

ABSTRACT

Title of Thesis: QUANTIFYING THE COST OF DATA
CENTER COOLING APPROACHES
INCLUDING SERVICE-LEVEL
AGREEMENTS

Edward Lee Honadel
Master of Science
2025

Thesis directed by: Professor Peter Sandborn
Department of Mechanical Engineering

This thesis developed a comprehensive cost model for evaluating and comparing the financial viability of different data center cooling systems. As data center demand and energy consumption rise, driven by AI and cloud computing, optimizing cooling systems is critical for reducing operational expenditures. Existing cost models overlook the financial penalties associated with downtime defined in Service-Level Agreements (SLAs). This research addresses this gap by creating a model that integrates capital and operational costs with a stochastic discrete-event simulator for maintenance, failures, repairs, and availability. By quantifying downtime penalties in addition to energy and capital expenditures, the model provides a more accurate total cost of ownership. The model calculates the Return on Investment (ROI) and Internal Rate of Return (IRR) to assess the economic justification for deciding when investing in advanced or redundant cooling technologies is economically justifiable. A case study on the CEETHERM lab data center at Georgia Tech shows that the specific method of calculating

penalties based on either cumulative downtime, continuous downtime, or number of interruptions can significantly shift the projected Return on Investment and breakeven time. The results demonstrate that penalties based on counting the number of interruptions during a contract period resulted in the largest system cost impact, which prevented the system from ever reaching financial breakeven, demonstrating that Service-Level Agreements based on the frequency of service interruptions pose a significantly higher financial risk to data center operators than those based on outage duration.

QUANTIFYING THE COST OF DATA CENTER COOLING APPROACHES
INCLUDING SERVICE-LEVEL AGREEMENTS

by

Edward Lee Honadel

Thesis submitted to the Faculty of the Graduate School of the
University of Maryland, College Park, in partial fulfillment
of the requirements for the degree of
Master of Science
2025

Advisory Committee:

Dr. Peter Sandborn, Chair

Dr. F. Patrick McCluskey

Dr. Diganta Das

© Copyright by
Edward Lee Honadel
2025

Acknowledgements

The author would like to express his thanks to Peter DeBock and Thomas Bress for their technical support of this research. This material is based upon work supported by ARPA-E (U.S. Department of Energy) under Award Number DE-AR0001755.

Disclaimer: This thesis contains an account of work sponsored by an agency of the United States Government. Neither the United States Government nor any agency thereof, nor any of their employees, makes any warranty, express or implied, or assumes any legal liability or responsibility for the accuracy, completeness, or usefulness of any information, apparatus, product, or process disclosed, or represents that its use would not infringe privately owned rights. Reference herein to any specific commercial product, process, or service by trade name, trademark, manufacturer, or otherwise does not necessarily constitute or imply its endorsement, recommendation, or favoring by the United States Government or any agency thereof. The views and opinions of authors expressed herein do not necessarily state or reflect those of the United States Government or any agency thereof.

Table of Contents

Acknowledgements.....	ii
List of Abbreviations	iv
Chapter 1: Introduction	1
1.1 Background on Data Centers	1
1.2 Data Center Business Types	3
1.3 Data Center Cost Modeling	4
1.4 Commercially Available Cost Models.....	8
1.5 Thesis Objective Statement.....	9
1.6 List of Tasks.....	9
Chapter 2: Cost Modeling.....	12
2.1 ROI and IRR Calculations	13
2.3 Maintenance Discrete-Event Simulator	16
2.3.1 Validating Discrete-Event Simulator.....	20
Chapter 3: Modeling Service-Level Agreements	22
3.1 Taxonomy of SLA Penalty Structures	22
3.1.1 Duration-Based Penalties.....	23
3.1.2 Frequency-Based Penalties	24
3.1.3 Mathematical Framework for Penalty Calculation.....	25
3.2.1 Duration-Based Penalty Models	25
3.2.2 Graduated Duration Penalty.....	26
3.2.3 Escalating Interrupt Penalty.....	27
Chapter 4: Case Study.....	28
4.1 Data Center Description – Reference Architecture	28
4.1.1 Parameters in Cost Model.....	29
4.1.2 Data Center Scaling Penalties.....	29
4.2 Case Study 1: The Impact of Varying SLA Structures.....	30
4.2.1 Case Study 1 Data Assumptions	31
4.2.2 Case Study 1 Results.....	32
4.3 Case Study 2: Direct Liquid Cooling vs. Traditional Air Cooling	35
4.3.1 Case Study 2 Data.....	35
4.3.2 Case Study 2 Results.....	36
4.4 Case Study 3: Baseline vs. Enhanced Redundancy	37
4.4.1 Case Study 3 Data.....	37
4.4.2 Case Study 3 Results.....	38
Chapter 5: Summary, Contribution, and Future Work	41
Appendix – Parametric Equipment Cost Models.....	44
References.....	46

List of Abbreviations

AI - Artificial Intelligence

BOM – Bill of Materials

CAPEX – Capital Cost

CPU – Central Processing Unit

CRAC – Computer Room Air Conditioning Unit

CRAH – Computer Room Air Handler

DES – Discrete-Event Simulator

DOE – Department of Energy

GPU – Graphics Processing Unit

HVAC – Heating, Ventilation and Air Conditioning

IoT – Internet of Things

IRR – Internal Rate of Return

MTBF – Mean Time Between Failure

MTTR – Mean Time To Repair

NPV – Net Present Value

OPEX – Operational Expenses

ROI – Return on Investment

SLA – Service-Level Agreement

TCO – Total Cost of Ownership

UPS – Uninterruptible Power Supply

Chapter 1: Introduction

1.1 Background on Data Centers

Demand for data center capacity is constantly increasing worldwide, primarily driven by the training and delivery of artificial intelligence (AI) and cloud computing. The need for data processing and storage is growing at an exponential rate due to the Internet of Things (IoT), data analytics, and the increasing complexity of computational workloads [1]. The construction and operation of data centers has become a multi-billion-dollar industry and a significant source of energy consumption. It is estimated that data centers currently account for 1-2% of global electricity usage, a figure that is projected to rise significantly in the coming decade [2].

Data centers are composed of IT equipment, storage systems, network infrastructure, power infrastructure, cooling systems and security systems. The cooling systems provide continuous heat removal from the IT equipment to prevent its temperature from exceeding operational limits, which could result in degraded performance or component failure, both of which potentially impact the data center's availability. The IT equipment, consisting of servers and storage systems, performs the core function of the data center by processing, storing, managing, and securely transmitting digital information. The cooling system is therefore a critical component that directly affects the data center's reliability and financial success [20]. The energy consumed by a cooling system can account for 30-50% of a data center's total electricity usage [1]. This makes optimizing the cooling system a target for reducing Operational Expenditures (OpEx) and environmental impact. As CPUs and GPUs become more advanced, the rack power densities in data centers continue to increase [3]. This trend is pushing traditional air-cooling methods to their limits, necessitating the exploration of more

advanced and efficient cooling systems.

A data center is a facility designed to house and support IT infrastructure, including servers with CPUs and GPUs, data storage drives, and networking infrastructure [5]. The facility provides the infrastructure for electrical distribution to the servers, which includes an uninterruptible power supply (UPS), environmental controls, and security. The performance of a data center is measured by its availability, or uptime [4]. As defined by the Uptime Institute, data centers are classified into four different tiers depending on the availability they can achieve (Table 1).

	Uptime per year	Downtime per year	Redundancy	Cost	Traditional Customer
Tier I Basic Capacity	99.671%	<28.8 hours	- None	\$	- Small businesses - Start-ups - Businesses with simple requirements
Tier II Redundant Capacity Components	99.741%	<22 hours	- Partial power and cooling redundancy	\$\$	- SMB
Tier III Concurrently Maintainable"	99.982%	<1.6 hours	- N+1 fault tolerance	\$\$\$	- Growing businesses - Large businesses - Government entities
Tier IV Fault Tolerant"	99.995%	<26.3 minutes	- 2N or 2N+1 - fully fault-tolerant	\$\$\$ \$	- Large enterprises - Businesses with an international reach

Table 1: Data center tiers and their downtime. The types of customers and redundancy used in the data center are also listed. [4]

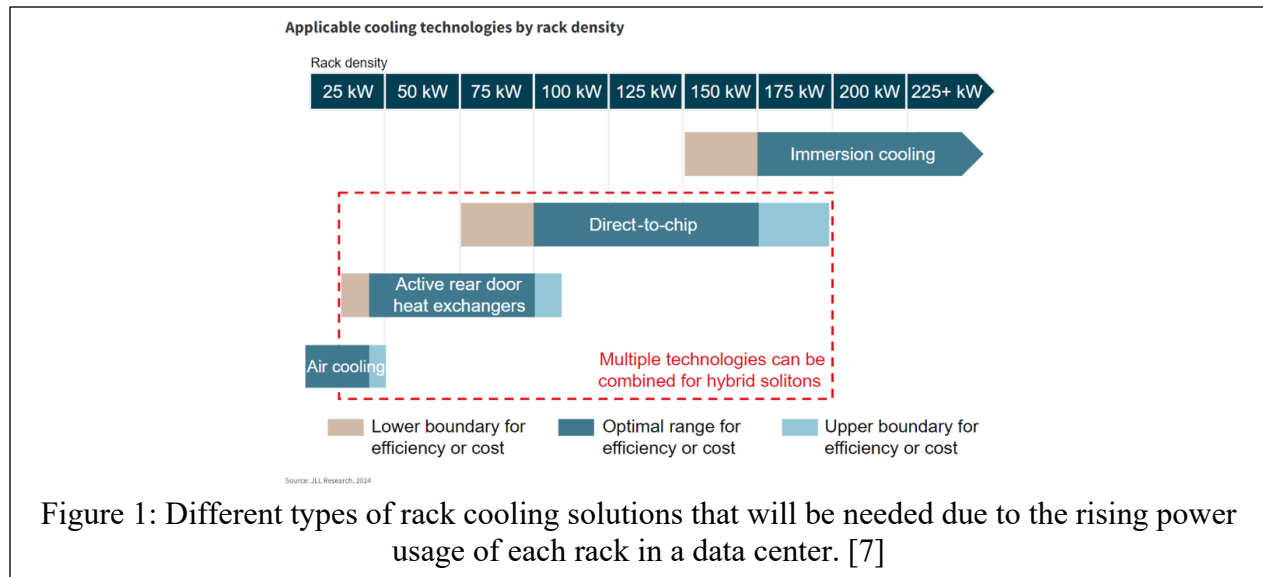
The thermal management system is important to achieving a high level of reliability. The heat generated by IT equipment must be efficiently captured and removed. Data center cooling systems can be broadly categorized into two interconnected loops: the primary and secondary cooling loops. The Primary Cooling Loop is responsible for generating the cooling medium and rejecting the collected heat to the outside environment. Key components of the primary loop include large-scale chillers, cooling towers, and pumps that circulate the water [5]. The design of the primary loop is a major capital investment and a significant driver of the facility's overall energy efficiency. The Secondary Cooling Loop operates inside the data hall and is responsible for transferring heat from the IT equipment to the primary loop. Traditional secondary loop

systems utilize Computer Room Air Conditioners (CRACs) or Computer Room Air Handlers (CRAHs). These units pull warm air from the hot aisle, which cools it by passing it over coils containing chilled water from the primary loop, and supply the resulting cold air to the cold aisle to be drawn into the servers [6]. As the advancements in artificial intelligence drive power usage in data centers, newer cooling solutions will be needed. As shown in Figure 1, as the rack density increases, newer solutions will need to be used to effectively cool newer power intensive CPUs and GPUs.

1.2 Data Center Business Types

Enterprise Data Centers: These are data centers that are owned and operated by a single organization for its internal IT needs. These businesses prioritize reliability and security over cost-of-service [8].

Colocation Data Centers: These facilities provide building, power, cooling, and security, for leased equipped space to multiple customers. Colocation providers compete on price, connectivity, and, most importantly, the assurance of uptime, which is guaranteed through a



Service-Level Agreement (SLA) [9].

