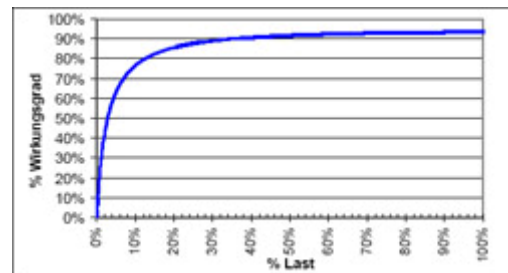


Steigerung der Effizienz großer USV-Systeme - Teil 1

Energieressourcen werden immer knapper und teurer. Daher kommt dem Wirkungsgrad bei der Spezifikation und Auswahl großer USV-Systeme eine immer höhere Bedeutung zu. Der Betrieb eines USV-Systems und die damit zusammenhängenden Energiekosten werden von drei Faktoren maßgeblich beeinflusst. Leider werden diese Faktoren bei der Spezifizierung der Systeme häufig nicht berücksichtigt. Für Unternehmen, die derartige Systeme einsetzen, ergeben sich aus dieser mangelnden Berücksichtigung des Wirkungsgrads erhebliche Mehrkosten. Die vorliegende Fachartikelreihe befasst sich mit den häufigsten Fehlern und Missverständnissen bei der Beurteilung des Wirkungsgrads von USV-Systemen. Dabei werden die einzelnen Wirkungsgradkurven erläutert, verglichen und auf ihre Kostenrelevanz hin untersucht.

Einführung

Traditionell liegt das Augenmerk bei der Spezifizierung und Auswahl von USV-Systemen nahezu ausschließlich auf der Systemverfügbarkeit, d. h. auf dem mittleren Ausfallzeitraum (MTBF) der Systeme, wie er von Herstellern und technischen Beratern genannt wird. Heutzutage sorgen jedoch zwei Faktoren dafür, dass neben der Verfügbarkeit zunehmend auch der Wirkungsgrad von USV-Systemen ins Blickfeld rückt: (1) die steigende Bedeutung der Gesamtbetriebskosten (TCO) über den Nutzungszeitraum des Systems hinweg und (2) Umweltinitiativen im öffentlichen und privaten Sektor wie beispielsweise entsprechende Zertifizierungen für die Gebäudeeffizienz oder auch nachfrageseitige Managementprogramme, die von Versorgungsunternehmen angeboten werden.



Der Wirkungsgradverlust von USV-Systemen hat zwei Hauptursachen: erstens die inhärenten Stromverluste der USV-Module und zweitens die Art und Weise, wie die Systeme implementiert werden (d. h. bedarfsgerechte Auslegung, Redundanz). Bei der Spezifikation von USV-Systemen wird häufig nur der vom Hersteller angegebene Wirkungsgrad betrachtet, der unter Best-case-Bedingungen erzielt wird. Dieser Wert ist irreführend, wie im Folgenden näher erörtert wird.

Anhand eines hypothetischen Beispiels soll hier aufgezeigt werden, welche Auswirkungen diese Vorgehensweise auf die Stromkosten eines Unternehmens haben kann. Dabei sollen zwei 1-MW-USV-Systeme von zwei verschiedenen Herstellern betrachtet werden. System 1 und System 2 weisen identische Herstellerangaben zum Wirkungsgrad auf (93 % bei Volllast), werden in einer 2N-Architektur betrieben, schlagen mit Stromkosten von 0,10 USD/kWh zu Buche und unterstützen eine Last von 300 kW. Im Normalfall wird daraus automatisch gefolgert, dass die beiden betrachteten Systeme identische Jahresstromkosten verursachen. Dies trifft jedoch so nicht zu. Mit Ausnahme von Notfall- oder Wartungsszenarien werden USVen in einer 2N-Konfiguration nie mit einem Lastpegel von 100 % betrieben, da jedes System in der Lage sein muss, die volle Last zu übernehmen, falls das andere System ausfällt. Aus diesem Grund darf die projektierte Höchstlast eines jeden Systems im Normalbetrieb 50 % nicht überschreiten. In der Realität erreichen 2N-Systeme selten eine Last von 50 % pro System. Eine Reihe von Feldstudien lassen vielmehr darauf schließen, dass 2N-Datencenter nur mit 20 bis 40 % ihrer 2N-Leistung gefahren werden¹. Im obigen Beispiel wird eine typische Last von 30 % veranschlagt, bei der jede USV eine Kapazität von 150 kW bietet. Pro USV in System 1 fallen Jahresstromkosten für die Verlustleistung in Höhe von 10.470 USD an. Dem stehen 28.322 USD pro USV in System 2 gegenüber. Da jedes System über zwei USV-Module verfügt, verdoppeln sich die Kosten auf 20.940 bzw. 56.644 USD pro Jahr. Dabei manifestieren sich die Kosten für die Verlustleistung in Form von Abwärme, die von einem Kühlsystem abgeleitet werden muss.

