



NAVAL POSTGRADUATE SCHOOL

MONTEREY, CALIFORNIA

THESIS

**CLOUD COMPUTING AND VIRTUAL DESKTOP
INFRASTRUCTURES IN AFLOAT ENVIRONMENTS**

by

Stefan E. Gillette

June 2012

Thesis Co-Advisors:

Douglas J. MacKinnon
Rachel Goshorn

Approved for public release; distribution is unlimited

THIS PAGE INTENTIONALLY LEFT BLANK

REPORT DOCUMENTATION PAGE			<i>Form Approved OMB No. 0704-0188</i>	
Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instruction, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302, and to the Office of Management and Budget, Paperwork Reduction Project (0704-0188) Washington DC 20503.				
1. AGENCY USE ONLY (Leave blank)		2. REPORT DATE June 2012	3. REPORT TYPE AND DATES COVERED Master's Thesis	
4. TITLE AND SUBTITLE Cloud Computing and Virtual Desktop Infrastructures in Afloat Environments			5. FUNDING NUMBERS	
6. AUTHOR(S) Stefan E. Gillette				
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Naval Postgraduate School Monterey, CA 93943-5000			8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING /MONITORING AGENCY NAME(S) AND ADDRESS(ES) N/A			10. SPONSORING/MONITORING AGENCY REPORT NUMBER	
11. SUPPLEMENTARY NOTES The views expressed in this thesis are those of the author and do not reflect the official policy or position of the Department of Defense or the U.S. Government. IRB Protocol Number __N/A__				
12a. DISTRIBUTION / AVAILABILITY STATEMENT Approved for public release; distribution is unlimited			12b. DISTRIBUTION CODE A	
13. ABSTRACT (maximum 200 words) The phenomenon of "cloud computing" has become ubiquitous among users of the Internet and many commercial applications. Yet, the U.S. Navy has conducted limited research in this nascent technology. This thesis explores the application and integration of cloud computing both at the shipboard level and in a multi-ship environment. A virtual desktop infrastructure, mirroring a shipboard environment, was built and analyzed in the Cloud Lab at the Naval Postgraduate School, which offers a potential model for the foundation of a cloud computing infrastructure in a network environment aboard ship. This research develops a Concept of Operations to propose how a cloud computing infrastructure may be employed and how it might operate in a multi-ship environment. This thesis' findings indicate that cloud computing, when combined with virtualization technologies, can improve interoperability via the loose coupling of systems, decrease network footprints via server consolidation, and increase elasticity of resources. Additionally, cloud computing may alleviate bandwidth constraints because data and information in a cloud network can be stored, shared, and accessed locally. This could also reduce if not eliminate reachback through satellites. Future efforts in this area of research may involve more rigorous testing, and opportunities toward improved security, as well as leveraging ever-improving cloud software.				
14. SUBJECT TERMS Cloud Computing, Virtualization, Virtual Technology, Virtual Desktop Infrastructure, Virtual Machine, Service Oriented Architecture, Afloat Architecture, Consolidated Afloat Network Enterprise Services, Thin Client, Zero Client			15. NUMBER OF PAGES 98	
			16. PRICE CODE	
17. SECURITY CLASSIFICATION OF REPORT Unclassified	18. SECURITY CLASSIFICATION OF THIS PAGE Unclassified	19. SECURITY CLASSIFICATION OF ABSTRACT Unclassified	20. LIMITATION OF ABSTRACT UU	

THIS PAGE INTENTIONALLY LEFT BLANK

Approved for public release; distribution is unlimited

**CLOUD COMPUTING AND VIRTUAL DESKTOP INFRASTRUCTURES IN
AFLOAT ENVIRONMENTS**

Stefan E. Gillette
Lieutenant Junior Grade, United States Navy
B.A., University of Washington, 2006

Submitted in partial fulfillment of the
requirements for the degree of

MASTER OF SCIENCE IN SYSTEMS TECHNOLOGY

from the

**NAVAL POSTGRADUATE SCHOOL
June 2012**

Author: Stefan E. Gillette

Approved by: Douglas J. MacKinnon
Thesis Co-Advisor

Rachel Goshorn
Thesis Co-Advisor

Albert Barreto
Second Reader

Dr. Dan Boger
Chair, Department of Information Sciences

THIS PAGE INTENTIONALLY LEFT BLANK

ABSTRACT

The phenomenon of “cloud computing” has become ubiquitous among users of the Internet and many commercial applications. Yet, the U.S. Navy has conducted limited research in this nascent technology. This thesis explores the application and integration of cloud computing both at the shipboard level and in a multi-ship environment. A virtual desktop infrastructure, mirroring a shipboard environment, was built and analyzed in the Cloud Lab at the Naval Postgraduate School, which offers a potential model for the foundation of a cloud computing infrastructure in a network environment aboard ship. This research develops a Concept of Operations to propose how a cloud computing infrastructure may be employed and how it might operate in a multi-ship environment. This thesis’ findings indicate that cloud computing, when combined with virtualization technologies, can improve interoperability via the loose coupling of systems, decrease network footprints via server consolidation, and increase elasticity of resources. Additionally, cloud computing may alleviate bandwidth constraints because data and information in a cloud network can be stored, shared, and accessed locally. This could also reduce if not eliminate reachback through satellites. Future efforts in this area of research may involve more rigorous testing, and opportunities toward improved security, as well as leveraging ever-improving cloud software.

THIS PAGE INTENTIONALLY LEFT BLANK

TABLE OF CONTENTS

I.	INTRODUCTION.....	1
A.	BACKGROUND	1
B.	PURPOSE.....	3
C.	METHODS	3
	1. Literature Review	4
	2. Virtual Desktop Infrastructures.....	4
	3. Cloud Computing Concept of Operations	5
D.	STRUCTURE.....	5
II.	LITERATURE REVIEW	7
A.	CLOUD COMPUTING.....	7
	1. Cloud Computing Defined	7
	a. Defense Information Systems Agency (DISA).....	8
	b. National Institute of Standards and Technology (NIST).....	8
	2. Cloud Computing Component Services.....	9
	a. Infrastructure-as-a-Service	9
	b. Platform-as-a-Service	9
	c. Software-as-a-Service	10
	d. Desktop-as-a-Service.....	11
	3. Cloud Computing Models	11
	a. Private Cloud.....	11
	b. Public Cloud.....	11
	c. Hybrid Cloud.....	12
B.	SERVICE-ORIENTED ARCHITECTURE (SOA).....	12
C.	VIRTUALIZATION.....	13
	1. Virtual Local Area Network	14
	2. Virtual Machine	14
	3. Hypervisor	15
	4. Paravirtualization	15
D.	BANDWIDTH AND LATENCY.....	16
	1. Bandwidth.....	16
	2. Latency.....	17
E.	COMMON INTERNET PROTOCOLS	18
	1. Transmission Control Protocol.....	18
	2. User Datagram Protocol.....	18
F.	COMMON VDI PROTOCOLS.....	19
	1. Remote Desktop Protocol.....	19
	2. Personal Computer over Internet Protocol	19
G.	THICK CLIENTS, THIN CLIENTS, AND ZERO CLIENTS	20
	1. Thick Clients.....	20
	2. Thin Clients	21
	3. Zero Clients	21
H.	DOD ENTERPRISE NETWORK PROGRAMS AND INITIATIVES....	21

1.	Integrated Shipboard Network System	22
2.	Information Technology for the Twenty-First Century.....	22
3.	Consolidated Afloat Networks and Enterprise Services	23
III.	VIRTUALIZATION MODELS	27
A.	NPS CLOUD LAB PHYSICAL INFRASTRUCTURE	29
1.	Blade Enclosures	29
a.	Blade Management Controller	30
2.	Blade Servers	31
3.	Server Operating Systems	31
B.	NPS CLOUD LAB VIRTUAL INFRASTRUCTURE	32
1.	VMware ESXi.....	32
2.	VMware vCenter Server	34
3.	VMware vSphere	35
a.	vMotion.....	36
b.	High Availability (HA).....	37
c.	Fault Tolerance.....	37
4.	Hosts, Clusters, and Resource Pools	37
a.	Hosts	37
b.	Clusters	38
c.	Resource Pools	38
5.	Virtual Networks.....	39
a.	vSphere Distributed Switch.....	39
6.	Storage	40
7.	VMware View.....	40
8.	Virtual Machines and the End User Interface	41
C.	SHIPBOARD VIRTUAL DESKTOP INFRASTRUCTURES.....	44
IV.	CLOUD COMPUTING CONCEPT OF OPERATIONS	49
A.	LINKING CLOUDS	49
1.	LTA Ships	51
2.	Fixed-wing Aircraft	52
3.	IP Routable Satellites.....	53
a.	Broadband Global Area Network.....	53
b.	Cisco Internet Routing in Space.....	54
4.	Worldwide Interoperability for Microwave Access.....	54
B.	CLOUD DATA STORAGE	56
C.	CLOUD-TO-CLOUD INTEROPERABILITY.....	59
1.	Levels of Information System Interoperability Maturity Model ..	60
D.	AFLOAT CLOUD COMMAND AND CONTROL STRUCTURE.....	64
V.	CONCLUSION	67
A.	SUMMARY	67
B.	FUTURE RESEARCH.....	70
	LIST OF REFERENCES.....	73
	INITIAL DISTRIBUTION LIST	77

LIST OF FIGURES

Figure 1.	Simple network diagram depicting a cloud as the Internet [3]	7
Figure 2.	Cloud computing component service layers [7]	10
Figure 3.	Server Virtualization [9]	14
Figure 4.	A Type-1 Hypervisor [11]	15
Figure 5.	The effect on bandwidth and latency as distance increases [15]	17
Figure 6.	Consolidation of legacy networks into CANES [24].....	24
Figure 7.	Server configuration without virtualization [10]	29
Figure 8.	Dell M1000e Blade Server Enclosure [30]	30
Figure 9.	Dell M1000e Blade Management Controller [30]	30
Figure 10.	Dell PowerEdge M610 Blade Server [31]	31
Figure 11.	Server configuration with virtualization [10]	33
Figure 12.	Screenshot of an ESXi DCUI Interface	34
Figure 13.	VMware vCenter Server hierarchy [32]	35
Figure 14.	The three layers of the vSphere software stack [10]	36
Figure 15.	Hosts, Clusters, and Resource Pools [10]	38
Figure 16.	Parent image with linked clones [10].....	41
Figure 17.	A Wyse P20 zero client device [34]	42
Figure 18.	NPS Cloud Lab Blade Chassis and VDI physical connections	43
Figure 19.	NPS Cloud Lab VDI physical and virtual connections	44
Figure 20.	Physical server consolidation [10]	45
Figure 21.	Screenshot of the vSphere Client manager screen.....	48
Figure 22.	Multi-ship Afloat Cloud Infrastructure	51
Figure 23.	LTA Ship acting as a data relay.....	52
Figure 24.	Illustration of Cisco's Internet Routing in Space servicing remote users [40].....	54
Figure 25.	A WiMAX Tower (left) and a WiMAX Antenna (right) [43], [44]	55
Figure 26.	Possible WiMAX infrastructure	55
Figure 27.	Possible data storage process afloat.....	59
Figure 28.	The Levels of Information System Interoperability Maturity Model [49]	61
Figure 29.	Notional force level (hub) cloud C2 structure	65
Figure 30.	Notional unit level (edge cloud) C2 structure.....	66

THIS PAGE INTENTIONALLY LEFT BLANK

LIST OF TABLES

Table 1.	PEO-C4I select steps of the Information Life Cycle Process [27]	57
----------	---	----

THIS PAGE INTENTIONALLY LEFT BLANK

LIST OF ACRONYMS AND ABBREVIATIONS

AOR	Area of Responsibility
API	Application Programming Interface
BGAN	Broadband Global Area Network
BIOS	Basic Input/output System
CANES	Consolidated Afloat Networks and Enterprise Services
CENTRIXS	Combined Enterprise Regional Information Exchange
CONOPS	Concept of Operations
COTS	Commercial off the Shelf
CPU	Central Processing Unit
CWSP	Commercial Wideband Satellite Program
DaaS	Desktop-as-a-Service
DCUI	Direct Console User Interface
DISA	Defense Information Systems Agency
DoD	Department of Defense
DSCS	Defense Satellite Communications System
EHF	Extremely High Frequency
ERN	Education and Research Network
GB	Gigabyte
Gbps	Gigabits per second
GHz	Gigahertz
GUI	Graphical User Interface
HA	High Availability
HDD	Hard Disk Drive
IaaS	Infrastructure-as-a-Service
INMARSAT	International Maritime Satellite Organization
IP	Internet Protocol
IRIS	Internet Routing in Space
ISNS	Integrated Shipboard Network System

IT	Information Technology
IT-21	Information Technology for the Twenty-First Century
Kbps	Kilobytes per second
LAN	Local Area Network
LCD	Liquid Crystal Display
LISI	Levels of Information System Interoperability
LTA	Lighter-Than-Air
LTE	Long-Term Evolution
Mbps	Megabytes per second
MB	Motherboard
MILSTAR	Military Strategic and Tactical Relay
NIC	Network Interface Card
NIST	National Institute of Standards and Technology
NPS	Naval Postgraduate School
OCONUS	Outside of the Continental United States
ONE-NET	OCONUS Navy Enterprise Network
OS	Operating System
PaaS	Platform-as-a-Service
PAID	Procedures, Application, Infrastructure, Data
PCoIP	Personal Computer over Internet Protocol
PEO-C4I	Program Executive Office – Command, Control, Communications, Computers, and Intelligence
POR	Program of Record
QoS	Quality of Service
RAM	Random Access Memory
SaaS	Software-as-a-Service
SAN	Storage Area Network
SCI	Secret Compartmented Information
SCSI	Small Computer System Interface
SHF	Super High Frequency

SOA	Service-Oriented Architecture
SPAWAR	Space and Naval Warfare Systems Command
TB	Terabyte
TCP	Transmission Control Protocol
UDP	User Datagram Protocol
UHF	Ultra High Frequency
USB	Universal Serial Bus
VDI	Virtual Desktop Infrastructure
VDS	Virtual Distributed Switch
VLAN	Virtual Local Area Network
VM	Virtual Machine
VTC	Video Teleconferencing
WiMAX	Worldwide Interoperability for Microwave Access

THIS PAGE INTENTIONALLY LEFT BLANK

ACKNOWLEDGMENTS

Gratitude must be given to the three people who helped me most with this thesis: Dr. Douglas MacKinnon, Dr. Rachel Goshorn, and Albert “Buddy” Barreto III. Dr. MacKinnon, your willingness to be my primary advisor is most appreciated. You gave me advice and guidance that had I not received would have made completing this thesis much more difficult. Your knowledge in how to conduct research and how to write a thesis is invaluable. Thank you.

Dr. Goshorn, if it were not for you I would not have chosen to write this thesis. Your enthusiasm in promoting thesis topics from PEO-C4I and in working with students is inspiring. The trip you made possible for myself and other students to visit the PEO-C4I/SPAWAR facilities in San Diego encouraged and excited me to write this thesis. I cannot thank you enough for taking the time to be one of my advisors.

Buddy, enthusiastic and inspirational are words that barely begin to describe you. You kept me going when at times I lost motivation and felt this thesis would never be completed. The devotion and interest you have in cloud computing and virtualization technologies is second to none and was a true inspiration in completing this thesis. There is no amount of thanks I can give for the time you spent with me in the NPS Cloud Lab and the numerous meetings in your office. You were always eager to help and to listen to my ideas. I appreciated working with you more than I can express, and I welcome any opportunity to do so again in the future. Thank you for everything.

THIS PAGE INTENTIONALLY LEFT BLANK

I. INTRODUCTION

Cloud computing is evolutionary in information technology and revolutionary as a business model. Prominent businesses such as Amazon, Facebook, Google, and Microsoft have adopted cloud computing as the information technology model for many of their services. Only within the last decade has the U.S. Federal Government recognized and invested in cloud computing, promoting the movement of its Federal IT enterprises to a cloud model. In 2010, former U.S. Chief Information Officer Vivek Kundra listed the move to cloud computing as point number three of the “25 Point Implementation Plan to Reform Federal Information Technology Management,” declaring that “beginning immediately, the Federal Government will shift to a “Cloud First” policy.” [1]

The Department of Defense and the Defense Information Systems Agency are actively researching cloud computing and are developing cloud computing infrastructures in some of their shore facilities, connecting to and becoming part of the Global Information Grid (GIG). Though cloud computing is gaining traction in government IT infrastructures ashore, the concept of implementing cloud computing infrastructures in afloat environments is a rather nascent idea. In fact, the Navy’s acquisition office for command, control, communications, computers, and intelligence (C4I) PEO-C4I has requested this thesis topic to help them with concepts for cloud computing afloat. This thesis advocates the integration of cloud computing in afloat U.S. Naval network environments.

A. BACKGROUND

The current afloat network infrastructure on a U.S. naval vessel uses the client-server model. In a client-server model, a client requests services from a server and a server processes the requests and returns the results to the client. Typically, in an IT enterprise there are a small number of servers fielding requests from several hundred or several thousand clients.

An Arleigh-Burke class destroyer provides an example of a naval vessel that has a small IT enterprise using the client-server model in an afloat network environment. Within the ship there is a central space where several servers are located (this space is usually referred to as the “radio room,” or just “radio”). Throughout the ship are approximately 250 workstations. Each workstation has a desktop computer (a client) that communicates with the servers over a network using Ethernet cables. Each desktop computer consists of a monitor, a mouse, a keyboard, and a mid-tower computer case containing the necessary hardware components a computer needs to operate, such as a processor, memory, a hard disk drive, and a network card. The desktop computers also have their own operating system and software programs that are installed on the computer.

Physical resources in the client-server model reside primarily in desktop computers. These resources are rarely ever completely used, resulting in resources that could be used by a computer in need of additional resources elsewhere on the network. Cloud computing is a solution for providing the access to these resources to computers in need on the network. This is accomplished by enabling on demand access to shared resource pools. Virtualization, currently non-existent in afloat network environments, is a method of providing scalable desktop computers and elastic resources when used as the foundation of a cloud computing infrastructure. Cloud computing, combined with virtualization, transitions away from the client-server model to a model in which applications are no longer installed and executed on a client device. Instead, applications are installed and executed on a server, streamed over a network, and presented to a user on a device with the minimal hardware necessary to present applications (such as zero clients, which require only a graphics card and a network interface card (NIC)).

Moving beyond the shipboard level, cloud computing infrastructures may also have potential benefits in large scale multi-ship environments, such as in a strike group. While afloat, information is sent to a ship via satellite over UHF, SHF or EHF frequencies, either from other ships or from shore facilities. Bandwidth afloat is increasingly becoming limited and as a result throughput to afloat naval assets is minimal, connectivity often unreliable, and data flow sluggish due to high-latency in the

satellite links. Cloud computing potentially offers tangible benefits to alleviating bandwidth constraints by utilizing a cloud accessible by multiple ships. Given the theoretical possibility that a strike group may become victim to a satellite denied environment, a cloud shared among ships may be a suitable alternative for sharing information and data.

B. PURPOSE

The U.S. Navy is undergoing vast changes in its computer networks, both afloat and ashore. Planned budget cuts may cause a reduction in the number of afloat assets and manpower. The Naval Networks Enterprise of 2016 calls for the reduction or elimination of many of the Navy's legacy systems, part of which is intended to reduce costs as a means of meeting budget constraints.

Cloud computing may increase bandwidth efficiency and would likely reduce costs. Virtual technology in a cloud network consolidates disparate resources into a small number of central servers, allowing multiple users to access resources that were once spread out over various networks and systems. Loose coupling of systems in a cloud computing environment allow for greater interoperability among networks and systems, ultimately leading to a smaller local area network footprint and a reduction in the amount of hardware and software onboard afloat assets; together this may reduce the number of manpower needed, resulting in cost reduction. Because data and information in a cloud network can be stored, shared, and accessed locally there exists the potential to reduce bandwidth usage. Cloud computing infrastructures integrated in afloat network environments have the potential to increase the elasticity of resources and would bring data to the tactical edge, reducing if not eliminating reach back through satellites.

