



Burgiss Applied Research: Working Paper

February 2018

Modeling Cash Flows for Private Capital Funds

Vishv Jeet and Luis O'Shea

In this paper we focus on predicting cash flows for private capital funds. We start by discussing the characteristics of cash-flow data that make these predictions challenging. We then examine several models for expected contributions and distributions, and evaluate their performance, both in-sample and out-of-sample. We find that for contributions, uncalled capital, in addition to age, is a useful predictor. Distributions are more difficult to predict; we find that disentangling the effect of a fund's performance on distributions produces better forecasts than other simpler approaches. We compare the models explored in this paper with those outlined by Takahashi and Alexander and find that they underperform our models by a wide margin. Finally, we draw some lessons regarding how to model cash flows, and how to measure model performance. We also make some observations regarding the intersection of risk and prediction with regard to cash flows.

1 Introduction

Private capital cash flows are uncertain in both their magnitude and timing. The uncertainty of capital calls (contributions) often forces Limited Partners (LPs) to keep their uncalled capital in low-return investments, such as treasury bills, that are both liquid and less risky.¹ Large amounts of uncalled capital sitting in low-return investments, waiting to be called, could be put to better use if contributions could be predicted to some extent or, better still, if distributions could be recycled to make contributions. A good cash-flow model for investing in private capital funds is extremely desirable as it can help LPs minimize the capital sitting in low-return investments, set realistic investment targets, assess financial viability of new or existing commitments, and enhance reporting.

In this paper we discuss private capital cash-flow models to predict the *average* amount of capital that funds will call (or distribute) over a future period, such as a quarter. We model contributions and distributions separately, since contributions, being relatively unaffected by the fund's performance, can be modeled more directly than distributions. Also, it is important to understand contributions independently of distributions since they are liabilities for LPs. A fundamental econometric challenge in modeling cash-flow data of private capital funds is that anywhere between 60 and 80 percent of quarterly observations are zeros, i.e., most quarters have no cash flows.² The presence of large number of zeros in the analysis data limits the explanatory power of econometric models. With these limitations in place, we build upon ideas from other studies such as modeling contributions as a function of the age and uncalled capital of a fund. We model distributions as a function of the age, uncalled capital, valuation, and cumulative distributions of a fund. We calibrate our models using splines-based linear regression and fund-level cash-flow data. We report the performance of our models in three prominent subclasses of funds, namely buyout, real estate, and venture capital in the US private capital market.

The remainder of this paper is organized as follows. Section 2 discusses literature related to cash-flow modeling of private capital funds. Section 3 provides a general overview of our methodology; sections 4 and 5 then discuss modeling and empirical comparison of models for contribution and distribution respectively. Section 6 compares the performance of our models with Takahashi and Alexander's model (calibrated to Burgiss data) for both contributions and distributions. Section 7 examines important econometric challenges and lessons in private capital cash-flows modeling and prediction. The paper concludes with section 8. The appendices discuss our dataset, notation, and mathematical details of the models.

2 Literature Review

Cash-flow modeling for private capital is largely unaddressed by the existing literature; a few exceptions are Buchner et al. (2010), de Malherbe (2004, 2005), Robinson and Sensoy (2016), and Takahashi and Alexander (2002). Historically, investing in private capital was governed by simple rules of thumb which were rendered ineffective as allocations to alternatives grew over time. Recognizing this, Takahashi and Alexander (2002) proposed a framework that modeled contributions and distributions as *rates* using uncalled capital and valuation as the respective denominators. However, their model is deterministic in nature. Buchner et al. (2010) and de Malherbe (2004, 2005) introduced stochastic models of the cash-flow dynamics for drawdowns and distribution. Most recently Robinson and Sensoy (2016) noted that most cash-flow variation at a point in time is diversifiable — either idiosyncratic to a given fund or explained by the funds age. They also observed

¹For LPs, not servicing capital calls in a timely fashion may result in the loss of reputation or relationships with General Partners (GPs) and perhaps the eventual default on the commitment.

²In the econometric literature this is frequently referred to as “zero-inflation.”

that private capital cash flows are procyclical in their co-movement with the public markets as the conditions in the public market affect exit opportunities for private capital investments. Another body of literature relevant to this paper is related to finite mixture models that frequently generate zero-valued observations for modeling zero-inflated data. These models have been successfully applied to count data (see Dalrymple et al. (2003) and Zuur et al. (2009)), however, their applications to continuous data is limited. Even for count data the applications of zero-inflation models are limited to only explaining the underlying data-generating process. The explanatory power of zero-inflation models does not necessarily extend to their predictive power—nevertheless, it is common to expect some similarity between the two. Recent literature (Shmueli 2010) has begun to recognize the distinction between the in-sample explanatory power and out-of-sample predictive power of econometric models.

3 Methodology

In this section we give a general introduction to the methodology used. More complete descriptions can be found in the sections devoted to modeling contributions (section 4) and distributions (section 5) as well as in appendices C and D that contain formal descriptions of the models. See appendix B for details of the data used to estimate the models.

The goal of this paper is to predict the expected (or average) cash flows a private capital investor will experience over a future horizon (such as a quarter, or a year). Since contributions are liabilities we model them separately from distributions rather than combining them into net cash flows. However, lurking beneath the seeming simplicity of this goal are a number of complications, some of which we will be returning to in future work (see section 7).

An initial question is whether to predict cumulative variables³ (such as cumulative contributions) or per-period variables (such as contributions in a given quarter). We consider both, but focus on the latter since we find that they lead to models whose performance is superior.

A second question is how to measure the performance of the model. We will use root-mean-square error (RMSE) since other measures result in undesirable model choices (see section 7 for further details). Cash-flow data is very noisy; consequently we also seek to minimize the effect of this noise to avoid drowning out model differences. We will achieve this by also looking at performance over periods longer than a quarter, namely one year. Finally, we evaluate these measures out-of-sample, by backtesting all our models.

A third question is what form of per-period cash flows to try to predict. For example one could predict the cash flows divided by fund size (termed contribution and distribution *fractions* in this document); alternatively one could predict the cash flow as a fraction of a number representing the size of its source (uncalled capital for contributions and fund valuation for distributions). We find that these two modeling choices can be subsumed into our models via the addition of suitable interaction terms.

Finally, note that the relationships we expect to find are non-linear. For example, fund age is clearly an important variable when explaining cash flows. Equally clear is that this relationship is not linear. For contributions it will peak early on, while for distributions it will reach a peak somewhere near the middle of the fund's life. We will repeatedly use the same flexible tool to incorporate this non-linearity into our models: splines. Briefly, the idea is that the space of splines can be spanned by some basis (basis splines). The regression can be carried out against the raw variable (such as age) composed with these basis splines. See appendix A for some examples of basis splines.

³See appendix F for more details.

4 Contributions

4.1 Empirical Properties of the Data

Clearly there are significant statistical regularities in contributions. The pattern is evident if one looks at cumulative contributions, as shown in figure 1. This figure also shows that while the median cumulative contribution follows a simple pattern, the dispersion around the median is high.

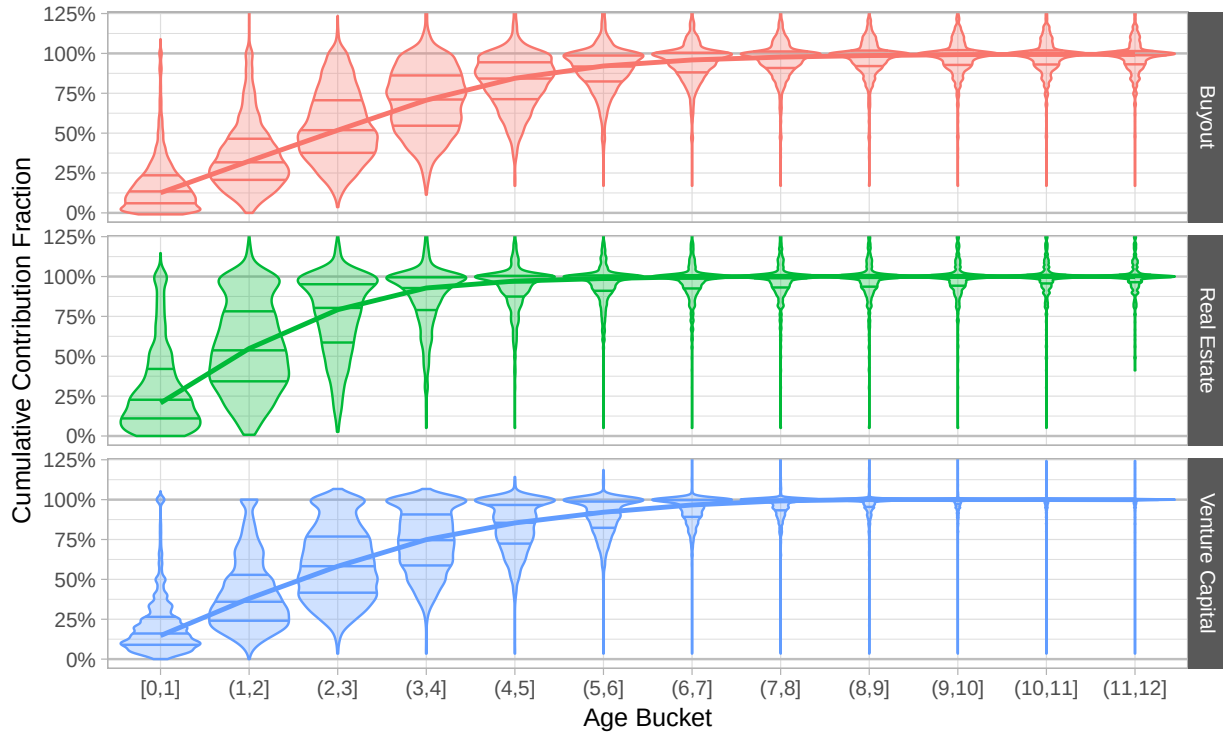


Figure 1: Violin plots of cumulative contribution fractions using fund-level data, the solid lines are smooth curves fitting the 50th percentiles across age groups

In addition to the cumulative contribution, we can also focus directly on per-period contributions. Figure 2 illustrates how contributions vary by fund age and also by amount of uncalled capital. In addition to a clear age-dependence (they decline with fund age) there is also a non-trivial dependence on uncalled capital; for example for 4-year old buyout funds the lower quartile (by uncalled capital) funds call at a third of the rate as there in the top quartile.

Finally, we examine the dependence on vintage; figure 3 plots how contribution fractions evolve with age by vintage. The vintages are grouped by half-decade in order to gather enough data within each group to generate a number that is not overly noisy. Financial crises (such as the global financial crisis (GFC)) have a significant effect on contributions (as a function of age); this is especially noticeable for real estate. Outside of crises the pattern of contributions seems reasonably stable. Note, too, that vintage, as a factor, is less easily incorporated into a predictive model since in general the information we would need to extract from a vintage is not available at the time of the prediction.

The probability density of contribution fractions is very far from normal, as can be seen in figures 4 and 5. In fact it seems to be an exponential distribution (or perhaps gamma distribution) that has been zero-inflated (see table 1). For venture capital there is inflation at multiples of 5% (see the bottom panel of figure 5).

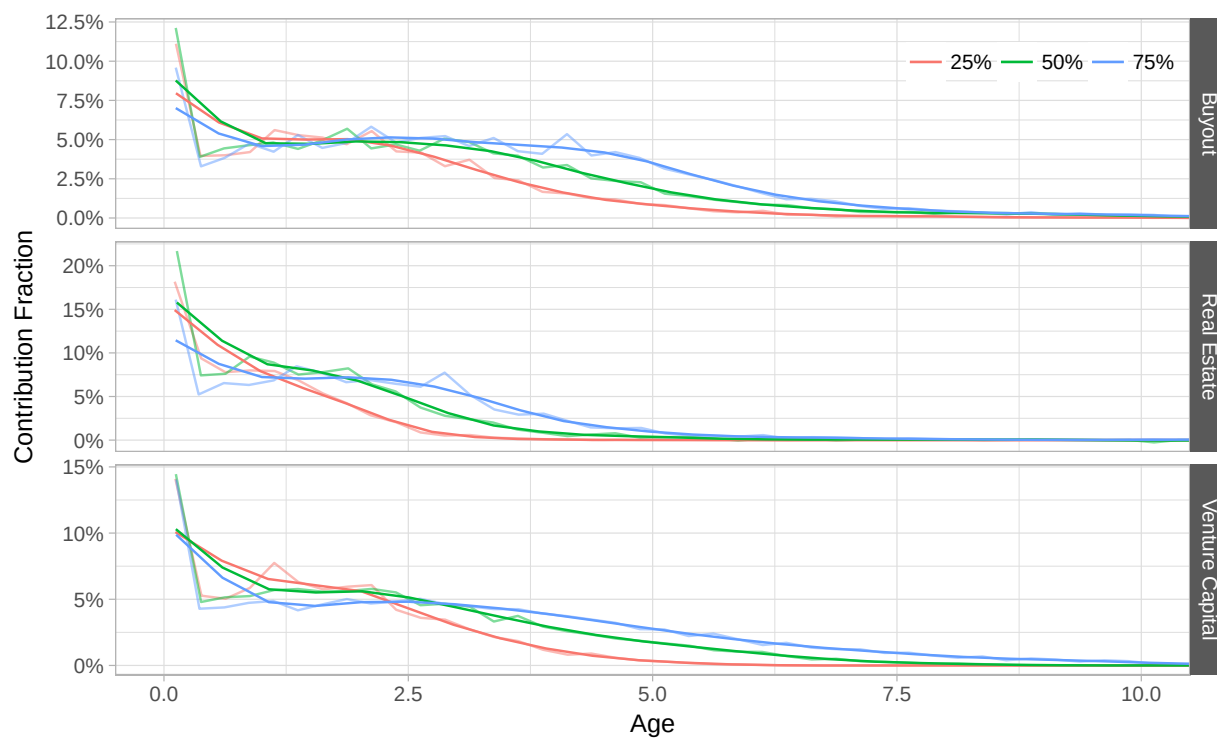


Figure 2: Empirical contribution fractions (equally-weighted) against age for funds grouped by uncalled capital quartile

