



**QUINT S-ORING with
QUINT POWER 20 A**

Creating failure tolerance with redundancy in 24 V DC systems

By Andrew Bruce, Associate Business Development Application Engineer – Power Supplies, Phoenix Contact USA

Introduction

Failure tolerance is an important design criterion for any automation system. In cases where losing power can have serious physical or financial repercussions, redundancy becomes necessary. While the word redundant strictly means “exceeding what is necessary,” in power systems, redundancy increases reliability.

The most basic form of redundancy is wiring two power supplies in parallel. If one power supply fails, the redundant power supply takes over. This simple concept can be elaborated on and dissected to provide a toolkit for full system redundancy. Designing a failure-tolerant system simply requires answering two questions:

- What can go wrong?
- What can be done about it?

INSIDE:

Introduction	1
Understanding failures	2
The failure tolerance toolbelt	3
Independent paths	4
System diversity	4
Monitoring and maintenance	5
Conclusion	6
References	7

Understanding failures

The first step in designing a redundant power system is to understand why redundancy would be necessary. In an ideal world, extra components are inefficient. Devices would perform forever, and the system would never lose power. However, in reality, failures will occur, and redundancy is about designing with that in mind.

Understanding failure mechanisms improves the ability to design a failure-tolerant system. Failures fall into three categories:

1. Early
2. Random
3. End of life (EOL)

Early failures occur relatively quickly when the component is initially placed under load. Generally, these failures are the result of imperfections in the manufacturing process. Manufacturers can largely eliminate early failures by performing a burn-in process and performing regular quality checks during production. A burn-in process entails putting components under stress after manufacturing and exposing errors. Pieces that fail the burn-in process are discarded or reworked.

Random failures are quantified by the mean time between failure (MTBF). This is derived from a formula based on the number, configuration, and type of components used in a device and their individual rates of failure. By sheer probability, random failures will occur, and some components are more likely to experience them. Complicated systems will have lower MTBFs because they have more components that can fail. Often manufacturers will list the MTBF of assembled products in their datasheets so users do not have to apply the complicated formula or derive this value experimentally.

Finally, end-of-life failure occurs when a component reaches the end of its useful life. This type of failure is expected and inevitable. Most industrial components will have a weak link in their design; that is, a component that reaches end-of-life first in the majority of cases. The lifetime of a device can be calculated by the lifetime of the weakest component. In the example of a car tire, an early failure would be if the tires failed on their first trip. A lifetime failure occurs when you have driven thousands of miles, and the tread is worn down. A random failure would be an accident like running over a nail.

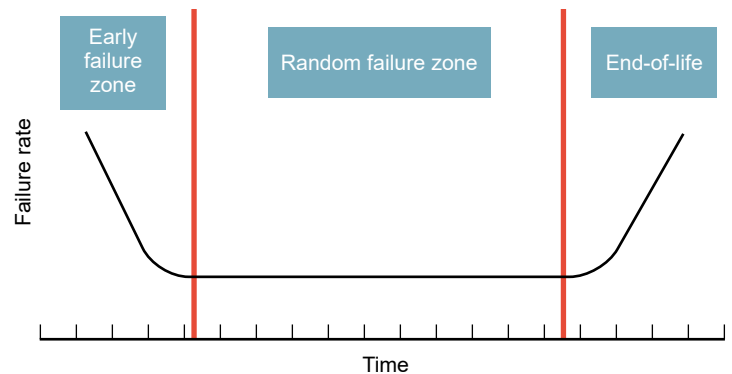


Figure 1: Failure rate curve

Figure 1 shows the failure rate of a component over its life. The failure rate is initially high, where missed imperfections in manufacturing are uncovered. The rate drops during the normal useful lifetime, which is described by MTBF. The rate of random failure can be reduced with properly installed surge protection devices (SPDs).

Transients, also known as surges, damage and weaken components and are the most common source of random failures. Protecting against these will have a significant impact on the rate of random failure. As the components reach the end of their useful life, failure rate will increase again and asymptote toward infinity.

As manufacturing becomes more precise, random failures are less likely to occur. Industrial components will often have a larger MTBF than lifetime. For example, the lifetime of the product is 10 years, but the MTBF is 50 years. When this is the case, product lifetime is more important to design around as end-of-life failures occur more frequently than random failures. Selecting a device with high-quality components will reduce the rate of failure from EOL.

The next step in understanding failures is to review what can happen when a component fails. In a 24 V power system, there are three primary failure modes:

- 1. Open failure.** The component stops functioning, and no current passes through it.
- 2. Short failure.** The component directly shorts to ground, and all available current flows to the short.
- 3. High failure.** This can occur in a power supply as a failure in regulation, where high internal voltages are passed to the output. Voltage-limiting devices can be built into the power supply or added externally.

