

Enhancing Cloud Security and Privacy: The Unikernel Solution

Alfred Bratterud
Dept. of Computer Science
Oslo and Akershus University
Oslo, Norway
Email: alfred.bratterud@hioa.no

Andreas Happe
Dept. Digital Safety & Security
Austrian Inst. of Tech. GmbH
Vienna, Austria
Email: andreas.happe@ait.ac.at

Bob Duncan
Computing Science
University of Aberdeen
Aberdeen, UK
Email: bobduncan@abdn.ac.uk

Abstract—In previous work, we proposed how a new approach using unikernels might be appropriate to help resolve the problem of cloud security and privacy, in which we identified 10 key management problems which needed to be addressed. In this paper, we outline the technical details of such an approach, identifying how this new approach can better address the issues involved.

Index Terms—Cloud security and privacy; management control; compliance; complexity

I. INTRODUCTION

In our previous paper in this series [1], we provided a high level account of ten security issues which management need to take account of when using cloud computing systems, and suggested how unikernel-based systems might address many of those issues, as we see in TABLE I below.

Issue	Description	Helped by:
1	The definition of security goals	Mgt and Unikernels
2	Compliance with standards	Mgt and Unikernels
3	Audit issues	Mgt and Unikernels
4	Management approach	Mgt and Unikernels
5	Technical complexity of cloud	Unikernels
6	Lack of responsibility and accountability	Mgt/CSPs
7	Measurement and monitoring	Unikernels
8	Management attitude to security	Mgt
9	Security culture in the company	Mgt
10	The threat environment	Unikernels can help to enforce good design / architecture decisions

TABLE I: Items addressed by Unikernels ©2016 [1]

While these security issues can be successfully addressed by other means, the reality, as evidenced by the recurring success of attackers, is that many companies are failing to apply the necessary rigour needed to resolve these issues in their existing approaches. Year after year, many attacks which are both simple and relatively inexpensive to defend against, continue to be exploited.

In this paper we introduce a framework of definitions and metrics for classifying unikernel systems which we later make use of in the design, testing and assessment of new unikernel based system architectures. In Section II, we outline some necessary definitions and preliminary observations on unikernels, and in Section III we consider 6 security observations relating to unikernels. The remainder of the paper is organized as follows: in Section IV, we review the relationship between unikernels and microkernels; in Section V, we discuss the implications of implementation language choice; in Section VI, we present well defined properties of unikernel systems; in Section VII, we outline our proposed solution. In Section VIII, we show how our proposed approach will address those key issues identified in TABLE: I. In Section IX, we consider some initial thoughts on attack vectors; and in Section X, we discuss our conclusions.

II. A PRECURSOR TO FORMAL WORK ON UNIKERNELS

Modern unikernel research, particularly concerning deployment in cloud, is still in its infancy and the literature currently available does not include much in the way of theoretical work or precise definitions, but rather takes a pragmatic approach [2]–[4]. One popularized definition states that “*Unikernels are specialised, single-address-space machine images constructed by using library operating systems*”; and that *Unikernels provide many benefits compared to a traditional OS, including improved security, smaller footprints, more optimisation and faster boot times*” [5]. Without further qualification, directly associating unikernels with security benefits can be hazardous. In particular, it is not at all clear what “machine images” or “operating system libraries” are, nor what it means to construct them, and hence it is unclear precisely how unikernels can be said to be more secure. It is our view that stricter definitions are both necessary and achievable. We propose a set of working definitions intended to make the following exposition

more precise, and to serve as a theoretical basis of a framework for unikernel based cloud computing.

To this end we can go back over four decades to early work on virtualization of mainframes. Early single address space operating systems such as Mungi or Opal and many others do not seem to have much traction today. In 1995 [6], there was some early work on the Exokernel system, with some updating over the years, but little widespread use. Microsoft work on library operating systems, Drawbridge [7], saw little use other than for research, although it later evolved into the Haven system [8], which was intended for use in the Azure cloud.

Definition II.1. Popek-Goldberg virtual machine. An environment created by the virtual machine monitor, which is functionally equivalent to the physical machine on a given hardware platform, as defined in [9]

Popek and Goldberg provide formal definitions of both Virtual Machine Monitor and Virtual Machine, which form the most well known, if not the only, formally defined virtualization platform which also directly corresponds to modern hardware and nomenclature. An x86 Popek-Goldberg virtual machine is a virtual machine running under x86 Popek-Goldberg compliant hardware virtualization, e.g. x86 hardware with *vt-x* extensions.

In order to give a precise definition of unikernels, we need to define their constituent parts. Our intention is not to provide definitions that encompass all the complexities or functionalities of unikernels, but rather the opposite; just enough to get a precise idea of what unikernels are and what they are not.

Definition II.2. Compiled program object. Symbol. An n -tuple of machine instructions from a Turing complete instruction set, e.g. Intel x86, or an arbitrary sequence of bytes b , i.e. $0 \leq b < 2^8$.