Hyperscale (Cloud) Data Centers: These are large facilities owned and operated by a single company to deliver cloud services to many customers. For hyperscalers, the key is massive scale, operational efficiency, and the lowest possible total cost of ownership (TCO) to remain competitive [10]. Hyperscalers own all the data center equipment and lease data center services to customers.

1.3 Data Center Cost Modeling

Existing cost models for data centers provide a framework for understanding the total cost of ownership (TCO), which encompasses both capital expenditures (CAPEX) and operational expenditures (OPEX). CAPEX includes the initial investment in fixed assets like land, buildings, and the core infrastructure for power and cooling. OPEX, covers the ongoing day-to-day costs of running the facility, such as electricity, maintenance, and staffing. Several models have been proposed that break down these costs, they generally include infrastructure, server hardware, networking, power, cooling and maintenance as the primary components.

A significant portion of a data center's TCO is attributed to the initial capital investment, with some estimates suggesting that CAPEX can account for 60-70% of the total costs in the first 5 years of operation [11]. However, over the lifetime of the data center, OPEX can grow to exceed the initial CAPEX, making operational efficiency a critical factor in long-term financial viability [11].

The cost of energy is a substantial and recurring operational expense, often representing 10-15% of the total cost of the data center [11]. This includes not only the direct cost of electricity from the grid but also the costs associated with the maintenance and depreciation of the power delivery and conditioning systems. Because a portion of the power consumed by IT

equipment is converted into heat, a robust cooling infrastructure is required, which in turn, also consumes a significant amount of power.

The 2007 Uptime Institute white paper by Koomey introduced a cost model that connected IT equipment choices to the cost of the data center [11]. It established a direct link between IT power consumption and facility costs by calculating the site infrastructure investment, including cooling and power systems. For operational expenditures (OpEx), the model calculated total electricity use by applying multipliers to the IT load to account for cooling and auxiliary systems and then it factors in a load percentage and average electricity price. Hamilton [12] presents a financial model designed to deconstruct the total monthly expenses of a “mega” data center. The monthly electricity bill is calculated based on the facility's critical IT power load, an average utilization percentage, the Power Usage Effectiveness (PUE) ratio, and the price per kilowatt-hour. In all of these models, the cooling system cost is determined by multiplying the IT CAPEX and OPEX cost by a set percentage without any actual cooling system detail.

In the 2017 paper "Economic analysis of data center cooling strategies," Cho et al. presents a detailed model to evaluate and compare the economic performance of seven different data center cooling strategies [13]. Unlike models that determine cooling costs as a fixed percentage or multiplier to IT capital and operational expenses, this research builds its analysis from the specific components and performance of the cooling systems themselves. The model is based on a life-cycle cost analysis over a 15-year period with 4 iterations and deconstructs the total cost into the following key components:

- Installation Cost (CAPEX): The initial investment is calculated based on the actual costs of the specific combination of components required for each cooling

strategy, such as chillers, cooling towers, pumps, and fans. The cost varies directly with the complexity and hardware requirements.

- Operational Cost (OpEx): This is broken down into multiple parts:
 - Energy Cost: This is the most significant factor in the model and is not determined by a simple PUE multiplier. Instead, the authors use detailed energy simulations to calculate the annual energy consumption of each specific cooling strategy. This simulated consumption is then multiplied by the unit price of electricity to determine the total energy cost.
 - Maintenance and Replacement Costs: Future costs are calculated using established annual maintenance rates and replacement periods for each individual component of the cooling system.
 - Carbon Dioxide Emission (CDE) Cost: The model also factors in the cost of carbon emissions, which is calculated based on the simulated energy consumption of each strategy.

In the 2022 paper "Optimal cost management of the CCHP based data center with district heating and district cooling integration in the presence of different energy tariffs," Keskin and Soykan introduce a comprehensive energy cost management model that treats the data center not as an isolated energy consumer, but as integrated within a regional energy network [14]. This model moves beyond static multipliers by using a dynamic optimization framework to actively manage energy sourcing and expenditure on an hourly basis. A key innovation of this model is the inclusion of revenue generation. The optimization framework decides when it is profitable to sell surplus energy. This includes selling excess electricity generated by the CCHP system back

to the grid, and, critically, selling captured waste heat and excess cooling capacity to regional district heating and cooling networks.

In the 2020 paper "Configuration Optimization Model for Data-Center-Park-Integrated Energy Systems under Economic, Reliability, and Environmental Considerations," Liu et al. develop a model to determine the optimal mix of energy technologies for a data center park [15]. The model uses Mixed Integer Linear Programming (MILP) to go beyond simple operational analysis and find the most economical equipment configuration from the initial planning stage, while also factoring in environmental impact and reliability requirements. The model uniquely quantifies the cost of high reliability. After determining an optimal baseline configuration, it calculates the additional initial investment required to add redundant equipment to meet specific availability targets (e.g., "four nines" or "five nines"). This creates a direct financial link between the cost of CAPEX and the desired level of system resilience.

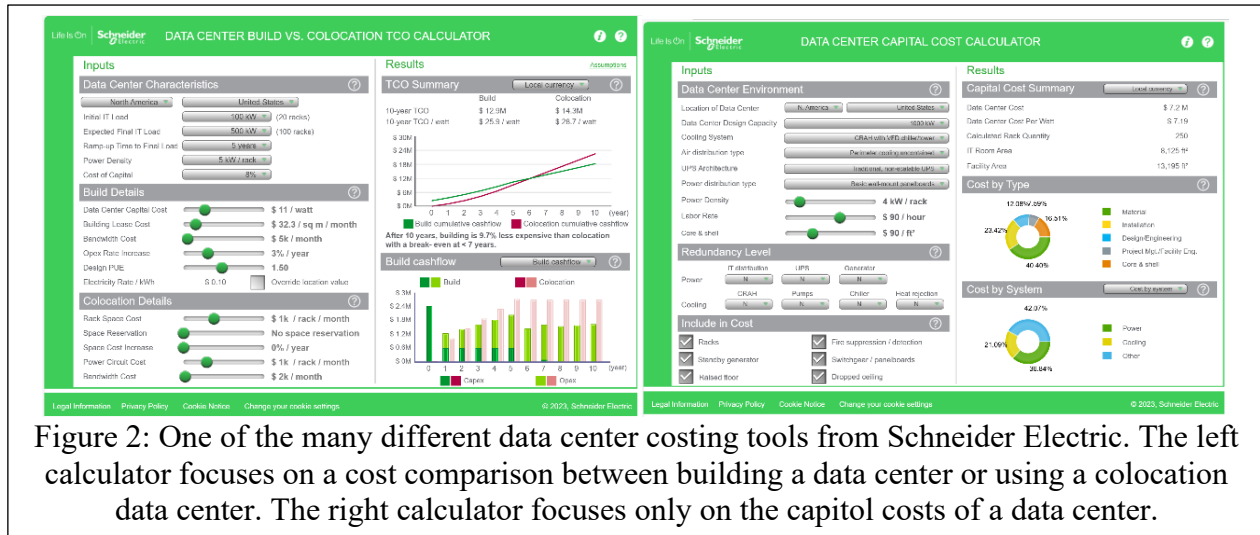
In the 2019 paper "Energy and economic assessment of major free cooling retrofits for data centers in Turkey," Gözcü and Erden present a financial model focused specifically on the economic viability of retrofitting existing data centers with various free cooling technologies [16]. Rather than modeling a new build, their approach evaluates the Total Cost of Ownership (TCO) of adding one of four different free cooling systems listed as direct air-side (ASE), indirect air-side (IASE), indirect evaporative (IEC), and indirect water-side (WSE) to a baseline 1 MW data center. The model's foundation is a detailed, hour-by-hour thermodynamic simulation which calculates the actual energy and water consumption for each potential retrofit based on local weather data. This detailed operational data is then fed into a comprehensive TCO model that is deconstructed into CAPEX and OPEX. This includes the cost of labor and replacement parts. The model projects these costs to increase annually based on historical

inflation and wage data in Turkey. The model then uses these cost streams to perform a comprehensive economic analysis over a 15-year lifetime, utilizing key financial metrics including Net Present Value (NPV), Modified Internal Rate of Return (MIRR), and the discounted payback period to compare the options.

1.4 Commercially Available Cost Models

Currently there are several commercially available models that perform cost modeling of data centers. These include: Schneider Electric data center models [17], Gordian RS Means [19], and The Green Grid [5].

Schneider Electric's data center cost models are divided into a set of independent models for different attributes of a data center [17][18]. For cooling systems, the Schneider Electric model allows the user to select a list of predefined cooling system designs. The cost comparison tool also focuses on building a data center compared to using a colocation service, which does not provide the flexibility to



Gordian RS Means focuses on the cost of building a data center and has an extensive list of components related to the building of a data center (CAPEX). This model does not provide

detail into the operational cost of a data center and does not detail on hardware in the cooling system's primary or secondary loop [19].

The Green Grid cost model provides an in-depth approach to the cooling system details. The Green Grid model allows the user to input parameters for air cooling and liquid cooling systems. An issue with this model is that the availability of the data center is not calculated and as a result it does not incorporate penalties that the data center may incur by failing to meet their customer requirements [13.2].

The models discussed above have several gaps, most notably the complete lack of models for the Service-Level Agreements (SLAs) and the potential penalties articulated within them. Additionally, while many models adjust CAPEX to account for redundancy, they often fail to consider the corresponding impact on the OPEX and the effect that newly added redundant systems have on downtime. The Mean Time Between Failures (MTBF) of cooling system components is also frequently omitted and the maintenance cost is considered a fraction of the data center capacity multiplied by a maintenance cost factor.

1.5 Thesis Objective Statement

The objective of this thesis is to develop a cost model that can be used to compare different data center cooling systems that includes capital cost, maintenance costs, energy costs, and penalties associated with downtime. Use the model to determine the relative cost of implementing alternative cooling systems and compare different SLA penalty models.

1.6 List of Tasks

The research outlined in this thesis was accomplished through the completion of four primary tasks:

1. Development of a Relative Cost Modeling Framework

The foundational task is to construct a financial framework for the relative comparison of data center cooling architectures. This framework was designed to evaluate the economic viability of a proposed cooling system against a designated reference system by calculating key relative performance indicators, including the Return on Investment (ROI) and Internal Rate of Return (IRR). The model accounts for the different capital expenditures (CAPEX) and operational expenditures (OPEX), energy consumption, and scheduled maintenance costs.

2. Creation of a Stochastic Discrete-Event Maintenance Simulator

To accurately model the impact of unscheduled downtime, a stochastic discrete-event simulator was developed. This simulation models the operational life cycle of cooling system components by incorporating probabilistic time-to-failure (MTBF) and time-to-repair (MTTR) distributions. A critical function of this simulator is its ability to account for various system redundancy configurations (e.g., N, N+1, 2N), thereby generating data on system availability, expected downtime frequency, and duration.

3. Modeling of Service-Level Agreement (SLA) Financial Penalties

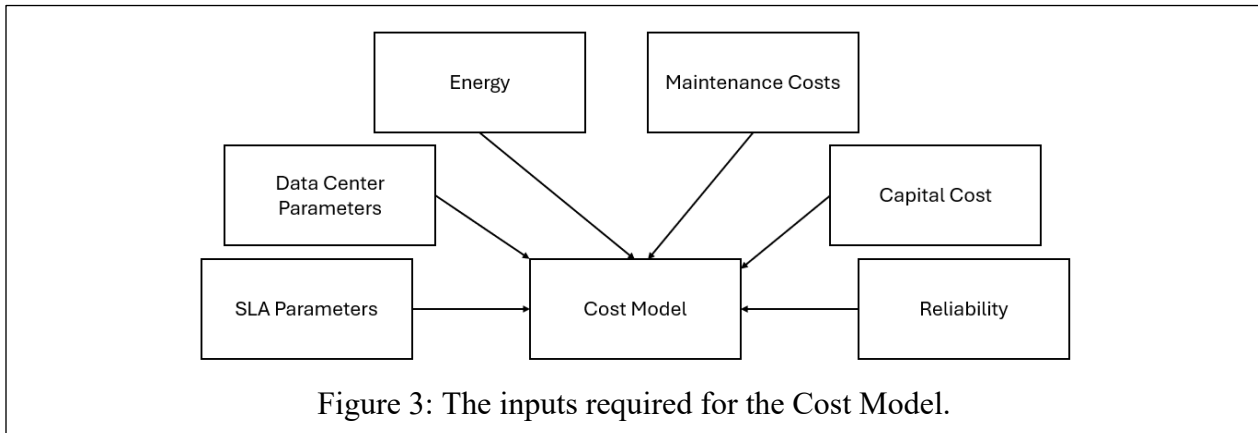
This task focuses on quantifying the financial consequences of downtime, a factor frequently omitted from standard Total Cost of Ownership (TCO) analyses. A flexible penalty model was developed to translate the downtime outputs from the discrete-event simulator into direct financial costs. The model is capable of representing various common SLA penalty structures, including those based on total cumulative downtime, the duration of a single continuous outage, and the frequency of service interruptions within a specified contract period.