C. METHODS

This discovery thesis required three distinct phases; research, experimentation, and modeling. The research phase was done through a literature review. Experimentation was conducted by building an actual virtual desktop infrastructure (VDI) at the Naval Postgraduate School (NPS). The modeling phase explored a cloud computing concept of operations (CONOPS) in afloat environments.

1. Literature Review

A literature review was carried out to provide knowledge in five subject areas specific to this thesis. The first subject area investigated the need for cloud computing in government IT infrastructures. Specifically, what interests the DoD has in cloud computing and the goals and objectives the DoD intends to achieve using cloud computing technologies and infrastructures.

The second subject area is cloud computing, a background on what cloud computing is and what it does for users of information technology. Common cloud component services and cloud deployment models are listed and defined.

The third subject area researched was virtualization, a key foundation of the cloud computing infrastructure proposed in this thesis. Virtual desktop infrastructures (VDI), a relatively new method of offering desktop computers to users, was researched in order to experiment with virtualization in the Naval Postgraduate School's Cloud Lab and to determine the applicability of a VDI into cloud computing infrastructures.

The fourth subject area researched and identified concepts and terms common in the field of computer networking, such as service-oriented architecture (SOA) and bandwidth and latency. Popular internet and VDI protocols are explained and various client devices popular in VDI and cloud infrastructures are discussed.

The last subject area examined current DoD network programs such as the Integrated Shipboard Network System and Information Technology for the 21st century, and DoD initiatives such as Consolidated Afloat Networks and Enterprise Services, which is under development by PEO-C4I. Identifying and researching network programs and initiatives allowed for the assessment of how current enterprise networks in afloat environments operate, the capabilities they offer, and the feasibility of integrating cloud computing into afloat environments.

2. Virtual Desktop Infrastructures

The experimentation phase of this thesis involved the construction of a VDI in the Cloud Lab at the Naval Postgraduate School (NPS). This was carried out in three steps,

the first being the identification and configuration of physical computer equipment on which to build a VDI. The second step was the actual building of a VDI using virtualization software products by VMware, Inc., a leader in commercial off-the-shelf virtualization technologies. [2] Testing the VDI using various end user client devices was the third step. Experimenting with a VDI allowed for a first-hand look at a virtual infrastructure as a possible foundation for a cloud computing infrastructure in afloat environments.

3. Cloud Computing Concept of Operations

Modeling a cloud computing CONOPS is the last phase of this thesis. The CONOPS ties together the concepts reviewed in the literature review phase with the experimentation of the VDI in the second phase through applying them to afloat environments. This thesis shows it is possible to model what a cloud computing infrastructure would look like in an afloat environment by understanding the concepts, definitions, and terms researched in the literature review phase and by applying a VDI as built in the NPS Cloud Lab.

Modeling a cloud infrastructure is best achieved through the development of a CONOPS and is a preliminary step required in the systems engineering design of any system. A CONOPS can be defined simply as how a system will be employed and how it will operate. To begin, a cloud computing infrastructure, using virtualization as a foundation on a single ship, is modeled. Then, the scale is expanded outward beyond the shipboard level to a multi-ship environment, such as a strike group, and offers several possible cloud infrastructures using various naval assets and communication methods.

D. STRUCTURE

Following the current chapter's introduction, Chapter II draws from the literature review and explains the concepts, definitions, and terminology relevant to this thesis. Current Navy network programs and initiatives are discussed and the implications and possibilities they present regarding their potential integration with cloud computing infrastructures are examined.

Chapter III covers the step-by-step process that was taken to design and build a VDI in the Cloud Lab. The virtualization software products that were used for the VDI are described and the features they offer to a cloud computing infrastructure are explained. Though the majority of the chapter outlines the VDI as it was built in the Cloud Lab, the end of the chapter provides a generic approach on to how to scale this and build a VDI for a naval asset in an afloat environment.

Chapter IV presents a model of what a cloud computing infrastructure would look like in a large scale multi-ship afloat environment. A cloud computing infrastructure based on virtualization at the single shipboard level is presented in Chapter III, followed by a cloud computing infrastructure in a large scale multi-ship environment in Chapter IV. This is achieved by developing a CONOPS limited in scope, specifically determining what a cloud computing infrastructure in an afloat environment would look like, how it would be employed, and how it would operate. Discussed are potential communication methods of linking clouds, cloud data storage, cloud-to-cloud interoperability, and an afloat cloud command and control structure.

To conclude, Chapter V summarizes the thesis and restates the benefits of moving to a cloud computing infrastructure in afloat environments. Suggestions for further research bring the chapter and the thesis to a close.

II. LITERATURE REVIEW

This chapter lists and defines various concepts, terms, and technologies common in the field of computer networking and in cloud computing. The Navy's current enterprise network programs and initiatives are also examined in this chapter.

A. CLOUD COMPUTING

The term cloud computing stems from the ubiquitous use of a cloud to represent the internet as depicted in a network diagram; Figure 1 serves as a simple example.

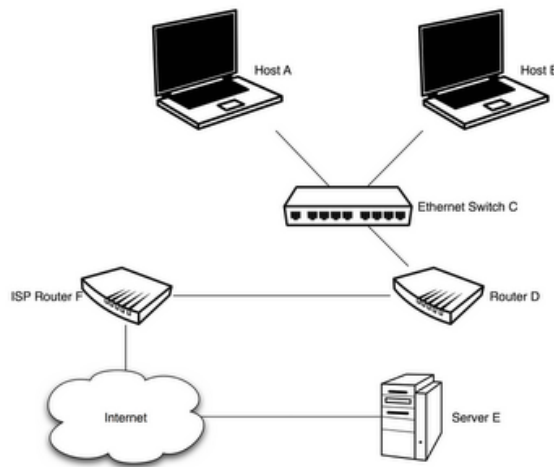


Figure 1. Simple network diagram depicting a cloud as the Internet [3]

1. Cloud Computing Defined

Cloud computing is the modern “buzzword” used to describe the delivery of computing in the form of services over a network. Often, a user or company does not own the resources or the equipment on which the services reside, but instead pays for the services like a utility. A simple analogy is an electricity grid. A user of electricity does not own the power plant where the electricity is produced, nor does the user own the power lines over which the electricity is transmitted. The user only pays for the amount of electricity used over a given period.

A cloud computing infrastructure operates much like a utility company when applied to a network environment. But instead of offering electricity, a cloud computing company can offer computing resources such as data storage or software applications. The hosting company's resources reside on their computer servers in a remote location. An end user can access and use the computing resources on his or her computer device, such as a desktop computer or smart phone.

There are more details to cloud computing than the example just given, which this thesis will show. Cloud computing also provides an enterprise with agility, cost savings, resource elasticity, and scalability. Defining cloud computing can be difficult and many definitions exist. Two of the most prominent and widely accepted definitions are given below:

a. *Defense Information Systems Agency (DISA)*

DISA defines cloud computing as “a means of enabling convenient, on-demand network access to a shared pool of configurable computing resources (e.g.; networks, servers, storage, applications and services) that can be rapidly provisioned and released with minimal management effort or service provider interaction.” [4] Additionally, DISA lists five key characteristics of cloud computing: On-Demand Self Service, Broad Network Access, Resource Pooling, Rapid Elasticity, and Measured Service. [4]

b. *National Institute of Standards and Technology (NIST)*

The National Institute of Standards and Technology published their final definition of cloud computing in September, 2011: “Cloud computing is a model for enabling ubiquitous, convenient, on-demand network access to a shared pool of configurable computing resources (e.g., networks, servers, storage, applications, and services) that can be rapidly provisioned and released with minimal management effort or service provider interaction. This cloud model is composed of five essential characteristics, three service models, and four deployment models.” [5] Like DISA, NIST lists the same five essential characteristics: On-Demand Self Service, Broad Network Access, Resource Pooling, Rapid Elasticity, and Measured Service.

2. Cloud Computing Component Services

Cloud computing can further be defined by the component services that together form a typical cloud computing infrastructure. Figure 2 illustrates the components services as they apply in a cloud computing architecture. Listed below are the most common component services provided in a cloud computing infrastructure.

a. Infrastructure-as-a-Service

The most basic service offering in a cloud infrastructure is Infrastructure-as-a-Service (IaaS). IaaS provides the power, storage, networks (including IP addresses), physical computers, and virtual machines and/or entire virtual infrastructures. Amazon Elastic Compute Cloud (EC2) and GoDaddy are two of the most popular companies offering IaaS services. NIST defines IaaS as

the capability provided to the consumer is to provision processing, storage, networks, and other fundamental computing resources where the consumer is able to deploy and run arbitrary software, which can include operating systems and applications. The consumer does not manage or control the underlying cloud infrastructure but has control over operating systems, storage, and deployed applications; and possibly limited control of select networking components (e.g., host firewalls). [5]

b. Platform-as-a-Service

Platform-as-a-Service (PaaS) delivers to the user a computing platform and provides the entire infrastructure needed to run applications over a network. [6] It is a platform on which developers can build custom web applications, requiring no downloading or installation of any kind for the end user. Simply put, PaaS is a platform that delivers applications over a network (like the Internet). Users may build their own applications or use applications already built. Google App Engine and Windows Azure are two popular companies offering PaaS services. PaaS is defined by NIST as

the capability provided to the consumer... to deploy onto the cloud infrastructure consumer-created or acquired applications created using programming languages, libraries, services, and tools supported by the provider. The consumer does not manage or control the underlying cloud infrastructure including network, servers, operating systems, or storage,

but has control over the deployed applications and possibly configuration settings for the application-hosting environment. [5]

c. Software-as-a-Service

The ability to access and use an application over a network is Software-as-a-Service. Before cloud computing, software was purchased by a company and installed on their computers. With SaaS, a company can instead use the software as necessary over the Internet and not have to purchase the software. Email providers like Hotmail and Gmail are examples of free applications hosted as a SaaS on a cloud infrastructure. Unlike PaaS, applications running as a SaaS are not open for development; i.e., a platform is not provided for application development. NIST defines SaaS as

the capability provided to the consumer... to use the provider's applications running on a cloud infrastructure². The applications are accessible from various client devices through either a thin client interface, such as a web browser (e.g., web-based email), or a program interface. The consumer does not manage or control the underlying cloud infrastructure including network, servers, operating systems, storage, or even individual application capabilities, with the possible exception of limited user-specific application configuration settings. [4]

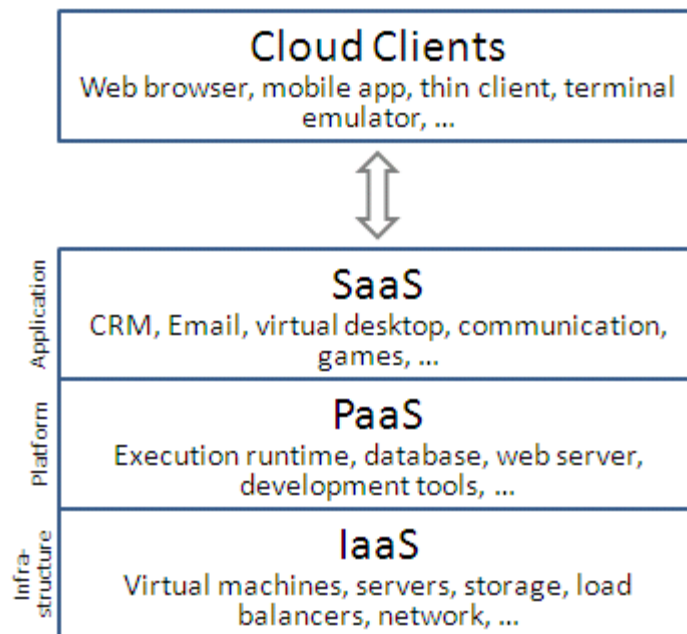


Figure 2. Cloud computing component service layers [7]

d. Desktop-as-a-Service

Desktop-as-a-Service is the delivery of hosted desktop services. The desktops are typically virtual desktops and are part of a larger VDI. This service works by delivering a virtual desktop over a network to an end user. The virtual desktop resides on a host server in a remote location. The end user accesses the virtual desktop on a client device (e.g. a desktop computer, a laptop, or a smart phone). Data created by the end user is saved to a datacenter where the host server resides. This service is similar to IaaS in that a complete user desktop is streamed over a network by a host server.

3. Cloud Computing Models

There are three primary cloud computing models (commonly known as “deployment models”). The private cloud, the public cloud, and the hybrid cloud are discussed below.

a. Private Cloud

A private cloud is a cloud in which all resources reside on site behind an organization’s firewall, leaving management of the cloud in the hands of the organization’s IT staff. Reasons for a private cloud are security, availability, and reliability. With a private cloud the organization purchases and maintains all of the hardware and software.

b. Public Cloud

Public clouds are the most common cloud models, the largest being the Internet. These clouds contain services offered to the general public and are owned by organizations renting out the services to customers. Services are offered over the Internet via web applications or web services. Often, companies with their own private cloud will connect to the Internet for various reasons (e.g., email exchange). This combination of public cloud of the Internet and a company’s private cloud results in a hybrid cloud.

c. Hybrid Cloud

A hybrid cloud infrastructure is the composition of two different types of clouds, such as the combination of a private cloud and a public cloud. The individual clouds within a hybrid cloud remain unique among the clouds within the hybrid infrastructure, “but are bound together by standardized or proprietary technology that enables data and application portability (e.g., cloud bursting for load balancing between clouds).” [5]

B. SERVICE-ORIENTED ARCHITECTURE (SOA)

The wide use of standard web services has prompted a shift in IT architectures to an architecture based on services. Service-oriented architecture attempts to address information systems as services and as such it has become a prominent architecture in IT enterprises and it has also become a programming paradigm. An SOA can be defined as

a strategic framework of technology that allows all interested systems, inside and outside of an organization, to expose and access well-defined services, and information bound to those services, that may be further abstracted to process layers and composite applications for solution development. [8]

SOA depends on the loose coupling of devices, the ability of a service to be interoperable with another service even if the services have no knowledge of the definitions of each other; the services only need to be able to logically recognize the other. This method of loose coupling and interoperability between services is also common in cloud computing architectures.

Like cloud computing, SOA operates on a paradigm in which services are at its core. Also similar to cloud computing are many of the characteristics in a SOA, including agility, granularity, interoperability, modularity, and reuse. The relationship between cloud computing and SOA is that cloud computing, which provides IT resources capable of being leveraged on demand, is a suitable IT infrastructure on which a SOA can operate. SOA is a “good approach to architecture that deals with the proper formation of the information systems using mechanisms that make them work and play well together.”

[8] Thus, when combined with cloud computing, SOA supplies an enterprise with the interfaces and architecture that link the enterprise with cloud services.

Simply put, SOA and cloud computing complement each other. An IT enterprise built on cloud computing requires governance, which includes policies and the right management tools; SOA provides such governance. Cloud computing is an instance of an architecture, whereas SOA is a pattern of architectures. In sum, “SOA is more holistic and strategic, meaning it deals with the complete enterprise including business drivers, whereas cloud computing is more tactical and is a way of solving a problem.” [8] An enterprise successfully integrating SOA and cloud computing is likely to be more effective and efficient than an enterprise that does not. [8]

C. VIRTUALIZATION

Virtualization is the partitioning of physical resources into multiple virtual machines. An example is the partitioning of a hard drive in order to allocate virtual storage to virtual machines, or the division of a physical server into multiple virtual servers. Virtualization is an enabler of resource sharing and resource allocation. In a cloud computing infrastructure virtualization offers scalability, greater agility of services, elasticity of resources, and improved infrastructure utilization. [9] Almost every component of IT can be virtualized, including servers, desktops, applications, local area networks, switches, routers, firewalls, and more.

Server virtualization is key to the consolidation of physical servers, which results in a reduction in an enterprises’ datacenter footprint and in costs savings via a reduction in the number of server hardware. Once server virtualization has been implemented a single physical server can then support multiple virtual machines (VM). Figure 3 shows a typical server virtualization structure, with the ability to scale to multiple virtual machines.

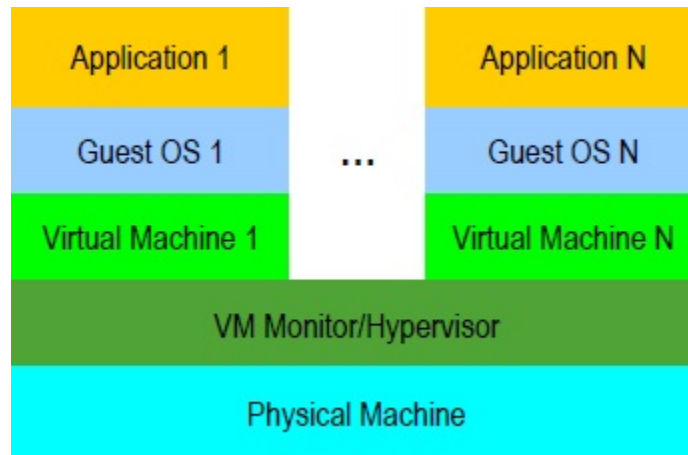


Figure 3. Server Virtualization [9]

In addition to server virtualization, virtual local area networks, VMs, and virtual machine monitors (also known as hypervisors) are a few of the most common virtual technologies.

1. Virtual Local Area Network

A virtual local area network (VLAN) replicates a physical LAN in that multiple hosts are linked together to form a network. VLANs have the same attributes found in a physical LAN, but end user devices can be grouped together on one physical machine or server via multiple network switches vice multiple physical locations sharing a single switch. [10] Though VLANs do not require the hardware (cables, switches, etc.) a physical LAN would, they do consume bandwidth from the network the VLAN is connected to.

2. Virtual Machine

A virtual machine is an isolated software application that runs its own operating systems and applications, acting like a physical computer. [10] It contains its own virtual processor, RAM, hard disk drive, and NIC. Operating systems and other computers on a network cannot tell the difference between a virtual machine and a physical machine. Applications that run on a physical computer can run on a virtual machine.

3. Hypervisor

A virtual machine monitor, or hypervisor, is the foundation of hardware virtualization and allows multiple operating systems to run on a computer. A hypervisor is the core (kernel) of a virtualization platform. It is installed on the physical server's hardware and serves as a virtual operating platform for multiple virtual operating systems. Its sole purpose is to run virtual operating systems. In terms of cloud computing it is a function of IaaS.

Hypervisors simplify the management of virtual machines via controlling resources by allocating what is needed to each operating system residing on the layer above the hypervisor. Hypervisors that run directly on top of a physical server's hardware are known as “bare-metal” hypervisors, or “Type-1” hypervisors. Figure 4 depicts an example of a Type-1 Hypervisor, with the hardware layer on bottom and the virtualized OS and applications above.

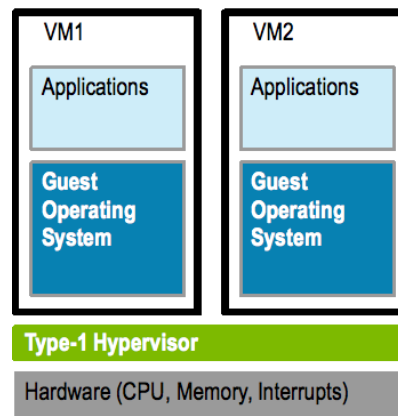


Figure 4. A Type-1 Hypervisor [11]

4. Paravirtualization

Virtualization requires that an entire system be emulated (BIOS, processor, NIC). A more efficient use of resources lies in a specific type of virtualization, called paravirtualization. With paravirtualization, multiple operating systems can run on hardware at the same time, just like with virtualization. The key difference lies in resources sharing, which paravirtualization does more efficiently because it does not emulate an entire system, but rather only certain abstractions. Generally, an operating

system performs better using paravirtualization vice full virtualization; however, flexibility is lost and security is reduced when using paravirtualization given that the OS may not be readily available for the required abstraction and that the OS has greater access to the underlying hardware. [12] As will be discussed later, the virtualization used in the NPS Cloud Lab and the virtualization advocated for integration into afloat cloud computing infrastructures is full virtualization.

