

Nonstandard Errors

ALBERT J. MENKVELD, ANNA DREBER, FELIX HOLZMEISTER,
 JUERGEN HUBER, MAGNUS JOHANNESSEN, MICHAEL KIRCHLER,
 SEBASTIAN NEUSÜß, MICHAEL RAZEN, UTZ WEITZEL, DAVID ABAD-DÍAZ,
 MENACHEM (MENI) ABUDY, TOBIAS ADRIAN, YACINE AIT-SAHALIA,
 OLIVIER AKMANSOY, JAMIE T. ALCOCK, VITALI ALEXEEV, ARASH ALOOSH,
 LIVIA AMATO, DIEGO AMAYA, JAMES J. ANGEL, ALEJANDRO T. AVETIKIAN,
 AMADEUS BACH, EDWIN BAIDOO, GAETAN BAKALLI, LI BAO,
 ANDREA BARBON, OKSANA BASHCHENKO, PARAMPREET C. BINDRA,
 GEIR H. BJØNNES, JEFFREY R. BLACK, BERNARD S. BLACK,
 DIMITAR BOGOEV, SANTIAGO BOHORQUEZ CORREA, OLEG BONDARENKO,
 CHARLES S. BOS, CIRIL BOSCH-ROSA, ELIE BOURI, CHRISTIAN BROWNLEES,
 ANNA CALAMIA, VIET NGÀ CAO, GUNTHER CAPELLE-BLANCARD,
 LAURA M. CAPERA ROMERO, MASSIMILIANO CAPORIN, ALLEN CARRION,
 TOLGA CASKURLU, BIDISHA CHAKRABARTY, JIAN CHEN,
 MIKHAIL CHERNOV, WILLIAM CHEUNG, LUDWIG B. CHINCARINI,
 TARUN CHORDIA, SHEUNG-CHI CHOW, BENJAMIN CLAPHAM,
 JEAN-EDOUARD COLLIARD, CAROLE COMERTON-FORDE, EDWARD CURRAN,
 THONG DAO, WALE DARE, RYAN J. DAVIES, RICCARDO DE BLASIS,
 GIANLUCA F. DE NARD, FANY DECLERCK, OLEG DEEV, HANS DEGRYSE,
 SOLOMON Y. DEKU, CHRISTOPHE DESAGRE, MATHIJS A. VAN DIJK,
 CHUKWUMA DIM, THOMAS DIMPFL, YUN JIANG DONG,
 PHILIP A. DRUMMOND, TOM DUDDA, TEODOR DUEVSKI,
 ARIADNA DUMITRESCU, TEODOR DYAKOV, ANNE HAUBO DYHRBERG,
 MICHAŁ DZIELIŃSKI, ASLI EKSI, IZIDIN EL KALAK, SASKIA TER ELLEN,
 NICOLAS EUGSTER, MARTIN D. D. EVANS, MICHAEL FARRELL,
 ESTER FELEZ-VINAS, GERARDO FERRARA, EL MEHDI FERROUHI,
 ANDREA FLORI, JONATHAN T. FLUHARTY-JAIDEE, SEAN D. V. FOLEY,
 KINGSLEY Y. L. FONG, THIERRY FOUCAULT, TATIANA FRANUS,
 FRANCESCO FRANZONI, BART FRIJNS, MICHAEL FRÖMMEL,
 SERVANNA M. FU, SASCHA C. FÜLLBRUNN, BAOQING GAN, GE GAO,
 THOMAS P. GEHRIG, ROLAND GEMAYEL, DIRK GERRITSEN,
 JAVIER GIL-BAZO, DUDLEY GILDER, LAWRENCE R. GLOSTEN,
 THOMAS GOMEZ, ARSENY GORBENKO, JOACHIM GRAMMIG,
 VINCENT GRÉGOIRE, UFUK GÜÇBILMEZ, BJÖRN HAGSTRÖMER,
 JULIEN HAMBUCKERS, ERIK HAPNES, JEFFREY H. HARRIS,
 LAWRENCE HARRIS, SIMON HARTMANN, JEAN-BAPTISTE HASSE,
 NIKOLAUS HAUTSCH, XUE-ZHONG (TONY) HE, DAVIDSON HEATH,
 SIMON HEDIGER, TERRENCE HENDERSHOTT, ANN MARIE HIBBERT,
 ERIK HJALMARSSON, SETH A. HOELSCHER, PETER HOFFMANN,
 CRAIG W. HOLDEN, ALEX R. HORENSTEIN, WENQIAN HUANG, DA HUANG,
 CHRISTOPHE HURLIN, KONRAD ILCZUK, ALEXEY IVASHCHENKO,
 SUBRAMANIAN R. IYER, HOSSEIN JAHANSHAHLOO, NAJI JALKH,
 CHARLES M. JONES, SIMON JURKATIS, PETRI JYLHÄ, ANDREAS T. KAECK,

This is an open access article under the terms of the [Creative Commons Attribution](#) License, which permits use, distribution and reproduction in any medium, provided the original work is properly cited.

DOI: 10.1111/jofi.13337

© 2024 The Authors. *The Journal of Finance* published by Wiley Periodicals LLC on behalf of American Finance Association.

GABRIEL KAISER, ARZÉ KARAM, EGLE KARMAZIENE, BERNHARD KASSNER,
 MARKKU KAUSTIA, EKATERINA KAZAK, FEARGHAL KEARNEY,
 VINCENT VAN KERVEL, SAAD A. KHAN, MARTA K. KHOMYN, TONY KLEIN,
 OLGA KLEIN, ALEXANDER KLOS, MICHAEL KOETTER,
 ALEKSEY KOLOKOLOV, ROBERT A. KORAJCZYK, ROMAN KOZHAN,
 JAN P. KRAHNEN, PAUL KUHLE, AMY KWAN, QUENTIN LAJAUNIE,
 F. Y. ERIC C. LAM, MARIE LAMBERT, HUGUES LANGLOIS, JENS LAUSEN,
 TOBIAS LAUTER, MARKUS LEIPPOLD, VLADIMIR LEVIN, YIJIE LI, HUI LI,
 CHEE YOONG LIEW, THOMAS LINDNER, OLIVER LINTON, JIACHENG LIU,
 ANQI LIU, GUILLERMO LLORENTE, MATTHIJS LOF, ARIEL LOHR,
 FRANCIS LONGSTAFF, ALEJANDRO LOPEZ-LIRA, SHAWN MANKAD,
 NICOLA MANO, ALEXIS MARCHAL, CHARLES MARTINEAU,
 FRANCESCO MAZZOLA, DEBRAH MELOSO, MICHAEL G. MI, ROXANA MIHET,
 VIJAY MOHAN, SOPHIE MOINAS, DAVID MOORE, LIANGYI MU,
 DMITRIY MURAVYEV, DERMOT MURPHY, GABOR NESZVEDA,
 CHRISTIAN NEUMEIER, ULF NIELSSON, MAHENDRARAJAH NIMALENDRAN,
 SVEN NOLTE, LARS L. NORDEN, PETER O'NEILL, KHALED OBAID,
 BERNT A. ØDEGAARD, PER ÖSTBERG, EMILIANO PAGNOTTA,
 MARCUS PAINTER, STEFAN PALAN, IMON J. PALIT, ANDREAS PARK,
 ROBERTO PASCUAL, PAOLO PASQUARIELLO, LUBOS PASTOR, VINAY PATEL,
 ANDREW J. PATTON, NEIL D. PEARSON, LORIANA PELIZZON,
 MICHELE PELLI, MATTHIAS PELSTER, CHRISTOPHE PÉRIGNON,
 CAMERON PFIFFER, RICHARD PHILIP, TOMÁŠ PLÍHAL, PUNEET PRAKASH,
 OLIVER-ALEXANDER PRESS, TINA PRODROMOU, MARCEL PROKOPCZUK,
 TALIS PUTNINS, YA QIAN, GAURAV RAIZADA, DAVID RAKOWSKI,
 ANGELO RANALDO, LUCA REGIS, STEFAN REITZ, THOMAS RENAULT,
 REX W. RENJIE, ROBERTO RENO, STEVEN J. RIDDIOUGH, KALLE RINNE,
 PAUL RINTAMÄKI, RYAN RIORDAN, THOMAS RITTMANNSSBERGER,
 IÑAKI RODRÍGUEZ LONGARELA, DOMINIK ROESCH, LAVINIA ROGNONE,
 BRIAN ROSEMAN, IOANID ROȘU, SAURABH ROY, NICOLAS RUDOLF,
 STEPHEN R. RUSH, KHALADDIN RZAYEV, ALEKSANDRA A. RZEŹNIK,
 ANTHONY SANFORD, HARIKUMAR SANKARAN, ASANI SARKAR,
 LUCIO SARNO, OLIVIER SCAILLET, STEFAN SCHARNOWSKI,
 KLAUS R. SCHENK-HOPPÉ, ANDREA SCHERTLER, MICHAEL SCHNEIDER,
 FLORIAN SCHROEDER, NORMAN SCHÜRHOFF, PHILIPP SCHUSTER,
 MARCO A. SCHWARZ, MARK S. SEASHOLES, NORMAN J. SEEGER,
 OR SHACHAR, ANDRIY SHKILKO, JESSICA SHUI, MARIO SIKIC,
 GIORGIA SIMION, LEE A. SMALES, PAUL SÖDERLIND, ELVIRA SOJLI,
 KONSTANTIN SOKOLOV, JANTJE SÖNKSEN, LAIMA SPOKEVICIUTE,
 DENITSA STEFANOVA, MARTI G. SUBRAHMANYAM, BARNABAS SZASZI,
 OLEKSANDR TALAVERA, YUEHUA TANG, NICK TAYLOR, WING WAH THAM,
 ERIK THEISSEN, JULIAN THIMME, IAN TONKS, HAI TRAN, LUCA TRAPIN,
 ANDERS B. TROLLE, M. ANDREEA VADUVA, GIORGIO VALENTE,
 ROBERT A. VAN NESS, AURELIO VASQUEZ, THANOS VEROUSIS,
 PATRICK VERWIJMEREN, ANDERS VILHELMSSON, GRIGORY VILKOV,
 VLADIMIR VLADIMIROV, SEBASTIAN VOGEL, STEFAN VOIGT,
 WOLF WAGNER, THOMAS WALTHER, PATRICK WEISS,
 MICHEL VAN DER WEL, INGRID M. WERNER, P. JOAKIM WESTERHOLM,
 CHRISTIAN WESTHEIDE, HANS C. WIKA, EVERT WIPPLINGER,

MICHAEL WOLF, CHRISTIAN C. P. WOLFF, LEONARD WOLK,
 WING-KEUNG WONG, JAN WRAMPELMEYER, ZHEN-XING WU, SHUO XIA,
 DACHENG XIU, KE XU, CAIHONG XU, PRADEEP K. YADAV, JOSÉ YAGÜE,
 CHENG YAN, ANTTI YANG, WOONGSUN YOO, WENJIA YU, YIHE YU,
 SHIHAO YU, BART Z. YUESHEN, DARYA YUFEROVA, MARCIN ZAMOJSKI,
 ABALFAZL ZAREEI, STEFAN M. ZEISBERGER, LU ZHANG, S. SARAH ZHANG,
 XIAOYU ZHANG, LU ZHAO, ZHUO ZHONG, Z. IVY ZHOU, CHEN ZHOU,
 XINGYU S. ZHU, MARIUS ZOICAN, and REMCO ZWINKELS*

ABSTRACT

In statistics, samples are drawn from a population in a data-generating process (DGP). Standard errors measure the uncertainty in estimates of population parameters. In science, evidence is generated to test hypotheses in an evidence-generating process (EGP). We claim that EGP variation across researchers adds uncertainty—nonstandard errors (NSEs). We study NSEs by letting 164 teams test the same hypotheses on the same data. NSEs turn out to be sizable, but

*Albert J. Menkveld is at Vrije Universiteit Amsterdam and Tinbergen Institute. Anna Dreber is at Stockholm School of Economics and University of Innsbruck. Felix Holzmeister is at University of Innsbruck. Juergen Huber is at University of Innsbruck. Magnus Johannesson is at Stockholm School of Economics. Michael Kirchler is at University of Innsbruck. Sebastian Neustiß is at Optiver Amsterdam. Michael Razen is at University of Innsbruck. Utz Weitzel is at Vrije Universiteit Amsterdam, Radboud University, and Tinbergen Institute. These first nine authors are the project coordinators. They conceptualized and designed the project, managed it, conducted the meta-analyses, and wrote the paper. Any errors are therefore their sole responsibility. The three authors affiliated with the Certification Agency for Scientific Code and Data (Cascad) conducted the reproducibility verification. The other authors all significantly contributed to the project by participating either as a member of a research team or as a peer evaluator. We have included the full list of authors names and their affiliations in Appendix A, because space in this footnote is limited. The views expressed here are those of the authors and do not represent the views of the Bank of England, the Federal Reserve Bank of New York, the Federal Reserve System, or any other of the institutions that the authors are affiliated with or receive financing from. The coordinators thank Andrew Chen; Amit Goyal; Campbell Harvey; Lucas Saru; Eric Uhlmann; and participants at the Microstructure Exchange 2021, Derivatives Forum Frankfurt 2022, Financial Intermediation Research Society (FIRS) 2022, Research in Behavioral Finance Conference (RBFC) 2022, Society for Experimental Finance (SEF) 2022, Society for Financial Econometrics (SoFiE) 2022 where the paper was runner-up for the best-paper prize, Vienna-Copenhagen Conference on Financial Econometrics 2022, and the Western Finance Association (WFA) 2022 for valuable comments. They further thank Adam Gill, Eugénie de Jong, Ingrid Löfman, and Elmar Nijkamp for research assistance. The coordinators are grateful for financial support from (Dreber) the Knut and Alice Wallenberg Foundation, the Marianne, Marcus Wallenberg Foundation, the Jan Wallander, Tom Hedelius Foundation, (Huber) an Austrian Science Fund grant P29362, (Huber and Kirchler) Austrian Science Fund SFBF63, (Johannesson) Riksbankens Jubileumsfond grant P21-0168, and (Menkveld) Dutch Research Council 016.Vici.185.068. Authors at Gothenburg University, Lund University, and Stockholm University gratefully acknowledge financial support from the Swedish House of Finance. The collection of data on human subjects was approved by the Institutional Review Board (IRB) of the School of Business and Economics at the Vrije Universiteit Amsterdam prior to the experiment (IRB reference number is SBE5/5/2021gwl260). The project coordinators have read *The Journal of Finance's* disclosure policy and have no conflicts of interest to disclose.

Correspondence: Albert J. Menkveld, School of Business and Economics, Vrije Universiteit Amsterdam, De Boelelaan 1105, 1081 HV Amsterdam, The Netherlands; e-mail: albertjmenkveld@gmail.com.

smaller for more reproducible or higher rated research. Adding peer-review stages reduces NSEs. We further find that this type of uncertainty is underestimated by participants.

IN THEIR RECENT BOOK, Kahneman, Sibony, and Sunstein (2021), (KSS) discuss variability in human judgment in terms of noise. To illustrate their analysis, they consider the setting of judges passing sentence. They decompose total variation in sentencing into two canonical components: level noise and pattern noise (chapter 6). Level noise captures the extent to which some judges are more lenient than others, while pattern noise captures variation in the sentences of the same judge hearing similar cases. In statistical terms, this distinction can be thought of as across-judge versus within-judge variation. Variation across judges is also referred as variation in judge fixed effects.

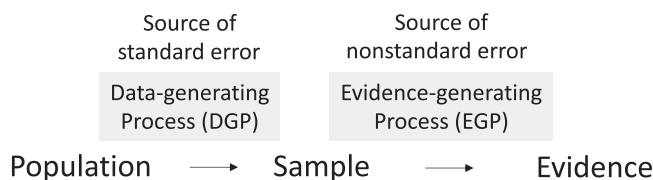
There are similarities to empirical science, where researchers analyze samples to test hypotheses. Within-researcher variation arises from sampling error. Resampling (or bootstrapping) yields different values of the estimator. The standard deviation (SD) of this distribution is referred to as standard error (SE) (Yule (1897)). SE is a source of uncertainty that researchers are well aware of, and typically account for, in conducting their tests.

Researchers are less aware of the additional uncertainty due to there not being a standard analysis path. Researchers vary in what they deem to be the most reasonable path in the “garden of forking paths” (Gelman and Loken (2014)). Conditional on the path, there is a well-defined estimator and SE. Conditional on the sample, however, estimates may vary across researchers as they might pick different paths.¹ We refer to this additional variation as nonstandard error (NSE). Note that the adjective *nonstandard* emphasizes the lack of a standard approach. In other words, if all researchers agree on one path being the most reasonable one, NSE is zero.

The schema below summarizes the overarching idea behind NSEs. Statisticians use the term data-generating process (DGP) to refer to the idea that samples are random draws from a population. Estimators therefore exhibit SE. Using similar language, one could say that researchers also engage in an evidence-generating process (EGP), whereby researchers potentially pick different analysis paths. Estimators therefore also exhibit NSE. Note that the latter type of error can be thought of as erratic as opposed to erroneous, in the sense that there simply is no right path in an absolute sense.²

¹ An important source of such variation is that researchers need to translate conceptual research questions into empirical research questions (Brezna et al. (2022)).

² Variation in estimates reported in meta studies is of both types. Under the polar cases, estimates vary because researchers do the analysis the exact same way but on different samples (SE) or because the sample is the same but the analysis differs (NSE). Mavroeidis, Plagborg-Møller, and Stock (2014) is a special case, as they conduct their meta study by applying all observed analysis paths on all samples. However, unlike us, they do not focus on distinguishing the two sources of variation explicitly. For a review of meta studies in finance, see Geyer-Klingenberg, Hang, and Rathgeber (2020).



To illustrate the idea, consider the following example. In microstructure, market efficiency is conceptually defined as the extent to which a price process resembles a random walk. Suppose that one is interested in estimating the trend in a measure of market efficiency. To estimate, say, the mean annual change in market efficiency, a researcher must decide how to measure market efficiency, at what frequency to sample the data, how to define outliers, etc. These decisions together form what we refer to as the analysis path.

Our objective is to measure and analyze NSEs. The four questions that we focus on are as follows:

1. How large are NSEs in finance?
2. Can NSEs be “explained” in the cross-section of researchers? In particular, are they smaller for papers by higher quality teams, papers with more reproducible results, and papers that score higher in peer evaluations?
3. Does peer feedback reduce NSEs?
4. Do researchers know the size of NSEs?

The motivation for these questions is the fact that NSEs are undesirable in the sense that they add uncertainty. Such uncertainty is particularly worrisome when some estimates are positive while others are negative. We therefore want to understand if higher quality coincides with tighter NSEs, and if feedback reduces NSEs.

Finding answers to the four questions above is extremely costly in terms of human resources. The core structure of an ideal experiment involves two sizable sets of representative researchers. A first set of researchers independently tests the same hypotheses on the same data and writes a short paper presenting the results. A second, nonoverlapping set of researchers evaluates these papers and provides feedback in a single-blind process.

We ran such an experiment under the #fincap tag (FINance Crowd Analysis Project), with 164 research teams (RTs) and 34 peer evaluators (PEs) participating, and each PE evaluating about 10 papers. The Deutsche Börse kindly made proprietary data available spanning 17 years of trading in Europe’s most actively traded instrument—the EuroStoxx 50 index futures. These data enabled researchers to test predefined RT hypotheses³ on several important market trends. This unique opportunity might explain why participation was exceptionally high (at least double that of similar experiments elsewhere,

³ We refer to these hypotheses as RT hypotheses to distinguish them from the hypotheses that we test when analyzing the #fincap results. The RT hypotheses are based on the four overarching questions (Section I.B).

as we discuss later in the introduction).⁴ A back-of-the-envelope calculation shows that the total human resources for #fincap span almost a single academic career: $(164 \times 2 \text{ months} + 34 \times 2 \text{ days} \approx 27 \text{ years})$.

