

TECHNOLOGY SUPPORT AND POST-ADOPTION IT SERVICE USE: EVIDENCE FROM THE CLOUD

German F. Retana

INCAE Business School
Alajuela, Costa Rica
gf@incae.edu

Chris Forman

Cornell University
Dyson School of Applied Economics and
Management, Cornell SC Johnson College of
Business, Ithaca, NY 14850
chris.forman@cornell.edu

Sridhar Narasimhan, Marius Florin Niculescu, D. J. Wu

Georgia Institute of Technology
Scheller College of Business
800 West Peachtree NW, Atlanta GA 30308
{sri.narasimhan, marius.niculescu, dj.wu}@scheller.gatech.edu

September 2012
This Version: May 2017

Abstract

Does a provider's technology support strategy influence its buyers' post-adoption IT service use? We study this question in the context of cloud infrastructure services. The provider offers two levels of support, basic and full. Under basic support, the provider handles simple service quality issues. Under full support, the provider also offers education, training, and personalized guidance through two-way interactions with buyers. Using unique data on public cloud infrastructure services use by 22,179 firms from March 2009 to August 2012, we find that buyers who receive full support use the service 34.85% more than other users. We further show that buyers who have full support deploy infrastructure services more efficiently, increasing the fraction of servers they run in parallel by 3.53 percentage points relative to those who do not. Furthermore, buyers who drop full support and switch back to basic support continue using 15.01% more of the service and have a proportion of servers running in parallel that is 2.82 percentage points higher compared to buyers who have never received full support. These findings provide suggestive evidence of buyer learning as a result of provider support.

Keywords: IT service, organizational learning, IT use, cloud computing, Infrastructure-as-a-Service, technology support, service strategies.

TECHNOLOGY SUPPORT AND POST-ADOPTION IT SERVICE USE: EVIDENCE FROM THE CLOUD

1. Introduction and Motivation

Businesses are increasingly shifting their information technology (IT) *infrastructure* from traditional on-premises deployment to the cloud. However, this shift is non-trivial from a technical standpoint. For example, many of the expected features of enterprise-grade servers, such as redundant components that ensure high availability and physical access to servers, are not present in the cloud. The cloud requires users to design for failure (Reese 2009) and consider how to keep an application running if any given server randomly disappears. Moreover, the cloud's scaling capabilities can only be exploited if the applications scale out horizontally (i.e., employ several servers performing functions in parallel) rather than vertically (i.e., increasing capacity of single servers). The former scaling method involves a greater degree of technical sophistication than the latter. Thus, it is difficult for some of the buyers of cloud infrastructure services to overcome these knowledge barriers *on their own*. While faster access to infrastructure and greater scalability attract firms to cloud infrastructure, it does not come as a surprise that the *lack of in-house resources and expertise* is reported as the prime challenge faced by firms attempting to place workloads in the cloud (RightScale 2016).

It has been well-documented in the Information Systems (IS) literature that firms' internal capabilities and technical know-how affect both the timing of new IT adoption (Attewell 1992; Bresnahan and Greenstein 1996; Forman et al. 2008) as well as the post-adoption usage (Parthasarathy and Bhattacharjee 1998; Zhu and Kraemer 2005; Zhu et al. 2006). In particular, firms are known to delay not only the adoption (initial purchase) but also the actual assimilation of a technology because of knowledge barriers (Åstebro 2004; Fichman and Kemerer 1997). While extant literature has shown the role of organizational learning in overcoming knowledge barriers (Attewell 1992; Chatterjee et al. 2002; Fichman and Kemerer 1997), much less is known about how providers' knowledge transfer strategies affect buyers' consumption of IT services. While a firm's own internal efforts in learning are known to be associated with its post-adoption use of business IT systems (Åstebro 2004), to our knowledge there is little quantitative evidence on how a *provider's* strategies to transfer knowledge to buyers affect the realized post-adoption consumption level for the offered service.

In this study, we take steps towards filling this research gap. In the context of cloud infrastructure services, we focus on a strategy that directly facilitates interactions and knowledge transfer between providers and buyers – the offering of *personalized* technology support – and seek to measure its impact on service use in terms of volume and efficiency. In our research setting, the provider's buyers use its hardware resources and choose between two levels of technology

support, *basic* or *full*. When receiving full support¹, buyers have access to *personalized* guidance and training, and thus have the opportunity to learn through *two-way* interactions with the provider from the latter's prior experience in deploying applications in the cloud. Through full support, the provider takes a proactive approach significantly beyond the level of basic support to lower the aforementioned knowledge barriers associated with the usage of cloud infrastructure services. Full support is different from pure outsourcing – where the provider does everything for the buyer and “takes the burden of learning off the back of a potential user” (Attewell 1992) – in that buyers remain in charge of implementing the business process and must acquire the knowledge to control and utilize the provider's resources. Table 1 summarizes some attributes of buyers who may opt to use or switch between both support levels.

Table 1. Characteristics of Buyers Employing or Switching between each Support Level		
Support Level Used	Basic	Full
Typical Buyer ^a	Buyers who know what they are doing. If they break something, they know how to fix it themselves.	Buyers who want to avoid the technicalities and want provider to help them with everything. They may not have a clear idea on what cloud infrastructure is or they know what it is yet do not know how to deploy applications on it.
Why switch to this level of support? ^a	Buyers feel they have learned enough to not need the provider's safety net. Their willingness to pay the full support premium has fallen. Buyer-side budget cuts may also force buyer to drop full support.	Buyers are in need of better skills and may find it more profitable to pay to access provider's knowledge than to invest in internal team. They may still have in-house IT staff, but not system administrators with cloud knowledge.
Examples of Software Applications ^a	Proprietary software: Buyers who have developed their own in-house applications, implying they have a strong development team, and run these apps in the cloud.	Standard or open source software: Buyers with standard retail applications (e.g., Magento) or some other CMS who want assistance in its deployment.
Common Industries ^b	IT Services, software, consulting, telecommunications.	e-commerce, education, financial services, non –profit.
Common Use Cases ^c	SaaS offering, test and development environment.	Corporate website, e-commerce site, social media site (blogging, social networking, etc.), marketing campaign or advertising, online publishing.
^a Insights on this table are the main points extracted from semi-structured interviews to provider executives, account managers, and support agents. ^b Insight based on optional single-selection sign-up survey item completed by basic support buyers and full support buyers who did not switch support level during their tenure. Common industries within each support level were identified as those where we find a significantly greater proportion of buyers belonging to that industry relative to the other support level. ^c Insight based on optional mark-all-that-apply sign-up survey item completed by basic support buyers and full support buyers who did not switch support level during their tenure. The survey item asked buyers their intentions for how they planned to use the service. Common use cases within each support level were identified as those where we find a significantly greater proportion of buyers selecting each use case relative to the other support level.		

¹ While *receiving* (i.e., to be a recipient) may be interpreted as a buyer actively and continually interacting with the provider, the term can also be defined as the capability of receiving a service (recipient, a., 2017). Hence, it fits our research setting given that buyers who have access to or receive full support may or may not actively engage with the provider in a continuous fashion.

To test these assertions, we collected unique data from a major global public cloud provider of infrastructure services (computing power and storage). Our panel data consist of 22,179 firms that used the provider's service at some point between March 2009 and August 2012. We find that buyers receiving full support use, on average, 34.85% more of the IT service relative to buyers who receive basic support. We also provide evidence that technology support helps buyers make better and more efficient use of the service by quantifying the effects that full support has on buyers' deployment of horizontally distributed and scalable architectures. Buyers increase the fraction of servers they run in a parallel and horizontally scalable architecture by 3.53 percentage points after they switch from basic to full support. Given that the mean proportion of servers running in parallel in our sample is only 12%, this is an economically significant change in behavior.

We also find evidence that full support's effects persist even if buyers eventually drop full support and switch back to basic support. Former full support buyers continue using, on average, 15.01% more of the service and have a proportion of servers running in parallel 2.82 percentage points higher compared to buyers who have never received full support. Since in our setting buyers who switch back to basic must redeploy their architectures on their own, the differentiated behavior between former full support buyers and those who never received support is suggestive of the durability of buyer learning.

Last, we extend these models to allow the effects of receiving and dropping full support to vary over time. The difference in service use between buyers receiving full and basic support continues to increase over time from the moment the former adopted full support. A potential reason for this is that buyers who have received full support are more capable of foreseeing new cloud deployment opportunities. Service usage diminishes little over time after the buyer drops full support, further suggesting buyer learning has taken place.

To alleviate concerns of reverse causality and omitted variable bias, we probe the robustness of our results through the use of matched samples, instrumental variables, and a generalized methods of moments (GMM) estimation approach. Our results continue to qualitatively hold under these alternative approaches, supporting our initial analysis and theory.

To our knowledge, this is the first study to quantitatively document how technology support can influence IT service use and offer suggestive evidence that technology support facilitates buyer learning. As such, our work not only informs the IS literature on post-adoption IT usage but also offers managers evidence of the importance of overcoming knowledge barriers to cloud infrastructure service use through *two-way* interactions with the provider.

2. Empirical Model

2.1. IT Service Use

We employ linear fixed effects dynamic panel data models to tease out the effects of receiving and dropping full support on IT service use. We model two dimensions of IT service use: volume and efficiency of use.

In our setting, the provider bundles server capacity in terms of memory (GB of RAM), processing power (number of virtual CPUs), and storage (GB space of local hard disk). The three attributes are highly correlated in the offer menu; a server with more of one attribute has more of the other two. Since the servers are priced based on the amount of memory they have, and memory is the basis for buyers' infrastructure sizing decisions, the amount of memory consumed over time is a direct measure of buyers' volume of use of the cloud infrastructure service. We compute the average GB of RAM used by a buyer per month and denote it as $Memory_{i,t}$. Then, given the strong positive skew in its distribution, following standard practice we compute $\ln Memory_{i,t} = \ln(Memory_{i,t} + 1)$ and use it as our first dependent variable.

As mentioned earlier, buyers who receive full support may learn from the provider through two-way interactions that enable them to make more efficient use of the cloud service. An advantage of our setting is that we can partially observe certain attributes of buyers' deployments, some of which are diagnostic in assessing how proficient a buyer is in making use of the infrastructure. If full support helps buyers use the service better, one would expect that they employ architectures that can scale more efficiently, albeit at the cost of increased complexity. We explain this assertion below.

Although the on-demand nature of the service along with its rapid elasticity provides firms with the opportunity to reduce idle computing capacity waste and eliminates the necessity of an up-front capital commitment in overprovisioning resources (Armbrust et al. 2010; Harms and Yamartino 2010), doing so requires firms to scale their infrastructure in a cost-efficient manner. There are two ways of growing an IT infrastructure: vertically and horizontally (Garcia et al. 2008; Michael et al. 2007; Reese 2009, p. 176). Scaling vertically is easy to execute since it only involves increasing the capacity of the single server performing a function. However, it does not allow the buyer to truly leverage the cloud's scalability. For example, growth is capped by the maximum server capacity available. In contrast, scaling horizontally—with several servers performing functions in parallel—is complex.² Although launching a single server is a trivial task for any system administrator, launching several of them in a horizontally scalable manner is non-trivial. However, horizontal scaling offers virtually unlimited growth potential plus it allows buyers to have a more resilient architecture through redundancies.

² The complexities include load balancing workloads and managing concurrent sessions across servers, among others (Casalicchio and Colajanni 2000; Cherkasova 2000, and interviews with cloud experts at IBM Thomas J. Watson Research Center, Yorktown Heights, New York, and a major technological research university).

