

Testing Capacity-Constrained Learning*

Andrew Caplin¹ Daniel Martin^{2,*} Philip Marx³
Anastasiia Morozova² Leshan Xu²

¹Department of Economics, New York University, New York, NY, USA.

²Department of Economics, University of California, Santa Barbara, Santa Barbara, CA, USA.

³Department of Economics, Louisiana State University, Baton Rouge, LA, USA.

*Corresponding author: danielmartin@ucsb.edu

ORCID: Caplin [0000-0002-1169-9027](#), Martin [0000-0001-6483-3923](#)

September 13, 2026

Abstract

We introduce a general test of capacity-constrained learning models. Learning has capacity constraints when the set of possible ways to learn is exogenously fixed, as in the widely used fixed-capacity versions of rational inattention (Sims 2003) and efficient coding (Woodford 2012). With such models, changes in incentives do not alter the extent of attention, only how individuals decide to allocate their scarce attention. We show that choice data are consistent with capacity-constrained learning if and only if they satisfy a *No Improving (Action or Attention) Switches (NIS)* condition. Based on existing experiments in which the incentives for being correct are varied, we find strong evidence that participants fail NIS for a wide range of standard perceptual tasks: identifying the proportion of ball colors, recognizing shapes, and counting the number of balls. We further show that violations of NIS occur systematically in response to higher incentives, suggesting that incentives often expand attention beyond what capacity-constrained models allow. However, we find that this is not true for all existing perceptual tasks in the literature, which offers insights into settings where we do or do not expect incentives to impact the extent of attention.

*We thank the Editor, three anonymous referees, Federico Echenique, Mark Dean, and Gerelt Tserenjimid for helpful comments. We also thank the Sloan Foundation for support under the “Cognitive Economics at Work” grant.

1 Introduction

Whether incentives alter attention has become a key question for a growing literature on the role of cognition in economics (e.g., Bronchetti et al. 2023).¹ An important aspect of this question is whether perceptual limits (“capacities” or “constraints”) are exogenously fixed or if perception can be improved when desired. For example, in the rational inattention literature, Sims (2003) assumes the bounds on attention are fixed, while Matejka and McKay (2015) assume the bounds are elastic and can be loosened at a cost. In the efficient coding literature, Woodford (2012) provides versions of optimal encoding with both fixed bounds and elastic bounds.² Whether or not perceptual capacity is fixed or variable has substantial implications for the economic impact of bounded cognition. Most importantly, if perceptual constraints are exogenously fixed, then changes in incentives will not alter the extent of attention, only how individuals decide to allocate their scarce attention.

The question of whether incentives impact the extent of attention also relates to an ongoing debate about whether incentivization matters in experiments. Various reasons have been given for incentivizing experiments (e.g., Plott 1986, Smith 1991), but most experiments on perception and attention in the psychology literature do not incentivize performance; implicitly treating incentives as irrelevant to cognition. However, this assumption has been challenged by papers showing the impact of rewards on visual perception in the psychology literature on psychometrics (e.g., Pessoa and Engelmann 2010).

We offer a method for answering the question of whether perceptual capacity is exogenously fixed, and hence whether incentives impact the extent of attention, by providing a sharp test of capacity-constrained learning models. In this class of models, a decision-maker can learn about the true state only through an exogenously fixed set of feasible learning strategies. Incentive changes can affect which feasible strategy they choose, but cannot expand the underlying set of learning strategies that is available to be used. Because feasibility is imposed as a fixed constraint on learning rather than a choice variable, this model class includes settings with exogenous perceptual limits, such as fixed-capacity rational inattention (Sims 2003) and efficient coding (Woodford 2012). Our contribution is to provide a test that directly addresses the questions above: because incentives can only expand attention if perceptual capacity is elastic, violations of our test can be interpreted as evidence that incentives matter for attention and perception.

¹Motivated by a long literature on subjective perception in cognitive science (e.g., Weber 1834), cognitive economics blends economics, psychology, and neuroscience methods to better understand the limits of human perception (Caplin 2025).

²Fixed-capacity efficient coding has been applied to risky choice (Khaw, Li, and Woodford 2021, Frydman and Jin 2022), investing (Charles, Frydman, and Kilic 2024), probability weighting (Frydman and Jin 2023), belief updating (Ba, Bohren, and Imas 2022), and play in games (Frydman and Nunnari 2021).

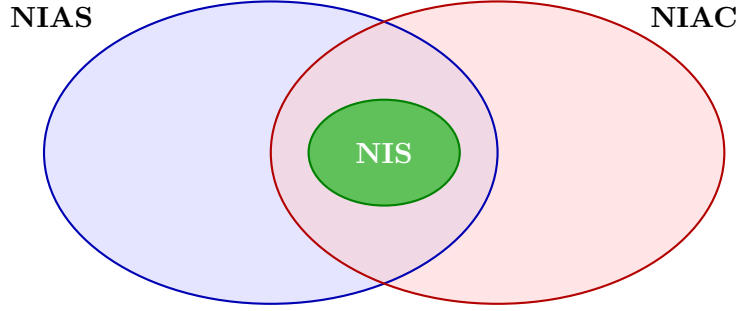


Figure 1: Relationship between NIS and existing conditions (NIAS and NIAC).

We show that a single condition, *No Improving (Action or Attention) Switches (NIS)*, is both necessary and sufficient for choices to be consistent with models of capacity-constrained learning. NIS sharpens the *No Improving Action Switches (NIAS)* condition, which requires optimal actions given information (Caplin and Martin 2015), and the *No Improving Attention Cycles (NIAC)* condition, which rules out profitable reallocations of attention across problems (Caplin and Dean 2015). In words, NIS states that *wholesale* switches of attention or actions, both within and across decision problems, should not improve utility. Incentives thus enter our framework not as an auxiliary consideration but as the central empirical tool: when NIS fails, it is because incentives induce systematic reallocations of attention that would not occur if perceptual capacity were fixed.

We apply our test of capacity-constrained learning to data from three existing experiments that vary the payoffs to being correct in visual perception tasks. We first look at experiments 1.2 and 2.2 of Dean and Neligh (2023), which vary the reward to correctly guessing the state in two tasks: the number of colored balls in a visual display (1.2) and the number of correct equations (2.2). Using the aggregated choices of participants, we find that NIAS and NIAC are satisfied for both experiments. However, NIS fails in both experiments, and the failure is statistically significant in Experiment 1.2 ($p < 0.01$).³ Importantly, NIS fails in a systematic way in both experiments: switching attention in the direction of higher incentives (switching to the attention used when incentives were higher) always leads to improvements in expected utility.

Next, we reexamine the data from the experiment of Caplin, Csaba, Leahy, and Nov (2020), which varies the payoff for correctly guessing which shape appears the most frequently. We find that NIAS is satisfied for all four difficulty levels of the task, but that NIS is not satisfied for any of the difficulty levels (statistically significant for all four).

³One possible reason that the violation of NIS in experiment 2.2 is not statistically significant is that performance is already high (over 80%) at very low incentives, which leaves little room for improvement with incentives. Another possible reason is that there are only two incentive levels, so there are few inequalities to check.

Finally, we re-analyze the data from the experiment of Dewan and Neligh (2020), which varies the payoffs using two possible prize sizes (\$10 and \$20). For the first task, correctly guessing the number of balls in a display, we again find that NIAS and NIAC are satisfied for both prize sizes, but NIS fails for both ($p < 0.01$ for each prize size). As in the other experiments, we find that for a given prize size, switches of attention in the direction of higher incentives always lead to improvements in expected utility.

However, for the second task of Dewan and Neligh (2020), which requires participants to judge the degree of an angle, we do not find strong evidence that NIS fails for either of the prize sizes ($p = 0.22$ and $p = 0.86$). In addition, unlike the other experiments, switches of attention in the direction of higher incentives (for a given prize size) often do not lead to improvements in expected utility.

These results suggest that certain task factors might dictate the suitability of capacity-constrained learning. Specifically, this model class might be more appropriate for tasks that are not responsive to additional effort.⁴ This can occur because (1) even at low incentives, an individual chooses to exert enough effort to reach their perceptual limit, and (2) additional effort would not meaningfully improve perception beyond this limit. This combination of forces causes individuals to reach the same level of “irreducible uncertainty” at all levels of incentives. Another possibility is that uncertainty could actually be reduced further, but individuals choose not to spend the additional effort required to reduce it further at standard incentive levels. Determining what class of perceptual and judgment tasks is insensitive to incentive changes is an open area for experimental work and will be central to assessing when models of capacity-constrained learning are suitable for environments with changing incentives.