4. Comparative Case Study and Sensitivity Analysis

The final task applies the model to a practical case study in order to demonstrate its capabilities. These case studies conduct a comprehensive financial comparison between a conventional data center cooling system and an alternative technology. The analysis combines all components of the model to project the total life-cycle cost, including capital investment, operational expenses, and potential SLA penalties. A sensitivity analysis is performed to identify the most critical variables influencing the economic outcome and to explicitly evaluate the financial trade-off between investing in higher redundancy and accepting a greater risk of downtime penalties.

Chapter 2: Cost Modeling

A cost model for a data center cooling system consists of many factors such as the capital cost, energy costs, maintenance costs (Figure 3). Capital costs represent the initial investment required to purchase and install the cooling infrastructure. These costs are influenced by the cooling architecture and the level of redundancy used within the data center.



Energy consumption is a major component of a data center's operational expenditure. IT and cooling consume the majority of the energy. In many cases, the data center is defined by its power budget, i.e., designers may start with the power budget and work from there to determine the size of the IT needed for the data center [11]. In fact, several of the referenced cost models (e.g., [10], [11]) are organized in this way, i.e., they start with the power and work backwards to determine the number of servers.

Maintenance costs include all expenses related to servicing, repairing, and replacing components of the cooling system to ensure its continued operation. These can include service contracts, replacement parts, and labor.

2.1 ROI and IRR Calculations

Relative cost modeling is a comparative approach that evaluates the financial viability of one investment against another. Instead of focusing on the absolute TCO of a single option, it calculates the differential costs between two systems. Since all the system costs that are the same for the two systems subtract out, only the costs that differ need to be determined. In the context of this thesis, a relative cost model will be developed to specifically compare a new, alternative cooling system against the original (or reference) system. This model will not only consider the upfront capital and ongoing maintenance and energy costs but will also incorporate the financial penalties associated with downtime, which can be a significant factor in the overall economic assessment. By calculating a Return on Investment (ROI) and Internal Rate of Return (IRR) based on these relative costs, the model will provide a clear and direct answer as to whether the investment in the new cooling technology is economically justifiable (and when it becomes economically advantages, i.e. the breakeven point). This approach moves beyond a simple TCO calculation and provides a more dynamic and decision-oriented tool.

The following subsections detail how the return on investment (ROI) and internal rate of return (IRR) are determined in the model.

2.1.1 $I_{new} > I_0$ Analysis Model

If the investment in the new cooling approach is larger than (or in addition to) the investment in a reference cooling system, then we use a commonly used cost-avoidance return on investment [1],

$$ROI = \frac{C_0 - C_{new}}{I_{new} - I_0} = \frac{\Delta NPV}{I_{new} - I_0} \quad (1)$$

where,

C_0 = total life-cycle cost using the reference approach

C_{new} = total life-cycle cost using the new approach

I_0 = investment cost in the reference approach

I_{new} = investment cost in the new approach

The total life-cycle costs and therefore, the ROI are time dependent. All cost contributions within C_0 and C_{new} are discounted to present value, I_0 and I_{new} are assumed to happen at time 0. ROI represents the net return over the entire time period of the investment.¹ The internal rate of return (IRR), which is the annual rate of return, can be calculated from the ROI using:

$$(1 + ROI)^1 = (1 + IRR)^{\text{entire investment period (in years)}} \quad (2)$$

Breakeven are the points in time when the $IRR = 0$. Note, there may be several points on the timeline where the rate of return transitions from negative to positive and vice versa. IRR is in general a complex number when $ROI < -1$. In odd investment periods (i.e., odd years), the real IRR solution can be calculated.² In even investment periods (i.e., even years), the IRR consists of only complex solutions.

The annual IRR at a particular time step is determined from the ROI (over all time steps) from

$$IRR = (1 + ROI)^{n/m} - 1 \quad (3)$$

where n is the number of periods per year and m is the current period (time step). For example, if the analysis has monthly time steps, then 1 period = 1 month, $n = 12$ periods/year, m = the current time step, and ROI is the rate of return over the entire m time steps.

The cost avoidance ROI is the ratio of the total life-cycle cost difference between the data center with the new cooling system and the data center with a reference cooling system, and the

¹ In this thesis we need a “cost avoidance” ROI because the cooling system does not generally earn revenue, rather it avoids future spend (on IT equipment maintenance and replacement). The heat removed may have value (which we consider), but the cooling system is not a net revenue earning portion of the data center.

² In this case the IRR solutions are composed of a real answer and one or more complex conjugates.

difference in the investments in the new and reference cooling systems. This cost avoidance *ROI* is useful because the numerator and denominator are both differences, therefore, only the life-cycle costs that differ between the new and reference approaches need to be modeled (all other costs subtract out). The life-cycle costs that have to modeled include: capital equipment for cooling, energy for cooling, maintenance of the cooling system, maintenance frequency of the IT equipment (if it is impacted by the new cooling approach), heat removal value from the IT equipment, and other data center properties that vary from the reference case. The model will also depend on discount rate and inflation rates.

2.1.2 $I_{new} < I_0$ Analysis Model

It is likely that some of the proposed cooling technology changes will result in overall decreases in the investment (in the cooling system) so that $I_{new} < I_0$. When $I_{new} < I_0$ a comparison with the reference case can still be made, but in this case a negative investment is made to get a return, so *ROI* and *IRR* are not the appropriate measures to use, i.e., there is no *ROI* or *IRR* in these cases. In these cases, an additional suggested measure is ΔNPV , which is the difference in the net present value between the new case and the reference case,

$$\Delta NPV = C_0 - C_{new} \quad (4)$$

where, NPV is time dependent and all cost contributions within C_0 and C_{new} are discounted to present value. Note, I_0 and I_{new} are not ignored, they are included within C_0 and C_{new} . Equation (4) is the life-cycle cost difference that appears in the numerator of the *ROI* in Equation (1).

Note the ΔNPV is equal to zero at the breakeven point (the point where $IRR = 0$).

2.2 Energy

The power associated with each component in the cooling system is determined using EnergyPlus, which is a comprehensive, open-source whole-building energy simulation program funded by the U.S. Department of Energy (DOE), [26]. It is widely used by architects, engineers, and researchers to model energy consumption for heating, cooling, ventilation, and lighting in buildings. For data centers, EnergyPlus can calculate heating and cooling loads necessary to maintain thermal setpoints, simulate the performance of HVAC systems, and determine the resulting energy consumption.[27] Its heat balance-based solution technique allows for a detailed and integrated simulation of the thermal interactions between the building envelope, internal loads, and the mechanical cooling systems, providing a powerful tool for optimizing data center cooling design and operational strategies.

2.3 Maintenance Discrete-Event Simulator

The cost model includes a discrete-event simulation (DES) that models the time-history of component failures and their corresponding maintenance events. The DES is embedded within a Monte Carlo model. The DES starts by sampling the time to failure distribution of each component in the system to determine the first failure event. When the first component fails, the system is assumed to be down until the failure is repaired, Figure 4. Once the repair is made the system is restored to operation. If the failed component has redundancy, then the system will continue to operate with the redundant component until either the redundant component fails or the original component is repaired, whichever comes first, Figure 5. The DES then determines the next component failure and repeats the process until the end of support for the system is reached. The downtime and discounted maintenance costs are accumulated.

The Discrete Event Simulator (DES) has the following hierarchy:

- **Data Center:** The Data Center is considered operational only if all its component groups are simultaneously operational.
- **Component Group:** An intermediate level comprising one or more identical components for the purpose of handling redundancy. The operational status of a component group is defined by a redundancy such as N, N+1 or 2N. This allows for component failures without the data center failing.
- **Component:** Each component's behavior is characterized by probabilistic time-to-failure distributions determined from the mean time-between-failure (MTBF) (or Weibull parameters if provided) and the time-to-repair determined from the mean time-to-repair (MTTR).

The input parameters for defining each component include:

- **Time-to-Failure Distribution:** Component failures are modeled using a three-parameter Weibull distribution, characterized by its shape (β), scale (η), and location (γ) parameters. If only an MTBF for a component is defined, it is converted to a three-parameter Weibull for use in the DES.
- **Time-to-Repair Distribution:** The MTTR for components is modeled using a Lognormal distribution.
- **Maintenance Cost:** A fixed cost associated with each repair completion event for a component.

The simulation progresses through a set of discrete points in time when significant events occur. Each Monte Carlo sample, representing the system's operational life over a specified duration, follows these steps:

At the start of each sample run, the simulation clock is set to zero. All components are initialized to an operational state with zero damage (this assumes that the data center starts as good as new at time zero). An initial failure event is sampled for each component from its respective Weibull distributions, and a corresponding failure event is scheduled. The initial operational status of each component group and the overall data center is then determined according to their redundancy configurations.

The simulation proceeds by:

1. **Identifying the Next Event:** The algorithm scans all scheduled events (future failures of components and future repair completions of failed components) across all component groups. The event with the chronologically earliest time is identified.
2. **Advancing the Clock:** The simulation clock is advanced directly to the time of this next event.
3. **Processing the Event:** Upon reaching the event time, the system state is updated based on the event type:
 - **Failure Event:** The affected component transitions from an "UP" to a "DOWN" state. Its next failure event is canceled (set to infinity). A MTTR is sampled from its Lognormal distribution, and a repair completion event is scheduled for the current clock time plus the sampled MTTR. A beta value of 1 is assumed for the Lognormal distribution sampling. The component's operational age is updated.
 - **Repair Completion Event:** The component transitions from a "DOWN" to an "UP" state. The associated maintenance cost is recorded. Assuming as good as

new repair, a new failure event is sampled from its Weibull distribution, scheduling a new failure event.

4. **Updating System State and Redundancy Logic:** Following any component state change, the redundancy and standby policy logic is invoked:

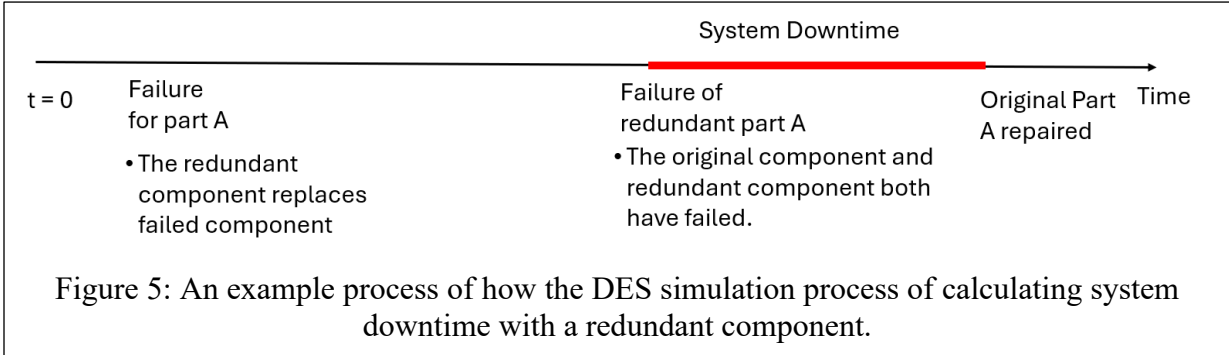
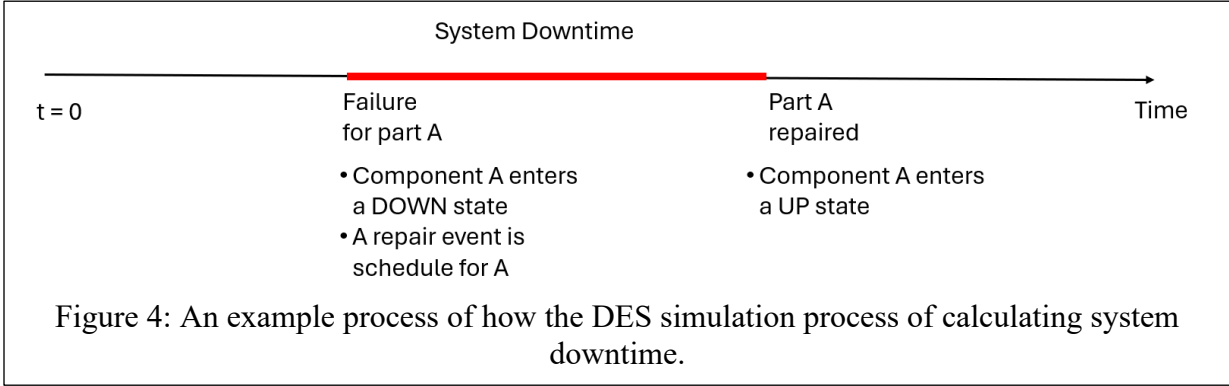
- The number of operational components within the affected group is re-evaluated.
- The component group's state is updated: if the count of "UP" components falls below the number of components required for the cooling system, the group transitions to "DOWN". Conversely, if it rises to component multiplier or above, it transitions to "UP".
- The overall data center's operational status is then re-assessed. If any component group is "DOWN", the data center becomes "DOWN". If all component groups are "UP", the data center is "UP".
- Downtime and uptime intervals for both component groups and the overall data center are logged based on these state transitions to determine.

