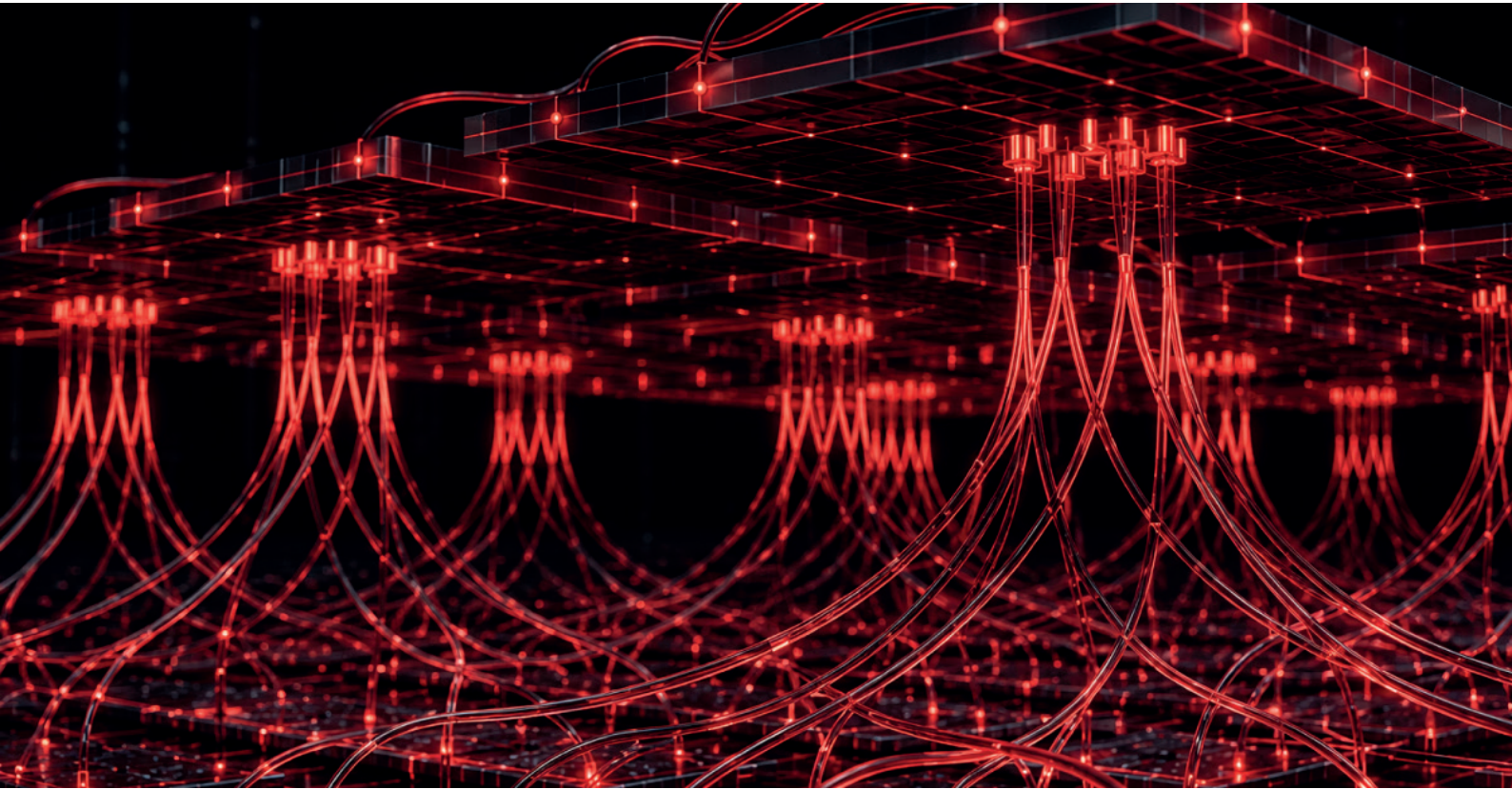




— TEIL 1

Grundlagen und Anforderungen

Verfügbarkeitsklassen, Risikoanalyse,
Redundanz und Stromversorgung.

Fabian Rieke

Für den Betrieb eines Rechenzentrums zählt vor allem eines: Es muss jederzeit laufen. Mit dem zunehmenden Einsatz künstlicher Intelligenz (KI) wachsen Leistungsdichte und Anforderungen, und damit die Verantwortung des Betreibers für Ausfallsicherheit und Rechtssicherheit. Fällt ein Rechenzentrum aus, drohen nicht nur Datenverluste, sondern auch Stillstand, wirtschaftlicher Schaden und Haftungsfragen.