Failures can occur from the strain of operation or from operator error. It is very easy for a human to make a wiring mistake and cause a short. This paper will primarily address component failures from operation, but it is important to mitigate human errors with clear labeling, intuitive design, and proper training.

The most reliable way of assessing the robustness of a 24 V system is to go through each component and visualize what would happen if each failure mode occurred. Several concepts can be applied in layers to work around these failures.

The failure tolerance toolbelt

There are five key concepts for creating a failure-tolerant and reliable system:

1. Component redundancy
2. Fault isolation
3. Independent paths
4. System diversity
5. Monitoring and maintenance

Component redundancy

When people use the word redundancy in reference to a 24 V DC power system, they are normally referring to component redundancy: having an additional component in parallel to perform the intended function should the first component fail. In the case of a 24 V DC system, this is done with power supplies. The concept of $n+1$ component failure tolerance means that one additional component is added to the minimum number needed to perform the function. In this way, any single component can fail without impacting the system. For example, for a 40-amp load, three parallel 20-amp power supplies can give $n+1$ redundancy, where $n=2$. Two power supplies are needed to output 40 amps, and the third is redundant in case one fails. To increase the failure tolerance of the system, $n+2$ or $n+3$ redundancy can be used.

Isolation

It is important to protect against each type of failure mode when designing a system. If one of these power supplies failed in a short, or if there was a wiring short on their outputs, current would go to the path of least resistance. All parallel power supplies would feed the short instead

of the load. With hypothetically infinite current, the voltage output would drop, or the power supplies would turn off their output to protect themselves. This can be managed by isolating the shorts in two ways: direction and disconnection. These two methods isolate components and prevent one fault from impacting the whole system.

The standard approach for current direction is to wire two power supplies to a diode module (Figure 2). A diode will only allow current to pass through it in one direction, toward the load.

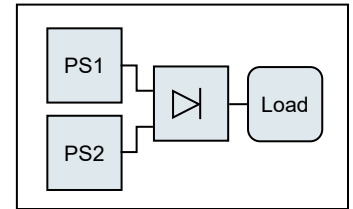


Figure 2: Diode module

Diodes are a very simple solution and have high MTBFs. However, diodes are not particularly efficient and generally do not have intelligent features.

Current direction can also be achieved using MOSFETs instead of diodes. MOSFETs offer up to 75% energy savings through less voltage drop and power dissipation. MOSFET-based modules are also more suitable for intelligent features, such as automatic current balancing and monitoring. The voltage difference between the two power supplies will determine which provides more current to the load. Two power supplies with the exact same voltage will split the current evenly. By adjusting internal resistance, a MOSFET-based module can compensate for a voltage difference between the power supplies, ensuring each supply shares the load evenly. This is known as automatic current balancing.

Current flow produces heat, and heat is the largest detriment to the useful life of a power supply. Balancing the current will evenly distribute the heat, increasing the lifetime and the reliability of the system. Some redundancy modules will also include voltage limitation circuits, protecting against high failures.

The second way to isolate faults is with disconnection. While many standards require circuit breakers, the strategic use of circuit breakers can also increase failure tolerance. Individual breakers on the AC side of the power supply will enable maintenance and disconnection in the case of a short upstream. Circuit breakers or fusing on the DC side for each load will prevent a single load from

taking down the whole system. If a load experiences a short, the breaker will open up and disconnect the load. This keeps the damage caused by the short isolated to that particular branch.

Since direction and disconnection address faults in different devices, most applications will benefit from a combination of the two.

Independent paths

In figure 2, both power supplies feed into a single module. Although diodes and MOSFETs are largely passive and have much lower failure rates than a power supply, the module itself is a single point of failure. This is where independent paths come into play. If the redundancy module fails, the system would need to be powered down for maintenance. To get around this, the extra power supply terminals can be redirected to an additional dual-channel module, creating a second, independent path for current to flow. This is often called “cross-wiring” and is shown in Figure 3. For this to work, the power supply must have multiple output terminals, or extra terminal blocks are needed. A path is considered independent if it does not rely on its redundant counterpart to operate.

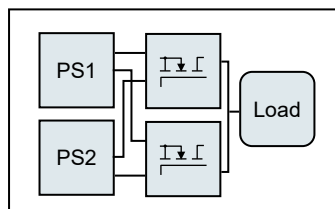


Figure 3: Cross-wiring

A single-channel module is effectively a traditional module split in half. Instead of feeding both power supplies into a single module to direct the current and prevent back-feeding, each

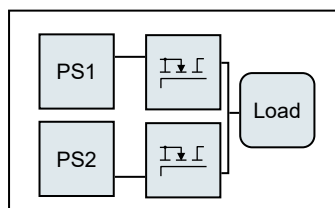


Figure 4: Single-channel redundancy

power supply has its own module. This elevates the failure tolerance from component-dependent to path-dependent. In the case of the cross-wired unit, we have multiple independent paths that can allow each component to fail independently of the others. The failure of the single-channel module also takes out the power supply upstream of it.