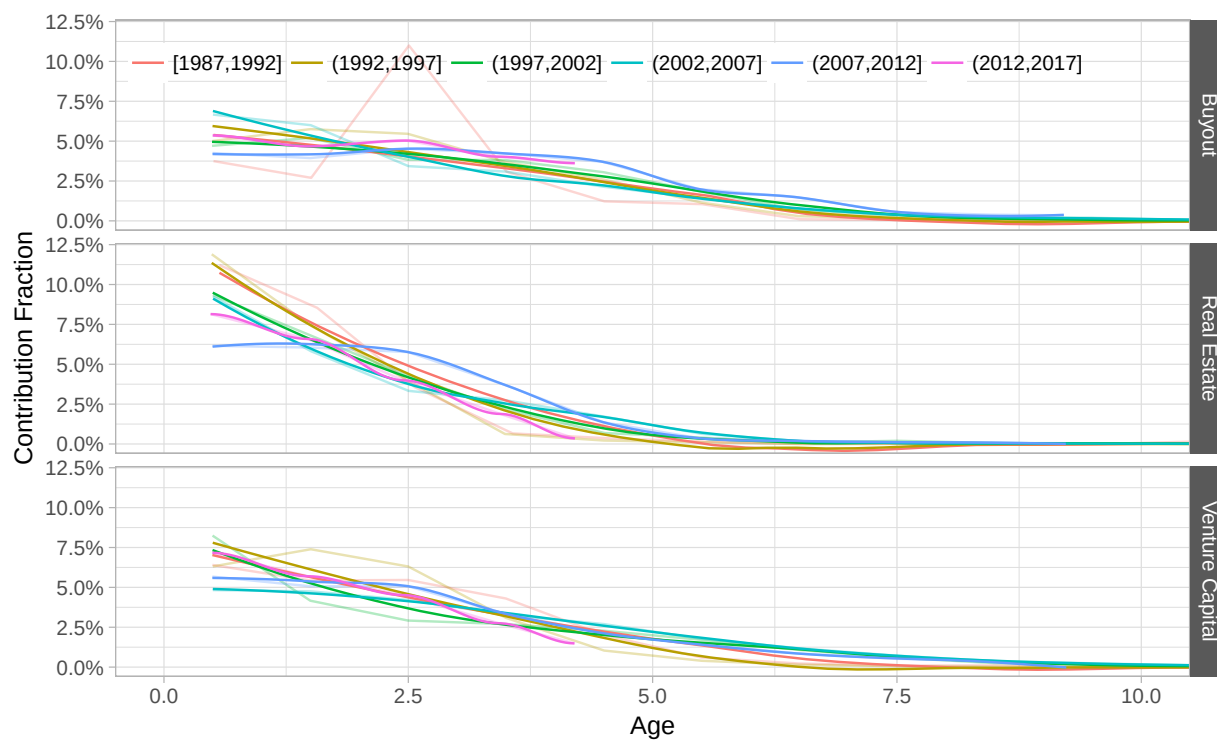


Figure 3: Empirical contribution fractions by vintage group

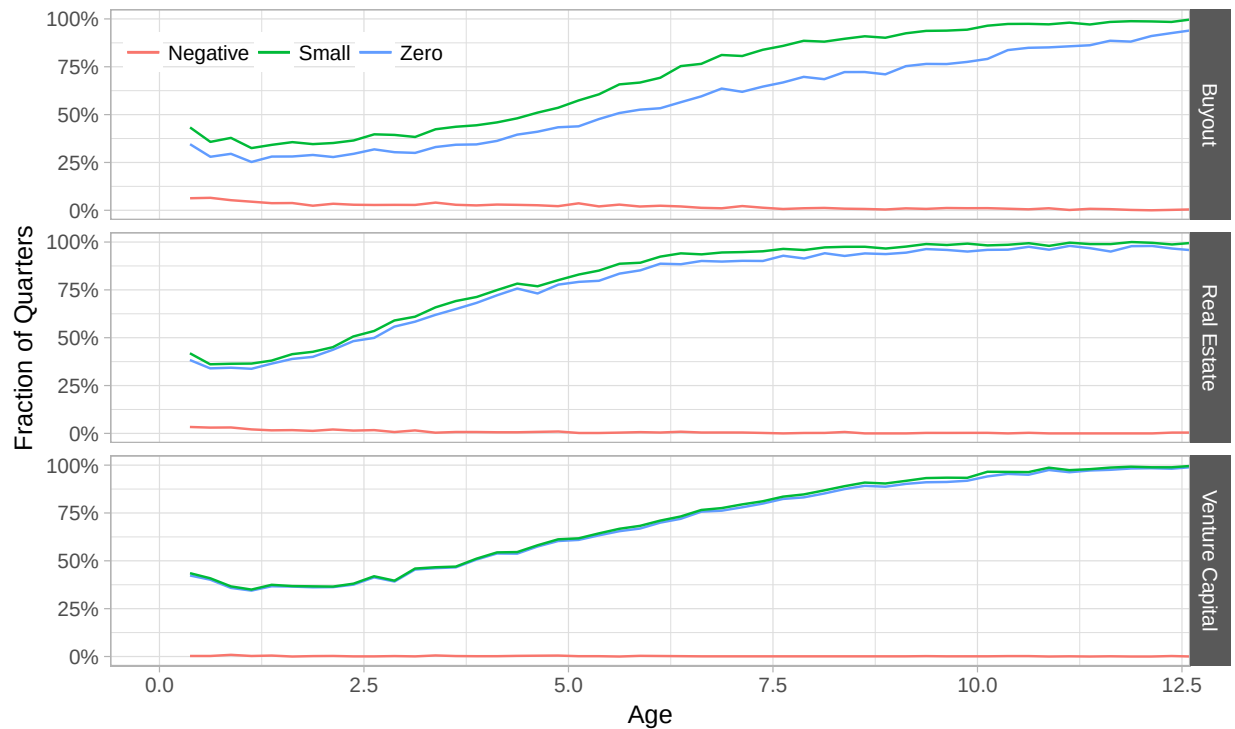


Figure 4: Fraction of calendar quarters, as a function of age, with contributions that are zero (or small or negative)

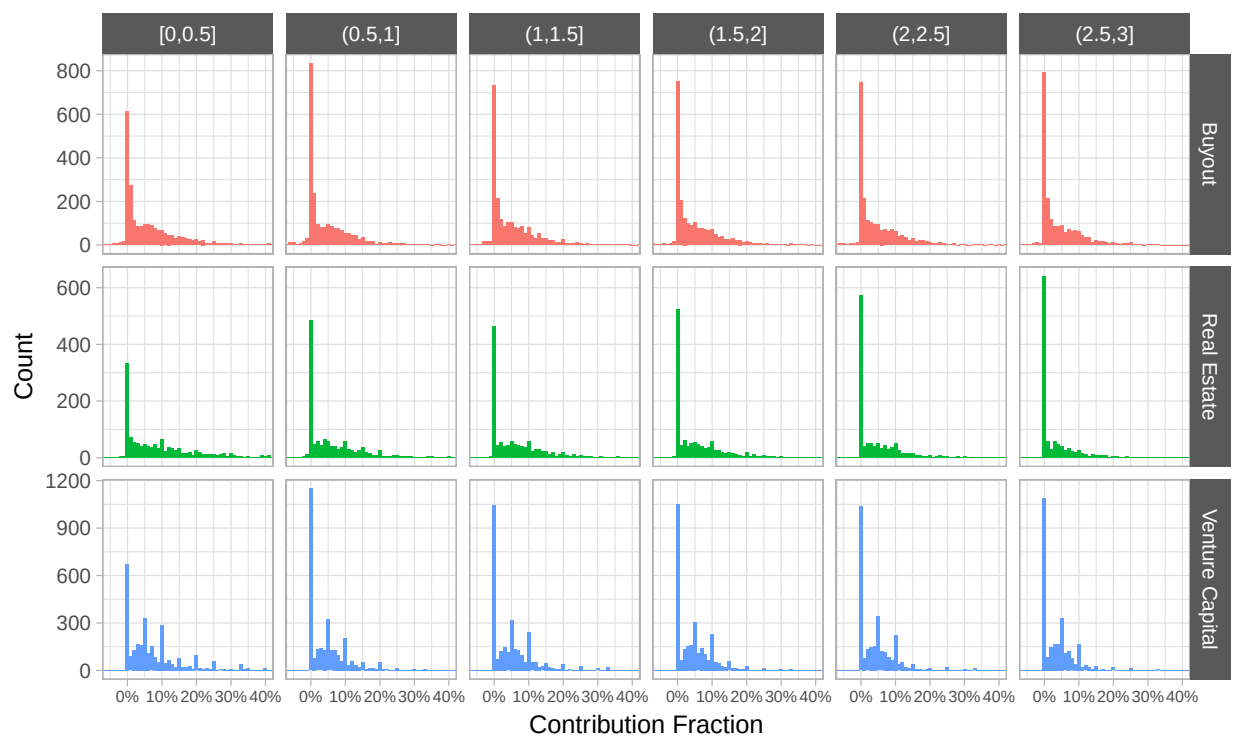


Figure 5: Histogram of contribution fractions grouped by subclass and fund age range

Table 1: Percentage of zeros in cash flow data across three subclasses

Subclass	Contribution	Distribution
Buyout	61%	67%
Real Estate	77%	60%
Venture Capital	76%	82%

4.2 Contribution Models

We model the *contribution fraction* of each fund in a quarter, defined as

$$\text{Contribution Fraction} = \frac{\text{Contribution}}{\text{Fund size}}.$$

Using the contribution fraction, we estimate the following models, each separately on each subclass.⁴

Zero A model that always predicts zero.

Constant A model that predicts the same contribution fraction, regardless of age or uncalled capital.

JustAge A model that predicts contribution fraction as a smoothly varying function of age.

NoInteractions A model that predicts contribution fraction in terms of a smoothly varying function of age and uncalled capital, but not their interactions.

WithInteractions A model that predicts contribution fraction in terms of a smoothly varying function of age, uncalled capital, and their interactions.

4.3 Predictions

Next we illustrate the predictions of the above models. These predictions are a function of what data is used to estimate the model. Later in this section we will backtest the models, namely ask the models to predict the next period using only data from the past. However here we pick a representative fit, namely the most recent one (i.e., we estimate the model using all available data).

Figure 6 illustrates the predictions of the simpler of the above models, namely those that do not employ uncalled capital as a variable. In order to plot the more complex models, we need to pick a level of uncalled capital corresponding to each age. In order to see the effect of this variable, we choose three such levels by estimating the 25th, 50th, and 75th percentiles of uncalled capital corresponding to each age;⁵ these data are displayed in figure 7. With these uncalled-capital quartiles as a function of age we can now plot the model predictions for the models with a dependency on uncalled capital. These predictions are displayed in figure 8. Comparing these results with figure 2 we see general agreement. The dependence on age is as expected, in addition the relationship with uncalled capital mirrors the data. For example, somewhere between 1 and 2 years the inevitable impact of uncalled capital becomes evident (namely that low uncalled capital implies lower contribution rates). Perhaps more interesting is that before that point the relationship is unclear, or even inverted (as appears to be the case for venture capital funds younger than 2 years⁶). Finally note the spike at age zero in all age-dependent models; this is due to the data equating the inception date of a fund with its first cash flow (presumably a contribution) thus guaranteeing a cash flow at age zero.

⁴See appendix C for the mathematical details of these models.

⁵The estimation of age-dependent quantiles is done via a spline quantile regression.

⁶Thus these funds exhibit a form of momentum: funds that call faster than average continue to call faster than average, until they run out of funds to call, of course.

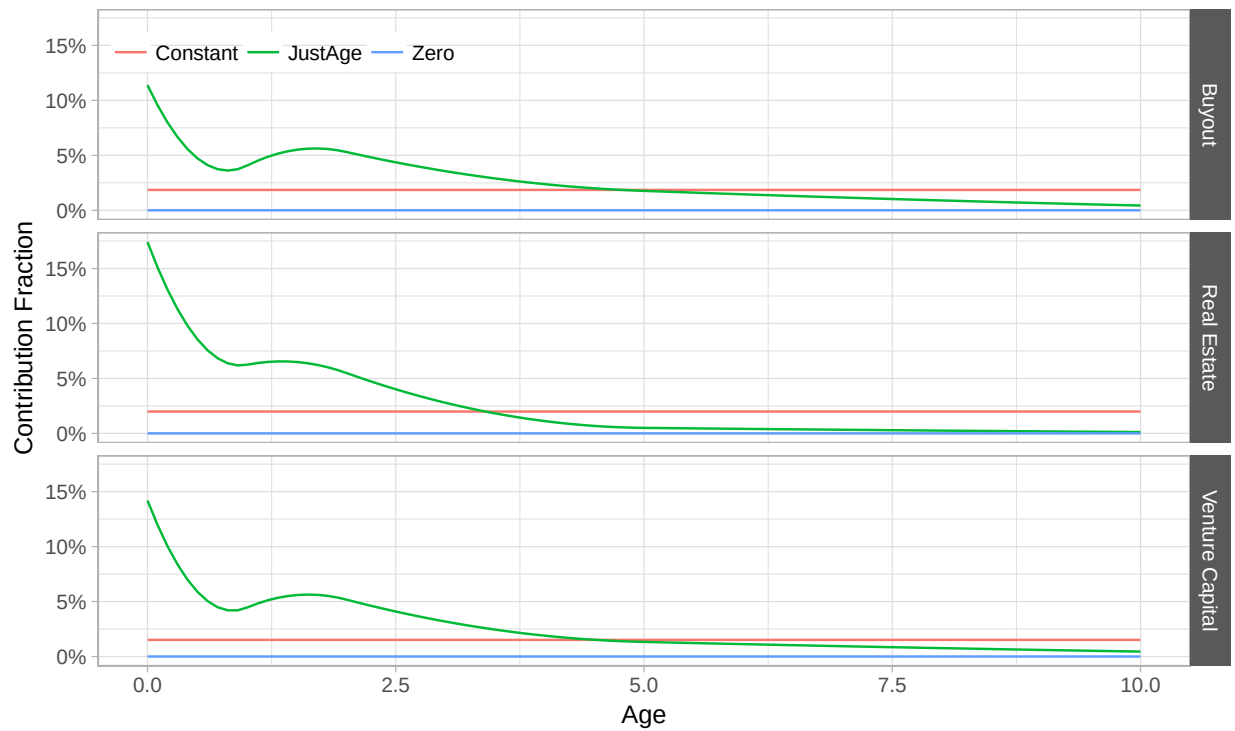


Figure 6: Predictions of models with no dependence on uncalled capital against age

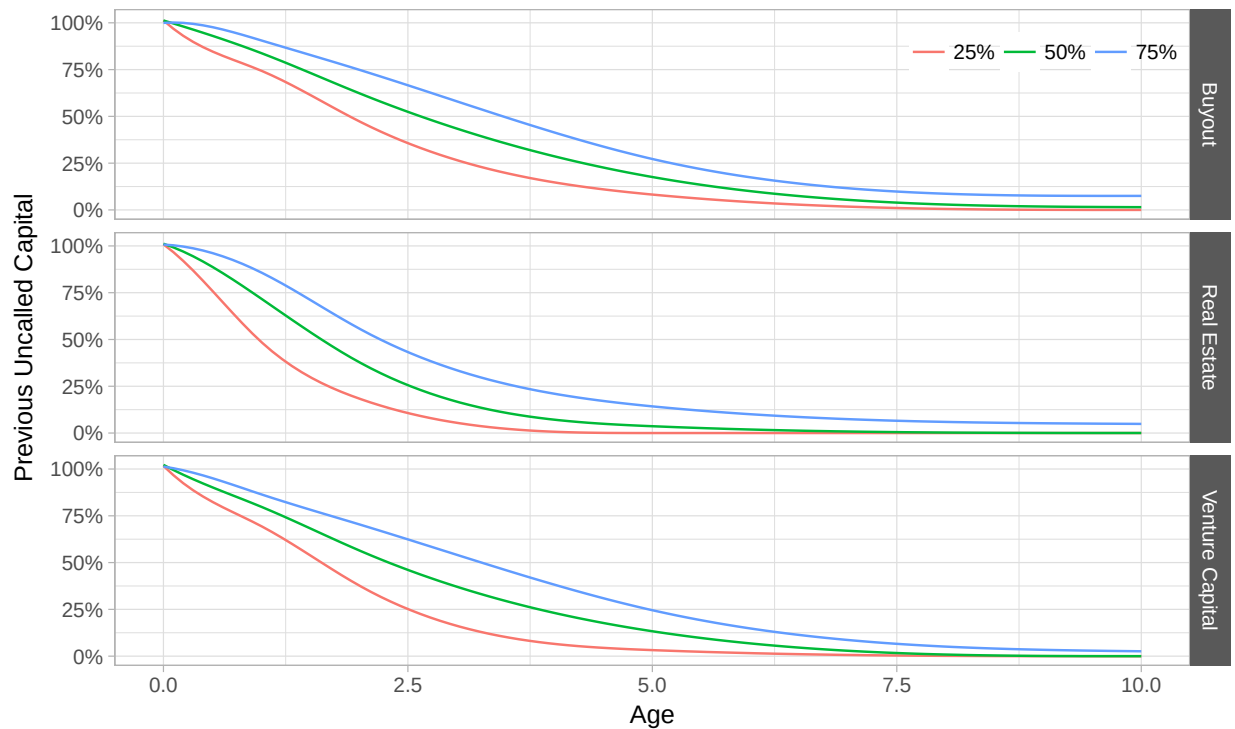


Figure 7: Empirical contribution fractions (fund-size-weighted) against age for funds grouped by uncalled capital quartile