Wie verfügbar ein Rechenzentrum sein muss, ist dabei keine Bauchentscheidung, sondern lässt sich klassifizieren, bewerten und planen. Dieser erste Teil legt die Grundlagen dafür, von der Klassifizierung über die Risikoanalyse bis zur Stromversorgung.

Klassifizierung: Wie sich die Verfügbarkeit eines Rechenzentrums messen lässt

Für die Planung sind drei Faktoren entscheidend: der Zweck des Rechenzentrums, das Sicherheitsniveau und die physische Größe. Beim Zweck wird zwischen Unternehmens-, Colocation-, Hosting- und Netzbetreiber-Rechenzentrum unterschieden.

Der Maßstab, an dem ein Rechenzentrum gemessen wird, ist seine Verfügbarkeit, also die Ausfallsicherheit und Erreichbarkeit aus Sicht des Nutzers. Der europäische Standard teilt sie in die Verfügbarkeitsklassen VK 1 bis 4 ein, der amerikanische Standard in die Kategorien Tier I bis IV.

Bei der Einstufung werden unter anderem folgende Parameter berücksichtigt:

- die durchschnittliche Erreichbarkeit bzw. Jahresverfügbarkeit,
- die Redundanz der Versorgungswege (z. B. Energie- und Kälteversorgung),
- die Möglichkeit zur Wartung im laufenden Betrieb.

Darüber hinaus wird der Single Point of Failure (SPOF) betrachtet, also ein einzelner Fehlerpunkt, dessen Ausfall den Gesamtbetrieb lahmlegen kann. Die Fehlertoleranz bewertet mögliche Szenarien, etwa ein Brandereignis oder Wärmestau bei Ausfall der Allgemeinen Stromversorgung.

Die Ausfallzeit wird auf Basis von 8.760 Stunden im Jahr berechnet. Für die Verfügbarkeitsklasse 1 bzw. Tier I wird eine Verfügbarkeit von 99,671 % angestrebt, was einer tolerierbaren Ausfallzeit von maximal 28,8 Stunden pro Jahr entspricht. Bei VK 4 bzw. Tier IV liegt der Zielwert bei 99,995 %, also einer maximalen Ausfallzeit von rund 0,4 Stunden pro Jahr (rund 26 Minuten).

Zur Einordnung: Die Tier-Klassifizierung (Tier I bis IV) ist ein Standard des Uptime Institute, einer US-amerikanischen Organisation für die Ausfallsicherheit von Rechenzentren ^[1]. Die **DIN EN 50600** ist dagegen eine europäische Norm, die die Infrastruktur von Rechenzentren systematisch beschreibt ^[2]. Sie wurde inhaltlich in die internationale Normenreihe **ISO/IEC 22237** überführt, bleibt in Deutschland aber in Gebrauch. Beide Systeme sind allerdings nicht deckungsgleich. Die genannten Prozent- und Stundenwerte stammen aus der Tier-Klassifizierung, denn die **DIN EN 50600** ordnet ihren Verfügbarkeitsklassen keine festen Prozentwerte zu. Die Zuordnung von VK 1 bis 4 zu Tier I bis IV dient daher nur der Orientierung.

Risikoanalyse: Die passende Verfügbarkeitsklasse bestimmen

Als Entscheidungshilfe, welche Verfügbarkeitsklasse die richtige ist, sollten Sie im Vorfeld eine Risikoanalyse durchführen. Als Grundlage dafür kann die **DIN EN IEC 31010 (VDE 0050-1):2024-12** (Vorgängerausgabe **DIN EN 31010:2010-11**) herangezogen werden ^[3].

Hilfreich ist eine Risikomatrix. Auf der Y-Achse tragen Sie die Auswirkung ab, also die Stärke oder Schwere negativer Vorfälle. Das kann ein Geldbetrag sein oder die erwartete Dauer des Verlustes der Verfügbarkeit eines Dienstes. Auf der X-Achse steht die Eintrittswahrscheinlichkeit, also die Wahrscheinlichkeit, dass das jeweilige Ereignis eintritt.

In diese Matrix gehören alle relevanten Faktoren, etwa Schäden durch Ausfall der Netzspannung, Vandalismus oder die geografische Lage. Dazu zählen auch Risiken wie Hochwasser, Gefahrstofflager in benachbarten Industrieanlagen oder stark frequentierte Verkehrswege.