Given the benefits of scaling horizontally, we use the fraction of servers running in parallel as a measure that proxies for buyer efficiency in service use. This measure varies separately from memory use, our first dependent variable. For example, a buyer can consume a large volume of memory with none of its servers running in parallel, in which case the fraction is zero, or alternatively it can consume a small volume with all of its servers running in parallel, which makes the fraction equal to 1. To compute this metric we scan the names of the servers used daily by buyers and count, to the extent possible, how many of them are performing the same functions.³ The monthly average fraction of servers running in parallel is captured in our second dependent variable, $FractionParallel_{i,t}$. A summary of these two dependent variables and all other variables used in this work is available in Table 2.

2.2. Effects of Full Support on Service Use

Our first model tests if receiving or having dropped full support is associated with greater volume and efficiency of use. Letting $y_{i,t} \in \{lnMemory_{i,t}, FractionParallel_{i,t}\}$, we have:⁴

$$y_{i,t} = \alpha + \beta FullStatus_{i,t} + \gamma FormerFullStatus_{i,t} + \sum_{s=1}^p \lambda_s y_{i,t-s} + \mu_i + \tau_t + v_{i,t} + \varepsilon_{i,t} \quad (1)$$

Subscripts i and t index individual buyers (firms) and time periods (months) respectively. $FullStatus_{i,t}$ is a binary variable that indicates if full support was adopted by buyer i by time t , and is equal to one in all periods after the buyer adopts full support. Thus, β identifies the effects on cloud use of receiving full support; we expect $\hat{\beta} > 0$. After adopting full support some buyers may opt to switch to basic support. $FormerFullStatus_{i,t}$ is a binary variable that signals if buyer i has dropped full support by the end of the focal month t but was using full support at the start of the focal month or in some prior month(s). The sum $\beta + \gamma$ identifies differences in use behavior between basic support buyers who received full support in the past and those who exclusively received basic support. If the effects of technology support are lasting then we should find that $\hat{\beta} + \hat{\gamma} > 0$.

³ We develop an algorithm (available upon request) that compares the names of the servers being run by each buyer at the end of every day during our sample and check if we find servers with names very similar to each other. Our assumption is that if we find two or more servers with very similar names, they will very likely be performing the same function in parallel (e.g., `web1.domain.com` and `web2.domain.com`).

⁴ We acknowledge that $FractionParallel_{i,t} \in [0,1]$ and that our model represents a linear approximation to a nonlinear model. However, in other settings with fractional dependent variables it has been shown that linear models offer similar estimates of the coefficients to those of more sophisticated models such as a fractional probit (Papke and Wooldridge 2008). More importantly, whereas nonlinear models such as the cross-sectional fractional probit can be used with unbalanced panel data (Wooldridge 2011), they are unable to accommodate fixed effects and the use of lagged values of our variables as instruments. The latter are key elements of our empirical model.

Table 2. Summary of Variables		
Variable	Role	Description
$Memory_{i,t}$	Metric	Average GB of RAM memory used by buyer i during month t .
$\ln Memory_{i,t}$	Dependent variable	$= \ln(Memory_{i,t} + 1)$
$FractionParallel_{i,t}$	Dependent variable	Average proportion of servers run in parallel by buyer i during month t .
$FullStatus_{i,t}$	Support choice status indicator	Indicates if full support was adopted by buyer i by time t . If buyer i received full support for the first time in time period z , then $FullStatus_{i,t} = 1_{\{t \geq z\}}$.
$FormerFullStatus_{i,t}$	Support choice status indicator	Indicates if buyer i dropped full support (i.e., switched to basic) by time t . If buyer i switched from full to basic support in period w , then $FormerFullStatus_{i,t} = 1_{\{t \geq w\}}$.
$AdoptFull_{i,t}$	Support choice indicator	Indicates if full support was adopted by buyer i at time t . If buyer i received full support for the first time in time period z , then $AdoptFull_{i,t} = 1_{\{t=z\}}$.
$DropFull_{i,t}$	Support choice indicator	Indicates if full support was dropped by buyer i at time t . If buyer i dropped full support (i.e., switched to basic) in period w , then $DropFull_{i,t} = 1_{\{t=w\}}$.
μ_i	Fixed effect	Buyer fixed effect. A vector of dummies - one dummy per buyer i .
τ_t	Fixed effect	Calendar time fixed effect. A vector of dummies - one dummy per calendar month t in the data.
$v_{i,t}$	Fixed effect	Buyer tenure time fixed effect. A vector of dummies - one dummy per each month in buyer i 's tenure (i.e., months since adoption of cloud service). ^b
$\varepsilon_{i,t}$	Error Term	Assumed to be correlated only within individual buyers, but not across them
$FullStatus_{i,t}^f$	Instrument	Fitted value of $FullStatus_{i,t}$ attained from probit models that use failure-related variables as covariates.
$FailOutageN_{i,t}$	Instrument ^a	Indicates if buyer i has suffered at least N service outage-related failures by time t .
$FailNetworkN_{i,t}$	Instrument ^a	Indicates if buyer i has suffered at least N network outage-related failures by time t .
$FailHardwareN_{i,t}$	Instrument ^a	Indicates if buyer i has suffered at least N physical hardware-related failures by time t .
^a Please see Online Appendix 0 for further details on the construction of all the exogenous failure-related instruments.		
^b This is an alternative to having a discrete integer valued buyer tenure term (e.g., $Tenure_{i,t}$). The $v_{i,t}$ vector allows to control for the possibility that buyers' use of the service may increase in a nonlinear fashion over time t .		

We additionally include lagged values of the dependent variables to control for persistence in use behavior, i.e., that buyers' use in prior periods may strongly influence their use in the focal period. This approach suffers from dynamic panel bias as it fails the strict exogeneity assumption common to fixed effects panel models (Nickell 1981; Roodman 2009a). We address this bias through System GMM estimation (Anderson and Hsiao 1981; Archak et al. 2011; Arellano and Bond 1991; Arellano and Bover 1995; Blundell and Bond 1998; Ghose 2009). We show our main results using $p = 2$ lags in Model (1) for comparability with our later System GMM estimations; however, our models are consistent (i.e., panels do not have unit roots) if we use fewer or more lags (e.g., $p = 1, 3$). We elaborate on our use of System GMM in Section 4.1.

Parameter μ_i is the buyer fixed effect and τ_t is a vector of calendar month fixed effects. We also include a vector of dummy variables, $v_{i,t}$, indicating in what month of its tenure a buyer is when month t starts. Finally, parameter $\varepsilon_{i,t}$ is our error term which we assume is correlated only within individual buyers, but not across them.

Our fixed effects model allows us to difference out unobserved time-invariant buyer-level heterogeneity that may influence both the choice of support level and IT use. We also run our models using a matched sample constructed using a coarsened exact matching (CEM) procedure (Blackwell et al. 2010) – details are included in Section 0. CEM reduces the dependence of our estimates on our model specification and also reduces endogeneity concerns when making causal inferences (Ho et al. 2007).

Despite these steps, our estimates may be influenced by time-varying unobserved factors correlated with the support decision and service use. To address this issue, we probe the robustness of our results to the use of instrumental variables. We use exogenous failure events experienced by buyers as an instrument for their decision to receive full support. We identify 3 types of failures: generalized outages across the cloud infrastructure service, network-related failures, and degraded performance issues due to hardware problems on the physical server host; please see Table 3 for their detailed descriptions. We employ a probit model that has the exogenous failures while receiving basic support as regressors to generate predicted values for $FullStatus_{i,t}$, which we denote $FullStatus_{i,t}^f$. We then use the fitted value, $FullStatus_{i,t}^f$, as our instrument in a standard two-stage least squares (2SLS) estimation (Angrist and Pischke 2009, pp. 142-144; Imbens and Wooldridge 2007).

Table 3. Description of exogenous failure incidents used as instruments	
Failure Type	Description of the Event
Service outage	Provider may suffer from generalized outages in different components of its service (e.g., memory leak in provider's cloud management system). Such generalized problems are announced in the provider's status webpage and/or announced to buyers.
Network-related failure	Some node in the provider's infrastructure, generally belonging to some buyer, is suffering from a distributed denial of service attack (DDoS) or some networking hardware device has failed.
Hardware-related failure	Buyer is suffering degraded service performance due to a hardware-related problem in the physical host in which the buyer's virtual machine runs.

These exogenous failures satisfy the criteria for an instrument. When unforeseeable problems occur (e.g., an unexpected failure on the provider's hardware), the support interactions that take place between buyers and the provider can serve as a signal to buyers of the value of full support. Basic support buyers who, because of the failure, gain experience using the service with a greater involvement from the provider, may be more likely to upgrade to full support than buyers who do not have such experiences with the provider. However, such interactions on their own are

unlikely to alter use of the provider's service. Since the failures are exogenous (i.e., can occur with equal probability to any server independent of the support choice), they are not directly related to the technical sophistication of the buyer. One potential concern is if failures are more prevalent for buyers who are using more servers or deploying a more complex architecture. The lags of our dependent variables in our model, which proxy for size and complexity, control for this.

2.3. Time Varying Effects of Full Support

Model (1) allows us to identify the overall effects of receiving each support level, but not how such effects may vary with the time elapsed since the switch between support types. Knowledge transfer is a potential cause for the change in buyer behavior, and, in turn, it is reasonable to expect that the amount of learning is linked in some way to the length of exposure (or lack thereof) to full support. To allow the marginal effect of switching to and from full support to vary in a flexible way over time, we employ lags of indicators of the adoption event, $AdoptFull_{i,t}$, and the switching to basic support event, $DropFull_{i,t}$. These variables are set to 1 only in the period when full support is adopted or when it is initially dropped, respectively. Thus, $AdoptFull_{i,t-j}$ indicates if buyer i adopted full support j periods ago (counting from period t), and $DropFull_{i,t-k}$ indicates if buyer i dropped full support k periods ago. We use these indicators in the following autoregressive distributed lag (ARDL) model (Greene 2008, pp. 681-689):

$$\begin{aligned}
y_{i,t} = & \alpha + \sum_{j=0}^q \beta_j AdoptFull_{i,t-j} + \beta_{\infty} FullStatus_{i,t-(q+1)} \\
& + \sum_{k=0}^r \gamma_k DropFull_{i,t-k} + \gamma_{\infty} FormerFullStatus_{i,t-(r+1)} \\
& + \sum_{s=1}^p \lambda_s y_{i,t-s} + \mu_i + \tau_t + \nu_{i,t} + \varepsilon_{i,t}.
\end{aligned} \tag{2}$$

We include $q = r = 12$ lags of $AdoptFull_{i,t}$ and $DropFull_{i,t}$ so that our model identifies the effects of adopting or dropping full support during the 12 months (1 year) following the event. Our results are consistent with those under a different number of lags. We then leverage the β_j and γ_k coefficients to estimate the dependent variables' impulse response function (Hamilton 1994, pp. 318-323) to the buyer's decision to adopt or drop full support. The dummy variable $FullStatus_{i,t-(q+1)}$ controls for the effect of having adopted full support more than q time periods ago; $FormerFullStatus_{i,t-(r+1)}$ does the same for having dropped full support more than r periods ago.