5. **Termination:** The event loop continues until the simulation clock reaches or exceeds the simulation duration.

This complete process constitutes a single Monte Carlo sample. The entire procedure is repeated for many samples to generate a statistical basis for analysis.

The DES makes several fundamental assumptions. Component failures are treated as independent events. The failure of one component does not directly influence the failure

probability of another component. All repairs of a component are assumed to restore the component to as good as new condition.



2.3.1 Validating Discrete-Event Simulator

The DES can be validated against closed-form availability calculations that assume idealized conditions. Assuming one component with an availability of a and a redundancy given by R . $R = 1$ is no redundancy (N), $R = 2$ assumes 1 redundant component, i.e., N+1 redundancy, etc. The availability of the one-component system is given by [28],

$$A = 1 - f(1 - a)^R \quad (5)$$

where f is the number of ways component can fail. This relation also assumes that the system is running in steady-state.

In order to test the DES, a single component is defined with an MTBF of 400 hours, an MTTR of 300 hours, and a redundancy of $N+1$. The DES is run for a long time (100 years) in order to minimize the good-as-new component starting point.

The availability of the calculations using equation 5 is 81.63%. To have the DES match the availability of equation 5 then the components are assumed to be repaired good as new. The redundant components are assumed to be hot redundant and therefore continue to experience operational wear and performance degradation even when not actively carrying load.

Chapter 3: Modeling Service-Level Agreements

Service-level agreements (SLA) are agreements between the data center and the data center customer where the data center agrees on a penalty structure if downtime occurs that results in a loss of operation for the data center customer. SLA's for data centers are not generally published, but a few examples are publicly available: [17,19,20,25,26].

In an SLA for a data centers include maintenance scheduling, downtime durations that qualify for monetary compensation, and support services for hardware and software.

Current work that discusses SLA penalties focuses mainly on IT or software failure within a data center. Generally, these research papers do not discuss the downtime penalties themselves but simply use them as a motivator for their research [21][22], i.e., there is no quantitative treatment of SLA penalties in existing data center cost analyses.

This type of penalty analysis could be critical in the initial planning of the data center. Having redundancy on every cooling component within the data center could reduce the downtimes significantly. Is this trade-off in both CAPEX and OPEX worth it to avoid paying the downtime penalties?

3.1 Taxonomy of SLA Penalty Structures

Service-Level Agreements in the data center industry are not standardized, but the penalty structures generally fall into distinct categories based on how downtime is measured and penalized. The specific terms of an SLA can alter the profitability of a data center. The same cooling system could be economically advantageous under one agreement and financially unviable under another. The penalties are typically triggered based on measurements of downtime within a defined contract period, such as one calendar month. The primary

classifications of these penalty structures are based on downtime duration and interruption frequency. [23][24][29][30]

Generally, the downtime for a data center is measured in one of three ways: total downtime, continuous downtime, or number of interruptions. These downtimes are measured in a contract measurement period defined in the SLA. In a total downtime contract, only the combined downtime for all events in the contract measurement period is used to calculate a penalty. In a continuous downtime contract, a continuous downtime that exceeds a threshold amount defined in the SLA is penalized. For an interrupts based contract, the penalty is assessed based on the number of separate interrupts experienced during the contract measurement period. Figure 6 shows a taxonomy of the types of SLAs and their penalties.

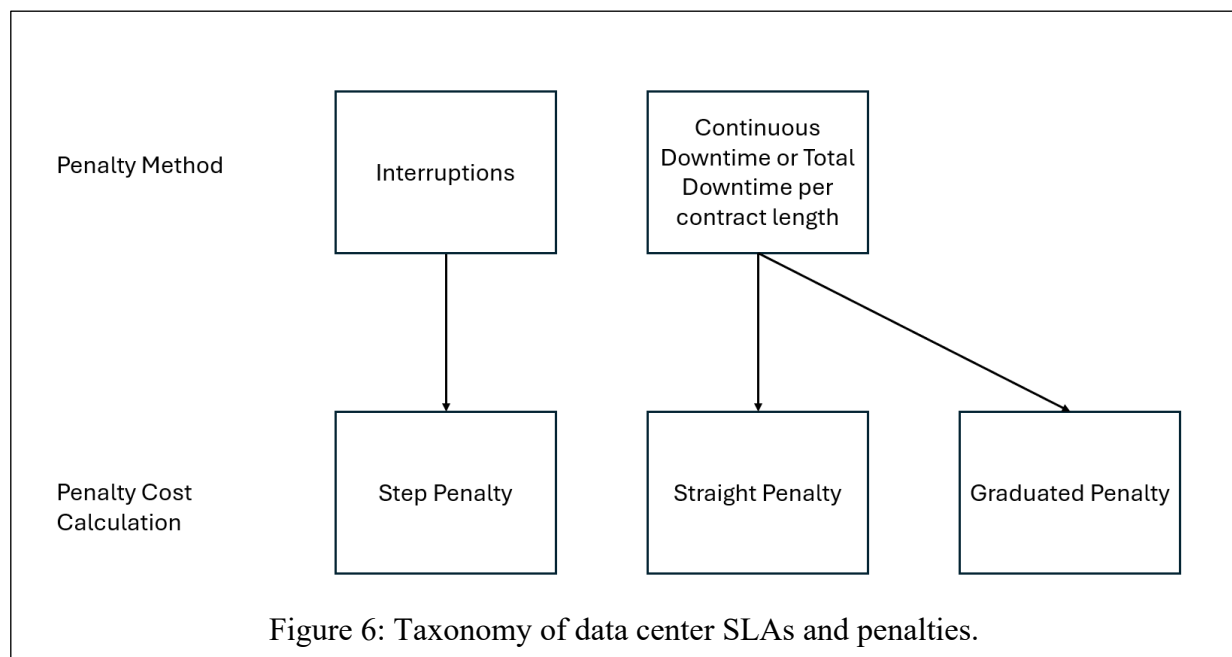


Figure 6: Taxonomy of data center SLAs and penalties.

3.1.1 Duration-Based Penalties

These penalties are tied to the amount of time a service is unavailable. They are designed to compensate the customer for the direct impact of an outage.

- **Total Cumulative Downtime:** This method sums up all independent downtime events within the contract measurement period. For example, four separate outages of two minutes each would result in eight minutes of total downtime. This model penalizes overall system unreliability, as frequent, minor disruptions can accumulate to trigger significant penalties.
- **Continuous Downtime:** This method targets single, prolonged outages. A penalty is only applied if a single, unbroken period of downtime exceeds a predefined threshold (e.g., five consecutive minutes). This model is less sensitive to minor system downtime but imposes severe costs for catastrophic failures that require lengthy repair times. This reflects the reality that a 30-minute continuous outage is often far more damaging to a customer's operations than thirty 1-minute outages.

3.1.2 Frequency-Based Penalties

Some customer operations are severely impacted by the interruption itself, regardless of duration, as it may force a restart of complex computational jobs. To address this, SLAs can penalize the number of service interruptions.

- **Number of Interrupts:** This model assesses a penalty based on the count of distinct downtime incidents within the contract measurement period. A system that fails for one minute five times in a month would be penalized more heavily than a system that fails for five continuous minutes once. This structure is often designed with escalating costs, where each subsequent interruption in a period is more costly than the last, i.e., heavily penalizing system instability.

3.1.3 Mathematical Framework for Penalty Calculation

To integrate SLA penalties into the life-cycle cost model, the downtime events generated by the discrete-event simulator must be translated into a direct financial cost. The penalty structures described previously can be implemented by incorporating common contractual features such as grace periods (a downtime allowance before penalties begin) and ceilings (a maximum penalty per unit time).

The following mathematical models, which directly correspond to the penalty functions implemented for this thesis analysis, define the specific calculations used to translate downtime into a financial cost.

3.2.1 Duration-Based Penalty Models

These models calculate penalties based on the length of an outage.

Linear Duration Penalty

This model applies a fixed-rate penalty for each hour of downtime that exceeds a contractual grace period, up to a maximum penalized duration. The cost is calculated as,

$$\text{Linear Penalty Cost per Unit Time} = \begin{cases} 0, & d \leq T_{\text{grace}} \\ pd, & T_{\text{grace}} < d \leq T_{\text{max}} \\ p(T_{\text{max}}), & d > T_{\text{max}} \end{cases} \quad (6)$$

where,

p = is the linear penalty rate.

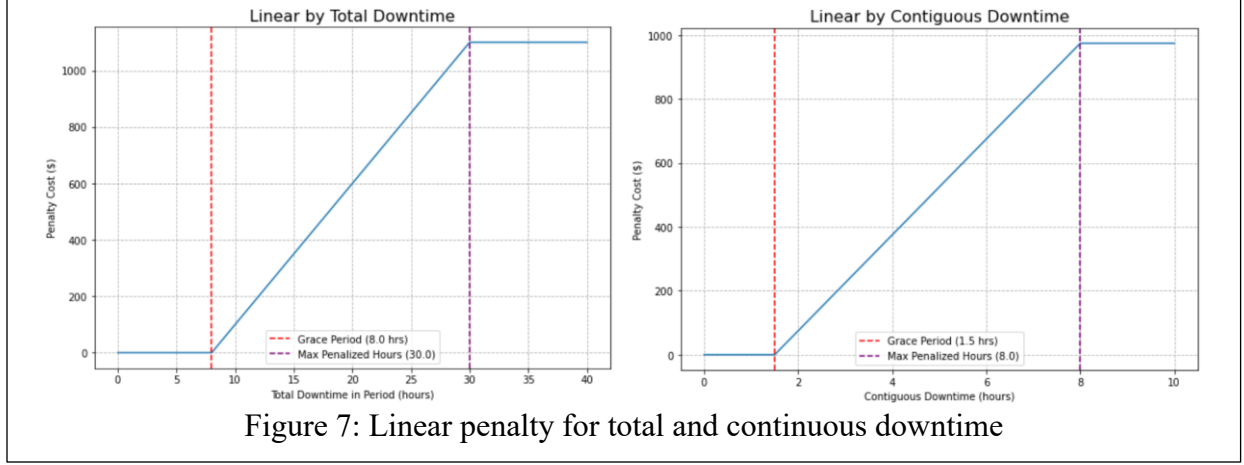
d = is the duration of the downtime being measured.

T_{grace} = is the grace period, or the downtime allowed before penalties are incurred.

T_{max} = is the penalty ceiling, the maximum duration for which penalties will be assessed.

Equation (6) is applicable to both the total downtime and continuous downtime models.

