

Zeitschrift: Bulletin Electrosuisse
Herausgeber: Electrosuisse, Verband für Elektro-, Energie- und Informationstechnik
Band: 98 (2007)
Heft: 13

Artikel: Wenn eine Anlage nicht ausfallen darf
Autor: Müller, Thomas
DOI: <https://doi.org/10.5169/seals-857454>

Nutzungsbedingungen

Die ETH-Bibliothek ist die Anbieterin der digitalisierten Zeitschriften. Sie besitzt keine Urheberrechte an den Zeitschriften und ist nicht verantwortlich für deren Inhalte. Die Rechte liegen in der Regel bei den Herausgebern beziehungsweise den externen Rechteinhabern. [Siehe Rechtliche Hinweise.](#)

Conditions d'utilisation

L'ETH Library est le fournisseur des revues numérisées. Elle ne détient aucun droit d'auteur sur les revues et n'est pas responsable de leur contenu. En règle générale, les droits sont détenus par les éditeurs ou les détenteurs de droits externes. [Voir Informations légales.](#)

Terms of use

The ETH Library is the provider of the digitised journals. It does not own any copyrights to the journals and is not responsible for their content. The rights usually lie with the publishers or the external rights holders. [See Legal notice.](#)

Download PDF: 16.05.2025

ETH-Bibliothek Zürich, E-Periodica, <https://www.e-periodica.ch>

Wenn eine Anlage nicht ausfallen darf

Hochverfügbare Automatisierungssysteme

«Doppelt genäht hält besser!» Die gängige Phrase meint, dass eine Doppelnaht zwei Teile auch noch zusammenhält, wenn ein Faden reisst. Auf Anlagen übertragen, würde das heissen, dass doppelt ausgelegte Infrastruktur zu einer besseren Verfügbarkeit führt. Aber stimmt diese Annahme wirklich? Um zunächst bei der Naht zu bleiben, können die zusätzlichen Einstiche der zweiten Naht zu einer Schwächung des Materials führen, sodass es entlang dieser reisst. Die richtige Ausführung ist entscheidend. Das Gleiche gilt für hochverfügbare Anlagen.

Die Automatisierung durchdringt immer mehr Bereiche des täglichen Lebens und übernimmt Aufgaben, die bis vor Kurzem undenkbar erschienen. Dazu gehören sicherheitskritische Aufgaben, die eine hohe Verfügbarkeit bedingen. Extreme Anwendungen sind Steuerungen von Systemen, die keinen sicheren Zustand kennen und

Thomas Müller

von denen Leben abhängen können (Fly/Drive by Wire). Aber auch Applikationen, bei denen es nicht direkt um Leben und Tod geht, ein Ausfall also «nur» einen hohen wirtschaftlichen Schaden nach sich ziehen würde, müssen hochverfügbar sein.

Einige Grundbegriffe

Interessiert die Zuverlässigkeit eines Systems, so betrachtet man die Zeitdauer bis zu einem Ausfall. Diese sogenannte MTTF (Mean Time To Fail) gibt die wahrscheinliche Zeitspanne an, in der das System unter gewissen Bedingungen seine Funktionen korrekt ausübt. Wenn wir beispielsweise davon ausgehen, dass eine Reifenpanne bei einer bestimmten Anzahl gefahrener Kilometer pro Tag im Mittel alle 2,7 Jahre auftritt, ein Reifen also eine MTTF von 1000 Tagen hat, so können wir umgekehrt sagen, dass wir bei einer Tagesausfahrt eine Wahrscheinlichkeit von einem Tausendstel haben, dass ein (bestimmter) Reifen defekt wird. Dafür, dass am selben Tag zwei Reifen defekt sind, ist die Wahrscheinlichkeit nur noch 1 Millionstel.

Die Zuverlässigkeit kann durch regelmässige Wartung beeinflusst werden. Um beim Reifenbeispiel zu bleiben: Neue Reifen sind weniger anfällig als alte, zumal der frische Gummi elastischer ist und das dickere Profil schwerer durchstochen werden kann.

