



WHITEPAPER

AI server hot-plugging with robust hot-swap controllers

Strategies for reliable and continuous operation in AI server data centers using hot-swap controllers and the role of eFuses in the future

Nitish Agarwal, Senior Application Engineer, Infineon Technologies Americas Corp.

Edition: v1.0
Date: 11/2024

www.infineon.com



Table of contents

Abstract	3
1 AI servers and the need for hot-plugging	4
1.1 Transition from 12 V to 48 V architecture	4
1.2 Hot-swap controllers overcome 48 V architecture's shortcomings	4
2 Events during the hot-plugging of a server tray into a server rack	5
2.1 The occurrence of the inrush current	5
2.2 Reducing the inrush current	6
2.3 A hot-swap controller's basic mechanism	6
2.4 Reduction of the inrush current by a hot-swap controller at startup	8
2.5 The importance of the SOA curve for hot-plugging	10
3 Introduction to digital hot-swap controllers	11
3.1 An example of a 4 kW hot-swap solution	11
3.2 Use of burst mode to turn on the FET instead of continuous mode	14
3.3 Eliminating the effect of contact debounce during hot-plugging	15
3.4 Protecting the server blade and backplane during fault events	15
3.5 Active monitoring and health reporting of the system	16
4 Future trends	17
References	18

Abstract

The rapidly evolving AI landscape has intensified the demand for continuous, high-performance computing. AI servers, comprising multiple AI processor blades operating in parallel are the backbone of this computational power. These servers often rely on a 48 V backplane architecture to meet growing power needs while minimizing energy losses. In this “always online” era, the challenge of maintaining 24/7 server uptime necessitates the ability to hot-plug or hot-swap AI server blades without disrupting operations.

This whitepaper delves into the necessity and mechanics of hot-plugging in AI servers, highlighting the critical role of hot-swap controllers. It begins with an introduction to the requirements for hot-plugging in AI server environments, followed by a detailed exploration of the events and challenges encountered during the hot-plugging process, including the occurrence and mitigation of inrush currents.

Various methods to handle inrush currents are discussed, leading to the introduction of hot-swap controllers and their basic control mechanisms. The paper explains how these controllers manage MOSFET operations within their safe operating area (SOA) to ensure reliability and prevent damage. It also covers the advantages of digital hot-swap controllers, which offer programmable SOA profiles for enhanced system protection.

Additionally, the paper addresses the elimination of contact debounce issues during hot-plugging and the comprehensive fault protection features offered by modern hot-swap controllers. The importance of active monitoring and telemetry for system health reporting is also emphasized. Finally, future trends, including integrated solutions like eFuses and the transition to higher voltage backplanes, are explored.

1 AI servers and the need for hot-plugging

In today's era, we see ourselves surrounded by AI in one way or the other, from virtual chatbots providing real-time tech support to autonomous vehicles taking us to our destination without any human intervention. AI is expected to perform highly complex computational algorithms that can think, learn, and act in a way that is similar to human intelligence, in real time or even exceed a human brain's prowess. These complex algorithms are processed either via one powerful processor or many processors working in parallel to compute vast amount of data in real time. Due to the advantages, cost savings, and reliability that AI has to offer, industries from a wide spectrum are moving towards implementing AI into their business models, which has increased the overall need for more and more AI processors to be operational to feed this growing demand.

Essentially, an AI processor is operated at a voltage level of 1 V or lower, but this voltage is converted from a 48 V backplane using a power converter. This power converter — together with the processor, network switches, and a memory — forms an AI server blade (or server tray). Multiple AI server blades are operated in parallel to compute the required data. The increasing dependency on AI makes it essential to keep these AI servers online 24/7 to avoid any interruption. Such a continuous on time can be achieved by operating several AI server blades in parallel to share the load forming an AI server rack.

To visualize the parallel arrangement of AI server blades, imagine the drawers of a large filing cabinet: the back of the cabinet into which these drawers (server tray) wheel in is a 48 V live backplane, and a cabinet is the server rack. The AI server blades are plugged into this 48 V backplane, supplied by the same power source which powers up the processors. At any point, if any one of the AI server blades has an issue, the load is shared by the other blades to keep the system running and online 24/7. In the meantime, the faulty blades can be wheeled out of the rack and be replaced by a new server blade, plugging into the live backplane without taking the entire system offline. This process of plugging the server blade into the live backplane is called “hot-plugging” or “hot-swapping”.

1.1 Transition from 12 V to 48 V architecture

The old architecture of a server had a 12 V backplane; however, with growing power demands, the 12 V architecture is not a feasible option anymore. This has forced the industry to transition towards the 48 V architecture. This higher power architecture reduces power losses by 16x which allows greater power extraction with reduced cooling costs. Equally important is the fact that these server blades have to stay immune to any potential damage caused by external factors such as voltage fluctuations or unwanted thermal events. The hot-swap controller also ensures that any fault in the server blades is isolated locally so that it does not propagate to the adjacent blades, which can potentially cause a system level failure.

Transitioning the backplane changes to the 48 V architecture, hot-plugging, and the protection requirements of these server blades add up costs rapidly. A study from Forbes [1] has highlighted that a server downtime can cost companies up to \$9000 per minute. For higher risk enterprises, the costs can exceed \$5 million an hour in certain scenarios.

1.2 Hot-swap controllers overcome 48 V architecture's shortcomings

All these shortcomings and requirements can be taken care of by a hot-swap controller. For today's server blades, it is crucial to have a robust and reliable hot-swap solution deployed, which will be the focus of the present paper. The paper is expected to assist engineers working on exploring different hot-swap solutions for their application.

After highlighting the need for a hot-swap controller during the hot-plugging operation and explaining the different methods of hot-plugging, the paper introduces the hot-swap controller and discusses the different operating regions of the MOSFETs and their safe operating area (SOA) charts. It then outlines the SOA control method for digital hot-swap controllers with an example of a 4 kW design along with a brief introduction of protection schemes and telemetry accuracy. Finally, it briefly discusses the future trends in hot-swap controllers.

2 Events during the hot-plugging of a server tray into a server rack

Every server blade that is hot-plugged into the live backplane has a capacitor as the first component for energy storage and filtering out voltage ripple and high-frequency noises. The capacitor at the time of the insertion of the power module is completely discharged and acts as a short or a low-impedance path from the input voltage V_{IN} to the ground (GND). If, at this moment, the server blade is hot-plugged into the live backplane, a huge amount of inrush current can be observed for a short period of time — strong enough to blow up the fuse or cause a significant voltage dip in the adjacent connected server trays — triggering a system-wide shut down.

To feed the growing power demands of AI servers, the power density of the power supply modules is also increasing, and this further increases the amount of capacitance installed on the system at the front end. This makes it critical to reduce this inrush current.

2.1 The occurrence of the inrush current

The following circuit shows the connection between an AI server blade and the backplane. The R_{INT} is the parasitic series resistance in the power path caused due to contact resistance and PCB trace resistance.