Definition II.3. Software library, symbol. A *Software library* is taken to be a collection of compiled program objects, $\{O_1, \dots, O_m\}$ and arbitrary byte-sequences (e.g. data) $\{D_{m+1}, \dots, D_n\}$ providing symbol resolutions, i.e. linkable objects, each corresponding to a *symbol* in the set $\{S_1, \dots, S_m, S_{m+1}, \dots, S_n\}$

The intuition here is to capture the idea of a library of compiled code, where functions and data are represented as binary objects in e.g. `libos.so`, `libos.a`, `libos.dll` etc., where functions and static data can be accessed via symbols available to a linker. For the purposes of this paper we assume the process of linking and symbol resolution are given and well defined.

Definition II.4. Software service. Objects, and optionally symbols, from a software library, with the addition of an entry point object O_0 corresponding to a symbol S_0 (e.g. `_start`) that may or may not be present in the service, providing the compiled code necessary for a program to be executable on a given hardware architecture.

Compiling an executable binary typically includes defining the entry point, e.g. `main` in ISO C, from where to start executing functions provided from a library. Compilers such as `gcc` will typically pre-pend some additional functionality through e.g. the `_start`-symbol, mapped to code for initializing the C-runtime, including calling global constructors, zero-initializing the `.bss`-segment etc. The symbols are required during link-time but can later be stripped out when they are mapped to memory addresses relative to the binary.

We can now give an operational definition of a library operating system. The term has held different meanings [2], [10] and we do not intend the following to be canonical, but merely one that will suffice to precisely define a unikernel for the purposes of our framework.

Definition II.5. Library operating system. For a software service SW where the objects $\{O_0, \dots, O_n\}$ form the set of objects necessary and sufficient for SW to run on a given hardware platform, a library operating system is a software library that can provide $\{O_0, \dots, O_n\}$

It is implied in this definition that a library operating system provides all the objects necessary to form a fully functional program, independent of any other software present on a system, except that which may or may not be presented through an instruction level interface, e.g. software that responds to a *trap* on the *Virtual Machine Monitor* in the Popek-Goldberg model.

1) *Definition of a unikernel:* In the context of virtual machines and cloud computing, it makes sense to describe the whole virtual machine as a unikernel [2], as there is in fact no classical boundary between kernel- and user space, and also because any combination of objects that can be pulled from the library operating system individually can be combined with a piece of software to form a unique whole. This piece of completely linked software will have full access to hardware on the same level as a classic kernel.

In the context of classical operating system kernels, however, the library operating system designed to produce unikernels may also be called a unikernel [11] in reference to “microkernel”, “nanokernel”, “monolithic kernel” etc. It could be argued that if all the contents of a virtual machine were to be considered a unikernel, there wouldn’t really be any point in using the word “kernel”.

The following definition is intended to be sufficiently flexible to allow both interpretations.

Definition II.6. Unikernel. Given a library operating system OS, a unikernel U is defined as $U \subseteq OS$ such that U is sufficient and necessary to provide complete linkage to some service S for a given hardware platform.

Using an inclusive subset allows both the whole library operating system and any subset to be called a unikernel.

Definition II.7. Unikernel machine image A software service $SW \cup U$ where U is the unikernel for SW , they both share the same address space, and with the addition of any facilities

necessary to start *SW* on a given well-defined virtualization platform e.g. a bootloader in the case of an x86 Popek-Goldberg virtual machine.

Definition II.8. Popek-Goldberg unikernel. A Popek-Goldberg virtual machine initialized with a unikernel machine image.

III. SIX SECURITY OBSERVATIONS IN UNIKERNEL-BASED SYSTEMS

We have identified 6 security observations which are exhibited by unikernel systems:

- Choice of service isolation mechanism;
- The concept of reduced attack surface;
- The use of a single address space, shared between service and kernel;
- No shell by default, and the impact on debugging and forensics;
- Micro services architecture and immutable infrastructure;
- Single thread by default.

A. Choice of service isolation mechanism

In the previous paper in this series, the argument was made for why classical virtualization is the preferred platform for secure cloud computing. While many alternatives exist, which are both practical and widely trusted, one cannot reason precisely about their security properties unless they are well defined. In this paper we make no judgements about their usefulness, but merely note that classical virtualization has had a precise foundation since 1974. We believe that the lack of similar models for other modes of virtualization is due to the fact that Popek-Goldberg virtualization exists at the instruction level, which is necessarily simple in nature as it must be implemented in physical circuitry. Other approaches typically rely on higher level software interfaces, and are thus harder to define precisely. Despite the simplicity of C, it still proves a hard nut to crack for the purposes of formal verification [12].

B. Reduced attack surface

Using the above definitions we can now define the attack surface of a system as the sum of all objects, in bytes.

Definition III.1. Attack surface. The number of bytes in a system, physically available for reading, writing or executing as instructions for a given hardware architecture.

Physical protection can be seen as a gray area when it comes to microcode, firmware and otherwise mutable hardware such as e.g. field-programmable gate arrays (FPGAs). This definition is intentionally kept general in order to allow further specifications to refine the meaning of “physically available” for a given context. The following example can serve to illustrate how the definition can be used for one of many purposes.