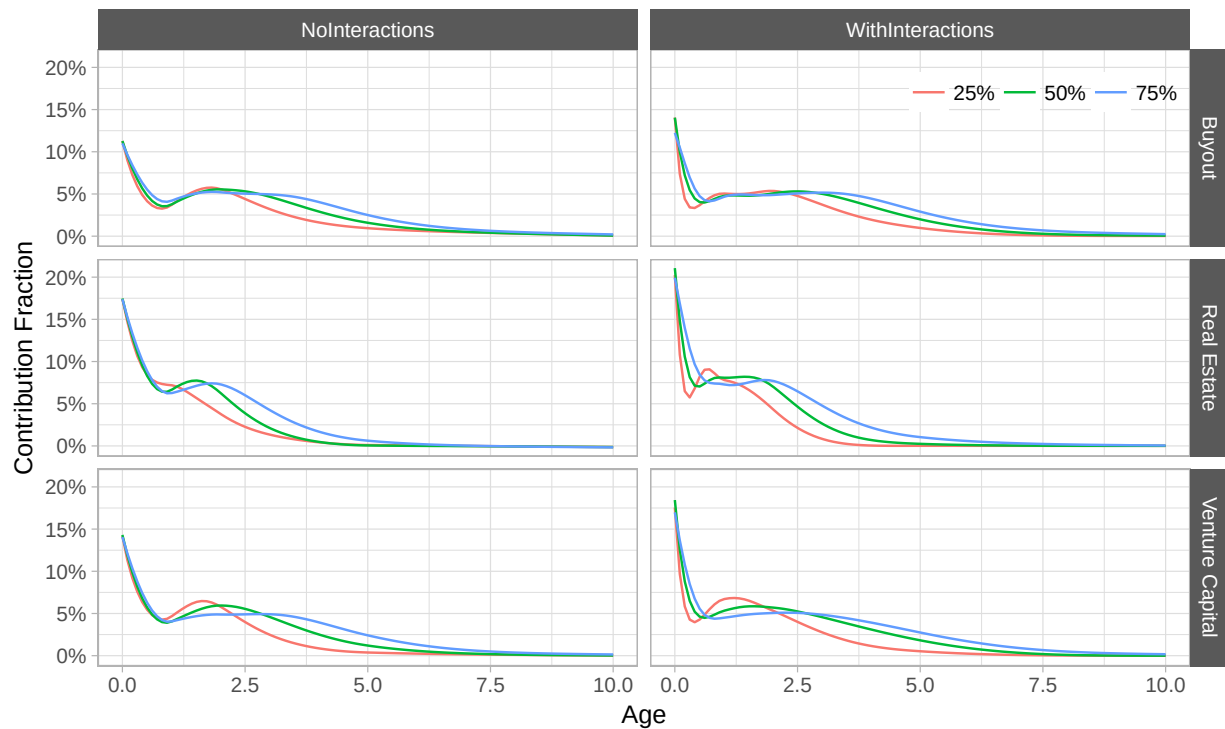


Figure 8: Predictions of models with dependence on uncalled capital against age; the uncalled capital is set at the percentile value indicated by the line color

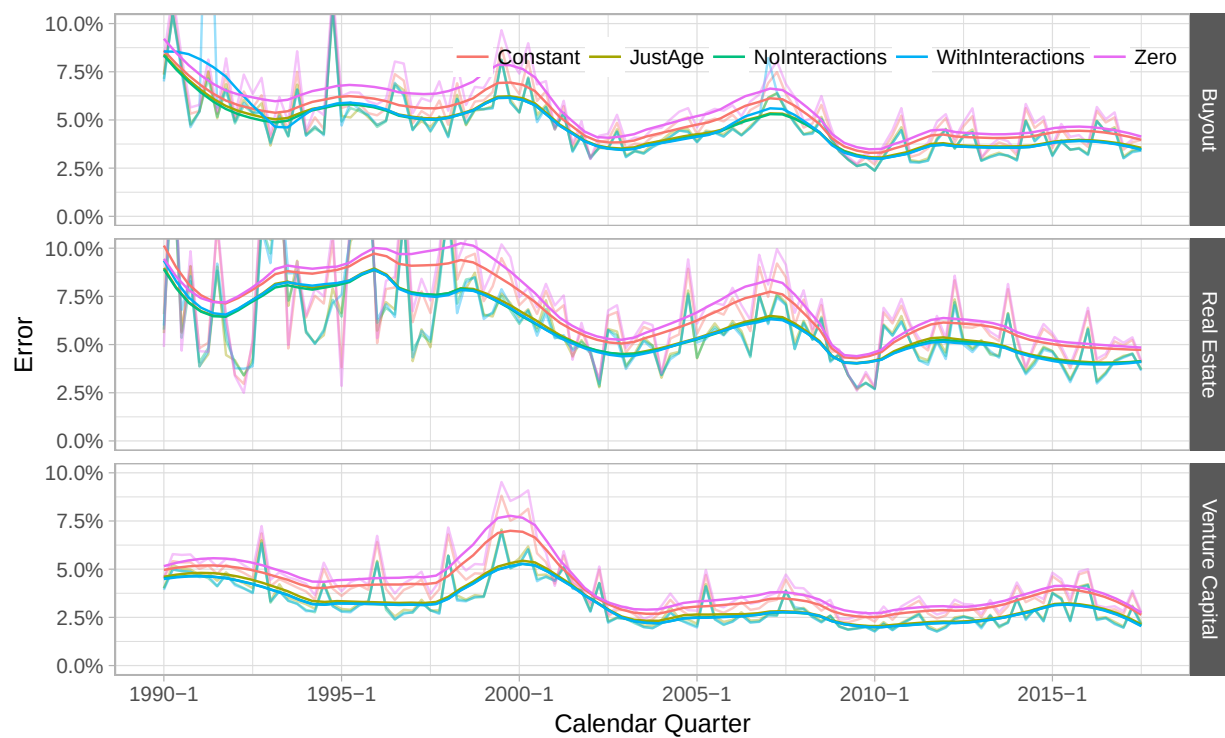


Figure 9: Backtesting performance of several models for predicting quarterly contributions

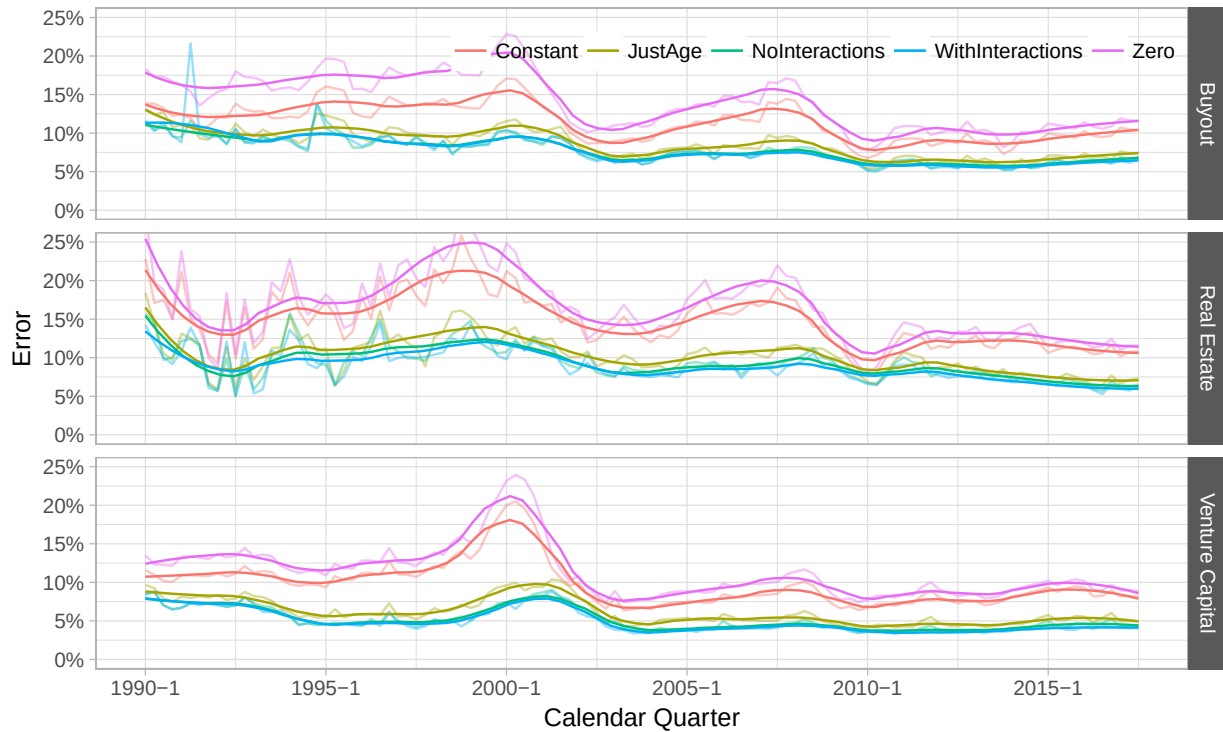


Figure 10: Backtesting performance of several models for predicting annual contributions

4.4 Backtesting Results

While in-sample measures of model performance are useful to understand how good a model is, ultimately what matters is out-of-sample performance. For time-series data this is usually carried out by *backtesting* the model. Here the out-of-sample data is from the future. For example when predicting the cash flows in (say) 2016 Q4 the model will only have access to data up to the end of 2016 Q3. Thus the backtesting procedure consists of repeatedly estimating each model using data strictly before each quarter and then predicting the cash flows for that quarter. Since data from far in the past is, presumably, less relevant than more recent data we also use exponentially decaying weights in the regression with a characteristic time of 1 year.⁷ We measure performance of the model by the root-mean-square error.

Figures 9 and 10 display the results of carrying out this backtesting for quarters from 1990 Q1 to 2017 Q2. Figure 9 uses quarterly data (i.e., the model observes quarterly data and predicts the next quarter); figure 10 uses annual data (i.e., the model is fit to past annual data and predicts the cash flows during the next year).

Starting with the quarterly backtest (figure 9) we see the expected ordering of the models (namely, **Constant** is worst, and either **NoInteractions** or **WithInteractions** is best). Note that while this ordering is guaranteed in-sample, it is not guaranteed out-of-sample (which is what backtesting constitutes). Also the most complex model (**WithInteractions**) is essentially over-fitting the data. Nevertheless it out-performs simpler models (such as **NoInteractions**) by a small margin. However evidence of this overfitting can occasionally be seen; for example if that model is estimated using a shorter half-life (which is similar to reducing the amount of data during estimation) then poor predictions become

⁷The characteristic time is how long it takes for the weights to decay to $1/e$ of their value for the most recent observations. A characteristic time of 1 year is equivalent to a half-life of about 0.6 years.

Table 2: Backtesting contribution models on annual data
The models were calibrated on, and predicted, annual cash flows. Reported statistics are out-of-sample R^2 since 1990; the numbers in parentheses are since 2010.

Subclass	Out-of-sample R^2 / %			
	Constant	JustAge	NoInteractions	WithInteractions
Buyout	30 (23)	62 (60)	69 (66)	70 (68)
Real Estate	19 (15)	61 (58)	68 (64)	71 (68)
Venture Capital	24 (19)	71 (71)	79 (79)	82 (82)

possible, especially early in the backtesting period.

Moving to annual backtesting (figure 10) we see the differences between models become more evident. In particular the superiority of models that in addition to age also use uncalled capital as a predictor becomes clear. Also of note is that even during the GFC the models out-performed a prediction of zero. This was one of the motivations for using weighted regressions. For example if one performs a regression with no weights (or a long half-life) then during the GFC contributions plummet and a prediction of zero is better than the model predictions. In a sense this is not surprising since the model does not use any variables that indicate in which “regime” it finds itself.

Finally note that while both the signal (as indicated by the Zero model) and the model errors are larger in the annual backtest, the relative size of the errors as function of the signal (this can be thought of as an out-of-sample version of $1 - R^2$) declines. The performance of the annual models is compared in table 2, using an out-of-sample R^2 measure (see appendix E for more details on the computation of out-of-sample R^2). As can be seen, the models that incorporate information from age and uncalled capital significantly improve their out-of-sample R^2 . Note also that the error time-series are more stable in the annual backtest than in the quarterly one. The reason for these effects is that we are performing temporal aggregation. As discussed later, our models are trying to predict the general rate of cash flows, as opposed to the precise timing of each individual cash flow. However, we measure performance relative to actual cash flows. One can consider the precise timing of individual cash flows as analogous to idiosyncratic noise, and this is diversified away by aggregation, in this case, temporally.

See later in this document (section 7) for a discussion of aggregation in general.

5 Distributions

Modeling distributions of private capital funds can be aided by separating out the effects of performance on them. For example, a drop in distributions could either be due to delayed exits or reduced asset prices. This makes distributions more complex compared to contributions that are directly scaled with the respective fund sizes (or total contributions) known *a priori*. Total distributions, on the other hand, would not be known until the fund is liquidated. To get around this difficulty we *estimate* the total distributions of each fund in a quarter as

$$\text{Total Distributions Estimate} = \text{Cumulative Distributions} + \text{Valuation} + \text{Uncalled Capital}.$$

The estimate of total distributions consists of three components: the sum of distributions that are already paid-out, the valuation that will potentially be paid-out, and the uncalled capital. Early in the life of a fund this estimate will be very similar to the fund size⁸ and later on it will be very

⁸This explains why the uncalled capital is included in the estimate.

similar to the actual total distributions (that would eventually be known). Similar to fund size for scaling contributions, an estimate of total distributions is a suitable choice for scaling per-period distributions, we define *distribution ratio* as

$$\text{Distribution Ratio} = \frac{\text{Distribution}}{\text{Total Distributions Estimate}}.$$

Distribution ratios are extremely similar to contribution fractions as they are completely abstracted from the size and performance of the fund; also their cumulative value approaches one as the underlying fund gets old (see figure 11). The concept of cumulative distribution ratios is also mentioned in Mathonet and Meyer (2008) as “repayment age.”