Figure 7 shows an example of both models.



3.2.2 Graduated Duration Penalty

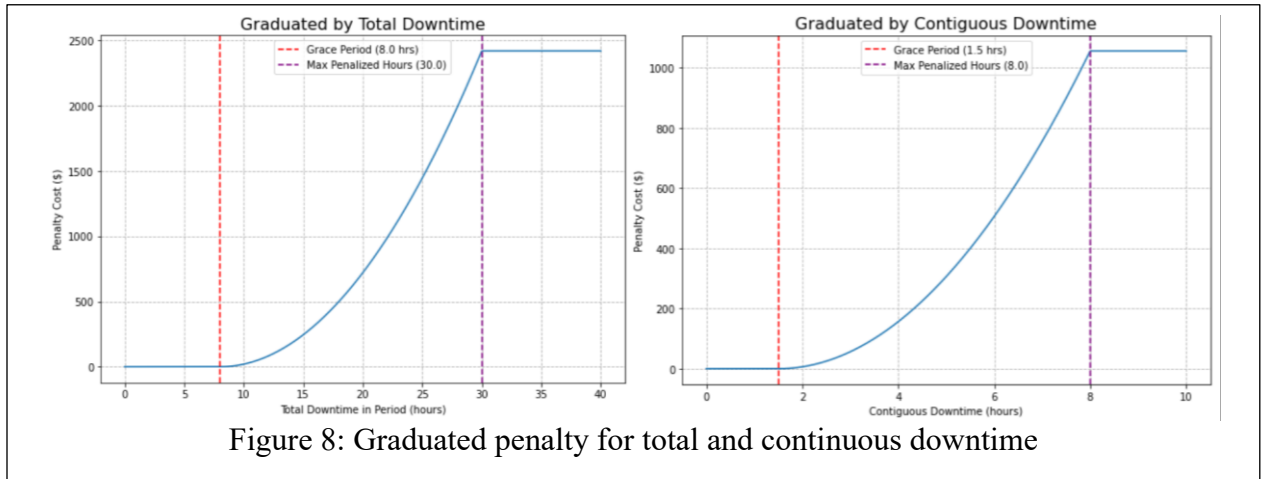
This model applies a quadratically increasing penalty for downtime that exceeds the grace period.

$$\text{Graduated Penalty Cost per Unit Time} = \begin{cases} 0, & d \leq T_{\text{grace}} \\ \frac{p(T_{\text{max}})}{T_{\text{max}} - T_{\text{grace}}} (d - T_{\text{grace}}), & T_{\text{grace}} < d \leq T_{\text{max}} \\ p(T_{\text{max}}), & d > T_{\text{max}} \end{cases} \quad (7)$$

where,

All other variables (d , T_{grace} , T_{max}) are as defined previously.

Figure 8 shows the graduated duration penalty applied to the total downtime and continuous downtime models.



3.2.3 Escalating Interrupt Penalty

This model applies to a penalty for each distinct interruption, with the cost increasing for each subsequent event up to a maximum number of penalized interruptions.

$$\text{Interrupt Penalty Cost per Unit Time} = \begin{cases} 0, & i \leq I_{\text{grace}} \\ c(i - I_{\text{grace}}), & I_{\text{grace}} < i \leq I_{\text{max}} \\ c(I_{\text{max}}), & i > I_{\text{max}} \end{cases} \quad (8)$$

where,

c = is the penalty per interruption.

i = is the number of interrupts measured.

I_{grace} = is the grace period, or the interrupts allowed before penalties are incurred.

I_{max} = is the penalty ceiling, the maximum interrupts for which penalties will be assessed.

Figure 9 shows the interrupt penalty model applied.

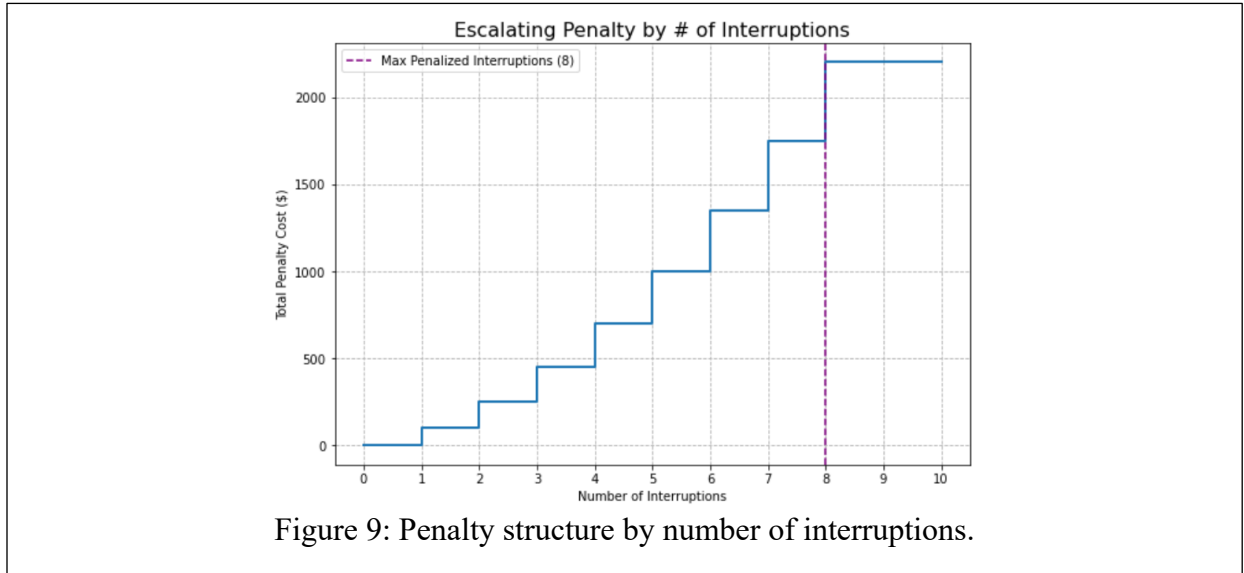


Figure 9: Penalty structure by number of interruptions.

Chapter 4: Case Study

This chapter presents a collection of case studies performed using the model described in Chapters 2 and 3.

4.1 Data Center Description – Reference Architecture

The data center considered in all the case studies is the CEETHERM lab data center at Georgia Tech [25].³ While this facility provides a detailed component-level configuration, its scale is not representative of a commercial data center. To bridge this gap and ensure the analysis reflects a commercially relevant context, a linear adjustment based on data center power load can be used to scale the financial impact of downtime.

The data center consists of a primary cooling loop, which is configured in a non-redundant layout. The major components include a single 30 kW cooling tower, a primary 150 kW chiller, and a 10 kW Computer Room Air Handling (CRAH) unit as shown in Figure 10. These components represent critical single points of failure for the entire system.

The data hall is configured with 10 server racks. These racks house a total of 40 servers, resulting in a density of 4 servers per rack. The cooling architecture within each rack consists of a dedicated 2 kW heat exchanger, a coolant reservoir, and a 12 kW pump to circulate the liquid coolant to the servers.

At the component level, each of the 40 servers employs a hybrid cooling strategy. Low-power components are cooled by a single 0.05 kW fan. The primary heat-generating components

³ This is the reference data center considered in the ARPA-E COOLERCHIPS program.

are managed by a direct-to-chip liquid cooling loop. This loop services four high-power chips, with each chip coupled to a dedicated liquid cold plate.

BOM Level	Name	Multiplier	Cost (dollars)	Redundancy	MTBF (hours)	MTTR (hours)	Cost per Maintenance Event (dollars)	Maintenance Type	Power (Watts)
	1 Cooling Tower		300,000		1504800.00	16.67	10,000		30
	1 Pump 1		20,000		78271.75	3.67	1,000		12
	1 Chiller		600,000		190872.00	14.52	15,000		150
	1 Pump 2		20,000		78271.75	3.67	1,000		12
	1 Distributor		8,000		100000.00	100.00	500		1
	1 Collector		8,000		100000.00	100.00	500		1
	1 CRAH		25,000		61701.73	6.22	2,000		10
	1 Rack	10							
	2 Heat Exchanger		15,000		85477.39	3.52	1,000		2
	2 Pipe 1		10,000		2127659.57	100.00	500		0
	2 Coolant Reservoir		5,000		2413680.00	0.50	300		0
	2 Pipe 2		10,000		2127659.57	100.00	500		0
	2 Pump 3		20,000		78271.75	3.67	1,000		12
	2 Server	40							
	3 Fan		100		150000		5		0.05
	3 Liquid Cooled Chip 1		500		100000		100		1.5
	3 Liquid Cold Plate 1		100		100000		100		0
	3 Liquid Cooled Chip 2		500		100000		100		1.5
	3 Liquid Cold Plate 2		100		100000		100		0
	3 Liquid Cooled Chip 3		500		100000		100		1.5
	3 Liquid Cold Plate 3		100		100000		100		0
	3 Liquid Cooled Chip 4		500		100000		100		1.5
	3 Liquid Cold Plate 4		100		100000		100		0

Figure 10: Georgia Tech CEETHERM Data Center.

4.1.1 Parameters in Cost Model

Several input parameters are held constant across all case studies to isolate the impact of the cooling architectures and SLA structures. These fixed values are:

- Electricity Cost: 0.145 \$/kWh
- Energy Inflation: 2.24% per year
- Inflation: 2.5% per year
- Discount Rate: 6% per year
- Simulation Time Step: Monthly

4.1.2 Data Center Scaling Penalties

To validate the linear scaling assumption applied to downtime penalties, calculations were performed comparing the baseline architecture to a scaled equivalent data center. In order

to calculate ROI and IRR for the reference system shown in Figure 10, the analysis system will have a modified energy and capital cost. Figure 11 illustrates two scenarios, Run 1 represents the baseline Georgia Tech data center utilizing the total cumulative downtime metrics. Run 2 models Georgia Tech data center scaled with a multiplier 10 times of Run 1. This case increases the capitol cost and energy cost by 10 times the original amount. The resulting ROI and IRR are nearly identical across both simulations after running the DES for 100 time histories. This consistency confirms that the financial impact of downtime penalties scales linearly within the model.

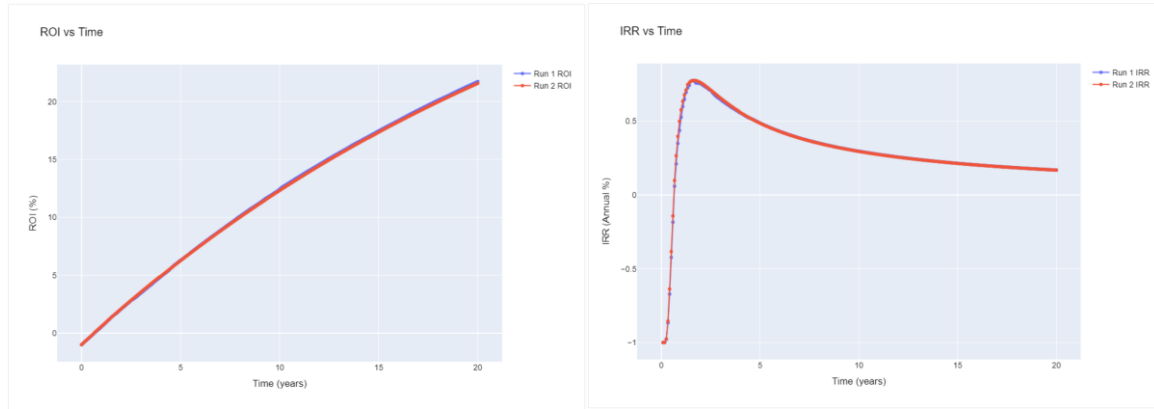


Figure 11: ROI and IRR of the Georgia Tech data center and a data center that is 10 times larger than Georgia Tech data center

4.2 Case Study 1: The Impact of Varying SLA Structures

This case study used the baseline Georgia Tech configuration to analyze its financial performance under three different SLA penalty structures: (a) a penalty based on total cumulative downtime per month, (b) a penalty based on any single continuous downtime event exceeding a threshold, and (c) a penalty based on the number of service interruptions per month. This analysis demonstrates how the financial risk profile of the same physical infrastructure can

change depending on the terms of the customer agreement, highlighting the importance of including SLA modeling in the design phase of the data center.