Sind die möglichen Ereignisse im Kontext Ihres Rechenzentrums bewertet, ermitteln Sie die Kosten der jeweiligen Ausfallzeiten. Daraus leiten Sie Maßnahmen ab, die in einem sinnvollen Verhältnis von Aufwand und Nutzen stehen.

Hinzu kommen mögliche gesetzliche Anforderungen. Sie greifen, wenn das Rechenzentrum als kritische Infrastruktur gilt, etwa bei Polizei- oder Feuer-

wehrendienststellen, Flughäfen und Flugsicherungen, Anlagen zum Bevölkerungsschutz. Für Rechenzentren ist der regulatorische Rahmen zuletzt deutlich gewachsen: Das NIS2-Umsetzungsgesetz (in Kraft seit Dezember 2025) verschärft die Anforderungen an die Cybersicherheit, das KRITIS-Dachgesetz (vom Bundestag im Januar 2026 angenommen) ergänzt sie um die physische Resilienz kritischer Anlagen. Zuständig sind das BSI für die IT-Sicherheit und das Bundesamt für Bevölkerungsschutz und Katastrophenhilfe (BBK) für den physischen Schutz.

Redundanzstufen: N, N+1 und 2N im Überblick

Für den Aufbau einer redundanten Energie- und Kälteversorgung sind die Stufen „N“, „N+1“ und „2N“ entscheidend.

„N“ beschreibt die erforderliche Anzahl der Komponenten, etwa USV oder Kälteanlage, sowie die Kapazität bei voller IT-Auslastung.

Bei einer Auslegung nach „N+1“ kommt eine zusätzliche Komponente hinzu, um den Ausfall oder die Wartung einer Anlageneinheit zu kompensieren. Häufig wird dieses Vorgehen bei modularen USV-Anlagen angewendet. Ein Beispiel: Benötigt eine Anlage eine Systemleistung von 28 kW und stellt ein Modul 20 kW bereit, werden statt zwei Modulen drei verbaut, um Ausfall oder Wartung abzufangen.

Dasselbe Prinzip findet sich in unserer Energieversorgung wieder, um Netz- und Versorgungssicherheit zu gewährleisten. Komponenten wie Transformatoren oder Leitungen werden in der Regel nicht am Grenzwert ihrer maximalen Auslastung betrieben, sondern mit einer Reserve, die kurzfristige Lastschwankungen und einzelne Ausfälle abfedert.

Eine Auslegung nach „2N“ bedeutet dagegen ein vollständig redundantes, gespiegeltes System mit zwei unabhängigen Verteilungssystemen. Selbst beim vollständigen Ausfall eines Systems übernimmt das redundante System die Last vollständig. Ausfallzeiten sinken dadurch deutlich.

- N: Mindestanforderung. Alles läuft genau mit der benötigten Kapazität, ohne Reserve.
- N+1: Mindestanforderung plus eine zusätzliche Komponente als Reserve.
- 2N: Zwei vollständig unabhängige Systeme mit jeweils der Kapazität „N“. Maximaler Ausfallschutz.

Stromversorgung: Primär, Sekundär und

Notstrom

Bei der Versorgung wird zwischen Primär-, Sekundär- und Notstromversorgung unterschieden. Primär- und Sekundärversorgung bezeichnen die Einspeisung über den bereitgestellten Haus- bzw. Mittelspannungsanschluss.

Die Notstromversorgung gliedert sich in zwei Bereiche: Sicherheitsstromversorgung und Ersatzstromversorgung.

Die Sicherheitsstromversorgung umfasst elektrische Anlagen und Stromquellen, die aufgrund baurechtlicher Anforderungen eingerichtet werden müssen, etwa im Rahmen eines Brandschutzkonzepts. Diese Anforderungen beruhen im Wesentlichen auf den Landesbauordnungen der Bundesländer. Die konkreten Forderungen stellen in der Regel die Bauaufsichtsbehörden im Rahmen der Baugenehmigung.

Bei der Versorgung der IT-Bereiche handelt es sich in der Regel um eine Ersatzstromversorgung, da die weiter betriebenen Bereiche keine baurechtlich sicherheitsrelevanten Einrichtungen im Sinne der Bauordnung sind. Diese Anlagen dienen in erster Linie dazu, wirtschaftliche Schäden durch längere Ausfallzeiten abzumildern.

