

Comparing UPS System Design Configurations

White Paper 75

Revision 4

by Kevin McCarthy, EDG2 Inc.
Victor Avelar, Schneider Electric

> Executive summary

There are five principle UPS system design configurations that distribute power from the utility source of a building to the critical loads of a data center. The selection of the appropriate configuration or combination thereof for a particular application is determined by the availability needs, risk tolerance, types of loads in the data center, budgets, and existing infrastructure. This paper will focus on these five configurations; the advantages and disadvantages of each are discussed. The impact on availability is addressed for each configuration and guidelines are provided for choosing the appropriate design.

Introduction

Although the public power distribution system is fairly reliable in most developed countries, studies have shown that even the best utility systems are inadequate to meet the needs of mission-critical applications. Most organizations, when faced with the likelihood of downtime, and data processing errors caused by utility power, choose to implement an uninterruptible power supply (UPS) system between the public power distribution system and their mission-critical loads. The UPS system design configuration chosen for the application directly impacts the availability of the critical equipment it supports. There are many variables that affect a system's availability, including human error, reliability of components, maintenance schedules, and recovery time. The impact that each of these variables has on the overall system's availability is determined to a large degree, by the configuration chosen.

Over time, many design engineers have tried to create the perfect UPS solution for supporting critical loads, and these designs often have names that do not necessarily indicate where they fall in the spectrum of availability. "Parallel redundant", "isolated redundant", "distributed redundant", "hot tie", "hot synch", "multiple parallel bus", "system plus system", "catcher systems, and "isolated parallel" are names that have been given to different UPS configurations by the engineers who designed them or by the manufacturers who created them. The problems with these terms are that they can mean different things to different people, and can be interpreted in different ways. Although UPS configurations found in the market today are many and varied, there are five that are most commonly applied. These five include: (1) capacity, (2) isolated redundant, (3) parallel redundant, (4) distributed redundant and (5) system plus system.

This paper explains these UPS system configurations and discusses the benefits and limitations of each. A system configuration should be chosen to reflect the criticality of the load. Considering the impact of downtime and the corporate risk tolerance will help in choosing the appropriate system configuration.

Guidelines are provided for selecting the appropriate configuration for a given application.

Scale, availability, and cost

Availability

The driving force behind the ever-evolving possibilities for UPS configurations is the ever-increasing demand for availability by data processing managers. "Availability" is the estimated percentage of time that electrical power will be online and functioning properly to support the critical load. An analysis in the **Appendix** quantifies the availability differences between the configurations presented in this paper. As with any model, assumptions must be made to simplify the analysis, therefore, the availability values presented will be higher than what is expected in an actual installation. Furthermore, the availability numbers are more a comparison tool than they are a predictor of any given system's performance. For the purposes of comparing the five common design configurations, a simple scale is provided in **Table 1** illustrating their availability ranking based on the results found in the appendix. After reviewing the explanations of the different configurations, this order should become evident.

Criticality / redundancy levels

All UPS systems (and electrical distribution equipment) require regular intervals of maintenance. The availability of a system configuration is dependent on its level of immunity to equipment failure, and the inherent ability to perform normal maintenance, and routine testing while maintaining the critical load. The Uptime Institute discusses this topic further in a document titled "Industry Standard Tier Classifications Define Site Infrastructure Perfor-

mance”¹. In addition to the Uptime Institute, TIA-942 also provides information on tiers.² The tiers described in the Uptime Institute document encompass the 5 UPS architectures mentioned in this paper and are also depicted in **Table 1**.

The following terms are sometimes used in describing the various tiers and drive both distributed redundant as well as system plus system configurations:

Concurrent maintenance – The ability to completely shut down any particular electrical component, or subset of components, for maintenance or routine testing without requiring that the load be transferred to the utility source.

Single point of failure – An element of the electrical distribution system that at some point will cause downtime, if a means to bypass it is not developed in the system. An N configuration system is essentially comprised of a series of single points of failure. Eliminating these from a design is a key component of redundancy.

Hardening – Designing a system, and a building, that is immune to the ravages of nature, and is immune to the types of cascading failures that can occur in electrical systems. The ability to isolate and contain a failure; for example, the two UPS systems would not reside in the same room, and the batteries would not be in the same room with the UPS modules. Circuit breaker coordination becomes a critical component of these designs. Proper circuit breaker coordination can prevent short circuits from affecting large portions of the building.

Hardening a building can also mean making it more immune to events such as hurricanes, tornadoes, and floods, as might be necessary depending on where the building is. For example designing the buildings away from 100 year flood plains, avoiding flight paths overhead, specifying thick walls and no windows all help to create this immunity.

Cost

As the configuration goes higher on the scale of availability, the cost also increases. **Table 1** provides approximate ranges of costs for each design. These costs represent the cost to build out a new data center and include not only the UPS architecture cost, but also the complete data center physical infrastructure (DCPI) of the data center. This consists of generator(s), switchgear, cooling systems, fire suppression, raised floor, racks, lighting, physical space, and the commissioning of the entire system. These are the up-front costs only and do not include operating costs such as maintenance contracts. These costs assume an average of 30 square feet (2.79 square meters) per rack, and are based on a range of power densities from 2.3 kW / rack to 3.8 kW / rack. The cost per rack will decrease as the size of the building increases, providing a larger footprint over which to spread costs and greater buying power from vendors.

¹ <http://uptimeinstitute.com/>

² TIA-942, *Telecommunications Infrastructure Standard for Data Centers*, April 2005

Table 1

Scale of availability and cost for UPS configurations

Configurations	Scale of availability	Tier class	Data center scale of cost (US\$)
Capacity (N)	1 = Lowest	Tier I	\$13,500 - \$18,000 / rack
Isolated redundant	2	Tier II	\$18,000 - \$24,000 / rack
Parallel redundant (N+1)	3		
Distributed redundant	4	Tier III	\$24,000 - \$30,000 / rack
System plus system (2N, 2N+1)	5 = Highest	Tier IV	\$ 36,000 - \$42,000 / rack

What is “N”?

UPS design configurations are often described by nomenclatures using the letter “N” in a calculation stream. For instance, a parallel redundant system may also be called an N+1 design, or a system plus system design may be referred to as 2N. “N” can simply be defined as the “need” of the critical load. In other words, it is the power capacity required to feed the protected equipment. IT equipment such as RAID (redundant array of independent disks) systems can be used to illustrate the use of “N”. For example, if 4 disks are needed for storage capacity and the RAID system contained 4 disks, this is an “N” design. On the other hand, if there are 5 disks and only 4 are needed for storage capacity that is an N+1 design.

Historically, the critical load power requirement has had to be projected well into the future in order to allow a UPS system to support loads for 10 or 15 years. Projecting this load has proven to be a difficult task, and justifiably so. In the 1990’s the concept of “Watts / Square Area” was developed in order to provide a framework for the discussion and the ability to compare one facility to the next. Misunderstanding exists with this measure of power simply by the fact that people can’t agree on what the square area is. More recently, with the trend of technology compaction, the concept of “Watts / Rack” has been used to drive the system capacity. This has proven to be more reliable as the quantity of racks in a space is very easy to count. Regardless of how the load “N” is chosen, it is essential that it be chosen from the onset to allow the design process to begin on the right track.

Scalable, modular UPS system designs now exist to allow the UPS capacity to grow as the IT “need” grows. For more information on this topic, refer to White Paper 37, [Avoiding Costs from Oversizing Data Center and Network Room Infrastructure](#).