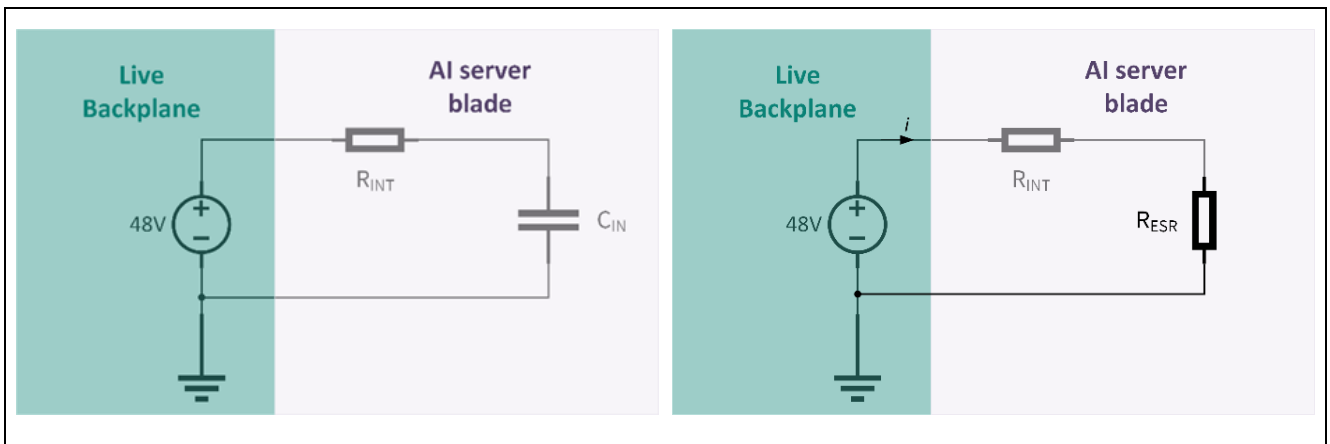


Figure 1 Typical front-end of a server blade's circuit (left) and the equivalent server blade impedance (right)

At $t = 0$, when the power modules are hot-plugged into the backplane, the capacitor C_{IN} is completely discharged and acts as a short between V_{IN} and GND, in theory. But in reality, it has a small equivalent series resistance (ESR) in the range of a few m Ω . For simplicity, let us consider it to be around 10 m Ω . Similarly, R_{INT} is also a very small value and is close to 1 m Ω .

As shown in Figure 1 (right) V_{IN} is connected to GND with an impedance of R_{INT} and R_{ESR} . Considering this to be a 48 V system, the input current (i) at startup can be calculated using the following equation:

$$i(t = 0) = \frac{V_{in}}{R_{INT} + R_{ESR}} = \frac{48}{0.001 + 0.01} = 4.363 \text{ kA}$$

Equation 1 Inrush current calculation

Even though this 4.3 kA of inrush current lasts for a small period of time, it can cause significant damage to the entire backplane as well as to the hot-plugged server blade.

2.2 Reducing the inrush current

As we have set the basis for the cause of the inrush current, we can now explore different techniques to lower it.

Upon taking a closer look at [Equation 1](#), we can observe that by increasing the resistance in the current's path, i.e., by adding a series resistance in between V_{IN} and C_{IN} , the current can be reduced to a desirable limit. But adding an extra resistance in the power path results in a voltage drop and adds to I^2R losses.

To demonstrate with an example, consider the total series resistance ($R_{INT} + R_{ESR} +$ the introduced resistance) between V_{IN} and GND to be $1\ \Omega$ and after turning on the power converter, the input current is around 10 A. With this added resistance, the inrush current at startup will come down to 48 A from 4.363 kA, but the resistor will dissipate power equal to 100 W ($I^2R = 10 \times 10 \times 1$). It also introduces a voltage drop of 10 V ($IR = 10 \times 1$), which significantly reduces the efficiency of the converter converting 48 V to 1 V. This results in added costs due to cooling and increased power consumption.

To reduce the adverse effect of the constant added resistance, there are Negative Temperature Coefficient (NTC) thermistors available in the market that offer a high resistance when they are cold, i.e., before the event of hot-plugging. When the server blade is taken off the shelf at room temperature and when the inrush current passes through these NTCs, it generates heat from the I^2R losses, which reduces the resistance. Thus, NTCs reduce the losses in the system at the steady state. But this resistance does not go down to a level where the losses would be insignificant. The resistances could be shunted after the short circuit event using a switch or relay, but this adds to the size, cost, and complexity of the system. Also, this mechanism offers an unreliable solution in the event of a fault when the server blades need to be isolated.

2.3 A hot-swap controller's basic mechanism

As explored above, adding an extra resistance in the power path is not a viable solution. A solution is needed where the input current to the capacitor can be controlled during the hot-plug event without causing a voltage drop or power losses during normal operation. [Figure 2](#) visualizes this idea by considering the backplane as an infinite water reservoir and the input capacitor as an empty water tank. To control the flow of water going into the tank from the reservoir, i.e., the flow of current, a control valve is added.

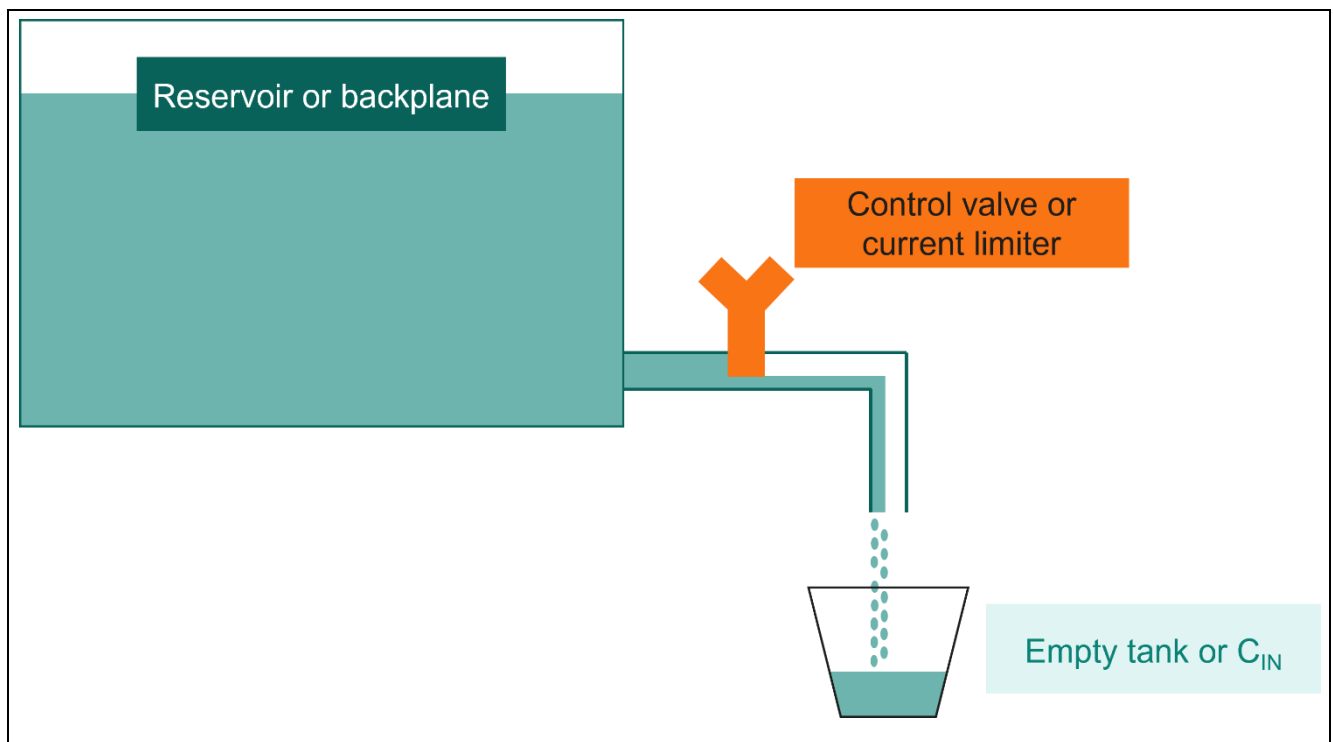


Figure 2 Capacitor inrush current analogy with water reservoir example

Without this control valve, the water will try to flow into the tank with the maximum possible flow rate, just like the inrush current. This is why the flow of water (the inrush current) must be regulated or controlled by a control valve.