Aufgrund der obigen Rechnung könnte jemand auf die Idee kommen, zur Erhöhung der Zuverlässigkeit stets einen defekten Reifen mitzuführen, um damit das Risiko eines Reifendefekts zu mindern. Die Rechnung stimmt aber nur, wenn die Ausfallwahrscheinlichkeiten der Komponenten gleich verteilt und unabhängig sind. Die Gleichverteilung ist bei einem defekten und einem ganzen Reifen natürlich nicht vorhanden. Aber auch wenn wir über ein Nagelbrett fahren, stimmt die Rechnung nicht,

da in diesem Fall die Unabhängigkeit nicht gegeben ist. Kann eine einzelne Ursache (wie hier das Nagelbrett) zu einer Reihe von Ausfällen führen, spricht man von Common Mode Failures.

In der Praxis ist die Verfügbarkeit eines Systems wichtig. Diese gibt an, wie gross der Zeitanteil ist, zu dem das System seine Funktion erfüllt (z.B. 99,99%) oder umgekehrt, wie hoch der Anteil ist, in dem das System nicht läuft (z.B. in Stunden pro Jahr). Für die Verfügbarkeit ist die Zuverlässigkeit ein wichtiger Faktor – aber nicht der einzige. Die Dauer bis zur Wiederinbetriebnahme nach einem Unterbruch, also die Reparaturzeit, ist genauso wichtig. Falls beim obigen Reifenbeispiel die Reparatur einen Tag dauert, ist die Verfügbarkeit 99,9%. Schafft man die Reparatur in einer Viertelstunde, ist sie 99,999%. Da die Reparaturzeit von vielen Faktoren abhängt, sind für die Verfügbarkeit neben den technischen Massnahmen die organisatorischen genauso wichtig.

Als technische Massnahme versucht man bei hochverfügbaren Systemen eine Fehlertoleranz zu erreichen, indem man sie redundant auslegt. Als redundant werden dabei alle Komponenten bezeichnet, die es nur braucht, damit im Fehlerfall die Funktion aufrechterhalten werden kann. Je nach Anforderungen sind in der Praxis unterschiedliche Grade von Redundanz (doppelte oder mehrfache Auslegung von Teilen oder der vollständigen Anlage) üblich.

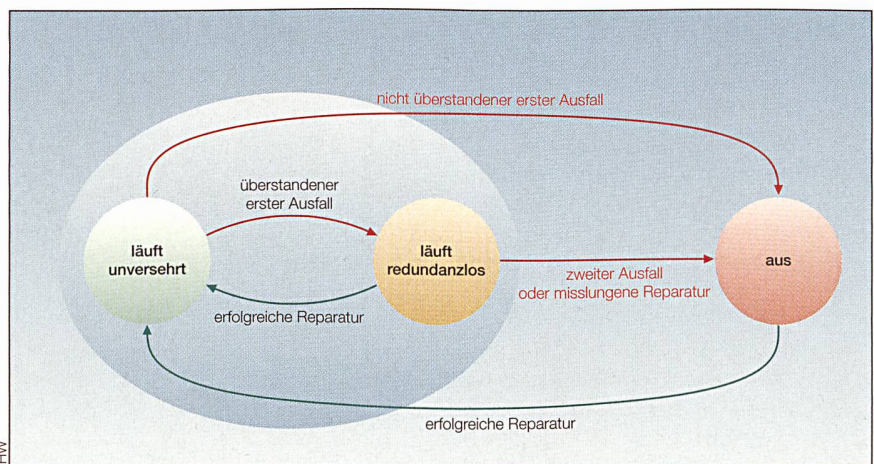


Bild 1 Zustände eines Systems mit einfacher Redundanz.

