

The role of analysts in unregulated financial markets: Evidence from Initial Coin Offerings *

Andreas Barth, Valerie Laturus, Sasan Mansouri and Alexander F. Wagner[†]

September 15, 2021

ABSTRACT

Initial Coin Offerings (ICOs) offer a potentially promising way for funding new ventures, but the market has imploded after an initial “hot” phase. This paper exploits the rich data available from this market to study the contribution of analysts to the functioning and failure of unregulated capital markets. The assessments of freelancing ICO analysts varies in quality and exhibits biases due to the reciprocal interactions of analysts with ICO team members. Ratings predict ICO success, but imperfectly. Even favorably rated ICOs tend to fail when a greater portion of their ratings reciprocate prior ratings, and the market capitalization 90 days after listing on an exchange is smaller for tokens with more reciprocal ratings. These findings suggest that the failure of ICOs is not uniform but is related to measures of conflicts of interest. Thus, information about the track record of analysts and their potentially conflicting activities proves to be valuable to investors. These results help assess the need for regulation of more recent forms of blockchain-based and decentralized finance.

JEL classification: G14, G24, L26, D82, D83

Keywords: Analysts, Asymmetric Information, Blockchains, FinTech, Initial Coin Offering (ICO)

*We are grateful to Christian Leuz, Evgeny Lyandres, Arthur Petit-Romec, and David Yermack for comments. Barth acknowledges financial support from the Chaire Fintech at University Paris Dauphine - PSL. Wagner acknowledges financial support from the University of Zurich Research Priority Program “Financial market regulation”. We declare that we have no relevant or material financial interests that relate to the research described in this paper.

[†]Barth, Laturus and Mansouri are with Goethe University Frankfurt, Wagner is with University of Zurich, CEPR, ECGI and SFI. ✉: Andreas.Barth@finance.uni-frankfurt.de ;Laturus@finance.uni-frankfurt.de ;mansouri@finance.uni-frankfurt.de ;alexander.wagner@bf.uzh.ch

1 Introduction

The question of how analysts contribute to the functioning of capital markets has been on the agenda of accounting and finance research for years (Bradshaw et al., 2017). While professional analysts in traditional financial markets are heavily regulated, little is known about the role of free-lancing analysts in unregulated financial markets. This gap is striking given the recently renewed boom of blockchain-based markets and decentralized finance (DeFi).

This paper uses the setting of Initial Coin Offerings (ICOs) – an unregulated financial market that experienced a massive rise and fall in the late 2010s – to investigate determinants and consequences of the quantitative and qualitative aspects of investment ratings issued by human experts (henceforth ICO analysts). Ratings predict the success of ICOs, but imperfectly. Even among ICOs with an average rating in the top quartile, more than 50% fail. The analysis in particular addresses potential conflicts of interest in ICO analyst ratings. We find that ICO analysts tend to reciprocate favorable ratings for their own ventures; however, the results also suggest that investors place lower emphasis on reciprocating ratings.

Initial Coin Offerings (ICOs) are token sale events on an own or existing blockchain that facilitate financing for an entrepreneurial venture. ICOs experienced an enormous boom in 2017-2018, but the volume of the market has declined massively since then. Tokens offerings are a potentially powerful instrument for new ventures to obtain crowdfunding-like resources (Goldstein et al., 2019; Lyandres, 2019; Chod and Lyandres, 2020; Lee and Parlour, 2020; Li and Mann, 2020; Lyandres et al., 2020; Gryglewicz et al., 2020). However, despite all the promises, the ICO market failed.

Understanding the workings and failure of this relatively new market and in particular studying the cross-section of ICOs and their analysts is of particular interest for at least three reasons. First, the ICO environment provides a relatively clean setup to investigate how analysts contribute to capital markets. The market is particularly interesting to study the role of information intermediaries because its regulation has only recently begun to clar-

ify. Initial Exchange Offerings (IEOs) and Security Token Offering (STOs) emerged recently as alternatives to ICOs. STOs need to be registered and approved by the U.S. Securities and Exchange Commission (SEC), but like ICOs they offer little investor protection. Understanding which ICOs failed despite potential monitoring by human professionals and the possible concomitant market discipline is important to clarify the motivation for further regulation of these newer versions of FinTech markets.

Second, like financial analysts, ICO analysts potentially suffer from conflicts of interests.¹ However, the conflicts in this case (i) are potentially more extreme and (ii) can be more directly identified than in the case of the typical security analysts. As for (i), ICO analysts do not only *provide ratings* for ICOs, but may also *run their own* ICOs. Thus, whenever an ICO analyst i provides a rating for an ICO j , there is a chance that a team member of this ICO j will rate for the ICO of analyst i at a later stage. As for (ii), most of the literature on financial analysts classifies analysts as “affiliated” (and thus potentially conflicted) if they belong to a bank that has or applies for an underwriting relationship with the firms on which they are reporting or if analysts want to get hired by the firm they analyze (“revolving door analysts”). These potential biases are largely hidden information, and particularly revolving door analysts can only be identified ex post their job change. By contrast, the ICO setting presents a situation where linkages are more direct and where investors can be aware of potential biases right away.

Third, non-professional analysts and their crowd forecasts have been shown to be important information intermediaries for equity investors (Chen et al., 2014; Jame et al., 2016; Drake et al., 2017; Campbell et al., 2019; Farrell et al., 2020; Da and Huang, 2020). However, we know little about the potential conflicts of interest that such analysts face and whether market participants consider the differential credibility and informativeness of these analyses in their investment decisions.

We collect data on 5,384 ICOs between 2017 and 2020 from the platform ICObench.com.

¹See, for example, Lin and McNichols (1998), Michaely and Womack (1999), and Chan et al. (2007) for evidence of biased financial analysts.

According to the web traffic statistics from Alexa Internet, ICObench.com was an important source of rating information for investors and was able to achieve a site visit rank of 3,644 during the peak of the ICO market (compared to a site visit rank of around 2,200 for the Financial Times, in the same period).² We identify 531 experts who issued a total of 13,834 ratings. Figure 1 illustrates some main results using binned scatter plots.

[Figure 1 about here]

We begin by investigating determinants of analysts’ ratings. We first find that analysts who, in the past, had issued very positive ratings for ICOs that did not succeed (i.e., analysts with large forecast errors) provide, on average, lower ratings in the future (see Panel (a) of Figure 1). ICObench.com also rank ICO analysts, which gives an equivalent setting to all-star financial analysts (Leone and Wu, 2002). We observe that “star analysts” are less optimistic and their ratings are, on average, lower. In addition to quantitative ratings, we also consider the length and linguistic tone of the reviews that accompany the evaluation, i.e., the qualitative nature of ICO analyst ratings. We observe that lower ratings often accompany longer reviews with a more negative tone. In all these analyses, we compare different ratings for the same ICO, which helps to rule out that these results are purely driven by the self-selection of analysts to certain ICOs.

Importantly, reciprocal ratings are special (see Panel (b) of Figure 1): the total rating score an analyst gives to an ICO j is higher if she received a rating in the past for her own ICO by any team member of coin j . This effect is stronger the higher the prior received rating was. These effects continue to hold when we compare analysts providing a rating to the same ICO in a given month. Comparing different assessments of the same analyst and for the same ICO allows also to rule out that the optimistic assessment is due to the high difficulty of forecasting tasks or due to a non-random match between founders of good ICOs that also serve as analysts.

²Alexa Internet identifies “ICO rating” as the main ‘Buyer Keyword’ for ICObench.com, that is, those people who were searching to buy a product or service landing on ICObench.com searched primarily for “ICO rating”.

Next, we analyze the explanatory power of ICO analyst ratings for the success of an ICO campaign, and for the failure of an ICO campaign despite strong analyst endorsement. We first confirm the result of prior work that investors appear to value the fact that a human analyst provided a rating for the ICO.³ Moreover, a better average quantitative ICO analyst rating translates into a higher probability that the ICO offering is completed and received funding and into a higher market capitalization 90 days after listing on an exchange.

However, while the unconditional failure rate of ICOs is about 64%, even 53.6% of ICOs with an average rating in the top quartile fail. Our main interest is in the characteristics of analysts or the ICO itself that lead to such disagreement between analysts’ advice and the market outcome.

The share of reciprocal ratings is an important determinant of failure despite high ratings: If ICO j receives a rating from many reciprocal analysts, i.e., analysts whose rating is a response to a rating they received from a team member of ICO j , the market is more likely to disregard analyst recommendations; see Panel (c) of Figure 1. Moreover, even among successful ICOs, the market capitalization 90 days after listing on an exchange is smaller for ICOs with more reciprocal ratings. There are two possible interpretations of this result. First, it is conceivable that, even though we control for a wide variety of factors presumably capturing variation in ICO quality, reciprocal ratings occur with “objectively” bad ICOs, i.e., they pick up some additional variation in quality. Second, investors may trust ICOs with more reciprocal ratings less (even when they may potentially be worth funding).⁴ Either way, the findings imply that investors do not blindly pile capital into highly rated ICOs.

Interesting patterns also emerge for the linguistic measures of the rating. The length and linguistic tone of the reviews that accompany the evaluation explain only little of the variation in the success of ICOs. However, the likelihood that an ICO fail despite high ratings increases with the positivity of the tone and complexity of the language in the reviews.

³This is in line with several studies that document the benefits of analyst coverage (Sufi, 2009; Demiroglu and Ryngaert, 2010; Mola et al., 2013; Crawford et al., 2012).

⁴Several studies discuss whether or not investors are sophisticated enough to detect biased ratings (Ellis, 1998; Baker and Mansi, 2002; Livingston et al., 2010; Hirth, 2014; Badoer et al., 2019).

Finally, the quantitative and qualitative ratings by human analysts do not systematically differ on average for ICOs that prove to be fraudulent. A higher share of reciprocal ratings is not associated with a higher fraud probability, suggesting that criminal intentions do not typically drive reciprocity. ICOs exhibiting fraud do show a larger dispersion of both rating scores and tone of rating reviews among analysts.

Overall, the results suggest that, the failure of ICOs was not uniform but was related to measures of conflicts of interest. Having access to information about the track record and potentially conflicting activities of analysts allowed ICO investors to respond to qualitative differences among analyst ratings in a differentiated way. Even easier access would arguably have further enhanced efficiency of capital allocation in this market. Information intermediaries and platforms collecting data about crypto analysts play an important role in the functioning of market discipline in unregulated markets.⁵

These results add to the literature in four important ways. First, the literature on financial analysts suggests that a close link between analysts and firms leads to superior information and better assessments (Bradley et al., 2017; Bae et al., 2008), but also highlights the problem of conflicts of interest in a similar spirit of “affiliated” analysts (e.g. Lin and McNichols, 1998; Michaely and Womack, 1999; O’Brien et al., 2005; Malmendier and Shanthikumar, 2007; Agrawal and Chen, 2008; Kadan et al., 2009) or revolving door analysts (Lourie, 2019; Kempf, 2020).⁶ However, there is scarce data on direct interactions of analysts with the firms they analyze. The data on ICO analysts provide distinct advantages in that respect, and by showing that investors do take differences among analysts into account we highlight that these data are of value to investors.

Second, the paper complements the literature on semi-professional analysts in equity

⁵Transparency about the background, characteristics, or track records of information providers and intermediaries has been identified as critically important in other settings. For example, Law and Mills (2019) highlight the importance of the transparency provided by the Financial Industry Regulatory Authority (FINRA) about brokers’ (criminal) backgrounds. In academia, too, market discipline based on transparent disclosure can work. For example, Leuz et al. (2021) find that medical scholars cite papers reporting research sponsored by drug companies less frequently.

⁶A similar conflict of interest is present for rating agencies (e.g. Bolton et al., 2012; Bar-Isaac and Shapiro, 2013; Baghai and Becker, 2017; Chu and Rysman, 2019).

markets (Chen et al., 2014; Drake et al., 2017). That literature recognizes the possibility of conflicts of interest if the semi-professional analyst is holding positions on the stock themselves, resulting in a subjective, distorted analysis (Campbell et al., 2019).⁷ While these studies focus on equity markets in which semi-professional analysts complement the information produced by professional analysts, one particular advantage of the ICO market, besides very detailed and structured information, is the absence of professional analysts.⁸

Third, the paper adds to the growing literature on the relationship between machine-generated evaluations and human expert ratings.⁹ In addition to human evaluations, many platforms set up machine-generated ratings. These ratings do not evaluate the content of an ICO, but are based on observable factors such as features of the ICO’s campaign and team.¹⁰ Importantly, we show that both ratings are informative about ICO success. However, many ICOs fail despite high ratings, which is why we analyze this discrepancy.

Finally, ICOs are (or were) a potentially powerful way to fund new ventures, not least because of the underlying distributed ledger-based technology and the platform’s special features (Bakos and Halaburda, 2019; Biais et al., 2019; Cong and He, 2019; Easley et al., 2019; Hinzen et al., 2020). This paper advances our knowledge of the failure of the ICO market. Usually, the sale of tokens or ICOs appear at a very early planning stage of a product’s or a firm’s life cycle and suffer from severe information asymmetry and adverse selection problems (Malinova and Park, 2018; Chod and Lyandres, 2020; Chod et al., 2020; Gan et al., 2020). As such, tokens have no intrinsic value at the time of the investment. Instead, they rather derive value from trust in future usage (Conley, 2017). Hence, the

⁷Campbell et al. (2019) use non-professional analysts’ disclosures of stock positions as an indicator of the analyst’s position, which may not be reported truthfully.

⁸There are of course many stocks that professional analysts do not cover. This lack of coverage is the analysts’ choice, however, and as such provides information to the market.

⁹For example, Aubry et al. (2019) use data on paintings auctioned to study the accuracy and usefulness of valuations generated by using a pricing algorithm based on neural networks. With data from a leading startup accelerator, Catalini et al. (2018) show that artificial intelligence can help humans to screen and evaluate information when there is an information overload.

¹⁰Automated algorithms that simply count disclosed information are usually applied. For example, a high number of social media platforms on which an ICO is present or being listed on several rating websites automatically improves the rating for the respective ICO (Boreiko and Vidusso, 2019).

literature has investigated both the supply side, i.e., choices by ICO entrepreneurs (Adhami et al., 2018; Amsden and Schweizer, 2018; Benedetti and Kostovetsky, 2018; Cerchiello and Toma, 2018; Roosenboom et al., 2020; Howell et al., 2020; Fisch, 2019; Ernst and Young, 2018; Chakraborty and Swinney, 2020; PwC, 2019; Deng et al., 2018; Davydiuk et al., 2019), and the demand side, i.e. choices by investors (Fahlenbrach and Frattaroli, 2020; Fisch et al., 2019; Fisch and Momtaz, 2020).¹¹ Little attention has been paid to the information providing intermediaries in between supply and demand, however, and the literature largely focuses on the governance role of whitepapers provided by the ICO team (Adhami et al., 2018; Feng et al., 2019; Giudici and Adhami, 2019; Zhang et al., 2019; Samieifar and Baur, 2020; Florysiak and Schandlbauer, 2019; Zetzsche et al., 2019).

