



Weniger Störgeräusche bei der feldorientierten Regelung

Toshibas Methode der symmetrischen PWM-Trägersignale macht hochfrequente Signale in der Vektorregelung überflüssig und sorgt für leisen Motorlauf.

Seite 14

Sieben Tipps zum Display-Design

Worauf Entwickler bei der Auswahl des Grafik-Displays für die Anwendung achten sollten.

Seite 20

Kryptoschlüssel sicher einsetzen

Vernetzte MCUs brauchen eine vor Manipulationen geschützte Verwaltung der Schlüssel.

Seite 36

Leistungshalbleiter detailliert messen

Dynamische Eigenschaften von Power-ICs mittels Doppelpulsverfahren vergleichen.

Seite 64

Über
9,6 Millionen
Produkte online
DIGIKEY.DE



ÜBER 9,6 MILLIONEN PRODUKTE ONLINE | ÜBER 1200 BRANCHENFÜHRENDE ANBIETER



SIE ENTWICKELN. WIR HELFEN. **DIGIKEY.DE**



Die weltweit größte Auswahl an elektronischen Komponenten für den sofortigen Versand™

*Für alle Bestellungen unter 50,00 € wird eine Versandgebühr von 18,00 € in Rechnung gestellt. Bei Bestellungen unter \$60,00 USD wird eine Versandgebühr von \$22,00 USD berechnet. Alle Bestellungen werden per UPS, Federal Express oder DHL für die Lieferung innerhalb von 1 bis 3 Tagen (abhängig vom endgültigen Bestimmungsort) versendet. Keine Bearbeitungsgebühren. Alle Preise werden in Euro oder US-Dollar angegeben. Digi-Key ist ein autorisierter Distributor für alle Lieferpartner. Neue Produkte werden täglich hinzugefügt. Digi-Key und Digi-Key Electronics sind eingetragene Marken von Digi-Key Electronics in den USA und anderen Ländern. © 2021 Digi-Key Electronics, 701 Brooks Ave. South, Thief River Falls, MN 56701, USA

ECIA MEMBER
Supporting The Authorized Channel

Dusch mich, aber mach' mich nicht nass!

Der akute Chipmangel in der europäischen Automobilindustrie führt ungeschönt vor Augen, wie abhängig Hightech-Standorte in Europa von Lieferanten in fernen Regionen dieser Welt sind. Dass wegen fehlender ICs aus Fernost tatsächlich hier Fertigungsbänder stillstehen und Belegschaften in Kurzarbeit geschickt werden, ist ein starkes Stück. Doch genau das war vorhersehbar, ist man beim ZVEI überzeugt. Hiesige Produzenten wollten von globalen Lieferketten und günstiger Produktion in Fernost profitieren, den Halbleiterproduzenten durch unzureichende Forecasts aber kaum Planungssicherheit zugestehen. Dann dürfe man sich nicht wundern, wenn letztere ihre Fertigungen auf andere Märkte ausrichten. Schließlich müssten die Chipproduzenten ihre etliche Milliarden Euro teuren Fabriken maximal auslasten, um profitabel sein zu können.

Die schlechte Nachricht: Die aktuelle Mangelsituation wird kein kurzfristiger Einmal-Effekt bleiben, ist ZVEI-Chef Dr. Gunther Kegel sicher. Die Fertigungskapazität ist nun einmal begrenzt. Ist diese für andere Abnehmer blockiert, können sich die hiesigen Kunden erst einmal hinten anstellen. Das ist die harte Realität. Längst gibt es daher Rufe nach mehr Souveränität für Europas Chipfertigung.

„Der Weg zur Technologieführerschaft in der Chipfertigung wäre kein Sprint, sondern ein Marathon.“



Michael Eckstein, Redakteur
michael.eckstein@vogel.de

Gleichzeitig will man die Vorzüge der Globalisierung nicht aufgeben. Das klingt ein wenig nach: „Dusch mich, aber mach' mich nicht nass!“ – zumindest lassen sich Protektionismus (light) und freie Märkte kaum unter einen Hut bringen. Die Alternative wäre: Die eigenen Produkte müssen weltweit führend und ausreichend verfügbar sein. In einem aktuellen Positionspapier hat der VDE dargelegt, wie der Produktionsanteil an Halbleitern in Europa massiv erhöht werden könnte: Mit der richtigen Fokussierung und langfristigen (Förder-)Programmen.

Keine Frage: Der Weg zur Technologieführerschaft in der Chipproduktion wäre kein Sprint, sondern ein Marathon. Dafür braucht man einen langen Atem, festen Willen, muss über einen langen Zeitraum viele Ressourcen in hartes Training investieren und auch Fehlschläge verkraften. Ist Europa bereit dafür?

Herzlichst, Ihr

FlowCAD

Work from Home Special

OrCAD PCB Designer Standard
inklusive 2 Jahre Updates/Hotline



Probleme beim Zugriff auf Software Lizenzen von zu Hause sind mit dem „Work from Home“ Programm von FlowCAD und Cadence kein Thema mehr.

Erhalten Sie OrCAD PCB Designer Standard und zwei Jahre Updates und Hotline zum Sonderpreis von 999,- €*!

* netto + ges. MwSt.

Damit kann jeder Entwickler, Bibliothekar und PCB Designer seine eigene Lizenz lokal im Home Office nutzen und

- Schaltpläne erstellen und mit PSpice simulieren
- Bauteile platzieren und PCBs routen
- Bibliothekselemente erstellen

Es muss keine online-Verbindung zur Firmen IT bestehen. Die Mitarbeiter sind produktiv und können sich ihre Arbeitszeit frei einteilen.



FlowCAD.com/wfh

OrCAD®

cādense®

ELEKTRISCHE ANTRIEBE

Weniger Störgeräusche bei der feldorientierten Regelung

Toshiba hat eine Methode zur feldorientierten Regelung entwickelt, die mit symmetrischen PWM-Trägersignalen arbeitet. Sie kommt ohne hochfrequente Signale aus, vermeidet Vibrationen und sorgt für einen ruhigeren Motorlauf.

14



ELEKTRONIKSPIEGEL

6 **Zahlen, Daten, Fakten**

8 **Aktuelles**

SCHWERPUNKTE

Elektrische Antriebe

TITELTHEMA

- 14 **Störgeräusche bei Motorregelungen vermeiden**
Synchrone PWM-Trägersignale für die einzelnen Motorphasen machen hochfrequente Signale überflüssig.

Mensch-Maschine-Interface

- 18 **Sicherheit bei Touchscreen-Eingabe umsetzen**
Richtig ausgelegte Funktions- und Bediensicherheit schützt vor gefährlichen Fehleingaben.

- 20 **Auswahlhilfe für Grafikdisplays**
Mit diesen Tipps finden Entwickler schnell ein Display, das die spezifischen Anforderungen ihres Projekts erfüllt.

Verbindungstechnik

- 56 **Leistungssteckverbinder für 60 A im Raster 8,5 mm**
Wie kleine Steckverbinder mit langer Lebensdauer unter rauen Betriebsbedingungen punkten.
- 58 **Hebelbedienbare PCB-Klemmen und -Steckverbinder**
Aktuelle Leiterplatten-Steckverbinder kombinieren den Push-in-Federanschluss mit praktischer Hebelbetätigung.

Messtechnik

- 62 **Skalierbare Encoder-Anwendung mit einem Chip**
Schaltungsblöcke mit reflexiver On-Chip-Sensorik passen sich unterschiedlichen Encoder-Durchmessern an.
- 64 **Dynamische Eigenschaften von Leistungshalbleitern**
Der Doppelpulstest hilft dabei, die dynamischen Parameter von Power-ICs zuverlässig zu vergleichen.

Bauteilebeschaffung

- 68 **Mit hochintegrierten Lösungen schnell zum Produkt**
SoCs und SiPs ermöglichen den raschen Zugang zu innovativen Halbleitertechnologien.
- 70 **Leistungsgrenzen optischer Module erweitern**
MEMS-basierende Oszillatoren ermöglichen sehr hohe Taktraten etwa für optische Übertragungsmodule.

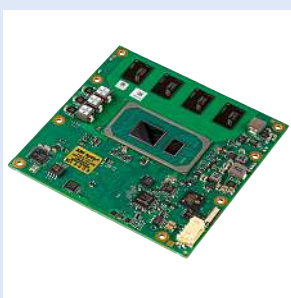
SONDERTEIL

EMBEDDED SYSTEMS DEVELOPMENT UND INTERNET OF THINGS

- 28 **Embedded-Servermodule für Edge-Rechenzentren**
- 30 **IoT- und KI-Plattform für Edge-Computing**
- 32 **FPGA-basierter Autopilot für Drohnen**
- 34 **Flexible System-on-Chip-Lösung für IoT-Endgeräte**



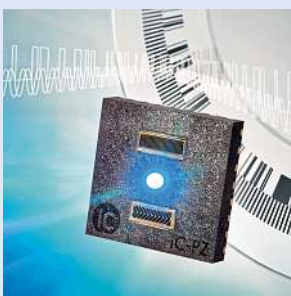
18 Touch-Bedienung
muss eindeutig sein



30 COM-Express-Module
für Edge-Server



52 NFC-Ladelösung für
kleine Endgeräte



62 Skalierbarer Encoder
mit nur einem Chip

- 36 Kryptographie-Schlüssel sicher injizieren und nutzen
- 40 Schnellere Software-Verifizierung und -Validierung
- 44 Multikernel-Betriebssystem für Automotive
- 48 ESE-Kongress auch digital auf Erfolgskurs
- 50 Asset-Management: Funktionsweise von LoRa Edge
- 52 Drahtlose NFC-Ladelösung für kleine Endgeräte

TIPPS UND SERIEN

- 13 Power-Tipp: PSR-Sperrwandler verbessern
- 60 Steckverbinder erklärt: Steckzyklen und Oberflächen
- 74 Zum Schluss: Fertigung in Zeiten von Lockdowns

RUBRIKEN

- 3 Editorial
- 12 Veranstaltungen
- 72 Impressum

RUTRONIK TECHTALK MEETS EMBEDDED WORLD

March 2 - 4, 2021 | online

**SEE
YOU!**

Online Registration*

www.rutronik.com/TechTalk_Embedded_World

* This invitation is subject to the approval of your superior



AUFGEMERKT



Bild: 8mb_ibm_disk_on_key / Gbuchana / CC BY-SA 3.0 / Wikimedia Commons

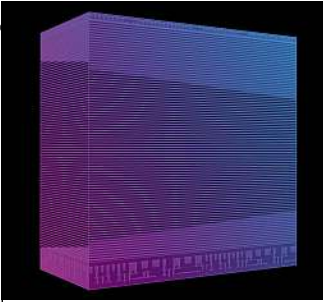
2000: Der erste USB-Stick

Shimon Shmueli oder Dov Moran – wer das erste USB-Speicherstäbchen erdacht hat, ist nicht vollständig geklärt. Die ersten Flash-Speicher mit USB-Anschluss im Hosentaschenformat brachte IBM Ende 2000 als DiskOnKey auf den Markt. Kapazität 8 bis 32 MByte. Die USB-Sticks bestehen nur aus wenigen Bauelementen: dem Flash-Chip und dem Controller-Chip mit Firmware, bei denen die Firma M-Systems zu den Pionieren gehörte. Deren Inhaber Dov Moran fertigte in Kfar Saba, Israel, die USB-

Speicher für IBM, später u.a. auch für Hewlett-Packard Apple. Als die Speicherkapazitäten der USB-Sticks so rasant stiegen wie die Preise sanken, war der Durchbruch geschafft. Der Verkaufshöhepunkt wurde 2012 und 2016 mit über 16 Millionen Sticks erreicht. Die praktischen USB-Sticks werden uns sicherlich auch in Zeiten der Cloud begleiten, aber die Stückzahlen dürften weiter sinken. Im Jahr 2006 verkaufte Dov Moran sein Unternehmen an den Konkurrenten SanDisk für rund 1,6 Milliarden US-Dollar. // KR

AUFGEDREHT: Raspberry Pi Pico

Bild: Micron Technologies



NAND-Chip mit 176 Lagen

176 NAND-Einzelchips, jeder nur ein Fünftel so dünn wie ein Blatt Papier, stapelt Micron zu einem Rekord-3D-Flash-Baustein. Damit zielt der US-Hersteller auf Mobil-, Automobil-, Consumer- und Rechenzentrums-Anwendungen. Unter dem Paket aus TLC-NAND-Speicherzellen hat Micron die Steuerelektronik integriert, so dass diese seitwärts keinen zusätzlichen Platz benötigt. So entsteht ein überaus kompakter Speicherblock mit einer Kapazität von zunächst 512 GBit. // ME

Chip

Mikrocontroller RP2040, ein ARM Cortex-M0+ Dual-Core mit 2 x 133 MHz Takt. Der Chip ist die erste Eigenentwicklung der Foundation.

USB

USB 1.1 PHY und Controller. Der Standard-Mikro-USB-Port kann sowohl im Device- als auch im Host-Modus verwendet werden.

Bootsel

Soll ein Device den Pico als USB-Gerät erkennen, muss der Taster während des Anschlusses des USB-Kabels gedrückt werden.

Pinleiste

2 x 20 I/Os als lötbare Wabenleiterplatte. Alternativ fassen die Löcher 2,54-mm-Stiftleisten. Pinbelegung auf der Geräteunterseite.

VBUS

VBUS auf Pin 40 adressiert die Micro-USB-Eingangsspannung, angeschlossen am Micro-USB-Port. Sie beträgt nominell 5 V.

VSYS

Haupteingangsspannung (Pin 39; 1,8-5,5 V). Das interne Schaltnetzteil erzeugt 3,3 V für RP2040; versorgt auch externe Schaltungen.



Bild: Raspberrypi.org

Raspberry Pi Pico ist ein leistungsstarkes Mikrocontroller-Board mit flexiblen digitalen Schnittstellen für rund 5 Euro. Das erste MCU-Board des Raspberry-Pi-Herstellers taktet mit bis zu 2 x 133 MHz und bietet 264 GByte SRAM und 2 MByte

Flash-Speicher. Der Stromversorgungs-Chip unterstützt Eingangsspannungen von 1,8 bis 5,5 V. Der Pico bietet etwa 2 x SPI, 2 x I²C, 2 x UART, 3-x-12-Bit-ADC und 16 x steuerbare PWM-Kanäle. Desweiteren einen Uhr- und Timer-Chip. // MK

AUFGE-SCHNAPPT

„Das Auto wird das komplexeste Digitalprodukt der Welt.“

Herbert Diess, Vorstandsvorsitzender der Volkswagen AG in einem Podcast vom November 2020.

Joe Kaeser tritt als Siemens-CEO ab

Seit 2013 hatte Joe Kaeser die Geschicke geleitet, am 3. Februar 2021 nahm der Siemens-Chef nun Abschied als CEO. In seine Zeit fielen Umstrukturierungen wie die Abspaltung der Medizintechnik-Tochter Siemens Healthineers. Nachfolger als CEO ist nun Roland Busch. Kaeser wird als Aufsichtsratschef beim Unternehmen bleiben. // SG



Bild: Siemens

6

RENESAS KAUFT DIALOG SEMICONDUCTOR

Für 6 Mrd. US-\$ (knapp 5 Mrd. Euro) übernimmt Renesas die deutsche Chipschmiede Dialog Semiconductor. Der Spezialist für Analogtechnik, Mixed Signal und Low Power hatte erst kürzlich aufgrund des Brexits seine Firmenzentrale nach Deutschland verlegt. Renesas befindet sich seit 2018 auf Einkaufstour: Prominente Übernahmen waren Intersil und IDT. Mit der aktuellen Akquise wollen die Japaner vor allem die Bereiche IoT, Industrie und Automotive auf dem europäischen Markt strategisch stärken.

ENERGIESPEICHER

Post-Lithium-Ionen-Batterien erfordern neue Kompetenzen



Bild: WWU - MEET

Coating: Das Beschichten der Elektroden ist einer von vielen Prozessschritten in der Zellfertigung.

Die Forschung zur Batteriezellherstellung gewinnt an Dynamik. Ein Forscherteam unter der Federführung von Wissenschaftlern der Westfälischen Wilhelms-Universität Münster (WWU) hat nun im Fachmagazin „Nature Energy“ einen Übersichtsartikel zu Fertigungsprozessen verschiedener Batterietypen veröffentlicht. Dieser legt nahe, dass die aktuell etablierte Lithium-Ionen-Batterie (LIB) den Markt wiederaufladbarer Hochenergiebatterien mittelfristig dominieren wird. Alternative Batterie-

technologien, insbesondere Feststoffbatterien, aber auch Lithium-Schwefel- oder Lithium-Luft-Batterien, werden zwar intensiv erforscht, jedoch noch nicht industriell in Großserie produziert. Ausgehend von zahlreichen aktuell entstehenden Produktionskapazitäten für LIB, würde eine Umstellung auf sogenannte Post-Lithium-Ionen-Batterien (PLIB) mit neuen Prozesstechnologien, Fertigungsumgebungen sowie Kompetenzen einhergehen und erfordere deshalb Milliardeninvestitionen.

Einzig die Produktion von Natrium-Ionen-Batterien ist in vielen Prozessschritten vergleichbar mit der von LIB. Die Prozesse weiterer PLIB wie Festkörper-, Lithium-Schwefel-, oder Lithium-Luft-Batterien unterscheiden sich deutlich von der Herstellung der Lithium-Ionen-Batterien: Schritte wie Elektrodenherstellung, Zellbau oder Zyklierung erfordern andere Techniken, Herstellungsumgebungen oder Maschinen. // TK

WWU Münster

MCCLEAN-REPORT ZUM WACHSTUM DES CHIPMARKTS

Speicher, Automotive-Chips und Embedded-MPUs 2021 enorm gefragt

DRAM und NAND-Flash werden im Jahr 2021 die am stärksten wachsenden Teilmärkte der Halbleiterbranche sein. Zu diesem Ergebnis kommt die jüngste Edition des McClean Reports von IC Insights, der regelmäßige Prognosen über die Entwicklungen der Halbleiterbranche entwirft.

Für DRAM wird für das Jahr 2021 ein Wachstum von 18 % in Aussicht gestellt, während NAND voraussichtlich um 17 % zulegen soll. Neben den ähnlichen zyklischen Schwankungen wie auch bei DRAM profitierten

2020 die NAND-Umsätze besonders stark von der andauernden Covid-19-Pandemie. Der Speichermarkt unterliegt generell enormen Schwankungen: So hatten DRAMs zwar 2018 zuletzt den ersten Platz im Ranking eingenommen, aufgrund stark fallender Preise aber 2019 den letzten Platz unter den 33 geführten Chip-Produktkategorien belegt.

Automotive: Application-Specific Analog und Automotive: Special Purpose Logic sollen voraussichtlich um je 16 % wachsen. Durch zusätzliche elektroni-

sche Systeme, Onboard-Konnektivität, Fortschritte beim autonomen Fahren und ein weltweit höherer Absatz von Elektro-Fahrzeugen soll laut IC Insights der durchschnittlichen Halbleiteranteil pro Neufahrzeug im Jahr 2021 mehr als 550 US-\$ betragen.

Auch der Markt für Embedded-MPUs soll Fahrt aufnehmen (15 % Wachstum). SoC-Anbieter wie Qualcomm, Samsung und MediaTek konzentrieren sich verstärkt auf 64-Bit-Embedded-Prozessoren, die Sicherheitsfunktionen und KI-Beschleuni-

gung durch maschinelles Lernen zusammen mit Grafik- und Videofunktionen für automatisierte Fahrzeuge oder Drohnen sowie IoT-Anwendungen integrieren. Weitere besonders starke Wachstumsfelder sind Display-Treiber, Anwendungsspezifische Analog-Bausteine für drahtgebundene (je 11 %) bzw. schnurlose Kommunikation (10 %), 32-Bit-MCUs sowie Logik-ICs für Computer und Peripheriegeräte (ebenfalls jeweils 10 %). // SG

IC Insights

AZURE QUANTUM PUBLIC PREVIEW

Öffentlich zugängliche Testversion für Quantencomputing via Cloud



Bild: Microsoft

Public Preview: Azure Quantum ist eine per Cloud zugängliche Entwicklungsplattform für Quantencomputer.

Microsoft hat mit der Public Preview seines Dienstes Azure Quantum ein offenes Ökosystem zum Umgang mit und der Entwicklung für Quantencomputer zur Verfügung gestellt.

Entwickler, Forschende und Firmenkunden erhalten so die Möglichkeit, sich über eine Cloud-Anbindung kostenlos mit einer Entwicklungs- und Bedienungsumgebung für Quantencomputer vertraut zu machen.

Hierfür stellt Microsoft eine offen zugängliche, umfangreiche Dokumentation bereit. Diese

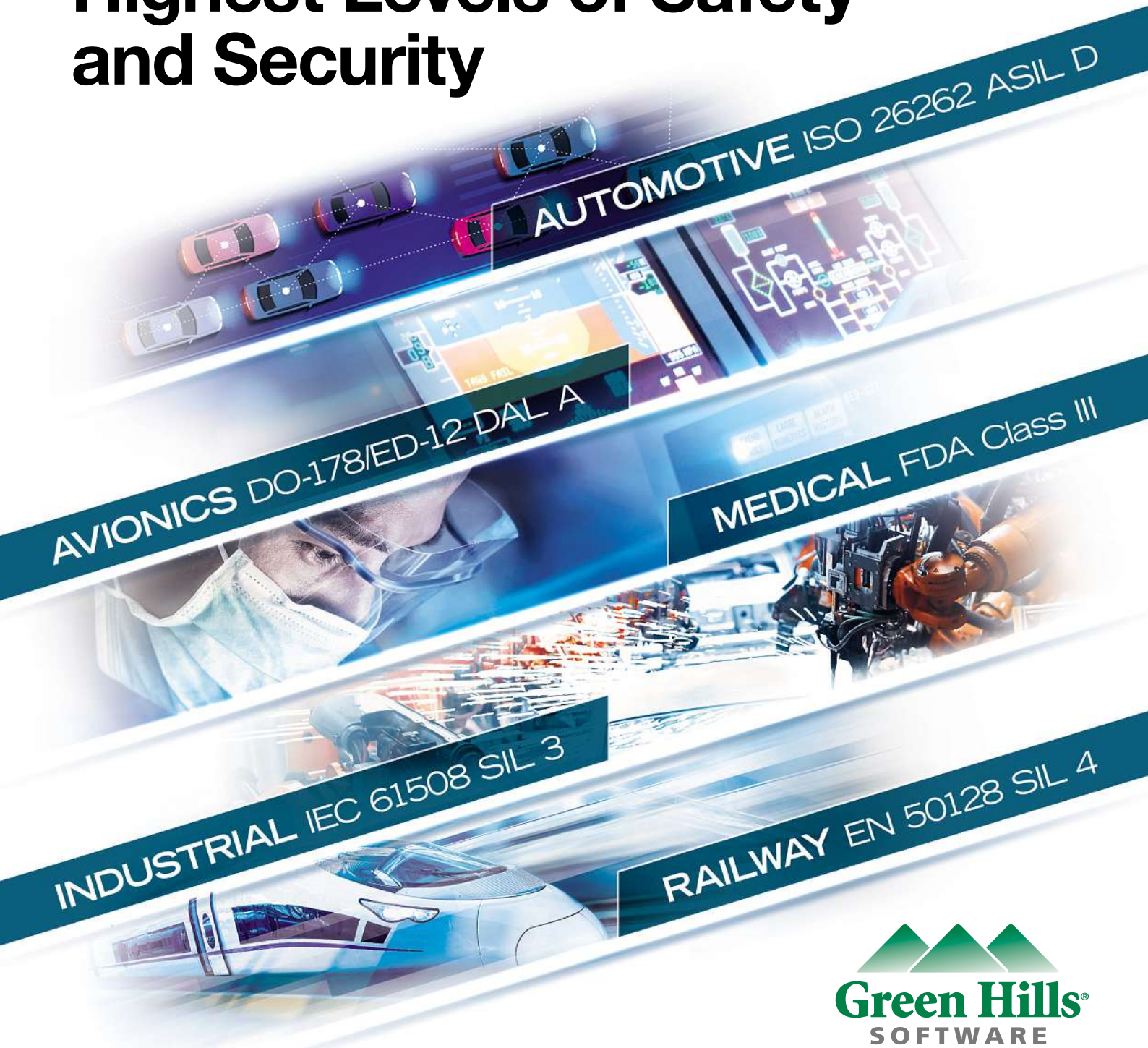
umfasst eine Schritt-für-Schritt-Einweisung in das Ökosystem und Tutorials für den Umgang mit der Plattform. Zur Programmierung dient die von Microsoft in einer Open-Source-Lizenz zur Verfügung gestellte Quantencomputer-Programmiersprache Q#. Für die Programmierung mittels der Azure Quantum Plattform wird zudem ein ebenfalls quelloffenes Entwicklungskit, von Microsoft als Quantum Development Kit (QDK) bezeichnet, benötigt. Das QDK lässt sich direkt in Entwicklungsumge-

bungen wie Visual Studio oder Jupyter Notebook einbinden.

Um Zugriff auf den direkten Test der Quantencomputing-Plattform zu bekommen ist ein Azure-Konto nötig. Zu Testzwecken stellt Microsoft Probeaccounts zur Verfügung, die mit maximal 200 US-\$ an Azure Credits belastet werden dürfen. Der Probezeitraum gilt 12 Monate und umfasst auch diverse kostenlose Produkte mit unbegrenzter Nutzungsdauer. // SG

Microsoft

Develop Your Software to the Highest Levels of Safety and Security



For 39 years, world-class companies have trusted Green Hills Software's integrated software platforms, engineering services, and certification experts as the foundation to develop and deploy next-generation embedded systems with confidence to the highest levels of safety and security.

Call us on **+49 228 4330 777** or contact us at ghs.com/go/contact



embeddedworld2021

Exhibition & Conference
... it's a smarter world

1-5 March 2021

DIGITAL

Meet Us Online:
ghs.com/go/ew-meet



ERSTE FERTIGUNG AUSSERHALB CHINAS

Huawei baut Fabrik für Mobilfunktechnik im Elsass

Der Mobilfunkkonzern Huawei wird im Businesspark Brumath in der französischen Region Grand Est nahe der deutschen Grenze eine Fabrik für Mobilfunkausrüstungen aufbauen. Konkret soll die „Huawei European Wireless Factory“ den Hauptteil der Mobilfunkbasisstationen für die europäischen Kunden des Unternehmens herstellen. Nach der Ankündigung des Projekts Ende 2019 durch Liang Hua, den Vorstandsvorsitzenden von Huawei, sollen die Bauarbeiten noch in diesem Jahr be-



Zentral: Huawei stärkt seine Präsenz in Europa mit einer neuen Fabrik für Mobilfunkausrüstungen.

ginnen. Die Eröffnung ist laut Huawei für 2023 vorgesehen. Die Fabrik wäre die erste Produktionsstätte dieser Art außerhalb Chinas. Auch der Aufbau von weiteren 110 neuen Forschungs- & Entwicklungsarbeitsplätzen in Irland bis Ende 2022 zeigt die Bedeutung des europäischen Marktes für den Konzern. Nach eigenen Angaben wird Huawei die lokale Wirtschaft unterstützen, „indem es bei der Planung und dem Bau dieses Werks mit französischen Unternehmen zusammenarbeitet“. Darüber hin-

aus plant das Unternehmen, 300 Experten in Frankreich einzustellen, hauptsächlich in der Region Grand Est. Die Belegschaft wird insgesamt 500 Mitarbeiter umfassen und Produkte im Wert von etwa 1 Mrd. Euro pro Jahr herstellen. Laut Huawei ist es ausdrückliches Ziel des Unternehmens, Nachhaltigkeit und Umweltschutz in den Mittelpunkt dieses Projekts zu stellen. Daran wird sich Huawei messen lassen müssen. // ME

Huawei

FACHMESSEN

PCIM Europe 2021 verschoben und die SMTconnect 2021 abgesagt



Lisette Hauser, Mesago: „Wir sind zuversichtlich, mit dem neuen Termin optimale Bedingungen für eine sichere Veranstaltung schaffen zu können“.

Die ursprünglich für den 04. bis 06.05.2021 geplante PCIM Europe Fachmesse und Konferenz wird auf den 31.08. bis 02.09.2021 verschoben. Im September bietet die PCIM Europe der Branche eine Plattform, um sich zu präsentieren und mit Partnern und Kunden persönlich auszutauschen, heißt es. Mit den PCIM Europe digital days wird ein ergänzendes digitales Format als Verlängerung der Messe- und Konferenzteilnahme offeriert. Der Community werden so noch mehr Möglichkeiten geboten,

sich untereinander zu vernetzen und sich über Entwicklungen und Produktneuheiten zu informieren (<https://pcim.mesago.com/events/de.html>). Veranstaltungsort für den neuen Termin bleibt die Messe Nürnberg.

Angeichts des anhaltenden hohen Infektionsgeschehens und die damit verbundenen Restriktionen hat die Mesago Messe Frankfurt sich auch dazu entschieden, die zeitgleich für den 04. bis 06.05.2021 geplante SMTconnect abzusagen. Gespräche mit Ausstellern, aber auch die

Rückmeldungen seitens Besucher, haben gezeigt, dass der Fokus der SMT-Community auf Präsenzveranstaltungen liegt. Diesem Interesse kommt der Veranstalter Mesago Messe Frankfurt nach, weshalb es keine digitale Fachmesse geben wird. Die SMTconnect soll in ihrer gewohnten Form vom 10. bis 12.05.2022 wieder in Nürnberg auf dem Messegelände stattfinden (aktuelle Infos unter www.smtconnect.com). // KU

Mesago Messe Frankfurt

LUFTFAHRT

Die nächste Generation von hybrid-elektrischen Flugzeugen

In der Luftfahrt wird intensiv an elektrischen Antrieben getüftelt. In dem Projekt CHYLA wird untersucht, wie hybrid-elektrische Antriebe in kleinen Flugzeugen der allgemeinen Luftfahrt bis hin zu großen Transportflugzeugen skaliert werden können. Die Technische Universität Braunschweig analysiert dazu zusammen mit der Delft University of Technology, welche Technologien für welche Flugzeugklasse geeignet sind und wie Technologien kombiniert werden können, um die Elektrifizierung von Ver-

kehrsflugzeugen an den Bedarf anzupassen.

In Braunschweig werden Simulations- und Optimierungsmethoden für verschiedene Technologien entwickelt, die in hybrid-elektrischen Flugzeugen zum Einsatz kommen. Zudem wird ein multidisziplinärer Rahmen für die Design-Optimierung erarbeitet, mit dem die Entwicklung hybrid-elektrischer Flugzeuge unter Berücksichtigung der Unsicherheiten bei den zugrunde liegenden Technologien verbessert werden sollen. Darü-



Ein Mittelklasse-Passagierflugzeug: mit neuartigen Flugwerk- und Energietechnologien, entwickelt im Exzellenzcluster SE2A.

Bild: TU Braunschweig / Stanislaw Karpuk

ber hinaus analysieren die Forscherinnen und Forscher Infrastrukturen, die für den Einsatz von hybrid-elektrischen Flugzeugen erforderlich sind, und die damit verbundenen Sicherheitsfragen.

An der TU Braunschweig wird also die Machbarkeit der nächsten Generation von hybrid-elektrischen Flugzeugen bewertet. Berücksichtigt werden dabei Unsicherheiten hinsichtlich der neuen Technologien, die verwendet werden. // TK

TU Braunschweig

ENERGIESPEICHER

Wasserstoffantriebe für E-Scooter und Co.

Wasserstoff gilt als Antrieb der Zukunft. Während bereits erste Wasserstoffautos über deutsche Straßen fahren, ist der bisher übliche Drucktank für E-Scooter jedoch nicht handhabbar. Ein Forscherteam am Fraunhofer-Institut für Fertigungstechnik und Angewandte Materialforschung IFAM in Dresden hat Powerpaste entwickelt, eine Paste, die auf Magnesiumhydrid basiert. In ihr lässt sich Wasserstoff sicher chemisch speichern, einfach transportieren und ohne teure Tankstellen-Infrastruktur nachtanken.

„Mit Powerpaste lässt sich Wasserstoff bei Raumtemperatur und Umgebungsdruck chemisch speichern und bedarfsgerecht wieder freisetzen“, erklärt Dr. Marcus Vogt, Wissenschaftler



Powerpaste: In ihr lässt sich Wasserstoff auf sichere Weise chemisch speichern, einfach transportieren und ohne teure Tankstellen-Infrastruktur nachtanken.

am Fraunhofer IFAM. Das ist auch dann unkritisch, wenn der Roller bei sommerlicher Hitze stundenlang in der Sonne steht, denn Powerpaste zersetzt sich

erst oberhalb von etwa 250 °C. Der Tankvorgang ist denkbar einfach: Statt eine Tankstelle anzusteuern, wechselt der Rollerfahrer einfach eine Kartusche

und füllt zusätzlich Leitungswasser in einen Wassertank – fertig. Das geht auch bequem zuhause oder unterwegs.

Ausgangsmaterial der Powerpaste ist pulverförmiges Magnesium. Bei 350 °C und fünf- bis sechsfachem Atmosphärendruck wird dieses mit Wasserstoff zu Magnesiumhydrid umgesetzt. Nun kommen noch Ester und Metallsalz hinzu. Um das Fahrzeug anzutreiben, befördert ein Stempel die Powerpaste aus der Kartusche heraus. Aus dem Wassertank wird Wasser zugegeben, es entsteht gasförmiger Wasserstoff. Der Clou: Nur die Hälfte des Wasserstoffs stammt aus der Powerpaste, die andere Hälfte liefert das Wasser zu. // TK

Fraunhofer IFAM

RIGOL
Possibilities and More

Jetzt: VNA-Modus für Echtzeit-Spektrumanalysatoren

UltraReal

Unsere Neuen:
Stark in Preis
und Leistung!

Vektor-Netzwerk-Analyse-Modus (VNA, Standard):

- S11-, S21- und Distanz-zu-Fehler-Messung (DTF)
- Smith-, Polar-, SWR- und Gruppenlaufzeit darstellbar

RTSA-Modus (Echtzeit):

- Bis zu 40 MHz Echtzeitbandbreite
- FMT, Density, PVT, Spektrogramm

EMI-Modus (Option)

ab € 7.895,-
plus MwSt.

RSA5032N / RSA5065N

- 9 kHz bis 3,2 oder 6,5 GHz Frequenzbereich

GPSA-Modus (Suchlauf):

- -165 dBm (typ.) mittlere Rauschanzeige (DANL)
- -108 dBc/Hz Phasenrauschen

VSA-Modus (Option)

RSA3015N / RSA3030N / RSA3045N

- 9 kHz bis 1,5 / 3 oder 4,5 GHz Frequenzbereich

GPSA-Modus (Suchlauf):

- -161 dBm (typ.) mittlere Rauschanzeige (DANL)
- -102 dBc/Hz Phasenrauschen

ab € 2.099,-
plus MwSt.

**3 Jahre Garantie –
verlängerbar!**

Registrieren
Sie sich
für unseren
Newsletter →



@

www.rigol.eu

RIGOL Technologies EU GmbH | Telefon +49 8105 27292-0 | info-europe@rigol.com

TECH-WEBINARE

www.elektronikpraxis.de/webinare

Live-Webinar am 23.02.2021

So bauen Sie robuste und intelligente industrielle Netzwerke

Immer mehr Funktionen, eine geringere Latenzzeit und ein robuster EMV-Schutz für eine lange Haltbarkeit auch in rauen Betriebsumgebungen: Dies alles sind Anforderungen, die moderne industrielle Kommunikationsnetzwerke stellen.

Im Webinar am **23.02.2021 um 10.00 Uhr**

- erfahren Sie mehr über die Grundlagen von industriellen Netzwerken,
- bekommen Sie Tipps für den Aufbau robuster und sicherer Netzwerke, welche die Herausforderungen der Protokollübertragung meistern,
- lernen Sie aktuelle, digital isolierte Transceiver-Lösungen kennen und
- erfahren Sie, welche Verbesserungen bei der Isolierung und der Leistung die neuesten isolierten RS-485- und CAN-Transceiver aufweisen.

Referent: Joachim Baumm, Semitron

WHITEPAPER

www.elektronikpraxis.de/whitepaper-elektronik

Moderne Rechenzentren effizient und sicher betreiben
www.elektronikpraxis.de/wp-43812/

Custom-SoCs schaffen Mehrwert für IoT und Industrie 4.0
www.elektronikpraxis.de/wp-43581/

Dossier: Mobilität von morgen
www.elektronikpraxis.de/wp-43821/

Host-Memory-Buffer für SSDs implementieren
www.elektronikpraxis.de/wp-43875/

Wie Luftfeuchte vor Elektrostatik und Viren schützt
www.elektronikpraxis.de/wp-43413/

VERANSTALTUNGEN

www.elektronikpraxis.de/event

Technologietag Leiterplatte

08. - 09. Juni 2021, Würzburg
www.leiterplattentag.de

19. EMS-Tag

10. Juni 2021, Würzburg
www.ems-tag.de

Praxisforum Elektrische Antriebstechnik

23. - 24. Juni 2021, Würzburg
www.praxisforum-antriebstechnik.de

Anwenderkongress Steckverbinder

05. - 07. Juli 2021, Würzburg
www.steckverbinderkongress.de

FPGA-Conference Europe

06. - 08. Juli 2021, München
www.fpga-conference.eu

Batterie Praxis

13. - 14. Juli 2021, Würzburg
www.batterie-praxis.de

SEMINARE

www.b2bseminare.de

Embedded Linux Woche

10. - 12. März 2021, Würzburg
www.b2bseminare.de/160

Embedded Machine Learning

19. März 2021, digital
www.b2bseminare.de/1112

Steckverbinder, das Rückgrat der Elektronik

22. - 23. März 2021, digital
www.b2bseminare.de/1105

C++11 und C++14

14. - 16. April 2021, Leipzig
www.b2bseminare.de/115

Partner und Veranstalter:



Regeleigenschaften für PSR-Sperrwandler verbessern

TIMOTHY HEGARTY *

Bild: TI

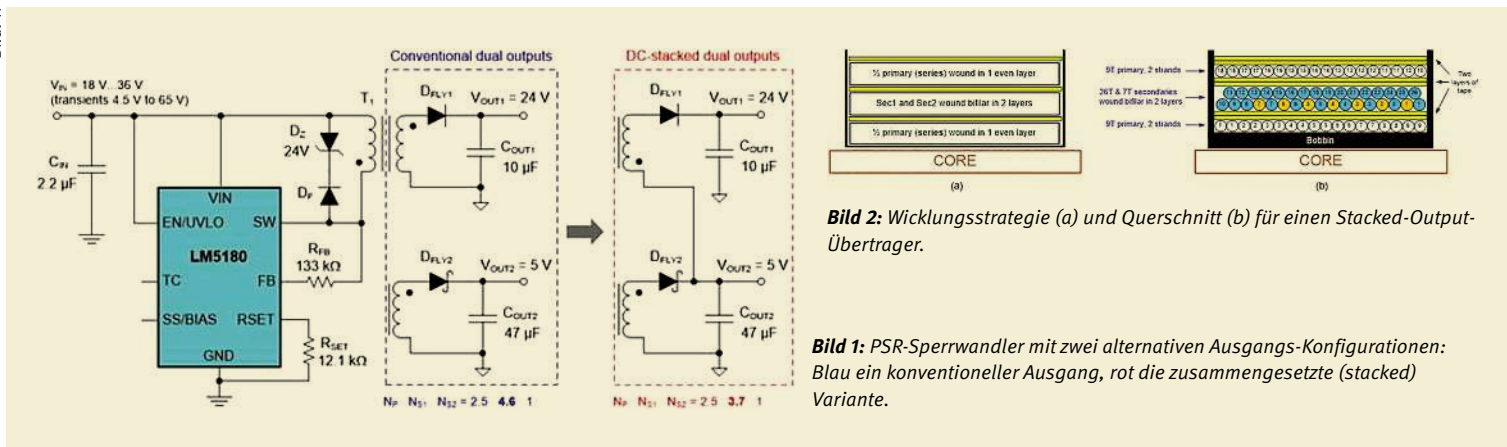


Bild 2: Wicklungsstrategie (a) und Querschnitt (b) für einen Stacked-Output-Übertrager.

Bild 1: PSR-Sperrwandler mit zwei alternativen Ausgangs-Konfigurationen: Blau ein konventioneller Ausgang, rot die zusammengesetzte (stacked) Variante.

Sperrwandler mit primärseitiger Regelung (Primary-Side Regulation, PSR) und mehreren Ausgängen empfehlen sich dank ihrer Zuverlässigkeit und Einfachheit für viele Anwendungsfälle. Da nur der Leistungsübertrager die Isolationsbarriere überquert, wird für die Regelung weder ein Optokoppler noch eine Hilfswicklung, ein Signalübertrager oder eine externe Referenz benötigt. Ein positiver Aspekt ist außerdem der geringe Gesamt-Bauteileaufwand, besonders dann, wenn mehrere isolierte Ausgänge erforderlich sind.

Ohne großen Aufwand lassen sich PSR-Sperrwandler für mehrere Ausgänge und mit flexibler sekundärseitiger Konfiguration entwickeln, unter anderem mit bipolaren Ausgängen von beispielsweise $\pm 12\text{ V}$ oder $+15\text{ V}/-8\text{ V}$, oder auch mit zwei positiven Ausgängen (z.B. $3,3\text{ V}/12\text{ V}$ oder $5\text{ V}/12\text{ V}$). Die Ausgänge sind die entweder mit gemeinsamer Masse oder als unabhängige Ausgänge mit unterschiedlichen Massepotenzialen realisiert.

Dazu ein Beispiel. Der auf dem LM5180 basierende PSR-Sperrwandler in Bild 1 stellt mit der linken, sekundärseitigen Schaltung

(blau umrandet) Ausgangsspannungen von 24 V (V_{OUT1}) und 5 V (V_{OUT2}) bereit. Die rechte, rot umrandete Ausgangsschaltung dagegen basiert auf dem Stacked-Output-Prinzip, wobei die obere Sekundärwicklung eine Spannung von 19 V liefert, die sich zu den 5 V des unteren Ausgangs addiert. Die Windungsverhältnisse sind jeweils angegeben.

Durch Erfassen der Kniepunktspannung zum richtigen Zeitpunkt lässt sich eine sehr genaue Regelung der Ausgangsspannung unter dem Einfluss von Netz-, Last- und Temperaturschwankungen erreichen. Allerdings ist die Entmagnetisierungszeit bei mehreren Ausgängen nicht eindeutig definiert, denn jede Wicklung hat ihr eigenes Entmagnetisierungs-Intervall, und die Spannung wird am Ende der letzten Entmagnetisierung erfasst.

Stacked Winding: Die Vorteile des Prinzips

Die Stacked-Winding-Variante kommt mit weniger Windungen aus, ihre Streuinduktivität ist geringer, und die für die Kreuzregelung wichtige Kopplung von einer Sekundärwicklung zur anderen ist genauer. Da überdies der negative Spannungsaussschlag an der Flyback-Diode geringer ist, werden die Gleichtakt-Störgrößen eingedämmt, und die beim Einschalten des Schalters an die Primärseite reflektierte, kapazitiv bedingte Stromspitze wird gemildert.

Die Streuinduktivität sowohl zwischen der Primär- und der Sekundärwicklung als auch zwischen den Sekundärwicklungen lässt sich mit verschiedenen Maßnahmen minimieren. Zunächst sollte die Primärwicklung (bzw. die Wicklung mit der höchsten Windungszahl) verschachtelt gewickelt werden. Außerdem sollten komplette Lagen bewickelt werden, um die Kopplung von Lage zu Lage zu optimieren (Auffüllen der Lagen mit Litzen).

Der Kern sollte ferner einen breiten Spulenkörper aufweisen, um mit möglichst wenig Lagen auszukommen. Die Streuinduktivität zwischen den Sekundärwicklungen lässt sich reduzieren, indem man die Wicklungen bifilar oder multifilar auf derselben Lage wickelt und Sekundärwicklungen mit niedrigen Strömen zwischen Lagen von Sekundärwicklungen mit hohen Strömen anordnet.