Capacity or “N” system

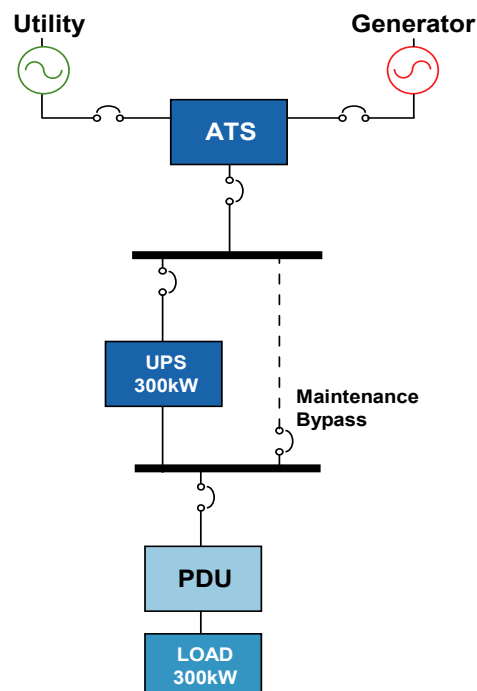
An N system, simply stated, is a system comprised of a single UPS module, or a paralleled set of modules whose capacity is matched to the critical load projection. This type of system is by far the most common of the configurations in the UPS industry. The small UPS under an office desk is an N configuration. Likewise, the 5,000 square foot (465 square meters) computer room with a projected design capacity of 400 kW is an N configuration whether it has a single 400 kW UPS, or two 200 kW UPS paralleled onto a common bus. An N configuration can be looked at as the minimum requirement to provide protection for the critical load.

Although both examples above are considered N configurations, the UPS module designs are different. Unlike the small UPS, systems above single-phase capacities (roughly 20 kW) have internal static bypass switches that allow the load to be transferred safely to the utility source if the UPS module experiences internal problems. The points at which a UPS transfers to static bypass are carefully selected by the manufacturer to provide the utmost protection for the critical load, while at the same time safeguarding the module itself against situations that could damage it. The following example illustrates one of these protective measures: It is common in three-phase UPS applications for the modules to have overload ratings. One of these ratings may state that the “module will carry 125% of its rated load for 10 minutes” or alternatively until the component reaches a given temperature. Once a 125% overload is detected, a module will start a timing routine where an internal clock begins a 10-minute countdown. When the timer expires, if the load has not returned to normal levels, the module will transfer the load safely to static bypass. There are many scenarios in which the bypass will be activated, and they are stated clearly in the specifications of a particular UPS module.

A way to augment an N configuration design is to provide the system with “maintenance” or “external” bypass capability. An external bypass would allow the entire UPS system (modules and static bypass) to be safely shut down for maintenance if and when that situation arises. The maintenance bypass would emanate from the same panel that feeds the UPS, and would connect directly to the UPS output panel. This, of course, is a normally open circuit that can only be closed when the UPS module is in static bypass. Steps need to be taken in the design to prevent the closing of the maintenance bypass circuit when the UPS is not in static bypass. When properly implemented into a system, the maintenance bypass is an important component in the system, allowing a UPS module to be worked on safely without requiring the shutdown of the load.

Most “N” system configurations, especially under 100 kW, are placed in buildings with no particular concern for the configuration of the overall electrical systems in the building. In general, buildings’ electrical systems are designed with an “N” configuration, so an “N” UPS configuration requires nothing more than that to feed it. A common single module UPS system configuration is shown in **Figure 1**.

Figure 1
Single module “capacity”
UPS configuration



Disadvantages of an “N” design

- Limited availability in the event of a UPS module break down, as the load will be transferred to bypass operation, exposing it to unprotected power
- During maintenance of the UPS, batteries or down-stream equipment, load is exposed to unprotected power (usually takes place at least once a year with a typical duration of 2 - 4 hours)
- Lack of redundancy limits the load's protection against UPS failures
- Many single points of failure, which means the system is only as reliable as its weakest point

Isolated redundant

An isolated redundant configuration is sometimes referred to as an “N+1” system, however, it is considerably different from a parallel redundant configuration which is also referred to as N+1. The isolated redundant design concept does not require a paralleling bus, nor does it require that the modules have to be the same capacity, or even from the same manufacturer. In this configuration, there is a main or “primary” UPS module that normally feeds the load. The “isolation” or “secondary” UPS feeds the static bypass of the main UPS module(s). This configuration requires that the primary UPS module have a separate input for the static bypass circuit. This is a way to achieve a level of redundancy for a previously non-redundant configuration without completely replacing the existing UPS. **Figure 2** illustrates an isolated redundant UPS configuration.

In a normal operating scenario the primary UPS module will be carrying the full critical load, and the isolation module will be completely unloaded. Upon any event where the primary module(s) load is transferred to static bypass, the isolation module would accept the full load of the primary module instantaneously. The isolation module has to be chosen carefully to ensure that it is capable of assuming the load this rapidly. If it is not, it may, itself, transfer to static bypass and thus defeat the additional protection provided by this configuration.

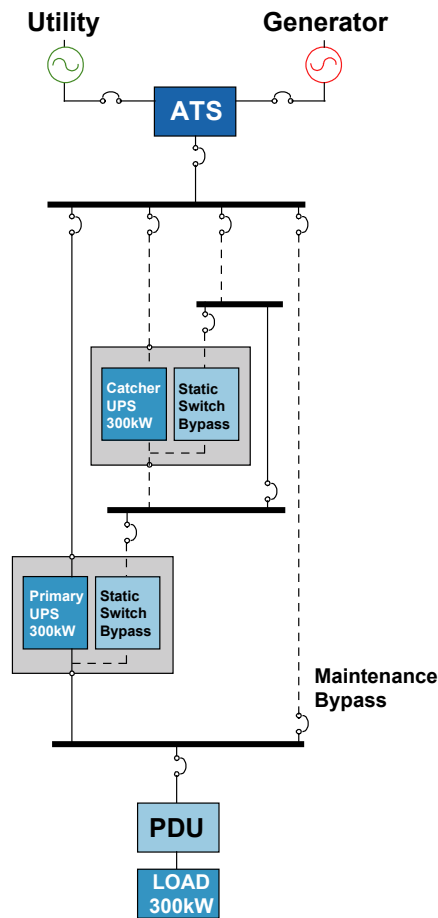
Service can be performed on either module by transferring the load to the other module. A maintenance bypass is still an important design feature, as the output single point of failure still exists. The entire system needs to be shutdown for 2 – 4 hours per year for system-level preventive maintenance. Reliability gains from this configuration are often offset by the complexity of the switchgear and associated controls. MTechnology Inc.³, consultants specializing in high reliability electric power systems, performed a comparative reliability analysis. Using the techniques of Probabilistic Risk Assessment (PRA), MTech developed quantitative models for both an isolated redundant UPS system and a non-redundant (capacity) system. The most basic fault tree analysis, which ignored contributions to failure arising from human error, component aging, and environmental effects, demonstrate that the isolated redundant system does not materially affect the probability of failure (unreliability.) Both systems had an unreliability of 1.8% per year of operation. The isolated redundant model resulted in 30 failure modes (minimal cut sets) vs. 7 for the capacity system. While the probability of the additional 23 failure modes is generally small, the analysis illustrates that adding complexity and additional components to the system invariably increases the number of potential failure modes. Mtech contends that when human errors and the effects of aging are considered, the case against the isolated redundant system is even stronger. The operation of the isolated redundant system is much more complex than in a non-isolated system, and the probability of human error very much higher. The benefits of the preventative maintenance procedures that are enabled by the isolated redundant designs do not withstand careful scrutiny. The primary beneficiaries of the isolated redundant UPS design are those who sell the original equipment and those who profit from servicing the additional

³ MTechnology, Inc; 2 Central Street, Saxonville, MA 01701; phone 508-788-6260; fax 508-788-6233

UPS modules. The customer's equipment does not benefit from higher reliability electric power.