Our paper contributes to a growing literature on testing economic models of cognition (e.g., Caplin, Csaba, Leahy, and Nov 2020; de Clippel and Rozen 2020; Dewan and Neligh 2020; Dean and Neligh 2023; Almog and Martin 2024; Almog, Gauriot, Page, and Martin 2024). This literature includes tests of models of costly attention and rational inattention, in which cognition is treated as an endogenous, resource-constrained input into choice (see Shenhav et al. 2017; Gabaix 2019; Maćkowiak, Matějka, and Wiederholt 2023 for reviews). It also includes experimental and psychometric approaches that use controlled tasks to identify and test restrictions implied by attention-based models (e.g., Dean and Neligh 2023; Caplin, Csaba, Leahy, and Nov 2020; Ambuehl, Ockenfels, and Stewart 2025). Our primary contribution to this literature is to examine the consistency of human choices with capacity-constrained learning models.

Related work emphasizes that economic behavior can also depend on the choice of at-

⁴The relationship between incentives and the effort required to complete judgment tasks is explored in Camerer and Hogarth (1999).

tention environments and attention-improving tools, and develops theory-driven tests using willingness-to-pay and incentive variation in such settings (e.g., Bronchetti et al. 2023; Gabaix 2014; De Oliveira, Denti, Mihm, and Ozbek 2017). Our primary contribution within this agenda is to provide and implement a sharp revealed-preference test for whether observed behavior is consistent with capacity-constrained learning, thereby isolating when incentive changes can be rationalized by reallocating fixed attention versus requiring an expansion of attentional capacity. This positions our test as not only the first general test of capacity-constrained learning, but also as a novel way of adjudicating the long-standing debate about whether incentives alter the extent of human attention.

Our test also has implications for non-human learners. Caplin, Martin, and Marx (2025) apply our test to the predictions of a state-of-the-art machine learning algorithm that is trained to identify pneumonia from chest x-rays. In the context of machine learning, capacity-constrained learning aligns with standard notions of how a machine learns, as this model class can be interpreted as the machine choosing among a feasible set of mathematical operations to best match the incentives provided by its loss function. Echoing the results in this paper, Caplin, Martin, and Marx (2025) find that a state-of-the-art machine learning algorithm violates NIS but not NIAS or NIAC in an experiment in which the loss function used to train the algorithm is varied.

The rest of the paper is as follows. Section 2 introduces NIS, our test of capacity-constrained learning. Section 3 proposes approaches to measuring the extent of NIS violations and the statistical testing of NIS. Section 4 provides the results of testing NIS using the data from existing experiments. Section 5 concludes.

2 Our Test of Capacity-Constrained Learning

In a capacity-constrained learning model, the decision-maker faces an objectively defined but uncertain state and must choose an action whose payoff depends on that state. Before acting, the decision-maker may acquire information, but only via a capacity constraint that is represented as a fixed feasibility set: a collection of learning strategies the person is capable of using. A learning strategy can be understood as a rule that determines what noisy evidence the person can generate about the state and therefore what posterior beliefs they can end up holding.

Capacity-constrained learning imposes three key requirements. First, the set of learning strategies available to the decision-maker is the same across decision problems and does not change with the payoff scale. Second, in each decision problem, the decision-maker selects whichever feasible learning strategy yields the highest expected payoff for that problem. Finally, once a signal is realized, the decision-maker updates beliefs in a Bayes-consistent

way and chooses an action that is optimal given those beliefs.

Consider a researcher who wants to determine if a decision-maker chose in line with capacity-constrained learning, but does not observe the information actually obtained by the decision-maker, their beliefs, or their attention. The following sections introduce a test the researcher can use to answer this question. A more precise definition of capacity-constrained learning and our representation result are provided in the Online Appendix.

2.1 Preliminaries

This section introduces the objects a researcher needs in order to apply our test using observed choice data. The researcher observes a finite set of experimental decision problems, each defined by a known set of possible states, a set of available actions, and a payoff rule that rewards actions as a function of the realized state. In particular, there is a finite set of possible states of the world $\omega \in \Omega$ distributed according to an objective prior,⁵ a finite global set of actions $a \in \mathcal{A}$ with $|\mathcal{A}| \geq 2$, and a finite global set $x \in X$ of prize specifications, with realizations $x(a, \omega)$ that depend on the realized state and action in a way known by the researcher. Thus, a decision problem $i \in D$ consists of an action set $A_i \subseteq \mathcal{A}$ and a prize specification $x_i \in X$.⁶

For example, Experiment 1.2 of Dean and Neligh (2023) (DN23 hereafter) has two equally likely states: out of the 100 balls presented on the screen, there can either be 49 red balls and 51 blue balls or 51 red balls and 49 blue balls (labeled states ω_1 and ω_2 , respectively). The global set of actions $\mathcal{A} = \{a_1, a_2\}$ is common across decision problems, and individuals are incentivized to match actions and states (choosing a_1 in state ω_1 and a_2 in state ω_2) according to a varying prize specification such that a correct guess leads to either 5, 40, 70, or 95 *probability points* (the probability of receiving a fixed cash amount) while an incorrect guess yields 0 probability points. This structure is typical of many experimental designs: the states and actions are fixed across treatments, while incentives vary. To summarize, the set of decision problems $i \in \{1, 2, 3, 4\}$ consists of common states and actions yielding varying probability point prizes, which are known to correspond to actions and states as follows:

⁵That is, we assume the distribution of states is known to the researcher and the subjects, as is standard in perceptual and counting tasks. The tests in this paper can be readily extended to settings where the objective prior varies, but such variation is not necessary for testing capacity-constrained learning.

⁶Our specification is similar to that of Dean and Neligh (2023), who vary both action sets and prizes in their experiments. To reconcile with Caplin and Dean (2015), note that a set of actions a and prizes x can always be reduced to a set of actions by redefining $a \leftarrow (a, x)$.

$$\Omega = \{\omega_1, \omega_2\}, \quad A_i = \{a_1, a_2\}, \quad x_i(a, \omega) = \begin{matrix} & \omega_1 & \omega_2 \\ a_1 & \begin{pmatrix} p_i & 0 \end{pmatrix} \\ a_2 & \begin{pmatrix} 0 & p_i \end{pmatrix} \end{matrix} \quad \text{where } p_i = \begin{cases} 5 & \text{for } i = 1, \\ 40 & \text{for } i = 2, \\ 70 & \text{for } i = 3, \\ 95 & \text{for } i = 4. \end{cases}$$

The decision-maker (DM) is assumed to have a utility function u over prizes. In all the experiments we examine, prizes take the form of “probability points” which directly determine the chances of receiving a fixed monetary amount. Under the expected utility framework, the utility of a prize that offers a given number of probability points is linear in the number of probability points, because these points scale the probability of a fixed outcome rather than the outcome itself. If we additionally assume that receiving cash yields more utility than not receiving it, then the utility of the prize is also strictly increasing in the number of probability points. Without loss of generality, we can normalize the utility of receiving the cash prize to 100 and not receiving the cash to 0, in which case the utility of the probability point prize is point identified: the utility of p probability points is exactly p .

The data required for our test are the joint frequencies with which subjects choose each action in each realized state, separately for each decision problem. We summarize these frequencies as state-dependent stochastic choice (SDSC) matrices $P_i(a, \omega)$. These include, for each decision problem $i \in D$, the joint distribution of actions and states $P_i(a, \omega)$ for all $a \in A_i$ and $\omega \in \Omega$. We denote as P the collection of P_i for every decision problem $i \in D$.

For example, in Experiment 1.2 of DN23, the set of aggregate (across-subject) joint distributions P is:

$$\begin{aligned} P_1 &= \begin{matrix} & \omega_1 & \omega_2 \\ a_1 & \begin{pmatrix} 0.37 & 0.20 \end{pmatrix} \\ a_2 & \begin{pmatrix} 0.13 & 0.30 \end{pmatrix} \end{matrix} & P_2 &= \begin{matrix} & \omega_1 & \omega_2 \\ a_1 & \begin{pmatrix} 0.38 & 0.17 \end{pmatrix} \\ a_2 & \begin{pmatrix} 0.12 & 0.33 \end{pmatrix} \end{matrix} \\ \\ P_3 &= \begin{matrix} & \omega_1 & \omega_2 \\ a_1 & \begin{pmatrix} 0.39 & 0.17 \end{pmatrix} \\ a_2 & \begin{pmatrix} 0.11 & 0.33 \end{pmatrix} \end{matrix} & P_4 &= \begin{matrix} & \omega_1 & \omega_2 \\ a_1 & \begin{pmatrix} 0.39 & 0.14 \end{pmatrix} \\ a_2 & \begin{pmatrix} 0.11 & 0.36 \end{pmatrix} \end{matrix} \end{aligned}$$

In summary, to apply our test, an experimentalist needs only (i) observed choices in each state and (ii) the payoff rule linking actions and states.

2.2 NIS: Capacity-Constrained Learning

The condition for a collection of SDSC data P to be consistent with capacity-constrained learning is given below. Importantly, our test relies only on observed choices and known payoff rules.