Bild 2 zeigt den Querschnitt durch einen Übertrager für das Design in Bild 1. Die Primärwicklung ist im Verhältnis 50:50 aufgeteilt, und die Sekundärwicklungen sind bifilar gewickelt.

Mit einem speziellen Rechen-Tool lässt sich die Auswahl und Optimierung der Übertrager-Windungsverhältnisse, der Magnetisierungs-Induktivität und der Wicklungswiderstände vereinfachen.

// KR

Texas Instruments



* Timothy Hegarty
... arbeitet als Applikationsingenieur im Geschäftsbereich „Buck Switching Regulators“ bei Texas Instruments in Phoenix / USA.



TITELSTORY

Die von Toshiba entwickelte Methode der symmetrischen PWM-Trägersignale macht hochfrequente Signale überflüssig und vermeidet so elektromagnetische Störungen, Schwankungen des Drehmoments und die daraus resultierenden Störgeräusche. Dazu wird die spezielle Fähigkeit eines integrierten Motortreibers genutzt sowie unter-

schiedliche aber synchrone PWM-Trägersignale für die einzelnen Motorphasen. Zusammen mit einer Vector Engine und synchron gekoppelten A/D-Konvertern sorgt diese Methode für eine sensorlose, feldorientierte Motorregelung, die beim Anlaufen und bei niedrigen Drehzahlen für optimales Drehmoment bei minimaler Geräuschentwicklung sorgt.

Deutlich weniger Störgeräusche bei der feldorientierten Regelung

Um die Probleme mit Vibrationen und Störgeräuschen bei der Messung der magnetischen Salienz zu beseitigen, hat Toshiba eine alternative Methode entwickelt, die ohne hochfrequente Signale auskommt.

PATRICK OSTERLOH *

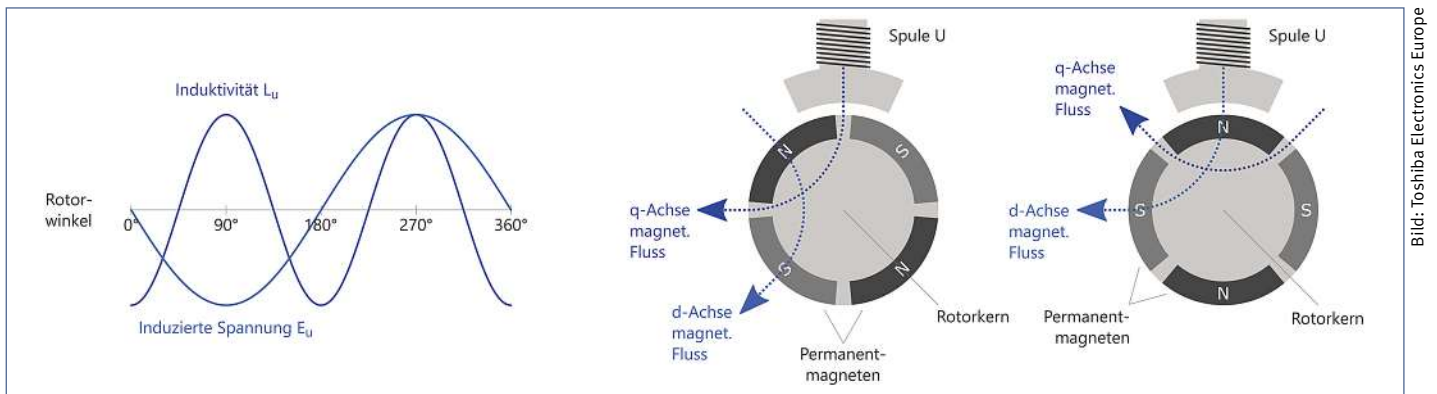


Bild 1: Durch die magnetische Salienz variiert die Induktivität der Statorspule mit der doppelten Rotationsfrequenz und kann so zur Bestimmung der Rotorposition benutzt werden.

Durch den Einsatz bürstenloser, elektrisch kommutierter Motoren anstelle von Bürstenmotoren lassen sich signifikante Energieeinsparungen bei gleichzeitig reduziertem Wartungsaufwand erzielen. Allerdings entfällt durch den Wegfall der mechanischen Bürsten auch die implizite Information über die richtigen Zeitpunkte zur Kommutierung, sodass auf alternative Ansteuerungstechniken zurückgegriffen werden muss. Idealerweise liefern Hall-Sensoren, Resolver oder Encoder die benötigte Information über die Stellung des Rotors, die dann zur Steuerung der Kommutierung benutzt wird. Es gibt aber auch viele Anwendungen, bei denen aus Kosten-, Zuverlässigkeits- oder Platzgründen der Einsatz solcher Messgeber nicht möglich ist. Im Laufe der Jahre sind viele innovative Lösungen für dieses Problem entstanden, die im Wesentlichen alle auf dem Prinzip der in den Statorspulen induzierten elektromotorischen Gegenkraft (Gegen-EMK) basieren.



* Patrick Osterloh
... ist Senior Manager bei Toshiba Electronics Europe.

Bei einem Permanentmagnet-Synchronmotor (PMSM) generieren die auf oder in dem Rotor angebrachten Permanentmagneten eine elektromotorische Gegenkraft in den Statorspulen, sobald sie die Statorpole passieren. Die von der Gegen-EMK in den Statorspulen erzeugte Spannung, deren Amplitude unmittelbar von der Rotationsgeschwindigkeit abhängt, kann mit Hilfe von A/D-Konvertern in den für die Motorregelung eingesetzten Mikrocontrollern gemessen werden, sodass sich auf die Position des Rotors zurückschließen lässt. Das Problem bei diesem Konzept liegt darin, dass bei Stillstand des Motors oder bei niedrigen Rotationsgeschwindigkeiten die von der Gegen-EMK erzeugte Spannung faktisch nicht oder nur sehr ungenau messbar ist.

Folglich kann die Position des Rotors (und seiner Magnete) im niedrigen Drehzahlbereich nicht unmittelbar ermittelt werden, wodurch die optimale Regelung des Motors in diesen Situationen erschwert wird. Natürlich ist es immer möglich, den Motor durch eine erzwungene Kommutierung (wie z.B. der Blockkommutierung) anlaufen zu lassen, bis die Rotationsgeschwindigkeit hoch genug ist, um eine ausreichende Gegen-EMK zu erzeugen, die dann zur Positionsbestimmung

genutzt werden kann. Dieser Ansatz funktioniert zwar, erzeugt aber erhebliche Schwankungen des Drehmoments, die durch den sich fortwährend verändernden Winkel zwischen den magnetischen Feldern von Stator und Rotor hervorgerufen werden und die Vibrationen und störende Geräusche mit sich bringen. In vielen Anwendungen wie Kompressoren, Pumpen oder Frequenzumrichtern, die in unmittelbarer Umgebung von Menschen eingesetzt werden, soll der mechanische Stress und die zusätzliche Lärmbelastigung vermieden werden, um die Zuverlässigkeit der Produkte und die Akzeptanz bei den Endverbraucher*innen zu steigern.

Eine kurze Exkursion zur feldorientierten Regelung

Hauptziel der bei der Ansteuerung von bürstenlosen DC-Motoren eingesetzten feldorientierten Regelung (Field Oriented Control, FOC) ist es, einen konstanten Winkel von 90° zwischen den magnetischen Feldern von Stator und Rotor herzustellen und beizubehalten, um so ein optimales und konstantes Drehmoment zu erreichen. Dazu werden vom Motortreiber drei verschiedene, sinusförmige Spannungen erzeugt, die an die drei Statorspulen angelegt werden. Diese

Bild: Toshiba Electronics Europe

Bild 2:

Bei der INFORM-Methode werden die hochfrequenten Signale nur während der Pausen der pulsweitenmodulierten Phasensignale angelegt, was immer noch zu unerwünschten Störgeräuschen führen kann.

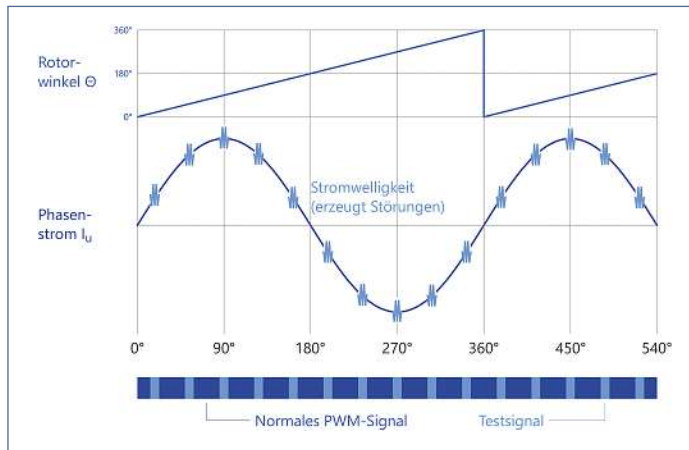


Bild: Toshiba Electronics Europe

ändert wird. Die Induktivität weist dabei ein Maximum auf, wenn Statorspule und die magnetische Flussachse des Rotors senkrecht zu einander stehen (bei 90° und 270°), und ein Minimum, wenn diese parallel zu einander stehen (bei 0° und 180°). Daraus ergibt sich ein vom Rotationswinkel abhängiger, sinusförmiger Verlauf der Spuleninduktivität, die mit der doppelten Frequenz des Rotors variiert (Bild 1). Dieser Effekt, der als magnetische Saliency (magnetic saliency) bezeichnet wird, ist durch Anlegen einer hochfrequenten und den Motorphasen überlagerten Spannung bestimmbar. Die sich dabei ergebenen Änderungen in den Motorströmen lassen sich messen, um so die Winkelposition des Rotors mit hinreichender Genauigkeit zu bestimmen.

Bild 3:

Durch Verwendung von drei unterschiedlichen, symmetrischen PWM-Trägersignalen kann die Veränderung der Spuleninduktivität durch differentielle Messung um den gemeinsamen Bezugspunkt bestimmt werden.

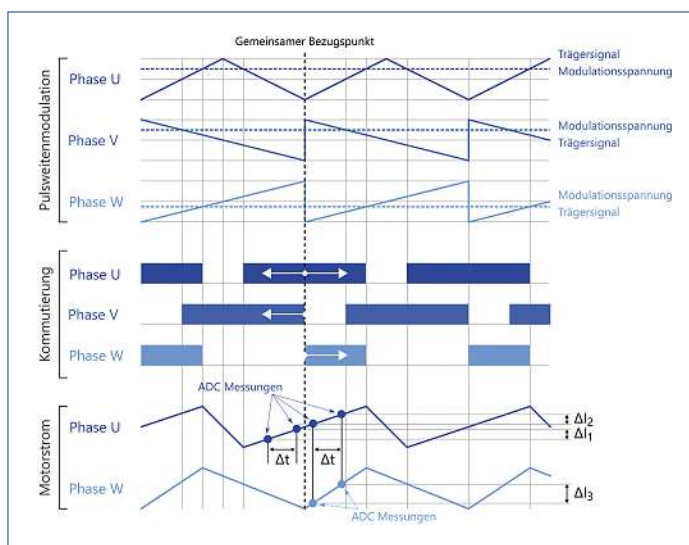


Bild: Toshiba Electronics Europe

Diese Methode hat aber auch Nachteile. Die zur Bestimmung der Induktivität erforderlichen, hochfrequenten Signale erzeugen nicht nur vermehrt elektromagnetische Störungen, sondern auch Schwankungen des Drehmoments, was wiederum Vibrationen und Störgeräusche nach sich zieht. Entwickler*innen von Motoranwendungen werden damit wieder mit genau den Problemen konfrontiert, die sie eigentlich durch den Einsatz der feldorientierten Regelung lösen wollten. Einen verbesserten Ansatz bietet die INFORM-Methode (Indirect Flux detection by Online Reactance Measurements), bei der die hochfrequenten Signale nicht permanent, sondern nur gelegentlich in den Pausen der PWM-Steuersignale angelegt werden. Aber auch dieser Ansatz kann die Schwankungen im Drehmoment und die sich daraus ergebenen Vibrationen und Störgeräusche nicht gänzlich vermeiden (Bild 2).

drei Phasen sind zu einander um 120° phasenverschoben, sodass ein gleichmäßig rotierendes Magnetfeld im Stator entsteht. Der Motortreiber selbst wird durch pulsweitenmodulierte Signale angesteuert, die von einem Mikrocontroller erzeugt werden. Moderne Mikrocontroller beinhalten daher optimierte Schaltungen, mit denen die Erzeugung der PWM-Signale mit geringem Software-Aufwand realisierbar ist.

Die Berechnung und der Umgang mit den zeitabhängigen, sinusförmigen Signalen ist aber sehr aufwendig. Die vom A/D-Konverter gemessenen Motorströme werden deshalb zunächst mit Hilfe der Clarke/Park-Transformationen in ein zeitunabhängiges, sich mit dem Rotor drehendes Bezugssystem umgewandelt, das aus den Achsen für den magnetischen Fluss (direct axis, abgekürzt d) und für das Drehmoment (quadrature axis, abgekürzt q) besteht. Mit PI-Reglern können so relativ einfach die Position, Geschwindigkeit und das Drehmoment des Rotors ermittelt, mit den Sollwerten verglichen und die für

den optimalen Arbeitspunkt erforderlichen Motorspannungen angepasst werden. Anschließend wird das Ergebnis wieder durch inverse Park/Clarke-Transformationen in das zeitabhängige, dreiphasige Bezugssystem des Stators überführt, bevor es dann in ein pulsweitenmoduliertes Signal umgewandelt wird, das den Motortreiber ansteuert.

Sensorlose Motorregelung bei niedrigen Drehzahlen

Herkömmliche sensorlose, feldorientierte Regelalgorithmen arbeiten (wie erwähnt) am besten, wenn die Rotationsgeschwindigkeit hoch genug ist, sodass die Gegen-EMK ein brauchbares Spannungssignal erzeugt. Bei niedrigeren Geschwindigkeiten muss auf andere Algorithmen zurückgegriffen werden, die alternative Methoden zur Bestimmung der Rotorposition nutzen. Ein fortschrittlicher Ansatz, der heute häufig zum Einsatz kommt, ist die Messung der Induktivität der Statorspulen, die durch das sich vorbeibewegende magnetische Feld des Rotors ver-

Pulsweitenmodulation mit symmetrischen Trägersignalen

Um die Probleme mit Vibrationen und Störgeräuschen bei der Messung der magnetischen Saliency zu beseitigen, hat Toshiba eine alternative Methode entwickelt, die ohne hochfrequente Signale auskommt, indem sie eine besondere Eigenschaft der in der neuen TXZ+-Mikrocontrollerfamilie enthaltenen PWM-Module ausnutzt. Während die PWM-Module der meisten Mikrocontroller das gleiche Trägersignal für die Ansteuerung aller drei Motorphasen benutzen, können die Bausteine aus der TXZ+-Familie individuelle Trägersignale für jede der drei Motorphasen verwenden. Jede einzelne Phase kann so gleichzeitig auf Trägersignale in Form eines Dreiecks, eines ansteigenden Sägezahns oder eines abfallenden Sägezahns für die Pulsweitenmodulation zurückgreifen, während alle drei Phasen einen gemeinsamen

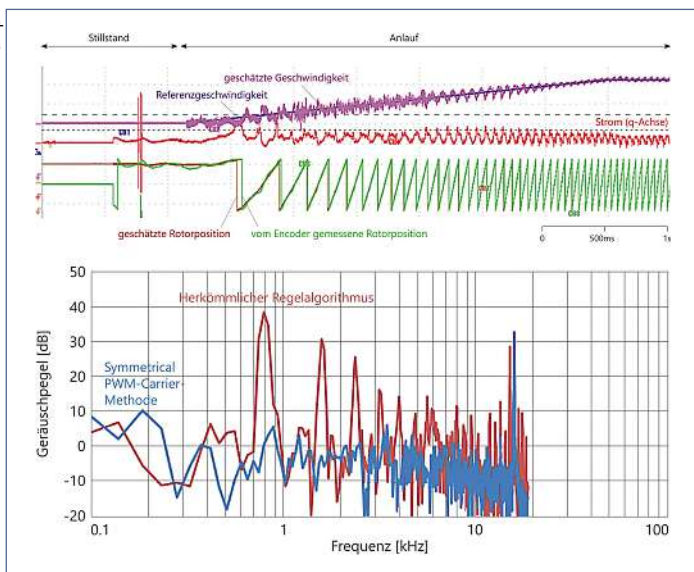


Bild 4: Rotorpositionsschätzung der Symmetric-PWM-Carrier-Methode im Vergleich zu der von einem Messgeber bestimmten Position (oben) und der beim Motoranlauf gemessene Geräuschpegel im Vergleich zu konventionellen, sensorlosen Motorregelungen (unten).

Bezugspunkt behalten. Da während der Strommessungen jeweils nur zwei der drei Phasen aktiv sind, kann der Motorstrom durch die differentielle Messung der zwei jeweils aktiven Phasen (z.B. U und V) vor und nach dem gemeinsamen Bezugspunkt bestimmt und die Veränderung der Spulen-Induktivität daraus ermittelt werden (Bild 3).

Die Bestimmung der Rotorposition mit Hilfe dieses auf symmetrischen PWM-Trägersignalen basierten Ansatzes hat sich als sehr genau herausgestellt. Darüber hinaus wurde bewiesen, dass mit dieser Methode ein wesentlich niedriger Geräuschpegel als bei den anderen hier vorgestellten Methoden zur Bestimmung der magnetischen Salienz erreicht werden kann (Bild 4).

Integrierte Mikrocontroller für die Motorregelung

Diese Methode der sensorlosen, feldorientierten Regelung kann durch Einsatz von Mikrocontrollern aus den zur Toshiba-TXZ+-Familie gehörenden M3H- und M4K-Gruppen realisiert werden. Ausgestattet mit Arm-Cortex-M-Prozessorkernen sind diese Bausteine speziell für die feldorientierte Regelung von PMSM- und BLDC-Motoren konzipiert. Entscheidend für die Implementierung des auf symmetrischen PWM-Trägersignalen basierenden Regelalgorithmus ist der in den M3H- und M4K-Bausteinen integrierte Advanced-Programmable-Motortreiber A-PMD. Neben den schon erwähnten, für jede Phase individuell auswählbaren Trägersignalen, beinhaltet dieses PWM-Modul auch Funktionen für die Synchronisierung mehrerer Motorkanäle und der dabei verwendeten A/D-Konverter. So kann sichergestellt werden, dass alle erforderlichen Strommessun-

gen zu dem richtigen Zeitpunkt durchgeführt werden. Außerdem sind im A-PMD auch Funktionen zur Kontrolle der Stromleitung, für die funktionale Sicherheit und für die Totzeit-Generierung enthalten. Eine weitere Vereinfachung der sensorlosen, feldorientierten Motorregelung wird durch die integrierte Advanced Vector Engine Plus (A-VE+) ermöglicht. Mit diesem für die feldorientierte Regelung spezialisierten Coprozessor können ein Großteil der benötigten Berechnungen in Hardware ausgeführt werden. Dazu zählen neben den (inversen) Clarke/Park-Transformationen auch PI-Regler, Space-Vector-Modulation sowie die Bereitstellung von praktischen und schnellen trigonometrischen Funktionen. Die A-VE+ verfügt auch über einen eigenen Taskscheduler, mit dessen Hilfe sie im Zusammenspiel mit A-PMD und A/D-Konvertern den überwiegenden Teil der sensorlosen, feldorientierten Regelaufgaben selbstständig ausführen kann.

Fazit: Die Methode der symmetrischen PWM-Trägersignale macht hochfrequente Signale überflüssig und kann so elektromagnetische Störungen, Schwankungen des Drehmoments und die daraus resultierenden Störgeräusche vermeiden. Dazu wird die Fähigkeit des A-PMD-Moduls genutzt, unterschiedliche aber synchrone PWM-Trägersignale für jede Motorphase zu verwenden. Zusammen mit dem A-VE+-Coprozessor und den synchronen A/D-Konvertern ermöglicht diese Methode eine sensorlose, feldorientierte Regelung, die auch beim Motoranlauf und bei niedrigen Drehzahlen für das optimale Drehmoment bei minimaler Geräuschentwicklung sorgt.

// KU

Toshiba Electronics Europe

IoT Multitalent



Wireless Gecko – SoCs und Module...

Überlegene Energie-Effizienz bei ARM® Cortex®-M4 mit Ultra-Low Power Modi. Sparsam im Verbrauch und vielseitig bei der Funkverbindung:

- ▶ Bluetooth Low Energy, Bluetooth 5
- ▶ Zigbee, Thread, Bluetooth Mesh
- ▶ WiFi
- ▶ Multi-Protokoll Support

Hier erhalten Sie **roten Support** für grüne Wireless SoC und Module:

www.glyn.de | wireless@glyn.de



GLYN
High-Tech Distribution

Sicherheit bei der Eingabe an einem Touchscreen umsetzen

Die Eingabe von Informationen an einem Touchscreen ist praktisch. Funktionssicherheit und Bediensicherheit schützen vor gefährlichen Fehleingaben und sorgen für die entsprechende Sicherheit.

EBERHARD SCHILL *



Bild: Kyocera

Kunden genau besprochen werden. Der Fokus liegt auf folgenden Komponenten:

■ **Backlight:** Es bestimmt, wie hell ein Display leuchtet. Helle Umgebungsbedingungen (im Freien) erfordern hohe Helligkeiten, wechselnde Beleuchtungsbedingungen eine umfassende Dimmbarkeit des Displays. Moderne Backlights bestehen aus LEDs, die zur Erhöhung der Ausfallsicherheit in mehreren Ketten/Strängen betrieben werden. Sollte ein LED-Strang ausfallen, wird das Display zwar dunkler, bleibt aber ablesbar.

■ **Optische Folien und Displayelektronik:** Für einen Betrieb in feuchter Umgebung müssen besonders stabile Polarisationsfilter benutzt und entsprechend verarbeitet werden, da sonst eindringende Feuchtigkeit diese für den Displaybetrieb essenziellen Folien zerstört. Im Extremfall kann sogar eine Versiegelung notwendig sein. Analog dazu ist es möglich, Elektronik-Komponenten oder Leiterbahnen zu verkapseln. In manchen Anwendungen werden Feuchtigkeitsabsorber im Gerätegehäuse vorgesehen.

■ **Mechanischer Aufbau:** Ein HMI in Arbeitsmaschinen muss dem mechanischen Aufbau und der Vibrationsfestigkeit Rechnung tragen. In einem Standard-Display sind alle Komponenten mit einem leichten Spiel verbaut, damit sich bei thermischer Ausdehnung keine Verspannungen zwischen Glas und Metallrahmen ergeben. Dieser Ausdehnungsspielraum kann zu einem klappern und Abrieb führen. Das Glas ist bei Industriedisplays meist federnd im Plastikrahmen des Backlights gelagert. Mechanischer Stress des Displays sollte vermieden werden, da die Anzeigequalität (Mura) beeinträchtigt wird und im Extremfall das Display sogar brechen kann. Bei EMV-kritischen Anwendungen muss außerdem darauf geachtet werden, dass sich Display und Elektronik in einem Metallgehäuse befinden. Die Display-Leiterplatte

Sicherheit bei der Eingabe: Landmaschinen oder Baumaschinen werden mittlerweile per Touchscreen bedient. Dabei darf es zu keiner Fehlbedienung kommen.

Über das Smartphone lassen sich heute im vernetzten Heim Lampen, Fernseher oder Heizungen steuern. Mit dem Smartphone oder Tablet jedoch ein Fahrzeug oder gar eine Maschine bedienen? Auf diese Idee käme wohl niemand. Touch-Funktionen sind hier nicht eindeutig und es kann zu einer Fehlbedienung kommen. Einige Funktionen werden nicht direkt beim ersten Touch ausgeführt. Auf Bedienterminals sind für sicherheitsrelevante Anwendungen mechanische Schalter vorhanden.

Schalter werden nur dargestellt, wenn sie gebraucht werden, Symbole und Hilfstexte

erleichtern den Ablauf. Maschinenzustände lassen sich grafisch oder per Kamera anzeigen, Regler und Schieber wie ihre analogen Pendanten einstellen. Idealerweise sollte eine Bedienterminal zu einhundert Prozent aus einem Display bestehen und dennoch verlässlich funktionieren.

Die zwei Arten von Sicherheit bei einer HMI-Schnittstelle

Zwei Arten von Sicherheit: Die **Funktionssicherheit** einer HMI-Schnittstelle hängt von der Hardware des Bedientdisplays ab. Für eine uneingeschränkte Funktionalität über den geplanten Lebenszyklus sind langlebige Komponenten notwendig, die für die spezifizierte Anwendung ausgelegt sind. Um eine passende Auswahl von Komponenten zu treffen, müssen in der Design-in-Phase die geplanten Einsatzbedingungen mit dem



* Eberhard Schill
... ist Manager Marketing und Distribution bei Kyocera in Dietzenbach, Deutschland.

sollte nach EMV-Gesichtspunkten ausgelegt sein. Meist sind diese Leiterplatten mehrschichtig. Offen liegende Flexleiterplatten, wie bei einem Laptop, müssen zusätzlich abgeschirmt werden.

■ **Displayalterung:** Kritisch ist ein Einbrennen des Bildes. Skalen oder Icons befinden sich während des gesamten Betriebszeitraumes unverändert an derselben Position. Neben einem Bildschirmschoner kann über Software das gesamte Bild nach gewisser Zeit regelmäßig um eine Pixelbreite horizontal und/oder vertikal verschoben werden.

Eingabe zurück an den Anwender melden

Die **Bediensicherheit** einer HMI-Schnittstelle hängt von einer ergonomisch gestalteten Softwareoberfläche ab. Industrieterminals sind Arbeitsgeräte und müssen leicht und ermüdungsfrei bedienbar sein. Dazu sollte die HMI gut in der Hand liegen und alle virtuellen oder mechanischen Tasten gut erreichbar sein. Informationen sollten gut ablesbar und strukturiert in der Mitte dargestellt werden. Bei der Anordnung sollten

häufig genutzte Bedienelemente so positioniert sein, dass sie gut für Zeigefinger oder Daumen zu erreichen sind und Unterarmen vermieden werden. Tastengröße und -abstand sollten eine Fehlbedienung möglichst ausschließen.

Ein Touchpanel speziell in der Industrie bleibt im täglichen Einsatz nicht sauber und lässt sich nicht mehr gut ablesen. Mit einer Fett abweisenden Oberflächen lässt sich das vermeiden, womit eine verschmierte Oberfläche reduziert wird. In Kombination mit einem entsprechenden Controller lässt sich ein Touchdisplay selbst unter Benetzung mit Flüssigkeiten mit mehreren Fingern oder des Handballens bedienen. Selbst Handschuhe sind möglich.

Wichtig ist eine zeitnahe Rückmeldung von Seiten der Schnittstelle, die eine erfolgte Eingabe bestätigt. Das erfolgt optisch oder akustisch und ist bei Maschinen mit einer gewissen Reaktionszeit wichtig, wo keine unmittelbar erkennbare Reaktion erfolgt. Wer bisher mechanische Schalter betätigt hat, tut sich schwer, nach der Berührung einer Glasoberfläche auf eine korrekte Eingabe zu vertrauen. Für das gewohnte hapti-

sche Feedback zurück zum Anwender hat der Hersteller Kyocera eine spezielle Technik entwickelt, die das simuliert. Die patentierte Haptivity Technologie erzeugt den gewohnten Tastenklick auf beliebigen Oberflächen, sodass sich auf einem Display dargestellte Tasten und Schalter so anfühlen, wie der Benutzer es von mechanischen Bedienelementen gewohnt ist. Dazu wird die berührte Oberfläche nach Überschreiten einer bestimmten Druckschwelle durch einen Piezo als Aktuator gezielt in eine kurze Vibration versetzt. Durch Variation der Ansteuerparameter des Piezos lässt sich der wahrgenommene Klick durch den Anwender in weiten Grenzen per Software einstellen.

Die einstellbare Auslösekraft einer Berührung verhindert Fehlbedienungen. Die technisch mögliche Implementierung mehrstufiger Schalter oder Slider hilft, die HMI benutzerfreundlicher zu gestalten. Außerdem helfen eingearbeitete Mulden in der Glasoberfläche, um den Finger zu führen und die Bedienoberfläche nahezu blind zu ertasten. // HEH

Kyocera



Mouser hat etwas, das anderen Bezugsquellen fehlt

Alles für Ihre Stücklisten

mouser.de/available-to-ship

Mouser ist ein autorisierter Distributor
für Halbleiter und elektronische Bauelemente.



Worauf bei der Wahl eines Grafikdisplays zu achten ist

Bei der Wahl eines elektronischen Displays sollten Entwickler einige Kriterien unbedingt berücksichtigen. Denn jede Anwendung hat ihre spezifischen Anforderungen.

THOMAS BILLER *

Grafikdisplay:

Displays für industrielle Anwendungen gibt es in vielen Ausführungen.



schen der Größe der Anzeige selbst, also Zeichenhöhe bzw. Sichtbereich, Viewing Area, und den Außenabmessungen, also den Einbaumaßen, der Anzeige unterschieden werden. Bei den gebräuchlichen Ablesentfernungen limitieren oftmals die Applikationen, das Gerät oder das Gehäuse die Anzeigengröße. So definieren beispielsweise

Handheld-Geräte festgelegte maximale Abmessungen oder die Geräte-/Einbaugröße für die Anzeige. Ist die Entfernung zwischen Anwender und Display hoch, dann ist eine bestimmte Mindestgröße der Anzeige notwendig.

Sind die äußeren Abmessungen noch nicht durch das zu entwickelnde Gerät vorgegeben, erfolgt die Dimensionierung der Anzeigengröße unter Berücksichtigung von Aspekten wie: vorgesehene Ablesentfernung, Auflösungsvermögen und Gesichtsfeld des Auges (Scharfsicht ohne Kopf- oder größere Augenbewegungen).

Da das Auflösungsvermögen des menschlichen Auges bis zu einer Winkelminute (1/60 Grad), das horizontale Gesichtsfeld ungefähr 40° und das vertikale ungefähr 30° betragen, ergeben sich folgende Werte:

■ Für normale Ablesentfernungen von 30 bis 50 cm: maximal Anzeigengröße 22 cm x 16 cm bis 36 cm x 27 cm, minimale Auflösung des Auges: 0,09 mm bis 0,15 mm.

■ Für größere Ablesentfernungen ab 5 m: maximale Anzeigengröße: 3,6 m x 2,7 m, minimale Auflösung des Auges: 1,5 mm.

Um die Anzeige mit einem Blick vollständig erfassen zu können, darf die maximale Anzeigengröße nicht überschritten werden. Damit Anwender die Zeichen einer alphanumerischen Darstellung problemlos erkennen, sollte die minimale Zeichengröße ein Mehrfaches, ungefähr ein Faktor 10 bis 20, der minimalen Auflösung des Auges betragen (Bild 1). Bei Grafikdisplays ist die mini-

Die Schritte bei der Wahl eines passenden elektronischen Displays sind zwar variabel, aber es kommt auf einige wichtige Punkte an. Sollen nur Zahlen bzw. Zahlen und Buchstaben auf dem Display abgebildet werden, dann kommt aus Kostengründen eine Sieben-Segment- oder Dot-Matrix-Anzeige infrage. Hier muss der Entwickler genau prüfen, ob eine solche Anzeige ausreichend ist. Sollen künftig auch Grafiken auf dem Display angezeigt werden, dann muss die Entscheidung Richtung Grafikdisplay gehen. Es kommt nicht allein auf den Inhalt an, der auf dem Display angezeigt

wird. Wichtig für die Wahl des Displays ist der Umfang der darzustellenden Informationen.

Bei alphanumerischen Anzeigen umfasst dieser Punkt die Anzahl der erforderlichen Zeichen und Zeilen, bei Grafik-Anzeigen die Anzahl der für eine klare Darstellung erforderlichen Bildpunkte. Zusammen mit der benötigten Auflösung legt der Entwickler damit bereits das spätere Format fest: Hoch- oder Querformat.

Wie sich das Display korrekt dimensionieren lässt

Wenn der Darstellungsumfang festgelegt ist, dann folgt für die Auswahl das Seitenverhältnis. Durch unterschiedliche Segment- bzw. Pixelgeometrien können die Hersteller der Displays verschiedene Erscheinungsbilder und Seitenverhältnisse für gleiche Zeichengrößen oder Auflösungen in gewissen Grenzen umsetzen. Damit die Anzeigengröße festgelegt werden kann, muss zunächst zwi-



* Thomas Biller
... arbeitete als Produkt-Manager bei Schukat electronic.

male Auflösung wichtig in Bezug auf die maximale Pixelgröße. Beträgt sie ein Vielfaches der minimalen Auflösung, wird die Grafik verpixelt wahrgenommen.

Ansteuerschnittstelle und Versorgungsspannung

Sobald feststeht, welcher Mikrocontroller die Anzeige ansteuert, ist in der Regel eine Schnittstelle vorgegeben: 8-Bit-Parallel, SPI oder I²C. Das schränkt die Display-Auswahl stark ein, denn eine einfache, langsamere Schnittstelle kann kein Farb- oder Grafikdisplay mit hoher Auflösung und Farbtiefe ansteuern. Bei den aktuellen Displays besteht oft die Wahl, diese über verschiedene Interfaces anzusteuern. Sofern der Entwickler bei der Schnittstellenwahl noch ungebunden ist, lässt sich am Ende des Designprozesses entscheiden, welche Schnittstelle am besten geeignet ist.

Ein anderer begrenzender Faktor für die Wahl des Displays ist die Versorgungsspannung. Für spezielle Anwendungen können bereits eine oder mehrere feste Versorgungsspannungen für alle Komponenten festgelegt

sein. Viele Anzeigen hingegen kommen mit einer der gängigen Versorgungsspannung von 3,3 V oder 5,0 V aus. Ob das Display Farben darstellen soll, kann durch das Corporate Identity eines Unternehmens vorgegeben sein. Grundsätzlich gilt es zu entscheiden, ob eine mehrfarbige Anzeige, also die gleichzeitige Darstellung verschiedener Farben, oder eine monochrome Anzeige, etwa eine reine Zahlen-/Textanzeige bzw. Dot-Matrix-Anzeige, gewünscht ist. Bei monochromen passiven LC-Displays stellt sich die Frage nach der positiven oder negativen

Darstellung, ob weiße Zeichen auf dunklem Hintergrund oder umgekehrt dargestellt werden sollen. Bei passiven LCD-Anzeigen und negativer Darstellung erscheinen die Zeichen bzw. Pixel in der Farbe des Backlights, das hier unbedingt erforderlich ist.

Soll zur Minimierung der Leistungsaufnahme bei einer negativen Darstellung auf eine LCD-Hinterleuchtung verzichtet werden, empfiehlt sich der Einsatz einer OLED-Anzeige. Hier sind die Zeichen/Pixel selbstleuchtend, der Leistungsverbrauch ist wesentlich geringer und der Kontrast besser als bei vergleichbaren LC-Displays. Allerdings fallen höhere Kosten an.

Bei mehrfarbigen Grafikdisplays wie 16,7M, 262k, 256 oder Farben besteht die Wahl zwischen verschiedenen Display-Techniken: mehrfarbige OLED-Anzeigen, passive

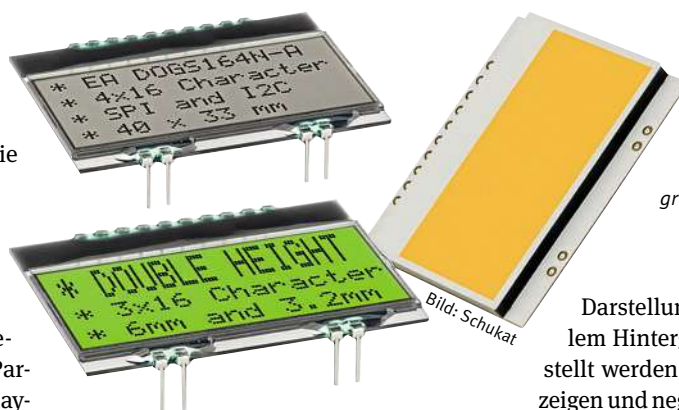


Bild 1:
Damit die Anzeige mit einem Blick vollständig erfasst werden kann, darf die maximale Anzeigengröße nicht überschritten werden.

Ungewöhnliche Display-Formate: Neue Bar-type TFTs

C O D I C O[®]

- Größen: 3.9" / 4.6" / 6.5"
- Auflösungen:
480*128 / 800*320 / 800*320 Pixel
- Seitenverhältnisse: 15:4 / 10:4 / 10:4
- Optional mit kapazitivem Touch
- Temperaturbereich: -20 bis +70°C



Kontakt: +43 1 86 305-0 | office@codico.com | www.codico.com/shop

YB Yeebo
International
Holdings Ltd.

Adobe Stock/© Zinnikon & Co. Grafikhoff

LCD-Anzeigen/CSTN, Aktiv-Matrix-Displays/TFT und mehr. Mit einer Grafikanzeige kann das Display gleichzeitig als Eingabeeinheit bzw. Human Machine Interface (HMI) dienen. Entscheidet sich der Entwickler für diese Option, spart das möglicherweise Kosten für alternative Eingabekomponenten. In diesem Fall fällt die Wahl auf ein Grafikdisplay mit Touchpanel (PCAP oder resistiv).

Auf die Umgebungsbedingungen kommt es an

Ebenfalls ausschlaggebend für die Wahl des Displays sind die zu erwartenden Umgebungsbedingungen: Staub und Feuchtigkeit, Betriebstemperaturbereich sowie vorherrschende Lichtsituation. Die Umgebungshelligkeit im Betrieb entscheidet darüber, ob grundsätzlich eine selbstleuchtende Anzeige wie LED oder OLED oder bei LC-Displays eine Hinterleuchtung (Backlight) erforderlich ist. Soll die Anzeige bei Dunkelheit oder schwachem Umgebungslicht ablesbar sein, schließt das rein reflektive LCD-Anzeigen aus. Das schränkt die Auswahl auf transflektive oder transmissive LC-Displays, wie STN, FSTN oder TFT mit Hinterleuchtung, sowie auf selbstleuchtende Anzeigetechniken wie LED/OLED ein. Fällt die Auswahl auf LCD-Anzeigen mit Hinterleuchtung, kann der Entwickler zwischen verschiedenen Hinterleuchtungs-Typen wie LED oder CCFL wählen. LED-Hinterleuchtungen punkten mit geringem Stromverbrauch und langer Lebensdauer.

Bei einer hohen Umgebungshelligkeit wie dem Einsatz im Freien sind eine ausreichende Helligkeit und ein hohes Kontrastverhältnis bei der Anzeige entscheidend. Hier eignen sich Anzeigen mit Helligkeiten von über



Bild 2:

Bei passiven LCD-Anzeigen und negativer Darstellung erscheinen die Zeichen bzw. Pixel in der Farbe des Backlights.

1000 cd/m² und einem Kontrastverhältnis von über 1000:1.

Eine kontrastreiche Darstellung bieten E-Paper- und OLED-Displays. Für den Einsatz im Gebäude genügen Standardanzeigen mit 300 bis 600 cd/m² und 500:1 aus.

Weitere Umweltbedingungen bestimmen unter anderem die Schutzklasse. Für den Einsatz im Automobilbau sollte der erweiterte Temperaturbereich von -30 bis 85 °C oder sogar weiter liegen. Standard-Displays sind von -20 bis 70 °C spezifiziert und eignen sich in Mitteleuropa für den Außeneinsatz.

Da ein Display bei ungünstigen Blickwinkeln dunkler und kontrastloser erscheint, sind zwei optische Charakteristiken zu beachten. Meist befindet sich der Betrachter mit den Augen ober- oder unterhalb der Display-Mitte. Schaut der Anwender überwiegend von unten, dann wählt man eine Anzeige, die vom Betrachten einer Uhr entlehnt ist: „Viewing Direction 6 o'clock“. Im umgekehrten Fall „Viewing Direction 12 o'clock“. Displays mit Full Viewing besitzen keine Vorzugsrichtung bezüglich des vertikalen Betrachtungswinkels. Den maximalen Betrachtungswinkel in horizontaler und vertikaler Richtung, in dem das Display noch akzeptabel ablesbar ist, spezifiziert Viewing Angle. Soll die Anzeige selbst aus extremen Winkeln gut lesbar sein, ist ein

Display mit dem größten Betrachtungswinkel in horizontaler Richtung die richtige Wahl. Gute TFT-Displays erreichen horizontale Betrachtungswinkel von bis zu ±85 Grad. Für die meisten Applikationen reicht ein Betrachtungswinkel von ±60° oder ±65°.

Die Leistungsaufnahme und Bildwechsel eins Displays

Die Display-Techniken unterscheiden sich in Bezug auf die Leistungsaufnahme teilweise erheblich. Vor allem bei portablen bzw. batteriebetriebenen Geräten entscheidet die Leistungsaufnahme. Sie steigt mit der Größe und Ausstattung. Bei monochromen Anzeigen lässt sich die Leistungsaufnahme reduzieren:

■ Soll im Dunkeln abgelesen werden, ist eine selbstleuchtende OLED-Anzeige einer LCD-Anzeige mit Backlight vorzuziehen. OLED-Anzeigen besitzen eine geringere Leistungsaufnahme und einen besseren Kontrast, sind aber kostenintensiver.

■ Für Displays in einer hellen Umgebung eignen sich reflektive LCD-Anzeigen oder E-Paper. Sie bieten ausreichend Kontrast und benötigen nur bei einem Bildwechsel Energie.

Bei der Response Time (Bildwechsel) sollten es für eine flüssige Videowiedergabe weniger als 15 ms sein. Aktuelle LCD-TFT-Displays auf dem Markt sind mit Reaktionszeiten von 5 bis 30 ms spezifiziert. // HEH

Schukat electronic

The Customizing Class

sales@dmotechnics.com
www.dmbtechnics.com

Industrielle

Displaylösungen

individuell design

Von der ersten Idee bis zur fertigen Lösung – exakt auf Ihre Wünsche abgestimmt.

✚ Ihr Experte für kundenspezifische Displays

TFT-DISPLAYS

Ablesbar bei weitem Blickwinkel

Bei einem herkömmlichen TFT-Display nimmt der Kontrast ab, je schräger man auf die Anzeige blickt. Ab einem bestimmten Winkel kippt der Kontrast vollständig (Gray Scale Inversion Effect). Electronic Assembly hat mit der AACs-Technik (All Angle Color Stability) ein IPS-Grafikdisplay entwickelt, bei dem Kontrast und Farben selbst bei einem extremen Blickwinkel nahezu unverändert bleiben. Es bietet eine Helligkeit von bis zu 1000 cd/m² und ist von 2 bis 7 Zoll (ca. 5 bis 18 cm) verfügbar. Seine Auflösung liegt laut Hersteller bei 240 x 320 Bildpunkten und die Dicke ist mit 2,2 mm angegeben. Alle Displays sind als kapazitiver multigestenfähiger Touchoberfläche erhältlich.

Die Panels sind dank eines Flexkabelanschlusses und passendem ZIFF-Stecker als Zubehör zu verbauen. Sie lassen sich außerdem über die klassische RGB-

Bild: Electronic Assembly



Schnittstelle mit 16 oder 24 Bit ansteuern, die kleineren Displays mit 2, 2,8 und 3,5 Zoll zusätzlich auch über SPI. Die Versorgungsspannung beträgt 3,3 V. Für den Dauereinsatz in der Industrie konzipiert, arbeiten die TFT-Displays zwischen -20 bis 70 °C. Der Hersteller garantiert dabei eine Lebensdauer von mindestens 50000 Betriebsstunden.

