

How Well Do LLMs Predict Human Behavior? A Measure of their Pretrained Knowledge*

Wayne Gao, Sukjin Han and Annie Liang

Discussion Paper 26/835
18 January 2026

School of Economics

University of Bristol
Priory Road Complex
Bristol
BS8 1TU
United Kingdom



How Well Do LLMs Predict Human Behavior?

A Measure of their Pretrained Knowledge*

Wayne Gao[†] Sukjin Han[‡] Annie Liang[§]

January 18, 2026

Abstract

Large language models (LLMs) are increasingly used to predict human behavior. We propose a measure for evaluating how much knowledge a pretrained LLM brings to such a prediction: its *equivalent sample size*, defined as the amount of task-specific data needed to match the predictive accuracy of the LLM. We estimate this measure by comparing the prediction error of a fixed LLM in a given domain to that of flexible machine learning models trained on increasing samples of domain-specific data. We further provide a statistical inference procedure by developing a new asymptotic theory for cross-validated prediction error. Finally, we apply this method to the Panel Study of Income Dynamics. We find that LLMs encode considerable predictive information for some economic variables but much less for others, suggesting that their value as substitutes for domain-specific data differs markedly across settings.

*For helpful discussions, we thank Tom Cunningham, Matt Gentzkow, Guido Imbens, and Pat Kline.

[†]Department of Economics, University of Pennsylvania

[‡]School of Economics, University of Bristol

[§]Department of Economics, Northwestern University

1 Introduction

Large language models (LLMs) are increasingly used in economics as predictive tools—both to generate synthetic responses in place of human subjects (Horton, 2023; Anthi et al., 2025), and to forecast economic outcomes directly (Hewitt et al., 2024a; Faria-e Castro and Leibovici, 2024; Chan-Lau et al., 2025). Their appeal in these roles is obvious: A pretrained LLM embeds a vast amount of information and can be deployed at negligible cost, often in settings where collecting new, domain-specific human data would be expensive or infeasible.

What remains unclear is how to assess the *quality* of these predictions. This paper proposes a measure that quantifies the domain-specific value of LLMs in an interpretable unit: the amount of human data they substitute for. Specifically, we ask how much human data would be required for a conventional model trained on that data to match the predictive performance of the pretrained LLM in that domain. We define the LLM’s *equivalent sample size* to be the smallest size of training data required for the model to match the LLM.

For some prediction tasks, LLMs may perform comparably to models trained on large datasets, making them a useful surrogate for data collection. For others, LLMs may be so poorly suited that even models trained on a small amount of domain-specific data quickly outperform them. Quantifying the LLM’s equivalent sample size can thus inform the need for data collection and the interpretation of LLM-based empirical results across domains.

Beyond introducing this measure, the paper develops a method for estimating the equivalent sample size and for conducting statistical inference. In this part of the paper, we construct a confidence interval for our proposed measure via sequential hypothesis testing, whose validity is shown by developing a novel asymptotic theory for cross-validated estimators. Finally, we illustrate the measure by evaluating LLM predictions of variables in the Panel Study of Income Dynamics, where we find substantial heterogeneity in the extent to which LLMs can proxy for observed data across prediction tasks.

The paper proceeds as follows. Section 2 formalizes the setting and introduces our proposed measure. We consider general prediction problems in which an analyst seeks to predict

an outcome Y given covariates X . A *prediction rule* is a mapping from covariates to outcomes, and the error of any such rule is evaluated by its expected loss under a specified loss function. The standard approach for learning a prediction rule is to collect a dataset of (X, Y) observations and train a model or machine learning algorithm on that data.

Large language models (LLMs) offer an alternative. Rather than training on domain-specific data, the analyst can describe the prediction problem to a pretrained LLM and request predictions of Y given X . While the LLM is itself trained on data, we focus on settings in which the analyst deploys a fixed, pretrained model without further domain-specific fine-tuning, as is common in the literature (see Section 1.1). We refer to the resulting mapping from covariates to predicted outcomes as the *LLM prediction rule*.

We then compare the error of the LLM prediction rule to that of a fixed comparator trained on datasets of increasing size. This comparator may be an off-the-shelf machine-learning method or a structural economic model; for brevity, we refer to it simply as an algorithm. Training the algorithm on datasets of different sizes yields a sequence of prediction rules that trace out an error curve. We define the *equivalent sample size* to be the smallest training sample size at which the algorithm’s error weakly improves upon that of the LLM prediction rule. This quantity can be interpreted as the amount of data required to train a model that is as informative as the pretrained LLM for the prediction task at hand.

Section 3 develops an estimator and inference procedure for the equivalent sample size. We estimate the LLM’s prediction error using its performance on the full sample, and estimate the error curve of the comparator algorithm via block-out cross-validation. We finally introduce a sequential testing procedure that compares the LLM and algorithm across an ordered sequence of training sizes and yields a one-sided confidence interval for the equivalent sample size.

Section 4 provides theoretical foundations for this inference approach and establishes its validity under the mild assumption that the algorithm’s error is weakly reducing in the size of its training data. Because the sequence of hypotheses comparing the algorithm and the LLM is nested, the sequential procedure controls coverage without requiring multiple-testing cor-

rections, thus avoiding the conservativeness that typically arises in multiple hypothesis testing. This feature is especially important in our setting, where we want to compare algorithm performance across a large range of training sample sizes. We finally derive a central limit theorem for the block-out cross-validated risk estimator under a fixed-training-size asymptotic regime, and construct a consistent estimator of its asymptotic variance. Together, these results establish that the proposed procedure delivers a valid one-sided confidence interval for the equivalent sample size with correct asymptotic coverage.

Section 5 applies our proposed measure to evaluate LLM prediction of variables in the Panel Study of Income Dynamics (PSID). The PSID is one of the longest-running longitudinal household surveys in the U.S., and is frequently used in empirical research. Using the 2021 wave of the Panel Study of Income Dynamics, we study a suite of prediction tasks in which respondents’ demographic, labor-market, and household covariates are used to forecast hourly wages, homeownership, drinking, and smoking. We construct LLM prediction rules by translating each covariate vector into a natural-language persona description and querying the model. To limit concerns about data leakage, we emphasize post-cutoff evaluation by using static open-source models with documented training cutoffs. The estimated equivalent sample size varies substantially across outcomes—e.g., approximately 20 observations for predicting hourly wages to 600 observations for predicting homeownership—highlighting substantial heterogeneity in the extent to which pretrained knowledge substitutes for domain-specific data.

Finally, Section 6 extends our evaluation framework to assessing LLM performance when prompted to infer counterfactual outcomes or treatment effects. Under an assumption of unconfoundedness, we define equivalent sample sizes for both the average potential outcome conditional on covariates and the conditional average treatment effect (CATE). Although the true CATE is unknown, a prediction framework with a transformed outcome enables us to calculate the risk for CATE for both the LLM and ML and thus the equivalent sample size.

1.1 Related Literature

LLMs as Human Surrogates. LLMs are increasingly used as *surrogate* experimental and survey subjects in place of human participants. Researchers prompt LLMs with experimental or survey scenarios and treat the resulting responses as proxies for human behavior. This approach has been applied in survey settings (Argyle et al., 2023; Jansen et al., 2023; Sun et al., 2024) and in experimental settings (Aher et al., 2023; Hewitt et al., 2024b; Park et al., 2024; Bohren et al., 2025). Related work explores broader applications of LLMs as simulated economic or social agents, including their use in market simulations and policy analysis (Horton, 2023; Manning et al., 2024).

Despite this growing adoption, there is little consensus on how to assess the quality or value of LLMs as human surrogates. Much of the existing literature evaluates performance in task-specific or descriptive terms—asking whether LLM outputs resemble human responses in a given domain—without a unifying metric that quantifies what is gained or lost by substituting LLMs for human data. At the same time, several papers caution against uncritical reliance on LLMs, documenting systematic biases, domain failures, and sensitivity to prompting choices (Gao et al., 2025; Ludwig et al., 2025). This paper contributes a general measure for evaluating the appropriateness of LLM predictions: the amount of human data the LLM effectively substitutes for.

Model Evaluation. Our paper is closely related to Huang et al. (2025), who also interpret the value of LLM outputs in terms of an equivalent number of human observations. In their setting, the effective sample size arises from an uncertainty-quantification problem: how many synthetic LLM responses are required to construct confidence intervals with valid coverage for human population parameters. In contrast, we treat the LLM as a fixed prediction rule and define the equivalent sample size as the amount of real, domain-specific data required for a supervised learning algorithm to match the LLM’s predictive performance.

Additionally, our approach is in the spirit of Erev et al. (2007)’s notion of the *equivalent number of observations* (*ENO*), which evaluates game-theoretic models by asking how many

subject observations would be required for the sample average of observed behavior to predict as well as the theory. The objects of comparison are however different: In their setting, the object being evaluated—a behavioral or equilibrium model—is itself estimated from data, whereas we treat the LLM as a fixed prediction rule that is not trained or tuned on the domain sample. Additionally, their comparator is the sample average of observed behavior, while ours is the prediction of a supervised learning algorithm. Thus our measure applies to a substantially broader set of prediction tasks.

Finally, our work is related more broadly to work on evaluating economic models (Fudenberg et al., 2022, 2026; Andrews et al., 2025), where here the “model” is the pretrained LLM. In particular, our benchmarking of LLM performance against that of a comparator algorithm is similar to Fudenberg et al. (2022), although our focus on the equivalent *sample size* of the comparator algorithm is entirely new.

Asymptotic Theory with Cross-Validation. A growing literature develops asymptotic theory for cross-validated prediction error to address the dependence induced by overlapping training samples across folds (Austern and Zhou, 2020; Bates et al., 2024; Fava, 2025; Lei, 2025). Our results in Section 3.4 for block-out cross validation contribute to this literature, and are in particular related to Lei (2025) and Fava (2025). The key difference in our analysis is that we consider an asymptotic regime where the training sample size is held fixed. Lei (2025) and Fava (2025) instead work with a standard asymptotic regime where the training sample size grows large, which is natural for typical machine learning applications. Inference for our proposed measure demands a different asymptotic regime, since we implement hypothesis testing with respect to a given number of training observations.

2 Framework

Section 2.1 introduces our framework: we study prediction problems in which an analyst observes a vector of covariates X and seeks to predict an outcome Y . Section 2.2 defines

a prediction rule obtained by providing the LLM with a text description of X and eliciting a prediction of Y . This approach contrasts with the standard machine learning paradigm, reviewed in Section 2.3, in which the analyst first collects a dataset of (x, y) observations and then trains a task-specific prediction algorithm on that data. Finally, Section 2.4 introduces our proposed measure: the smallest training sample size for which a domain-specific model or algorithm achieves predictive performance comparable to that of the LLM.

2.1 Setup

A *prediction problem* is a pair (X, Y) where $X \in \mathcal{X}$ is a covariate vector and $Y \in \mathcal{Y}$ is an outcome of interest.

Example 1. *The covariate vector X describes a worker’s educational attainment, years of experience, occupation, industry, geographic location, and past earnings. The outcome $Y \in \mathbb{R}_+$ is the worker’s wage in the following year.*

Example 2. *The covariate vector X describes an individual’s age, educational attainment, marital status, and race. The outcome $Y \in [0, 1]$ is the probability with which the individual owns a house.*

A *prediction rule* is any map $f : \mathcal{X} \rightarrow \mathcal{Y}$ that predicts the outcome Y given X , where the LLM-based prediction rule is one special choice of f (see Section 2.2). As is standard, we measure the accuracy of a prediction $\hat{y}$ given true outcome y using a loss function $\ell : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$, where $\ell(y, y')$ is the inaccuracy of predicting y' when the true outcome is y (e.g., $\ell(y, y') = (y' - y)^2$ for regression or $\ell(y, y') = 1\{y' \neq y\}$ for classification). Define the *error* of prediction rule f to be

$$e(f) \equiv E_P[\ell(f(X), Y)]$$

i.e., the expected loss when predictions are made using f , where P denotes the true (unknown) joint distribution of (X, Y) . This is also known as the *risk* of f .

2.2 LLM-Based Prediction

Our procedure and its guarantees generalize to any fixed prediction rule f , but we are in particular interested in the prediction rule f_{LLM} that corresponds to asking a pretrained LLM to predict the outcome. The covariate vector X and desired outcome Y are described in natural language and given to the LLM, and the LLM’s response is taken as the predicted outcome.

Here is an example prompt:

Prompt 1. *Here is your persona: You are a 37-year-old male living in Tennessee in 2020. You work in architecture and engineering occupations in manufacturing industry. You have 4 years of work experience in the current job. You have completed 14 years of education. You are married and have 1 child. Your father has 12 years of education. Your mother has 12 years of education. Your race is white. Your health is excellent. You are a Protestant.*

Given your persona, do you own a house or apartment or do you rent? If you own, say 1; if you rent, say 5; if you neither own nor rent, say 8.

The prompt text is fixed up to placeholders for individual covariate values, so the structure of the prompt is constant and varies only in the realized covariate values. Let

$$e_{\text{LLM}} \equiv e(f_{\text{LLM}})$$

denote the error of this prediction rule.

Throughout we treat the LLM-based rule f_{LLM} as fixed, and do not model the background data on which it is trained. In practice, researchers often do not retrain foundation models, but instead deploy a given pretrained system. Our object of interest is therefore the predictive performance of this fixed system.

2.3 Domain-Specific Learning

We compare the above approach to prediction rules that are obtained by training a model or algorithm on domain-specific data. This approach first specifies a set of prediction rules $\mathcal{F}$, and then uses data to select among them. Formally, an *algorithm* a is a map

$$a : \mathcal{D} \rightarrow \mathcal{F}$$

from the set of all finite datasets $\mathcal{D} \equiv \cup_{N \geq 1} (\mathcal{X} \times \mathcal{Y})^N$ to the set of prediction rules $\mathcal{F}$.¹

A typical choice of a would be an economic model or a machine-learning procedure. In the first case, $\mathcal{F}$ represents an interpretable class of economically motivated prediction rules, and $a(D)$ denotes the rule obtained by estimating the model's parameters on the dataset D using a specified estimation method. In the second case, $\mathcal{F}$ is a potentially rich class of predictors—such as ensembles of decision trees—and $a(D)$ is the fitted predictor produced by training the algorithm on D , including any regularization or tuning steps. In both cases, the algorithm maps data into a concrete prediction rule whose out-of-sample performance typically improves with the size of the training dataset, providing a natural benchmark against which to evaluate the predictive value of a fixed LLM.

For any sample size $N \in \mathbb{N}_+$, let

$$f_N \equiv a(D_N)$$

be the prediction rule obtained by training the algorithm on a random sample $D_N \equiv \{(x_i, y_i)\}_{i=1}^N$ of N observations (x_i, y_i) drawn i.i.d. from P . Since D_N is random, the prediction rule f_N is as well. We define

$$e_N^a \equiv E[e(f_N)] = E_{D_N} [E_{(x,y)}[\ell(a(D_N), (x, y))]]$$

¹All of our procedures and results extend if we instead model the algorithm as a map $a : \mathcal{D} \times [0, 1] \rightarrow \mathcal{F}$, where the additional argument in $[0, 1]$ represents an exogenous source of randomness, allowing the algorithm's output to be stochastic—for example, due to random initialization, subsampling, or internal randomization.

to be the expected prediction error of the algorithm trained on N datapoints. That is, e_N^a integrates both over the randomness of the sample D_N used to train f_N and the randomness of the new observation $(X, Y) \sim P$ on which f_N is evaluated.