Wenn das Kühlsystem für die Ableitung von 1 kW erzeugter Wärme insgesamt 400 Watt benötigt, entstehen zusätzliche jährliche Kosten in Höhe von 8.376 bzw. 22.651 USD². Im obigen Beispiel erreicht ein typisches Datencenter eine Nutzungsdauer von 10 Jahren; dies resultiert in Gesamtkosten für den USV-Stromverbrauch in Höhe von 293.165 bzw. 793.021 USD (siehe Tabelle 1). Wie kann es dazu kommen, dass sich zwei auf den ersten Blick identische USV-Systeme bei der Verlustleistung nahezu um den Faktor 3 unterscheiden?

USV-System	USV-Kosten für Verlustleistung	Kühlkosten	Jähl. Kosten durch geringen Wirkungsgrad	Zusatzkosten durch geringen Wirkungsgrad (210 Jahre)
USV-System 1	\$20.940	\$8.376	\$29.317	\$293.165
USV-System 2	\$56.644	\$22.651	\$79.302	\$793.021

Die Antwort auf diese Frage liegt in den unterschiedlichen Wirkungsgradkurven der beiden Systeme und in ihrer unterschiedlichen Lastanpassung. Schon eine Verbesserung des Wirkungsgrads einer USV von nur 5 % kann – je nach Last – zu einer Verringerung der Stromkosten zwischen 18 und 84 % führen. Diese Tatsache wird weiter unten anhand von zwei aktuellen USV-Systemen näher erläutert.

Mit Blick auf die heutigen Anforderungen an Wirkungsgrad und Umweltverträglichkeit stehen USV-Herstellern drei Faktoren zur Verfügung, mit denen sie groß dimensionierte USV-Systeme optimieren können: **Technologie, Topologie und Modularität**. Gemeinsam können diese Faktoren dazu beitragen, die elektrischen Verluste von USVen in Form von Wärmeenergie (kW) wirksam zu verringern. Diese technische Dokumentation erläutert das Prinzip der Wirkungsgradkurve und geht auf häufige Fehler bei der Beurteilung des USV-Wirkungsgrades ein. Es macht darüber hinaus deutlich, wie sich mit den Faktoren Technologie, Topologie und Modularität der Wirkungsgrad von USV-Systemen optimieren lässt. Eine ausführliche Erörterung zum Thema Datencenter-Wirkungsgrad kann dem APC White Paper Nr. 113, „Elektrisches Wirkungsgradmodell von Datencentern“ entnommen werden.