Es braucht einen Schiedsrichter

Aber was hat eine Verdoppelung aller Komponenten für einen Einfluss? Zunächst verdoppeln sich die Kosten, und da die Zuverlässigkeit der einzelnen Komponenten gleich bleibt, wird auch die Anzahl der Komponentenausfälle im gleichen Masse steigen. Für ein redundantes System benötigt man zusätzliche Teile: Schiedsrichter (Voter), die Fehler erkennen und entscheiden, welches Teilsystem das Sagen hat. Diese kosten ebenfalls und sind natürlich auch störungsbehaftet, wodurch die Bilanz noch etwas schlechter wird.

Wie Bild 1 zeigt, müssen, bezogen auf die Verfügbarkeit, drei Zustände unterschieden werden. Nach dem Anlauf, der hier nicht betrachtet wird, läuft das System zunächst ohne Fehler im Zustand «läuft unversehrt». Fällt jetzt eine Komponente aus, so müssen zwei Fälle unterschieden werden:

Der erste Fall betrifft die Störung aller redundanten Komponenten durch einen Common Mode Failure (Beispiel Nagelbrett) oder aber eine nicht redundante Komponente. Dieser hoffentlich seltene Fall führt dazu, dass das System seine Funktion nicht mehr erfüllen kann (Zustand «aus»). In klassischen Fehlkonstruktionen sind die Voter für solche Ausfälle verantwortlich. Diese komplexen Schaltungen sind architekturbedingt oft nicht redundant. Sind sie nun auch noch unzuverlässig, kann die Verfügbarkeit der Anlage schlechter sein, als wenn sie nicht doppelt ausgelegt worden wäre.

Im zweiten Fall fällt nur eine der redundanten Komponenten aus. Das System erfüllt seine Funktion immer noch, verfügt aber über keine Fehlertoleranz mehr. Fällt in diesem Zustand «läuft redundanzlos» zusätzlich eine Komponente des aktiven Teilsystems aus, geht die Anlage ebenfalls in den Zustand «aus». Die Verfügbarkeit hängt also davon ab, wie schnell das defekte Teilsystem repariert werden kann. Je länger es dauert, bis der Fehler erkannt und behoben wird, desto grösser ist die Chance eines zweiten Ausfalls.

Wer installiert das Ersatzteil?

Hier kommt dem Diagnose- und Alarmsystem eine entscheidende Rolle zu. Aber auch die klarste Fehlermeldung und der beste Alarm nützen nichts, wenn niemand da ist, um diesen zu bearbeiten, oder wenn das erforderliche Ersatzteil nicht verfügbar respektive ebenfalls defekt ist. Die innerbetrieblichen Abläufe und Einsatzpläne sowie die Logistik der Ersatzteile haben einen direkten Einfluss auf die Verfügbarkeit einer Anlage. Wird beim Reparieren versehentlich das aktive Teilsystem beschädigt

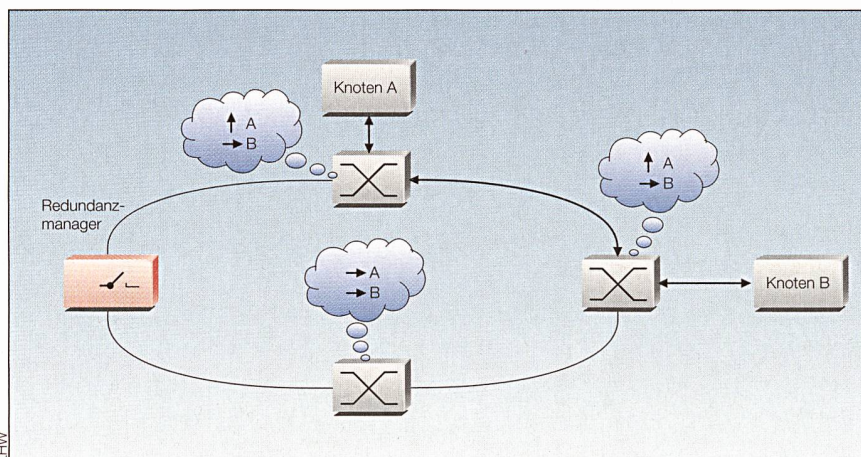


Bild 2 Beim Medium Redundancy Protocol (MRP) schliesst ein Redundanzmanager den Ring.

oder ausgeschaltet, führt dies ebenfalls zum Ausfall der gesamten Anlage. Somit ist klar, dass die Kompetenz des Wartungspersonals einen entscheidenden Einfluss hat.