Electronic Assembly

INDUSTRIE-DISPLAYS

KOE von 7 bis 12,3 Zoll im Angebot

Im Programm der Distec sind TFT-Displays für den Einsatz in der Industrie des taiwanischen Herstellers Kaohsiung Opto-Electronics (KOE), einem Unternehmen der Japan Display Inc. (JDI Group). KOE hat das Design, die Entwicklung und die Produktion der kleinen und mittelgroßen TFT-Displays 2012 vom renommierten Display-Pionier Hitachi übernommen. Distec

startet mit Displays in den Größen von 7 bis 12,3 Zoll (17,78 bis 31,242 cm). Die robusten Displays eignen sich besonders für den Einsatz in Umgebungen wie beispielsweise Automation, Industrie 4.0, Digital Signage, öffentlicher Verkehr, Landwirtschaft und Bauwesen.

Die TFT-Displays aus der Serie Rugged+ sind widerstandsfähig gegen Vibration und Schock. Den weiten Arbeitstemperaturbereich gibt der Hersteller von -40 bis 85 °C. Damit eignen sich die Displays für einen Betrieb bei Kälte oder Hitze. Außerdem bieten die Displays eine hohe Helligkeit von bis zu 1000 cd/m² mit integrierter Dimmkontrolle und einen weiten Winkel dank IPS-Technik. Damit lässt sich das Display auch bei ungünstigen Winkel ablesen.

Bild: Distec



Distec

...since 1984

LCD LED

Not only a project, it's a Partnership!

TOUCH LCD KEYPADS

TFT OLED

KEYPADS

COLOUR UP

YOUR LIFE

www.display-elektronik.de

Display Elektronik GmbH · Am Rauner Graben 15 · D-63667 Nidda
Tel. 0 60 43 - 9 88 88 - 0 · Fax 0 60 43 - 9 88 88 - 11

NEWSLETTER: www.display-elektronik.de/newsletter.html

SCHURTER
ELECTRONIC COMPONENTS

Elektronische Komplettlösungen

- Spezialist für alle HMI Technologien
- nach Medizinnorm ISO 13485 zertifiziert
- kompetenter Partner für Komplettsysteme
- Anbieter von kundenspezifischen Lösungen
- weltweite Präsenz mit Kompetenzzentren
- Service über den gesamten Produktlebenszyklus

info.de@schurter.com | +49 7642 6820
schurter.com

SCHURTER
ELECTRONIC COMPONENTS

REFLEKTIVES DISPLAY

Display mit 31,5 Zoll für Anwendungen im Freien



Bild: Data Modul

Ein reflektives Display im Format 31,5 Zoll und Full-HD-Auflösung von Sharp bietet Data Modul. Eine Besonderheit ist der vom Hersteller angegebene geringe Stromverbrauch. Beim Abspielen von Videos werden nur rund 2 % des Energiebedarfs zu vergleichbaren transmissiven Displays benötigt. Bei diesem Display-Typ wird kein Reflektor verwendet. Sie sind nur mit Beleuchtung ablesbar, dafür sehr hell. Somit verbessert sich bei zunehmendem Umgebungslicht die Ablesbarkeit des Displays.

Der Kontrast bleibt umgebungsunabhängig konstant.

Ein zusätzliches Leuchtmittel mit 4 cd reicht aus, um Informationen auch nachts vom Display abzulesen. Das Display lässt sich mit Batterien, Akkumulatoren oder Solarpaneels betreiben. Damit ist eine Installation selbst an Orten ohne feste Energieanbindung möglich. Dank seiner Einbautiefe und der geringen Abwärme eignet sich das Display für Anwendungen im Freien.

Data Modul

EMPFÄNGER MIT 60 GHZ

4K-Signale unkomprimiert zum Display übertragen

Spezielle Displays in den Größen 13,3 und 15,6 Zoll mit einer 4K-Auflösung und einem integrierten 60-GHz-Empfänger sind bei Holitech im Angebot. Über den Wireless-Standard lassen sich drahtlos unkomprimierte 2D-Signale in 4K-Auflösung mit einer Farbtiefe von bis zu 48 Bit und Bildwiederholraten von bis zu 240 Hz übertragen.

Somit haben Anwender die Möglichkeit, Bild- und Toninformationen ohne Komprimierung im Raum auf eine Distanz von bis zu 10 Metern zu übertragen. Die

Quelle kann beispielsweise ein Smartphone oder ein Tablet sein. Über den angeschlossenen Sender erfolgt die Datenübertragung latenzfrei. Dabei ist beim Sender keine weitere Spannungsversorgung notwendig, er wird über die angeschlossene Quelle versorgt und misst 15 mm x 12 mm.

Mögliche Anwendungen sind Videokonferenzen über das Smartphone, wo das Bild drahtlos auf das Display übertragen wird.

Holitech Europe



Bild: Holitech

MOBILE ANWENDUNGEN

Display für Low-Power-Anwendungen auswählen

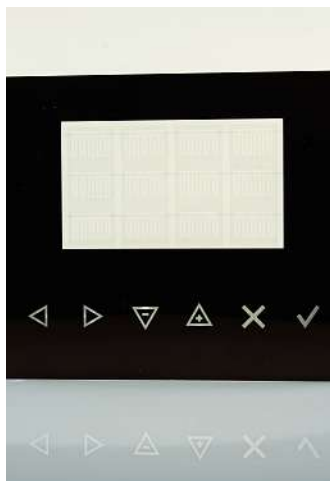


Bild: DMB

Für portable Anwendungen müssen Entwickler bei der Systemdefinition zuerst auf die notwendige Energieversorgung achten. Denn diese bestimmt, wie lange das System ohne einen Batteriewechsel oder Aufladen läuft. Bei Applikationen wie einem E-Bike-Display sollten es mehrere Stunden sein.

Ist die Batterielebensdauer festgelegt, wird die Display- und die eingesetzte Touch-Technik festgesetzt. Speziell beim kapazitiven Touchsensor unterscheidet sich die Stromaufnahme je

nach Betriebsmodus. Bei einigen Touchcontrollern lässt sich die Leistungsaufnahme über die Software steuern. Außerdem hat der Betriebsmodus-Wechsel Einfluss auf die Leistungsaufnahme. Es müssen also sowohl Touch-Technik als auch Betriebsmodus optimiert sein. Neben den Formfaktor ist es der Kostenfaktor, der die Komponentenwahl beeinflusst. Für eine lange Batterielebensdauer ist die spätere Funktion wichtig.

DMB

Wir sind die Taktgeber.

31 Hersteller

Frequenzbestimmender Bauelemente. Wir sind der Spezialdistributor für Resonatoren, Schwingquarze, Oszillatoren (XO, VCXO, TCXO, OCXO) und Real-Time-Clock-Module.

30.000+ Produkte

als auch kundenspezifische Sonderlösungen.

14 Spezialisten

an Ihrer Seite – wir beraten Sie herstellerunabhängig, technologieoffen und technisch kompetent.

Wir helfen Ihnen bei der Auswahl des für Sie richtigen Produktes und begleiten Ihr Projekt vom Erstmuster bis zur Serienfertigung. Rufen Sie uns an.

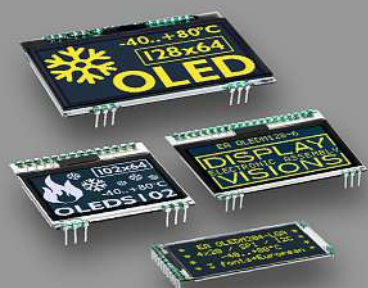
The Design-In Company.

☎ +49 4103 1800-0 ✉ info@wdi.ag 🌐 www.wdi.ag
WDI AG · Industriestrasse 21 · 22880 Wedel (Holstein)

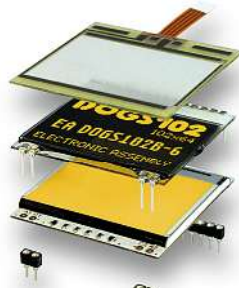
wdi ag



TFT Displays von 2" bis 10,1"
mit und ohne Touch oder in IPS
Technologie und bis zu 1000cd/m²



OLED für die Industrie. Pins.
-40...+80°C, SPI Interface.



COG Text- und Grafikdisplays
1x8 Zeichen bis 240x128 Pixel.



**MESSE
ANGEBOT**
69,- €*

2.8" HMI DEMOPACK

Überprüfen Sie die Luftqualität  und zeichnen Sie mit dem Farbdisplay Temperatur und Feuchte auf. Passen Sie Ihr Touchpanel mit unserem WYSIWYG-Editor individuell an.

**EA DEMOPACK-CLIMA:
netto zzgl. MWSt. Versandkostenfrei in Deutschland.
Dieses Angebot ist bis zum 15.03.2021 gültig.*

WELT DER DISPLAYS

ELECTRONIC ASSEMBLY bietet eine breite Palette hochwertiger Displays für den industriellen Einsatz. Wir haben auch für Sie das passende Display als HMI mit Touch, OLED, LCD oder TFT.



Intelligente HMI-Displays mit Touch und USB. Auch als Stand-alone. Mit WYSIWYG-Editor. Industrietauglich.

BESUCHEN SIE UNS

 **embeddedworld2021**
Exhibition & Conference
... it's a smarter world

DIGITAL



Registrieren Sie sich über den QR-Code und nutzen Sie den Gutscheincode.

ew21456801

ELEKTRONIK PRAXIS

www.elektronikpraxis.de

Wissen.
Impulse.
Kontakte.



Embedded Systems Development und IoT

Servermodule für Rechenzentren

Der Bedarf an Server-
technologie und Rechen-
zentren am Edge steigt
enorm.

Seite 28

FPGA-basierter Drohnen-Autopilot

Schneller zum Ziel mit
FPGAs am Beispiel einer
automatischen Steuerung
für Drohnen.

Seite 32

Verifizierung vs. Validierung

Für sicherheitskritische
Software gibt es strenge
Industriestandards. Wo sind
die Unterschiede?

Seite 40

Funktionsweise von LoRa Edge

WLAN- und Satelliten-
signale einfach auswerten
mit softwaresteuerbaren
Funk-Frontends.

Seite 50

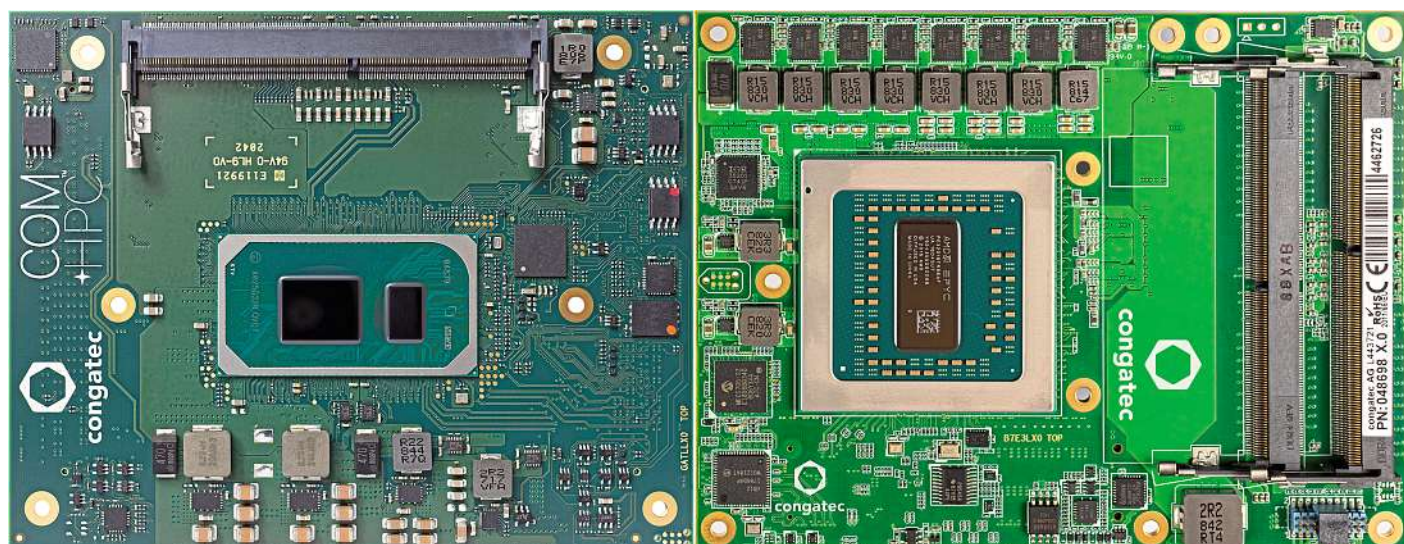


Bild: congatec

Die Congatec-Module COM-Express-Type 7 und COM-HPC: sie sind prädestiniert für die Entwicklung modularer Edge-Server-Technologie.

Embedded-Servermodule für Edge-Rechenzentren

Die Forderungen nach immer geringerer Latenz und Reduzierung des energiehungrigen Datentraffics über lange Strecken lässt den Bedarf nach Servertechnologie und Rechenzentren am Edge enorm steigen.

ANDREAS BERGBAUER*

Der Bedarf nach Servertechnologie und Rechenzentren am Edge steigt enorm. Da sich diese Anwendungen oft in rauen Umgebungen finden und die Infrastruktur für einen langen Zeitraum ausgelegt ist, sind robuste, langzeitverfügbare Servertechnologien gefragt. Server-on-Modules nach den COM-HPC- und COM-Express-Standards bieten sich als Designgrundlage an.

Einer der größten Teilmärkte ist die Telekommunikation. Allen voran die 5G-Technologien, die aufgrund der hohen Bandbreiten mehr Performance vor Ort vorhalten müssen. Hier werden auch mehr Basisstationen benötigt, da die Anzahl der verwaltbaren Devices bei voller Bandbreitenausnutzung sinkt und auch die Funkreichweiten dieser Basisstationen gegenüber bisherigen Mobilfunkstandards geringer ist. Der IT&Telekom-

Markt für Edge-Datenzentren wird einzig vom Markt der Co-Location-Applikationen überflügelt, in dem es um „ins Netz“ ausgelagerte Rechenzentren geht, die sich Unternehmen teilen. Beide Segmente halten mit 54 % mehr als die Hälfte des Gesamtmarkts.

Edge-Rechenzentren halten in vielen Bereichen Einzug

Der Bedarf an Edge-Rechenzentren wird im Bank- und Finanzsektor steigen, da digitale Transaktionen dank schnellerer Abrechnungsterminals und dem Bedarf nach kontaktlosen Technologien zunehmen werden. Auch der öffentliche Sektor mit seinen steigenden und immer komplexeren Sicherheitsanforderungen sowie der Trend zur Smart City begünstigen die Edge-Servertechnologie. Ebenso der Energiesektor und das Gesundheitswesen mit zunehmend bandbreitenstarken, möglichst latenzfreien Applikationen sowie die für Deutschland so wichtigen Märkte des autonomen Fahrens und der industriellen Fertigung. Am Edge herrschen jedoch andere Bedingungen als in

bisher wohlklimatisierten Serverräumen. Insbesondere kleinere Edge-Datencenter – teils auch Mini-, Micro- oder gar Pico-Rechenzentren genannt – werden vielfach im freien Feld in robust schützenden Outdoor-Schränken oder Containern installiert. Diese sind Wind und Wetter ausgesetzt und müssen umfassend vor Hitze, Kälte und Klimastürzen sowie Feuchtigkeit, Staub und – sofern Offshore eingesetzt – mitunter auch vor Salzwasser geschützt werden.

Die Ashra – der amerikanischstämmige Berufsverband für Heizungs-, Kühlungs-, Lüftungs- und Klimaanlagebau, nach dessen Standards und Richtlinien sich viele Rechenzentren beim Serverraummanagement richten – hat deshalb jüngst in einem Bulletin des Technical Committees (TC) 9.9 Vorschläge gemacht, wie Ausfälle und Probleme verhindert werden können. Neben der Maßgabe, dass man Umgebungsbedingungen mit den Spezifikationen der Komponenten abgleichen sollte, wird hier im Wesentlichen aber nur darauf hingewiesen, dass man für Service und Wartungszwecke nicht mal



* Andreas Bergbauer
... ist Senior Product Line Manager bei der congatec AG.

so eben die Türen öffnen sollte, denn so könne sehr heiße, kalte, feuchte oder staubige Luft zu schnell eindringen. Infolgedessen sollte der maximal erlaubte Temperaturwechsel bei höchstens 20 K innerhalb einer Stunde und maximal 5 K in 15 Minuten liegen. Der Temperaturwechsel solle zudem auch nur kontinuierlich und nicht plötzlich erfolgen. Es scheint, als wolle man hier die schonenden Bedingungen für Serverräume auch in die Wüste oder auf die Offshore-Bohrinseln und Windparks übertragen.

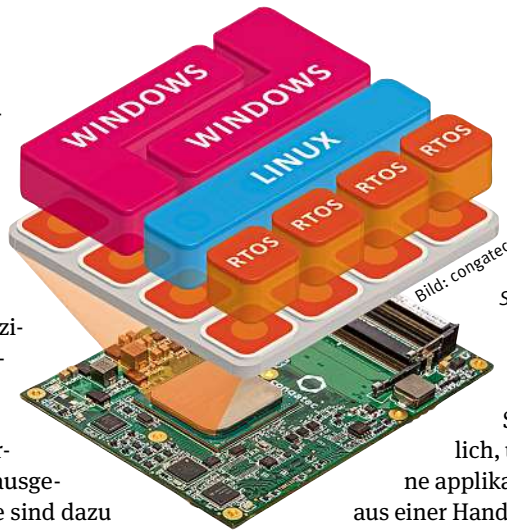
Ist es aber nicht sinnvoller, Edge-Server so zu entwickeln, dass man sie nicht wie ein rohes Ei behandeln muss und dass man infolge auch für die Klimatisierung nicht mehr ganz so viel Geld und Energie aufwenden muss, weil die Hardware in einem deutlich erweiterten Temperaturbereich betrieben werden kann und auch Temperaturschocks übersteht? Dies insbesondere dort, wo nicht mehrere Racks, sondern nur einige wenige Server redundant als Fog-Server-Farm arbeiten, um mit geringstmöglicher Latenz Echtzeitprozesse noch zuverlässig steuern zu können? Diesen Bedarf haben Prozessorsesteller wie AMD und Intel erkannt. Deshalb kommen zunehmend mehr Entry-Class-Serverprozessoren auf den Markt, die im Grunde ähnlich spezifiziert sind wie Embedded-Prozessoren. Sie sind für den Einsatz in rauen Umgebungen als aufgelötete Prozessoren mit BGA-Sockel ausgelegt und können weitere Temperaturbereiche als die für Standard-IT abdecken. Sie erhalten zudem auch einen Support über sieben Jahre und mehr, was es OEM in diesem Segment erleichtert, eigene Produkte auf den Markt zu bringen und langfristig pflegen zu können.

Die Embedded Computing Community hat sich darüber hinaus ebenfalls dem Design robuster Fog- und Edge-Server gestellt und

mit der COM-Express-Type-7-Spezifikation sowie mit dem neuen Standard COM-HPC Server zwei Modul-Spezifikationen geschaffen, die exakt auf diese neue Prozessorgenerationen ausgelegt wurden. Sie sind dazu prädestiniert, insbesondere die vielen kleinen verteilten Edge- und Fog-Server-Installationen zu entwickeln, die man beispielsweise im Automobilsektor für die vielen kleinen Serverracks am Wegesrand der Autobahnen benötigt. Durch den Einsatz von Real-Time-Hypervisor-Technologien ermöglichen sie dabei ein perfektes Performance-Balancing. Zudem reduzieren sie erheblich die Gesamtbetriebskosten, da mit ihnen die Leistungsskalierung der nächsten Fog-Generation durch einen einfachen Wechsel des Prozessormoduls möglich wird, anstatt das gesamte Rackmount-System austauschen zu müssen.

Nutzer wollen individuelle Rechenkapazitäten

Neben einem robusten Design liegen die Herausforderungen bei Echtzeit-Edge-Computing-Anwendungen auch darin, das beste Setup sowohl für die Fog-Dienste als auch für die über zeitsensitive Netzwerke verbundenen Edge-Geräte zu finden. Da an den Edges viele Aufgaben zu bewältigen sind, benötigen OEM-Kunden und professionelle Endanwender einen individuellen Mix der Rechenkapazitäten. Modularität ist also so-



Module von Congatec: sie unterstützen die Echtzeit-Hypervisor-Technologie, wie die von Echtzeitsystemen, zur Konsolidierung mehrerer Applikationen auf einem System.

wohl auf der Hardware- als auch auf der Softwareebene erforderlich, um perfekt zugeschnittene applikationsfertige Plattformen aus einer Hand liefern zu können.

Modularität auf der Hardwareebene ist die Kernkompetenz von Congatec. Die firmeneigene Hypervisor-Software für echtzeitfähige virtuelle Maschinen rundet das Plattformangebot für Fog-Server auf Softwareebene ab und schafft die Grundlage für weitere Entwicklungsschritte der OEMs zum Aufbau perfekt zugeschnittener Rugged Fogs. Das Unternehmen will sein Angebot deshalb in Zukunft durch die Unterstützung von Lösungspartnern für Vision, KI, VR, AR, Big Data Analytics und dedizierte Edge-Computing-Services erweitern. Zudem sind auch Konfigurationen mit virtuellen Maschinen in der Evaluierung, um IoT-Gateway- und Sicherheitsfunktionen zur Erkennung von Schwachstellen, Angriffen und Anomalien oder kryptographische Funktionen anbieten zu können. Diese könnten sogar die Standards FIPS 140-2 Level 3 oder BSI Common Criteria EAL5 von Hochsicherheitsanwendungen erreichen, was nicht unerheblich ist, denn mit zunehmender Digitalisierung sind die Daten in den Systemen auch das Öl unserer Zeit und damit viel Geld wert, was entsprechend hohen Schutz erfordert. // MK

Congatec

**HAMMOND
MANUFACTURING®**



**1557 Kunststoffgehäuse
aus Polycarbonat IP68 und ABS IP66**

Mehr erfahren: hammfg.com/1557

Kontaktieren Sie uns, um ein kostenloses Bewertungsmuster anzufordern.
eusales@hammfg.com • + 44 1256 812812



Industrielle IoT- und KI-Plattform für das Edge Computing

Die COM-Express-Type-6-Modulfamilie MSC C6C-TLU basiert auf Intels Core-Prozessoren der 11. Generation. Avnet Integrated informiert über die Einsatzgebiete seiner hochleistungsfähigen Module.

Die Modulfamilie MSC C6C-TLU eint hohe Computing Performance mit KI-Beschleunigung und integrierten Echtzeitfähigkeiten für den Einsatz in industriellen Netzwerken. Einzelne Modulvarianten meistern den erweiterten Temperaturbereich und raue Umgebungsbedingungen.

Einer der vielversprechenden Wachstumsmärkte der Zukunft sind innovative Edge-Computing-Lösungen, die eine Vielzahl an Technologien, Komponenten und Applikationen mit unterschiedlichen Anforderungen kombinieren. Jede Anwendung erfordert eine aufgabenspezifisch optimierte Hardware mit modularer Prozessorleistung und unterschiedlichen Schnittstellen. Die entscheidenden Bausteine sind eine leistungsstarke Computing Performance, eine hohe Datensicherheit, eine einfache Netzwerkfähigkeit und ein Betrieb im erweiterten Temperaturbereich. Durch den Einsatz von Künstlicher Intelligenz und Deep-Learning-Technologien können auf der Basis neuronaler Netzwerke und mit Hilfe großer Datenmengen die Edge-Computing-Anwendungen kontinuierlich optimiert und beschleunigt werden. Oft ist

es wichtig, dass die KI-Beschleunigung direkt vor Ort in der Edge, also nicht erst in der Cloud oder in High Performance Computing-Lösungen, zur Verfügung steht.

Ein breites Einsatzgebiet sind Computer-Vision-Systeme zur zuverlässigen Erkennung und Identifizierung von Objekten und Personen und anschließenden Bildverarbeitung und Analyse. Die leistungsfähigen Systeme kombinieren Edge- oder Cloud-basierende Computertechnologien, Software und KI mit smarten Kameramodulen. Die Anwendungen reichen von der Robotik, intelligenten HMI-Terminals über Zutrittskontrollsysteme bis hin zu medizintechnischen Monitoring-Systemen und POS/POI-Geräten. Für eine schnelle Realisierung anspruchsvoller Edge-Computing-Anwendungen sorgen eine neue Generation an Prozessor- und Modultechnologien, welche die benötigte Rechen-, Grafik- und Videoleistung bieten. Zudem spielen zahlreiche integrierte Funktionen wie eine hohe Netzwerkfähigkeit und schnelle Schnittstellen eine zentrale Rolle und bilden – zu einer einzigen Plattformlösung vereint – die Ausgangsbasis für unterschiedliche Anwendungen.

Speziell für Edge-Computing-Lösungen bietet Avnet Integrated die hochleistungsfähige COM-Express-Type-6-Modulfamilie MSC C6C-TLU, die auf der aktuellen 11. Generation von Intel-Core-Prozessoren (früherer Codename „Tiger Lake UP3“) basiert (Bild 1). Die High-End-Module zeichnen sich im Vergleich zu den Vorgängermodellen, die einen Intel-Core-Prozessor der 8. Generation integrieren, durch einen deutlichen Performance-Zuwachs aus. Mit den sofort einsetzbaren Embedded-Modulen steht die neue Prozessortechnologie frühzeitig für die Entwicklung vielfältiger Endsysteme in kurzer Zeit mit optimierter Time-to-Market zur Verfügung.

Vorteile der engen Partnerschaft mit Intel

„Dank der engen Partnerschaft mit Intel sind wir in der Lage, zeitgleich mit der Einführung der neuen Prozessorgeneration unsere leistungsfähige Modulfamilie mit Intels Core-Prozessoren der 11. Generation vorzustellen“, freut sich Jens Plachetka, Manager Board Platforms bei Avnet Integrated. „Als einer der führenden COM-Hersteller positionieren wir die Embedded-Baugruppen besonders für Anwendungen, wo eine dezentrale Rechenleistung und eine hohe Konnektivität gefragt sind. Mit der erneuten Erweiterung unseres Produktangebots können wir neue Märkte erschließen und unsere Kunden vom Design-In bis hin zur kompletten Systemlösung umfassend unterstützen.“ Neben den Modulen MSC C6C-TLU stellt Avnet Integrated das komplette Ecosystem mit passenden Design Tools wie Starter Kit und Board Support Package sowie umfangreiche Service-Leistungen wie Design-in-Support, Carrier Design Review und Temperatur-Screening zur Verfügung. Wie alle Modulfamilien werden die Boards in den Technology-Campussen von Avnet Integrated entwickelt und in hochautomatisierten Produktionsstätten im eigenen Hause gefertigt. Damit ist die Langzeitverfügbarkeit der Module von

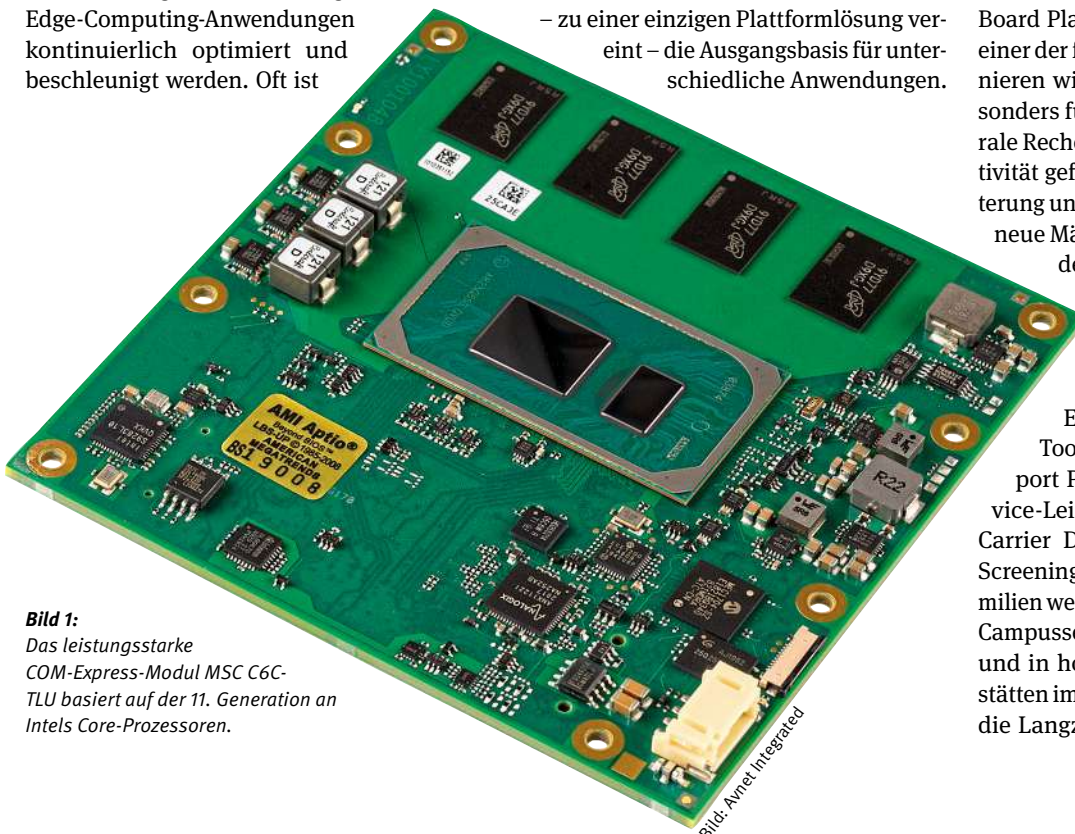


Bild 1:
Das leistungsstarke
COM-Express-Modul MSC C6C-
TLU basiert auf der 11. Generation an
Intels Core-Prozessoren.

Bild: Avnet Integrated

mindestens 15 Jahren ab Produkteinführung gesichert.

Die leistungsfähige COM-Express-Type-6-Modulfamilie MSC C6C-TLU wird mit Core-Prozessoren der 11. Generation von Intels Embedded Roadmap bestückt. Die energieeffizienten Prozessoren mit zwei bzw. vier Cores und acht Threads mit Taktraten bis 2,2 GHz (4,4 GHz max. Turbo) haben eine TDP (Thermal Design Power) zwischen 12 und 28 W. Dank neuer Prozessorarchitektur und der fortschrittlichen 10-nm-Prozesstechnologie bietet die Modulfamilie eine gute Balance zwischen Performance und Verlustleistung.

Intels Core-Prozessoren der 11. Generation nutzen die schnelle Grafikprozessoreinheit (GPU) Intel Iris Xe mit bis zu 96 Ausführungseinheiten, die bis zu vier unabhängige Displays mit bis zu zwei Kanälen für 8K- bzw. mit bis zu vier Kanälen für 4K-Video sowie zwei Video-Decoder mit bis zu 40 simultanen 1080p-Streams bei 30 Bildern pro Sekunde meistert. Dank einer Hardware-Beschleunigung und Virtualisierung können etliche IoT-spezifische Aufgaben gleichzeitig laufen.

Um eine hohe Speicherbandbreite zu erreichen, sind die Embedded-Module MSC C6C-TLU mit schnellem, 4 bis zu 32 GB großem Dual-Channel-LPDDR4X-4266-Speicher erhältlich. Speziell für kritische Anwendungen kann zur Verbesserung der Sicherheit und Zuverlässigkeit im Speicher optional eine In-band-ECC-Fehlerkorrektur (IB ECC) implementiert werden. Die IB ECC speichert ECC-Korrektur-Bits in internen Caches und externen reservierten Speicherblöcken. Die ECC-Prüfung ist für einzelne Speicherbereiche des externen DDR-Speichers separat aktivierbar. Der Vorteil ist, dass kein zusätzlicher externer ECC-Speicher erforderlich ist.

Die COM-Express-Type-6-Modulfamilie MSC C6C-TLU stellt bis zu neun PCIe Gen3 und Gen4 Lanes, USB-4- und USB-3.1-Schnittstellen, GPIOs und ein 1-Gb-Ethernet-Port zur Verfügung. Zwei SATA-6-GB/s-Kanäle stehen für Massenspeicher bereit. An Grafikschnittstellen sind drei DisplayPort/HDMI-Interfaces und ein LVDS- und Embedded-Display-Port-Interface vorhanden. Die Baugruppen sind mit einem Trusted Platform Module TPM 2.0 bestückt.

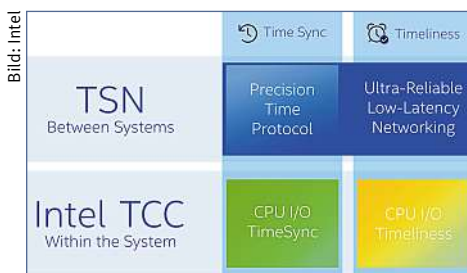


Bild 2: Die COM-Express-Modulfamilie MSC C6C-TLU unterstützt Echtzeitfunktionalitäten wie TSN und Intel TCC.

Speziell für moderne Netzwerkanwendungen bieten die COM-Express-Standardmodule besondere Echtzeitfähigkeiten wie TSN (Time-Sensitive Networking). Unterschiedliche Prozesse innerhalb eines Netzwerkes können so mit hoher zeitlicher Deterministik ablaufen und z.B. immer wieder neu gestartet werden. Das Ziel ist, Maschinen und komplexe Anlagen über alle Industriebusse hinweg sicher zu vernetzen und synchron zu steuern um ein Echtzeit-Computing zu ermöglichen. Die TSN-Spezifikation erweitert die Datenübertragung über Standard-Ethernet-Netzwerk um eine neue Funktionalität, die einen sicheren Transfer in Echtzeit mit geringer Latenzzeiten und hoher Verfügbarkeit ermöglicht. Um strenge Timing-Vorgaben und geringe Latenzzeiten auch unter harten Echtzeitbedingungen sicherzustellen, wird neben TSN auch das Software-Protokoll Intel Time Coordinated Computing (Intel TCC) zur Minimierung von Jitter unterstützt (Bild 2).

Was die High-End-Modulfamilie MSC C6C-TLU auszeichnet

Die High-End-Modulfamilie MSC C6C-TLU ist mit besonders hochwertigen Komponenten bestückt und dank ihres robusten Designs auch für Anwendungen mit hohen Anforderungen geeignet, z.B. beim Einsatz in rauen industriellen Umgebungen oder im Außenbereich. Avnet Integrated bietet hochleistungsstarke Modulvarianten im industriellen Umgebungstemperaturbereich von -40 bis 85 °C und für den zuverlässigen 24/7-Dauerbetrieb an. Voraussetzung dafür ist, dass der eingesetzte 11.-Gen.-Intel-Core-Prozes-

sortyp für die erweiterten Betriebsbedingungen spezifiziert ist. Dank des 10-nm-Fertigungsprozesses, der durch eine kleine Verlustleistung und geringen thermischen Stress der Bauteile gekennzeichnet ist, ist die Intel-Core-Technologie erstmals für den industriellen Einsatz im erweiterten Temperaturbereich erhältlich. Besonderes Augenmerk wird auf das Kühlungskonzept der Komponenten gelegt, das neben der CPU auch den on-board-Speicher berücksichtigt.

Die Produktfamilie MSC C6C-TLU ist unempfindlich gegen Schock und Vibrationen, da u.a. die Speicher auf das Board gelötet sind. Optional besteht die Möglichkeit, die Module durch eine Schutzbeschichtung (Conformal Coating) zuverlässig gegen Umwelteinflüsse wie Verschmutzung, Feuchtigkeit, Tau- oder Salzbeschlagung zu schützen.

Kombiniert mit Software Development Kits (SDKs) wie Intels Distribution-Toolkit OpenVINO ist die COM-Express-Type-6-Modulfamilie MSC C6C-TLU für unterschiedliche High-End-Anwendungen im industriellen Umfeld geeignet, z.B. für komplexe HMI-Lösungen, KI- und Industrial-IoT-basierende Systeme und echtzeitfähige Maschinensteuerungen. In der kollaborativen Robotik sind Anwendungen realisierbar, bei denen autonome Roboter und Personen in einem dynamischen Umfeld gleichzeitig arbeiten können. Auf der neuen Modulgeneration basierende intelligente Kamerasysteme erkennen auch komplizierte Muster und/oder Objekte und verarbeiten die aufgenommenen Daten dank hoher Prozessorleistung bereits vor Ort in der Edge. Die nächste Generation moderner Medizingeräte ist gekennzeichnet durch hochauflösende Displays und weitreichende KI-basierenden Diagnosemöglichkeiten. Ein weiteres Einsatzgebiet der COM-Express-Modulfamilie sind professionelle, vernetzbare Audio- und Videoanlagen, u.a. für komplexe Public-Signage- und Infotainment-Lösungen. Hohe Temperaturunterschiede und widrige Witterungsverhältnisse müssen öffentliche Überwachungs- und Sicherheitssysteme sowie Verkehrsmanagementsysteme verkraften können. // MK

Avnet Integrated

compmall
► Industrie-PCs



Industrie-PCs von compmall ...für Ihr Projekt passend gemacht!

Lieferkonzepte
bis zu 10 Jahre

www.compmall.de

FPGA-basierter Autopilot für Drohnen

Schneller zum Entwicklungsziel mit FPGAs: Das Beispiel einer automatischen Steuerung für Drohnen zeigt, wie sich mit FPGAs die Time-to-Market für innovative Produkte verkürzen lässt.

MARCO KÜMMERLING *

Elektronische Steuerungen sind heute die Gehirne vieler Produkte und ermöglichen eine enorme Vielfalt an Variationen und maßgeschneiderter Lösungen. Dies gelingt nur auf Basis eines gut strukturierten Designs der elektronischen Steuerung, was über alle Branchen hinweg eine nicht zu unterschätzende Herausforderung an die entwickelnden Unternehmen und Dienstleister ist.

Wer leistungsfähige und robuste Geräte in kurzer Zeit realisieren will greift häufig auf lokal programmierbare Logik zu – etwa FPGAs (Field Programmable Gate Arrays). Je nach Aufgabenstellung kommen FPGAs allein oder in Kombination mit frei program-

mierbaren Controllerlösungen zum Einsatz. Die Bausteine bieten bemerkenswerte Vorteile, wenn es um die Verarbeitung vieler Datenkanäle und die Ausführung von Algorithmen geht. Aufgrund ihrer internen Struktur können sie Signale im Gegensatz zur sequentiellen Verarbeitung in Controllern und Prozessoren getaktet parallel erfassen und verarbeiten, so dass eine hohe Bandbreite zur Verfügung steht und auch Signale im GHz-Bereich verarbeitet werden können. Zudem ist bei einem getakteten Design die Verarbeitung der Prozesse deterministisch, was die Bewertung des Designs in Sicherheitsanwendungen vereinfacht und die Entwicklungskosten solcher Anwendungen verringern kann.

Auch die geringe Fehleranfälligkeit von FPGAs spielt eine wesentliche Rolle. Die zahlreichen Punkt-zu-Punkt-Verbindungen stellen sicher, dass Variablen nicht von der falschen Quelle beschrieben werden. Die

Wahrscheinlichkeit von Fehlern bei der Interprozess-Kommunikation ist gering, die Taktfrequenz kann deutlich geringer ausgelegt werden als bei einer CPU, und ein unterlagertes Betriebssystem als potentielle Störungsquelle ist nicht vorhanden.

FPGAs verarbeiten Prozesse deterministisch

Mit ihren Eigenschaften sind FPGA optimal geeignet zum Erfassen und Verarbeiten großer Datenmengen mit hoher Geschwindigkeit. Besonders erfolgreich kommen sie zum Beispiel in folgenden Bereichen zum Einsatz.

- Videodatenverarbeitung
- Bilddatenverarbeitung
- Künstliche Intelligenz
- komplexe Prozessverarbeitung in der Automation
- im industriellen IoT-Bereich oder
- für Safety-Anwendungen

Entwickler profitieren von der hohen Integrationsdichte und Flexibilität der Bausteine; können somit Kosten senken und Entwicklungszeiten verkürzen. Möglich wird dies – neben der hohen Anzahl von Logikelementen – durch integrierte Memory- und DSP-Blöcke, die für Multiplikation, FIR- und FFT-Filter und ähnliche rechenintensive Algorithmen direkt verwendet werden können. Auch viele schnelle serielle Transceiver mit über 10 GBit/s für Gigabit Ethernet oder PCI Express sind in vielen FPGAs bereits integriert. Es ist zu erwarten, dass die Vielfalt der integrierten Blöcke und Verarbeitungsmethoden in Zukunft weiter steigen wird, womit ein noch größeres Anwendungsfeld für die programmierbaren Bausteine erschlossen werden kann.

Der Hersteller Microsemi bietet mit seinen Polarfire-FPGA noch eine Reihe weiterer Vorteile, wie eine niedrige statische Stromaufnahme und aufgrund ihrer Konstruktion eine hohe Robustheit gegenüber externen Störeinflüssen. Die Unternehmen Arrow Electronics als Auftraggeber und IMG Electronic &



* Marco Kümmerling
... ist Entwicklungsingenieur
bei IMG Electronic & Power Systems
in Nordhausen.

Bild 1:

Das Blockschaltbild des SupervisorPilot zeigt den redundanten Aufbau der FPGA-basierten automatischen Steuerung.

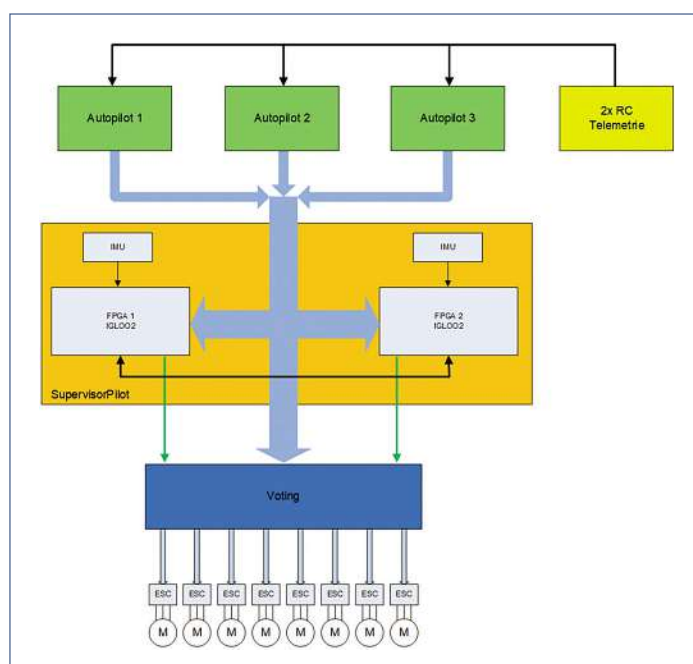


Bild: IMG Electronic & Power Systems

Power Systems als Dienstleister haben in Kooperation das FPGA Development Board „Everest“ entwickelt, das die Entwicklungszeit enorm verkürzen kann. Zentrales Element ist der Polarfire FPGA MPF300T mit 300K Logikelementen. Um ein breites Spektrum an Anwendungen abdecken zu können, beinhaltet das Board zwei DDR3-Speicherblöcke mit 1 GByte respektive 2 GByte Speicherkapazität, drei Gigabit-Ethernet-Schnittstellen, ein SFP+ Interface für z.B. 10 Gigabit Ethernet, eine 4-Lane PCIe-Schnittstelle sowie einen HDMI-Ausgang.

Zum Programmieren und Debuggen ist ein FlashPro 5 Programmer mit auf dem Board integriert. Für eigene Hardwareerweiterungen bietet das Board einen FMC-HPC-Connector und einen PMOD-Connector. Das Board wird mit 12 V DC versorgt, wobei die Spannung extern oder über das PCIe Interface eingespeist werden kann.

Dieses Development-Board entstand aus der Praxis für die Praxis, so dass bereits viele Designbeispiele vorhanden sind. Dazu zählen ein Videoverarbeitungssystem, ein Web-Server mit 1 Gigabit Ethernet und eine Anwendung mit KI (Künstliche Intelligenz) sowie Spezialanwendungen, für die keine Standardkomponenten verfügbar sind. Dazu zählen auch Anwendungen mit dem RISC-V-Open-Source-Prozessor. Auf der Webseite der IMG Electronic & Power Systems GmbH sind eine Reihe konkreter Designbeispiele für das Development Board verfügbar, die den Einstieg in die FPGA-Welt erleichtern und als Grundlage für eigene Projekte dienen können. Bei solchen Projekten unterstützt IMG als technischer Dienstleister die Industrie, ihre Innovationen schnell und kostengünstig auf den Markt zu bringen. Neben kompetentem FPGA-Design, Programmierung und Schulung bietet das Unternehmen umfassende technologische Dienstleistungen in der Hard- und Softwareentwicklung, für Embedded Systems, Batterie-Management-Systeme und in der Leistungselektronik. Mit angeschlossener Fertigung und hauseigenem Prüflabor kann von der Idee bis zum Serienprodukt die gesamte Wertschöpfungskette abgebildet werden.