The analogy of control valve can be drawn to controlling the flow of current through a MOSFET by controlling its gate voltage. As shown in Figure 3, when the gate voltage of the MOSFET is below V_{th} (turn-on threshold), the MOSFET operates in the cut-off region and does not allow any current to flow through it, thus blocking the current flow (inrush current) while inserting the server blade into the backplane.

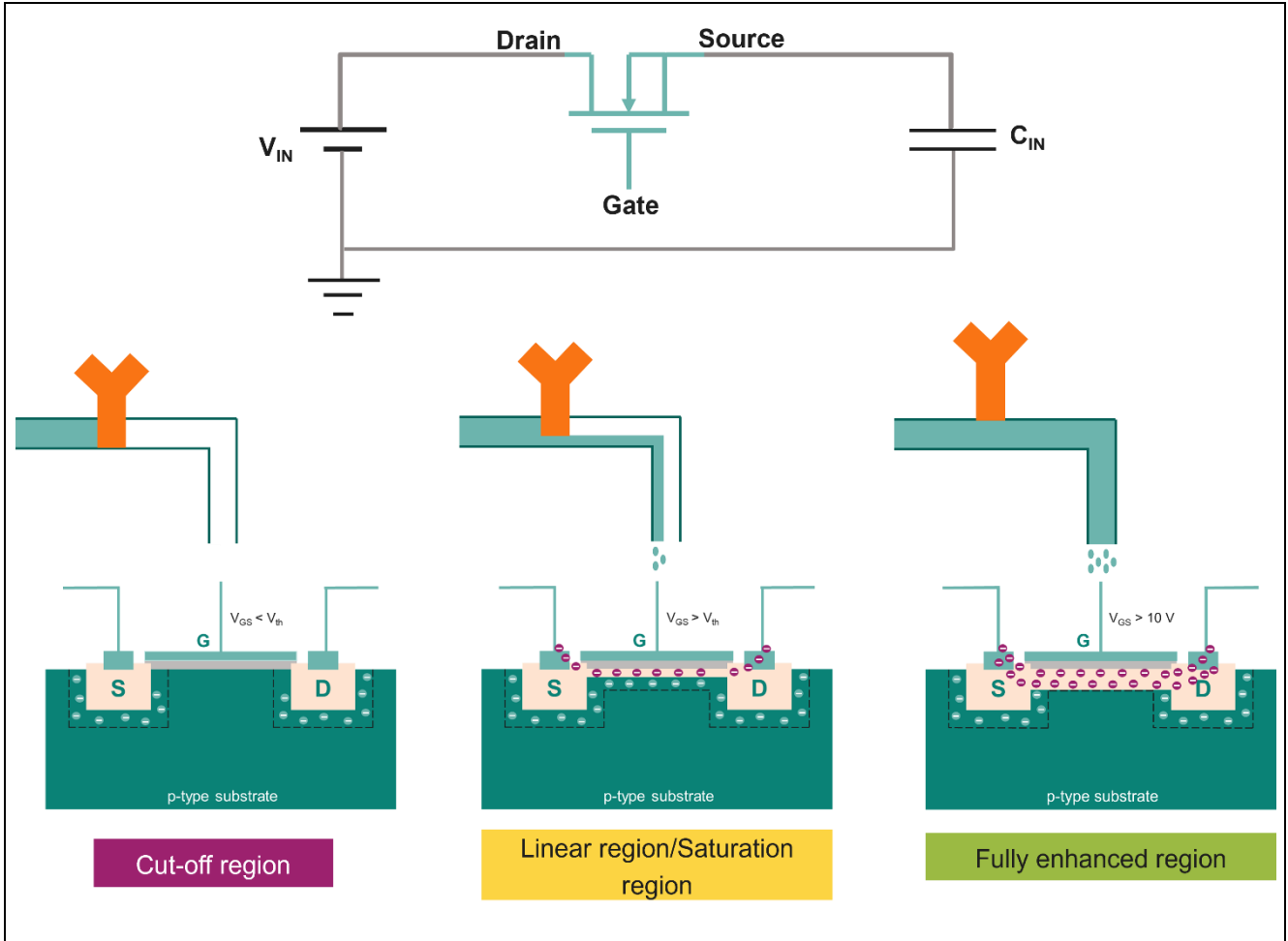


Figure 3 Capacitor inrush current analogy with water reservoir example

As the voltage on the gate, i.e., V_{GS} (gate to source voltage) exceeds the turn-on voltage, the current flows through the FET in a controlled manner. The amount of current that flows through the FET is dependent on the voltage present on the gate with reference to the source. As the voltage on the gate increases, the current through the FET also increases. This region of operation is called the linear region.

Regarding the hot-plugging event, V_{DS} (drain-to-source voltage) is high enough that it has a negligible effect on the drain current, which will be solely dependent on V_{GS} . This region of operation is called the saturation region. In the saturation region, as long as V_{GS} stays constant, the current through the MOSFET also stays constant. As the gate voltage is increased, more current is allowed to flow through the MOSFET. This charges the input capacitor and raises the voltage at the source; thus the system reaches a point where V_{DS} comes down to a level low enough for the MOSFET to be in the ohmic region. Here, the current through the MOSFET is dependent on the $R_{DS(on)}$ (resistance between the source and the drain) of the MOSFET. This impedance is further dependent on the gate voltage and the typical $R_{DS(on)}$ value is reached at a V_{GS} of 10 V for most MOSFETs.

The MOSFETs that are selected for hot-swap applications are in the 80 V to 100 V range, with a majority of the designs using 100 V MOSFETs. The reasoning behind this choice is explained later. These MOSFETs typically have their $R_{DS(on)}$ values in the range of 1.5 m Ω to 3.5 m Ω , and come in different packages depending on the system space, power, and thermal requirements. As the power starts scaling up, to reduce the thermal stress on the MOSFETs, identical MOSFETs are paralleled together to allow current-sharing between them and reduce the thermal stress. Thermal stress is directly proportional to the $R_{DS(on)}$ of the FETs, and the parallel arrangement of FETs is beneficial in this regard, as it divides each FET's $R_{DS(on)}$ by the number of MOSFETs in parallel.