Statistical Framework. We define the NSE for a particular RT hypothesis as the interquartile range (IQR) in estimates across researchers. The reason for picking a robust dispersion measure instead of SD is that the distribution of the SD could exhibit fat tails and thus be prone to outliers. #fincap itself is a case in point as will become clear. The distribution of estimates across researchers tends to the distribution of researcher fixed effects (RFEs), in the sense of each researcher picking his preferred analysis path. Importantly, a distribution of RFEs could be any distribution. Using a robust dispersion measure is therefore a prudent choice.⁵

Statistical inference in #fincap needs to account for multiple hypothesis testing (MHT) (Bonferroni (1936), Šidák (1967)). In particular, the critical values for individual tests need to account for multiple teams testing the same hypothesis. Put simply, if individual tests are performed at a 5% level, then the probability of at least one being significant over multiple tests (weakly) exceeds 5%. Harvey, Liu, and Zhu (2016) illustrate how to adjust levels in asset pricing tests. In his presidential address, Harvey (2017) emphasizes that MHT affects all of finance. We follow in his footsteps when applying MHT in #fincap.

Finally, to address the four questions of interest, we need to analyze how NSEs covary with quality measures as well as how they change across stages. Since NSE is defined in terms of quantiles, we use quantile regressions to conduct this analysis (Koenker and Bassett Jr. (1978)). Note that ordinary least squares (OLS) models conditional means only, and therefore it is inappropriate for an analysis of dispersion. In addition to the first and third quartiles, we also model the median, the first decile, and the ninth decile to obtain a more complete view of the distribution, including results on the interdecile range (IDR).

Summary of Our Findings. We first show that the group of #fincap participants is representative of the academic community in empirical finance/liquidity. About one-third of the 164 RTs have at least one member with publications in the top-three finance or top-five economics journals.⁶ For the group of PEs, this share is 85%. Similarly, 52% of RTs consist of at least one

⁴ #fincap was presented to all participants by means of a dedicated website (<https://fincap.academy>) and a short video (<https://youtu.be/HPtnus0Yu-o>).

⁵ The intuition is as follows. If the number of researchers tends to infinity, then the distribution of estimates tends to the distribution of RFEs, plus sampling errors. If, in addition, the sample size tends to infinity, then the distribution of estimates tends to the distribution of RFEs (because, for each analysis path, the group mean for this path tends to the RFE associated with this path). This distribution can be any distribution and might therefore exhibit fat tails. Section I.C provides a statistical framework.

⁶ Finance: *The Journal of Finance*, *Journal of Financial Economics*, and *Review of Financial Studies*. Economics: *American Economic Review*, *Econometrica*, *Journal of Political Economy*, *Quarterly Journal of Economics*, and *Review of Economic Studies*.

associate or full professor, and for the group of PEs, this share is 88%. On a scale from 1 (low) to 10, the average self-ranked score on experience with empirical finance is 8.1 for RTs and 8.4 for PEs. For experience with market liquidity, the average self-ranked score is 6.9 for RTs and 7.8 for PEs.

The evidence on the four overarching questions of interest is as follows. First, the dispersion in estimates across RTs is sizable. All six RT hypotheses had to be tested by proposing a measure and computing the average annual percentage change. The first RT hypothesis, for example, was “Market efficiency has not changed over time.” The median estimate across RTs is -1.1% with an NSE (IQR) of 6.7% (i.e., 6.7 percentage points). The IDR is 27.3% .⁷ The dispersion for the other RT hypotheses is similar in magnitude, albeit smaller for RT hypotheses that arguably involve fewer decisions along the analysis path (e.g., testing for a trend in market share).

Statistical tests show that, for all RT hypotheses, at least a few estimates are significant (at a family level of 0.5%).⁸ This number ranges from 6 (out of 164) for RT-H6 to 125 for RT-H3. We further test the null hypothesis of no dispersion in RFEs. We reject the null for all RT hypotheses. NSEs are therefore statistically significant for all RT hypotheses.

It is worth noting that the uncertainty due to NSEs is similar in magnitude to that due to SEs. For RT-H1, for example, the median SE across RTs is 2.5% . For a Gaussian distribution, this implies an IQR of $1.35 \times 2.5\% = 3.4\%$, which compares to an NSE of 6.7% .

Second, the quantile regressions show that higher quality tends to coincide with smaller NSEs. A one-SD increase in reproducibility significantly reduces NSEs by 25.0% and a one-SD increase in PE rating significantly reduces them by 33.3% . In contrast, a one-SD increase in team quality significantly raises NSEs by 2.8% . This effect, however, is small in economic magnitude. If IDR is used instead of IQR, then a one-SD increase in quality significantly reduces IDRs for all quality measures to 13.3% , 17.9% , and 11.9% , respectively. Overall, higher quality seems to make extreme values less likely.

Third, peer feedback significantly reduces NSEs. The peer-feedback process involves multiple stages. We find that each stage reduces NSEs, albeit insignificantly, while the reduction across all four stages taken together is a significant 47.2% . This number for IDRs is also significant and amounts to an even larger decline of 68.2% .

Fourth, RTs mostly underestimate the dispersion in estimates across RTs, which we test in an incentivized belief survey. Such underestimation might well be why NSEs never attracted much attention until recently.

Finally, we dig deeper to discover what drives dispersion in estimates. A particularly useful tool for such analysis is a multiverse analysis (Liu et al.

⁷ This RT hypothesis further illustrates the importance of robust statistics. One RT reports an estimate of $+74,491\%$. This extreme outlier causes the mean and SD to be 446.3% and $5,817.5\%$, respectively.

⁸ We use the conservative significance levels advocated by Benjamin et al. (2018): 0.5% for significance and 5% for weak significance. They refer to the latter as “suggestive evidence.”

(2021)). For key forks on the analysis path, the multiverse reveals how sensitive the distribution of estimates is to decisions at each particular fork.

It turns out that many of the key forks in #fincap add substantial noise. For RT-H1 on market efficiency, for example, which frequencies teams choose for their variance ratio calculations matters. Some teams compare seconds to minutes, others days to months. A comparison of higher frequencies tends to find a decline in market efficiency, whereas for lower frequencies some find an increase in market efficiency.

The multiverse further reveals that Jensen's inequality can cause large dispersion. If a researcher is interested in assessing an N -period (long-term) trend in X_t , and estimates it based on one-period observations, this could add substantial noise (Blume (1974)). Consider, for example, the expectation of a product of two independent and identically distributed relatives, where a relative is defined as X_t/X_{t-1} . Jensen's inequality implies that the expectation of this product is larger than the product of the expected relatives. The multiverse shows that the noise that this adds can be particularly large for teams that sample at a daily frequency to estimate an annual trend and use relatives instead of, for example, log-differences or a trend-stationary approach.

Contribution to the Literature. The issue of variability in the research process is not new. Leamer (1983), for example, was troubled by the “fumes which leak from our computing centers,” and called for studying “fragility in a much more systematic way.”

Replication studies echo his concern as they typically find much weaker effects and less statistical strength (Ioannidis, 2005, Open Science Collaboration, 2015, Camerer et al., 2016, 2018). This is potentially the result of p -hacking, whereby researchers try analysis paths until insignificant results turn significant.⁹ We caution, however, that poor replication could also be demand- instead of supply driven. This is the case when journals prefer to publish papers with low p -values. Munafò et al. (2017) survey the various threats to credible empirical science and propose several solutions.

The literature on replicability in finance is young but growing rapidly. Examples include McLean and Pontiff (2016), Hou, Xue, and Zhang (2018), Linnainmaa and Roberts (2018), Chordia, Goyal, and Saretto (2020), Harvey and Liu (2020), Ben-David, Franzoni, and Moussawi (2021), Black et al. (2021), Chen (2021), Mitton (2022), Jensen, Kelly, and Pedersen (2023), and Pérignon et al. (2023).¹⁰ None of these replication studies focuses on explaining the dispersion of estimates in a cross-section of researchers, or on the impact of peer feedback. We are the first to run an experiment in which this can be done in a clean way. Our objective is to study dispersion in estimates, short of a potential bias due to p -hacking. By design, there is no need to p -hack for #fincap researchers, because anyone who completed all stages of the project had been guaranteed

⁹ The p -value is the probability of observing an effect that is at least as large as the estimated effect, under the null hypothesis that there is no effect.

¹⁰ Following up on our work, Soebhag, van Vliet, and Verwijmeren (2023) and Walter, Weber, and Weiss (2023) study NSEs in asset pricing as a result of portfolio sort decisions.

coauthorship. Similarly, PEs were guaranteed coauthorship to ensure thoughtful feedback.

We are the first in finance to run an experiment to study dispersion in estimates, but we are not the first in science. Silberzahn et al. (2018) pioneered the multianalyst study by letting multiple teams test whether soccer referees are more likely to draw red cards for players with a darker skin color. Other examples are Botvinik-Nezer et al. (2020) in neuroscience, Huntington-Klein et al. (2021) in economics, and Breznau et al. (2021) and Schweinsberg et al. (2021) in sociology. We innovate relative to these studies by explaining dispersion in estimates with quality attributes, by adding peer feedback stages, and by soliciting beliefs on dispersion *ex ante*. A further strength of our study is the large cross-section of RTs: $N = 164$. This sample is more than twice the size of any of the other multi-analyst samples.

The remainder of the paper is organized as follows. Section I provides an in-depth discussion of the project design.¹¹ It further presents the hypotheses associated with the four overarching questions, and develops an appropriate statistical framework to test them. Section II presents our results. Section III concludes.

I. Project Design and Hypotheses

This section first presents details of the #fincap experiment. It then presents hypotheses based on the four overarching questions. It finishes by discussing an appropriate statistical framework.

A. Project Design

Before starting the #fincap experiment, we filed a pre-analysis plan (PAP) with the Open Science Foundation (OSF; <https://osf.io/h82aj/>). The original version of this paper contains the results of the analysis outlined in the PAP. This original version remains available as Tinbergen Institute Discussion Paper TI 2021-102/IV. Subsequent feedback from various presentations and from reviewers at the *The Journal of Finance* have led to the results presented here. Relative to the PAP, we now use robust methods to cope with unanticipated extreme outliers, we account for multiple testing, and we add a multiverse analysis to deepen insights. Appendix B reconciles the current results with those in the original version.

In a nutshell, the #fincap experiment is about multiple RTs independently testing the same hypotheses on the same sample. We refer to these hypotheses as RT hypotheses and to the sample as the RT sample. This labelling distinguishes the hypotheses that RTs test from the hypotheses that *we* test based on the results generated by RTs and PEs (Section I.B).¹²

¹¹ The design of #fincap follows the guidelines for multianalyst studies proposed by Aczel et al. (2021).

¹² RTs and PEs were recruited mostly by alerting appropriate candidates through suitable channels (e.g., the <https://microstructure.exchange/>). To inform them about #fincap, we created an on-

The RT sample is a plain-vanilla trade sample for the EuroStoxx 50 index futures with the addition of a principal-agent flag.¹³ For each side to a trade (i.e., buy and sell), we therefore know whether the exchange members traded for their own account or for a client. The sample runs from 2002 through 2018 and contains 720 million trade records. These index futures are among the world's most actively traded index derivatives. They give investors exposure to Europe, or, more precisely, to a basket of euro-area blue-chip equities. With the exception of over-the-counter activity, all trading is done through an electronic limit-order book (see, for example, Parlour and Seppi (2008), for details on limit-order book markets).

The RT hypotheses are all statements about annual trends in the following market characteristics (with the null being no change):

- RT-H1: market efficiency,
- RT-H2: realized bid-ask spread,
- RT-H3: share of client volume in total volume,
- RT-H4: realized bid-ask spread on client orders,
- RT-H5: share of market orders in all client orders, and
- RT-H6: gross trading revenue of clients.

The RT hypotheses are presented only briefly here to conserve space. The full presentation of RT-H1, for example, characterizes informationally efficient prices as a random walk. Appendix C motivates and discusses all RT hypotheses in detail. For the purpose of our analysis that follows, we highlight two points. First, the RT hypotheses are chosen to address first-order questions in the field of empirical finance/liquidity. These questions were used to market #fincap and convince appropriate candidates to join the project. Second, we ask for trends expressed as percentage changes to make them invariant to choice of unit (e.g., whether they are expressed in thousands or not).

Note that, by design, there is considerable variation across RT hypotheses in the level of abstraction. RT-H1, for example, corresponds to the relatively abstract notion of market efficiency. In contrast, RT-H3 corresponds to the share of client volume in total volume, which should be relatively straightforward to calculate because, in the RT sample, each buy and sell trade is flagged agent (client) or principal (proprietary).

line repository: <https://fincap.academy>. The repository remains largely unaltered (except for, for example, adding FAQs).

¹³ Trade records contain the following fields: datetime, expiration, buy-sell indicator, size, price, aggressor flag, principal-agent flag, and a full- or partial-execution flag. Note that each side to a trade becomes a record, where the aggressor is the side whose incoming, say, buy order is matched with a resting sell order of the other side. The record is labeled *principal* if the exchange member trades for his own account, and *agent* when he trades for a client. More details on the sample are in Section II of the Internet Appendix. The Internet Appendix is available in the online version of the article on The Journal of Finance website.

RTs are asked to test these RT hypotheses by estimating an average yearly change for a self-proposed measure.¹⁴ They are further asked to report SEs for these estimates. We compute the ratio of the two, which we refer to as the implied *t*-value, or *t*-value for short.

RTs write a short academic paper in which they present and discuss their findings. These papers are evaluated by PEs who were recruited outside the set of researchers who registered as RTs. RT papers were randomly and evenly assigned to PEs in such a way that each paper is evaluated twice, and each PE evaluates nine or 10 papers. PEs score the papers by providing an overall rating as well as a rating per RT hypothesis. They do so in a single-blind process: PEs see the names of RTs, but not vice versa.¹⁵ The reason for single-blind instead of double-blind is to incentivize RTs to exercise maximum effort.

PEs are asked to motivate their scores in a feedback form where they are encouraged to add constructive feedback. RTs receive this feedback unabridged, and are allowed to update their results based on it. Importantly, the design of #fincap was common knowledge to all because it had been available on a dedicated website before registration opened (see footnote 4).

More specifically, #fincap consists of the following four stages:

- Stage 1 (January 11 to March 23, 2021). RTs receive the detailed instructions along with access to the RT sample. They conduct their analysis and submit their results (short paper plus code). We emphasized in our emails and on the project website that RTs should work in *absolute secrecy* so as to ensure independence across RTs.
- Stage 2 (May 10 to May 28, 2021). RTs receive feedback from two anonymous PEs and are allowed to update their analysis based on it. They are asked to report their findings in the same way they did in stage 1.
- Stage 3 (May 31 to June 18, 2021). RTs receive the five best papers based on the average raw PE score. The names of the authors of these five papers were removed before distributing the papers.¹⁶ Similar to stage 2, all RTs are allowed to update their analysis and resubmit their results.
- Stage 4 (June 20 to June 28, 2021). RTs report their final results, this time not constrained by delivering code that produces them. In other words, RTs are allowed to Bayesian update their results (i.e., estimates and SE) using all of the information that has become avail-

¹⁴ RTs are asked to express their results in annualized terms. To some, this was not clear. We therefore notified everyone of the following clarification that we added to the FAQ section on <https://fincap.academy>: “Research teams are asked to report annualized estimates (and the corresponding standard errors); research teams are not required, however, to consider only annualized data.”

¹⁵ In our analysis, we remove PE fixed effects by demeaning (see Section I.B).

¹⁶ If two papers were tied in terms of their average score, then following the PAP, we picked the one that had the highest reproducibility score provided by Cascad. For more information on Cascad, see the statement of H2 in Section I.B.

able to them, in particular the five best papers. They could, for example, echo the results of one of these papers, simply because of an econometric approach that they believe is superior but that is beyond their capacity to code. This stage was added to remove all constraints and see how far the RT community can get in terms of reaching consensus.

The stages subsequent to the first one mimic the feedback researchers get from various interactions with peer researchers in the research process *before* a first journal submission. Stage 2 mimics, for example, immediate feedback from colleagues over lunch, during seminars, or in coffee breaks at conferences. Stage 3 mimics indirect feedback by means of seeing competitive papers that gain a lot of visibility through endorsements by colleagues or that are presented in seminars or at conferences. Stage 4 solicits a final estimate whereby researchers are allowed to attach weight to estimates of others who, for example, they believe implement a superior methodology that they are unable to code themselves. We emphasize that all of these stages are designed in a way to keep the full dynamics of a referee process at a scientific journal out of scope.¹⁷

B. Hypotheses

Before running the experiment, we translated the project's four overarching questions into a set of preregistered hypotheses. These hypotheses all center on the dispersion in estimates across RTs. Our main measure is the IQR, which we refer to as NSE. All hypotheses are stated as null hypotheses and tests are two-sided.

The first three sets of hypotheses focus on how NSEs relate to various quality measures:

H1: The NSE of stage 1 estimates does not covary with team quality.

Team quality is proxied by the largest common factor in various candidate proxies for team quality. We prefer an appropriately weighted average over simply adding all proxies to maximize statistical power in the regressions. More specifically, we define team quality as the first principal component of the following standardized series:¹⁸

- (a) *Top publications*: The team has at least one top-three publication in finance or one top-five publication in economics (0/1) (see footnote 6).
- (b) *Expertise in the field*: Average of self-assessed experience in market liquidity and empirical finance (scale from 0 to 10).

¹⁷ Studying such dynamics would require a different experiment that involves “publishing” papers, including the names of the authors. Note that we do reveal the best five papers (according to PEs) to all RTs in stage 4, but the authors of these papers remain hidden. Our focus is narrowly on the pure findings and beliefs of the RTs, avoiding any possible corruption by “the publication game.”

¹⁸ An important advantage of a principal component analysis (PCA) is that the weighting is data driven, thus avoiding subjective weights. Note that even the five proxies that enter were picked ex ante in the PAP filed at OSF. The PCA results will be discussed in Section II.B.1.

- (c) *Experience with big data*: The team has worked with samples at least as large as the sample they analyze in #fincap (0/1).
 - (d) *Academic seniority*: At least one team member holds an associate or full professorship (0/1).
 - (e) *Team size*: The team size attains its maximum of two members (0/1).
- H2: The NSE of stage 1 estimates does not covary with the reproducibility score.
- This score measures the extent to which RT estimates are reproducible from RT code. The scoring was done by Cascad. Cascad is a non-profit certification agency created by academics with the support of the French National Science Foundation (CNRS) and a consortium of French research institutes. The objective of Cascad is to provide researchers with a way to credibly signal the reproducibility of their research (used by, for example, the *American Economic Review*).¹⁹
- H3: The NSE of stage 1 estimates does covary with the average PE rating (RT hypothesis level).
- To remove a possible PE fixed effect, we use demeaned PE ratings in all of our analysis.

The next hypothesis corresponds to convergence in estimates across the four stages.

- H4: The NSE does not change across all feedback stages.

The final hypothesis focuses on RT *beliefs* about the dispersion in estimates across RTs.

- H5: The average belief of RTs on the dispersion in estimates across RTs, is correct.
- The dispersion predictions were solicited in terms of the SD measure.

C. Statistical Framework

To formalize the analysis of NSEs in a statistical sense, consider a set of researchers indexed by $j \in \{1, \dots, J\}$. All researchers are given the same sample of size K . Researchers are asked to estimate the mean of a particular object (e.g., the per-year change in market efficiency). All researchers independently decide on the optimal analysis path and estimate the mean accordingly. Collectively, let these estimates, $X_1, \dots, X_J$, be distributed as

$$X_j = e_j + \varepsilon_j, \quad (1)$$

¹⁹ Cascad rates reproducibility on a five-category scale: RRR (perfectly reproducible), RR (practically perfect), R (minor discrepancies), D (potentially serious discrepancies), and DD (serious discrepancies). For #fincap, Cascad converted their standard categorical rating to an equal-distance numeric rating: RRR, RR, R, D, and DD become 100, 75, 50, 25, and 0, respectively.

where e_j is a researcher-specific mean, an RFE, and ε_j is a sampling error. The Central Limit Theorem (CLT) implies that, for large K , ε_j is approximately normal with mean zero and variance $\sigma_{j,K}^2 = \sigma_j^2/K$, where σ_j^2 is the path-specific variance of residuals.