Building a classic VM using Linux implies simply installing Linux, and then installing the software on top. Any reduction in attack surface must be done by removing unneeded software and kernel modules (e.g. drivers). Take TinyCore Linux as an

example of a minimal Linux distribution and assume that it can produce a machine image of 24MB in size.

Given this intuition, let L be a collection of compiled program objects for x86, such that $L = \{O_1, \dots, O_{2400}\}$, i.e. all the objects provided by TinyCore Linux, totalling 24MB in size, and for simplicity assume the objects are uniformly sized, 1Kb each. Adding a 1MB software service *SW* which require $\{O_0, \dots, O_{1000}\}$ to be executable, we get an attack surface of 25MB, regardless of how many objects of L were actually needed by *SW*. Assuming a *library operating system* existed that could provide $\{O_0, \dots, O_{1000}\}$, a *unikernel* would by definition II.6, provide exactly those objects, forming a 2MB sized unikernel machine image. (2MB + 512 bytes of bootloader code in an x86 hw-VM). Hence the attack surface of the unikernel VM is reduced by 92%. Conversely, the Linux VM could be said to add $23/1 * 100 = 2300\%$ of unnecessary code, which we will refer to as *bloat* or *increased attack surface* depending on context.

C. The use of a single address space

The main objective for postulating a single address space is to imply single process or singular purpose. In a classic kernel, the need for multiple address spaces is prompted by the need to run multiple processes, which must be kept separate to ensure consistency and integrity among them. A classic kernel will typically rely on virtual memory implemented in hardware to ensure process isolation and to provide a controlled virtual to physical address translation. Popek-Goldberg virtualization relies directly on this general concept without further extension. Aside from the performance degradation often seen in nested address translation, we take the view that introducing virtual memory and multiple processes inside a Popek-Goldberg virtual machine needlessly complicates an already complex system. In particular it lays upon the virtual machine the added responsibility of creating and maintaining a process-kernel boundary.

Definition III.2. Single address space. For a computer system, a *single address space* is defined as an interval of positive integers $[a_0, \dots, a_n]$ where n is a power of two, representing the total addressable memory of a system in a given state, such that dereferencing any address $*a_x$ from anywhere inside the system would access the same physical memory cell.

The intuition is that virtual memory is not employed inside the system, effectively eliminating the possibility of running several disjoint processes. We are not making any assumptions or requirements as to whether or not all addresses are in fact accessible, e.g. physically present or readable / writable / executable, merely that they point to the same location if any. Note that virtual memory can and will be employed on the virtual machine monitor, to protect one VM from another, but further nesting of virtual memory would violate the single address space principle.

D. No shell by default and the impact on debugging and forensics

One feature of unikernels that immediately makes it seem very different from classical operating systems is the lack of a command line interface. This is however a direct consequence of the fact that classical POSIX-like CLI's are run as a separate process (e.g. `bash`) with the main purpose of starting other processes. Critics might argue that this makes unikernels harder to manage and “debug”, as one cannot “log in and see what’s happened” after an incident, as is the norm for system administrators. We take the position that this line of argument is vacuous; running a unikernel rather corresponds to running a single process with better isolation, and in principle there is no more need to log in to a unikernel than there is to log in to e.g. a web server process running in a classical operating system.

It is worth noting that while unikernels by definition are a single address space virtual machine, with no concept of classical processes, a text based CLI can easily be provided (e.g. IncludeOS does provide an example) — the commands just wouldn’t start processes, but rather call functions inside the program. From a security perspective we take the view that this kind of ad-hoc access to program objects should be avoided. While symbols are very useful for providing a stack trace after a crash or for performance profiling, stripping out symbols pointing to program objects inside a unikernel would make it much harder for an attacker to find and execute functions for malicious and unintended purposes. Our recommendation is that this should be the default mode for unikernels in production mode.

We take the view that logging is of critical importance for all systems, in order to provide a proper audit trail. Unikernels however simply need to provide the logs through other means, such as over a virtual serial port, or ideally over a secure networking connection to a trusted audit trail store.

Lastly it is worth mentioning that unikernels in principle have full control over a contiguous range of memory. Combined with the fact that a crashed VM by default will “stay alive” as a process from the VMM perspective, and not be terminated, this means that in principle the memory contents of a unikernel could be accessed and inspected from the VMM after the fact, if desired. Placing the audit trail logs in a contiguous range of memory could then make it possible to extract those logs also after a failure in the network connection or other I/O device normally used for transmitting the data. Note that this kind of inspection requires complete trust between the owner of the VM and the VMM (e.g. the cloud tenant and cloud provider). Our recommendation would be not to rely on this kind of functionality in public clouds, unless all sensitive data inside the VM is encrypted and can be extracted and sent to the tenant without decrypting it.