To the best of our knowledge, only three previous papers examine ICO analysts (Aggarwal et al., 2019; Bourveau et al., 2019; Lee et al., 2019).¹² All three document that ICOs with higher expert assessments are more successful. Our baseline results confirm this finding, but we focus on the striking fact more than 50% of the ICOs with the highest quartile of ratings fail. We show that accounting for the heterogeneity among analysts is important. For example, we exploit the specific feature of the market that ICO analysts provide ICO ratings, while also often running their own ICOs. We show that reciprocal ratings are biased, but also that investors discount such reciprocal ratings. We also uncover several other predictors of failure despite praise from analysts.

The rest of the paper is organized as follows. Section 2 presents the data and descriptive statistics. Section 3 describes the results, and Section 4 concludes.

¹¹There is also a literature on price dynamics of tokens (Cong et al., 2020a,b; Lee and Parlour, 2020; Li and Mann, 2020) as well as studies of asset pricing properties of coins on secondary markets and post-ICO performance (Dittmar and Wu, 2019; Hu et al., 2019; Fisch and Momtaz, 2020; Lyandres et al., 2020). See Li and Mann (2019) for a review of recent literature advances in ICO research.

¹²Momtaz (2020) uses information on ICO analysts, too. However, the paper does not focus on ICO analysts per se, but uses their evaluations as a proxy for one dimension of project quality.

2 Data and descriptive statistics

2.1 Sample and data source

We collect data on ICOs, ICO ratings and ICO experts from the platform ICObench.com. Our sample consists of 5,384 ICOs (of which 2,378 were rated by at least one expert) and spans the time period from the start of ICObench.com in 2017 to February 2020. ICOs in our sample were launched in 127 different countries, of which the USA, Singapore and the UK have the highest market shares.¹³

2.2 ICO analysts

In contrast to regulated financial analysts, ICO analysts are not certified. However, they have to apply for expert status on a platform, in our case ICObench.com. In their application, experts are required to describe their level of experience in crypto assets and motives to rating ICOs. The platform confirms the analysts after reviewing their credentials. The selection is relatively stringent. As of March 2020, the ICObench.com platform hosts more than 111,000 community members of which only 531 have expert status and thus the ability to provide ratings.

ICObench.com ranks the analysts based on several factors like profile completeness and analysts' consistency in contributions to the platform.¹⁴ This provides an analogy to the widely used all-star rankings of financial analysts. We collect these rankings over time and flag whether an analyst is among the top 30 analysts, i.e., within approximately the top 5%. The dummy variable $StarAnalysts_{ij}$ equals one if analyst i is listed among the top 30 list

¹³The compiled dataset is of comparable size to data used in other empirical ICO studies. For example, Benedetti and Kostovetsky (2018) use a sample of 4,003 ICO campaigns from five websites. Most information was retrieved from ICObench.com and ICOrating.com. Florysiak and Schandlbauer (2019) analyze 4,053 ICOs from ICObench.com. Deng et al. (2018) hand-collected a sample of 4,489 ICOs. Recently, Lyandres et al. (2020) cover the largest data set from the ICO universe with 7,514 ICO projects merged from various websites. Note that our sample period also covers the time after the collapse of the ICO market.

¹⁴The expert weight is calculated based on a profile score, a rating score, a time score, an acceptance score, and a contribution score. See <https://icobench.com/faq> for a detailed description.

prior to evaluating ICO j .¹⁵

Interestingly, many ICO analysts are involved in one or more ICO campaigns themselves.¹⁶ This dependent network structure of analysts offers a unique setting for investigating the role of human experts in crowdfunding markets (see Section 2.4).

2.3 Ratings

We identify 531 experts on ICObench.com who rate for 2,378 ICOs. Each analyst rate an average of 29.64 ICOs, resulting in a total number of 13,834 ratings. Experts can provide a rating for three dimensions of an ICO - team, vision and product - with each dimension being scored from 1 (poor) to 5 (best). The *TotalRating* is defined as the sum of these three individual ratings i.e., an integer in the interval $[3, 15]$.

For all ratings, we collect the date when the analyst issued the rating. The main analysis only considers ratings issued before ICO completion (or cancellation), which helps prevent look-ahead bias. However, our findings also hold when we include all ratings. Analysts have the opportunity to modify their ratings: when this happens, users can only see the updated rating score as well as two dates - the date of the first rating and the date of the update, but not the full history.¹⁷ This paper considers the modification date as the date for the rating and flag a modified rating by analyst i to ICO j with a dummy variable $Modified_{ij}$.

When issuing a rating, the analyst gives a score and, typically, justifies the decision by writing a review. We collect all reviews and calculate linguistic measures from these texts. Based on the Loughran and McDonald (2011) dictionaries, we calculate the tone of the language, defined as the difference between positive and negative words to total words, as well as the uncertainty of the language, defined as the count of uncertain words divided by

¹⁵We find that star analysts were not only active on ICObench, but also among the most active users on other ICO websites, e.g., ICOholder.com.

¹⁶Note that if ICO analysts become part of the ICO project by advising the team members, they lose the ability to rate their own ICO. We found that 4 analysts rated an ICO project before becoming an advisory team member.

¹⁷There is a well-documented phenomenon of “walking down” forecasts in the literature on sell-side analysts. The absence of access to the rating history of analysts on ICObench.com prevents us from studying this phenomenon in the ICO context.

total words. We further control for the complexity of the reviews, measured by the Gunning (1952) Fog index, which is a function of the number of words per sentence (length of a sentence) and the share of complex words (words with more than two syllables) relative to total words.¹⁸

For some analyses, we aggregate the analyst-ICO information to the ICO level. More precisely, we count the total number of analysts who rated ICO j in $\#Analysts_j$. We further aggregate all analyst ratings to ICO j in the variable $TotalRating_j$ by averaging all ratings that ICO j received from all analysts that cover this ICO. Finally, we proxy the lack of consensus among analysts that provide a rating for ICO j with $AnalystDispersion_j$, defined as the standard deviation of all ratings for ICO j .

Figure 2 presents the number of ratings in a given month over time of the newly announced ICOs, the number of ratings by analysts who registered in the same month, as well as the Bitcoin price in US dollar. While the number of new ratings went up hand-in-hand with the number of ICOs to the peak of the Bitcoin price in January 2018, the number of ratings exploded thereafter and only recently has converged again to the number of announced ICOs. Figure 2 further shows that the surging demand for information about crypto assets was met by the increase in the supply of analysts.

[Figure 2 about here]

Figure 3 shows the monthly averaged $TotalRating$ as well as the analysts' rating dispersion, measured by the standard deviation of ratings within an ICO in a given month. We observe that the average total rating of experts is overall very positive with a small decrease in the rating score around the Bitcoin price drop in 2018. Analyst dispersion remains at a relatively constant level over the sample period. It only slightly increases around the time when the Bitcoin price was low at the end of 2018, but decreases as the Bitcoin price rises at the end of 2019 again.

¹⁸The Gunning (1952) Fog index is defined as $Fog = 0.4 \cdot \left(\frac{TotalWords}{TotalSentences} + \frac{ComplexWords}{TotalWords} \right)$.

[Figure 3 about here]

Complementing the assessment of human experts, many platforms have set up machine-generated ratings. Instead of evaluating an ICO’s quality directly, these ratings based on the availability of information *about* the ICO. The idea is that more transparency indicates higher trustworthiness and quality of the ICO. For every ICO in our database, we collect the machine-generated rating by ICObench.com, which is called “Benchy”. The Benchy bot provides a higher rating for higher transparency on team and event information. Moreover, Benchy uses factual data such as “presence of the social media links” and “the level of activity on them”, see <https://icobench.com/faq>. Benchy re-evaluates each ICO profile at least once daily and issues a rating ranging between 1 (poor) and 5 (best). Only the most recent evaluation is observable, not the history of Benchy ratings.

While all ICOs listed on the platform ICObench.com automatically receive a machine-generated rating from the Benchy bot, 2,378 out of 5,384 ICOs listed on this website were also rated by ICO analysts. On average, the ICOs with(out) an analyst rating have a Benchy rating of 3.2 (2.7) out of 5.

2.4 Reciprocal ratings

A specific feature of the market is that ICO analysts also participate in ICOs. We identify those experts that are involved in one or several ICO projects by collecting each expert’s self-description of experiences and achievements from the ‘About’-section of their profile page on ICObench.com. Table 1 shows the distribution of ICO projects among analysts. Out of the 531 experts in our sample, 329 have been involved in at least one ICO, with some analysts being very active in launching ICOs.

[Table 1 about here]

We use this information to flag whether a rating of analyst i to ICO j is a response to a rating that analyst i received from any team member of coin j at any point in time, and

generate the indicator variable $ReciprocalRating_{ij}$ as follows:

$$ReciprocalRating_{ij} = \begin{cases} 1, & \exists TotalRating_{j'i'}^{ij} \text{ where } \mathbf{i} \in \Omega_{i'}, j' \in \Omega_j \\ 0, & \text{else} \end{cases}$$

where Ω_j refers to the set of all team members of the ICO j . Table 2 represents a hypothetical illustration of how we define this variable. $ReciprocalRating_{ij}$ thus flags whether any member of ICO j has provided a rating of any ICOs that expert i is associated with. Reciprocal ratings are not directly flagged by ICObench.com, but users can easily obtain the information given the available links to the analyst’s associated ICOs and the timeline of the ratings provided on ICObench.com.

Whenever $ReciprocalRating_{ij}$ indicates reciprocity, we additionally identify the level of the reciprocal rating, i.e., the $TotalRating$, as well as the three components $TeamRating$, $VisionRating$ and $ProductRating$ by any member of ICO j to the ICO with which expert i is associated. The level of the reciprocal rating is labeled $ReceivedTotalRating_{ij}$.

[Table 2 about here]

2.5 ICO outcome variables

We consider multiple measures of ICO success, some related to the initial completion, some related to the medium-term performance. Specifically, as a short-term measure for ICO success we construct a dummy variable $Success$, which takes the value of 1 if the ICO-related coin successfully completes the offering and receives funding. For these ICOs, we collect information on the dollar amount raised during the campaign from ICObench.com, ICOmarks.com, tokendata.io and ICOdata.io. Tokens were classified as failed when we could not find the amount raised nor any success information on the above-mentioned web pages. In total, we identify 1,932 successful ICOs.

In addition, we generate a more medium-term success measure, the logarithmic market

value of the token 90 days after listing on an exchange from CoinMarketCap.com. We observe the market capitalization information only for a subset of ICOs in our sample, either because CoinMarketCap.com does not cover the exchange where the token was listed or because the project failed. Therefore, we either use the market capitalization 90 days post exchange listing for the restricted sample of successfully listed ICOs or set the market capitalization to zero for projects without any information on CoinMarketCap.com, assuming a failure of these projects (Howell et al., 2020).

Figure 4 shows the time trend of successful ICOs. ICOs became popular at the beginning of 2017. While only 29 ICO tokens were on sale before then, the number increased to 1,127 ICOs within one year with around 94 offerings per month and a 53% success rate. The market peaked in 2018, with 3,360 ICOs in total and a success rate of 33%. In 2019, around 64 ICO offerings were sold per month, of which 25% were successful on average. Thus, flow of ICOs continues, albeit at a lower level, even after the sharp decline of cryptocurrency prices and the corresponding decline of enthusiasm towards ICOs.

[Figure 4 about here]

Finally, we also collect information about scams, i.e., ICOs that were launched with the intention to defraud investors. To do so, we use the marker ‘Scam or Other Issues’ for dead coins listed on Coinopsy.com, as well as information from Deadcoins.com, a message board where users post about scams. Some of these ICOs can also be found in the U.S. Securities and Exchange Commission (SEC) press releases, especially when they fine ICO companies for fraudulent practices.¹⁹ With this (likely conservative) method, 234 ICOs were flagged as scams in our data.

2.6 Forecast errors

Combining the ICO success variable and the analyst rating score allows us to construct, an ex-post forecast error measure for each rating. As the outcome of an ICO is either success

¹⁹See, e.g., <https://www.sec.gov/news/press-release/2019-259>.

or failure, we define the forecast error of a rating as the distance to the highest (lowest) possible rating in case of success (failure):

$$ForecastError_{ij} = \begin{cases} 15 - TotalRating_{ij}, & \text{if ICO succeeded} \\ TotalRating_{ij} - 3, & \text{if ICO failed} \end{cases}$$

If an analyst gives a successful ICO a total rating of 15, their rating was fully precise, resulting in a *ForecastError* measure of 0. If the ICO had failed, however, the forecast error of this rating would flag a 12, such that $ForecastError \in \{0, \dots, 12\}$.²⁰

Figure 5 shows the monthly averaged *TotalRating*, the average forecast error and the number of successful ICOs (as a share of total ICOs) over time. In addition, we plot the monthly average forecast error separately for ratings where analysts were too optimistic and too pessimistic, respectively.²¹ Interestingly, ratings become less precise over time. This is driven by overoptimistic analysts.

[Figure 5 about here]

In our regression analysis, we use an analyst-specific measure of the forecast error that takes the entire history of an analyst’s ICO-specific *ForecastError_{ij}* into account. We recursively average the *ForecastError_{ij}* of analyst *i* over all of their issued ratings up to ICO *j* using an expanding window. We denote this variable *ForecastError_i^j*.

2.7 ICO characteristics

For every ICO in the sample, we collect data on the campaign characteristics that have been found in the literature to indicate the perceived quality of an ICO by investors (Amsden and

²⁰While this *ForecastError* measure is not immediately available for investors on ICObench.com, one can easily view the entire timeline of an analysts’ ratings with a link to detailed information on the rated ICO.

²¹An analyst is too optimistic (pessimistic) if her rating of a failing (successful) ICO was larger (smaller) than 3 (15). Thus, we calculate the monthly average optimistic (pessimistic) forecast error over the lower (upper) cases of the *ForecastError_{ij}* definition.