2.4 Proposed Measure

Our proposed measure of the LLM’s *equivalent sample size* is the minimum amount of domain-specific data an experimenter would need to collect so that the algorithm a , when trained on that data, produces predictions at least as accurate as the LLM’s.

Definition 2.1. *Fix a prediction problem (X, Y) . The a -equivalent sample size (a -ESS) of f_{LLM} is*

$$N^* \equiv N_a^* = \min \{N \in \mathbb{N}_+ : e_N^a \leq e_{\text{LLM}}\}.$$

i.e., the smallest training sample size N^ such that the predictive performance of f_{N^*} weakly improves upon that of the LLM. If no such N exists, we set $N^* = \infty$.*

The equivalent sample size depends on the benchmark algorithm a , as well as on the choice of LLM, its prompting protocol, and the prediction problem. Thus it should be interpreted as measuring the LLM’s predictive ability in a specific prediction problem and relative to the specified algorithm, rather than as an absolute measure of the LLM’s predictive content. We encourage the use of algorithms a that are well-suited to the prediction problem and thus serve as more demanding benchmarks.

At one extreme, $N^* = 1$ means that the LLM provides no predictive advantage relative to the algorithm a in this prediction problem, even when the algorithm is trained on a single data point. This case corresponds (for example) to an LLM that is very poorly suited to the prediction task, or which relies on spurious correlations in its original training data.

At the other extreme, the equivalent sample size can also be infinite, meaning that no amount of domain-specific data allows the benchmark algorithm to match the LLM’s predictive performance. This could occur if a is too restrictive to approximate the LLM’s mapping

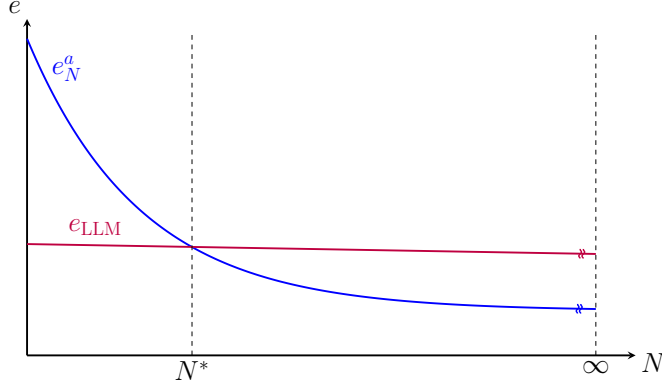


Figure 1: a -Equivalent Sample Size (N^*)

even asymptotically—for example, if a is an economic model that imposes strong structural restrictions.

In general, we expect the LLM’s equivalent sample size to be finite. That is, while the LLM’s pretraining enables it to make useful predictions even in the absence of any new data, a sufficiently flexible algorithm trained on data from the relevant prediction task eventually outperforms it.

In our application we will report figures such as Figure 1, which plots the constant error e_{LLM} against e_N^a as N varies. The intersection point defines the a -equivalent sample size N^* .

3 Estimation and Inference: Procedure

This section describes our procedure for estimating and conducting inference for the a -equivalent sample size N^* . Recall that N^* is defined as the smallest training sample size N given which the machine learning algorithm weakly outperforms the LLM-based prediction rule:

$$N^* \equiv \min \{N : e_N^a \leq e_{\text{LLM}}\}.$$

Section 3.1 describes our approach to estimating the error of the LLM-based prediction rule, e_{LLM} . Section 3.2 describes a block-out cross-validation procedure for estimating the error curve of the machine learning algorithm, $\{e_N^a\}_{N \geq 1}$. These can be combined to compute

a plugin estimator for N^* , as we describe in Section 3.3. Finally, Section 3.4 develops a sequential hypothesis-testing procedure that delivers a one-sided confidence interval (CI) for N^* . To aid readability, this section focuses on the procedure itself, while Section 4 provides theoretical guarantees for these procedures.

3.1 Estimating the LLM Benchmark

We estimate the error of the LLM-based prediction rule, e_{LLM} , as its performance on the full data sample $\{(X_i, Y_i)\}_{i=1}^n$. For each observation X_i , obtain a prediction $\hat{Y}_i^{\text{LLM}} = f_{\text{LLM}}(X_i)$ by querying the LLM using the fixed prompt template described in Section 2.2. The empirical error of the LLM is then given by

$$\hat{e}_{\text{LLM}} \equiv \frac{1}{n} \sum_{i=1}^n \ell(\hat{Y}_i^{\text{LLM}}, Y_i). \quad (3.1)$$

3.2 Estimating the Error Curve of the ML Algorithm

We next estimate the algorithm’s error curve by training it on disjoint training blocks of a given size and evaluating each fitted predictor on the held-out complement. Averaging the resulting out-of-sample losses produces an estimate of the algorithm’s expected error as a function of the size of the training sample. Throughout this section, we drop dependence of e_N^a on the algorithm a to ease notation.

Given a dataset of size n , fix a candidate training size N_k for $k \in \{1, \dots, K\}$. Set $B_k \equiv \lfloor n/N_k \rfloor$ to be the number of blocks and $M_n \equiv n - N_k$ to be the size of the test set. For simplicity, assume $n = B_k N_k$; otherwise, discard the last $n - B_k N_k = O(1)$ observations.²

Partition the index set $\{1, \dots, n\}$ into B_k disjoint blocks $S_{k,1}, \dots, S_{k,B_k}$, each of size N_k . For concreteness, take $S_{k,b} = \{(b-1)N_k + 1, \dots, bN_k\}$ for $b = 1, \dots, B_k$. For each block

²Discarding these $B_k N_k$ observations does not affect the later asymptotic theory.

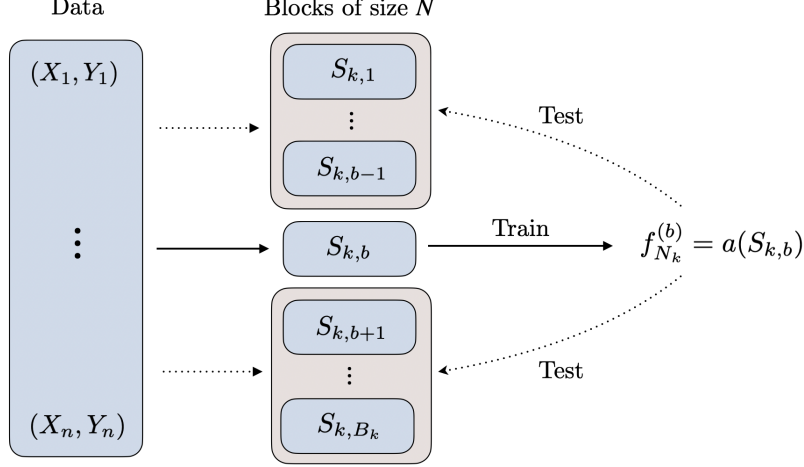


Figure 2: Depiction of block-out cross-validation: The data is split into B_k blocks each of size N_k . In each round, one block is used for training and the remaining are used for testing.

$b \in \{1, \dots, B_k\}$, define the *block b training set* to be

$$\mathcal{D}_{\text{train}}^{(k,b)} \equiv \{(X_i, Y_i) : i \in S_{k,b}\}$$

and the *block b test set* to be

$$\mathcal{D}_{\text{test}}^{(k,b)} \equiv \{(X_i, Y_i) : i \notin S_{k,b}\}$$

as depicted in Figure 2. Because the observations (X_i, Y_i) are i.i.d. and the sets $S_{k,b}$ form a partition of $\{1, \dots, n\}$, the training blocks $\mathcal{D}_{\text{train}}^{(k,1)}, \dots, \mathcal{D}_{\text{train}}^{(k,B_k)}$ are independent and each has distribution P^{N_k} .

Recall that each training set is of size N_k while each test set is of size $M_n \equiv n - N_k$. In each round b , we fit the learning algorithm on the block b training set to obtain the block-wise predictor

$$f_{N_k}^{(b)} \equiv a(\mathcal{D}_{\text{train}}^{(k,b)}),$$

and evaluate its empirical error on the test complement:

$$\hat{e}_{k,b} \equiv \frac{1}{M_n} \sum_{i \notin S_{k,b}} \ell(f_{N_k}^{(b)}(X_i), Y_i).$$

The block-out CV estimator of e_{N_k} is finally

$$\hat{e}_{N_k}^{CV} \equiv \frac{1}{B_k} \sum_{b=1}^{B_k} \hat{e}_{k,b}. \quad (3.2)$$

i.e., the average test error across the choice of training block.

3.3 Plugin Estimator for the ESS

Given the estimates (3.1) and (3.2), the plug-in estimator of the α -equivalent sample size is defined as

$$\hat{N}^* \equiv \min \{N_k : \hat{e}_{N_k}^{CV} \leq \hat{e}_{LLM}, k \in \{1, \dots, K\}\}$$

i.e., the smallest sample size at which the estimated LLM error $\hat{e}_{LLM}$ exceeds the estimated error $\hat{e}_N^{CV}$.

The next section defines a sequential estimator $\hat{N}_\alpha^*$ that incorporates sampling uncertainty and delivers a lower confidence bound for N^* with confidence level $1 - \alpha$. Under standard regularity conditions, both estimators are asymptotically equivalent, but $\hat{N}_\alpha^*$ is conservatively biased in finite samples to ensure correct coverage.

3.4 One-Sided Confidence Interval for the ESS

To develop a one-sided CI for N^* , we compare the predictive performance of the machine learning algorithm and the LLM across an increasing sequence of training sample sizes, and identify the smallest training size at which we can no longer conclude that the LLM outperforms the machine learning algorithm at the desired confidence level.

Let $\{N_k\}_{k=1}^K$ be an increasing sequence of candidate training sizes, with $N_1 < N_2 < \dots <$

N_K . For each $k \in \{1, \dots, K\}$, consider the one-sided hypotheses

$$H_{0,k} : e_{N_k} \leq e_{\text{LLM}}, \quad H_{1,k} : e_{N_k} > e_{\text{LLM}}.$$

Each hypothesis $H_{0,k}$ compares the predictive performance of the machine learning algorithm trained on N_k observations to that of the LLM, with the null asserting that the machine learning algorithm weakly outperforms the LLM at that training size. To test $H_{0,k}$, we use the test statistic

$$T_{N_k,n} \equiv \frac{\hat{e}_{N_k}^{\text{CV}} - \hat{e}_{\text{LLM}}}{\hat{\sigma}_k / \sqrt{n}}, \quad (3.3)$$

where $\hat{\sigma}_k^2$ is an estimator of the variance of $\hat{e}_{N_k}^{\text{CV}} - \hat{e}_{\text{LLM}}$, as detailed in (4.6) or (4.7) in Section 4.3 below. We reject $H_{0,k}$ at level α when $T_{N_k,n} > c_\alpha$ where c_α is the critical value.

We will be interested in the error curve $\{e_{N_k}\}_{k \geq 1}$. Theoretically it is possible that e_N is not monotone in N , so that increasing the training size *reduces* expected performance. This behavior is somewhat counterintuitive, and we expect that it is not descriptive of most economic models and machine learning algorithms that are considered in practice. Our subsequent results will impose the following assumption.

Assumption 3.1 (Monotonicity). e_N is nonincreasing in N .

Given a monotone error curve, we can conduct inference on N^* using the following sequential testing procedure at overall level α :

Algorithm 3.1. *Proceed as follows:*

Step 1. Test $H_{0,1}$ at level α . If $H_{0,1}$ is rejected, proceed to Step 2. Otherwise, stop and set $\hat{N}_\alpha^* = 1$.

$\vdots$

Step k . Test $H_{0,k}$ at level α . If $H_{0,k}$ is rejected, proceed to Step $k+1$. Otherwise, stop and set $\hat{N}_\alpha^* = N_{k-1} + 1$.

$\vdots$

Step K . Test $H_{0,K}$ at level α . If $H_{0,K}$ is rejected, conclude $N^* > N_K$ (and set $\hat{N}_\alpha^* = N_K + 1$). Otherwise, set $\hat{N}_\alpha^* = N_{K-1} + 1$.

The procedure sequentially increases the training sample size of the machine-learning algorithm and tests, at each step, whether its predictive performance weakly exceeds that of the LLM. As long as the null hypothesis is rejected, the procedure continues to larger training sizes. At the first step k at which the corresponding null $H_{0,k}$ cannot be rejected, the procedure stops and sets the estimated equivalent sample size to $N_{k-1} + 1$, the smallest sample size exceeding the largest training size for which rejection occurred.³ If all null hypotheses up to $H_{0,K}$ are rejected, the procedure concludes that the equivalent sample size exceeds N_K .

This stopping rule yields a valid lower confidence bound for the equivalent sample size N^* . In particular, suppose Assumption 3.1 holds and the level α test in each step of Algorithm 3.1 is valid. Then we show below that

$$\text{CI}_{N^*} \equiv \left[\hat{N}_\alpha^*, \infty \right) \quad (3.4)$$

is a one-sided $(1 - \alpha)$ -confidence interval for N^* . That is, $\Pr(\hat{N}_\alpha^* \leq N^*) \geq 1 - \alpha$, i.e., with probability at least $1 - \alpha$, one can conclude that the true equivalent sample size is no smaller than $\hat{N}_\alpha^*$. In practice, this allows one to say “ N^* is at least as large as $\hat{N}_{.05}^*$ with 95% confidence,” so that we have a minimal guarantee on the value of the LLM prediction rule. The validity of the tests in Algorithm 3.1 and the CI (3.4) are established in Section 4.

Remark 3.1 (Alternative CI Representation). *The CIs for N^* can be expressed in terms of CIs for the excess risk of the ML algorithm relative to the LLM (i.e., $e_{N_k}^a - e_{\text{LLM}}$). Recall that the algorithm stops at the first N_k where the null hypothesis $H_{0,k} : e_{N_k}^a \leq e_{\text{LLM}}$ is not rejected. By the duality between hypothesis tests and CIs, non-rejection at level α corresponds to the value 0 being contained in the one-sided CI for the difference $e_{N_k}^a - e_{\text{LLM}}$.*

³We set $\hat{N}_\alpha^*$ to $N_{k-1} + 1$ rather than N_k to account for the possible coarseness of the grid (i.e., $\{N_1, \dots, N_K\} \subsetneq \{1, \dots, K\}$). When one employ the finest grid (i.e., $N_k = K$), then simply $N_{k-1} + 1 = N_k$.