Figure 2

Isolated redundant UPS configuration



Advantages of an isolated redundant design

- Flexible product choice, products can be mixed with any make or model
- Provides UPS fault tolerance
- No synchronizing needed
- Relatively cost effective for a two-module system

Disadvantages of an isolated redundant design

- Reliance on the proper operation of the primary module's static bypass to receive power from the reserve module
- Requires that both UPS modules' static bypass must operate properly to supply currents in excess of the inverter's capability
- The secondary UPS module has to be able to handle a sudden load step when the primary module transfers to bypass. (This UPS has generally been running with 0% load for a long period of time. Not all UPS modules can perform this task making the selection of the bypass module a critical one).
- Switchgear becomes complex and costly when catcher UPS supports multiple primary UPS

- Higher operating cost due to a 0% load on the secondary UPS, which draws power to keep it running
- A two module system (one primary, one secondary) requires at least one additional circuit breaker to permit choosing between the utility and the other UPS as the bypass source. This is more complex than a system with a common load bus and further increases the risk of human error.
- Two or more primary modules need a special circuit to enable selection of the reserve module or the utility as the bypass source (static transfer switch)
- Single load bus per system, a single point of failure

Parallel redundant or “N+1”

Parallel redundant configurations allow for the failure of a single UPS module without requiring that the critical load be transferred to the utility source. The intent of any UPS is to protect the critical load from the variations and outages in the utility source. As the criticality of data increases, and the tolerance for risk diminishes the idea of going to static bypass, and maintenance bypass is seen as something that needs to be further minimized. N+1 system designs still must have the static bypass capability, and most of them have a maintenance bypass as they still provide critical capabilities.

A parallel redundant configuration consists of paralleling multiple, same size UPS modules onto a common output bus. The system is N+1 redundant if the “spare” amount of power is at least equal to the capacity of one system module; the system would be N+2 redundant if the spare power is equal to two system modules; and so on. Parallel redundant systems require UPS modules identical in capacity and model. The output of the modules is synchronized using an external paralleling board or in some cases this function is embedded within the UPS module itself. In some cases the paralleling function also controls the current output between the modules.

The UPS modules communicate with each other to create an output voltage that is completely synchronized. The parallel bus can have monitoring capability to display the load on the system and the system voltage and current characteristics at a system level. The parallel bus also needs to be able to display how many modules are on the parallel bus, and how many modules are needed in order to maintain redundancy in the system. There are logical maximums for the number of UPS modules that can be paralleled onto a common bus, and this limit is different for different UPS manufacturers. The UPS modules in a parallel redundant design share the critical load evenly in normal operating situations. When one of the modules is removed from the parallel bus for service (or if it were to remove itself due to an internal failure), the remaining UPS modules are required to immediately accept the load of the failed UPS module. This capability allows any one module to be removed from the bus and be repaired without requiring the critical load to be connected to straight utility.

The 5,000 square foot (465 square meters) computer room in our N configuration example would require two 400 kW UPS modules, or three 200 kW UPS modules paralleled onto a common output bus to become redundant. The parallel bus is sized for the non-redundant capacity of the system. So the system comprised of two 400 kW modules would have a parallel bus with a rated capacity of 400 kW.

In an N+1 system configuration there is an opportunity for the UPS capacity to grow as the load grows. Capacity triggers need to be set up so that when the percentage of the capacity in place reaches a certain level, (acknowledging that delivery times for some UPS modules can be many weeks or even months), a new redundant module should be ordered. The larger the UPS capacity, the more difficult a task this can become. Large UPS modules weigh thousands of pounds and require special rigging equipment in order to set them into place. There would typically be a spot reserved in the UPS room for this module. This type

of deployment needs to be well planned as placing a large UPS module into any room comes with some risk.

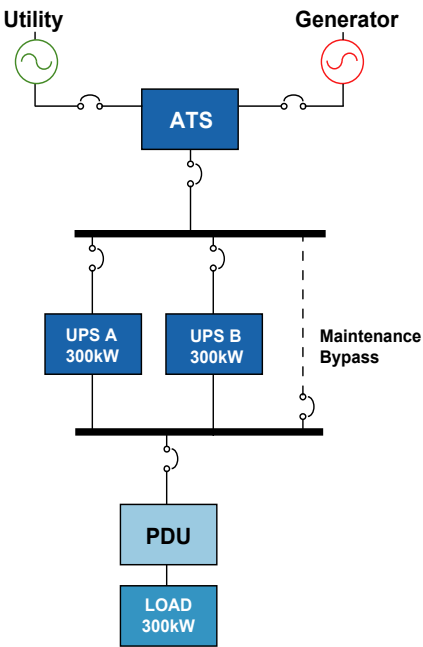
System efficiency can be an important factor to consider in the design of redundant UPS systems. Lightly -loaded UPS modules are typically less efficient than a module that is loaded closer to its capacity. **Table 2** shows the typical running load for a system using various UPS module sizes, all feeding a 240 kW load. As can be seen in the table, the module size chosen for a particular application can seriously affect the system efficiency. The efficiency of any particular UPS at low loads varies from manufacturer to manufacturer, and should be investigated during a design process.

Table 2
N + 1 configurations

UPS modules in parallel	Mission critical load	Total UPS system capacity	% each UPS module is loaded
2 x 240 kW	240 kW	480 kW	50%
3 x 120 kW	240 kW	360 kW	66%
4 x 80 kW	240 kW	320 kW	75%
5 x 60 kW	240 kW	300 kW	80%
2 x 240 kW	240 kW	480 kW	50%

Figure 3 depicts a typical two module parallel redundant configuration. This figure shows that even though these systems provide protection of a single UPS module failure, there still remains a single point of failure in the paralleling bus. As with the capacity design configuration, a maintenance bypass circuit is an important consideration in these designs in order to allow the UPS modules to be shut down for maintenance periodically.

Figure 3
Parallel redundant (N+1) UPS configuration



Advantages of an “N+1” design

- Higher level of availability than capacity configurations because of the extra capacity that can be utilized if one of the UPS modules breaks down
- Lower probability of failure compared to isolated redundant because there are less breakers and because modules are online all the time (no step loads)
- Expandable if the power requirement grows. It is possible to configure multiple units in the same installation
- The hardware arrangement is conceptually simple, and cost effective

Disadvantages of an “N+1” design

- Both modules must be of the same design, same manufacturer, same rating, same technology and configuration
- Still single points of failure upstream and downstream of the UPS system
- The load may be exposed to unprotected power during maintenance if the service extends beyond a single UPS module, or its batteries. If service is required in the parallel board or the parallel controls or down-stream equipment, the load will be exposed to un-protected power. Lower operating efficiencies because no single unit is being utilized 100%
- Single load bus per system, a single point of failure

Distributed redundant

Distributed redundant configurations, also known as tri-redundant, are commonly used in the large data center market today especially within financial organizations. This design was developed in the late 1990s in an effort by an engineering firm to provide the capabilities of complete redundancy without the cost associated with achieving it. The basis of this design uses three or more UPS modules with independent input and output feeders. The independent output buses are connected to the critical load via multiple PDUs. In some cases STS are also used in this architecture. From the utility service entrance to the UPS, a distributed redundant design and a system plus system design (discussed in the next section) are quite similar. Both provide for concurrent maintenance, and minimize single points of failure. The major difference is in the quantity of UPS modules that are required in order to provide redundant power paths to the critical load, and the organization of the distribution from the UPS to the critical load. As the load requirement, “N”, grows the savings in quantity of UPS modules also increases.

Figures 4, 5, and 6 illustrate a 300 kW load with three different distributed redundant design concepts. **Figure 4** uses three UPS modules in a distributed redundant design that could also be termed a “catcher system”. In this configuration, module 3 is connected to the secondary input on each STS, and would “catch” the load upon the failure of either primary UPS module. In this catcher system, module 3 is typically unloaded.