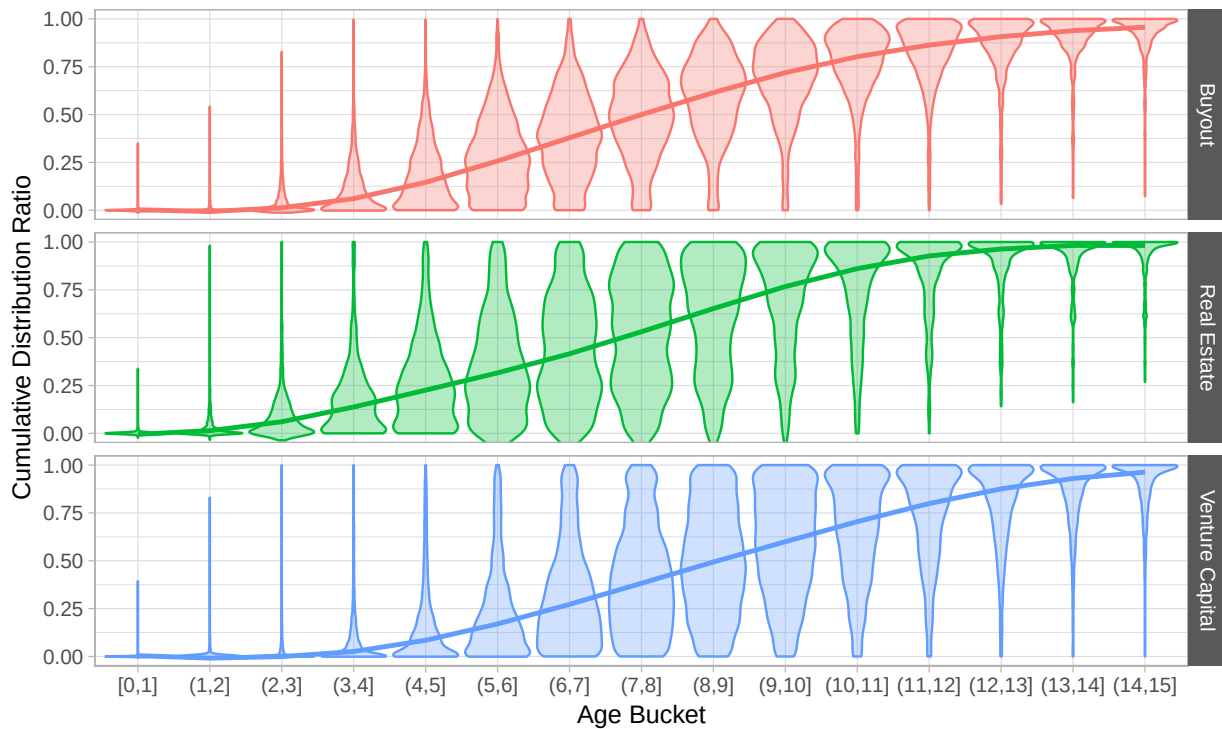


Figure 11: Violin plots of cumulative distribution ratios using fund-level data, the solid lines are smooth curves fitting the 50th percentile across age groups.

Figure 11 reveals that the distribution ratio is clearly a function of age but in the cross-section of funds there is large deviation from the median in each age bucket. To explain the variation from the median, we consider the *valuation ratio* as an independent explanatory variable,⁹ defined as

$$\text{Valuation Ratio} = \frac{\text{Valuation}}{\text{Total Distribution Estimate}}.$$

To examine the explanatory power of the valuation ratio we plot the mean distribution ratio as a function of age in the quartiles of valuation ratio. Figure 12 suggests that along with age, valuation

⁹Uncalled capital and the sum of distributions that are already paid-out can also potentially serve as explanatory variables, but in our experience they do not improve the model performance significantly, perhaps because they are very highly correlated with age, which is already included in our models. These factors may provide independent sources of information if included as investment speed or performance respectively but this is not explored in this paper.

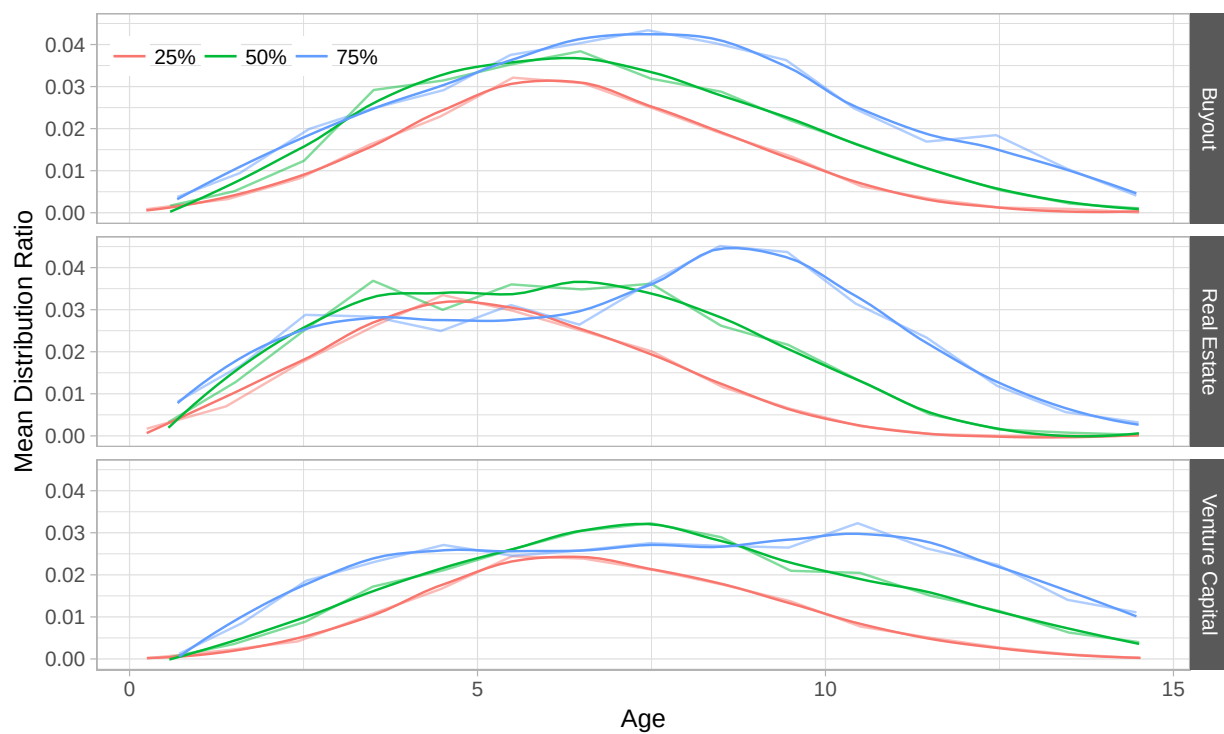


Figure 12: Empirical distribution ratio (equally-weighted) against age group by valuation ratio quartiles

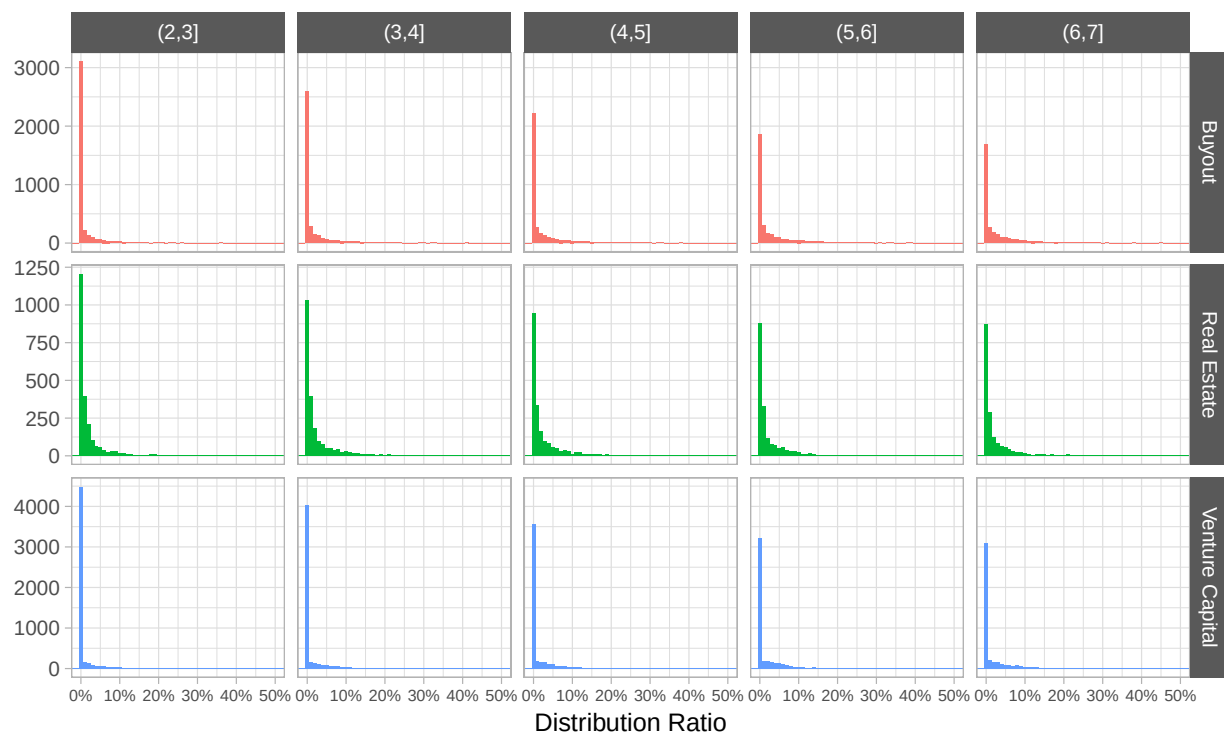


Figure 13: Histograms of distribution ratio in various age buckets

ratio also explains the mean distribution ratio. This will become clearer when we compare models with and without valuation ratio as an independent variable; but before we explore modeling, it is important to discuss another important aspect of distributions data, namely that it is also zero-inflated. This is shown in figure 13 and also summarized in table 1. For a massively zero-inflated data (such as venture capital, see table 1) zero will be the correct prediction most of the time. A model that always predicts zero is included in our models as it can serve as a benchmark for assessing the performance of other models.

5.1 Distribution Models

Using distribution ratios, we estimate the following models, each separately on each subclass.¹⁰ Our models are designed to predict only the timing of distributions; any effect of the performance of fund on distributions is not captured by these models.¹¹

Zero A model that always predicts zero.

Constant A model that predicts the same distribution ratios, regardless of age or valuation ratio.

JustAge A model that predicts distribution ratios as a smoothly varying function of age.

NoInteractions A model that predicts distribution ratios in terms of a smoothly varying function of age and valuation ratio, but not their interactions.

WithInteractions A model that predicts distribution ratios in terms of a smoothly varying function of age, valuation ratio, and their interactions.

5.2 In-sample Predictions

Figure 14 shows the model predictions for the distribution ratio against age for the first three models which do not include valuation ratio as regressor. A quick in-sample comparison of **JustAge** model with the empirical profile in figure 12 confirms that age has significant predictive power, and when combined with valuation ratio the **WithInteraction** model produces plots nearly identical to the empirical profile (compare figure 12 and the right panel of figure 15).

5.3 Backtesting Results

For an out-of-sample analysis of the distribution models we run backtests in the same framework as defined in section 4.4 and use the same characteristic time of 1 year to generate exponentially decaying weights. Again we compare the out-of-sample performance of models using root-mean-square error in *distribution fraction* space, which we defined as

$$\text{Distribution Fraction} = \frac{\text{Distribution}}{\text{Fund Size}}.$$

The predicted distribution ratios are converted to predicted distribution fractions as follows,

$$\text{Predicted Distribution Fraction} = \text{Predicted Distribution Ratio} \times \frac{\text{Total Distributions Estimate}}{\text{Fund Size}}.$$

¹⁰See appendix D for the mathematical details of these models.

¹¹This is not to suggest that the distributions forecasts cannot be improved by incorporating information about the performance of fund. In fact forecasting distributions could be a two-step process in which separate forecasts about timing and performance are combined.

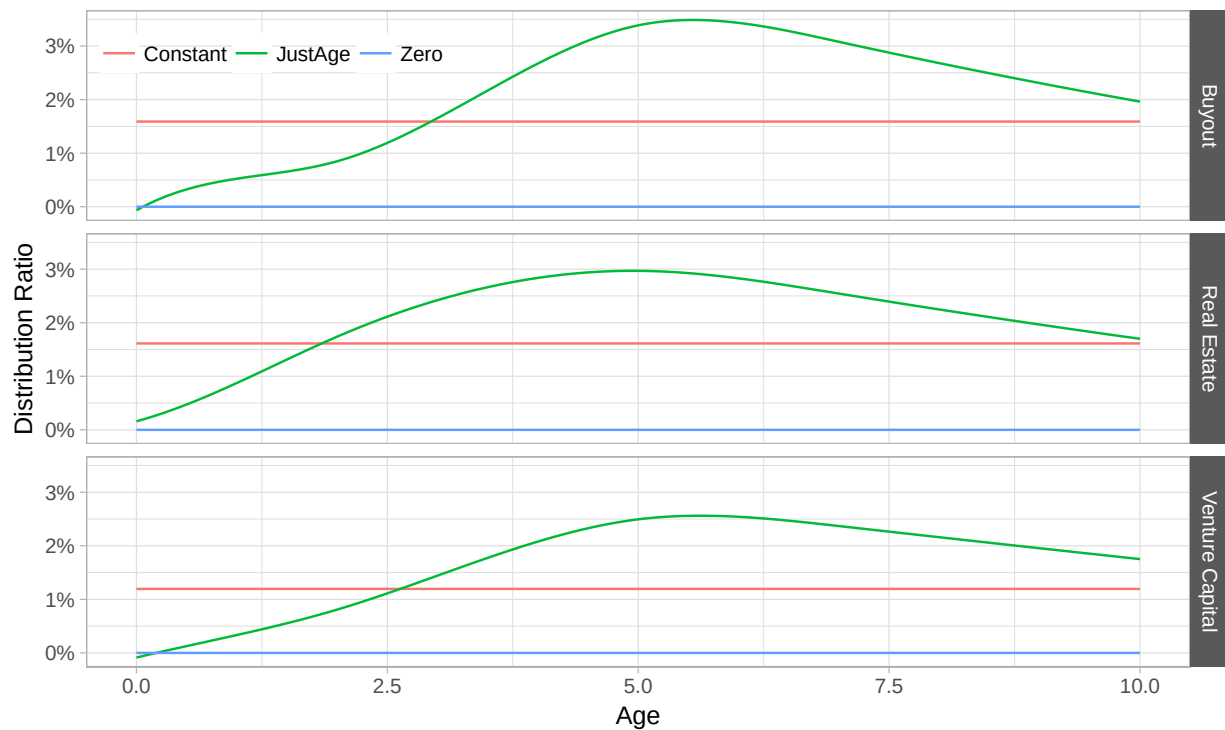


Figure 14: Model predictions (not including models that also use valuation ratio as a regressor) of distribution ratio against of age