Condition 1 (No Improving (Action or Attention) Switches (**NIS**)). *Utility function u satisfies NIS for P if and only if for any set of decision problems $i, j \in D$,*

$$\sum_{a \in A_i} \sum_{\omega \in \Omega} P_i(a, \omega) u(x_i(a, \omega)) \geq \sum_{a \in A_j} \max_{\hat{a} \in A_i} \sum_{\omega \in \Omega} P_j(a, \omega) u(x_i(\hat{a}, \omega)) \quad (1)$$

In words, this condition states that *wholesale* switches of attention and actions should not improve utility according to the baseline prize specification. Intuitively, NIS states that reusing any feasible pattern of learning observed in one decision problem should not improve expected utility in another with the same payoff rule. NIS involves both checks between decision problems (when $i \neq j$), where both actions and choice probabilities can change, and checks “within” each decision problem (when $i = j$), where only actions can change. Thus, NIS involves checking $N^2 + N$ inequalities. The Online Appendix describes how NIS relates to two existing tests of Bayesian learning (NIAS and NIAC).

2.2.1 Example: Dean and Neligh (2023)

To illustrate NIS we turn to the decision problems, prizes, and aggregate SDSC data from Experiment 1.2 of DN23, which are given in Section 2.1. Restricting for simplicity to just the first two decision problems (1, 2) in which the number of probability points was 5 and 40 respectively, the data P is:

$$P_1 = \begin{matrix} & \omega_1 & \omega_2 \\ \begin{matrix} a_1 \\ a_2 \end{matrix} & \begin{pmatrix} 0.37 & 0.20 \\ 0.13 & 0.30 \end{pmatrix} \end{matrix} \quad P_2 = \begin{matrix} & \omega_1 & \omega_2 \\ \begin{matrix} a_1 \\ a_2 \end{matrix} & \begin{pmatrix} 0.38 & 0.17 \\ 0.12 & 0.33 \end{pmatrix} \end{matrix}$$

The NIS condition makes four comparisons: switches of actions within (decision) problem 1, switches of actions within problem 2, switches of choice probabilities and actions from problem 1 to 2, and switches of choice probabilities and actions from problem 2 to 1. Concretely, the test checks whether outcomes could be improved by either reassigning actions within a given problem, or by reusing the learning reflected in choices under one problem and applying it to the other problem, in both directions. In this case, switches of choice probabilities from problem 1 to 2 generate a violation of NIS, which we will show below. Practically, this means that the pattern of choices in problem 2 (the problem that is identical

to problem 1 except for in terms of prize) reveals a feasible way of learning and applying that same learning to problem 1 would improve expected utility. Under fixed feasibility, such an improvement should not be possible, so NIS is violated.

As a first step, we compute the expected utility implied by the observed choice frequencies in decision problem 1, treating these frequencies as revealing how the subject learned about the state. The value of the data in decision problem 1 is

$$\begin{aligned} & \underbrace{P_1(a_1, \omega_1)u(x_1(a_1, \omega_1))}_{0.37 \times 5} + \underbrace{P_1(a_1, \omega_2)u(x_1(a_1, \omega_2))}_{0.20 \times 0} \\ & + \underbrace{P_1(a_2, \omega_1)u(x_1(a_2, \omega_1))}_{0.13 \times 0} + \underbrace{P_1(a_2, \omega_2)u(x_1(a_2, \omega_2))}_{0.30 \times 5} \\ & = 3.34. \end{aligned}$$

We next ask whether the learning revealed in decision problem 2 could have been reused in decision problem 1. To do this, we evaluate the choice frequencies from problem 2 using the payoffs and optimal actions from problem 1, which results in

$$\begin{aligned} & \underbrace{P_2(a_1, \omega_1)u(x_1(a_1, \omega_1))}_{0.38 \times 5} + \underbrace{P_2(a_1, \omega_2)u(x_1(a_1, \omega_2))}_{0.17 \times 0} \\ & + \underbrace{P_2(a_2, \omega_1)u(x_1(a_2, \omega_1))}_{0.12 \times 0} + \underbrace{P_2(a_2, \omega_2)u(x_1(a_2, \omega_2))}_{0.33 \times 5} \\ & = 3.55. \end{aligned}$$

Because the optimal actions remain the same and choice probabilities on the diagonal increase, switching results in a net gain of utility (0.21). Thus, there would have been an improving switch in utility by using the implied learning in decision problem 2 in decision problem 1. As a result, NIS is not satisfied for this data. Since the existing experiments vary payoffs, NIS effectively tests whether attention expands in response to incentives, and NIS failing effectively means the failure of the capacity-constrained learning model.

3 Taking NIS to Data

For the experiments that we study, NIS makes the stringent prediction that the percentage correct has to be identical across treatments. When evaluating whether data are consistent with this prediction, two distinct questions arise. The first is a question of *magnitude*: if a violation occurs, how large is it? The second is a question of *statistical inference*: given that we observe only a finite sample, can we confidently conclude that the population violates NIS? A violation can be large but statistically insignificant (if the sample is small or noisy), or small but statistically significant (if the sample is large enough to detect minor deviations). We address both questions in turn. Section 3.1 proposes two Improvability Indices (IDI

and IEI) that quantify the economic magnitude of NIS violations. Section 3.2 develops a statistical test whose p -value quantifies our confidence that observed violations are not due to sampling noise.

3.1 Improvability Indices

A known challenge in testing axioms with choice data is that axioms are either satisfied or not, which is fairly stark. Instead, when they fail it might be useful to know the extent of the failure.

In the same spirit as a literature that measures the extent of violations of rationality for deterministic models of decision-making (Afriat 1973; Varian et al. 1991; Echenique, Lee, and Shum 2011; Apesteguia and Ballester 2015; Dean and Martin 2016), we propose two measures for the extent of violations of rationality for stochastic models of decision-making.⁷ The measures are stated here for NIS, but could be extended to cover other conditions, such as NIAS and NIAC.

The first measure we propose is the *Improvability Difference Index (IDI)*. This index indicates the largest difference between the actual expected utility obtained in a decision problem and what could be achieved by switching actions and attention. To aid comparability across decision problems, it is normalized by the maximum achievable expected utility within a decision problem, so takes a value between 0 and 1.⁸

Technically, IDI is the larger of 0 and the (normalized) difference between the expected utility under the chosen attention in a decision problem (the left-hand side of the NIS condition 1) and the expected utility from using the attention chosen in a different decision problem (the right-hand side of the NIS condition 1):

$$\max_{i,j \in D} \frac{\sum_{a \in A_j} \max_{\hat{a} \in A_i} \sum_{\omega \in \Omega} P_j(a, \omega) u(x_i(\hat{a}, \omega)) - \sum_{a \in A_i} \sum_{\omega \in \Omega} P_i(a, \omega) u(x_i(a, \omega))}{\sum_{\omega \in \Omega} \max_{\hat{a} \in A_i} \mu(\omega) u(x_i(\hat{a}, \omega))} \quad (2)$$

NIS passes if and only if IDI is 0, and higher values of IDI indicate that NIS is violated by more of the maximum expected utility.

Second, we propose the *Improvability Efficiency Index (IEI)*, which is the fraction of expected utility that needs to be shaved from the right-hand side of the NIS condition (the expected utility from using the attention chosen in a different decision problem) to make it

⁷Another approach to measuring the extent of violations of rationality for stochastic choice data is proposed by Ok and Tserenjigmid (2023), who measure deviations from (possibly incomplete) preference maximization.

⁸If we interpret the expected utility gain from switching to the learning from another decision problem as something that a third party could extract, then IDI is in a similar spirit to the Money Pump Index (Echenique, Lee, and Shum 2011), especially when adapted to the cyclical switches of NIAC.

smaller than the left-hand side of the NIS condition (expected utility obtained in a decision problem). Technically, we calculate IEI by finding the largest value of $\epsilon \in [0, 1]$ such that the following inequality holds for all sets of decision problems $i, j \in D$:

$$\sum_{a \in A_i} \sum_{\omega \in \Omega} P_i(a, \omega) u(x_i(a, \omega)) \geq \epsilon \sum_{a \in A_j} \max_{\hat{a} \in A_i} \sum_{\omega \in \Omega} P_j(a, \omega) u(x_i(\hat{a}, \omega)) \quad (3)$$

Clearly, NIS passes if and only if IEI is 1, and lower values of IEI indicate that a higher fraction of expected utility has to be shaved off of the RHS to make NIS pass. This can be interpreted as the utility that has to be shaved from every state on the RHS to make NIS pass.⁹

The closest measure to IEI is the Critical Cost Efficiency Index (CCEI) (Afriat 1973), which measures the budget set reduction necessary for the Generalized Axiom of Revealed Preference (GARP) to be satisfied for deterministic choice. Like IDI and IEI, the CCEI measures the largest violation in a set of choice data.