3. Data and Sample Construction

Our data set includes 79,619 buyers that used the provider's services at some point between March 2009 and August 2012. When using the cloud infrastructure services, buyers under basic support only pay hourly rates contingent on server capacity and operating system. Buyers under full support pay a fixed price premium per server-hour used plus an additional fixed monthly fee (which is prorated on a daily basis). There are no sign-up or termination fees for the servers' usage or the full support service. Online Appendix A offers further details about the provider's cloud infrastructure services, their pricing, and the corresponding levels of technology support. To isolate the causal effects of full support, we restrict our baseline sample to buyers who are likely to have similar usage profiles over time, but for their adoption of full support. We exclude buyers who use the service very little or who do not change their cloud architecture configuration (i.e., do not resize their infrastructure).⁵ These buyers have very different time-varying profiles relative to full support buyers and, although we exclude them ex ante, it is likely that they would have also been excluded later by our CEM procedures. After these restrictions, our baseline sample includes 22,179 buyers and 368,606 buyer-month observations. Table 4 provides descriptive statistics of the cloud use time-varying variables in our baseline sample, both in aggregate as well as contingent on buyers' support choice ($FullStatus_{i,t}$); difference in means t-tests for all variables are significant at the 1% level.

Table 4. Descriptive Statistics of Time-Varying Variables (Baseline sample, 22,179 buyers)												
Support Type Used	Full or Basic				$FullStatus_{i,t} = 0$				$FullStatus_{i,t} = 1$			
Observations	368,606				309,544				59,062			
Variable	Mean	S.D.	Min	Max	Mean	S.D.	Min	Max	Mean	S.D.	Min	Max
$Memory_{i,t}$	7.88	31.37	0	2,284.54	7.26	30.92	0	2,284.54	11.11	33.41	0	1,917.40
$\ln Memory_{i,t}$	1.35	1.04	0	7.73	1.30	1.01	0	7.73	1.62	1.15	0	7.56
$FractionParallel_{i,t}$	0.12	0.27	0	1	0.12	0.26	0	1	0.13	0.28	0	1
$FullStatus_{i,t}$	0.16	0.37	0	1	0	0	0	0	1	0	1	1
$FormerFullStatus_{i,t}$	0.01	0.09	0	1	0	0	0	0	0.05	0.22	0	1

In addition to the buyers' cloud use data, we have collected data from a survey administered to buyers upon sign-up of a new account from which we identify the buyers' total employment and their intended use case for the cloud service. After joining the survey data with the cloud usage data, we match buyers who exclusively receive basic support (controls) to buyers who start with basic support and later upgrade to full support (treated) across six different attributes: (1) pre-upgrade volume of IT use (i.e., memory use), (2) pre-upgrade efficiency of IT use or architecture

⁵ We exclude buyers who only received basic support and averaged 512 MB RAM/hour or less during their first 6 months (excluding 1st month) or made no adjustments to the size of their infrastructure during their first 6 months (excluding 1st month). An infrastructure resizing occurs in any launch, halt, or resizing of a server in the buyers' cloud infrastructure. We do not consider their behavior during their 1st month in our threshold because most buyers are setting up their infrastructure during this time. Results without excluding these buyers are consistent with our main findings.

complexity (i.e., fraction of servers running in parallel), (3) pre-upgrade frequency of cloud infrastructure resizing (i.e., how often buyers launch a server, halt a server, or resize an existing one), (4) operating system of preference, (5) employment, (6) and intended use case for the cloud infrastructure service. The matching process yields our CEM-based sample of 1,525 buyers. Further details regarding the sign-up survey and the construction of our CEM-based sample are included in Online Appendix B.

Finally, we have also collected data on the timing and content of all support interactions through online live chat sessions and support tickets between the buyers and the provider, starting from October 2009. We provide further details on these data when we describe our instrumental variables procedure.

4. Results

4.1. Effects of Technology Support on IT Use

The results for Model (1) are shown in Table 5. We show results for both dependent variables and with both the baseline and CEM samples. The estimates across the two samples are very consistent with each other so hereafter we leave the results with the baseline sample as reference and discuss the results with the CEM sample. Moreover, we show results employing two lags of the dependent variables as covariates ($p = 2$) for consistency with our System GMM model results below, though the results are consistent if we use fewer or more lags (e.g., $p = 1, 3, 4$).

Table 5. Main Results				
Column	(1)	(2)	(3)	(4)
Dependent Variable	$y_{i,t} = \ln Memory_{i,t}$		$y_{i,t} = FractionParallel_{i,t}$	
Sample	Baseline	CEM	Baseline	CEM
$FullStatus_{i,t}$	0.297*** (0.008)	0.299*** (0.023)	0.032*** (0.002)	0.035*** (0.006)
$FormerFullStatus_{i,t}$	-0.143*** (0.018)	-0.159*** (0.045)	-0.008*** (0.004)	-0.007*** (0.011)
$y_{i,t-1}$	0.953*** (0.006)	0.967*** (0.017)	0.897*** (0.005)	0.885*** (0.020)
$y_{i,t-2}$	-0.149*** (0.005)	-0.190*** (0.014)	-0.165*** (0.004)	-0.127*** (0.018)
Observations	324,406	25,298	324,406	25,298
Buyers	21,573	1,525	21,573	1,525
R ²	0.773	0.791	0.637	0.657
Upgrade change	34.57%	34.85%	3.24	3.53
Downgrade change	16.68%	15.01%	2.47	2.82
$\hat{\beta} + \hat{\gamma} = 0$ test p-value	0.000	0.001	0.000	0.016
All regressions include monthly calendar (τ_t) and tenure dummies ($v_{i,t}$). Robust standard errors, clustered on buyers, in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$. In columns(1) and (2) upgrade and downgrade changes are computed as $e^{\hat{\beta}} - 1$ and $e^{\hat{\beta} + \hat{\gamma}} - 1$ respectively. In columns(3) and (4) upgrade and downgrade changes are computed as $\hat{\beta} \times 100$ and $(\hat{\beta} + \hat{\gamma}) \times 100$ respectively.				

The results in column (2) indicate that buyers who receive full support use 34.85% more memory than buyers who receive basic support. In column (4), we find that full support is associated with an increase of 3.53 percentage points in the proportion of servers running in a parallel and horizontally scalable manner. Together, these results suggest the provider's support has a positive and significant influence on buyer IT service use. Furthermore, column (2) suggests that buyers who drop full support continue consuming about 15.01% more relative to buyers who have never received full support. Similarly, former full support buyers have a proportion of servers functioning in parallel about 2.82 percentage points higher than those who never received full support. These outcomes are suggestive of buyer learning, especially if we consider that because of technical limitations in the service offering (during our observation period), buyers must redeploy their servers after they downgrade. Thus, the observed post-downgrade behavior is the outcome of buyers acting on their own without personalized guidance from the provider.

Instrumental Variables Approach: We use the support interaction data to identify when buyers suffer from exogenous failures when using the cloud service. These unforeseeable exogenous shocks force the buyer to interact with the provider, which serves as a useful signal of the provider's service capabilities. Buyers may discover that by interacting more closely with the provider they can reduce their total cost of solving cloud-related problems. This motivates them to upgrade from basic to full support so that they can continue to have similar interactions. As noted above, the number of failures may be correlated with the number of servers a buyer is employing. It may also be correlated with the complexity of the architecture employed, however we believe this to be a less significant concern given the nature of the failures we consider. We instrument using lagged failures rather than concurrent failures and also include lags of our dependent variables as controls to mitigate the risk of the potential correlation between the failures and our metrics of IT use. Finally, although the failures can serve as instruments for the full support adoption decision, we lack an appropriate instrument for the switching to basic support decision. Hence, in this section we employ a reduced version of Model (1) that omits $FormerFullStatus_{i,t}$, letting $FullStatus_{i,t}$ represent buyer behavior post-upgrade irrespective of a potential downgrade decision.

The vectors of variables $FailOutageN_{i,t}$, $FailNetworkN_{i,t}$, and $FailHardwareN_{i,t}$, correspond to the types of exogenous failures described in detail in Table 3 and Online Appendix 0. They consist of dummies that are turned on if buyers have experienced at least N failures of each corresponding type by time t . In this section we comment on our results using 2 dummies for each failure type (i.e., $N = 1, 2$), yet our results are consistent using 1 or 3 of them.

Given our binary endogenous variable, we follow the approach suggested by Imbens and Wooldridge (2007) and Angrist and Pischke (2009, pp. 142-144) and first include the vector of failure-related indicators in a probit model with $FullStatus_{i,t}$ as dependent variable. We use each failure type independently in columns (1) through (3) in Part A of Table 6, and all 3 types of failures in column (4). The results suggest that, as proposed, all failure types are positively associated with buyers' likelihood of using full support. We use the probit model to generate fitted

values of $FullStatus_{i,t}$, which we denote as $FullStatus_{i,t}^f$. The descriptive statistics of the fitted values are in Part B of Table 6. Next, we use $FullStatus_{i,t}^f$ as our instrument for $FullStatus_{i,t}$ in a 2SLS estimation procedure.

Table 6. Probit for $FullStatus_{i,t}$ and Descriptive Statistics for $FullStatus_{i,t}^f$				
Column	(1)	(2)	(3)	(4)
Failure Types	Outage	Network	Hardware	All 3
Part A. Coefficients of Probit with $FullStatus_{i,t}$ as dependent variable				
$FailOutage1_{i,t-1}$	0.871*** (0.052)			0.472*** (0.057)
$FailOutage2_{i,t-1}$	0.619*** (0.097)			0.135 (0.117)
$FailNetwork1_{i,t-1}$		1.169*** (0.088)		0.771*** (0.101)
$FailNetwork2_{i,t-1}$		1.558*** (0.227)		1.181*** (0.255)
$FailHardware1_{i,t-1}$			1.084*** (0.043)	0.925*** (0.045)
$FailHardware2_{i,t-1}$			0.438*** (0.076)	0.302*** (0.081)
Constant	-0.674 (0.681)	-0.674 (0.681)	-0.956 (0.801)	-0.909 (0.775)
Observations ^a	26,629	26,629	26,629	26,629
Pseudo-R2	0.115	0.103	0.146	0.164
Part B. Descriptive Statistics of $FullStatus_{i,t}^f$				
Mean	0.099	0.099	0.099	0.099
Std. Dev.	0.094	0.086	0.111	0.117
Min	0.000	0.000	0.000	0.000
Max	0.913	0.956	0.900	0.988
Probit regressions in Part A include monthly calendar (τ_t) and tenure dummies ($v_{i,t}$). ^a Number of observations in these models is slightly larger than in others (e.g., 26,629 here vs. 25,298 elsewhere) since these models only employ 1 lag of the covariates whereas other models employ 2 lags of the dependent variables (used as covariates), thus altering the number of observations available. Robust standard errors, clustered on buyers, in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.				

The first stage results for $lnMemory_{i,t}$ are reported in Part B of Table 7. The values of the F-statistic for the excluded instruments range between 65.73 and 79.61, and in all cases are significant at the 1% level. The second stage results are reported in Part A of the same table. Column (1) is the fixed effects specification of Model (1) (no instruments) excluding $FormerFullStatus_{i,t}$ and is included for comparison purposes. The results across all models with instruments and those in column (1) are highly consistent with each other. They suggest the adoption of full support increases the volume of service consumption by 35.53% to 49.99%.