Systeme, deren Redundanz auf Umschaltung vom aktiven auf ein passives Teilsystem basiert (Stand-by-Lösungen wie beispielsweise ein Reserverad), sind problematisch, da Ausfälle der redundanten Komponenten oft erst bemerkt werden, wenn es zu spät ist (Wann haben Sie das letzte Mal das Reserverad Ihres Autos kontrolliert?). Besser sind funktionsbeteiligte Redundanzsysteme, bei denen sämtliche Teilsysteme im Normalbetrieb gleichzeitig arbeiten (Doppelräder statt Reserverad).

Speziell kritisch sind Anlagen, die hohe Anforderungen an eine kontinuierliche Bedienung durch das Automatisierungssystem stellen. Der längste Unterbruch, den die Anlage tolerieren kann, wird als Gnadenzeit bezeichnet. Diese ist vom Prozess abhängig und reicht von mehreren Sekunden in der Gebäude- oder Prozessautomation bis in den Submillisekundenbereich.

Beispiele dafür sind die Traktionskontrolle von Fahrzeugen sowie Steuerungen von Robotern oder Unterstationen. Während Stand-by-Lösungen immer eine gewisse Umschaltzeit benötigen, ist die funktionsbeteiligte Redundanz im Idealfall unterbrechungsfrei.

In diesen sensiblen Bereichen werden meist eigens für diese Anforderungen entwickelte Steuerungen eingesetzt. Für weniger kritische Anwendungen sind Off-the-Shelf-Lösungen in Form von redundanten speicherprogrammierbaren Steuerungen und Peripherie verfügbar. Bei den verbreiteten Lösungen (z.B. von Siemens) können die Zentraleinheiten mit speziellen Synchronisationsbaugruppen gekoppelt werden. Bei der zweikanalig geschalteten Peripherie wird jedes Sensorsignal auf zwei Eingänge geführt; die Ausgangssignale werden entsprechend über Dioden auf die Aktoren geführt.

Um das Risiko eines Totalausfalls zu minimieren, wird die Intelligenz gerne verteilt, indem mehrere Steuerungen mit dezentraler Peripherie eingesetzt werden. Heutige

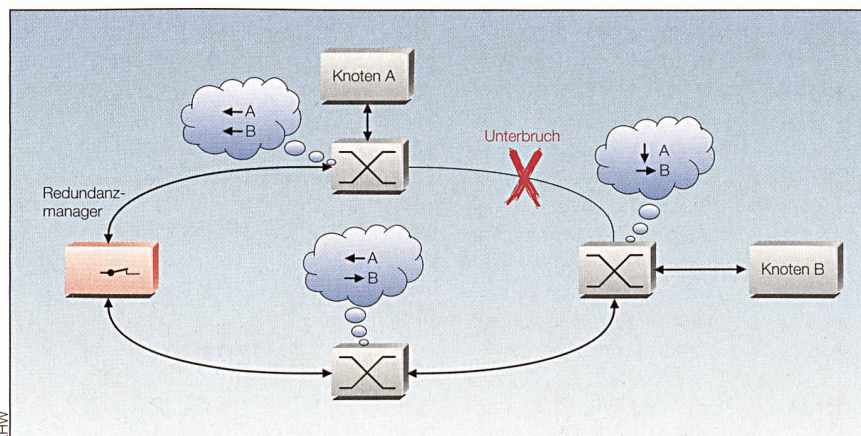


Bild 3 Bei einem Unterbruch im Netz leitet der Redundanzmanager die Datenpakete weiter.