Projekt: Ausfallsicherer Autopilot für Drohnen

Eine sicherheitsrelevante Anwendung von FPGAs sei hier am Beispiel des Projekts „SupervisorPilot“ dargestellt. Dabei handelt es sich um ein System, das den sicheren und zuverlässigen Betrieb von kleinen unbemannten Flugobjekten, sogenannten Drohnen, sicherstellen soll. Mikrodrohnen und kleine Drohnen (S-UAS = „small unmanned

aerial system“) verfügen in der Regel über keinen serienmäßigen und redundanten System-Ausfallschutz, der mit den Ansprüchen der bemannten Luftfahrt vergleichbar wäre. Dies begründet die Nachfrage nach einer technischen Lösung. Bei kritischen, industriellen Anwendungen oder beim Überflug von Menschen ist ein Absturz unter allen Umständen zu vermeiden. Zur Kollisionsverhinderung ist das Erfassen von Hindernissen essenziell; zudem ist das Wissen um die eigene Position auch für Notfall-Routinen wie das „Coming Home“ erforderlich.

Bis zu drei redundante Autopiloten parallel

Die Aufgabenstellung zu diesem Projekt umfasste im Wesentlichen das Erhöhen der Ausfallsicherheit und Verfügbarkeit von unbemannten Luftfahrzeugen durch den Einsatz von bis zu drei redundant arbeitenden Standard-Autopiloten. Hierzu erfolgt die Bewertung (Voting) der Motorsteuersignale von den Autopiloten und die Weiterleitung der Signale eines verfügbaren Autopiloten zu den ESC (Electronic Speed Control) der Motoren durch zwei FPGAs. Zusätzlich erhält der „SupervisorPilot“ über ein Interface-Modul weitere Statusdaten von den Autopiloten (z.B. Bestätigung Position = OK, Kontakt Fernsteuerung = OK). Die Auswertung dieser Daten beeinflusst das Voting, so dass sichergestellt wird, dass immer ein Autopilot mit vollständig positivem Status aktiv ist. Nur dessen Daten und Signale werden dann auch vom Voter an die Telemetrie und aktive Fernsteuerungen übermittelt.

Parallel bewertet ein eigenes Inertialsystem (IMU) an beiden FPGAs die aktuelle Lage des unbemannten Luftfahrzeugs. Es erfolgt eine permanente Kommunikation der FPGAs untereinander zum Vergleich der Bewertungsergebnisse inklusive Ausfallerkennung eines FPGAs. Integriert ist zudem die Option, eine Fernsteuerung auszuwählen – beispielsweise für Schüler- / Lehrer-Betrieb, oder zur Übernahme eines unbemannten Luftfahrzeugs in kritischen Flugzonen z.B. durch Sicherheitsbehörden. Schließlich bietet das System die Möglichkeit – je nach kundenspezifischer Konfiguration der Drohne – zur Einleitung geeigneter Maßnahmen bei Ausfall einer oder mehrerer Komponenten bis hin zur manuellen oder automatischen Auslösung des Notfallsystems (Fallschirm und/oder Airbags). Das Projekt befindet sich derzeit in der Baumusterphase, die Erprobung fliegender Prototypen ist für diese Jahr (2021) geplant.

// ME

IMG Electronic & Power Systems

Embedded Hardware

- universelle und effiziente Kühlrippengehäuse zur Entwärmung von Embedded Mainboards
- optimal angepasste Kühlkörperlösungen durch präzise Fräsbearbeitungen
- effektive Wärmespreizung mittels im Kühlelement verpresster Kupferflächen
- kundenspezifische Anfertigungen



Mehr erfahren Sie hier:
www.fischerelektronik.de

Fischer Elektronik GmbH & Co. KG

Nottebohmstraße 28
 58511 Lüdenscheid
 DEUTSCHLAND
 Telefon +49 2351 435-0
 Telefax +49 2351 45754
 E-Mail info@fischerelektronik.de

Flexible „System on Chip“-Lösung für Endgeräte

Endgeräte für das IoT sollen intelligenter werden. Eine sinnvolle Kombination von FPGA, System-on-Chip und Mikrocontroller kann die erforderliche Rechenleistung energieeffizient bereitstellen.

HARALD WERNER *

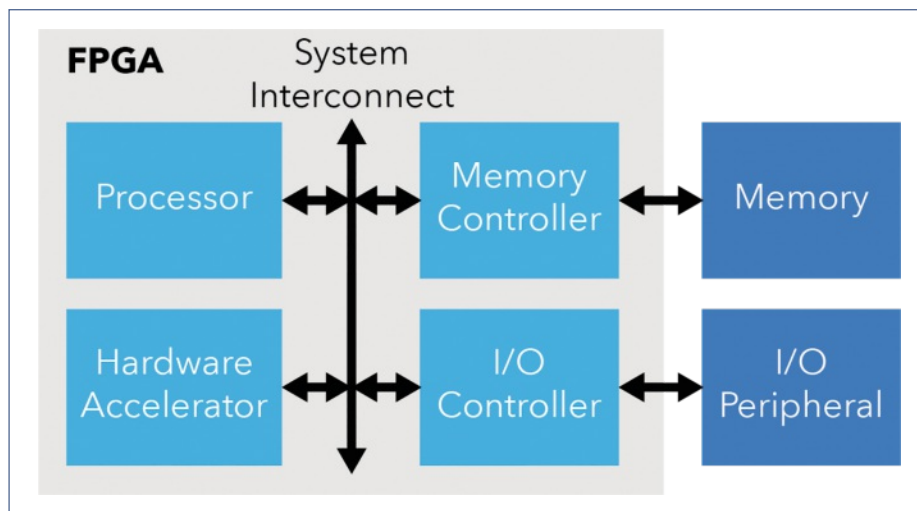


Bild: Efinix

system übernehmen und Standard-Algorithmen verarbeiten, die wenig Rechenleistung benötigen. Dort wo seine Rechengeschwindigkeit nicht ausreicht, lassen sich Prozesse in das FPGA auslagern, das sie über massiv-parallele Verarbeitung beschleunigt. Über den Memory Controller kann der Prozessor auf die zu verarbeitenden Daten zugreifen oder diese auch speichern. Über die Peripherie-Komponenten werden die Daten zur Verarbeitung in das System gebracht. Die Herausforderung, ein solches System zu designen, liegt in der Komplexität sowie den benötigten Ressourcen und den damit verbundenen Kosten. Der Vorteil der optimierten Lösungen von Efinix besteht darin, dass sie applikationsspezifische I/O-Peripherien sowie Schnittstellen als fertige Lösungen zur direkten Implementierung anbietet. Auch das Interface zwischen Prozessor und externem DDR-Speicher wird als Plug & Play-Komponente zur Verfügung gestellt. Als Prozessor wird ein Open-Source-RISC-V-Prozessor benutzt, der ein Hardware-/Software-Co-Design ermöglicht. Dieser Ansatz erlaubt eine schnelle Implementierung des Designs und damit einen schnelleren Verfügbarkeit eines Produkts auf dem Markt.

Als erste vorgefertigte Lösung dient hier ein Kamera-Design mit einem Kamerasensor, der über die MIPI-CSI2-Schnittstelle an das Trion FPGA angeschlossen ist. Über eine I2C-Schnittstelle, die über einen APB3-Bus an den RISC-V-Prozessor angeschlossen ist, kann der MIPI-CSI2-Sensor über den RISC-V-Prozessor konfiguriert werden. Die Daten aus dem MIPI-CSI2-Interface können über eine Vorverarbeitung im Eingangsmodul verarbeitet werden. Die verarbeiteten Daten werden über ein standardisiertes FIFO-Interface in den DMA geschrieben. Aus dem DMA kann der Speicher angesprochen werden der über den System Interconnect Bus ebenfalls mit dem Prozessor (RISC-V) verbunden ist. Der Hardware-Beschleuniger ist ebenfalls mit dem DMA verbunden. Im Beschleuniger ist

Bild 1: Blockschaltbild des System-on-Chip-Bausteins. Neben der programmierbaren Logik sind ein Prozessor, Hardware-Beschleuniger, Speicher- und I/O-Controller integriert.

Endgeräte sollen immer intelligenter werden, unter anderem um den Datenaustausch zwischen ihnen und einer Cloud zu optimieren. Alle Funktionen, die ein Endgerät selbst erledigen kann, müssen nicht in die Cloud ausgelagert werden. Die Produkte werden damit auch zuverlässiger, da keine Abhängigkeit von einer Cloud-Verbindung besteht. Andererseits setzt dies voraus, dass die Geräte über mehr Rechenleistung verfügen. Gleichzeitig dürfen diese Geräte aber nicht mehr Strom verbrauchen – schließlich sollen sie ihre Batterie nicht stärker belasten und möglichst wenig Wärme erzeugen. Es gibt hier die verschiedensten Ansätze mit FPGAs, festverdrahteten SoCs oder auch Mikroprozessoren, die aber meist an Ihre Grenzen stoßen. Während eine Lösung einfach zu realisieren ist, aber nicht

genug Rechenleistung hat, verbraucht die andere Lösung zu viel Strom oder ist wirtschaftlich nicht sinnvoll. Doch die Vorteile der Einzellösungen lassen sich kombinieren.

SoC mit applikations-spezifischen Funktionsblöcken

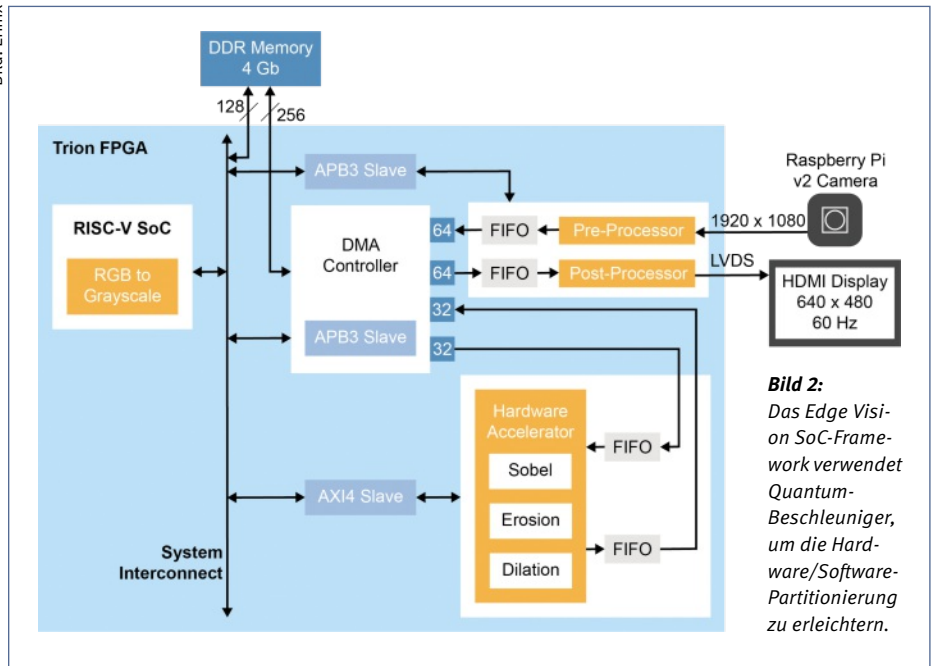
System-on-Chip (SoC) verbinden viele Anwender mit Bausteinen, die Prozessoren, programmierbare Logik und I/O-Blöcke auf einem Siliziumchip vereinen. Dieser Ansatz ist sicherlich eine gute Lösung, wenn man sehr schnellen Datenaustausch zwischen FPGA und Prozessor benötigt, mit sehr geringer Latenzzeit bei gleichzeitiger hoher Geschwindigkeit des Prozessors. Im Prinzip besteht ein System on Chip aus einem Prozessor mit einem oder mehreren Cores, einem Memory Controller, Peripherie-Komponenten sowie einem Hardware-Beschleuniger. Um diese Komponenten flexibel zu halten, hat Efinix sein Edge Vision SoC entwickelt. Es ist möglich, diese Lösung auf die Zielapplikation zu optimieren. Der Mikroprozessor kann die Kontrolle über das Gesamt-



* Harald Werner
... ist European Sales Director von Efinix mit Sitz in Allershausen

hier z.B. ein Sobel-Filter mit Erosion und Dilation implementiert, der in der Lage ist, die Kanten hervorzuheben, was zu einer besseren Kantenerkennung führt. Die Daten kommen aus dem DMA über ein FIFO und werden vom Hardware Beschleuniger auch wieder über ein FIFO in den DMA geschrieben. Der Prozessor steuert den Ablauf und hat entsprechenden Zugriff auf die Daten im DMA bzw. Speicher. Die abgearbeiteten Daten werden dann über das Ausgabemodul an ein Display mit LVDS Schnittstelle angeschlossen. Hier können auch die entsprechenden Verarbeitungen und die Anpassung an das Display erfolgen. Das Eingangs-Modul, Ausgabe-Modul sowie das Hardware-Beschleuniger-Modul sind in einem standardisierten Wrapper eingebettet, so dass sich der Entwickler keine Gedanken um die Anbindung an DMA, oder die verschiedenen Module machen muss.

Bild: Efinix



RISC-V Prozessor in Ruby-Konfiguration

Weitere Funktionen wie RGB 2 Gray können im RISC-V Prozessor erfolgen. Das Design wurde aus den vorgefertigten Modulen aufgebaut, die über Highspeed-Schnittstellen zum DMA sowie über Low-speed-Verbindungen zum Prozessor verbunden sind. Nur dort, wo hohe Datentransferraten erforderlich sind, kommen schnelle Verbindungen zum Einsatz. Für Kontrollsignale werden nur langsamere, ressourcensparende Verbindungen benötigt und verwendet. Über diesen Ansatz können schnell die verschiedensten Designs verbindungsoptimiert erstellt werden. Nicht alle Komponenten greifen direkt auf den externen Speicher zu. Somit kann dort auch nicht so schnell ein Engpass entstehen. Zum Realisieren solcher Systeme bietet Efinix die Trion-FPGA-Familie sowie verschiedene Funktionsblöcke an. Diese Bausteinfamilie besteht aus FPGAs mit bis zu 120K Logikelementen und einem eingebauten Speicher-Controller, der DDR3, LPDDR3, LPDDR2 bis 1066 MBit/s unterstützt. Auch die MIPI-CSI2-Schnittstellen sind fest verdrahtet und unterstützen ein MIPI-CSI2-Interface bis 4x 1,5 GBit/s. Die in den FPGAs verwendete Quantum-Technologie ermöglicht geringe Verlustleistungen sowie schnelle Systemfrequenzen. Die Tabelle zeigt die im Edge-Vision-SoC-Design verwendeten Blöcke. Der RISC-V Prozessor, hier in der Ruby-Konfiguration, arbeitet mit bis zu 100 MHz Systemfrequenz. Für die neue Trion-Titanium-FPGA-Familie ist eine Systemgeschwindigkeit von mehr als 250 MHz, sowie die Unterstützung von DDR4/LPDDR4 vorgesehen. Der RISC-V-Prozessor ist vorkonfiguriert

Bild: Efinix

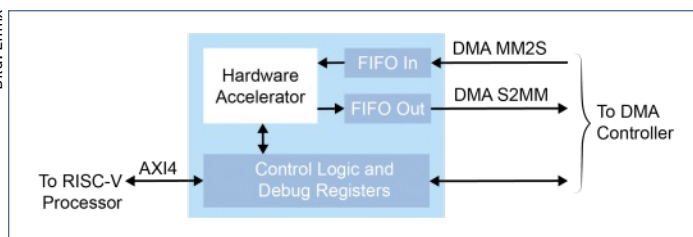


Bild 3: Das Eingangs- und das Ausgabe-Modul sowie das Hardware-Beschleuniger-Modul sind in einem standardisierten Wrapper eingebettet.

Bild: Efinix

BUILDING BLOCK	LE	FF	ADD	LUT	MEMORY BLOCK (M5K)	MULTIPLIER
Edge Vision SoC (Total)	23097	12409	4151	13610	171	4
RubySoC	–	5987	1263	6779	77	4
DMA Controller	–	5029	1418	5576	52	0
Camera	–	516	883	582	18	0
Display	–	180	117	112	10	0
HW Accel	–	607	444	402	14	0

Tabelle: Die im Design verwendeten Ressourcen im Überblick. Bild: Efinix

riert und kann zusammen mit der Entwicklungsplattform kostenlos von der Efinix Webseite heruntergeladen werden.

Mit solch einer flexiblen Lösung lassen sich Endgeräte effizienter gestalten bei gleichzeitiger Optimierung der Verlustleistung und der entstehenden Kosten. Das beschriebene System ist für eine Kamera ausgelegt, kann aber mit nur geringem Aufwand auch auf andere Systeme adaptiert werden. Zum Beispiel können für KI-Applikationen die erforderlichen Neuronalen Netze in ein solches Hardware-Beschleuniger-Modul aus-

gelagert werden. Der Prozessor kann über Standardaufrufe diesen Ablauf kontrollieren und ggf. zusätzliche Informationen zu den KI-Ergebnissen hinzufügen bzw. die Ergebnisse interpretieren. Weitere Applikationsbeispiele werden bald verfügbar sein. Diese Lösung ist sicherlich nicht direkt mit festverdrahteten SoC-Lösungen zu vergleichen, stellt aber eine gute Alternative dar, wenn Verlustleistung und Wirtschaftlichkeit eine wichtige Rolle spielen.

// ME

Efinix

Kryptographie-Schlüssel sicher injizieren und nutzen

Eine manipulationsgeschützte Schlüsselverwaltung ab Werk ist essenziell für Controller, die in vernetzten Geräten zum Einsatz kommen. Renesas hat dafür die Secure Crypto Engine entwickelt.

GIANCARLO PARODI *

Die zunehmende Nutzung vernetzter Geräte führt zu einer verstärkten Nachfrage nach Daten- und Informationssicherheit. So müssen etwa die Kommunikationskanäle und die lokalen Speicher der Geräte gesichert werden, um die Vertraulichkeit und Integrität der Benutzerdaten zu schützen. Entwickler von Embedded Devices

stehen vor der Frage, wie sie die Möglichkeiten ihrer Anwendungen weiterentwickeln können, um diese neuen Anforderungen zu erfüllen. Gesucht ist auch ein geeigneter Migrationspfad für vorhandene, aber veraltete Lösungen. Der Umgang mit Krypto-Schlüsseln ist dabei ein Problem, besonders deren Verwaltung in der Produktion – etwa wenn Fertigungsanlagen bislang ohne spezifische Security-Maßnahmen gearbeitet haben.

Ein Ansatz besteht darin, alle erforderlichen kryptografischen Funktionen mit Hilfe einer Software-Bibliothek zu implementieren, die entweder kommerziell unterstützt

oder von der Open-Source-Gemeinschaft übernommen wird. Dies ist der wohl einfachste Ansatz, hat jedoch einige Nachteile. Typischerweise verbraucht eine solche Lösung eine beträchtliche Menge an CPU-Ressourcen, was den Stromverbrauch und die Systemlatenz erhöht. Je nach Prozessortyp und Taktfrequenz nehmen viele Operationen mit asymmetrischer Kryptographie wie dem Setup einer TLS-verschlüsselten Verbindung zum Aufbau eines sicheren Kanals zwischen dem Gerät und einem Server recht viel Zeit in Anspruch, wodurch sich die Verbindungsleistung verlangsamt.

Das Integrieren eines Sicherheitscontrollers (Secure Element, SE) scheint das Problem zu lösen. Ein SE ist ein geschlossenes, isoliertes Subsystem und bietet höchstmöglichen physischen Manipulationsschutz. Entwickler müssen sich nicht um die kryptographische Implementierung kümmern, was den Aufwand reduziert. Es gibt aber auch Schwachstellen. Der zusätzliche Chip kostet Geld. Die Verbindung zur MCU über einen seriellen Bus wie I2C oder SPI ist ein Single-Point-of-Failure mit geringer Bandbreite. Ein SE hat nur wenig Speicherplatz für Benutzerschlüssel. Oft ist der – meist geringe – Durchsatz der implementierten Krypto-Algorithmen nicht dokumentiert. Eine AES-Beschleunigung etwa zum Verarbeiten von Massendaten (z.B. eines verschlüsselten Firmware-Updates) ist möglicherweise überhaupt nicht verfügbar. Hinzu kommt: Die im Chip implementierte Funktionalität lässt sich später nicht mehr aktualisieren. Und: Das Einfügen des Schlüssels in das SE erfordert möglicherweise zusätzliche Tools, die in den Produktionsfluss integriert werden müssen. Programmierhäuser könnten einen Teil des Aufwands abkürzen und einen sicheren Programmierservice anbieten, was aber ebenfalls die Kosten treiben würde.

Eine Alternative ist das sichere Subsystem der Renesas Secure Crypto Engine (SCE, Bild 1). Bei zum SE vergleichbarer Funktion stellt



* Giancarlo Parodi
... ist Principal Engineer bei Renesas Electronics

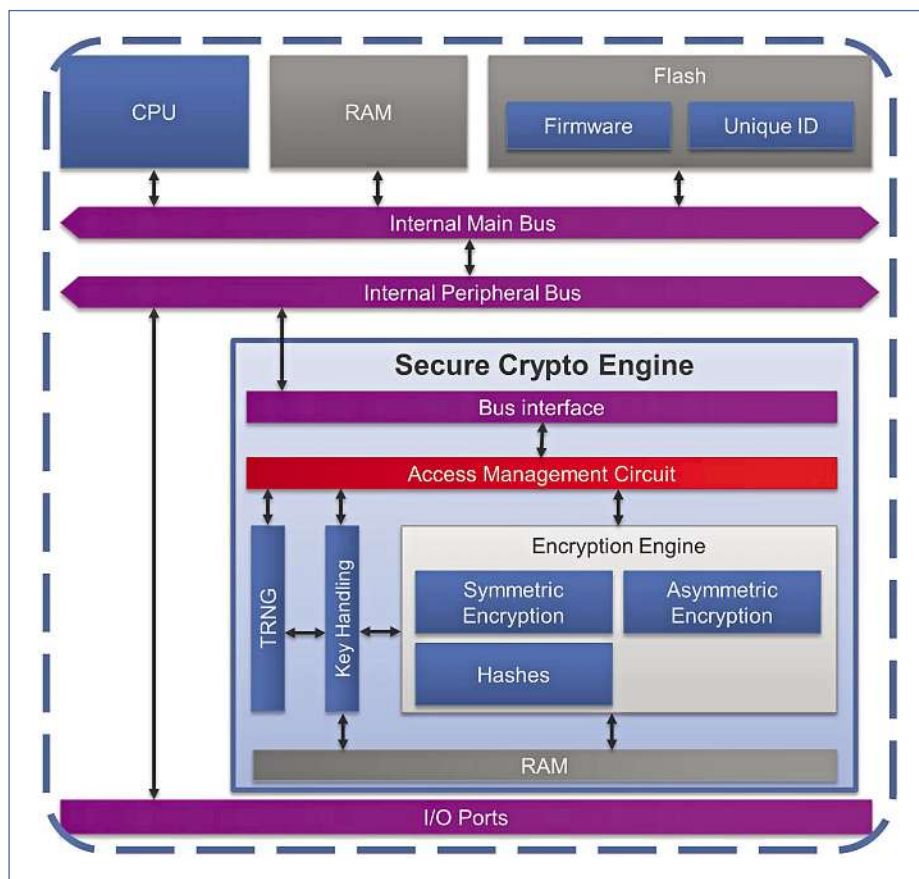


Bild 1: Implementierung der Secure Crypto Engine mit eigenem RAM und Kryptographie-Prozessor.

es viel mehr Rechenleistung bereit. Es unterstützt AES-, RSA-, ECC- und Hashing-Algorithmen sowie eine echte Zufallszahlengenerierung (TRNG; True Random Number Generation) innerhalb einer isolierten On-Chip-Umgebung. Die SCE ermöglicht eine sichere, unbegrenzte Schlüsselspeicherung unter Verwendung eines werkseitig programmierten 256-Bit Hardware Unique Key (HUK) zur Ableitung von Keyschlüsseln, die zum Wrapping (Verschlüsseln) der Kundenschlüssel verwendet werden können. Das Kryptosubsystem Secure Crypto Engine (SCE9) ist vollständig in der MCU enthalten und durch einen Access Management Circuit geschützt, der die Kryptoengine bei illegalen Zugriffsversuchen abschalten kann. Es führt alle Klartext-Krypto-Operationen unter Verwendung seines eigenen dedizierten internen Speicherbereichs durch, auf den kein CPU- oder DMA-zugänglicher Bus zugreifen kann. Die erweiterten Schlüsselspeicher- und Schlüsselverarbeitungsfunktionen der SCE9 stellen sicher, dass keine Klartextschlüssel jemals offengelegt oder im CPU-zugänglichen RAM oder nichtflüchtigen externen Speicher gespeichert werden.

Das Renesas-Subsystem Secure Crypto Engine

Die RA-Mikrocontroller-Familie umfasst die SCE in allen konnektivitätsfähigen Gruppen wie die RA6- und RA4-Derivate. Der RA6M4 ist die erste MCU der RA-Familie in der neuen Generation von Security-fokussierten MCUs von Renesas. Die implementierten Security-Features kombiniert mit umfassender Peripherie-IP sowie Funktions- und Pin-Kompatibilität zwischen den MCU-Serien machen die MCUs der RA-Familie zur geeigneten Wahl für nahezu jede vernetzte Embedded-Lösung.

Bei einer Security-fokussierten Anwendung besteht die Herausforderung schließlich nach wie vor darin, wie die Schlüsselbereitstellung durchgeführt werden kann, ohne den Schlüsselinhalt während der Produktionsphase offenzulegen. Ein Ansatz ist die Verwendung eines Hardware-Security-Moduls (HSM), eines Programmier-Tools, das die Anwendungsschlüssel sicher verwalten und mit einer Provisioning-Backbone-Infrastruktur kommunizieren kann. Eine andere Möglichkeit ist die Verwendung eines Programmierdienstes in einem Highspeed-Programmier-Equipment, das viele ähnliche Bausteine in einem Batch für hochvolumige Anwendungen bereitstellen könnte.

Alle MCUs von Renesas unterstützen bereits das werkseitige Programmieren über serielle und USB-Schnittstellen – und damit

einen zuverlässigen und kostengünstigen Mechanismus für die Massenproduktion. Über diese Schnittstelle lassen sich Anwendungen samt ihrer Parameter programmieren. Renesas hat nun die integrierten Funktionen seiner Mikrocontroller der RA6M4- und Folgeserien so erweitert, dass sie das Bereitstellen von Schlüsseln über die gewohnten Programmierschnittstellen unterstützen. Ein dediziertes Befehlsprotokoll ermöglicht es dem Anwender, einen Installationsschlüssel und einen verschlüsselten Nutzlast (der Anwendungsschlüssel wird injiziert) über die Boot-Firmware an die SCE-Grenze zu übergeben (Bild 2).

Das SCE-Subsystem enthält einen Hardware-Root-Key (HRK), mit dem sich Key Encryption Keys (KEKs) ableiten lassen. Mit diesen KEKs lassen sich Benutzerschlüssel verschlüsselt innerhalb der SCE übertragen. In der sicheren Umgebung wird für die Verwendung während des Installationsprozesses entschlüsselt. Wir bezeichnen den aus dem Verfahren KDF (HRK) resultierenden KEK als HRK*. Der HRK selbst ist abgesichert und außerhalb der SCE nicht zugänglich. Auch der HRK* existiert nur innerhalb der SCE. Wie bereitet ein Anwender also seinen Anwendungsschlüssel vor, der über die Boot-F/W-Schnittstelle im Werk eingespeist wird? Bild 3 zeigt diesen Vorgang. Der erste Schritt ist das Generieren eines Installationsschlüssels (grüner Schlüssel) und des Anwendungsschlüssels (gelber Schlüssel). Der Installationsschlüssel ist ein symmetrischer AES-Schlüssel, der Anwendungsschlüssel kann, muss aber nicht symmetrisch sein. Er muss lediglich ein von der SCE unterstütztes Format haben. Dieser Vorgang muss in einer sicheren Umgebung erfolgen. Der Anwendungsschlüssel ist die Nutzlast und wird mit dem Installationsschlüssel verschlüsselt. Natürlich muss auch der Installationsschlüssel verschlüsselt werden. Ansonsten kann jeder den Anwendungsschlüssel dechiffrieren, der die Protokollkommunikation abfängt oder die Installationsdaten kopiert, die in dem Programmier-Equipment der Fabrikanlage gespeichert sind.

Der Schlüsseldienst kommt kostenlos

Aber woher weiß der Anwender von der HRK*, um den Installationsschlüssel damit zu verschlüsseln? Das kann er nicht. Wie also teilt Renesas das Geheimnis, ohne es mitzuteilen? Um den Installationsschlüssel mit der HRK* zu verschlüsseln, ohne die inhaltlichen Informationen über die HRK selbst preiszugeben, bietet Renesas seinen Kunden einen kostenlosen Service an. Dieser Service

RELAX
It's all rugged.



CompactPCI Serial®
Modular Architecture

Embedded Blue®
Boxed Solutions

With:
Intel® Core i7™
Intel® Atom™
ARM® v8 Networking SoC

EKF Elektronik GmbH

+49 (0) 2381 68900
www.ekf.com • sales@ekf.de

The diagram illustrates the internal components of an integrated circuit. A dashed box encloses the main components: Boot firmware (green), CPU (grey), Secure Subsystem (yellow), and Memory (grey). Below the dashed box is the Programming Interface (green). A Communication channel (indicated by a double-headed arrow) connects the Programming Interface to the CPU. Bidirectional arrows show communication between Boot firmware and CPU, between CPU and Secure Subsystem, between CPU and Memory, and between Secure Subsystem and Memory. A long arrow at the bottom points from the Programming Interface towards the Secure Subsystem and Memory.



The diagram illustrates the key management process for Renesas MCUs. It shows a user (laptop icon) interacting with an MCU (chip icon). The process involves installing a user key and storing it wrapped by the MCU-unique key (HUK). A legend defines the key types: MCU HUK*, Install Key, User Key, and MCU-unique Key*. A note states: 'No plaintext user keys are transferred or exposed. Install key lifetime ends after user key is installed. User key is stored wrapped by MCU-unique key (HUK).'

Renesas Electronics

MIKROCONTROLLER AUF ARM-CORTEX-M-BASIS

Energieeffiziente 32-Bit-Controller

Toshiba Electronics erweitert sein Angebot stromsparender und leistungsfähiger 32-Bit-Mikrocontroller der TXZ+-Reihe auf Arm Cortex-M-Basis. Die neuen Produkte sind für Antriebssteuerungen, Haushaltsgeräte, Mensch-Maschine-Schnittstellen und vernetzte IoT-Infrastruktur optimiert. Es gibt Modelle mit Flash-Speichergrößen bis 2 MByte und einer Lebensdauer von 100.000 Schreibzyklen in vielfältigen Gehäuseoptionen. Jede MCU verfügt über einen mehrkanaligen 12-Bit A/D-Wand-



ler (ADC) und Schnittstellen wie CAN, Ethernet und USB. Integriert sind auch Sicherheitsmechanismen wie Secure Boot.

Toshiba Electronics

AUTOMOTIVE-CONTROLLER

Erste 8-Bit-MCUs für CAN-FD

Microchip stellt die nach eigenen Angaben ersten 8-Bit-MCUs für CAN-FD-Netzwerke vor. Mit den Bausteinen der Serie PIC18-Q84 sollen Automotive-Entwickler Systemfunktionen flexibler erweitern können. Die MCUs werden dafür durch umfangreiche Core-unabhängige Peripherie (CIPs; Core Independent Peripherals) unterstützt, die eine Vielzahl von Aufgaben erledigt, ohne dass ein Eingreifen der CPU erforderlich ist. Dies spart Zeit und Kosten beim Anschluss von Systemen an ein CAN-FD-Netz-



werk. Zum Übertragen von Sensordaten sind beim Einsatz von PIC18-Q84-MCUs keine Gateways erforderlich.

Microchip

AUF KI UND SICHERHEIT GETRIMMT

Controller für AIoT-Anwendungen



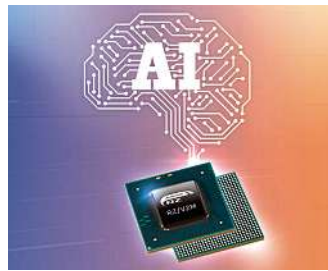
Auf den Bedarf des AIoT-Marktes hat Espressif seine neue 240-MHz-Dual-Core-MCU ESP32-S3 zugeschnitten. Der Baustein bietet laut Hersteller ohne externe Komponenten alle erforderli-

chen Sicherheitsanforderungen für die Entwicklung sicher vernetzter Geräte. Er unterstützt AES-XTS-basierte Flash-Verschlüsselung und RSA-basierten Secure Boot. Darüber hinaus verfügt der ESP32-S3 über ein Modul für digitale Signaturen und ein HMAC-Modul, die ähnliche Funktionen bieten wie das Secure Element. Der Baustein stellt per „World Controller“ zwei isolierte Ausführungsumgebungen bereit.

Espressif / Macnica

ECHTZEIT-INFERENZ AM EDGE

KI-Prozessor behält kühlen Kopf



Renesas hat seine RZ/V-Serie von Mikroprozessoren (MPUs) vorgestellt. Diese ist mit dem DRP-AI IP (Dynamically Reconfigurable Processor) ausgestattet – einem für Bildverarbeitung optimierten

KI-Beschleuniger. Der RZ/V2M ist das erste Produkt der Serie und bietet eine Kombination aus Echtzeit-KI-Inferenz und hoher Energieeffizienz von typisch 4 W für Embedded-Anwendungen in Industrie, Infrastruktur und Einzelhandel. Der Baustein benötigt keine Kühlkörper und Lüfter und eignet sich daher auch für den Einsatz in kompakten Geräten, was die KI-Integrationsmöglichkeiten in Embedded-Anwendungen erweitert.

Renesas Electronics

MADE IN GERMANY

proboxx - Zusammenspiel aus Design und Funktion

→ innovativer Befestigungsmechanismus für Wand- und Profilmontage

- Leergehäuse
- AS-Interface
- Funk (sWave®)

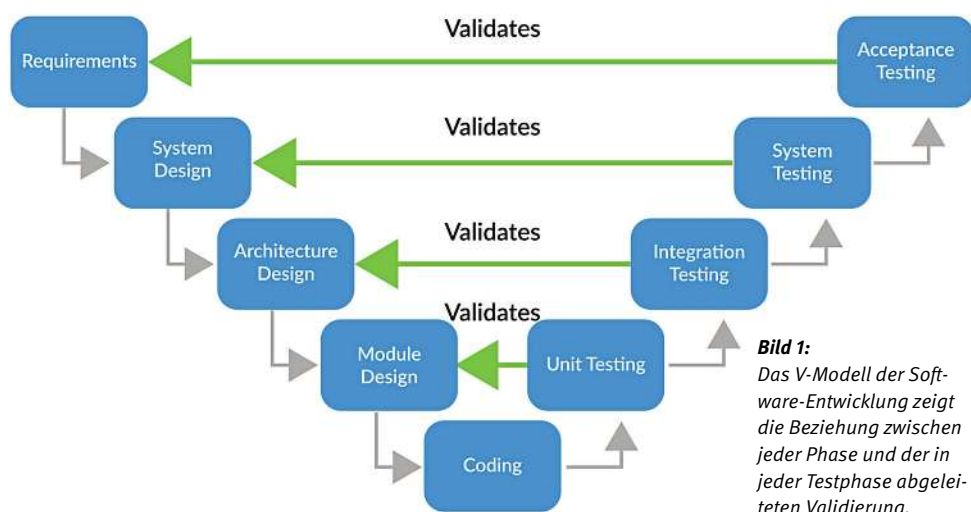
- M12-Gehäuse
- IO-Link

SCHLEGEL®
ELEKTROKONTAKT
www.schlegel.biz

Verifizierung vs. Validierung in eingebetteter Software

Für sicherheitskritische Software gibt es strenge, oft auch bestimmte Industriestandards, festgeschriebene Ansätze zur Verifizierung und Validierung. Was ist der Unterschied?

RICARDO CAMACHO *



dukten beinhaltet die Validierung, dass die Kunden das Produkt selbst sehen, ausprobieren und testen. Bei Software hingegen besteht sie in der Ausführung der Software und der Demonstration, dass sie funktioniert. Dies beinhaltet typischerweise:

- Codeausführung zum Nachweis der korrekten Funktionalität.
- Ausführung in Zielumgebungen.
- Stress-, Performance-, Penetrations- und andere nicht-funktionale Tests.
- Direkte und oft ausgeführte Abnahmetests bei Kunden.

■ Nutzung von Artefakten aus Verifikationsprozessen, um die Traceability von Anforderungen auf die Endfunktionalität zu veranschaulichen, insbesondere für bestimmte Sicherheitsfunktionen.

Es ist wichtig, nicht nur die Unterschiede zwischen Verifizierungs- und Validierungszielen zu verstehen sondern auch, dass die Softwareentwicklung beides braucht. Weil diese Aktivitäten einen Großteil des Aufwands bei der Softwareentwicklung ausmachen, sind Unternehmen stets bestrebt, sie zu rationalisieren, ohne die Sicherheit oder die Qualität zu beeinträchtigen.

Herangehensweise an Verifizierung an Validierung

Das V-Modell zeigt den Ansatz für eine formale Verifizierung und Validierung, der bei der Entwicklung von sicherheitskritischer Software zur Verwendung kommt (Bild 1). Es veranschaulicht die Aktivitäten in jeder Entwicklungsphase und die Beziehungen zwischen ihnen. In jeder Testphase werden vollständigere Teile der Software validiert und mit der Phase verglichen, die sie definiert. Es gibt Möglichkeiten, um Agile, DevOps und CI/CD in diese Art der Produktentwicklung einzubinden.

Die Verifizierung hingegen stellt sicher, dass jeder Schritt vollständig und korrekt ausgeführt wird (vgl. Bild 2). Die Verifizierung umfasst Reviews, Durchläufe, Analy-

Die offiziellen Definitionen von Validierung und Verifizierung stehen im IEEE Standard Glossary of Software Engineering Terminology.

■ **Verifizierung:** Der Prozess, bei dem festgestellt wird, ob die Produkte einer bestimmten Phase des Software-Entwicklungszyklus die in der vorangegangenen Phase festgelegten Anforderungen erfüllen oder nicht.

■ **Validierung:** Der Prozess der Evaluierung von Software am Ende des Software-Entwicklungsprozesses, um die Übereinstimmung mit den Software-Anforderungen sicherzustellen.

Barry Boehms „Verifying and Validating Software Requirements and Design Specifications“ bringen es kurz und knapp auf den Punkt. Verifizierung heißt demnach: „Baue

ich das Produkt richtig?“. Validierung bedeutet dagegen: „Baue ich das richtige Produkt?“

Das Ziel von Verifizierung und Validierung

Das ultimative Ziel lautet, das richtige Produkt zu bauen. Mehr noch, es geht darum, sicherzustellen, dass das Produkt die Qualität, die Sicherheit und den Schutz bietet, um zu gewährleisten, dass es das richtige Produkt bleibt. Die Verifizierung stellt als Teil des Software-Entwicklungsprozesses sicher, dass die Arbeit korrekt ist. Dazu umfasst sie in der Regel eine Konformität mit Industriestandards als Garantie, dass Prozess und Artefakte den Richtlinien entsprechen; Reviews, Durchläufe, Inspektionen; Statische Analyse und andere Aktivitäten an Artefakten, die während der Entwicklung entstehen; und Durchsetzung von Architektur-, Design- und Programmierstandards.

Die Validierung ist der Nachweis, dass das Endprodukt die Anforderungen – sie umfassen Funktionalität, Zuverlässigkeit, Leistung und Sicherheit – erfüllt. Bei physischen Pro-



* Ricardo Camacho
... ist Senior Technical Marketing Manager bei Parasoft in Carlsbad, Kalifornien.

sen, Nachvollziehbarkeit, Test- und Codeabdeckung und andere Aktivitäten, um sicherzustellen, dass der Prozess und das Produkt korrekt erstellt werden. So ist etwa die Ausführung von Unit-Tests eine Validierungsaktivität, die Sicherstellung der Nachvollziehbarkeit, der Codeabdeckung und des Testfortschritts der Unit-Tests eine Verifizierung. Als Hauptaufgabe stellt letztere sicher, dass die gelieferten Artefakte aus der vorherigen Phase der Spezifikation entsprechen und mit den Richtlinien übereinstimmen.

Namensverwirrung zwischen Verifizierung und Validierung

Unternehmen verwenden Verifizierung und Validierung oft leicht unterschiedliche. Es ist üblich, die Validierung als eine Aktivität zu betrachten, die gegen Ende der Entwicklung stattfindet, bei der die Teams das Produkt dem Kunden oder Testern präsentieren und die Erfüllung aller Anforderungen nachweisen. In der Embedded-Entwicklung kommt immer noch das Konzept des Factory Acceptance Test (FAT) zum Einsatz. Ein weiterer Unterschied in der Auffassung ist DO-178B/C und verwandter Standards für die Entwicklung von sicherheitskritischer Avionik-Software, die die Validierung als solche nicht kennen.

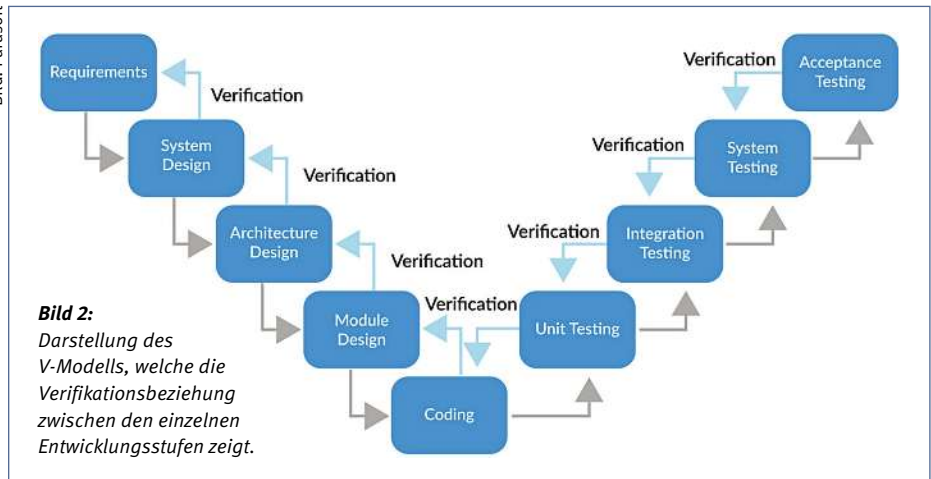
Hybride DevOps-Pipelines für sicherheitskritische Software

In vielen Embedded-Software-Unternehmen ist die Implementierung eines vollständig agilen Prozesses nicht mit den Auflagen der Sicherheitsstandards der Branche kompatibel. Artefakte, Code, Testergebnisse und Dokumentation haben oft vorgeschriebene und festgelegte Liefertermine. Der Fortschritt basiert auf solchen Meilensteinen. Teams können hybride Modelle nutzen, um mit iterativen und agilen Methoden intern Meilensteine für Deliverables zu erreichen (Bild 3).