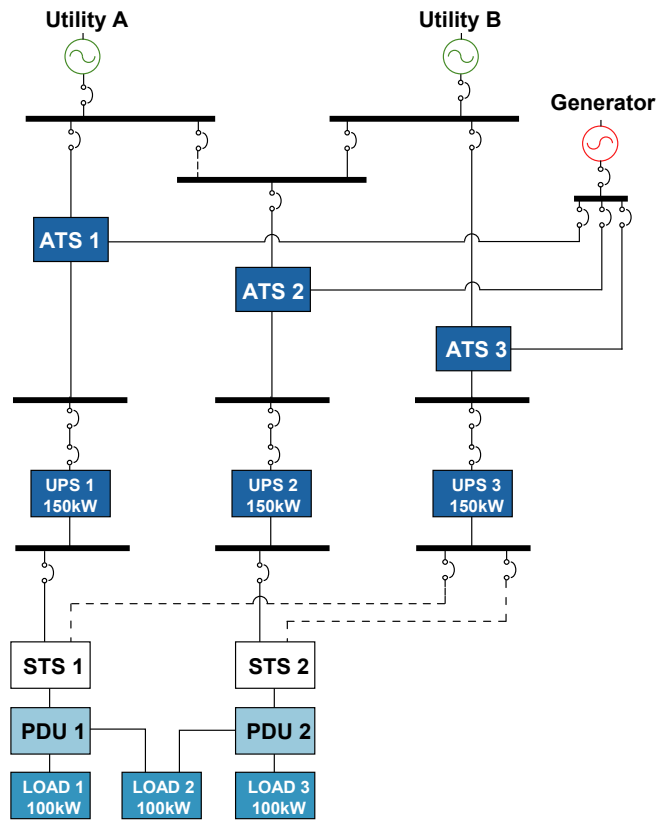


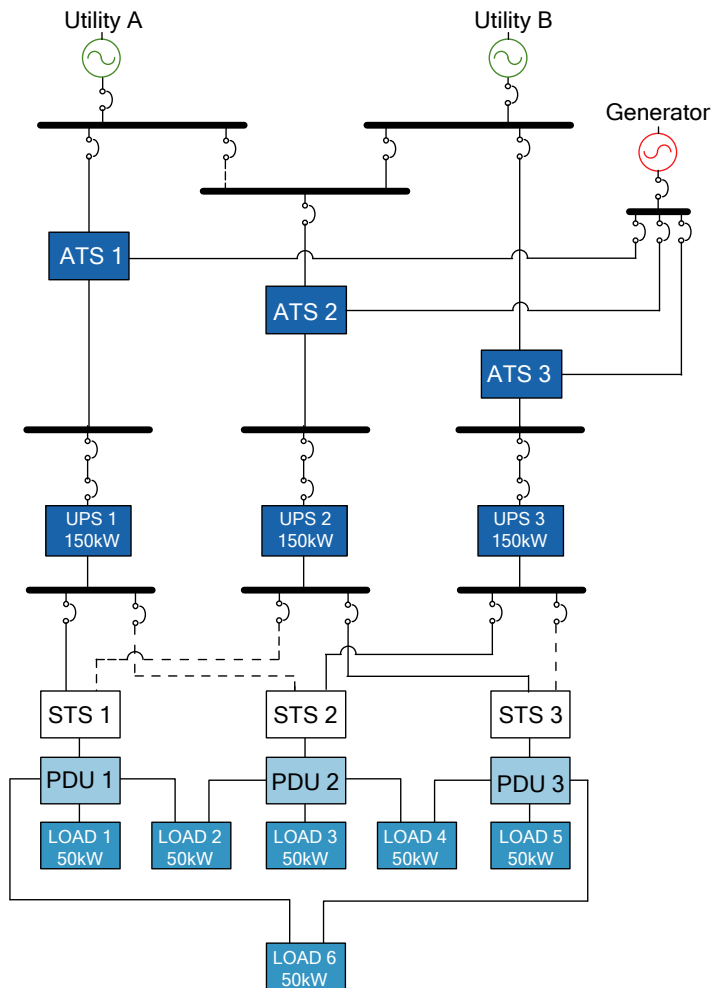
Figure 4

*Distributed redundant
"catcher" UPS configuration*

Figure 5 depicts a distributed redundant design with three STS and the load evenly distributed across the three modules in normal operation. The failure on any one module would force the STS to transfer the load to the UPS module feeding its alternate source.

Figure 5

Distributed redundant UPS configuration (with STS)

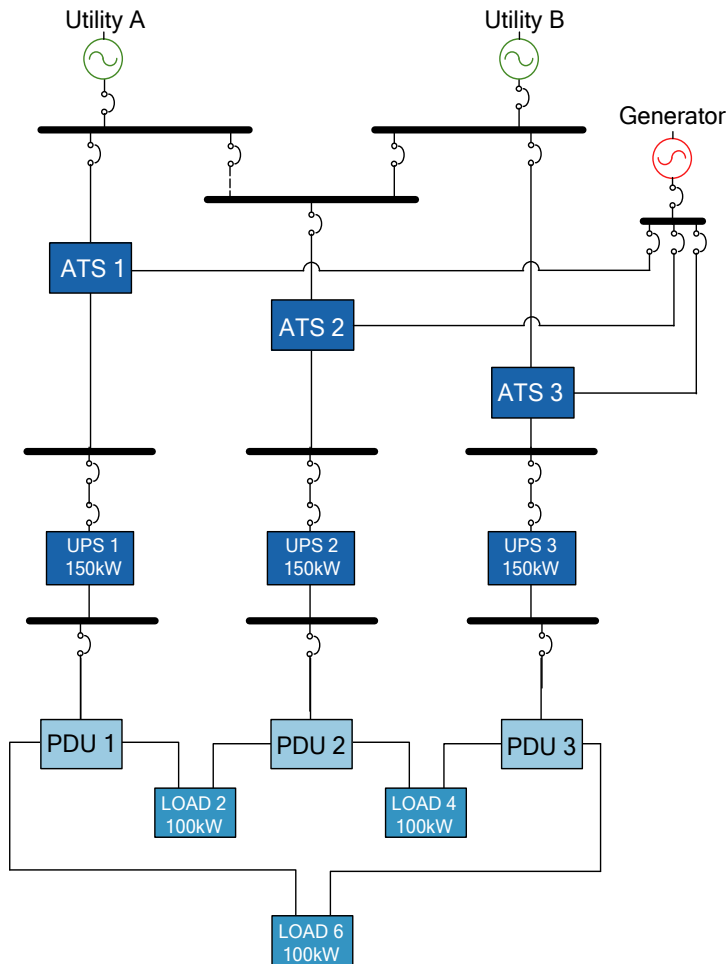


Evident in both of these one lines is the difference between distributing power to dual-corded loads and single-corded loads. The dual-corded loads can be fed from two STS units or no STS units, while the single-corded loads need to be fed from a single STS. As the quantity of single-corded loads in data centers today are becoming fewer and fewer it is becoming more practical, and less costly to apply multiple, small, point of use transfer switches close to the single-corded loads. In cases with 100% dual-corded loads this configuration could be designed without STS units as shown in **Figure 6**. This design is typically known as a tri-redundant and uses no static transfer switches.

Whereas **Figure 5** depicts a 1N STS design, large institutions with extreme electrical system reliability requirements use redundant STS as a means of isolating electrical maintenance activities from critical IT loads. For example, four “layered” events would need to occur to drop a dual-corded server during UPS maintenance. First the transfer to UPS static bypass would need to fail followed by the side “A” STS, then the side “B” UPS, and finally the side “B” STS. This “layering” approach provides small incremental reliability gains compared to the large expenses that come along with them – law of diminishing returns. Ultimately the best redundancy is geographical redundancy whereby redundant data centers are built in two distant locations. However, it is currently difficult for financial institutions to implement geo-redundancy since they must have secure and instant access to all their data.

Figure 6

Tri-redundant UPS configuration (no STS)



Overall, distributed redundant systems are usually chosen for large multi-megawatt installations where concurrent maintenance is a requirement and space is limited. UPS module savings over a 2N architecture also drive this configuration. Other industry factors that drive distributed redundant configurations are as follows:

Static transfer switch (STS) – An STS has two inputs and one output. It typically accepts power from two different UPS systems (or any other type of sources), and provides the load with conditioned power from one of them. Upon a failure of its primary UPS feeders the STS will transfer the load to its secondary UPS feeder in about 4 to 8 milliseconds, and thus keep the load on protected power at all times. This technology was developed in the early 1990's, has been improved over time, and is commonly used in distributed redundant configurations.