Formally, let LB_k denote the lower bound of the asymptotic one-sided $(1 - \alpha)$ CI for the excess risk, constructed from the test statistic (3.3).

$$LB_k \equiv (\hat{e}_{N_k}^{CV} - \hat{e}_{LLM}) - z_{1-\alpha} \frac{\hat{\sigma}_k}{\sqrt{n}}.$$

Assume the finest grid (i.e., $N_k = k$) for simplicity. The sequential estimator $\hat{N}_\alpha^*$ is simply the smallest sample size where this lower bound is non-positive (indicating that we cannot rule out that the ML model has matched the LLM):

$$\hat{N}_\alpha^* = \min \{N_k : LB_k \leq 0\} = \min \{N_k : \hat{e}_{N_k}^{CV} - z_{1-\alpha} \hat{\sigma}_k / \sqrt{n} \leq \hat{e}_{LLM}\}.^4 \quad (3.5)$$

This expression holds regardless of whether the LLM's error is treated as fixed or estimated, provided the variance estimator $\hat{\sigma}_k$ correctly accounts for the variability of both terms.

4 Estimation and Inference: Theory

This section presents theoretical results justifying our proposed procedure in Section 3. It is organized as follows: First, Section 4.1 proves the validity of the sequential testing procedure described in Section 3.4, taking as given the validity of the individual tests of the hypotheses $H_{0,k} : e_{N_k}^a \leq e_{LLM}$. Section 4.2 then establishes a central limit theorem for the block-out CV error estimator of each $e_{N_k}^a$, which yields pointwise confidence intervals for the error curve $\{e_{N_k}^a\}_{k \geq 1}$. Finally, Section 4.3 uses this central limit theorem to establish the validity of individual test of each $H_{0,k}$.

4.1 Validity of Sequential Hypothesis Testing

The validity of Algorithm 3.1 in terms of overall size control is established as follows.

Theorem 4.1 (Validity of Sequential ESS Inference). *Suppose that Assumption 3.1 holds.*

⁴For a generic coarser grid, $\hat{N}_\alpha^* = \max \{N_k : LB_k > 0\} + 1$.

Assume that for each k , the test of $H_{0,k}$ employed in Algorithm 3.1 has asymptotic size controlled by α , i.e.,

$$\limsup_{n \rightarrow \infty} \Pr(\text{reject } H_{0,k} \mid H_{0,k} \text{ is true}) \leq \alpha. \quad (4.1)$$

Then, the sequential testing procedure is valid at overall level α and the resulting CI for N^* has a correct coverage:

$$\liminf_{n \rightarrow \infty} \Pr(N^* \geq \hat{N}_\alpha^*) \geq 1 - \alpha.$$

The proof of Theorem 4.1 is short but reveals a structural feature of the sequential test that avoids the size distortions typical of multiple hypothesis testing. We therefore present it in the main text.

Proof. Define $k^* \equiv \min \{k : e_{N_k} \leq e_{\text{LLM}}\}$ to be the index of the equivalent sample size in our sequence $\{N_1, \dots, N_K\}$. Note that $N_{k^*-1} + 1 \leq N^* \leq N_{k^*}$ by the coarseness of the grid $\{N_k\}$.

Write $\underline{N}_k \equiv N_{k-1} + 1$ with $N_0 \equiv 0$. We will show that $[\hat{N}_\alpha^*, \infty)$ is a valid $(1 - \alpha)$ -CI for $\underline{N}_{k^*}$, i.e.,

$$\liminf_{n \rightarrow \infty} \Pr(\underline{N}_{k^*} \geq \hat{N}_\alpha^*) \geq 1 - \alpha. \quad (4.2)$$

Then, since $\underline{N}_{k^*} \leq N^*$ and consequently $\{\underline{N}_{k^*} \geq \hat{N}_\alpha^*\} \subseteq \{N^* \geq \hat{N}_\alpha^*\}$, we can conclude that $[\hat{N}_\alpha^*, \infty)$ is also valid for N^* based on

$$\liminf_{n \rightarrow \infty} \Pr(N^* \geq \hat{N}_\alpha^*) \geq \liminf_{n \rightarrow \infty} \Pr(\underline{N}_{k^*} \geq \hat{N}_\alpha^*) \geq 1 - \alpha. \quad (4.3)$$

We now prove (4.2). By Assumption 3.1, the sequence of null hypotheses $H_{0,k} : e_{N_k} \leq e_{\text{LLM}}$ has a nested truth structure:

(S1) For all $k < k^*$, we have $e_{N_k} > e_{\text{LLM}}$, so the null $H_{0,k}$ is *false*.

(S2) For all $k \geq k^*$, we have $e_{N_k} \leq e_{\text{LLM}}$, so the null $H_{0,k}$ is *true*.

Thus, H_{0,k^*-1} is the *last false null hypothesis* in the sequence, while H_{0,k^*} is the *first true null hypothesis* in the sequence.

Now, let $\hat{k}$ denote the random index at which Algorithm 3.1 terminates. The CI lower bound is defined as $\hat{N}_\alpha^* = \underline{N}_{\hat{k}}$. The algorithm stops at step k if it fails to reject $H_{0,k}$. Therefore, the event $\{\hat{k} = k\}$ implies that $H_{0,k}$ was not rejected, while all preceding hypotheses $H_{0,1}, \dots, H_{0,k-1}$ were rejected. Consequently, the event $\{\hat{N}_\alpha^* > \underline{N}_{k^*}\} = \{\underline{N}_{\hat{k}} > \underline{N}_{k^*}\}$ is equivalent to $\{\hat{k} > k^*\}$.

For the algorithm to go beyond step k^* (i.e., for $\hat{k} > k^*$), it must have rejected the hypothesis H_{0,k^*} . Conversely, if H_{0,k^*} is rejected, then $\hat{k} > k^*$. Therefore,

$$\Pr(\hat{N}_\alpha^* > \underline{N}_{k^*}) = \Pr(\text{reject } H_{0,k^*}).$$

Since H_{0,k^*} is true by (S2), the rejection of H_{0,k^*} constitutes a Type I error (a false rejection). By the condition of the theorem, the test for each individual step controls the asymptotic size at level α . Specifically, for the step k^* , we have $\limsup_{n \rightarrow \infty} \Pr(\text{reject } H_{0,k^*}) \leq \alpha$.

Combining these results, we obtain the bound on the coverage probability:

$$\begin{aligned} \liminf_{n \rightarrow \infty} \Pr(N^* \geq \hat{N}_\alpha^*) &= 1 - \limsup_{n \rightarrow \infty} \Pr(\hat{N}_\alpha^* > N^*) \\ &\geq 1 - \limsup_{n \rightarrow \infty} \Pr(\text{reject } H_{0,k^*}) \geq 1 - \alpha. \end{aligned}$$

which implies the desired conclusion in view of (4.3). $\square$

An important feature of this procedure is that it controls the family-wise error rate (FWER) (i.e., the probability of rejecting any true nulls) without requiring Bonferroni-type corrections (e.g., [Holm, 1979](#)). This is because the sequence of hypotheses are ordered (nested) and the sequential testing procedure stops at the first non-rejection. Thus we only test true null hypotheses if we have correctly rejected all preceding false nulls, so the probability of making *any* false rejection is bounded simply by the probability of rejecting the *first* true null H_{0,k^*} . This property is particularly valuable given that we will typically want to evaluate a wide range of sample sizes.

Remark 4.1 (Impact of the Coarse Grid $\{N_k\}_{k=1}^K$). Let $\hat{N}_{\alpha, \text{fine}}^*$ denote the cutoff $\hat{N}_\alpha^*$ obtained under the “finest grid” $N_k = k$. Then $\hat{N}_{\alpha, \text{fine}}^*$ coincides with the smallest integer k at which we reject the null in Algorithm 3.1, which would produce a “sharp” $(1 - \alpha)$ CI for N^* in the form of $[\hat{N}_{\alpha, \text{fine}}^*, \infty)$. For a generic coarse grid $\{N_k\}$, the cutoff $\hat{N}_\alpha^*$ is conservative against LLM in the sense of $\hat{N}_\alpha^* \leq \hat{N}_{\alpha, \text{fine}}^*$. Equivalently, we may view $[\hat{N}_{\alpha, \text{fine}}^*, \infty)$ as a sharp CI for $\underline{N}_{k^*} \leq N^*$, which results in conservative one-sided coverage for N^* .

The validity of our sequential test, as established in Theorem 4.1, requires a valid test statistic at each step. Recall in Section 3.4 that we use the studentized test statistic (3.3) for the boundary of the null hypothesis, $\tilde{H}_{0,k} : e_{N_k} = e_{\text{LLM}}$, i.e., the null where the machine learning algorithm and the LLM have equal expected error. Since this is the least favorable configuration within the null—i.e., the case where false rejection is most likely—a level- α test of $\tilde{H}_{0,k}$ is also a valid level- α test of the original null hypothesis $H_{0,k}$. An asymptotic one-sided test rejects $H_{0,k}$ at level α if $T_{N_k, n} > z_{1-\alpha}$, where $z_{1-\alpha}$ is the $(1 - \alpha)$ -quantile of the standard normal distribution.⁵ The validity of this individual test follows from Proposition 4.1 in Section 4.3, which is based on the Central Limit Theorem and consistent variance estimation in Section 4.2.

4.2 Central Limit Theorem under Block-Out Cross Validation

This section develops a Central Limit Theorem for risk estimators using the block-out CV procedure described in Section 3.2. Our established CLT in Theorem 4.2, and the corresponding consistent variance estimator in Theorem 4.3 are established under very mild conditions without the requirement of stability conditions (Lei, 2025), and may be of independent interest beyond the scope of this paper.

Specifically, we consider a setting where researchers have access to a single large data set with sample size n . For each candidate training size N_k , we estimate the out-of-sample error $e_{N_k}^a$ of any learning algorithm a via block-out CV. We work with an asymptotic framework

⁵We reject the null (that ML is as good as LLM) when the ML error is significantly *higher* than the LLM error. Thus, the rejection region is the right tail.

where $n \rightarrow \infty$ while N_k remains fixed.⁶ These CIs in turn produce CIs for N^* as discussed in Remark 3.1.

We work under the following sampling framework and regularity conditions. Let $Z \equiv (X, Y) \sim P$ denote the distribution of observed data.

Assumption 4.1 (Random Sampling). *A random sample of data $D_n \equiv (Z_1, \dots, Z_n)$ is observed, with $Z_i \equiv (X_i, Y_i) \sim_{i.i.d.} Z \equiv (X, Y) \sim P$.*

For each given $N \leq n$ and a sample D_N of size N , we define a generic learning algorithm as a measurable map $a : \mathcal{Z}^N \rightarrow \mathcal{F}$, where $\mathcal{F}$ is a class of measurable functions $f : \mathcal{X} \rightarrow \mathcal{Y}$. For any training sample $D_N \in \mathcal{Z}^N$, recall that we write $f_N = a(D_N)$.

We consider a generic loss function ℓ that satisfies the following conditions:

Assumption 4.2 (Loss and Moments). *The loss function $\ell : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$ is measurable and, for $D_N \sim P^N$ and an independent $Z = (X, Y) \sim P$,*

$$\mathbb{E}[|\ell(f_N(X), Y)|^{2+\delta}] < \infty \quad \text{for some } \delta > 0.$$

Recall that $e(f_N) \equiv \mathbb{E}_Z[\ell(f_N(X), Y) \mid D_N]$ and $e_N^a \equiv \mathbb{E}_{D_N}[e(f_N)]$. In addition, for each realization $z = (x, y) \in \mathcal{Z}$, define

$$e_N^a(z) \equiv \mathbb{E}_{D_N}[\ell(f_{D_N}(x), y)],$$

where the expectation is over an independent copy $D_N \sim P^N$ with z fixed. Note that $e_N^a = \mathbb{E}_Z[e_N^a(Z)]$. Under Assumption 4.2, $e(f_N)$, $e_N^a(z)$, and e_N^a are well-defined with $\mathbb{E}[e(f_N)^2]$, $\mathbb{E}[e_N^a(Z)^2] < \infty$. We now present our main result.

Theorem 4.2 (CLT for block-out CV Risk Estimator). *Suppose Assumptions 4.1-4.2 hold*

⁶In Appendix B, we consider an alternative asymptotic framework where $N_k \rightarrow \infty$ proportionally to n (i.e., B is fixed), and establish the validity of our approach under an alternative set of assumptions.

and N_k is fixed while $n \rightarrow \infty$. Then

$$\sqrt{n}(\widehat{e}_{N_k}^{CV} - e_{N_k}^a) \xrightarrow{d} \mathcal{N}(0, \sigma_{N_k}^2),$$

where the asymptotic variance $\sigma_{N_k}^2$ is given by

$$\sigma_{N_k}^2 \equiv N_k V_{N_k, \text{train}} + V_{N_k, \text{test}} + 2N_k C_{N_k}, \quad (4.4)$$

with $V_{N_k, \text{train}} \equiv \text{Var}(e(f_{N_k}))$, $V_{N_k, \text{test}} \equiv \text{Var}(e_{N_k}^a(Z))$, and $C_{N_k} \equiv \text{Cov}(e(f_{N_k}), e_{N_k}^a(Z_1))$, where $Z_1 \in \mathcal{D}_{N_k}$.

Importantly, we note that the asymptotic variance of the block-out CV estimator features three components: $V_{N_k, \text{train}}$ that captures the training-sample randomness contribution, $V_{N_k, \text{test}}$ that captures the testing-sample randomness contribution, and C_{N_k} that reflects the fact that each observation appears both in one training block and in the test sets of all other blocks, a key feature of cross-validation approaches.

4.2.1 Asymptotic Variance Estimation

We now construct an estimator $\widehat{\sigma}_{N_k}^2$ of the asymptotic variance $\sigma_{N_k}^2$ by plugging in estimators for each component of the three-term decomposition in (4.4), as described below.

Estimator of $V_{N_k, \text{train}}$ For each block b , define the block-wise population risk $e_{b, \text{train}} \equiv \mathbb{E}[\ell(f_{N_k}^{(b)}(X), Y) \mid \mathcal{D}_{\text{train}}^{(k, b)}]$. The $(e_{b, \text{train}})_{b=1}^{B_k}$ are i.i.d. with mean $e_{N_k}^a$ and variance $V_{N_k, \text{train}}$. We estimate $V_{N_k, \text{train}}$ by the sample variance of the block-wise CV risks:

$$\widehat{V}_{N_k, \text{train}} \equiv \frac{1}{B_k - 1} \sum_{b=1}^{B_k} (\widehat{e}_{k, b} - \widehat{e}_{N_k}^{CV})^2.$$

Estimator of $V_{N_k, \text{test}}$ Each i belongs to exactly one training block and appears in the test sets of all other $B_k - 1$ blocks. Define the block-averaged test loss for observation i :

$$\hat{\mu}_i \equiv \frac{1}{B_k - 1} \sum_{b: i \notin S_{k,b}} \ell(f_{N_k}^{(b)}(X_i), Y_i), \quad i = 1, \dots, n.$$

As $B_k \rightarrow \infty$, $\hat{\mu}_i \xrightarrow{p} e_{N_k}^a(Z_i)$. We estimate $V_{N_k, \text{test}}$ by:

$$\hat{V}_{N_k, \text{test}} \equiv \frac{1}{n - 1} \sum_{i=1}^n (\hat{\mu}_i - \bar{\mu})^2, \quad \bar{\mu} \equiv \frac{1}{n} \sum_{i=1}^n \hat{\mu}_i.$$

Estimator of C_{N_k} For each block b , define the block-wise summaries

$$\hat{e}_b \equiv \hat{e}_{k,b}, \quad \hat{m}_b \equiv \frac{1}{N_k} \sum_{i \in S_{k,b}} \hat{\mu}_i, \quad b = 1, \dots, B_k.$$

Let $\bar{r} \equiv B_k^{-1} \sum_{b=1}^{B_k} \hat{e}_b$ and $\bar{m} \equiv B_k^{-1} \sum_{b=1}^{B_k} \hat{m}_b$. We estimate the cross term by:

$$\hat{C}_{N_k} \equiv \frac{1}{B_k - 1} \sum_{b=1}^{B_k} (\hat{e}_b - \bar{r})(\hat{m}_b - \bar{m}).$$

Combining the above, we construct the asymptotic variance estimator as

$$\hat{\sigma}_{N_k}^2 \equiv N_k \hat{V}_{N_k, \text{train}} + \hat{V}_{N_k, \text{test}} + 2N_k \hat{C}_{N_k}, \quad (4.5)$$

whose validity is established in the following theorem.