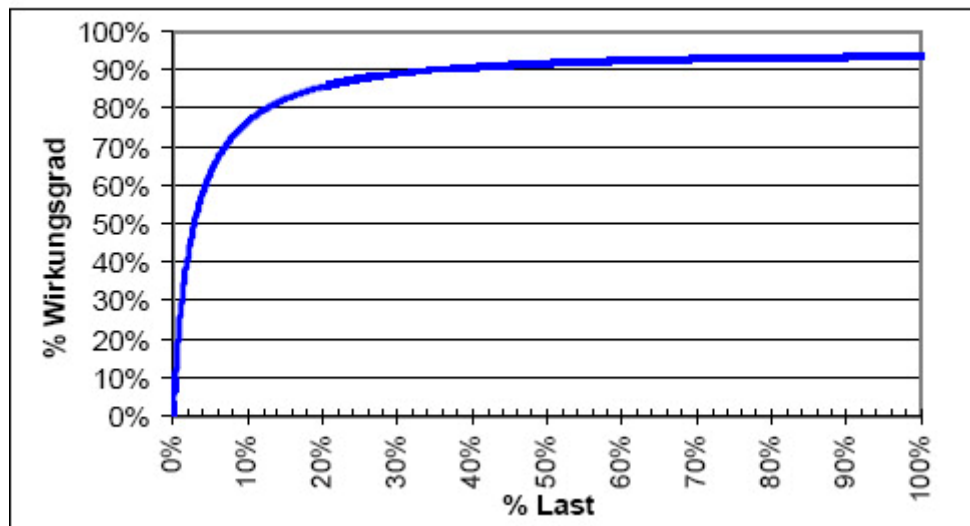
Die USV-Wirkungsgradkurve

Wenn ein USV-Datenblatt lediglich eine einzige Kennzahl für den Wirkungsgrad nennt, handelt es sich nahezu immer um einen Wert, der bei einer Last von 100 % (Nennlast) und unter verschiedenen anderen vorteilhaften Systembedingungen erreicht wird. Dazu zählen vollständig geladene Batterien, eine Nenneingangsspannung der USV und der Verzicht auf den Anschluss bzw. die Installation optionaler Eingangstransformatoren und -filter. Der Grund, weshalb USV-Hersteller den Wirkungsgrad bei Volllast angeben, ist ganz einfach: Es ist der beste Wirkungsgrad, der sich mit der jeweiligen USV erzielen lässt. **Leider lässt sich dieser Wert in der Praxis nur äußerst selten realisieren, da eine Last von 100 % nahezu nie gegeben ist.** Wenn eine USV auf Basis des vom Hersteller genannten Wirkungsgrads spezifiziert wird, ist das genauso, als würde man ein Auto kaufen, dessen niedrigster Benzinverbrauch für eine Geschwindigkeit von 120 km/h berechnet wurde, und dieses Auto dann ausschließlich im Stadtverkehr fahren. Sehr viel sinnvoller ist es hier, den Wirkungsgrad bei einer Last von rund 30 % zugrunde zu legen, d. h. bei der durchschnittlichen Last, mit der die meisten mittleren bis großen Datencenter betrieben werden. Hierfür muss zunächst klar sein, was eine USV-Wirkungsgradkurve ist und wie sie erzeugt wird.

Abbildung 1 zeigt die grundlegende Form einer USV-Wirkungsgradkurve. Dabei entspricht der höchste Punkt der Kurve dem höchsten Wirkungsgrad (y-Achse) und der höchsten Last (x-Achse). Je nach USV wird dieser Punkt bei einer Last von 20 bis 30 % erreicht, nach deren Unterschreiten der Wirkungsgrad scharf abfällt. In der abgebildeten Kurve beträgt der maximale Wirkungsgrad 93 %. Um eine USV unter Verwendung eines realistischen

Lastpegels zu spezifizieren, muss ihr Wirkungsgrad bei einer üblichen Last von beispielsweise 30 % ermittelt werden. In der unten stehenden Kurve wird bei diesem Lastpegel ein Wirkungsgrad von 89 % erzielt. In Datacentern mit redundanten USV-Systemen (2N) ist ein noch stärker ausgeprägter Abfall des Wirkungsgrads zu beobachten, da sich die Last auf beide USVen verteilt, so dass sich ein Wirkungsgrad von nur noch 82 % ergibt. Diese Auswirkung der Redundanz wird weiter unten noch näher erläutert.