E. Micro services architecture and immutable infrastructure.

Micro services is a relatively new term founded on the idea of separating a system into several individual and fully disjoint services, rather than continuously adding features and

capabilities to an ever growing monolithic program. Being single threaded by default unikernels naturally imply this kind of architecture; any need for scaling up beyond the capabilities of a single CPU should be done by spawning new instances. While classical VM’s require a lot of resources and impose a lot of overhead, minimal virtual machines are very lightweight. As demonstrated in [13] more than 100,000 instances could be booted on a single physical server and [11] showed that each virtual machine, including the surrounding process require much less memory than a single “Hello World” Java program running directly on the host.

An important feature of unikernels in the context of micro services is that each unikernel VM is fully self contained. This also make them immune to breaches in other parts of the service composition, increasing the resilience of the system as a whole. Add to this the idea of *optimal mutability* (defined below), and each unikernel-based micro service can in turn be as immutable as is physically possible on a given platform. In the next paper in this series we expand upon these ideas and take the position that composing a service out of several micro services, each as immutable as possible, enables overall system architects and decision makers to focus on a high level view of service composition, not having to worry too much about the security of their constituent parts. We take the view that this kind of separation of concerns is necessary in order to achieve scalable yet secure cloud services.

F. Single threaded by default

While the above definitions do not impose any restrictions on whether or not a unikernel can run several concurrent threads or multiple CPU cores, it is well known that concurrency is a major source of errors accounting for a significant number of vulnerabilities. IncludeOS and MirageOS are both examples of unikernels that are single threaded by default. Efficiency is achieved by event based asynchronous interfaces with no blocking calls. While pre-emptive interrupt handling and concurrency using shared memory are necessary for certain workloads, we take the view that single threaded concurrency free services are by nature less complex and thus less error prone. It is also well known that threaded applications perform worse inside virtual machines than single threaded applications due to the extra layer of context switches necessary to schedule threads inside the VM as well as outside.

Our recommendation is to keep unikernels single threaded by default and rather achieve concurrency by adding more instances, to the extent possible. In a modular library OS one can add threading and re-entrant versions of libraries as optional components without causing bloat or increased complexity to unikernels not requiring concurrency.

IV. RELATIONSHIP TO MICROKERNELS

While there exists a rich fauna of operating system kernel types, the most well known distinction is between monolithic kernels and micro kernels. For this reason we’ll briefly explain how unikernels fit in this spectrum. Microkernel operating systems are absolutely minimal in the sense that nothing

that doesn't have to be in the kernel is. However, most implementations such as the L4 are still A) multi-process; B) not library operating systems; and C) will typically have a classical style command line etc., which would make it almost orthogonal to our purpose as it addresses other issues (L4 is focussed mainly on fast IPC). That being said, they have an advantage over classic kernels when it comes to: A) attack surface (it can run many programs, but it does not have to); and B) complexity. The simplicity of the microkernel is what made it possible to do formal verification of the Haskell implementation of L4 - and that is a major security benefit.

Our position is that unikernels have the potential to incorporate the "small and simple" from microkernels, while still adding new security features — in particular: 1) The library operating system approach, which guarantees a minimal amount of unnecessary code is introduced; 2) the single-purpose approach; and 3) it is single-threaded by default. This provides a further means of simplification (parallel programming is notoriously error-prone), while also strongly encouraging micro service architecture, which increase resilience of the system as a whole.

V. CHOICE OF IMPLEMENTATION LANGUAGE

Definition V.1. Independent systems language. A Turing complete programming language with facilities to utilize the whole instruction set for a given hardware architecture, including writing arbitrary data to arbitrary addresses.

C and C++ are examples of independent systems languages for most modern hardware architectures e.g. x86: the `asm` keyword (i.e. "inline assembly") makes the full instruction set available to the programmer, including privileged instructions such as `hlt`, `in` and `out`, and the pointer data type and (unsafe) type conversion allows arbitrary data to be written to arbitrary addresses. Type safe languages such as javascript, OCaml, Haskell and Python are not independent systems languages by design; type safety can be immediately violated by e.g. type coercion. The requirement for being an independent systems language is thus incompatible with type safety. To bridge this incompatibility, unikernels written in type safe languages must necessarily contain a portion of code written in an independent systems language. In most cases, such as with MirageOS, this is done in C.

VI. WELL DEFINED PROPERTIES OF UNIKERNEL SYSTEMS

Based on the previous definitions we can now provide a framework of well-defined properties of unikernel-based systems:

- **Service isolation:** A well-defined and absolute isolation mechanism such as Popek-Goldberg virtualization is preferable as 0 bytes of code needs to be shared between services during runtime. Enforcement is performed by hardware at the instruction level. Following closely is Xen PVH, which is mostly hardware virtualization, but some shared code. Paravirtualization shares a thin yet fairly complex set of software bindings and hardware

is only used for classical process isolation. One way to quantify this property is the amount of software, in bytes, shared by each service on the virtual machine monitor. Microcode / firmware would be a grey area, but that would be common to all current isolation mechanisms.