D. BANDWIDTH AND LATENCY

Bandwidth and latency are data transmission characteristics. The terms have different meanings in different fields of study. In this section they are defined as they apply to the field of computer networking.

1. Bandwidth

Generally speaking, bandwidth is the measure of the amount of information that can be transmitted over a communications medium. For example, the optical carrier classification OC-1 is capable of transmitting 51.84 Mbps on optical fiber. [15] In other words, the OC-1 optical fiber has a bandwidth of 51.48 Mbps. More ascribable to afloat environments is the bandwidth provided by satellite communications. For example, the Defense Satellite Communications System (DSCS), which uses the SHF medium, has a typical bandwidth allocating between 4Mbps and 5.5Mbps per satellite footprint, but only 256Kbps to 2.048Mbps are allocated to afloat units. [12] Other afloat assets have even lower bandwidth.

Bandwidth is highly important to afloat units because they are constrained to a very limited amount of bandwidth while at sea. Increase in bandwidth is almost always the number one reason for technology upgrades to old communications satellite programs and new defense satellite programs are always aimed at increasing bandwidth to afloat units. Cloud computing infrastructures may alleviate bandwidth constraints at sea by reducing or eliminating the need for information to travel via satellite between afloat units.

2. Latency

Latency is a measurement of time, referring specifically to the delay between the transmission of a signal and its eventual receipt. [14] As distance increases between two devices, so too does the latency. Likewise, as bandwidth decreases, the amount of time it takes for information to reach its destination increases. This situation is exacerbated when large amounts of data are being sent over a medium with small bandwidth, a scenario common in afloat environments.

Potential errors can occur due to latency if data sent is not received over a given period. For example, a receiving node expects a communication over a certain time frame. The node may assume it has received all of the data in the communication when the time frame has expired, although more data may still be en route. Figure 5 illustrates the effect distance has on bandwidth and latency.

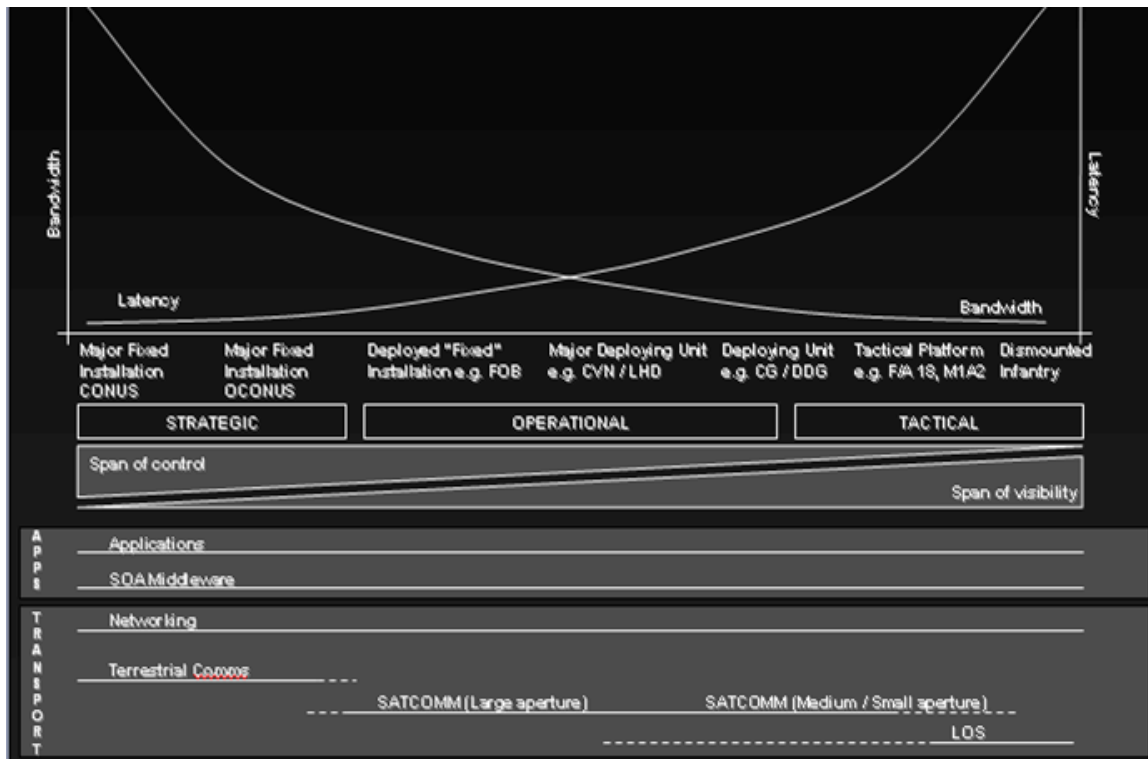


Figure 5. The effect on bandwidth and latency as distance increases [15]

E. COMMON INTERNET PROTOCOLS

A multitude of logical communications protocols exist in computer networking. The most common are those that comprise the Internet Protocol Suite and enable the transfer of information between computers on the Internet and on intranets. IP (Internet Protocol), HTTP (Hypertext Transfer Protocol), and FTP (File Transfer Protocol) are examples of the most popular protocols. IP is the primary method of transmitting data across networks. HTTP is the foundation of and primary method of transferring data across the World Wide Web. FTP is the standard method of transferring files between hosts. Computers will use either the TCP (Transmission Control Protocol) or UDP (User Datagram Protocol) protocol to transfer information between a server and remote clients on a network. [15] Both protocols have their own unique benefits and are discussed below.

1. Transmission Control Protocol

Transmission Control Protocol is a connection oriented protocol and focuses on reliability and accuracy in the delivery of “packets” (packets are the units that carry data across networks). This is accomplished via a three-way handshake. For example, before data is sent a server will initiate a transmission by sending a synchronization packet request to a client device. The client device will respond with a synchronization-acknowledge packet (a confirmation packet). The initiating device (the server) then responds with an acknowledge packet, establishing a connection. Once the connection has been established, the initiating device is free to deliver its data. This exchange illustrates the guarantee that data is being received by the client device given that a connection was first established. Though TCP provides reliability, latency increases as a result because of the three-way handshake.

2. User Datagram Protocol

User Datagram Protocol is a connectionless oriented protocol that focuses less on reliability and more on the speedy delivery of “datagrams” (datagrams are the units that carry data across networks, similar to the packets in the TCP protocol). With UDP there is no handshake or guarantee of delivery. Datagrams are simply sent one-way over a

network. Thus, UDP is a more efficient method of delivering information and reducing latency, but it lacks the reliability found in TCP. Due to the nature of its efficiency in speed, UDP is the preferred method in transmitting large volumes of information that is time-sensitive. Video, VTC, and live audio transmissions in particular benefit from UDP.

F. COMMON VDI PROTOCOLS

Protocols used to exchange information on a VDI are designed to provide a user with a graphical user interface (GUI). This is because, in a VDI, processes are being executed on the servers in a datacenter and not at the client device. A user only needs to interact with the presentation of applications on a client device (monitor, smart phone, etc.). Remote Desktop Protocol and Personal Computer over Internet Protocol are two common VDI protocols.

1. Remote Desktop Protocol

The Remote Desktop Protocol (RDP) is a protocol developed by Microsoft. Designed to support different types of network topologies and LAN protocols, RDP provides remote display and input capabilities over network connections. [16] Users on a remote client, such as a zero client, connect to a server or VM via RDP. Connection initialization, capabilities negotiation, and the transfer input between a remote client and a server are specific functions of RDP. [17]

Devices using RDP send information using TCP, ensuring data integrity. When viewing video files, however, it is often the case a user will experience a mismatch in the synchronization of the audio with the video. Video feeds and VTCs are becoming more popular on afloat units as bandwidth increases with every new generation of satellites. Thus, RDP may not always be suitable for data transmission on networks at sea, particularly in a VDI environment. A solution may be in the Personal Computer over Internet Protocol.

2. Personal Computer over Internet Protocol

Personal Computer over Internet Protocol, or PCoIP, is a proprietary protocol owned by the Teradici Corporation and is a relatively new method of delivering a desktop

over internet protocol. Based on display compression, the PCoIP protocol “compresses, encrypts, and encodes the entire computing experience at the data center and transmits it ‘pixels only’ across any standard IP network to stateless PCoIP zero clients.” [18]

The design of PCoIP is intended to “deliver a user’s desktop from a centralized host PC with an immaculate, uncompromised end-user experience across standard IP networks; including full DVI dual monitor video, complete USB compatibility, and high-definition audio.” [19] Characteristics of PCoIP include optimization for minimal bandwidths for low bandwidth situations; use of existing IP networks; compatibility with all operating systems; compatibility with all PC applications; and the ability to adapt to changing network conditions and use less bandwidth when a network is congested. [19]

Devices using PCoIP send information using either TCP or UDP. This can be advantageous depending on what type of data is being sent over a network. In situations where large volumes of data are being transmitted, such as in a VTC, then the option to select UDP is very beneficial.

G. THICK CLIENTS, THIN CLIENTS, AND ZERO CLIENTS

There are three types of client devices in client-server architectures; thick clients, thin clients, and zero clients. The differences between each of them and their advantages and disadvantages are discussed below.

1. Thick Clients

Typical network enterprises, particularly those that do not implement cloud architectures, use thick clients. A thick client (also called a “fat client”) is a complete computer desktop, equipped with the necessary hardware components to run applications on the computer. This type of client is typical of client-server networks and functions for the most part independent of a central server. A thick client has installed on it its own OS and applications and processes are executed on the client and data is stored locally on the client. Downsides of thick clients are that they must be managed individually and their resources are not elastic across a network.

2. Thin Clients

Thin clients are slimmed down versions of thick clients, containing only the bare minimum hardware components necessary to present applications via a GUI. A thin client has installed on it its own OS or at a minimum the software necessary for the thin client to operate, but applications reside on a server connected to the thin client. Thus, a thin client is server dependent; most processes are executed on the server and data is stored on the server or at a remote location (other than the thin client). A thin client is designed to be a low-end terminal, and as such little management is required of the device itself. Thin clients also practice resource elasticity, given they only use what they need from the resources supplied over the network by remote servers in a datacenter. (A datacenter is simply a facility or area in one of an organization's buildings that houses the primary processing computers and data storage of an organization's IT enterprise.)

3. Zero Clients

Zero clients are small in size, as the name would suggest, and do not contain an OS, a CPU, or memory. A display adaptor and a NIC are typically the only hardware components in a zero client. Zero clients can be defined as a device that "only connects a monitor and peripherals (mouse, keyboard, USB devices) back to a VDI or similar infrastructure in the data center." [20] All processing and storage is done at the servers in the data center, making the zero client's sole purpose the presentation of applications via a GUI.

Given the nature of their design, zero clients are optimal end user devices in a VDI. In a network environment using a VDI, benefits of zero clients include lower operating costs and easier management and deployment. [20] Like thin clients, zero clients also use only the resources they need, enabling resource elasticity.

H. DOD ENTERPRISE NETWORK PROGRAMS AND INITIATIVES

The U.S. Navy operates several enterprise networks both on land and at sea. In 2011 the Naval Enterprise Networks (NEN) Program Office (PMW205) was established to manage the Navy's three largest enterprise networks; the OCONUS Navy Enterprise

Network (ONE-NET), the Navy-Marine Corps Intranet (NMCI), and the Next Generation Enterprise Network (NGEN). These Programs of Record are shore based enterprise-wide IT networks that offer end-to-end information and telecommunication services. They provide common computing environments for both the Non-secure IP Router Network (NIPRNet) and the Secure IP Router Network (SIPRNet). [21] Currently, none of these enterprise networks utilize cloud computing infrastructures. (PEO-C4I is looking into integrating cloud computing into Navy enterprise networks and is discussed at the end of this chapter).

Enterprise networks afloat include the Integrated Shipboard Network System (ISNS) and Information Technology for the Twenty-First Century (IT-21). These afloat network systems also do not use cloud infrastructures. They are discussed below.

1. Integrated Shipboard Network System

The Integrated Shipboard Network System is a ship's primary LAN infrastructure. Both secret and unclassified networks are provided by ISNS. In addition to a network infrastructure, basic network information distribution services are also provided. Operating on a client-server model, ISNS is the backbone of an afloat unit's network infrastructure and is an integral part of IT-21. Considered a legacy system, ISNS is one of many IT systems expected to be consolidated into the Navy PEO-C4I's Consolidated Afloat Networks and Enterprise Services (CANES) initiative, discussed later in this section.

2. Information Technology for the Twenty-First Century

IT-21 was established after the introduction of Transmission Control Protocol/Internet Protocol (TCP/IP) on afloat units became the standard method of sending and receiving data. An information transfer strategy, IT-21 is a "formalization of architecture and capabilities to provide TCP/IP services to afloat units." [14] IT-21 provides IP network connectivity capable of data, video, and voice, and provides access to NIPRNet, SIPRNet, and other Navy intranets between afloat units. IT-21 also supports internet chat relay programs such as Microsoft Chat (MS Chat) and mIRC, which became

prevalent in the fleet after TCP/IP was accepted as the preferred method for coordination between afloat units.

To communicate off ship to other afloat units or shore installations IT-21 uses approved acquisition programs. These programs include, but are not limited to, the SHF based Commercial Wideband Satellite Program (CWSP), the SHF based Defense Satellite Communications System (DSCS), and the EHF based Military Strategic Relay Satellites (MILSTAR).

IT-21 implementation into the fleet began in the early 2000's and the program is currently comprised of over 60 legacy networks. The legacy systems and the technology supporting IT-21 have become operationally obsolete and are expected to be replaced by the CANES initiative.

3. Consolidated Afloat Networks and Enterprise Services

The Consolidated Afloat Networks and Enterprise Services initiative, under development by PEO-C4I, will provide a common computing environment on afloat units and is expected to replace and/or consolidate several of the Navy's aging network programs (see Figure 6), such as the Combined Enterprise Regional Information Exchange (CENTRIXS), ISNS, IT-21, Submarine Local Area Network (SUBLAN), and SCI networks. A primary function of CANES will be to provide a single support framework for C4I applications. According to PEO-C4I, "CANES will take advantage of the new business model of open architecture, Service Oriented Architecture (SOA), and rapid COTS insertion, in order to bring fiscal savings to the Navy, as well as operational agility to the warfighter." [22] The PEO-C4I Masterplan of 2012 states that CANES will reach full deployment on surface platforms by 2021. [23]

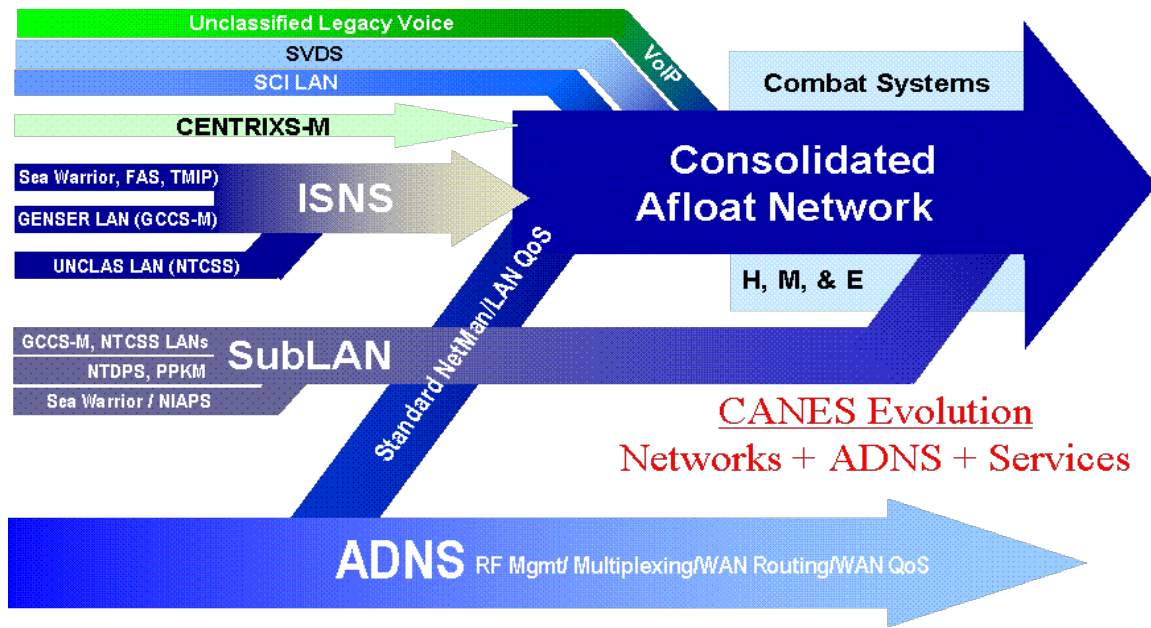


Figure 6. Consolidation of legacy networks into CANES [24]

Furthermore, CANES is an effort for “reducing server footprints and migrating existing shipboard hardware into a centralized, managed process, replacing ISNS and parts of IT-21. It will also provide the Fleet with an ability to collaborate and share information across the warfighter domain with reach-back to the assisting shore establishments.” [25] These objectives will be met by the implementation of two subprograms, the Common Computing Environment (CCE) and the Afloat Core Services (ACS). The CCE is considered the hardware portion of CANES and the ACS represents the software portion of CANES.

The primary goals of CANES are to “1) reduce the number of networks through the use of mature, certified, cross domain technologies; 2) reduce the infrastructure footprint and associated costs for hardware afloat; and 3) provide increased capability to meet current and projected warfighter requirements.” [26] Additionally, CANES will decouple applications and services software from independent hardware stacks, and instead they will be hosted on a common interoperable environment. [26]

The integration of SOA and the consolidation of legacy systems are big steps towards a cloud computing architecture. The focus on delivering services and the

decoupling of hardware and software are characteristics of CANES, and it is these same characteristics that are also at the core of what cloud computing offers an enterprise. CANES, however, will not be using a cloud computing architecture; but, it does stand as a precursor to cloud computing. Should cloud computing eventually be integrated into afloat network environments, much of the necessary foundation will already be in place given that CANES is built on a SOA. As stated earlier, SOA and cloud computing are complementary to each other. Therefore, the advocacy for the integration of cloud computing into afloat enterprise networks is both appropriate and applicable given the SOA design of CANES and the implementation of CANES in the very near future.

The Navy is aware of the tangible benefits of cloud computing, as is evident in the request by PEO-C4I to define a technical framework for cloud computing at the tactical edge. (Cloud computing is an NPS thesis area of interest for PEO-C4I, as well.) In the Framework for Cloud Computing at the Tactical Edge report [27], three strategic challenges are listed:

- How to govern, manage, process, and exploit dramatic increases in C4ISR data ashore and afloat
- How to boost IT efficiency/utilization while lowering costs
- How to align IT acquisition with fleet operational needs and rapidly deliver incremental capabilities

To address these strategic challenges, PEO-C4I suggests using cloud computing and is researching the integration of cloud computing into the US Navy. The Framework for Cloud Computing at the Tactical Edge report specifically discusses “architectural considerations, an operational concept, and content distribution management within the Navy’s C4ISR afloat/ashore environments.” [27] A portion of this report is examined further in Chapter IV to address cloud data storage afloat. [27]

THIS PAGE INTENTIONALLY LEFT BLANK

III. VIRTUALIZATION MODELS

The foundational building block of the cloud computing model for afloat shipboard networks presented in this thesis is virtualization. Virtualization eliminates the one-server, one-application model and replaces the client-server model with a virtual desktop infrastructure (VDI). In a VDI conventional desktop computers are replaced with thin or zero clients. A thin client or zero client, which contain very little hardware, are what provides a user with a GUI on which to access a virtual machine (VM). VMs reside on a host computer over a network. Hosts, which are typically servers, can have thousands of virtual machines. A VDI, when coupled with cloud computing services, is a powerful enterprise level infrastructure capable of leveraging elastic resources.