2.4 Reduction of the inrush current by a hot-swap controller at startup

Fundamentally, a hot-swap controller controls the gate voltage of the MOSFET to regulate the current being supplied to the capacitor at startup. Additionally, the hot-swap controller is responsible for ensuring that the MOSFET does not get damaged in this process and stays within its design limits. The design limits are defined by the SOA curves found in the [datasheet \[2\]](#) of the MOSFET. The SOA curves show the amount of drain current the MOSFET can handle with the applied V_{DS} across it. As an example, the SOA curve of Infineon's OptiMOS™ 5 Linear FET 2 IPT017N10NM5LF2 [3], with improved SOA technology, is shown below:

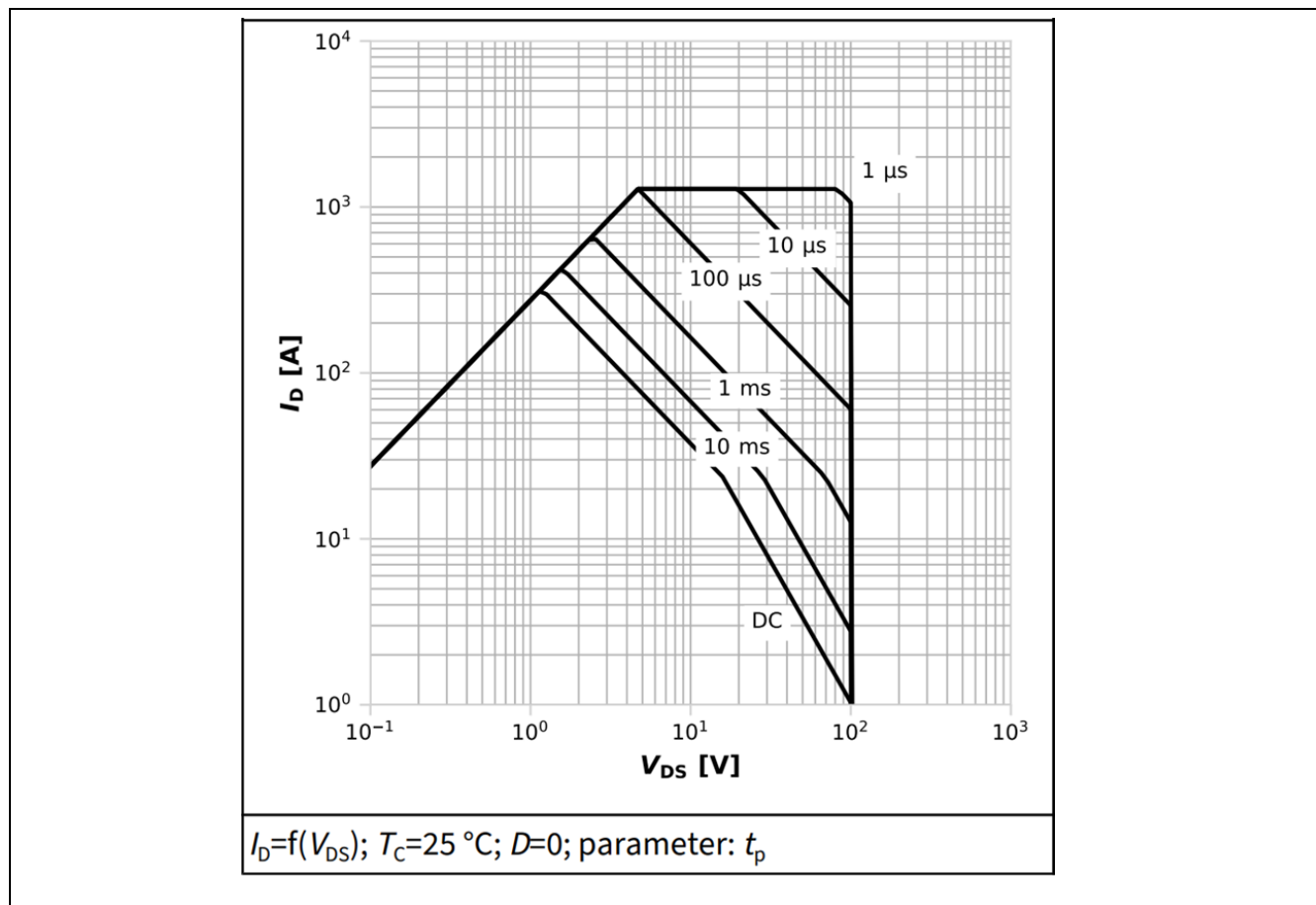


Figure 4 Safe operating area (SOA) curve of OptiMOS™ 5 Linear FET 2 IPT017N10NM5LF2 FET at 25°C

From the SOA curve, we can see that as V_{DS} decreases, the current across the MOSFET (drain current) increases. Also, the current handling capability of the FET increases further if the duration of the current pulse across the MOSFET is reduced. As shown in [Figure 4](#), when the FET is operated on the 1 ms SOA line, the current handling capability increases. For hot-swapping applications, it is generally considered to be safe to design for operation at the DC SOA line.

Note that the SOA curves from [Figure 4](#) are at 25°C but in server applications, the ambient temperature ranges from 45°C to 55°C. Moreover, in cases where the FET operating in normal conditions is turned off for some time due to a fault or some other factors and then turned back on, the instantaneous temperature around the FET will not be at an

ambient temperature, but at a rather high temperature. It is safe to consider the temperature at this point to be close to 85°C to 95°C. In general, the designs are done considering SOA of the FET at 95°C.

As mentioned, the SOA data in a datasheet is measured at 25°C. So, to calculate the approximate data at a different desired temperature, use [Equation 2](#).

$$SOA \text{ Current } (@T_{a2}) = \frac{SOA \text{ Current } (@25^{\circ}C) \times (T_{jmax} - T_{a2})}{(T_{jmax} - 25)}$$

$$SOA \text{ Current } (@95^{\circ}C) = \frac{SOA \text{ Current } (@25^{\circ}C) \times (175 - 95)}{(175 - 25)}$$

Equation 2 MOSFET SOA current extraction at desired junction temperature

The data calculated for OptiMOS™ 5 Linear FET 2 IPT017N10NM5LF2 at 95°C on the DC SOA is shown in [Figure 5](#).