3.2 Statistical Testing

NIS is a property of the population-level joint distribution P of states and choices, which is not directly observable in experimental data. Instead, experiments generate a finite sample drawn from P , inducing an empirical distribution $\hat{P}$. This gap creates the need for statistical testing.

A number of methods have been used to evaluate the statistical significance of NIAS and NIAC. For example, in Dewan and Neligh (2020), NIAS is tested using bootstrapping. If no more than 5% of samples for each action for a given subject fail, then that subject fails to reject NIAS. In Dewan and Neligh (2020), NIAC is tested using regression. They run a linear weighted least squares regression of correctness on incentive level and then perform a one-sided t -test on the coefficient of the incentive level.

This paper evaluates significance using the Wald test. This test calculates the distance between our estimated parameter values $P(a, \omega)$ across actions and states, and the set of values that satisfy the null hypothesis of NIS, and then compares it with a χ^2 distribution to determine the p -value. A key step in the Wald test is calculating the covariance matrix of the difference between the right-hand side and the left-hand side of the NIS inequality, which is

$$\sum_{a \in A_i} \sum_{\omega \in \Omega} P_i(a, \omega) u(x_i(a, \omega)) - \sum_{a \in A_j} \max_{\hat{a} \in A_i} \sum_{\omega \in \Omega} P_j(a, \omega) u(x_i(\hat{a}, \omega)). \quad (4)$$

⁹This is reminiscent of the ϵ -contour set of Echenique and Pourbabaee (2024), which is all acts that are a $\frac{\epsilon}{1-\epsilon}$ % improvement over another act.

This is achieved using the Delta Method, which is also applied in Dean and Neligh (2023) to test NIAC in Experiment 1.2. This approach assumes only the asymptotic normality of the estimators. For example, in the common case where actions are binary, we observe $\hat{P}(a|\omega)$, which represents a draw from a binomial distribution. This binomial distribution has a mean of $P(a|\omega)$, with the number of trials determined by the frequency of state ω . With a large number of trials, this distribution is approximately normal, validating the assumption. Additionally, the Delta Method relies on the first-order approximation of the Taylor expansion of expression (4). However, since in our case the expression is a linear function of $\hat{P}(a|\omega)$, the first-order approximation of the Taylor expansion is the exact value. This is another reason the Delta Method is especially suitable for our case.

In both Dean and Neligh (2023) and Caplin, Csaba, Leahy, and Nov (2020), the prizes, states, and action spaces are all binary and the states are ex ante equally likely.¹⁰ In that case, the NIS expression (4) reduces to

$$P_i(a_1|\omega_1) + P_i(a_2|\omega_2) - \max_{\hat{a} \in \{a_1, a_2\}} \{P_j(\hat{a}|\omega_1) + (1 - P_j(\hat{a}|\omega_2))\} \geq 0$$

for all $i, j \in D$. This implies

$$P_i(a_1|\omega_1) + P_i(a_2|\omega_2) - \max_{\hat{a} \in \{a_1, a_2\}} \{P_j(\hat{a}|\omega_1) + (1 - P_j(\hat{a}|\omega_2))\} = 0 \quad (5)$$

for all $i \neq j \in D$, and

$$P_i(a_1|\omega_1) + P_i(a_2|\omega_2) - P_i(a_1|\omega_2) - P_i(a_2|\omega_1) \geq 0 \quad (6)$$

for all $i \in D$. As discussed in the Online Appendix, the latter condition (6) within decision problem is equivalent to NIAS in the restricted setting. Therefore, assuming that NIAS is satisfied and plugging into the preceding equation (5) yields the implication that:

$$P_i(a_1|\omega_1) + P_i(a_2|\omega_2) = P_j(a_1|\omega_1) + P_j(a_2|\omega_2) \quad (7)$$

for all $i, j \in D$. Equation (7) represents a subset of the NIS inequalities in our restricted setting, taking NIAS as given. Intuitively, it requires that accuracy be independent of incentives.¹¹ This is easily tested using an off-the-shelf Wald test for multivariate equality. If this relaxed test is rejected, then so is the more stringent NIS condition nesting NIAS. Otherwise, it is possible to test the full set of NIS conditions (5) and (6) using a Wald test for mixed joint hypotheses with equalities and inequalities (Kodde and Palm 1986). In what follows, we focus on the simplified implication (7), given that NIAS is typically satisfied pointwise.¹²

¹⁰In Dewan and Neligh (2020), the state space and action space have size 5. The expression is slightly longer but follows the same principle.

¹¹Intuitively, this is a two-sided strengthening of an analogous one-sided implication of NIAC in (Dean and Neligh 2023, Sec 3.C), namely that subjects become no less accurate as incentives increase.

¹²A single exception is the Angles experiment of Dewan and Neligh (2020) with a \$10 incentive.

4 Existing Experiments

In this section, we re-examine the data from existing experiments of Dean and Neligh (2023), Caplin, Csaba, Leahy, and Nov (2020), and Dewan and Neligh (2020). We focus on these three studies because they jointly satisfy the criteria necessary to implement our test cleanly. First, each study employs a match-the-state structure with a well-defined, finite state space: binary-state environments in DN23 and CCLN, and a five-state environment in DN20. Second, each varies incentives across multiple levels while holding the action set and payoff mapping fixed, which is precisely the variation NIS exploits. Third, each provides the joint distribution of realized states and chosen actions; the minimal data requirement for our test. Finally, all three studies have publicly available datasets and use samples large enough to support statistical inference. Together, these features make them well-suited for the first empirical application of NIS.

Table 1 and Figure 2 provide an overview of the analysis. The empirical results are reached in two steps. First, for each dataset we spell out the inequalities implied by NIS: these specify how accuracy should compare across incentive conditions if perceptual capacity is fixed. Second, we test these predictions against the observed data. A rejection of NIS occurs whenever the observed pattern of accuracies contradicts these inequalities. That is, when different incentives yield systematically and significantly different accuracy than predicted under NIS. With this structure in place, the reader can interpret each set of results as a direct test of the NIS inequalities in the context of the relevant experiment. Therefore, when Table 1 and Figure 2 specify that NIS is rejected, it is rejecting precisely the hypothesis that a change in incentives did not alter the capacity of participants in this experiment. As Table 1 and Figure 2 show, NIS is rejected in most tasks we study, including DN23 (ball counts and equations), CCLN (shapes), and DN20 (dots). However, NIS is not rejected in DN20’s angle-judgment task, suggesting that this type of perceptual problem is less responsive to incentives. This contrast highlights the main takeaway: while incentives systematically expand attention in many standard perceptual settings, some tasks exhibit limits that are insensitive to payoff variation. Figure 2 further shows that (i) there is a negative relationship between IEI and IDI and (ii) IEI/IDI is strictly smaller/larger for tasks in which NIS is rejected.¹³

¹³Table 1 reports the joint p-value aggregated across all NIS inequalities in both directions, as described in Section 3.2. Whenever NIAS holds throughout the relevant decision problems (DN23, CCLN, and DN20’s Dots task) this test reduces to a regular two-sided Wald test. Where NIAS fails for at least one decision problem, as it sometimes does for DN20’s Angles task, we instead use the Kodde and Palm (1986) test. Individual pairwise p-values, reported in Tables 2–6, are one-sided, testing the increasing-incentive direction only.

Experiment	# of incentive levels	NIAS	NIAC	NIS rejected ($\alpha = 0.05$)	NIS joint p-value	IDI	IEI	N
DN23 1.2	4	Pass	Pass	Yes	< 0.01	0.09	0.89	10400
DN23 2.2	2	Pass	Pass	No	0.38	0.02	0.98	5500
CCLN	6	Pass	Fail	Yes	< 0.01	0.07	0.91	26048
CCLN Difficulty 1	6	Pass	Fail	Yes	0.04	0.05	0.92	7136
CCLN Difficulty 2	6	Pass	Fail	Yes	< 0.01	0.08	0.89	6240
CCLN Difficulty 3	6	Pass	Fail	Yes	0.01	0.07	0.90	6656
CCLN Difficulty 6	6	Pass	Fail	Yes	< 0.01	0.08	0.90	6016
DN20 Dots	4	Pass	Pass	Yes	< 0.01	0.24	0.66	8100
DN20 Dots \$10	4	Pass	Pass	Yes	< 0.01	0.23	0.67	4100
DN20 Dots \$20	4	Pass	Pass	Yes	< 0.01	0.25	0.66	4000
DN20 Angles	4	Fail	Fail	No	0.58	0.02	0.96	8100
DN20 Angles \$10	4	Fail	Fail	No	0.22	0.03	0.94	4100
DN20 Angles \$20	4	Pass	Fail	No	0.86	0.02	0.96	4000