Benefits of virtualization include datacenter automation, a reduction in capital costs by increasing energy efficiency, maximization of hardware resources, and ease of enterprise desktop management. [28] Additionally, administrators can focus on one physical machine that hosts a multitude of virtual machines vice having to manage multiple servers and hundreds or thousands of physical desktop computers throughout an organization.

The cloud computing infrastructure as envisioned in a shipboard environment in this thesis is not possible without the implementation of virtualization because the creation, replication, and distribution of virtual machines to end users would not be possible. Currently, shipboard IT and network infrastructures onboard ships operate on a client-server model, generally with one room (radio) containing the central servers that process and store all data on a ship, with various smaller sized servers scattered throughout the ship that run specific applications. Depending upon the size of the ship there can be several hundred or several thousand desktop computers and/or laptops. These desktop computers contain their own hardware, such as a processor, a HDD, a CD-ROM drive, a NIC card, etc., and rely on the shipboard network for connectivity to the servers.

Often, it is the case that the hardware in the desktop computers and servers are not used to its full potential. Virtualization allows for maximizing the potential use of hardware resources that would otherwise go unused or become underutilized. Maximizing physical resources is one of the prime benefits of virtualization. Additionally, virtualization helps consolidate physical resources and it can isolate servers from each other while operating on the same physical server. This is achieved by creating a virtualization layer and virtual servers.

Once a virtualization layer has been created, the hardware can be partitioned and assigned to discrete virtual operating systems. These virtual operating systems operate independently of each other and are unaware other OSs exist on the same physical machine. The virtualization layer allows for administrators to easily control and configure the VDI. Resource pooling solves the issue of underutilized hardware, and it is easily managed by administrators without having to take a physical server offline. For example, if it has been determined that one of many virtual servers does not maximize its potential of allocated resources, then those resources may be transferred to another server in need of those resources by changing the configuration without taking offline the physical server on which the virtual server resides. If a physical server needs to be taken offline for maintenance or repairs, then the virtual servers residing on that physical server can be moved seamlessly to another physical server. This can be done manually or configured to be done automatically.

The virtualization software that was used for this thesis is VMware. [29] VMware is a frontrunner in the virtualization industry and has pioneered many of the virtualization products available on the market. Their software can run on the most common operating systems, such as Windows, Linux, and Mac OS X, while their enterprise software hypervisors run on bare-metal (i.e., directly on top of hardware, requiring no operating system).

This chapter discusses the VDI built in the NPS Cloud Lab. It serves as a simplistic example of how a virtual desktop infrastructure could be built and integrated into a shipboard environment. The sections are presented in an order that corresponds to the steps that were taken to design and build a VDI from beginning to end. The specific

VMware software products used in the Cloud Lab are likewise introduced at the appropriate step that would require their installation, rather than introducing them all at the same time in the beginning of this chapter.

A. NPS CLOUD LAB PHYSICAL INFRASTRUCTURE

Figure 7 depicts a standard server setup without virtualization. This consists of the server (the physical machine) with its assorted hardware (CPU, HDD, RAM, etc.), an OS installed on top of the hardware, and applications installed on top of the OS. A user would access data on the server by logging in to a desktop with an available connection to the server. This is an example of a typical client-server network infrastructure, which is the current infrastructure in use onboard naval ships.

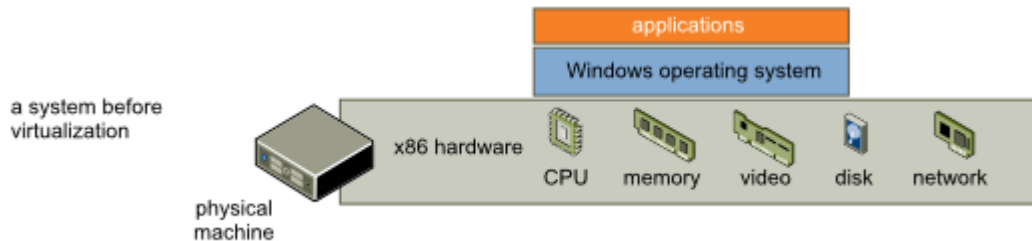


Figure 7. Server configuration without virtualization [10]

As stated earlier, a ship has several servers. To maximize physical resources and to reduce space, the physical servers can be consolidated by using blade enclosures and blade servers.

1. Blade Enclosures

The first step necessary in converting to a virtual infrastructure is to consolidate physical servers. Usually, there are several server racks onboard a ship. Each rack contains either a single, stand-alone server or a small number of servers. These server racks can be consolidated using blade enclosures and blade servers, which together are a suitable solution for server consolidation. Blade enclosures are stripped down versions of rack mounted servers and are modular in shape. The modular enclosure (chassis) reduces the use of physical space as well as energy by housing multiple blade servers. The enclosures provide power, cooling, networking, interconnects, and management

controllers. In terms of physical space, server consolidation using blade enclosures provides a much smaller footprint compared to the server racks on ships today.

In building the virtual infrastructure in the NPS Cloud Lab a Dell PowerEdge M1000e Blade Enclosure (see Figure 8) was used. These blade chassis can hold up to 16 half-height or eight full-height blade servers, has space for six power supplies, contains six bays for I/O modules, and can accommodate up to 256 processor cores and 4TB of RAM.



Figure 8. Dell M1000e Blade Server Enclosure [30]

a. Blade Management Controller

Integrated into the Dell M1000e Blade Enclosure is a Chassis Management Controller (see Figure 9). These are small LCD devices situated on the front of the blade enclosure. They provide powerful management for the entire blade enclosure and include real-time power management and monitoring, flexible security, and status and inventory alerts for the chassis and any blade servers installed in the chassis.



Figure 9. Dell M1000e Blade Management Controller [30]

2. Blade Servers

Blade servers offer a robust and scalable enterprise platform suitable for the infrastructure required of cloud computing and virtualization. Small in size, they are easily added or removed from a blade chassis as required. A blade server contains all of the typical hardware components found in a computer, but many of the components have been reduced in size or removed altogether in order to save space. However configured, blade servers contain the necessary components required to support the most demanding computing and networking infrastructures. They are also equipped with internal USB connections for embedded hypervisors, capable of delivering virtualization when using virtualization software such as VMware.

For the VDI and cloud computing infrastructure in this thesis, five Dell PowerEdge M610 Blade Servers were used (see Figure 10). A M610 Blade Server has a memory capacity of up to 192GB of RAM and two CPU sockets each capable of six cores. A M610 Blade Server can effectively scale I/O bandwidth via end-to-end 10GbE, important for alleviating bandwidth constraints.

A total of five M610 Blade Servers were installed in the M1000e Blade Enclosure in the NPS Cloud Lab. These servers were labeled AEGIS-1, AEGIS-2, AEGIS-3, AEGIS-4, and AEGIS-5. Each blade server was equipped with two dual core Intel Xeon E5540 CPUs running at 2.53GHz, 24 GB of memory, and 131GB of hard disk space.



Figure 10. Dell PowerEdge M610 Blade Server [31]

3. Server Operating Systems

Prior to building a virtual infrastructure a server operating system must be installed on one of the physical servers. A server operating system provides management tools and other features suitable for enterprise level infrastructures. Such features include

the ability to reconfigure and update hardware or software without restarting the server, flexible networking capabilities, automation capabilities, and tighter system security. A server operating system can also serve as the central management point for all physical servers installed in a blade chassis, and when using virtual software it can serve as the central management point for every virtual server installed on the physical servers.

For the Cloud Lab server operating system, Windows Server 2008 Revision 2 was used. Windows Server 2008 R2 is designed to increase the reliability and flexibility of private cloud infrastructures. Windows Server 2008 R2 was installed on server AEGIS-1. Of the entire infrastructure, the server OS is the only program installed directly on top of the physical hardware that is not virtualization software.

B. NPS CLOUD LAB VIRTUAL INFRASTRUCTURE

The physical servers and a server OS completed the physical infrastructure on which the virtual infrastructure was to be installed. From this point on, only a few virtualization programs are installed as programs on the server OS and are used solely for the management of the virtual infrastructure on each physical server (AEGIS-1 – AEGIS-5). With the physical infrastructure in place, the next step was to install the virtual servers that would eventually house virtual machines. A virtual server shares many of the same characteristics as a physical server, and is tasked with managing the resources for virtual machines. Virtual servers allow for the consolidation of underutilized physical servers, leading to reductions in datacenter space, power, cooling and administration requirements. [9] The VMware ESXi hypervisor was the virtual server used in the Cloud Lab.

1. VMware ESXi

On each physical server is installed an operating system or at a minimum a kernel (also known as a hypervisor). A kernel shares functionality similar to an OS and is the bridge between the hardware level and the software level. VMware developed a kernel aptly named VMkernel, which is embedded within their hypervisor product called VMware ESXi, an enterprise level software hypervisor. ESXi is a bare-metal hypervisor

that runs on its own kernel; i.e., it runs directly on a physical server's hardware. Figure 11 shows a server configuration using ESXi.

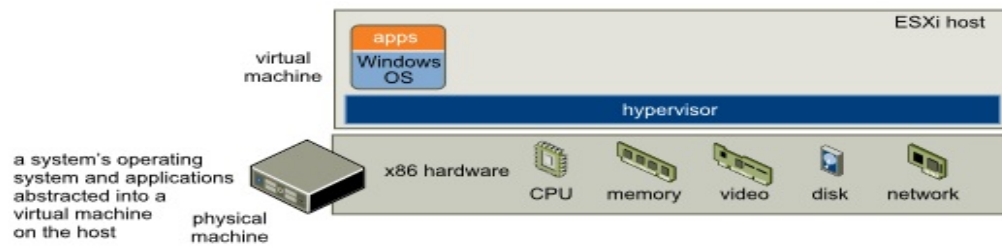


Figure 11. Server configuration with virtualization [10]

ESXi is the first virtualization layer of the VDI and is designed to specifically support running multiple VMs. It does this by abstracting physical resources such as processing, memory, and storage into multiple VMs. The hypervisor footprint is small, requiring less than 32MB of the server's physical HDD, while the entire ESXi image is only 750MB. (As a comparison, Windows Server 2008 R2 requires 10GB of space.) As a side benefit, smaller footprints lead to a lower overall attack surface and thus better security.

Access to ESXi is via a Direct Console User Interface (DCUI), a screenshot of which is shown in Figure 12. The DCUI interface is similar to that of a BIOS menu used to modify a computer's boot configuration and provides basic configuration, such as setting the administrator password and network configuration.

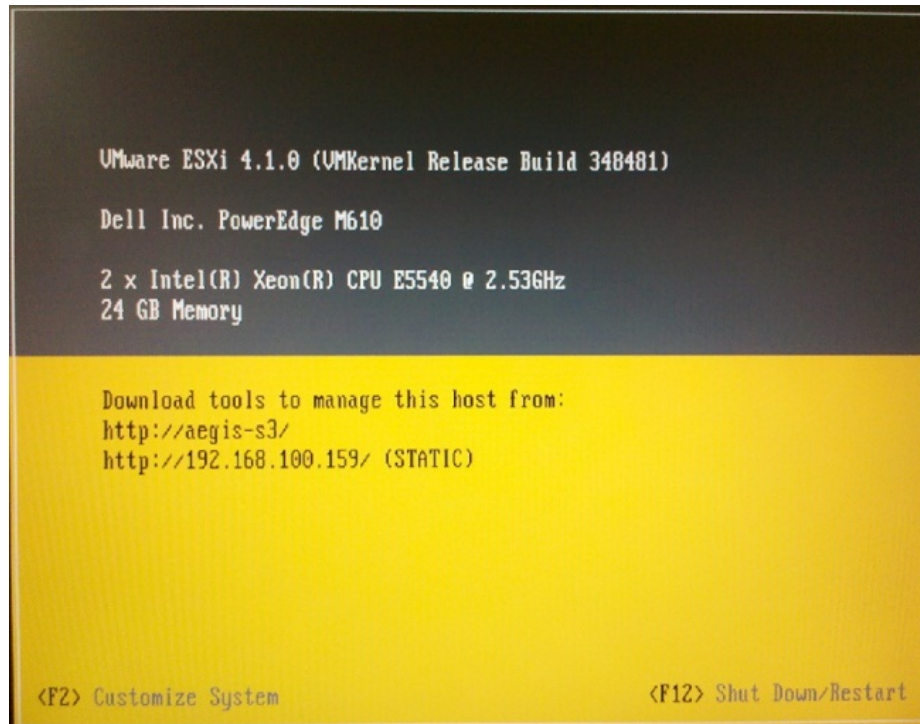


Figure 12. Screenshot of an ESXi DCUI Interface

An instance of VMware ESXi Version 4.1 was installed on each of the remaining four physical blade servers in the NPS Cloud Lab: AEGIS-2, AEGIS-3, AEGIS-4, and AEGIS-5. ESXi was the only program installed directly onto these blade servers and the DCUI is the only screen in which a user would interact with the hypervisor. The rest of the virtual infrastructure was constructed and managed from a program called VMware vCenter Server. VMware vCenter Server is explained next.

2. VMware vCenter Server

VMware vCenter Server is a simple and efficient way to manage a virtual infrastructure and it provides a scalable and extensible platform that forms the foundation for virtualization management. [32] A central point for the configuration and management of virtualized IT environments, vCenter Server provides datacenter services such as access control and performance monitoring.

In the NPS Cloud Lab vCenter Server was installed as a program running on the Windows Server 2008 R2 server OS on AEGIS-1. From vCenter Server, each instance of

ESXi installed on the other AEGIS blade servers could be managed and configured. Figure 13 shows the vCenter server hierarchy, in which vCenter Server manages over each instance of vSphere platforms and their virtual machines.

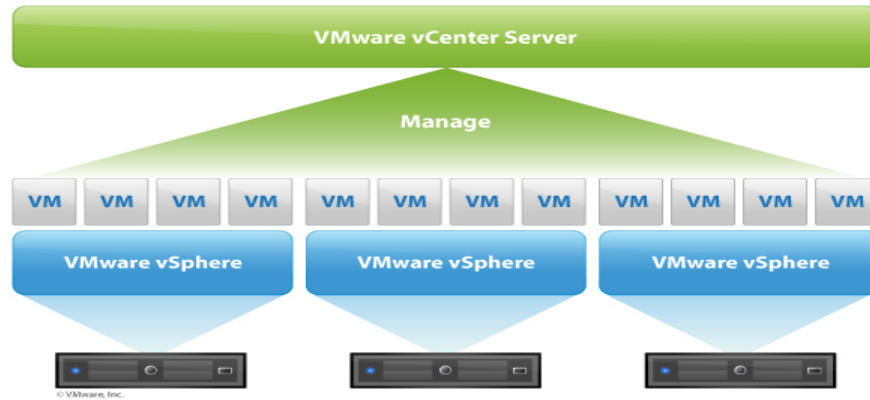


Figure 13. VMware vCenter Server hierarchy [32]

3. VMware vSphere

VMware vSphere is the VMware infrastructure platform designed specifically for cloud computing when combined with other VMware products, such as vCenter Server and ESXi. Specifically, vSphere “manages large collections of infrastructure, such as CPUs, storage, and networking, as a seamless and dynamic operating environment, and also manages the complexity of a datacenter.” [10] As it relates to the cloud computing paradigm of shared resources, vSphere “virtualizes and aggregates the underlying physical hardware resources across multiple systems and provides pools of virtual resources to the datacenter.” [10]

VMware vSphere is not a single program, rather it is a software stack; i.e., it is a composition of virtualization, management, and interface layers (illustrated in Figure 14). Describing the layers as it pertains to the setup in the Cloud Lab, VMware ESXi is the virtualization layer and vCenter Server is the management layer, each comprising a level of the vSphere software stack. The interface layer is the last layer to be built, not materializing until the last step of the virtualization infrastructure was completed. It

appears at the end user segment in which a user accesses a virtual machine via a thin client or a zero client.

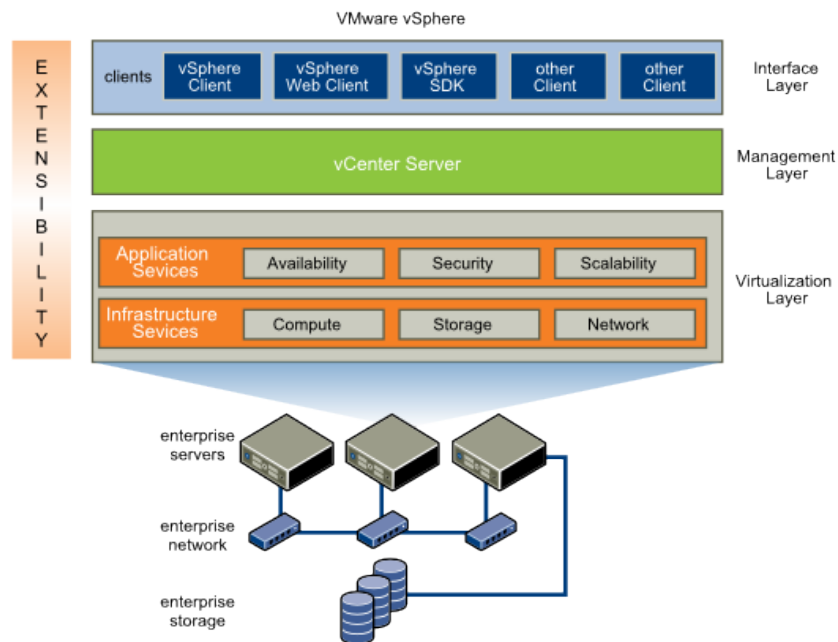


Figure 14. The three layers of the vSphere software stack [10]

In addition to ESXi and vCenter Server there are several key features of vSphere that when leveraged provide access to crucial resources when needed. These features together form one of the core characteristic of the cloud computing paradigm – elasticity. Each of the following vSphere features were installed and running on the ESXi servers in the Cloud Lab.

a. vMotion

When this feature is installed a powered-on virtual machine can be migrated from one physical server to another with zero downtime. In the event a server needs to be taken down for maintenance, any virtual machine in use by a user will be seamlessly transferred to another server. Continuous service availability and complete transaction integrity are guaranteed. This prevents any user from having to wait to continue working.

b. High Availability (HA)

In the event a server fails the High Availability feature will automatically restart any affected virtual machines on another server. Similar to the vMotion feature, having HA provides users continuous access to a virtual machine.

c. Fault Tolerance

With this feature enabled a secondary copy of a virtual machine is created and maintained. Continuous availability is provided in the event the original virtual machine becomes unavailable. Fault tolerance is achieved by combining HA with distributed resource sharing.

4. Hosts, Clusters, and Resource Pools

Important to allocating resources are hosts, clusters, and resource pools, which are quantities of shared computing, memory, and storage resources; Figure 15 graphically shows the formation of hosts, clusters, and resource pools. By aggregating resources a uniform set of elements (an element being a host, cluster, or resource pool) are created in the virtual environment. [33] The management of these elements is done via vSphere and can be managed like a shared utility and dynamically provisioned out to the virtual machines demanding resources. This sharing of resources is an example of elasticity and on demand resource sharing, key characteristics of cloud computing.

In the Cloud Lab, several hosts (the AEGIS servers) and a multiple of clusters and resource pools were used. Each of these elements is described below, followed by the setup of the elements as used in the virtual infrastructure.

a. Hosts

A host is simply a computer that utilizes virtual software to run VMs. Hosts provide the resources such as processing power and memory to the VMs. They also provide the VMs with access to storage and networks. A host can support hundreds or thousands of VMs, depending on the amount of hardware a host has.

Hosts can be added to the vCenter Server for management. Five hosts were added to vCenter Server; AEGIS-1 through AEGIS-5, which are the servers for the entire VDI in the Cloud Lab.

b. Clusters

Clusters represent the aggregate computing power and resources of several hosts sharing the same network and storage arrays. Although a cluster draws resources from several components, it acts and can be managed as a single entity. [10]

c. Resource Pools

Resources of a host or of a cluster can be divided into resource pools. Resource pools are self-contained and can be isolated from other resource pools on the virtual infrastructure. VMs are assigned resources from resources pools. The benefit here is that, for example, one group of VMs assigned to a set of users whose work requires a small amount of resources can be assigned a resource pool equating to their need, whereas a different group of users may need a larger amount of resources to complete their work and have their VMs assigned to a resource pool with the requisite resources needed. Each resource pool can draw from multiple hosts (physical servers), using an aggregate amount of the total resources available.