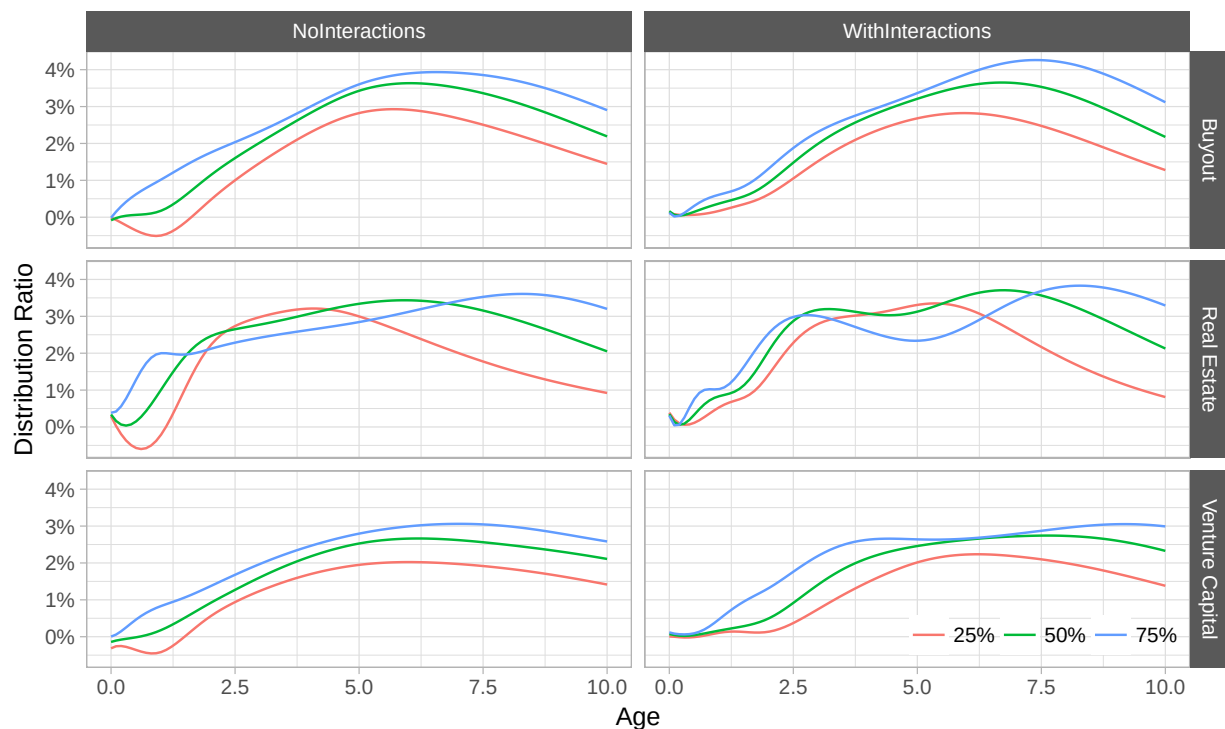


Figure 15: Models predictions (models that also use valuation ratio as a regressor) of distribution ratio against of age; the valuation ratio is set at the percentile value indicated by the line color

Table 3: Backtesting distribution models on annual data

The models were calibrated on, and predicted, annual cash flows. Reported statistics are out-of-sample R^2 since 1990; the numbers in parentheses are since 2010.

Subclass	Out-of-sample R^2 / %			
	Constant	JustAge	NoInteractions	WithInteractions
Buyout	24 (22)	45 (45)	48 (49)	46 (51)
Real Estate	29 (22)	43 (41)	45 (44)	45 (46)
Venture Capital	26 (-9 [†])	46 (25)	51 (35)	53 (38)

[†] Attributed to the fact that since 2010 Zero has been a better model than Constant.

Note that measuring the error in fraction space makes them equally-weighted, i.e., insensitive to fund sizes. The results of backtesting the distribution models are plotted in figures 16 and 17, where we compare the performance of all five models quarter-by-quarter from the first quarter of 1990 to the second quarter of 2017. Figure 16 shows backtesting results produced using quarterly distributions data and predictions; figure 17 uses rolling annual distributions data and predictions. These models only predict the timing of distributions, so spikes in the errors are inevitable as funds regularly (and unpredictably) experience large capital gains or losses. Similarly large-scale market events like the dot-com crash of 2002 and the GFC of 2008 are not predictable. When such events have a negative impact on market performance Zero is almost guaranteed to outperform more complex models.

Looking at figure 16 it seems that all models perform quite similarly, and during crises Zero can even outperform the rest. In general other models modestly improve upon Zero and with a closer look one can rank them. Constant is clearly the worst and WithInteractions the best but is nearly identical to NoInteractions. WithInteractions model is also prone to overfitting especially when data is sparse (note large spikes for buyout funds in figure 16). Similarly NoInteractions is only slightly better than JustAge. Moving to figure 17, where the models are compared with temporally aggregated data (namely annual data), the ordering of models (worst to best) does not change but the differences among the models become clearer as a results of temporal aggregation. Note also, as a result of temporal aggregation, both prediction and observed data become less spiky. The performance of these models on annual data is also summarized in table 3. The out-of-sample R^2 s are about 45–50% for the JustAge, NoInteractions, and WithInteractions models; this represents a large improvement over the Constant model. Looking at table 3 and also figures 16 and 17, one might justly wonder how little information funds valuations (as an independent variable) provide as the difference in the JustAge and NoInteractions is barely noticeable. One way to explain this is to notice that the JustAge model is already exploiting the information in fund valuations for computing distribution ratios as the dependent variable and converting the distribution-ratio forecast into distribution-fraction forecast. It is possible that beyond this fund valuations have little or no further information to add. This will become clearer in section 5.4 where we compare the JustAge model with another model in which fund valuations are explicitly used as an independent variable and not used to construct the dependent variable.

5.4 Separating the Effect of Performance

When modeling distributions, a natural question to ask is what would happen if we did not separate the effect of fund performance on distributions. To explore this question, we compare the WithInteractions model with an analogous fraction-based model, which we call WithInteractions-Fraction (see appendix D for more details). Figure 18 compares the out-of-sample performance

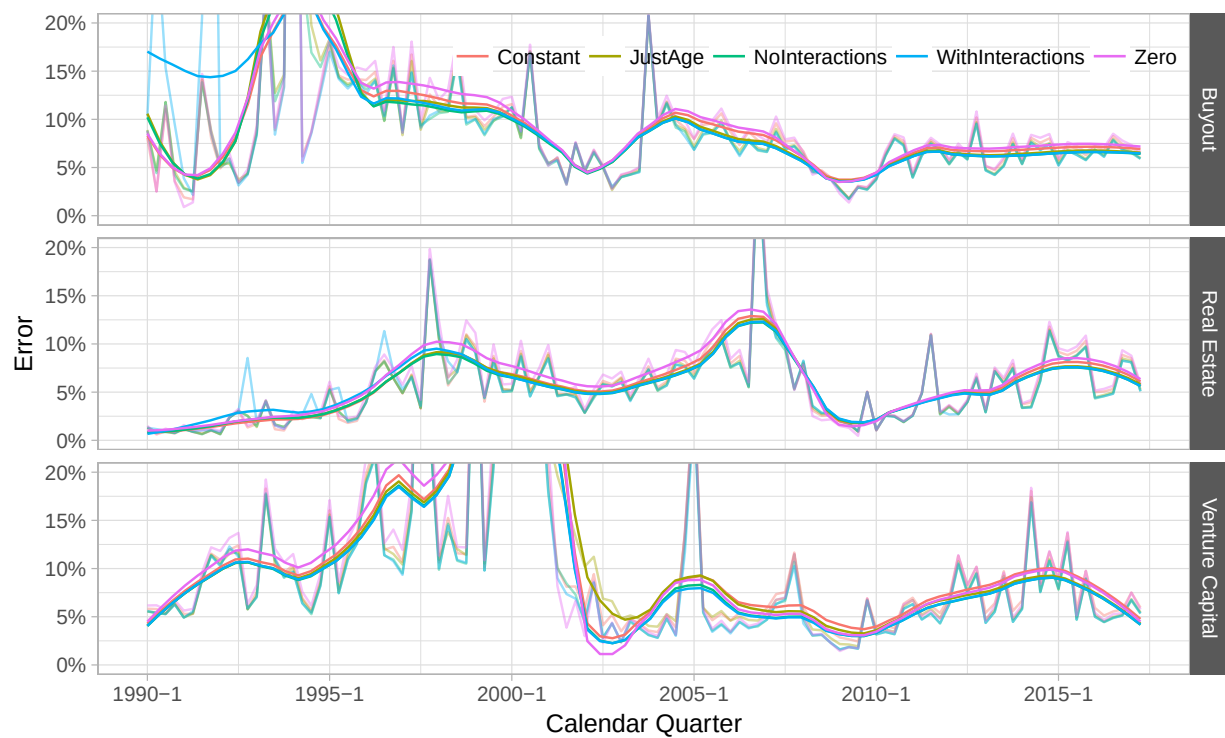


Figure 16: Performance of several models for predicting quarterly distributions

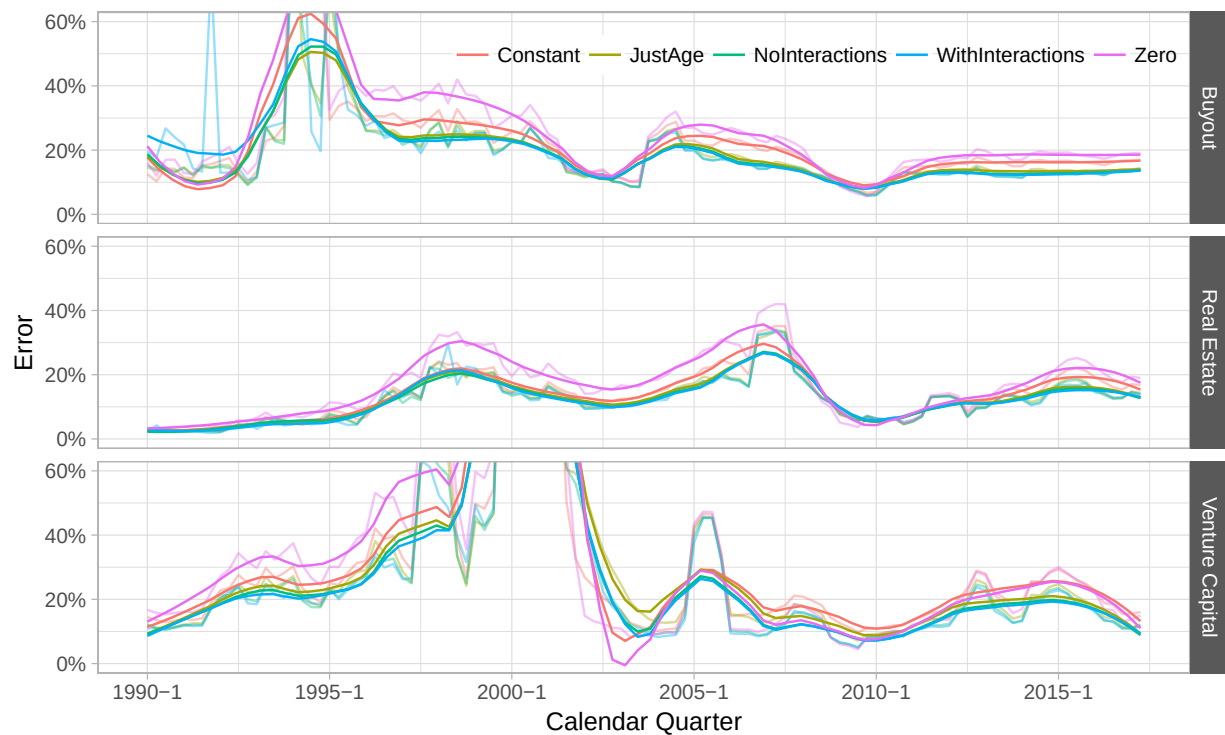


Figure 17: Performance of several models for predicting annual distributions

of the two modeling approaches for the last twelve years in all three subclasses of funds. Except during the GFC, the ratio-based model always outperformed the fraction-based model for buyout and real estate. For venture capital the two approaches are pretty competitive but in the last 5 years or so the ratio-based model has outperformed. Since the fraction-based model does not try to separate out the effect of performance on distributions it can outperform the ratio-based model during the broad market boom or bust. During the GFC of 2008 the fraction-based model marginally outperformed the ratio-based model because it could “see” the impact of negative return on fund valuations. However, since it does not separate the performance effects it is also prone to overfitting distribution fractions (dependent variable) which may have large swings (notice a spike in the case of real estate funds during 2010 in figure 18). Aside from the empirical evidence, in order to understand why the ratio-based model is a better choice, we turn to in-sample analysis and compare R^2 (see table 4) of both models in the fraction space (this requires recomputing R^2 of the ratio-based model in the fraction space). The ratio-based model achieves a greater fit for both buyout and real estate subclasses and slightly underfit for venture capital which explains why the two approaches are competitive for venture capital.

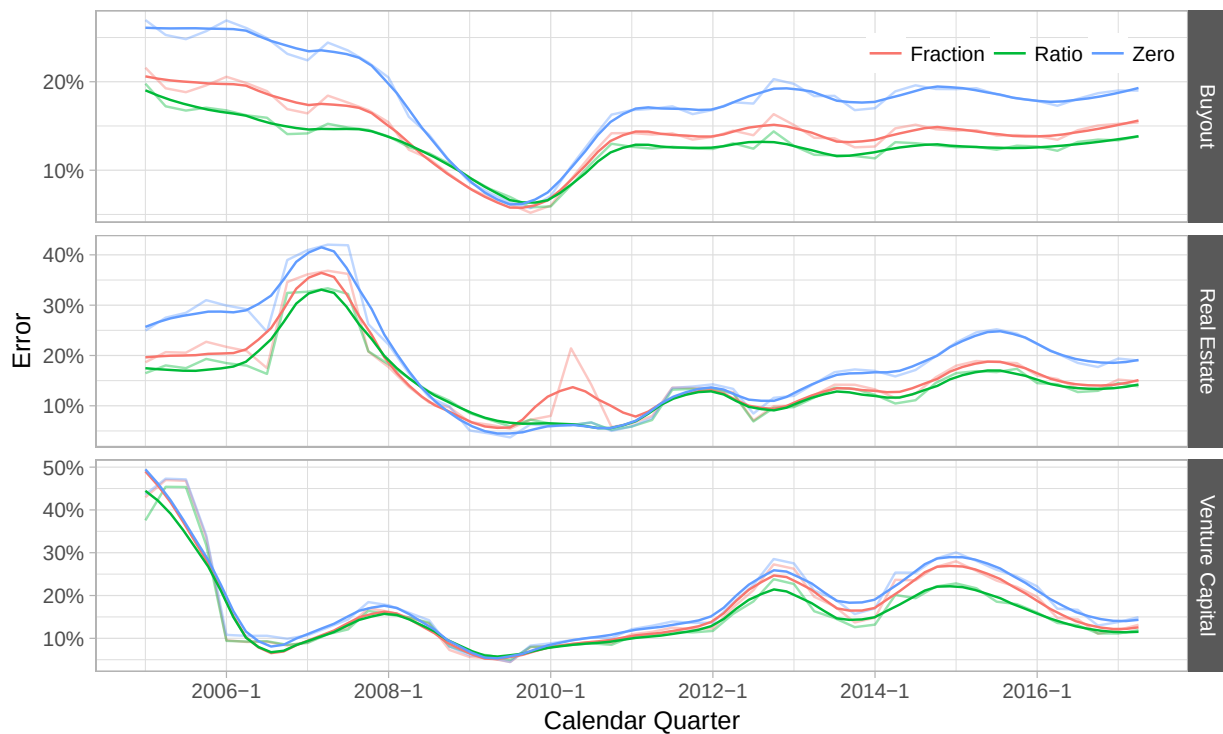


Figure 18: Backtesting-based (annual) comparison of the WithInteractions and WithInteractions-Fraction models