Schweizer, 2018; Burns and Moro, 2018; Howell et al., 2020; Roosenboom et al., 2020). For many characteristics, we generate binary indicators that flag whether an ICO exhibits the respective feature. The dummy variable *Presale* equals one if an ICO offers coins at pre-sale stage and zero otherwise. The *Bonus* and *Bounty* dummies equal 1 if there were discounts on the token sale or incentives to boost social media presence, respectively. The dummy *MVP* flags the availability of a minimum viable product or whether a product prototype was in place. The dummy *KYC* equals one if investors need to validate their identity by signing up to a whitelist to get access to the token sale.²² The dummy *IEO* indicates the use of a centralized token launch platform provided by a cryptocurrency exchange. The *RetentionRatio* is the retained share of total token supply (in percent); it captures the “skin in the game” of ICO team members. We collect the information whether the ICO has its own webpage on the forum Bitcointalk.org and/or Facebook to discuss and promote their project idea. Finally, we include *LengthWhitePaper*, the natural logarithm of (1 + total words of the white paper) as a proxy for the informative value of the ICO whitepaper. We set this variable to zero if no white paper could be found. In addition to these campaign-specific characteristics, we collect the year-month information of the date when the ICO was launched.

2.8 Descriptive statistics

Table 3 shows descriptive statistics of the key variables of rating and ICO characteristics. All variables are defined in Table A1.

[Table 3 about here]

In our sample, the analyst’s average *TotalRating* is 11, where the *ProductRating* is slightly more pessimistic than the other two dimensions *TeamRating* and *VisionRating*.

²²Note that ICObench.com provides information on two different KYC procedures. One KYC symbol means the identity verification of ICObench.com profiles, while the second flags the identification and registration process of investors to receive access to the token sale. We use the second KYC throughout the paper.

Of all ratings, 12% are flagged *ReciprocalRating*, which are found to be somewhat more positive with an average *ReceivedTotalRating_{ij}* of 13. On the ICO dimension, we observe a success rate of 36%. In terms of dollar amount raised, EOS, Telegram, and Bitfinex were the most successful ICOs in our sample. The scam rate is 4.3%. Each ICO is covered by 2.6 analysts, on average, and 44% of all ICOs are covered by at least one analyst. The ICOs for Sharpay (94), Truegame (82) and WePower (64) had the largest number of analysts covering them.

3 Empirical Analysis

Section 3.1 analyzes the determinants of an analysts' rating (both quantitative and qualitative) of an ICO. Section 3.2 in turn considers whether investors consider differences in the reliability of analyst ratings.

3.1 What determines analyst ratings?

3.1.1 Baseline results

We model the rating of analyst i for ICO j as a function of analyst characteristics, as indicated in the following equation:

$$\begin{aligned}
 Rating_{ij} = & \beta_0 + \beta_1 \cdot Benchy_j + \beta_2 \cdot StarAnalysts_{ij} + \beta_3 \cdot ForecastError_i^j \\
 & + \beta_4 \cdot Modified_{ij} + \beta_5 \cdot X_j + Month_{ij} + \alpha_j + \epsilon_{ij},
 \end{aligned} \tag{1}$$

$Rating_{ij}$ denotes the respective rating score that analyst i gives to ICO j for the different rating categories team, vision and product (on a scale from 1–5), as well as the *TotalRating_{ij}* score as the sum of the three categories (on a scale from 3–15). The vector X_j contains the ICO campaign characteristics as described in Subsection 2.7. Time trends of ratings are absorbed by $Month_{ij}$ dummies. α_j denotes ICO fixed effects. We allow for a potential

serial correlation of ratings within each analyst and within each ICO and employ two-way clustering of standard errors (Cameron et al., 2011) at the analyst and ICO dimension.

[Table 4 about here]

Table 4 summarizes the results of this analysis. Column (1) shows that machine-generated and human expert ratings point in the same direction qualitatively, i.e., ICOs with higher machine-generated ratings receive a higher rating score by human analysts on average. Moreover, in the cross-section of analysts, columns (2) and (3), we find a statistically significant negative coefficient on $ForecastError_i^j$, implying that analysts with historically higher forecast errors give on average lower ratings. The negative relationship between the past forecast errors and the rating remains also in the within ICO estimation, as column (4) shows. Analysts listed within the top 30 analysts on ICObench.com are more critical and issue lower ratings on average.

Furthermore, in line with the literature, the coefficients of the control variables suggest that analysts consider the characteristics of the underlying ICO (Deng et al., 2018; Bourveau et al., 2019; Roosenboom et al., 2020). In general, we find that ICOs with a pre-sale event, with a KYC feature and an IEO feature receive better ratings. Moreover, analysts perceive it as a good signal when founders retain a higher share of the tokens themselves.

3.1.2 Reciprocal ratings

When ICO analysts issue new ratings, do these ratings depend on ratings that their own affiliated ICOs previously received? To answer this question, we run regressions, as specified in the following equation:

$$\begin{aligned}
 Rating_{ij} = & \beta_0 + \beta_1 \cdot ReciprocalRating_{ij} + Analyst \times Month_{ij} \\
 & + ICO \times Month_{ij} + \epsilon_{ij},
 \end{aligned} \tag{2}$$

where $ReciprocalRating_{ij}$ indicates a dummy that flags whether analyst i received a rating from a team member of ICO j . We include $Analyst \times Month$ and $ICO \times Month$ dummies to exploit only the analyst and ICO pairing within the month of the rating. These fixed effects detect the variation previously established in Table 4, and they help to rule out that the results were driven by a non-random match between founders of good ICOs that also serve as analysts. Comparing different assessments of the same analyst and for the same ICO allows also to differentiate whether analysts behave in a deliberately optimistic biased manner or whether the optimistic assessment is due to the high difficulty of forecasting tasks. For reciprocal ratings, we also analyze whether the level of the prior rating predicts the level of the reciprocal rating.²³ We again employ two-way clustering at the analyst and ICO dimension.

[Table 5 about here]

Table 5 shows that ratings indeed contain a reciprocal element. Column (1) indicates a positive association between the total rating an analyst gives to an ICO and the $ReciprocalRating_{ij}$ dummy. More specifically, the total rating score is around 0.25 points higher when the analyst is in a position to respond to a prior rating. Additionally, within the sample of reciprocal ratings, column (2) shows that ratings are more positive, the higher the prior received rating was. In other words, analysts reciprocate positive ratings. For example, column (2) shows that each one-unit (one standard deviation) increase for the previous rating leads an analyst to issue a total rating around 0.08 (0.15) higher. Note that this result holds within ICO-time and analyst-time combinations, i.e., comparing ratings by two (otherwise) identical analysts, where one analyst previously received a rating by a team member of coin j and the other one did not. This reciprocal rating behavior is similar to the *quid pro quo* between hedge funds and sell-side equity analysts described in Klein et al. (2019).

²³We use $Analyst \times Quarter$ and $ICO \times Quarter$ dummies for the analysis because of the restricted sample size.

Columns (3)–(8) analyze the three different rating categories team, vision and product separately. The coefficient for the *ReciprocalRating_{ij}* dummy is positive and significant for all three categories, indicating that, on average, analysts give a higher rating for the team, vision and product of ICO *j* if any team member of ICO *j* rated them. However, the actual score is significant only for the team dimension. That is, an analyst rates the team component of ICO *j* higher if she received a more favorable team rating from a team member of ICO *j*. By contrast the scores for vision and product are not significant. These findings are intuitive, as the team category constitutes a “soft factor”. The results also indicate a relatively personal nature of the reciprocity.

3.1.3 Linguistic characteristics of rating reviews

When issuing ratings, analysts often justify the rating scores with written reviews. We next analyze whether more optimistic ratings are special in terms of the linguistic nature of the written review. The literature on earnings conference calls uses the number of words spoken by analysts as a proxy for the question difficulty, so analysts who ask lengthier questions are regarded as more critical (Merkley et al., 2017). Correspondingly, we investigate whether the rating score correlates with the length of the written text or with the linguistic tone of the review. Moreover, we investigate whether the relationship between the rating score and the review length and tone differs for reciprocal versus non-reciprocal ratings. This idea follows Cohen et al. (2020), who document that biased analysts ask easier questions. We run regressions for the overall sample as well as for reciprocal and non-reciprocal ratings separately as specified in the following equation:

$$\begin{aligned} Linguistic\ Measure_{ij} = & \beta_0 + \beta_1 \cdot TotalRating_{ij} + Analyst \times Time_{ij} \\ & + ICO \times Time_{ij} + \epsilon_{ij}, \end{aligned} \tag{3}$$

where $Linguistic Measure_{ij}$ indicates interchangeably the length of the rating review measured by the (natural logarithm of the) number of words and the ratio of positive words minus negative words to total words in the review. As before, we employ two-way clustering by analysts and ICOs.

[Table 6 about here]

Table 6 shows the results. In Panel A column (1), we find a negative relationship between the rating score and the length of the review, suggesting that more negative ratings come with a more detailed explanation. In column (2), we include $Analyst \times Month$ and $ICO \times Month$ dummies to rule out that the results are driven by a non-random match between analyst characteristics (e.g. mood) and the quality of the rated ICO. For the review tone in Panel B, columns (1) and (2), we find that analysts use more positive terminology when reviewing an ICO that they score higher.

When investigating whether the relationship between the rating score and the review length and tone differs for reciprocal versus non-reciprocal ratings, we find that lower rating scores are justified with even lengthier reviews for reciprocal ratings, with a statistically significant difference to the coefficient for non-reciprocal ratings.²⁴ The relationship between review tone and rating score does not differ noticeably between reciprocal and non-reciprocal ratings.

3.1.4 Order of ratings

The literature on security analysts has documented herding behavior among analysts and shows that their buy or sell recommendations have a significant positive influence on the next analysts' recommendations (Welch, 2000). Thus, reciprocal analysts' scores may impact investors as well as other analysts when they cover the ICO at an early stage. We therefore analyze whether analysts provide reciprocal ratings faster and move earlier for ICOs where

²⁴Because of the restricted sample size for the sample of reciprocal ratings, we use $Analyst \times Quarter$ and $ICO \times Quarter$ dummies for these analyses.

they issue more positive ratings. We generate a variable that counts the rank of rating arrival per ICO j from analyst i , i.e., whether analyst i was the first, second, third, ... last analyst that rated for ICO j . We relate the order of the rating coverage to the *ReciprocalRating_{ij}* dummy, as indicated in the following equation:

$$\begin{aligned} OrderRank_{ij} = & \beta_0 + \beta_1 \cdot TotalRating_{ij} + \beta_2 \cdot ReciprocalRating_{ij} + \beta_3 \cdot StarAnalysts_j \\ & + \beta_4 \cdot ForecastError_i^j + Month_{ij} + \alpha_i + \alpha_j + \epsilon_{ij}. \end{aligned} \quad (4)$$

We again absorb any ICO and analyst characteristics with ICO and analyst fixed effects and control for time trends by *Month_{ij}*. As before, we use two-way clustered standard errors at the analyst and ICO dimension.

[Table 7 about here]

The results are shown in Table 7. In line with the literature on the analyst coverage of stocks (Demiroglu and Ryngaert, 2010), we first find that analysts who give favorable ratings tend to issue their rating early. Second, star analysts tend to move first and rate the same ICO earlier than their less experienced peers. Third, reciprocal ratings tend to be issued early. In particular, in the chronological sequence of ratings given to an ICO j , a reciprocal analyst appears to issue her rating on average 1.3 positions earlier than a non-reciprocal analyst.

3.2 Are ICOs with higher ratings more successful, and which ICOs fail despite high ratings?

So far, we have established several important determinants of ICO analyst ratings, with analyst-specific factors, such as experience, prior forecasting ability and reciprocal status, playing a major role in addition to objective differences among the ICOs. Now, we investigate

whether ICOs with higher ratings are indeed more successful. First, we establish baseline results for (unconditional) ICO success, but our main interest is in explaining when investors deviate from the ICO analyst consensus, that is, the ICO success probability conditional on an extreme positive (or negative) rating outcome. We also consider whether the factors that explain such deviations predict scams.

3.2.1 Ratings and ICO success

Table 8 presents descriptive statistics for the relationship between ratings and ICO success. Panel A indicates that ICOs are more likely to be successful when it motivates analysts to rate it. In Panel B, we tabulate success statistics for groups of the quantitative rating score. The probability of receiving funding, the market capitalization 90 days after listing and the average dollar amount raised is higher for ICOs with more positive ratings, though the relationship is not strictly monotonic for the average dollar amount raised.

While these results highlight that successful ICOs have, on average, higher ratings, there are numerous cases in which ICOs were either unsuccessful despite positive ratings or successful despite negative ratings. To quantify this phenomenon, we define for each ICO j a *Disagreement_j* dummy as a conditional success outcome. More precisely, the *Disagreement_j* dummy equals one if (i) the average *TotalRating_j* of an ICO is greater than or equal to 13 but the ICO is unsuccessful, or (ii) the average total rating is less than or equal to 5 and the ICO is successful. In our sample, this *Disagreement_j* dummy is one in 413 of 2,378 rated ICOs (17%).²⁵ While the unconditional failure rate of ICOs is about 64%, 53.6% of ICOs with an average rating in the top quartile fail.

These mismatches between ratings and ICO success do not occur randomly. To illustrate, in Panel C, we tabulate the disagreement dummy against the occurrence of reciprocal ratings. We observe that the ICO outcome does not correspond to what one would expect

²⁵Note that disagreement most often concerns the case where the rating is high but the ICO fails. We focus mainly on these cases when analyzing disagreement. There are only very few cases of successful ICOs with an average poor rating (N=44).

given the ratings level if reciprocal analysts cover the ICO. ICOs that receive very favorable recommendations fail much more frequently if the reciprocal rating share is positive than if none of the ratings is reciprocal. Moreover, Panel C also shows that the market capitalization 90 days after listing is lower for ICOs with a reciprocal rating.

[Table 8 about here]

In order to formally analyze ICO success in a regression framework, we first explain the unconditional success of ICO j by characteristics of participating analysts and other variables in a logit regression:

$$\begin{aligned} Success_j = & \beta_0 + \beta_1 \cdot TotalRating_j + \beta_2 \cdot \#Analysts_j + \beta_3 \cdot ReciprocalRatingShare_j \\ & + \beta_4 \cdot StarAnalysts_j + \beta_5 \cdot PreviousRatings_j + \beta_6 \cdot AnalystDispersion_j \quad (5) \\ & + \beta_7 \cdot Benchy_j + \beta_8 \cdot Z_j + \beta_9 \cdot X_j + Month_j + \epsilon_j. \end{aligned}$$

$Success_j$ indicates the success dummy as described in Section 2 as well as the market capitalization of the token 90 days after listing on an exchange.²⁶ X_j again represents the controls as in Equation 1. Additionally, we control for linguistic measures with Z_j , which contains the average tone, uncertainty and complexity as well as the length of all rating reviews written about ICO j . We further include a dummy for each month of the sample, $Month_j$ to absorb time trends common to all ICOs. In regressions with market capitalization as the dependent variable these fixed effects are arguably particularly important to abstract from general market developments and focus on the cross-section of ICOs. Note that the sample size drops when including the $AnalystDispersion_j$ measure, as it is only defined for ICOs with at least two ratings.