Table 7. Two-Stage Least Squares Estimation for Volume of IT Use					
Column	(1)	(2)	(3)	(4)	(5)
Failure Types	None	Outage	Network	Hardware	All 3
Part A. Second Stage Results for $\ln Memory_{i,t}$					
$FullStatus_{i,t}$	0.304*** (0.023)	0.405*** (0.141)	0.389** (0.170)	0.353*** (0.105)	0.367*** (0.096)
$\ln Memory_{i,t-1}$	0.943*** (0.024)	0.931*** (0.030)	0.933*** (0.034)	0.937*** (0.028)	0.935*** (0.027)
$\ln Memory_{i,t-2}$	-0.171*** (0.020)	-0.170*** (0.020)	-0.170*** (0.020)	-0.170*** (0.020)	-0.170*** (0.020)
Observations	25,298	25,297	25,297	25,297	25,297
Buyers	1,525	1,524	1,524	1,524	1,524
R ²	0.782	0.733	0.733	0.734	0.734
Upgrade change ($e^{\hat{\beta}} - 1$)	35.53%	49.99%	47.59%	42.27%	44.37%
Part B. First Stage Regression of Fitted $FullStatus_{i,t}^f$ on Real $FullStatus_{i,t}$					
$FullStatus_{i,t}^f$		0.400*** (0.087)	0.437*** (0.066)	0.423*** (0.074)	0.439*** (0.061)
$\ln Memory_{i,t-1}$		0.124*** (0.010)	0.126*** (0.010)	0.122*** (0.010)	0.121*** (0.010)
$\ln Memory_{i,t-2}$		-0.017*** (0.007)	-0.014*** (0.007)	-0.018*** (0.007)	-0.021*** (0.007)
Observations ^a		25,297	25,297	25,297	25,297
Buyers ^a		1,524	1,524	1,524	1,524
R ²		0.112	0.112	0.123	0.130
First Stage F-Statistic		65.73	70.66	71.13	79.61
All regressions include monthly calendar (τ_t) and tenure dummies ($v_{i,t}$). Robust standard errors, clustered on buyers, in parentheses.					
^a 2SLS models with instruments in columns (2) through (4) drop 1 singleton buyer with a single observation.					
* $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.					

We turn to Table 8 for the results with $FractionParallel_{i,t}$ as the dependent variable. The first stage results again suggest $FullStatus_{i,t}^f$ is positively associated with $FullStatus_{i,t}$. The F-statistics have lower values than before but remain statistically significant at the 1% level. Moving to Part A of the table, we again have in column (1) results based on Model (1) omitting the downgrade indicator. With the exception of column (3) that only employs the network-related failures as instrument and where $FullStatus_{i,t}$ does not have a statistically significant coefficient, the remaining specifications are consistent with each other. They suggest that receiving full support is associated with an increase of 3.01 to 7.89 percentage points in the proportion of servers running in parallel. In sum, the results of the 2SLS estimations are consistent with those of our prior specification in Table 5.

Table 8. Two-Stage Least Squares Estimation for Efficiency of IT Use					
Column	(1)	(2)	(3)	(4)	(5)
Failure Types	None	Outage	Network	Hardware	All 3
Part A. Second Stage Results for $FractionParallel_{i,t}$					
$FullStatus_{i,t}$	0.030*** (0.005)	0.079*** (0.035)	0.021 (0.044)	0.056** (0.023)	0.047** (0.022)
$FractionParallel_{i,t-1}$	0.901 (0.020)	0.893*** (0.021)	0.903*** (0.022)	0.897*** (0.021)	0.898*** (0.021)
$FractionParallel_{i,t-2}$	-0.144*** (0.018)	-0.145*** (0.018)	-0.144*** (0.018)	-0.145*** (0.018)	-0.145*** (0.018)
Observations	25,298	25,297	25,297	25,297	25,297
Buyers	1,525	1,524	1,524	1,524	1,524
R ²	0.661	0.640	0.645	0.643	0.644
Upgrade change ($\hat{\beta} \times 100$)	3.01	7.89	2.14	5.56	4.75
Part B. First Stage Regression of Fitted $FullStatus_{i,t}^f$ on Real $FullStatus_{i,t}$					
$FullStatus_{i,t}^f$		0.545*** (0.097)	0.498*** (0.069)	0.547*** (0.076)	0.554*** (0.063)
$FractionParallel_{i,t-1}$		0.153*** (0.030)	0.168*** (0.031)	0.153*** (0.030)	0.151*** (0.030)
$FractionParallel_{i,t-2}$		-0.011 (0.020)	0.001 (0.020)	-0.009 (0.019)	-0.017 (0.019)
Observations ^a		25,297	25,297	25,297	25,297
Buyers ^a		1,524	1,524	1,524	1,524
R ²		0.035	0.028	0.051	0.062
First Stage F-Statistic		21.38	28.83	28.18	37.18
All regressions include monthly calendar (τ_t) and tenure dummies ($v_{i,t}$). Robust standard errors, clustered on buyers, in parentheses.					
^a 2SLS models with instruments in columns (2) through (4) drop 1 singleton buyer with a single observation.					
* $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.					

System GMM Estimation and Endogenous Switching Decisions: While the failure events identified through the support interactions are completely unexpected to the buyer, their exogeneity can still be questioned if buyers with a greater number of servers or more complex architectures are more likely to suffer failures in any of their servers. To address this concern, we employ System GMM estimation methods that consider IT use and support choice as endogenous and use their lagged values as their instruments. The use of system GMM will also allow us to address concerns that our use of a lagged dependent variable violates the strict exogeneity assumption in panel data models.⁶

We first find the minimum number of lags of the dependent variables that we can use while finding a valid specification that passes the Hansen (1982) J test of overidentifying restrictions and does not suffer from serial correlation (Arellano and Bond 1991). The tests' outcomes are similar for both dependent variables and suggest the minimum number of lags we can use of each is two ($p = 2$). Then, we find the minimum number of lags of the covariates that we can use as

⁶ In addition to estimating system GMM, we note that the number of time periods in our regressions is large and that any bias from violating the strict exogeneity assumption asymptotically goes to zero as the number of time periods increases (Hsiao 2003).

instruments to avoid the problem of overfitting the model with too many instruments (Roodman 2009b). Finally, we run our models first relying solely on the lags of the covariates as instruments and then augmenting our instrument matrix with the exogenous failure-based instruments used in column (4) of Table 6. The results of these estimations are shown in Table 9; the table's footer has details of the specific instruments used in each model.

Table 9. System GMM Estimation Results				
Column	(1)	(2)	(3)	(4)
Dependent Variable	$y_{i,t} = \ln Memory_{i,t}$		$y_{i,t} = FractionParallel_{i,t}$	
Failure-based IVs	No	Yes	No	Yes
$FullStatus_{i,t}$	0.076*** (0.028)	0.088*** (0.030)	0.020** (0.009)	0.021** (0.009)
$FormerFullStatus_{i,t}$	0.124 (0.108)	0.100 (0.099)	0.048 (0.034)	0.046 (0.033)
$y_{i,t-1}$	0.813*** (0.078)	0.800*** (0.075)	0.763*** (0.054)	0.762*** (0.056)
$y_{i,t-2}$	0.152 (0.087)	0.159 (0.085)	0.055 (0.057)	0.055 (0.057)
Observations	25,298	25,298	25,298	25,298
Buyers	1,525	1,525	1,525	1,525
Total Number of IVs	259	265	455	461
Hansen J Statistic p-value	0.654	0.381	0.592	0.399
Upgrade change	7.94%	9.25%	1.95	2.06
Downgrade change	22.13%	20.68%	6.73	6.68
$\hat{\beta} + \hat{\gamma} = 0$ test p-value	0.054	0.048	0.032	0.030
All regressions include monthly calendar (τ_t) and tenure dummies ($v_{i,t}$). Results for $\ln Memory_{i,t}$ in columns (1) and (2) have AR(2), and hence use the 2nd lag of the first difference of all covariates as their instruments for the levels equation. They also use the 3rd lag of $\ln Memory_{i,t}$ and $FullStatus_{i,t}$ as well as the 3rd to 8th lags of $FormerFullStatus_{i,t}$ as instruments for the differences equation. Upgrade and downgrade changes are computed as $e^{\hat{\beta}} - 1$ and $e^{\hat{\beta} + \hat{\gamma}} - 1$ respectively. Results for $FractionParallel_{i,t}$ in columns (3) and (4) have AR(2), and hence use the 2nd lag of the first difference of all covariates as their instruments for the levels equation. They also use the 3rd to 8th lags of $\ln Memory_{i,t}$ and $FullStatus_{i,t}$ as well as the 3rd to 7th lags of $FormerFullStatus_{i,t}$ as instruments for the differences equation. Upgrade and downgrade changes are computed as $\hat{\beta} \times 100$ and $(\hat{\beta} + \hat{\gamma}) \times 100$ respectively. Models in columns (2) and (4) augment the instruments matrix by considering the same vector of exogenous failure-related instruments shown in in columns (4) of Table 4. Robust standard errors using Windmeijers' (2005) finite sample correction. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.				

Columns (1) and (2) show the results for the volume of IT use. The coefficients for $FullStatus_{i,t}$ suggests an increase in memory usage of 7.94% to 9.25%. The coefficients for $FormerFullStatus_{i,t}$ do not show evidence of a change in buyer consumption after the downgrade decision, and the volume of usage is statistically greater than that of buyers who had never received full support. Columns (3) and (4) present the results for the efficiency of IT use. The upgrade to full support is associated with an increase of 1.95 to 2.06 percentage points in the fraction of servers running in parallel. Meanwhile, the downgrade action, if anything, appears to be associated with a further increase in the proportion. Overall, the outcomes are qualitatively similar to those attained before.

4.2. Time Varying Effects of Full Support

The estimation results for Model (2) employing both the baseline and CEM samples are shown in Table 10. The coefficients do not change much if we employ a different number of lags for the support indicators (q and r) or the dependent variables (p). Nevertheless, in this model computing the marginal effects of adopting and dropping full support on $\ln Memory_{i,t}$ and $FractionParallel_{i,t}$ is not straightforward. $AdoptFull_{i,t}$ and $DropFull_{i,t}$ (and their lags) influence our dependent variables in two ways. First, they influence behavior directly through the coefficient estimates on those variables. Second, they influence behavior indirectly through their effects on the lags of the dependent variables. Thus, while the results in Table 10 indicate that the direct effects of $AdoptFull_{i,t}$ and $DropFull_{i,t}$ may decrease in absolute value over time, the marginal effects (i.e., the combination of the direct and indirect effects) cannot be read off the coefficient estimates in the table.