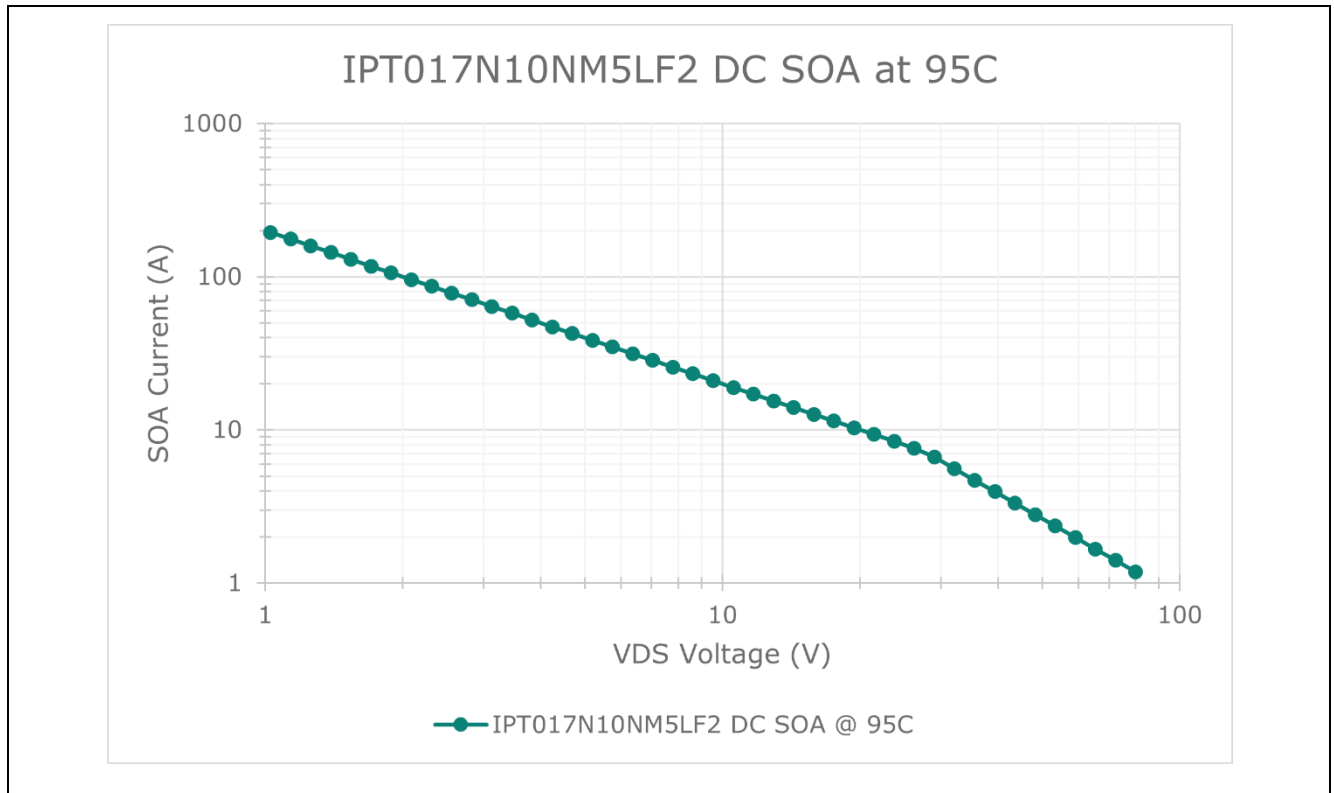


Figure 5 Calculated SOA curve of OptiMOS™ 5 Linear FET 2 IPT017N10NM5LF2 at 95°C

The typical operating voltage range of AI and datacom backplane is in between 40 V to 60 V, so it is essential to consider the SOA current of a FET from $V_{DS} = 60$ V down to 1 V.

2.5 The importance of the SOA curve for hot-plugging

When a server blade is first plugged into the backplane, the input voltage is V_{IN} , i.e., the drain voltage is V_{IN} and the FET is fully off (cut off region) and the source voltage is 0 V and thus, $V_{DS} = V_{IN}$. Let us consider the input voltage, $V_{IN} = 50$ V which puts V_{DS} also at 50 V. The maximum current that the MOSFET can sustain is close to 2.6 A at $V_{DS} = 50$ V as shown in [Figure 5](#). If the current in this condition goes above the specified limit, it can cause significant damage to the MOSFET. So it is essential for the controller to regulate the V_{GS} voltage to keep the drain current at or below the 2.6 A level.

As the current starts flowing through the MOSFET, the input capacitor C_{IN} starts charging. In other words, the voltage of the capacitor (voltage at the source of the MOSFET) starts increasing, which reduces the V_{DS} voltage. Hence, the MOSFET can be operated at an even higher SOA current.

3 Introduction to digital hot-swap controllers

To keep up with the growing power and reliability needs of an AI server, the hot-swap controllers are also evolving in today's market. The most notable evolution is that of the digital hot-swap controllers, where the entire SOA profile of the FET can be programmed, ensuring that the FET stays within its safe operating region at all times, which improves the reliability and lifetime of the overall system.

The algorithm of such a digital hot-swap controller works in the following way:

1. The SOA current profile of the MOSFET is programmed inside the controller from $V_{DS} = 80\text{ V}$ down to $V_{DS} = 1\text{ V}$. Mainly, the DC SOA profile is programmed inside the controller at the desired temperature.
2. At the event of hot-plugging, the FET's V_{DS} is the same as V_{IN} . So the controller refers to the internal lookup table to find the corresponding permissible drain current and sets that value as the target current through the FET (drain current).
3. The gate voltage (V_{GS}) is then ramped up slowly by the controller and the current flowing through the FET is measured. When this current reaches the set target current, the controller regulates the gate voltage in order to maintain the programmed current level.
4. As the current starts flowing through the FET, the capacitor C_{IN} starts charging, increasing the output voltage (i.e., the voltage on the source pin of the FET). This reduces the V_{DS} voltage across the FET.
5. Once the V_{DS} is reduced, the controller adjusts the target regulation current and the gate voltage is adjusted accordingly to reach a new target current.
6. Once the capacitor C_{IN} is charged to the input voltage and the inrush event is completed, it is safe to pass the load current through the FET. The hot-swap controller at this instant sends out a Power Good signal to the power converter, instructing that it is now safe to turn on.
7. This power converter then powers up the processor and puts the AI server blade online.

3.1 An example of a 4 kW hot-swap solution

An example of a 4 kW hot-swap solution is shown in Figure 6 which has been used in power delivery boards. The solution consists of four OptiMOS™ 5 Linear FET 2 IPT017N10NM5LF2 FETs and a XDP™ XDP710-002 hot-swap controller [4]. The region of operation for this board is in the voltage range from 40 V to 60 V. Hence, the nominal load current for 4000 W will be $4000/40 = 100\text{ A}$, which is why four FETs were selected to be operated in parallel for this design, after thermal calculations. As discussed in the previous section, the SOA profile of the FET at 95°C is programmed into the controller.

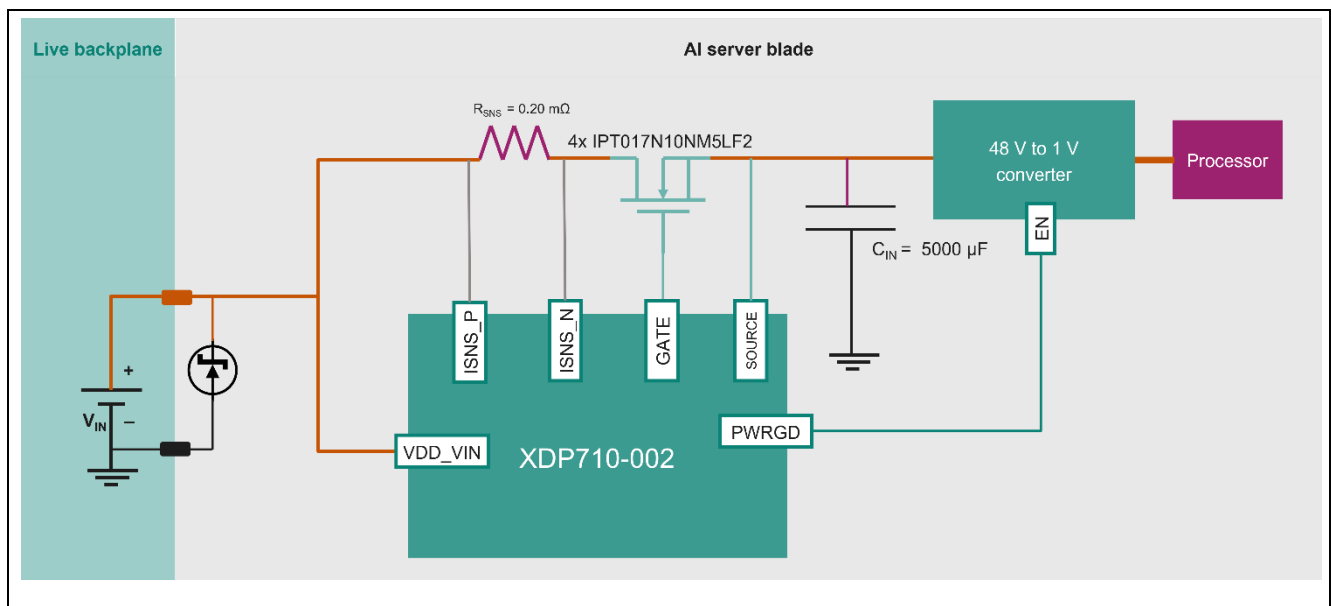


Figure 6 4 kW hot-swap digital controller solution for a 5000 μF capacitance

Note that a weaker SOA profile is programmed into the controller to allow for more margin in the system, as shown in Figure 7.