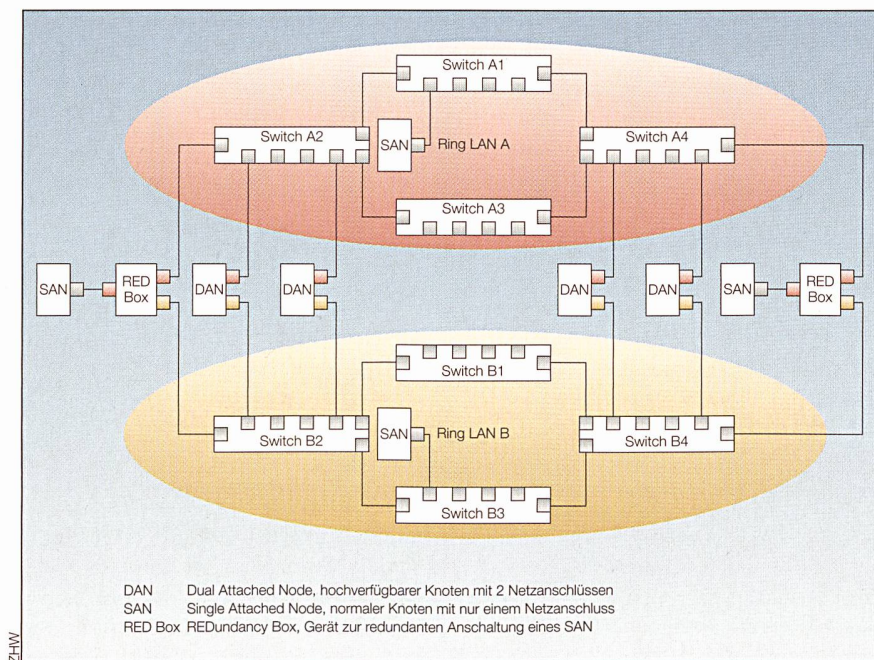


Bild 4 Beim Parallel Redundancy Protocol (PRP) sendet jedes Endgerät über zwei unabhängige Netzwerke.

Automationssysteme sind oft verteilte Systeme, und somit kommt dem Netzwerk eine entscheidende Rolle zu.

Netzwerkredundanz

Während bei der eigentlichen Steuerung die Umschaltzeiten noch relativ leicht erreicht werden können, müssen bei den Netzwerken spezielle Vorkehrungen getroffen werden. Es ist der IEC gelungen, eine Norm für hochverfügbare Automationsnetzwerke zu produzieren, die unabhängig von höheren Protokollen ist.

Die Norm IEC 62439 beschreibt drei Verfahren:

- Das Medium Redundancy Protocol (MRP) verwendet eine Ringtopologie mit 200 ms – 500 ms Umschaltzeit und erlaubt den Anschluss von normalen Geräten ohne doppelte Netzwerkanschlüsse.
- Das Parallel Redundancy Protocol (PRP) basiert auf einem Doppelnetzwerk mit null Umschaltzeit für Geräte mit je einem Anschluss zu den zwei unabhängigen, parallel betriebenen Netzwerken.
- Das Cross-Network Redundancy Protocol (CRP) verwendet ebenfalls ein Doppelnetzwerk, benötigt aber 2 s Umschaltzeit, und Geräte brauchen je einen Anschluss zu den zwei gekoppelten Netzwerken.

Bereits die Norm IEEE-802.1d beschreibt, wie zur Verbesserung der Netzwerkzuverlässigkeit zwischen den Switches redundante Pfade verwendet werden können. Damit die Pakete nicht endlos im Kreis herumgeschickt werden, legt das

Spanning Tree Protocol den Transportweg fest. Dabei sperren die Switches alle redundanten Pfade bis auf einen. Im Fehlerfall wird der Prozess wiederholt und (falls möglich) der ausgefallene Pfad durch einen anderen ersetzt.