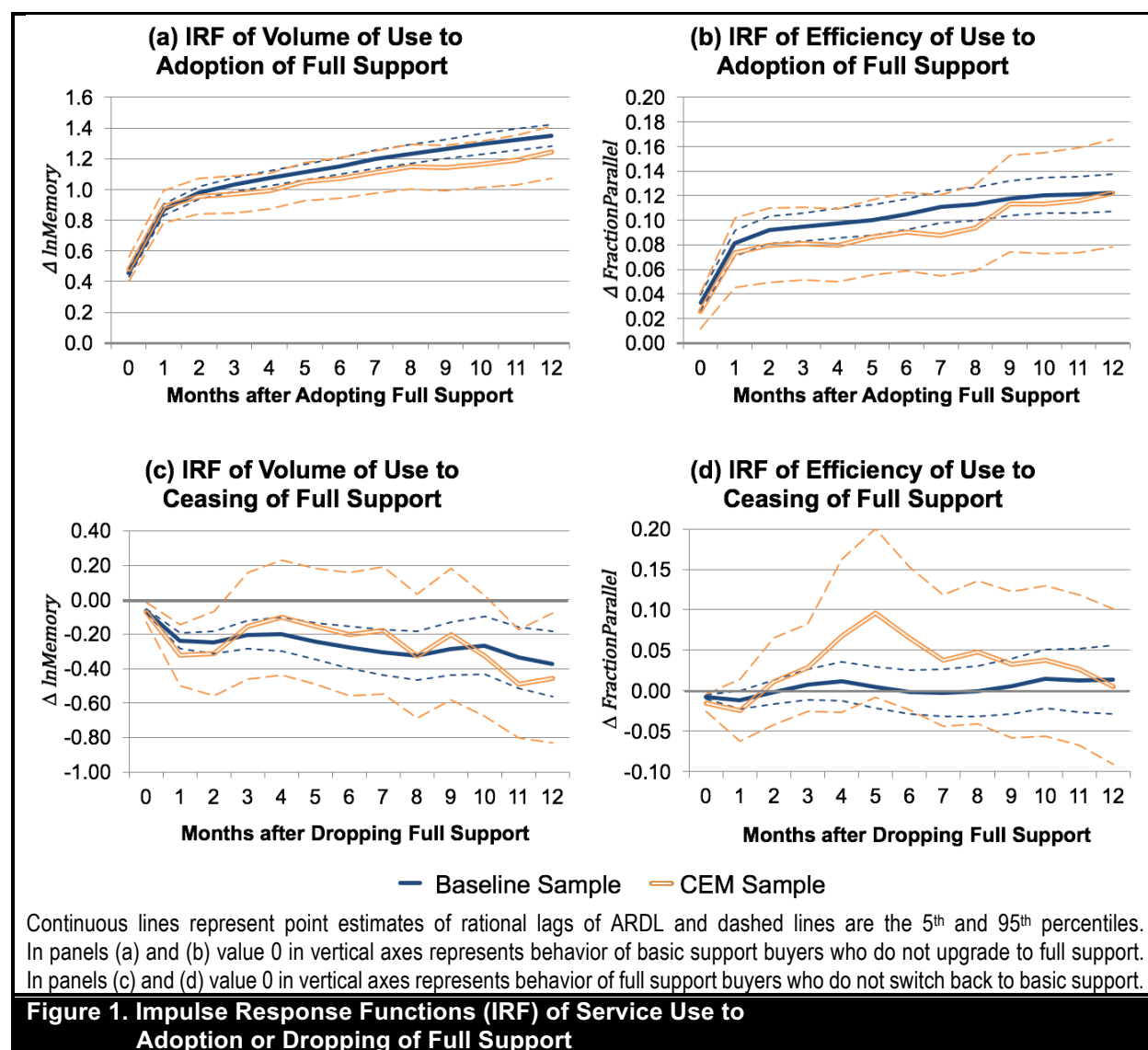
Table 10. Results with Lags of Full Support Adoption Indicators				
	(1)	(2)	(3)	(4)
Dependent Variable	$y_{i,t} = \ln Memory_{i,t}$		$y_{i,t} = FractionParallel_{i,t}$	
Sample	Baseline	CEM	Baseline	CEM
$AdoptFull_{i,t}$	0.458*** (0.017)	0.482*** (0.047)	0.033*** (0.004)	0.026*** (0.009)
$AdoptFull_{i,t-1}$	0.432*** (0.015)	0.418*** (0.041)	0.052*** (0.005)	0.049*** (0.012)
$AdoptFull_{i,t-2}$	0.217*** (0.010)	0.182*** (0.033)	0.025*** (0.003)	0.016*** (0.007)
$AdoptFull_{i,t-3}$	0.222*** (0.009)	0.201*** (0.029)	0.025*** (0.003)	0.020*** (0.007)
$AdoptFull_{i,t-4}$	0.234*** (0.009)	0.221*** (0.028)	0.028*** (0.002)	0.018*** (0.008)
$AdoptFull_{i,t-5}$	0.236*** (0.009)	0.265*** (0.029)	0.028*** (0.002)	0.026*** (0.007)
$AdoptFull_{i,t-6}$	0.248*** (0.009)	0.233*** (0.031)	0.031*** (0.002)	0.024*** (0.007)
$AdoptFull_{i,t-7}$	0.259*** (0.009)	0.259*** (0.028)	0.033*** (0.003)	0.018*** (0.007)
$AdoptFull_{i,t-8}$	0.256*** (0.009)	0.262*** (0.034)	0.031*** (0.003)	0.027*** (0.007)
$AdoptFull_{i,t-9}$	0.263*** (0.010)	0.228*** (0.028)	0.034*** (0.003)	0.041*** (0.011)
$AdoptFull_{i,t-10}$	0.270*** (0.010)	0.263*** (0.031)	0.033*** (0.003)	0.024*** (0.007)
$AdoptFull_{i,t-11}$	0.272*** (0.011)	0.267*** (0.041)	0.032*** (0.003)	0.030*** (0.011)
$AdoptFull_{i,t-12}$	0.275*** (0.011)	0.296*** (0.032)	0.034*** (0.003)	0.033*** (0.007)
$FullStatus_{i,t-13}$	0.291*** (0.010)	0.274*** (0.031)	0.035*** (0.003)	0.030*** (0.008)

Table 10 continues on next page.

$DropFull_{i,t}$	-0.066 ^{***} (0.006)	-0.074 ^{***} (0.034)	-0.008 ^{***} (0.002)	-0.015 ^{***} (0.007)
$DropFull_{i,t-1}$	-0.176 ^{***} (0.028)	-0.251 ^{***} (0.109)	-0.004 (0.007)	-0.010 (0.022)
$DropFull_{i,t-2}$	-0.029 (0.021)	-0.015 (0.082)	0.007 (0.005)	0.031 (0.023)
$DropFull_{i,t-3}$	-0.002 (0.027)	0.094 (0.171)	0.007 (0.007)	0.014 (0.017)
$DropFull_{i,t-4}$	-0.041 (0.023)	-0.013 (0.054)	0.004 (0.008)	0.043 (0.047)
$DropFull_{i,t-5}$	-0.082 ^{***} (0.022)	-0.083 (0.056)	-0.005 (0.007)	0.038 (0.031)
$DropFull_{i,t-6}$	-0.075 ^{***} (0.031)	-0.070 (0.075)	-0.003 (0.007)	-0.012 (0.017)
$DropFull_{i,t-7}$	-0.077 ^{***} (0.027)	-0.012 (0.069)	-0.000 (0.007)	-0.007 (0.017)
$DropFull_{i,t-8}$	-0.075 ^{***} (0.026)	-0.193 ^{***} (0.094)	0.001 (0.007)	0.024 (0.022)
$DropFull_{i,t-9}$	-0.019 (0.030)	0.087 (0.155)	0.006 (0.007)	-0.005 (0.019)
$DropFull_{i,t-10}$	-0.041 (0.040)	-0.195 ^{***} (0.050)	0.009 (0.008)	0.015 (0.018)
$DropFull_{i,t-11}$	-0.124 ^{***} (0.041)	-0.209 ^{***} (0.057)	0.001 (0.015)	-0.003 (0.015)
$Drop$	-0.091 ^{***} (0.037)	-0.038 (0.121)	0.005 (0.011)	-0.013 (0.018)
$FormerFullStatus_{i,t-13}$	-0.091 ^{***} (0.027)	-0.152 (0.118)	-0.007 (0.007)	0.008 (0.013)
$y_{i,t-1}$	0.953 ^{***} (0.006)	0.971 ^{***} (0.017)	0.897 ^{***} (0.005)	0.925 ^{***} (0.020)
$y_{i,t-2}$	-0.144 ^{***} (0.005)	-0.181 ^{***} (0.014)	-0.165 ^{***} (0.004)	-0.167 ^{***} (0.017)
Observations	324,406	25,298	324,406	25,298
Buyers	21,573	1,525	21,573	1,525
R-Squared	0.774	0.795	0.638	0.670
All regressions include monthly calendar (τ_t) and tenure dummies ($v_{i,t}$). Robust standard errors, clustered on buyers, in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.				

To show the time-varying effects of support we plot the impulse response functions (IRFs) of the dependent variables to the switch in the support type (i.e., a unit change in a binary variable) (Hamilton 1994, pp. 318-323). Specifically, we estimate and plot $\frac{\Delta \hat{y}_{i,t}}{\Delta x_{i,t-j}}$ ($y_{i,t} \in \{\ln Memory_{i,t}, FractionParallel_{i,t}\}, x_{i,t} \in \{AdoptFull_{i,t}, DropFull_{i,t}\}$) over time j to show how current service usage is influenced by the switch events from j periods ago. We describe the estimation procedure in detail in Online Appendix D.

In Figure 1, we show the IRFs of the dependent variables with respect to the adoption and dropping of full support. Panels (a) and (b) suggest that, following the adoption of full support, buyers' service consumption grows month-over-month compared to similar buyers who only receive basic support. The same is true of the proportion of servers that run in parallel. The results hold in both in the baseline and CEM samples. Therefore, the difference in service use between buyers receiving full and basic support continues to increase over time from the moment the former adopted full support.



While panels (a) and (b) show that full support positively impacts usage growth over time, panels (c) and (d) tell a somewhat different story when buyers drop full support. Panels (c) and (d) show that after dropping full support, buyers do not radically alter their usage compared to buyers who continue using full support. When looking at the CEM sample, for most part, the 90% confidence interval of the estimates of the IRF contains value 0, which represents the behavior of full support buyers who do not downgrade. The upper bound of the confidence interval falls below 0 at the very beginning and towards the end of the time frame of the IRF evaluation in panel (c), indicating a mild decline in use. When looking at the complexity of deployment (panel (d)) we do not find any statistically significant evidence of behavioral change after dropping full support. When looking at the baseline sample, panel (c) shows a slight decline in use after dropping full support, but panel (d) does not show a significant change in complexity of the deployment. Along with our prior results related to buyer behavior after dropping full support (i.e., $\hat{\beta} + \hat{\gamma}$ from Model (1)), the figure suggests buyers continue leveraging what they have learned from the provider even after ending two-way communications.

5. Conclusion

To our knowledge, this note provides the first empirical evidence of how a service provider's technology support influences a buyer's post-adoption IT use. We show that enhanced technology support increases volume and efficiency of usage, and also provide suggestive evidence that buyer learning from the provider may be responsible for these patterns. Our results call for a more complete and fully integrated theory of organization learning and post-adoption usage in the IS literature.

Our study has important managerial implications. From the provider's perspective, our results highlight the impact of full support on user behavior. Before our study, the provider who is the subject of our study was unsure of the precise economic value of offering full support (costs were understood but the impact on the revenue stream was unclear). A rough estimate of the profit gains for the provider from having a buyer under basic support vs. full support suggests the switch yields at least 147% increase in profits after considering revenue and support cost increases (see Table 11). Due to their commoditization, cloud services have been perceived as being fully self-service, on-demand offerings with minimal necessity for interactions between buyers and service providers (Mell and Grance 2011). For example, Amazon, the largest provider of cloud infrastructure services, initially did not offer technology support. Our findings suggest that the buyers' continuous access to full support has significant, quantifiable and sustainable business value. Thus, our research adds to other recent findings about the value of service and support in the cloud setting. For example, Retana et al. (2016) show that proactively providing customers with information about the value of a service during the customer onboarding process decreases both customer attrition and the number of costly support interactions.