Finally we return to the discussion started in section 5.3 concerning the predictive power of fund valuations for forecasting distributions. We compare the JustAge model with the WithInteractions-Fraction model. Both of these models exploit the same information in fund valuations, albeit differently. Figure 19 plots the comparison of two models using a rolling annual backtest. It seems that JustAge is uniformly better than the WithInteractions-Fraction model except during a crisis. For real estate and venture capital the JustAge model is usually competitive, and often beats the WithInteractions-Fraction model. With this it is no surprise that the JustAge model is already using the information in fund valuations which is why we should not expect a large improvement over the

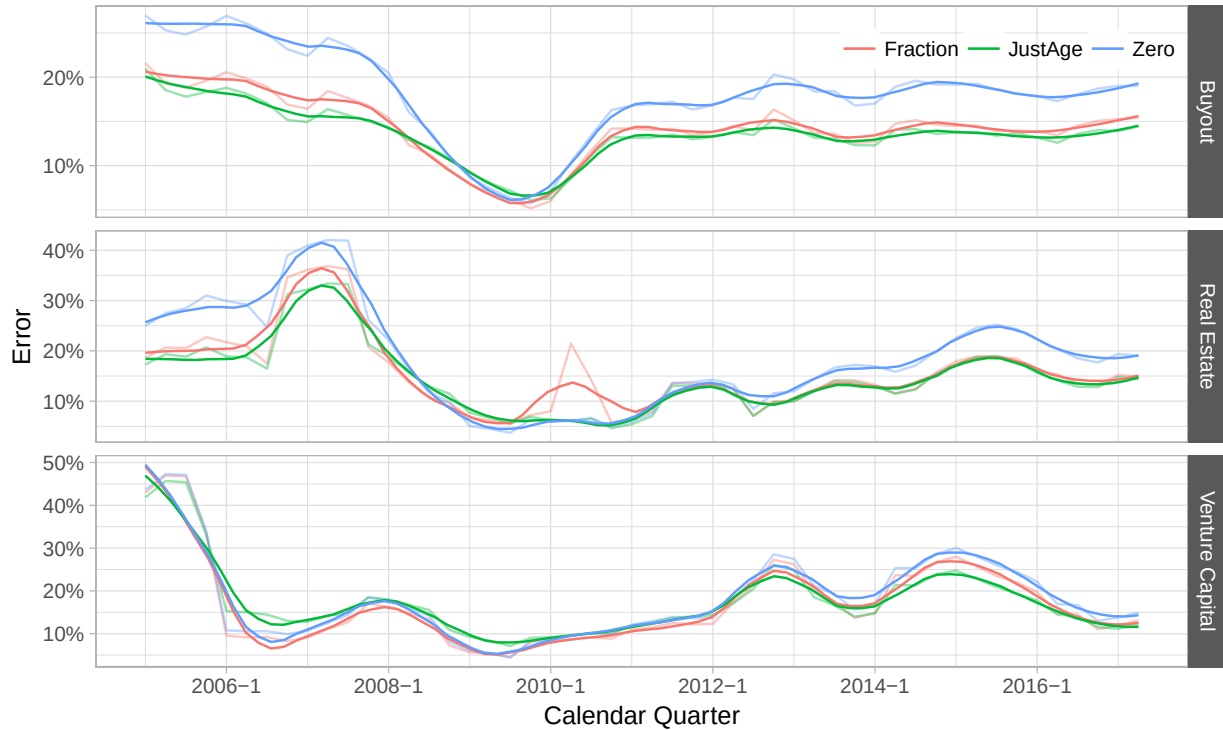


Figure 19: Backtesting-based (annual) comparison of the JustAge and WithInteractions-Fraction models

Table 4: Comparison of R^2 in fraction space

Subclass	Modeling Approach	
	Ratio-based	Fraction-based
Buyout	40%	36%
Real Estate	38%	27%
Venture Capital	25%	27%

JustAge model by the NoInteractions or WithInteractions models.

6 Comparison with Takahashi and Alexander

This section compares the models we have explored in this paper with the approach outlined in Takahashi and Alexander (2002).

The Takahashi and Alexander model (henceforth, **TA**) for contributions is arguably an instance of what we referred to above as an age-dependent, uncalled-capital-dependent, with-interactions model, albeit of a particularly simple form. The model assumes contributions are proportional to uncalled capital, and that the proportionality constant (termed the *contribution rate*) can vary with the age of the fund. The contribution rate is assumed to be a step function which is constant in the first year, constant in the second year and constant thereafter. Thus the model has three parameters (if we hold constant the ages at which the step function changes level, which we do, following the original paper). For distributions the model has even less parameters. It assumes distributions are proportional to fund value (or NAV) and that the proportionality constant (termed the *distribution*

rate) is simply the fund age as a fraction of the fund's life raised to some fixed power (called the *bow* parameter).¹²

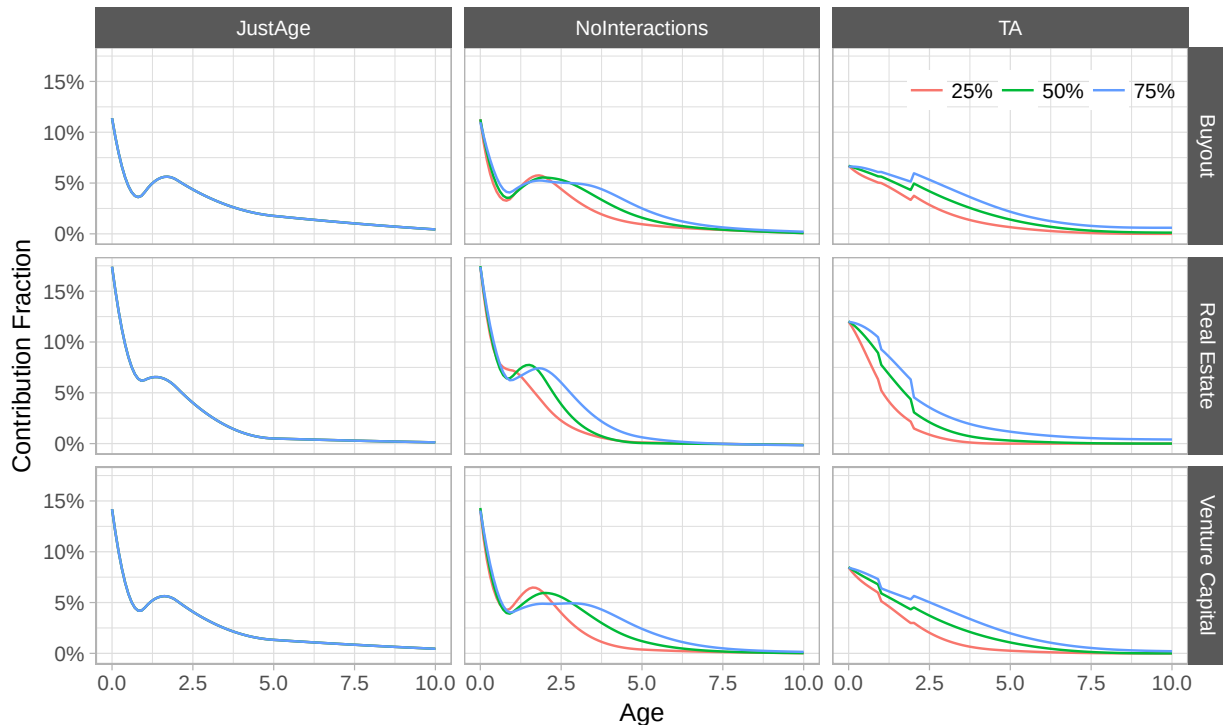


Figure 20: Comparison of predictions from the TA contribution model with those of other contribution models discussed in this paper

In figure 20 we see the predictions of TA for contributions (calibrated to all our data), where we have set the level of uncalled capital at three different levels (determined by quartiles). It is instructive to compare these model outputs with the empirical data shown in figure 2. For very young funds (age close to 0) TA underpredicts contributions. For funds between about 1 to 2 years of age it exhibits too much dependence on the amount of uncalled capital, while for funds that are about 4 years of age it, perhaps, exhibits insufficient dependence on the amount of uncalled capital. In fact note that this answers a natural question regarding the patterns of contributions to funds: to what degree are their contributions best modeled as being proportional to the amount of uncalled capital? TA assumes this proportionality; the models in this paper allow for more complex (or less complex) behaviors. What the data indicates (and the backtesting, below, supports) is that funds behave in non-proportional way for the first few years of their life (contra TA) although later (after about 3 years) they do exhibit a dependence on uncalled capital, as, indeed, they must. An interesting question regarding TA is how to set its various parameters. In the original paper (Takahashi and Alexander 2002) they suggest several different sets of parameters suitable for venture capital funds; however those parameters were based on the authors' experience with a relatively small set of venture capital funds two decades ago. Instead we calibrate the model's parameters to our universe of data (adhering to our backtesting methodology¹³). We show in figure 21 a backtesting-based comparison

¹²The TA model also has an assumed growth rate as a parameter, which we set to a suitable value. Takahashi and Alexander (2002) also includes a minimum distribution rate (termed *yield*) which we ignore.

¹³This is a secondary reason for not using the parameters from the original paper: doing so would result in a look-ahead bias for results before the year 2000.

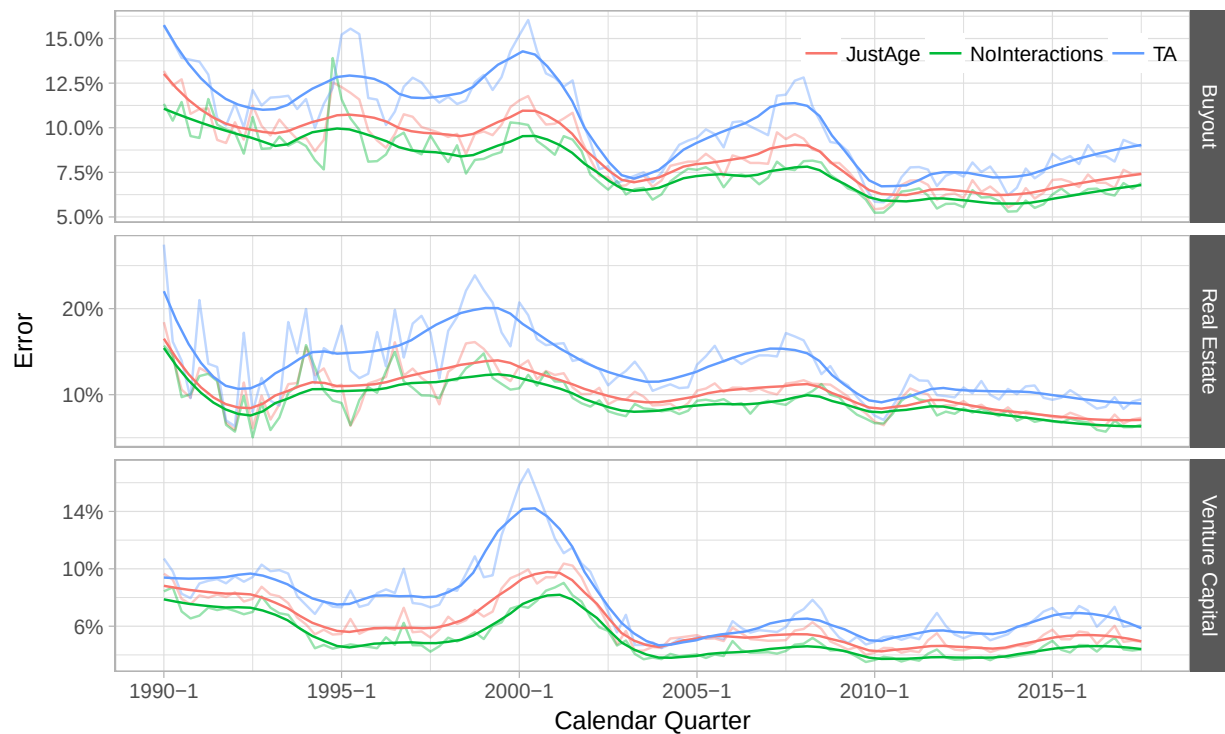


Figure 21: Backtesting-based (annual) comparison of the TA model for contributions with various other models from this paper

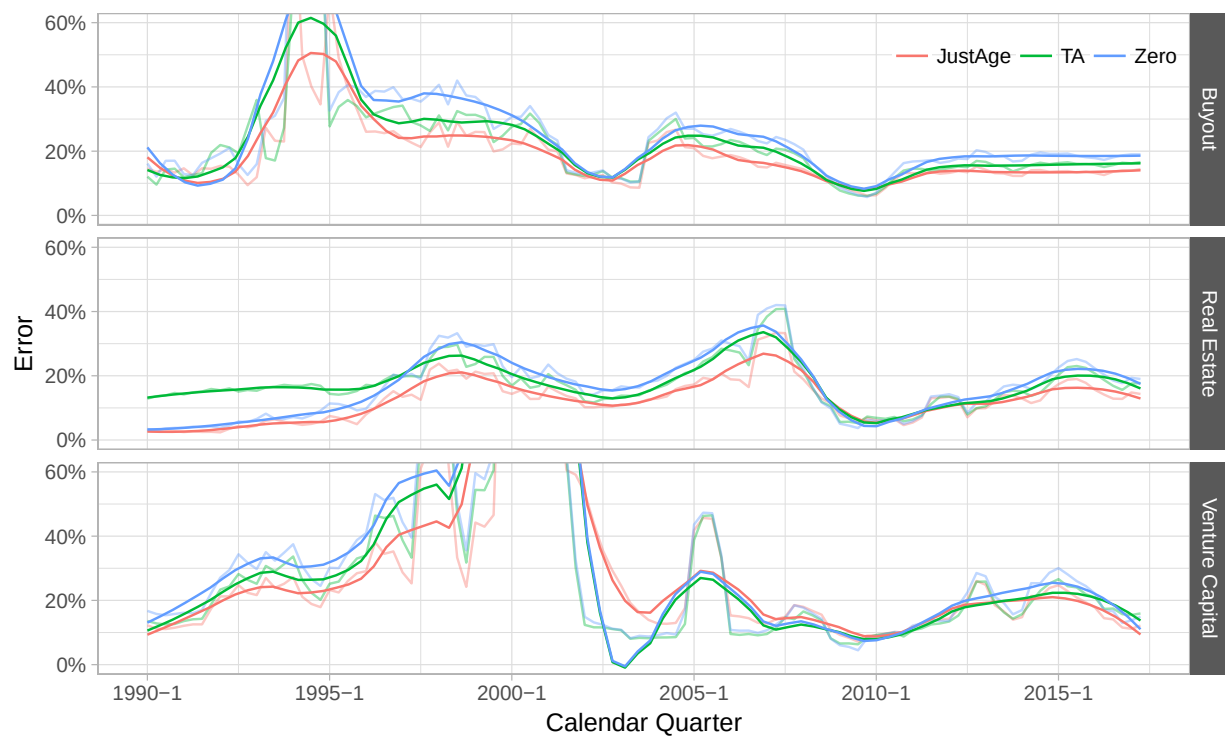


Figure 22: Backtesting-based (annual) comparison of the JustAge and TA model for distributions, which uses a bow of 2.0, annual growth of 13%, and a life span of 20 years for private capital funds

Table 5: Backtesting comparison of some models from this paper with Takahashi and Alexander. The models were calibrated on, and predicted, annual cash flows. Reported statistics are out-of-sample R^2 since 1990; the numbers in parentheses are since 2010.