Der Grund, dies in einer Diskussion über Verifizierung und Validierung anzusprechen, ist, dass sich viele Vorteile der kontinuierlichen Integration und des Testens auf komplexe sicherheitskritische Anwendungen anwenden lassen. Ein Teil davon ist, die Verifizierung und Validierung so früh wie möglich in den Entwicklungsprozess vorzulegen („Shift-Left“). Zum Beispiel gibt es keinen Grund, Unit-Tests zu verzögern, bis alle Units kodiert sind, oder mit der Überprüfung des Codes mit statischer Analyse zu warten, bis er für die Integration bereit ist.

Ähnlich sollten Entwickler mit Integrations-tests beginnen, sobald die Unit-getesteten Komponenten fertig sind. Die Kombination von Automatisierung mit einem konti-

Bild: Parasoft



nuierlichen, iterativen Ansatz bietet enorme Vorteile für die Software-Validierung und -Verifizierung. Moderne Anbieter wie Parasoft verfügen über Tools, die das Testen, die Qualität und die Sicherheit in allen Phasen des SDLC ansprechen (siehe auch Bild 4).

Die Verifizierung umfasst die Arbeit, die sicherstellt, dass jede Phase der Entwicklung die Spezifikation des vorherigen Schritts er-

füllt. Bezogen auf die Software-Codierung und -Tests bedeutet Verifizierung die Sicherstellung, dass der Code das Moduldesign und letztendlich das High-Level-Design und die darüber liegenden Anforderungen erfüllt.

Darüber hinaus gewährleistet die sie die Erfüllung der Anforderungen auf Projektebene. Dazu gehört das Einhalten von Industriestandards, Risikomanagement, Rückverfolg-

Anzeige

Universal Debug Engine®
Multicore • Debugging • Trace • Test Automation

AURIX / TriCore
Arm Cortex-M/R/A
SPC5 / MPC5xxx
RH850 / R-Car
XE166 / XC2000

pls Development Tools
www.pls-mc.com

Bild 3:
Das Agile V-Model von Parasoft.

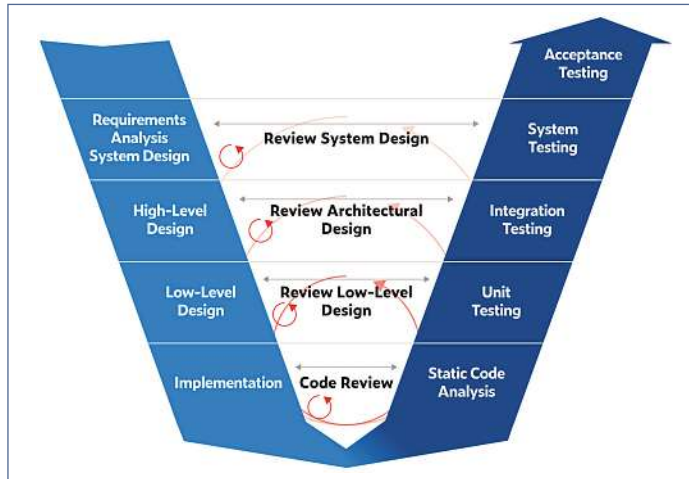


Bild: Parasoft

ponenten, Tests und aufgezeichneten Abweichungen in Beziehung setzen.

Schnellere Validierung durch Testautomatisierung

Validierung ist der Nachweis, dass ein Produkt seine Anforderungen erfüllt, wobei die Ausführung von Code entweder isoliert für Unit-Tests oder in verschiedenen Phasen der Integration erforderlich ist. Für die Entwicklung von Embedded Software bedeutet die Automatisierung dieser Testreihen eine enorme Zeitersparnis.

Validierung verlangt die Ausführung auf der Zielhardware. Die Optimierung von Regressionstests schöpft verfügbare Ressourcen, Mitarbeiter und Hardware bestmöglich aus. Hier profitieren Entwickler von Testautomatisierungstools, weil diese die Abhängigkeit von manuellen Tests reduzieren – unter Beibehaltung der Nachvollziehbarkeit und Code-Abdeckung aller Ergebnisse:

- Die Automatisierung aller Testsuiten minimiert das manuelle Testen und reduziert den Testengpass durch begrenzte Hardwareverfügbarkeit.
- Die Target- und Host-basierte Testausführung unterstützt je nach Bedarf verschiedene Validierungstechniken.
- Shift-left- (Vorverlegte) Tests beginnen, sobald Teams Code entwickeln. Hierbei kommen Unit-Testing-Frameworks zum Einsatz und sobald der Code fertig ist, werden automatisch Softwaresätze zum Testen generiert. Unterstützung für testgetriebene Entwicklung und kontinuierliches Testen ist mit dem Reifegrad des Prozesses im Unternehmen verfügbar.
- Verwalten von Änderungen mit der intelligenten Testausführung, um sich auf Tests für den geänderten Code und die betroffenen Abhängigkeiten zu konzentrieren.
- Rückverfolgbarkeit in zwischen Code, Tests, Ergebnissen der statischen Analyse und Anforderungen sowie Unterstützung für unternehmensweite Application Lifecycle Management (ALM)-Tools.

Weil Validierung und Verifizierung einen großen Teil der Ressourcen in der Embedded-Produktentwicklung verschlingen, haben sich moderne Toolanbieter Lösungen überlegt. So können Entwicklungsteams etwa mit der Software-Testautomatisierungssuite von Parasoft ein das Testen in die frühen Phasen der Entwicklung vorverlegen. Zugleich stellt es es die Traceability, die Aufzeichnung der Testergebnisse, die Details zur Codeabdeckung, die Berichterstellung und die Compliance-Dokumentation sicher. // SG

Bild: Parasoft

barkeit und Metriken (Codeabdeckung und Konformität).

Beschleunigen der Verifizierung mittels Automatisierungstools

Mit Software-Test-Automatisierungstools können Entwickler die Verifizierung beschleunigen, indem sie die vielen mühsamen Aspekte der Aufzeichnung, Dokumentation, Berichterstattung, Analyse und des Reportings automatisieren:

- Frühestmöglicher Einsatz der statischen Analyse, um Qualität und Sicherheit schon beim Programmieren zu gewährleisten. Zudem schützt die statische Analyse vor zukünftigen Fehlern und Schwachstellen und verringert so die nachgelagerten Auswirkungen von Bugs, die bei Inspektionen und Tests übersehen werden.
- Das Automatisieren der Konformität mit Programmierstandards, um den manuellen Aufwand zu reduzieren und Code-Inspektionen zu beschleunigen.
- Rückverfolgbarkeit in beide Richtungen

für alle Artefakte, um sicherzustellen, dass der Code und die Tests die Anforderungen erfüllen.

- Code- und Testabdeckung, um sicherzustellen, dass alle Anforderungen implementiert sind, und die Implementierung wie erforderlich getestet wird.
- Berichte und Analysen, um die Entscheidungsfindung zu unterstützen und den Fortschritt zu verfolgen. Die Entscheidungsfindung muss auf den von den automatisierten Prozessen gesammelten Daten basieren.
- Automatisierte Dokumentationserstellung aus Analysen und Testergebnissen, um das Einhalten von Prozessen und Standards zu unterstützen.
- Automatisierte Konformität mit Standards, um den Aufwand und die Komplexität zu senken, indem die sich am meisten wiederholenden und langwierigen Prozesse automatisiert werden. Zudem lassen sich die Projekthistorie verfolgen und Ergebnisse mit den Anforderungen, Softwarekom-

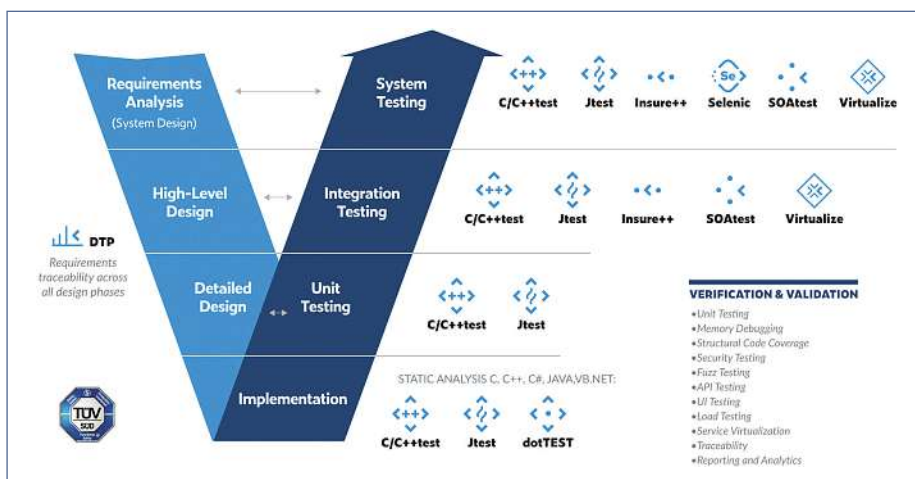
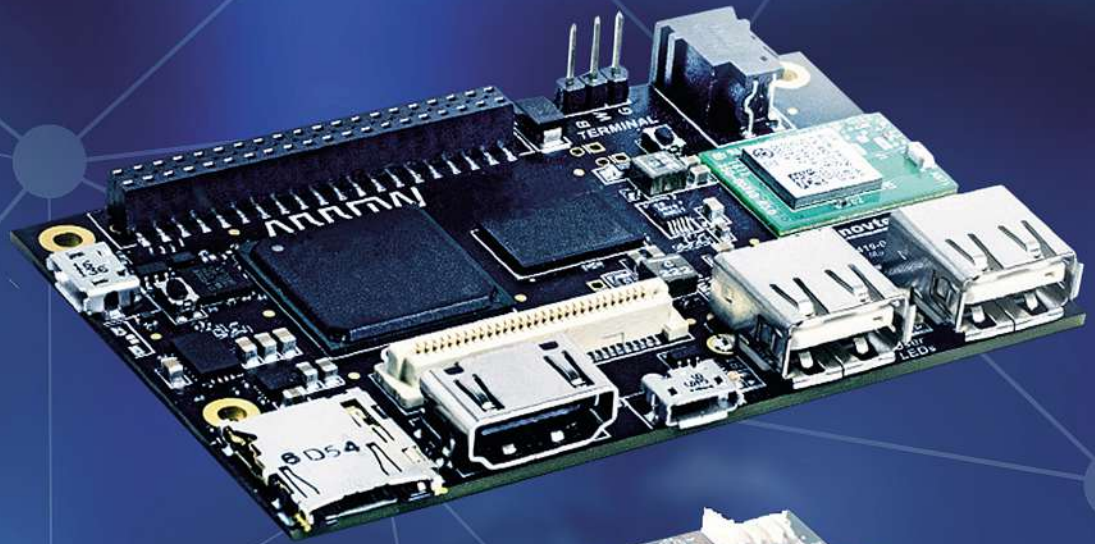


Bild 4: Moderne Tools lassen sich in jede Phase des V-Modells im SDLC einfügen.

Parasoft



Teilen Sie Ihr Wissen

 **FPGA-Conference**
Europe

6. - 8. Juli 2021

München-Dornach

Call for Papers – Jetzt Vortrag einreichen!

Werden Sie Referent auf Europas größter Konferenz mit dem Fokus auf Field Programmable Gate Arrays. Geben Sie der Branche mit Ihrem Vortrag neue Impulse, und profitieren Sie als Referent von den wertvollen Kontaktmöglichkeiten auf der Konferenz und der begleitenden Fachausstellung.

www.fpga-conference.eu

Eine Veranstaltung von Marken und Partnern der



VOGEL COMMUNICATIONS
GROUP

ELEKTRONIK
PRAXIS

 **PLC2**

Ein Automotive-Betriebssystem für neue Herausforderungen

Zukünftige Fahrzeugsysteme müssen harte Echtzeitfähigkeit mit hoher Rechenleistung kombinieren. Das erfordert eine wesentliche Änderung auf der Betriebssystemebene.

MASAKI GONDO *

Vernetzte, autonome und nachhaltige Mobilität ist eine verlockende Vision, die alle Automobil-Innovatoren stark unter Druck setzt. Bordsysteme kommender Fahrzeuge müssen hohe Rechenleistung erbringen und gleichzeitig harte Echtzeitfähigkeit zur Verfügung stellen. Dies kann durch die Hardware trotz performanter Mehrkernprozessoren nur teilweise erfüllt werden. Damit ein Automotive-Betriebssystem eine verteilte Mikrokern-Architektur ordentlich unterstützen kann, ist eine grundlegende Änderung des OS Betriebssystems nötig.

Disruptive Entwicklungen im Automobilmarkt, wie Konnektivität, autonomes und

automatisiertes Fahren, Mobilität als Dienstleistung und Elektrifizierung des Antriebsstrangs, stellen Hochleistungsanforderungen an die großen Fahrzeugrechner, die aber gleichzeitig stromsparend sein müssen. Diese Hochleistungsrechner müssen mit einer Anzahl kleinerer Steuergeräte koexistieren, die die traditionellen Karosserie-, Chassis-, Antriebs- und andere elektrische Systeme des Fahrzeugs steuern. In der Vergangenheit gab es Premium-Autos mit bis zu 100 Steuergeräten, die jeweils für eine bestimmte Funktion optimiert waren.

Diese elektrische/elektronische (E/E) Architektur der Fahrzeuge ändert sich gerade, da immer umfangreichere Funktionen eingeführt werden. Verteilte Steuergeräte verschmelzen zu größeren Domänencontrollern, um die Komplexität der Fahrzeugverkabelung, das Gewicht und vor allem rapide steigenden Entwicklungskosten zu kompensieren.

Zunehmend aggregieren zentralisierte Hochleistungsrechner mehrere Steuergeräte über Domänengrenzen hinweg. Dies wiederum führt zu einer neuen Kategorie von Prozessoren, die mit vielen, heterogenen Kernen die leistungsfähigste und energiesparendste Lösung darstellen, um gleichzeitig die alten, bereits vorhandenen Funktionen zu integrieren und die neuen, sehr rechenintensiven Aufgaben, zur bewältigen.

Zusätzlich erhöhen insbesondere Konnektivität und das automatisierte Fahren die Nachfrage nach neuen Sicherheitsstandards. Mittlerweile ist es eine Frage der nationalen Sicherheit, kriminelle Angreifer daran zu hindern autonome Fahrzeuge zu hacken. Etablierte Sicherheitsstandards wie ISO 26262 reichen für neue Anwendungsfälle, wie z.B. dem autonomen Fahren, nicht aus. Es werden gerade neue Standards wie SOTIF (Safety Of The Intended Functionality) und



* Masaki Gondo
...ist CTO und Vice President bei eSol.

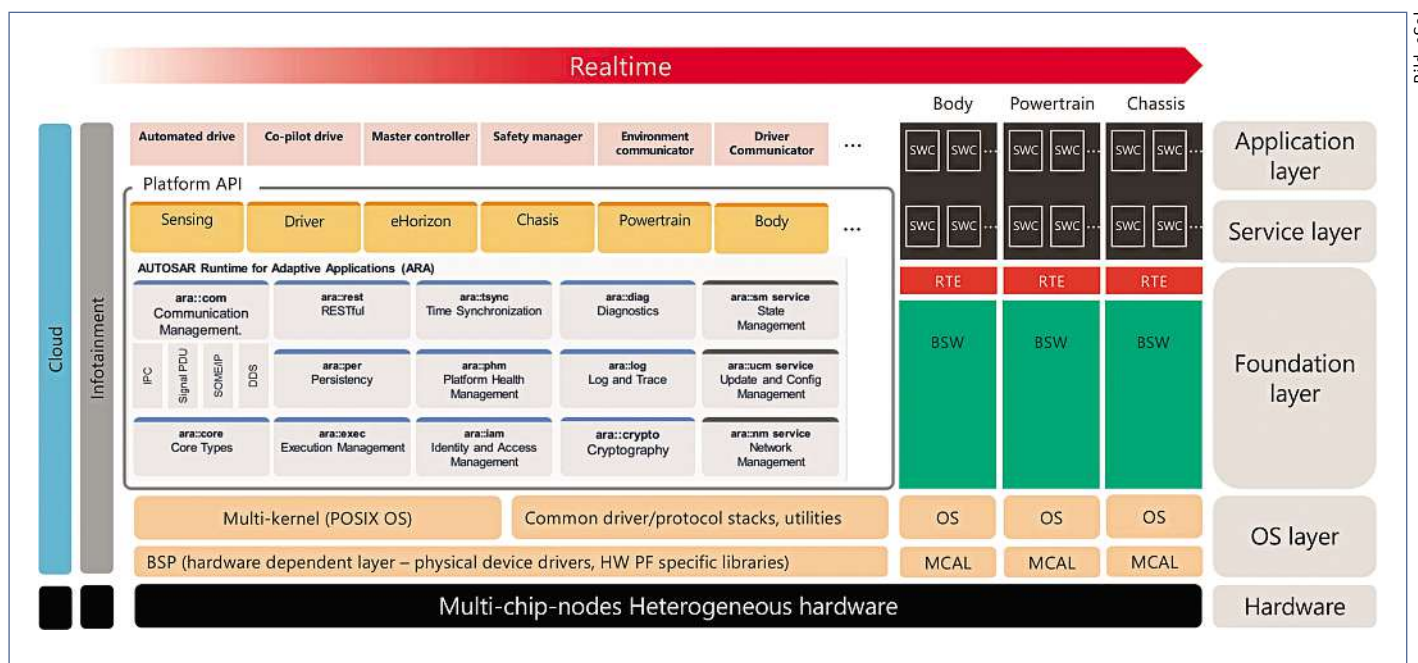


Bild: eSol

Bild 1: Diverse Software-Architekturen erfordern Software-Standards mit einem skalierbaren Ansatz.

UL4600 entwickelt, um diesen Anwendungen gerecht zu werden.

Um die Dinge noch schwieriger zu machen, erfordert die öffentliche Akzeptanz zum Thema „sicheres autonomes Fahren“, eine Diskussion über die Begriffe „Festgelegten Betriebsbereich“ (Operational Design Domain, ODD) und „Überwachung des Verkehrsgeschehens und der Fahrzeugsteuerung“ (Object Event Detection and Response, OEDR) von autonomem Fahrsystemen (Automated Driving Systems, ADS). Diese Diskussion ist aber sehr OEM-spezifisch und eng an dessen zukünftiger Fahrzeugspezifikation gebunden und wird deswegen von allen OEMs nicht öffentlich geführt.

Um den verschärften Sicherheitsherausforderungen gerecht zu werden, benötigen OEMs und Zulieferer komplette Hardware- und Software-Architekturen, die ihnen bei der Überwindung dieser Hürden helfen. Diese Architekturen müssen hoch skalierbar sein damit differenzierte Modellreihen gefertigt werden können. In der Automobilproduktion großer OEMs kann die Unterscheidung zwischen einem A-Segment- und einem F-Segment-Modell nicht nur durch einen Softwareschalter umgesetzt werden. Um die Kosten bei großen Modellfamilien unter Kontrolle zu halten, muss die Hardware und Software der E/E-Architektur skalierbar sein.

Erweiterung der Software-Architektur

Im Grunde genommen haben die OEMs die Aufgabe, Systeme zu entwickeln, die eine noch nie dagewesene Kombination aus Hochleistung, Energieeffizienz, Echtzeit (Determinismus), Kosteneffizienz und Sicherheit bieten. Und als ob die Skalierung heterogener Steuergeräte nicht schon schwierig genug wäre, wird diese Herausforderung noch dadurch weiter verschärft, dass diese Systeme für die harten Anforderungen der automobilen Massenproduktion geeignet sein müssen.

Um alle diese Herausforderungen an die E/E-Architektur des Fahrzeugs zu bewältigen, kann man nicht nur die Hardware-Plattform verbessern, sondern muss auch die Software-Plattform mit einbeziehen. Dies gilt insbesondere für die Basis aller Softwarefunktionen, dem Betriebssystem, das die Anwendungssoftware mit der Hardware verbindet.

In der Automobilindustrie wird heute eine Vielzahl von Ansätzen und Software-Plattformen verwendet. Neben Plattformen, die wie AUTOSAR auf Automotive-Standards basieren, gibt es auch Open-Source-Software-Plattformen (OSS), z.B. Linux, sowie

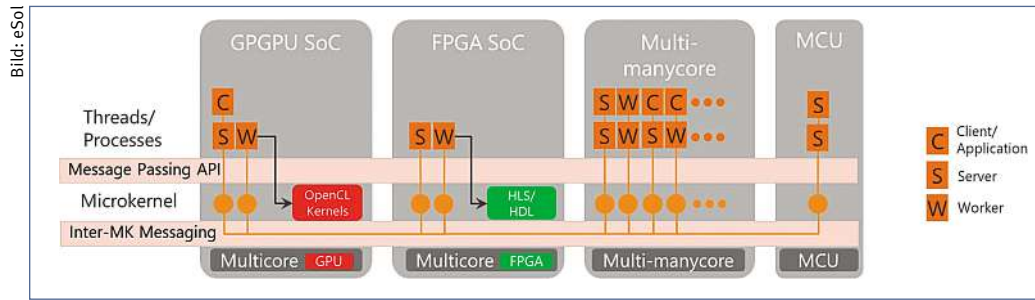


Bild 2: Mit einem Multikernel OS kann eine verteilte Software Architektur implementiert werden.

verschiedene proprietäre Plattformen, die von OEMs, Tier 1s und unabhängigen Anbietern entwickelt und gepflegt werden. Alle unterscheiden sich typischerweise in Bezug auf Softwareschichten und Funktionsbereiche, die sie abdecken.

Bild 1 beschreibt eine Softwareplattform für die Automobilindustrie, die auf einer serviceorientierten Architektur (SOA) basiert. Sie ist für den Einsatz in E/E-Architekturen mit Domänencontrollern oder den neuen zonenbasierten Architekturen geeignet.

Die SOA bietet eine flexible Plattform, um neue autonome Fahrfunktionen wie Fahrspurerkennung, Objekterkennung und Fahrerstatusüberwachung als Software-dienste bereitzustellen, auf die dann eine übergeordnete Softwareebenen über Standardschnittstellen zugreifen kann. Darüber hinaus bietet die SOA Transparenz in Bezug auf die Zuordnung, da der Standort des Servers, der einen Dienst zur Verfügung stellt, unabhängig von dessen Nutzung ist. Dies ist der Schlüssel zur Anwendung des Modells des verteilten Rechnens. Außerdem sorgt die SOA für Transparenz bei der Implementierung, indem sie ein konsistentes Interface beibehält, unabhängig davon, wie der Service innerhalb des Servers umgesetzt wird.

Die gezeigte Architektur basiert auf dem AUTOSAR Adaptive Platform Standard (AUTOSAR AP), der die Software der unteren Schicht standardisiert und eine „geplante Dynamik“ ermöglicht. Im Gegensatz zu einer vollständig offenen SOA-Architektur, handelt es sich hierbei um ein Konzept, dass eine große Anpassungsfähigkeit erlaubt, ohne dabei die Handhabung sicherheitskritischer Prozesse zu beeinträchtigen. Geplante Dynamik wird durch verschiedene Maßnahmen erreicht, wie z.B. dass alle Prozesse während der Systemintegration registriert werden müssen und die Privilegien zum Starten von Prozessen eingeschränkt werden. Auch die Kommunikation zwischen Anwendungsprozessen und externen Entitäten wird nach strengen Richtlinien verwaltet, die bereits während der Systemintegration festgelegt

werden. Mit Hilfe einer standardisierten Prozess- und Werkzeugkette muss das Softwareteam alle Prozesse, ihre Rechte, ihre Interaktionen (d.h. die Bereitstellung oder Nutzung eines Dienstes), ihre Sicherheitsrechte usw. beschreiben. Ein einfaches Beispiel: in einem Auto könnte es drei Fahrkonfigurationen wie „manuelles Fahren“, „autonomes Fahren“ und „automatisches Parken“ geben. In jeder dieser drei Konfigurationen dürfen jeweils unterschiedliche Anwendungen mit unterschiedlich aktiven Überwachungsdiensten laufen.

Störfreiheit und Interprozess-Kommunikationen

Eines der zentralen Konzepte der funktionalen Sicherheit ist die „zeitliche Störfreiheit“ (Freedom from interference, FFI). Die AUTOSAR Adaptive Plattform bietet eine gute Grundlage für eine SOA-Architektur, die insbesondere FFI abdeckt, da verschiedene Funktionalitäten einer Plattform in meist unabhängige Services aufgeteilt und wiederum von meist unabhängigen Anwendungen genutzt werden. Das bietet zwar eine gute logische Software-Architektur, aber die eigentliche Garantie für FFI muss durch Hardware-Mechanismen wie die Memory Management Unit (MMU) des Prozessors gewährleistet werden. Das Betriebssystem virtualisiert diesen Mechanismus in Form von Betriebssystemprozessen, die die Instanzen der Dienste und Anwendungen darstellen. Die FFI ist für die Sicherheit unerlässlich. Wie man in Bild 1 sehen kann, laufen viele der Komponenten als Prozesse. Diese müssen tatsächlich sehr häufig miteinander interagieren. Wenn z.B. ein Anwendungsprozess einen Service nutzen will, der als separater Prozess läuft, müssen das beide vorher miteinander kommunizieren. Das Betriebssystem stellt für diese Interaktionen besondere Funktionen zur Interprozess-Kommunikation (inter-process communication, IPC) zur Verfügung. Allerdings müssen zwei Prozesse, die ursprünglich aus Sicherheitsgründen voneinander isoliert wurden, diese

Sicherheitsgrenze nun überschreiten, um miteinander Daten austauschen zu können. Das benötigt jetzt deutlich mehr Verarbeitungszyklen als eine prozessinterne Kommunikation. Und das entwickelt sich schnell zu einem erheblichen Problem für die Systemleistung, wenn die gesamte Software auf einem Steuergerät integriert wird.

Durch den rapide zunehmenden Einsatz von Mehrkernprozessoren, wird die schnelle Kommunikation zwischen den Kernen immer wichtiger. Auf Chip-Ebene haben sich bereits Hochgeschwindigkeits-Netzwerke (high-speed network-on-chip, NoC) durchgesetzt, die eine schnelle und latenzarme Kommunikation zwischen den Kernen unterstützen. So führen all diese Aspekte bei Hochleistungsrechner im Fahrzeug zu einer Nachfrage von schnellen IPCs, bei der die gesamte Software auf einer großen Anzahl von miteinander kommunizierenden Prozessorkernen arbeiten kann. Hier werden die traditionellen Single-Core-Betriebssysteme immer weniger in der Lage sein, alle Teile des Systems angemessen zu unterstützen, um die benötigte Leistung zur Verfügung zu stellen.

Multikernel-OS für Multicore-Hardware

Eine Multikernel-Betriebssystem-Architektur, auch als verteiltes Mikrokernel-Betriebssystem bezeichnet (untere Schicht von Bild 1), ist am besten dazu geeignet, eine große Anzahl miteinander verbundener Kerne und Prozesse zu bedienen. Da sich die Fahrzeugarchitekturen immer weiter entwickeln werden, hat eSOL ein solches Multikernel-Betriebssystem mit dem Namen eMCOS (embedded multi-core OS) entwickelt, um die erforderliche Leistung und Skalierbarkeit zu bieten.

Zusätzlich zur Skalierbarkeit für kleine oder große Funktionsgruppen, bietet dieses verteilte Mikrokernel-Betriebssystem auch schnelle und deterministische Reaktionen für Echtzeitsteuerungsanwendungen, wie z.B. für die verschiedenen Stufen des autonomen Fahrens. Das Betriebssystem kann auf vielfältige Weise skaliert werden: Anwendungen können zwischen den Mikrokerneln verbunden werden und Entwickler können die Anwendungsprogrammierungsschicht an ihren jeweiligen Zweck anpassen.

Das Multikernel OS (verteilt Mikrokernel OS) unterscheidet sich von herkömmlichen Mikrokernel Oses. Da keine globalen Kernel-Blockierungen mehr notwendig sind, um die Echtzeitfähigkeit während der IPCs zu gewährleisten, stellt diese Architektur sicher, dass die Parallelität und damit die maximale Leistungsfähigkeit erhalten bleibt.

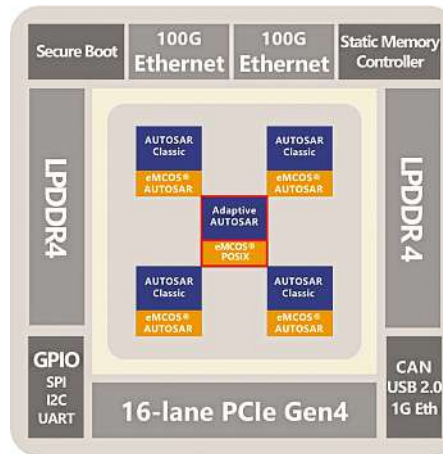


Bild 3: Mit eMCOS können AUTOSAR Adaptive Platform und Classic Platform auf einem einzigen Chip gehostet werden.

eMCOS ermöglicht durch einen besonderen Zeitplanungsmechanismus auf zwei Ebenen sowohl einen harten Echtzeit-Determinismus als auch Hochleistungsperformanz auf Basis von Lastverteilung zwischen den Kernen. Vergleichstests zeigen, dass diese Architektur auf Multi-Core- oder Multi-Chip-Hardware viel mehr Leistung bringt als bestehende Betriebssystem Architekturen, die auf Techniken wie SMP, AMP und deren Mischkonfiguration setzen.

Von eMCOS werden standardmäßig Multi-Prozess-POSIX- und AUTOSAR-Programmschnittstellen unterstützt, und es gibt spezielle APIs für Funktionen wie Distributed Shared Memory (DSM), Fast-Messaging, NUMA-Speicherverwaltung, Thread-Pool und andere, die nicht im Rahmen dieses Artikels behandelt werden können.

Bild 2 zeigt ein Beispiel wie die verteilte Mikrokernel-Architektur ein heterogenes Hardwaresystem aus GPGPU, FPGA, Many-Core und sogar MCU unterstützt. Ein solches SOA System unterscheidet nicht, ob die Software auf einem einzigen Chip läuft oder durch die Kommunikation von mehreren Chips realisiert wurde. Es spielt auch keine Rolle ob die Chips dabei über ein Netzwerk zwischen verschiedenen ECUs kommunizieren.

Unterstützung von Manycore-Hardware

Zusätzlich zu den vielen Varianten von MCU-Architekturen, Arm-basierten Mehrkernprozessoren, Dual-Cores bis hin zu Octa-Cores wird eMCOS auch in sogenannten Manycore-Prozessoren eingesetzt. Ein Beispiel dafür sind die neuesten Coolidge MPPA Chips (Massively Parallel Processor Array) von Kalray. Diese Chips besitzen 80 Kerne

und können sowohl in den neuen E/E-Architekturen als auch in den schnell wachsenden Bereichen des Edge-Computings und der KI eingesetzt werden: moderne Rechenzentren, Netzwerke (5G), Luft- und Raumfahrt, Gesundheitsausrüstung, Industry 4.0, Drohnen und Roboter.

Das eMCOS Betriebssystem lässt auf dem Coolidge Chip gleichzeitig AUTOSAR Classic Platform (CP) und AUTOSAR Adaptive Platform (AP) laufen (Bild 3), wodurch Anwendungen hohe Geschwindigkeit und Echtzeit-Determinismus mit geringem Stromverbrauch kombinieren können. Auf den äußeren I/O-Clustern laufen mehrere eMCOS AUTOSAR und AUTOSAR CP Instanzen, während auf den zentralen Rechenclustern eine Instanz von AUTOSAR AP auf eMCOS POSIX läuft. Die CP- und AP-Instanzen kommunizieren über die schnelle eMCOS IPC Schnittstelle.

In Zukunft werden neue Chiptechnologien die Hardware noch weiter in Richtung heterogener Manycore-Computer verändern. Aussichtsreiche Kandidaten sind hier 3D-Stapel-Chiplets. Hierbei werden mehrere Chips übereinandergestapelt und durch schnelle Durchkontaktierungen aus Silizium verbunden, um immer mehr Rechenleistung zur Verfügung zu stellen und dabei gleichzeitig die Kosten zu reduzieren.

Die Wahl der OS-Plattform ist entscheidend für die Zukunft

Konnektivität, autonomes Fahren, Mobilitätsdienste und Elektrofahrzeuge sind wichtige Automobilrends, die die Erwartungen an fahrzeuginterne Hochleistungsrechner und verbesserte E/E-Architekturen erhöhen. Gleichzeitig kommen zusätzliche Sicherheits Herausforderungen auf uns zu.

OEMs brauchen jede Hilfe, um mit diesen gestiegenen Anforderungen umgehen zu können und gleichzeitig noch eine akzeptable Markteinführungszeit für neue Features und Modelle zu erreichen. Die Software-Plattform ist ein entscheidendes Element, um alle diese Herausforderungen zu bewältigen. Skalierbarkeit ist unerlässlich und Software, die auf einer serviceorientierten Architektur (SOA) basiert, ist dabei der bewährteste Ansatz. Ein Multikernel OS (verteilt Mikrokernel OS) ist am besten geeignet, um zukünftige Hardwarearchitekturen auf Basis von Multicore-Prozessoren zu unterstützen und gleichzeitig eine SOA mit einer sehr schnellen Nachrichtenübermittlung zwischen den CPU-Kernen oder Steuergeräten zu unterstützen.

// SG

eSol

WIND RIVER STUDIO

Intelligente Netzwerksysteme

Wind River Studio ist eine Cloud-native Plattform für die Entwicklung, die Bereitstellung, den Betrieb und die Wartung von geschäftskritischen intelligenten Systemen. Die Software ermöglicht eine flexible und effiziente Bereitstellung, Verwaltung und den Betrieb von intelligenten Distributed 5G-Edge-Clouds über eine Managementkonsole („Single Pane of Glass“). Die Kubernetes- und Container-basierte Architektur sorgt für Bereitstellung und Management verteilter Edge-Netzwerken im großen

Maßstab. Die enthaltene Netzwerkanalysetechnologie erlaubt die Datenerfassung, Überwachung und Berichterstellung, um die Verwaltung und den Betrieb eines verteilten Cloud-Netzwerks zu optimieren. Zudem erlaubt die Core-to-Edge-Orchestrierung eine Automatisierung von verteilten Cloud-Netzwerken und Anwendungsdiensten, einschließlich der Verwaltung von Containern, Netzwerkelementen und Edge-Geräten.

Wind River

M-VISOR

Sichere Virtualisierung für MCUs

Der Embedded-Virtualisierungshypervisor μ -visor sorgt für die betriebs- und datensichere Konsolidierung mehrerer MCU-basierter Systeme in einem einzigen Multicore-RH850/U2A-Design. Dies senkt Kosten, Größe und Stromverbrauch und vereinfacht zugleich das Sicherheitsdesign des Systems. Durch die hardwarebasierten Virtualisierungsfunktionen des RH850/U2A von Renesas ermöglicht μ -visor den gleichzeitigen Betrieb mehrerer virtueller Maschinen. Es unterstützt eine Vielzahl

von Scheduling- und Core-Management-Strategien, um verschiedene Automotive-Anwendungsfälle zu bedienen, bei geringem Overhead. μ -visor unterstützt die kommenden Anforderungen an Cybersicherheit im Automotive-Bereich nach ISO/SAE 21434, UNECE WP.29 und für funktionale Sicherheit nach ISO 26262 ASIL-D. Die Lösung ergänzt die bewährte INTEGRITY-Multivisor-Virtualisierung für Anwendungsprozessoren.

Green Hills

PARASOFT C/C++TEST 2020.2

Schnellere Testautomatisierung

Parasoft C/C++test 2020.2 birgt zahlreiche Neuerungen zur Optimierung automatisierter Testabläufe. Beim Source Code Management soll die Integration mit Git künftig neue Code-Verletzungen in Sekunden oder Minuten statt Stunden melden. Dies soll die Feedback-Schleife für Entwickler drastisch verkürzen. Gleichzeitig soll der Abgleich von Code-Analyse und SCM-Baselines dafür sorgen, dass sich Entwickler auf relevante Verletzungen und Code-Änderungen konzentrieren können.

Auf Basis der Dashboard-Berichtslösung Parasoft DTP hat der Hersteller außerdem eine Erweiterung für Visual Studio Code erstellt. Diese soll es ermöglichen, statische Analyseergebnisse aus Remote-Server-Sitzungen innerhalb von Sekunden oder Minuten zu importieren. Zusätzliche Neuerungen betreffen die Möglichkeit von In-Datei-Unterdrückungen von Warnmeldungen sowie den Installationsprozess.

Parasoft

Du suchst einen
passenden
Job in der
Elektronik-
branche?



jobs.elektronikpraxis.de

Hier findest du ihn!

Der exklusive Stellenmarkt für alle, die in der Elektronikwelt oder angrenzenden Hightech-Branchen beruflich durchstarten wollen.

jobs.elektronikpraxis.de

ELEKTRONIK
PRAXIS

ist eine Marke der

VOGEL COMMUNICATIONS
GROUP

CLOUDANBINDUNG LEICHT GEMACHT

Kompakter Edge-PC als Gateway für das IIoT

spectra.de/
Edge-PC

Der Edge-PC Spectra PowerBox 100-IoT bringt die Daten von der Prozessebene in die Private oder Public Cloud.

Wir begleiten Sie auf Ihrem Weg in die industrielle Zukunft. Fragen Sie uns!

spectra
Industrie-PC & Automation

Corona kann den Drang zur Software-Weiterbildung nicht bremsen

Mit über 1100 Teilnehmern kann sich der coronabedingt erstmals virtuell ausgetragene ESE Kongress nicht vor den Rekordjahren der Präsenzveranstaltung verstecken.

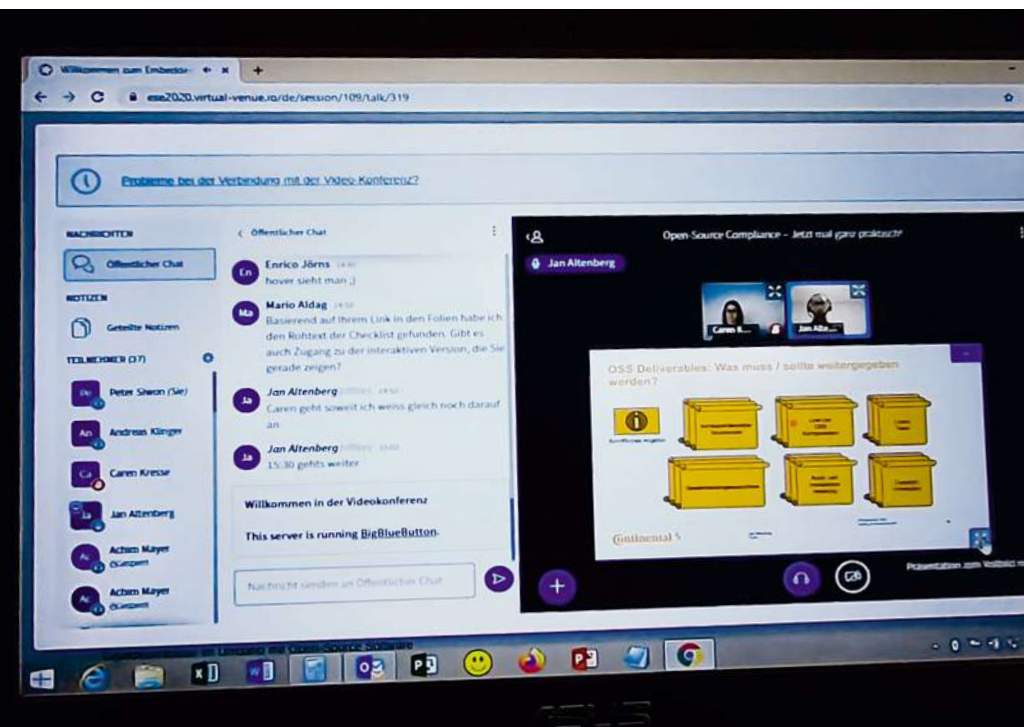


Bild: Peter Siwon

„Wir sind begeistert von der interessierten und aktiven Teilnahme über alle Tage des Kongresses hinweg. Die Embedded-Branche hat wieder bewiesen, dass sie schnell auf neue Rahmenbedingungen und unvorhergesehene Situationen reagieren kann – eine Eigenschaft, die auch in Embedded-Projekten täglich gefragt ist“, resümiert Ingo Pohle, Geschäftsführer des Mitveranstalters Microconsult.

„Der ESE Kongress 2020 hat sich wie ein Chamäleon dieser Herausforderung gestellt: Statt der beliebten und größten Embedded-Präsenzveranstaltung in Europa – dem Schmelztiegel aller Embedded-Themen – hat das ESE-Team die Migration zu einem digitalen Kongress erfolgreich gemeistert. Ein grandioser Abschluss für ein sonst schwieriges Jahr, der die Weichen für Projekte in 2021 stellt.“

ESE digital: Neue Konzepte und Erkenntnisse für die Zukunft

Das digitale Konzept, das die Konferenz in diesem Jahr aufgrund der aktuellen Pandemielage kurzfristig auf die Beine stellen musste, brauchte auch Neuerungen und Vorzüge mit sich, die Teilnehmer mit Freude auf- und wahrnahmen.

So wurde etwa die kostenlose Pre-Conference, die das digitale Format in diesem Jahr erstmals ermöglicht hat, sehr geschätzt. Und die Möglichkeit für Teilnehmer, nach dem Ende des ESE Kongress Aufzeichnungen sämtlicher Vorträge im Video-on-Demand-Archiv abrufen zu können, wurde mit Begeisterung aufgenommen. Hier sind einige positive Erfahrungen gewonnen worden, die sicher auch in künftigen Live-Kongressen eine Rolle spielen werden.

Dennoch hoffen die Veranstalter, dass 2021 der ESE Kongress wieder, so Corona es zulässt, vor Ort und in Person stattfinden kann – womöglich angereichert durch ein hybrides Konzept mit den Learnings aus der digitalen Veranstaltung 2020. // SG

ELEKTRONIKPRAXIS

Virtual Venue: ESE-Kongress-erfahrene Referenten vermittelten auch in der digitalen Variante der Konferenz den Teilnehmern handfestes Praxiswissen.

Auch Deutschlands größter Fachkongress rund um das Thema Embedded Software Engineering musste aufgrund der andauernden Corona-Pandemie im Jahr 2020 auf eine rein digital ausgetragene Veranstaltung ausweichen. Das tat dem Andrang allerdings kaum Abbruch: Insgesamt 1112 Teilnehmer fanden sich auf der virtuellen Kongressplattform des ESE Kongress 2020 ein, um das umfangreiche Konferenzprogramm mit 15 Kompaktseminaren, 96 Vorträgen, 3 Keynotes, Live-Moderation und -Diskussion sowie Direct Chats mit Referenten, Branchenkollegen, Eventpartnern und Sponsoren zu nutzen.

Damit konnte sich der ESE Kongress auch im digitalen Format nicht nur vom Umfang, sondern auch den Teilnehmerzahlen her mit den Vorjahren messen: 2019 und 2018 hatten die Präsenzveranstaltungen in Sindelfingen

die Rekordmarke von 1200 Teilnehmern durchbrochen. Trotz andauernder Corona-Krise bleibt die Embedded-Software-Branche weiter stark, der Bedarf nach Austausch zu Trends, Tipps und Tricks ungebremst.

Schnelles Reagieren auf ungewohnte Veränderungen

Eine digitale Konferenz kann natürlich die körperliche, familiäre Nähe, die die ESE Kongresse in den letzten 12 Jahren ausgezeichnet hatten, nicht ersetzen. Dennoch lag es dem Veranstalterteam am Herzen, durch Studio-Atmosphäre und Live-Moderation den Geist der physischen Veranstaltung auch im digitalen zu erhalten. Die Teilnehmer beteiligten sich engagiert an den Seminarangeboten, verfolgten aufmerksam die Vorträge und nutzten auch die Gelegenheit, nachzufragen und sich auszutauschen.

Der Embedded Software Engineering Kongress 2021: 29. November bis 3. Dezember Sehen wir uns?

13808

Rückblick 2020 und Informationen auf www.ese-kongress.de

ESE Kongress – Ideen entwickeln, Profis treffen, Lösungen finden.