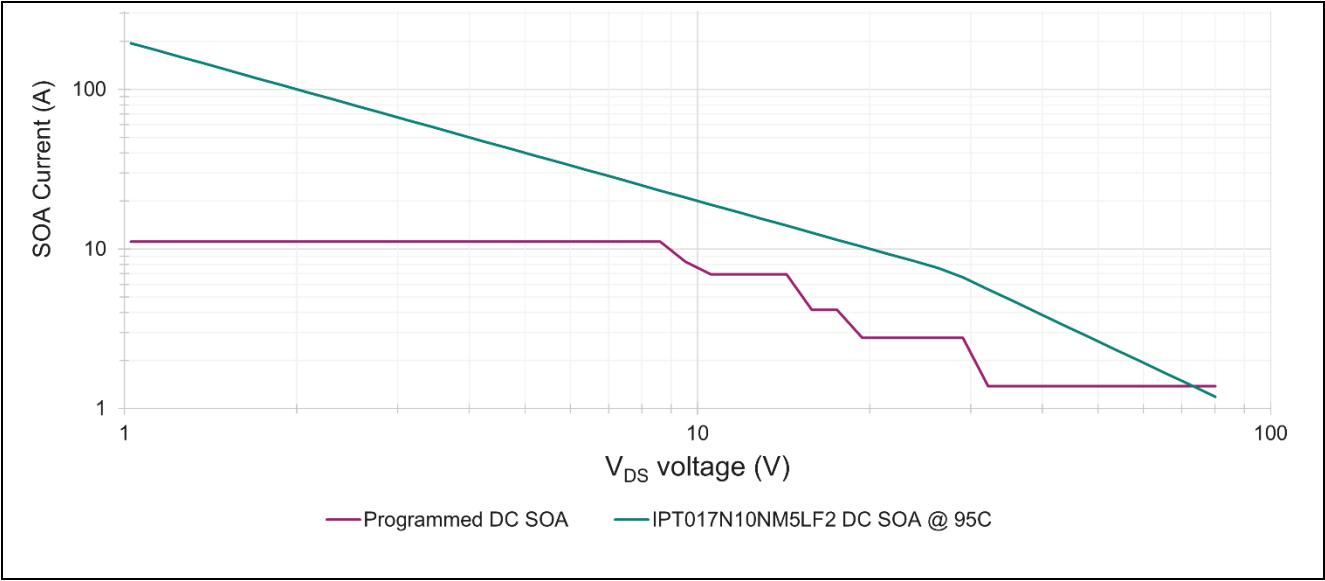


Figure 7 SOA profile programmed in the XDP™ XDP710-002 versus the FET's SOA capability

It is also important to note that the worst case is always considered. During the inrush phase, i.e., when the multiple FETs in parallel are operated in the linear region, the entire inrush current is flowing through a single FET and the other FETs are still operating in the cut-off region. This condition is quite common as the gate turn-on thresholds of FETs connected in parallel vary depending on the placement as well as the parasitics and the FET's intrinsic properties. Thus, the FET with the lowest gate turn-on threshold will get into the saturation region and the other FETs stay in the cut-off region.

The turn-on waveform in Figure 8 shows that the inrush current has been kept below the level of 1.5 A at the instant of time where $V_{DS} = 60$ V, and the gate is regulated to maintain the current through the FET at a constant level. The output voltage rises gradually, which reduces the V_{DS} voltage across the FET, enabling the controller to let more current flow through the FET by raising the gate voltage. At all times, the FET is operated at its safe level, based on the SOA profile programmed in the controller. This 4 kW hot-swap controller design is designed to handle up to 105 A of current in regular operation, and the inrush for this system is clipped below 12 A by the controller. As the controller is digital, the programmed SOA can be fine-tuned more conservatively or aggressively to adjust the inrush time, i.e., the time for output voltage to rise and for the maximum current allowed to flow across the FET at the time of turn-on.

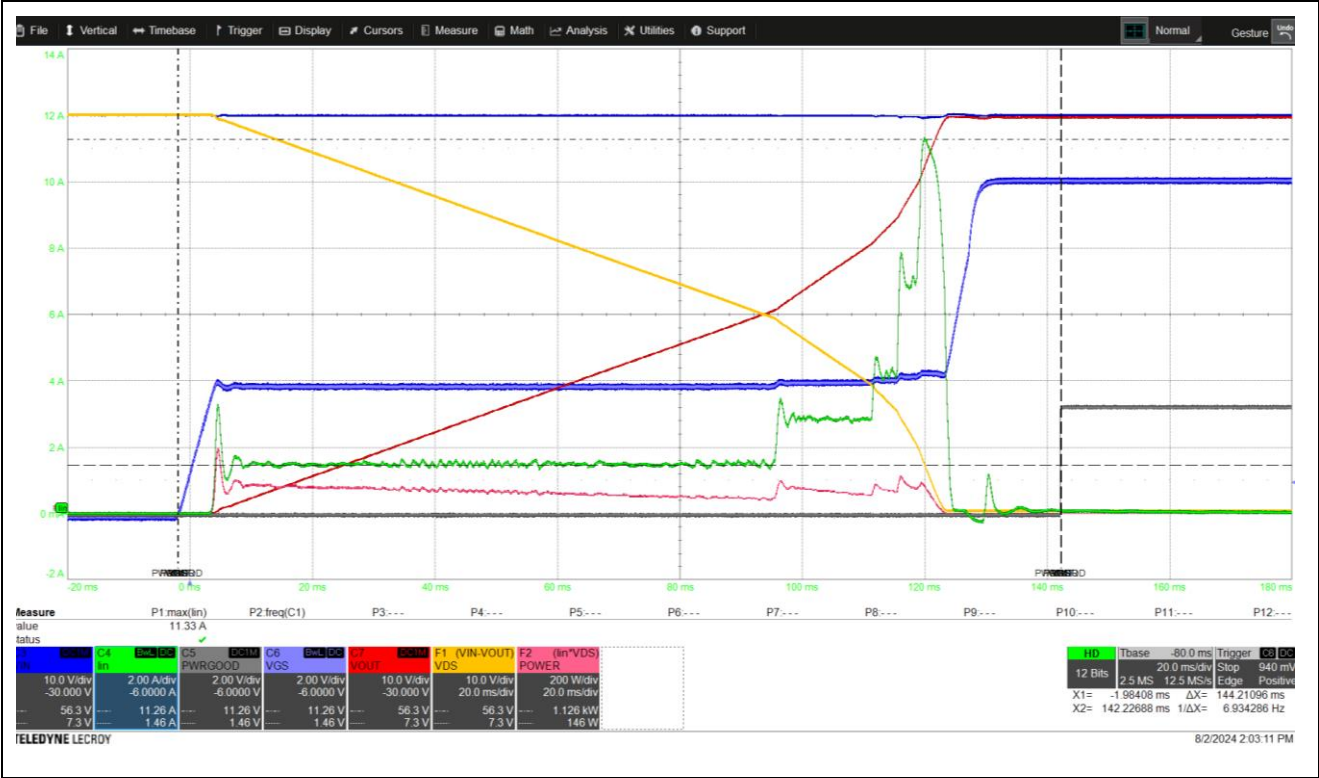


Figure 8 Typical turn-on waveform with XDP™ XDP710-002 charging a 5000 μ F capacitance

3.2 Use of burst mode to turn on the FET instead of continuous mode

While operating on the DC SOA line of the FET, it can be observed that power is continuously dissipated across the FET equal to $V_{DS} \times I_D$. This can become a reliability concern for some applications when the output capacitor size starts increasing further, which increases the overall time the MOSFET stays in the linear region. To overcome this issue, the 1 ms SOA line of the MOSFET can be used, where the FET will be turned on for only 1 ms to charge the capacitor (C_{IN}), turned off for some time to cool down the FET, and then turned back on again for 1 ms to charge the capacitor more.