Subclass	Out-of-sample R^2 / %					
	Contribution Models			Distribution Models		
	Constant	TA	JustAge	Constant	TA	JustAge
Buyout	30 (23)	47 (45)	62 (60)	24 (22)	25 (27)	45 (45)
Real Estate	19 (15)	35 (37)	61 (58)	29 (22)	14 (17)	43 (41)
Venture Capital	24 (19)	55 (55)	71 (71)	26 (-9 [†])	9 (15)	46 (25)

[†]Attributed to the fact that since 2010 Zero has been a better model than Constant.

of the TA model, using annual data, to some of the models for contributions explored in this paper. As can be seen that even JustAge (which exhibits no dependence on the amount of uncalled capital) outperforms the TA model for all quarters and for all subclasses. Note that before the GFC the underperformance of TA was very significant. In recent years the difference has declined somewhat, although for buyout it seems to be on the rise again. The TA model for contributions is compared with the other models explored in this paper using out-of-sample R^2 in table 5. The TA model is only a slight improvement over the Constant model but significantly underperforms JustAge (for all other models, see table 2).

Turning to a similar backtest of the TA model for distributions in figure 22 we see that again even a simple model (JustAge) almost always outperforms it (in this case, the exception is during the dot-com crash for venture capital). The out-of-sample R^2 comparison of the TA model for distributions is provided in table 5. As can be seen, in summary, TA significantly underperforms JustAge and other models for distributions as well (see table 3).

7 Lessons from Modeling Cash Flows

It turns out that modeling fund-level cash flows, from a purely statistical perspective of calibrating a linear model with some explanatory variables, is relatively straightforward. They say the devil is in the details; in this case the devil is in the choice of question, and in the omissions. This section addresses these issues.

7.1 Aggregation: Temporal versus Cross-Sectional

Quarterly cash-flow data (contributions or distributions) contains a large amount of idiosyncratic noise (we describe this more precisely below). Because of this, two models, one of which is much better than the other, may appear quite similar. Consequently it would be useful, for purely modeling reasons, to diminish this idiosyncratic noise. One way to achieve this is to aggregate the data in some way. In this section we discuss predicting cash flows over a period longer than a quarter (temporal aggregation) as well as predicting the cash flows of a portfolio of funds, rather than a single fund (cross-sectional aggregation).

We start with a discussion of what we called idiosyncratic noise by way of a simple example.

An Illustrative Example Consider tossing a coin 100 times, obtaining 52 heads, and estimating the (true) probability that the coin comes up heads based on that data. Suppose we ignore the

categorical nature of the outcome (heads or tails) and instead treat it as a continuous variable, where heads is 1 and tails is 0. And suppose we treat this as regression (with no independent variables) so all we need to estimate is the intercept. The natural estimate is 0.52 (the “good” estimate); let’s compare this with 0.6 (the “bad estimate”). If we compare the RMSEs then the good estimate gets a score of 0.4996 and the bad 0.5060, which might lead one to conclude that the two estimates are almost indistinguishable (and perhaps the bad estimate is in fact better since it is a round number). Before jumping to any conclusions let us take a second look at the data. Since the data is distributed binomially we can compute the relative likelihood of observing the given data as a function of our two estimates; from this perspective the good estimate makes the observations 3.7 times more likely than the bad estimate. Thus treating the data as a continuous variable creates a notion of error (RMSE) which is swamped by “coin noise”. In fact, the real issue is that in this example we only seek to model the *probability* of heads, but the RMSE numbers assume we are trying to predict the actual coin toss itself, and thus produce a performance measure that is mostly made up of errors that we do not consider to be errors at all, and which serve to mask the differences between models (in this case estimates of the probability of heads).

Indeed if funds always called capital (and distributed proceeds) in precise units of (say) 5% of the fund size, then the above analogy would be very close to reality. Alas, cash flows are more complicated and we must treat them as a continuous variable (because they are). Thus, like the coin example, our performance measures hide important differences between models. One way to highlight these differences is via aggregation. For example, in the case of the coin example one could group the 100 tosses into two groups of 50 tosses each (say 25/50 heads in the first and 27/50 heads in the second). This would also serve to make the difference between the two models very clear. A second strategy would be to ask many people to toss 100 coins and then average the number of heads in all their first tosses, all their second tosses, etc. These two strategies correspond to temporal and cross-sectional aggregation.

Temporal Aggregation In sections 4 and 5 we looked at cash flows over a periods of a quarter and then a year. The latter is an example of temporal aggregation, and indeed we found that the differences between models became much clearer in this case. However, there is a limit to temporal aggregation; anything more than about one year seems irrelevant (from a practical perspective). Thus it seems natural to turn to the second form of aggregation.¹⁴

Cross-Sectional Aggregation In our setting cross-sectional aggregation corresponds to forming portfolios of funds and predicting the total cash flows of the portfolio. Not only does this appear to be useful if one is trying to (statistically) distinguish models, but it is also the natural perspective of *users* of such a model since they would typically invest in a number of funds and be interested in the aggregated cash flows.

If one compares the performance of models on portfolios, however, something disappointing occurs. With the exception of the **Zero** model, all models become almost indistinguishable. Note that this includes the **Constant** model! The reason for this is that our models are linear and as a result the average prediction for several funds is the prediction of the average fund. Thus forming portfolios of funds makes the average fund in the portfolio become more and more alike to the average fund in any other portfolio. The **Constant** model more-or-less assumes that all funds are like an average fund, and hence all models converge to that simple model.

¹⁴A third kind of aggregation (also temporal in nature) is discussed in appendix F, where we discuss regressions in levels.

Thus statistically, cross-sectional aggregation is not a useful technique to distinguish models.¹⁵ A natural question is whether a good model is even necessary if one is primarily interested in a large well-diversified portfolio. We turn to this in the next section (section 7.3).

7.2 How to Measure Performance?

In this document we always measure the performance of a model based on its root-mean-square error. For example, when working with contributions we take the difference between the predicted contribution fraction and the observed contribution fraction, square it, compute the mean over some group of such errors, and take the square root. Typically we group by subclass and calendar quarter. The RMSE or the L^2 norm is given by

$$L^2(\{e_1, \dots, e_n\}) = \sqrt{\frac{1}{n}(e_1^2 + \dots + e_n^2)}.$$

This norm is, implicitly, what is used when a model is estimated via OLS. However the L^2 norm is sensitive to extreme values.¹⁶ An alternative norm is the L^1 norm:

$$L^1(\{e_1, \dots, e_n\}) = \frac{1}{n}(|e_1| + \dots + |e_n|),$$

which has the advantage of being less sensitive to outliers. As a crude approximation, one can consider the sample mean as minimizing the L^2 norm, while the sample median minimizes the L_1 norm. For cash-flow data, this makes the L^1 norm almost useless, since the median is often zero. Thus if models are compared via the L^1 norm, it is hard to beat the model Zero!

7.3 Cash Flows at Risk

Throughout this document we have assumed that the primary variable of interest was the *expected* cash flows. What does this mean? Roughly it means that if one had a very large set of identical funds, then the average contribution (or distribution) during the next quarter for that set is the expected contribution (or distributions). Suppose two-thirds of these identical funds would call zero and one third would call \$1M, which gives an expected contribution of \$0.33M per fund. Suppose an investor was invested in just one fund and was interested in how much cash would have to be available in order to service those calls. Would \$0.33M suffice? Clearly not. In fact with probability 1/3 such an investor would be short by \$0.66M. What this example shows is that the seemingly innocuous assumption of estimating expected cash flows, may not be well-aligned with the needs of investors. Instead, a better estimate would be a statement of the following sort: *the probability that one's contributions in the next quarter will be \$1.23M or less is 95%*. We term a quantity like \$1.23M a *cash flow at risk*. This notion is analogous to the notion of value at risk (VaR), the latter being commonly used when measuring market risk. However, the highly non-normal probability distribution of cash flows make this notion more challenging than in market risk. We intend to return to this topic in future research.

¹⁵This is probably too strong a statement. One possible approach to using portfolios to distinguish models would be to form portfolios that are deliberately alike in terms of their funds. For example one portfolio could be composed of only old funds, or only funds with a relatively large amount of uncalled capital.

¹⁶A more precise version of this statement is that it is an estimator that assumes the data is normally distributed. For other distributions it may be a poor estimator.

8 Conclusion

In this paper we attempted to model and predict cash flows for private capital using a splines-based linear regressions framework. We modeled contributions and distributions separately since they require somewhat different modeling approaches, and given that contributions are liabilities it is useful to understand them in isolation.

We explored a number of different models for contributions and found that most of the variation in cash flows can be explained in terms of age, but that nevertheless uncalled capital is a worthwhile addition to the set of regressors. A model with all interactions between age and uncalled capital was only marginally better than one without interactions, and does not seem justified in view of the danger of overfitting. With annual data we obtained an out-of-sample R^2 of about 60% (sometimes almost 80%).

Distributions were best modeled in terms of distribution ratios (as opposed to distribution fractions). Again, age was the most explanatory variable. In this case the next most explanatory variable—valuation ratio—was of limited value. In general, distributions are far less predictable than contributions and we obtained an out-of-sample R^2 of about 40%.

We also compared our models with those of Takahashi and Alexander (2002) and found that our models (even relatively simple ones) outperformed those of Takahashi and Alexander by a wide margin. This seemed to be due to Takahashi and Alexander underpredicting contributions for very young funds as well as misspecifying the dependence on the amount of uncalled capital.

Finally we drew a number of noteworthy conclusions regarding the modeling of cash flows. First, although RMSE is sensitive to extreme values, it is not possible to measure model performance using more robust measures (such as mean absolute error) since the latter tend to favor models that always predict zero. Second, in order to isolate model differences between models it is useful to perform aggregation. Indeed temporal aggregation makes it easier to distinguish models, however cross-sectional aggregation (forming portfolios) makes all models alike, and hence is useless as a technique to distinguish models. Third, the above issues draw attention to the question of whether predicting *expected* cash flows is the most natural model output. In fact, it could be argued that cash flows (especially contributions) should be approached through the lens of risk management and as a result we should be computing something akin to “cash flows at risk.” We intend to pursue this topic in subsequent research.

References

- Buchner, A., C. Kaserer, and N. Wagner (2010). “Modeling the cash flow dynamics of private equity funds: Theory and empirical evidence”. In: *The Journal of Alternative Investments* 13.1, pp. 41–54.
- Dalrymple, M. L., I. Hudson, and R. P. K. Ford (2003). “Finite mixture, zero-inflated Poisson and hurdle models with application to SIDS”. In: *Computational Statistics & Data Analysis* 41.3, pp. 491–504.
- de Malherbe, E. (2004). “Modeling private equity funds and private equity collateralised fund obligations”. In: *International Journal of Theoretical and Applied Finance* 7.03, pp. 193–230.
- (2005). “A model for the dynamics of private equity funds”. In: *The Journal of Alternative Investments* 8.3, pp. 81–89.
- Mathonet, P.-Y. and T. Meyer (2008). “J-Curve exposure: Managing a portfolio of venture capital and private equity funds”. In: John Wiley & Sons. Chap. 7, pp. 109–113.
- Robinson, D. T. and B. A. Sensoy (2016). “Cyclicality, performance measurement, and cash flow liquidity in private equity”. In: *Journal of Financial Economics* 122.3, pp. 521–543.
- Shmueli, G. (2010). “To explain or to predict?” In: *Statistical science* 25.3, pp. 289–310.
- Takahashi, D. and S. Alexander (2002). “Illiquid Alternative Asset Fund Modeling”. In: *The Journal of Portfolio Management* 28.2, pp. 90–100.
- Zuur, A. F., E. N. Ieno, N. J. Walker, A. A. Saveliev, and G. M. Smith (2009). “Zero-truncated and zero-inflated models for count data”. In: *Mixed effects models and extensions in ecology with R*. Springer, pp. 261–293.

A Regression Formula Notation

In this section we describe the notation used in this document to define regressions. The philosophy behind this notation is to compactly list all the *variables* that should appear on the RHS of the regression, without naming the associated coefficient.¹⁷

Suppose we have random variables y, a, b . Thus our data will consist of N observations of these variables $\{(y_i, a_i, b_i) \mid i = 1, \dots, N\}$. A simple model for this data might be

$$y_i = \alpha + \beta a_i + \epsilon_i.$$

This model explains y in terms of two variables: a and a constant (which we can think of as 1). A compact way of listing all the explanatory variables (and leaving the coefficients unnamed) is

$$y \sim 1 + a.$$

One can think of the “+” operator as taking the union of two sets of variables (namely 1 and a) to create a new set of variables (namely $\{1, a\}$). This can easily be extended to more variables and functions of variables. For example $y \sim 1 + a + a^2 + b + \exp b$ is a compact formula equivalent to

$$y_i = \alpha + \beta_1 a_i + \beta_2 a_i^2 + \beta_3 b_i + \beta_4 \exp b + \epsilon_i.$$

An important set of functions that can be applied to variables are *basis splines*. For example, suppose a is a variable taking on values between 0 and 10, then $\mathcal{B}(a; \text{degree} = 3; \text{knots} = 2)$ represents

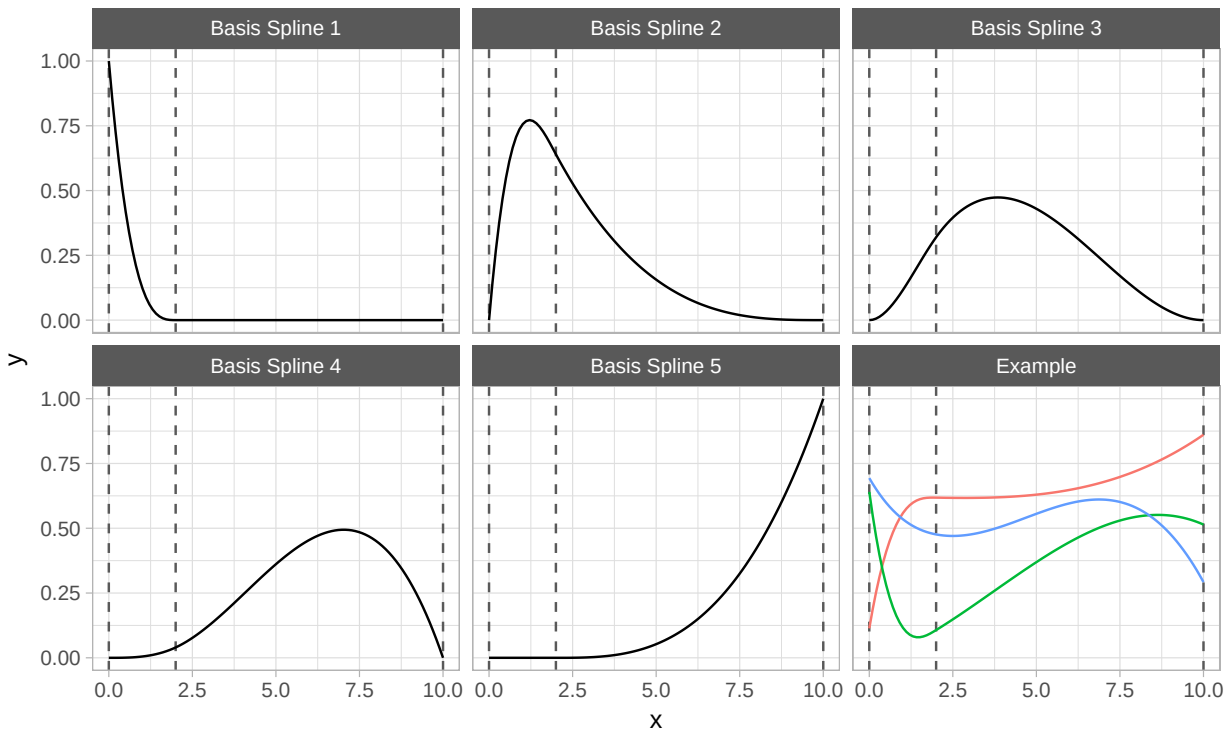


Figure 23: Cubic basis splines with knots at 0, 2, and 10 (indicated by dashed lines); the last facet, labeled “Example,” plots three linear combinations of the basis splines in the previous facets

¹⁷This notation is inspired by, and essentially the same as, the formula notation used in the statistical programming languages S and R.

cubic (i.e., degree 3) basis splines with knots at 0, 2 and 10; these basis splines are illustrated in figure 23. Note that $\mathcal{B}(a; \dots)$ represents a *set* of variables. Thus $y \sim 1 + \mathcal{B}(a; \text{degree} = 3; \text{knots} = 2)$ represents a regression formula with a total of six terms (five cubic basis splines of a and an intercept).