- **VM Slimness, Bloat and attack surface:** For a software service SW requiring objects A, B and C to form a machine image, a system (e.g. unikernel library) that can produce a virtual machine containing exactly $\{SW, A, B, C\}$ without $\{D, E, F, \dots\}$ provides *optimal slimness*. Conversely, the amount of code added to the VM in addition to $\{SW, A, B, C\}$ adds *bloat*. As an example, if $\{SW, A, B, C\}$ was 10MB in size and an operating system added 5 MB that would be 50% of bloat. Linux-based virtual machines would easily be 2000% bloated for single-purpose virtual machines, e.g. using a trimmed down Linux micro-core of 24MB used to run a 1MB service, 96% of the machine image would be operating system. IncludeOS instances are typically 1 MB of OS, so if the service could run with IncludeOS that would be 50/50 software and OS, plus a few percent overweight (one typically would not use all the code included, even if one included only the objects needed, unless the objects themselves are each absolutely minimal). Given that 1MB was sufficient to add the required operating system parts, wrapping it in a Linux VM would literally make for 2300% of bloat.
- **System mutability:** To what extent is it possible to change the system once launched? This is hard to quantify, but we propose the following set of properties as "optimal immutability" which a system should strive for:
 - 1) All data that can be read-only is
 - 2) All executable code is write-protected
 - 3) Write-able areas of memory (i.e. the heap / working memory) is not executable

Enforcement of these rules should be implemented at the lowest level. In IncludeOS, this kind of protection cannot really be enforced on current platforms. Type-safe language unikernels, such as Mirage, have a certain degree of language-level protection, but only in the parts of the unikernel not written in C. We propose a future work on hypervisors where we provide an interface for specifying which parts of the VM that should be read-only, execute- and read/write (but not execute), when the system boots. This way, the hypervisor at ring -1 can set up memory segments inside the VM before it starts, denying even the VM itself the ability to modify read-only parts of memory. Having the CPU enforce these rules will make it useless to inject code into a VM, if one found a way to do it, as jumping to that code would trigger a hardware trap.
- **Possibility of internal system misuse:** To what extent does the operating system allow parts of the code to be used for unintended purposes? Having a terminal makes several commands available for "general purpose" or "ad-

hoc use” of the code embedded into the system. Not having a terminal, or other similar means of allowing ad-hoc function calls, greatly reduces or entirely removes this possibility.

VII. OUR PROPOSED SOLUTION

By default, in the interests of usability, conventional systems open many more ports than may be needed to run a system. An open port, especially one which is not needed, is another route in for the attacker. We also take the position that the probability of vulnerabilities being present in a system increases proportionally to the amount of executable code it contains. Having less executable code inside a given system will reduce the chances of a breach and also reduce the number of tools available for an attacker once inside. As Meireles [14] said in 2007 “... *while you can sometimes attack what you can’t see, you can’t attack what is not there!*”. Given the success with which the threat environment continually attacks business globally [15]–[19], it is clear that many companies are falling down on many of the key issues we have highlighted in Section I. It is also clear that a sophisticated and complex solution is unlikely to work. Thus we must approach the problem from a more simple perspective.

A. Service isolation

A fundamental premise for cloud computing is the ability to share hardware. In private cloud systems, hardware resources are shared across a potentially large organization, while on public clouds, hardware is shared globally across multiple tenants. In both cases, isolating one service from the other is an absolute requirement.

The simplest mechanism for service isolation is simply *process isolation* in classic kernels, relying on hardware supported virtual memory e.g. provided by the now pervasive x86 protected mode. While process isolation has been used successfully in mainframe setups for decades, access to terminals with limited user privileges has also been the context for classical attack vectors such as stack smashing, root-kits etc., the main problem being that a single kernel is being shared between several processes and that gaining root access from one terminal would give access to everything inside the system. As a result, much work was done in the sixties and seventies to find ways to completely isolate a service without sharing a kernel. This work culminated with the seminal 1974 paper by Gerald J. Popek and Robert P. Goldberg [20] where they present a formal model describing the requirements for complete instruction level virtualization, i.e. *hardware virtualization*.

While hardware virtualization was in wide use on e.g. IBM mainframes from that time, it wasn’t until 2005 that the leading commodity CPU manufacturers, Intel and AMD, introduced these facilities into their chips. In the meantime, paravirtualization had been re-introduced as a workaround to get virtual machines on these architectures, notably in [21]. While widely deployed and depended upon, the Xen project has recently been evolving its paravirtualization interface towards using

hardware virtualization in e.g. PVH [22] stating that “*PVH means less code and fewer Interfaces in Linux/FreeBSD: consequently it has a smaller TCB and attack surface, and thus fewer possible exploits*” [23].

Another isolation mechanism is operating system-level virtualization with containers, e.g. LXC popularized in recent years by Docker, where each container represents a userspace operating environment for services that all share a kernel. The mechanism for isolating one container from another is classical process isolation, augmented with software controls such as cgroups and Linux namespaces. While containers do offer less overhead than classic virtual machines, a good example where containers make a lot of sense would be trusted in-house clouds, i.e. Google is using containers internally for most purposes [24]. We take the position that hardware virtualization is the simplest and most complete mechanism for service isolation, with the best understood foundations as formally described by Popek and Goldberg, and that this should be the preferred isolation mechanism for secure cloud computing.