Theorem 4.3 (Studentized CLT). *Suppose Assumptions 4.1 and 4.2 hold and N_k is fixed while $n \rightarrow \infty$. Let $\hat{\sigma}_{N_k}^2$ be defined by (4.5). Then:*

1. $\hat{\sigma}_{N_k}^2 \xrightarrow{p} \sigma_{N_k}^2$ as $n \rightarrow \infty$.
2. If $\sigma_{N_k}^2 > 0$, then the studentized statistic satisfies

$$\frac{\sqrt{n}(\hat{e}_{N_k}^{CV} - e_{N_k}^a)}{\hat{\sigma}_{N_k}} \xrightarrow{d} \mathcal{N}(0, 1).$$

4.3 A Valid Test of $H_{0,k} : e_{N_k}^a \leq e_{\text{LLM}}$

We finally show how to construct a valid hypothesis test of $H_{0,k} : e_{N_k}^a \leq e_{\text{LLM}}$ based on Theorem 4.3, which can then be used to implement the sequential testing procedure in Algorithm 3.1.

We base our test on the difference between the estimated risks:

$$\hat{\Delta}_k \equiv \hat{e}_{N_k}^{CV} - \hat{e}_{\text{LLM}}.$$

Since both terms are computed on the same dataset, they are correlated. To account for this correlation within the fixed- N_k asymptotic framework, we apply the general block-out CV theory directly to the *loss difference*. Define the loss difference for a given predictor f and observation $Z = (X, Y)$ as

$$\delta(f, Z) \equiv \ell(f(X), Y) - \ell(f_{\text{LLM}}(X), Y).$$

Note that the estimator $\hat{\Delta}_k$ is algebraically equivalent to the block-out CV estimator applied to this difference loss function.⁷ Applying Theorem 4.3 to the loss difference δ , we obtain the following result.

Proposition 4.1 (Test Validity). *Define the studentized test statistic:*

$$T_{\text{diff}, N_k, n} \equiv \frac{\hat{e}_{N_k}^{CV} - \hat{e}_{\text{LLM}}}{\hat{\sigma}_{\text{diff}, k} / \sqrt{n}}. \quad (4.6)$$

where $\hat{\sigma}_{\text{diff}, k}^2$ estimates the asymptotic variance of the difference. Under Assumptions 4.1 and 4.2, under the boundary of the null hypothesis ($e_{N_k}^a = e_{\text{LLM}}$):

$$T_{\text{diff}, N_k, n} \xrightarrow{d} \mathcal{N}(0, 1) \quad \text{as } n \rightarrow \infty.$$

⁷Note that $\hat{e}_{\text{LLM}}$ is the sample mean of the LLM loss. Since the LLM prediction rule is fixed, this is identical to the block-out CV estimator for a fixed function f_{LLM} .

provided that $\sigma_{\text{diff},k}^2 > 0$. Consequently, the one-sided test that rejects $H_{0,k}$ when $T_{N_k,n} > z_{1-\alpha}$ has asymptotic size α .

Remark 4.2 (Conservative Alternative). *In practice, implementing the exact difference variance estimator $\hat{\sigma}_{\text{diff},k}^2$ requires computing covariances between ML and LLM errors. A simpler alternative is to use the sum of their individual variances, i.e.,*

$$T_{\text{con},N_k,n} \equiv \frac{\hat{e}_{N_k}^{CV} - \hat{e}_{\text{LLM}}}{\sqrt{(\hat{\sigma}_{N_k}^2 + \hat{\sigma}_{\text{LLM}}^2)/n}}. \quad (4.7)$$

Note that $\sigma_{\text{diff},k}^2 = \sigma_{N_k}^2 + \sigma_{\text{LLM}}^2 - 2\text{Cov}(\sqrt{n}\hat{e}_{N_k}^{CV}, \sqrt{n}\hat{e}_{\text{LLM}})$. Assuming that the predictions of the ML model and the LLM are positively correlated (which is typical, as “hard” observations tend to yield high losses for both models), the covariance term is positive. In this case, the sum of variances overestimates the true variance ($\hat{\sigma}_{N_k}^2 + \hat{\sigma}_{\text{LLM}}^2 > \hat{\sigma}_{\text{diff},k}^2$), making the test statistic $T_{N_k,n}$ smaller and the test conservative (size $< \alpha$).

5 Applications

5.1 Data and Prediction Problems

The Panel Study of Income Dynamics (PSID) is one of the longest-running longitudinal household surveys in the United States, following individuals and families since 1968. It collects detailed information on economic, social, and health outcomes and is widely used in empirical research and policy analysis. We use the 2021 PSID wave and evaluate LLM prediction of four outcomes Y : hourly wage ($Y \in \mathbb{R}_+$), homeownership ($Y \in \{\text{own}, \text{rent}\}$), smoking ($Y \in \{0, 1\}$), and drinking ($Y \in \{0, 1\}$). In each prediction problem, the covariate vector X denotes the respondent’s characteristics, including age, sex, race, education, occupation, hours worked, labor market experience, self-reported health, state of residence, marital status, number of children, parents’ education, religion, and the other outcomes listed above except for the one being predicted.

For the continuous-valued outcome (hourly wage), we set the loss function to be squared-error loss $\ell(y, y') = (y' - y)^2$, although for readability we report the *root mean-squared error*

$$\text{RMSE}(f) = \sqrt{e(f)}$$

rather than the error $e(f)$. For the other discrete outcomes (homeownership, drinking, and smoking) we set the usual 0-1 loss $\ell(y, y') = 1\{y' \neq y\}$, so error is the misclassification rate.

5.2 LLM Prediction Rules

We generate LLM predictions as follows. For each test observation $i = 1, \dots, n$, we transform the covariate vector X_i into a natural-language *persona* description and then prompt the LLM to predict Y_i . Importantly, we do not fine-tune the LLM on PSID data and do not use few-shot examples drawn from PSID. Recall Prompt 1.

A natural concern is potential data leakage (Ludwig et al., 2025), namely that PSID data could have been included—directly or indirectly—in LLMs’ training corpora. While access to PSID data requires registration and human verification through the PSID Data Center, we nonetheless implement two safeguards: (i) we evaluate a static LLM with a documented training data descriptions and cutoffs: OpenAI’s ChatGPT-4 (cutoff: April 2023),⁸ and (ii) we use the 2021 wave of PSID released on October 2024,⁹ which postdates the reported knowledge cutoffs of the LLMs. (See Appendix D for supplementary results considering additional LLMs whose cutoffs also predate the 2021 PSID release.)

5.3 Domain-Specific Learning

We compare the LLM’s predictions to those of supervised ML models trained directly on PSID data. We consider two popular ML algorithms, which are briefly reviewed here.

⁸We access these models via an API. Based on interviews with current and former OpenAI employees, we expect that when ChatGPT is accessed via the API it is not subject to further fine-tuning (e.g., through reinforcement learning from human feedback), and that the reported knowledge cutoff is reliable.

⁹<https://www.icpsr.umich.edu/web/sbeccc/studies/39190/versions/V1>

Lasso regression and Logit L1. Lasso is a regularized linear prediction method. Different from standard OLS, it seeks to maximize fit to the data subject to a penalty for model complexity, evaluated as the sum of the absolute value of coefficients. Lasso regression thus encourages small coefficients, and often leads to sparse linear models in which many coefficients are set equal to zero. Logit L1 applies the same ℓ_1 -regularization principle in a logistic regression framework for binary outcomes.

Random forests. Random forests are a flexible prediction method built by averaging many decision trees. Each component tree is trained on a different subsample of the data and uses a different subset of covariates, and the final prediction aggregates across them. This procedure allows the algorithm to capture nonlinear relationships and interactions without requiring these features to be specified in advance.

For each ML model, we compute the error curve, namely, the test error as a function of the training sample size N . We tune hyperparameters (e.g., the lasso penalty parameter, the depth and leaf size for RF) for every N using k -fold CV; see Section C in the Appendix for the implementation details. These error curve estimates are subsequently compared to the LLM benchmark to compute the ESS.

5.4 Results

We estimate error curves for each prediction problem and compute the equivalent sample size of the LLM. Figures 3–6 present the full estimated error curves, the estimated equivalent sample size $\hat{N}^*$, and confidence intervals for all estimated quantities. The reported confidence intervals for the block-out CV errors are calculated as hybrid CIs: they are based on the fixed- N asymptotic theory (Theorem 4.2) when $N \leq 400$ and the fixed- B asymptotic theory (Corollary B.1 in Section B) otherwise.¹⁰ The confidence intervals for N^* are calculated based on (3.5) in Remark 3.1. Table 1 more succinctly presents our confidence intervals for

¹⁰This choice is made based on calibrated simulation to ensure desirable coverage performance.

the equivalent sample size.

Table 1: Equivalent sample size across prediction tasks

Outcome	LLM Error	Algorithm	One-Sided CI
Hourly Wage	12.50	Lasso	$[20, \infty)$
		RF	$[17, \infty)$
Home Ownership	0.24	Logit L1	$[676, \infty)$
		RF	$[590, \infty)$
Drinking	0.4	Logit L1	$[59, \infty)$
		RF	$[41, \infty)$
Smoking	0.23	Logit L1	$[1, \infty)$
		RF	$[1, \infty)$

We find substantial heterogeneity in the LLM’s predictive performance across outcomes. The LLM performs extremely well in predicting homeownership: at the 95% confidence level, its equivalent sample size is at least 590 observations when benchmarked against a random forest and at least 676 observations when benchmarked against L1-penalized logit.

The LLM’s performance is weaker, but still meaningful, for predicting drinking behavior. At the same confidence level, the LLM’s equivalent sample size is at least 41 observations relative to the random forest and at least 59 observations relative to L1-penalized logit.

Predictive performance deteriorates further for hourly wages. Here, we can only conclude that the LLM’s equivalent sample size is at least 17 observations when compared to the random forest and at least 20 observations when compared to L1-penalized logit.

Finally, the LLM provides essentially no predictive value for smoking behavior: we fail to reject the null hypothesis that conventional machine-learning algorithms outperform the LLM even when trained on a single observation.

Taken together, these results illustrate that the value of LLM predictions is highly task-specific—ranging from a close substitute for hundreds of observations in some settings to providing no meaningful predictive content in others. This underscores the importance of evaluating the LLM’s predictive performance across domains.

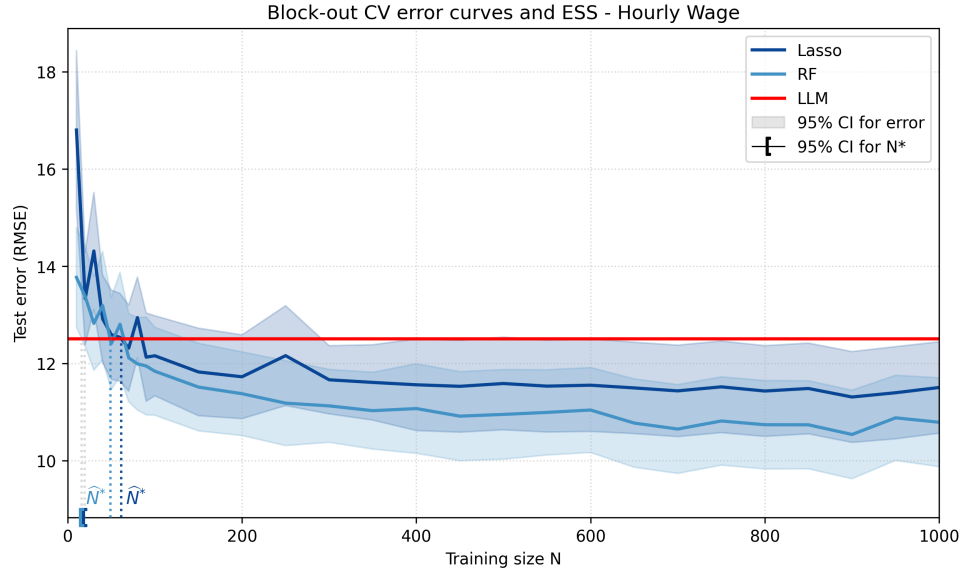


Figure 3: RMSE across training sample sizes (Y : Hourly wage)

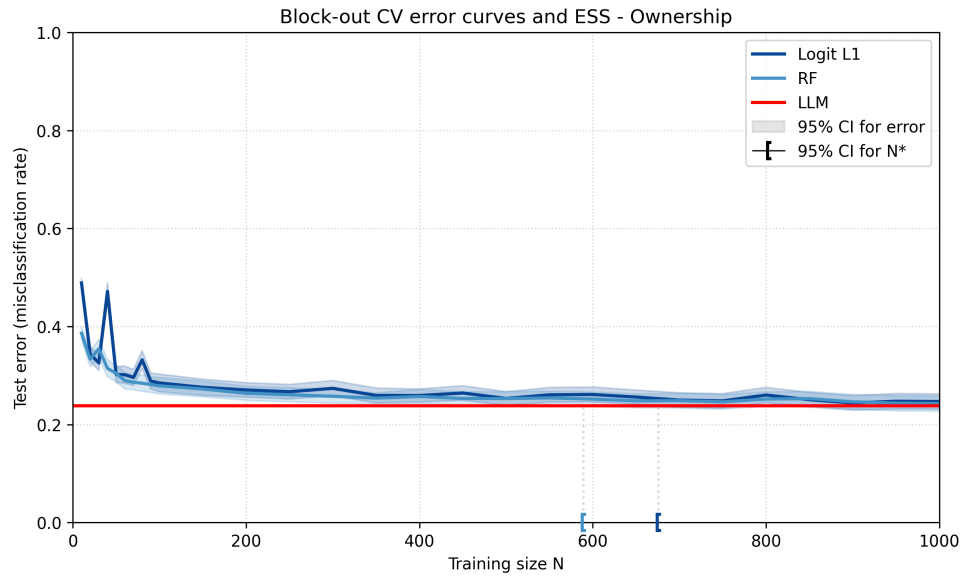


Figure 4: Misclassification rate across training sample sizes (Y : Homeownership)