Die Fehlererkennung erfolgt durch spezielle Meldungen, die die Switches im Abstand von 2 s versenden. Bleibt eine solche Meldung aus, wird das Netzwerk rekonfiguriert. Während der Rekonfiguration werden alle Ports für den normalen Datenverkehr gesperrt, was zu einem bis 30 s dauernden Unterbruch führen kann.

Wenn zwei Sekunden zu lang sind

30 s sind für die meisten Anwendungen zu lang. Aus diesem Grund wurde das Rapid Spanning Tree Protocol (RSTP) definiert. Dabei wird bei einer Rekonfiguration der Verkehr (soweit möglich) auf den bisherigen Pfaden weitergeleitet, und parallel dazu werden die Alternativpfade bestimmt. Erst wenn der neue Baum bestimmt ist, wird umgeschaltet. Die Unterbruchzeit beträgt mit RSTP typisch 2 s. Da diese Zeit nicht garantiert werden kann (z.B. wenn Edge Ports nicht als solche erkannt werden), wird RSTP in der obigen IEC-Norm nur für unkritische Automationssysteme empfohlen.

Die garantierte Umschaltzeit von maximal 500 ms ist denn auch ein wesentlicher Vorteil von MRP. Dieses basiert auf dem Hiper-Ring, einem von Hirschmann und Siemens entwickelten und 1999 vorgestell-

ten Protokoll zur Ringredundanz. Entsprechende Produkte sind seit einiger Zeit auf dem Markt und in der Praxis erprobt. MRP erlaubt es, Einzelausfälle in einer einfachen Ringtopologie zu kompensieren. Da keine vermaschten Topologien unterstützt werden, ist MRP deterministisch und wesentlich einfacher als RSTP.

MRP verwendet einen Redundanzmanager (RM), der den Ring schliesst. Der Redundanzmanager ist üblicherweise eine logische Funktion in einem Switch, wird hier aber als separate Komponente dargestellt. Im Normalbetrieb überprüft er durch spezielle Testpakete die Durchgängigkeit des Rings. Er leitet aber keine Pakete weiter und verhindert damit, dass diese endlos im Ring zirkulieren.

Fällt ein Switch oder eine Leitung aus, werden die auf einem Port ausgesendeten Testpakete am anderen nicht mehr empfangen. Der Redundanzmanager leitet von nun an die Pakete in beide Richtungen weiter und informiert die Switches über die Topologieänderung, sodass diese ihre Pakete nicht auf die unterbrochene Strecke geben, sondern über den Redundanzmanager verschicken.

Die besten Resultate erreicht man mit PRP, dem Parallel Redundancy Protocol. Dieses beruht auf einem funktionsbeteiligten Redundanzkonzept (ohne Umschal-

Institute of Embedded Systems (InES)

Das InES an der Zürcher Fachhochschule Winterthur ist mit seinen rund 40 Mitarbeitern spezialisiert auf industrielle Kommunikation. Die Schwerpunkte sind hochverfügbare Kommunikationssysteme, präzise Zeitsynchronisation PTP gemäss IEEE 1588 sowie Wireless-Technologien. Insbesondere bietet InES fertige Lösungen für die hochverfügbaren Kommunikationssysteme MRP und PRP gemäss der Norm IEC 62439:

- Portabler MRP-Software-Stack für normale Knoten sowie Redundanzmanager
- Portabler Software-Stack für das PRP mit einer Adaption für den Linux-Kernel 2.6
- FPGA-basierte Hardwarelösung für die PRP-RED-Box

Andererseits werden auch kunden-spezifische Redundanzlösungen entwickelt. Schulung, Support und andere Dienstleistungen des InES helfen, seine Lösungen möglichst effektiv in ihre Produkte zu integrieren.