Table 1: Summary of results for aggregate data from all experiments

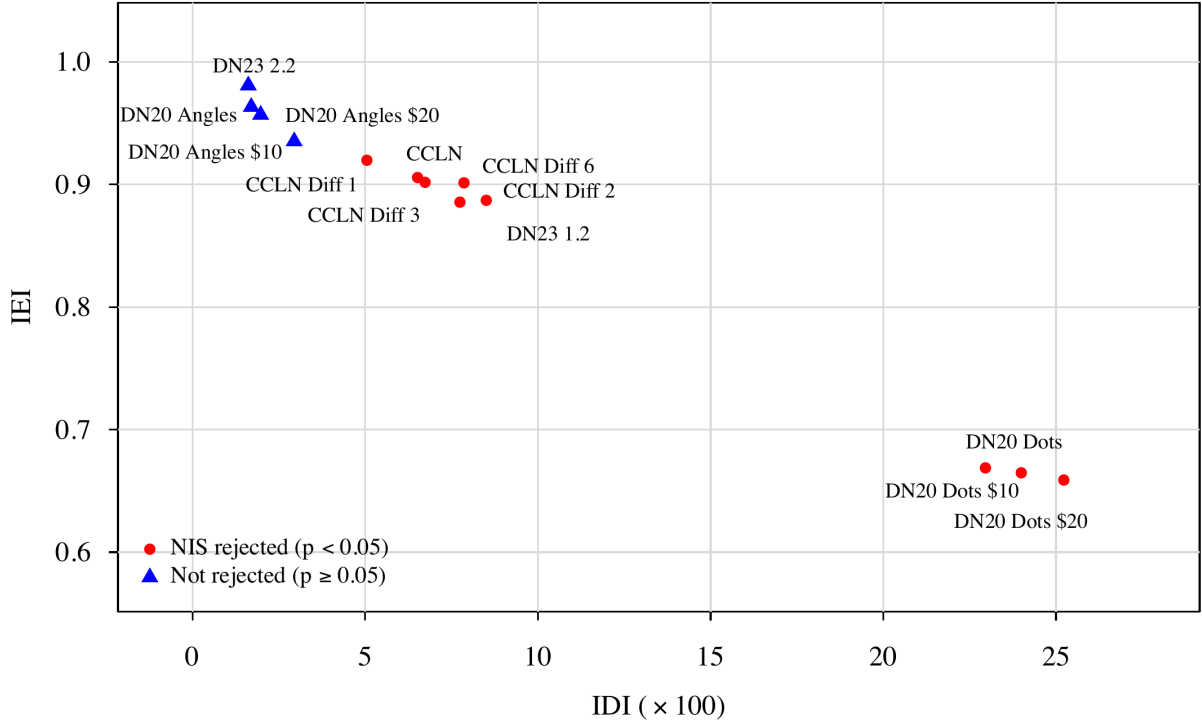


Figure 2: IDI vs. IEI across experiments (DN20 Angles and DN20 Angles \$20 coincide exactly). Color/shape indicates whether the NIS test is rejected at $\alpha = 0.05$.

4.1 Dean and Neligh (2023): Dots & Equations

4.1.1 Design

As described in Section 2.1, participants in Experiment 1.2 of Dean and Neligh (2023) (DN23 hereafter) face a choice between two actions (a or b) with two equally likely states (represented by 49 blue and 51 red balls (R) or 49 red and 51 blue balls (B)). When the state is R (more red balls), the “correct” action is a , and when the state is B (more blue balls), the correct action is b . In our notation, this is analogous to an incentive to choose actions to match the state, i.e. choosing a_k when the state is ω_k for $k = R, B$. For a randomly selected choice, participants are given probability points (for a prize of \$40) if they chose the correct action and the number of probability points varies across decision problems: either 5, 40, 70, or 95 points (4 incentive levels), as can be seen in Table 1. Fifty-two participants face 50 repetitions of each of these four decision problems. Experiment 2.2 has a similar design, but the 55 subjects face seven equations instead of dots with two states of the world being either three or four of the seven equations being correct. The incentive structure remained broadly the same with only the 5 probability points payment and the 95 probability points payment (2 incentive levels, as can be seen in Table 1).

4.1.2 Results

As can be inferred from the joint distribution matrices presented in Section 2.1, subjects improve their accuracy in the direction of increasing incentives, i.e., using the higher incentive-level information structure for a lower incentive-level problem. Formally, we test pairwise NIS inequalities for each ordered pair of incentive levels. With four incentive levels, there are sixteen such ordered pairs in total: six in the direction of increasing incentives, six in the direction of decreasing incentives, and four within the same incentive level. We find that all six comparisons in the direction of increasing incentives yield an improving switch of attention, while none of the remaining comparisons do. The magnitude of these violations in aggregate leads us to reject NIS with a joint p -value less than 0.01 (see Table 1). Table 2 unpacks these results using the aggregate data from DN23 Experiment 1.2 for each incentive level pair, with p -values corresponding to the one-sided null hypothesis that the NIS RHS (performance given capacity at the higher incentive level) is no greater than the LHS (performance given capacity at the lower incentive level) for each incentive-level switch. These violations of NIS arise systematically because higher incentives lead subjects to deploy more attention, generating accuracy gains inconsistent with fixed-capacity models. Put differently, the test fails here not at random but precisely where incentives are raised, highlighting their role in shaping perceptual effort.

If we take the RHS to be the best possible result that can be achieved with the subject’s

Incentive level		NIS LHS	NIS RHS	NIS inequality fails?	p-value
Lower	Higher				
5	40	3.34	3.56	Yes	0.03
5	70	3.34	3.60	Yes	0.02
5	95	3.34	3.77	Yes	< 0.01
40	70	28.46	28.84	Yes	0.33
40	95	28.46	30.12	Yes	0.05
70	95	50.46	52.71	Yes	0.11

Table 2: NIS inequalities in direction of increasing incentives (aggregate data from experiment 1.2 of DN23)

current level of learning and LHS to be the actual level achieved, the difference between them can be used to calculate the Improvability Index (IDI), as explained in Section 3.1. Our measure of IDI is the maximum improvement over all decision problems, normalized by the best possible outcome in each decision problem. The largest normalized improvement relative to baseline comes from the test comparing the performance at the lowest level of incentives (5) with the learning structure from the highest level of incentives (95). The improvement is 0.43 probability points, which corresponds to an increase in 8.5% of the maximum that could be achieved in that decision problem.

The results from Experiment 2.2 of DN23 show a similar pattern: $RHS > LHS$ in Table 3, indicating that switching attention in the direction of increasing incentives would improve expected utility. However, the violation is small in magnitude ($IDI = 1.6\%$) and statistically insignificant ($p = 0.19$), so we cannot reject the null hypothesis that NIS holds in the population. Two factors likely explain this. First, Experiment 2.2 has only two incentive levels, yielding fewer inequalities to test and limiting statistical power. Second, accuracy exceeds 80% even at the lowest incentive level, leaving little room for incentives to expand attention further; a ceiling effect that dampens the scope for detectable violations.

DP1	DP2	LHS	RHS	NIS inequality fails?	p-value
5	95	4.13	4.21	Yes	0.19

Table 3: NIS inequalities in direction of increasing incentives (aggregate data from experiment 2.2 of DN23)

4.2 Caplin, Csaba, Leahy, and Nov (2020): Polygons

4.2.1 Design

In this experimental task, subjects are shown 24 geometric objects at once. Each one of these objects is a polygon that has either 7, 8, 9 or 10 sides. Subjects are asked to make a binary decision indicating whether they believe that there are more 7-sided polygons (action “heptagon”) or 9-sided polygons (action “nonagon”). Again, in our notation, this is analogous to an incentive to choose the action to match the state, i.e. choosing a_k when the state is ω_k for $k = 7, 9$. The number of 7- and 9-sided polygons thus determines the difficulty level of the task: the smaller the difference, the harder the task. Caplin, Csaba, Leahy, and Nov (2020) (CCLN hereafter) vary between participants the difference to be either 1, 2, 3, or 6 (the difficulty level is fixed for a given participant). The 8- and 10-sided polygons are thus decoys. Subjects who make a correct decision in a randomly chosen round are rewarded with probability points of winning \$10 (they receive either 0, 1, 2, 4, 8, 16 or 32 points). To ensure credible probabilistic rewards on MTurk, subjects stopped a computer’s built-in clock, with the last two digits providing a uniform random draw; if this number was below their earned probability points score-based threshold minus 100 (e.g., drawing the number 67 with a score of 172), they won the prize (72% probability in our example), otherwise, they received nothing. Each subject faces 40 rounds in total: 8 at 0 points, 8 at 1 point, 8 at 2 points, 6 at 4 points, 5 at 8 points, 3 at 16 points, and 2 at 32 points. Given that NIS makes no predictions about optimal actions when there are 0 points, we exclude choices under that incentive level.