4.2.1 Case Study 1 Data Assumptions

To analyze the financial impact of downtime, three distinct, SLA penalty structures were modeled. These structures are designed to be representative of common industry practices discussed in Chapter 3.

1. **Total Cumulative Downtime:** This structure penalizes the data center based on the total minutes of downtime accumulated over a one-month period.
 - **Grace Period:** 8 minutes per month
 - **Penalty:** \$5,000 for every minute of downtime beyond the grace period.
 - **Ceiling:** 10 hours per month
2. **Continuous Downtime:** This structure imposes a penalty only when a single, uninterrupted outage exceeds a specified duration.
 - **Threshold:** 8 minutes
 - **Penalty:** \$5,000 for every minute exceeding the threshold.
 - **Ceiling:** 10 hours per month
3. **Number of Interruptions:** This structure penalizes the data center for the frequency of downtime events, regardless of their duration.
 - **Grace Period:** 1 interruption per month
 - **Penalty:** \$100,000 per interruption
 - **Ceiling:** 30 interruptions per month

4.2.2 Case Study 1 Results

Before analyzing the specific financial risks imposed by varying SLA structures, it is necessary to establish a baseline profile for the cooling system. This baseline assumes all downtime penalties are zero.

In Figures 11 and 12, Run 1 establishes the baseline scenario with no downtime penalties. The remaining runs demonstrate the impact of the three different SLA structures: Run 2 applies the Total Cumulative Downtime penalty, Run 3 represents the Continuous Downtime penalty, and Run 4 depicts the Number of Interrupts penalty.

In Run 1 of Figures 12 and 13, both the ROI and IRR curves intersect the x-axis at the exact same point in time. This is the breakeven point, where the discounted operational savings offset the upfront capital expenditure. Beyond this point, the ROI curve trends upward but exhibits a concave trajectory, as time progresses rather than following a linear path. This non-

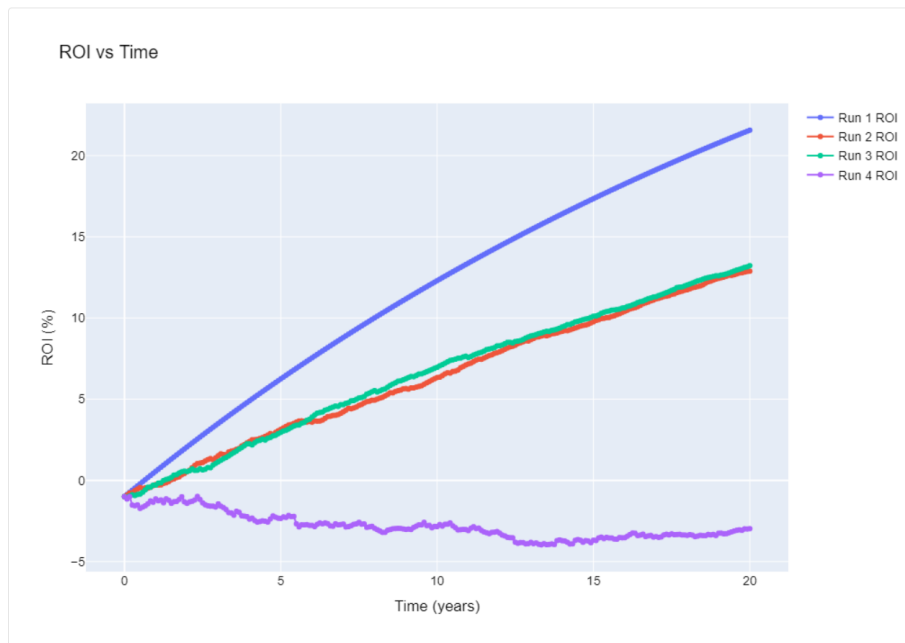


Figure 12: ROI baseline and for the three different scenarios listed in Section 4.2.1

linear behavior is a direct consequence of the non-zero discount rate, which is 0.06 %/year in this case. As the operational timeline extends into the future, the present value of the fixed annual energy and maintenance savings diminishes. The incremental growth of the Net Present Value slows down, causing the slope of the ROI to flatten.

The IRR curve displays a distinct behavior where it rises sharply immediately following the breakeven point, reaches a maximum peak, and subsequently decreases. The rapid initial rise reflects the ratio of cost avoidance to investment duration. The annualized nature of the IRR metric causes the rate to decline, as constant annual avoidance are averaged over a longer period. The specific time at which the IRR peaks is economically significant, representing the point of maximum efficiency. While operating the system beyond this peak continues to yield a financial

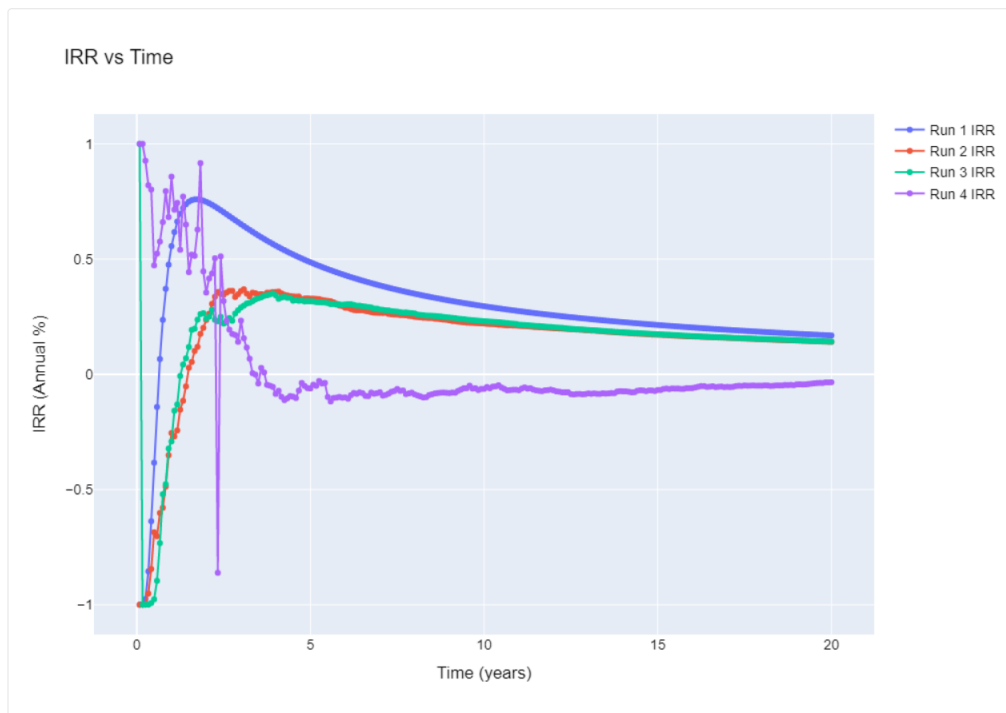


Figure 13: IRR baseline and for the three different scenarios listed in section 4.2.1

advantage, the rate of return relative to the time invested begins to diminish.

For the total cumulative downtime scenario, Figure 12 shows breakeven point is reached at 0.515 years for Run 1. The total downtime model (Run 2) has a breakeven year of 0.684 years. This SLA structure is sensitive to frequent, small disruptions that accumulate over time. A cooling system with recurring issues would be heavily penalized under this model. Under the continuous downtime model (Run 3), the breakeven point is extended to 0.633 years. This structure is less concerned with minor, short-duration interruptions and instead heavily penalizes prolonged, continuous outages. The number of interruptions scenario (Run 4) never reaches a positive ROI value and as a result does not have a breakeven year. This penalty structure is not influenced by the duration of an outage and instead punishes any interruption to service. This penalty method is for operations that are severely impacted by the need to restart processes after any downtime event. Run 4 is very erratic due to the many downtime events that last a couple of hours and heavily penalize these events compared to total downtime and continuous downtime penalty methods.

Figure 13 provides a sensitivity analysis of the breakeven time as a function of increasing SLA penalty costs for total downtime penalty. The mean breakeven year exhibits a flat trajectory as the penalty cost increases. This behavior is due to the reference and analysis architectures sharing the same redundancy and reliability values, resulting in statistically similar availability profiles. Because both systems incur the same frequency and duration of downtime, the associated SLA penalties rise symmetrically for both (Figure 14 represents a validation exercise, i.e., the characteristic should be flat). The standard deviation increases as the penalty size increases due to stochastic nature of the Discrete-Event Simulator. This analysis will be

revisited in Case study 3 where the reference and analysis cases do not have the same redundancy.

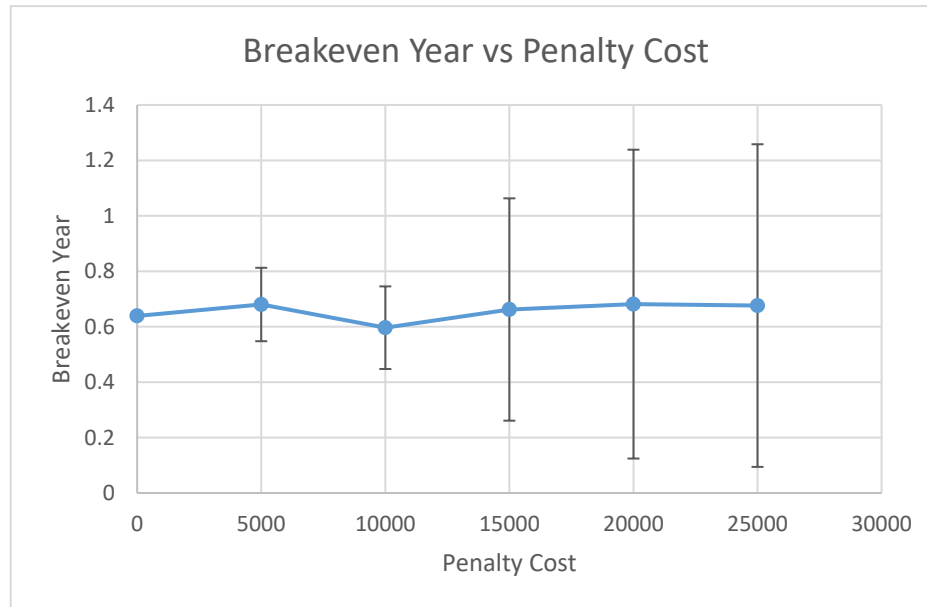


Figure 14: Sensitivity analysis for breakeven year vs. penalty cost

4.3 Case Study 2: Direct Liquid Cooling vs. Traditional Air Cooling

The baseline configuration utilizes a hybrid approach with a Computer Room Air Handling (CRAH) unit at the facility level and direct liquid cooling at the chip level. This study compares the baseline system to a traditional air-cooled system of equivalent cooling capacity. This will require developing a representative Bill of Materials (BOM) for the air-cooled alternative. The comparison will highlight trade-offs in capital cost, energy efficiency, and system reliability between the two distinct cooling philosophies.

4.3.1 Case Study 2 Data

To achieve the same level of heat rejection without direct liquid cooling, this system requires a significantly larger and more powerful CRAH unit. The primary cooling loop (chiller, cooling tower, pumps) remains the same to isolate the comparison to the secondary loop. All components related to in-rack liquid cooling (Liquid Cooled Components) are removed.

BOM Level	Name	Multiplier	Cost (dollars)	Redundancy	MTBF (hours)	MTTR (hours)	Cost per Maintenance Event (dollars)	Power (Watts)
1	Cooling Tower		650,000		1504800.00	16.67	30,000	45
1	Pump 1		40,000		78271.75	3.67	1,000	40
1	Chiller		1,200,000		190872.00	14.52	50,000	400
1	Pump 2		40,000		78271.75	3.67	1,000	40
1	Distributor		8,000		100000.00	100.00	500	1
1	Collector		8,000		100000.00	100.00	500	1
1	CRAH		800,000		61701.73	6.22	2,000	150
1	Rack	10						
2	Server	40						
3	Fan		100		150000		5	0.05
3	Server Components		1,000		300000.00		50	1.5

Figure 15: Georgica Tech Data Center modified for air cooling.

