcomes (life satisfaction, subjective loneliness, and affect), (ii) relationship quality and perceived social connectedness outcomes (partner, family, and friendship satisfaction, as well as perceived social support), and (iii) objective social interaction outcomes (number of meals eaten alone and frequency of in-person meetings with friends or relatives).

4.5 Heterogeneity

Pre-specified moderators include baseline high/low subjective loneliness (UCLA-3 below 5, the median at baseline in our pilot study), remote working frequency, gender, age groups (18–24, 25–34, 35–44, 45–54, 55+), and relationship status (single vs couple). We estimate treatment heterogeneity by interacting the treatment dummy with each moderator in our main ANCOVA specification (controlling for the baseline outcome value).

4.6 Missing Data and Attrition

Our main analyses will use complete cases. We will test whether survey completion (endline response) is correlated with treatment assignment; in the pilot, no such correlation was detected.

We will also implement, as robustness checks:

1. Mean imputation with missing-value indicators,
2. Lee-bound trimming to account for differential attrition.

5 Data, Ethics, and Registration

Data collection. Two online surveys (baseline and ~28 days later). No conversation content is collected; participants privately record their compliance.

Privacy and ethics. Participants consent online; the study does not access or store AI conversation content. We will adhere to all IRB-approved procedures for privacy, data security, and participant protection.

Trial registration. The study will be registered on the AEA RCT Registry. This PAP will be uploaded alongside the registration, and any deviations from this plan will be transparently documented in a post-analysis note.

6 Deliverables and Reporting

We will publish a replication package (code and de-identified data);

Administrative

JEL: C93, I31, I12, O33

Keywords: Artificial Intelligence; Subjective Well-Being; Loneliness; Social Connectedness; Work From Home; WTA; Survey Experiment.