Note that sampling errors are likely to correlate across researchers so that, collectively, the estimates are approximately distributed as

$$\underset{(J \times 1)}{X} = \underset{(J \times 1)}{e} + \underset{(J \times 1)}{\varepsilon}, \text{ where } \underset{(J \times 1)}{\varepsilon} \sim N\left(\underset{(J \times 1)}{0}, \underset{(J \times J)}{\Sigma}\right), \quad (2)$$

where Σ is a positive semidefinite matrix. The off-diagonal elements of Σ are expected to be mostly positive since, if for a particular sample draw X_i is above its (unconditional) mean e_i , then X_j is likely also above its mean e_i .²⁰

Nonstandard error. NSE is defined as the IQR in estimates,

$$\text{NSE} := Q_{0.75}(x) - Q_{0.25}(x), \quad (3)$$

where x denotes a realization of the random vector X and $Q_\alpha(x)$ is the α^{th} quantile of x . Note that NSE tends to the IQR of RFEs when J and K both tend to infinity:

$$\text{NSE} \xrightarrow{J,K \rightarrow \infty} Q_{0.75}(e) - Q_{0.25}(e). \quad (4)$$

We reiterate that the distribution of RFEs (i.e., the distribution of e) could be any distribution. It is therefore prudent to pick a robust dispersion measure, which is why we use IQR instead of SD. The latter tends to get dominated by the size of extreme outliers.²¹

Testing for NSE. We test for significance of NSEs by testing whether there is any dispersion in RFEs. We do so by testing the null hypotheses

$$H_0 : e_j = \nu, \quad \forall j \in \{1, \dots, J\}, \quad (5)$$

where ν is the median RFE. Since X_j is an estimator of e_j , these hypotheses can be tested by conducting multiple tests, one for each X_j where $j \in \{1, \dots, J\}$. Each individual test is a two-sided one that verifies whether X_j is statistically different from ν . (Note that the critical values used in the test need to account for MHT, which we discuss below.) In the implementation, we set ν equal to the median estimate. If in at least one of these tests the null is rejected, then dispersion is nonzero and we consider NSEs to be statistically significant.²²

²⁰ For example, consider the case of estimating the mean of a distribution. If two researchers estimate this mean by taking the sample average, but one winsorizes the sample and the other does not, then a particular sample draw with unusually high values likely yields above-mean estimates for both researchers.

²¹ #fincap is a case in point. For RT-H4, one team reports an estimate of $-6,275,383\%$, whereas the estimates of other teams range from $-2,897\%$ to 870% . The SD based on all estimates is $490,024\%$, but it is only 245% if one omits the outlier.

²² Two more technical points merit discussion. First, we prefer the median over the mean to have a robust location parameter. The asymptotic variance of the mean is smaller than that of the

Conceptually, the distribution of X could be obtained by bootstrapping. Such a procedure is infeasible, however, because it requires that researchers rerun their analysis for every new draw of the sample. Instead, we use MHT results to make the testing procedure feasible.

Before turning to MHT, we pause for a moment and take stock of what is available to us. The #fincap sample consists of estimates x_j , along with their SE s_j . This is useful, but misses information on the covariance among all possible pairs of estimates across researchers.

To account for multiple testing, we rely on well-developed statistical theory. If one aims to test at a level of 5% for a family of N tests, then individual tests should be performed with a $(5/N)\%$ critical value, if the test statistics are mutually independent (Bonferroni (1936), Šidák (1967), Harvey, Liu, and Zhu (2016)).²³

In summary, we propose an NSE test where the null hypothesis is that there is no dispersion in RFEs. We use a Bonferroni adjustment of significance levels to account for multiple testing. The test is conservative, because Bonferroni assumes independence. As pointed out in footnote 20, estimates are likely to correlate across researchers, in which case, the *effective* number of tests is likely to be smaller than the actual number of tests. In the implementation, we add a trivial extension in which correlations between estimates are calibrated based on our multiverse analysis (Section II.C). We close the section by discussing an alternative test and pointing out a caveat.

Alternative test. Note that a natural alternative to the proposed test is to simply test whether IQR is statistically different from zero. We did not use this shortcut because our focus is on whether there is any dispersion at all in estimates across researchers. Although we use IQR to express dispersion in a single number, the underlying interest is whether the distribution in estimates is nondegenerate.

Caveat. We point out a potential caveat. The procedure to obtain a conservative test on RFEs implicitly assumes that SEs reported by researchers are consistent estimators of the true SEs. This may not be true if (some) researchers

median for Gaussian distributions, but typically not for distributions with fat tails. The reason is that the former depends on variance and thus on extreme outliers, whereas the latter does not: σ^2/N and $1/(4Nf(m))$, respectively, where N is the sample size, f is the density function, and m is the median. Figure IA.I in the Internet Appendix shows that, in #fincap, the variance of the median is an order of magnitude smaller than the variance of the mean. Second, the proposed test assumes that sampling error is negligible for the median estimate as an estimator for the median RFE, because randomness in the median estimate is ignored. Figure IA.I illustrates that the variance of the median estimate is indeed negligible for #fincap. (If it were not negligible, then one could resort to bootstrapping to establish critical values.)

²³ If the N test statistics are independent, then the probability of at least one significant result is $(1 - (1 - \alpha)^N)$. For example, for $\alpha = 0.05$ and $N = 10$, this probability is 40%. Šidák (1967) proposes that the significance level of the individual tests be adjusted to $\alpha' = 1 - (1 - \alpha)^{1/N}$. A Taylor expansion of α' around zero yields $\alpha' \approx \alpha/N$, which is known as the Bonferroni correction (Bonferroni (1936)).

report nonrobust SEs. Nonrobust SEs tend to be smaller because they ignore commonalities. To the extent that this holds, NSE tests tend to turn significant more often. NSEs themselves, however, remain consistent estimators.²⁴

II. Results

This section presents our findings. The results are based on a balanced sample of 164 out of 168 RTs that completed all stages of the project. The first subsection presents various summary statistics and tests whether NSEs are statistically significant. The second subsection tests our hypotheses. The third subsection digs deeper by means of a multiverse analysis. The fourth and final subsection discusses alternative explanations.

A. Summary Statistics

Table I summarizes our stage 1 sample by means of three sets of statistics, organized in three panels.²⁵ Panel A summarizes the characteristics of the #fincap community. It consists of 164 RTs and 34 PEs. Maximum RT size is two members, which is the size of 79% of RTs.

The statistics testify to the high quality of the #fincap community. The percent of RTs that have at least one publication in a top finance or economics journal is 31% (see footnote 6 for the list of journals). For PEs, this figure is 85%. The percent of RTs who have at least one member who is tenured at the associate or full professor level is 52%. For PEs, this figure is 88%. Feedback seems to come from more established scholars, which likely mirrors reality.

RT members and PEs cover the global academic finance community reasonably well (see Figure IA.II in the Internet Appendix). RT members reside in 34 countries, with most residing in the United States (51 out of 293). PEs reside in 13 countries with, again, most residing in the United States (13 out of 34). The strong skew toward the United States is not surprising given that the more senior, well-published finance scholars are predominantly affiliated with U.S. universities.

Most RTs and PEs appear to have the appropriate background for testing the RT hypotheses on the RT sample. Their average self-reported scores on having experience in the field of empirical finance is 8.1 for RTs and 8.4 for PEs on a scale from 0 (low) to 10. For experience with market liquidity, these average scores are 6.9 for RTs and 7.8 for PEs. There is considerable variation around these averages as the SDs range from 1.7 to 2.4. When it comes to working with samples as large as the RT sample (720 million trade records) most RTs and PEs seem up to it as 65% of RTs have worked with samples at

²⁴ Unfortunately, we do not have precise information on the SEs reported in #fincap, because not all RTs provide detailed information on how they calculate SEs.

²⁵ Table IA.I through IA.III in the Internet Appendix repeat Panel C of Table I for the other stages. Panel A is the same for all stages, and Panel B is available only for stage 1 results, since only these results are evaluated by peers and scored by Cascad on reproducibility.

Table I
Summary Statistics

This table presents summary statistics. Standard deviations are in parentheses.

Panel A: Quality of the #fincap community						
	Research Teams (RTs)		Peer Evaluators (PEs)			
Fraction with top finance/econ publications (see footnote 6)	0.31		0.85			
Fraction including at least associate/full professor	0.52		0.88			
Experience empirical-finance research (low-high, 1–10)	8.1 (1.7)		8.4 (1.8)			
Experience market-liquidity research (low-high, 1–10)	6.9 (2.4)		7.8 (2.3)			
Relevant experience (average of the above two items)	7.5 (1.3)		8.1 (1.7)			
Fraction with “big data” experience (>#fincap sample)	0.65		0.88			
Fraction teams consisting of two members (maximum team size)	0.79					
Number of observations	164		34			
Panel B: Quality of the analysis of RTs						
	RTs					
Reproducibility score according to Cascad (low-high, 0-100)	64.5 (43.7)					
Paper quality as judged by PEs (low-high, 0-10)	6.2 (2.0)					
Panel C: Dispersion across teams of stage 1 results: Estimates, SEs, and <i>t</i> -values						
	RT-H1 Efficiency	RT-H2 RSpread	RT-H3 Client Volume	RT-H4 Client RSpread	RT-H5 Client- MOrders	RT-H6 Client GTR
Estimate (yearly change, %)						
Mean	446.3	−1,093.4	−3.5	−38,276.1	−3.5	−87.1
SD	5,817.5	14,537.2	9.4	490,024.2	37.6	728.5
Min	−171.1	−186,074.5	−117.5	−6,275,383.0	−452.9	−8,254.5
Q(0.10)	−23.7	−6.9	−3.8	−6.7	−1.6	−192.1
Q(0.25)	−6.2	−3.6	−3.5	−2.1	−0.6	−18.2
Median	−1.1	−0.0	−3.3	0.1	−0.0	0.0
Q(0.75)	0.5	3.9	−2.4	3.8	0.2	3.2
Q(0.90)	3.7	21.5	−0.1	20.4	1.0	56.5
IQR (i.e., NSE)	6.7	7.5	1.2	5.9	0.8	21.4
IDR	27.3	28.4	3.7	27.1	2.5	248.5
Max	74,491.1	4,124.0	8.7	870.2	69.5	1,119.0
Standard error						
Mean	468.7	1,195.3	3.7	38,302.0	6.2	148.2
SD	5,810.6	14,711.9	29.5	489,929.5	40.1	526.0
Min	0.0	0.0	0.0	0.0	0.0	0.0
Q(0.10)	0.1	0.2	0.1	0.2	0.1	0.0
Q(0.25)	0.5	1.1	0.3	1.2	0.2	0.7
Median	2.5	5.0	1.4	4.4	1.0	9.7
Q(0.75)	9.3	13.9	2.0	14.3	2.4	77.1
Q(0.90)	44.7	39.6	2.2	31.2	3.1	235.4
IQR	8.8	12.8	1.7	13.1	2.2	76.4
IDR	44.6	39.4	2.1	31.0	3.1	235.4
Max	74,425.5	188,404.1	378.8	6,274,203.0	463.7	4,836.2

(Continued)

Table I—Continued

Panel C: Dispersion across teams of stage 1 results: Estimates, SEs, and *t*-values

	RT-H1 Efficiency	RT-H2 RSpread	RT-H3 Client Volume	RT-H4 Client RSpread	RT-H5 Client- MOrders	RT-H6 Client GTR
<i>t</i> -value						
Mean	−3.6	35.3	−47.1	24.3	−5.7	−2.0
<i>SD</i>	28.4	541.2	269.9	406.0	60.1	21.2
Min	−322.3	−764.6	−2,770.6	−852.6	−631.6	−191.7
Q(0.10)	−4.7	−5.7	−37.4	−3.5	−2.3	−1.7
Q(0.25)	−1.9	−1.5	−11.5	−1.0	−0.6	−1.0
Median	−0.7	−0.1	−1.8	0.1	0.0	0.0
Q(0.75)	0.3	0.8	−1.6	1.0	0.8	0.7
Q(0.90)	1.7	1.5	−0.3	1.6	1.7	1.2
IQR	2.2	2.3	9.9	1.9	1.3	1.7
IDR	6.4	7.2	37.1	5.2	3.9	2.9
Max	51.6	6,880.5	29.5	5,119.5	89.6	100.6

least as large and 88% of PEs. In sum, we believe that the group of #fincap participants is reasonably representative of the academic community in empirical finance/liquidity.

Panel B of Table I shows that the average quality of the RT analysis is solid, but the dispersion is large. The average reproducibility score is 64.5 on a scale from 0 (low) to 100 (see footnote 19). This is high when benchmarked against other studies on reproducibility (Colliard, Hurlin, and Pérignon (2021)). The accompanying SD is 43.7, which shows that there is large variation across RTs: Most code reproduces either close to perfectly or barely at all. The paper quality scores provided by PEs show a similar pattern, albeit with less dispersion. The average score across RTs is 6.2 on a scale from 0 (low) to 10, with an SD of 2.0.

Panel C provides descriptive statistics on the distribution of results across RTs, both by RT hypothesis and by type of result (estimate, SE, and *t*-value). Since our focus is on dispersion in estimates across RTs, we relegate discussion of RT medians to Appendix C. More specifically, Appendix C discusses the RT hypotheses in depth and summarizes what RTs as a group seem to find, with a focus on the across-RT median instead of the across-RT IQR (i.e., the NSE).

Perhaps the most salient feature of the evidence in Panel C is that there is substantial variation across RTs for both all RT hypotheses and all types of results. Panel A in Figure 1 illustrates this result for estimates. For RT-H1 on market efficiency, for example, the median estimate across RTs is −1.1% with an IQR of 6.7%. Even for RT-H3, which is a seemingly straightforward calculation of market share, the dispersion is sizable: The IQR is 1.2% with a median of −3.3%. The variation in IQR across RT hypotheses suggests that the more abstract a hypothesis is, the larger the IQR is. Finally, the figure

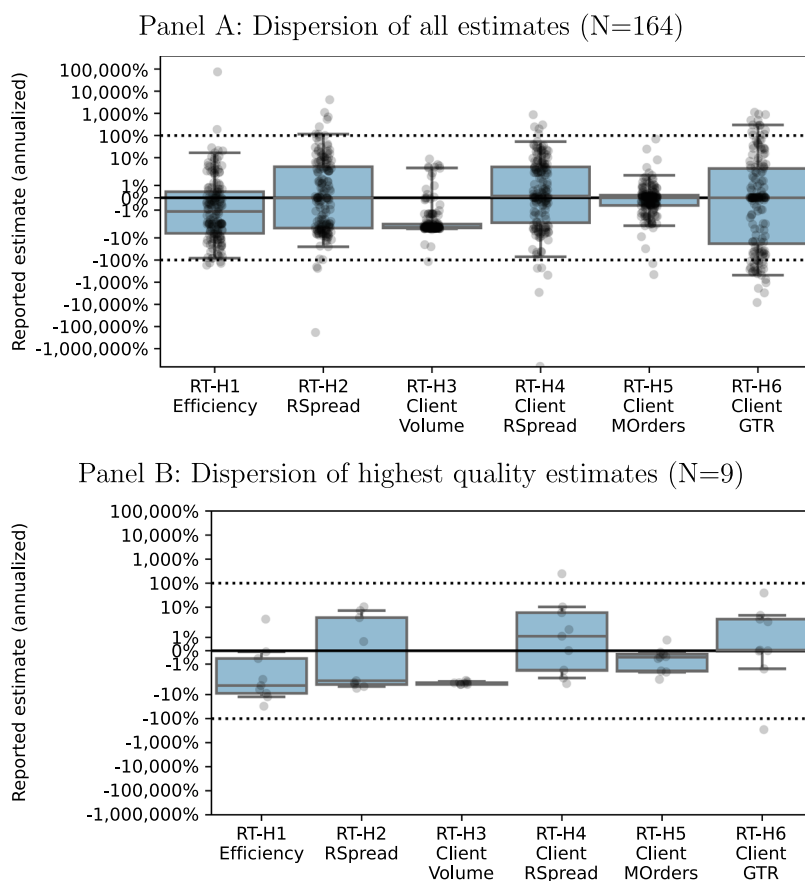


Figure 1. Dispersion of stage 1 estimates across research teams (RTs). This plot illustrates the dispersion of stage 1 estimates across RTs. These estimates all focus on a trend in the sample, expressed in terms of a yearly percentage change. The six box plots correspond to the six trends that the RTs were asked to estimate. The boxes depict the first and third quartiles. The horizontal line in the box corresponds to the median. The whiskers depict the 2.5% and 97.5% quantiles. All estimates are also plotted individually as gray dots. (Color figure can be viewed at [wileyonlinelibrary.com](https://onlinelibrary.wiley.com/terms-and-conditions))

illustrates that there are extreme outliers for all RT hypotheses, which motivates our analysis in terms of robust statistics.

NSE test results. Is the dispersion in estimates statistically significant? Table II presents the NSE test results. The null of no dispersion in RFE is rejected for all RT hypotheses at the 0.5% (family) significance level. The conservative Bonferroni adjustment in Panel A yields at least 11 estimates that are individually significantly different from the median (RT-H6), and at most 38 significant differences (RT-H3). There are significant estimates both above and below the median for all RT hypotheses.

Table II
Nonstandard Error Test

This table tests for the presence of NSEs in stage 1. It does so by testing whether estimates provided by researchers deviate from the median across researchers. Critical values of the individual tests are raised to achieve the desired significance level at the family of tests. The number of significantly negative tests and significantly positive tests is reported in brackets. The reported family p -value is the probability that out of all test statistics, at least one is larger than the reported value, under the null of a multivariate normal with means equal to the realized #fincap medians, and a covariance matrix with squared SEs (reported by the RTs) on the diagonal and off-diagonals that are either zero (Bonferroni) or based on the multiverse analysis (Section II.C).

Panel A: Multiple tests (Bonferroni)

	Reject No-NSE at 0.5%?	p -Value of Family Test	Mean (SD) Correlation Test Statistics	Effective Number of Tests
RT-H1	Yes (8, 25)	< 0.0001	0.00 (0.00)	164
RT-H2	Yes (24, 10)	< 0.0001	0.00 (0.00)	164
RT-H3	Yes (13, 25)	< 0.0001	0.00 (0.00)	164
RT-H4	Yes (22, 4)	< 0.0001	0.00 (0.00)	164
RT-H5	Yes (13, 10)	< 0.0001	0.00 (0.00)	164
RT-H6	Yes (8, 3)	< 0.0001	0.00 (0.00)	164

Panel B: Multiple tests (based on multiverse analysis)

	Reject No-NSE at 0.5%?	p -Value of Family Test	Mean (SD) Correlation Test Statistics	Effective Number of Tests
RT-H1	Yes (8, 26)	< 0.0001	0.03 (0.21)	77
RT-H2	Yes (24, 10)	< 0.0001	0.05 (0.22)	81
RT-H3	Yes (13, 26)	< 0.0001	0.22 (0.34)	21
RT-H4	Yes (22, 4)	< 0.0001	0.08 (0.24)	67
RT-H5	Yes (13, 10)	< 0.0001	0.20 (0.34)	31
RT-H6	Yes (8, 3)	< 0.0001	0.02 (0.21)	86

If, instead of assuming zero correlation across test statistics as in Bonferroni, one calibrates them based on bootstrapping from the multiverse analysis (Section II.C), results change as shown in Panel B. The implied “effective” number of tests is much lower than the 164 tests used in Bonferroni—it ranges from 21 (RT-H3) to 86 (RT-H6). The factor by which significance levels are adjusted is therefore up to almost seven times smaller than what Bonferroni suggests (i.e., $164/24 = 6.8$). The result is that more differences, indeed, become significant. The increases are moderate though, with at most two more differences becoming significant.

In sum, the statistics presented thus far show that there is substantial dispersion across RTs, in terms of not only their estimates, but also their team quality, reproducibility score, and in PE rating. In the next subsection, we use this dispersion to test the first three hypotheses, that is, we ask: is there more

dispersion in estimates for lower quality teams, for results that are harder to reproduce, or for lower quality papers?

B. Hypothesis Tests

The results on the three sets of hypotheses are discussed in the following three subsections. SEs in the quantile regressions account for correlation in residuals by adding RT hypothesis fixed effects and by clustering per RT across all stages.