In this manner, during the turn-on phase of the FET, it is not feeding the load current but merely charging the input capacitor (C_{IN}) of the server blade, so that the energy supplied to charge the output capacitor is not lost. The controller can then slowly charge the output capacitor by providing some time for the MOSFET to cool down. As shown in the turn-on waveforms in [Figure 9](#), the same 4 kW design as before is operated in burst mode. Here, the FETs are operated at the 1 ms SOA line during the high V_{DS} region. When V_{DS} is down to a low value, the operation transitions to continuous mode, compensating for the lost turn-on time in burst mode. This method provides ample time for the MOSFET to cool down (approximately 9 ms in this example), which increases the reliability of the entire system.

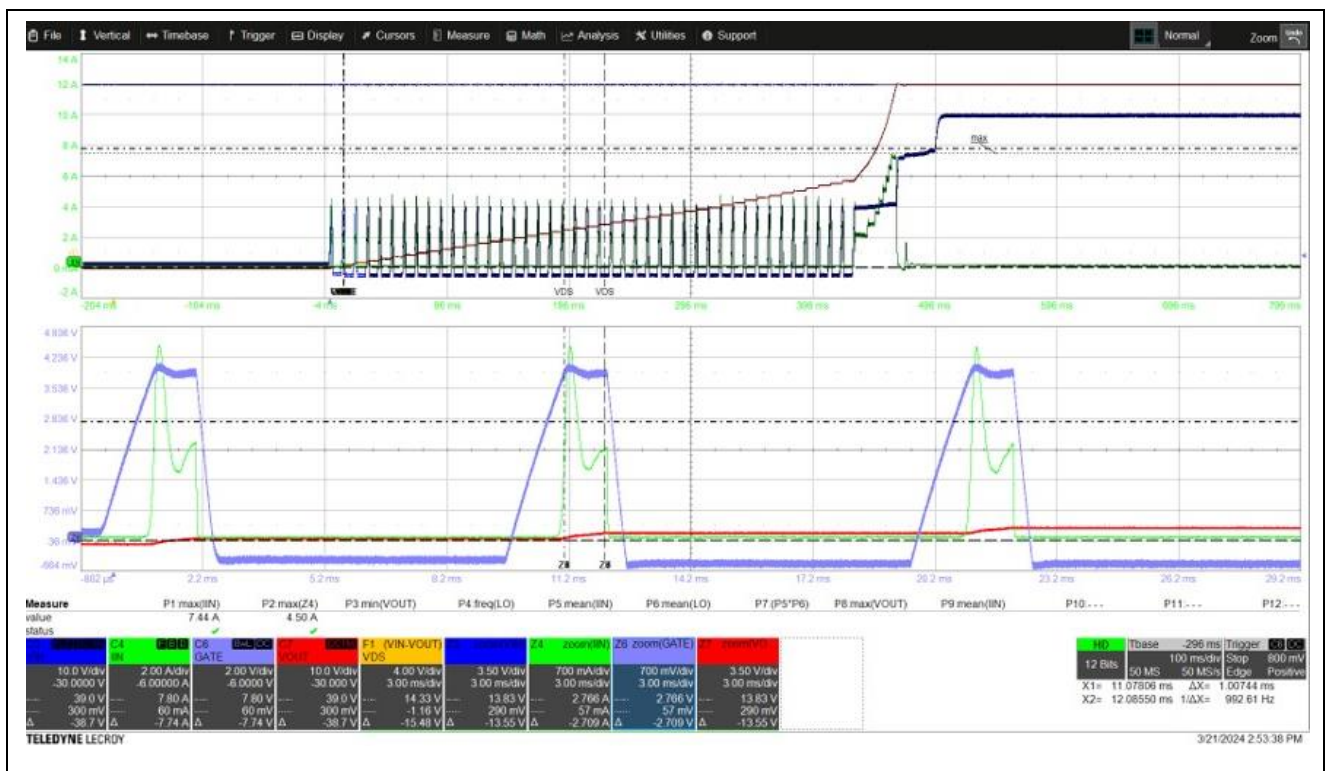


Figure 9 Turn-on with burst mode

3.3 Eliminating the effect of contact bounce during hot-plugging

Whenever the server blade is hot-plugged into the backplane, it is done as a mechanical connection. It is always difficult to make a connection of the input to the server blade without any bouncing. This can be visualized with an example of trying to plug in a laptop's battery charger into an AC socket, where sometimes a small arcing can be observed due to the bouncing of the input connection.

The bouncing during hot-plugging can be prevented by providing the gate pulse to the inrush FET after the server blade has been plugged-in securely into the backplane. Most applications make use of a voltage-sense pin for this purpose. This voltage-sense pin connects to the same backplane, but it is the last pin to make a connection with the backplane. It is also connected to the “Enable” pin of the hot-swap controller. Only when the voltage on the “Enable” pin is above a certain threshold for a programmed amount of time, the hot-swap controller will turn on the inrush FETs. If not, C_{IN} stays disconnected from the backplane. The connections are shown in [Figure 10](#).