B. Why Use Unikernels?

Using hardware virtualization as the preferred isolation mechanism requires an operating system to be embedded into the virtual machine. IaaS cloud providers will typically offer virtual machine images running a classical general purpose operating system, such as Microsoft Windows and one or more flavours of Linux, possibly optimized for cloud by e.g. removing device drivers that are not needed. While specialized Linux distributions can greatly reduce the memory footprint and attack surface of a virtual machine, general purpose multi-process operating systems will, by design, contain a large amount of functionality that is simply not needed by one single service. We take the position that virtual machines should be specialized to a high degree, each forming a single purpose micro service, to facilitate a resilient and fault tolerant system architecture, which is also highly scalable.

We argue that the unikernel approach offers the potential to meet all our needs, while delivering a much reduced attack surface, yet providing exactly the performance we require. An added bonus will be the reduced operating footprint, meaning a more green approach is delivered at the same time.

C. How Does This Compare to a Conventional System?

Looking at what Frederick P. Brooks Jnr. suggests in [25] “Because ease of use is the purpose, this ratio of function to conceptual complexity is the ultimate test of system design. Neither function nor simplicity alone defines a good design”, we can see where modern software systems are missing the point. The more complex a system becomes, the more overhead is introduced, leading to greater complexity and ultimately unnecessary bloat, draining performance, and exposing vulnerabilities. Conventional cloud systems tend to be over-complicated, unnecessarily bloated, and therefore expensive to scale. Unikernels, on the other hand in [5], “Unikernels are specialized, single-address-space machine images constructed

by using library operating systems”, meaning they are exactly the right size to carry out their given task — no larger, and no smaller.

Our proposed approach, using unikernels, limits/enforces the software architect to use a given pattern (event-based computing using the single-responsibility-principle, service-oriented architectures, separation of data and processing, and modularity) — which is very good from a software design point of view. We are trying to get people to use “best-of-breed” patterns, and thus develop better software through this limitation.

VIII. HOW DOES THIS ADDRESS OUR TEN KEY CONCERNS?

As we saw in the introduction, we identified 10 key security issues which have been identified and which need to be addressed. Of these, we believe our proposed unikernel solution can help us address seven of these issues, namely:

- 1 The definition of security goals;
- 2 Compliance with standards;
- 3 Audit issues;
- 4 Management approach;
- 5 Technical complexity of cloud;
- 7 Measurement and monitoring;
- 10 The threat environment.

A. The Definition of Security Goals

We propose to build in a number of basic sensible security goals to the system, so that these goals are included by design. Naturally, it will be possible to accommodate additional goals, where the user identifies those as appropriate for their needs.

B. Compliance with Standards

Compliance with standards is generally achieved through some form of assurance [26], which generally can be achieved by a compliance process or by audit. Audit is expensive if done well, thus compliance through the use of checklists is the usual method chosen, but this process brings weaknesses with it [27]. By tightening up information flows within the system, and by providing rigorous audit trails, our proposed solution seeks to maximise assurance, thus leading to compliance in a much more accurate and cost effective way.

C. Audit Issues

There are many audit issues which need to be addressed [28], not least surrounding the use of the humble audit trail [29]. In a forthcoming paper, we outline in more detail how our proposed system will tackle this key issue with a much more rigorous approach.

D. Management Approach

Cloud ecosystems involve far more actors than conventional systems, and it is important to recognise that many of these actors will have differing agendas [30]. Our proposed approach will seek to minimise the impact of third party actors by reducing the opportunity for these actors to adversely influence the effectiveness of the security approach.

E. Technical Complexity of Cloud

There is no doubt that distributed systems are highly complex, and that cloud ecosystems are, by their nature, far more complex [31]. We propose to tackle this issue through simplification of how the system is constructed, in order to minimise the attack surface.

F. Measurement and monitoring

As an adjunct to a provable level of security [32], it is necessary to measure and monitor what is happening with a system. To that end, our proposed system will, by default, provide a considerable armoury of measurement and monitoring capabilities, which will allow users to be satisfied of the level of security they have achieved, and will continue to achieve through use of the system.

G. The threat environment

This is a major and very worrying issue, which continues to evolve day by day. Our proposed approach will seek to tackle this through minimising the attack surface, minimising access routes to attackers, and generally making life difficult for the attackers. This is an issue that will bear much ongoing scrutiny by the research community in order to try to keep ahead of the attackers.

IX. INITIAL THOUGHTS ON PENETRATION TESTING

Penetration testers often refer to the OWASP foundation Top 10 report, see Table II below for details of the most used attack techniques. In its current 2013 installation, two vulnerabilities—A5-Security Misconfiguration and A9-Using Known Vulnerable Components—are directly related to the rich landscape of available server-side functions which commonly are neither minimized nor properly configured. Recent years have given rise to opinionated frameworks, i.e. frameworks that guide developers with sensible security defaults. Their security measures efficiently reduce threats from common attack vectors, e.g. A1-Injection or A2-Cross-Site Scripting, but those frameworks themselves can introduce vulnerabilities, as OWASP noted with its introduction of A9 as “the growth and depth of component based development has significantly increased the risk of using known vulnerable components”.