B.1. Covariates for Stage 1 Dispersion (H1–H3)

The first set of hypotheses relates NSEs to various quality variables. One of these is team quality, which we measure using the first principal component (PC1) of five standardized quality proxies (see H1 in Section I.B). PC1 explains 38.3% of total variance. Moreover, it loads positively on all quality proxies—it loads strongest on publications and weakest on big data, but, importantly, it loads positively on all of them. Table I.A.IV in the Internet Appendix provides detailed results on the PCA.

Table III summarizes results of the stage 1 quantile regressions, where the dependent variables are the 10th, 25th, 50th, 75th, and 90th percentiles of the distribution in estimates across RTs. Figure 2 depicts the results by showing how a one-SD increase in each covariate affects IQR (i.e., NSE) and IDR. Taken together, these results allow us to test the first three hypotheses that relate quality variables to dispersion in estimates.

First, we find that higher team quality coincides with somewhat larger IQR but with smaller IDR. The effect of team quality on the 25th percentile is not significant except for the 75th percentile, for which it is significantly positive. The economic magnitude is small, however, as can be seen in Figure 2. A one-SD increase in quality raises IQR by only $(0.032 - 0.004) \times 7.2 = 0.2$ percentage points (pps), where 7.2 is the average IQR across hypotheses (see Panel C of Table I). This increase of 0.2 pps implies a relative increase of 2.8%.²⁶ In contrast, a one-SD increase in team quality *reduces* IDR by 6.7 pps (−11.9%, since average IDR is 56.3). This result is a significant increase in the first decile and a significant reduction in the ninth decile. These findings suggest that higher quality teams are less likely to report extreme estimates.

If one replaces team quality by the five quality variables on which it is based, a more nuanced picture emerges (Table I.A.V in the Internet Appendix). Among the statistically significant and sizable relationships, we find that a one-SD increase in academic seniority (i.e., an associate/full professor in the team) reduces IQR by 1.4 pps (−19.4%) and a one-SD increase in team size reduces

²⁶ A direct test on IQR, instead of separate tests on the 25th and 75th percentiles, requires jointly modeling these percentiles. Such multivariate modeling, combined with clustering on errors, is a nontrivial econometric challenge. Univariate modeling with clustering, in contrast, is relatively standard. We use a python package to run these regressions: pyqreg.

Table III
Stage 1 Quantile Regressions

This table presents results of quantile regressions that characterize how the distribution of stage 1 estimates covaries with various quality metrics. These metrics are team quality, reproducibility score, and (demeaned) peer rating. The three quality variables have been standardized and, subsequently, multiplied by the IQR for each RT hypothesis. The coefficients therefore measure the result of a one-standard-deviation change, expressed in terms of IQR units. * and ** correspond to significance at the 5% and 0.5% level, respectively.

	Q(0.10)	Q(0.25)	Q(0.50)	Q(0.75)	Q(0.90)
Team quality (standardized/scaled)	0.597** (0.030)	0.004 (0.014)	0.002 (0.007)	0.032** (0.012)	-0.325** (0.030)
Reproducibility score (standardized/scaled)	0.473** (0.033)	0.109** (0.014)	-0.001 (0.007)	-0.142** (0.011)	-0.555** (0.028)
Average rating (standardized/scaled)	0.766** (0.034)	0.230** (0.014)	-0.001 (0.007)	-0.097** (0.011)	-0.626** (0.028)
Dummy RT-H1 Efficiency	-29.592** (0.813)	-6.099** (0.340)	-1.132** (0.166)	0.939** (0.269)	9.057** (0.708)
Dummy RT-H2 RSpread	-15.933** (0.849)	-3.930** (0.342)	-0.017 (0.166)	3.674** (0.268)	22.451** (0.705)
Dummy RT-H3 Client Volume	-5.629** (0.836)	-3.789** (0.339)	-3.319** (0.166)	-2.386** (0.268)	0.221 (0.721)
Dummy RT-H4 Client RSpread	-12.089** (0.837)	-2.437** (0.340)	0.162 (0.166)	4.161** (0.266)	19.619** (0.704)
Dummy RT-H5 Client MOrders	-2.479** (0.837)	-0.744* (0.339)	-0.001 (0.166)	0.297 (0.268)	1.625* (0.721)
Dummy RT-H6 GTR	-194.457** (0.806)	-21.385** (0.337)	0.022 (0.167)	5.137** (0.268)	65.203** (0.679)
#Observations	984	984	984	984	984

IQR by 0.9 pps (-12.5%). A one-SD increase in top publications, however, *increases* IQR by 1.9 pps (+26.4%). These three variables are positively correlated, which explains why we find that the (aggregate) team variable has a relatively small effect on IQR. For IDR, the effects are of the same sign but larger in magnitude, at -19.4, -7.0, and +6.1 pps respectively (-34.4%, -12.4%, and +10.8%). Note that now the negative effects really dominate, which explains why IDR covaries negatively with team quality. In sum, these findings suggest that well-published scholars seem to disagree more, but this effect is offset by the presence of a senior scholar or a second team member.²⁷

Second, all percentiles covary significantly with reproducibility except for the median. The 10th and 25th percentiles covary positively and the 75th and the 90th percentiles covary negatively. Figure 2 shows that these changes are sizable. A one-SD increase in reproducibility reduces IQR by 1.8 pps (-25.0%)

²⁷ The finding of more disagreement among well-published scholars deserves further study. These scholars might excel simply because they are extraordinarily creative and therefore more likely to see new and idiosyncratic analysis paths that they believe are most appropriate, or they might delegate more of the analysis to (junior) research assistants resulting in more variation in analysis paths. Related, Harvey (2017, p. 1414) discusses delegation in the context of *p*-hacking.

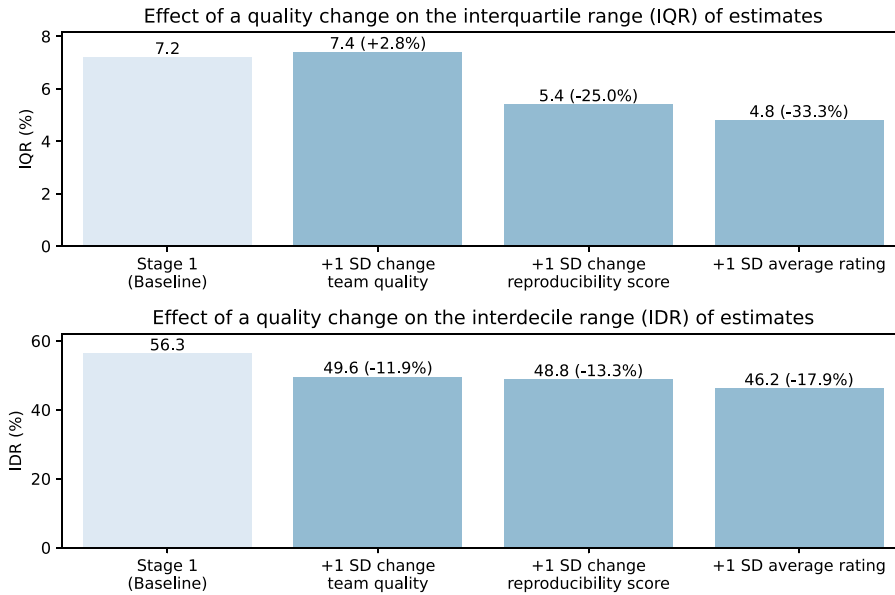


Figure 2. Dispersion in estimates related to quality measures. This figure plots how the dispersion in stage 1 estimates covaries with various quality measures. The top graph uses the interquartile range (IQR) as a dispersion measure and the bottom graph uses the interdecile range (IDR). The quality variables are team quality, reproducibility score, and the rating by PEs. The IQR and IDR estimates are taken from Table III, where relative changes are averaged across RT hypotheses. The baseline level is the average dispersion across RT hypotheses. (Color figure can be viewed at [wileyonlinelibrary.com](https://onlinelibrary.wiley.com/doi/10.1111/j.1469-7610.2020.02501.x))

and IDR by 7.5 pps (-13.3%). In sum, better reproducibility lowers overall dispersion.

Third, the results for paper quality mirror those of reproducibility, albeit are a bit stronger in magnitude. The 10th and 25th percentiles covary significantly positively, and the 75th and 90th percentiles covary significantly negatively. A one-SD increase in paper quality reduces IQR by 2.4 pps (-33.3%) and IDR by 13.6 pps (-17.9%). Higher rated papers exhibit less dispersion in estimates.

In summary, the evidence on the first three hypotheses is such that the null of no covariation is rejected for all three. Generally, higher quality is associated with less dispersion in estimates.

B.2. Convergence across Stages? (H4)

The analysis of first-stage results shows that dispersion in estimates is sizable and statistically significant. We next ask whether peer feedback creates convergence. In particular, we examine whether dispersion in estimates declines in the three subsequent stages in which teams get feedback from peers. This is the focus of the fourth hypothesis.

Table IV presents results of quantile regressions to explain the dispersion in estimates in all four stages (thus far, only stage 1 has been analyzed). To

Table IV
All-Stages Quantile Regressions

This table presents results of quantile regressions that characterize how the distribution of estimates varies across all stages of the #fincap project. The stage dummies have been multiplied by the (stage 1) IQR for each RT hypothesis. The coefficients therefore measure the effect in terms of IQR units. Standard errors account for correlation in residuals by adding RT hypothesis fixed effects and by clustering by RT across all stages. * and ** correspond to significance at the 5% and 0.5% level, respectively.

	Q(0.10)	Q(0.25)	Q(0.50)	Q(0.75)	Q(0.90)
Dummy Stage 2 - Dummy Stage 1	2.44* (1.18)	0.07 (0.14)	−0.00 (0.01)	−0.06 (0.06)	−0.73 (0.64)
Dummy Stage 3 - Dummy Stage 2	0.94* (0.41)	0.15 (0.09)	0.00 (0.01)	−0.09 (0.05)	−0.73 (0.40)
Dummy Stage 4 - Dummy Stage 3	0.21* (0.09)	0.06* (0.03)	0.00 (0.01)	−0.04 (0.03)	−0.25* (0.11)
Dummy Stage 4 - Dummy Stage 1	3.59** (1.23)	0.28* (0.14)	−0.00 (0.01)	−0.19** (0.05)	−1.71** (0.50)
RT-hypotheses dummies	Yes	Yes	Yes	Yes	Yes
#Observations	3,936	3,936	3,936	3,936	3,936

account for heterogeneity in dispersion across RT hypotheses, the explanatory variables are stage dummies that are multiplied by stage 1 (estimate) IQR separately for each hypothesis RT hypothesis. The coefficients therefore measure a stage effect, expressed in IQR units.

The evidence rejects the null hypothesis of no convergence across all stages. All changes across consecutive stages are positive for the 10th and 25th percentiles, and negative for the 75th and 90th percentiles. The majority, however, are insignificant. In contrast, the *total* change across stages is significant for all of the percentiles considered at the 5% level and for all but one at the 0.5% level. Taken together, these results show that there is significant convergence from the first to the last stage, but a decomposition across the various stages lacks significance.

Figure 3 illustrates that the convergence is sizable. Panel A shows that the total decline in IQR is 3.4 pps (−47.2%). The decline appears evenly distributed across the stages, although this decomposition is mostly insignificant. Panel B shows that the total decline in IDR is even larger, at 38.4 pps (−68.2%). More than half of the effect appears to occur from the first to the second stage, where RTs receive anonymized feedback from two PEs. However, this result is only weakly significant, since only the increase in the first decile is weakly significant (i.e., at the 5% level, not at the 0.5% level).

B.3. Are RT Beliefs on Dispersion in Estimates Accurate? (H5)

The fifth and final hypothesis focuses on whether RTs are aware of the dispersion in estimates across teams. Beliefs have been solicited in an incentivized way. All teams were asked to predict SDs in estimates across

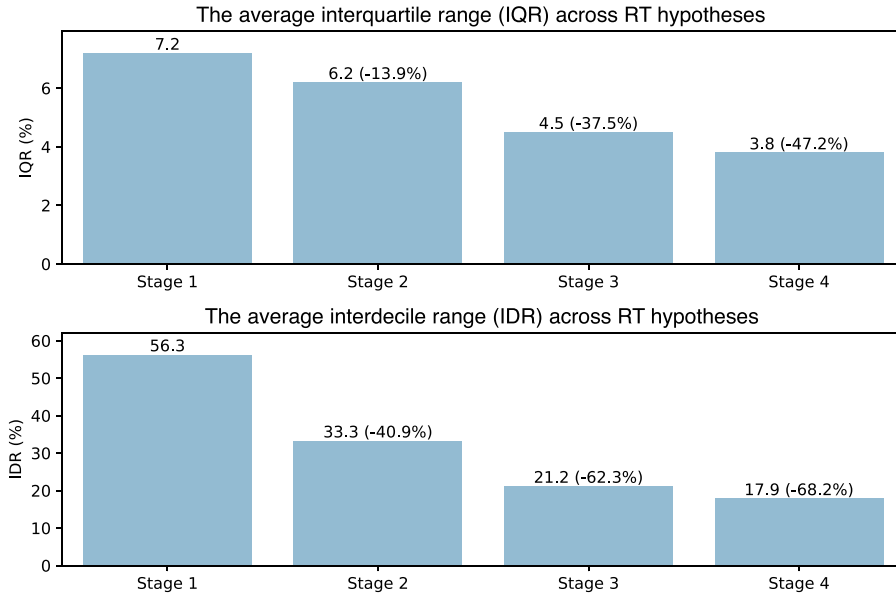


Figure 3. Dispersion in estimates related to feedback stages. This figure plots how the dispersion in estimates changes across feedback stages. Stage 1 is the baseline stage, which is the stage before any feedback. The top graph uses the interquartile range (IQR) as a dispersion measure, whereas the bottom graph uses the interdecile range (IDR). The IQR and IDR values are based on the estimates in Table IV, where relative changes are averaged across all RT hypotheses. The baseline level is the average dispersion in stage 1 estimates across RT hypotheses. (Color figure can be viewed at [wileyonlinelibrary.com](https://onlinelibrary.wiley.com/doi/10.1111/j.1468-2020.11406.2020))

teams.²⁸ We randomly selected 20% of all RTs and paid each of them \$300 if one of their predictions (randomly drawn) was within 50% of the realized SD. Details on the reward scheme are in the instruction sheet they received before reporting their beliefs (see Section III.I in the [Internet Appendix](#)). The hypothesis pertains to stage 1 estimates, because beliefs are solicited for this stage only.

As H5 is stated in terms of the average belief being correct, testing it requires a test on the equality of means: the mean belief about SDs in estimates across teams, and the SDs of these estimates in the population. We measure the relative distance between beliefs and realizations using the test statistic D ,

$$D = \frac{1}{6n} \sum_{i,j} \left(\frac{BeliefOnSD_{ij} - RealizationOfSD_j}{RealizationOfSD_j} \right), \quad (6)$$

²⁸ In retrospect, we should have (also) asked for an IQR prediction, because SD is highly sensitive to extreme outliers (see footnote 21). To assess whether RTs might have overlooked such outliers, we compare their SD predictions with realized SDs, both on the full sample and on a trimmed sample.

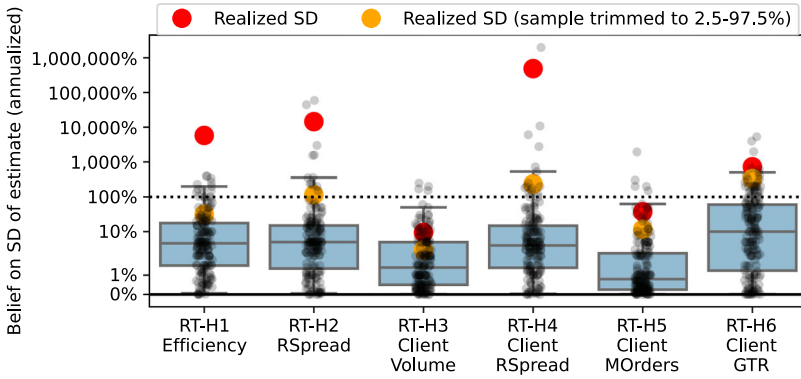


Figure 4. Research team beliefs on dispersion stage 1 estimates. This plot illustrates the dispersion in beliefs across RTs, for all six RT hypotheses. All teams were asked to predict the SD in estimates across all RTs. The boxes depict the first and third quartiles. The horizontal line in the box corresponds to the median. The whiskers depict the 2.5% and 97.5% quantiles. All estimates are also plotted individually as gray dots. The red dots show the realized SD in estimates across RTs. The orange dots do the same, but are based on a 2.5% to 97.5% trimmed sample. (Color figure can be viewed at wileyonlinelibrary.com)

where $BeliefOnSD_{ij}$ is the belief of team i on the SD in estimates across teams for RT hypothesis j , and $RealizationOfSD_j$ is the realized SD for this RT hypothesis in the raw sample.²⁹ The distribution of D under the null of equal means is obtained by bootstrapping. For details on the bootstrap procedure, we refer the reader to Appendix E.

Figure 4 plots the distribution of beliefs on SDs, along with realized SDs depicted by red dots. It illustrates that the vast majority of teams underestimate dispersion in estimates. The IQR denoted by the boxes is consistently below the red dot, which implies that at least 75% of the teams underestimate the dispersion.

One might think that teams simply overlook the extreme values that make realized SDs explode. This does not appear to be the case, however, because even if one trims the estimates by removing the top and bottom 2.5%, the IQR box stays below these “trimmed” realized SDs, depicted by orange dots. The only exception is RT-H3, for which the orange dot is just within the top of the box.

The formal test results are in Table IA.VI of the Internet Appendix. Pooling across all RT hypotheses, the test statistic shows that the predicted SD is 71.7% below the realized SD. This underestimation is significant at the 0.5% level. Similar results hold for all RT hypotheses individually except RT-H3,

²⁹ The benefit of a relative measure as opposed to an absolute one is that (i) it is easy to interpret as it allows for statements of RTs over- or underestimating by some percentage and (ii) it accounts for level differences across hypotheses (e.g., under the null of accurate beliefs, a uniform distribution of beliefs on the support 0.09 to 0.11 will exhibit the same dispersion as a uniform distribution of beliefs on 900 to 1100).

for which the underestimation is insignificant. Its value was also lowest, with only 9.0% underestimation. RT-H3 is an hypothesis on market shares that is arguably relatively straightforward to test. In summary, the vast majority of tests show significant underestimation and we therefore firmly reject the null that beliefs on the dispersion in estimates are accurate.

C. Digging Deeper: A Multiverse Analysis

NSEs in #fincap are significant and sizable. Why? Can we somehow identify which forks on the analysis path cause most of the dispersion? More specifically, can we rank key forks on the path according to the degree of refraction they cause in the light the sample sheds on the research question at hand? We employ a multiverse analysis to address these questions.

Steege et al. (2016) coined the term multiverse analysis to emphasize that data construction involves multiple decisions. The sample that enters the analysis, therefore, is a function of the set of reasonable choices, with the sample becoming a (p. 702) “many worlds or *multiverse* of data sets.” A particular result of an analysis then becomes a distribution of results (because samples vary). We generalize this approach by adding decision forks for the part of the analysis that follows the sample construction (e.g., the choice of econometric model).

The strength of a multiverse analysis is that it reveals how sensitive an estimate is to a particular fork on the analysis path. It does so by studying how much the estimates refract when varying across all reasonable alternatives at the fork. For example, let there be N reasonable analysis paths. Now suppose there are $k \leq N$ reasonable alternatives at the j^{th} fork. Then, the N estimates associated with the N paths are sorted into k sets, depending on the alternative selected at the fork. The degree to which the results differ across the k sets determines how sensitive results are to the j^{th} fork. We measure the degree to which k distributions differ by a k -sample Anderson-Darling (AD) test. Appendix F discusses the AD test in detail, including why it fits our application particularly well. AD is a standard option in the Boba software that we use (Liu et al. (2021)).³⁰

To make the multiverse feasible, we identify key forks on the analysis path and, for each fork, we ask RTs to select the alternative they picked among a set of predefined alternatives. This is done by means of a questionnaire that all participants filled out after the experiment. The choice of forks and the alternatives at each fork is informed by the short papers RTs wrote for #fincap. The discretization of the decision space enables us to project the large space of realized analysis paths onto a manageable space of “representative” paths. Table V provides an overview of all forks for the six RT hypotheses. It lists the alternatives at each fork, along with the fraction of RTs that picked them (depicted in Figure IA.V of the Internet Appendix).