In addition to the unconditional success of ICOs, we investigate the success conditional on having received very high or very low ratings. Thus, we run the following logit regression

²⁶In the Appendix, we alternatively use the dollar amount raised during the ICO as a measure of success.

on the ICO level:

$$\begin{aligned}
Disagreement_j = & \beta_0 + \beta_1 \cdot \#Analysts_j + \beta_2 \cdot ReciprocalRatingShare_j \\
& + \beta_3 \cdot StarAnalysts_j + \beta_4 \cdot PreviousRatings_j \\
& + \beta_5 \cdot AnalystDispersion_j + \beta_6 \cdot Benchy_j \\
& + \beta_7 \cdot Z_j + \beta_8 \cdot X_j + Month_j + \epsilon_j,
\end{aligned} \tag{6}$$

where X_j is the same set of controls as in Equation 1 and $Month_j$ dummies absorb common time trends. We again control for some linguistic measures with Z_j , which contains the average tone, uncertainty, complexity and length of all rating reviews written on ICO j .

[Table 9 about here]

[Table 10 about here]

As a baseline result, the regressions confirm that ratings are predictive on average. The likelihood of an ICO being successful either as measured by initial listing, as in columns (1) through (3), or in terms of the market capitalization 90 days after the listing, as in columns (4) through (6), is higher if that ICO motivates analysts to rate it. This result holds even after controlling for a wide variety of ICO characteristics.²⁷ Moreover, as Table 9 shows, the level of that rating matters as well. A more positive $TotalRating_j$ is associated with a higher probability of success, and a higher market capitalization. The rating level is not significant in column (6), where we consider only those ICOs that were successful and for which capitalization data are available on CoinMarketCap.com. The machine-generated rating $Benchy$ is predictive as well, indicating that ICOs are, on average, more likely to be successful the more easily publicly available information there is about it.²⁸ Importantly,

²⁷This finding is in line with the general literature on analysts and rating agencies, which indicates that the market appreciates analyst coverage (Demiroglu and Ryngaert, 2010) and the existence of ratings (Sufi, 2009).

²⁸This result is in line with the finding that investors value the dissemination of corporate news releases via robots, even when that information in principle is already available (Blankespoor et al., 2018).

the human ratings remain significant determinants of success throughout, indicating that investors use the two rating types for different information. As for control variables, we observe similar results as the prior literature.²⁹

Our main interest is in the role of the heterogeneity among analysts, and in what predicts ICO failure despite high ratings. First, Table 9 shows that ICOs with a higher share of reciprocal ratings experience lower market capitalizations 90 days after being listed. While the reciprocal share does not explain the binary ICO success indicator unconditionally, it does correlate significantly with failure (success) conditional on high (low) ratings (see Table 10). Table A3 shows that the effect emerges largely from failed ICOs despite high ratings.

There are two possible interpretations of this negative association of the reciprocal share and the short-term and medium-success of ICOs. First, we note that we control for a wide variety of factors presumably capturing variation in ICO quality. However, it is still possible that reciprocal ratings are correlated with some additional unobserved variation in ICO quality. The second interpretation is that, a matter of principle, investors trust ICOs with more reciprocal ratings less, even when these ratings do not suffer from a conflict of interest.

While these interpretations are not mutually exclusive, an additional test provides further insight. For each ICO, we calculate the difference between the average reciprocal and non-reciprocal ratings. We then divide the sample into cases where the average reciprocal rating is higher than or equal to the average of non-reciprocal ratings, and cases where reciprocal ratings were lower than the non-reciprocal ones. In the former case reciprocal ratings influence the overall ICO rating to a large extent, whereas in the latter case reciprocal ratings are less likely to bias the overall ICO rating. If investors dislike reciprocal ratings in general, we would expect the reciprocal rating share to be a significant determinant of disagreement

²⁹For example, we find a positive coefficient for Bitcointalk and negative coefficients for the Bounty and Bonus dummies. Successful ICOs tend to have longer white papers. For reasons of space, we only tabulate two control variables that received little prior attention in the literature: MVPs and IEOs. The use of crypto-exchange launchpads for Initial Exchange Offerings (IEOs) positively correlates with the two success variables. Somewhat surprisingly, ICOs with a minimum viable product (MVP) feature a lower probability of success and were only able to attract a lower dollar amount of funding. This unexpected result might be due to a non-regulated definition of minimum viable products. For example, drafts of codes on GitHub.com that are open to a discussion by other GitHub users were classified as MVP.

in both cases. Columns 3 and 4 of Table 10 present the results. It is noteworthy that the share of reciprocal analysts matters only for the conditional success for those cases where reciprocal ratings are at least as positive as non-reciprocal ratings. A caveat is that these regressions are based on relatively small samples (because they are only available for the subsample with reciprocal ratings). That said, they provide some suggestive evidence that investors are not concerned with reciprocal ratings per se, but rather that positive reciprocal ratings provide an additional signal of poor quality of an ICO.

Consider next the linguistic measures of the rating. This feature also does not predict ICO success per se (at least once controlling for the quantitative rating). Nevertheless, linguistic measures still provide insight. Table 10 shows that the likelihood of failure for a highly-rated ICO increases as the positivity of the tone and complexity of the language increase. Similarly, ICO failure despite high average ratings occurs more frequently when the analysts were more positive in ratings prior to their rating of ICO j .

Table 9 suggests that star analyst coverage is not predictive for ICO success. Again, however, Table 10 reveals that highly rated ICOs fail less frequently when many star analysts cover them.

Finally, analyst dispersion is also relevant for ICO success only if the average view of analysts is very positive. Interestingly, and at first surprisingly, when analysts' rating are highly dispersed, ICOs are less likely to fail. Intuitively, the combination of high average ratings and high dispersion occurs when there are several extremely positive and some negative views. The very positive ratings then carry the day. This is similar in spirit to the apparent anomaly that stocks with high dispersion of analyst opinions have high prices and, thus, lower future returns (Diether et al., 2002).³⁰

Overall, the results highlight that even when a characteristic is not related unconditionally to ICO success, it is not irrelevant for understanding ICO success. Specifically, factors such

³⁰This interpretation of analyst dispersion has been challenged in equity markets. For example, Avramov et al. (2009) show that the analyst dispersion anomaly is driven by a small fraction of firms with very high credit risk.

as the reciprocity of ratings, analyst dispersion and the presence of star analysts explain deviations from the outcome given a very high level of ratings. Thus, while it is perhaps unsurprising that average ratings predict the success of an ICO campaign, our key result is that the detailed characteristics of the ICO ratings and those who provide them contain important additional information.

3.2.2 Ratings and ICO scams

We have established that analyst ratings help predict ICO success, but that investors tend to disregard reciprocal ratings. Does the latter result occur because ICOs with a higher fraction of reciprocal ratings are more likely to be fraudulent? To answer, we rerun the regression from Equation 5, but replace the success dummy with a dummy that equals one if the ICO was detected to intentionally defraud investors.

[Table 11 about here]

As Table 11 shows, we find no correlation between the share of reciprocal analysts and fraudulent ICOs. Also, the level of machine-generated ratings or the level of human analyst ratings do not help identify fraudulent ICOs. It still pays for investors to consider the human analyst assessments, however. In particular, ICOs with more dispersion among analysts both in the quantitative and in the qualitative rating tend to be fraudulent.

4 Conclusion

The intersection of new technologies and financial markets (FinTech) holds great promise. One relatively recent phenomenon in this space is the opportunity for new ventures to engage in Coin/Token Offerings, a new form of financing. Yet, despite the problem of asymmetric information looms large in these markets, there was a tremendous rise of ICOs followed by a market crash. This paper studies the role of information intermediaries, human experts that may help ameliorate this asymmetric information problem, in unregulated financial markets.

While the rise and fall of the ICO market is interesting in its own right, ICO analysts show many interesting parallels to equity analysts or rating agencies. Particularly noteworthy are potential conflicts of interest, and how investors interpret them. The advantage of the ICO setting is that detailed data on links between analysts and securities they rate are available. For example, we document that an ICO analyst i , when rating an ICO j , tends to issue a rating that depends on the rating that their own affiliated ICO had previously received from team members of ICO j . However, there is a higher probability that an ICO fails, even when it has very favorable ratings, when more of those ratings are reciprocal. ICOs with a high share of reciprocal ratings also tend to have a lower market capitalization 90 days after listing on an exchange.

Thus, while the prior literature shows that human experts' average ratings predict the success of an ICO campaign (over and above machine-generated ratings), our key result is that understanding ICO success requires looking at past averages and studying the detailed characteristics of the ratings and those who provide them. After all, failure is frequent even among the most highly rated ICOs. Reciprocal ratings and highly positive reviews may be correlates of problems of highly rated ICOs not reflected in the many characteristics we control for; alternatively, investors may trust ICOs with more reciprocal and optimistic analysts less despite them being worth funding. Either way, the findings suggest that investors do not blindly pile capital into highly rated ICOs.

A necessary precondition for such investors to take such a differentiated approach to investments is the availability of information about the track record and potentially conflicting activities of analysts. Thus, the availability of information appears to support the ICO market's allocative role in society. While these results obtain on this largely unregulated market, the general insight that investors seem to value information about analysts is likely to be relevant for other markets as well.

Finally, our results suggest that the failure of the unregulated ICO market was largely attributed to measures of conflicts of interest. Therefore, to ensure the functioning of newly

emerging and largely unregulated financial markets as for example decentralized finance, IEOs or STOs, there needs to be at least some regulation that prevents such conflicts of interest.

References

- Adhami, S., Giudici, G., and Martinazzi, S. (2018). Why do businesses go crypto? An empirical analysis of Initial Coin Offerings. *Journal of Economics and Business*, 100(C):64–75.
- Aggarwal, R., Hanley, K. W., and Zhao, X. (2019). The role of information production in private markets: Evidence from Initial Coin Offerings. Working paper.
- Agrawal, A. and Chen, M. A. (2008). Do analyst conflicts matter? evidence from stock recommendations. *The Journal of Law & Economics*, 51(3):503–537.
- Amsden, R. and Schweizer, D. (2018). Are blockchain crowdsales the new ‘gold rush’? Success determinants of Initial Coin Offerings. Working paper.
- Aubry, M., Kräussl, R., Manso, G., and Spaenjers, C. (2019). Machine learning, human experts, and the valuation of real assets. Working paper.
- Avramov, D., Chordia, T., Jostova, G., and Philipov, A. (2009). Dispersion in analysts’ earnings forecasts and credit rating. *Journal of Financial Economics*, 91(1):83–101.
- Badoer, D. C., Demiroglu, C., and James, C. M. (2019). Ratings quality and borrowing choice. *The Journal of Finance*, 74(5):2619–2665.
- Bae, K.-H., Stulz, R. M., and Tan, H. (2008). Do local analysts know more? A cross-country study of the performance of local analysts and foreign analysts. *Journal of Financial Economics*, 88(3):581–606.
- Baghai, R. and Becker, B. (2017). Non-rating revenue and conflicts of interest. *Journal of Financial Economics*, 127.
- Baker, H. K. and Mansi, S. A. (2002). Assessing credit rating agencies by bond issuers and institutional investors. *Journal of Business Finance & Accounting*, 29(9-10):1367–1398.
- Bakos, Y. and Halaburda, H. (2019). The role of cryptographic tokens and ICOs in fostering platform adoption. Working paper.
- Bar-Isaac, H. and Shapiro, J. (2013). Ratings quality over the business cycle. *Journal of Financial Economics*, 108(1):62 – 78.
- Benedetti, H. and Kostovetsky, L. (2018). Digital tulips? Returns to investors in Initial Coin Offerings. Working paper.
- Biais, B., Bisière, C., Bouvard, M., and Casamatta, C. (2019). The blockchain folk theorem. *Review of Financial Studies*, 32(5):1662–1715.
- Blankespoor, E., deHaan, E., and Zhu, C. (2018). Capital market effects of media synthesis and dissemination: Evidence from robo-journalism. *Review of Accounting Studies*, 23(1):1–36.

- Bolton, P., Freixas, X., and Shapiro, J. (2012). The credit ratings game. *The Journal of Finance*, 67(1):85–111.
- Boreiko, D. and Vidusso, G. (2019). New blockchain intermediaries: Do ICO rating websites do their job well? *Journal of Alternative Investments*, 21(4):67–79.
- Bourveau, T., De George, E. T., Ellahie, A., and Macciocchi, D. (2019). Information intermediaries in the crypto-tokens market. Working paper.
- Bradley, D., Gokkaya, S., and Liu, X. (2017). Before an analyst becomes an analyst: Does industry experience matter? *The Journal of Finance*, 72(2):751–792.
- Bradshaw, M., Ertimur, Y., O’Brien, P., et al. (2017). Financial analysts and their contribution to well-functioning capital markets. *Foundations and Trends (R) in Accounting*, 11(3):119–191.
- Burns, L. and Moro, A. (2018). What makes an ICO successful? An investigation of the role of ICO characteristics, team quality and market sentiment. Working paper.
- Cameron, A. C., Gelbach, J. B., and Miller, D. L. (2011). Robust inference with multiway clustering. *Journal of Business and Economic Statistics*, 29(2):238–249.
- Campbell, J., DeAngelis, M., and Moon, J. (2019). Skin in the game: personal stock holdings and investors’ response to stock analysis on social media. *Review of Accounting Studies*, 24(C):731–779.
- Catalini, C., Foster, C., and Nanda, R. (2018). Machine intelligence vs. human judgement in new venture finance. Working paper.
- Cerchiello, P. and Toma, A. M. (2018). ICOs success drivers: A textual and statistical analysis. Working paper.
- Chakraborty, S. and Swinney, R. (2020). Signaling to the crowd: Private quality information and rewards-based crowdfunding. *Manufacturing & Service Operations Management*, 1(C):1–15.
- Chan, L. K. C., Karceski, J., and Lakonishok, J. (2007). Analysts’ conflicts of interest and biases in earnings forecasts. *Journal of Financial and Quantitative Analysis*, 42(4):893–913.
- Chen, H., De, P., Hu, Y. J., and Hwang, B.-H. (2014). Wisdom of crowds: The value of stock opinions transmitted through social media. *The Review of Financial Studies*, 27(5):1367–1403.
- Chod, J. and Lyandres, E. (2020). A theory of ICOs: Diversification, agency, and information asymmetry. *Management Science*, forthcoming.
- Chod, J., Trichakis, N., and Yang, A. S. (2020). Platform tokenization: Financing, governance, and moral hazard. Working paper.