Stromverteilung: Unterbrechungsfrei umschalten zwischen Stromquellen

Bei der Stromverteilung sorgen Umschaltssysteme dafür, dass kritische Lasten bei einem Ausfall auf eine andere Quelle wechseln. Zwei Technologien sind gebräuchlich, der automatische Lastumschalter (ATS) und das statische Transfersystem (STS).

ATS (Automatic Transfer Switches)

Ein automatischer Lastumschalter (ATS) schaltet die Versorgung automatisch von der ausgefallenen Hauptstromquelle auf eine alternative Quelle um, etwa auf ein Netzersatzaggregat. Dazu überwacht er fortlaufend Spannung und Frequenz der Primärquelle und leitet die Umschaltung ein, sobald diese ausfällt oder die Grenzwerte verlässt. Typische Einsatzfälle sind Netz-/Netz- und Netz/Generator-Anwendungen ^[4].

Der ATS arbeitet elektromechanisch. Die Umschaltung ist deshalb mit einer kurzen Unterbrechung der Versorgung verbunden und für Systeme ausgelegt, in denen eine solche Unterbrechung tolerierbar ist. In Rechenzentren wird diese Lücke in der Regel durch eine USV überbrückt. ATS arbeiten dabei

vorgelagert und sichern die Redundanz zwischen Netzeinspeisung und Notstromaggregaten ab. Als Bezugsnorm gilt die **DIN EN 60947-6-1 (VDE 0660-114)**.

STS (Statisches Transfersystem)

Statische Umschaltssysteme (STS) schalten kritische Lasten unterbrechungsfrei von der ausgefallenen Stromquelle auf eine alternative Quelle um^[5].

Im Gegensatz zum ATS verwenden STS-Halbleiter wie Thyristoren, um zwischen zwei Stromquellen umzuschalten. Die Umschaltung erfolgt praktisch verzögerungslos in wenigen Millisekunden. Diese Geschwindigkeit ist für kritische Anwendungen entscheidend, die selbst extrem kurze Unterbrechungen der Stromversorgung nicht tolerieren. Daher eignet sich das STS besonders für Bereiche, in denen die Kontinuität der Stromversorgung höchste Priorität hat, etwa im Bank- und Finanzwesen, im Gesundheitswesen oder in Rechenzentren. Als Bezugsnorm gilt die **DIN EN 62310** (Static transfer systems), die sich von der **DIN EN 60947-6-1** für elektromechanische ATS abgrenzt.

STS lassen sich auch direkt in die Racks eines Rechenzentrums integrieren. Sie sind eine kompakte und effiziente Lösung für das Energiemanagement. So bleibt die Zuverlässigkeit der Infrastruktur gewährleistet, und der verfügbare Platz wird optimal genutzt. Im Zusammenspiel sichern STS nachgelagert die unterbrechungsfreie Versorgung von Sammelschienen und Power Distribution Units (PDU), während ATS die Quellen vorgelagert umschalten.

Fazit und Ausblick

Die beste IT-Technik bleibt wirkungslos, wenn die versorgende Infrastruktur nicht fachgerecht geprüft und instandgehalten wird. Genau hier setzt unsere Arbeit an. Auf Basis unserer praktischen Erfahrung zeigen wir als Ensmann Consulting, wie sich die Supportbereiche eines Rechenzentrums prüfen lassen, ohne die IT-Komponenten außer Betrieb nehmen zu müssen. Wenn Sie wissen möchten, wie ausfallsicher die versorgende Infrastruktur Ihres Rechenzentrums ist, unterstützen wir Sie bei Prüfung und Bewertung. Sprechen Sie uns an.

Wie die einzelnen Versorgungswege zusammenspielen und welche Topologie welche Verfügbarkeit ermöglicht, zeigt der zweite Teil dieser Serie: Netzaufbau und Topologie.

Quellen & Verweise

^[1] Uptime Institute, Tier-Klassifizierung (Tier Standard) für die Ausfallsicherheit von Rechenzentren.

^[2] DIN EN 50600, europäische Norm für Einrichtungen und Infrastrukturen von Rechenzentren.

^[3] DIN EN IEC 31010 (VDE 0050-1):2024-12, Norm zu Verfahren der Risikobeurteilung (Vorgängerausgabe DIN EN 31010:2010-11).

^[4] Socomec – Automatische Lastumschalter (ATS)

^[5] Socomec – Statische Transfersysteme (STS)

Ensmann Consulting GmbH

Aubachstraße 22
56410 Montabaur

+49 2602 83997-0

info@ensmann.com