Resource pools are also dynamic. When a resource is not used by a VM that resource is free to be shared among other VMs. Resource pools can be nested and organized hierarchically, allowing for dynamic reconfiguration among any other running VM. [33]

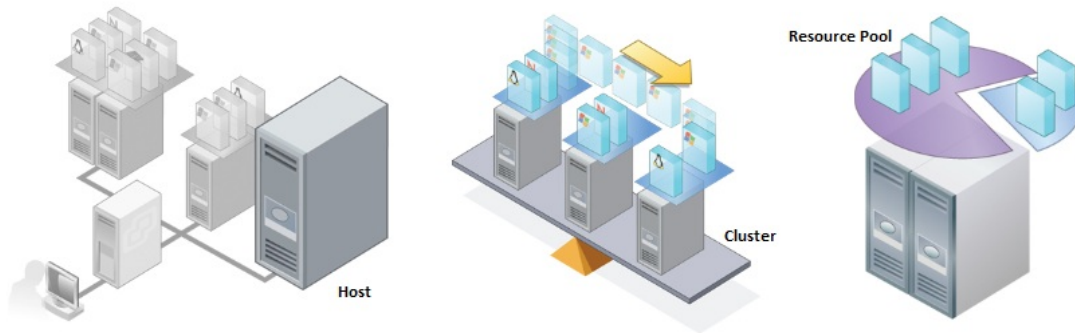


Figure 15. Hosts, Clusters, and Resource Pools [10]

5. Virtual Networks

The network portion of a virtual infrastructure forms the backbone along which information transits and involves both physical and virtual components. The virtual infrastructure created in the NPS Cloud Lab had its own virtual network, equipped with virtual switches and virtual ports, and was connected to the physical NPS ERN domain. Because the VDI was connected to the NPS ERN domain, access to virtual machines was possible from anywhere via a computer terminal with the VDI client software installed.

A virtual network is very similar to a physical network. A physical network consists of a number of physical computers connected via a switch. Switches manage traffic between computers and form the core of a physical network. Switches have several ports, which connect to a computer or another switch. Ports can be configured to match the needs of a computer connected to it.

Whereas a physical network has several computers connected to each other via a switch, a virtual network has several virtual machines running on a single physical computer and is managed by a virtual switch, which is created during the initial construction of a virtual infrastructure. VMware vSphere contains its own switches, called vSphere Distributed Switches.

a. vSphere Distributed Switch

A vSphere Distributed Switch (VDS) is similar to a physical network switch in that it connects network segments within Local Area Networks (LANs). The virtual switch detects the virtual machines connected to the switches' virtual ports and one VDS can span multiple ESXi hosts. The benefits include a reduction of network maintenance activities and an increase in network capacity and the resulting increase in efficiency enables VMs to maintain consistent network configuration as they move among hosts. [33]

Each virtual port on a VDS can be grouped together, forming a port group. These groups specify configuration options, including bandwidth limitations, and define how connections are made through a VDS to a network. The virtual network connects to a physical network by connecting a VDS to a physical network using uplink adapters.

6. Storage

The storage setup used for the virtual infrastructure in the Cloud Lab included the hard disk drives on the AEGIS servers and the hard disk drives that comprised the Storage Area Network (SAN). Storage available on the AEGIS servers was reserved specifically for applications installed directly on the server (such as ESXi). Storage for users of virtual machines resided on the SAN. A SAN is a dedicated network providing access to storage devices accessible by a server. Suitable for enterprise networks, a SAN can consolidate storage space in much the same way blade servers consolidate servers. For example, a desktop computer has its own internal HDD. When zero clients replace desktop computers, which do not have internal HDDs, all storage resides on a SAN. A SAN is one physical unit containing multiple hard drives, much like how a blade enclosure contains multiple blade servers.

The SAN used in the NPS Cloud Lab contained 16 SCSI HDDs. Each HDD had 500GB of storage space. All 16 HDDs combined equaled 8TB of storage. The SAN was partitioned into three separate resource pools, with two resource pools each having 2TB of storage and the third resource pool having 1.5TB.

7. VMware View

VMware View was the last virtual software program to be installed in the VDI. There are two versions of View, View Manager and View Client. View Manager is a universal client solution that lets an administrator manage every aspect of a deployed virtual machine. Operating systems, hardware, applications, and independent users can be managed from anywhere, regardless of where a user accesses a virtual machine. View Client is the end user segment in which a user accesses a virtual machine from a pool of available virtual machines, and is discussed in the next section of this chapter.

Of all the VMware applications used in the NPS Cloud Lab, VMware View essentially offers a unique cloud computing service: delivery of user desktops as a service. View accomplishes this by delivering desktops and applications securely to remote clients anywhere, resulting in a highly available, agile cloud service.

View has a powerful tool that makes the creation of virtual machines easy – View Composer. View Composer is designed specifically for making virtual machines in large-scale environments. In View Composer, a customizable template can be created and tailored to the needs of an enterprise. Once this template is created it becomes a parent (standard) virtual machine image. This template then allows for the rapid and automatic replication of itself to create multiple virtual machines (see Figure 16). By using linked clones, updates need only be applied to the parent image. All other virtual machines that are an image of the parent image can then be updated in a streamlined manner.

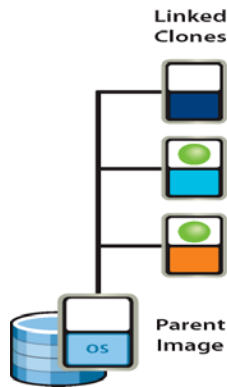


Figure 16. Parent image with linked clones [10]

8. Virtual Machines and the End User Interface

The last step in building the virtual desktop infrastructure was to create VMs using View Manager and to test them via an end user interface. To create the VMs a template was first made tailored to the needs of the audience that would be using the VMs on the NPS campus (mainly NPS students).

With the template constructed the replication process was simple to execute. After the desired number of VMs was entered (50, in this case), View automatically created the VMs based off the customized template. Each VM was linked to an account in the Active Directory, which allowed a user with an NPS account to log-in to a terminal with the VMware View Client installed and access a VM. View Client is the end user's gateway to a virtual machine and is the software portion of the interface layer of the vSphere software stack.

The hardware portion of the interface layer can be any computer client device such as a laptop, a desktop computer, a mobile device (smart phones), or a thin client or zero client. In the Cloud Lab PCoIP zero client devices, manufactured by Wyse, were used. The virtual machines were delivered using the PCoIP protocol and presented to the end user on a thin client device or a zero client device. All virtual desktops were centrally located and managed in the data center by vSphere and View Manager, while the end user accessed a virtual machine through View Client on the user's client device.

The VDI was tested using Wyse P20 zero client devices (see Figure 17). Wyse is a global leader in cloud client computing devices. The P20 is a small zero client designed to be used with VMware View and contains very little hardware; 128MB RAM, a small processor, video card, and a NIC. Several audio, video, USB, and network ports are integrated into the P20, offering all the connections necessary for a mouse, keyboard, and monitor, as well as the connection to the host servers that are home to the VMs. Overall, the P20 is compact, energy efficient, and “delivers a rich user experience while resolving the challenges of provisioning, managing, maintaining and securing enterprise desktops.” [34]

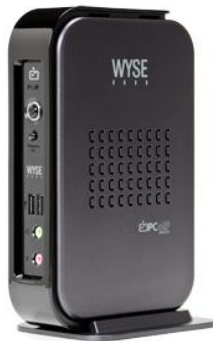


Figure 17. A Wyse P20 zero client device [34]

The VDI was successfully accessed, not only from the Wyse P20 zero client, but also from standard desktop computers and laptop computers, located on and off the NPS campus. The PCoIP protocol was implemented successfully; VMs were delivered over the internet to the client devices and ThinApps and standard applications streamed uninterrupted from the host servers in the datacenter to the VMs. Figure 18 is an

illustration of the VDI's physical connections between the Blade Chassis and Blade Servers, the Zero Clients, and the ERN and the Internet. Figure 19 depicts the overall VDI, including each AEGIS server virtual layer composition, and the physical and virtual connections between the AEGIS servers and the View Manager Workstation, View Client Devices (zero clients), the ERN, and the Internet.

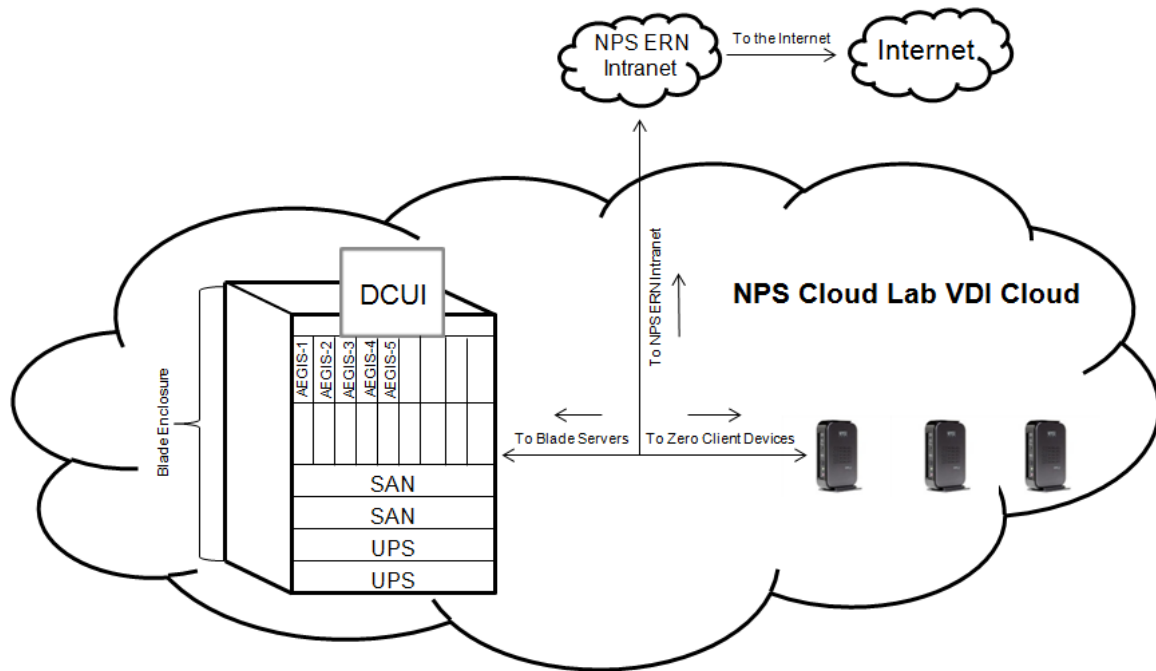


Figure 18. NPS Cloud Lab Blade Chassis and VDI physical connections

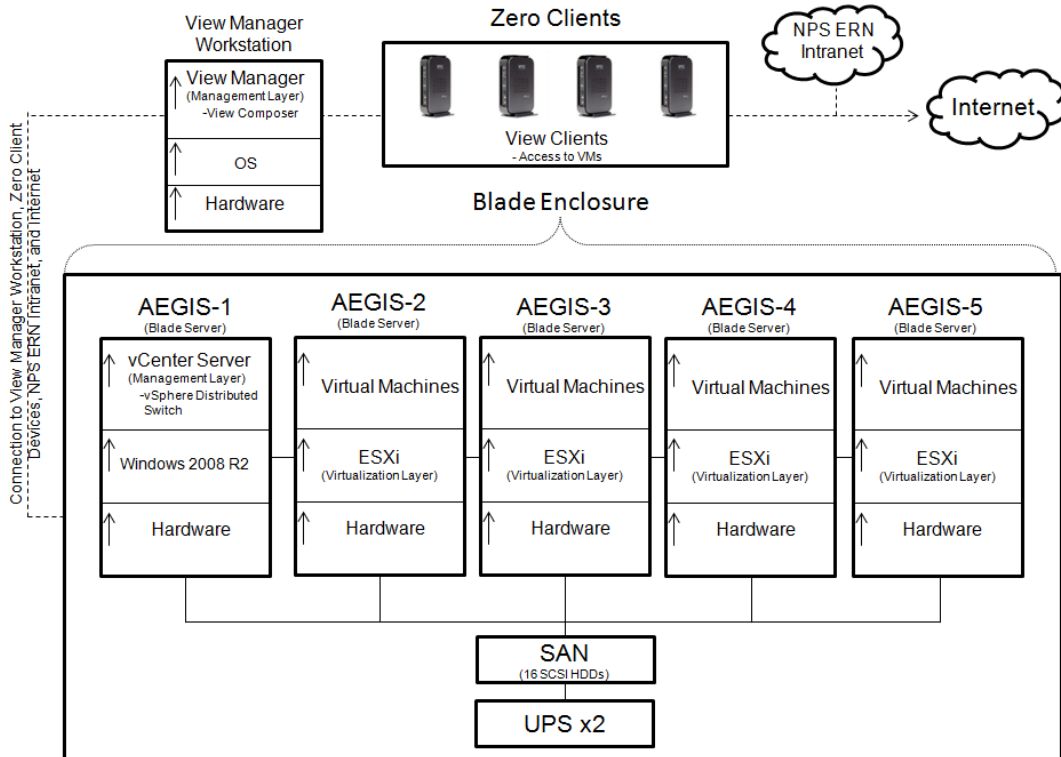


Figure 19. NPS Cloud Lab VDI physical and virtual connections

C. SHIPBOARD VIRTUAL DESKTOP INFRASTRUCTURES

Based on the VDI created in the Cloud Lab it is possible to conceptually scale it to construct what a VDI would look like onboard a ship. The major differences lie in the size of the infrastructure, the hardware requirements, and the level of application the VDI would be subjected to. The VDI proposed here is not for tactical applications, such as the AEGIS weapons system on an Arleigh-Burke class destroyer or on a Ticonderoga class cruiser. Rather, it is designed for use with non-tactical applications, such as standard email programs and servers, Microsoft Office products, administrative programs, and applications like those delivered by NIAPS (NIAPS – or Navy Information Application Product Suite – delivers “maintenance, logistics, administrative, training and management applications to users at sea.” [35]) Tactical applications could also be considered conceptually as another implementation of the NPS Cloud Lab VDI.

The technology standards required to run the VDI as built in the Cloud Lab are not different from those currently found on shipboard network infrastructures. The only

exception is the relatively new logical internet protocol PCoIP, which is integrated into VMware software and requires nothing more than using zero client devices capable of utilizing the protocol (such as the Wyse P20 zero client used in the Cloud Lab).

The physical requirements necessary to run a VDI on a ship would be relative to the size of the ship and the number of users onboard. The requisite amount of CPUs, HDDs, RAM, and other components for the host servers on which a VDI would reside would have to be determined prior to consolidation of the ship's existing servers. One method would be to total the hardware of the rack-mounted servers onboard the ship and use the resulting calculated figures as the minimum amount of hardware necessary to run the ship's networking infrastructure using blade servers. The physical resources for the VDI would then be added by expanding the number of blade servers after the ship's servers have been consolidated. Figure 20 illustrates the consolidation of physical servers into one physical server, with an instance of the ESXi hypervisor and virtual machines installed.

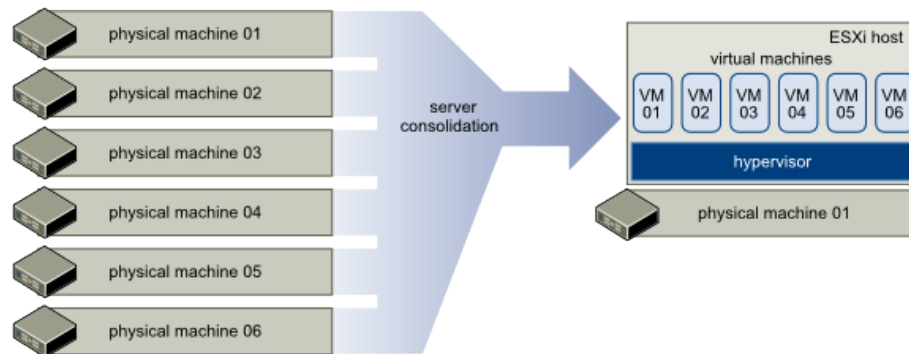


Figure 20. Physical server consolidation [10]

Requirements for individual virtual machines depend on the number of users. A ship with 300 personnel would require 300 VMs. However, a spare VM should be available to every user in the event a spare VM is needed. Therefore 600 VMs is a fair number of VMs for a ship with 300 personnel on board. To adequately determine the resources needed for 600 virtual machines a test VM should be created first. This test VM should have installed on it all of the applications that would be used onboard the ship. Stress tests should then be conducted to determine the physical resources necessary for

the VM to operate efficiently and satisfy user needs. After the test VM has met these conditions the physical resources that were required to sufficiently run the VM can be multiplied by the number of desired VMs to determine the overall requirement of physical resources to run the VDI.

For example, if 3GB of RAM and 300Hz of processing power was the minimum required to successfully run the test VM, then 600 x 3GB of RAM and 600 x 300Hz of processing power is necessary to support 600 VMs (or 1,800GB (1.8TB) of RAM and 180,000Hz of processing power - the equivalent of 60 3GHz CPUs). A small number of blade servers are more than capable of handling this amount of physical resources.

Storage requirements depend on the combined size of the applications plus the amount of storage administrators want to allocate to a user. The size of a SAN for 600 users can be calculated in the same fashion as the previous example to determine RAM and processing power. If, for example, 10GB of storage were to be allocated to each user and there are 300 users, then 3,000GB (3TB) would be a sufficient amount of storage. For comparison, the SAN in the Cloud Lab had the equivalent of 5,500GB (5.5TB) of storage.

It is important to remember that one of the prime reasons for moving to a cloud architecture with a virtualized infrastructure is to capitalize on elasticity. While it is unlikely 100% of resources will be used at any given time, additional resources to the 100% must be available if the need arises. A VDI integrated into a shipboard network infrastructure designed for 300 users but with the capacity to handle 600 users should be able to accommodate a sudden spike in demand for resources. However, ships are not equipped with a number of workstations equivalent to the number of personnel on board. Though a ship such as an Arleigh-Burke class destroyer may have 300 personnel, there are approximately 250 workstations throughout the ship. In this case, it is only possible for 250 of the 600 VMs to be used at any given time, and even then that is only if every workstation was being used, which is unlikely.

The replacement of desktop computers with zero clients like the Wyse P20 would account for much of the creation of elastic resources. Desktop computers contain their

own hardware and therefore their own resources. Those resources are rarely, if ever, fully utilized. A VM, on the other hand, pulls its resources from resource pools on the host servers and only uses the resources necessary to run the VM itself and any applications installed on it. Any remaining unused resources in the resource pools are available for other VMs to use. The resources on desktop computers are wasted because it cannot be used by another computer if there is a demand; i.e., desktop computers are not scalable. In a VDI resources are scalable; they are available to other VMs when in demand. If demand for resources is low enough to meet a set threshold then the host server(s) can automatically shutdown to preserve power consumption.

Based on the examples given in this section and the VDI built in the NPS Cloud Lab as discussed in this chapter, a model shipboard VDI requires a small number of phases to build and integrate into a shipboard network infrastructure. There are primarily three phases. The first is to consolidate servers currently existing onboard a ship and to add the requisite number of servers on which a virtual infrastructure would reside. In this first phase, analysis of the requisite number of servers must be done. Blade chassis and blade servers are suitable solutions for server consolidation. The second step is to construct a VDI and install on it the applications to be used in the enterprise. Using the VDI powered by VMware, as built in the Cloud Lab, offers a simple approach. VMware's software products are designed specifically for enterprises seeking a cloud computing infrastructure with virtualization at its core. The last step is to replace a ship's desktop computers with zero clients. Zero clients reduce the footprint of physical space throughout a ship and are easily managed.