- Chu, C. S. and Rysman, M. (2019). Competition and strategic incentives in the market for credit ratings: Empirics of the financial crisis of 2007. *American Economic Review*, 109(10):3514–55.
- Cohen, L., Lou, D., and Malloy, C. J. (2020). Casting conference calls. *Management Science*, 66(11):5015–5039.
- Cong, L. W. and He, Z. (2019). Blockchain disruption and smart contracts. *The Review of Financial Studies*, 32(5):1754–1797.
- Cong, L. W., Li, Y., and Wang, N. (2020a). Token-based platform finance. Working paper.
- Cong, L. W., Li, Y., and Wang, N. (2020b). Tokenomics: Dynamic adoption and valuation. *Review of Financial Studies*, forthcoming.
- Conley, J. P. (2017). Blockchain and the economics of crypto-tokens and Initial Coin Offerings. Working paper.
- Crawford, S. S., Roulstone, D. T., and So, E. C. (2012). Analyst initiations of coverage and stock return synchronicity. *The Accounting Review*, 87(5):1527–1553.
- Da, Z. and Huang, X. (2020). Harnessing the wisdom of crowds. *Management Science*, 66(5):1847–1867.
- Davydiuk, T., Gupta, D., and Rosen, S. (2019). De-crypto-ing signals in Initial Coin Offerings: Evidence of rational token retention. Working paper.
- Demiroglu, C. and Ryngaert, M. (2010). The first analyst coverage of neglected stocks. *Financial Management*, 39(2):555–584.
- Deng, X., Lee, Y. T., and Zhong, Z. (2018). Decrypting coin winners: Disclosure quality, governance mechanism and team networks. Working paper.
- Diether, K. B., Malloy, C. J., and Scherbina, A. (2002). Differences of opinion and the cross section of stock returns. *The Journal of Finance*, 57(5):2113–2141.
- Dittmar, R. and Wu, D. A. (2019). Initial Coin Offerings hyped and dehyed: An empirical examination. Working paper.
- Drake, M. S., Thornock, J. R., and Twedt, B. J. (2017). The internet as an information intermediary. *Review of Accounting Studies*, 22(2):543–576.
- Easley, D., O’Hara, M., and Basu, S. (2019). From mining to markets: The evolution of bitcoin transaction fees. *Journal of Financial Economics*, 134(1):91–109.
- Ellis, D. (1998). Different sides of the same story: Investors? And issuers? Views of rating agencies. *The Journal of Fixed Income*, 1:35–45.
- Ernst and Young (2018). Initial Coin Offerings (ICOs): The class of 2017 – one year later. EY study.

- Fahlenbrach, R. and Frattaroli, M. (2020). ICO Investors. *Financial Markets and Portfolio Management*, pages 1–59.
- Farrell, M., Jame, R., and Qiu, T. (2020). The cross-section of non-professional analyst skill. Working paper.
- Feng, C., Li, N., Wong, M., and Zhang, M. (2019). Initial Coin Offerings, blockchain technology, and white paper disclosures. Working paper.
- Fisch, C. (2019). Initial Coin Offerings (ICOs) to finance new ventures. *Journal of Business Venturing*, 34(1):1–22.
- Fisch, C., Masiak, C., Vismara, S., and Block, J. (2019). Motives and profiles of ICO investors. *Journal of Business Research*, forthcoming.
- Fisch, C. and Momtaz, P. P. (2020). Institutional investors and post-ICO performance: An empirical analysis of investor returns in Initial Coin Offerings (ICOs). *Journal of Corporate Finance*, 64:101679.
- Florysiak, D. and Schandlbauer, A. (2019). The information content of ICO white papers. Working paper.
- Gan, R. J., Tsoukalas, G., and Netessine, S. (2020). Initial Coin Offerings, speculators, and asset tokenization. Working paper.
- Giudici, G. and Adhami, S. (2019). The impact of governance signals on ICO fundraising success. *Journal of Industrial and Business Economics*, 46(2):283–312.
- Goldstein, I., Gupta, D., and Sverchkov, R. (2019). Initial Coin Offerings as a commitment to competition. Working paper.
- Gryglewicz, S., Mayer, S., and Morellec, E. (2020). Optimal financing with tokens. Working paper.
- Gunning, R. (1952). *The technique of clear writing*. McGraw-Hill, New York.
- Hinzen, F. J., John, K., and Saleh, F. (2020). Bitcoin’s fatal flaw: The limited adoption problem. Working paper.
- Hirth, S. (2014). Credit rating dynamics and competition. *Journal of Banking & Finance*, 49(C):100 – 112.
- Howell, S. T., Niessner, M., and Yermack, D. (2020). Initial Coin Offerings: Financing growth with cryptocurrency token sales. *The Review of Financial Studies*, 33(9):3925–3974.
- Hu, A. S., Parlour, C. A., and Rajan, U. (2019). Cryptocurrencies: Stylized facts on a new investible instrument. *Financial Management*, 48(4):1049–1068.
- Jame, R., Johnston, R., Markov, S., and Wolfe, M. C. (2016). The value of crowdsourced earnings forecasts. *Journal of Accounting Research*, 54(4):1077–1110.

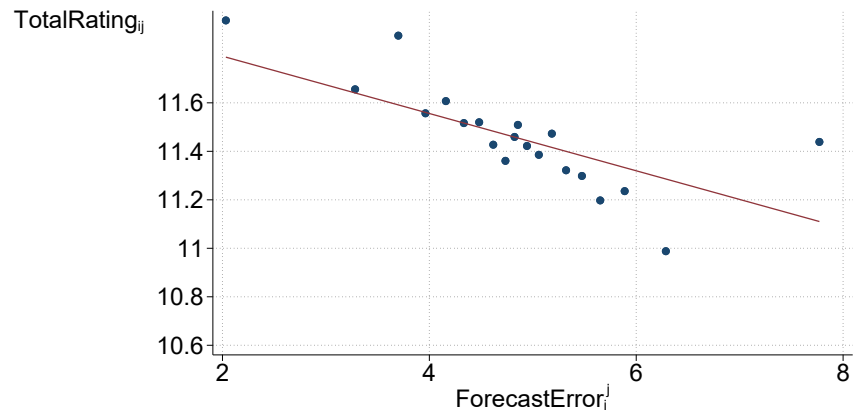
- Kadan, O., Madureira, L., Wang, R., and Zach, T. (2009). Conflicts of Interest and Stock Recommendations: The Effects of the Global Settlement and Related Regulations. *The Review of Financial Studies*, 22(10):4189–4217.
- Kempf, E. (2020). The job rating game: Revolving doors and analyst incentives. *Journal of Financial Economics*, 135(1):41–67.
- Klein, A., Saunders, A., and Wong, Y. T. F. (2019). Is there a quid pro quo between hedge funds and sell-side equity analysts? *The Journal of Portfolio Management*, 45(6):117–132.
- Law, K. K. F. and Mills, L. F. (2019). Financial gatekeepers and investor protection: Evidence from criminal background checks. *Journal of Accounting Research*, 57(2):491–543.
- Lee, J., Li, T., and Shin, D. (2019). The wisdom of crowds in FinTech: Evidence from Initial Coin Offerings. Working paper.
- Lee, J. and Parlour, C. A. (2020). Consumers as financiers: Consumer surplus, crowdfunding, and Initial Coin Offerings. Working paper.
- Leone, A. J. and Wu, J. (2002). What does it take to become a superstar? Evidence from institutional investor rankings of financial analysts. Working paper.
- Leuz, C., Jakab, L., Malani, A., and Muhn, M. (2021). Do conflict of interests disclosures work? evidence from citations in medical journals. Working paper.
- Li, J. and Mann, W. (2019). Initial Coin Offering: Current research and future directions. In *Handbook of Alternative Finance*. Palgrave-MacMillan, forthcoming.
- Li, J. and Mann, W. (2020). Initial Coin Offerings and platform building. Working paper.
- Lin, H. and McNichols, M. F. (1998). Underwriting relationships, analysts’ earnings forecasts and investment recommendations. *Journal of Accounting and Economics*, 25(1):101–127.
- Livingston, M. B., Wei, J. D., and Zhou, L. (2010). Moody’s and S&P ratings: Are they equivalent? Conservative ratings and split rated bond yields. *Journal of Money, Credit and Banking*, 42(7).
- Loughran, T. and McDonald, B. (2011). When is a liability not a liability? *Journal of Finance*, 66(1):35–65.
- Lourie, B. (2019). The revolving door of sell-side analysts. *The Accounting Review*, 94(1):249–270.
- Lyandres, E. (2019). Product market competition with crypto tokens and smart contracts. Working paper.
- Lyandres, E., Palazzo, B., and Rabetti, D. (2020). ICO success and post-ICO performance. Working paper.
- Malinova, K. and Park, A. (2018). Tokenomics: When tokens beat equity. Working paper.

- Malmendier, U. and Shanthikumar, D. (2007). Are small investors naive about incentives? *Journal of Financial Economics*, 85(2):457–489.
- Merkley, K., Michaely, R., and Pacelli, J. (2017). Does the scope of the sell-side analyst industry matter? An examination of bias, accuracy, and information content of analyst reports. *The Journal of Finance*, 72(3):1285–1334.
- Michaely, R. and Womack, K. L. (1999). Conflict of interest and the credibility of underwriter analyst recommendations. *Review of Financial Studies*, 12(4):653–686.
- Mola, S., Rau, P. R., and Khorana, A. (2013). Is there life after the complete loss of analyst coverage? *The Accounting Review*, 88(2):667–705.
- Momtaz, P. P. (2020). Initial Coin Offerings. *PLOS ONE*, 15(5):1–30.
- O’Brien, P., McNichols, M. F., and Hsiou-Wei, L. (2005). Analyst impartiality and investment banking relationships. *Journal of Accounting Research*, 43(4):623–650.
- PwC (2019). 5th ICO / STO report. PwC study.
- Roosenboom, P., van der Kolk, T., and de Jong, A. (2020). What determines success in Initial Coin Offerings? *Venture Capital*, 22(2):1–23.
- Samieifar, S. and Baur, D. G. (2020). Read me if you can! An analysis of ICO white papers. *Finance Research Letters*, page 101427.
- Sufi, A. (2009). The real effects of debt certification: Evidence from the introduction of bank loan ratings. *The Review of Financial Studies*, 22(4):1659–1691.
- Welch, I. (2000). Herding among security analysts. *Journal of Financial Economics*, 58(3):369–396.
- Zetzsche, D. A., Buckley, R. P., Arner, D. W., and Föhr, L. (2019). The ICO Gold Rush: It’s a Scam, It’s a Bubble, It’s a Super Challenge for Regulators. *Harvard International Law Journal*, 60(2):267–315.
- Zhang, S., Aerts, W., Lu, L., and Pan, H. (2019). Readability of token whitepaper and ICO first-day return. *Economics Letters*, 180(C):58–61.

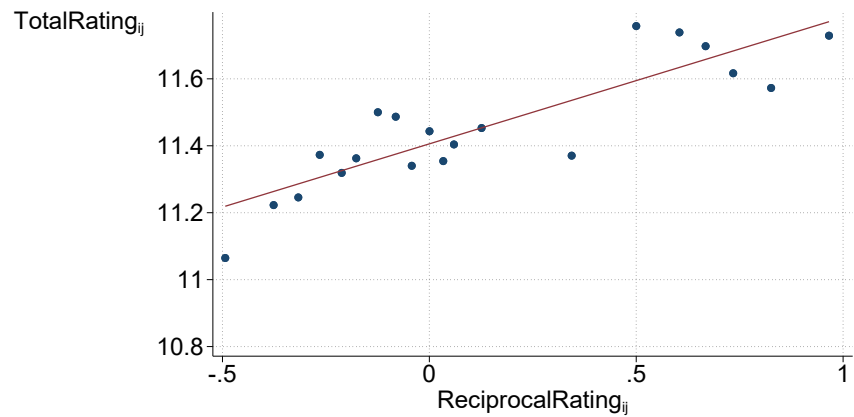
Figure 1: Summary of the main results

The figure shows binned scatter plots summarizing the main results. Panels (a) and (b) use within ICO variation, i.e., ICO fixed effects are absorbed. All variables are defined in Table A1.

(a) Optimistic analysts become more careful



(b) Reciprocal ratings are more favorable



(c) Even ICOs with high average ratings fail frequently, and especially so when many ratings are reciprocal

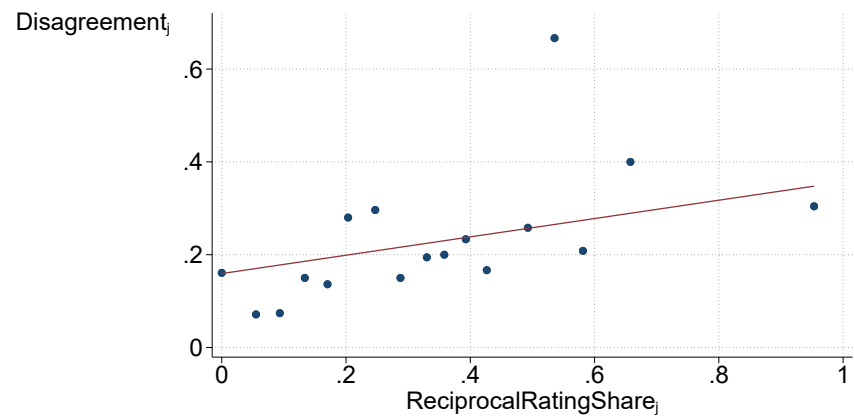


Figure 2: Number of ratings in a month and the Bitcoin price in \$

This figure presents the number of ratings in a given month over time, the number of the newly announced ICOs, the number of ratings by analysts who registered in the same month, as well as the Bitcoin price in \$.