Der Embedded Software Engineering Kongress ist die größte deutschsprachige Veranstaltung, die sich ausschließlich der Entwicklung von Geräte-, Steuerungs- und Systemsoftware für Industrie, Kfz, Telekom sowie Consumer- und Medizintechnik widmet. Vom 29. November bis 3. Dezember 2021 trifft sich die Embedded-Software-Branche bereits zum 14. Mal – wenn nicht in Sindelfingen, dann digital im Netz. Wir freuen uns auf Sie!

Der Call for Papers für Referenten startet Mitte März.

Danke an alle Eventpartner 2020:

Arm, Axivion, BlackBerry QNX, ELEKTRONIKPRAXIS, embeff, emmtrix Technologies, froglogic, GitHub, Green Hills Software, Hitex, IAR Systems, Kernkonzept, LDRA, LieberLieber Software, Logic Technology, macio, MathWorks, MicroConsult, oose Innovative Informatik, Parasoft, Pengutronix, PLS Programmierbare Logik & Systeme, PROTOS, Razorcat Development, RTI Real-Time Innovations, Sodus Willert, Tasking, Vector Informatik, Xilinx



Embedded Software Engineering Kongress

2021

29.11. bis 3.12.2021 in Sindelfingen

Hauptsponsoren 2020

axivion
stopping software erosion



Veranstalter

**ELEKTRONIK
PRAXIS**

 **MICROCONSULT**

Asset-Management: Funktionsweise von LoRa Edge

Per Software steuerbare Funk-Frontends vereinfachen das Auswerten von WLAN- und Satellitensignalen – und ermöglichen die energieeffiziente Standortbestimmung von Gütern.

PEDRO PACHUCA *



Bild: Semtech
Band bis zu den Satellitensignalen von BeiDou und GPS. Mit SDR ist es möglich, nur bestimmte Teile des eingehenden Signals zu verarbeiten – das spart Rechenressourcen. Die lassen sich nutzen, um mehr Informationen aus einem verrauschten Signal zu gewinnen – was aber häufig die Stromaufnahme erhöht. Ein SDR kann sein Energieverbrauch dynamisch so abstimmen, dass die Batterie möglichst lange hält.

WLAN- oder Satellitensignal? Der Standort entscheidet

Situationsabhängige Signalverarbeitung ist die zentrale Anforderung an ein Geolocation-Tag, das im Innen- und Außenbereich eingesetzt werden soll. Soll es seinen Standort ermitteln, muss das Tag zuerst feststellen, welche Ortungstechnik die besten Ergebnisse liefert. Befindet sich das Tag im Freien und besteht eine gute Sichtlinie zu Satelliten, sollte es die GNSS-Signale leicht erkennen. LoRa Edge nutzt dafür einen stromsparenden Scanmodus, den ein externer Controller aktiviert – möglicherweise nach einer Zeit der Ruhephase. Bis zu 0,65 s lang verarbeitet die LoRa-Edge-Firmware die auf GNSS-Frequenzbändern erwarteten Signale. Nur wenn eins mit einem Signal-Rausch-Verhältnis von mehr als -134 dB erkannt wird, versucht der Empfänger eine weitere Verarbeitung. Gelingt es, in dieser Zeit ein ausreichend starkes Signal zu finden, ändert die Empfänger-Firmware ihre Verarbeitung in Algorithmen mit höherer Empfindlichkeit. Es versucht dann bis zu acht Satelliten mit einer Signalstärke von mehr als -141 dB zu finden. Die geringere Empfindlichkeit während der anfänglichen Suche begründet sich darin, dass die Wahrscheinlichkeit, genügend Satelliten zu finden, um einen Fix zu erhalten, äußerst gering ist, wenn das Signal der stärksten in Betracht kommenden Quelle unterhalb dieses anfänglichen Grenzwerts liegt.

Sind genügend Satelliten in Sicht, erhält der Empfänger ausreichend Daten, um eine

Spezialisiert: LoRa Edge ermöglicht energieeffizientes Güter-Tracking auf Basis des LoRaWAN-Protokolls

Anwendungen wie das Echtzeit-Asset-Management beruhen auf der Nachverfolgung elektronischer Etiketten (Tags), die mit einer einzigen Batterieladung Monate, wenn nicht Jahre in Betrieb sein können. Low-Voltage-Designs und hocheffiziente Prozessor-Cores machen dies möglich. Oft ist eine ganzheitliche Bewertung des Stromverbrauchs einer Anwendung erforderlich. LoRa Edge bietet diesen ganzheitlichen Entwicklungsansatz. Die Technik basiert auf direkter Demodulation und Cloud-Anbindung und will so eine hohe Energieeffizienz, niedrige Kosten und einfache Handhabung garantieren. Ein häufiges Problem bei der Standorterfassung auf Basis herkömmlicher

Techniken ist, dass mehrere Funk-Frontends auf Leiterplattebene erforderlich sind. Der Grund: Der Integrator muss – sofern er nicht sicher ist, dass die Tags nur in einer streng kontrollierten Umgebung verwendet werden – Positionsinformationen von mehreren Quellen abrufen. Während GNSS-Dienste gut im Freien funktionieren, erfolgt die Lokalisierung in Innenräumen oft per Positionstriangulation über WLAN-Signale. Eine weitere HF-Schnittstelle übernimmt die stromsparende Funkkommunikation.

LoRa Edge geht einen anderen Weg: Es verwendet Software-programmierbare Funktechnik (SDR; Software-Defined Radio), um drei Frontends in einer Einheit zu integrieren. Signale von drei Antennen werden über rauscharme Verstärker in einen A/D-Wandler (ADC) geleitet, der direkt einen digitalen Demodulator speist. Dieser Aufbau ermöglicht es, mehrere Signale zu verarbeiten – von der LoRa-Kommunikation im Sub-1GHz-



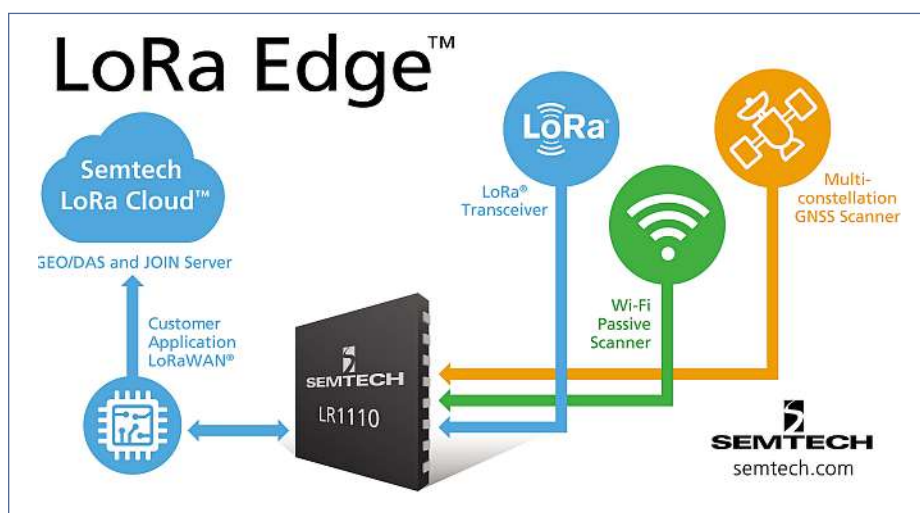
* Pedro Pachuca
... ist Director Wireless Products von Semtech.

genaue Positionsbestimmung innerhalb von 1,65 s zu erzielen. Nach der Signalerfassung kann der Empfänger die Verarbeitung beenden. Herkömmliche GNSS-Empfänger empfangen und verarbeiten Signale weiterhin direkt, was viel Energie verbraucht. LoRa Edge versucht nicht, die empfangenen Satellitendaten lokal zu verarbeiten. Stattdessen werden die Datenelemente zu einer Nachricht zusammengefasst und an einen Cloud-Dienst übertragen. Mit der dort verfügbaren Rechenkapazität werden die empfangene Satellitennachrichten in eine genaue Geolokalisierung umgewandelt. Steht GNSS nicht zur Verfügung oder stellt die Host-Controller-Firmware fest, dass eine andere Art von Korrektur erforderlich ist, kann LoRa Edge auf die Dekodierung der Signale von der 2,4-GHz-Antenne umschalten. Wie bei der GNSS-Implementierung versucht die HF-Engine nicht, die Daten vollständig zu dekodieren und zu verarbeiten. Sie konzentriert sich nur auf die Elemente, die der Remote-Cloud-Dienst benötigt, um eine genaue Geolokalisierung zu bestimmen. Dabei nutzt sie die Struktur des WLAN-Protokolls.

Die HF-Engine muss keine Daten an WLAN-Router übertragen, sondern nur die Umgebung scannen. Im WLAN-Scanmodus sucht der Empfänger 802.11b/g/n-Signale im 2,4-GHz-Band. Die Empfänger-Firmware kann geeignete Pakete auswählen, indem sie auf die von WLAN-Routern verwendete Header wartet, bevor diese Nutzdaten übertragen. Sobald die ersten Bytes empfangen werden, demoduliert die Firmware das Signal und erfasst Bytes, bis sie über eine vollständige MAC-Adresse des Zugriffspunkts verfügt. Danach beendet sie den Scanvorgang, speichert die Adresse und den zugehörigen Signalstärkewert und schaltet das HF-Frontend ab, um Strom zu sparen.

Meist muss der Host die MAC-Adressen mehrerer Zugriffspunkte erfassen, um eine genaue Standortbestimmung per WLAN zu erhalten. Dazu kann der Host-Controller den passiven Scan-Modus mehrmals aktivieren, bis er über genug Informationen verfügt. Um keinen Strom in Bereichen mit schlechtem WLAN-Zugang zu verschwenden, kann die HF-Engine einen Timeout-Modus nutzen und den Empfänger automatisch deaktivieren, wenn kein gültiges Paket übertragen wurde – bis der Host-Controller beschließt, es erneut zu versuchen. Dieser kann warten oder in den GNSS-Modus wechseln, falls er sich in großer Entfernung zu WLAN-Routern und an einem Außenstandort befindet.

Sobald der Host über eine Liste mit MAC-Adressen und Signalstärken verfügt, kann der Host die Daten wie bei den GNSS-Daten



3-fach-Funk: Das Software-konfigurierbare Funk-Frontend des Semtech-Chips kann Sub-1GHz-, Satelliten- und WLAN-Signale verarbeiten.

zur Umwandlung in eine Geolokalisierung an die Cloud weiterleiten. Das Ergebnis ist ein geolokalisierbares Tag, das weniger Strom verbraucht als herkömmliche Designs und die Batterielebensdauer von einigen Monaten auf zwei bis drei Jahre verlängert.

Dank SDR kann ein LoRa-Edge-Gerät seine Empfangsantennen auch zum Senden verwenden. Hat der Empfänger die Erfassung der GNSS-Daten abgeschlossen, kann der Host-Controller die HF-Engine in den Funkmodus schalten. Nach dem Senden kann die HF-Engine wieder in den Empfangsmodus übergehen oder in einem Standby-Modus mit geringem Stromverbrauch verbleiben, um bis zu einem geplanten Zeitpunkt auf den Empfang von Anweisungen oder auf eine Antwort von einem Remote-Server zu warten.

Ende-zu-Ende-Verschlüsselung der Anwendungsdaten

Der Integrator beziehungsweise Dienstleister entscheidet, wohin Datenpakete gesendet werden. LoRa Edge greift dazu auf die Sicherheitsmerkmale des LoRaWAN-Protokolls zurück. Integrierte Sicherheit ist ein entscheidendes Merkmal von LoRaWAN. Im Gegensatz zu den meisten anderen für das IoT genutzten Protokollen implementiert LoRaWAN eine Ende-zu-Ende-Verschlüsselung für Anwendungsdaten. Dies erfolgt zusätzlich zu einer Verschlüsselungsschicht auf Netzwerkebene, die den Zugriff nicht autorisierter Knoten verhindert.

Die Inbetriebnahme umfasst eine Anfrage an einen Join-Server, der Authentifizierungsroutinen ausführt und die Anmeldedaten des Geräts bzw. Systems mithilfe von AES-Standardprotokollen überprüft. Nach dieser Authentifizierung erstellen der Join-Server

und das System gemeinsam Sitzungsschlüssel zum Schutz der Netzwerknachrichten. Die Systeme können dann ähnliche Verfahren nutzen, um sich bei den eigenen Anwendungsservern des Nutzers zu authentifizieren. Auf diese Weise müssen Anwendungen und der Netzbetreiber keine Schlüssel gemeinsam nutzen.

Die Unterscheidung zwischen Netzwerk- und Anwendungsdiensten ist für die Cloud-Geolocation-Dienste ebenso wichtig wie für andere Anwendungen. Das Design von LoRa Cloud und LoRa Edge stellt sicher, dass standortbezogene Anfragen vom eigenen Anwendungsserver eines Kunden gestellt werden, anstatt dass das Gerät selbst die Anfrage auf Netzwerkebene stellt. Soll eine Geolokalisierung an ein Tag zurückgemeldet werden, wird dies auf Anwendungsebene durch das System des Nutzers gehandhabt. In vielen Fällen müssen die Daten jedoch nicht im Gerät selbst gespeichert werden – sie können in der Cloud gespeichert und nur bei Bedarf verteilt werden.

Gleichzeitig bietet LoRa Edge einen praktischen Mechanismus zum Speichern der für den Netzwerk- und Anwendungszugriff erforderlichen Schlüssel. Ein geschützter Speicherbereich wird mit den Schlüsseldaten programmiert, die für den Beitritt zu einem LoRaWAN-Netzwerk beim Start verwendet werden. Dieser Speicherbereich unterstützt auch die Möglichkeit, benutzerdefinierte Schlüssel für Nutzeranwendungen zu speichern. Die On-Chip-Logik führt alle sicheren und verschlüsselungsbezogenen Vorgänge aus, die für den Zugriff auf LoRaWAN-Funktionen erforderlich sind. // ME

Semtech

Bild: Semtech

Per NFC: Drahtlose Ladelösung für kleine Consumer-Geräte

Der bewährte Qi-Standard für drahtloses, induktives Laden ist für kleine Wearables wie Hörgeräte, Fitness-Tracker usw. wenig geeignet. Lösungen auf Basis des „Near Field Communication“-Standards schon.

ALESSANDRO GOITRE *

Die Integration drahtloser Ladevorrichtungen in Consumer-Geräte ist durch den Qi-Ladestandard des Wireless Power Consortium (WPC) beflügelt worden. Qi ist aktuell die verbreitetste Technik zum Laden von Smartphones. Diese Entwicklung beschleunigt sich mit der zunehmenden Verfügbarkeit von Lademöglichkeiten etwa in Mittelkonsolen im Auto, Haushaltsgeräten oder im Handel. Das führt zu steigender Bekanntheit und Akzeptanz bei den Verbrauchern, und führende Hersteller anderer Ar-

ten von Geräten, über Mobiltelefone hinaus, sehen das Potenzial für drahtloses Laden. Die Hersteller von Consumer-Geräten sind nun bereit, auch alternative Techniken zum drahtlosen Laden bei kleinen Formfaktoren in Erwägung zu ziehen. Die großen Vorteile beim Ersatz kabelgebundener Ladegeräte durch das drahtlose Laden sind:

- **Höhere Zuverlässigkeit:** Jede Kabelverbindung ist eine potenzielle Störungsquelle. Miniatursteckverbinder sind besonders empfindlich für Schäden durch mechanische Beanspruchung, z.B. durch Biegen.

- **Mehr Gestaltungsmöglichkeiten:** Consumer-Geräte und Wearables haben relativ kleine Oberflächen. Der Entwickler kann

die Form des Geräts und die Nutzung seiner Oberfläche optimieren, wenn er keinen Steckverbinder unterbringen muss.

- **Höhere IP-Schutzart:** Durch den Wegfall mechanischer Verbindungen entfallen Stellen, an denen Wasser, Schweiß und Staub eindringen können.

Aus verschiedenen Gründen ist die Qi-Technik zum Laden sehr kleiner Wearables, z.B. Armbänder von Aktivitätstrackern, drahtlose Ohrhörer oder intelligente Brillen, nur wenig geeignet. Eine vielversprechendere Alternative ist NFC – die Technik, die hinter kontaktlosen Bezahl- und Ticket-Systemen steht. Der Grund: Neben der Datenkommunikation ist NFC für Energy Harvesting ausgelegt, also die Fähigkeit des Empfängers, aus dem Sendesignal eines Lesers Energie zu gewinnen. Das heißt, dass jedes NFC-fähige Gerät mit der Möglichkeit zum Energy Harvesting ohne zusätzliche Antenne und weitere Komponenten drahtlos geladen werden kann. Außerdem erfordert das drahtlose Laden per NFC, im Gegensatz zu anderen Ladetechniken, keine perfekte Ausrichtung der Antennen von Ladegerät und Empfänger. Ein NFC-Ladegerät arbeitet mit hoher Effizienz auf der Ladeseite selbst dann, wenn die beiden Antennen um die halbe Antennengröße gegeneinander versetzt sind.

NFC, eine sinnvolle Ergänzung zum drahtlosen Laden per Qi

Trotz der in fast allen Smartphones verfügbaren NFC-Funktionen – die als Ladestation für andere, kleinere Geräte fungieren könnten – ist das Laden per NFC nur in wenigen kleinen Geräten implementiert worden. Der Grund dafür ist, dass die Sendeleistung von gewöhnlich 1 bis 1,5 W an der Antenne die Möglichkeiten zum Laden von Geräten mit größeren Akkus oder kleineren NFC-Antennen beschränkt. Nun verspricht jedoch ein Durchbruch bei der Entwicklung von NFC-Systemen eine Verdoppelung der Leistung, die über eine NFC-Verbindung übertragen

* Alessandro Goitre
... ist Produktmarketingdirektor von Panthronics in Graz.



Bild: Panthronics

Bild 1: Für das drahtlose Laden kleiner Elektronikgeräte ist der Qi-Standard überdimensioniert. NFC-Schaltungen bieten sich als kompakte Lösungen an.

Bild: Panthronics

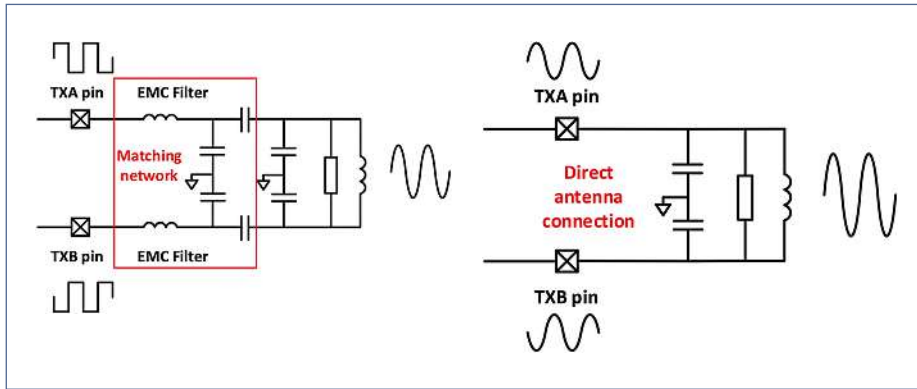


Bild 2: Anders als herkömmliche NFC-Sender zum drahtlosen Laden benötigt der PTX100W kein EMV-Filter und nur wenige Bauteile zur Antennenanpassung.

werden kann, bei gleichzeitiger Verringerung der Bauteileanzahl, der Materialkosten und der Systemgröße. Dieser Artikel beschreibt, wie sich diese NFC-Fähigkeit zum schnellen Laden auf die Entwicklung neuer Consumer-Produkte auswirken könnte.

Das drahtlose Laden per Qi hat sich als erfolgreiche Technik für Geräte wie Smartphones erwiesen. Deren große Akkus – üblicherweise 3.000 mAh oder mehr – erfordern die von Qi gebotene Ladeleistung von über 15 W. Das WPC unterstützt die Bemühungen, die Leistung von Qi zu erhöhen, um noch größere Geräte wie Tablets und Notebooks drahtlos laden zu können. Doch die Faktoren, die die Qi-Technik für den Einsatz bei diesen Anwendungen mit hoher Leistung qualifizieren – z.B. die Notwendigkeit einer großen Antenne im Empfänger – verhindern ihren Einsatz in kleineren Produkten wie Wearables. Wenn ein zu ladendes Gerät nur 1 W statt 15 W benötigt, ist ein Qi-System im Allgemeinen zu sperrig und zu teuer. Außerdem sind seine Fähigkeiten zur Datenkommunikation beschränkt. Im Gegensatz hierzu ist NFC geeignet zum drahtlosen Laden von Geräten, bei denen Flexibilität ein wesentlicher Faktor ist.

■ Es gibt bereits eine Basis von Milliarden NFC-fähiger Geräte bei den Verbrauchern, sodass sowohl die Verbraucher als auch die Entwicklungsingenieure mit dieser Technik vertraut sind. Die meisten Smartphones könnten als Ladegeräte für NFC-Geräte funktionieren.

■ Die NFC-Technik unterstützt eine bidirektionale Kommunikation. Je nach dem verwendeten Kommunikationsprotokoll und den Fähigkeiten des Geräts auf der Gegenseite unterstützt NFC Kommunikationsgeschwindigkeiten bis 848 kBits/s. Wenn NFC ohnehin als Technik für die Kommunikation des Geräts eingesetzt wird, kann drahtloses Laden das Funktionsangebot

des Geräts und damit seine Attraktivität für den Verbraucher erhöhen, ohne dass zusätzliche Materialkosten anfallen.

■ NFC lässt eine Fehlausrichtung der Antennen von Sender und Empfänger zu, ohne die Effizienz der Leistungsübertragung zu beeinträchtigen.

■ NFC ermöglicht kompakte Lösungen mit kleinerer Antenne auf der Sender- und Empfängerseite.

■ NFC lässt sich einfach implementieren, denn die Ladefähigkeit ist in den Standardprotokollen des NFC-Forums bereits enthalten. Beim drahtlosen Laden per NFC kann die Feldstärke des HF-Felds erhöht werden, um die Leistungsübertragung zwischen zwei kompatiblen Geräten zu maximieren.

Bislang nur wenige NFC-Ladelösungen – warum?

NFC wird bisher zum drahtlosen Laden wenig eingesetzt. Das ist hauptsächlich deshalb so, weil NFC-Sender auf einer konventionellen Architektur aufbauen, die die Ausgangsleistung an der Sendeantenne auf < 1,5 W begrenzt. In dieser Architektur erzeugt der NFC-Sender ein rechteckförmiges Ausgangssignal (siehe Bild 2). Diese Rechteck-Architektur ist für die Hersteller von NFC-Komponenten eine beliebte Wahl, denn sie lässt sich schaltungsseitig leicht umsetzen. Doch das rechteckförmige Ausgangssignal muss zur Übertragung über die Antenne in ein Sinussignal umgewandelt werden, um zu verhindern, dass die elektromagnetischen Emissionen die Grenzwerte überschreiten. Hierzu ist ein EMV-Filter aus mehreren externen Komponenten erforderlich.

Beim NFC-Laden hat dies zwei erhebliche Nachteile: Hohe Leistungsverluste im EMV-Filter des Senders reduzieren die Ausgangsleistung und damit die Eingangsleistung beim Empfänger. Die höhere Impedanz der Antenne wegen der Toleranzen der diskreten

SAVE
THE
DATE

Entwicklerkonferenz für Intelligent Edge und KI

21. – 22. September 2021
Sindelfingen

Get Ready for the Next Decade

Experten behandeln wichtige Aspekte,
erklären Methoden, Werkzeuge und
Best Practices.

www.intelligent-edge.de

Eine Veranstaltung von ELEKTRONIK PRAXIS

einer Marke der VOGEL COMMUNICATIONS GROUP

Bild 3:
Musterschaltung
mit dem NFC-Sender
PTX100W zum Laden
einer Smartwatch.

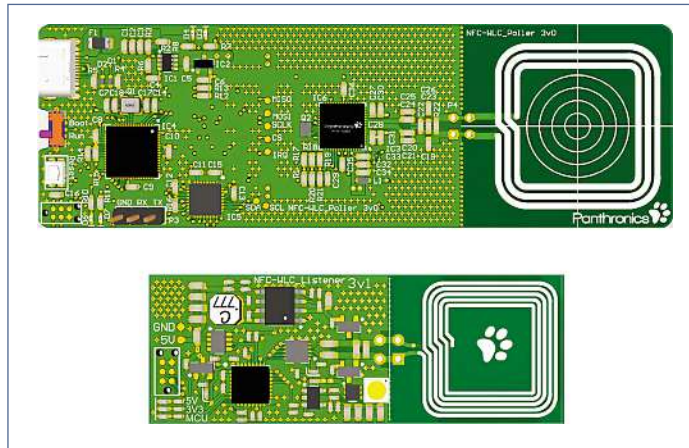


Bild: Panthronics

nics-Architektur mehr Leistung als eine herkömmliche NFC-Architektur liefert, ist die niedrigere Impedanz bei der Antennenanpassung, die durch eine Schaltung ohne EMV-Filter möglich wird. Eine Anwendung mit PTX100W kann mit einer Antennenanpassimpedanz von unter 5Ω ausgelegt werden, während die Impedanz mit EMV-Filter mindestens doppelt so hoch ist.

Sehr wichtig ist auch, dass der Wegfall des EMV-Filters und weiterer externer Bauteile sonst benötigten Platz auf der Leiterplatte spart. Damit verringern sich Kosten, Größe und Komplexität des Boards im geladenen Gerät, was für die Hersteller von Wearables ein wichtiger Vorteil ist. Die folgenden Anwendungsbeispiele demonstrieren die Lade-fähigkeit des PTX100W.

Ladeschaltung in einer Smartwatch: Das Sender-Board enthält den PTX100W. Die Abmessungen der Antennen von Sender und Empfänger erlauben die Integration in eine Smartwatch oder einen Fitness-Tracker (siehe Bild 3). Unter diesen Bedingungen überträgt der PTX100W bis zu 1 W an den Akku des Empfängers, gemessen mit dem Leistungssensor auf der Empfängerplatine. Eine Fehlausrichtung zwischen den Antennen von Sender und Empfänger hat nur geringe Auswirkungen auf die übertragene Gesamt-leistung.

NFC-Laden von Earbuds: Bei dieser Anwendung sind die Abmessungen der Antenne noch geringer, was die übertragene Leistung weiter reduziert. Bild 4 zeigt die mit dem PTX100W als Sender beim Akku der Earbuds ankommende Leistung. Selbst bei einem Antennenabstand bis zu 5 mm liegt die zur Antenne der Earbuds übertragene Leistung immer noch über der maximalen Leistung zum Laden eines typischen Earbud-Akkus.

Höhere Sendeleistung erlaubt schnelleres Laden per NFC

Mit der Einführung des neuen NFC-Lade-Controllers PTX100W bietet Panthronics eine neue Lösung mit hoher Leistung durch die innovative Architektur mit sinusförmigem Signal. Während die Hersteller von Consumer-Geräten das drahtlose Laden per NFC wegen der langen Ladezeiten selbst bei kleinen Akkus nur langsam akzeptiert haben, überwindet die Einführung des PTX100W dieses Problem, denn sie ermöglicht eine Leistung an der Antenne von bis zu 2,5 W und verkürzt die typische Ladezeit bei Akkus bis 400 mAh auf die Zielvorgaben der Hersteller von unter drei Stunden.

// ME

Panthronics

Bauteile im EMV-Filter beschränkt die mögliche Ausgangsleistung des NFC-Senders. In der Praxis heißt das, dass die in heutigen drahtlosen Ladegeräten eingesetzten NFC-Sender an der Antenne auf eine Ausgangsleistung von bestenfalls 1,5 W und maximal 500 mW beim Empfänger beschränkt sind. Bei diesen niedrigen Leistungen wird die Zeit zum Laden eines Geräts unverhältnismäßig lang. Glücklicherweise gibt es einen anderen Weg, dessen Vorzüge das Ergebnis einer völlig anderen Architektur im Sender und im Empfänger sind.

Neue Architektur mit sinusförmigem Ausgangssignal

Die neue Architektur – das Ergebnis einer von Panthronics entwickelten patentierten Umsetzung in Silizium – erzeugt ein sinusförmiges Signal am Senderausgang (siehe Bild 2). Damit benötigt der NFC-Kreis kein verlustbehaftetes EMV-Filter und die Antenne kann direkt am Senderausgang angeschlossen werden (DIRAC). Die wichtigsten Vorteile dieser Architektur sind die genaue

Umkehrung der Nachteile herkömmlicher NFC-Controller:

- Die Verluste werden reduziert, da das EMV-Filter und die meisten Bauteile zur Anpassung wegfallen.
- Durch die Sinusform des Signals entfallen verschiedene Kondensatoren und Induktivitäten mit großen Toleranzen und es kann eine Antennenanpassschaltung mit viel niedrigerer Impedanz verwendet werden, wodurch sich die Ausgangsleistung des Senders erhöht.
- Die Verringerung der Bauteileanzahl vereinfacht außerdem die Systemanpassung und beseitigt Produktstreuungen durch die großen Toleranzen der Bauteile für die Anpassung.
- Gleichmäßige übertragene Leistung als Funktion des Volumens
- Konstant optimierte Systemanpassung als Funktion der Verschiebung

Beim Betrieb mit 5 V liefert eine Panthronics-Lösung mit dem NFC-Lade-Controller PTX100W bis zu 2,5 W an der Antenne. Einer der Hauptgründe dafür, dass die Panthro-

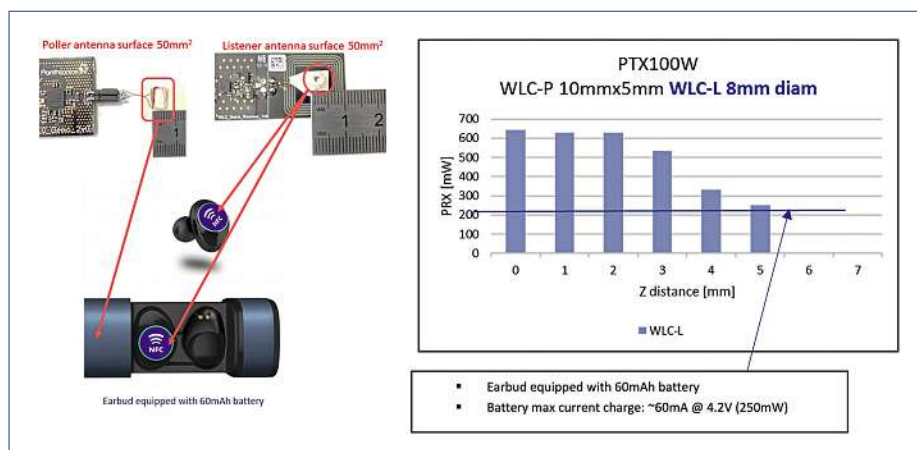


Bild: Panthronics

Bild 4: Musterschaltung zum NFC-Laden eines Earbud-Akkus. Die Grafik zeigt die Abnahme der übertragenen Leistung bei einer Fehlausrichtung der Antennen von Sender und Empfänger.

JEDEM DRITTEN EMS-UNTERNEHMEN DROHT DAS AUS

Call for Paper zum 19. Würzburger EMS-Tag am 10. Juni 2021

EMS-Führungskräfte, Manager von OEMs und Inhouse-Fertigern sowie Experten der Zulieferindustrie treffen sich seit vielen Jahren auf dem Würzburger EMS-Tag. Der praxisorientierte Managementevent besteht aus einem Seminartag und einer Abendveranstaltung am Vortag zum entspannten Meinungsaustausch unter Referenten und Teilnehmern. Der EMS-Tag gilt als eine der wichtigsten Veranstaltungen der Electronics-Manufacturing-Service-Branche. Geschäftsführer und Führungskräfte von EMS-Providern, OEMs und Inhouse-Fertigern sowie deren Zulieferern treffen sich seit vielen Jahren mit unabhängigen Experten, um über aktuelle Themen und Problemstellungen der Branchen zu diskutieren.

Wenn Sie einen Vortrag halten möchten, dann reichen Sie bitte Ihren Themenvorschlag sowie einen Vortrags-Abstract bis zum 15. März 2021 ein unter www.ems-tag.de. Der Programm-ausschuss wird die Agenda zum 19. Würzburger EMS-Tag zusammenstellen und Ihnen eine Rückmeldung geben, ob Ihr Vortrag in das Programm aufgenommen werden kann.

Bild: VCG



EMS-Tag 2020: Trotz oder gerade wegen Corona kamen im September mehr Teilnehmer denn je nach Würzburg bzw. ins Ausweichquartier Mainfrankensäle Veitshöchheim.

Weitere Details können Sie auch gerne direkt telefonisch mit dem Chefredakteur von ELEKTRONIKPRAXIS Johann Wiesböck besprechen: Tel. +49931-4183081 oder per E-Mail: johann.wiesboeck@vogel.de. Themenwünsche, Referententipp und Diskussionsbeiträge ehemaliger und künftiger Teilnehmer sind ebenfalls herzlich willkommen – Einsendeschluss 15. März 2021. Das Programm von 2020 – zur Orientierung – und alle weiteren

Infos zu dieser Managementkonferenz finden Sie unter www.ems-tag.de.

Rückblick: „Jedem dritten EMS-Unternehmen droht in den nächsten Jahren das Aus.“ Dies war eine der heiß diskutierten Thesen auf dem EMS-Tag 2020: Der Markt für Electronic Manufacturing Services habe zwar ein großes Potenzial – nur nicht für jeden. Etwa ein Drittel der EMS-Firmen könnte in den nächsten Jahren vom Markt verschwinden.

Branchenkenner, Marktakteure und Spezialisten zeigten beim 18. Würzburger EMS-Tag, wie sich die EMS-Unternehmen verändern und aufstellen müssen, um im Wettbewerb Schritt zu halten.

„Wer glaubt, er brauche nur die Coronakrise auszusitzen, hat nichts verstanden“, warnte Branchenkenner Dieter Weiss, in4ma. Seit drei Jahren wertet die in4ma-Marktforschung die betriebswirtschaftlichen Daten aller EMS-Firmen in Europa akribisch aus. Das Ergebnis: Der 2019 eingeläutete Paradigmenwechsel wird den EMS-Markt verändern und sich nun noch beschleunigen.

Drei Indizien sprechen dafür. Erstens: Die Fertigungstiefe der EMS-Firmen erhöht sich und damit das erforderliche Investitionsvolumen. Zweitens: Die großen EMS-Firmen wachsen schneller als die kleinen und erhöhen ihren Marktanteil. Drittens: Der Preisdruck nimmt zu und die Gewinnmargen sinken konstant. Als Faustformel gilt: Um wettbewerbsfähig zu sein, darf die Summe aus Material- und Personalkosten nicht größer sein als 90%. // JW

www.ems-tag.de

mes THE CONNECTOR

**MANCHE VERBINDUNGEN
SIND EINFACH DICHTER,
ALS SIE SICH
VORSTELLEN KÖNNEN.**

+ zum Beispiel der IP68-Rundsteckverbinder SP13 von Weipu. www.mes-electronic.de

Leistungssteckverbinder für 60 A im Raster 8,5 mm

Kona-Steckverbinder wurden für raue Betriebsbedingungen entwickelt sowie auf eine lange Lebensdauer sowie eine einfache Handhabung getrimmt. Was bieten die kleinen Stecker?



Leistungssteckverbinder:

Harwins leistungsfähigster Steckverbinder ist die Kona-Reihe mit einer Strombelastbarkeit pro Kontakt von 60 A und einer Länge von 7 mm.

Bild: Harwin

Die Verriegelung mit dem „Mate Before Lock“-Mechanismus (Zusammenstecken eines Steckerpaares und Verriegelung mit Rändelschrauben, wobei die Verriegelungssequenz erst nach dem vollständigen Zusammenstecken der Stecker gestartet wird) verhindert Schäden beim Stecken. Ein Kontaktindikator stellt sicher, dass die Komponenten korrekt ausgerichtet sind. Durch die Ummanthelung der einzelnen Kontakte und einen Verpolschutz ist falsches Stecken nicht möglich.

Die Befestigung erfolgt werkzeuglos über Rändelschrauben aus Edelstahl, die an der Buchse angebracht sind. Die Steckverbinder lassen sich damit schnell und einfach stecken.

Mögliche Anwendungen der kleinen Leistungsstecker

Die Leistungssteckverbinder sind für Anwendungen wie das Laden von Batterien konzipiert, bei denen der Strom nicht auf mehrere Kontakte aufgeteilt werden soll. Zu den Hauptanwendungen zählen laut Hersteller die Überwachung und das Management von Batterien bzw. Akkus in Elektrofahrzeugen, die Leistungssteuerung in Elektrofahrzeugen, Roboterantriebe, Servosteuerungen, UAV (unbemannte Luftfahrzeuge) und Satelliten.

Unser Fazit: Klein und oho. Harwin hat (auf Kundennachfrage) seine High-Rel-Steckverbinder um eine Variante mit einer Stromtragfähigkeit von 60 A pro Kontakt erweitert, die einige Marktpotenziale erschließen wird. Die kleinen Leistungssteckverbinder zeichnen sich insbesondere durch ihre hohe Vibrationsfestigkeit in dieser Klasse aus. Durchdacht sind ebenfalls die Kontaktgeometrie und die Verriegelung, bewährt aus den Vorgänger-Versionen.

Der Name Kona ist übrigens aus dem Hawaianischen entlehnt. Dort steht er für das Wort „Dame“ und zeigt eine edle Herkunft, Beliebtheit und Vorherrschaft an. // KR

Harwin

Harwin hat mit der Steckverbinderfamilie „Kona“ seine hochzuverlässigen Steckverbinder im Raster 8,5 mm um eine Variante für einen Dauerstrom von 60 A erweitert. Das Steckerpaar besteht aus einem vertikalen Durchgangssteckverbinder für die Leiterplatte sowie einer Buchse, die Kabel mit einem Durchmesser von 8 AWG (8,34 mm²) aufnehmen kann.

Die Steckverbinder eignen sich für Leiterplatten mit einer Dicke von 1,60 bis 3 mm. Das Gewicht der größten Version (4 Kontakte) mit einer Länge von 7 mm liegt bei 25 g. Kona ist eine Weiterentwicklung der Reihe Datamate Mix-Tek, die weitere Anwendungsgebiete erschließen soll.

Mit einer maximalen Nennspannung von 3 kV pro Kontakt (1 min) und einem Arbeitstemperaturbereich von -65 bis 150°C eignen sich die Leistungssteckverbinder für den Einsatz unter rauen Betriebsumgebungen. Die Lebensdauer wird mit 250 Steckzyklen angegeben. Angeboten werden die Steckver-

binder im einreihigen Gehäuse (Versionen mit 2, 3 und 4 Kontakten). Der Kontaktwiderstand wird mit 2 mOhm angegeben.

Das Herz des Steckverbinders: die Kontakte

Die Steckverbinder basieren auf dem kompakten 6-poligen Beryllium-Kupfer-Stiftkontakt des Herstellers, der in der hochmodernen HQ-Anlage in Portsmouth / UK gefertigt wird. Eine Hartgold-Auflage von 0,76 bis 1 µm mit 98% Reinheit garantiert eine kontinuierliche elektrische Leitfähigkeit auch bei starken Erschütterungen und Vibrationen.

Die Steckverbinder wurden für 12 Stunden auf Vibrationskräfte von 20 G getestet (10 bis 2000 Hz). Weiterführende Prüfungen sind geplant.

Die Ausgasung erfüllt die Anforderungen der NASA und ESA für Vakuum- oder Weltraumanwendungen. Die Auszugskraft (pro Kontakt) wird mit 5 N angegeben, die Einsteckkraft mit 50 N.

THE FUTURE CODE

Experience the Transformation of Industry

29. – 30. Juni 2021

Vogel Convention Center Würzburg
Hybrides Event

We empower you to transform your business!

Die Event-Plattform „The Future Code“ bringt Digitalisierungs- und Industrie-Experten zusammen und bietet Zündstoff für die zentralen Themen der Digitalen Transformation. Dabei können Sie frei entscheiden, ob Sie am Event Online oder in Präsenz teilnehmen!

www.thefuturecode.de/anmeldung

Hebelbedienbare Leiterplattenklemmen und -Steckverbinder

Die aktuelle Leiterplatten-Steckverbinder der Baureihen LPT(A) und LPC(H) kombinieren die Zuverlässigkeit des Push-in-Federanschlusses mit der Bedienfreundlichkeit der Hebelbetätigung.

THORSTEN ROSIN *

Die Welt steht vor einem der größten Umbrüche in der Geschichte der Menschheit – fossile Energieträger sollen in den kommenden Jahren reduziert und durch Wind-, Wasser, Solar- oder Bioenergie substituiert werden. Der Strom aus den „erneuerbaren“ Energien wird direkt verbraucht, in E-Fuels (englisch: electrofuels; synthetische Kraftstoffe) umgewandelt oder in großen Lithium-Ionen-Batterien gespeichert.

In den elektronischen Geräten wie Steuerungen, Stromversorgungen, Wechselrich-

tern und Wallboxen ist die Leiterplattenklemme ein wichtiger Bestandteil – sie bildet zwischen Leiter und Leiterplatte eine Schnittstelle, ohne die kein Gerät auskommt.

Bei der Auswahl der geeigneten Schnittstelle stehen Hersteller solcher Baugruppen unterschiedlichen Anforderungen gegenüber. Dabei gelangen auch Aspekte wie Normenkonformität der Zielmärkte sowie Anwenderfreundlichkeit zunehmend in den Fokus. Hier zeigen die Leiterplattenklemmen und -Steckverbinder der Baureihen LPT(A) und LPC(H) ihre Stärken (Bild 1).

Neuer Komfort für den Leiteranschluss

Neben dem seit Jahrzehnten bewährten Schraubanschluss hat sich in den letzten Jahren der Push-in-Federanschluss für den komfortablen und sicheren Leiteranschluss als zeitsparende Alternative etabliert.

Diese Technologie, die der Blomberger Hersteller maßgeblich weiterentwickelt, macht den Leiteranschluss erheblich komfortabler. Zur Verfügung steht heute ein

breites Produktprogramm von Leiterplattenklemmen und -Steckverbindern in den Querschnittsbereichen von 0,14 mm² bis zu 25 mm² (Bild 2).

Das Prinzip ist einfach: Ein Betätigungshebel ermöglicht Anwendern, den Klemmraum werkzeuglos zu öffnen und zu schließen. Durch Öffnen des Betätigungshebels wird die integrierte Feder gespannt und der Klemmraum geöffnet. Beim Schließen des Hebels wird die Feder wieder entspannt, wodurch sie sich selbstständig in ihre Ursprungslage bewegt und damit die Kontaktkraft dauerhaft sicherstellt.

Die Form der Kabeleinführtrichter verhindert das Abspießen einzelner Litzen und ermöglicht den werkzeuglosen, schnellen und zuverlässigen Anschluss von flexiblen Leitern ohne Aderendhülse. Dank Push-in-Technik lassen sich flexible Leiter mit Aderendhülsen sowie starre Leiter auch bei geschlossenem Hebel direkt anschließen.

Der farbig abgesetzte und intuitiv bedienbare Hebel hat zwei wesentliche Vorteile: Zum einen signalisiert die Hebelposition von außen sicht- und fühlbar die definierten



* Dipl.-Ing. (FH) Thorsten Rosin
... arbeitet als Produktmanager Device Connector bei Phoenix Contact in Blomberg.

Ohne Werkzeug:
Leiterplattenklemmen und -Steckverbinder der durchgängigen Serien LPT und LPC verbinden die hohe Bedienfreundlichkeit der Hebelbetätigung mit dem zuverlässigen Push-in-Federkraftanschluss.