Table 11. Estimate of Net Profit Gain from Full Support vs. Basic Support per Buyer

Item	Support Type		Units	Calculation
	Basic	Full		
Support Costs				
Number of Chats	0.366	0.702	Quantity / month	Mean number of chats / month
Cost of a Chats	\$2.73	\$5.24	\$ / month	Quantity × \$7.46 ^a
Number of Tickets ^b	0.117	0.650	Quantity / month	Mean number of tickets / month
Cost of a Tickets	\$4.31	\$23.95	\$ / month	Quantity × \$36.83 ^a
Cost of Support	\$7.04	\$29.19	\$ / month	Costs of Chats + Cost of Tickets
Cloud Server Profits				
Estimated Usage ^c	1,440.0	1,941.8	GB RAM/month	For full, median usage × 1.3485
Server Hourly Rate ^d	\$0.045	\$0.090	\$ / GB RAM / hour	Based on AWS pricing.
Estimated ARPU ^e	\$64.80	\$174.77	\$ / month	Estimated Usage × Hourly Rate
Estimated Profits	\$51.84	\$139.81	\$ / month	ARPU × 80% ^f
Difference in Profits				
Net Profits	\$44.80	\$110.63	\$ / month	Server Profits – Support Costs
Net Profits Gains (abs.)	\$65.83		\$ / month	\$110.63 – \$44.80
Net Profits Gains (%)	147%		%	\$110.63 / \$44.80 – 1
^a These are the estimated costs per chat session and ticket given to us by the provider.				
^b We only count buyer-initiated (inbound) tickets. We exclude (outbound) announcements by provider through tickets.				
^c Median usage under basic support is 2 GB RAM/hour; we multiply by 720 hours/month to get monthly usage under basic support. For full support we consider a 34.85% increase in usage from estimate in column (2) of Table 5.				
^d During our sample period, Amazon Web Services' (AWS) Elastic Compute Cloud (EC2), the public IaaS with the largest market share and thus with the dominant price-setting position, offered small 1.7 GB RAM servers at \$0.08/hour and medium 3.75 GB RAM servers at \$0.16/hour (source: aws.amazon.com). Based on these rates, we compute the mid-point price for 1 GB RAM / hour at \$0.045. This is the price under basic support. For full support, even though the provider adds \$0.12 to the hourly rate, we only add \$0.045 to attain a conservative estimate. We also ignore the fixed monthly fee charged by the provider to buyers under full support. See Online Appendix A for more details on pricing.				
^e Average Revenue per User.				
^f The provider estimates their server-related variable costs are around 20%. These include server and datacenter depreciation expenses, datacenter rent, power and cooling, and non-infrastructure related items like credit card fees and bad debt expenses.				

Our research has included a range of analyses used to isolate the effects of full support on IT service use. However, as in any empirical study, our research has limitations. In particular, while our analyses of user decisions to use horizontally scalable architectures provide suggestive evidence of learning, we do not directly observe learning using our current research design. Further, while we have sought to address concerns related to omitted variable bias through a range of approaches, we recognize the inherent challenges of identification given the nature of decisions we study and our use of observational data.

Such limitations offer exciting opportunities for future research. For example, better data on the productivity of service use would help researchers to more precisely isolate the effects of provider interactions on buyer learning. More broadly, we believe that future work should use transactional data such as ours to gauge the impact of other buyer interactions with third parties, such as traditional outsourcing firms and individuals in online communities of practice, to assess their impact on the manner and effectiveness with which firms use IT. We hope our findings will encourage additional work in this important area.

References

- Anderson, T.W., and Hsiao, C. 1981. "Estimation of Dynamic Models with Error Components," *Journal of the American Statistical Association* (76:375), pp 598-606.
- Angrist, J.D., and Pischke, J.-S. 2009. *Mostly Harmless Econometrics: An Empiricist's Companion*. Princeton, NJ: Princeton University Press.
- Archak, N., Ghose, A., and Ipeirotis, P.G. 2011. "Deriving the Pricing Power of Product Features by Mining Consumer Reviews," *Management Science* (57:8), August 1, 2011, pp 1485-1509.
- Arellano, M., and Bond, S. 1991. "Some Tests of Specification for Panel Data: Monte Carlo Evidence and an Application to Employment Equations," *The Review of Economic Studies* (58:2), pp 277-297.
- Arellano, M., and Bover, O. 1995. "Another Look at the Instrumental Variable Estimation of Error-components Models," *Journal of Econometrics* (68:1), pp 29-51.
- Armbrust, M., Fox, A., Griffith, R., Joseph, A.D., Katz, R., Konwinski, A., Lee, G., Patterson, D., Rabkin, A., Stoica, I., and Zaharia, M. 2010. "A View of Cloud Computing," *Communications of the ACM* (53-58:4), pp 50-58.
- Åstebro, T. 2004. "Sunk Costs and the Depth and Probability of Technology Adoption," *The Journal of Industrial Economics* (52:3), pp 381-399.
- Attewell, P. 1992. "Technology Diffusion and Organizational Learning: The Case of Business Computing," *Organization Science* (3:1), pp 1-19.
- Blackwell, M., Iacus, S., King, G., and Porro, G. 2010. "cem: Coarsened Exact Matching in Stata," *Stata Journal* (9:4), pp 524-546.
- Blundell, R., and Bond, S. 1998. "Initial Conditions and Moment Restrictions in Dynamic Panel Data Models," *Journal of Econometrics* (87:1), pp 115-143.
- Bresnahan, T.F., and Greenstein, S. 1996. "Technical Progress and Co-invention in Computing and in the Uses of Computers," *Brookings Papers on Economic Activity: Microeconomics*, 1996, pp 1-83.
- Chatterjee, D., Grewal, R., and Sambamurthy, V. 2002. "Shaping Pp for E-Commerce: Institutional Enablers of the Organizational Assimilation of Web Technologies," *MIS Quarterly* (26:2), pp 65-89.
- Fichman, R.G., and Kemerer, C.F. 1997. "The Assimilation of Software Process Innovations: An Organizational Learning Perspective," *Management Science* (43:10), pp 1345-1363.
- Forman, C., Goldfarb, A., and Greenstein, S. 2008. "Understanding the Inputs into Innovation: Do Cities Substitute for Internal Firm Resources?," *Journal of Economics & Management Strategy* (17:2), pp 295-316.

- Garcia, D.F., Rodrigo, G., Entrialgo, J., Garcia, J., and Garcia, M. 2008. "Experimental Evaluation of Horizontal and Vertical Scalability of Cluster-based Application Servers for Transactional Workloads," in: *8th International Conference on Applied Informatics and Communications (AIC'08)*. Rhodes, Greece: World Scientific and Engineering Academy and Society (WSEAS), pp. 29-34.
- Ghose, A. 2009. "Internet Exchanges for Used Goods: An Empirical Analysis of Trade Patterns and Adverse Selection," *MIS Quarterly* (33:2), pp 263-291.
- Hamilton, J.D. 1994. *Time Series Analysis*. Princeton: Princeton University Press.
- Hansen, L.P. 1982. "Large Sample Properties of Generalized Method of Moments Estimators," *Econometrica* (50:4), pp 1029-1054.
- Harms, R., and Yamartino, M. 2010. "The Economics of the Cloud," Microsoft, <http://www.microsoft.com/en-us/news/presskits/cloud/docs/the-economics-of-the-cloud.pdf>.
- Ho, D., Imai, K., King, G., and Stuart, E. 2007. "Matching as Nonparametric Preprocessing for Reducing Model Dependence in Parametric Causal Inference," *Political Analysis* (15:3), pp 199-236.
- Hsiao, C. 2003. *Analysis of Panel Data*. Cambridge, UK: Cambridge University Press.
- Imbens, G., and Wooldridge, J.M. 2007. "Control Function and Related Methods," in: *NBER Summer Institute - What's New in Econometrics?* <http://www.nber.org/minicourse3.html>.
- Mell, P., and Grance, T. 2011. "The NIST Definition of Cloud Computing," National Institute of Standards and Technology Information Technology Laboratory (ed.). Gaithersburg, MD.
- Michael, M., Moreira, J.E., Shiloach, D., and Wisniewski, R.W. 2007. "Scale-up x Scale-out: A Case Study using Nutch/Lucene," *Parallel and Distributed Processing Symposium (IPDPS 2007)*. *IEEE International*, pp. 1-8.
- Nickell, S. 1981. "Biases in Dynamic Models with Fixed Effects," *Econometrica* (49:6), pp 1417-1426.
- Papke, L.E., and Wooldridge, J.M. 2008. "Panel data methods for fractional response variables with an application to test pass rates," *Journal of Econometrics* (145:1), pp 121-133.
- Parthasarathy, M., and Bhattacharjee, A. 1998. "Understanding Post-Adoption Behavior in the Context of Online Services," *Information Systems Research* (9:4), pp 362-379.
- recipient, a. 2017. *Oxford English Dictionary Online*. <https://en.oxforddictionaries.com/definition/recipient>.
- Reese, G. 2009. *Cloud Application Architectures: Building Applications and Infrastructure in the Cloud*. O'Reilly Media.

- Retana, G.F., Forman, C., and Wu, D.J. 2016. "Proactive Customer Education, Customer Retention, and Demand for Technology Support: Evidence from a Field Experiment," *Manufacturing and Service Operations Management* (18:1), pp 34-50.
- RightScale. 2016. "RightScale 2016 State of the Cloud Report." Retrieved May 24, 2016, 2016, from <http://www.rightscale.com/blog/cloud-industry-insights/cloud-computing-trends-2016-state-cloud-survey>
- Roodman, D. 2009a. "How to Do xtabond2: An Introduction to Difference and System GMM in Stata," *Stata Journal* (9:1), pp 86-136.
- Roodman, D. 2009b. "A Note on the Theme of Too Many Instruments," *Oxford Bulletin of Economics & Statistics* (71:1), pp 135-158.
- Wooldridge, J.M. 2011. "Fractional Response Models with Endogenous Explanatory Variables and Heterogeneity," in: *Stata Conference Chicago 2011*. Chicago, IL.
- Zhu, K., and Kraemer, K.L. 2005. "Post-Adoption Variations in Usage and Value of E-Business by Organizations: Cross-Country Evidence from the Retail Industry," *Information Systems Research* (16:1), pp 61-84.
- Zhu, K., Kraemer, K.L., and Xu, S. 2006. "The Process of Innovation Assimilation by Firms in Different Countries: A Technology Diffusion Perspective on E-Business," *Management Science* (52:10), pp 1557-1576.

ONLINE APPENDIX FOR

TECHNOLOGY SUPPORT AND POST-ADOPTION IT SERVICE USE: EVIDENCE FROM THE CLOUD

A.Provider Cloud Infrastructure Service and Technology Support Offerings

This appendix offers additional details to those presented in section 0 of the manuscript in relation to the research context and the provider's service characteristics. In our particular setting, the cloud provider has recognized that the novelty of the service plus the complexities involved in deploying distributed architectures that best leverage the cloud's scalability may pose significant knowledge barriers to buyers attempting to use the service. In response to this, the provider offers them the option to contract and receive full support. We discuss first the pricing and terms of the cloud infrastructure service offering, and then elaborate on what characterizes full support.