A best practice for redundant dual path architectures is to isolate both paths so that they are independent of each other so that a failure on one side can not propagate to the other side. The use of static transfer switches in dual path architectures prevents the isolation of both redundant paths. Therefore, it is critical to base STS selection on a thorough investigation of static switch design and field performance. There are many options in an STS configuration and several grades of STS reliability on the market to consider. In **Figure 5**, the STS is upstream of the PDU (on the higher voltage side). Improvements in STS logic and design have improved the reliability of this configuration. Placing the STS on the lower voltage side of two PDUs is more reliable but is also much more expensive because twice as many PDUs must be purchased, and the STS will be at a lower voltage making it have a much higher

current rating. This configuration is discussed in greater detail in White Paper 48, [Comparing Availability of Various Rack Power Redundancy Configurations](#).

Single-corded loads - When the environment consists of single-corded equipment, each piece of IT equipment can only be fed from a single STS or rack mount transfer switch. Bringing the switch closer to the load is a prerequisite for high availability in redundant architectures as demonstrated in White Paper 48. Placing hundreds of single-corded devices on a single large STS is an elevated risk factor. Deploying multiple smaller switches feeding smaller percentages of the loads would mitigate this concern. In addition, distributed rack-mount transfer switches do not exhibit the failure modes that propagate faults upstream to multiple UPS system as is the case with larger STS. For this reason, the use of rack-based transfer switches is becoming more common, particularly when only a fraction of the load is single-corded. White Paper 62, [Powering Single-Corded Equipment in a Dual Path Environment](#) discusses the differences between STS and rack mount transfer switches in greater detail.

To use STS since the IT equipment would not detect the short transfer time upon failure of “A” or “B”.

Dual-corded loads – Dual-corded loads are becoming more the standard as time progresses, therefore the use of an STS is becoming less necessary. The loads can simply be connected to two separate PDUs which are fed from separate UPS systems. A common concern is that dual-corded loads may experience downtime if one of the power paths fails. This could happen in cases where both cords are mistakenly plugged into the same power path. Some customers claim to have experienced downtime of IT equipment when they unplug one of the redundant cords and attribute this to a defective power supply(s). These examples are often cited as a reason to use STS since the IT equipment would not detect the short transfer time upon failure of “A” or “B”.

Multiple source synchronization - When STS units are employed in a data center, it is preferable for the two UPS feeds to be in synchronization. Without synchronization control, it is possible for UPS modules to be out of phase, especially when they are running on battery. Many STS today have the ability to transfer sources that are not synchronized. This capability should make a multiple source synchronizer unnecessary.

A solution to prevent an out of phase transfer is to install a synchronization unit between the two UPS systems, allowing them to synchronize their AC output. This is especially critical when the UPS modules have lost input power and are on battery operation. The synchronization unit makes sure that all UPS systems are in sync at all times, so during a transfer in the STS, the power will be 100% in phase, thus preventing an out of phase transfer and possible damage to downstream equipment. Of course, adding a synchronization unit on the output of independent UPS systems allows for the possibility of a failure that can simultaneously drop all UPS systems.

Advantages of a distributed redundant design

- Allows for concurrent maintenance of all components if all loads are dual-corded
- Cost savings versus a 2(N+1) design due to fewer UPS modules
- Two separate power paths from any given dual-corded load’s perspective provide redundancy from the service entrance
- UPS modules, switchgear, and other distribution equipment can be maintained without transferring the load to bypass mode, which would expose the load to unconditioned power. Many distributed redundant designs do not have a maintenance bypass circuit.

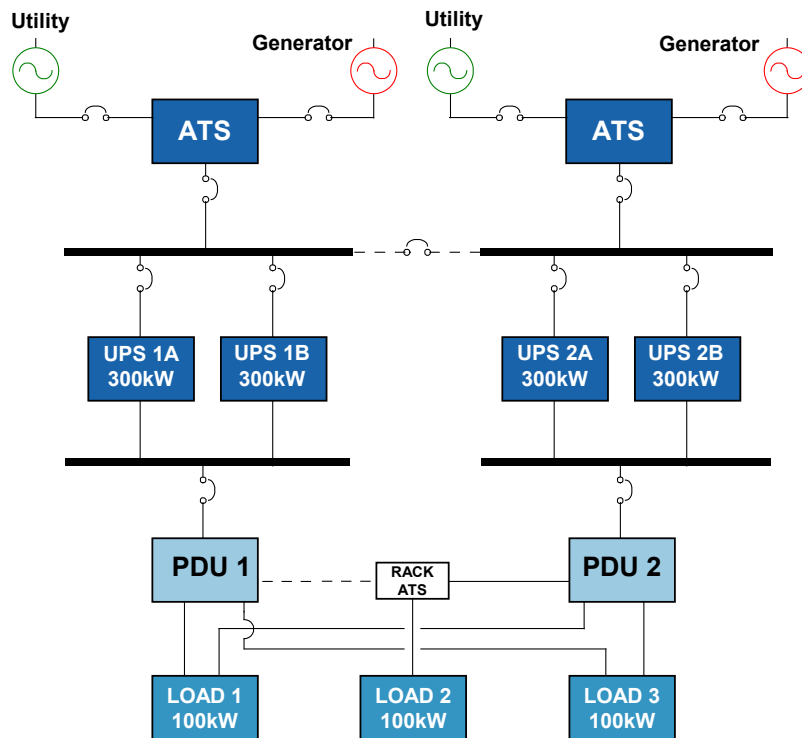
Disadvantages of a distributed redundant design

- Relatively high cost solution due to the extensive use of switchgear compared to previous configurations
- Design relies on the proper operation of the STS equipment which represents single points of failure and complex failure modes
- Complex configuration; in large installations that have many UPS modules and many static transfer switches and PDUs, it can become a management challenge to keep systems evenly loaded and know which systems are feeding which loads.
- Unexpected operating modes: the system has many operating modes and many possible transitions between them. It is difficult to test all of these modes under anticipated and fault conditions to verify the proper operation of the control strategy and of the fault clearing devices.
- UPS inefficiencies exist due to less than full load normal operation

System plus system redundant

“System plus system”, “isolated parallel”, “multiple parallel bus”, “double-ended”, “ $2(N+1)$ ”, “ $2N+2$ ”, “ $[(N+1) + (N+1)]$ ”, and “ $2N$ ” are all nomenclatures that refer to variations of this configuration. With this design, it now becomes possible to create UPS systems that may never require the load to be transferred to the utility power source. These systems can be designed to wring out every conceivable single point of failure. However, the more single points of failure that are eliminated, the more expensive this design will cost to implement. Most large system plus system installations are located in standalone, specially designed buildings. It is not uncommon for the infrastructure support spaces (UPS, battery, cooling, generator, utility, and electrical distribution rooms) to be equal in size to the data center equipment space, or even larger.

This is the most reliable, and most expensive, design in the industry. It can be very simple or very complex depending on the engineer’s vision and the requirements of the owner. Although a name has been given to this configuration, the details of the design can vary greatly and this, again, is in the vision and knowledge of the design engineer responsible for the job. The $2(N+1)$ variation of this configuration, as illustrated in **Figure 7**, revolves around the duplication of parallel redundant UPS systems. Optimally, these UPS systems would be fed from separate switchboards, and even from separate utility services and possibly separate generator systems. The extreme cost of building this type of facility has been justified by the importance of what is happening within the walls of the data center and the cost of downtime to operations. Many of the world’s largest organizations have chosen this configuration to protect their critical load.