2013 Code	Threat
A1	Injection Attacks
A2	Broken Authentication and Session Management
A3	Cross Site Scripting (XSS)
A4	Insecure Direct Object References
A5	Security Misconfiguration
A6	Sensitive Data Exposure
A7	Missing Function Level Access Control
A8	Cross Site Request Forgery (CSRF)
A9	Using Components with Known Vulnerabilities
A10	Unvalidated Redirects and Forwards

TABLE II
OWASP TOP TEN WEB VULNERABILITIES — 2013
[33]

The Unikernel approach implicitly minimizes infrastructure — runtime environment, and libraries as well as operating system shells — and thus reduces exposure to attack vectors A5 and A9. In addition, their single-process paradigm enforces beneficial architecture design decisions that yields systems with clearer separation-of-concerns. Given the rise of opinionated frameworks, we envision a web-development framework that de-constructs high-level workflows into separate unikernels, structures communication between those, and provides sensible security defaults. We assume that such a system of unikernels can solve complex web-application workflows in a secure manner without negatively impacting developer’s productivity during development and debugging. We address this area in much more depth in our next paper.

X. CONCLUSIONS

In this paper, we have introduced a framework of definitions and metrics for classifying unikernel systems. We have started the process of developing a formal approach to describing our framework, and have considered how such a theoretical framework might provide a more secure approach to the challenging issues of cloud security and privacy.

We have proposed a novel means of significantly reducing the attack surface for a cloud based system, removing in the process many classic attack vectors. We consider the architecture of the proposed system and its resilience to attack in much more depth in our next paper.

This will be followed by the need to look at, and solve, the challenge presented by audit trail issues. This, in turn, raises the need to address secure internal communication, access logging and log storage, and provision for the need to maintain a strong forensic trail. Once we have the security basics we are looking for in place, we can then turn our attention to a robust approach to privacy.