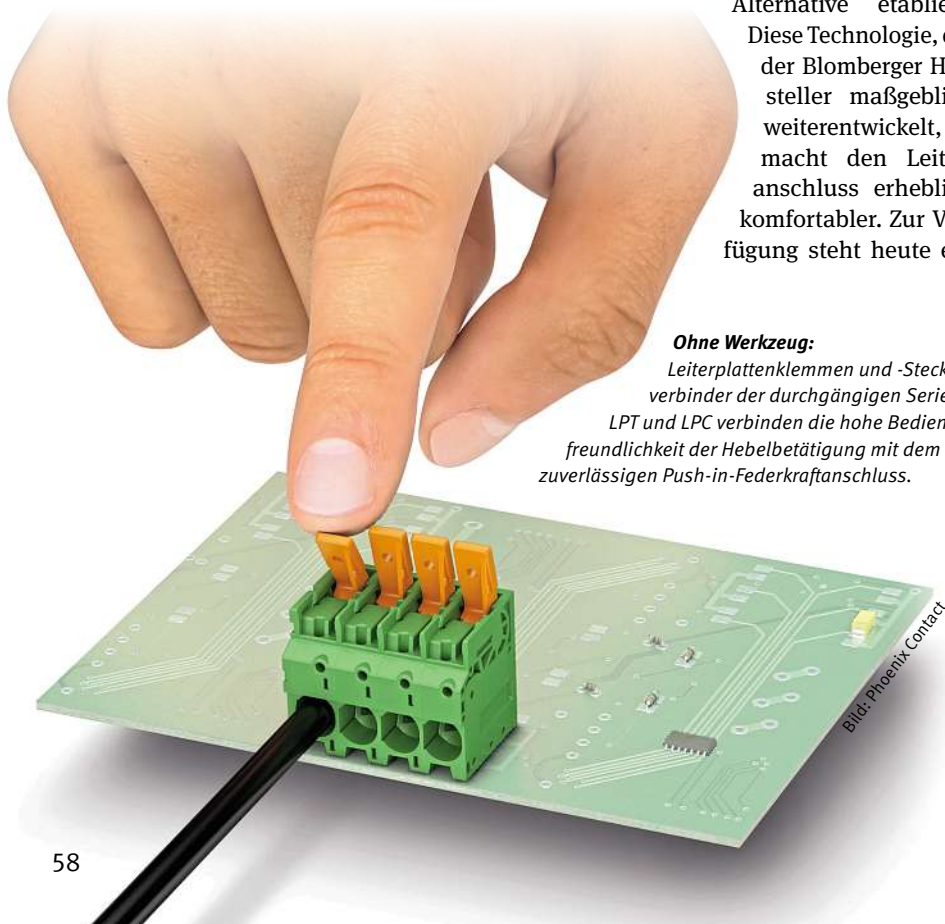


Bild 1:
Auch der Geräteanschluss wird durch die Hebelbedienung deutlich vereinfacht, wie an dem Applikationsbeispiel Photovoltaik zu sehen ist.

Bild: Phoenix Contact

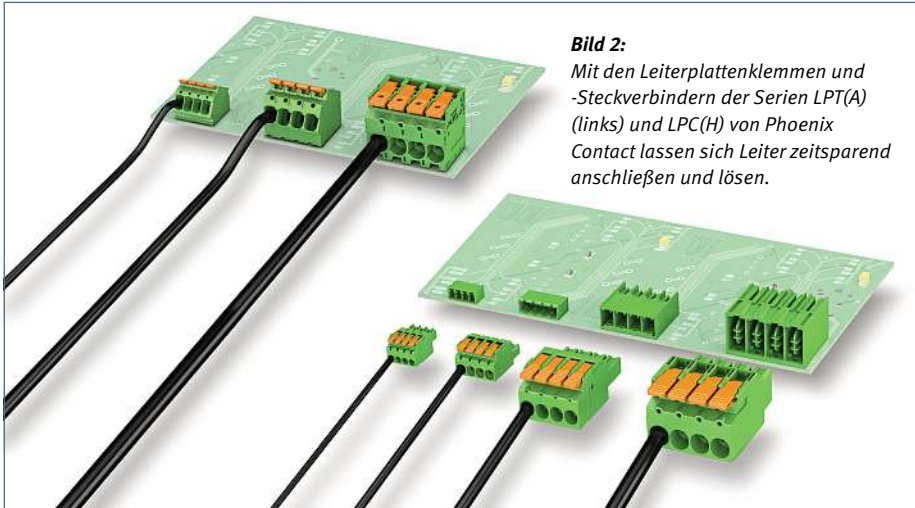


Bild 2:
Mit den Leiterplattenklemmen und -Steckverbindern der Serien LPT(A) (links) und LPC(H) von Phoenix Contact lassen sich Leiter zeitsparend anschließen und lösen.

Zustände des Klemmraums. So fallen nicht ordnungsgemäß geschlossene Klemmräume und somit fehlerhafte Verbindungen sofort auf.

Zum anderen ist die Kontaktkraft durch die Feder vorprogrammiert und stets gleich. Der Anwender kann sich sicher sein, dass die eingeführten Leiter zuverlässig und langzeitstabil durch das Umlegen des Hebels kontaktiert werden. Hierdurch wird eine potenzielle Fehlerquelle – etwa ein falsches Anzugsdrehmoment beim Schraubanschluss – sicher ausgeschlossen.

Beide Eigenschaften machen die Leiterplattenklemmen und -Steckverbinder zur Lösung für Anwendungen, bei denen der Anschluss selbsterklärend und zeitsparend erfolgen soll.

Initialkosten während der Installation und Inbetriebnahme sowie mögliche Folgekosten für Fehlersuche und Instandsetzung werden durch die Hebelbedienung auf ein Minimum reduziert. Die integrierten Prüfabgriffe der neuen hebelbedienbaren Leiterplattenklemmen und -Steckverbinder erleichtern auch nach der Verkabelung und Inbetriebnahme der Baugruppen eine einfache Prüfung im eingebauten und verkabelten Zustand.

Universelle Anwendung in allen Bereichen

Normative und zulassungsrelevante Anforderungen nehmen bereits in der frühen Entwicklungsphase Einfluss auf das Geräte-Design. Besonders wenn das Gerät international vermarktet werden soll, müssen neben den internationalen Normen – wie EN/IEC – auch die amerikanischen Normen erfüllt werden.

Insbesondere bei der Isolationskoordination wird den Geräteschnittstellen hohe Beachtung geschenkt. Die hierfür verwendeten

Komponenten finden sich auf der sogenannten „list of critical components“. Die Leiterplattenklemmen LPT(A) 6 und LPT 16 sind ebenso wie die Leiterplatten-Steckverbinder LPC(H) 6 und LPC 16 nach dem Standard UL 1059 uneingeschränkt für Spannungen bis zu 600 V geeignet.

Damit ist eine universelle Anwendung als sogenannter „field wiring terminal block“ in allen Bereichen gestattet. Sie bieten zudem eine erweiterte Fingerberührsicherheit von drei Millimetern (nach IEC/UL 61800-5-1). Damit erfüllen sie den für 400-V-TN-Systeme geforderten Schutz gegen direktes Berühren und erlauben den Geräteeinsatz ohne zusätzliche Abdeckungen.

Die Leiterplattenklemme LPT(A) 2,5 und die Steckverbinder LPC 1,5 / 2,5 sind ebenfalls nach dem Standard UL 1059 zugelassen. Sie weisen eine Spannung von 300 V nach Use-group B auf. Der Zulassungsprozess wird durch die Konformität zu den nationalen und internationalen Normen vereinfacht und ermöglicht so eine schnellere Markteinführung neuer Geräte auf den diversifizierten Märkten in Europa, den USA und Asien.

Fazit: Die Leiterplattenklemmen der Serie LPT(A) und die Leiterplatten-Steckverbinder der Serie LPC(H) sind im industriellen Umfeld eine weitere Alternative zu anderen Leiterplattenklemmen und -Steckverbindern mit Push-in und Schraubanschluss. Denn sie ermöglichen dem Anwender eine intuitive und zeitsparende Installation.

Das sorgt für eine hohe Zuverlässigkeit im Feld und stellt einen wirtschaftlichen Betrieb der gesamten Anlage sicher. Damit wird auch der Weg in die „All Electric Society“ auf kleinster Komponentenebene einfacher und sicherer. // KR

Phoenix Contact

IoT-fähig

Jetzt Termin vereinbaren!



Smart und Virtuell

Produkte digital transformieren

Wir machen aus Ihren Ideen globale IoT-Produkte:

- Hard- und Softwareentwicklung
- Industrial Engineering
- Produktion: Leiterplatten, Kabel, Gehäuse
- Komplettgeräteeinbau
- Einbindung in Cloud-Services (ADVANTECH WisePaaS)



Innovation Hub

Lernen Sie unsere IoT-Integration kennen! Werksführung und Innovation Hub zeigen Ihnen alle Möglichkeiten der IoT-Integration Ihrer Ideen und Produkte.

Lacon

Lacon Electronic GmbH

Hertzstraße 2
85757 Karlsfeld
www.lacon.de



Steckverbinder kurz erklärt: Steckzyklen und Oberflächen

Steckverbinder werden hinsichtlich ihrer Steckzyklen klassifiziert. Wir erklären, was damit gemeint ist und welche Eigenschaften durch die Oberflächenbeschichtung bestimmt werden.

UTE NIEMANN *

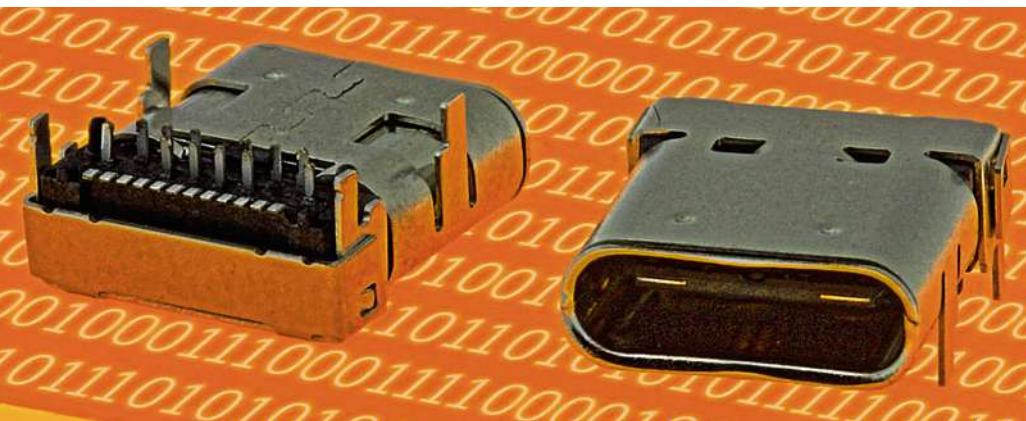


Bild: W + P

Steckzyklen: Die USB-Steckverbinder der Serie 8.320 von W+P mit selektiv vergoldeten Kontakten sind bis zu 10.000 Steckzyklen spezifiziert.

Die wesentliche Funktion von Steckverbindern ist der Anschluss eines Kabels an eine Leiterplatte, die Verbindung zweier Kabel oder zweier Leiterplatten. Derartige Verbindungen müssen entsprechend zuverlässig, kompakt und in der Regel lösbar sein.

Diese Trennbarkeit ist der Hauptgrund für die Verwendung von Steckverbindern: Wird beispielsweise ein Standort unabhängiger Einsatz gefordert oder die dezentrale Fertigung mit zentraler Endmontage oder eine Wartung, werden lösbare Steckverbinder mit der gewünschten Anzahl an Steckzyklen applikationsspezifisch ausgewählt.

Was verbirgt sich hinter dem Steckzyklus?

Der Kennwert „Steckzyklus“ bezieht die maximal ausführbaren Steck- und Ziehvorgänge in einem Kontaktsystem, welche von mehreren Einflussfaktoren bestimmt wer-

den: der Kontaktform, dem Kontaktwerkstoff und den Einsatzbedingungen. Idealerweise sollten aus mechanischer Sicht Steck- und Ziehkräfte minimal sein, damit nur ein geringer Materialverschleiß erfolgt. Ist der Verschleiß gravierender, verändern sich die elektrischen Eigenschaften und der Korrosionsschutz.

In einer typischen Kontakt-Architektur ist der Kontakt-Basiswerkstoff meistens eine Kupferlegierung. Darunter befindet sich in der Regel eine sogenannte Nickelsperrschicht, die bei Zinnoberflächen auch die Whiskerbildung minimiert. Darüber ist eine Schicht entweder aus Gold, Zinn, Silber bis hin zu Speziallegierungen wie Palladium-Nickel aufgebracht. Die Top-Beschichtung sichert letztendlich die elektrische Leistungsfähigkeit der Kontakte, mindestens für die geforderte Anzahl von Steckzyklen.

Die verschiedenen Arten der Oberflächenveredelung

Die Beschichtung der Kontakte erfolgt galvanisch oder chemisch. Gold beispielsweise ist bei entsprechender Schichtdicke (ab ca. 0,8 µm) mit geeigneter Kontaktgeometrie sowie Führung für eine sehr hohe Steckzyk-

lenzahl geeignet. Zudem können Steckverbinder mit dieser Goldschichtdicke unter nahezu allen Umgebungsbedingungen eingesetzt werden.

Die Kontaktoberfläche ist dem Verschleißprozess direkt ausgesetzt. Hier spielt die Härte der Beschichtung eine wesentliche Rolle.

Unterschiedliche Schichten: Flash- und Hartgold

Abhängig von der Anwendung erhalten unterschiedliche Oberflächen Gold-Beschichtungen bis ca. 1,27 µm. Alternativ werden Zinn und Silber verwendet, zu Zeiten von günstigen Palladium-Preisen auch dieses in einer Palladium-Nickel-Legierung.

Flash-Gold weist Schichtdicken im Bereich um 0,1 µm auf. Diese Beschichtung dient ausschließlich dem Korrosionsschutz während der Lagerung oder des Transports der Steckverbinder.

Bei D-Sub-Steckverbindern, modularen Steckverbindern sowie bei Anwendungen, die hohe Steckzyklen in rauer Industrieumgebung garantieren müssen, können die Hartgold-Schichtdicken höher sein, in der Regel bis zu 1,27 µm.

Die Oberflächenveredelung der Kontakte ist letztendlich eine wirtschaftliche Optimierungsaufgabe (Edelmetallkosten gegenüber der wirksamen minimalen Schichtdicke und Beschichtungswerkstoff). Denn die Oberflächenbeschichtung sorgt im Kontaktsystem für eine verlässliche Verbindung während der gesamten Produktlebensdauer.

Bei D-Sub Steckverbindern werden Steckzyklen üblicherweise in Güteklassen (GK) angegeben. Güteklasse 3 bezeichnet mehr als 50 Steckzyklen, Güteklasse 2 mehr als 200 und Güteklasse 1 mehr als 500 Steckzyklen.

Die beschriebenen Informationen gelten als Überblick und sind möglicherweise bei den verschiedenen Herstellern unterschiedlich.

// KR



* Ute Niemann
... arbeitet im Marketing bei W+P in Bünde.

W + P

HIGH-SPEED-VERKABELUNG

Glasfaser-Steckverbinder der nächsten Generation für 400 GBit/s

Da die Anforderungen an die Bandbreiten immer schneller steigen, werden 200-/400-Gbit/s-Übertragungsraten in Hochleistungs-Rechenzentren immer alltäglicher. Das bedingt erweiterte Funktionen auf Rack-Ebene. Auch Panduit hat für diese Anwendungen den Duplex-Glasfaser-Steckverbinder CS auf den Markt gebracht.

Dabei handelt es sich um die nächste Generation von Glasfaser-Steckverbindern des CS-Konsortiums, die das Rechenzentrum für 200-G bzw. 400-G-Anwendungen fit machen sollen.

Die kompakte Größe des Steckverbinders verdoppelt die Kanalkapazität von Racks, erweitert die Optionen beim Breakout-Modus und ermöglicht eine höhere Effizienz für Implemen-



Bild: Panduit

ermöglicht einen einfachen Zugang auch bei dicht gepackten Anwendungen. Beschädigungen benachbarter Steckverbinder beim Stecken oder Trennen werden verringert. Dadurch können CS-Steckverbinder auf einem Patchpanel auch vertikal verdichtet werden.

Mit CS-Steckverbindern sind bis zu 216 Verbindungen (LC-Stecker: 144) bei gleichem Platz möglich. Der Steckverbinder reduziert damit die Grundfläche des Bauteils auf der Leiterplatte in der Hyperscale-, Multi-Tenant- und großen Rechenzentren. Auch die Länge des Steckverbinders von der Ferrulenspitze bis zum Steckerfuß ist wesentlich kürzer als beim LC-Stecker.

Panduit

tierungen von 25 bis 200 G im Rack.

Komplexe Anwendungen wie KI, maschinelles Lernen (ML) und Deep Learning (DL) treiben die Anforderung voran, eine höhere Datendichte auf das Rack zu bringen. Die CS-Steckverbinder bieten eine 50 % höhere Port-Dichte pro Rack Unit und die

doppelte Dichte in QSFP-DD- und OSFP-Transceivern im Vergleich zu LC-Steckverbindern.

Das Kabelmanagement im Rack mit einem einheitlichen Kabel (Verkabelungssystem HD Flex des Herstellers) reduziert die Überlastung in den Pathways und die Push-Pull-Einsteck-/Ausziehfunktion des CS-Steckers

Hohe Leistung **UND** Zuverlässigkeit?

Sie wollen eine höhere Strombelastbarkeit für jeden Kontakt unserer hochzuverlässigen Steckverbinder?

Sie wollen eine einfache Anwendung und optimalen Einsatz unter Vibrationsbedingungen?

Sie wollen den hohen Qualitätsstandard von Harwin?

Wir sind ganz Ohr.

- 60A pro Kontakt
- Stoß-/Vibrationsfest bis 100G
- Edelstahl-„Mate-Before-Lock“-Fixierung
- Betriebstemperaturbereich bis 150°C



harwin.com/kona

KONA

HIGH RELIABILITY | HIGH POWER

HARWIN
Connect with confidence

Skalierbare Encoder-Anwendung mit nur einem Chip

Die Integration von Mixed-Signal-Schaltungsblöcken mit reflexiver On-Chip-Sensorik bietet hohe Messgenauigkeit und ist flexibel. Somit passt sich der Baustein individuellen Encoder-Durchmessern an.

SILVAN ETTLE *

Integrierte Encoder-Entwicklungen werden maßgeblich durch die steigenden Anforderungen aus der Automatisierung und Robotik bestimmt. Dazu sind flexible Durchmesser für den Einsatz in unterschiedlichen Gelenken eines Roboterarms genauso wichtig, wie die einfache Montage in kompakten Bauräumen und Hohlwellen. Außerdem sind hohe Auflösungen sowie moderne Kommunikationskonzepte gefragt. Hier kombiniert iC-Haus mit seiner Serie iC-PZ die Vorteile des reflexiven Messverfahrens mit der aktuellen Digitaltechnik. Möglich sind hohe mechanische Toleranzen und die gewünschten geringen Aufbauhöhen bei einer hohen Messpräzision.



* Silvan Ettle
... ist bei iC-Haus für Vertrieb und Applikation für Encodersysteme zuständig.

Der Einsatz der Encoder ist flexibel: Miniatur-Encoder, Hohlwelle oder Linearführung. Mit der zum Patent angemeldeten FlexCode-Technik lassen sich die Abmessungen der Maßverkörperung nahezu beliebig wählen. Ein speziell hierzu entwickelter Algorithmus stellt sicher, dass die aus der Projektion eines Pseudo-Random-Codes extrahierte absolute Positionsinformation bei individuellen Codelängen eindeutig interpretiert wird.

So lassen sich Messsysteme mit Scheibendurchmessern von wenigen Millimetern bis zu meterlangen Linearskalen mit nur zwei Chipvarianten der Serie iC-PZ realisieren. Die Designs mit FlexCode werden bei iC-Haus entwickelt, simuliert und in Form von CAD-Daten bereitgestellt. Eine optimierte Maßverkörperung aus Aluminium, Stahl, Glas, Polycarbonat oder anderen Basismaterialien lässt sich individuell fertigen. Viele unterschiedliche Prozesstechniken sind für die

Herstellung der Maßverkörperung geeignet. So lassen sich die reflektierenden Code-Strukturen beispielsweise durch Ätzen, Lasern, oder lithographische Verfahren auf das Basismaterial applizieren.

Eine gleichbleibend hohe Messqualität

Auch bei der Montage ist die Serie flexibel. Der Sensor kann von $\pm 0,5$ mm (tangential), $\pm 0,4$ mm (radial) und $\pm 2^\circ$ (Verdrehung) relativ zur Maßverkörperung positioniert werden. Innerhalb des Arbeitsbereiches sorgen die im Chip integrierten automatischen Kalibrierfunktionen für eine gleichbleibend hohe Messqualität des Encoders. Zusätzlich lässt sich die aufbaubedingte Exzentrizität einer Codescheibe vom Sensor ermitteln und weitgehend kompensieren. So wird dank einer intelligenten Signalverarbeitung der Einfluss der Mechanik auf die Messgenauigkeit signifikant reduziert. Somit sinken die

Skalierbarer Encoder:
Mit der Serie iC-PZ kombiniert
Hersteller iC-Haus bewährte Analog-
mit Digitaltechnik.

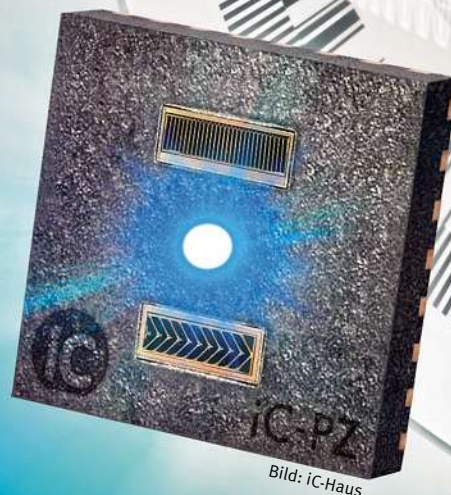


Bild: iC-Haus



Serie iC-PZ: Wurde vor allem durch die Automatisierung und Robotik getrieben. Hier kommt es auf flexible Durchmesser für die unterschiedlichen Gelenke des Roboters an.

Fertigungskosten, während der Anwender von den präzisen Messergebnissen profitiert.

Die Serie iC-PZ basiert auf einem Mixed-Signal-Halbleiter-Chip, der in CMOS-Technik gefertigt wird. Im Zentrum des hochintegrierten System-on-Chip (SoC) ist eine blaue LED montiert. Oberhalb bzw. unterhalb der LED befinden sich die optischen Sensorflächen, die mit Glas bedeckt plan abschließen. Die im Chip integrierte LED-Regelung ist auf die Empfindlichkeit der Fotodioden abgestimmt, sodass nicht nur die Lichtintensität, sondern auch die Signalgüte konstant bleibt. Durch dieses Zusammenspiel gewinnt das System zudem an Flexibilität in der Aufbauhöhe.

Der Luftspalt zwischen Sensor und Maßverkörperung kann um bis zu einem Millimeter variieren, ohne Einfluss auf die Signalqualität zu haben. Das SoC im gemoldeten Standard QFN-Gehäuse arbeitet zuverlässig innerhalb einer Umgebungstemperatur von -40 bis 125 °C und eignet sich damit auch für thermisch anspruchsvolle Einsätze. Mit einer Kantenlänge von fünf Millimeter findet das quadratische Gehäuse in kompakten Bauräumen Platz und kann mit Standard SMD-Technik assembliert werden.

14-Bit-Interpolator errechnet Positionswort

Das Messsystem basiert auf zwei synchron abgetasteten optischen Kanälen. Die Projektion reflektierender Strukturen auf zwei Spuren der Maßverkörperung wird von den gegenüberliegenden Fotodioden erfasst. Die innere Spur auf der Scheibe liefert die absolute Positionsinformation als Pseudo Random Code. Mit äquidistanten Strichen auf der äußeren Spur wird ein inkrementelles Signal generiert. Der integrierte 14-Bit-Interpolator verrechnet beide Informationen zu einem hochauflösenden Positionswort. Die hohe Interpolierbarkeit ist den fein aufeinander abgestimmten Systemkomponenten und Funktionsblöcken zu verdanken. Eine Codescheibe mit einem Durchmesser von

26 mm wird mit rund 0,3" aufgelöst. Linear-systeme mit der Serie iC-PZ bieten eine Auflösung von bis zu 12,5 nm.

Auf drei konfigurierbaren Ports können verschiedene Schnittstellen parallel betrieben werden. Eine beliebige Kombination aus ABZ, Sin/Cos, UVW, SPI sowie BiSS/SSI ist möglich. Die jeweilige Konfiguration wird über Pins geschaltet, sodass grundsätzlich weder Design noch Hardware angepasst werden müssen. Das Positionsformat der seriellen Schnittstellen ist individuell einstellbar. Dank FlexCount-Technik ist zudem die Auflösung der Quadratursignale frei wählbar. Die umfangreichen Konfigurationsmöglichkeiten eignen sich besonders für das Design unterschiedlicher Varianten innerhalb einer Produktfamilie mit den Vorteilen der Skalierbarkeit, von denen Entwicklung, Produktion und Logistik beim Encoder-Hersteller gleichermaßen profitieren.

Betriebszustände und Statusmeldungen überwachen

Die integrierten Sensoren der Serie iC-PZ werden vielfältig ausgewertet. Angeordnet in einer BiSS-Kette werden nicht nur synchron abgetastete Positionswerte zuverlässig und in Echtzeit bereitgestellt, selbst umfangreiche Diagnoseinformationen können in den BiSS-Stream eingebunden werden. So lassen sich Betriebszustände und Statusmeldungen bequem überwachen. Zusätzliche Informationen von externer Sensorik, wie beispielsweise Feuchte und Temperatur, können über den im iC-PZ integrierten I²C-Master abgefragt und an die übergeordnete Steuerung übermittelt werden.

Die Datenbasis für umfangreiches Condition Monitoring wird so direkt vom Sensor bereitgestellt. Zusätzliche Mikrocontroller oder FPGAs im Encoder, die speziell für diesen Zweck programmiert werden müssen, können entfallen.

// HEH

iC-Haus

VOLLER EINSATZ, RUNDUM GESCHÜTZT

Produkte aller
Sicherheitsbereiche aus
einer Hand, schnell und
zuverlässig geliefert.

PERSÖNLICHE
SCHUTZAUSRÜSTUNG
MASCHINENSICHERHEIT
BETRIEBSSICHERHEIT
ELEKTRISCHE SICHERHEIT



de.rs-online.com

Dynamische Eigenschaften von Leistungshalbleitern

Der Doppelpulstest hilft dabei, wenn die dynamischen Eigenschaften eines Leistungshalbleiters verglichen werden sollen. Für den Vergleich der dynamischen Parameter sind wichtige Punkte zu beachten.

MICHAEL ZIMMERMANN, RYO TAKEDA, BERNHARD HOLZINGER UND MIKE HAWES *

Bei der Bewertung von Leistungstransistoren für ein Leistungswandler-Design müssen Entwickler den richtigen Baustein auswählen. Für den Test eignet sich ein standardisierter Doppelpulstest (DPT), um Leistungstransistoren mit großem

Bandabstand (WBG) zu charakterisieren. Das Testsystem muss die parasitären Spannungen und Ströme klein halten und schließlich von System zu System konsistent halten. Doch wie muss ein Standard-DPT-System designet werden, damit die Ergebnisse zwi-

schen mehreren Testsystemen korreliert werden können? Bei genauerem Hinsehen und unter Berücksichtigung der ständig steigenden Schaltgeschwindigkeiten von WBG-Bausteinen, gibt es jedoch viele externe parasitäre Komponenten, die im System berücksichtigt werden müssen.

Einfluss der parasitären Anteile auf die Schaltparameter

Viele externe parasitäre Anteile, insbesondere die drei Hauptschleifen – Power-Schleife, Gate-Schleife und Zwischenkreis-Schleife – werden oftmals wegen eines Überschwingers (Ringing), das sie in die Signale einbringen, berücksichtigt. Parasitäre Störungen haben jedoch auch einen großen Einfluss auf die extrahierten Schaltparameter. Externe Einflüsse wie die gemeinsame Induktivität von Power-Schleife und Gate-Schleife L_{SS} , der externe Gate-Widerstand R_G , die externe Gate-Induktivität L_G und Lastinduktionsparasitäten beeinflussen die Schaltgeschwindigkeit eines Leistungshalbleiters. Zusätzlich wird die gemessene Schaltenergie durch die parasitäre Kapazität C_{ind} der Lastinduktivität und die parasitäre Induktivität L_{shunt} des Strom-Shunts beeinflusst.

Die Bandbreite des Shunt beeinflusst die Schaltenergie und das parasitäre Element L_{shunt} führt zu einer höheren gemessenen Stromspitze beim Einschalten und verstärkt alle hochfrequenten Komponenten des gemessenen Stromsignals. Mit der Charakterisierung des verwendeten Shunts lässt sich der Einfluss von L_{shunt} minimieren, sodass er für einen Vergleich bei entsprechender Kompensation nicht berücksichtigt werden muss. Datenblätter von Leistungshalbleitern enthalten nur wenige Informationen über Doppelpulstest-Systeme. Somit ist es für Entwickler schwierig, die notwendigen

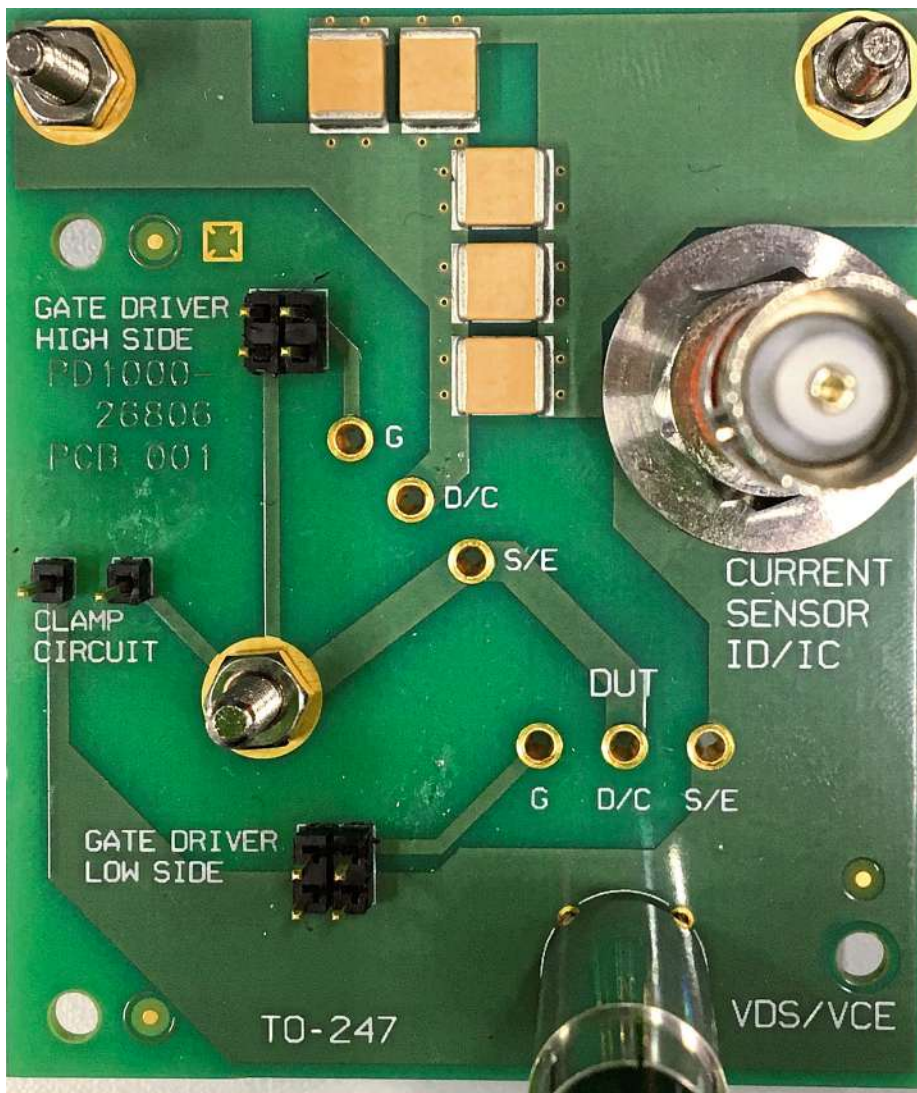


Bild: Keysight Technologies

Standard-Testboard: Das TO247 bietet dem Entwickler die Möglichkeit, sowohl bei Änderungen an Bauteilen und am Gate-Treiber flexibel zu sein.

* Michael Zimmermann, Ryo Takeda, Bernhard Holzinger und Mike Hawes ... arbeiten bei Keysight Technologies.

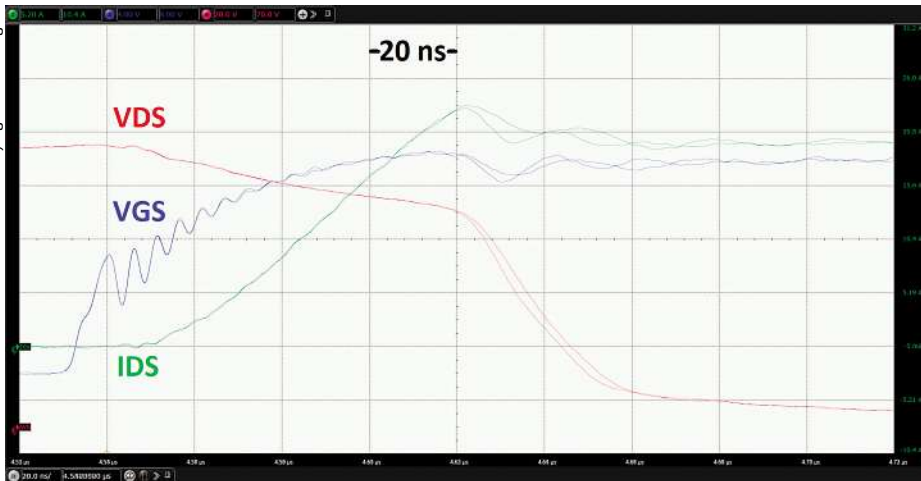


Bild 1: Vergleich der Einschaltsignale zweier Induktivitäten mit unterschiedlichen C_{ind}

Schaltparameter abzuleiten. Zusätzlich zu den Testparametern wie V_{DS} , I_{DS} oder V_{GS} werden nur R_G und L_{ind} angegeben. Alle genannten Parameter lassen sich leicht kontrollieren und ändern.

Einige Datenblätter zeigen die parasitäre Kapazität C_{ind} der Lastinduktivität und die Gesamtinduktivität der Power-Schleife an. C_{ind} ist ein wichtiger Parameter, da er eine zusätzliche Kapazität parallel zum High-Side-Bauelement einführt. Während des Einschaltens verursacht die zusätzliche Kapazität höhere gemessene Spitzenströme als der tatsächliche Sperrerrholungsstrom des Bausteins und erhöht daher die Schaltenergie während des Einschaltens. Das Layout der Vorrichtung selbst kann ebenfalls parasitäre Kapazitäten verursachen, die eine ähnliche Wirkung haben und niemals spezifiziert werden. Dieser Effekt zeigt das Bild 1, das Einschaltkurvenformen für $V_{DS} = 100\text{ V}$ und $I_{DS} = 20\text{ A}$ zeigt. Die Messung mit einer Induktivität mit höherem C_{ind} zeigte eine höhere und etwas längere Stromspitze und eine verzögerte Flanke des V_{DS} -Abfalls. Das erhöht die Schaltenergie um 4,5%. Wichtig ist, sowohl C_{ind} als auch die durch das Layout eingeführte parasitäre Kapazität minimal zu halten.

Ein Blick auf die Gesamtinduktivität der Power-Schleife

Die Gesamtinduktivität der Power-Schleife ist wichtig, da sie einen Spannungsabfall in V_{DS} während der Einschaltflanke von i_D erzeugt:

$$V_{DS,droop} = L_{powerloop} \times (di/dt)$$

Der Spannungsabfall wird bei schnell schaltenden Geräten mit steilen Stromflanken signifikant und muss bei Berechnungen der Anstiegszeit berücksichtigt werden. Die Induktivität der Power-Schleife muss in ihre Komponenten L_{SS} und L_{DS} zerlegt werden, da

sie das System und die Messergebnisse auf unterschiedliche Weise beeinflussen. L_{SS} verlangsamt die Schaltgeschwindigkeit, da es eine negative Rückkopplung in der Gate-Schleife erzeugt. Dadurch wird die Steilheit des Ausgangsstroms für Ein- und Ausschaltübergänge verlangsamt. Die langsamere Stromflanke verringert das Überschwängen. Im Gegensatz dazu ist L_{DS} die Hauptquelle des Überschwängens bei V_{DS} , I_{DS} , V_{GS} , hat aber praktisch keinen Einfluss auf die Schaltgeschwindigkeit. Es ist wichtig, L_{SS} und L_{DS} zu kennen, da ihr Einfluss gegensätzliche Auswirkungen hat.

Der Signalversatz und sein Einfluss auf das System

Der Zeitversatz an den Eingängen des Oszilloskops beeinflusst ebenfalls die Messung, lässt sich aber minimieren. Dieser Signalversatz ist wichtig für Parameter, die von zwei verschiedenen Signalen abhängen, wie Schaltenergie und Verzögerungszeit. Entscheidend ist die Position des Tastkopfs. Sie können entweder nahe oder weit entfernt vom Prüfling platziert werden. Dadurch ändert sich die Leiterbahnlänge zum Tastkopf. Aufgrund der Signallaufzeit kommt es zu einer zusätzlichen Abweichung im System. Die Änderung der Position von V_{GS} um einige Zentimeter verändert die Verzögerungszeit um 0,5 ns (Bild 2). Bei langsam schaltenden Bauteilen mit Verzögerungszeiten $>50\text{ ns}$ ist das nicht signifikant, jedoch führen bei schnelleren Bausteinen Verzögerungszeiten von $\leq 5\text{ ns}$ zu einem Fehler von 10% oder mehr.

Aktuell verwenden Halbleiterhersteller selbstgebaute DPT-Systeme, um Schaltparameter zu erhalten. Viele Hersteller sind nicht in der Lage, die Ergebnisse von System zu System zu korrelieren. Für vergleichbare Er-

Schneller zur Marktreife dank
innovativer Messtechnik



- Oszilloskope 50MHz → 1GHz
- Signal Generatoren 10MHz → 500MHz
- Labornetzteile → 200W
- Elektronische Lasten 200 → 300W
- Multimeter 4½ → 6½ Stellen

**Siglent HF-Messgeräte: Helfer
bei der Entwicklung fehlerfreier
und zuverlässiger Kommunikation**



Spektrum- und Vektornetzwerkanalysator

- EMI-Pre-Compliance
- Modulationsanalyse
- Distance-to-Fault Analyse

HF-Signalquellen

- mit/ohne IQ-Modulation



Bild: Keysight Technologies

Bild 2: Überlagerte Einschaltssignale mit verschiedenen Positionen des Tastkopfes V_{GS} (blau).

gebnisse zwischen mehreren Systemen gibt es zwei Möglichkeiten, die WBG-Leistungstristoren zu testen. Im ersten Ansatz charakterisieren die Hersteller die Parasitäten des DPT-Systems und schließen sie in die Testbedingungen der spezifizierten Eigenschaften ein. Jedoch ist es schwierig, Informationen über alle Parasitäten von Layout bis Induktivität zu kennen und auszutauschen. Einige lassen sich kaum oder sogar überhaupt nicht mit hoher Präzision messen. Selbst dann nicht, wenn alle Bedingungen mit hoher Genauigkeit bekannt wären und angegeben würden (Tabelle).

Vergleich von Messergebnissen verschiedener DPT-Systeme

Der Vergleich von Messergebnissen aus zwei verschiedenen DPT-Systemen ist schwierig. Eine entsprechende Tabelle mit allen notwendigen Testbedingungen wäre hierbei sehr umfangreich. Die beiden zu untersuchenden Bausteine wurden bei gleicher Prüfspannung, gleichen Strom und gleichen Gate-Widerstandswert aber auf unterschiedlichen DPT-Systemen gemessen. Baustein A zeigt deutlich ein langsames Einschalten und eine höhere Schaltenergie zu Baustein B. Einer der Haupteinflüsse von Schaltgeschwindigkeit und -energie ist Parameter L_{SS} . Das zur Charakterisierung von Bauelement A verwendete Testsystem weist im Vergleich zum Testsystem von B eine signifikant höhere L_{SS} auf. Ohne Simulationen und Analyse ist es schwierig, Bausteine zu vergleichen und zu wissen, welcher Baustein in der Zielanwendung schneller ist und weniger Schaltenergie verbraucht. Für vergleichbare Ergebnisse werden die parasitären Anteile des DPT-Systems konstant gehalten. Dazu ist ein Standard-DPT-System erforderlich. Ein gut konzipiertes DPT-System reduziert mögliche

menschliche Fehler auf ein Minimum. Damit bleiben alle Parasitäten und andere Einflussfaktoren weitgehend konstant. Es lassen sich die Geräteparameter mehrerer Geräte ermitteln und der Vergleich von Geräten verschiedener Hersteller ist einfacher.

Für ein zuverlässiges und flexibles System

Für konstante Parasitäten und ein zuverlässiges sowie flexibles System bietet es sich an, Komponenten auf eine Leiterplatte zu löten. Allerdings wird Flexibilität eingebüßt. Bei Änderungen muss das alte Bauteil ausgelötet und ein neues eingelötet werden. Die Leiterplatte nimmt Schaden und die Messergebnisse sind nicht mehr vergleichbar. Auf eine spezielle DUT-Schnittstellenkarte mit festen Sockel-Steckverbindern lassen sich die Bauteile ohne Löten wechseln. Bei den Gate-Treibern bietet eine separate, austauschbare Platine mit unterschiedlichen R_G -Werten, die auf DUT-Platine gesteckt werden kann. Die Gate-Treiberplatine ermöglicht eine wiederholbare und konsistente Verbindung mit der DUT-Platine.

PARAMETER	BAUSTEIN A	BAUSTEIN B
V/I	600 V/20 A	600 V/20 A
R_G	0 Ω	0 Ω
$t_{\text{delay,on}}$	43 ns	39 ns
t_{rise}	34 ns	32 ns
E_{on}	563 μJ	547 μJ
L_{DS}	10 nH	15 nH
L_{SS}	10 nH	5 nH

Tabelle: Vereinfachtes Beispiel, das die Schwierigkeit bei der Vergleichbarkeit von Daten zeigt, die auf verschiedenen Testsystemen erfasst wurden.

Für verschiedene Gate-Widerstände sollte entweder eine Gate-Treiberplatine mit schaltbaren Widerständen oder mit austauschbaren Widerständen verwendet werden. Ein Schalter in der Gateschleife trägt zur parasitären Gate-Induktivität bei und ist daher nicht empfehlenswert. Gleiches für austauschbare Through-Hole-Widerstände. Die parasitäre Induktivität ist viel höher als bei oberflächenmontierbaren Widerständen. Außerdem können die Widerstände leicht verwechselt werden. Eine kleine und konstante Gate- und Power-Schleife lässt sich mit einem Standard-DUT-Board-Design inklusive austauschbaren Standard-Gate-Treiberplatinen aufrechterhalten.

Um die parasitäre Induktivität L_{DCLink1} und L_{DCLink2} vom Zwischenkreis-Kondensator zum DUT klein zu halten, ist ein Kompromiss nötig. Große Kondensatoren sind unpraktisch. Sie sollten auf eine separate Hauptkondensator-Platine montiert werden. Allerdings erhöht das L_{DCLink1} und L_{DCLink2} . Stattdessen enthält jede DUT-Testplatine Entkopplungskondensatoren, um die Power-Schleife möglichst klein zu halten und dabei nicht bei der Flexibilität einzubüßen. Über gut gewählte Werte der Entkopplungs-Kondensatoren wird der Einfluss größerer L_{DCLink} -Werte gesenkt. Bei konstanten parasitären Induktivitäten in der Gate- und Power-Schleife hat die Hauptinduktivität den größten Einfluss.

■ Das Layout oder Design der Induktivität beeinflusst den Gleichstromwiderstand und die parasitäre Kapazität. Für vergleichbare Ergebnisse sollte für jedes System das gleiche Design verwendet werden.