³⁰ The Boba software is available at <https://github.com/uwdata/boba>.

Table V
Analysis Paths

This table summarizes all analysis paths by spelling out all forks and all alternatives at these forks. It further presents the empirical distribution of decisions at all forks.

RT-Hypothesis	Fork	Fork Description	Alternatives	Frequency
All	1	Remove open/close	No	79%
All	2	Days excluded	Yes, 30 minutes	21%
All	3	Outlier treatment	None	81%
			Settlement weeks	19%
			None	65%
			Winsorize measure at 2.5 and 97.5 percentile ^a	20%
			Trim measure at 2.5 and 97.5 percentile ^a	14%
All	4	Frequency analysis	Daily	37%
			Weekly	1%
			Monthly	21%
			Annual	41%
All	5	Model	Trend stationary (regression with linear trend)	35%
			Log difference (trivial regression, i.e., intercept only)	5%
			Relative difference (trivial regression)	60%
1	6	Measure	Variance ratio (low frequency in numerator)	63%
			Autocorrelation (R^2 of AR model for returns)	37%
1	7	Measure frequencies	Second to minute	18%
			One to five minutes	26%
			Five to thirty minutes	34%
			Day to week	13%
			Day to month	10%
2,4,5	6	Tick test or aggressor flag	Aggressor flag (available only for part of the sample)	84%
			Tick test	16%
2,4	7	Post-trade value	Price five minutes after trade	81%
			Price 10 minutes after trade	6%
			Price 30 minutes after trade	13%

(Continued)

Table V—Continued

RT-Hypothesis	Fork	Fork Description	Alternatives	Frequency
2,4	8	Aggregation	Equal-weighted average	47%
			Trade-size-weighted average	53%
3	6	Units...	Volume expressed in #contracts	70%
			Volume expressed in euro	30%
6	6	Reference price	Last trade price in the day	62%
			Last trade price one day later	1%
			Volume-weighted-average-price (VWAP) full-day	24%
6	7	Mean or median	VWAP based on last five trades in the day	0%
			Mean	96%
			Median	4%
6	8	Handle non-negatives	Translate and transform ($\varepsilon = 0.001$)	14%
			Translate and transform ($\varepsilon = 1$)	7%
6	9	Retain negative-trend sign	Set to missing	79%
			Yes	79%
			No	21%

^aWinsorization is applied at the frequency of analysis (fork 4).

For each fork, we also asked RTs to rate the fit between the alternative they picked from the set and what they actually did in #fincap. Their average rating ranges between 4.0 for RT-H6 and 4.4 for RT-H3 on a scale from 1 “Far from what we did” to 5 “Very close to what we did” (see Figure IA.IV in the [Internet Appendix](#)). We therefore believe that the multiverse analysis is representative of the #fincap analysis itself.

A multiverse analysis is powerful but resource intensive. The table illustrates that the analysis becomes very large very quickly. For RT-H6, for example, the nine forks generate $2 \times 2 \times 3 \times 4 \times 3 \times 4 \times 2 \times 3 \times 2 = 6,912$ possible paths. Not all possible paths are equally reasonable, however, and the #fincap data help us select the most reasonable ones. The result is a weighted multiverse, with untraveled paths get tiny zero weight. The other paths get weights proportional to the number of teams that picked the path. The vast majority of paths, however, was picked by only one team, so the size of the multiverse is slightly less than 164 (the actual number varies across RT hypotheses).

The analysis is done for the original sample as well as for 1,000 bootstrapped samples. These additional samples are needed to estimate the correlations in test statistics across paths. These correlations are used to adjust significance levels when accounting for MHT. This is used in assessing whether NSEs are statistically significant, and whether individual estimates are statistically significant (see Panel B in Table II and Table IA.VII, where the latter is in the [Internet Appendix](#), respectively). Each RT hypothesis therefore requires processing the 720 million trade records almost 164,000 times.³¹

Results. Figure 5 illustrates that the multiverse is able to generate dispersion in estimates that is on par with the dispersion in reported estimates. The box plots for reported estimates are drawn in gray, overlaid by the multiverse box plots in color. The large dispersion in multiverse estimates is remarkable, since they are based on a few decisions only.³²

Figure 6 illustrates how sensitive the distribution of estimates is to variation across alternatives at the various forks. The plots reveal that two common strong refractors are the (econometric) model choice and the sampling frequency. A well-known force that drives a wedge between high- and

³¹ To keep the multiverse analysis feasible or tractable, we optimized the code by identifying commonalities across paths and use these to economize on loops. For example, for a particular day, realized spread calculations can iterate once over all trades to obtain realized spreads both for the path that retains all trading and the path that excludes the first and last 30 minutes of trading. Efficient coding further involves identifying opportunities for parallel processing. The multiverse analysis has been implemented on Snellius, a national supercomputer available to Dutch scientists (128 cores and 200 GB internal memory). With all of this help, the code took a few days instead of a few months to run for each RT hypothesis.

³² The multiverse models only a few forks and hence its estimates are unlikely to accurately predict reported estimates. The explanatory power of regressions with reported #fincap estimates as dependent variables and multiverse estimates as explanatory variables is low. The larger point of the multiverse analysis is to illustrate that, for a subset of forks, variation across paths can generate large NSEs. It further allows researchers to drill down and identify the forks that generate most of the dispersion in estimates.

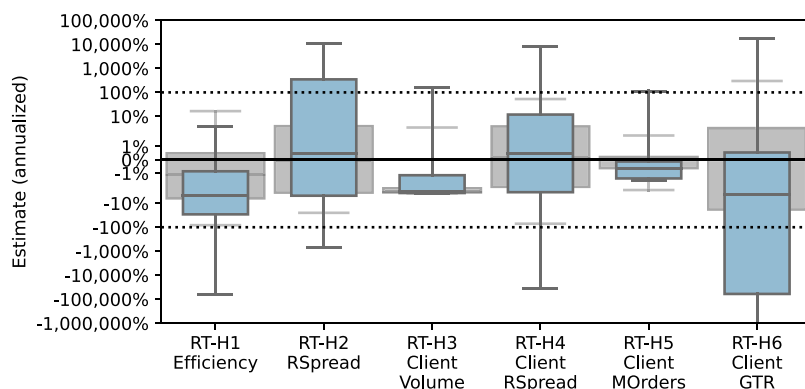


Figure 5. Dispersion in stage 1 estimates of multiverse analysis. This plot illustrates the dispersion in stage 1 estimates obtained from the multiverse analysis. The dispersion in *reported* estimates appears in gray and corresponds to Panel A in Figure 1. The boxes depict the first and third quartiles. The horizontal line in the box corresponds to the median. The whiskers depict the 2.5% and 97.5% quantiles. (Color figure can be viewed at [wileyonlinelibrary.com](https://onlinelibrary.wiley.com))

low-frequency relatives is Jensen's inequality (Blume (1974)),

$$\underbrace{\Pi_{t=1}^T E(M_t)}_{\text{Expected high frequency relative}} < \underbrace{E(\Pi_{t=1}^T M_t)}_{\text{Expected low frequency relative}}, \quad (7)$$

if $M_t \in \mathbb{R}^+$ are identical independently distributed random variables, since $f(x) = x^T$ is a convex function. First-order Taylor expanding the left-hand side around one and subtracting one from both sides yields

$$T(E(M_t) - 1) \lesssim E(\Pi_{t=1}^T M_t) - 1. \quad (8)$$

If there are T high-frequency periods in a low-frequency period, then T times the average high-frequency return is expected to be lower than the average low-frequency return. Figure 7 illustrates the effect of this inequality. The three right-most bars illustrate how, for the relative-change model, the median annualized return is $-23,000\%$ for data sampled at the daily frequency, -200% for the monthly frequency, and only -4.56% for the yearly frequency. The left-most six bars that correspond to the trend-stationary or log-difference model do not show such discrepancy across frequencies. The reason is that both of these models are linear and hence do not suffer from Jensen's inequality. The trend-stationary model features a linear trend and, in a log-difference model, the log of a product of relatives becomes a sum of log relatives.

Figure 6 further highlights some idiosyncratic sensitivities. For RT-H1, for example, the second-most sensitive fork is the frequencies that are picked to assess the deviation from a random walk. Further analysis reveals that when comparing high frequencies, such as one-second to one-minute returns,

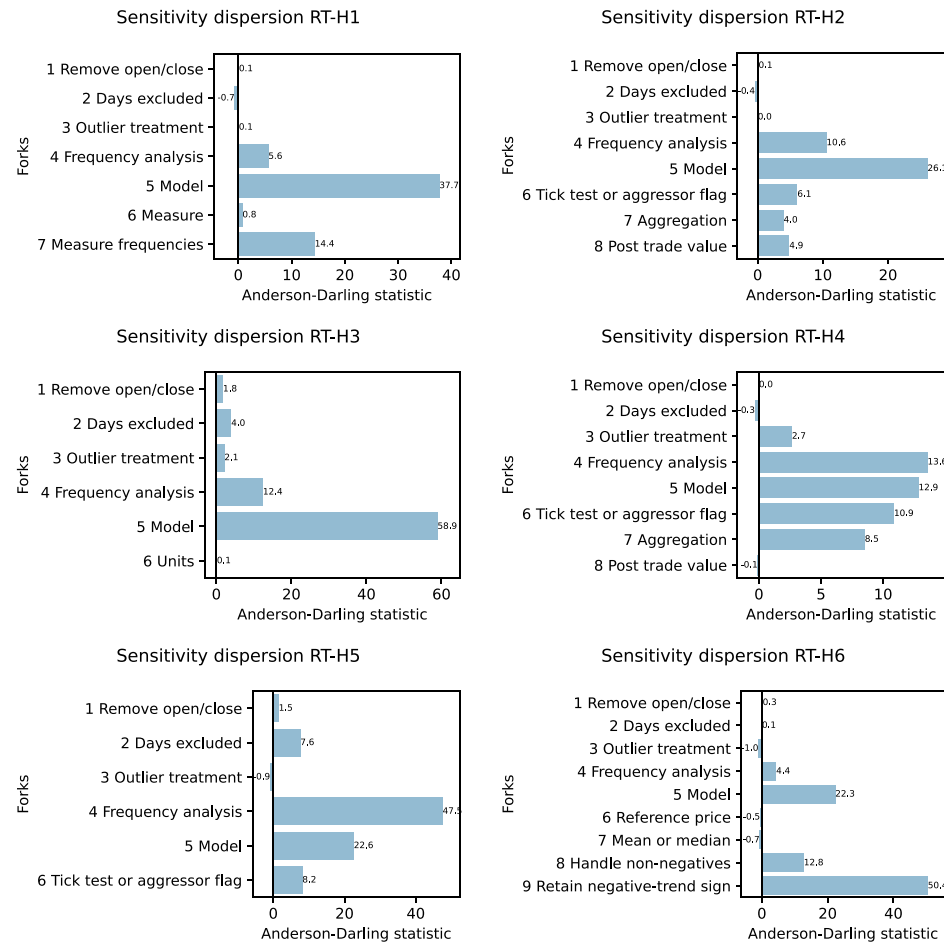


Figure 6. Fork sensitivity of estimates in multiverse analysis. This figure plots how sensitive the distribution in estimates is to the alternatives available at a fork in the multiverse analysis. The sensitivity is measured by the standardized Anderson-Darling test statistic. Higher values of the statistic imply that distributions become more dissimilar across alternatives at the fork. (Color figure can be viewed at wileyonlinelibrary.com)

almost all analyses exhibit a decline in market efficiency. However, when comparing low frequencies, such as daily to monthly returns, about half of the analyses show an increase in market efficiency whereas the other half show a decline.

Another example is the retain-negative-sign fork, which is the most sensitive one for RT-H6. The decision that each team had to make is whether a negative number that becomes more negative yields a positive percentage change or a negative percentage change. The first case emphasizes that a (negative) number becomes magnified, whereas the second emphasizes a negative trend (i.e.,

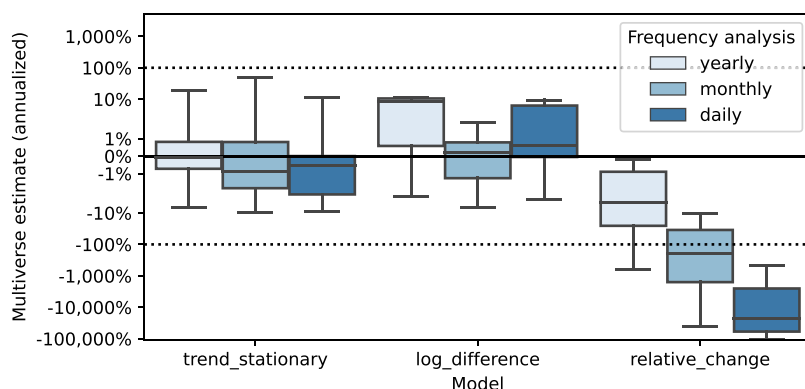


Figure 7. Sensitivity of estimates in multiverse analysis of RT-H1. This plot illustrates how the distribution of RT-H1 estimates depends on two influential forks in the multiverse analysis: (i) the model and (ii) the frequency of the analysis. Distributions are obtained by bootstrapping 1,000 times from the original sample for each analysis path. To avoid clutter, the weekly frequency is dropped since it is used by only one team (out of 164). (Color figure can be viewed at [wileyonlinelibrary.com](https://onlinelibrary.wiley.com/doi/10.1111/jofa.13337))

“retain a negative sign”). We find that 21% of the teams picked the first option and 79% picked the second one. It is not surprising that mapping an estimate from the positive to the negative domain causes strong refraction in estimates. This is an example of how a decision that each team might have thought was a trivial one (in the sense that there is only one option) can generate NSE.

D. Alternative Explanations

Having presented our results, it is useful to discuss alternative explanations. In Section II.A above, we argue that the group of #fincap participants is representative of the academic community in empirical finance/liquidity. One might nonetheless wonder to what extent the sizable NSEs are due to the presence of inexperienced researchers testing unsuitable hypotheses with little effort. We believe this is unlikely to be the case for the following three reasons.

Experience. Aware of this potential pitfall, we selectively approached researchers (for RTs and PEs) whom we knew were sufficiently experienced in the field. When signing up, they ticked a box indicating that they understood that participating in #fincap requires research experience in empirical finance/liquidity and in the analysis of large data sets. Ticking the box further means that they acknowledge that one of the team members held a PhD in finance or economics. After ticking the box, researchers had to specify in an open text box why they believe they meet these requirements. We parsed the content of this box to ensure that the team qualifies before accepting them into #fincap (see Figure A in the Internet Appendix for the sign-up sheet).

Hypotheses. We proceeded with care when designing the RT hypotheses. Early versions were shared with senior scholars, and their feedback helped us fine-tune the RT hypotheses. We therefore feel confident that the RT hypotheses are suitable and well-motivated hypotheses to test with the RT sample (see Figure E in the [Internet Appendix](#) for the RT instruction sheet, which shows how RT hypotheses were presented to RTs).

Related to the suitability question, one might wonder whether vagueness of an RT hypothesis might be a viable alternative explanation for sizable NSEs. To address this concern, we included a very precise RT hypothesis: RT-H3 on client volume share. The results for RT-H3 show that NSEs can be sizable, even for relatively precise hypotheses. It is true, however, that NSEs tend to be lower for the more precise RT hypotheses.

Effort. We incentivized RTs to exert effort by providing them with the following information (before they signed up): the deadlines of the various stages so that they could plan for them; their *nonanonymized* paper would be evaluated by senior peer reviewers; the top-five (anonymized) papers would be announced to all others³³, and only those who complete all stages become coauthors. In addition to these incentives, we believe that most scientists are propelled by an intrinsic motivation to do good research.

Looking back, we have several reasons to believe that researchers did indeed exert serious effort. First, only four out of 168 RTs failed to complete all stages; 123 out of 168 teams (73.2%) handed in their stage 1 report at least a day early, and none of the teams seriously breached any deadline. Second, the average reproducibility score was 64.5 on a scale from 0 (low) to 100, which is high in comparison to what has been reported in other reproducibility studies (Colliard, Hurlin, and Pérignon (2021)). Finally, the average paper quality was 6.2 on a scale from zero (low) to 10. As for PEs, we also believe they exerted serious effort, because all who signed up as a PE completed their reviews on time.

III. Conclusion

Researchers need to make many decisions when testing hypotheses on a particular sample—pick an appropriate measure, treat outliers, select a statistical model, etc. If researchers are not perfectly aligned on these decisions, then their estimates are likely to differ. This potential dispersion in estimates therefore adds uncertainty to an estimate reported by a *single* team, as other teams might have reported other estimates based on the same data.

We measure dispersion in estimates across researchers robustly with an IQR, which we refer to as NSE. We study NSEs in an experiment in which 164 teams test the same six RT hypotheses on the same sample. We find NSEs to be

³³ Individuals obtain “ego utility” from positive views about their ability to do well, and they exert more effort (or take more risks) when they are informed about their rank in nonincentivized competitions (Köszegi (2006), Tran and Zeckhauser (2012), Kirchler, Lindner, and Weitzel (2018)).

substantial, even for a relatively straightforward market-share hypothesis. For this RT hypothesis, we find NSE to be 1.2% around a median of -3.3% . A more opaque RT hypothesis on market efficiency yields larger variation with an NSE of 6.7% around a median of 1.1%. We further find that NSEs are smaller for more reproducible and higher quality papers as rated by peers.

A multiverse analysis based on key forks sheds light on how important each fork is in generating dispersion in estimates. It turns out that many forks add substantial dispersion in estimates. Two particularly powerful ones are sampling frequency and the statistical model. Using a nonlinear model at high frequency to estimate a low-frequency trend can therefore add substantial noise (Jensen's inequality).

NSEs being substantial is worrisome. An encouraging result, however, is that peer feedback reduces NSEs by half. We further note that #fincap NSEs are likely to be an upper bound for real-world dispersion in (published) estimates. First, published papers have likely gone through more stages of feedback. Second, papers submitted to a journal enter further review stages only if referees judge them to be of high enough quality. Published results might also be affected by *p*-hacking, which is a selective process and thus likely further reduces dispersion, and potentially introduces bias. Overall, we believe the full empirical research publication process deserves further scrutiny. Notwithstanding, the strong peer feedback that we document encourages the profession to think of more ways to interact to reduce NSEs.

Finally, our multiverse analysis provides guidance on what threshold to use in individual tests when accounting for multiple testing. Bonferroni assumes independence among test statistics and adjusts significance levels by the number of tests, 164 in the case of #fincap. Bootstrapped multiverse results show that there is substantial correlation among test statistics and finds adjustment factors that range between 13 and 91 (depending on the RT hypothesis). The threshold for two-sided testing at the 5% level should therefore be at least $\Phi(1 - 0.025/13) = 2.9$. This is in line with the 3.0 lower bound recommended by Harvey, Liu, and Zhu (2016) for factor tests in asset pricing.