Finally, given two *sets* of variables we can take their Cartesian product, creating all pairwise interactions as well as the uninteracted variables. For example $y \sim 1 + \{a, a^2\} \times \{b, b^2\}$ is the same as $y \sim 1 + a + b + a^2 + b^2 + ab + ab^2 + a^2b + a^2b^2$.

B Data used for Model Estimation and Backtesting

All computational results in this paper are based on private capital data in the Burgiss Manager Universe (BMU) as of 2017 Q2.¹⁸ Our data consisted of United States equity funds denominated in USD (excluding funds of funds) from three subclasses of funds, namely buyout, real estate, and venture capital. These funds are distributed over a broad range of vintage years from 1980 to 2017 although we only used data from vintage year 1990 onwards. For each fund we have a complete history of quarterly valuations as well as all cash flows (the latter with date-level precision).

C Modeling Contributions

We model the *contribution fraction* of each fund in a quarter, defined as

$$c_{i,q} = \frac{C_{i,q}}{S_i}$$

where i indexes the set of funds in our estimation universe, S_i is the larger of the fund size and the sum of all contributions made to fund i (thus $\sum_q c_{i,q} \leq 1$), and $C_{i,q}$ are the contributions to fund i in quarter q .

Contribution Models

Zero

$$c_{i,q} \sim 0.$$

Constant

$$c_{i,q} \sim 1.$$

JustAge

$$c_{i,q} \sim 1 + \mathcal{B}(a_{i,q})$$

where $a_{i,q}$ is the age of fund i in quarter q , and $\mathcal{B}(\cdot)$ represent a set of basis splines (see appendix A). For the age variable we use quadratic splines with knots at 1, 2, and 5 years.

NoInteractions

$$c_{i,q} \sim 1 + \mathcal{B}(a_{i,q}) + \mathcal{B}(u_{i,q-1})$$

where $u_{i,q-1}$ is the uncalled capital fraction for fund i one quarter before quarter q (so $u_{i,q-1} = 1 - \sum_{q' < q} c_{i,q'}$). The splines for the age variable are the same as in the previous model. The splines for uncalled capital are cubic splines with a knot at 1/2.

¹⁸The Burgiss Manager Universe is a research-quality dataset comprised of nearly 40 years of daily cash flows and valuations for over 7,200 private capital funds, representing more than \$5 trillion of capital committed across the globe in various Private Equity, Private Debt and Real Asset strategies. The dataset covers a full spectrum of strategies across all geographies, and BMU data is representative of global institutional investor experience because it is sourced entirely from limited partners, avoiding the natural biases associated with other data sourcing models.

WithInteractions

$$c_{i,q} \sim 1 + \mathcal{B}(a_{i,q}) \times \mathcal{B}(u_{i,q-1})$$

See the previous section for an explanation of “ $\times$ ”. Both sets of splines are defined as in the previous models.

D Modeling Distributions

We compute the estimate of *total distributions* for each fund (indexed by i) in a quarter (indexed by q), defined as

$$T_{i,q-1} = \sum_{t=1}^{q-1} D_{i,t} + V_{i,q-1} + U_{i,q-1}$$

where $D_{i,t}$ is the distribution in quarter t and $V_{i,q-1}$ is the fund valuation and $U_{i,q-1}$ is the uncalled capital in the previous quarter. Using $T_{i,q-1}$, we define the *distribution ratio* and *valuation ratio* as follows, respectively:

$$d_{i,q} = \frac{D_{i,q}}{T_{i,q-1}}$$

and

$$v_{i,q-1} = \frac{V_{i,q-1}}{T_{i,q-1}}.$$

The analogous quantities in the fraction space are computed using fund size S_i , instead of $T_{i,q-1}$, as follows, respectively:

$$\delta_{i,q} = \frac{D_{i,q}}{S_i}$$

and

$$\nu_{i,q} = \frac{V_{i,q-1}}{S_i}.$$

Distribution Models The models for distributions do not have an intercept because distributions start with zero.

Zero

$$d_{i,q} \sim 0.$$

Constant

$$d_{i,q} \sim 1.$$

JustAge

$$d_{i,q} \sim \mathcal{B}(a_{i,q}).$$

For the age variable we use quadratic splines with knots at 2, 5, and 7 years.

NoInteractions

$$d_{i,q} \sim \mathcal{B}(a_{i,q}) + \mathcal{B}(v_{i,q-1}).$$

The splines for the age variable are the same as in the previous model. The splines for valuation ratio are specified with no knots.

WithInteractions

$$d_{i,q} \sim \mathcal{B}(a_{i,q}) \times \mathcal{B}(v_{i,q-1}).$$

Both sets of splines are defined as in the previous models.

WithInteractions-Fraction

$$\delta_{i,q} \sim \mathcal{B}(a_{i,q}) \times \mathcal{B}(\nu_{i,q-1}).$$

This model is defined in the *fraction* space, and uses the same sets of splines as defined in the previous models.

E Computing Out-of-sample R^2

During the backtesting of a model we generate a large number of predictions, $\hat{c}_i$, arising from many funds and many periods. Corresponding to these predictions we have the realized cash flows, c_i . We compute an out-of-sample R^2 measure of the performance of the model as follows

$$R^2 = 1 - E^2,$$

where

$$E = \frac{L^2(\{\hat{c}_1 - c_1, \dots, \hat{c}_n - c_n\})}{L^2(\{c_1, \dots, c_n\})}$$

and n is the total number of predictions.

F Regression in Levels

All the regressions in sections 4 and 5 directly model the cash flows we are trying to predict: the variable we are interested in is the contribution (or distribution) in some period, and the model tries to explain that variable in terms of others. Note that these cash flows are very far from normally-distributed; they are often zero (see figure 4) and when they are not zero they may be best modeled by an exponential (or gamma) distribution. Consequently, a tempting alternative is to model the evolution of *cumulative* variables (such as cumulative contribution fraction and cumulative distribution ratio; see figures 1 and 11). Once we have a model that predicts a cumulative variable then we can recover the cash flow by simply taking the difference between the current value of the cumulative variable and the predicted value. However, it turns out that models based on this idea generally underperform the more direct models described above. We believe the reason for this is straightforward: during estimation these models try to fit the cumulative variables; this, in effect, is similar to fitting the differences but with an age-dependent weighting scheme. This weighting scheme is absent when we measure the performance of the model, and hence the predictions are simply not optimal from the perspective of how we judge the models.

A variation on the above idea is to focus on the probability-density function (PDF) of the cumulative variable and model how it evolves with age. For concreteness suppose we are interested in contributions to buyout funds. Suppose further that we are interested in the contributions that will be made to a one-year-old fund in the next year. One could compute the PDF of contribution fractions for funds with age 1 and age 2, and compute the transition matrix between these two states. This would tell us not only what the expected cumulative contribution fraction would be in one year for our fund, but also the full PDF of such contributions. Subtracting off the current cumulative contribution fraction would yield the actual cash flows. Since contribution fractions have age-dependent medians, a simple transformation to this variable is to convert them into percentile

ranks (among funds of similar age). This new variable will be uniformly distributed between 0 and 1. An illustration of what we are trying to model is shown in figure 24.

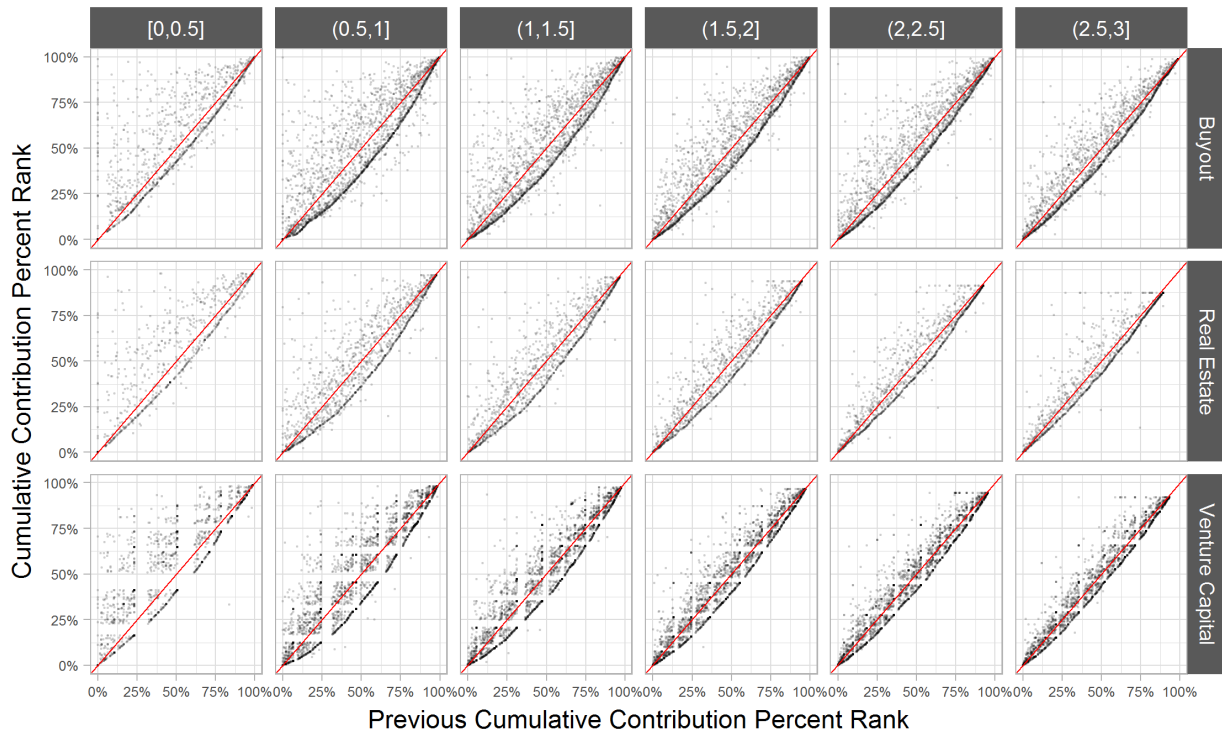


Figure 24: Transitions from the percentile rank of cumulative contribution fraction in the previous quarter to the current quarter

In this figure some interesting structure is apparent. For example, for venture capital the banding is a result of the fact that venture funds often call capital in multiples of 5% of the fund size. The dark bow-shaped line (which dips below the 45-degree line, in red) is the result of contributions often being zero, which result in the fund drifting downwards in its percentile rank. In terms of modeling, the simplest approach is to bucket the data by age and by percentile rank, then one can estimate a finite-dimensional transition matrix from one age bucket to the next.¹⁹ This approach introduces a tension between two forms of precision: the wider the buckets the more precisely the transition matrices can be estimated (since there will be more observations in each possible transition), however wide buckets lead to less precise percentile ranks (since we only know its value is somewhere in that bucket). Models based on bucketing underperformed the more traditional regression-based models. However, an approach based on bucketing is almost certainly not optimal; it is relatively model-free, but it does not take advantage of any “smoothness” in the data. For example, as one moves from one age bucket to the next in figure 24 it is clear that the data also changes smoothly; in addition as one moves along the horizontal or vertical axis the densities also vary smoothly. Thus a better approach would be to model these densities using some kind of smooth surface. We feel that this approach, with some work, would be an interesting alternative to the regression-based models, although there are some significant numerical challenges to fitting these surfaces robustly.

¹⁹An interesting and natural question is to what degree this process is Markovian. Having computed transition matrices it is a straightforward question to address; it boils down to how well the product of two consecutive transition matrices equals the corresponding two-step matrix. We do not discuss this question further in this document.

Notice and Disclaimer

This document and all of the contents (Content) within is the property of The Burgiss Group, LLC or its affiliates (collectively, “Burgiss”). The Information may not be reproduced or redistributed in whole or in part without prior written permission from Burgiss.

The Content is the confidential information of Burgiss, is only for internal use by the Clients of Burgiss (Clients), and may not be shared with third parties. The Content may not be used to create derivative works or be used to create any financial instruments or products and may not be used to provide investment consulting advice without prior written permission from Burgiss.

The Content is provided “AS IS” and any use of the Content is at Clients own risk. BURGISS MAKES NO EXPRESS OR IMPLIED WARRANTIES OR REPRESENTATIONS WITH RESPECT TO THE CONTENT (OR THE RESULTS TO BE OBTAINED BY THE USE THEREOF), AND TO THE MAXIMUM EXTENT PERMITTED BY APPLICABLE LAW, AND DISCLAIMS ALL IMPLIED WARRANTIES (INCLUDING, WITHOUT LIMITATION, ANY IMPLIED WARRANTIES OF ORIGINALITY, ACCURACY, TIMELINESS, NON-INFRINGEMENT, COMPLETENESS, MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE) WITH RESPECT TO ANY OF THE CONTENT.

Without limiting any of the foregoing and to the maximum extent permitted by applicable law, in no event shall Burgiss have any liability regarding any of the Content for any direct, indirect, special, punitive, consequential (including lost profits) or any other damages even if notified of the possibility of such damages.



111 River Street, 10th Floor
Hoboken, NJ 07030

p: +1 (201) 427.9600
f: +1 (201) 795.9237
w: burgiss.com