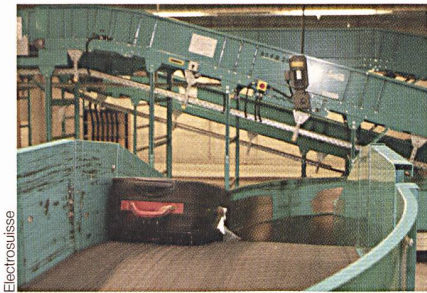


Bild 5 Die Gepäcksortieranlage im Flughafen Kloten darf nicht ausfallen.

tung) und bietet eine völlig unterbrechungs-freie Bewältigung von Ausfällen.

Zwei separate Netzwerke

Für das PRP werden zwei unabhängige Netzwerke parallel betrieben (LAN A und LAN B). Die LAN-Topologie ist grundsätzlich frei wählbar und wäre sogar mit anderen Redundanzkonzepten kombinierbar. Da es zwischen den beiden Netzwerken keine Verbindung gibt, sind sie ausfallunabhängig. Angeschlossene hochverfügbare Knoten besitzen zwei Netzwerkanschlüsse, wobei je einer mit dem LAN A und einer mit dem LAN B verbunden wird.

Alle Pakete werden immer über beide Anschlüsse respektive über beide LAN gesendet. Der Empfänger nimmt bei den Anschlüssen die Pakete entgegen und ignoriert das später eintreffende Paket. PRP kann in einer zusätzlichen Softwareschicht oder als vorgeschaltete Hardwarelösung (RED-Box) realisiert werden. Je nach An-

forderungen müssen dazu in jedem Paket zusätzliche Informationen (Redundancy Control Trailer) mitgeschickt werden, was die maximale Grösse der Nutzinformation einschränkt.

In der Automatisierungstechnik werden erste kommerzielle Pilotanwendungen der PRP-Redundanzlösung realisiert. Eine Automatisierungsanlage von ABB verwendet für den Leitstand ein System auf Basis des Betriebssystems Windows XP. Über zwei parallele Netzwerke ist er mit 12 hochverfügbaren Feldgeräten verbunden. Wegen deren hohen Echtzeitanforderungen laufen diese unter dem Betriebssystem Vxworks. PRP kann auf beiden Plattformen als reine Softwarelösung (in Form eines ladbaren Treibers) implementiert werden und erfüllt die zeitlichen Anforderungen der IED von kleiner als 5 ms Unterbruch im Fehlerfall.

Gepäcksortieranlage des Flughafens Zürich

Dass Steuerungs- und Netzwerkredundanz nicht ausreicht, zeigt die neue Gepäcksortieranlage des Flughafens Zürich. Auf der höchsten Ebene arbeiten die neue und die alte Anlage in verschiedenen Gebäuden in Lastteilung. Bei einem Totalausfall kann die eine Anlage bis zu einem bestimmten Grad die Aufgaben der anderen mittragen. Aber auch die neue Gepäcksortieranlage ist weitgehend redundant aufgebaut. Mechanisch ist dies durch übereinanderliegende Förderanlagen und Sorter gelöst. Die Energieversorgung erfolgt durch je eine Mittelspannungseinspeisung vom

Balsberg und von Kloten aus. Bei so viel Aufwand ist es nicht erstaunlich, dass auch die Steuerungen sowie das Netzwerk von der Betriebsebene bis hin zu Feldbussen und den IOs redundant ausgelegt sind. Auf den höheren Automatisierungsebenen werden Glasfaserringe eingesetzt, an die verschiedene SPS angekoppelt werden. Diese Siemens S7-400 verfügen über dezentralisierte Peripherieeinheiten (ET 200), die über Profibus-DP verbunden sind. Anlagenteile, die mechanisch redundant sind, werden nur von einer einfachen Steuerung bedient. Ist die physische Redundanz nicht gegeben, kommt dagegen eine redundante Master/Slave-SPS zum Einsatz.