To address the independent path issue caused by traditional diode modules, some manufacturers have designed single-channel modules (Figure 4). A

single-channel module is effectively a traditional module split in half. Instead of feeding both power supplies into a single module to direct the current and prevent back-feeding, each power supply has its own module. This elevates the failure tolerance from component-dependent to path-dependent. In the case of the cross-wired unit, we have multiple independent paths that can

The concept of current paths is useful in identifying potential weak points in a redundant system. Starting from the upstream power source and working down to the load, identify the path that current will flow given a short-circuit failure or an open-circuit failure for every wire and component. This approach will ensure power has an independent path to travel given any single-point fault. However, some types of failures can affect multiple points at once, creating the need for system diversity.

System diversity

System diversity is the technical equivalent of “Don’t put all of your eggs in one basket.” If all of a system’s power originates from a single utility and that utility has a fault, no amount of component redundancy can prevent downtime. A utility failure is a common-component failure, meaning it will cause any component relying on it to fail. There are multiple types of common-component failures, but each can be addressed with system diversity.

Source failures

Utility power is generally brought to a facility in two- or three-phase AC systems. The first step in adding diversity is to tap each redundant power supply off of a different phase. If one phase experiences a failure, the second has a chance of staying active. However, both phases can fail from a common cause, and some buildings will have separate, independent utilities supplying power. One utility can completely fail, and the second will have an independent path to supply the loads. A second way to diversify the source power is to include a variety of physical sources. Natural disasters can affect all utilities, and offline alternative power is useful. This can be a battery-powered uninterruptible power supply (UPS), solar power, or an on-site generator. Extremely critical cases may warrant a combination.

Manufacturing failures

Manufacturers can have a variety of problems supplying specific part numbers. Large orders can drain stock, manufacturing downtimes can occur, or political changes in the country of origin can affect delivery. While these are not common, they can be mitigated. Having backup component types, for instance, a diode module and a MOSFET module that will both work in the system, can be useful if

the manufacturing line for one goes down. Additionally, companies with manufacturing in different countries are incorporating a style of redundancy that can improve their ability to deliver products.

Component design failures

Even after extensive testing, components can have design flaws that go undiscovered. For example, a solder joint on a component could be vulnerable to rapid heating and cooling. While it is very unlikely the solder joint fails in every component at the same time, different design flaws could be more regular. If every current path has the same component in it, and this component is vulnerable to a design flaw, one event can cause all current paths to fail. Having a diverse set of components in each current path can prevent downtime from this type of common-mode fault. For instance, one path coming from utility power and a power supply, one path coming from a battery bank with a DC/DC converter, and one path coming from a solar panel and a charge controller virtually eliminate the possibility of failure. Diversity and redundancy at every level will only increase the robustness of a system.

Monitoring and maintenance

If components are not replaced after failure, redundancy can only delay system downtime. Replacing components requires both awareness of faults and the ability to remove components without impacting current flow. Many industrial components will have built-in monitoring capabilities, such as in the MOSFET-based redundancy module example earlier, but if this is not the case, external components can be added.

Monitoring is particularly necessary in redundant systems because initial component failures will not impact the loads. There will not be the obvious failure indication of everything grinding to a halt. A system will still work if a redundant power supply has failed, but if the operator is not aware of the failure, a second failure will cut power to the system.

The level of monitoring implemented can range significantly. On the simple end of the spectrum, the alarm built into most components can change states on failure, indicating maintenance is required. In applications where uptime is critical, more complex solutions that warn when the system is at risk of failure can be used. This is called preventive

maintenance. Instead of responding to problems, baseline parameters for each operating state are established, and deviations are monitored and addressed.

There are some cases where detailed analog monitoring can catch faults that binary, component-failure monitoring cannot. For example, slight changes in the load can also create a loss of redundancy for two power supplies. If the load increases past the capacity of one power supply, there is no component failure tolerance. These types of changes can result from increased mechanical load, like clogged filters and low oil, or from additional loads haphazardly being added to the system. If one power supply were to fail, the other would be overloaded, and the output would shut off. Without sufficient monitoring, redundancy can be lost without any indication, leaving the system vulnerable to failure. This is an example where analog or intelligent monitoring is needed.

Another factor to consider in monitoring is the existing infrastructure. Some facilities have high-level protocols and supervisory control and data acquisition (SCADA) systems in place. Monitoring products that use the same protocol can be easily integrated into the existing system. Other facilities may not have any type of visualization software to give detailed analysis of the system, and simple binary indications like LEDs and alarms are more suitable.