REFERENCES

- [1] B. Duncan, A. Bratterud, and A. Happe, “Enhancing Cloud Security and Privacy: Time for a New Approach?” in *INTECH 2016*, Dublin, 2016, pp. 1–6.
- [2] A. Madhavapeddy, R. Mortier, C. Rotsos, D. Scott, B. Singh, T. Gazagnaire, S. Smith, S. Hand, and J. Crowcroft, “Unikernels: Library Operating Systems for the Cloud,” *ASPLOS ’13 Proc. eighteenth Int. Conf. Archit. Support Program. Lang. Oper. Syst.*, vol. 48, pp. 461–472, 2013.
- [3] A. Madhavapeddy and D. J. Scott, “Unikernels: Rise of the Virtual Library Operating System,” *Commun. ACM*, vol. 57, no. 1, pp. 61–69, 2014.
- [4] A. Kantee, “The Rise and Fall of the Operating System,” *Login.*, pp. 6–9, 2015.
- [5] Unikernel.org, “Unikernels,” 2015. [Online]. Available: www.unikernel.org
- [6] D. R. Engler, M. F. Kaashoek, and Others, *Exokernel: An operating system architecture for application-level resource management*. ACM, 1995, vol. 29, no. 5.
- [7] D. E. Porter, S. Boyd-Wickizer, J. Howell, R. Olinsky, and G. C. Hunt, “Rethinking the library OS from the top down,” in *ACM SIGPLAN Not.*, vol. 46, no. 3. ACM, 2011, pp. 291–304.
- [8] A. Baumann, M. Peinado, and G. Hunt, “Shielding applications from an untrusted cloud with haven,” *ACM Trans. Comput. Syst.*, vol. 33, no. 3, p. 8, 2015.
- [9] G. J. Popek and R. P. Goldberg, “Formal Requirements for Virtualizable Third Generation Architectures,” *ACM SIGOPS Oper. Syst. Rev.*, vol. 7, no. 4, p. 112, 1973.
- [10] M. F. Kaashoek, K. Mackenzie, D. R. Engler, G. R. Ganger, H. M. Briceño, R. Hunt, D. Mazières, T. Pinckney, R. Grimm, and J. Jannotti, “Application Performance and Flexibility on Exokernel Systems,” *ACM SIGOPS Oper. Syst. Rev.*, vol. 31, no. 5, pp. 52–65, 1997.
- [11] A. Bratterud, A.-A. Walla, H. Haugerud, P. E. Engelstad, and K. Begnum, “IncludeOS: A Minimal, Resource Efficient Unikernel for Cloud Services,” *2015 IEEE 7th Int. Conf. Cloud Comput. Technol. Sci.*, pp. 250–257, 2015.
- [12] G. Klein, K. Elphinstone, G. Heiser, J. Andronick, D. Cock, P. Derrin, D. Elkaduwe, K. Engelhardt, R. Kolanski, M. Norrish, T. Sewell, H. Tuch, and S. Winwood, “seL4: Formal verification of an OS kernel,” *Proc. ACM SIGOPS 22nd Symp. Oper. Syst. Princ.*, pp. 207–220, 2009.
- [13] A. Bratterud and H. Haugerud, “Maximizing Hypervisor Scalability Using Minimal Virtual Machines,” *2013 IEEE 5th Int. Conf. Cloud Comput. Technol. Sci.*, pp. 218–223, 2013.
- [14] P. Meireles, “Narkive Mailinglist Archive,” 2007. [Online]. Available: <http://m0n0wall.m0n0.narkive.com/OI4NbHQq/m0n0wall-virtualization>
- [15] PWC, “Information Security Breaches Survey 2010 Technical Report,” pp. 1–22, 2010. [Online]. Available: <http://www.pwc.co.uk/>
- [16] PWC, “UK Information Security Breaches Survey - Technical Report 2012,” London, Tech. Rep. April, 2012. [Online]. Available: www.pwc.comwww.bis.gov.uk
- [17] PWC, “2014 Information Security Breaches Survey: Technical Report,” Tech. Rep., 2014.
- [18] Verizon, N. High, T. Crime, I. Reporting, and I. S. Service, “2012 Data Breach Investigations Report,” Verizon, Tech. Rep., 2012.
- [19] Verizon, “2014 Data Breach Investigations Report,” Tech. Rep. 1, 2014. [Online]. Available: <http://www.verizonenterprise.com/resources/reports/rp{\}Verizon-DBIR-2014{\}en{\}xg.pdf>
- [20] G. J. Popek and R. P. Goldberg, “Formal Requirements for Virtualizable Third Generation Architectures,” *Commun. ACM*, vol. 17, no. 7, pp. 412–421, 1974.
- [21] P. Barham, B. Dragovic, K. Fraser, S. Hand, T. Harris, A. Ho, R. Neugebauer, I. Pratt, and A. Warfield, “Xen and the art of virtualization,” *SIGOPS Oper.Syst.Rev.*, vol. 37, no. 5, pp. 164–177, 2003.
- [22] D. Chisnall, “Xen PVH: Bringing Hardware to Paravirtualization,” *Inf. IT*, 2014.
- [23] X. Project, “Xen Project Software Overview,” 2016. [Online]. Available: <http://wiki.xen.org/wiki/Xen{\}Project{\}Software{\}Overview{\}#PVH>
- [24] A. Verma, L. Pedrosa, M. Korupolu, D. Oppenheimer, E. Tune, and J. Wilkes, “Large-scale cluster management at Google with Borg,” *Proc. Tenth Eur. Conf. Comput. Syst. - EuroSys ’15*, pp. 1–17, 2015.
- [25] F. P. Brooks Jr, *The Mythical Man-Month: Essays on Software Engineering, Anniversary Edition, 2/E*. Pearson Education India, 1995.
- [26] B. Duncan and M. Whittington, “Compliance with Standards, Assurance and Audit: Does this Equal Security?” in *Proc. 7th Int. Conf. Secur. Inf. Networks*. Glasgow: ACM, 2014, pp. 77–84.
- [27] B. Duncan and M. Whittington, “Reflecting on whether checklists can tick the box for cloud security,” in *Proc. Int. Conf. Cloud Comput. Technol. Sci. CloudCom*, vol. 2015-Febru, no. February. Singapore: IEEE, 2015, pp. 805–810.
- [28] B. Duncan and M. Whittington, “Enhancing Cloud Security and Privacy: The Cloud Audit Problem,” in *Cloud Comput. 2016 Seventh Int. Conf. Cloud Comput. GRIDs, Virtualization*. Rome: IEEE, 2016, pp. 119–124.
- [29] B. Duncan and M. Whittington, “Enhancing Cloud Security and Privacy: The Power and the Weakness of the Audit Trail,” in *Cloud Comput. 2016 Seventh Int. Conf. Cloud Comput. GRIDs, Virtualization*. Rome: IEEE, 2016, pp. 125–130.
- [30] B. Duncan and M. Whittington, “Company Management Approaches Stewardship or Agency: Which Promotes Better Security in Cloud Ecosystems?” in *Cloud Comput. 2015*. Nice: IEEE, 2015, pp. 154–159.
- [31] B. Duncan, D. J. Pym, and M. Whittington, “Developing a Conceptual Framework for Cloud Security Assurance,” in *Cloud Comput. Technol. Sci. (CloudCom), 2013 IEEE 5th Int. Conf. (Volume 2)*. Bristol: IEEE, 2013, pp. 120–125.
- [32] B. Duncan and M. Whittington, “The Importance of Proper Measurement for a Cloud Security Assurance Model,” in *2015 IEEE 7th Int. Conf. Cloud Comput. Technol. Sci.*, Vancouver, 2015, pp. 1–6.
- [33] OWASP, “OWASP Top Ten Vulnerabilities 2013,” 2013. [Online]. Available: <https://www.owasp.org/index.php/Category:OWASP{\}Top{\}Ten{\}Project>