One of the essential characteristics of cloud infrastructure services is that they are offered on-demand (Mell and Grance 2011). Buyers only pay for what they use, and nothing else: there are no sign-up fees, no minimum spending requirements, no periodical subscription fees and – since buyers can choose not to use their service as well – there are no contract termination penalties either. Moreover, in the particular case of our provider, the computing resources are offered to buyers at fixed hourly rates that increase in server size or capacity, generally in a linear fashion. Servers' capacity is defined in terms of memory (GB of RAM), processing power (number of virtual CPUs), and local storage (GB space of local hard disk). During our observation period, the three capacity metrics tend to vary together as a bundle, meaning that more of one is generally associated with more of the other two, yet prices are set and buyers usually make infrastructure sizing decisions in terms of memory. Prices also vary depending on the operating system chosen for a server (e.g., Windows servers cost more than Linux servers), yet such heterogeneity does not alter our main findings. The results considering operating system heterogeneity are available upon request.

Buyers in our context can launch as many servers and of any size they want, when they want. However, as is discussed in section 2.1, there are important technical challenges in deploying horizontally scalable configurations where several cloud servers work in parallel. These challenges may in turn limit buyers' ability to use many servers at once. Finally, there are no usage caps, with the only exceptions to this being that the provider may have limited hardware installed at its data centers or may take security measures to prevent misuse of its service (e.g., spamming). In other words, for legit buyers, there is no pre-defined cap or limit to how much they can choose to use the service.

The provider complements its infrastructure offering with full support, which is offered for a fixed price premium per server-hour used plus an additional fixed monthly fee. For instance, instead of paying \$0.10 per hour for a 2GB RAM Linux server under basic support, a full support buyer would pay \$0.12 more, i.e., \$0.22 per hour. Similarly, for the 4GB RAM server priced at \$0.20 per hour under basic support, the full support buyer would pay \$0.32 per hour. The monthly fee is paid as a monthly subscription, which is a fee high enough to deter buyers with very low willingness to pay (i.e., bloggers that use a single very small server). There are no sign-up or termination fees for the full support service. The only explicit switching cost from one support level to another is technical rather than monetary: when downgrading from full support to basic support, because of technical limitations in the service offering (during our observation period), buyers must redeploy their servers on their own under the new support regime. The redeployment will involve launching new servers with virgin operating systems (i.e., “out of the box”), and then installing and configuring their business applications on them.

A prime goal of full support is to educate buyers on how to best use the cloud infrastructure service and adapt it to their idiosyncratic business needs. When receiving full support, buyers receive personalized guidance and training, and thus have the opportunity to learn directly from the provider’s prior experience in deploying applications in the cloud. Buyers not willing to pay the price premiums will only receive a basic level of support that has limited scope in the sense that it is intended to aid buyers with issues concerning account management or overall performance of the infrastructure service. For example, while a full support buyer may be personally guided step by step on how to deploy a web server through phone conversations, live chat sessions or support tickets, basic support buyers will be referred to a knowledge base. Similarly, if a server failed, which happens much more frequently than in traditional datacenter settings given the commodity hardware employed and the multi-tenant architecture (i.e., multiple organizations’ virtual servers are hosted in the same and shared physical server), the provider would work together with full support buyers in solving the issues, while basic support users would only be notified about the failure, if anything. Thus, basic support buyers do not have fluid access to external knowledge from the provider and have to rely mostly on their internal capabilities to make use of the service.

B. Description of CEM Procedure and Sample Construction

B.1. Overview of CEM Procedure

We run our models on subsamples defined using a coarsened exact matching (CEM) procedure (Blackwell et al. 2010; Iacus et al. 2012). For matching purposes, we consider buyers who adopted full support at any point in their tenure as treated and those that relied exclusively on basic support as controls. Matching reduces endogeneity concerns (Ho et al. 2007), and CEM has been used extensively in recent work to improve the identification of appropriate control groups in difference-in-differences estimation (e.g., Azoulay et al. 2011; Azoulay et al. 2010; Furman et al. 2012).

CEM is particularly convenient for our setting because it is a nonparametric procedure that does not require the estimation of propensity scores. This is useful because, aside from the exogenous failures, we have limited data that would allow us to directly predict the likelihood of full support. Each unique vector formed by combinations of the coarsened covariates describes a stratum. Since the number of treated and control observations in each strata may be different, observations are weighted according to the size of their strata (Iacus et al. 2012). The differences in means between the treated and the controls across the various matching variables are almost all statistically significant. However, once we apply the CEM weights the samples are perfectly balanced and any mean differences are eliminated. All our regressions with the CEM-based sample employ these weights. When exact matching is possible, such that for every treated observation there is a control observation identical to the first one across all possible covariates except for the treatment, a simple difference in means of the dependent variables would provide an estimate of the causal effect of interest. Nonetheless, since it is nearly impossible to use exact matching in observational data and thus there is always a concern about the influence of omitted variables, we continue using our fixed effects panel data model to control for them.

We match buyers based on six attributes: (1) pre-upgrade volume of IT use (i.e., memory use), (2) pre-upgrade efficiency of IT use (i.e., fraction of server running in parallel), (3) pre-upgrade frequency of cloud infrastructure resizing (i.e., how often buyers launch a server, halt a server, or resize an existing one), (4) operating system of preference, (5) employment, (6) and intended use case for the cloud infrastructure service. The first four attributes are derived directly from firms' observed usage of the cloud service. The latter two attributes come from an optional sign-up survey.

The survey is optional and administered as part of the online sign-up web form; the response rate is 43.4%, and we have not found systematic differences between respondents and non-respondents. The survey was first administered in June 2010, and we have all buyers' responses until February 2012. Although there can only be one survey response per account, since buyers can have multiple accounts, we may also have multiple responses per buyer. In our data we have 6,152 survey responses from 5,565 different buyers in the baseline sample, 431 of which

changed their response to at least one item across their surveys. However, for 42.3% of the buyers with varying responses the time gap between the survey responses is too short (i.e., less than 3 months) as to suggest that the variance is due to changes in firms' sizes or goals. Given this, we do not rely on variance across responses for our analysis and rather only consider the 5,134 buyers that either have a single survey response or that have consistent responses across all their submissions. Further, we have not considered firm attributes in the survey as controls in our models since they do not vary over time and thus would be absorbed by the firm fixed effect.

For the matching process, we only consider treated buyers who started using the cloud service with basic support and upgraded to full support later on. This allows us to match the upgraders to controls based on their usage behavior before they adopted full support, had the controls adopted full support in the same month of their tenure. This approach, which is similar to the one implemented by Azoulay et al. (2010) and Singh and Agrawal (2011), ensures to the extent possible that treated firms do not exhibit differential usage behavior before they adopt full support relative to controls. Among the 5,134 buyers for which we have all this data (i.e., they answered the sign-up survey), 1,259 are treated and 3,875 are potential controls. Using the six criteria described above, we develop a weighted matched subsample. As part of our research we ran our models with varying permutations of matching criteria which, in addition to the six already mentioned, included buyer industry. Our results were consistent across the various subsamples and are available upon request.

B.2. CEM Matching Criteria

Six different attributes of firms were used to match treated and controls. In this section we describe each of them. They are summarized in Table B.1.

Table B.1. Summary of Matching Criteria used in CEM Procedure			
Name	Description	# of Categories	Categories
Memory Use	Memory usage in GB of RAM in months before upgrade	9	<0.5, 0.5-1, 1-2, 2-4, 4-8, 8-16, 16-32, 32-64, >64
Architecture Complexity	Fraction of servers running in parallel in months before upgrade	5	0.00, 0.00-.25, .25-.50, .50-.75, >.75
Adjustments	Frequency of infrastructure resizing in months before upgrade	5	0, 1-2, 3-9, 10-43, >43
OS Preference	Primary OS before upgrade	6	Linux, Windows, RedHat, SQL, Mix
Employment	Employment	5	0-10, 11-50, 51-100, 101-250, >250
Use Case	General use cases (can have more than 1)	5	High variance, low variance, back office, hosting, test & development

IT Use, Architecture Complexity and Frequency of Infrastructure Sizing Adjustments: In

regards to overall use (i.e., memory use) and frequency of infrastructure resizing, when creating our baseline sample we had already discarded basic support users with very small and/or rather static deployments over the early periods of their tenure. We excluded buyers who averaged 512 MB RAM/hour or less during their first 6 months (excluding 1st month) or made no adjustments to size of their infrastructure during their first 6 months (excluding 1st month). Nonetheless, even among the remaining buyers there is considerable variation in these two variables.

The average memory usage, fraction of servers running in parallel (as a proxy for the architecture complexity of the deployment), and the frequency of infrastructure resizing actions used to match treated and controls were computed as follows. Assume that a given treated buyer adopted the service with basic support in some period t_0 and switched from basic to full support in a later time period, t_{fs} , $t_{fs} > t_0$. Then, we consider the set of controls (i.e., buyers who exclusively used basic support) who also adopted the service in month t_0 and used the service (i.e., did not churn) at least up to t_{fs} . This ensures all buyers were using the service during the same calendar time frame and have very similar tenure by period t_{fs} . For the treated group and all these controls, we compute the average memory usage, fraction of servers running in parallel, and frequency of scaling actions in the periods during which all buyers were using basic support: from t_0 up to $t_{fs}-1$. Finally, we use this metric, which represents their pre-upgrade behavior, to match buyers.

For average memory usage, we set our cutoff points at standard server sizes: 512MB, 1GB, 2GB, 4GB, 8GB, 16GB, 32GB and 64GB of RAM. For the fraction of servers running in parallel, we opted to create 5 bins. Its distribution has a strong mass at zero, which justified the first bin with all observations with zero value. Afterwards we built 0.25-width bins which are close to the 25th, 50th and 75th percentiles of the non-zero values. For frequency of infrastructure resizing, we base our cutoff points on percentiles of the distribution: the 25th percentile is a single change to the size of the deployment, the 50th percentile is 3 changes, the 75th percentile is 9 changes, and the 95th percentile is 43 changes. In total, as shown in Table B.1, we have 9 categories for memory use, 5 categories for fraction of servers running in parallel, and 5 categories for the frequency of infrastructure resizing to match on.

Operating System (OS) Preference: During the time span of our data, the provider offered its servers running 4 different OS and we observe which OS each individual server used:

1. Linux: Several distributions, although we do not observe which.
2. Windows: Several versions of Windows Server, although we do not observe which.
3. Red Hat Enterprise Linux.
4. SQL Server: This is really a Windows Server running SQL Server, yet it was offered under its own price scheme and hence is considered another operating system for this exercise.

Even though there were multiple OS available, as we show in Table B.2, most buyers either exclusively or at least primarily used a single OS. To determine if a buyer is a user of a particular

OS, we computed the proportion of the total amount of GB RAM-hours consumed by each buyer over its observed tenure that were consumed of each of the 4 different OS. Then, using arbitrary yet high thresholds (e.g., from 85% up to 100%), we flag a customer as user of a certain OS if the proportion of service use with that OS is greater or equal than the defined threshold. Using these proportions of workloads under each OS and varying thresholds, we populated each column of Table B.2 as follows:

- Threshold: Indicates the percentage of total usage using a specific OS used to flag a buyer as a user of that OS.
- Linux, Windows, Red Hat and SQL: Indicate the proportion of buyers who used at least as much as the threshold of their total usage under each corresponding OS. For example, 57.35% of buyers in the baseline sample used at least 99% of all their GB RAM-hours on Linux servers.
- Only 1: Given a certain threshold, it shows the proportion of total buyers that used only a single OS. The column is the sum of the 4 different OS columns to the left.
- Mixed: Given a certain threshold, it shows the proportion of total buyers that used a mix of more than a single OS. The “Only 1” column and this column add up to 100%.

The main takeaway from Table B.2, and in particular from the “Only 1” column, is that most customers primarily use a single OS. For instance, 66.96% of buyers in the baseline sample ran all their servers using a single OS, and 80.13% ran at least 95% of their workloads using a single OS.