Initial submission: November 11, 2021; Accepted: February 14, 2023
Editors: Stefan Nagel, Philip Bond, Amit Seru, and Wei Xiong

Appendix A: Authors and Affiliations

The full list of authors and affiliations did not fit into the initial footnote. This appendix makes up for this omission by providing the full list of authors and their affiliations:

Albert J. Menkveld is at Vrije Universiteit Amsterdam and Tinbergen Institute. Anna Dreber is at Stockholm School of Economics and University of Innsbruck. Felix Holzmeister is at University of Innsbruck. Juergen Huber is at University of Innsbruck. Magnus Johannesson is at Stockholm School of Economics. Michael Kirchler is at University of Innsbruck. Sebastian Neusüß is at Optiver Amsterdam. Michael Razen is at University of Innsbruck. Utz

Weitzel is at Vrije Universiteit Amsterdam, Radboud University, and Tinbergen Institute. David Abad-Díaz is at University of Alicante. Menachem (Meni) Abudy is at Bar-Ilan University. Tobias Adrian is at International Monetary Fund. Yacine Ait-Sahalia is at Princeton University. Olivier Akmansoy is at CNRS and Cascad. Jamie T. Alcock is at University of Birmingham. Vitali Alexeev is at University of Technology Sydney. Arash Aloosh is at Neoma Business School. Livia Amato is at University of Chicago Booth School of Business. Diego Amaya is at Wilfrid Laurier University. James J. Angel is at Georgetown University. Alejandro T. Avetikian is at Pontificia Universidad Católica de Chile. Amadeus Bach is at University of Mannheim. Edwin Baidoo is at Tennessee Technological University. Gaetan Bakalli is at Emlyon Business School. Li Bao is at University Toulouse Capitole. Andrea Barbon is at University of St. Gallen and Swiss Finance Institute. Oksana Bashchenko is at SFI at University of Lausanne. Parampreet C. Bindra is at University of Innsbruck. Geir H. Bjønnes is at BI Norwegian Business School. Jeffrey R. Black is at University of Memphis. Bernard S. Black is at Northwestern University. Dimitar Bogoev is at EDF Energy London. Santiago Bohorquez Correa is at Universidad EAFIT. Oleg Bondarenko is at University of Illinois at Chicago. Charles S. Bos is at Vrije Universiteit Amsterdam. Ciril Bosch-Rosa is at Technische Universität Berlin. Elie Bouri is at Lebanese American University. Christian Brownlees is at Universitat Pompeu Fabra and Barcelona School of Economics. Anna Calamia is at TBS Business School. Viet Nga Cao is at Monash University. Gunther Capelle-Blancard is at Paris School of Business. Laura M. Capera Romero is at Vrije Universiteit Amsterdam. Massimiliano Caporin is at University of Padova. Allen Carrion is at University of Memphis. Tolga Caskurlu is at University of Amsterdam. Bidisha Chakrabarty is at Saint Louis University. Jian Chen is at Queen's University. Mikhail Chernov is at UCLA, NBER, and CEPR. William Cheung is at Waseda University. Ludwig B. Chincarini is at University of San Francisco, USCF Investments, and Wedge Capital. Tarun Chordia is at Emory University, Visiting Professor Deakin University. Sheung-Chi Chow is at Australian National University. Benjamin Clapham is at Goethe University Frankfurt. Jean-Edouard Colliard is at HEC Paris. Carole Comerton-Forde is at University of Melbourne and CEPR. Edward Curran is at Macquarie University. Thong Dao is at Nottingham Trent University. Wale Dare is at HEC Liège - University of Liège. Ryan J. Davies is at Babson College. Riccardo De Blasis is at Marche Polytechnic University. Gianluca F. De Nard is at University of Zurich, New York University, and OLZ AG. Fany Declerck is at University Toulouse Capitole and Toulouse School of Economics. Oleg Deev is at Masaryk University. Hans Degryse is at KU Leuven. Solomon Y. Deku is at Nottingham Trent University. Christophe Desagre is at ICHEC Brussels Management School and UCLouvain. Mathijs A. van Dijk is at Erasmus University Rotterdam. Chukwuma Dim is at George Washington University. Thomas Dimpfl is at University of Hohenheim. Yun Jiang Dong is at Queen's University. Philip A. Drummond is at The Brattle Group. Tom Dudda is at Technische Universität Dresden. Teodor Duevski is at HEC Paris. Ariadna Dumitrescu is at ESADE Business School, Universitat Ramon Llull.

Teodor Dyakov is at EDHEC Business School. Anne Haubo Dyhrberg is at Wilfrid Laurier University. Michał Dzielinski is at Stockholm University. Asli Eksi is at Salisbury University. Izidin El Kalak is at Cardiff University. Saskia ter Ellen is at International Monetary Fund and Norges Bank. Nicolas Eugster is at University of Queensland. Martin D. D. Evans is at Georgetown University. Michael Farrell is at University of Wisconsin-Milwaukee. Ester Felez-Vinas is at University of Technology Sydney. Gerardo Ferrara is at Bank of England. El Mehdi Ferrouhi is at Ibn Tofail University. Andrea Flori is at Politecnico di Milano. Jonathan T. Fluharty-Jaidee is at West Virginia University. Sean D. V. Foley is at Macquarie University. Kingsley Y. L. Fong is at University of New South Wales. Thierry Foucault is at HEC Paris and CEPR. Tatiana Franus is at Bayes Business School. Francesco Franzoni is at Università della Svizzera italiana and Swiss Finance Institute. Bart Frijns is at Open Universiteit. Michael Frömmel is at Ghent University. Servanna M. Fu is at University of Essex. Sascha C. Füllbrunn is at Radboud University and Institute for Management Research. Baoqing Gan is at University of Technology Sydney. Ge Gao is at Beijing Sport University. Thomas P. Gehrig is at University of Vienna. Roland Gemayel is at King's College London. Dirk Gerritsen is at Utrecht University. Javier Gil-Bazo is at Universitat Pompeu Fabra, Barcelona School of Economics, and UPF Barcelona School of Management. Dudley Gilder is at Cardiff University. Lawrence R. Glosten is at Columbia University. Thomas Gomez is at Utrecht University. Arseny Gorbenko is at Monash University. Joachim Grammig is at University of Tübingen and Centre for Financial Research Cologne. Vincent Grégoire is at HEC Montréal. Ufuk Güçbilmez is at University of Glasgow. Björn Hagströmer is at Stockholm University. Julien Hambuckers is at HEC Liège - University of Liège. Erik Hapnes is at nan. Jeffrey H. Harris is at American University. Lawrence Harris is at University of Southern California. Simon Hartmann is at Vienna University of Economics and Business. Jean-Baptiste Hasse is at Aix-Marseille University and UCLouvain. Nikolaus Hautsch is at University of Vienna and Vienna Graduate School of Finance. Xue-Zhong (Tony) He is at Xi'an Jiaotong-Liverpool University. Davidson Heath is at University of Utah. Simon Hediger is at University of Zurich. Terrence Hendershott is at UC Berkeley. Ann Marie Hibbert is at West Virginia University. Erik Hjalmarsson is at University of Gothenburg. Seth A. Hoelscher is at Missouri State University. Peter Hoffmann is at European Central Bank. Craig W. Holden is at Indiana University. Alex R. Horenstein is at Miami Herbert Business School. Wenqian Huang is at Bank for International Settlements. Da Huang is at Northeastern University. Christophe Hurlin is at University of Orléans and Cascad. Konrad Ilczuk is at Konrad Ilczuk - Consulting. Alexey Ivashchenko is at Vrije Universiteit Amsterdam. Subramanian R. Iyer is at University of New Mexico. Hossein Jahanshahloo is at Cardiff University. Naji Jalkh is at Saint Joseph University. Charles M. Jones is at Columbia Business School. Simon Jurkatis is at Bank of England. Petri Jylhä is at Aalto University. Andreas T. Kaeck is at University of Sussex. Gabriel Kaiser is at Vienna University of Economics and Business and University of Luxembourg. Arzé Karam is at Durham University. Egle

Karmaziene is at Vrije Universiteit Amsterdam, Tinbergen Institute, and Swedish House of Finance. Bernhard Kassner is at University of Munich (LMU). Markku Kaustia is at Aalto University. Ekaterina Kazak is at University of Manchester. Fearghal Kearney is at Queen's University Belfast. Vincent van Kervel is at Pontificia Universidad Católica de Chile. Saad A. Khan is at HEC Montréal. Marta K. Khomyn is at The University of Adelaide. Tony Klein is at Queen's University Belfast and University of Barcelona. Olga Klein is at University of Warwick. Alexander Klos is at Kiel University. Michael Koetter is at Halle Institute for Economic Research and Otto von Guericke University Magdeburg. Aleksey Kolokolov is at University of Manchester. Robert A. Korajczyk is at Northwestern University. Roman Kozhan is at University of Warwick. Jan P. Krahnen is at Goethe University Frankfurt and Leibniz Institute SAFE. Paul Kuhle is at Universidad Autónoma de Madrid. Amy Kwan is at University of New South Wales. Quentin Lajaunie is at University of Orléans and Square Research Center. F. Y. Eric C. Lam is at Hong Kong Institute for Monetary and Financial Research and Hong Kong Monetary Authority. Marie Lambert is at HEC Liège - University of Liège. Hugues Langlois is at HEC Paris. Jens Lausen is at Goethe University Frankfurt. Tobias Lauter is at Leibniz University Hannover. Markus Leippold is at University of Zurich and Swiss Finance Institute. Vladimir Levin is at University of Luxembourg. Yijie Li is at UBS. Hui Li is at La Trobe University. Chee Yoong Liew is at UCSI University. Thomas Lindner is at Vienna University of Economics and Business and University of Innsbruck. Oliver Linton is at University of Cambridge. Jiacheng Liu is at Purdue University. Anqi Liu is at University of Sydney. Guillermo Llorente is at Universidad Autónoma de Madrid. Matthijs Lof is at Aalto University. Ariel Lohr is at Arizona State University. Francis Longstaff is at UCLA. Alejandro Lopez-Lira is at University of Florida. Shawn Mankad is at North Carolina State University. Nicola Mano is at Università della Svizzera italiana and Swiss Finance Institute. Alexis Marchal is at EPFL. Charles Martineau is at University of Toronto. Francesco Mazzola is at ESCP Business School. Debrah Meloso is at TBS Business School. Michael G. Mi is at University of Sydney and NGS Super. Roxana Mihet is at SFI at University of Lausanne and CEPR. Vijay Mohan is at RMIT University. Sophie Moinas is at Toulouse Capitole University and Toulouse School of Economics. David Moore is at Loyola Marymount University. Liangyi Mu is at Queen's University Belfast and University of Manchester. Dmitriy Muravyev is at Michigan State University. Dermot Murphy is at University of Illinois at Chicago. Gabor Neszveda is at John von Neumann University and MNB Institute. Christian Neumeier is at Allianz Global Investors GmbH. Ulf Nielsson is at Copenhagen Business School, Danish Finance Institute, and University of Iceland. Mahendrarajah Nimalendran is at University of Florida. Sven Nolte is at Radboud University. Lars L. Norden is at Stockholm University. Peter O'Neill is at University of New South Wales. Khaled Obaid is at California State University - East Bay. Bernt A. Ødegaard is at University of Stavanger. Per Östberg is at University of Zurich and Swiss Finance Institute. Emiliano Pagnotta is at Singapore Management University. Marcus Painter is at Saint

Louis University. Stefan Palan is at University of Graz. Imon J. Palit is at RMIT University. Andreas Park is at University of Toronto. Roberto Pascual is at University of the Balearic Islands. Paolo Pasquariello is at Ross School of Business, University of Michigan. Lubos Pastor is at University of Chicago Booth School of Business. Vinay Patel is at University of Technology Sydney. Andrew J. Patton is at Duke University. Neil D. Pearson is at University of Illinois at Urbana-Champaign and Canadian Derivatives Institute. Loriana Pelizzon is at Leibniz Institute for Financial Research SAFE, Goethe University Frankfurt, and Ca' Foscari University of Venice and CEPR. Michele Pelli is at University of Zurich and Swiss Finance Institute. Matthias Pelster is at University of Duisburg-Essen and european center for financial services (ecfs). Christophe Pérignon is at HEC Paris and Cascad. Cameron Pfiffer is at University of Oregon. Richard Philip is at University of Sydney. Tomáš Plíhal is at Masaryk University. Puneet Prakash is at Missouri State University. Oliver-Alexander Press is at Copenhagen Business School. Tina Prodromou is at University of Wollongong and School of Business. Marcel Prokopczuk is at Leibniz University Hannover. Talis Putnins is at University of Technology Sydney and Stockholm School of Economics in Riga. Ya Qian is at Deutsche Bank AG. Gaurav Raizada is at IIM Ahmedabad and iRageCapital, Mumbai. David Rakowski is at University of Texas at Arlington. Angelo Ranaldo is at University of St. Gallen and Swiss Finance Institute. Luca Regis is at University of Torino and Collegio Carlo Alberto. Stefan Reitz is at Kiel University. Thomas Renault is at Université Paris 1 Panthéon-Sorbonne. Rex W. Renjie is at Vrije Universiteit Amsterdam. Roberto Reno is at ESSEC Business School. Steven J. Riddiough is at University of Toronto. Kalle Rinne is at University of Luxembourg. Paul Rintamäki is at Aalto University. Ryan Riordan is at Queen's University and University of Munich (LMU). Thomas Rittmannsberger is at Technical University Munich, School of Management. Iñaki Rodríguez Longarela is at Stockholm University and UiT - The Arctic University of Norway. Dominik Roesch is at State University of New York at Buffalo. Lavinia Rognone is at University of Edinburgh Business School. Brian Roseman is at Oklahoma State University. Ioanid Roşu is at HEC Paris. Saurabh Roy is at Université du Québec à Montréal. Nicolas Rudolf is at University of Lausanne. Stephen R. Rush is at University of Economics HCMC and Liminal Markets. Khaladdin Rzayev is at University of Edinburgh Business School, Koç University, and London School of Economics. Aleksandra A. Rzeźnik is at York University. Anthony Sanford is at HEC Montréal. Harikumar Sankaran is at New Mexico State University. Asani Sarkar is at Federal Reserve Bank of New York. Lucio Sarno is at University of Cambridge. Olivier Scaillet is at University of Geneva and Swiss Finance Institute. Stefan Scharnowski is at University of Mannheim. Klaus R. Schenk-Hoppé is at University of Manchester. Andrea Schertler is at University of Graz. Michael Schneider is at Deutsche Bundesbank and Leibniz Institute for Financial Research SAFE. Florian Schroeder is at Macquarie University. Norman Schürhoff is at SFI at University of Lausanne and CEPR. Philipp Schuster is at University of Stuttgart. Marco A. Schwarz is at Düsseldorf Institute for

Competition Economics and CESifo. Mark S. Seasholes is at Arizona State University. Norman J. Seeger is at Vrije Universiteit Amsterdam and Tinbergen Institute. Or Shachar is at Federal Reserve Bank of New York. Andriy Shkilko is at Wilfrid Laurier University. Jessica Shui is at Federal Housing Finance Agency. Mario Sikic is at University of Zurich. Giorgia Simion is at Vienna University of Economics and Business and Vienna Graduate School of Finance. Lee A. Smales is at University of Western Australia. Paul Söderlind is at University of St. Gallen. Elvira Sojli is at University of New South Wales. Konstantin Sokolov is at University of Memphis. Jantje Sönksen is at University of Tübingen. Laima Spokeviciute is at Cardiff University. Denitsa Stefanova is at University of Luxembourg. Marti G. Subrahmanyam is at NYU Stern and NYU Shanghai. Barnabas Szasz is at ELTE, Eotvos Lorand University. Oleksandr Talavera is at University of Birmingham. Yuehua Tang is at University of Florida. Nick Taylor is at University of Bristol. Wing Wah Tham is at University of New South Wales. Erik Theissen is at University of Mannheim. Julian Thimme is at Karlsruhe Institute of Technology. Ian Tonks is at University of Bristol. Hai Tran is at Loyola Marymount University. Luca Trapin is at University of Bologna. Anders B. Trolle is at Copenhagen Business School and Danish Finance Institute. M. Andreea Vaduva is at Universidad Carlos III de Madrid. Giorgio Valente is at Hong Kong Institute for Monetary and Financial Research and Hong Kong Monetary Authority. Robert A. Van Ness is at University of Mississippi. Aurelio Vasquez is at ITAM. Thanos Verousis is at University of Essex. Patrick Verwijmeren is at Erasmus University Rotterdam. Anders Vilhelmsson is at Lund University. Grigory Vilkov is at Frankfurt School of Finance and Management. Vladimir Vladimirov is at University of Amsterdam. Sebastian Vogel is at Erasmus University Rotterdam. Stefan Voigt is at University of Copenhagen and Danish Finance Institute. Wolf Wagner is at Erasmus University Rotterdam. Thomas Walther is at Utrecht University and Technische Universität Dresden. Patrick Weiss is at Reykjavik University and Vienna University of Economics and Business. Michel van der Wel is at Erasmus University Rotterdam. Ingrid M. Werner is at The Ohio State University and CEPR. P. Joakim Westerholm is at University of Sydney. Christian Westheide is at University of Vienna and Leibniz Institute for Financial Research SAFE. Hans C. Wika is at University of Minnesota and Norges Bank. Evert Wipplinger is at Vrije Universiteit Amsterdam. Michael Wolf is at University of Zurich. Christian C. P. Wolff is at University of Luxembourg. Leonard Wolk is at Vrije Universiteit Amsterdam. Wing-Keung Wong is at Asia University, China Medical University Hospital, and The Hang Seng University of Hong Kong. Jan Wrampelmeyer is at Vrije Universiteit Amsterdam and Tinbergen Institute. Zhen-Xing Wu is at Zhongnan University of Economics and Law. Shuo Xia is at Halle Institute for Economic Research and Leipzig University. Dacheng Xiu is at University of Chicago Booth School of Business. Ke Xu is at University of Victoria. Caihong Xu is at Stockholm University. Pradeep K. Yadav is at University of Oklahoma. José Yagüe is at University of Murcia. Cheng Yan is at University of Essex. Antti Yang is at Cornerstone Research and Erasmus University

Rotterdam. Woongsun Yoo is at Central Michigan University. Wenjia Yu is at Aalto University. Yihe Yu is at State University of New York at Buffalo. Shihao Yu is at Columbia University. Bart Z. Yueshen is at INSEAD. Darya Yuferova is at Norwegian School of Economics (NHH). Marcin Zamojski is at University of Gothenburg. Abalfazl Zareei is at Stockholm University. Stefan M. Zeisberger is at Radboud University and University of Zurich. Lu Zhang is at University of Luxembourg. S. Sarah Zhang is at University of Manchester. Xiaoyu Zhang is at Vrije Universiteit Amsterdam and Tinbergen Institute. Lu Zhao is at Southwestern University of Finance and Economics. Zhuo Zhong is at University of Melbourne. Z. Ivy Zhou is at University of Wollongong. Chen Zhou is at Erasmus University Rotterdam. Xingyu S. Zhu is at Stockholm School of Economics. Marius Zoican is at University of Toronto. Remco Zwinkels is at Vrije Universiteit Amsterdam and Tinbergen Institute.

Appendix B: Reconciliation with PAP Results

The original version of the paper—*NSEs*—contains results of the analysis outlined in the PAP. This original version is available as Tinbergen Institute Discussion Paper TI 2021-102/IV. Most tables and figures have not changed.³⁴

The tables that have changed are Tables 3 and 4 in the original version. The reason is that these are the only two regression tables. In the original version, we estimate a heteroskedasticity model using OLS. The dependent variable is log squared error. However, OLS estimates are notoriously sensitive to extreme outliers, which turn out to be a feature of the #finap sample (see footnote 21 or Figure 1). Quantile regressions are robust to the presence of extreme outliers and are therefore more appropriate for the analysis of our sample. Moreover, they model the entire distribution instead of just a conditional mean (as emphasized in the introduction). In the remainder of this appendix, we compare results across the two tables in the original version and the current version to reconcile previous with current findings.³⁵

Table 3 in the original version has become Table III in the current version. These tables both relate dispersion in estimates to quality variables in order to test the first hypothesis. In the original version, most results are insignificant. The only significance is for reproducibility when using a 2.5% to 97.5% winsorized sample. The coefficient of -0.24 implies that a 10% increase in reproducibility coincides with a reduction in the SD of estimates of $1/2 \times 0.24 \times 10\% = 1.2\%$ (the coefficient $1/2$ converts variance to SD; see footnote 21 in original paper). In the current version, the first quartile (Q1) covaries significantly *positively* with reproducibility and paper quality, whereas the third quartile covaries significantly *negatively* with these factors. They, therefore, co-vary significantly negatively with IQR. A 10% increase in reproducibility

³⁴ More specifically, tables 1, 2, and 5, figures 1, 2, 3, 4, and 5, have not changed. In the current version, they appear as Tables I, IA.IV, and IA.VI, and Figures IA.II, 1a, 1b, IA.III, and 4, respectively, where the IA prefix indicates that they are in the Internet Appendix.

³⁵ We have changed the statistical methodology guided by the feedback of *The Journal of Finance* referees. One could say that this “peer feedback” likely reduced the NSE of our results.

coincides with a reduction in IQR of $10\% \times (0.109 + 0.142) \times 0.44 = 1.1\%$.³⁶ Note that this effect is in the same ballpark as the 1.2% in the original paper.