4.2.2 Results

Table 4 provides the results aggregating across difficulty levels for the task in CCLN. The Online Appendix provides the results disaggregated by task difficulty level. The NIS test outcomes follow the general pattern we saw previously in DN23: the test generally fails in the direction of increasing incentives (with a joint p-value < 0.01). Once again, the violations are concentrated when comparing low- to high-incentive treatments, reinforcing that incentives consistently expand effective perceptual capacity. This pattern directly ties the empirical failures of NIS to the influence of incentives, rather than to random noise or task artifacts.

Figure 3 disaggregates the IDI and IEI measures (described in detail in Section 3) by task difficulty level, complementing the aggregate results in Table 4. Each point in the figure represents a pairwise comparison between two incentive levels, with the x-axis showing the percent increase in reward from the lower to the higher incentive level. The two panels plot the IDI (left) and IEI (right) for each such comparison, where larger NIS violations correspond to higher IDI values and lower IEI values. The best-fit lines are estimated

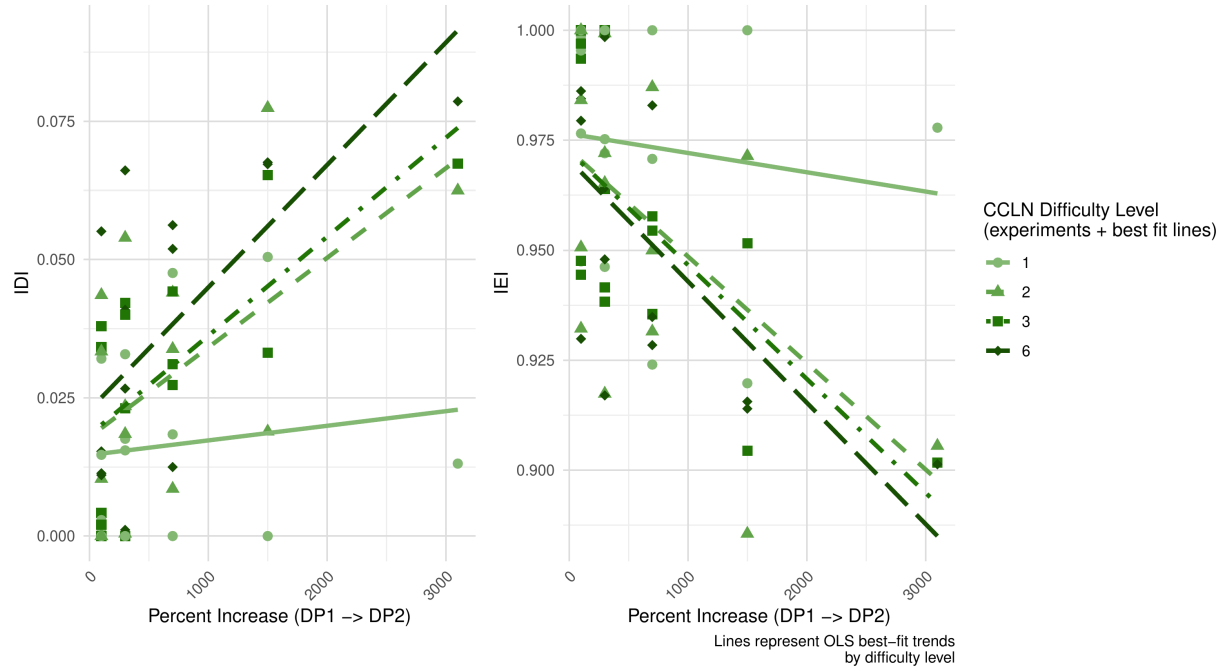


Figure 3: Incentive Responsiveness of NIS Violations by Task Difficulty (CCLN)

Notes: Each point represents a pairwise incentive comparison (DP1 $\rightarrow$ DP2), with the x-axis showing the percent increase in offered reward. The left panel plots IDI (larger values indicate greater violations of NIS) and the right panel plots IEI (smaller values indicate greater violations). Lines of best fit are shown separately for each difficulty level (1, 2, 3, and 6), where higher numbers correspond to easier tasks.

Steeper slopes indicate greater responsiveness of attentional capacity to incentive increases.

DP1	DP2	LHS	RHS	NIS inequality fail?	p-value
1	2	0.63	0.66	Yes	< 0.01
1	4	0.63	0.66	Yes	< 0.01
1	8	0.63	0.66	Yes	< 0.01
1	16	0.63	0.69	Yes	< 0.01
1	32	0.63	0.68	Yes	< 0.01
2	4	1.31	1.32	Yes	0.33
2	8	1.31	1.32	Yes	0.37
2	16	1.31	1.38	Yes	< 0.01
2	32	1.31	1.36	Yes	0.03
4	8	2.64	2.64	No	0.52
4	16	2.64	2.76	Yes	< 0.01
4	32	2.64	2.72	Yes	0.06
8	16	5.27	5.52	Yes	< 0.01
8	32	5.27	5.44	Yes	0.05
16	32	11.05	10.87	No	0.79

Table 4: NIS inequalities in direction of increasing incentives (aggregate data pooled across difficulty levels in CCLN)

separately for each difficulty level, allowing the reader to assess how the relationship between incentive increases and the magnitude of NIS violations varies across task difficulties. While not every individual pairwise comparison yields a statistically significant violation, the joint test confirms that NIS is rejected at every difficulty level (statistically significant in all cases, as reported in Table 1). Crucially, as Figure 3 illustrates, the slope of the best-fit line flattens as tasks become harder (moving from Level 6 to Level 1): easier tasks exhibit a stronger expansion of attention in response to incentives, while harder tasks are less responsive. This gradient directly anticipates the DN20 angles result we find: a perceptual judgment so difficult that it appears largely insensitive to incentive variation altogether. Taken together, these findings suggest that the degree to which incentives can expand attentional capacity is mediated by task difficulty, with harder tasks approaching a ceiling beyond which additional effort yields diminishing perceptual returns.

4.3 Dewan and Neligh (2020): Dots & Angles

4.3.1 Design

Dewan and Neligh (2020) (DN20 hereafter) use two perceptual tasks with a similar structure. In one version of the task, subjects are shown a random arrangement of dots on the screen and are asked to determine the correct number of dots, between 38 and 42 inclusive, with each number being equally likely. The second task involves an identification of the degrees of an angle presented on the screen: 35, 40, 45, 50 or 55 degrees. Thus, both of these decision problems had five possible actions and five possible states. Again, in our notation, this is analogous to an incentive to choose the action to match the state. Each subject was exposed to 100 variations of each task with varying rewards. As is standard, the reward for both tasks was provided in the form of probability points ranging from 1 to 100 for a prize of \$10 in one group of sessions (for 41 subjects) or \$20 in another group of sessions (for 40 subjects) for a total of 8 sessions and 81 subjects. Experimental earnings were based on two randomly selected tasks, one from each half of the experiment, where only correct answers were rewarded. The incentive level of each selected task determined the probability of winning a monetary prize. Since each subject performed each task at a given incentive level only once, we group the observations at the incentive quartile level.

4.3.2 Results

Aggregate results from the dots task in DN20 adhere to the expected pattern: the six tests in the direction of increasing incentives do not pass the NIS condition. Among the analyzed experiments, the dots task also ranks high on the IDI scale. The largest relative improvement occurs when switching from using the learning attributed to the 1st quartile of incentives (lowest) to the 4th quartile of incentives (highest) and constitutes an aggregate improvement of 24.0%.

However, the angles task, a perceptually more difficult task that cannot be verified through effort, shows less responsiveness to the incentives: the joint probability matrices have little to no change in the direction of increasing incentives. The results of the NIS tests reflect this difficulty: in the direction of increasing incentives, all six tests show that participants could not have improved significantly by using a learning structure from the higher incentive tasks. The IDI results also confirm this: while the greatest possible improvement here still comes in the direction of increasing incentives, the magnitude is close to 2%.

Figure 4 plots the IDI and IEI measures for both tasks in DN20, disaggregated by task type and prize size. This allows us to contrast the two tasks: for the dot-counting task, the best-fit line slopes steeply upward (IDI) and downward (IEI) as incentives increase, while

DP 1 (incentive quartile)	DP 2 (incentive quartile)	LHS	RHS	NIS inequality fail?	p-value
1st	2nd	6.18	7.05	Yes	< 0.01
1st	3rd	6.18	8.26	Yes	< 0.01
1st	4th	6.18	9.30	Yes	< 0.01
2nd	3rd	20.62	24.16	Yes	< 0.01
2nd	4th	20.62	27.20	Yes	< 0.01
3rd	4th	40.05	45.09	Yes	< 0.01

Table 5: NIS inequalities in direction of increasing incentives (aggregate data pooled across prize sizes from the Dot task of DN20)

for the angle-judgment task, both lines are nearly flat across the full range of incentive variation. This visual pattern maps directly onto the statistical results in Tables 5 and 6. In the dots task, all six pairwise NIS inequalities in the direction of increasing incentives are violated at conventional significance levels ($p < 0.01$ for all comparisons), with the largest relative improvement of 24.0% arising when switching from the lowest to the highest incentive quartile. By contrast, in the angles task, the NIS inequalities are not systematically violated in the direction of increasing incentives, with p-values ranging from 0.08 to 0.80 and the RHS occasionally falling below the LHS, suggesting no meaningful expansion of attentional capacity with higher stakes.