Figure 7*2(N+1) UPS configuration*

The cost of this configuration is affected by how “deep and wide” the design engineer deems is necessary to take the system duplication efforts to meet the needs of the client. The fundamental concept behind this configuration requires that each piece of electrical equipment can fail or be turned off manually without requiring that the critical load be transferred to utility power. Common in 2(N+1) design are bypass circuits that will allow sections of the system to be shut down and bypassed to an alternate source that will maintain the redundant integrity of the installation. An example of this can be seen in **Figure 7**: the tie circuit between the UPS input panelboards will allow one of the utility service entrances to be shut down without requiring one of the UPS systems to be shut down. In a 2(N+1) design, a single UPS module failure will simply result in that UPS module being removed from the circuit, and its parallel modules assuming additional load. Maintenance bypass is not a benefit in many of these designs since they have a complete system “bypass”.

In this example, illustrated in **Figure 7**, the critical load is 300 kW, therefore the design requires that four 300 kW UPS modules be provided, two each on two separate parallel buses. Each bus feeds the necessary distribution to feed two separate paths directly to the dual-corded loads. The single-corded load, illustrated in **Figure 7**, shows how a transfer switch can bring redundancy closer to the load. However, tier IV power architectures require that all loads are dual-corded, including electrical feeds to air conditioning equipment.

Companies that choose system plus system configurations are generally more concerned about high availability than the cost of achieving it. These companies also have a large percentage of dual-corded loads. In addition to the factors discussed in the distributed redundant section, other factors that drive this design configuration are as follows:

Static transfer switch (STS) – With the advent of dual-cord capable IT equipment, these devices along with their undesirable failure modes can be eliminated with a significant increase in system availability.

Single-corded loads – To take full advantage of the redundancy benefits of system plus system designs, single-corded loads should be connected to transfer switches at the rack level. The benefits of doing so are illustrated in White Paper 48, [Comparing Availability of Various Rack Power Redundancy Configurations](#).

Advantages of a system plus system design

- Two separate power paths allows for no single points of failure; Very fault tolerant
- The configuration offers complete redundancy from the service entrance all the way to the critical loads
- In 2(N+1) designs, UPS redundancy still exists, even during concurrent maintenance
- UPS modules, switchgear, and other distribution equipment can be maintained without transferring the load to bypass mode, which would expose the load to unconditioned power
- Easier to keep systems evenly loaded and know which systems are feeding which loads.

Disadvantages of a system plus system design

- Highest cost solution due to the amount of redundant components
- UPS inefficiencies exist due to less than full load normal operation
- Typical buildings are not well suited for large highly available system plus system installations that require compartmentalizing of redundant components

Choosing the right configuration

How then does a company choose the path that is right for them? The considerations for selecting the appropriate configuration are:

- Cost / impact of downtime – How much money is flowing through the company every minute, how long will it take to recover systems after a failure? The answer to this question will help drive a budget discussion. If the answer is \$10,000,000 / minute versus \$1,000,000 / hour the discussion will be different.
- Budget – The cost of implementing a 2(N+1) design is significantly more, than a capacity design, a parallel redundant design, or even a distributed redundant. As an example of the cost difference in a 3.2 MW data center, a 2(N+1) design may require ten 800 kW modules (five modules per parallel bus; two parallel busses). A distributed redundant design for this same facility requires only six 800 kW modules.
- Types of loads (single vs. dual-corded) – Dual-corded loads provide a real opportunity for a design to leverage a redundant capability, but the system plus system design concept was created before dual-corded equipment existed. The computer manufacturing industry was definitely listening to their clients when they started making dual-corded loads. The nature of loads within the data center will help guide a design effort, but are much less a driving force than the issues stated above.
- Types of IT architecture – Virtualization and drastic improvements in network bandwidth and speed have opened up the possibility of failing over an entire data center to another location with little to no latency. This has brought into question the notion that the highest availability data centers are those with highly redundant power and cooling architectures. As virtualization technology matures, two remote data centers with 1N redundancy are likely to be more available than a single highly redundant data center.

- Risk tolerance – Companies that have not experienced a major failure are typically more risk tolerant than companies that have not. Smart companies will learn from what companies in their industry are doing. This is called “benchmarking” and it can be done in many ways. The more risk intolerant a company is, the more internal drive their will be to have more reliable operations, and disaster recovery capabilities.
- Availability performance – How much downtime can the company withstand in a typical year for a particular data center? If the answer is none, then a high availability design should be in the budget. However, if the business can shut down every night after 10 PM, and on most weekends, then the UPS configuration wouldn’t need to go far beyond a parallel redundant design. Every UPS will, at some point, need maintenance, and UPS systems do fail periodically, and somewhat unpredictably. The less time that can be found in a yearly schedule to allow for maintenance the more a system needs the elements of a redundant design.
- Reliability performance – The higher the reliability of a given UPS, the higher the probability that the system will continue working over time. For more information on reliability performance see White Paper 78, [*Performing Effective MTBF Comparisons for Data Center Infrastructure*](#).
- Maintainability performance – Simply having high reliability doesn’t prevent a failure from significantly affecting downtime. The amount of time to repair a system is heavily dependent on system design and skill level of the service technician. It is important to identify design attributes that increase repair time while decreasing the chance for human error.
- Maintainability support performance – The “effectiveness of an organization in respect of maintenance support”.⁴ One of the best ways to assess this criteria is to look at the experience other companies have had with a specific service organization.

The last four bullets can be rolled up into a term called dependability. Dependability, as defined by the International Electrotechnical Vocabulary 191-01-22 (IEV), is the “ability to perform as and when required”.⁵ While not a quantitative term, dependability combines the important factors that should be designed into a UPS and other critical systems that support a data center.

Table 3 is a useful starting point for selecting the right UPS system design configuration for a particular application. For designs with no or little redundancy of components, periods of downtime for maintenance should be expected. If this downtime is unacceptable, then a design that allows for concurrent maintenance should be selected. By following the questions in the flowchart, the appropriate system can be identified.

⁴ <http://std.iec.ch/iev/iev.nsf/display?openform&ievref=192-01-29> (accessed March 11, 2016)

⁵ <http://std.iec.ch/iev/iev.nsf/display?openform&ievref=192-01-22> (accessed March 11, 2016)

Table 3*Design configuration selection*

Configurations	Historical uses	Reasons for use
Capacity (N)	Small businesses Businesses with multiple local offices Businesses with geographically redundant data centers	Reduce capital cost and energy cost Support for lower criticality applications Simple configuration and installation Ability to bring down load for maintenance
Isolated redundant	Small to medium businesses Data centers typically below 500 kW of IT capacity	Improved fault tolerance over “1N” Ability to use different UPS models Ability to increase future capacity
Parallel redundant (N+1)	Small to large businesses with data centers typically below 500 kW of IT capacity	Improved fault tolerance over “1N” Ability to increase future capacity
Distributed redundant catcher	Large businesses with data centers typically above 1 MW of IT capacity	Ability to use different UPS models Ability to add more capacity Reduced UPS expense vs. 2N
Distributed redundant with STS	Large enterprises with data centers greater than 1 MW	Concurrent maintenance capability Reduced UPS expense vs. 2N
Distributed redundant without STS i.e. tri-redundant	Large collocation providers	Reduced UPS expense vs. 2N Increased savings over designs with STS
System plus system 2N, 2(N+1)	Large multi-megawatt data centers	Complete redundancy between side A & B Easier to keep UPS systems evenly loaded

Conclusion

The power infrastructure is critical to the successful operation of a data center's equipment. There are various UPS configurations that can be implemented, with advantages and limitations of each. By understanding the business's availability requirements, risk tolerance, and budget capability, an appropriate design can be selected. As demonstrated in the analysis of this paper, 2(N+1) architectures fed directly to dual-corded loads provide the highest availability by offering complete redundancy and eliminating single points of failure.



About the authors

Kevin McCarthy is a Vice President and Corporate Technical Officer at EDG2, a global engineering design and program management firm specializing in the innovative design and engineering of mission critical facilities such as data centers and trading floors. He has designed over 1,500 data centers, ranging in size from 100 square feet to 1,000,000 square feet, totaling over 5 million square feet of data center space. Kevin holds a Bachelor's degree in Electrical Engineering from The Ohio State University, with a minor in Computer Science. He is a regular speaker at national Data Center conferences, and is a member of 7x24, AFCOM, and the Washington Builders Conference.