Table 4 in the original version has become Table IV in the current version. In the original version, the unwinsorized sample shows a weakly significant decline in dispersion of estimates across all stages. The effect is also relatively small in magnitude since the SD decline is only 9%. With extreme outliers removed in the 2.5% to 97.5% winsorized sample, the decline becomes both significant and larger in magnitude. The SD now declines by 53.5% across all stages. The results in the current version show that Q1 of the estimate distribution increases significantly across all stages and Q3 decreases significantly. The result is a decline of 47.2% (depicted in Figure 3). Again, the numbers in both versions are in the same ballpark.

Appendix C: RT Sample, RT Hypotheses, and Results

This appendix presents the RT hypotheses in detail and the test results of #fincap RTs as a group. The instruction sheet itself is available as Figure E in the [Internet Appendix](#). We start by providing the context that motivates the RT hypotheses.

A. Context

Electronic order matching systems (automated exchanges) and electronic order generation systems (algorithms) have changed financial markets over time. Investors used to trade through broker-dealers by paying dealer ask prices when buying, and accepting dealer bid prices when selling. The wedge between these bid and ask prices, the bid-ask spread, was a useful measure of trading cost, and often still is.

Today, investors more commonly trade in electronic limit-order markets (as is the case for EuroStoxx 50 futures). They still trade at bid and ask prices. They do so by submitting so-called market orders and marketable limit orders. However, investors can now also quote bid and ask prices themselves by submitting (nonmarketable) standing limit orders. Investors also increasingly use agency algorithms to automate their trades. Concurrently, exchanges have been continuously upgrading their systems to better serve their clients. Has market quality improved, in particular when taking the perspective of nonexchange members: (end-user) clients?

B. RT Hypotheses and Test Results

The RT hypotheses and results are discussed based on estimates in the final stage of the project (available as Table IA.III in the [Internet Appendix](#)). We therefore focus our discussion on the results that RTs settle on after receiving all feedback. What do RTs find after having shown some convergence across

³⁶ The square root of the average variance of reproducibility (demeaned by RT hypothesis) is 0.44.

the stages? Further, consistent with the main text, we focus our discussion on robust location and dispersion statistics, namely, the median and IQR, respectively. We note that such discussion is meaningful because Table IA.VII in the Internet Appendix shows that, for all RT hypotheses, the null of a zero trend is rejected at the 0.5% significance level. This significance level is used for all tests in the remainder of the subsection.

(The first two hypotheses focus on all trades.)

RT-H1. Assuming that informationally efficient prices follow a random walk, did market efficiency change over time?

Null hypothesis: Market efficiency has not changed over time.

Findings. The median estimate is -1.1% with an IQR of 2.6% . The third quartile is -0.2% and the vast majority therefore finds a negative trend in efficiency. The Bonferroni tests show that 31 RTs find a significant negative trend against only four who find a significant positive trend. The decline seems modest as the across-RT median³⁷ is -1.1% per year. The small changes add up, though, to a total change in the 2002 to 2018 sample of approximately $(0.989^{17} - 1) = -17.1\%$. This might reflect a trend of declining depth in the market, possibly due to new regulation in the aftermath of the global financial crisis of 2007 to 2008. Postcrisis regulation constrains the supply of liquidity by sell-side banks (e.g., Bao, O'Hara, and Zhou (2018), Jovanovic and Menkveld (2022)). If these banks incur higher inventory costs as a result, then in equilibrium one observes larger transitory price pressure thus reducing market efficiency (e.g., Pastor and Stambaugh (2003), Hendershott and Menkveld (2014)). In the interest of brevity, we discuss all remaining hypotheses in the same way.

RT-H2. Did the (realized) bid-ask spread paid on market orders change over time? The realized spread could be thought of as the gross-profit component of the spread as earned by the limit-order submitter.

Null hypothesis: The realized spread on market orders has not changed over time.

Findings. The median estimate is -2.3% with an IQR of 4.3% . The third quartile is -0.1% and the vast majority therefore finds a negative trend in realized spread. The tests show that 38 RTs find a significant negative trend, whereas only three RTs find a significant positive trend. The median decline of 2.3% per year implies a 32.7% decline over the full sample. This trend might be due to the arrival of high-frequency market-makers who operate at low costs. They do not have the deep pockets that sell-side banks have, but they will offer liquidity for regular small trades by posting near the inside of the market. Their arrival is typically associated with a tighter bid-ask spread, but not necessarily with better liquidity supply for large orders (e.g., Jones (2013), Angel, Harris, and Spatt (2015), Menkveld (2016)).

³⁷ The across-RT median includes all RTs, including those who report insignificant results.

(The remaining hypotheses focus on agency trades only.)

RT-H3. Did the share of client volume in total volume change over time?

Null hypothesis: Client share volume as a fraction of total volume has not changed over time.

Findings. The median estimate is -2.9% with an IQR of 1.7% . The ninth decile is -1.1% , which shows that almost all RTs report a negative trend. The tests show that 123 RTs find a significant negative trend against only two RTs documenting a significant positive trend. A median decline of 2.9% per year implies a total decline of 39.4% for the full sample. Intermediation therefore seems to have increased, which should surprise those who believe that the arrival of agency algorithms enables investors to execute optimally themselves, thus reducing the need for intermediation.³⁸

RT-H4. On their market orders and marketable limit orders, did the realized bid-ask spread that clients paid change over time?

Null hypothesis: Client realized spreads have not changed over time.

Findings. The median estimate is -0.2% with an IQR of 2.4% . The third quartile, however, is positive, suggesting that a modest majority finds a negative trend. The tests show a bit stronger evidence for a negative trend, because 15 RTs find it to be significantly negative against only eight who find a significant positive trend. The median decline of 0.2% per year translates to a 3.3% decline for the full sample. The decline in client realized spread is therefore only about a tenth of the total realized spread decline, which suggests that market orders of intermediaries benefited most from the general decline in realized spread.

RT-H5. Realized spread is a standard cost measure for market orders, but to what extent do investors continue to use market and marketable limit orders (as opposed to nonmarketable limit orders)?

Null hypothesis: The fraction of client trades executed via market orders and marketable limit orders has not changed over time.

Findings. The median estimate is 0.0% with an IQR of 0.6% . A significantly negative trend is found by 13 RTs, whereas nine find a significantly positive trend. The results seem rather balanced between a negative and a positive trend. The results therefore appear to suggest that clients neither increased nor decreased their share of market orders. One might have expected a decrease because an increase in the use of agency algorithms should allow them to execute more through nonmarketable limit orders as opposed to market or-

³⁸ We verified with Deutsche Börse that this change is not purely mechanical in the sense that, in the sample period, many institutions became an exchange member and, with it, the status of their volume changes from agency to principal.

ders or marketable limit orders. The benefit of execution via a nonmarketable limit order is that one earns half the bid-ask spread as opposed to paying it.

RT-H6. A measure that does not rely on the classic limit- or market-order distinction is gross trading revenue (GTR). Investor GTR for a particular trading day can be computed by assuming a zero position at the start of the day and evaluating an end-of-day position at an appropriate reference price. Relative investor GTR can then be defined as this GTR divided by the investor's total (euro) volume for that trading day. This relative GTR is, in a sense, a realized spread. It reveals what various groups of market participants in aggregate pay for (or earn on) their trading. It transcends market structure as it can be meaningfully computed for any type of trading in any type of market (be it trading through limit orders only, through market orders only, through a mix of both, or in a completely different market structure).

Null hypothesis: Relative GTR for clients has not changed over time.

Findings. The median estimate is 0.0% with an IQR of 1.1%. Three RTs find a significantly positive trend and another three find a significantly negative one. The significance, therefore, is rather weak and balanced. We cautiously conclude that GTR has stayed at mostly the same level throughout the sample.

Appendix D: Explanatory Variables for Error Variance

A. Team Quality

The quality measures for RTs are based on the survey that participants filled out upon registration (see Figure A in the [Internet Appendix](#)). To keep the regression model both concise and meaningful, we reduce the ordinal variable “current position” and the logarithmic interval-based variable “size of largest dataset worked with” to binary variables. The academic position variable is equal to one if a researcher is an associate or full professor. The data set variable is equal to one if the researcher has worked with data sets that contained at least 100 million observations, because the #fincap sample contains 720 million observations. We aggregate these binary variables to RT level by taking the maximum across the team members.

As for self-assessed experience, we asked for both empirical finance and market liquidity, which we deem equally relevant for testing the RT hypotheses. Thus, and because of the anticipated high correlation, we use the average of these two measures to obtain the individual score. In the interest of consistency, we then aggregate to the team level by taking the maximum across the team members.

B. Workflow Quality

We proxy for workflow quality with an objectively obtained score of code quality provided by Cascad (see footnote 19). The scale ranges from 0 (serious discrepancies) to 100 (perfect reproducibility).

C. Paper Quality

Papers are rated by an external group of PEs. They rate the analyses for each RT hypothesis individually, as well as the paper in its entirety (see Section III.J in the [Internet Appendix](#)). The ratings range from 0 (very weak) to 10 (excellent). Each paper is rated by two PEs and the paper rating is the average of the two (after removing a PE fixed effect as discussed in Section I.A above).

Appendix E: Bootstrap Procedure for Belief Statistic D

The distribution of D under the null of equal means is obtained by bootstrapping as follows. For each RT hypothesis, we subtract the difference between the average belief on SD and the observed SD from the beliefs,

$$AdjBeliefOnSD_{ij} = BeliefOnSD_{ij} - \left[\left(\frac{1}{n} \sum_i BeliefOnSD_{ij} \right) - RealizationOfSD_j \right] \quad (E1)$$

In this new sample with adjusted beliefs, the average belief about dispersion equals the observed dispersion, by construction. This sample is input into the bootstrapping procedure, which iterates through the following steps 10,000 times:

1. As we have n RTs, in each iteration we draw n times from the new sample, with replacement. Each draw picks a particular RT and stores its beliefs and its results for all six of the RT hypotheses. The result of these n draws therefore is a simulated sample that has the same size as the original sample.
2. The simulated sample is used to compute the test statistic D in (6). This statistic for iteration k , a scalar, is stored as D_k .

The bootstrap procedure yields 10,000 observations of the test statistic under the null. For a significance level of 0.005, the statistic observed in the #fincap sample is statistically significant if it lands below the 25th lowest simulated

statistic or above the 25th highest simulated statistic. Its p -value is³⁹

$$2 \min(\text{EmpiricalQuantileFincapStatistic}, 1 - \text{EmpiricalQuantileFincapStatistic}). \quad (\text{E2})$$

Appendix F: AD Test

The sensitivity of dispersion to a particular fork is measured by a k -sample AD test (Scholz and Stephens (1987)). This test was designed to verify whether k separate samples are drawn from the same distribution. The AD test statistic T_{k-1} measures the distance between the empirical distribution functions of k separate samples. It does not rely on parametric assumptions, and thus it is particularly attractive for our application as distributions are unknown *ex ante*. In the case of independence, the percentiles of the asymptotic distributions are known (Scholz and Stephens (1987), Table 1 with $m = k - 1$). The test statistic T_{k-1} converges to a standard normal for k tending to infinity.

The AD approach builds on tests previously proposed by Kolmogorov, Smirnov, Cramér, and von Mises. It adds a weight function to allow the researcher to attach differential importance to various portions of the distribution function (Anderson and Darling (1952)). It nests the Cramér-von Mises ω^2 statistic, which is based on equal weighting. The AD default weighting equalizes the sampling error across the (empirical) support of the distribution function (Anderson and Darling (1954), p. 767). It effectively attaches more weight to the tails of the distribution. Scholz and Stephens (1987, p. 919) argue that among alternatives, the AD test statistic has attractive small-sample (i.e., small k) properties.

REFERENCES

- Aczel, Balazs, Barnabas Szaszi, Gustav Nilsson, Olmo R. van den Akker, Casper J. Albers, Marcel A.L.M. van Assen, Jojanneke A. Bastiaansen, Dan Benjamin, Udo Boehm, Rotem Botvinik-Nezer, Laura F. Bringmann, Niko A. Busch, Emmanuel Caruyer, Andrea M. Cataldo, Nelson Cowan, Andrew Delios, Noah N.N. van Dongen, Chris Donkin, Johnny B. van Doorn, Anna Dreber, Gilles Dutilh, Gary F. Egan, Morton Ann Gernsbacher, Rink Hoekstra, Sabine Hoffmann, Felix Holzmeister, Juergen Huber, Magnus Johannesson, Kai J. Jonas, Alexander T. Kindel, Michael Kirchler, Yoram K. Kunkels, D. Stephen Lindsay, Jean-Francois Mangin, Dora Matzke, Marcus R. Munafò, Ben R. Newell, Brian A. Nosek, Russell A. Poldrack, Don van Ravenzwaaij, Jörg Rieskamp, Matthew J. Salganik, Alexandra Sarafoglou, Tom Schonberg, Martin Schweinsberg, David Shanks, Raphael Silberzahn, Daniel J. Simons, Barbara A. Spellman, Samuel St-Jean, Jeffrey J. Starns, Eric L. Uhlmann, Jelte Wicherts, and Eric-Jan Wagenmakers, 2021, Consensus-based guidance for conducting and reporting multi-analyst studies, *eLife* 10, 1–13.

³⁹ Note that the procedure accounts for within-RT correlations (i.e., including possible nonzero correlations among a particular RT's results and the beliefs that it reports). The reason the procedure accounts for these correlations is that the bootstrap uses block-sampling whereby, when an RT is drawn, all of its beliefs and all of its estimates are drawn. One therefore only assumes independence across RTs, which holds by construction given the design of #fincap.

- Anderson, Theodore W., and Donald A. Darling, 1952, Asymptotic theory of certain “goodness of fit” criteria based on stochastic processes, *Annals of Mathematical Statistics* 23, 193–212.
- Anderson, Theodore W., and Donald A. Darling, 1954, A test of goodness of fit, *Journal of the American Statistical Association* 49, 765–769.
- Angel, James J., Lawrence E. Harris, and Chester S. Spatt, 2015, Equity trading in the 21st century: An update, *Quarterly Journal of Finance* 5, 1–39.
- Bao, Jack, Maureen O’Hara, and Xing (Alex) Zhou, 2018, The Volcker rule and corporate bond market making in times of stress, *Journal of Financial Economics* 130, 95–113.
- Ben-David, Itzhak, Francesco Franzoni, and Byungwook Kim Rabi Moussawi, 2021, Competition for attention in the ETF space, Working paper, Ohio State University.
- Benjamin, Daniel J., James O. Berger, Magnus Johannesson, Brian A. Nosek, E.-J. Wagenmakers, Richard Berk, Kenneth A. Bollen, Björn Brembs, Lawrence Brown, Colin Camerer, David Cesarini, Christopher D. Chambers, Merlise Clyde, Thomas D. Cook, Paul De Boeck, Zoltan Dienes, Anna Dreber, Kenny Easwaran, Charles Efferson, Ernst Fehr, Fiona Fidler, Andy P. Field, Malcolm Forster, Edward I. George, Richard Gonzalez, Steven Goodman, Edwin Green, Donald P. Green, Anthony G. Greenwald, Jarrod D. Hadfield, Larry V. Hedges, Leonhard Held, Teck Hua Ho, Herbert Hoijtink, Daniel J. Hruschka, Kosuke Imai, Guido Imbens, John P.A. Ioannidis, Minjeong Jeon, James Holland Jones, Michael Kirchler, David Laibson, John List, Roderick Little, Arthur Lupia, Edouard Machery, Scott E. Maxwell, Michael McCarthy, Don A. Moore, Stephen L. Morgan, Marcus Munafó, Shinichi Nakagawa, Brendan Nyhan, Timothy H. Parker, Luis Pericchi, Marco Perugini, Jeff Rouder, Judith Rousseau, Victoria Savalei, Felix D. Schönbrodt, Thomas Sellke, Betsy Sinclair, Dustin Tingley, Trisha Van Zandt, Simine Vazire, Duncan J. Watts, Christopher Winship, Robert L. Wolpert, Yu Xie, Cristobal Young, Jonathan Zinman, and Valen E. Johnson, 2018, Redefine statistical significance, *Nature Human Behavior* 2, 6–10.
- Black, Bernard S., Hemang Desai, Kate Litvak, Woongsun Yoo, and Jeff Jiewei Yu, 2021, Specification choice in randomized and natural experiments: Lessons from the regulation sho experiment, Working paper, Northwestern University.
- Blume, Marshall E., 1974, Unbiased estimators of long-run expected rates of return, *Journal of the American Statistical Association* 69, 634–638.
- Bonferroni, Carlo E., 1936, Teoria statistica delle classi e calcolo delle probabilità, *Pubblicazioni del R Istituto Superiore di Scienze Economiche e Commerciali di Firenze* 8, 1–62.
- Botvinik-Nezer, Rotem, Felix Holzmeister, Colin F. Camerer, Anna Dreber, Juergen Huber, Magnus Johannesson, Michael Kirchler, Roni Iwanir, Jeanette A. Mumford, R. Alison Adcock, Paolo Avesani, Blazej M. Baczkowski, Aahana Bajracharya, Leah Bakst, Sheryl Ball, Marco Barilari, Nadège Bault, Derek Beaton, Julia Beitner, Roland G. Benoit, Ruud M.W.J. Berkers, Jamil P. Bhanji, Bharat B. Biswal, Sebastian Bobadilla-Suarez, Tiago Bortolin, Katherine L. Bottenhorn, Alexander Bowering, Senne Braem, Hayley R. Brooks, Emily G. Brudner, Cristian B. Calderon, Julia A. Camilleri, Jaime J. Castrellon, Luca Cecchetti, Edna C. Cieslik, Zachary J. Cole, Olivier Collignon, Robert W. Cox, William A. Cunningham, Stefan Czoschke, Kamalaker Dadi, Charles P. Davis, Alberto De Lucas, Mauricio R. Delgado, Lysia Demetriou, Jeffrey B. Dennison, Xin Di, Erin W. Dickie, Ekaterina Dobryakova, Claire L. Donnat, Juergen Dukart, Niall W. Duncan, Joke Durnez, Amr Eed, Simon B. Eickhoff, Andrew Erhart, Laura Fontanesi, G. Matthew Fricke, Shiguang Fu, Adriana Galván, Remi Gau, Sarah Genon, Tristan Glatard, Enrico Glerean, Jelle J. Goeman, Sergej A.E. Golowin, Carlos González-García, Krzysztof J. Gorgolewski, Cheryl L. Grady, Mikella A. Green, Joao F. Guassi Moreira, Olivia Guest, Shabnam Hakimi, J. Paul Hamilton, Roeland Hancock, Giacomo Handjaras, Bronson B. Harry, Colin Hawco, Peer Herholz, Gabrielle Herman, Stephan Heunis, Felix Hoffstaedter, Jeremy Hogeveen, Susan Holmes, Chuan-Peng Hu, Scott A. Huettel, Matthew E. Hughes, Vittorio Iacovella, Alexandru D. Iordan, Peder M. Isager, Ayse I. Isik, Andrew Jahn, Matthew R. Johnson, Tom Johnstone, Michael J.E. Joseph, Anthony C. Juliano, Joseph W. Kable, Michalis Kassinosopoulos, Cemal Koba, Xiang-Zhen Kong, Timothy R. Koscik, Nuri Erkut Kucukboyaci, Brice A. Kuhl, Sebastian Kupek, Angela R. Laird, Claus Lamm, Robert Langner, Nina Lauharatanahirun, Hongmi Lee, Sangil Lee, Alexander Leemans, Andrea Leo, Elise Lesage, Flora Li, Monica Y.C. Li, Phui Cheng Lim, Evan N. Lintz, Schuyler W. Liphardt,