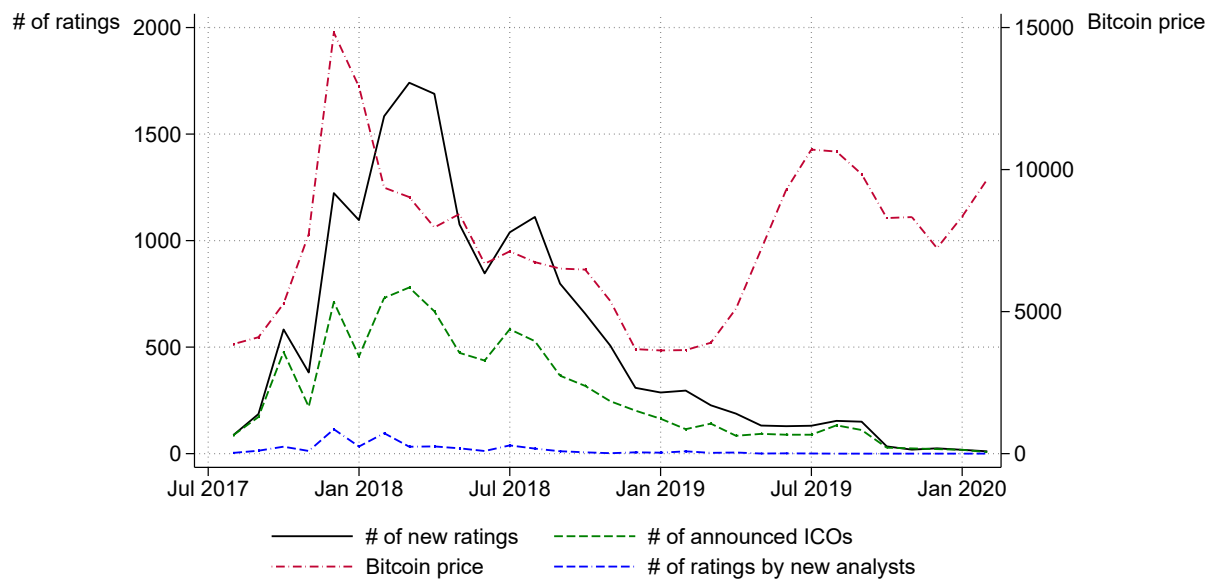


Figure 3: ICO analyst ratings and rating dispersion over time

This figure plots average total rating and analysts' rating dispersion (left axis) as well as the Bitcoin price in \$ (right axis).

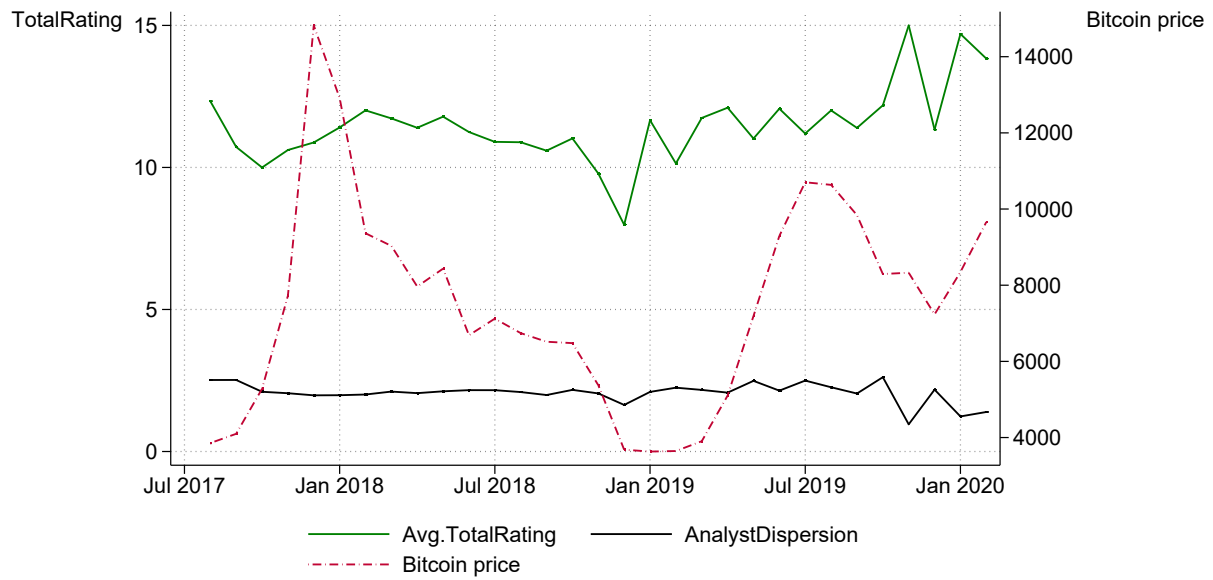


Figure 4: Successful and unsuccessful ICOs over time

The figure shows the number of ICOs over time, distinguishing between successful and failed ICOs. An ICO is labeled successful if the related coin successfully completes the offering and receives funding. In total, we identify 5,384 ICOs of which 1,932 ICOs succeeded.

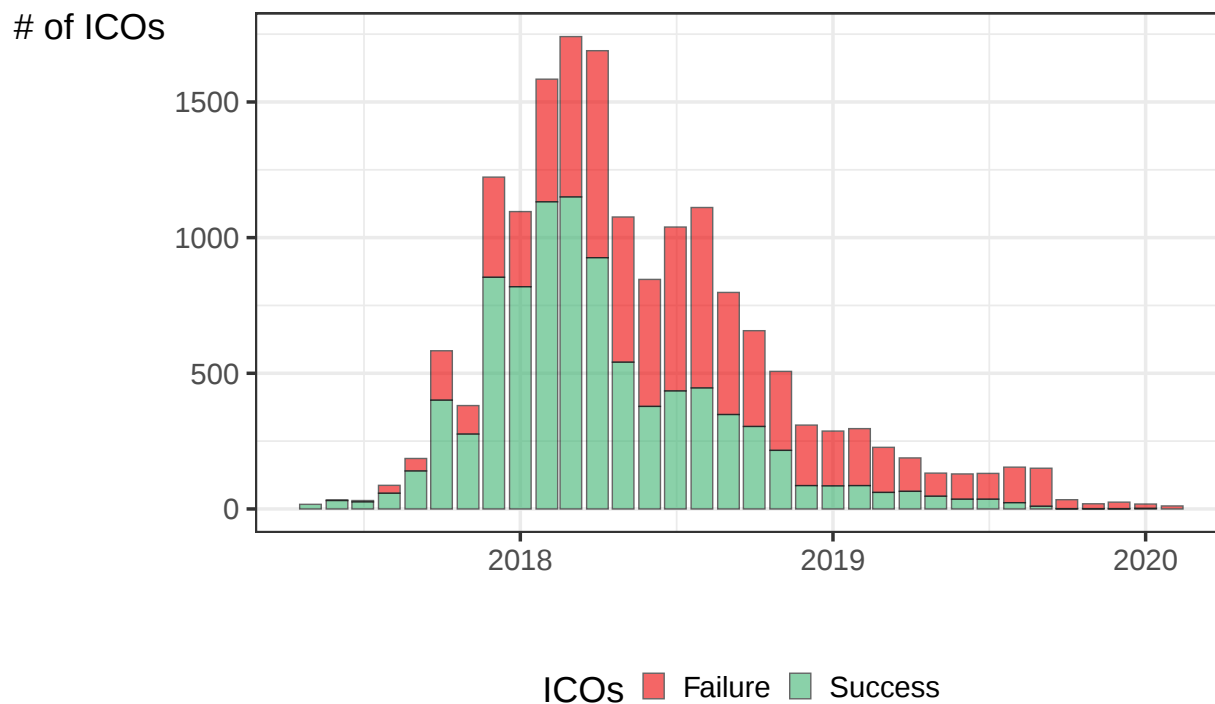


Figure 5: Forecast error over time

This figure shows the average forecast error and the number of successful ICOs (as a share of total ICOs) in a given month. The average forecast error is further split into $ForecastErrorOptimistic_i$ and $ForecastErrorPessimistic_i$ to capture the monthly averaged forecast error separately for the ratings when the analyst was too optimistic and pessimistic, respectively.

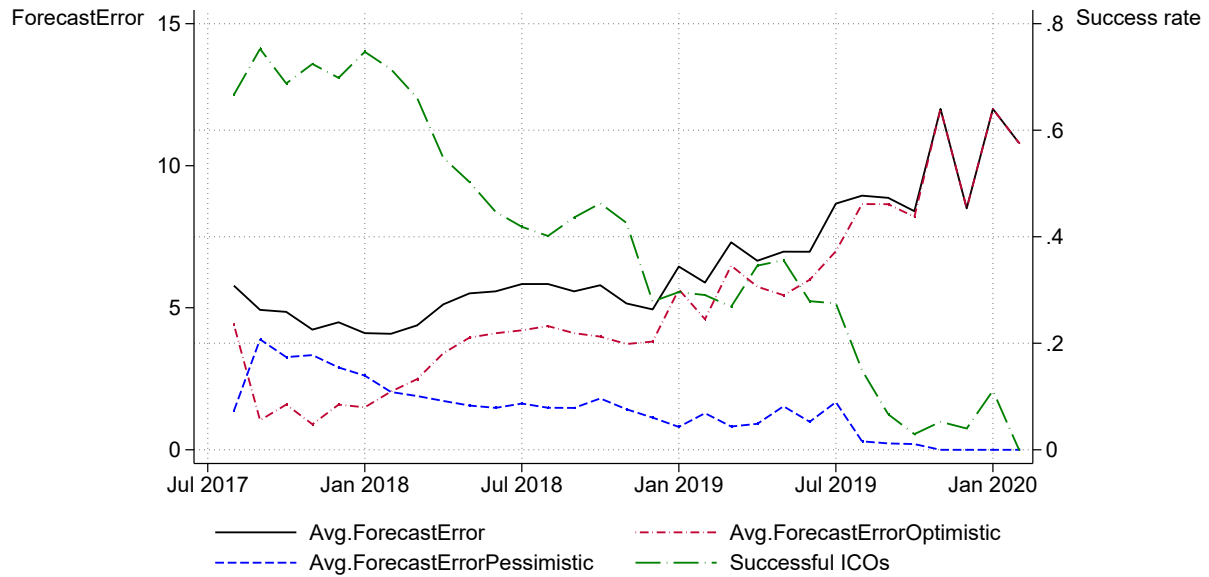


Table 1: ICO affiliation of analysts from the platform ICObench.com

This table tabulates the distribution of ICO projects among analysts. The total number of analysts in our sample is 531. The list of associated ICOs for each analyst is available on their webpage in ICObench.com

Number of associated ICOs	Count
0	230
1	127
2	52
3	29
4	24
5	15
6	9
7	13
8	3
9	8
≥ 10	49
Total number of analysts	531

Table 2: ICO analyst networks: An example

This table presents a hypothetical example of our data set. In Panel A, we show the team members of the three ICOs in the sample, namely, “A-Tokens” where Adam and Ashley are among the team members, “Bethereum” where the team includes Barbara and Benjamin, and “CryptoPay” with Cora and Chris in the team. In Panel B, we outline a hypothetical rating history. For example, in October 2017, Ashley (member of A-Tokens) provides a rating of 12 for Bethereum. In December 2017, Chris (member of CryptoPay) provides a rating of 15 for Bethereum. For this rating, we set *ReciprocalRating* equal to one because, in a month before that, in November 2017, Benjamin (member of Bethereum) gave a rating of 14 for CryptoPay, with which Chris is affiliated. Hence, we consider the rating given in December 2017 as a reply to the rating received in November 2017.

A. ICOs and members:

A-Tokens	Bethereum	CryptoPay	...
Adam Ashley	Barbara Benjamin	Cora Chris	

B. Ratings:

Date	Analyst	provides a rating for:	TotalRating	ReciprocalRating	ReceivedTotalRating if <i>ReciprocalRating</i> = 1
1)Oct 2017	Ashley	Bethereum	12	0	-
2)Nov 2017	Benjamin	CryptoPay	14	0	-
3)Dec 2017	Chris	Bethereum	15	1	14
4)Jan 2018	Adam	CryptoPay	9	0	-
...					

Table 3: Descriptive statistics

This table shows descriptive statistics of the variables used in the analysis. The variables are sorted alphabetically within each panel. The sample consists of 5,384 ICOs listed in ICObench.com, of which, 2,378 received 13,834 ratings in total. All variables are defined in Table A1.

	N	Min	P25	Mean	P50	P75	Max	Std. Dev.
<i>A. Rating characteristics</i>								
<i>ForecastError_i^j</i>	12,460	0	4.1	4.8	4.8	5.6	12	1.5
<i>Modified_{ij}</i>	13,834	0	0	.13	0	0	1	.33
<i>OrderRank_{ij}</i>	7,639	1	6	14	11	18	94	12
<i>ProductRating_{ij}</i>	13,834	1	3	3.6	4	5	5	1.1
<i>ReceivedProductRating_{ij}</i>	1,754	1	4	4.1	4	5	5	.74
<i>ReceivedTeamRating_{ij}</i>	1,754	1	4	4.3	4	5	5	.67
<i>ReceivedTotalRating_{ij}</i>	1,754	3	12	13	12	14	15	1.8
<i>ReceivedVisionRating_{ij}</i>	1,754	1	4	4.2	4	5	5	.7
<i>ReciprocalRating_{ij}</i>	13,834	0	0	.13	0	0	1	.33
<i>ReviewLength_{ij}</i>	9,165	1.1	3.4	3.8	3.9	4.3	7.9	.97
<i>ReviewTone_{ij}</i>	9,165	-.75	-.043	-.014	-.0086	.018	.67	.075
<i>StarAnalyst_{ij}</i>	13,834	0	0	.27	0	1	1	.44
<i>TeamRating_{ij}</i>	13,834	1	3	3.9	4	5	5	1.1
<i>TotalRating_{ij}</i>	13,834	3	10	11	12	14	15	3
<i>VisionRating_{ij}</i>	13,834	1	3	3.9	4	5	5	1.1
<i>B. ICO characteristics</i>								
<i>AmountRaised_j</i>	5,339	0	0	5.3	0	14	22	7.3
<i>AnalystDispersion_j</i>	1,638	0	1.2	2	1.9	2.6	8.5	1.3
<i>Bench_j</i>	5,339	.1	2.4	2.9	2.9	3.5	5	.75
<i>Bitcointalk_j</i>	5,339	0	0	.57	1	1	1	.49
<i>Bonus_j</i>	5,339	0	0	.14	0	0	1	.35
<i>Bounty_j</i>	5,339	0	0	.28	0	1	1	.45
<i>Disagreement_j</i>	2,378	0	0	.17	0	0	1	.38
<i>ForecastError_j</i>	2,322	0	4.3	5	4.9	5.7	11	1.1
<i>IEO_j</i>	5,339	0	0	.052	0	0	1	.22
<i>KYC_j</i>	5,339	0	0	.49	0	1	1	.5
<i>LengthWhitePaper_j</i>	5,339	0	0	1.2	0	0	11	3.1
<i>MVP_j</i>	5,339	0	0	.2	0	0	1	.4
<i>MarketCap_j^{all}</i>	5,339	0	0	1.9	0	0	22	5.1
<i>MarketCap_j</i>	813	0	12	12	15	16	22	6.5
<i>Presale_j</i>	5,339	0	0	.53	1	1	1	.5
<i>PreviousRatings_j</i>	2,322	3	11	11	11	12	15	1.4
<i>ReciprocalRatingShare_j</i>	2,378	0	0	.072	0	0	1	.19
<i>RetentionRatio_j</i>	4,224	0	30	46	45	60	100	21
<i>ReviewComplexity_j</i>	1,883	4.6	11	12	12	14	59	3.2
<i>ReviewLength_j</i>	1,883	1.1	3.6	4	4	4.5	7.3	.81
<i>ReviewToneDispersion_j</i>	1,240	0	.03	.057	.046	.073	.55	.046
<i>ReviewTone_j</i>	1,883	-.67	-.04	-.02	-.015	.003	.29	.056
<i>ReviewUncertainty_j</i>	1,883	0	0	.016	.011	.021	.33	.022
<i>Scam_j</i>	5,339	0	0	.043	0	0	1	.2
<i>StarAnalysts_j</i>	2,378	0	0	.31	.22	.5	1	.34
<i>Success_j</i>	5,339	0	0	.35	0	1	1	.48
<i>#Analysts_j</i>	5,339	0	0	2.6	0	2	94	6.1
<i>Facebook_j</i>	5,339	0	1	.78	1	1	1	.41

Table 4: Rating determinants

This table presents linear regression results for Equation 1. The dependent variable is the total rating score that an analyst gave an ICO. All variables are defined in Table A1. t -statistics are given in parentheses. Standard errors are clustered at the ICO and analyst level. ***, **, * indicate significance at the 1%, 5% and 10% levels.