Victor Avelar is a Senior Research Analyst at Schneider Electric. He is responsible for data center design and operations research, and consults with clients on risk assessment and design practices to optimize the availability and efficiency of their data center environments. Victor holds a Bachelor's degree in Mechanical Engineering from Rensselaer Polytechnic Institute and an MBA from Babson College. He is a member of AFCOM and the American Society for Quality.



[Avoiding Costs from Oversizing Data Center and Network Room Infrastructure](#)

White Paper 37



[Comparing Availability of Various Rack Power Redundancy Configurations](#)

White Paper 48



[Powering Single-Corded Equipment in a Dual Path Environment](#)

White Paper 62



[Performing Effective MTBF Comparisons for Data Center Infrastructure](#)

White Paper 78



[Browse all white papers](#)

whitepapers.apc.com



[Browse all TradeOff Tools™](#)

tools.apc.com



Contact us

For feedback and comments about the content of this white paper:

Data Center Science Center
dcsc@schneider-electric.com

If you are a customer and have questions specific to your data center project:

Contact your Schneider Electric representative at
www.apc.com/support/contact/index.cfm

Appendix - availability analysis

An availability analysis is done in order to quantify the availability difference between the five configurations presented in this paper. The details of the analysis are provided below.

Availability analysis approach

Schneider Electric's Availability Science Center uses an integrated availability analysis approach to calculate availability levels. This approach uses a combination of Reliability Block Diagram (RBD) and State Space modeling to illustrate the power availability at the outlet for these five configurations. RBDs are used to represent subsystems of the architecture, and state space diagrams, also referred to as Markov diagrams, are used to represent the various states the electrical architecture may enter. For example when the utility fails the UPS will transfer to battery. All data sources for the analysis are based on industry-accepted third parties such as IEEE and RAC. These statistical availability levels are based on independently validated assumptions.

Joanne Bechta Dugan, Ph.D., Professor at University of Virginia-

"[I have] found the analysis credible and the methodology sound. The combination of Reliability Block Diagrams (RBD) and Markov reward models (MRM) is an excellent choice that allows the flexibility and accuracy of the MRM to be combined with the simplicity of the RBD."

Data used in analysis

The data used to model the components is from third party sources. In this analysis the following key components are included:

1. Terminations
2. Circuit breakers
3. UPS systems
4. PDU
5. Static transfer switch (STS)
6. Generator
7. Automatic transfer switch (ATS)

The PDU is broken down into three basic subcomponents: circuit breakers, step-down transformer and terminations. The subpanel is evaluated based on one main breaker, one branch circuit breaker and terminations all in series.

Assumptions used in the analysis

It's critical that the reader correctly interpret the availability values of the five configurations. In order to carry out an availability analysis of complex systems, assumptions must be made to simplify the analysis. Therefore the availabilities presented here will be higher than what is expected in an actual installation. **Table A1** lists the basic assumptions used in this analysis.

Table A1*Assumptions of analysis*

Assumption	Description
Failure rates of components	All components in the analysis exhibit a constant failure rate. This is the best assumption, given that the equipment will be used only for its designed useful life period. If products were used beyond their useful life, then non-linearity would need to be built into the failure rate.
Repair teams	For “n” components in series it is assumed that “n” repair persons are available.
System components remain operating	All components within the system are assumed to remain operating while failed components are repaired.
Independence of failures	These models assume construction of the described architectures in accordance with Industry Best Practices. These result in a very low likelihood of common cause failures and propagation because of physical and electrical isolation. This assumption does not entirely apply to the distributed redundant architectures, because the Static Transfer Switch can cause two of the three UPS to fail thereby causing the failure of the entire architecture. This common cause failure was modeled for the two distributed redundant architectures.
Failure rate of wiring	Wiring between the components within the architectures has not been included in the calculations because wiring has a failure rate too low to predict with certainty and statistical relevance. Also previous work has shown that such a low failure rate minimally affects the overall availability. Major terminations have still been accounted for.
Human error	Downtime due to human error has not been accounted for in this analysis. Although this is a significant cause of data center downtime, the focus of these models is to compare power infrastructure architectures, and to identify physical weaknesses within those architectures. In addition, there exists a lack of data relating to how human error affects the availability.
Power availability is the key measure	This analysis provides information related to power availability. The availability of the business process will typically be lower because the return of power does not immediately result in the return of business availability. The IT systems typically have a restart time which adds unavailability that is not counted in this analysis.
Definition of failure per IEEE Std 493-1997 (Gold Book) IEEE Recommended Practice for the Design of Reliable Industrial and Commercial Power Systems	Any trouble with a power system component that causes any of the following to occur: •Partial or complete plant shutdown, or below-standard plant operation •Unacceptable performance of user's equipment •Operation of the electrical protective relaying or emergency operation of the plant electrical system •De-energization of any electric circuit or equipment

Failure and recovery rate data

Table A2 includes the values and sources of failure rate $\left(\frac{1}{MTTF}\right)$ and recovery rate

$\left(\frac{1}{MTTR}\right)$ data for each subcomponent, where MTTF is the mean time to failure and MTTR is mean time to recover.

Table A2*Components and values*

Component	Failure rate	Recovery rate	Source of data	Comments
Raw utility	3.887E-003	30.487	EPRI - Data for utility power was collected and a weighted average of all distributed power events was calculated.	This data is highly dependent on geographic location.
Diesel engine generator	1.0274E-04	0.25641	IEEE Gold Book Std 493-1997, Page 406	Failure rate is based on operating hours. 0.01350 failures per start attempt per Table 3-4 pg 44.
Automatic transfer switch	9.7949E-06	0.17422	Survey of Reliability / Availability - ASHRAE paper # 4489	Used to transfer electrical source from utility to generator and visa versa.
Termination, 0-600 V	1.4498E-08	0.26316	IEEE Gold Book Std 493-1997, Page 41	Used to connect two conductors.
6 terminations	8.6988E-08	0.26316	6 x IEEE value Computed from value by IEEE Gold Book Std 493-1997, Page 41	Upstream of the transformer, one termination exists per conductor. Since there are 2 sets of terminations between components a total of six terminations are used.
8 terminations	1.1598E-07	0.26316	8 x IEEE value Computed from value by IEEE Gold Book Std 493-1997, Page 41	Downstream of the transformer, one termination exists per conductor plus the neutral. Since there are 2 sets of terminations between components a total of eight terminations are used.
Circuit breaker	3.9954E-07	0.45455	IEEE Gold Book Std 493-1997, Page 40	Used to isolate components from electrical power for maintenance or fault containment. Fixed (including Molded case), 0-600A
PDU transformer, stepdown >100 kVA	7.0776E-07	0.00641	MTBF is from IEEE Gold Book Std 493-1997, Page 40, MTTR is average given by Marcus Transformer Data and Square D.	Used to step down the 480 VAC input to 208 VAC outputs, which is required for 120 VAC loads.
Static transfer switch	4.1600E-06	0.16667	Gordon Associates, Raleigh, NC	Failure rate includes controls; recovery rate was not given by ASHRAE for this size STS, so the value used is from the 600-1000A STS
UPS no bypass 150 kW	3.64E-05	0.125	Failure rate is from Power Quality Magazine, Feb 2001 issue, recovery rate data is based on assumption of 4 hours for service person to arrive, and 4 hours to repair system	UPS without bypass. MTBF is 27,440 hrs without bypass per MGE "Power Systems Applications Guide"

State space models

State space models were used to represent the various states in which each of the six architectures can exist. In addition to the reliability data, other variables were defined for use within the state space models (**Table A3**).