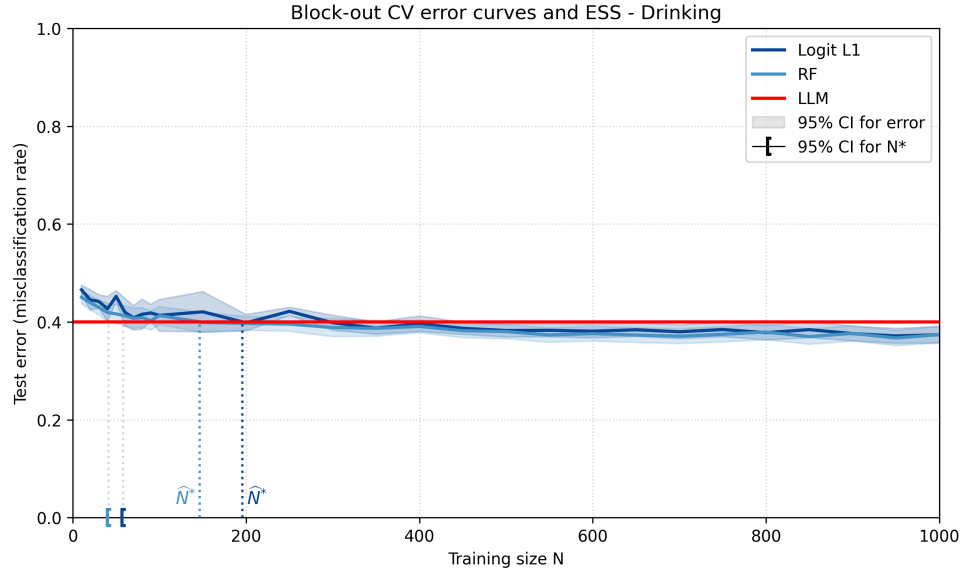


Figure 5: Misclassification rate across training sample sizes (Y : Drinking)

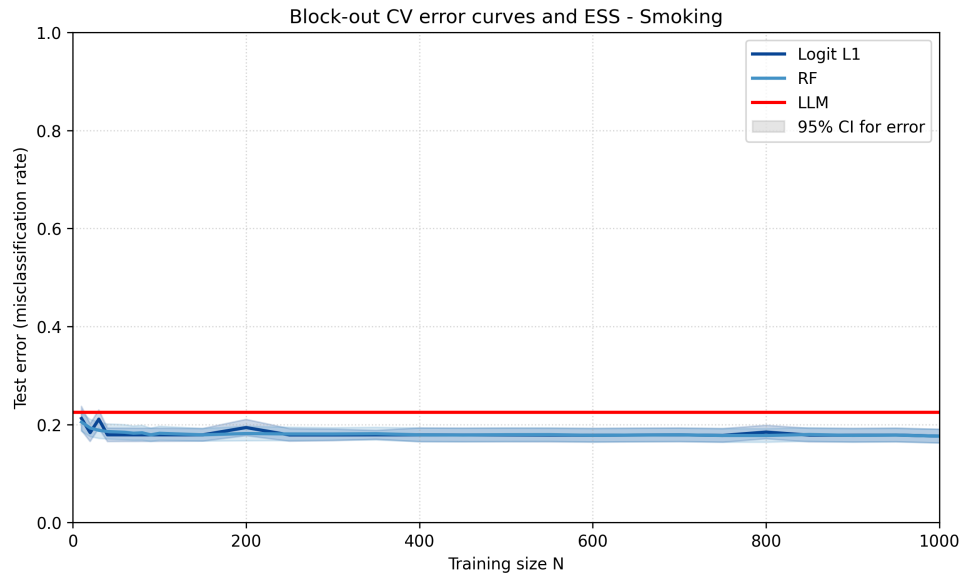


Figure 6: Misclassification rate across training sample sizes (Y : Smoking)

6 Extension: ESS for Treatment Effects

LLMs are considered useful tools for causal inference (Kiciman et al., 2023; Wang et al., 2024; Han, 2025; Liu et al., 2025). In light of this line of work, we extend our evaluation framework for LLMs that are employed for treatment effect estimation.

Let $\mu(X) \equiv \mathbb{E}_P[Y \mid X]$. Under squared loss $\ell(y, \hat{y}) = (y - \hat{y})^2$, the population risk admits the standard bias–variance-type decomposition

$$\begin{aligned} e(f) &\equiv \mathbb{E}_P[(Y - f(X))^2] = \mathbb{E}_P[(Y - \mu(X))^2] + \mathbb{E}_P[(\mu(X) - f(X))^2] \\ &\equiv e_{\text{irreducible}} + e_\mu(f), \end{aligned}$$

where the equality holds as the cross term is zero by the law of iterated expectations and the irreducible error is achieved as $\mu(X)$ uniquely minimizes $\mathbb{E}_P[(Y - f(X))^2]$. Therefore, any ESS defined with respect to $e(f)$ would be equivalent to an ESS defined with respect to $e_\mu(f)$.

Now consider the estimation of the conditional average treatment effect (CATE) with data (X, T, Y) distributed as P , where T is an indicator for an additional treatment. Assume unconfoundedness: $T \perp (Y_0, Y_1) \mid X$ where Y_t is the counterfactual outcome for $t \in \{0, 1\}$. Then, the conditional mean of Y_t is identified as $\mu_t(X) \equiv \mathbb{E}_P[Y_t \mid X] = \mathbb{E}_P[Y \mid X, T = t]$. Within treatment arm d and under squared loss $\ell(y, \hat{y}) = (y - \hat{y})^2$, the conditional risk admits the standard bias–variance-type decomposition:

$$\begin{aligned} e_t(f) &\equiv \mathbb{E}_P[(Y_t - f(X))^2] \\ &= \mathbb{E}_P[(Y_t - \mu_t(X))^2] + \mathbb{E}_P[(\mu_t(X) - f(X))^2] \\ &\equiv e_{\text{irreducible},t} + e_{\mu_t}(f), \end{aligned}$$

by an argument analogous to above. Therefore, any ESS defined with respect to $e_t(f)$ would be equivalently an ESS defined with respect to $e_{\mu_t}(f)$. Moreover, $e_t(f)$ can be computed

from observed data for each $T = t$, as $\mathbb{E}_P[(Y_t - f(X))^2] = \mathbb{E}_P[\mathbb{E}_P[(Y - f(X))^2 \mid X, T = t]]$ by unconfoundedness.

For the LLM, we can estimate $e_t(f_{\text{LLM}})$ by applying the prompt strategy in Section 2.2 to the test sample restricted to $T = t$ (i.e., to the treated or control test sample).¹¹ For ML trained on each sample of size N , define

$$e_{N,t}^a \equiv \mathbb{E}_{D_N} e_t(f_N) = \mathbb{E}_{D_N} e_t(a(D_N)), \quad D_N \equiv \{(x_i, t_i, y_i)\}_{i=1}^N,$$

where $e_{N,t}^a$ can be evaluated via block-out cross-validation within the $T = t$ subsample. Then, for $t \in \{0, 1\}$ the treatment-specific ESS is defined as

$$N_t^* \equiv \min\{N : e_{N,t}^a \leq e_t(f_{\text{LLM}})\}.$$

Define the CATE $\tau(X) \equiv \mathbb{E}_P[Y_1 - Y_0 \mid X]$ and the CATE risk

$$e_\tau(g) \equiv \mathbb{E}_P[(\tau(X) - g(X))^2].$$

For both the LLM and ML, using prediction rules (f_1, f_0) corresponding to $(e_{\mu_1}(f_1), e_{\mu_0}(f_0))$ only yields the upper bound on $e_\tau(f_1 - f_0)$ due to the nonlinearity of the loss function, and thus defining ESS for τ through $f_1 - f_0$ is elusive.¹² We thus proceed as follows. Instead of prompting the LLM to predict Y_t , we prompt it to directly “predict” $\tau(X)$. An example prompt can be as follows:

Prompt 2. *Here is your persona: You are a 37-year-old male living in Tennessee in 2020. You work in architecture and engineering occupations in manufacturing industry. You have 4 years of work experience in the current job. You have completed 14 years of education. You are married and have 1 child. Your father has 12 years of education. Your mother has 12 years of education. Your race is white. Your health is excellent. You are Protestant.*

¹¹This approach relates to counterfactual generation using LLMs in clinical trials (Wang et al., 2024).

¹²Specifically, we only have $e_\tau(f_1 - f_0) = \mathbb{E}_P[(\tau(X) - (f_1(X) - f_0(X)))^2] \leq 2\{e_{\mu_1}(f_1) + e_{\mu_0}(f_0)\}$.

Finally, you are a subject in a randomized controlled trial in which the treatment is the provision of health insurance, specifically Medicaid.

Given your persona, what would be the effect of the randomized provision of health insurance on the number of your primary care visits?

Let $g_{\text{LLM}}(x)$ denote such a prediction. For the benchmark algorithm, let $\tilde{a} : \mathcal{D} \rightarrow \mathcal{G}$ map data to a CATE function $g \in \mathcal{G}$ and define

$$e_{N,\tau}^{\tilde{a}} \equiv \mathbb{E}_{D_N} e_{\tau}(\tilde{a}(D_N)).$$

Then define the CATE-ESS as

$$N_{\tau}^* \equiv \min\{N : e_{N,\tau}^{\tilde{a}} \leq e_{\tau}(g_{\text{LLM}})\}.$$

In practice, because τ is unknown, we proceed via the transformed outcome. Under unconfoundedness, it can be shown that¹³

$$\mathbb{E}_P[\tilde{Y} \mid X] = \tau(X),$$

where

$$\tilde{Y} \equiv Y \frac{T - \pi(X)}{\pi(X)(1 - \pi(X))}, \quad \pi(X) \equiv \mathbb{P}[T = 1 \mid X].$$

Thus,

$$\begin{aligned} e_{\tau}(g) &= \mathbb{E}_P[(\tau(X) - g(X))^2] = \mathbb{E}_P[(\tilde{Y} - g(X))^2] - \mathbb{E}_P[(\tilde{Y} - \tau(X))^2] \\ &= \tilde{e}(g) - \tilde{e}_{\text{irreducible}}, \end{aligned}$$

¹³Observe that

$$\begin{aligned} E[\tilde{Y}|X] &= P(T = 1|X)E[\tilde{Y}|X, T = 1] + P(T = 0|X)E[\tilde{Y}|X, T = 0] \\ &= \pi(X) \left[\frac{E[Y_1|X]}{\pi(X)} \right] + (1 - \pi(X)) \left[\frac{-E[Y_0|X]}{1 - \pi(X)} \right] = E[Y_1|X] - E[Y_0|X], \end{aligned}$$

where unconfoundedness is applied in the second equality.

where $\tilde{e}_{\text{irreducible}}$ does not depend on g . Consequently the ESS does not change if we replace e_τ by $\tilde{e}$:

$$N_\tau^* = \tilde{N}^* \equiv \min\{N : \tilde{e}_N^a \leq \tilde{e}(g_{\text{LLM}})\}, \quad \tilde{e}_N^a \equiv \mathbb{E}_{D_N} \tilde{e}(\tilde{a}(D_N)).$$

In practice, $\tilde{e}(g_N)$ can be estimated by computing $\tilde{Y}$ (e.g., with known $\pi(X)$ in a randomized control trial), and the ML model can simply predict $\tilde{Y}$ based on X .

7 Conclusions

LLMs have become popular prediction tools in empirical social science. We measure the value of LLMs’ pretrained knowledge using the *equivalent sample size* (ESS): the training sample size required for an algorithm trained on domain-specific data to match the LLM benchmark performance. We propose an inference framework for ESS based on sequential hypothesis tests and a new asymptotic theory for cross-validation error under block-out CV. The framework can help researchers quantify and communicate the capability (or limitation) of LLMs in domain-specific prediction tasks.

References

- AHER, G. V., R. I. ARRIAGA, AND A. T. KALAI (2023): “Using large language models to simulate multiple humans and replicate human subject studies,” in *International conference on machine learning*, PMLR, 337–371.
- ANDREWS, I., D. FUDENBERG, L. LEI, A. LIANG, AND C. WU (2025): “The Transfer Performance of Economic Models,” Working paper.
- ANTHIS, J. R., R. LIU, S. M. RICHARDSON, A. C. KOZLOWSKI, B. KOCH, E. BRYNJOLFSSON, J. EVANS, AND M. S. BERNSTEIN (2025): “Position: LLM Social Simulations Are a Promising Research Method,” in *Proceedings of the 42nd International Conference on Machine Learning (ICML)*, position Paper Track.

- ARGYLE, L. P., E. C. BUSBY, N. FULDA, J. R. GUBLER, C. RYTTING, AND D. WINGATE (2023): “Out of one, many: Using language models to simulate human samples,” *Political Analysis*, 31, 337–351.
- AUSTERN, M. AND W. ZHOU (2020): “Asymptotics of Cross-Validation,” *arXiv preprint arXiv:2001.11111*.
- BATES, S., T. HASTIE, AND R. TIBSHIRANI (2024): “Cross-validation: what does it estimate and how well does it do it?” *Journal of the American Statistical Association*, 119, 1434–1445.
- BOHREN, J. A., P. HULL, AND A. IMAS (2025): “Systemic Discrimination: Theory and Measurement,” *The Quarterly Journal of Economics*, 140, 1743–1799.
- CHAN-LAU, J. A., T. L. QUACH, AND A. R. SHI (2025): “Forecasting Federal Funds Rates with AI: LSTM, GRU, and Large Language Model Approaches,” Working paper (wp/25-10), ASEAN+3 Macroeconomic Research Office (AMRO), september 15, 2025.
- EREV, I., A. E. ROTH, R. L. SLONIM, AND G. BARRON (2007): “Learning and Equilibrium as Useful Approximations: Accuracy of Prediction on Randomly Selected Constant Sum Games,” *Economic Theory*, 33, 29–51.
- FARIA-E CASTRO, M. AND F. LEIBOVICI (2024): “Artificial Intelligence and Inflation Forecasts,” *Federal Reserve Bank of St. Louis Review*, 106, 1–14.
- FAVA, B. (2025): “Training and Testing with Multiple Splits: A Central Limit Theorem for Split-Sample Estimators,” *arXiv preprint arXiv:2511.04957*.
- FUDENBERG, D., W. GAO, AND A. LIANG (2026): “How Flexible Is That Functional Form? Measuring the Restrictiveness of Theories,” *The Review of Economics and Statistics*, forthcoming.
- FUDENBERG, D., J. KLEINBERG, A. LIANG, AND S. MULLAINATHAN (2022): “Measuring the Completeness of Economic Models,” *Journal of Political Economy*, 130, 956–990.