Table B.2. Proportion of Buyers Using each OS under Different Thresholds						
Threshold	Proportion of Buyers using OS					
	Linux	Windows	Red Hat	SQL	Only 1	Mixed
100%	52.38%	10.47%	2.31%	2.13%	66.96%	33.04%
99%	57.35%	13.29%	2.81%	2.36%	75.49%	24.51%
95%	60.07%	14.78%	3.09%	2.52%	80.13%	19.87%
90%	61.60%	15.89%	3.24%	2.65%	83.05%	16.95%
85%	62.79%	16.78%	3.42%	2.83%	85.49%	14.51%

To put this proportion into perspective, recall the median buyer in the sample consumes an average of 0.5 GB RAM (or 512 MB RAM) per hour over its tenure. Thus, over a month, a median buyer consumes $0.5 \text{ GB RAM/h} \times 24 \text{ h/day} \times 30 \text{ days/month} = 360 \text{ GB RAM/month}$. If a buyer uses the same OS for at least 95% of its workload, then it is using some other OS for at most 18 GB RAM during a month. This level of usage is equivalent to running a very small, 256 MB RAM (0.25 GB RAM) server for 3 days of the month (i.e., 72h). Even for a threshold of 85%, the remaining 15% is 54 GB RAM during the month, or 9 days of a very small 256 MB RAM server.

We feel that such levels of usage (e.g., a very small server during 9 days per month) are inconsistent with running production applications, even if they are only used for short time spans.

Even a small blog would run in that 256 MB RAM server but for an entire month (i.e., 30 days), not 9 days, and any standard application will traditionally at least need a 512 MB RAM server, twice as large as this one. In other words, we are confident that customers who use at least 85% or more of their workloads on a single OS can be characterized as users of that OS. We find this is the case for 85.49% of the buyers in the baseline sample and 89.44% of the buyers in the CEM subsample described below. For our matching process, we employ the 85% threshold to flag buyers as users of each of the 4 different OS or a mix of any of the 4, resulting in 5 categories.

Employment: The employment and intended use case data are collected from the sign-up survey. The proportions of buyers falling into the relevant categories within these two attributes are shown in Table B.3. For the employment cutoff points, we broadly rely on the ranges used in the survey. Among the buyers with consistent survey responses across all their accounts, 65.60% indicated they have 10 or fewer employees. Another 19.75% indicated they have between 11 and 50 employees. We subdivide the remaining 15% of buyers in three bins each accounting for roughly 5% of our sample: from 51 to 100, from 101 to 250, and more than 250.

Table B.3. Proportion of Buyers per Category			
Buyer Role	All Buyers	Controls	Treated
Number of Buyers	5,134	3,875	1,259
Employment			
10 or less	65.60%	69.19%	54.57%
11 to 50	19.75%	18.68%	23.03%
51 to 100	5.03%	4.36%	7.07%
101 to 250	3.66%	3.02%	5.64%
More than 250	5.96%	4.75%	9.69%
Use Case			
High Use Uncertainty	46.34%	46.86%	44.72%
Low Use Uncertainty	59.14%	57.34%	64.65%
Back Office Applications	18.85%	19.48%	16.92%
Hosting	9.17%	9.29%	8.82%
Test & Development	29.31%	32.26%	20.25%

Intended Use Case: The intended use case is collected by a multiple choice question (i.e., “Mark all that apply”) that asked buyers to “Please indicate what solution(s) you intend to use [the cloud infrastructure service] for.” The 20 options available to buyers are very specific, and finding matches across such specific use cases would be extremely hard. Instead, we group the specific use cases into 3 more general use cases based on two dimensions: if the use case is related to back office or front office applications, and, in the latter case, if it is likely that the volume of usage for the use case is predictable or not. Our first general use case, which we call “High Use Uncertainty”, includes customer-facing websites that are prone to unpredictable variance in their volume of usage. Examples of such use cases are social media sites, online gaming sites, online publishing sites, rich media sites (e.g., audio or video), and other Software-as-a-Service (SaaS) offerings. Our second general use case, “Low Use Uncertainty”, includes customer-facing websites used for

regular operation of the firm that have steady or at least predictable use levels. Examples are corporate websites, collaboration platforms, online portals, and e-commerce sites. We chose to include e-commerce sites in this general use case since, although it may have a high variance, seasonality makes the peaks and valleys of the demand fairly predictable. Finally, our “Back Office Applications” general use case includes applications or systems used internally for business operations. Examples are a company’s intranet and systems used for accounting, customer relationship management, human resources, supply chain management, or backup. We additionally consider web hosting services and running test and development environments as additional general use cases. Altogether, we have 5 general use cases.

B.3. CEM Sample

To construct our CEM sample we started off with our entire dataset which has a total 79,619 different buyers. However, as noted in the main text we do not employ all buyers in our baseline sample. For the baseline sample we have excluded buyers who (1) only received basic support and (2) averaged 512 MB RAM/hour or less during their first 6 months (excluding 1st month) or (3) made no adjustments to size of their infrastructure during their first 6 months (excluding 1st month). We do not consider their behavior during their 1st month in our threshold because most buyers are setting up their infrastructure during this time. The baseline sample has 22,179 buyers.

From the baseline sample, we can only include in our CEM sample those buyers for which we have a survey response, which are 5,134, and for whom we observe at least two months of tenure. This is so that we can match buyers based on their behavior in the period before upgrading from basic to full support. This leaves us 4,200 buyers for the matching process.

The CEM procedure leaves in our sample 1,525 buyers, of which 1,303 are controls who exclusively used basic support and 222 are treated buyers who started with basic support and upgraded to full support. There are on average 5.86 control buyers per treated buyer.

C. Support Interactions and Construction of Instruments

The content of the support interactions between the provider and its buyers was used to identify three types of exogenous failures experienced by buyers. The descriptions of the failure events can be found in Table 3 in the manuscript. The following are the keywords and phrases used to identify each of these types of support interactions. All support interactions that matched some keyword or phrase were visually examined to rule out false positives.

Table C.1. Keywords and Phrases Searched for Support Interactions Coding		
Failure Type	Variable Name	List of keywords or phrases used for identification of failure
Service outage	<i>FailOutage</i>	Providers' service status URL, cloud status, outage, scheduled maintenance, undergoing maintenance
Network-related failure	<i>FailNetwork</i>	Server does not respond to ARP requests, faulty switch, network issue in our data center, lb in error state, load-balancer hardware nodes, DDoS
Hardware-related failure	<i>FailHardware</i>	Hardware failure, degraded hardware, drive failing, drives failing, server outage, host failure, server is down, server down, is hosted on has become unresponsive, problem with our server, host server, physical host, physical hardware, physical machine, host machine, failing hardware, hardware failure, imminent hardware issues, migrate your cloud server to another host, queued for move, issue on the migrations, host server of your cloud servers

Once we identified the occurrence of the failures through the coding process, we calculated the accumulated number of them occurring over time for each buyer and each failure type. Letting $F \in \{FailOutage, FailNetwork, FailHardware\}$ represent a type of support interaction, whenever the count reached N incidents we turn the corresponding FN indicator on. For example, variable $FailOutage2_{i,t}$ will be equal to 1 if buyer i has accumulated at least 2 support interactions that have been coded as type *FailOutage* by month t .

D. Impulse Response Functions

An impulse response function represents the response of a dependent variable to a (one-time) unit change in some covariate while all other variables dated t or earlier are held constant (Hamilton 1994, pp. 318-323). In our case, we compute and plot difference quotients $\frac{\Delta \hat{y}_{i,t}}{\Delta x_{i,t-j}}$ ($y_{i,t} \in \{lnMemory_{i,t}, FractionParallel_{i,t}\}$, $x_{i,t} \in \{AdoptFull_{i,t}, CeaseFull_{i,t}\}$) over time to show how current memory usage or current fraction of servers running in parallel is influenced by adoption or dropping of full support j periods ago.

These difference quotients are the coefficients of the associated rational lag model (Greene 2008, pp. 683-686). The rational lags identify the effect that each lag of the covariate, on its own, has on the dependent variable. Since we used two lags ($p = 2$) of the dependent variables in our estimations of Model (2), the early lags (i.e., $j \leq 1$) of $AdoptFull_{i,t-j}$ or $DropFull_{i,t-j}$ have a specific formulation, while the later lags (i.e., $j \geq 2$) follow a recursive form. The approach is very similar to that of example 20.4 in Greene (2008, pp. 685-686). The rational lag coefficients, which we denote $\hat{\delta}_j$, are computed as follows (we show $\hat{\beta}_j$ coefficients as those multiplying the $AdoptFull_{i,t}$ lags, but algebra is identical for the $\hat{\gamma}_k$ coefficients multiplying lags of $DropFull_{i,t}$):

$$\hat{\delta}_0 = \frac{\Delta \hat{y}_t}{\Delta x_t} = \hat{\beta}_0$$

$$\hat{\delta}_1 = \frac{\Delta \hat{y}_t}{\Delta x_{t-1}} = \hat{\beta}_1 + \hat{\lambda}_1 \hat{\delta}_0$$

$$\hat{\delta}_2 = \frac{\Delta \hat{y}_t}{\Delta x_{t-2}} = \hat{\beta}_2 + \hat{\lambda}_1 \hat{\delta}_1 + \hat{\lambda}_2 \hat{\delta}_0$$

$$\hat{\delta}_j = \frac{\Delta \hat{y}_t}{\Delta x_{t-j}} = \hat{\beta}_j + \hat{\lambda}_1 \hat{\delta}_{j-1} + \hat{\lambda}_2 \hat{\delta}_{j-2}, \quad 2 \leq j \leq r.$$

As an additional step, we use Monte Carlo simulation and draw 100,000 random samples of the vectors of the fitted $\hat{\beta}_j$ and $\hat{\lambda}_s$ coefficients using their estimates and their variance-covariance matrix. For each draw we compute and record the fitted rational lags $\hat{\delta}_j$, and use their distributions to estimate their 90% confidence interval. In Figure 1, the dashed lines represent the 5th and 95th percentiles of the distribution of each $\hat{\delta}_j$.

References

- Azoulay, P., Graff Zivin, J.S., and Sampat, B.N. 2011. "The Difusion of Scientife Knowledge Across Time and Space:Evidence from Professional Transitions for the Superstars of Medicine." National Bureau of Economic Research.
- Azoulay, P., Graff Zivin, J.S., and Wang, J. 2010. "Superstar Extinction," Quarterly Journal of Economics (25), pp 549-589.
- Blackwell, M., Iacus, S., King, G., and Porro, G. 2010. "cem: Coarsened Exact Matching in Stata," Stata Journal (9:4), pp 524-546.
- Furman, J.L., Jensen, K., and Murray, F. 2012. "Governing Knowledge in the Scientific Community: Exploring the Role of Retractions in Biomedicine," Research Policy (41), pp 276-290.
- Ho, D., Imai, K., King, G., and Stuart, E. 2007. "Matching as Nonparametric Preprocessing for Reducing Model Dependence in Parametric Causal Inference," Political Analysis (15:3), pp 199-236.
- Iacus, S.M., King, G., and Porro, G. 2012. "Causal Inference without Balance Checking: Coarsened Exact Matching," Political analysis (20:1), pp 1-24.
- Singh, J., and Agrawal, A. 2011. "Recruiting for Ideas: How Firms Exploit the Prior Inventions of New Hires," Management Science (57:1), pp 129-150.