Table A3

State space models

Variable	Value	Source of data	Comments
PbypassFailSwitch	0.001	Industry average	Probability that the bypass will fail to switch successfully to utility in the case of a UPS fault.
Pbatfailed	0.001	Gordon Associates - Raleigh, NC	Probability that the UPS load drops when switching to battery. Includes controls.
Tbat	7 minutes		Battery runtime remained the same for all configurations.
Pgenfail_start	0.0135	IEEE Gold Book Std 493-1997, Page 44	Probability of generator failing to start. Failure Rate is based on operating hours. 0.01350 failures per start attempt per Table 3-4 pg 44. This probability accounts for ATS as well.
Tgen_start	0.05278	Industry average	Time delay for generator to start after a power outage. Equates to 190 seconds.

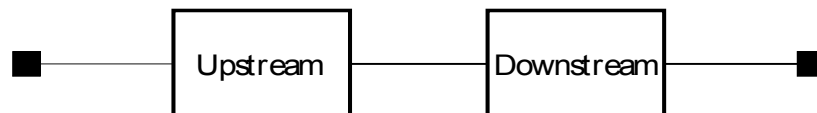
Availability model description

The purpose of this section is to provide an overview of how the analysis was carried out for the “capacity” configuration. **Figures A1** through **A3** represent the availability model for the “capacity” configuration of **Figure 1**. Models for the remaining UPS configurations were created using the same logic.

Figure A1 describes the series relationship between the upstream and downstream portions of the “capacity configuration”. The “Upstream” block represents everything between the utility and the UPS inclusive. The “Downstream” block represents everything after the UPS including all components up to and including the transformer output breaker.

Figure A1

Top level RBD representing upstream and downstream paths



Inside the “Power Input” block resides the Markov diagram used to calculate the availability of upstream components that feed the downstream components. The blocks at the top of **Figure A2** represent the individual components of the bypass, UPS system, generator, automatic transfer switch (ATS) and utility respectively. The failure and recovery rates from

these blocks feed the Markov diagram which results in an overall availability for the entire “Upstream” block.

Figure A2

Upstream Markov diagram

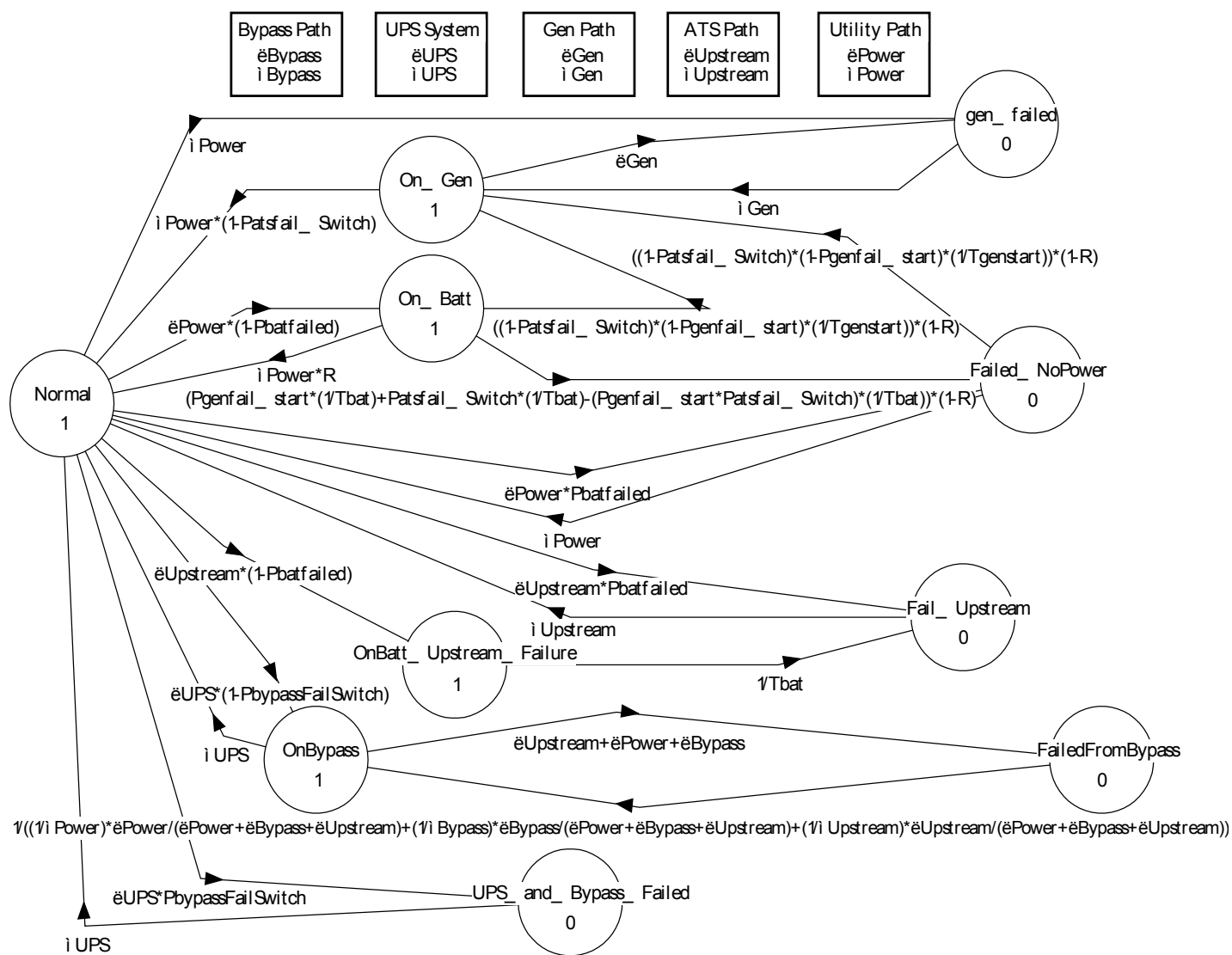
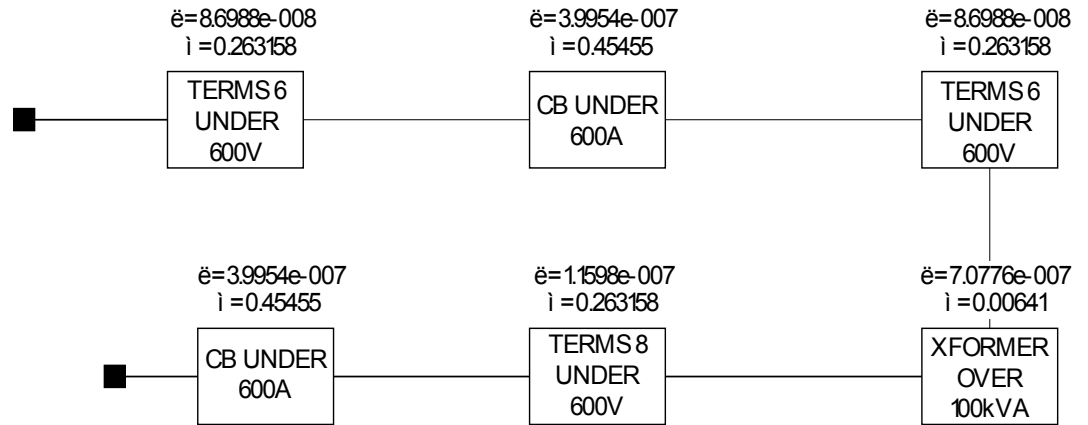


Figure A3 represents the components that make up the “Downstream” block in **Figure A1**. For the distributed redundant configurations (**Figures 4 and 5**), the STS would be added at the beginning of this component string.

Figure A3*Downstream diagram*

Results

Table A4 illustrates the results of the analysis for all 5 UPS configurations. Note that two versions of the distributed redundant UPS configuration were modeled (“catcher” and “with STS”).

Table A4*Results of analysis*

UPS configuration	Figure number	Availability
Single module “capacity” UPS configuration	1	99.92%
Isolated redundant UPS configuration	2	99.93%
Parallel redundant (N+1) UPS configuration	3	99.93%
Distributed redundant “catcher” UPS configuration	4	99.9989%
Distributed redundant UPS configuration	5	99.9994%
2(N+1) UPS configuration	7	99.99997%