- GAO, Y., D. LEE, G. BURTCH, AND S. FAZELPOUR (2025): “Take caution in using LLMs as human surrogates,” *Proceedings of the National Academy of Sciences*, 122, e2501660122.
- HAN, S. (2025): “Mining causality: Ai-assisted search for instrumental variables,” *arXiv preprint arXiv:2409.14202*.
- HEWITT, L., A. ASHOKKUMAR, I. GHEZAE, AND R. WILLER (2024a): “Predicting Results of Social Science Experiments Using Large Language Models,” Working paper.
- (2024b): “Predicting results of social science experiments using large language models,” *Preprint*.
- HOLM, S. (1979): “A simple sequentially rejective multiple test procedure,” *Scandinavian journal of statistics*, 65–70.
- HORTON, J. J. (2023): “Large language models as simulated economic agents: What can we learn from homo silicus?” Tech. rep., National Bureau of Economic Research.
- HUANG, C., Y. WU, AND K. WANG (2025): “How Many Human Survey Respondents is a Large Language Model Worth? An Uncertainty Quantification Perspective,” *arXiv preprint arXiv:2502.17773*.
- JANSEN, B. J., S.-G. JUNG, AND J. SALMINEN (2023): “Employing large language models in survey research,” *Natural Language Processing Journal*, 4, 100020.
- KICIMAN, E., R. NESS, A. SHARMA, AND C. TAN (2023): “Causal reasoning and large language models: Opening a new frontier for causality,” *Transactions on Machine Learning Research*.
- LEI, J. (2025): “A Modern Theory of Cross-Validation through the Lens of Stability,” *arXiv preprint arXiv:2505.23592*.
- LIU, X., P. XU, J. WU, J. YUAN, Y. YANG, Y. ZHOU, F. LIU, T. GUAN, H. WANG, T. YU, ET AL. (2025): “Large language models and causal inference in collaboration:

- A comprehensive survey,” *Findings of the Association for Computational Linguistics: NAACL 2025*, 7668–7684.
- LUDWIG, J., S. MULLAINATHAN, AND A. RAMBACHAN (2025): “Large language models: An applied econometric framework,” Tech. rep., National Bureau of Economic Research.
- MANNING, B. S., K. ZHU, AND J. J. HORTON (2024): “Automated Social Science: Language Models as Scientist and Subjects,” *National Bureau of Economic Research*.
- PARK, J. S., C. Q. ZOU, A. SHAW, B. M. HILL, C. CAI, M. R. MORRIS, R. WILLER, P. LIANG, AND M. S. BERNSTEIN (2024): “Generative agent simulations of 1,000 people,” *arXiv preprint arXiv:2411.10109*.
- ROSENTHAL, H. P. (1970): “On the subspaces of $L_p(p, 2)$ spanned by sequences of independent random variables,” *Israel Journal of Mathematics*, 8, 273–303.
- SUN, S., E. LEE, D. NAN, X. ZHAO, W. LEE, B. J. JANSEN, AND J. H. KIM (2024): “Random silicon sampling: Simulating human sub-population opinion using a large language model based on group-level demographic information,” *arXiv preprint arXiv:2402.18144*.
- WANG, Y., T. FU, Y. XU, Z. MA, H. XU, B. DU, Y. LU, H. GAO, J. WU, AND J. CHEN (2024): “TWIN-GPT: digital twins for clinical trials via large language model,” *ACM Transactions on Multimedia Computing, Communications and Applications*.

A Block-Out CV under Fixed- N Asymptotics

This appendix contains proofs of the main asymptotic results stated in Section 4.2. Throughout, we fix the training size N_k and let $n \rightarrow \infty$. For notational simplicity, we suppress the subscript k and write N for N_k , B for B_k , S_b for $S_{k,b}$, etc.

Recall the following definitions. The block-out CV estimator is

$$\widehat{e}_N^{\text{CV}} = \frac{1}{B} \sum_{b=1}^B \widehat{e}_b, \quad \widehat{e}_b = \frac{1}{M_n} \sum_{i \notin S_b} \ell(f_N^{(b)}(X_i), Y_i),$$

where $M_n = n - N$ and $f_N^{(b)} = a(\mathcal{D}_{\text{train}}^{(b)})$. For each block b , define the block-wise error

$$e_{b,\text{train}} \equiv \mathbb{E}[\ell(f_N^{(b)}(X), Y) \mid \mathcal{D}_{\text{train}}^{(b)}].$$

By construction, $(e_{b,\text{train}})_{b=1}^B$ are i.i.d. with $\mathbb{E}[e_{b,\text{train}}] = e_N^a$ and $\text{Var}(e_{b,\text{train}}) = V_{N,\text{train}}$.

Proof of Theorem 4.2. Decompose the CV estimator as

$$\widehat{e}^{\text{CV}} = A_n + B_n, \quad A_n \equiv \frac{1}{B} \sum_{b=1}^B e_{b,\text{train}}, \quad B_n \equiv \frac{1}{B} \sum_{b=1}^B (\widehat{e}_b - e_{b,\text{train}}).$$

First, we consider the training contribution A_n . The training blocks $\mathcal{D}_{\text{train}}^{(b)} = (Z_{(b-1)N+1}, \dots, Z_{bN})$ for $b = 1, \dots, B$ are i.i.d, and thus $e_{b,\text{train}}$ are i.i.d. with

$$\mathbb{E}[e_{b,\text{train}}] = \mathbb{E}_{D_N}[e(f_N)] = e_N^a, \quad \text{Var}(e_{b,\text{train}}) = V_{N,\text{train}}.$$

By the classical CLT,

$$\sqrt{B}(A_n - e_N^a) \xrightarrow{d} \mathcal{N}(0, V_{N,\text{train}}).$$

Since $B = n/N$ under our assumption that $n = BN$, we have

$$\sqrt{n}(A_n - e_N^a) = \sqrt{N} \cdot \sqrt{B}(A_n - e_N^a) \xrightarrow{d} \mathcal{N}(0, NV_{N,\text{train}}).$$

Second, we consider the testing contribution B_n . Consider the centered loss at $z \equiv (x, y)$,

$$\psi(D_N, z) \equiv \ell(f_{D_N}(x), y) - e(f_N),$$

Then for each block b and each $i \notin S_b$,

$$\ell(f_N^{(b)}(X_i), Y_i) - e_{b,\text{train}} = \psi(\mathcal{D}_{\text{train}}^{(b)}, Z_i).$$

Thus

$$B_n = \frac{1}{BM_n} \sum_{b=1}^B \sum_{i \notin S_b} \psi(\mathcal{D}_{\text{train}}^{(b)}, Z_i).$$

For each pair (b, i) with $i \notin S_b$, the training set $\mathcal{D}_{\text{train}}^{(b)}$ and the test point Z_i are independent. Therefore, the projection onto the test point is

$$\mathbb{E}[\psi(\mathcal{D}_{\text{train}}^{(b)}, Z_i) \mid Z_i] = e_N^a(Z_i) - e_N^a.$$

In the meanwhile, the projection onto a training point Z_j (where $j \in S_b$) involves taking the expectation over the test point Z_i conditional on the training set, which equals zero by the construction of the centered kernel ψ : $\mathbb{E}[\psi(\mathcal{D}_{\text{train}}^{(b)}, Z_i) \mid \mathcal{D}_{\text{train}}^{(b)}] = 0$.

We define T_n as the first-order Hájek projection of B_n :

$$T_n \equiv \sum_{j=1}^n \mathbb{E}[B_n \mid Z_j].$$

Since each observation Z_j serves as a test point for exactly $B - 1$ blocks (specifically, all blocks b such that $j \notin S_b$), we obtain:

$$T_n = \sum_{j=1}^n \frac{B-1}{BM_n} (e_N^a(Z_j) - e_N^a) = \frac{B-1}{BM_n} \sum_{i=1}^n (e_N^a(Z_i) - e_N^a).$$

Define the remainder $\mathcal{R}_n \equiv B_n - T_n$. Clearly,

$$\frac{B-1}{BM_n} = \frac{1}{n} + O(n^{-2}),$$

with the implicit constant depending only on N .

The remainder $\mathcal{R}_n$ is a higher-order term in the sense of U/V-statistics. Note that the term B_n constitutes an incomplete V-statistic of order $N + 1$ with kernel

$$h_N(z_1, \dots, z_N, z_{N+1}) \equiv \psi((z_1, \dots, z_N), z_{N+1}).$$

and T_n is the first-order Hájek projection of B_n . The remainder $\mathcal{R}_n \equiv B_n - T_n$ represents the degenerate component of the statistic. By construction of the projection, $\mathcal{R}_n$ is uncorrelated with any linear function of the observations, and $\mathbb{E}[\mathcal{R}_n] = 0$. To bound its variance, we observe that $\mathcal{R}_n$ can be expressed as a sum of degenerate kernels $\eta(\cdot)$ over the block structure:

$$\mathcal{R}_n = \frac{1}{BM_n} \sum_{b=1}^B \sum_{i \notin S_b} \eta(\mathcal{D}_{\text{train}}^{(b)}, Z_i),$$

where η satisfies the degeneracy properties $\mathbb{E}[\eta \mid Z_i] = 0$ (orthogonal to test point) and $\mathbb{E}[\eta \mid \mathcal{D}_{\text{train}}^{(b)}] = 0$ (orthogonal to training set). Consequently, in the expansion of the second moment $\mathbb{E}[\mathcal{R}_n^2]$, cross-terms $\mathbb{E}[\eta(\mathcal{D}_{\text{train}}^{(b)}, Z_i)\eta(\mathcal{D}_{\text{train}}^{(b')}, Z_j)]$ vanish if the two terms share the same block ($b = b'$ but $i \neq j$) due to the second degeneracy property. Cross-terms with disjoint blocks ($b \neq b'$) also vanish unless the indices overlap in a specific way. The only non-vanishing terms arise from configurations where training and test indices overlap across different blocks. Given the disjoint block design, the number of such non-vanishing terms scales as $O(n^2)$, while the normalization factor is $(BM_n)^{-2} = O(n^{-2})$. This implies the variance scaling:

$$\mathbb{E}[\mathcal{R}_n^2] = O(n^{-2}).$$

By Chebyshev's inequality, for any $\epsilon > 0$, $\mathbb{P}(|\sqrt{n}\mathcal{R}_n| > \epsilon) \leq \frac{n\mathbb{E}[\mathcal{R}_n^2]}{\epsilon^2} \rightarrow 0$, ensuring that $\sqrt{n}\mathcal{R}_n \xrightarrow{p} 0$. Third, we consider the covariance between A_n and B_n . Since $\sqrt{n}\mathcal{R}_n \xrightarrow{p} 0$ and $\sqrt{n}(A_n - e_N^a) = O_p(1)$, we have

$$\text{Cov}(\sqrt{n}(A_n - e_N^a), \sqrt{n}\mathcal{R}_n) \rightarrow 0.$$

Hence, it suffices to study the covariance between A_n and T_n . Using the definition of T_n ,

$$\text{Cov}(A_n, T_n) = \frac{B-1}{B^2 M_n} \sum_{b=1}^B \sum_{i=1}^n \text{Cov}(e_{b,\text{train}}, e_N^a(Z_i)).$$

If $i \notin S_b$, then $e_{b,\text{train}}$ depends only on $\{Z_j : j \in S_b\}$ whereas $e_N^a(Z_i)$ depends only on Z_i and an independent training sample, so $e_{b,\text{train}}$ and $e_N^a(Z_i)$ are independent and their covariance is 0. If $i \in S_b$, then the pair $(e_{b,\text{train}}, e_N^a(Z_i))$ has the same distribution as $(e(f_N), e_N^a(Z_1))$ with $D_N = (Z_1, \dots, Z_N) \sim P^N$. Hence,

$$\text{Cov}(e_{b,\text{train}}, e_N^a(Z_i)) = C_N \quad \text{for all } b \text{ and } i \in S_b.$$

Observing that there are exactly $BN = n$ pairs (b, i) with $i \in S_b$, we deduce

$$\text{Cov}(A_n, T_n) = \frac{B-1}{B^2 M_n} \cdot BN \cdot C_N = \frac{B-1}{B} \cdot \frac{N}{M_n} \cdot C_N = \frac{N}{n} C_N.$$

and consequently

$$\text{Cov}(\sqrt{n}(A_n - e_N^a), \sqrt{n}T_n) = n \text{Cov}(A_n, T_n) \rightarrow NC_N.$$

Third, we consider the covariance between A_n and B_n . Since $\sqrt{n}\mathcal{R}_n \xrightarrow{p} 0$ and $\sqrt{n}(A_n - e_N^a) = O_p(1)$, we have

$$\text{Cov}(\sqrt{n}(A_n - e_N^a), \sqrt{n}\mathcal{R}_n) \rightarrow 0.$$

Hence, it suffices to study the covariance between A_n and T_n . Using the definition of T_n ,

$$\text{Cov}(A_n, T_n) = \frac{B-1}{B^2 M_n} \sum_{b=1}^B \sum_{i=1}^n \text{Cov}(e_{b,\text{train}}, e_N^a(Z_i)).$$

If $i \notin S_b$, then $e_{b,\text{train}}$ depends only on $\{Z_j : j \in S_b\}$ whereas $e_N^a(Z_i)$ depends only on Z_i and an independent training sample, so $e_{b,\text{train}}$ and $e_N^a(Z_i)$ are independent and their covariance

is 0. If $i \in S_b$, then the pair $(e_{b,\text{train}}, e_N^a(Z_i))$ has the same distribution as $(e(f_N), e_N^a(Z_1))$ with $D_N = (Z_1, \dots, Z_N) \sim P^N$. Hence,

$$\text{Cov}(e_{b,\text{train}}, e_N^a(Z_i)) = C_N \quad \text{for all } b \text{ and } i \in S_b.$$

Observing that there are exactly $BN = n$ pairs (b, i) with $i \in S_b$, we deduce

$$\text{Cov}(A_n, T_n) = \frac{B-1}{B^2 M_n} \cdot BN \cdot C_N = \frac{B-1}{B} \cdot \frac{N}{M_n} \cdot C_N = \frac{N}{n} C_N + O(n^{-2}).$$

and consequently

$$\text{Cov}(\sqrt{n}(A_n - e_N^a), \sqrt{n}T_n) = n \text{Cov}(A_n, T_n) \rightarrow NC_N.$$

Finally, combining the analysis above, we deduce

$$\sqrt{n}(A_n - e_N^a) \xrightarrow{d} \mathcal{N}(0, NV_{N,\text{train}}), \quad \sqrt{n}B_n \xrightarrow{d} \mathcal{N}(0, V_{N,\text{test}}),$$

and the joint limit is bivariate normal with covariance $\lim_n \text{Cov}(\sqrt{n}(A_n - e_N^a), \sqrt{n}B_n) = NC_N$.

Thus

$$\sqrt{n}(\widehat{e}^{\text{CV}} - e_N^a) = \sqrt{n}(A_n - e_N^a) + \sqrt{n}B_n \xrightarrow{d} \mathcal{N}(0, \sigma_N^2),$$

with $\sigma_N^2 = NV_{N,\text{train}} + V_{N,\text{test}} + 2NC_N$. □

The proof of Theorem 4.3 relies on showing that each component of the plug-in variance estimator is consistent. We establish this through three lemmas.