Importantly, these results reflect within-subject incentive variation: subjects were exposed to all incentive levels within their session, so the differences in behavior across quartiles reflect responses to the incentive they faced. This contrast between the two tasks mirrors the difficulty gradient documented in Figure 3 for the CCLN experiment, where harder tasks (Level 1) showed a flatter slope of incentive responsiveness relative to easier tasks (Level 6). The angle-judgment task can be understood as an extreme case of this pattern: a perceptual judgment so resistant to effortful improvement that it behaves as though perceptual capacity is effectively fixed, consistent with the capacity-constrained learning model. Unlike dot-counting (where additional effort translates into measurable accuracy gains) judging angles appears to impose a hard perceptual limit that incentives alone cannot overcome.

Thus, NIS is a useful condition for demarcating the kinds of problems in which the returns to cognitive effort are negligible (relative to the costs). An interesting avenue for future work would be to verify the condition in settings where incentives may not be monotonically increasing.

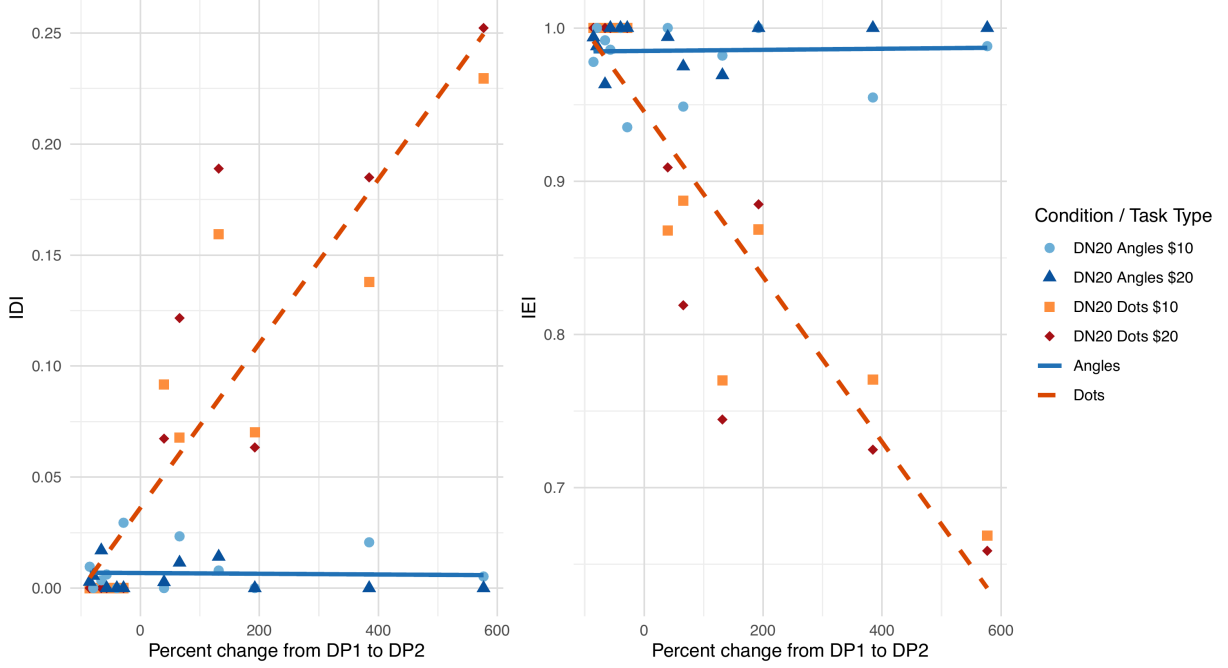


Figure 4: Incentive Responsiveness of NIS Violations by Task Type and Session Prize (DN20)

Notes: Each point represents a pairwise incentive comparison (DP1 $\rightarrow$ DP2), with the x-axis showing the percent increase in offered reward. The left panel plots IDI (larger values indicate greater violations of NIS) and the right panel plots IEI (smaller values indicate greater violations). Points and lines of best fit are shown separately for the dot-counting task (orange) and angle-judgment task (blue), each disaggregated by prize size (\$10 and \$20). The steep slope for the dots task indicates that attentional capacity expands substantially with incentives, while the near-flat slope for the angles task suggests that perceptual capacity is largely unresponsive to incentive variation, consistent with NIS not being rejected for that task.

DP 1 (incentive quartile)	DP 2 (incentive quartile)	LHS	RHS	NIS inequality fail?	p-value
1st	2nd	5.84	5.69	No	0.78
1st	3rd	5.84	5.94	Yes	0.29
1st	4th	5.84	5.78	No	0.63
2nd	3rd	16.62	17.37	Yes	0.08
2nd	4th	16.62	16.89	Yes	0.32
3rd	4th	28.80	28.00	No	0.80

Table 6: Aggregate NIS test results of the Angle task of DN20

5 Discussion

Our findings suggest that the failures of NIS observed across multiple datasets are not random but instead track systematic changes in incentives. In DN23, CCLN, and DN20 (dots), violations occur precisely when incentives increase, consistent with the view that higher stakes expand the effective capacity of attention. By contrast, the absence of violations in DN20 (angles) shows that not all perceptual tasks are equally incentive-responsive: some limits appear “hard” and insensitive to payoff variation. This heterogeneity implies that the role of incentives in shaping attention is conditional, operating strongly in domains where effortful attention can augment perception, but weakly where perceptual constraints are rigid. Taken together, these results clarify the contribution of the NIS test: it serves both as a diagnostic for capacity-constrained learning and as a tool for mapping when and how incentives shape the scope of human attention.

One of the key questions raised by the application of NIS in this paper is how to characterize the class of tasks that fall within the capacity-constrained learning paradigm, and those that do not. By using NIS as a diagnostic tool, an experimentalist can assess whether performance in a given task is consistent with a fixed capacity constraint, or instead responds to incentives in a way that suggests an expansion of attention. Doing so requires experimental designs in which the underlying states are well defined and observable to the analyst, the action set and payoff mapping are held fixed across incentive treatments, and incentives vary only in scale. Within-subject designs with randomized exposure to incentive levels are particularly useful, as they allow differences in observed behavior to be interpreted as changes in learning rather than differences in subject composition, and they provide the repeated observations needed to estimate state-dependent choice frequencies with precision.

A closely related direction for future research is to explicitly manipulate the learning feasibility set itself, for example by introducing additional noise, imposing time limits, or varying task difficulty. While the analysis of Caplin, Csaba, Leahy, and Nov (2020) shows that NIS is rejected across all difficulty levels in their task, that design was not intended to test whether increased difficulty shifts tasks into or out of the capacity-constrained regime. Experiments that independently vary incentives and feasibility, while maintaining a fixed mapping between states, actions, and payoffs, would therefore be particularly informative in clarifying which task characteristics are conducive to fixed-capacity behavior and which are not.

Our test may also apply to strategic settings with state-dependent utility, where players face uncertainty about an exogenous state of the world. Previous work applying rational inattention models to games includes de Clippel and Rozen (2020), Spurlino (2022), and Almog and Martin (2024). In such settings, NIS can be used to test whether players’ perceptual constraints are fixed or respond to incentives. Relatedly, Caplin and Martin (2015) establish

a connection between rational inattention and quantal response equilibrium (QRE): NIAS is satisfied by the choices that arise from QRE if and only if the noise is small. This suggests that NIS and related conditions can help distinguish models of noisy perception (consistent with Bayesian learning) from models of noisy response (such as standard QRE). However, our framework does not directly apply to games (or individual decision problems) that do not have an exogenous state and prior, such as the complete-information strategic settings studied by Goeree and Holt (2001) and the endogenous-state dynamic decision problem studied by Houser, Keane, and McCabe (2004). Extending NIS to such environments, where the relevant “state” is endogenous behavior, is a potential direction for future research.

Statements and Declarations

Funding: The authors acknowledge support from the Sloan Foundation under the “Cognitive Economics at Work” grant.

Use of artificial intelligence: OpenAI Codex (GPT-5) was used in August of 2026 to review the manuscript and replication code. The authors take full responsibility for the final content and accuracy of both.

Competing interests: The authors declare no competing interests.