A VDI designed to meet user needs and properly integrated into a ship's network infrastructure offers a cloud computing enterprise network foundation that is agile, elastic, and scalable. Resources are only utilized when in demand and are preserved when not in demand. Desktops are streamed to users over the shipboard network, delivering a unique cloud service – desktop-as-a-service. Management is centralized in one server that presides over the entire VDI. Conducting maintenance on physical servers and applying patches and updates to the VDI can be done without affecting the availability of VMs to

users. In sum, a VDI incorporating these beneficial characteristics is a suitable foundation for a shipboard level cloud computing architecture.

Figure 21 shows a screenshot of the vSphere Client, which provides status of the overall condition of the VDI, including the availability of resources and the state of the virtual machines (powered on/off). Note the amount of processing power and RAM available in the resource pool for the 50 virtual machines residing on the AEGIS servers.

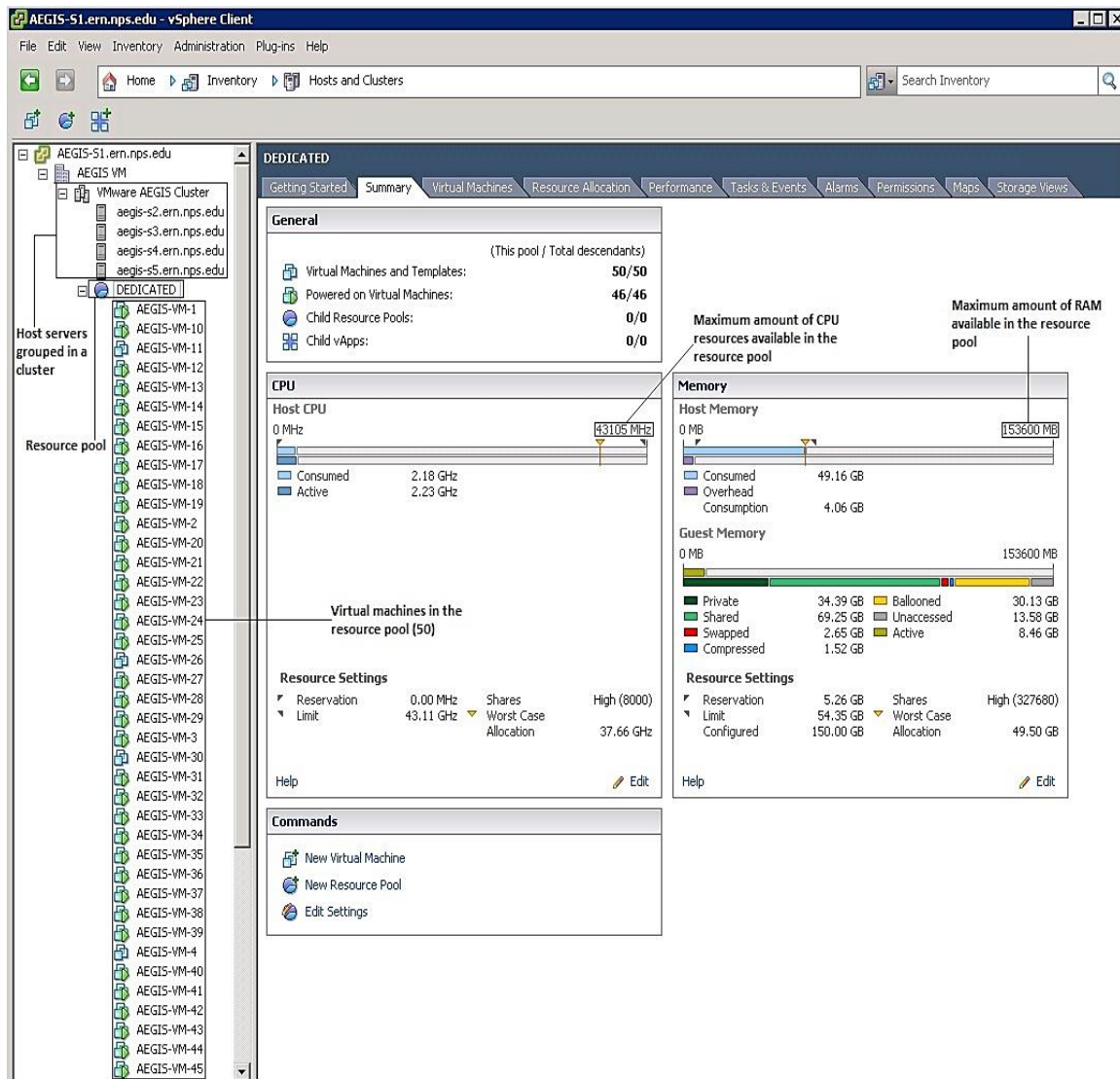


Figure 21. Screenshot of the vSphere Client manager screen

IV. CLOUD COMPUTING CONCEPT OF OPERATIONS

The previous chapter focused on virtual desktop infrastructures as the foundation of a cloud computing architecture on a shipboard level. A VDI was built in the Cloud Lab at NPS to give insight into what a virtual infrastructure looks like, the steps to design and implement it, how it operates, and the requirements to operate it. This chapter looks at the implementation of cloud computing in afloat multi-ship environments. For the physical cloud computing infrastructure in a multi-ship environment it is assumed that each ship has its own private cloud, outfitted with the virtual desktop infrastructure designed in the previous chapter.

An underlying yet important benefit of cloud infrastructures in multi-ship environments is the potential minimization of satellite communications use by ships in a strike group, which may result in a reduction in bandwidth consumption. Though the push and pull of data between a strike group and shore facilities via satellites is for the most part unavoidable (especially when operating in distant blue waters), the ability to link each ship's cloud with every other ship in a strike group offers an alternative to reliance on satellites. Mobility of cloud services and the mobility of clouds themselves (i.e., the ships) is naturally key to the afloat cloud infrastructure. To support this “cloud on the move”¹ concept, assets linking shipboard clouds in the afloat environment will be a crucial role if data is to be shared between clouds. The next section proposes several methods of linking clouds.

A. LINKING CLOUDS

Numerous possible configurations for an afloat cloud infrastructure in multi-ship environments exist given the various assets and communication methods within the DoD's inventory. This section identifies and discusses configurations using assets capable of acting as links for clouds afloat, such as lighter-than-air (LTA) ships, fixed

¹ “Cloud on the move” was coined by Albert Barreto while conducting research for this thesis in the NPS Cloud Lab. Credit is given to him here for the phrase.

wing aircraft, IP routable satellites, and Worldwide Interoperability for Microwave Access (WiMAX) communications.

The current typical strike group composition consists of an aircraft carrier, an oiler, a submarine, and a combination of a small number of surface combatants that can include cruisers, destroyers, and frigates. In a cloud infrastructure the ship with the best communications equipment and technology would serve as the “hub” of the cloud network, which in the case of a strike group the role would be fulfilled by the aircraft carrier. Each additional ship in the strike group then becomes a “node” or “edge cloud.” Though each ship has its own unique cloud, when linked together they form one composite cloud capable of sharing data.

When ships are within Line-of-Sight (LOS) data can be shared via UHF LOS or even HF, but these communication methods have very low data rates and would not meet the large bandwidth requirements to exchange data from cloud to cloud. When ships operate beyond LOS data is shared over satellite communications (SATCOM). In order to reduce satellite saturation and free bandwidth for other uses, alternatives to satellites are needed.

Figure 22 depicts a basic example of individual shipboard edge clouds linking to a hub. In this instance, the hub is the uplink and downlink of data from satellites. The hub receives data both from shore facilities via satellite and from each edge cloud via communications other than satellite. Upon receiving data, the hub processes and disseminates the data to the rest of the strike group, similar to a satellite broadcast. In a way, the hub conducts the push and the pull of data for the strike group. The benefit here would be that ships other than the hub in the strike group do not have to utilize a satellite, therefore freeing satellite bandwidth that would have been consumed by each ship. The infrastructure just described is one of many possible infrastructures. Listed below are technologies that may be suitable for linking clouds in a multi-ship environment.

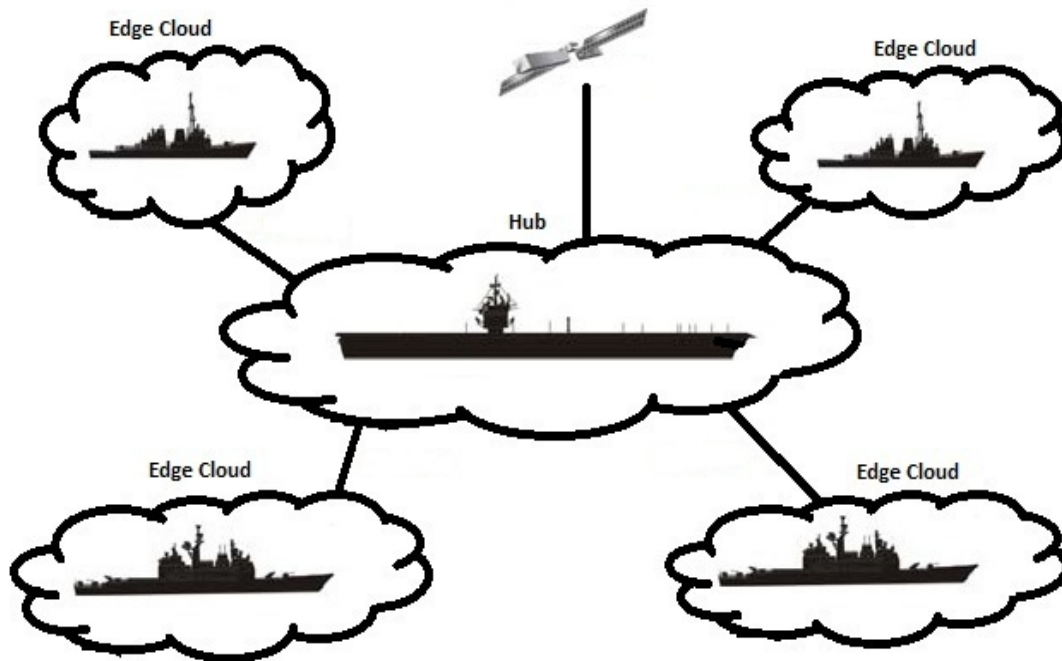


Figure 22. Multi-ship Afloat Cloud Infrastructure

1. LTA Ships

Lighter-Than-Air ships are lightweight, inexpensive, helium filled dirigibles. They have been in use in the US military for decades, but only recently has a case been made to reintroduce them into the fleet as a useful asset in maritime operations. LTA ships are capable of heavy strategic lift, require very little manning, have lengthy on station durations (up to 30 days), and most important in today's military they are considerably cheap relative to other ship programs of record. [36]

More relevant to this thesis is the potential LTA ships have in cloud computing infrastructures in afloat environments. An LTA ship, when outfitted with the requisite equipment, can act as a node in a cloud network, as illustrated in Figure 23. As a node, an LTA ship can serve as a forwarding station, linking assets in a strike group when they are beyond LOS; though, the LTA ship itself would have to be within LOS of the ships. Operating at high-speeds (70-90 knots) several thousands of feet in the air, able to withstand unfavorable weather conditions, and a low level of vulnerability to enemy fire,

LTA ships are a versatile solution to extending a cloud infrastructure and linking nodes within a cloud network over great distances. The same technologies found on satellites can also be installed on an LTA ship, able to accommodate the physical small size of satellite systems. Additionally, LTA ships can “readily accommodate common C4I architecture for maximum interoperability,” “enhance data sharing across multiple applications and domains,” and can be “scaled/tailored to suit lift, endurance, sensor and communication requirements.” [36]

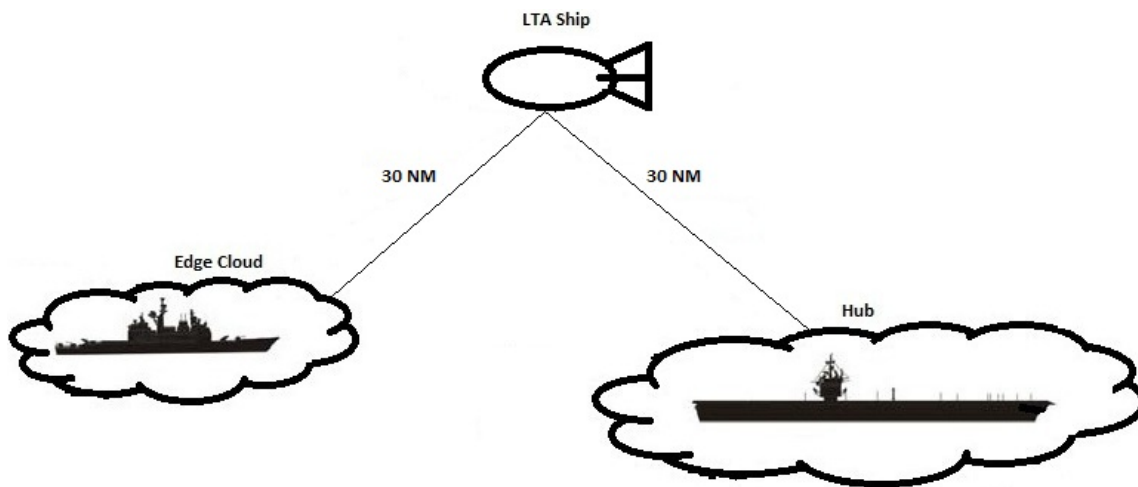


Figure 23. LTA Ship acting as a data relay

2. Fixed-wing Aircraft

The use of fixed-winged aircraft, particularly unmanned aerial vehicles, is a possible solution to extending edge clouds beyond LOS. These systems have capabilities similar to LTA ships and have very high data rates, but their on-station duration is relatively short given they must refuel often. A current POR likely capable of meeting the requirements to link shipboard clouds to a hub is the Broad Area Maritime Surveillance Unmanned Aircraft System (BAMS). BAMS provides “persistent maritime Intelligence, Surveillance, and Reconnaissance (ISR) data collection and dissemination capability,” and it functions as a communications relay. [37]

BAMS and other maritime aircraft such as the P-3 Orion or P-8 Poseidon can act as a communications relay of data from cloud to cloud. Unlike LTA ships, these data

relays would likely be used in high-intensity situations where LTA ships would not be suitable for use. Except for in those cases, the role of a data relay would be a secondary mission for fixed-wing aircraft unless the aircraft is designed specifically to link distributed cloud nodes at sea.

3. IP Routable Satellites

Though one of the primary reasons for a mobile cloud infrastructure is to operate in an environment devoid of satellites, IP routable satellite systems are a unique alternative to other satellites because they use the IP protocol. Many cloud services are web based services and are already designed to function using the IP protocol suite. Two recent and emerging technologies in this category are the Broadband Global Area Network system and the Internet Routing in Space satellite system by Cisco Satellite Solutions.

a. Broadband Global Area Network

The Broadband Global Area Network, or BGAN, is a global satellite Internet network system. The system utilizes the Internet Protocol suite to deliver internet services to end users. Typical applications include email, internet and intranet access, secure VPN, VoIP, file transfer, and remote surveillance, among others. [38] BGAN is a viable and suitable solution to linking nodes in an afloat cloud infrastructure for several reasons.

First, BGAN terminals are portable devices small enough to fit inside a backpack. Compact terminal size is a plus when considering network footprint reduction, especially when they are to be installed onboard ships or aircraft. Second, the IP streaming rates range from 32kbps to 492kbps. Although these speeds are less than desirable, a BGAN satellite would suffice given the only connection they are providing is a relay of data between cloud nodes (ships). Third, BGAN is delivered by satellites that are part of the INMARSAT (International Maritime Satellite Organization) satellite system, which has been in use by the US military for decades. Lastly, BGAN benefits from the security built into INMARSAT, a satellite system that has proven its security over time.

b. Cisco Internet Routing in Space

Cisco's Internet Routing in Space (IRIS) systems, like BGAN, offers Internet routing via satellites. Cisco claims their satellite solutions deliver “scalable service offerings for enterprise networking, broadband and backhaul applications.” [39] These satellite solutions may also be utilized as a data relay station between mobile cloud nodes (ships) in a strike group. A recent Joint Capabilities Technology Demonstration (JCTD) of the Cisco IRIS program was successful in providing end-to-end security and connectivity. [40]

IRIS runs on the Intelsat satellite system. Capable of delivering up to 50Mbps [41], these secure satellites are powerful and may adequately fulfill the role of a data relay in an afloat cloud infrastructure. Figure 24 provides an illustration of the IRIS system servicing various afloat and ashore end-users.



Figure 24. Illustration of Cisco's Internet Routing in Space servicing remote users [40]

4. Worldwide Interoperability for Microwave Access

Possibly the best method of transferring data between nodes in an afloat cloud infrastructure is via WiMAX (Worldwide Interoperability for Microwave Access). WiMAX is part of the 4G (Fourth Generation) mobile communications standards, capable of providing ultra-broadband Internet access over wireless communications at speeds of up to 70Mbps. [42] Of particular importance to afloat cloud infrastructures is the distance WiMAX offers, up to 30 miles.

For the relay of data to work between nodes in a strike group WiMAX antennas or towers would have to be installed on each ship (see Figure 25). For example, one WiMAX tower (capable of providing coverage to areas as large as 3,000 square miles [43]) would be installed on the hub, and smaller receive/transmit antennas on each ship. An example of this configuration is depicted in Figure 26.

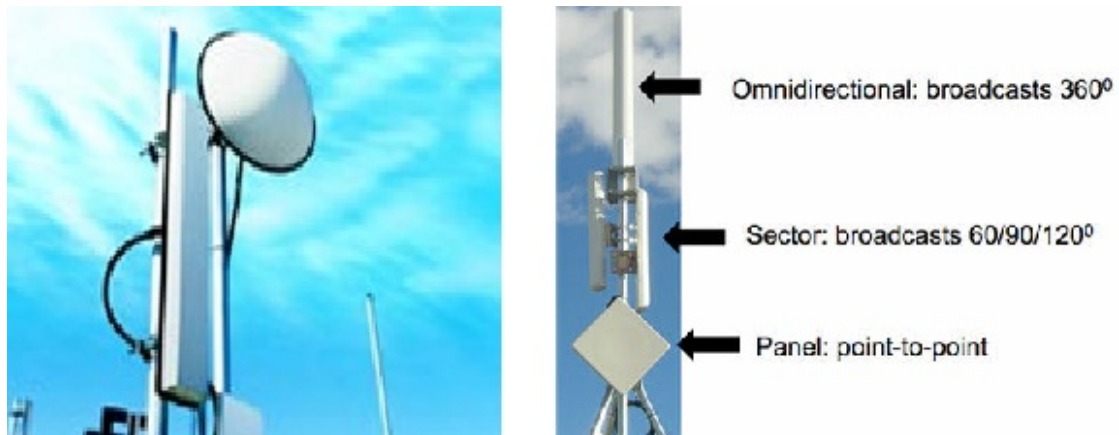


Figure 25. A WiMAX Tower (left) and a WiMAX Antenna (right) [43], [44]

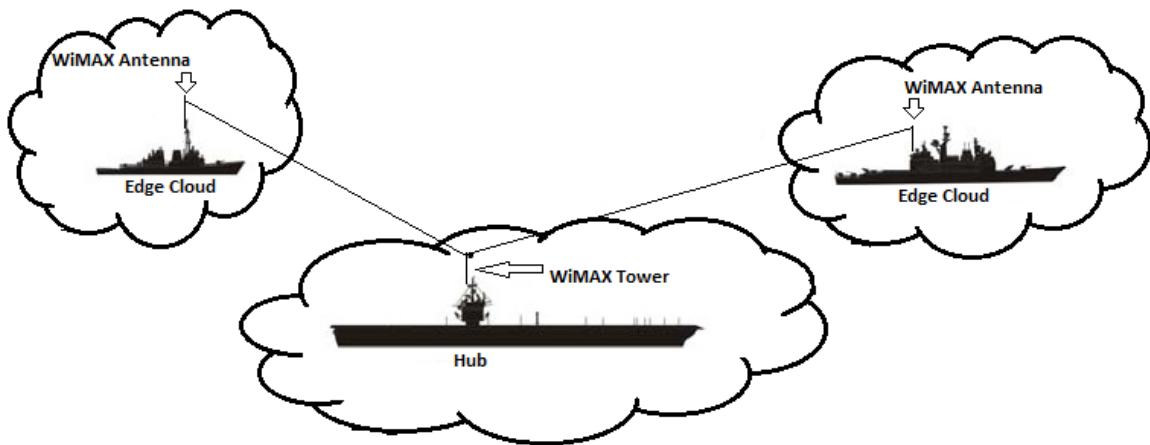


Figure 26. Possible WiMAX infrastructure

A message released in April, 2012 titled “Naval Air Systems Command plans 4G cell service aboard ships” is a promising step towards the integration of WiMAX in afloat network architectures. [45] The Navy plans on testing 4G LTE (Long-Term Evolution) broadband communications capabilities in 2013 on three ships. Key features include an initial throughput of 8 to 15 Mbps with an eventual throughput of 1Gbps as technology

matures, a range of 30 miles, and most relevant to an afloat cloud architecture the 4G systems “will support the exchange of a variety of broadband data while ships are underway, freeing up expensive and limited satellite bandwidth.” [45]

Whether it is the use of LTA ships or WiMAX, ships in an afloat cloud network should use technologies other than satellites to link the cloud infrastructure between ships. This should be done primarily for two reasons; (1) the latency and bandwidth characteristics found in technologies such as LTA ships or WiMAX are superior to those found in satellites, and (2) current and future warfare doctrine assumes the U.S. military will be fighting in a degraded or denied satellite environment.