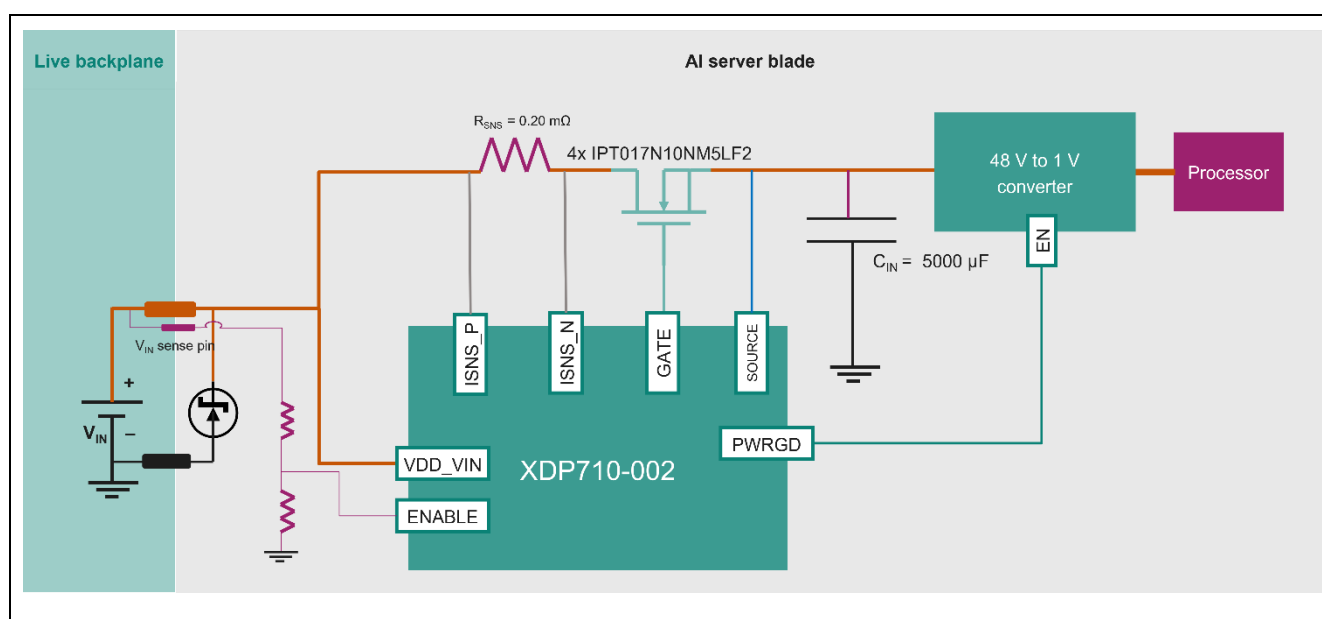


Figure 10 “Enable” pin connection to avoid bouncing during hot-plugging

3.4 Protecting the server blade and backplane during fault events

It is essential to protect the backplane also from the fault events that can occur inside the server blade, as they can potentially cause significant damage to the other servers connected in parallel. In such events, the fault should be detected and the server blade should be isolated from the backplane as quickly and safely as possible by opening the inrush FETs. At the same time, if any external factors such as when the backplane's voltage quality varies or a surge event occurs, it is equally important to protect the server blade from these events. This timely detection and isolation of the faults can save millions of dollars for a server farm in the long run and keep the system healthy and reliable.

Various fault protections are offered by the hot-swap controller, namely, short-circuit current, overcurrent, overvoltage, undervoltage, overtemperature, and FET faults protection. These protection features are essential to keep the system healthy by isolating the faulty power modules from the backplane in time without causing any damage to it, at the same time isolating the power modules from any potential power quality issue in the backplane. Also, when a fault occurs, the system is alerted about the fault via dedicated fault pins and the type of fault can be identified by reading the fault register via the PMBus protocol. These faults can be self-clearing, which turns the FET back on when the fault condition is cleared or can be latched off, which essentially needs an intervention from the host system to turn the server blade back on. Lastly, the controller can be configured in an auto-retry mode where the hot-swap controller attempts to turn the system on repeatedly. If the fault condition does not exist after a turn-on from fault, it keeps the system running; otherwise it turns the system off and tries again.

Notable here is the short-circuit detection, as this fault event can occur because of dropping a screwdriver or screw during maintenance, but is the most catastrophic of all the fault events. It can cause significant damage to the backplane power quality, modules connected in parallel, and the AI server blade itself. Thus, it is important to detect and isolate this as soon as possible. With today's controllers, isolation detection takes less than 1 μ s using fast comparators and strong gate pull-down to take the FET into the cut-off region. But this fast isolation gives rise to another issue of a high input voltage spike. When the short-circuit event is isolated so fast, the input current is changed from several hundreds of amperes to zero amperes in a few nanoseconds. Since the server rack backplane is connected with long cables (several feet long) to feed the power in, these cables add an inductance in the power path. So when the current transitions so fast from a very high value to a low value, it induces a high voltage spike on the backplane. This voltage can go easily above 100 V and thus a TVS diode is placed before the MOSFETs to clamp (or shunt the excessive energy) the excessive voltage. But this clamping is far higher than the 60 V nominal operating range, being closer to 80 V – 85 V. For this reason, the FETs used in hot-plugging are rated for 100 V.

To compliment the fault protection, warning alerts can also be set which will alert the system beforehand to make necessary adjustments in order to avoid entering fault conditions. This prevents the server blade from being disconnected from the backplane and thus providing uninterrupted supply.

3.5 Active monitoring and health reporting of the system

The active monitoring module of the hot-swap controller is responsible for providing accurate, real-time telemetry of input voltage, output voltage, input current, power, temperature, and energy along with the status of faults and warnings. All this information can be read by the host using the PMBus communication protocol, which is accepted industry-wide. It is essential to have <1 percent telemetry accuracy error to ensure that the system is operating within the desired limits. Moreover, the feature of capturing the peaks and valleys helps identify potential fault events.

From improving the reliability of the system and the ability to identify the root cause of a fault, to deploying a potential fix to eliminate the reason for a fault, the industry has been foreseeing the necessity of having a black-box feature in the hot-swap controller. The black-box feature will record telemetry data before, during, and after a fault event, which can later be accessed from the controller's memory for analysis.

4 Future trends

As the volume due to market demand increases, while higher performance is simultaneously expected, system designers are facing continued challenges to meet these requirements. For a robust and compact design, the industry is now looking for an integrated solution, i.e., eFuse. eFuses integrate linear MOSFETs, hot-swap controllers, current-sensing shunts, and temperature sensors together into one package, which reduces the size of the hot-swap solution considerably. For a 1.2 kW system, one eFuse in a footprint of 7 mm x 7 mm offers a very compact solution compared to a discrete hot-swap solution. These eFuses can be paralleled together for high power designs, since the controller, FET, and temperature sensors are integrated together under one package. This ensures that at startup, the current is shared reliably amongst the eFuses unlike a typical a hot-swap controller implementation where a single gate driver drives several MOSFETs in parallel.

Considering the future performance and reliability demands of AI servers, an eFuse provides an even more reliable and efficient solution with integrated temperature sensor placed close to the MOSFET, which can provide the real-time die temperature of the MOSFET without any latency. This ensures that if the MOSFET's temperature exceeds the programmed limit at any time, i.e., during turn-on or regular operation, the MOSFET can be turned off without causing any damage.

It is expected that the power consumption of the AI server blade will increase further to 8 kW or even 12 kW, which trends towards higher voltage of the backplane, rising to 400 V. This would require a new class of hot-swap controllers and MOSFETs to support this requirement. It is equally important that the solutions are even more robust and reliable, as 400 V DC is some serious voltage, and a failure at this voltage could be catastrophic.

References

- [1] Forbes: *The True Cost Of Downtime (And How To Avoid It)*; [Available online](#)
- [2] Infineon Technologies AG: *Datasheet - OptiMOS™ 5 Linear FET 2 IPT017N10NM5LF2*; [Available online](#)
- [3] Infineon Technologies AG: *OptiMOS™ 5 Linear FET 2 IPT017N10NM5LF2, single N-channel Linear FET 100 V, 1.7 mΩ, 321 A in TOLL package*; [Available online](#)
- [4] Infineon Technologies AG: *χDP™ χDP710-002 hot-swap controller*; [Available online](#)

Published by
Infineon Technologies AG
Am Campeon 1-15, 85579 Neubiberg
Germany

© 2024 Infineon Technologies AG.
All rights reserved.

Public

Version: V1.0_EN
Date: 11/2024



Stay connected!



Scan QR code and explore offering
www.infineon.com

Please note!

This Document is for information purposes only and any information given herein shall in no event be regarded as a warranty, guarantee or description of any functionality, conditions and/or quality of our products or any suitability for a particular purpose. With regard to the technical specifications of our products, we kindly ask you to refer to the relevant product data sheets provided by us. Our customers and their technical departments are required to evaluate the suitability of our products for the intended application.

We reserve the right to change this document and/or the information given herein at any time.

Additional information

For further information on technologies, our products, the application of our products, delivery terms and conditions and/or prices, please contact your nearest Infineon Technologies office (www.infineon.com).

Warnings

Due to technical requirements, our products may contain dangerous substances. For information on the types in question, please contact your nearest Infineon Technologies office.

Except as otherwise explicitly approved by us in a written document signed by authorized representatives of Infineon Technologies, our products may not be used in any life-endangering applications, including but not limited to medical, nuclear, military, life-critical or any other applications where a failure of the product or any consequences of the use thereof can result in personal injury.