Once faults are detected, maintenance needs to be performed. Depending on the failed component and the type of fault, it will be either repaired or replaced. How maintenance is handled will vary greatly depending on the industry and application. Some industries have planned downtime when all critical maintenance takes place. In these cases, the goal of a robust, redundant system is to survive until the planned downtime.

In applications where perpetual uptime is expected, the term “hot-swappable” becomes important. To hot-swap a component is to replace it under load. There are safety requirements and standards regulating how and when this can be done to ensure the safety of the operators. These guidelines must be followed whenever maintenance is performed under load. From the feasibility perspective, some design considerations can allow for uninterrupted power. These ideas are already introduced but can be applied differently from the maintenance perspective.

The two considerations are alternative paths and proper disconnection. The redundant system design should have already addressed the alternative paths. When a part is replaced,

current should have another way of getting to the load. Power needs to be disconnected from the component so it can be safely wired. Live wires carry a high probability of accidental shorting. Pluggable terminals can be a safe method for disconnection and can reduce the chance of errors upon rewiring the new component into the system. Alternatively, well-placed circuit breakers can isolate components for replacement.

Assessing redundancy needs

The first step in applying these concepts and actually designing a redundant power system is to determine how critical the system is. There are some applications where losing power is acceptable. In other applications, losing 24 V power can cost hundreds of thousands of dollars every hour. If losing power is not costly or dangerous, then it may be unwarranted to go all out to ensure independent current paths, disconnects, monitoring, and source diversity for every type of failure. However, in the situation where a loss of power can mean significant downtime or risk of human injury, the cost of a second power supply or a battery backup and management system is negligible.

Things to consider:

- Cost of power loss
- Likelihood of failures
- Safety implications of failure

Most applications will be somewhere in the middle. From there, the best approach is to design around the most likely and the most catastrophic failures. The most likely failures would be indicated by components with the lowest expected lifetime and lowest calculated MTBF. Often, lifetime is shorter than MTBF, so it is more likely to cause failure. These components should have redundant counterparts. The most common failure mode when the system is up and running (as opposed to on startup) is failing low, where the component stops passing current and is seen as an open circuit.

Utility power is also a likely source of failure. Natural disasters, routine lightning storms, and rogue squirrels pose a threat. Any system that needs to avoid downtime should take utility power loss into consideration. Adding source diversity with a UPS or a generator will lead to greater uptime.

Instead of qualitatively assessing how to apply redundancy, a quantitative approach can be taken. Comparing the cost of implementing redundant components against the cost and probability of downtime can provide an easy method to establish the budget for redundancy.

Budget for redundancy = \$(Cost of downtime) x (likelihood of failure) x (MTTR)

The cost of downtime can include the wages of the technician, loss of a batch, and cost of product that would have been produced if the system was running.

The likelihood of failure can be calculated by the summation of individual component MTBFs in the following equation. In cases where the lifetime of the product is lower than the MTBF, lifetime is an appropriate substitute value.

$$mtbf(c_1; \dots; c_n) = \left(\sum_{k=1}^n \frac{1}{mtbf(c_k)} \right)^{-1}$$

("Mean time between failures.")

MTTR is the mean time to repair, or the average time it takes to get the system back up and running. A remote system that requires a technician to drive to the site will have a higher MTTR than a system with staff and stock on site. If replacement products are not kept on hand, the manufacturing lead times should also be considered.

Conclusion

Evaluating redundancy is a more in-depth process than what is apparent at first glance. However, the two questions posed in the introduction summarize this process:

- What can go wrong?
- What can be done about it?

The concepts outlined here can be applied in most scenarios to determine how these questions should be answered. It is a process of viewing each component of a system at different levels of analysis, evaluating the possible failure mechanisms at that level, and designing a network of diverse paths for current to flow. Taking external factors like monitoring and maintenance into account further answers what can be done about potential failures. Ultimately, no system will reach 100% failure tolerance, but applying these redundancy concepts will bring the system to a tolerable level for the application.

References

“Mean time between failures.” Wikipedia. Wikimedia Foundation. 4 November 2020. https://en.wikipedia.org/wiki/Mean_time_between_failures

ABOUT PHOENIX CONTACT

Phoenix Contact develops and manufactures industrial electrical and electronic technology products that power, protect, connect, and automate systems and equipment for a wide range of industries. Phoenix Contact GmbH & Co. KG, Blomberg, Germany, operates 50 international subsidiaries, including Phoenix Contact USA in Middletown, Pennsylvania.

For more information about Phoenix Contact or its products, visit www.phoenixcontact.com, call technical service at **800-322-3225**, or e-mail info@phoenixcon.com.