Technik alleine genügt nicht

Die Methoden zur Realisierung von hochverfügbaren Systemen sind bekannt. Für einen weiten Anwendungsbereich sind Off-the-Shelf-Produkte verfügbar. Eine genaue Analyse der Anforderungen ist jedoch unumgänglich. Es gibt keine Standardlösung für alle Problemstellungen. Und technische Massnahmen allein genügen nicht: Diagnose, Alarmierung, Schulung des Wartungspersonals und Materiallogistik spielen eine entscheidende Rolle, damit bei einem Ausfall einer Komponente die Redundanz möglichst rasch und sicher wiederhergestellt werden kann. In diesem Sinn ist die Fehlertoleranz kein Ersatz für Qualität und Wartung.

Angaben zum Autor

Thomas Müller ist Professor an der Zürcher Hochschule Winterthur (ZHAW), Leiter des Instituts für Embedded Systems InES. Der Artikel enthält Informationen aus Vorträgen von Hubert Kirmann, ABB Schweiz AG, und Guido Egloff, Siemens Schweiz AG, aus der ITG-Tagung vom 25. Januar 2007.
Zürcher Hochschule Winterthur, 8401 Winterthur, thomas.mueller@zhwin.ch

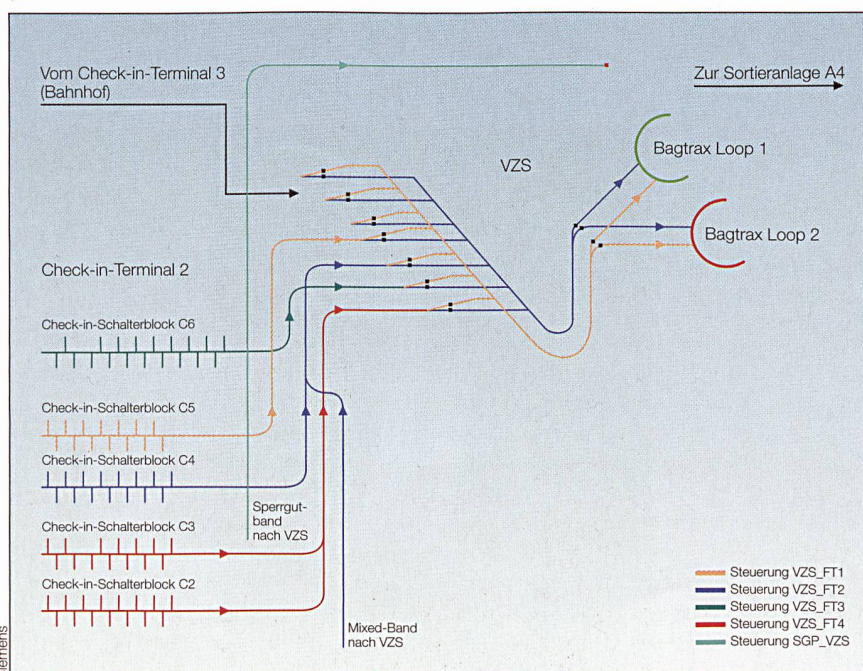


Bild 6 Bei der Gepäcksortieranlage ist auch die Mechanik redundant ausgelegt.

Résumé

Quand une installation n'a pas le droit de tomber en panne

Systemes d'automatisation à haute disponibilité. «Deux sécurités valent mieux qu'une!» Car si l'une fait défaut, l'autre est toujours là. Au niveau des installations, cela signifierait qu'une infrastructure redondante apporterait une meilleure disponibilité. Mais est-ce bien vrai? Si une double couture est censée mieux tenir, il n'en reste pas moins vrai que les piqûres supplémentaires de la seconde peuvent affaiblir le tissu, provoquant une déchirure. C'est l'exécution correcte qui est déterminante. Même chose pour les installations à haute disponibilité.