Lemma A.1. *Under Assumptions 4.1, 4.2, and $n = BN$,*

$$\widehat{V}_{N,\text{train}} \xrightarrow{p} V_{N,\text{train}}.$$

Proof of Lemma A.1. Write $\widehat{e}_b = e_{b,\text{train}} + U_b$ where

$$U_b \equiv \frac{1}{M_n} \sum_{i \notin S_b} Y_i, \quad \text{with} \quad Y_i \equiv \ell(f_N^{(b)}(X_i), Y_i) - e_{b,\text{train}}.$$

Conditional on the training set $\mathcal{D}_{\text{train}}^{(b)}$, the variables $\{Y_i\}_{i \notin S_b}$ are independent with mean zero. By Assumption 4.2, they satisfy $\mathbb{E}[|Y_i|^{2+\delta}] < \infty$ for some $\delta > 0$.

We recall Rosenthal's inequality (Rosenthal, 1970): Let $X_1, \dots, X_m$ be independent zero-mean random variables with finite p -th moments for $p \geq 2$. Then there exists a constant C_p depending only on p such that

$$\mathbb{E} \left[\left| \sum_{j=1}^m X_j \right|^p \right] \leq C_p \left(\sum_{j=1}^m \mathbb{E}[|X_j|^p] + \left(\sum_{j=1}^m \mathbb{E}[X_j^2] \right)^{p/2} \right).$$

Applying this to U_b with $p = 2 + \delta$ and $m = M_n$, and noting that $\sum \mathbb{E}[|Y_i|^{2+\delta}] = M_n \mathbb{E}[|Y_1|^{2+\delta}]$ and $\sum \mathbb{E}[Y_i^2] = M_n \text{Var}(Y_1)$, we obtain:

$$\begin{aligned} \mathbb{E}[|U_b|^{2+\delta} \mid \mathcal{D}_{\text{train}}^{(b)}] &= \frac{1}{M_n^{2+\delta}} \mathbb{E} \left[\left| \sum_{i \notin S_b} Y_i \right|^{2+\delta} \mid \mathcal{D}_{\text{train}}^{(b)} \right] \\ &\leq \frac{C_{2+\delta}}{M_n^{2+\delta}} \left(M_n \mathbb{E}[|Y_1|^{2+\delta}] + (M_n \text{Var}(Y_1))^{1+\delta/2} \right) \\ &= O(M_n^{-(1+\delta)} + M_n^{-(1+\delta/2)}) \\ &= O(n^{-(1+\delta/2)}). \end{aligned}$$

To bound the maximum error across blocks, we apply the union bound and Markov's inequality:

$$\Pr \left(\max_{1 \leq b \leq B} |U_b| > \epsilon \right) \leq \sum_{b=1}^B \frac{\mathbb{E}[|U_b|^{2+\delta}]}{\epsilon^{2+\delta}} \leq B \cdot \frac{O(n^{-(1+\delta/2)})}{\epsilon^{2+\delta}}.$$

Since $B = O(n)$, the probability is bounded by $O(n^{-\delta/2})$, which converges to zero as $n \rightarrow \infty$.

Thus, $\max_b |U_b| \xrightarrow{P} 0$.

Finally, note that $(e_{b,\text{train}})$ are i.i.d. with variance $V_{N,\text{train}}$, so the sample variance

$$S_B^2(e_{1,\text{train}}, \dots, e_{B,\text{train}}) \xrightarrow{p} V_{N,\text{train}}$$

by the standard LLN. Since the sample variance S_B^2 is continuous in e_b 's and $\max_b |\widehat{e}_b - e_{b,\text{train}}| = \max_b |U_b| \xrightarrow{p} 0$, we have

$$S_B^2(\widehat{e}_1, \dots, \widehat{e}_B) - S_B^2(e_{1,\text{train}}, \dots, e_{B,\text{train}}) \xrightarrow{p} 0$$

and thus $S_B^2(\widehat{e}_1, \dots, \widehat{e}_B) \xrightarrow{p} V_{N,\text{train}}$. □

Lemma A.2. *Under Assumptions 4.1, 4.2, and $n = BN$,*

$$\widehat{V}_{N,\text{test}} \xrightarrow{p} V_{N,\text{test}}.$$

Proof of Lemma A.2. Recall that $\widehat{\mu}_i = \frac{1}{B-1} \sum_{b:i \notin S_b} \ell(f_N^{(b)}(X_i), Y_i)$. Conditional on Z_i , the summands are independent and identically distributed with mean $e_N^a(Z_i)$. Let $v_N(z) \equiv \text{Var}(\ell(f_N(x), y))$ denote the conditional variance of the loss given a test point $z = (x, y)$. The conditional variance of the average is:

$$\text{Var}(\widehat{\mu}_i \mid Z_i) = \frac{v_N(Z_i)}{B-1}.$$

By the law of total expectation, the mean squared error is

$$\mathbb{E}[(\widehat{\mu}_i - e_N^a(Z_i))^2] = \mathbb{E}[\text{Var}(\widehat{\mu}_i \mid Z_i)] = \frac{\mathbb{E}[v_N(Z)]}{B-1}.$$

Under Assumption 4.2, $\mathbb{E}[v_N(Z)]$ is finite. Since $B \asymp n$ in the fixed- N regime, the right-hand side converges to zero uniformly in i as $n \rightarrow \infty$. Thus, $\widehat{\mu}_i$ converges to $e_N^a(Z_i)$ in L^2 .

To establish the consistency of the sample second moment, we use the decomposition:

$$\frac{1}{n} \sum_{i=1}^n \hat{\mu}_i^2 = \frac{1}{n} \sum_{i=1}^n e_N^a(Z_i)^2 + \frac{1}{n} \sum_{i=1}^n (\hat{\mu}_i^2 - e_N^a(Z_i)^2).$$

The first term is an average of i.i.d. variables with finite mean (by Assumption 4.2). By the standard LLN, it converges in probability to $\mathbb{E}[e_N^a(Z)^2]$.

For the second term, we apply the Cauchy-Schwarz inequality to bound the L^1 norm of the difference:

$$\begin{aligned} \mathbb{E}|\hat{\mu}_i^2 - e_N^a(Z_i)^2| &= \mathbb{E}|(\hat{\mu}_i - e_N^a(Z_i))(\hat{\mu}_i + e_N^a(Z_i))| \\ &\leq \|\hat{\mu}_i - e_N^a(Z_i)\|_2 \cdot \|\hat{\mu}_i + e_N^a(Z_i)\|_2. \end{aligned}$$

We established above that $\|\hat{\mu}_i - e_N^a(Z_i)\|_2 \rightarrow 0$. The second factor is bounded because $\|\hat{\mu}_i\|_2 \rightarrow \|e_N^a(Z_i)\|_2 < \infty$. Thus, the expected absolute difference converges to zero. By Markov's inequality, the average of the differences converges to zero in probability.

Combining these results:

$$\frac{1}{n} \sum_{i=1}^n \hat{\mu}_i^2 \xrightarrow{p} \mathbb{E}[e_N^a(Z)^2], \quad \frac{1}{n} \sum_{i=1}^n \hat{\mu}_i \xrightarrow{p} \mathbb{E}[e_N^a(Z)] = e_N^a.$$

Consequently,

$$\hat{V}_{N,\text{test}} = \frac{n}{n-1} \left(\frac{1}{n} \sum_{i=1}^n \hat{\mu}_i^2 - \bar{\mu}^2 \right) \xrightarrow{p} \mathbb{E}[e_N^a(Z)^2] - (e_N^a)^2 = V_{N,\text{test}}.$$

□

Lemma A.3. Under Assumptions 4.1, 4.2, and $n = BN$,

$$\hat{C}_N \xrightarrow{p} C_N.$$

Proof of Lemma A.3. Define the errors $\Delta e_b = \hat{e}_b - e_{b,\text{train}}$ and $\Delta m_b = \hat{m}_b - \tilde{m}_b$. From

the proof of Lemma A.1, we established the uniform bound $\max_b |\Delta e_b| \xrightarrow{p} 0$. Consequently, the mean squared error for the training component vanishes:

$$\frac{1}{B} \sum_{b=1}^B (\Delta e_b)^2 \leq \left(\max_b |\Delta e_b| \right)^2 \xrightarrow{p} 0.$$

For the test component $\widehat{m}_b$, we rely on the L^2 convergence from Lemma A.2. Note that $\widehat{m}_b - \widetilde{m}_b = \frac{1}{N} \sum_{i \in S_b} (\widehat{\mu}_i - e_N^a(Z_i))$. By Jensen's inequality and the fact that $|S_b| = N$ is fixed:

$$(\Delta m_b)^2 = \left(\frac{1}{N} \sum_{i \in S_b} (\widehat{\mu}_i - e_N^a(Z_i)) \right)^2 \leq \frac{1}{N} \sum_{i \in S_b} (\widehat{\mu}_i - e_N^a(Z_i))^2.$$

Averaging over all B blocks, we obtain

$$\frac{1}{B} \sum_{b=1}^B (\Delta m_b)^2 \leq \frac{1}{BN} \sum_{b=1}^B \sum_{i \in S_b} (\widehat{\mu}_i - e_N^a(Z_i))^2 = \frac{1}{n} \sum_{i=1}^n (\widehat{\mu}_i - e_N^a(Z_i))^2.$$

which, by Lemma A.2, converges to 0 in probability. Thus, $\frac{1}{B} \sum_{b=1}^B (\Delta m_b)^2 \xrightarrow{p} 0$.

Now, express the difference in sample covariances (ignoring the $B/(B-1)$ correction which approaches 1) as:

$$\begin{aligned} \widehat{C}_N - \widetilde{C}_N &\approx \frac{1}{B} \sum_{b=1}^B (\widehat{e}_b \widehat{m}_b - e_{b,\text{train}} \widetilde{m}_b) - \widetilde{\bar{e}} \widetilde{\bar{m}} + \bar{e}_{\text{train}} \bar{\widetilde{m}} \\ &= \frac{1}{B} \sum_{b=1}^B [\Delta e_b \widetilde{m}_b + e_{b,\text{train}} \Delta m_b + \Delta e_b \Delta m_b] - (\text{mean terms}). \end{aligned}$$

Using the Cauchy-Schwarz inequality, we bound the cross-terms. For example:

$$\left| \frac{1}{B} \sum_{b=1}^B e_{b,\text{train}} \Delta m_b \right| \leq \sqrt{\frac{1}{B} \sum_{b=1}^B e_{b,\text{train}}^2} \sqrt{\frac{1}{B} \sum_{b=1}^B (\Delta m_b)^2}.$$

The first factor is bounded (by the LLN applied to $e_{b,\text{train}}$), and the second factor converges to zero as shown above. Similar bounds apply to terms involving Δe_b (which converges

uniformly) and the product $\Delta e_b \Delta m_b$. The mean terms vanish similarly. Therefore, $|\widehat{C}_N - \widetilde{C}_N| \xrightarrow{p} 0$.

Finally, since $(e_{b,\text{train}}, \widetilde{m}_b)$ are i.i.d. vectors across blocks b , the standard LLN implies $\widetilde{C}_N \xrightarrow{p} \text{Cov}(e_{b,\text{train}}, \widetilde{m}_b) = C_N$. Thus, $\widehat{C}_N \xrightarrow{p} C_N$. $\square$

Proof of Theorem 4.3. By Lemmas A.1, A.2, and A.3,

$$N\widehat{V}_{N,\text{train}} \xrightarrow{p} NV_{N,\text{train}}, \quad \widehat{V}_{N,\text{test}} \xrightarrow{p} V_{N,\text{test}}, \quad N\widehat{C}_N \xrightarrow{p} NC_N.$$

From the variance decomposition (4.4), we obtain

$$\widehat{\sigma}_N^2 = N\widehat{V}_{N,\text{train}} + \widehat{V}_{N,\text{test}} + 2N\widehat{C}_N \xrightarrow{p} NV_{N,\text{train}} + V_{N,\text{test}} + 2NC_N = \sigma_N^2,$$

establishing part (1).

For part (2), Theorem 4.2 gives $\sqrt{n}(\widehat{e}^{\text{CV}} - e_N^a) \xrightarrow{d} \mathcal{N}(0, \sigma_N^2)$. Since $\widehat{\sigma}_N^2 \xrightarrow{p} \sigma_N^2 > 0$ by part (1), by the Slutsky's Theorem, we have

$$\frac{\sqrt{n}(\widehat{e}^{\text{CV}} - e_N^a)}{\widehat{\sigma}_N} \xrightarrow{d} \mathcal{N}(0, 1).$$

$\square$

B Block-Out CV under Fixed- B Asymptotics

The asymptotic results in Section 4.2 rely on the regime where the training size N is fixed while the number of blocks $B = n/N \rightarrow \infty$ (the fixed- N asymptotic regime). While valid for time-series rolling windows or massive datasets, in many econometric applications, the number of fold splits B is small and fixed (e.g., $B = 5$ or 10), while the sample size within each block N is large.

We define *fixed- B asymptotic regime* as the asymptotic framework where B is fixed and $N \rightarrow \infty$ (implying $n \rightarrow \infty$). Under this regime, the estimator $\widehat{e}^{\text{CV}}$ aggregates losses from

B highly stable predictors. The variance estimator $\hat{\sigma}_{N_k}^2$ based on the fixed- N asymptotic regime can perform poorly here because it relies on estimating the between-block covariance structure from very few (B) observations.

In this section, we derive a Central Limit Theorem (CLT) for the fixed- B asymptotic regime by exploiting the algorithmic stability of the learner as $N \rightarrow \infty$. This approach leverages recent results on cross-validation stability (Lei, 2025), specifically the behavior of risk gradients for deterministic centering.

Recall that the block-out cross-validation estimator is given by:

$$\hat{e}^{\text{CV}} = \frac{1}{n} \sum_{i=1}^n \tilde{\ell}_i, \quad \text{where} \quad \tilde{\ell}_i = \frac{1}{B-1} \sum_{b \neq b(i)} \ell(f_N^{(b)}(X_i), Y_i),$$

where $b(i)$ denotes the block containing observation i . Note that $\tilde{\ell}_i$ is the average loss of observation i evaluated by the $B-1$ models that did not include i in their training set. Let $e_N^a = \mathbb{E}[\ell(f_{D_N}(X), Y)]$ be the population prediction error for a training set of size N .

To establish asymptotic normality centered at e_N^a , we require conditions on the stability of the learning algorithm. Following Lei (2025), we distinguish between the stability of the loss function and the stability of the risk (expected loss). Define the conditional risk of a model trained on dataset D_N as $R(D_N) = \mathbb{E}_Z[\ell(f_{D_N}(X), Y)]$.