Each ship, equipped with its own private cloud, will need a connection that is fast (low latency) and has a bandwidth large enough that is capable of delivering large amounts of data between clouds. The technologies listed above are potential solutions to linking clouds, allowing for data to be shared between clouds. Just as important as how data is shared is where data is stored, discussed in the next section.

B. CLOUD DATA STORAGE

Where data is stored will dictate what an afloat cloud infrastructure looks like more so than the communications method used. The method of communication simply links clouds and delivers the data from its storage location. Where and how data is stored depends on several factors, such as the technology used, the size of the datacenter, the volume of data, the sensors collecting data, etc. One of the five essential characteristics of cloud computing is on-demand self-service, part of which includes provisioning network storage. Using a VDI like the one built in Chapter III provides each cloud its own data center capable of provisioning resources, including storage. This is acceptable for each ship individually, but information that needs to be shared among a strike group will need to be stored in a location accessible by all clouds.

The PEO-C4I Framework for Cloud Computing at the Tactical Edge report contains the DON Data/Information Life Cycle, which defines the process of the entire life cycle of data, from identifying the need for data to data disposal. [27] This detailed data process can be used to help determine where data should reside in an afloat cloud

computing multi-ship environment. Table 1 lists key steps in the process, providing insight as to the requirements of where data may be stored.²

Life Cycle Phase	Definition
Identify Need	Identify the potential container(s) for the specific data/information need for storage, processing, transport, and management
Identify Architecture	Identify options for data cache/staging/storage location(s)
Task Collection/ Creation	Identify the application(s), system(s), service(s), and/or network(s) that will use, process, store, and/or transport the created data/information
Collect/Create	Determine the accessibility requirement for the data/information asset(s)
Disseminate	Transmit, publish-and-subscribe, or broadcast data/information assets to applicable customers (applications, systems, services, or personnel)
Store	Store the data/information asset
Dispose	Destroy data/information asset(s) that are no longer needed to support the established requirement or have exceeded their time period for retention

Table 1. PEO-C4I select steps of the Information Life Cycle Process [27]

Based on the definitions of the life cycle phases listed above there are two phases most responsible for dictating where data storage should take place. The first is Identify Need. Where data will be stored will depend on which edge cloud actually needs the identified data. The second is Task Collection/Creation. Several factors can be incorporated into this phase. For example, if the information needed is intelligence based then the data must be processed at a node that has the systems to do the processing. In a strike group, intelligence data processing takes place on the aircraft carrier. If the raw data initially came from the sensors on a destroyer then the data is transmitted from the destroyer, processed at the hub (aircraft carrier), and made available in a data center residing on the hub. The raw data, now processed into usable intelligence information, may be needed by a cruiser to launch a surface missile. Now the data must be made available to the cruiser, which would be pushed from the hub.

² The table is a redacted form of the table from [27], Appendix D.

Generally, data would reside on the hub because the hub will have the largest data center(s) and the hub has the systems to process raw data into intelligence. Edge clouds would collect data from their sensors and transmit it to the hub, and the hub in turn processes and stores this data, making it available to the entire strike group. Figure 27 illustrates this process. (For concepts on metadata formats for multi-INT data extraction from sensors, that could be stored in these clouds, see “Metadata Standards for Multi-INT Fusion in Distributed ISR Systems” by LCDR Trevor Day, a former NPS student.)

Data created within each edge cloud will be stored internally in its own data center. That is, basic data such as user accounts and records, emails, PowerPoints, and any other typical data encountered on the regular day to day operations of a ship that can be processed via the CRUD computer programming functions (Create, Read, Update, and Delete).

Data received from outside the hub’s organic/inorganic sensors (i.e., from a shore facility via satellite) – including the edge clouds – will be ingested and stored in a data center in the hub. The hub can then transmit the data to the requisite edge clouds as needed over whichever method of communication links the strike group is operating on (such as WiMAX). Again, having the hub as the sole satellite receiving node is done to reduce satellite use.

This does not imply other ships will not use satellites. When in a multi-ship environment, however, having one ship with the capability to receive data, store it, and disseminate it to linked edge clouds will inhibit satellite saturation. This is achieved by making data available in one centralized location (the hub), creating a situation in which other ships will not have to access the data through a satellite.

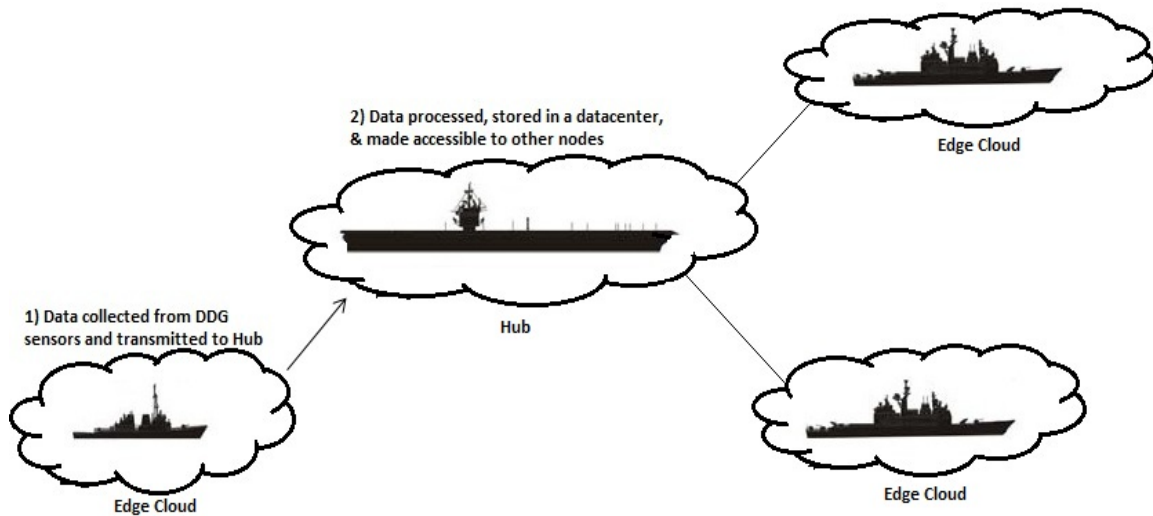


Figure 27. Possible data storage process afloat

C. CLOUD-TO-CLOUD INTEROPERABILITY

System-to-system interoperability has been and continues to be a challenge for the Navy and the DoD. New information systems may not be designed to integrate with other systems, resulting in “stove-pipe” systems. A stove-pipe system does not share data or functionality with other systems. As a result data may have to traverse several networks or nodes in order to reach its destination, or the receiving system may have to rely on a completely separate system to convert data into a format readable by the intended recipient system. The route data travels or the number of times data must be converted as it travels from origin to destination can be an unnecessary and preventable waste of time and resources.

The sharing of data between information systems can be referred to as portability. Portability of data between systems is directly related to the mobility of the same data between systems. Interoperability between systems dictates the portability and mobility of data. Cloud computing is highly dependent on the portability and mobility of data, hence the necessity of loose coupling between systems in any cloud computing infrastructure. To ensure the sharing of data from one cloud node to another cloud node in a multi-ship environment, cloud to cloud interoperability needs to be addressed.

1. Levels of Information System Interoperability Maturity Model

A relevant model designed to determine and classify levels of interoperability between systems is the Levels of Information System Interoperability (LISI) Maturity Model by the DoD C4ISR Working Group. [47] The LISI model is a “formal Reference Model, an assessment implementation of the interoperability maturity model, and a structured process for improving interoperability between varied information systems.” [48] As such, the LISI model provides the “how to” in systems architecture development and in linking systems. [49]

Four attributes of the LISI model define how data is shared between systems: procedures, application, infrastructure, and data (aka PAID, see Figure 28). The procedures attribute refers to the degree of interoperability as a result of policies, processes, and the compliance of technical and system architecture standards. [50] The application attribute defines the ability of software applications to work on different systems, while the infrastructure attribute defines the degree of connectivity between systems and applications. The data attribute “reflects the flexibility of the data format and the richness of the information being exchanged across systems and domains.” [50]

Furthermore, the ISIL model defines five maturity levels of system to system interoperability. They are listed below. [49]

Level 0: Isolated. Systems have no connection to each other.

Level 1: Connected. Systems are linked electronically, generally peer-to-peer.

Level 2: Functional. Systems operate on a distributed local network.

Level 3: Domain. Systems are capable of being connected via a WAN.

Level 4: Enterprise. Systems are capable of using a distributed global information space across multiple domains.

A cloud computing infrastructure set in an afloat environment would have to fulfill the maturity level of three, at a minimum. At this level, a strike group would be enabled with data sharing over a WAN, capable of sharing data over distances greater than LOS and allowing multiple users to access shared data. Optimally, afloat cloud

architectures would strive to become a robust Level 4 maturity level, being able to connect to DoD and Naval network enterprises, in particular the Global Information Grid. If afloat cloud nodes are to ride on the (future) CANES infrastructure, then Level 4 should be a goal to achieve.

The LISI model was written in 1997 and at that time cloud computing was still just an idea. A revised LISI model would likely need to be completed in order to address the revolutionary service delivering technology that is cloud computing. An IEEE paper published in 2011 titled “Cloud to Cloud Interoperability” addresses the applicability of the LISI model to cloud computing and makes several suggestions on how the LISI model can be updated to assess the maturity of cloud-to-cloud interoperability. [50]

Nature of Operational Information Interaction	Corresponding Interoperability Level		Implications			
			P	A	I	D
Cross-Domain Interactive Manipulation	Enterprise	4	Enterprise Level	Interactive	Multiple Topologies	Enterprise Model
Shared Applications & Databases	Domain	3	Domain Level	Groupware	World Wide Networks	Domain Model
Complex Media Exchange	Functional	2	Program Level	Desktop Automation	Local Networks	Program Model
Simple Electronic Exchange	Connected	1	Local/Site Level	Standard System Drivers	Simple Connection	Local
Manual Gateway	Isolated	0	Access Control	N/A	Independent	Private

Figure 28. The Levels of Information System Interoperability Maturity Model [49]

In [50], Barreto, et al., agree that portability and mobility are “prime indicators” of the degree of interoperability between clouds, and believe open standards for VMs and cloud-to-cloud APIs and advancement in virtualization technologies will be required for successful cloud to cloud interoperability. These statements are applicable to the afloat

cloud environment, particularly considering that each ship's onboard cloud infrastructure is built on a virtual desktop infrastructure.

The difficulties in making clouds interoperable are not trivial. Each ship will have its own unique cloud; no two ships in the Navy are exactly the same. An example is the Aegis Baseline, the software version of the Aegis weapons system. Often, ships will have varying versions of the Baseline. This can cause complications when ships are operating in the same AOR and the Aegis systems try to communicate. Like the Aegis Baseline, ships with their own unique cloud will likely encounter problems when attempting to share data with other clouds that are not similar.

As cloud computing becomes more prominent in business, industry, the military, etc., organizations will develop their own cloud enterprises that have their own unique characteristics and use specific standards. Granted, cloud computing is typically designed around web services, intended for the primary purpose of leveraging the Internet, and therefore should all use the same standards. The issue lies within each organization and the standards, equipment, technologies, etc., that they will use within their organization. Variations within each of these categories will cause interoperability issues like those exhibited in the example of the Aegis Baseline.

Listed below are suggested attributes that should be characteristic in any cloud to ensuring cloud to cloud interoperability [50]:

- Portability of data, i.e., data in a format readable by all systems within a cloud architecture. Loose coupling or the encapsulation of data may be suitable methods of portability.
- Mobility of data, or “the ability to move a live computer workload from one host to another without losing client connections or in-flight state.” [50]
- Applying SOA principles to help cloud to cloud integration. CANES will be of great assistance in this process considering it will be based on SOA.

- A well-defined cloud security policy across all clouds within a strike group domain.
- Establishment of formal trust relationships between clouds.
- A uniform tool set to provision services automatically and to manage VM instances.

In [50], additions to the LISI model's PAID attributes are suggested to further develop cloud to cloud interoperability:

- Procedures: Uniform security and privacy policies. These must be applied consistently across all cloud boundaries.
- Application: The ability of software applications to “work and migrate seamlessly across cloud boundaries.” [50] Additionally, a set QoS would have to be maintained.
- Infrastructure: This attribute helps maintain cloud mobility and the degree of cloud-services provisioning, management, monitoring, reporting and auditing.
- Data: Shift from application-centric to data-centric view of information processing. Data in the cloud must be treated as artifacts, artifacts being the embodiments of data. Users can manipulate this data, and the benefit of data as artifacts is that it provides a uniform manner in which systems can access the data.

Characteristics like those listed above will be necessary to support cloud to cloud interoperability in an afloat cloud computing infrastructure. The point of utilizing cloud architectures in an afloat environment is to provide cloud services capable of sharing data between ships. To do so, data must have the ability to be accessed and manipulated by the systems in the entire cloud enterprise (manipulated meaning data is able to meet the basic CRUD computer programming functions). Furthermore, ships in a strike group should be equipped with the same information systems, use the same standards and policies throughout the cloud enterprise, and offer the same cloud services. Doing so will

help to achieve semantic and syntactic cloud interoperability specifically, and will promote cloud to cloud interoperability in general.

To some extent ships will have many of the standards and resources to store cloud data when utilizing a VDI like the one built in Chapter III. Those VDIs contain a virtualized pool of storage in its datacenter and offer resource agility, scalability, and elasticity. Hard disk drives can easily be added or removed and hard disk space will scale when necessary to meet increasing volumes of information.

Linking clouds, cloud to cloud interoperability, and cloud data storage rely on physical computer and network systems. These systems would not operate without the human role of governance. The next section discusses the command and control architecture that would be required in an afloat cloud computing environment.

D. AFLOAT CLOUD COMMAND AND CONTROL STRUCTURE

Governance in cloud computing includes people and processes, not just technology. To ensure cloud technology operates correctly, the support of the people who will build, control, and monitor the cloud computing infrastructure is important. [8] Likewise, governance includes management-level visibility into the cloud infrastructure in afloat environments, important for the oversight of the people implementing the processes. Governance is necessary for designing and implementing policies, and to provide “command, control, discovery, and monitoring as well as design and development support for [cloud] services.” [8]

Governance of afloat cloud environments will be determined by a command and control (C2) structure. Each strike group has an organizational C2 structure delineating the chain of command, encompassing all ships belonging to the strike group. Before a strike group deploys, specific Concept of Operations and Operational Plans are promulgated by the applicable authority for each warfare area, designating certain responsibilities to assets in the strike group. A portion of these responsibilities establishes warfare commanders under the Composite Warfare Commander (CWC) concept. One of the warfare commanders is the Force Over-the-Horizon track coordinator, or FOTC.

As stated in NTTP 6-02, the FOTC is the key process in which “all information is evaluated at a central node before being disseminated.” [51] In an afloat cloud environment the FOTC participates in C4I planning and would be the appropriate warfare coordinator responsible for the operation of the cloud infrastructure. Traditionally, the FOTC resides on the aircraft carrier in a strike group, which is suitable given the aircraft carrier is also the hub in the afloat cloud infrastructure as presented in previous sections.

Overseeing the intelligence and network operations of the strike group is the N2/N6 staff officer. The N2/N6, like the FOTC, participates in C4I planning and is embarked on the aircraft carrier. Together, the N2/N6 and the FOTC would be the senior officers responsible for the afloat cloud network. Policy would generally be promulgated by the N2/N6, while the FOTC would execute the policy while the strike group is operational. At the operational level onboard the aircraft carrier (the hub), the communications officer (COMMO) and the Information Systems Officer (ISO) would together implement cloud policy on behalf of the FOTC. Figure 29 depicts a notional standard C2 structure at the force level.

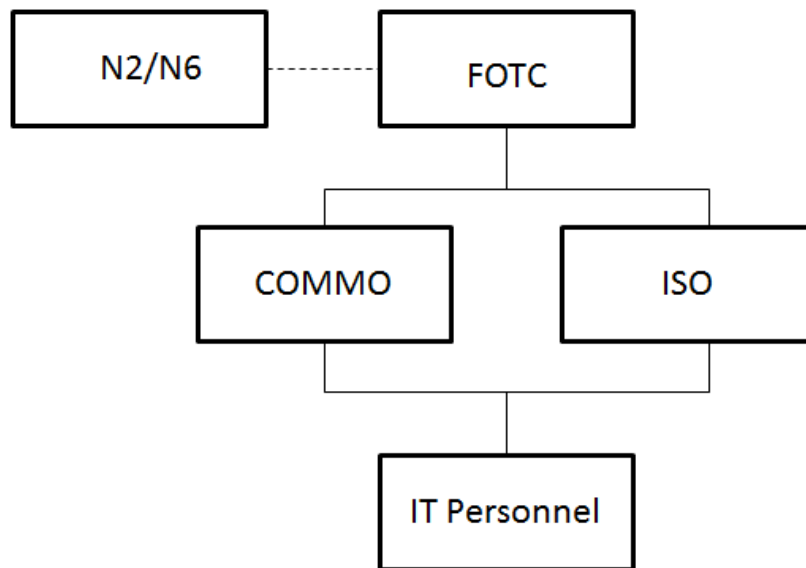


Figure 29. Notional force level (hub) cloud C2 structure

The C2 structure at the unit level (Figure 30) would be very similar to the force level structure. Unit level ships such as a DDG or CG do not have a FOTC or an N2/N6

staff officer embarked as there is only one of each in the entire strike group. Unit level ships will have a COMMO responsible for overall operations in the radio room, including the ship's internal networks and all communication channels and systems (HF, UHF, SHF, etc.). Often, an ISO is billeted on a unit level ship. In this case the ISO assists the COMMO, either sharing responsibilities tasked to the COMMO or fulfilling roles normally assigned to the COMMO. IT personnel work for both the COMMO and the ISO, with the COMMO being the senior officer.