Abb. 1 – USV-Wirkungsgradkurve

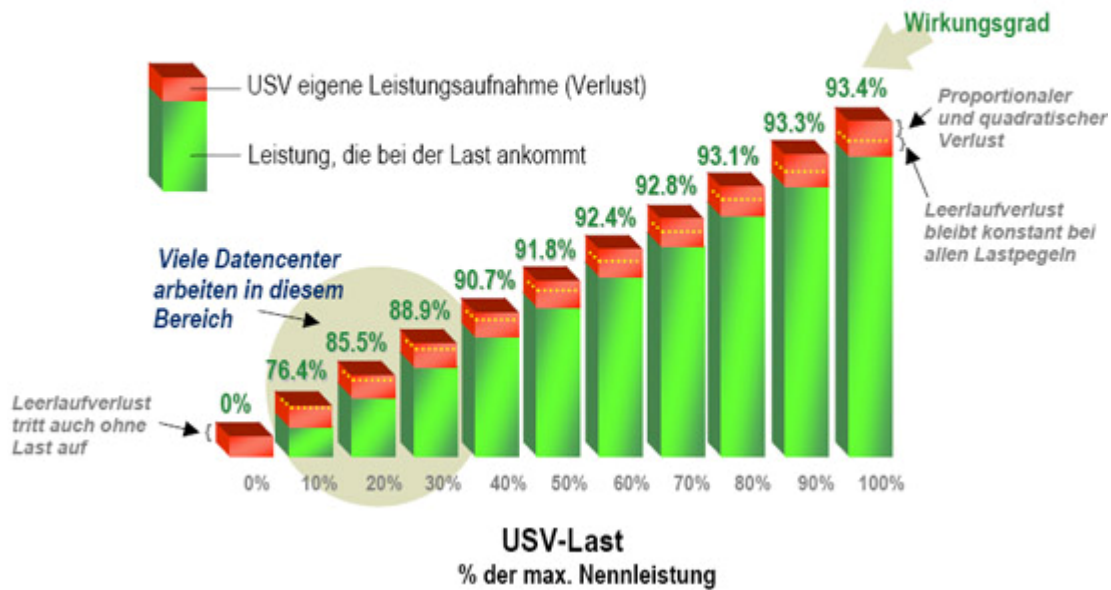


Erzeugen einer USV-Wirkungsgradkurve

Eine Wirkungsgradkurve wird erzeugt, indem zunächst die der USV bereitgestellte Leistung (Eingangsleistung) und die von USV für die Last aufgebrachte Leistung (Ausgangsleistung) gemessen werden. Diese Messungen erfolgen bei verschiedenen Lastpegeln (üblicherweise bei 25 %, 50 %, 75 % und 100 %). Eine weitere Messung findet bei einer Last von 0 % statt, um den Eigenstrombedarf (Leerlaufverlust) der USV zu bestimmen. Ausgehend von diesen Messungen werden die Verluste berechnet, indem die Eingangsleistung von der Ausgangsleistung subtrahiert wird. Die Verluste werden in einem Diagramm abgetragen; anschließend wird über die so erhaltenen Punkte eine Trendlinie gelegt. Die Trendlinie ergibt eine Gleichung, anhand derer alle übrigen Punkte für jeden beliebigen Lastprozentsatz ermittelt werden können. Nachdem alle Leistungsverluste berechnet wurden, wird die Wirkungsgradkurve erzeugt, indem das Verhältnis von Ausgangsleistung zu Eingangsleistung für die verschiedenen Lastpegel abgetragen wird.

Abbildung 2 veranschaulicht das Prinzip der Wirkungsgradkurve aus Abbildung 1, indem der Verbleib der Eingangsleistung aufgeschlüsselt wird.

Abb. 2 – Detaillierte Darstellung der USV-Eingangsleistung und ihrer Nutzung



In der Abbildung stehen die grünen Balken für die Gesamtleistung, die der IT-Last zur Verfügung steht, während die roten Balken für die internen USV-Verluste stehen, die die Wirkungsgradkurve aus Abbildung 1 ergeben. Bei einem idealen Wirkungsgrad wird die gesamte Eingangsleistung auf der Ausgangsseite dem Datacenter zur Verfügung gestellt, so dass bei allen Lastpegeln ausschließlich ein grüner Balken zu zeichnen wäre (d. h. es gibt keine Leistungsverluste). In diesem Fall hätte die Wirkungsgradkurve die Form einer horizontalen Linie (100 % für alle Lastpegel). Wie aus den oben eingezeichneten roten Balken hervorgeht, wird ein Teil der Eingangsleistung jedoch von der USV selbst benötigt. In einer USV fallen drei Arten von Verlusten an: Leerlaufverlust, proportionaler Verlust und quadratischer Verlust.