Assumption B.1 (Algorithm Stability). *The algorithm satisfies:*

$$\|\nabla_1 R\|_{L_2} := \sqrt{\mathbb{E} [(\nabla_1 R(D_N))^2]} = o(N^{-1}) \text{ as } N \rightarrow \infty, \quad (\text{B.1})$$

where $\nabla_1 R(D_N) = R(D_N) - R(D_N^{(1)})$ denotes the *Perturb-One Stability*, i.e., the change in conditional risk when one training point is replaced by an independent copy.¹⁴

Discussion on Stability. Assumption B.1 corresponds to the condition in Proposition 4.11.(1) of Lei (2025), where the variability induced by training data vanishes relative to the

¹⁴See Definition 2.3 in Lei (2025) for details.

evaluation noise. [Lei \(2025\)](#) also provides arguments¹⁵ that in machine learning contexts, the risk function usually enjoys much better stability than the loss function, making risk stability requirement in our Assumption [B.1](#) plausible.

Theorem B.1 (CLT under Fixed- B Asymptotics). *Under Assumption [B.1](#), as $N \rightarrow \infty$ with B fixed:*

$$\sqrt{n}(\hat{e}^{\text{CV}} - e_N^a) \xrightarrow{d} \mathcal{N}(0, \tau^2).$$

Proof of Theorem [B.1](#). We adapt Proposition 4.11.(1) in [Lei \(2025\)](#), along with the preceding analysis and lemmas [Lei \(2025\)](#), to our block-out CV design. We employ the Hájek projection technique for deterministic centering as formalized in [Lei \(2025\)](#). Let $Z_i = (X_i, Y_i)$ and define the statistic $L_n = \sum_{i=1}^n \tilde{\ell}_i$. The first-order Hájek projection is:

$$\hat{L}_n = \sum_{j=1}^n \mathbb{E}[L_n \mid Z_j] - (n-1)\mathbb{E}[L_n].$$

We decompose the conditional expectation $\mathbb{E}[L_n \mid Z_j]$ into two components:

1. Test component $j = i$: Z_j acts as the test point for $\tilde{\ell}_j$. Since $\tilde{\ell}_j$ uses models trained on blocks not containing j , Z_j is independent of the training sets. Thus:

$$\mathbb{E}[\tilde{\ell}_j \mid Z_j] = l_N(Z_j) \equiv \mathbb{E}[\ell(f_{D_N}(X), Y) \mid Z = Z_j].$$

2. Training component $j \neq i$: Z_j acts as a training point for observation i , which implies j is in a training block used for $\tilde{\ell}_i$. Hence, it affects the model prediction. In our block design, Z_j is in the training set for all $B-1$ blocks that do not contain j . The sum of these influences corresponds to the function $g_N(Z_j)$ defined in [Lei \(2025, Eq. 40\)](#).

The variance of the projection is $\text{Var}(\hat{L}_n) \approx n\text{Var}(l_N(Z_1) + g_N(Z_1))$, where g_N captures the sum of training influences. By the Efron-Stein inequality, the variance of the training

¹⁵See the arguments after Theorem 4.7 (pp.54-55) in [Lei \(2025\)](#).

influence scales with the stability of the risk:

$$\text{Var}(g_N(Z_1)) \leq C (N \|\nabla_1 R(D_N)\|_{L_2})^2.$$

By Assumption [B.1](#) (Risk Stability), we have $N \|\nabla_1 R\|_{L_2} \rightarrow 0$. Consequently, the total asymptotic variance is dominated by the test term:

$$\frac{1}{n} \text{Var}(\widehat{L}_n) \rightarrow \text{Var}(l_N(Z_1)) = \tau^2.$$

Standard results on U-statistics imply the remainder $L_n - \widehat{L}_n$ is $o_p(\sqrt{n})$. Applying the standard CLT to the sum of i.i.d. variables $l_N(Z_i)$ yields the result. $\square$

In the fixed- B asymptotic regime, since the training variance component vanishes, the asymptotic variance τ^2 is simply the variance of the prediction error on a fresh test point. This can be estimated by the sample variance of the aggregated losses:

$$\widehat{\tau}_B^2 = \frac{1}{n-1} \sum_{i=1}^n \left(\widetilde{\ell}_i - \widehat{e}^{\text{cv}} \right)^2. \quad (\text{B.2})$$

Corollary B.1 (Studentized CLT). *Under the conditions of Theorem [B.1](#), $\widehat{\tau}_B^2 \xrightarrow{p} \tau^2$, and*

$$\frac{\sqrt{n}(\widehat{e}^{\text{cv}} - e_N^a)}{\widehat{\tau}_B} \xrightarrow{d} \mathcal{N}(0, 1).$$

C Details of ML Implementation

Outcomes and feature construction. We apply the following pre-processing to all ML methods we use. For earnings as the target outcome, we winsorize it prior to analysis at the 1st and 99th percentiles to reduce sensitivity to extreme outliers.^{[16](#)} The continuous predictors are standardized, and the categorical predictors are handled as follows: (i) rare categories are collapsed into a single “other” level when their frequency falls below a minimum-count

¹⁶We alternatively train the ML models on $\log(1 + Y)$ while calculating the error on the original level scale.

threshold (e.g., 50), and (ii) the resulting categorical variables are converted to indicators (with one category dropped as a reference group). In addition, we augment X with other outcomes as predictors (continuous outcomes enter as standardized numeric features; categorical outcomes enter via indicators), excluding the current target and respecting the earnings/wage exclusion noted above.

Learning-curve design and block-out CV. Fix a training size N . After randomly permuting the analysis sample once for reproducibility, we partition the data into $B = \lfloor n/N \rfloor$ disjoint blocks of size N . For each block $b = 1, \dots, B$, we fit the ML model using only the observations in block b as the training set, generate predictions for all observations, and evaluate loss on the complement of block b . Averaging the out-of-block losses across b yields the block-out CV estimate of predictive risk at training size N . Repeating this procedure over a grid of N values produces the error curve.

Model classes and per- N hyperparameter tuning. For each N , we tune hyperparameters on a tuning subset of size N drawn at random from the analysis sample, and then apply the selected hyperparameters in the block-out CV loop.¹⁷ For regression (continuous Y), we consider lasso and random forest. The lasso penalty parameter is selected by `LassoCV` using 5-fold CV. The random forest uses 300 trees and is tuned by 3-fold CV over a small grid of structural parameters (maximum depth in $\{\text{None}, 10, 20\}$ and minimum leaf size in $\{1, 5\}$), with performance scored in terms of RMSE. For classification (discrete Y), we consider ℓ_1 -penalized logistic regression (logit lasso) and random forest classification. Logistic regression is fit using the `saga` solver and tuned over a grid of inverse-regularization parameters $C \in \{0.01, 0.1, 1, 10\}$ using stratified CV; random forest classification uses 300 trees and is tuned over the same depth/leaf grid as in regression using stratified CV.

For neural networks used in the preliminary analysis in Section D, we fit a shallow fully-

¹⁷This approach is useful in stabilizing the error curve across N , although to faithfully follow the definition of e_N^a , we may want to tune hyperparameters for each training block in the block-out CV loop. Still, the current approach may be acceptable as the tuning is low-dimensional and held fixed across CV folds for a given N .

connected feedforward network with two ReLU hidden layers (64 and 32 units) and an output layer that is linear for regression and softmax for classification (with the number of outputs set to the number of classes observed in the training data). Networks are trained with the Adam optimizer using MSE for regression (100 epochs) and categorical cross-entropy for classification (10 epochs, with one-hot encoded labels).

Small- N safeguards for classification. For very small N , stratified CV and/or block-level training sets can become degenerate due to severe class imbalance. We therefore implement explicit safeguards: (i) the number of folds in stratified CV is chosen adaptively and never exceeds the minimum class count (capped at 3 folds), and (ii) if a block-level training set contains only one class, we default to a majority-class prediction rule for that block.

Uncertainty quantification. We report uncertainty measures derived from the variance estimator introduced in Section 4: the cross-block variability of observation-level out-of-block losses yields an estimated variance component $\hat{\sigma}_k^2$, which is converted into standard errors of the CV risk via $\sqrt{\hat{\sigma}_k^2/n}$. For $\hat{\sigma}_k^2$, as mentioned in the main text, we use a hybrid approach based on the fixed- N asymptotic theory (Theorem 4.2) when $N \leq 400$ and the fixed- B asymptotic theory (Corollary B.1 in Section B) otherwise. For RMSE, we apply a delta-method transformation from MSE to RMSE to obtain confidence intervals.

D Additional Results for the Applications

D.1 Other Prediction Methods

Here we compare LLM performance to ML performance at a fixed training size for the ML models: 800 observations. In addition to ChatGPT-4, we report performance for Meta’s Llama 3 models (Llama3-8B cutoff: March 2023; Llama3-70B cutoff: December 2023). We find that predictive performance is broadly similar across LLMs and, separately, across machine-learning algorithms, while varying systematically across outcomes. This pattern

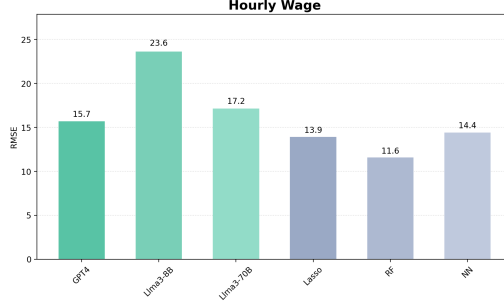


Figure 7: Root Mean Squared Error (lower is better) ($N = 800$)

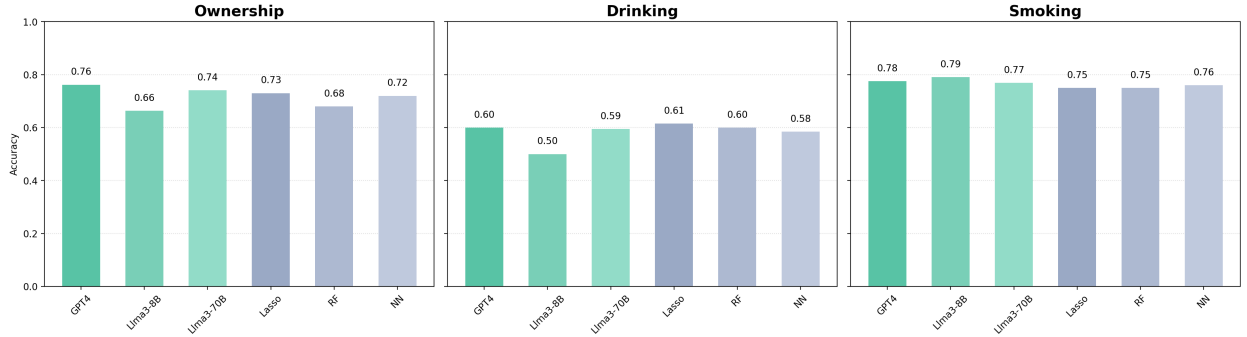


Figure 8: Accuracy (higher is better) ($N = 800$)

suggests that our main findings are not driven by the choice of a particular LLM or comparator algorithm.

D.2 Other Error Metrics

In the main text, prediction error for discrete outcomes is measured using the 0-1 loss function. Misclassification rate can be mechanically low for outcomes that are very skewed, which might mask differences between the LLMs and ML algorithms that would otherwise emerge. As a robustness check, we thus compare performance in terms of *balanced accuracy* and *F1*, defined as follows.

For discrete (binary) outcomes, let $Y \in \{0, 1\}$ denote the true label and $f(X) \in \{0, 1\}$ the predicted label. Writing TP, TN, FP, and FN for the usual entries of the confusion matrix,

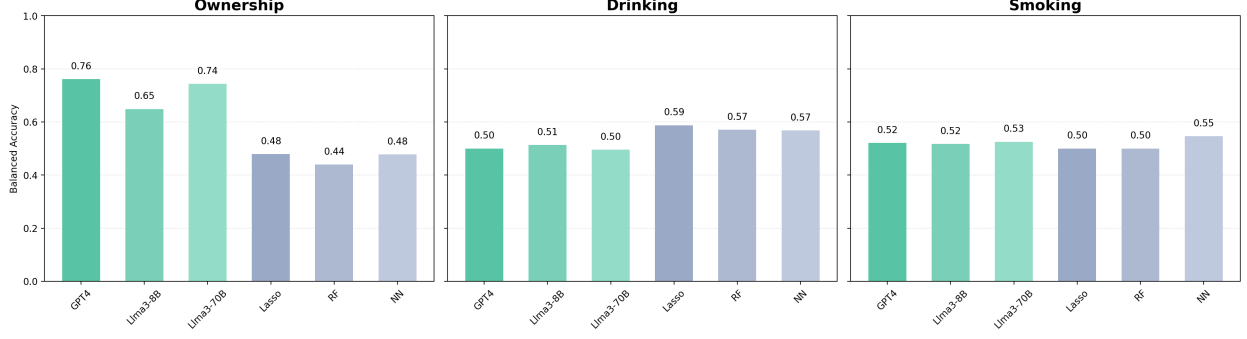


Figure 9: Balanced Accuracy (higher is better) ($N = 800$)

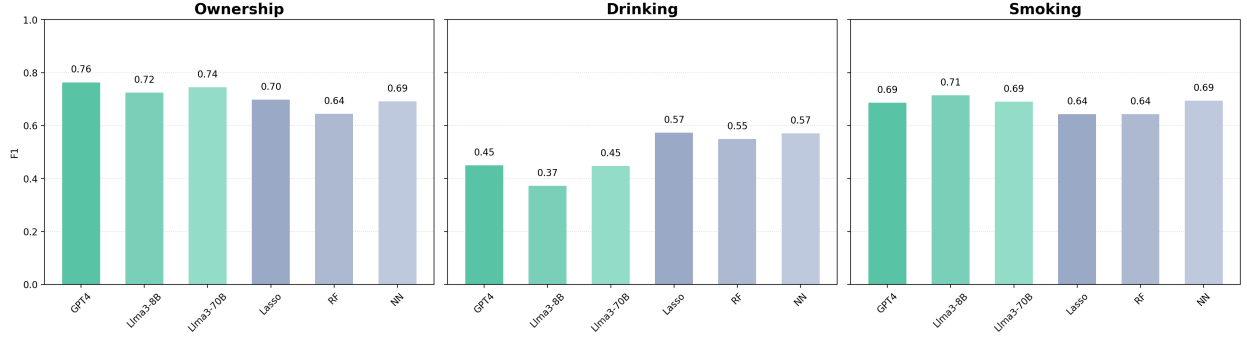


Figure 10: F1 (higher is better) ($N = 800$)

the *balanced accuracy* is defined as the average of the class-wise recall,

$$\text{BA}(f) = \frac{1}{2} \left(\frac{\text{TP}}{\text{TP} + \text{FN}} + \frac{\text{TN}}{\text{TN} + \text{FP}} \right).$$

The *F1 score* is defined as the harmonic mean of precision and recall,

$$\text{F1}(f) = \frac{2 \text{TP}}{2 \text{TP} + \text{FP} + \text{FN}} = \frac{2 \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}},$$

and emphasizes performance on the positive class by jointly penalizing false positives and false negatives. Figures 9 and 10 replicate our classification results in Appendix D.1 using these error metrics. We find that the performances across domains are qualitatively similar to those in Figure 8.