- Annabel B. Losecaat Vermeer, Bradley C. Love, Michael L. Mack, Norberto Malpica, Theo Marins, Camille Maumet, Kelsey McDonald, Joseph T. McGuire, Helena Melero, Adriana S. Méndez Leal, Benjamin Meyer, Kristin N. Meyer, Glad Mihai, Georgios D. Mitsis, Jorge Moll, Dylan M. Nielson, Gustav Nilsson, Michael P. Notter, Emanuele Olivetti, Adrian I. Onicas, Paolo Papale, Kaustubh R. Patil, Jonathan E. Peelle, Alexandre Pérez, Doris Pischedda, Jean-Baptiste Poline, Yanina Prystauka, Shruti Ray, Patricia A. Reuter-Lorenz, Richard C. Reynolds, Emiliano Ricciardi, Jenny R. Rieck, Anais M. Rodriguez-Thompson, Anthony Romy, Taylor Salo, Gregory R. Samanez-Larkin, Emilio Sanz-Morales, Margaret L. Schlichting, Douglas H. Schultz, Qiang Shen, Margaret A. Sheridan, Jennifer A. Silvers, Kenny Skagerlund, Alec Smith, David V. Smith, Peter Sokol-Hessner, Simon R. Steinkamp, Sarah M. Tashjian, Bertrand Thirion, John N. Thorp, Gustav Tinghög, Loreen Tisdall, Steven H. Tompson, Claudio Toro-Serey, Juan Jesus Torre Tresols, Leonardo Tozzi, Vuong Truong, Luca Turella, Anna E. van 't Veer, Tom Verguts, Jean M. Vettel, Sagana Vijayarajah, Khoi Vo, Matthew B. Wall, Wouter D. Weeda, Susanne Weis, David J. White, David Wisniewski, Alba Xifra-Portas, Emily A. Yearling, Sangsuk Yoon, Rui Yuan, Kenneth S. L. Yuen, Lei Zhang, Xu Zhang, Joshua E. Zosky, Thomas E. Nichols, Russell A. Poldrack, and Tom Schonberg, 2020, Variability in the analysis of a single neuroimaging dataset by many teams, *Nature* 582, 84–88.
- Breznau, Nate, Eike Mark Rinke, Alexander Wuttke, Muna Adem, Jule Adriaans, Amalia Alvarez-Benjumea, Henrik K. Andersen, Daniel Auer, Flavio Azevedo, Oke Bahnsen, Dave Balzer, Gerrit Bauer, Paul C. Bauer, Markus Baumann, Sharon Baute, Verena Benoit, Julian Bernauer, Carl Berning, Anna Berthold, Felix S. Bethke, Thomas Biegert, Katharina Blinzler, Johannes N. Blumenberg, Licia Bobzien, Andrea Bohman, Thijs Bol, Amie Bostic, Zuzanna Brzozowska, Katharina Burgdorf, Kaspar Burger, Kathrin Busch, Juan Carlos-Castillo, Nathan Chan, Pablo Christmann, Roxanne Connelly, Christian S. Czymara, Elena Damian, Alejandro Ecker, Achim Edelmann, Maureen A. Eger, Simon Ellerbrock, Anna Forke, Andrea Forster, Chris Gaasendam, Konstantin Gavras, Vernon Gayle, Theresa Gessler, Timo Gnambs, Amélie Godefroidt, Max Grömping, Martin Groß, Stefan Gruber, Tobias Gummer, Andreas Hadjar, Jan Paul Heisig, Sebastian Hellmeier, Stefanie Heyne, Magdalena Hirsch, Mikael Hjerm, Oshrat Hochman, Andreas Hövermann, Sophia Hunger, Christian Hunkler, Nora Huth, Zsófia S. Ignácz, Laura Jacobs, Jannes Jacobsen, Bastian Jaeger, Sebastian Jungkunz, Nils Jungmann, Mathias Kauff, Manuel Kleinert, Julia Klinger, Jan-Philipp Kolb, Marta Kotczyńska, John Kuk, Katharina Kunißen, Dafina Kurti Sinatra, Alexander Langenkamp, Philipp M. Lersch, Lea-Maria Löbel, Philipp Lutscher, Matthias Mader, Joan E. Madia, Natalia Malancu, Luis Maldonado, Helge-Johannes Marahrens, Nicole Martin, Paul Martinez, Jochen Mayerl, Oscar J. Mayorga, Patricia McManus, Kyle McWagner, Cecil Meeusen, Daniel Meierrieks, Jonathan Mellon, Friedolin Merhout, Samuel Merk, Daniel Meyer, Leticia Micheli, Jonathan Mijs, Cristóbal Moya, Marcel Neunhoeffer, Daniel Nüst, Olav Nygård, Fabian Ochsenfeld, Gunnar Otte, Anna Pechenkina, Christopher Prosser, Louis Raes, Kevin Ralston, Miguel Ramos, Arne Roets, Jonathan Rogers, Guido Ropers, Robin Samuel, Gregor Sand, Ariela Schachter, Merlin Schaeffer, David Schieferdecker, Elmar Schlueter, Regine Schmidt, Katja M. Schmidt, Alexander Schmidt-Catran, Claudia Schmiedeberg, Jürgen Schneider, Martijn Schoonvelde, Julia Schulte-Cloos, Sandy Schumann, Reinhard Schunck, Jürgen Schupp, Julian Seuring, Henning Silber, Willem Slegers, Nico Sonntag, Alexander Staudt, Nadia Steiber, Nils Steiner, Sebastian Sternberg, Dieter Stiers, Dragana Stojmenovska, Nora Storz, Erich Striessnig, Anne-Kathrin Stroppe, Janna Teltemann, Andrey Tibajev, Brian Tung, Giacomo Vagni, Jasper Van Assche, Meta van der Linden, Jolanda van der Noll, Arno Van Hootegem, Stefan Vogtenhuber, Bogdan Voicu, Fieke Wagemans, Nadja Wehl, Hannah Werner, Brenton M. Wiernik, Fabian Winter, Christof Wolf, Yuki Yamada, Nan Zhang, Conrad Ziller, Stefan Zins, Tomasz Żółtak, and Hung H.V. Nguyen, 2021, Observing many researchers using the same data and hypothesis reveals a hidden universe of uncertainty, Working paper, University of Bremen.
- Breznau, Nate, Deadric T. Williams, Katrin Auspurg, Josef Bruederl, Felix Holzmeister, Gustav Nilsson, Balázs Aczel, Cory J. Clark, and Eric Uhlmann, 2022, Is the variability in social science research results produced by different researchers explicable? An adversarial

- collaboration and joint effort to parse the dispersion in estimates, Working paper, University of Bremen.
- Camerer, Colin F., Anna Dreber, Eskil Forsell, Teck-Hua Ho, Jürgen Huber, Magnus Johannesson, Michael Kirchler, Johan Almenberg, Adam Altmejd, Taizan Chan, Emma Heikensten, Felix Holzmeister, Taisuke Imai, Siri Isaksson, Gideon Nave, Thomas Pfeiffer, Michael Razen, and Hang Wu, 2016, Evaluating replicability of laboratory experiments in economics, *Science* 351, 1433–1436.
- Camerer, Colin F., Anna Dreber, Felix Holzmeister, Teck-Hua Ho, Jürgen Huber, Magnus Johannesson, Michael Kirchler, Gideon Nave, Brian A. Nosek, Thomas Pfeiffer, Adam Altmejd, Nick Buttrick, Taizan Chan, Yiling Chen, Eskil Forsell, Anup Gampa, Emma Heikensten, Lily Hummer, Taisuke Imai, Siri Isaksson, Dylan Manfredi, Julia Rose, Eric-Jan Wagenmakers, and Hang Wu, 2018, Evaluating the replicability of social science experiments in nature and science, *Nature Human Behaviour* 2, 637–644.
- Chen, Andrew Y., 2021, The limits of p -hacking: Some thought experiments, *Journal of Finance* 76, 2447–2480.
- Chordia, Tarun, Amit Goyal, and Alessio Saretto, 2020, Anomalies and false rejections, *Review of Financial Studies* 33, 2134–2179.
- Colliard, Jean-Edouard, Christophe Hurlin, and Christophe Pérignon, 2021, The economics of research reproducibility, Working paper, HEC Paris.
- Gelman, Andrew, and Eric Loken, 2014, The statistical crisis in science, *American Scientist* 102, 460–465.
- Geyer-Klingeborg, Jerome, Markus Hang, and Andreas Rathgeber, 2020, Meta-analysis in finance research: Opportunities, challenges, and contemporary applications, *International Review of Financial Analysis* 7, 1–15.
- Harvey, Campbell R., 2017, Presidential address: The scientific outlook in financial economics, *Journal of Finance* 72, 1399–1440.
- Harvey, Campbell R., and Yan Liu, 2020, False (and missed) discoveries in financial economics, *Journal of Finance* 75, 2503–2553.
- Harvey, Campbell R., Yan Liu, and Heqing Zhu, 2016, ...and the cross-section of expected returns, *Review of Financial Studies* 29, 5–68.
- Hendershott, Terrence, and Albert J. Menkveld, 2014, Price pressures, *Journal of Financial Economics* 114, 405–423.
- Hou, Kewei, Chen Xue, and Lu Zhang, 2018, Replicating anomalies, *Review of Financial Studies* 33, 2019–2133.
- Huntington-Klein, Nick, Andreu Arenas, Emily Beam, Marco Bertoni, Jeffrey R. Bloem, Pralhad Burli, Naibin Chen, Paul Grieco, Godwin Ekpe, Todd Pugatch, Martin Saavedra, and Yani Stopnitzky, 2021, The influence of hidden researcher decisions in applied microeconomics, *Economic Inquiry* 59, 944–960.
- Ioannidis, John P.A., 2005, Why most published research findings are false, *PLoS Medicine* 2, 696–701.
- Jensen, Theis Ingerslev, Bryan Kelly, and Lasse Pedersen, 2023, Is there a replication crisis in finance, *Journal of Finance* 78, 2465–2518.
- Jones, Charles M., 2013, What do we know about high-frequency trading? Working paper, Columbia University.
- Jovanovic, Boyan, and Albert J. Menkveld, 2022, Equilibrium bid-price dispersion, *Journal of Political Economy* 130, 426–461.
- Kahneman, Daniel, Olivier Sibony, and Cass R. Sunstein, 2021, *Noise: A Flaw in Human Judgment* (William Collins, London).
- Kirchler, Michael, Florian Lindner, and Utz Weitzel, 2018, Rankings and risk-taking in the finance industry, *Journal of Finance* 73, 2271–2302.
- Koenker, Roger, and Gilbert Bassett Jr. 1978, Regression quantiles, *Econometrica* 46, 33–50.
- Köszegi, Botond, 2006, Ego utility, overconfidence, and task choice, *Journal of the European Economic Association* 4, 673–707.
- Leamer, Edward E., 1983, Let's take the con out of econometrics, *American Economic Review* 73, 31–43.

- Linnainmaa, Juhani T., and Michael R. Roberts, 2018, The history of the cross-section of stock returns, *Review of Financial Studies* 31, 2606–2649.
- Liu, Yang, Alex Kale, Tim Althoff, and Jeffrey Heer, 2021, Boba: Authoring and visualizing multi-verse analyses, *IEEE Transactions on Visualization and Computer Graphics* 27, 1753–1763.
- Mavroeidis, Sophocles, Mikkel Plagborg-Møller, and James H. Stock, 2014, Empirical evidence on inflation expectations in the new Keynesian Phillips curve, *Journal of Economic Literature* 52, 124–188.
- McLean, R. David, and Jeffrey Pontiff, 2016, Does academic research destroy stock return predictability?, *Journal of Finance* 71, 5–31.
- Menkveld, Albert J., 2016, The economics of high-frequency trading: Taking stock, *Annual Review of Financial Economics* 8, 1–24.
- Mitton, Todd, 2022, Methodological variation in empirical corporate finance, *Review of Financial Studies* 35, 527–575.
- Munafò, Marcus R., Brian A. Nosek, Dorothy V.M. Bishop, Katherine S. Button, Christopher D. Chambers, Nathalie Percie du Sert, Uri Simonsohn, Eric-Jan Wagenmakers, Jennifer J. Ware, and John P.A. Ioannidis, 2017, A manifesto for reproducible science, *Nature Human Behaviour* 1, 1–9.
- Open Science Collaboration, 2015, Estimating the reproducibility of psychological science, *Science* 349, 943–951.
- Parlour, Christine A., and Duane J. Seppi, 2008, Limit order markets: A survey, in Arnoud W.A. Boot, and Anjan V. Thakor, eds., *Handbook of Financial Intermediation and Banking* (Elsevier Publishing, Amsterdam, Netherlands).
- Pastor, Lubos, and Robert F. Stambaugh, 2003, Liquidity risk and expected returns, *Journal of Political Economy* 111, 642–685.
- Pérignon, Christophe, Olivier Akmansory, Christophe Hurlin, Anna Dreber, Felix Holzmeister, Jürgen Huber, Magnus Johannesson, Michael Kirchler, Albert J. Menkveld, Michael Razen, and Utz Weitzel, 2023, Computational reproducibility in finance: Evidence from 1,000 tests, Working paper, HEC Paris.
- Scholz, Fritz W., and Michael A. Stephens, 1987, K-sample Anderson-Darling tests, *Journal of the American Statistical Association* 82, 918–924.
- Schweinsberg, Martin, Michael Feldman, Nicola Staub, Olmo R. van den Akker, Robbie C.M. van Aert, Marcel A.L.M. van Assen, Yang Liu, Tim Althoff, Jeffrey Heer, Alex Kale, Zainab Mohamed, Hashem Amireh, Vaishali Venkatesh Prasad, Abraham Bernstein, Emily Robinson, Kaisa Snellman, S. Amy Sommer, Sarah M.G. Otner, David Robinson, Nikhil Madan, Raphael Silberzahn, Pavel Goldstein, Warren Tierney, Toshio Murase, Benjamin Mandl, Domenico Viganola, Carolin Strobl, Catherine B.C. Schaumans, Stijn Kelchtermans, Chan Naseeb, S. Mason Garrison, Tal Yarkoni, C.S. Richard Chan, Prestone Adie, Paulius Alaburda, Casper Albers, Sara Alspaugh, Jeff Alstott, Andrew A. Nelson, Eduardo Ariño de la Rubia, Adbi Arzi, Štěpán Bahník, Jason Baik, Laura Winther Balling, Sachin Banker, David AA Baranger, Dale J. Barr, Brenda Barros-Rivera, Matt Bauer, Enuh Blaise, Lisa Boelen, Katerina Bohle Carbonell, Robert A. Briers, Oliver Burkhard, Miguel-Angel Canela, Laura Castrillo, Timothy Catlett, Olivia Chen, Michael Clark, Brent Cohn, Alex Coppock, Natália Cugueró-Escofet, Paul G. Curran, Wilson Cyrus-Lai, David Dai, Giulio Valentino Dalla Riva, Henrik Danielsson, Rosaria de F.S.M. Russo, Niko de Silva, Curdin Derungs, Frank Dondelinger, Carolina Duarte de Souza, B. Tyson Dube, Marina Dubova, Ben Mark Dunn, Peter Adriaan Edelsbrunner, Sara Finley, Nick Fox, Timo Gnamb, Yuan Yuan Gong, Erin Grand, Brandon Greenawalt, Dan Han, Paul H.P. Hanel, Antony B. Hong, David Hood, Justin Hsueh, Lilian Huang, Kent N. Hui, Keith A. Hultman, Azka Javaid, Lily Ji Jiang, Jonathan Jong, Jash Kamdar, David Kane, Gregor Kappler, Erikson Kaszubowski, Christopher M. Kavanagh, Madian Khabsa, Bennett Kleinberg, Jens Kourous, Heather Krause, Angelos-Miltiadis Kryptos, Dejan Lavbič, Rui Ling Lee, Timothy Leffel, Wei Yang Lim, Silvia Liverani, Bianca Loh, Dorte Lønsmann, Jia Wei Low, Alton Lu, Kyle MacDonald, Christopher R. Madan, Lasse Hjorth Madsen, Christina Maimone, Alexandra Mangold, Adrienne Marshall, Helena Ester Matskewich, Kimia Mavon, Katherine L. McLain, Amelia A. McNamara, Mhairi McNeill, Ulf Mertens, David Miller, Ben Moore, Andrew Moore, Eric

- Nantz, Ziauddin Nasrullah, Valentina Nejkovic, Colleen S. Nell, Andrew Arthur Nelson, Gustav Nilsson, Rory Nolan, Christopher E. O'Brien, Patrick O'Neill, Kieran O'Shea, Toto Olita, Jahna Otterbacher, Diana Palsetia, Bianca Pereira, Ivan Pozdnyakov, John Protzko, Jean-Nicolas Rey, Travis Riddle, Amal (Akmal) Ridhwan Omar Ali, Ivan Ropovik, Joshua M. Rosenberg, Stephane Rothen, Michael Schulte-Mecklenbeck, Nirek Sharma, Gordon Shotwell, Martin Skarzynski, William Stedden, Victoria Stodden, Martin A. Stoffel, Scott Stoltzman, Subashini Subbaiah, Rachael Tatman, Paul H. Thibodeau, Sabina Tomkins, Ana Valdivia, Gerrieke B. Druiff-van de Woestijne, Laura Viana, Florence Villessèche, W. Duncan Wadsworth, Florian Wanders, Krista Watts, Jason D. Wells, Christopher E. Whelpley, Andy Won, Lawrence Wu, Arthur Yip, Casey Youngflesh, Ju-Chi Yu, Arash Zandian, Leilei Zhang, Chava Zibman, and Eric Luis Uhlmann, 2021, Same data, different conclusions: Radical dispersion in empirical results when independent analysts operationalize and test the same hypothesis, *Organizational Behavior and Human Decision Processes* 165, 228–249.
- Šidák, Zbyně, 1967, Rectangular confidence regions for the means of multivariate normal distributions, *Journal of the American Statistical Association* 62, 626–633.
- Silberzahn, Raphael, Eric L. Uhlmann, Dan P. Martin, Pasquale Anselmi, Frederik Aust, Eli Awtrey, Štěpán Bahník, Feng Bai, Colin Bannard, Evelina Bonnier, Rickard Carlsson, Felix Cheung, Garret Christensen, Russ Clay, Maureen A. Craig, Anna Dalla Rosa, Lammertjan Dam, Mathew H. Evans, Ismael Flores Cervantes, Nathan Fong, Monica Gamez-Djokic, Andreas Glenz, Shauna Gordon-McKeon, Tim J. Heaton, Karin Hederes, Moritz Heene, Alicia J. Hofelich Mohr, Fabia Högden, Kent Hui, Magnus Johannesson, Jonathan Kalodimos, Erikson Kaszubowski, Deanna M. Kennedy, Ryan Lei, Thomas A. Lindsay, Silvia Liverani, Christopher R. Madan, Daniel Molden, Eric Molleman, Richard D. Morey, Laetitia B. Mulder, Bernard R. Nijstad, Nolan G. Pope, Bryson Pope, Jason M. Prenoveau, Floor Rink, Egidio Robusto, Hadiya Roderique, Anna Sandberg, Elmar Schlüter, Felix D. Schönbrodt, Martin F. Sherman, S. Amy Sommer, Kristin Sotak, Seth Spain, Christoph Spörlein, Tom Stafford, Luca Stefanutti, Susanne Tauber, Johannes Ullrich, Michelangelo Vianello, Eric-Jan Wagenmakers, Maciej Witkowiak, Sanguk Yoon, and Brian A. Nosek, 2018, Many analysts, one data set: Making transparent how variations in analytic choices affect results, *Advances in Methods and Practices in Psychological Science* 1, 337–356.
- Soebhag, Amar, Bart van Vliet, and Patrick Verwijmeren, 2023, Non-standard errors in asset pricing: Mind your sorts, Working paper, Erasmus University Rotterdam.
- Steege, Sara, Francis Tuerlinckx, Andrew Gelman, and Wolf Vanpaemel, 2016, Increasing transparency through a multiverse analysis, *Psychological Science* 11, 702–712.
- Tran, Anh, and Richard Zeckhauser, 2012, Rank as an inherent incentive: Evidence from a field experiment, *Journal of Public Economics* 96, 645–650.
- Walter, Dominik, Rüdiger Weber, and Patrick Weiss, 2023, Non-standard errors in portfolio sorts, Working paper, Vienna Graduate School of Finance.
- Yule, G. Udny, 1897, On the significance of Bravais' formulae for regression, *Proceedings of the Royal Society of London* 60, 477–489.

Supporting Information

Additional Supporting Information may be found in the online version of this article at the publisher's website:

**Appendix S1: Internet Appendix.
Replication Code.**