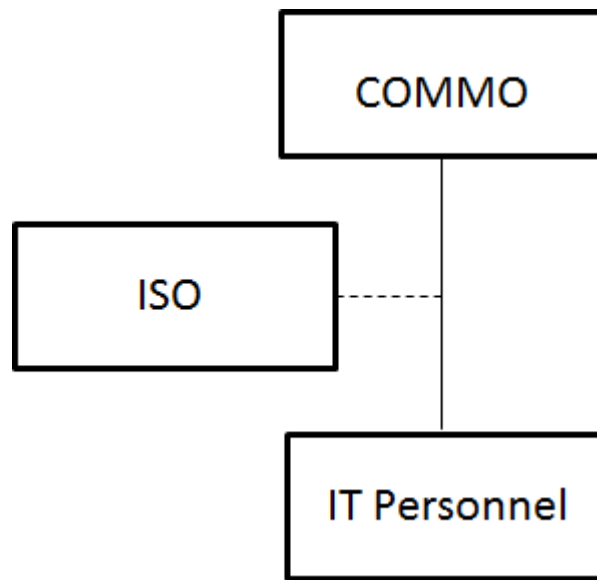


Figure 30. Notional unit level (edge cloud) C2 structure

In sum, governance is necessary to control and monitor the cloud services that make up an afloat cloud architecture. At the force level, the FOTC and the N2/N6 share responsibilities of the afloat cloud architecture, while at the unit level the COMMO is responsible. The implementation of SOA into the cloud architecture will reflect the policies and control of the FOTC and COMMO. In enterprise architecture, governance means “control, or the ability to mandate the use of standards and approaches,” while in the world of SOA governance means “designing, building, testing, and implementing policies for services and monitoring their use.” Governance in afloat cloud computing should encompass these definitions, and the FOTC and COMMO should adhere to them.

V. CONCLUSION

A. SUMMARY

Overall, this thesis expressed the need for cloud computing in afloat naval environments because of the benefits cloud computing offers and because the Federal Government has mandated a “Cloud First” policy regarding federal IT enterprises. The discussion on the need for cloud computing in afloat environments specifically stems from the request by PEO-C4I for research on cloud computing afloat.

Following the need for cloud computing, a literature review was carried out to understand concepts, definitions, and terminology relevant to this thesis. Current Navy network programs and initiatives were discussed and the implications and possibilities they present regarding their potential integration with cloud computing infrastructures were examined. The design and implementation of a VDI in the NPS Cloud Lab provided a generic approach on how to build a VDI for a naval asset in an afloat environment. The VDI was then scaled out to include a cloud computing CONOPS for a multi-ship environment.

This thesis proposed several ideas on what a cloud computing infrastructure may look like in an afloat environment. The first part focused on server consolidation and virtual desktop infrastructures, which hosts desktop operating systems on virtual machines. By consolidating servers into blade servers the physical network footprint is reduced as a result. Blade servers are robust and scalable enterprise platforms and are suitable for cloud computing enterprises.

Virtual desktop infrastructures utilize virtualization to create virtual machines. Virtual machines are then accessible to users via devices such as thin clients or zero clients, which replace conventional desktop computers. Virtual machines are hosted on servers, such as blade servers that can host several thousand virtual machines. Virtualization benefits include datacenter automation, reduction in capital costs via the increase in energy efficiency, maximization of hardware resources, and ease of enterprise desktop management.

The VDI built in the Cloud Lab at NPS serves as a notional model for integration onto naval ships, replacing the typical client-server model. With a VDI, aggregated resources form a uniform set of elements that can be managed like a shared utility and dynamically provisioned out to virtual machines demanding resources. Resources are elastic, provided to users when needed.

A VDI designed to meet user needs and properly integrated into a ship's network infrastructure offers a cloud computing enterprise network foundation that is agile, elastic, and scalable. Desktops are streamed to users over the shipboard network, delivering a unique cloud service – desktop-as-a-service, and management is centralized in one server that presides over the entire VDI. In short, a VDI is a suitable foundation for a shipboard level cloud computing architecture.

The second part of this thesis moved beyond the shipboard level and focused on afloat cloud computing infrastructures in a multi-ship environment. Several ideas were proposed on how to link clouds, including non-satellite and satellite technologies. One of the prime benefits of using cloud computing at sea is that bandwidth, particularly satellite bandwidth, is not much of a problem if data is in a database within the strike group. In the absence of satellites, however, a viable method of linking clouds is critical. This is even more so the case when ships are beyond LOS. As suggested in this thesis, WiMAX is a promising solution to linking clouds. WiMAX has a large bandwidth capacity and is very fast, especially when the underpinning infrastructure is 4G.

Linking clouds is a part of determining what an afloat cloud infrastructure looks like, but where data is actually stored will be more of an infrastructure determinant. Each ship will have its own datacenter, outfitted with large pools of virtualized storage resources in the VDI. Data sharing is a crucial part of afloat naval operations and making it accessible is a critical piece of the data storage process. A single point of ingest and dissemination of data from shore facilities and from satellites at the hub reduces satellite use from other ships in the strike group. Another important piece of the data process is the systems that process data into usable intelligence. The hub is therefore a suitable solution for being the primary place of storage given its large datacenters and that it houses the systems necessary to process data.

Clouds within a strike group must be interoperable for data sharing to take place. This thesis suggested the Levels of Information Systems Interoperability (LISI) Maturity Model as a relevant model for systems interoperability. Several attributes and characteristics from the LISI Maturity Model are listed and defined. Portability, mobility, SOA principles, trust relationships, and uniform tool sets are just some of the many characteristics that individual and perhaps dissimilar clouds should adopt when they are to integrate and form one large, interoperable cloud computing environment.

The last section regarding afloat cloud architectures discussed a possible afloat cloud command and control structure. Cloud computing governance does not just cover system governance, but also encompasses the human factor. When an afloat cloud infrastructure is built there will be a need for human operators. Promulgation of C4I planning for a strike group is conducted at the force level. Officers included at this level are the N2/N6 (Intelligence/C4I) and the FOTC. Together, these authority positions are suitable for creating and executing policies and plans relative to an afloat cloud computing infrastructure. At the unit level, the COMMO and the ISO (when applicable) would be the practical operators of the cloud infrastructure given that they are responsible for all other computer networks and communication systems on a ship.

Cloud computing presents a paradigm shift in the way IT services are provided over networks. When cloud computing is combined with virtualization the benefits to an IT enterprise are many, including:

- Consolidation of disparate resources
- Consolidation of servers leading to a reduction in the network footprint
- Greater system to system interoperability via loose coupling
- Resource elasticity and system scalability

Benefits of particular importance to the Navy include:

- Cost reduction in hardware, software, manning, and more (\$B in savings)
- Energy efficiency
- Potential increase in bandwidth efficiency and a reduction in satellite saturation
- Brings data to the tactical edge, practically eliminating reach back

The Department of Defense currently offers cloud services through DISA to shore facilities. Experiments with cloud computing in afloat environments have so far been limited, but there have been and continues to be experiments and the move toward cloud computing in afloat environments is gaining traction. PEO-C4I and SPAWAR are actively researching cloud computing for the purpose of using cloud architectures at sea. This thesis maps to several NPS thesis topics PEO-C4I has provided by the NPS C4I chair. In 2009 a company called Dataline LLC demonstrated that a “standard shipboard communications infrastructure could be used to manage a commercial cloud computing infrastructure as a service (IaaS) platform.” [52]

The near future role out of the CANES initiative, which is based on SOA, places cloud computing only steps away from fleet implementation. IaaS will be provided by CANES, while PaaS and desktops-as-a-service can be provided by a VDI such as the one built in this thesis. With the CANES initiative and the tangible benefits listed above, afloat cloud computing network enterprises should sound very promising to the Navy.

B. FUTURE RESEARCH

Utilization of cloud computing is growing, expanding into many types and areas of companies, organizations, governments and militaries. This thesis focused on cloud computing in afloat naval environments. Limited in scope, several aspects of cloud computing were not research or discussed. One aspect vital to military operations is security, which was not discussed in this thesis. Security in cloud computing is a topic deserving of in-depth research. Implementing cloud computing introduces new security vulnerabilities that need to be understood and addressed. When researching for this thesis much of the reference material included polls asking organizations to list their concerns about moving their IT infrastructure to a cloud based enterprise. In almost every case security topped the list.

This thesis proposed several concepts and implemented a VDI, but did not include physical testing of all that was proposed. Field experiments on linking clouds using any of the technologies listed in this thesis (such as LTA ships, WiMAX, BGAN) provides

plenty of opportunities for future research. The same is true for testing cloud-to-cloud interoperability.

The VDI built in the NPS Cloud Lab did not use several of the virtualization products or cloud computing products offered by VMware. Products such as vCloud Director, vCloud Connector, or the newest version of vSphere ESXi (version 5.0) were not used. Use of these products likely offers better cloud computing and virtualization performance and more configuration options with greater benefits than were achieved with the products and versions used. Scaling the VDI to include tactical applications is another future research need, especially if cloud computing becomes the standard in afloat network environments.

Cyber defense and cyber warfare applications will likely need research in cloud computing architectures and cloud computing concepts of operations. Ultimately, understanding future cloud infrastructures developed in industry is an area of research to determine how to best leverage cloud computing in the DoD.

THIS PAGE INTENTIONALLY LEFT BLANK

LIST OF REFERENCES

- [1] V. Kundra, "25 Point Implementation Plan to Reform Federal Information Technology Management," Office of the U.S. Chief Information Officer, Washington, D.C., 2010.
- [2] VMware, Inc., "VMware Awards," [Online]. Available: <http://www.vmware.com/company/news/awards.html>. [Accessed 28 May 2012].
- [3] M. Lindsey, "On Network Diagrams," 10 December 2007. [Online]. Available: <http://www.200ok.info/2007/12/on-network-diagrams.html>. [Accessed 14 March 2012].
- [4] "Cloud Computing & Enterprise Services," Defense Information Systems Agency, [Online]. Available: <http://www.disa.mil/Services/Computing/Cloud-Computing-and-Enterprise-Services>. [Accessed 8 January 2012].
- [5] P. Mell and T. Grance, "Publications," National Institute of Standards and Technology, September 2011. [Online]. Available: <http://csrc.nist.gov/publications/PubsSPs.html#800-145>. [Accessed 10 January 2012].
- [6] J. W. Rittinghouse and J. F. Ransome, Cloud Computing: Implementation, Management, and Security, Boca Raton, FL: CRC Press, 2010, p. 49.
- [7] "Cloud Computing Layers," 1 February 2012. [Online]. Available: http://en.wikipedia.org/wiki/File:Cloud_computing_layers.png. [Accessed 4 March 2012].
- [8] D. S. Linthicum, Cloud Computing and SOA Convergence in Your Enterprise: a Step-by-step Guide, Upper Saddle River, NJ: Pearson Education, Inc., 2010.
- [9] J. Metzler, "Virtualization: Benefits, Challenges, and Solutions," Riverbed Technology, San Francisco, 2011.
- [10] VMware, Inc., "vSphere 5 Documentation Center," [Online]. Available: http://pubs.vmware.com/vsphere-50/index.jsp?topic=/com.vmware.vsphere.introduction.doc_50/GUID-F7A7E6C0-FA25-4806-8921-0438F1B2AEAE.html. [Accessed 3 December 2011].
- [11] G. Heiser, "Much Ado About Type-2," 14 October 2010. [Online]. Available: <http://www.ok-labs.com/blog/entry/much-ado-about-a-type-2/>. [Accessed 20 January 2012].
- [12] R. McCarty, "Paravirtualization explained," [Online]. Available: <http://searchservirtualization.techtarget.com/tip/Paravirtualization-explained>. [Accessed 13 April 2012].

- [13] T. Dean, *Network + Guide to Networks*, Fifth ed., Boston, MA: Course Technology, Cengage Learning, 2010.
- [14] *Information Warfare and Information Professional Basic Officer Course (CORE Material)*, 10-08 ed., Corry Station, FL: Center for Information Dominance.
- [15] C. Suggs, PEO C4I Cloud WG, San Diego, CA, 2011, p. 7.
- [16] Microsoft, "Remote Desktop Protocol," 3 February 2012. [Online]. Available: [http://msdn.microsoft.com/en-us/library/windows/desktop/aa383015\(v=vs.85\).aspx](http://msdn.microsoft.com/en-us/library/windows/desktop/aa383015(v=vs.85).aspx). [Accessed 12 March 2012].
- [17] Microsoft, "Remote Desktop Protocol (RDP)," [Online]. Available: <http://www.microsoft.com/about/legal/en/us/intellectualproperty/iplicensing/programs/remotedesktopprotocol.aspx>. [Accessed 11 March 2012].
- [18] Teradici, "PCoIP Technology Explained," [Online]. Available: <http://www.teradici.com/pcoip/pcoip-technology.php>. [Accessed 3 March 2012].
- [19] "PCoIP Technology User Guide," Chung Ho City, 2008.
- [20] F. Ohlhorst, "Replacing Desktop PCs with Zero-Client Solutions," 25 March 2009. [Online]. Available: <http://www.channelinsider.com/c/a/Networking/Replacing-Desktop-PCs-with-ZeroClient-Solutions-179077/>. [Accessed 9 March 2012].
- [21] SPAWAR, "OCONUS Navy Enterprise Network (ONE-Net) is...", [Online]. Available: <http://www.public.navy.mil/spawar/PEOEIS/NEN/ONE-Net/Pages/default.aspx>. [Accessed 9 March 2012].
- [22] SPAWAR, "Consolidated Afloat Networks and Enterprise Services (CANES)," [Online]. Available: <http://www.public.navy.mil/spawar/productsServices/Pages/ConsolidatedAfloatNetworksandEnterpriseServicesCANES.aspx>. [Accessed 9 March 2012].
- [23] PEO C4I, "PEO C4I Masterplan," PEO C4I, San Diego, 2011.
- [24] J. Sprague, Naval C4I/IT Seminar Consolidated Afloat Networks and Enterprise Services (CANES), San Diego, CA, 2009, p. 6.
- [25] Department of the Navy Chief Information Officer, *Naval Networking Environment (NNE) 2016 Strategic Definition, Scope and Strategy Paper*, 1.1 ed., 2008.
- [26] Consolidated Afloat Network Enterprise Services (CANES), EXHIBIT R-2a, RDT&E Project Justification, 2009, p. 3.
- [27] PEO-C4I, "Technical Framework for Cloud Computing at the Tactical Edge," PEO-C4I, San Diego, 2011.

- [28] VMware, Inc., "Infrastructure Virtualization and Management," [Online]. Available: <http://www.vmware.com/virtualization/>. [Accessed 15 March 2012].
- [29] VMware, Inc., "VMware," VMware, Inc., [Online]. Available: <http://www.vmware.com/>.
- [30] Dell Inc., "PowerEdge M1000e," 2010.
- [31] Dell, Inc., "Dell PowerEdge M610," [Online]. Available: <http://www.dell.com/us/enterprise/p/powerededge-m610/pd>. [Accessed 8 November 2011].
- [32] VMware, Inc., "VMware vCenter Server," [Online]. Available: <http://www.vmware.com/products/vcenter-server/overview.html>. [Accessed 3 December 2011].
- [33] VMware, Inc., "VMware vSphere 4 - ESX and vCenter Server," [Online]. Available: http://pubs.vmware.com/vsphere-4-esx-vcenter/index.jsp?topic=/com.vmware.vsphere.intro.doc_41/c_hosts_clusters_and_resource_pools.html. [Accessed 8 December 2011].
- [34] Wyse Technology, Inc., "Wyse P20," [Online]. Available: <http://www.wyse.com/products/cloud-clients/zero-clients/P20>. [Accessed 7 January 2012].
- [35] U.S. Navy, "NIAPS Frequently Asked Questions," [Online]. Available: http://www.public.navy.mil/spawar/PEOEIS/NAVY311/Pages/NIAPS_Info_Center_FAQs-Basics.html#Q1. [Accessed 11 February 2012].
- [36] J. Herbert R. Race, "The Role of Airships in a Transformational Navy," 2004.
- [37] "MQ-4C BAMS UAS," Northrop Grumman Corporation, [Online]. Available: <http://www.as.northropgrumman.com/products/bams/index.html>. [Accessed 30 April 2012].
- [38] Inmarsat Inc., "BGAN Global Voice and broadband data," Inmarsat, [Online]. Available: http://www.inmarsat.com/Services/Land/Services/High_speed_data/BGAN.aspx?language=EN&textonly=False. [Accessed 30 April 2012].
- [39] Cisco Systems, Inc., "Cisco Satellite Solutions (IRIS)," [Online]. Available: <http://www.cisco.com/web/strategy/government/space-routing.html>. [Accessed 1 May 2012].

- [40] Cisco Systems, Inc., "IRIS JCTD - US DoD Evaluation: Success," Cisco Systems, Inc., [Online]. Available: http://www.cisco.com/web/strategy/docs/gov/iris_jctd_website.pdf. [Accessed 1 May 2012].
- [41] Intelsat, "Intelsat," [Online]. Available: <http://www.intelsat.com/>. [Accessed 1 May 2012].
- [42] HowStuffWorks, Inc, "How WiMAX Works," [Online]. Available: <http://computer.howstuffworks.com/wimax2.htm>. [Accessed 1 May 2012].
- [43] HowStuffWorks, Inc, How WiMAX Works, [Online]. Available: <http://computer.howstuffworks.com/wimax1.htm>. [Accessed 1 May 2012].
- [44] GoingWimax.com, "What is a WiMAX Antenna?," [Online]. Available: <http://www.goingwimax.com/what-is-a-wimax-antenna-4165/>. [Accessed 1 May 2012].
- [45] B. Brewin, "Naval Air Systems Command plans 4G cell service aboard ships," 2012.
- [46] T. Day, Naval Postgraduate School, Monterey, 2012.
- [47] C4ISR Architecture Working Group Interoperability Panel, "Levels of Information Systems Interoperability (LISI)," Department of Defense, Washington, DC, 1998.
- [48] S. Chiu, "Can Level of Information Systems Interoperability (LISI) Improve DOD C4I Systems' Interoperability?," Naval Postgraduate School, Monterey, 2001.
- [49] C4ISR Architecture Working Group, "C4ISR Architecture Framework Version 2.0," Department of Defense, 1997.
- [50] S. Dowell, A. Barreto III, J. Michael and M.-T. Shing, "Cloud to Cloud Interoperability," IEEE, Albuquerque, 2011.
- [51] Department of the Navy, "C4I Infrastructure NTTP 6-02," Office of the Chief of Naval Operations, 2005.
- [52] K. Jackson, "The Mil & Aero Blog," 19 November 2009. [Online]. Available: <http://www.pennwellblogs.com/mae/2009/11/guest-blog-navy-demonstrates-value-of.html>. [Accessed 22 July 2011].

INITIAL DISTRIBUTION LIST

1. Defense Technical Information Center
Ft. Belvoir, Virginia
2. Dudley Knox Library
Naval Postgraduate School
Monterey, California
3. RDML Jerry Burroughs
SPAWAR PEO C4I
San Diego, California
4. Pat Sullivan
SPAWAR PEO C4I Executive Director
San Diego, California
5. CAPT John Pope
SPAWAR PEO C4I Principal Deputy
San Diego, California
6. Robert Wolborsky
SPAWAR Chief Technology Officer
San Diego, California
7. CAPT Joel Beel
Commanding Officer, SPAWAR Systems Center Pacific
San Diego, California
8. CAPT Bryan Lopez
Executive Officer, SPAWAR Systems Center Pacific
San Diego, California
9. Robert Parker (CAPT, USN, ret.)
SPAWAR PEO C4I, APEO for S&T
San Diego, California
10. Charles Suggs
SPAWAR PEO C4I, DPEO Technical Director & Technical Integration
San Diego, California
11. Jerry Almazan
SPAWAR PEO C4I, PMW 120
San Diego, California

12. Delores Washburn
SPAWAR PEO C4I, PMW 160
San Diego, California
13. Alexander Vasel
SPAWAR PEO C4I, PMW 160
San Diego, California
14. Kurt Fisko
SPAWAR PEO C4I, PMW 170
San Diego, California
15. Ruth Youngs Lew
SPAWAR PEO C4I, PMW 790
San Diego, California
16. Patrick Garcia
SPAWAR PEO C4I, PMW 150
San Diego, California
17. Nicholas Gizzi
SPAWAR PEO C4I, PMW 150
San Diego, California
18. Dan Boger
Naval Postgraduate School
Monterey, California