Data availability: Data and reproduction materials are available at <https://doi.org/10.7910/DVN/1EXCE4>. This paper reanalyzes previously published and public data.

Ethics approval and informed consent: This paper reanalyzes previously published and public data and does not report newly collected human-subjects data.

References

- Afriat, Sydney N (1973). “On a system of inequalities in demand analysis: an extension of the classical method”. In: *International economic review*, pp. 460–472.
- Almog, David, Romain Gauriot, Lionel Page, and Daniel Martin (2024). “AI oversight and human mistakes: Evidence from Centre Court”. In: *arXiv preprint arXiv:2401.16754*.
- Almog, David and Daniel Martin (2024). “Rational inattention in games: experimental evidence”. In: *Experimental Economics*, pp. 1–28.
- Ambuehl, Sandro, Axel Ockenfels, and Colin Stewart (2025). “Who Opt In? Composition Effects and Disappointment from Participation Payments”. In: *The Review of Economics and Statistics* 107.1, pp. 78–94. ISSN: 0034-6535. DOI: [10.1162/rest_a_01268](https://doi.org/10.1162/rest_a_01268). eprint: https://direct.mit.edu/rest/article-pdf/107/1/78/2058750/rest_a_01268.pdf. URL: https://doi.org/10.1162/rest_a_01268.
- Apesteguia, Jose and Miguel A Ballester (2015). “A measure of rationality and welfare”. In: *Journal of Political Economy* 123.6, pp. 1278–1310.
- Ba, Cuimin, J Aislinn Bohren, and Alex Imas (2022). “Over-and underreaction to information”. In: *SSRN 4274617*.
- Bronchetti, Erin T et al. (2023). “Is Attention Produced Optimally? Theory and Evidence From Experiments With Bandwidth Enhancements”. In: *Econometrica* 91.2, pp. 669–707.
- Camerer, Colin F and Robin M Hogarth (1999). “The effects of financial incentives in experiments: A review and capital-labor-production framework”. In: *Journal of risk and uncertainty* 19, pp. 7–42.
- Caplin, Andrew (2025). “Data Engineering for Cognitive Economics”. In: *Journal of Economic Literature* 63.1, pp. 164–196. DOI: [10.1257/jel.20241351](https://doi.org/10.1257/jel.20241351).
- Caplin, Andrew, Dániel Csaba, John Leahy, and Oded Nov (2020). “Rational inattention, competitive supply, and psychometrics”. In: *The Quarterly Journal of Economics* 135.3, pp. 1681–1724.
- Caplin, Andrew and Mark Dean (2015). “Revealed preference, rational inattention, and costly information acquisition”. In: *American Economic Review* 105.7, pp. 2183–2203.
- Caplin, Andrew and Daniel Martin (2015). “A testable theory of imperfect perception”. In: *The Economic Journal* 125.582, pp. 184–202.
- Caplin, Andrew, Daniel Martin, and Philip Marx (2025). “Modeling Machine Learning: A Cognitive Economic Approach”. In: *Journal of Economic Theory* 224, p. 105970. DOI: [10.1016/j.jet.2025.105970](https://doi.org/10.1016/j.jet.2025.105970).
- Charles, Constantin, Cary Frydman, and Mete Kilic (2024). “Insensitive investors”. In: *The Journal of Finance* 79.4, pp. 2473–2503.
- de Clippel, Geoffroy and Kareen Rozen (2020). *Communication, Perception, and Strategic Obfuscation*. Working Paper.

- De Oliveira, Henrique, Tommaso Denti, Maximilian Mihm, and Kemal Ozbek (2017). “Rationally inattentive preferences and hidden information costs”. In: *Theoretical Economics* 12.2, pp. 621–654.
- Dean, Mark and Daniel Martin (2016). “Measuring rationality with the minimum cost of revealed preference violations”. In: *Review of Economics and Statistics* 98.3, pp. 524–534.
- Dean, Mark and Nathaniel Neligh (2023). “Experimental tests of rational inattention”. In: *Journal of Political Economy* 131.12, pp. 3415–3461.
- Dewan, Ambuj and Nathaniel Neligh (2020). “Estimating information cost functions in models of rational inattention”. In: *Journal of Economic Theory* 187, p. 105011.
- Echenique, Federico, Sangmok Lee, and Matthew Shum (2011). “The money pump as a measure of revealed preference violations”. In: *Journal of Political Economy* 119.6, pp. 1201–1223.
- Echenique, Federico and Farzad Pourbabaei (2024). “Individual and Collective Welfare in Risk Sharing with Many States”. In: *arXiv preprint arXiv:2401.07337*.
- Frydman, Cary and Lawrence J. Jin (2022). “Efficient coding and risky choice”. In: *The Quarterly Journal of Economics* 137.1, pp. 161–213.
- (2023). *On the Source and Instability of Probability Weighting*. Working Paper 31573. National Bureau of Economic Research.
- Frydman, Cary and Salvatore Nunnari (2021). *Coordination with cognitive noise*. Working Paper. CESifo Working Paper.
- Gabaix, Xavier (2014). “A Sparsity-Based Model of Bounded Rationality”. In: *The Quarterly Journal of Economics* 129.4, pp. 1661–1710. ISSN: 0033-5533. DOI: [10.1093/qje/qju024](https://doi.org/10.1093/qje/qju024). eprint: <https://academic.oup.com/qje/article-pdf/129/4/1661/30631668/qju024.pdf>. URL: <https://doi.org/10.1093/qje/qju024>.
- (2019). “Chapter 4 - Behavioral inattention”. In: *Handbook of Behavioral Economics - Foundations and Applications 2*. Ed. by B. Douglas Bernheim, Stefano DellaVigna, and David Laibson. Vol. 2. Handbook of Behavioral Economics: Applications and Foundations 1. North-Holland, pp. 261–343. DOI: <https://doi.org/10.1016/bs.hesbe.2018.11.001>. URL: <https://www.sciencedirect.com/science/article/pii/S2352239918300216>.
- Goeree, Jacob K and Charles A Holt (2001). “Ten little treasures of game theory and ten intuitive contradictions”. In: *American Economic Review* 91.5, pp. 1402–1422.
- Houser, Daniel, Michael Keane, and Kevin McCabe (2004). “Behavior in a dynamic decision problem: An analysis of experimental evidence using a Bayesian type classification algorithm”. In: *Econometrica* 72.3, pp. 781–822.
- Khaw, Mel Win, Ziang Li, and Michael Woodford (2021). “Cognitive imprecision and small-stakes risk aversion”. In: *The review of economic studies* 88.4, pp. 1979–2013.

- Kodde, David A and Franz C Palm (1986). “Wald criteria for jointly testing equality and inequality restrictions”. In: *Econometrica: journal of the Econometric Society*, pp. 1243–1248.
- Maćkowiak, Bartosz, Filip Matějka, and Mirko Wiederholt (2023). “Rational Inattention: A Review”. In: *Journal of Economic Literature* 61.1, pp. 226–73. DOI: [10.1257/jel.20211524](https://doi.org/10.1257/jel.20211524). URL: <https://www.aeaweb.org/articles?id=10.1257/jel.20211524>.
- Matejka, Filip and Alisdair McKay (2015). “Rational inattention to discrete choices: A new foundation for the multinomial logit model”. In: *American Economic Review* 105.1, pp. 272–98.
- Ok, Efe A and Gerelt Tserenjigmid (2023). “Measuring Stochastic Rationality”. In: *arXiv preprint arXiv:2303.08202*.
- Pessoa, Luiz and Jan B Engelmann (2010). “Embedding reward signals into perception and cognition”. In: *Frontiers in neuroscience* 4, p. 17.
- Plott, Charles R (1986). “Rational choice in experimental markets”. In: *Journal of Business*, S301–S327.
- Shenhav, Amitai et al. (2017). “Toward a Rational and Mechanistic Account of Mental Effort”. In: *Annual Review of Neuroscience* 40, pp. 99–124.
- Sims, Christopher A (2003). “Implications of Rational Inattention”. In: *Journal of Monetary Economics* 50.3, pp. 665–690.
- Smith, Vernon L (1991). “Rational choice: The contrast between economics and psychology”. In: *Journal of Political Economy* 99.4, pp. 877–897.
- Spurlino, Eric (2022). *Rationally Inattentive and Strategically (Un) sophisticated: Theory and Experiment*. Tech. rep. Working paper.
- Varian, Hal R et al. (1991). *Goodness-of-fit for revealed preference tests*. Working Paper. Department of Economics, University of Michigan Ann Arbor.
- Weber, Ernst Heinrich (1834). *De Pulsu, resorptione, auditu et tactu: Annotationes anatomicae et physiologicae...* CF Koehler.
- Woodford, Michael (2012). “Prospect theory as efficient perceptual distortion”. In: *American Economic Review* 102.3, pp. 41–46.