4.3.2 Case Study 2 Results

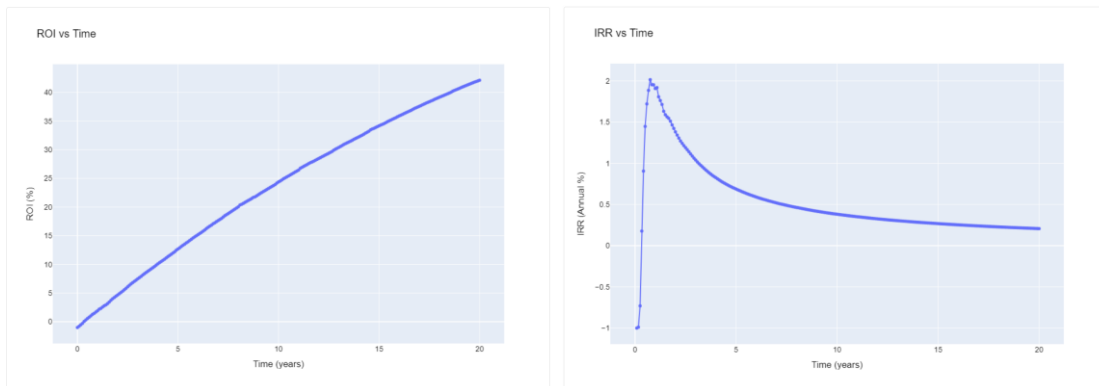


Figure 16: Liquid Cooling vs Traditional Air-Cooling ROI and IRR with the same downtime penalties

The comparison between direct liquid cooling and traditional air cooling reveals a significant financial advantage for liquid cooling, with a breakeven point occurring at 0.233 years.

While the initial capital expenditure for a direct liquid cooling system can be higher due to the need for specialized components like cold plates and a more complex fluid distribution network, this is often offset by substantial operational savings. Direct liquid cooling is significantly more energy-efficient than air cooling. By bringing the cooling medium directly to the heat-generating components, heat transfer is far more effective, reducing the energy required to maintain optimal operating temperatures. This translates to lower electricity costs, a major component of a data center's operational expenditures. Additionally, the increased cooling capacity of liquid cooling allows for greater server density within the same physical footprint. While both have potential points of failure, the reduced strain on a more efficient liquid cooling system can lead to a lower frequency of certain types of failures, thereby reducing downtime and associated SLA penalties.

4.4 Case Study 3: Baseline vs. Enhanced Redundancy

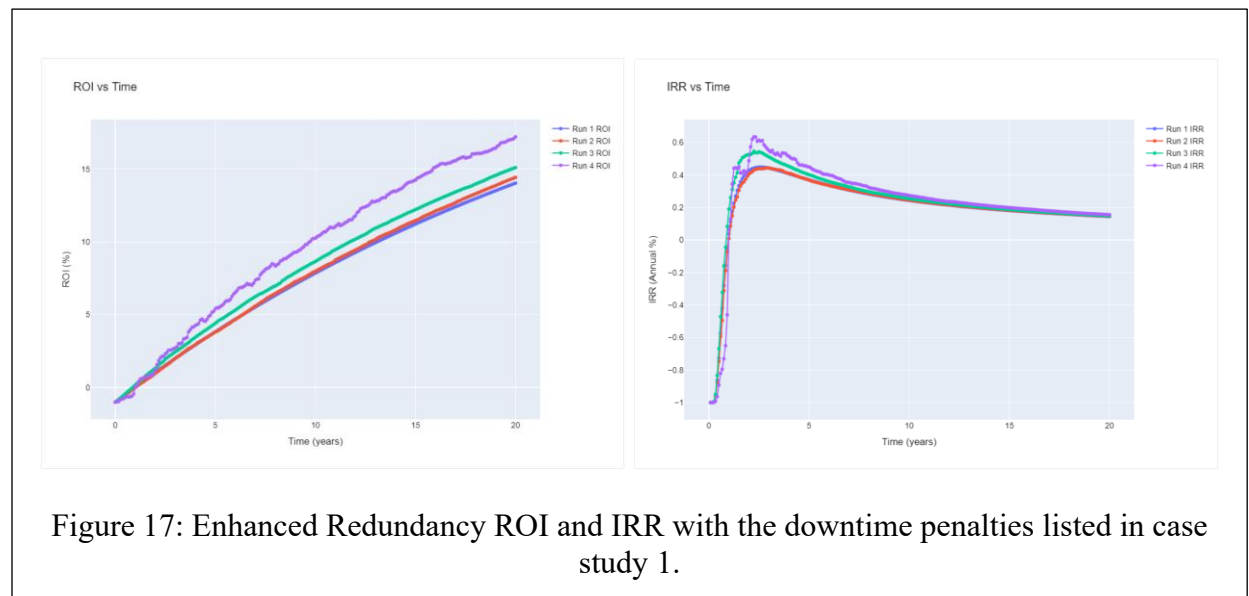
This study compares the baseline configuration against a modified version with N+1 redundancy added to a critical, single-point-of-failure component. The analysis will directly address the central question of whether the significant increase in initial capital cost for the redundant chiller is financially justified by the reduction in predicted system downtime and the associated SLA penalties over the data center's life cycle.

4.4.1 Case Study 3 Data

The baseline system is the standard Georgia Tech data center configuration, which features a single primary cooling loop with no redundancy for its major components. A failure of the single 150 kW chiller results in an immediate and complete failure of the data center's cooling system. The enhanced system is identical to the baseline, with one critical modification: the addition of a second, identical 150 kW chiller. This creates an N+1 redundancy configuration.

4.4.2 Case Study 3 Results

In Figure 16, Run 1 establishes the baseline scenario with no applicable downtime penalties. The remaining runs demonstrate the impact of the three different SLA structures: Run 2 applies the Total Cumulative Downtime penalty, Run 3 represents the Continuous Downtime penalty, and Run 4 depicts the Number of Interrupts penalty.



The analysis of adding N+1 redundancy to the primary chiller presents a clear trade-off between upfront capital investment and long-term operational resilience, with a breakeven point at 0.963 years.

While the inclusion of a redundant chiller increases the initial CAPEX, it significantly enhances overall system reliability. The breakeven point occurs earlier due to the significant reduction in downtime penalties when compared against the results in Case Study 1. In all modeled scenarios, the redundant configuration yields a higher overall ROI, as the financial savings from avoided SLA penalties far outweigh the cost of the additional hardware. In the baseline non-redundant configuration, a failure of the single chiller results in a complete loss of cooling and an immediate data center outage. The associated SLA penalties for such an event can be substantial, as can the reputational damage and potential loss of customers. With N+1 redundancy, it allows for uninterrupted operation while the primary unit is repaired.

The breakeven at 0.963 years indicates that the financial benefits of avoiding downtime, in the form of averted SLA penalties, begin to outweigh the initial cost of the redundant chiller in less than a year. This makes a compelling case for the investment in redundancy, particularly for data centers that have stringent uptime requirements and financial penalties for outages. This case study demonstrates that while redundancy comes at a cost, the financial protection it affords against the high cost of downtime can make it a profitable investment over the life of the data center.

Figure 18 presents a sensitivity analysis of breakeven time as a function of increasing SLA penalty costs, specifically utilizing the continuous downtime penalty model. In contrast to the previous sensitivity analysis (Figure 14), this graph demonstrates a clear inverse relationship: as the magnitude of the penalty increases, the time required to reach the breakeven point decreases. This trend is driven by the disparity in reliability between the non-redundant reference system and the N+1 redundant analysis system.

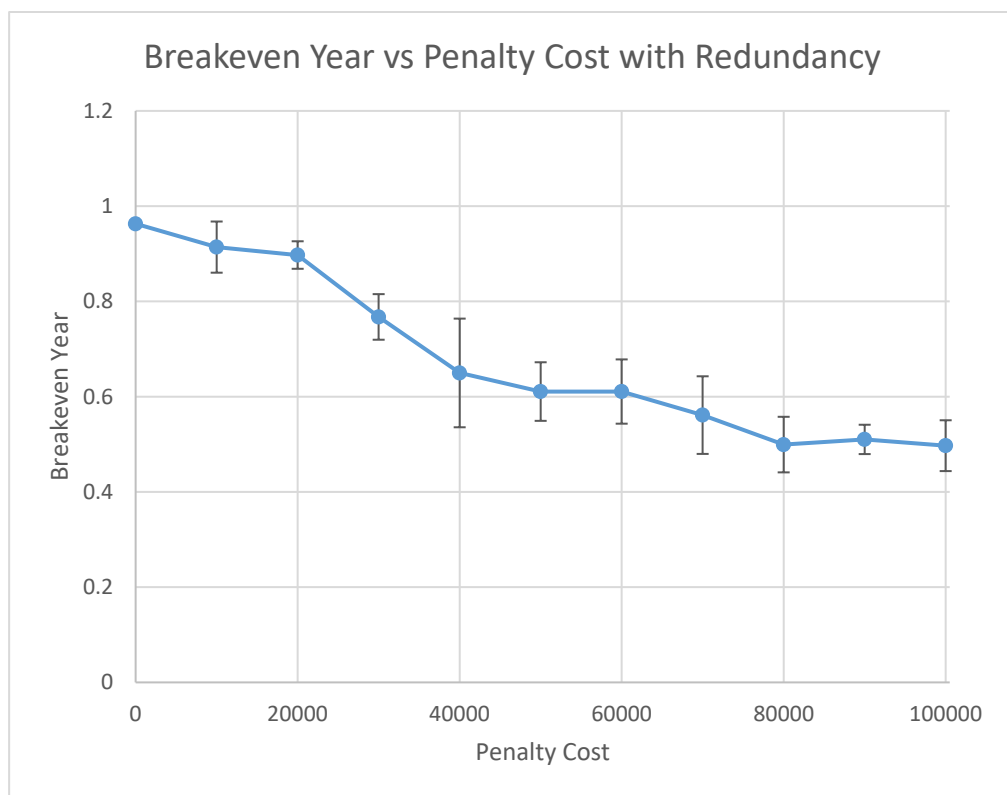


Figure 18: Sensitivity analysis for breakeven year vs. penalty cost with redundancy applied to the analysis system

Chapter 5: Summary, Contribution, and Future Work

This thesis developed a comprehensive cost model for the evaluation and comparison of data center cooling systems. Recognizing the escalating energy consumption and operational costs in the data center industry, driven by the demands of artificial intelligence and cloud computing, this research provides a tool for optimizing a critical component of data center infrastructure. A key innovation of this work is the incorporation of Service-Level Agreement (SLA) penalties into the total cost of ownership analysis. Existing cost models have largely overlooked the significant financial risks associated with downtime as defined in SLAs. By creating a model that combines capital and operational expenditures with a stochastic discrete-event simulator for maintenance, failures, and repairs, this research quantifies the financial impact of downtime penalties alongside energy and capital costs.

The model calculates key financial metrics, including Return on Investment (ROI) and Internal Rate of Return (IRR), to provide clear economic justification for investments in advanced or redundant cooling technologies. A series of case studies based on a real-world data center configuration demonstrated the model's utility. These case studies illustrated the sensitivity of financial outcomes to different SLA penalty structures, the economic advantages of direct liquid cooling over traditional air cooling, and the financial justification for investing in enhanced redundancy to mitigate the risk of costly downtime. The findings underscore the importance of a holistic approach to data center cost modeling that accounts for the complex interplay between capital investment, operational efficiency, and the financial consequences of service interruptions.

Contributions

- **Development and Integration of SLA Penalties into a Data Center Cost Model:** This thesis developed several general SLA penalty models and incorporated them into a novel cost modeling framework that quantitatively assesses the financial penalties associated with Service-Level Agreements. Resulting in a more accurate and realistic assessment of the total cost of ownership for different data center cooling systems.
- **Development of a Stochastic Maintenance and Failure Simulator:** A discrete-event simulator was created to model the probabilistic nature of component failures and repairs. This allows for the generation of realistic downtime scenarios to support the calculation of expected SLA penalties. The simulator's ability to account for various redundancy configurations enables a direct comparison of the financial benefits of increased system resilience.
- **Development of a Comprehensive Framework for Comparative Economic Analysis:** By calculating ROI and IRR, it offers a clear and actionable basis for decision-making regarding technology adoption and investment in redundancy, moving beyond current TCO comparisons.
- **The First Comparative Analysis of SLA Penalty Structures:** The SLA penalty structures were applied to an actual data center to assess their relative impacts.