Dep. variable:	<i>TotalRating_{ij}</i>				
	(1)	(2)	(3)	(4)	(5)
<i>Bench_{yj}</i>	1.518*** (10.43)	1.598*** (11.35)		0.996*** (7.01)	
<i>StarAnalyst_{ij}</i>		-0.720*** (-3.58)	-0.699*** (-3.41)	-0.701*** (-3.47)	-0.454** (-2.47)
<i>ForecastError_i^j</i>		-0.158*** (-3.61)	-0.156*** (-3.70)	-0.165*** (-4.12)	-0.081** (-2.31)
<i>MVP_j</i>			0.264* (1.79)	-0.084 (-0.58)	
<i>IEO_j</i>			0.850*** (3.98)	0.623*** (2.71)	
<i>Presale_j</i>			0.311** (2.40)	0.242** (1.99)	
<i>KYC_j</i>			0.912*** (5.36)	0.548*** (3.42)	
<i>Bounty_j</i>			0.258** (1.98)	0.179 (1.45)	
<i>Bonus_j</i>			0.145 (1.20)	0.148 (1.29)	
<i>RetentionRatio_j</i>			0.006** (2.04)	0.005* (1.72)	
<i>LengthWhitePaper_j</i>			0.012 (0.69)	-0.007 (-0.39)	
<i>Bitcointalk_j</i>			0.070 (0.36)	-0.026 (-0.14)	
<i>Facebook_j</i>			0.312 (1.26)	0.302 (1.24)	
Observations	13834	12460	11257	11257	11698
R^2	0.121	0.146	0.067	0.099	0.528
MonthRating Dummies	No	No	No	No	Yes
ICO FE	No	No	No	No	Yes

Table 5: Reciprocal ratings

This table presents linear regression results for Equation 2. The dependent variable is the total rating score that an analyst gave an ICO. In columns (1), (3), (5) and (7), regressions include all the ratings in the sample. In columns (2), (4), (6) and (8), we restrict the sample to the reciprocal ratings (*ReciprocalRating* = 1). All specifications include *Analyst* and *ICO* fixed effects multiplied by dummies for the time of the rating (i.e., *Analyst* $\times$ *Month* and *ICO* $\times$ *Month* fixed effects in odd columns and *Analyst* $\times$ *Quarter* and *ICO* $\times$ *Quarter* fixed effects in even columns). All variables are defined in Table A1. *t*-statistics are given in parentheses. Standard errors are clustered at the ICO and analyst level. ***, **, * indicate significance at the 1%, 5% and 10% levels.

Dep. variable:	<i>TotalRating_{ij}</i>		<i>TeamRating_{ij}</i>		<i>VisionRating_{ij}</i>		<i>ProductRating_{ij}</i>	
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>ReciprocalRating_{ij}</i>	0.252** (2.47)		0.062* (1.80)		0.074* (1.76)		0.117*** (2.85)	
<i>ReceivedTotalRating_{ij}</i>		0.080* (1.74)						
<i>ReceivedTeamRating_{ij}</i>				0.117*** (3.13)				
<i>ReceivedVisionRating_{ij}</i>						0.065 (1.14)		
<i>ReceivedProductRating_{ij}</i>								0.000 (0.01)
<i>Modified_{ij}</i>		-1.114*** (-3.54)		-0.381*** (-3.02)		-0.298** (-2.53)		-0.427*** (-3.46)
Observations	10354	1302	10354	1302	10354	1302	10354	1302
<i>R</i> ²	0.757	0.682	0.717	0.621	0.692	0.666	0.708	0.647
<i>Analyst</i> $\times$ <i>Time_{ij}</i> Dummies	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
<i>ICO</i> $\times$ <i>Time_{ij}</i> Dummies	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes

Table 6: Linguistic nature of rating reviews

This table presents linear regression results for Equation 3. The dependent variable in Panel A is *ReviewLength*, defined as the natural logarithm of the total number of words in a review, and in Panel B *ReviewTone*, defined as the ratio of positive words minus negative words to total words in the review. We restrict the sample to reciprocal ratings (*ReciprocalRating* = 1) in column (3) and to non-reciprocal ratings (*ReciprocalRating* = 0) in column (4). We include *Analyst* and *ICO* fixed effects multiplied by dummies for the month of ratings (i.e., *Analyst* $\times$ *Month* and *ICO* $\times$ *Month* fixed effects) in column (2). As in Table 5, we can only include the interaction of ICO (analyst) and quarter dummies when restricting the sample to reciprocal ratings in column (3), i.e., *Analyst* $\times$ *Quarter* and *ICO* $\times$ *Quarter* fixed effects. In order to compare the coefficients for reciprocal and non-reciprocal ratings, we also include *Analyst* $\times$ *Quarter* and *ICO* $\times$ *Quarter* fixed effects in column (4). All variables are defined in Table A1. *t*-statistics are given in parentheses. Standard errors are clustered at the ICO and analyst level. ***, **, * indicate significance at the 1%, 5% and 10% levels.

Dep. Variable:	Panel A			
	<i>ReviewLength_{ij}</i>			
	(1)	(2)	(3)	(4)
<i>TotalRating_{ij}</i>	-0.056*** (-5.42)	-0.041*** (-5.41)	-0.090*** (-5.10)	-0.036*** (-5.37)
Observations	9165	6206	866	6119
R^2	0.033	0.825	0.800	0.786
<i>Analyst</i> $\times$ <i>Time_{ij}</i> Dummies	No	Yes	Yes	Yes
<i>ICO</i> $\times$ <i>Time_{ij}</i> Dummies	No	Yes	Yes	Yes

Dep. Variable:	Panel B			
	<i>ReviewTone_{ij}</i>			
	(1)	(2)	(3)	(4)
<i>TotalRating_{ij}</i>	0.006*** (10.09)	0.006*** (7.51)	0.009*** (3.34)	0.005*** (8.17)
Observations	9165	6206	866	6119
R^2	0.062	0.537	0.522	0.477
<i>Analyst</i> $\times$ <i>Time_{ij}</i> Dummies	No	Yes	Yes	Yes
<i>ICO</i> $\times$ <i>Time_{ij}</i> Dummies	No	Yes	Yes	Yes

Table 7: Order of rating issuance

This table presents linear regression results for Equation 4. The dependent variable is the order rank of the rating for an ICO. A lower value of the variable indicates that analyst i issued the rating for ICO j earlier. All specifications include month dummies of the analyst rating. The sample is restricted to ICOs with more than ten ratings. All variables are defined in Table A1. t -statistics are given in parentheses. Standard errors are clustered at the ICO and analyst level. ***, **, * indicate significance at the 1%, 5% and 10% levels.

Dep. variable:	<i>OrderRank_{ij}</i>					
	(1)	(2)	(3)	(4)	(5)	(6)
<i>TotalRating_{ij}</i>	-0.104 (-1.65)		-0.085 (-1.36)	-0.182*** (-2.82)		-0.177*** (-2.75)
<i>ReciprocalRating_{ij}</i>		-1.764*** (-3.17)	-1.725*** (-3.08)		-1.249** (-2.56)	-1.213** (-2.48)
<i>StarAnalyst_{ij}</i>	-1.109*** (-3.00)	-0.855*** (-2.60)	-0.886*** (-2.69)			
<i>ForecastError_i^j</i>	-0.028 (-0.32)	-0.027 (-0.31)	-0.036 (-0.42)	-0.130 (-1.18)	-0.133 (-1.22)	-0.122 (-1.12)
Observations	6829	6829	6829	6767	6767	6767
R^2	0.672	0.674	0.674	0.709	0.709	0.709
MonthRating Dummies	Yes	Yes	Yes	Yes	Yes	Yes
ICO FE	Yes	Yes	Yes	Yes	Yes	Yes
Analyst FE	No	No	No	Yes	Yes	Yes

Table 8: Ratings and ICO success: Descriptive evidence

This table presents descriptive statistics for the relationship between ratings and ICO success. Panel A shows the success of ICOs that human analysts did or did not cover. Panel B links ICO success to the quantitative rating score. Panel C shows investor disagreement for ICOs with or without any reciprocal rating. Market capitalization displays the ICO value on the 90th day post listing. It is set to zero for all ICOs that were not listed on CoinMarketCap.com.

Panel A

Analyst Coverage	Total #	Funded #	in %	Ln(MarketCap) avg. in \$	AmountRaised avg. in \$	Ln(AmountRaised) avg. in \$
No	2,961	858	28.97	1.59	19,326,718	4.27
Yes	2,378	1,033	43.44	2.19	12,692,225	6.60
Total	5,339	1,891	35.42	1.86	15,704,397	5.34

Panel B

TotalRating Score	Total #	Funded #	in %	Ln(MarketCap) avg. in \$	AmountRaised avg. in \$	Ln(AmountRaised) avg. in \$
3	211	46	21.80	0.72	23,140,030	3.10
4–6	250	52	20.80	0.79	6,313,606	2.98
7–9	486	174	35.80	1.98	14,316,201	5.44
10–12	1,071	545	50.89	2.52	12,164,440	7.75
13–15	754	392	51.99	3.49	31,994,976	8.14

Panel C

Reciprocal Rating	Ln(MarketCap) avg. in \$	Total #	Disagreement #	Disagreement in %	Disagreement with Avg. Rating ≥ 13 #	Disagreement with Avg. Rating ≥ 13 in %
Yes*	1.49	415	97	23.37	96	23.13
No**	2.33	1,963	316	16.10	272	13.85

* $ReciprocalRatingShare_j > 0$

** $ReciprocalRatingShare_j = 0$

Table 9: Ratings and ICO success

This table presents marginal effects of logit regressions and coefficients of linear regressions for Equation 5. The dependent variable is the *Success* dummy in columns (1)-(3) and the market capitalization 90 days after listing on an exchange in columns (4)-(6). In columns (4)-(5), we set the dependent variable to zero for all ICOs that were not listed on CoinMarketCap.com. The controls for which coefficients are not shown for space reasons include the retention ratio and dummies for pre-sale, bonus/bounty options, *KYC*, Bitcointalk, and Facebook. All specifications include month dummies. Because the logit model predicts failure perfectly in some months, we lose a few observations from the inclusion of month fixed effects. All variables are defined in Table A1. *t*-statistics based on robust standard errors are reported in parentheses. ***, **, * indicate significance at the 1%, 5% and 10% levels.

	<i>Success_j</i>			<i>MarketCap_j</i>		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>TotalRating_j</i>	0.110*** (5.47)	0.130*** (3.99)	0.072* (1.65)	0.150*** (2.58)	0.186*** (2.62)	0.293 (0.92)
<i># Analysts_j</i>	0.026*** (3.51)	0.020** (2.56)	0.025*** (2.76)	0.054*** (2.77)	0.048** (2.18)	0.043* (1.85)
<i>StarAnalysts_j</i>		-0.105 (-0.40)	-0.165 (-0.46)	-0.543 (-0.86)	-0.207 (-0.29)	4.002 (1.44)
<i>ReciprocalRatingShare_j</i>		-0.110 (-0.35)	-0.111 (-0.31)	-3.306*** (-5.05)	-2.622*** (-3.30)	-4.393** (-2.25)
<i>PreviousRatings_j</i>		-0.020 (-0.27)	-0.067 (-0.75)	0.072 (0.50)	-0.126 (-0.73)	-0.446 (-0.74)
<i>AnalystDispersion_j</i>		-0.003 (-0.07)	-0.060 (-0.97)	-0.166* (-1.70)	-0.257** (-2.45)	-0.754 (-1.32)
<i>Bench_j</i>	0.682*** (8.15)	0.774*** (7.15)	0.852*** (5.37)	1.222*** (5.31)	1.103*** (3.53)	0.380 (0.49)
<i>ReviewToneDispersion_j</i>			-0.626 (-0.30)		-1.221 (-0.23)	-2.577 (-0.29)
<i>ReviewTone_j</i>			0.952 (0.48)		-7.314 (-1.49)	3.976 (0.34)
<i>ReviewUncertainty_j</i>			-5.779 (-1.11)		2.752 (0.20)	19.623 (1.18)
<i>ReviewComplexity_j</i>			-0.001 (-0.02)		0.008 (0.12)	0.030 (0.20)
<i>ReviewLength_j</i>			0.014 (0.11)		-0.525 (-1.64)	-0.997 (-1.35)
<i>MVP_j</i>			-0.374** (-2.11)		-0.610* (-1.67)	-0.597 (-0.56)
<i>IEO_j</i>			0.752** (2.42)		1.568** (2.21)	1.685 (0.98)
<i>LengthWhitePaper_j</i>			0.040* (1.87)		0.167*** (3.18)	0.056 (0.61)
Observations	2372	1612	1111	1629	1128	225
(Pseudo) <i>R</i> ²	0.146	0.154	0.186	0.122	0.164	0.361
Controls	No	No	Yes	No	Yes	Yes
MonthICO Dummies	Yes	Yes	Yes	Yes	Yes	Yes

Table 10: ICO outcomes that deviate from what ratings predict

This table presents marginal effects of logit regressions for Equation 6. The dependent variable is the *Disagreement* dummy which equals one if (i) analysts give an average $TotalRating_j \geq 13$ and the ICO fails, or if (ii) analysts give an average $TotalRating_j \leq 5$ and the ICO succeeds. In column (3), we restrict the sample to cases where the reciprocal ratings are on average greater than or equal to the average of non-reciprocal ratings for the same ICO. In column (4), we restrict the sample to ICOs where the average reciprocal rating is lower than the average of non-reciprocal ratings. All analyst variables are average values of every analyst that rates the ICO. Control variables for which coefficients are not shown for space reasons include the retention ratio and dummies for pre-sale, bonus/bounty options, *KYC*, *Bitcointalk*, and *Facebook*. All specifications include month dummies. All variables are defined in Table A1. *t*-statistics based on robust standard errors are reported in parentheses. ***, **, * indicate significance at the 1%, 5% and 10% levels.

Dep. variable:	<i>Disagreement_j</i>			
	(1)	(2)	(3)	(4)
<i># Analysts_j</i>	0.007 (0.83)	0.008 (0.91)	0.004 (0.20)	-0.001 (-0.03)
<i>StarAnalysts_j</i>	-0.762** (-2.11)	-0.609 (-1.55)	0.092 (0.08)	-2.598 (-1.56)
<i>ReciprocalRatingShare_j</i>	1.028*** (2.80)	0.885** (2.18)	2.697* (1.87)	0.900 (0.46)
<i>PreviousRatings_j</i>	0.295*** (3.24)	0.211** (2.08)	0.753** (2.46)	1.495** (2.22)
<i>AnalystDispersion_j</i>	-0.458*** (-7.07)	-0.474*** (-6.37)	-0.698*** (-2.82)	-0.504 (-1.45)
<i>Bench_j</i>	0.055 (0.48)	-0.124 (-0.83)	-1.065 (-1.53)	0.166 (0.22)
<i>ReviewTone_j</i>		8.626*** (4.09)	3.343 (0.59)	9.535 (0.85)
<i>ReviewUncertainty_j</i>		-0.010 (-0.00)	32.704 (1.61)	9.264 (0.32)
<i>ReviewComplexity_j</i>		0.064* (1.90)	0.139 (1.08)	-0.247 (-1.50)
<i>ReviewLength_j</i>		-0.072 (-0.49)	-0.421 (-0.82)	0.873 (1.26)
<i>MVP_j</i>		0.106 (0.53)	0.544 (0.86)	0.486 (0.63)
<i>IEO_j</i>		-0.722** (-2.29)	-0.549 (-0.39)	0.450 (0.38)
<i>LengthWhitePaper_j</i>		-0.053** (-2.11)	-0.060 (-0.78)	-0.111 (-1.55)
Observations	1615	1222	188	116
Pseudo R^2	0.194	0.201	0.235	0.270
Controls	No	Yes	Yes	Yes
<i>Month_j Dummies</i>	Yes	Yes	Yes	Yes

Table 11: ICO scams

This table presents marginal effects of logit regressions analogous to Equation 5, with the dependent variable being the *Scam* dummy. Control variables for which coefficients are not shown for space reasons include the retention ratio and dummies for pre-sale, bonus/bounty options, *KYC*, Bitcointalk, and Facebook. All specifications include month dummies. All variables are defined in Table A1. *t*-statistics based on robust standard errors are reported in parentheses. ***, **, * indicate significance at the 1%, 5% and 10% levels.

Dep. variable:	<i>Scam_j</i>		
	(1)	(2)	(3)
<i>TotalRating_j</i>	0.059 (1.31)	-0.009 (-0.13)	0.066 (0.55)
<i># Analysts_j</i>	0.027*** (2.82)	0.027*** (2.73)	0.031** (2.18)
<i>StarAnalysts_j</i>		-0.272 (-0.42)	-0.918 (-0.91)
<i>ReciprocalRatingShare_j</i>		-0.157 (-0.22)	-0.066 (-0.08)
<i>PreviousRatings_j</i>		0.248 (1.61)	0.517** (2.04)
<i>AnalystDispersion_j</i>		0.274** (2.34)	0.564*** (4.09)
<i>Benchy_j</i>	-0.217 (-1.20)	-0.386** (-2.05)	-0.208 (-0.58)
<i>ReviewToneDispersion_j</i>			5.808** (2.01)
<i>ReviewTone_j</i>			-1.859 (-0.53)
<i>ReviewUncertainty_j</i>			-9.649 (-1.11)
<i>ReviewComplexity_j</i>			-0.063 (-1.02)
<i>ReviewLength_j</i>			-0.132 (-0.41)
<i>MVP_j</i>			0.615 (1.62)
<i>IEO_j</i>			-0.606 (-0.80)
<i>LengthWhitePaper_j</i>			-0.075 (-1.26)
Observations	2127	1466	966
Pseudo <i>R</i> ²	0.057	0.066	0.191
Controls	No	No	Yes
<i>Month_j Dummies</i>	Yes	Yes	Yes

Appendix

Table A1: Variable definitions

Variable	Definition
$\#Analysts_j$	Number of analysts that rate an ICO.
$AmountRaised_j$	Natural logarithm of one plus \$ amount raised by an ICO.
$AnalystDispersion_j$	Standard deviation of ratings within an ICO.
$Benchy_j$	Machine-generated rating created by ICObench.com.
$Bitcointalk_j$	Dummy variable that equals one if the ICO is discussed on the forum bitcointalk.org.
$Bonus_j$	Dummy variable that equals one for ICOs with a quantity discount at the token sale or a discount program for early-bird investors.
$Bounty_j$	Dummy variable that equals one for ICOs with incentives to promote social media presence.
$Disagreement_j$	Dummy variable that equals one if (i) on average, analysts recommend buying ($TotalRating_j \geq 13$) and the ICO fails, or (ii) on average, analysts recommend selling ($TotalRating_j \leq 5$) and the ICO succeeds.
$Facebook_j$	Dummy variable that equals one if an ICO has a Facebook page.
$ForecastError_{ij}$	The distance of the total rating from the highest (lowest) possible rating in the case of ICO success (failure).
$ForecastErrorOptimistic_i$	The distance of the highest possible rating score to the actual total rating of analyst i , defined as $15 - RatingTotal$, if the ICO was unsuccessful, and averaged over all ICOs j .
$ForecastErrorPessimistic_i$	The distance of the total rating of analyst i to the lowest possible rating score, defined as $RatingTotal - 3$, if the ICO was successful, and averaged over all ICOs j .
$ForecastError_i^j$	A recursive average of all previous forecast errors for any analyst i up to the rating issuance date for ICO j .
$ForecastError_j$	A recursive average of the previous forecast errors of all analysts covering ICO j up to the rating issuance date.
IEO_j	Dummy variable that equals one for ICOs conducted on the platform of a cryptocurrency exchange (Initial Exchange Offerings).
KYC_j	Dummy variable that equals one for ICOs where investors are required to sign up to a whitelist using their wallet address to receive access to the ICO sale (Know Your Customer).

<i>LengthWhitePaper</i>	The natural logarithm of (1 + total words of the white paper), set to 0 if no white paper could be found.
<i>MarketCap_j^{listed}</i>	The natural logarithm of market capitalization 90 days after listing on an exchange from CoinMarketCap.com.
<i>MarketCap_j</i>	In a modified version of the variable, we set the value to zero if no information about the ICO was found on CoinMarketCap.com.
<i>Modified_{ij}</i>	Dummy variable that equals one if the rating for ICO <i>j</i> was modified by analyst <i>i</i> at any point in time.
<i>Month_j</i>	Dummy variable for each month, indicating the month when an ICO was launched.
<i>Month_{ij}</i>	Dummy variable for each month, indicating the month when a rating was given.
<i>MVP_j</i>	Dummy variable that equals one for ICOs with a prototype. This can be a version of a new product with sufficient features to satisfy early adopters (minimum viable product) or drafts of code on Github.com that are open to discussion by other GitHub users.
<i>Presale_j</i>	Dummy variable that equals one if an ICO features a token sale event that runs prior to the official ICO campaign.
<i>OrderRank_{ij}</i>	The order rank of the rating by analyst <i>i</i> issued for ICO <i>j</i> in a given month.
<i>PreviousRatings_j</i>	Average past TotalRating of all analysts that provide a rating for ICO <i>j</i>
<i>ReceivedTeamRating_{ij}/</i> <i>ReceivedVisionRating_{ij}/</i> <i>ReceivedProductRating_{ij}/</i> <i>ReceivedTotalRating_{ij}</i>	Level of the rating when ReciprocalRating dummy equals 1, i.e., level of rating that the analyst of ICO <i>j</i> received for their own ICO from any team member of ICO <i>j</i> prior to the rating issuance date.
<i>ReciprocalRating_{ij}</i>	Dummy variable that equals one for reciprocal ratings. A rating is reciprocal when the corresponding analyst was a team member of another ICO project that previously received a rating by one of the team members of this new ICO. Table 2 represents a hypothetical illustration of our variable composition.
<i>ReciprocalRatingShare_j</i>	Share of reciprocal analysts that provide a rating for ICO <i>j</i> .
<i>RetentionRatio_j</i>	The percentage of token supply that is retained by the ICO members, and not available for sale.
<i>ReviewComplexity_j</i>	The complexity of an analyst's review text, measured by the Gunning (1952) Fog index, and averaged together on ICO level.
<i>ReviewLength_{ij}</i>	The natural logarithm of the number of total words in an analyst review. For the <i>ReviewLength_j</i> , we measure the natural logarithm of the average review text lengths for ICO <i>j</i> .

$ReviewTone_{ij}$	The tone of the analyst review text. Using the Loughran and McDonald (2011) <i>Positive</i> and <i>Negative</i> word-lists, the tone of a text is defined as the difference between the count of positive and negative words divided by the total words.
$ReviewTone_j$	The tone averaged across all analysts' review texts for ICO j .
$ReviewToneDispersion_j$	The standard deviation of $ReviewTone_{ij}$ within an ICO.
$ReviewUncertainty_j$	The uncertainty of the analysts' review texts, averaged together on ICO level. Using the Loughran and McDonald (2011) <i>Uncertainty</i> word-list, the uncertainty of a text is defined as the count of uncertain words divided by the total words.
$Scam_j$	Dummy variable that equals one for ICO projects that intentionally defraud investors.
$Success_j$	Dummy variable that equals one for ICOs that completed the token sale and collected (at least \$1) funding.
$StarAnalysts_{ij}$	Dummy variable that equals one when ICO j was rated by one of the top 30 analysts i according to a ranking on ICObench.com.
$StarAnalysts_j$	Share of the top 30 analysts that provide a rating for ICO j .
$TeamRating_{ij}/$ $VisionRating_{ij}/$ $ProductRating_{ij}$	Rating score for team/ vision/ product of an ICO, ranging from 1 (lowest) to 5 (highest).
$TotalRating_{ij}$	The sum of team, vision and product ratings for the respective ICO, ranging from 3 to 15.
$TotalRating_j$	Average rating of ICO j by all analysts.

Table A2: Ratings and ICO success: An alternative success measure

This table presents linear regression results for Equation 5. The dependent variable is the natural logarithm of the amount raised by an ICO in columns. The controls for which coefficients are not shown for space reasons include the retention ratio and dummies for pre-sale, bonus/bounty options, *KYC*, *Bitcointalk*, and *Facebook*. All specifications include month dummies. All variables are defined in Table A1. *t*-statistics based on robust standard errors are reported in parentheses. ***, **, * indicate significance at the 1%, 5% and 10% levels.

Dep. variable:	<i>AmountRaised_j</i>		
	(1)	(2)	(3)
<i>TotalRating_j</i>	0.325*** (6.11)	0.390*** (4.74)	0.246** (2.14)
<i># Analysts_j</i>	0.097*** (5.11)	0.072*** (3.57)	0.078*** (3.52)
<i>StarAnalysts_j</i>		-0.371 (-0.45)	-0.302 (-0.28)
<i>ReciprocalRatingShare_j</i>		-0.402 (-0.39)	-0.340 (-0.30)
<i>PreviousRatings_j</i>		-0.103 (-0.48)	-0.201 (-0.75)
<i>AnalystDispersion_j</i>		-0.017 (-0.12)	-0.140 (-0.80)
<i>Benchy_j</i>	1.858*** (8.53)	2.062*** (7.38)	2.182*** (5.58)
<i>ReviewToneDispersion_j</i>			-2.660 (-0.45)
<i>ReviewTone_j</i>			2.451 (0.45)
<i>ReviewUncertainty_j</i>			-16.625 (-1.05)
<i>ReviewComplexity_j</i>			0.016 (0.15)
<i>ReviewLength_j</i>			-0.042 (-0.10)
<i>MVP_j</i>			-1.252** (-2.39)
<i>IEO_j</i>			1.985** (2.32)
<i>LengthWhitePaper_j</i>			0.117* (1.85)
Observations	2378	1629	1128
<i>R</i> ²	0.190	0.208	0.247
Controls	No	No	Yes
<i>Month_j Dummies</i>	Yes	Yes	Yes

Table A3: ICO outcomes that deviate from what ratings predict

This table presents marginal effects of logit regressions for Equation 6. The dependent variable is the *Disagreement* dummy, which equals one if analysts recommend buying (average *TotalRating_j* $\geq$ 13) and the ICO fails. All analyst variables are average values over all analysts that rate the ICO. Control variables for which coefficients are not shown for space reasons include the retention ratio and dummies for pre-sale, bonus/bounty options, *KYC*, Bitcointalk, and Facebook. All specifications include month dummies. All variables are defined in Table A1. *t*-statistics based on robust standard errors are reported in parentheses. ***, **, * indicate significance at the 1%, 5% and 10% levels.

Dep. variable:	<i>Disagreement_j</i>	
	(1)	(2)
<i># Analysts_j</i>	0.042*** (2.99)	0.056*** (3.30)
<i>StarAnalysts_j</i>	-1.346*** (-3.01)	-0.886* (-1.83)
<i>ReciprocalRatingShare_j</i>	2.033*** (4.34)	1.799*** (3.21)
<i>PreviousRatings_j</i>	0.532*** (4.25)	0.452*** (3.35)
<i>AnalystDispersion_j</i>	-0.678*** (-7.67)	-0.733*** (-7.13)
<i>Bench_j</i>	0.597*** (3.76)	0.277 (1.28)
<i>ReviewTone_j</i>		17.694*** (4.27)
<i>ReviewUncertainty_j</i>		-10.057 (-1.02)
<i>ReviewComplexity_j</i>		0.135*** (3.01)
<i>ReviewLength_j</i>		-0.015 (-0.08)
<i>MVP_j</i>		0.017 (0.07)
<i>IEO_j</i>		-0.540 (-1.23)
<i>LengthWhitePaper_j</i>		-0.039 (-1.10)
Observations	776	640
Pseudo R^2	0.263	0.308
Controls	No	Yes
<i>Month_j Dummies</i>	Yes	Yes