Future Work

- **Integration with Real-Time Monitoring and Predictive Maintenance:** The discrete-event simulator could be integrated with real-time data from data center monitoring systems. This would enable the model to be used for predictive maintenance, allowing

operators to anticipate potential failures and take pre-emptive action to avoid downtime.

Machine learning algorithms could be employed to identify patterns in operational data that are indicative of impending failures.

- **Inclusion of Environmental and Sustainability Metrics:** The current model focuses primarily on financial costs. Future iterations could incorporate environmental metrics, such as carbon footprint and water usage, to provide a more comprehensive assessment of the overall impact of different cooling solutions.

Appendix – Parametric Equipment Cost Models

The CAPEX model for the cooling equipment requires that a cost be computed for all equipment in the cooling system bill of materials. The cooling equipment list is obtained from EnergyPlus. EnergyPlus also “auto-sizes” the equipment. From the equipment type and its size, the costs of determined from the following parametric models:

Chiller: The size of the chiller capacity is measured in tons. Depending on the chiller capacity the \$/ton value will change [31], The design size reference capacity in watts is taken from the EnergyPlus sizing file. This sizing capacity in W is then converted into tons. Depending on the capacity of the chiller the \$/ton value will change as shown in the table below.

tons	\$/ton
$x < 50$	1500
$50 < x < 150$	700
$x > 150$	450

Table 1: x is the capacity of the chiller in tons. The \$/ton changes depending on the capacity of the chiller.

$$\text{Chiller Cost} = (\text{tons})(\$/\text{ton}) \quad (1)$$

Cooling Tower: The cooling tower like the chiller has the cost adjusted based on the capacity measured in tons [32]. The nominal capacity of the cooling tower in watts is taken from the Energy plus sizing file. Using the capacity of the cooling tower in watts the capacity can be converted into tons,

tons	\$/ton
$x < 10$	1500
$10 < x < 50$	1250
$x > 50$	700

Table 2: x is the capacity of the cooling tower in tons. The \$/ton changes depending on the capacity of the cooling tower.

$$\text{Cooling Tower Cost} = (\text{tons})(\$/\text{ton}) \quad (2)$$

Pumps: The cost of the pump is determined by its power consumption [33]. The pump models rely on calculations from the early 2000s, and to adjust for inflation, a Chemical Engineering Plant Cost Index (CEPCI) is applied to the pump costs. The power consumption of the pumps in watts will come from EnergyPlus,

$$\text{Pump Cost} = \left(9500 \frac{\text{Power Consumption}}{23} \right)^{0.79} \text{CEPCI} \quad (3)$$

Air Economizer: The air economizer model is based on the size of the data center in square feet [34]. The size of the data center in square meters will be the conditioned area of the building from Energy plus. The size of the data center is converted from square meters to square feet. It was later learned that the conditioned area of a building is different from the electrically active floor area. The electrically active floor area will need to be determined from the Energy plus file

$$\text{Air Economizer Cost} = 132.4(\text{size of data center in sq ft}) \quad (4)$$

Computer Room Air Handling Unit (CRAH): The current CRAH model uses the expected capacity of the data center to determine the cost [35]. The data center capacity in Kilowatts is a value that is currently required to be inputted into the program. It is expected that this value will be an input in either the energy or thermal model,

$$CRAH\ Cost = 240(data\ center\ capacity\ in\ kW) \quad (5)$$

Fluid: The current fluid model assumes that there is no loss of fluid throughout the system. The system volume will be taken from Energy plus and will be in cubic meters. The system volume is converted into cubic feet. The cost of the fluid used is currently set to the cost of water, however this value can be adjusted in the default parameter settings as shown in figure 3.

$$Fluid\ Cost = \frac{264.172(System\ volume)}{(12)(10000)} (Cost\ of\ Fluid) \quad (6)$$

In addition to the equipment, we also cost the cooling system infrastructure, i.e., piping and ducting. The size of the piping and ducting system is assumed to be 4-inch diameter for pipes and 4 square feet for the ducting. These values can be adjusted in the default parameters settings and a new piping and ducting.

Piping [36]: Using the diameter an area of the pipe can be calculated that can be used in equation 10. The volume of the piping in cubic meters is taken from the EnergyPlus sizing file. The volume of the piping is then converted into cubic feet,

$$Piping\ Cost = (\frac{80}{pipe\ area})(volume)CEPCI \quad (7)$$

Ducting [36]: Since the ducting size is inputted as an area, no changes need to be made. The volume of ducting in cubic meters is taken from the EnergyPlus sizing file. The volume of the piping is then converted into cubic feet,

$$Ducting\ Cost = (\frac{700}{duct\ area})(volume)CEPCI \quad (8)$$

References

- [1] Srivathsan, B., Sorel, M., & Sachdeva, P. (2024, October 29). *AI power: Expanding Data Center capacity to meet growing demand*. McKinsey & Company.
<https://www.mckinsey.com/industries/technology-media-and-telecommunications/our-insights/ai-power-expanding-data-center-capacity-to-meet-growing-demand>
- [2] Masanet, E., Shehabi, A., Lei, N., Smith, S., & Koomey, J. (2020). *Recalibrating global data center energy use estimates*. *Science*, 367(6481), 984-986.
https://datacenters.lbl.gov/sites/default/files/Masanet_et_al_Science_2020.full_.pdf
- [3] Dunaway, J. (2023, May 24). *AI and the data center: Driving greater power density*. Coresite. <https://www.coresite.com/blog/ai-and-the-data-center-driving-greater-power-density>
- [4] CoreSite. (2021, April 27). *Breaking down data center tiers & classifications*. CoreSite. <https://www.coresite.com/blog/breaking-down-data-center-tiers-classifications>
- [5] Avelar, V., et al. (2012). *PUE: A Comprehensive Examination of the Metric*. The Green Grid. White Paper #49.
- [6] Koomey, J. (2011) *Growth in Data Center Electricity Use 2005 to 2010*. Analytics Press.
- [7] Cushman & Wakefield. (2025). *2025 Global Data Center Outlook*.
<https://www.cushmanwakefield.com/en/insights/americas-data-center-update>
- [8] Ebrahimi, K., Jones, G. F., & Fleischer, A. S. (2015). A review of data center cooling technologies, operating conditions and the corresponding modernization opportunities. *Energy Conversion and Management*, 106, 624-644.

- [9] ASHRAE. (2015). *Thermal Guidelines for Data Processing Environments* (4th ed.). ASHRAE TC 9.9.
- [10] Cho, J., & Kim, Y. *A review of the state-of-the-art in data center cooling*. Energy Conversion and Management. vol. 132, pp. 123-142, 2017.
- [11] Koomey, Jonathan & Turner, Pitt & Stanley, John & Taylor, Bruce. (2007). *A Simple Model for Determining True Total Cost of Ownership for Data Centers*.
- [12] Hamilton, J. (2008, November 28). *Cost of Power in Large-Scale Data Centers*. Perspectives. <https://perspectives.mvdirona.com/2008/11/cost-of-power-in-large-scale-data-centers/>
- [13] Cho, K., Chang, H., Jung, Y., & Yoon, Y. (2017). *Economic analysis of data center cooling strategies*. Applied Thermal Engineering, 125, 229-239.
- [14] Keskin, B., & Soykan, G. (2022). *Optimal cost management of the CCHP based data center with district heating and district cooling integration in the presence of different energy tariffs*. Energy, 239, 122137.
- [15] Liu, Z., Zhang, J., Wu, Y., & Li, Z. (2020). *Configuration optimization model for data-center-park-integrated energy systems under economic, reliability, and environmental considerations*. Applied Energy, 275, 115383.
- [16] Gözcü, S. S., & Erden, H. S. (2019). *Energy and economic assessment of major free cooling retrofits for data centers in Turkey*. Energy and Buildings, 205, 119330.
- [17] Schneider Electric. (n.d.). *Data Center Build vs. Colocation TCO Calculator*. Schneider Electric. <https://www.se.com/ww/en/work/solutions/system/s1/data-center-and-network-systems/trade-off-tools/data-center-build-vs-colocation-tco-calculator/>

- [18] Schneider Electric. (n.d.). *Data Center Capital Cost Calculator*. Schneider Electric.
<https://www.se.com/ww/en/work/solutions/system/s1/data-center-and-network->
- [19] Gordian. 2025. *RSMeans Data*. RSMeans. <https://www.rsmeans.com/>.
- [20] Data Center Knowledge. (n.d.). *The Biggest Threats to Data Center Uptime – and How to Overcome Them*. Data Center Knowledge.
<https://www.datacenterknowledge.com/uptime/the-biggest-threats-to-data-center-uptime-and-how-to-overcome-them>
- [21] Ibrahim, Abdallah. (2015). *Service Level Agreement Assurance in Cloud Computing Data Centers*. DOI:10.13140/RG.2.1.4597.4486.
- [22] Alsaaidah, Y. A., Muhammed, A., Ala'anzy, M. A., Othman, M., & Abdullah, A. (2025). *Empowering cloud providers: optimised locust-inspired algorithm for SLA violation mitigation in green cloud computing*. Computing, 107.
<https://doi.org/10.1007/s00607-025-01527-7>
- [23] 365 Data Centers. (n.d.). *Service Level Agreement (SLA)*.
<https://365datacenters.com/sla/>
- [24] NTT Global Data Centers. (n.d.). *Service Level Agreement*.
<https://services.global.ntt/en-US/pdf-viewer/deea976d-ddaa-40cf-9944-05a127cdf398>
- [25] Hu, B., Patel, D., & Phan, L. (2021). *Guidance for CFD modeling of data centers* (Final Report No. RP-1675). American Society of Heating, Refrigerating and Air-Conditioning Engineers (ASHRAE).
- [26] U.S. Department of Energy. (n.d.). *EnergyPlus*. <https://energyplus.net/>

- [27] Big Ladder Software. (n.d.). *What is EnergyPlus?* <https://bigladdersoftware.com/epx/docs/8-4/getting-started/what-is-energyplus.html>
- [28] AWS. (n.d.). *Availability with redundancy - availability and beyond: Understanding and improving the resilience of distributed systems on AWS.* (n.d.).
<https://docs.aws.amazon.com/whitepapers/latest/availability-and-beyond-improving-resilience/availability-with-redundancy.html>
- [29] Managed Server. (n.d.). *Penalty Clauses and Damages in Hosting Service SLAs.*
<https://www.managedserver.eu/criminal-clauses-and-compensation-for-damages-in-hosting-services-slabs/>
- [30] Databank. (n.d.). *The Critical Role Of Service Level Agreements (SLAs) In Ensuring Data Center Reliability.* <https://www.databank.com/resources/blogs/the-critical-role-of-service-level-agreements-slas-in-ensuring-data-center-reliability/>
- [31] Fisher, J. 2019. *How to Manually Calculate Chiller Capacity for Your Process.*
<https://www.conairgroup.com/resources/resource/chiller-sizing-part-2-the-old-school-way/>.
- [32] Cooling Tower Products. 2024. *Cooling Tower Costs 2024.*
<https://www.coolingtowerproducts.com/cooling-tower-cost-2024/>
- [33] Woods, D. R., (2007). *Rules of Thumb in Engineering Practice* (pp. 383-384). Wiley-VCH Verlag GmbH & Co. KGaA. DOI:10.1002/9783527611119
- [34] Energy Star. (n.d.). *Use an Air-side Economizer.*
https://www.energystar.gov/products/data_center_equipment/16-more-ways-cut-energy-waste-data-center/use-air-side-economizer

- [35] Shehata, H., 2014. *Data Center Cooling: CRAC/CRAH redundancy, capacity, and selection metrics*. <https://journal.uptimeinstitute.com/data-center-cooling-redundancy-capacity-selection-metrics/>
- [36] Woods, D. R., (2007). *Rules of Thumb in Engineering Practice* (pp. 387). Wiley-VCH Verlag GmbH & Co. KGaA. DOI:10.1002/9783527611119