■ Eine Induktivität erzeugt ein Magnetfeld. Eine veränderte Position der Induktivität führt zu unterschiedlichen Auswirkungen des in das System induzierten Magnetfeldes. Die Position der Induktivität muss feststehen.

■ Die Kabellänge und die Kabelpositionierung der Induktivität führt ebenfalls zu Parasitäten. Ein langes Kabel erhöht den parasitären Einfluss. Die Induktivität sollte fest positioniert und fest angeschlossen sein.

Sowohl die IEC (International Electrotechnical Commission) als auch JEDEC arbeiten seit Jahrzehnten an der Definition von Test- und Charakterisierungsstandards. Mit den schnelleren WBG-Bausteinen sind Teststandards schwieriger zu entwickeln. Mit dem Wide Bandgap Power Electronic Conversion Semiconductor Committee sollen Standards für schnelle WBG-Leistungshalbleiter entwickelt werden.

// HEH

Keysight Technologies

OSZILLOSKOP FÜR EINE RACKINSTALLATION

Multiples Messgerät mit bis zu 512 analogen Eingängen

Die Oszilloskop-Serie DS8000-R von Rigol basiert auf der UltraVision-II-Architektur mit einem Phoenix-Chip-Set und zwei eigenentwickelten ASICs für das analoge Front-End und das Signal Processing übernehmen. Umgeben sind die Chips von einem Xilinx Zync-7000-SoC, Dual-Core ARM-9-Prozessoren, dem Linux +Qt-Betriebssystem, einem High-Speed-DDR-System-Speicher und einem QDR-II-Display-Speicher.

Bei der leicht erweiterbaren R-Serie mit Bandbreiten von 350 MHz bis 2 GHz (optionale Bandbreitenerweiterung) und einer maximalen Abtastrate von 10 GS/s eignen sich die Oszilloskope, um Signale schnell zu erfassen und zu analysieren. Dabei kommt dem Anwender die Halb-19"-Bauform (1HE) zusammen mit einem Modul zur Synchronisation zu gute: Es lassen sich mehrere Einschübe als multiples Messgerät mit bis zu maximal 512 analoge Eingangskanäle kaskadieren. Das dürfte gerade in Forschung und Entwicklung interessant sein. Ein Grundmodul hat vier analoge Eingangskanäle, sowie einen externen Trigger-Eingang und einen Arbiträr-Generatorausgang mit 25 MHz. Das Messgerät kann sowohl als Laborgerät als auch in ein Rack-System integriert verwendet werden. Typische Anwendungen



Oszilloskop: Das DS8000-R bietet bis zu 512 analoge synchronisierte Eingangskanäle.

sind automatisierte Tests in Fabriken, Protokollanalysen für serielle Busse in der Fahrzeugelektronik, Messen elektronischer Schaltungen, multiple Überwachungen für Forschungszwecke, Schaltleistungsmessungen sowie -analysen in der Leistungselektronik.

Um große Datenmengen zu erfassen und zu verarbeiten steht für alle Kanäle eine Speichertiefe von 500 MPts bereit. Mit einer Signalerfassungsrate von bis zu 600.000 Wfms/s können Entwickler sehr schnelle Fehlerimpulse erfassen. Die Echtzeit-Aufzeichnung und Wiedergabe von Signalen gibt Rigol mit bis zu

450.000 Frames an. Auch diese Gerätevariation bietet für die Frequenzanalyse den Einsatz einer hochauflösenden FFT mit 1 Mio. Abtastwerten an. Die Oszilloskop-Variante wurde mit der neuen integrierten Messmethode mit Echtzeit Augendiagramm und Jitter-Analyse-Software, sowie der Darstellung des Jitter-Trends speziell für die digitale Analyse deutlich erweitert.

Vielfältigste Trigger-, Mathematik- und Darstellungsmöglichkeiten (erweiterte FFT von 1 Mio. Punkten, Masken-Test und Power-Analyse) sind wie alle üblichen seriellen Bus-Protokoll-Analyse- und Triggerfunktionen

erhältlich. In das Messgerät integriert sind Spannungsmesser, Frequenzzähler und ein optionaler 2-Kanal arbiträrer Funktionsgenerator. Sollen mehr als vier Kanäle genutzt werden, lassen sich alle Einzelgeräte über eine Synchronisationseinheit miteinander kombinieren. Die Einheit lässt sich über die optische 10-Gbit-Schnittstelle 10GE SFP+ an einen Netzwerkrouter anbinden. Daten können mit hoher Geschwindigkeit auf einen PC übertragen werden.

Zur Auswertung der Messdaten auf dem PC bietet Rigol seine Software ULTRADAQ LITE. Entwickler können mit dem Software-Entwickler-Kit (SDK), das als Open Source angeboten wird und sich individuell anpassen lässt. Das DS8000-R lässt sich außerdem in rauen Industrie-Umgebungen mit Temperaturen von bis zu -40 °C betreiben. Weitere verschiedene Schnittstellen wie USB-Host, USB-Device, HDMI, LAN, und USB-GPIB (mit einem Adapter) sind verfügbar. Die Oszilloskope lassen sich zusätzlich über Maus und Tastatur oder über Web-Control steuern und nutzen. Höhere Bandbreiten, serielle Dekodierung und die 1-Kanal-Arb-Generator-Funktion sind per Software-Upgrade möglich.

// HEH

Rigol

DIE EINZIGEN 12-BIT OSZILLOSKOPE MIT 350 MHz – 2 GHz

Bis zu 16 analoge Kanäle, 32 digitale Kanäle und 5 Gpts Speicher.

WaveRunner 8000HD

HD
4096



TELEDYNE LECROY
Everywhere you look™

Tel. 06221-82700 | teledynelecroy.com/wr8000hd

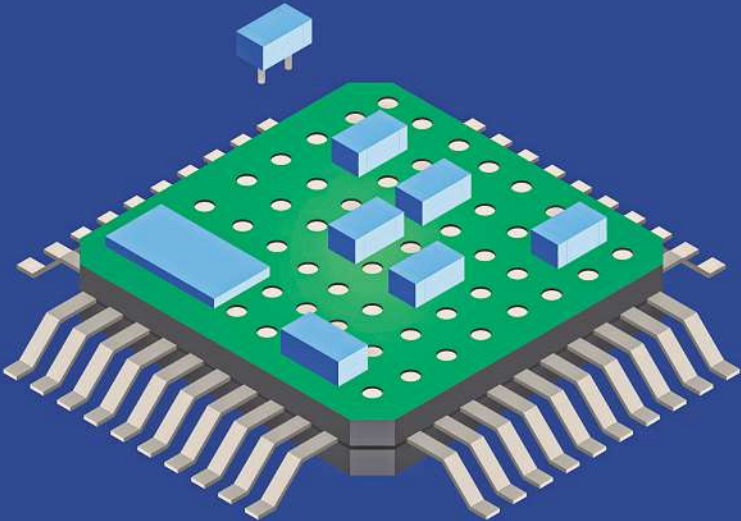
Schneller zum Produkt mit hochintegrierten Lösungen

Bauelemente wie System-on-Chip (SoC) oder System-in-Package (SiP) vereinen verschiedene Funktionen in sich und ermöglichen den raschen Zugang zu innovativen Halbleitertechnologien.

FRANK-STEFFEN RUSS *

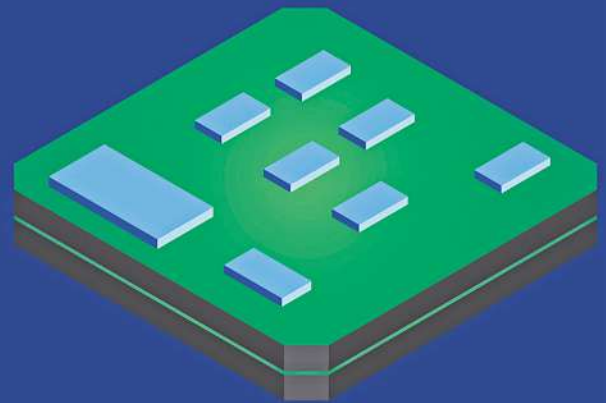
Bild: EBV

System-on-Chip (SoC)



Geringe Herstellungskosten
Geringer Platzbedarf
Geringer Energiebedarf

System-in-Package (SiP)



Kürzere Entwicklungszeit
Geringere Entwicklungskosten
Einfachere Hetero-Integration digitaler und analoger Komponenten

SoC und SiP im Vergleich: System-on-Chip (SoC) und System-in-Package (SiP) bieten für verschiedene Endmarktanwendungen unterschiedliche Vorteile.

Die Halbleitertechnologie ist heute die wichtigste Schlüsseltechnologie für Innovationen und Digitalisierung. Neue Entwicklungen in der Chip-Technologie machen mikroelektronische Systeme immer günstiger, leistungsfähiger und energieeffizienter. Intelligente Sensorlösungen und maßgeschneiderte Kommunikationstechnologien bilden die Basis des Internets der

Dinge – so entstehen neue Produkte, Geschäftsmodelle und Services, die Lösungen für die Herausforderungen unserer Gesellschaft versprechen.

Die Zahl der mit dem Internet verbundenen Geräte, einschließlich der Maschinen, Sensoren und Kameras, die das Internet der Dinge (IoT) ausmachen, wächst in einem stetigen Tempo. Eine neue Prognose der In-



* Frank-Steffen Russ
... ist Senior Director Market Segments bei EBV Elektronik.

ternational Data Corporation (IDC) schätzt, dass es im Jahr 2025 41,6 Milliarden angeschlossene IoT-Geräte oder „Dinge“ geben wird, die 79,4 Zettabytes an Daten erzeugen.

„Die Welt um uns herum wird immer ‚sensibilisiert‘ und bringt neue Ebenen der Intelligenz und Ordnung in persönliche und scheinbar zufällige Umgebungen. Die Geräte des Internet der Dinge sind ein integraler Bestandteil dieses Prozesses“, sagt David Reinsel, Senior Vice President, IDC’s Global DataSphere.

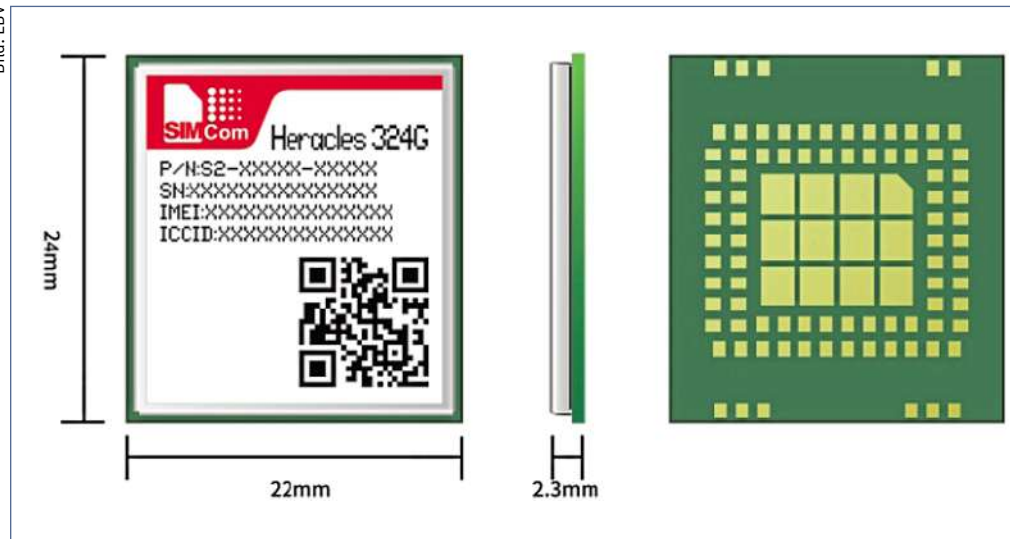
Mit der Digitalisierung der Welt ziehen immer kürzere Produktlebenszyklen bei gleichzeitig zunehmender Variantenzahl, wie man sie in der IT-Branche schon lange gewöhnt ist, jetzt auch in alle anderen Branchen ein. Doch die Entwicklung eines neuen Hardwareprodukts ist aufwendig und birgt viele Unwägbarkeiten, die die Entwicklungszeit erheblich verlängern, den Start verzögern und Unternehmen viel Geld kosten könnten. Die Programmierung von IoT-Bausteinen zum Sammeln und Senden von Sensormesswerten an die Cloud dauert zum Beispiel in der Regel Monate und erfordert hoch qualifizierte Fachleute.

Große Unternehmen können ihre Entwicklungsabteilungen mit dem entsprechenden Know-how und der erforderlichen personellen Bandbreite ausstatten. Aber Start-ups und Einzelunternehmer können auf diese Ressourcen nicht zurückgreifen.

Verschiedene Funktionen integriert

Dennoch bringen gerade kleine dynamische Unternehmen zunehmend innovative IoT-Geräte und smarte Produkte auf den Markt. Das ist auch der Halbleiterindustrie zu verdanken. Sie bietet heute nicht mehr nur reine Mikroprozessoren an, sondern liefert auch vorgefertigte Bausteine, die bereits über viele Funktionalitäten für verschiedene Anwendungen verfügen. Solche Bausteine werden als System-on-Chip (SoC) oder System-in-Package (SiP) beschrieben. Ein System-on-Chip vereint sämtliche Funktionen eines Systems auf einem Chip, zum Beispiel CPU, Signalprozessor, Grafikprozessor, Sicherheitselement sowie Konnektivität. So lassen sich kleinere Baugrößen und eine höhere Performance erreichen sowie Kosten und Energieverbrauch senken. Darüber hinaus kann ein SoC eine höhere Designsicherheit auf Firmware- und Hardwareebene bieten. Besonders die Reduzierung der Kosten durch SoCs sorgt zudem dafür, dass sich heute nahezu jedes Gerät beziehungsweise Produkt mit Intelligenz ausrüsten lässt. Während das System-on-Chip ein einzelner Chip

Bild: EBV



Das Telekommunikationsmodul Heracles: vereint ein komplettes Quad-Band GSM/GPRS-Modul und eine Prepaid-SIM-Karte.

ist, der die komplette Elektronik enthält, besteht ein System-in-Package aus einzelnen Chips, die in einem Package untergebracht sind. Dabei verfügen sie jeweils über spezifische Funktionalitäten. Das Ergebnis sind dreidimensionale Chips, die zu einer wesentlichen Platzersparnis und zu niedrigeren Montagekosten führen. System-in-Package-Lösungen eignen sich insbesondere für kundenspezifische Lösungen, die auch bei kleinen und mittleren Stückzahlen wirtschaftlich gefertigt werden können. Ein weiterer Vorteil ist, dass sich das Package ideal an die Anwendungsumgebung anpassen lässt. „SiPs bieten mehrere Vorteile ...“, so Santosh Kumar von dem Marktforschungsunternehmen Yole. „Dazu gehören die Reduzierung des Formfaktors, eine erhöhte Leistung, funktionale Integration mit Schutz vor elektromagnetischen Interferenzen, Designflexibilität im Vergleich zu SoCs und niedrigeren Kosten.“

Module aus vorkonfigurierten Baugruppen

Egal, ob SoC oder SiP – will ein Hersteller sein Produkt mit einer intelligenten Funktion ausstatten oder in das Internet der Dinge einbinden, benötigt er immer noch hochintegrierte Technologien, deren Eigenentwicklung für viele Unternehmen zu aufwendig, langwierig und teuer ist. Eine Lösung bieten hier komplett vorkonfigurierte Baugruppen. Sie verfügen nicht nur über die entsprechenden Halbleiterelemente, sondern über alle benötigten Komponenten, um eine bestimmte Funktion zu erfüllen.

Ein Beispiel dafür ist das Telekommunikationsmodul Heracles von EBV Elektronik. Es

vereint ein komplettes Quad-Band-GSM/GPRS-Modul und eine Prepaid-SIM-Karte. Es bietet Abdeckung und nahtlosen Zugang zum Mobilfunknetz von Orange sowie zu Tier-1-Roaming-Netzwerken in 33 europäischen Ländern. Durch die Vorintegration der Konnektivität in der elektronischen Designphase wird der Design- und Herstellungsprozess für die Objekthersteller stark vereinfacht. Die Lösung eignet sich bestens für alle Hersteller von IoT-Objekten – sei es für Automobil-, Tracking-, Mess-, Industrie- oder Wearables-Anwendungen.

Mit IT-Security und integriertem Mikroprozessor

Seit Ende 2020 ist das Modul zudem mit IT-Security und einem integrierten Mikroprozessor verfügbar, mit dem sich zum Beispiel schnell ein Sensorknoten im Narrow-Band-IoT beziehungsweise LTE-M realisieren lässt. Die Verfügbarkeit von sehr preisgünstigen Computern wie dem Raspberry Pi 4 in Form von Modulen ermöglicht es Herstellern, Start-ups und sogar softwareorientierten Unternehmen, ihre Prototypen und Proofs of Concept bis zur industriellen Produktion zu skalieren. Kunden können mit dem einsatzbereiten Raspberry Pi 4 Board-Prototypen erstellen und Software entwickeln. Anschließend können sie ihr eigenes System entwerfen, unter Verwendung des RPI-Computermoduls und aller für ihre spezifische Anwendung erforderlichen Peripheriegeräte, Kommunikations- und Schnittstellen. Das ist eine wirkliche Demokratisierung der Spitzentechnologie.

// MK

EBV



Hyperscale-Rechenzentren: einer der größten Treiber für einen höheren optischen Durchsatz.

Bild: Endrich

MEMS-Oszillatoren erweitern die Leistungsgrenzen optischer Module

5G wird die Kommunikationstechnik beflügeln. Einen hohen Stellenwert nehmen dabei optische Übertragungsmodule ein, die höchste Taktraten erfordern. Ein Fall für MEMS-Oszillatoren.

PARKER TRAWEEK *

Die Bereitstellung von 5G-Netzwerken wird enorme Fortschritte in der Kommunikation ermöglichen – eine 10-fach größere Bandbreite und eine 50-fache Reduzierung der Latenz. Um so massive Verbesserungen zu erzielen, werden verschiedene Technologien in rasantem Tempo weiterentwickelt, einschließlich Geräten und Komponenten, die in Rechenzentren verwendet werden. Ein Beispiel sind optische Transceiver, die für das Verbinden und Übersetzen von über Lichtwellenleiter übertragenen Daten in elektrische Signale innerhalb des Rechenzentrums verantwortlich sind.

Um den enormen Anstieg des Datenverkehrs zu bewältigen, verdoppeln sich die Übertragungsraten von optischen Modulen, beziehungsweise vervierfachen sich in einigen Fällen sogar. Heute werden üblicherweise Module mit Datenraten von 100 GBit/s verwendet. Der Einsatz von Modulen mit 400

GBit/s steigt jedoch rapide, und Varianten mit 800 GBit/s werden derzeit entwickelt. 400-GBit/s- und 800-GBit/s-Netzwerke mit höherer Kapazität stellen höhere Anforderungen an die optischen Module und die darin enthaltenen Oszillatoren.

Diese Geräte müssen eine größere Funktionalität mit dichterem Design, geringerer Leistung pro Bit und Jitter als ihre Vorgänger aufweisen.

5G erfordert große Datenmengen

Hyperscale-Rechenzentren sind einer der größten Treiber für einen höheren optischen Durchsatz. 5G erfordert die Übertragung und Berechnung großer Datenmengen. Um dies zu ermöglichen, müssen Rechenzentren optische Module mit höherer Kapazität verwenden. Hyperscale bezieht sich auf die vollständige Kombination von Hardware und Einrichtungen, mit denen sich eine verteilte Computing-Umgebung auf bis zu Tausende von Servern erweitern lässt. Bei Hyperscale

geht es darum, im Computing eine massive Skalierung zu erzielen – in der Regel für Big Data oder Cloud-Computing. Eine Hyperscale-Infrastruktur ist für horizontale Skalierbarkeit ausgelegt und erlaubt ein hohes Maß an Leistung, Durchsatz und Redundanz, das für Fehlertoleranz und Hochverfügbarkeit sorgt. Häufig werden beim Hyperscale-Computing massiv skalierbare Serverarchitekturen und virtuelle Netzwerke eingesetzt.

Die für den Betrieb von Rechenzentren erforderliche Energie ist enorm und deren Ausbau teuer. Einige Branchenexperten erwarten, dass Rechenzentren bis zum Jahr 2030 bis zu 8% des weltweiten Stromverbrauchs ausmachen. Von optischen Modulen wird erwartet, dass sie den Durchsatz mit wenig zusätzlichem Stromverbrauch erheblich verbessern. Rechenzentren erweitern neben anderen Datenkommunikationsanwendungen mit hoher Bandbreite die Grenzen der optischen Modultechnologie und stellen im weiteren Sinne höhere Anforderungen an die Oszillatortechnologie.

* Parker Traweck
... ist Product Marketing Engineer bei SiTime.

Die Rolle eines optischen Moduls besteht darin, eingehende optische Signale in elektrische Signale umzuwandeln und ausgehende elektrische Signale für den Transport ohne Fehler in das optische Format umzuwandeln. Dies wirft das komplexe Problem auf, die beiden Zeitbereiche – den des optischen Netzwerks und den des Chipsatzes auf der Hostplatine – zu synchronisieren. Dies macht das genaue Timing zu einem der kritischsten Faktoren innerhalb eines optischen Moduls. Die Komponente, die für die Überbrückung der Zeitlücke verantwortlich ist und daher als Re-Timer bezeichnet wird, erfordert einen Referenztakt, der mit zunehmender Datenrate von 100 auf 400 und 800 GBit/s einen zunehmend geringeren Jitter aufweisen muss.

Mit Beginn des Einsatzes von 400-Gigabit-Modulen wird der Phasenjitter des Referenzoszillators immer kritischer. RMS-Phasenjitter wird typischerweise durch Integrieren von Phasenrauschen über Offsetfrequenzen von 12 kHz bis 20 MHz berechnet. Der Differenzialoszillator SiT9501 von SiTime weist ein Phasenrauschen von -87 dBc/Hz bei einer Offsetfrequenz von 100 Hz und -170 dBc/Hz bei einer Offsetfrequenz von 400 MHz auf. Bei Integration führt das enge Phasenrauschen zu einem RMS-Phasenjitter von 70 Femtosekunden ($1 \text{ fs} = 10^{-15} \text{ s}$) bei einer Taktfrequenz von 156,25 MHz. Der Oszillator-RMS-Phasenjitter quantifiziert die Variation einer Taktflanke. RMS-Phasenjitter in Referenztakten, die optische Module ansteuern, ist besonders wichtig, da er den Jitter im seriellen Datenstrom, der durch das Modul fließt, verstärkt und Fehler verursachen kann, wenn dieser Jitter zu groß ist. Da sich der Durchsatz von 400 auf 800 GBit/s verdoppelt, sollte sich der Jitter im Signal proportional um den Faktor Zwei verringern, um eine ähnliche Zeitspanne beizubehalten.

Störimpulse im Phasenrauschen

Ein weiterer wichtiger Faktor, der bei der Berechnung des Phasenjitters berücksichtigt werden muss, ist das Vorhandensein von Störgeräuschen (Störimpulsen) im Phasenrauschen. Auf den ersten Blick scheint das Phasenrauschen zwischen einem SiT9501-MEMS-Oszillator und einem Quarz-PLL-basierenden Oszillator vergleichbar zu sein, doch bei näherer Betrachtung werden die „Spurs“ im Quarz-basierten Phasenregelkreis-Oszillator (PLL) deutlich (Bild 3). Das Phasenrauschen des SiT9501-Oszillators weist keine „Spurs“ auf, was zu einem RMS-Phasenjitter von nur 70 fs führt. Umgekehrt hat der Quarzoszillator einen Gesamt-RMS-

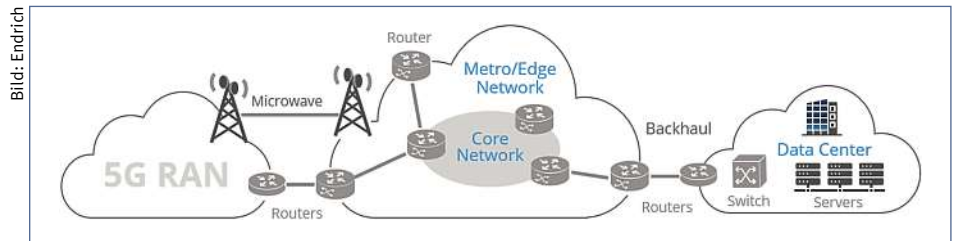


Bild 1: Optische Module werden an jedem Punkt des optischen Backbones – vom Fronthaul bis zum Backhaul – mit Transceivern mit hoher Datenrate verwendet, die in Metro-Netzen und Rechenzentren erforderlich sind (RAN = Radio Access Network).

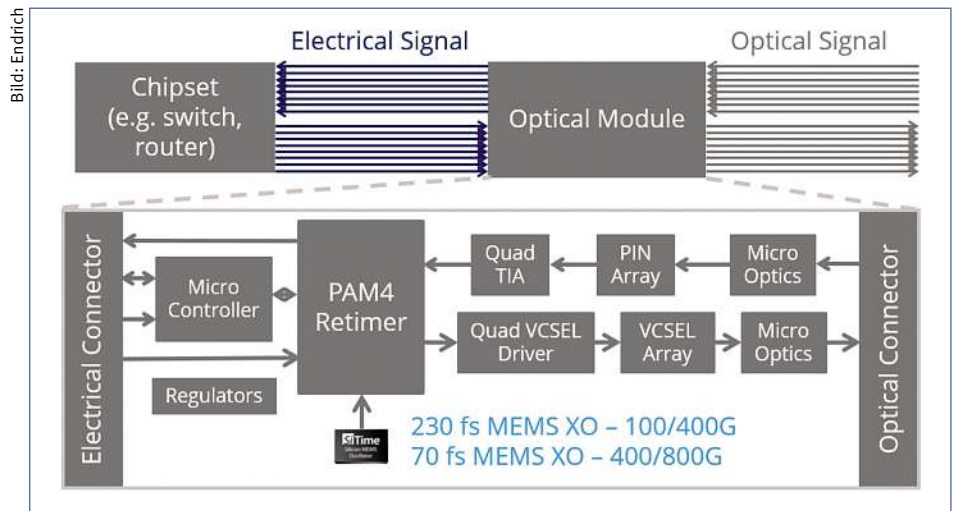


Bild 2: Blockdiagramm eines optischen Moduls mit einem SiTime-MEMS-Oszillator mit geringem Jitter, der den PAM4-Retimer taktet (PAM = Pulsamplitudenmodulation).

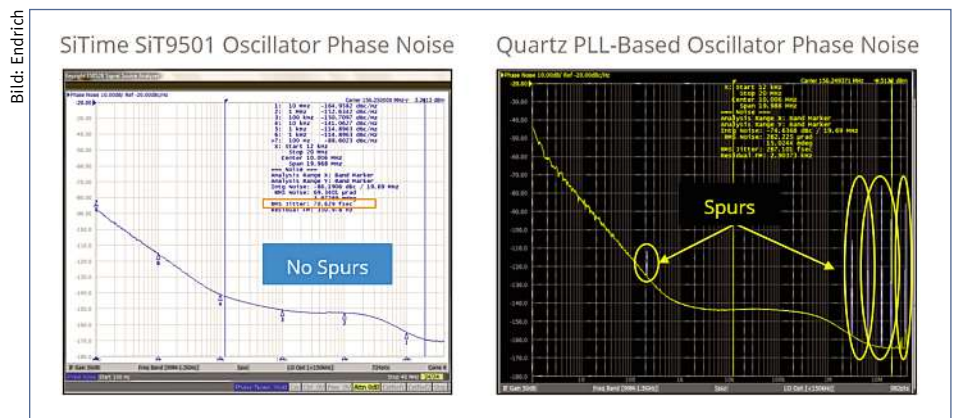


Bild 3: Vergleich des Phasenrauschens zwischen einem SiT9501-MEMS-Oszillator (RMS-Jitter: 70,629 fs; keine Störimpulse „Spurs“) und einem Quarz-PLL-basierenden Oszillator mit „Spurs“.

Phasenjitter von 267 fs. Ohne Berechnung der „Spurs“ beträgt der RMS-Phasenjitter des Quarzoszillators nur 90 fs, was bedeutet, dass die „Spurs“ 60% des gesamten Jitters ausmachen. Die fortschrittliche Integer-N-PLL-Technologie von SiTime ermöglicht dichtes Phasenrauschen und geringeren Jitter ohne „Spurs“. Da moderne optische Module die Datenraten um das Zwei- bis Vier-

fache erhöhen sollen, müssen die im Modul enthaltenen Komponenten diese Verbesserungen liefern, ohne ihren Platzbedarf zu erhöhen. Der SiT9501-Differenzialoszillator von SiTime ist die optimale Lösung für Designs von 400- und 800-GBit/s, da bei kleineren Größen mit nur 70 fs RMS-Phasenjitter keine Kompromisse bei der Leistung erforderlich sind. Darüber hinaus integriert der

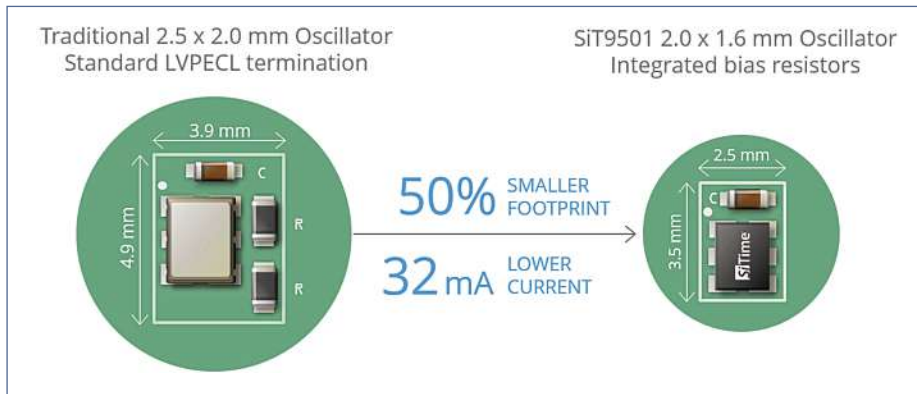


Bild 4: Vergleich der Grundfläche und Stromverbrauchs eines herkömmlichen AC-gekoppelten LVPECL-Layouts mit einem 2520-Oszillator (links) und des Layouts eines 2016-SiT9501-MEMS-Oszillators mit integrierten LVPECL-Source-Bias-Widerständen (rechts).

Oszillator SiT9501 (in einem Gehäuse mit 2,0 mm x 1,6 mm) Source-Bias-Widerstände, was zu einer Reduzierung des gesamten Platzbedarfs um 50% im Vergleich zu den derzeit meist verwendeten Quarzoszillatoren mit 2,5 mm x 2,0 mm führt. Der Oszillator SiT9501 integriert zudem On-Chip-Spannungsregler, die das Stromversorgungsrauschen filtern und die Leistungsintegrität bei Moduldesigns verbessern. Die Reduzierung des Timing-Footprints mit solchen integrierten Funktionen und der geringen Gehäusegröße

ist wichtig, da mehr als die Hälfte des optischen Moduls von der Laser-Baugruppe und der zugehörigen Elektronik verbraucht wird und nur wenig Platz für die Signalverarbeitung und den Datenpfad bleibt. Durch die Platzersparnis können Modulhersteller weitere Funktionen nutzen.

Um den strengen Strombegrenzungen für optische Module zu begegnen, führt das Entfernen der beiden Vorspannungswiderstände zu einem um 32 mA geringeren Stromverbrauch bei einem AC-gekoppelten Ausgang.

Mit dem SiT9501 wird auch die Flex-Swing-Technologie eingeführt, mit welcher der Differenzspannungshub werkseitig auf einzigartige Weise programmiert werden kann, um die Anforderungen an den Differential-Eingangshub eines beliebigen Chipsatzes zu erfüllen. Mit Flex-Swing können Ingenieure Niederspannungs-Chipsätze mit nicht standardmäßigen Spannungsschwankungen aufnehmen. Durch die Anpassung an die genauen Anforderungen des Chipsatzes kann der typische Abschluss beseitigt werden, wodurch der Strom mit einem DC-gekoppelten LVPECL-Ausgang um bis zu 16 mA reduziert wird. Die Entwicklung von optischen Modulen zu Datenraten von 400 und 800 Gbit/s, die von neuen Technologien angetrieben werden, erfordert Leistungssprünge ohne Erhöhung der Größe und des Stromverbrauchs. Dies wiederum treibt Oszillatoren an, energieeffizienter zu sein, weniger Platz zu verbrauchen und weniger Jitter zu erzeugen. Mit Innovationen wie integrierten Vorspannungswiderständen und programmierbarem Spannungshub reduziert der Differentialoszillator SiT9501 den gesamten Platzbedarf und den Stromverbrauch mit nur 70 fs RMS-Phasenjitter.

// MK

Endrich

**ELEKTRONIK
PRAXIS**

www.elektronikpraxis.de

ISSN 0344-1733

Kommunikationsdaten unserer Ansprechpartner:

E-Mail-Code: (bitte Schreibweise von Umlauten beachten): <vorname>.<name>@vogel.de; Telefon: +49-931-418-(4-stellige-Durchwahl)

Impressum

ABONNENTENSERVICE

DataM-Services GmbH

Franz-Horn-Straße 2, 97082 Würzburg
Tel. +49-931-41 70-4 62, Fax -4 94
vogel@datam-services.de, www.datam-services.de

REDAKTION

Leser-, Redaktionsservice:

Tel. +49-931-418-2333
fachmedien@vogel.de

Chefredakteur:

Johann Wiesböck (jw), Tel. -30 81

Redakteure:

Michael Eckstein (me), Tel. -30 96
Sebastian Gerstl (sg), Tel. -30 98
Hendrik Härter (heh), Tel. -30 92
Gerd Kucera (ku), Tel. -30 84
Thomas Kuther (tk), Tel. -30 85
Margit Kuther (mk), Tel. -30 99
Kristin Rinortner (kr), Tel. -30 86

Freie Mitarbeiter:

Anna-Lena Gutberlet (ag),
Richard Oed (ro)

Redaktionsanschrift:

München: Rablstr. 26, 81669 München, Tel. -30 87, Fax -30 93
Würzburg: Max-Planck-Str. 7/9, 97082 Würzburg
Tel. -24 77, Fax -27 40

Konzeption/Layout: Ltg. Daniel Grimm, Tel. -22 47

**ELEKTRONIKPRAXIS ist Organ des Fachverbandes
Elektronik-Design e.V. (FED). FED-Mitglieder erhalten
ELEKTRONIKPRAXIS im Rahmen ihrer Mitgliedschaft.**

Unternehmens- und Firmennamen:

Unternehmens- und Firmennamen schreiben wir gemäß Duden wie normale Substantive. So entfallen z.B. Großbuchstaben und Mittelinitialie in Firmennamen.

SALES

Chief Sales Officer:

Benjamin Wahler
Tel. -21 05, sales@vogel.de

Auftragsmanagement:

Tel. -20 78, auftragsmanagement@vogel.de

MARKETING

Produkt Marketing Manager:

Christian Jakob
Tel. -30 78, customer@vogel.de

VERTRIEB

Bezugspreis:

Einzelheft 12,90 EUR. Abonnement Inland: jährlich 249,00 EUR inkl. MwSt. Abonnement Ausland: jährlich 280,20 EUR (Luftpostzuschlag extra). Alle Abonnementpreise verstehen sich einschließlich Versandkosten (EG-Staaten ggf. +7% USt.).

Verbreitete Auflage:

Angeschlossen der Informationsgemeinschaft
zur Feststellung der Verbreitung von Werbeträgern –
Sicherung der Auflagenwahrheit.

Aktuelle Zahlen: www.iww.de

Datenbank:

Die Artikel dieses Heftes sind in elektronischer Form kostenpflichtig über die Wirtschaftsdatenbank GENIOS zu beziehen: www.genios.de



**VOGEL COMMUNICATIONS
GROUP**

Vogel Communications Group GmbH & Co. KG

Max-Planck-Str. 7/9 in 97082 Würzburg
Tel.: 0931/418-0, www.vogel.de

Beteiligungsverhältnisse:

Persönlich haftende Gesellschafterin:
Vogel Communications Group Verwaltungs GmbH
Max-Planck-Str. 7/9 in 97082 Würzburg
Kommanditisten:
Dr. Kurt Eckernkamp, Dr. Kurt Eckernkamp GmbH,
Nina Eckernkamp, Klaus-Ulrich von Wangenheim,
Heiko Lindner, Axel von Kaphengst

Geschäftsführung:

Matthias Bauer (Vorsitz)
Günter Schürger

Druck:

Vogel Druck und Medienservice GmbH
97204 Höchberg

Copyright:

Vogel Communications Group GmbH & Co. KG

Nachdruck und elektronische Nutzung:

Wenn Sie Beiträge dieser Zeitschrift für eigene Veröffentlichungen wie Sonderdrucke, Websites, sonstige elektronische Medien oder Kundenzeitschriften nutzen möchten, fordern Sie gerne Informationen über support.vogel.de an.

Durchstarten 2021 – gemeinsam aus der Krise

In dieser Interview-Reihe geben unsere Leserinnen und Leser Einblicke in die Herausforderungen und Chancen der Corona-Pandemie in ihrem Unternehmen und verraten, was sie aus dem Krisenjahr 2020 gelernt haben. Lassen Sie uns im neuen Jahr 2021 gemeinsam durchstarten!

Johann Wiesböck: Wir haben mit Ralf Welter als Leser von ELEKTRONIKPRAXIS über die Situation im Unternehmen und seine persönlichen Erkenntnisse aus der Coronakrise gesprochen. Was ihm wichtig ist, lesen Sie online – siehe Link.



Das ganze Interview können Sie unter
www.elektronikpraxis.de/durchstarten
nachlesen.



Ralf Welter: Der Geschäftsführer von Taiwan Semiconductor Europe leitet seit 2017 die europäische Zentrale in Zorneding bei München mit aktuell 25 Mitarbeitern. Seine Zielsetzung ist es, mit allen Mitarbeiter*innen mindestens einmal die Woche zu sprechen – und nicht nur über berufliche Dinge. Auch regelmäßige Telefonate mit Kunden sind ihm sehr wichtig. Aufgefallen ist ihm, dass viele Ansprechpartner*innen sich heute sogar etwas mehr Zeit für ein Gespräch nehmen und die Ergebnisse zum Teil auch besser sind als noch in der Zeit vor Corona.

So fertigen wir Elektronik in Lockdown-Zeiten



Wie gelingt der Spagat zwischen Einhaltung der Corona-Regeln und Schutz der Mitarbeiter bei höchst möglicher Produktivität? Mit Verantwortung auf allen Ebenen und steten Änderungen im Ablauf nach Plan.

Gerd Ohl, Geschäftsführung Limtronik: „Wir haben keine Umbesetzungen oder Freistellungen vorgenommen. Statt dessen haben wir sogar Neubesetzungen auf den Weg gebracht. Ebenfalls haben wir mit der Auswahl von neuen Auszubildenden für den Sommer begonnen, um uns zukunftsorientiert aufzustellen.“

In vielen Branchen ist Home-Office-Arbeit in Corona- bzw. Lockdown-Zeiten das bevorzugte Modell. Dieses Modell zu realisieren ist jedoch in der Fertigung nicht möglich. Deshalb haben wir uns bereits im Januar 2020 mit der Thematik Corona auseinandergesetzt und sehr früh entsprechende Maßnahmen eingeleitet, damit Covid-19 weder unsere Produktions- und Lieferprozesse noch die Gesundheit unserer Mitarbeiter gefährdet. Über das Jahr hinweg mussten selbstverständlich stetig Anpassungen und Optimierungen vorgenommen werden. Das war teils ein harter Weg – und der ist noch lange nicht zu Ende.

Limtronik entwickelt als Experte für Electronic Manufacturing Services und Joint-Development-Manufacturing-Partner spezifische Lösungen im Kundenauftrag. Dazu begleiten wir die Kunden von der Produktentwicklung bis zum serienreifen Endprodukt. Die ersten Herausforderungen im Zuge der Pandemie begannen Anfang 2020 mit anfänglichen Lieferengpässen für Bauteile. Die Lieferungen aus Asien erfolgten recht schnell wieder zu fast normalen Lieferzeiten. Doch ein Nebeneffekt ist allerdings, dass die Transportkosten angestiegen sind. Eine weitere Herausforderung stellt sich jetzt, weil unsere Lieferanten Corona-bedingte Preisanpassungen adressiert haben. Schwieriger geworden war und ist auch die Zusammenarbeit mit europäischen Lieferanten, da diese auf Grund massiver Umsatzeinbrüche ihre Kapazitäten anpassen mussten. Um aber die Produktivität weiterhin auf unserem gewohnten Niveau zu halten, haben wir zahlreiche innerbetriebliche Veränderungsprozesse sofort umgesetzt.

Mit Ausbruch des Coronavirus in Deutschland hat Limtronik einen Pandemieplan erstellt, der auf den Empfehlungen des „Deutscher Industrie- und Handelskammertag (DIHK)“ basiert. Dieser wird konsequent umgesetzt, stetig auf Aktualität geprüft und entsprechend ergänzt. Zu den Vorkehrungen zählen unter anderem die

Einhaltung der selbstverständlichen AHA-Regeln und diverse Desinfektionsmaßnahmen. Beispielsweise arbeitet das Team von Limtronik seit März 2020 mit Pausen zwischen den Schichten, damit eine Desinfizierung der Arbeitsplätze gewährleistet werden kann. Des Weiteren wird die Einhaltung der AHA-Regeln an den Arbeitsplätzen auch kontrolliert, Pausen werden versetzt durchgeführt und Masken zur Verfügung gestellt. Auch Schnelltests haben wir eingeführt.

Wir versuchen weiterhin alles, um unsere Produktivität auch unter diesen erschwerten Bedingungen auf möglichst hohem Niveau zu halten. Allerdings ist von Produktivitätseinbußen durch die neuen Pausenzeiten auszugehen. Der Anlauf der Maschinen nach einer Pause ist dabei immer der entscheidende Faktor. Neben den uns bereits bekannten Faktoren ergeben sich im weiteren Verlauf immer neue Anforderungen, die eine hohe Flexibilität und stete Anpassungsprozesse von uns fordern.

Eine neue Herausforderung für uns als Elektronikfertiger stellt aktuell die Regelung der Ausgangssperren für bestimmte Landkreise dar. Da Limtronik auch systemrelevante Komponenten baut, erstellte das Unternehmen entsprechende Nachweise der Arbeitsbedingungen. So möchten wir Mitarbeiter gegen Geldbußen oder Ähnlichem absichern, wenn diese sich um 22:15 Uhr nach der Schicht auf dem Nachhauseweg befinden.

Generell sind sich alle Beteiligten im Hause ihrer Verantwortung bewusst. Bisher gab es zwar einige Verdachtsfälle, aber nur eine einzige Corona-Infektion. Dies ist unter anderem unseren durchdachten Maßnahmen zu verdanken. Zu Jahresbeginn, nach den Feiertagen, testeten wir jeden Mitarbeiter vor dem Betreten des Betriebes per Schnelltest. Dies wurde einvernehmlich von der Belegschaft aufgenommen. So können wir sofort auf mögliche Ansteckungen reagieren und einer Infektionskette entgegenwirken. // KU

05. – 07. Juli 2021 | VCC | Würzburg

Anwenderkongress **Steck**verbinder



**SAVE
THE DATE!**

Praxisorientierte Lösungen für den Einsatz und das Design moderner Steckverbinder

Europas größter Fachkongress zum Thema ist der Pflichttermin für alle, die Steckverbinder entwickeln oder einsetzen und interessante Kontakte knüpfen möchten. Das erwartet Sie: Referenten aus Industrie und Forschung, praxisorientierte Lösungen, Workshops, Grundlagenseminare und Networking mit Experten aus der Industrie.

www.steckverbinderkongress.de

ONOFFONOFFONOFFONOFFONOFFONOFFON
 OFFONOFFONOFFONOFFONOFFONOFFONOFF
 ONOFFONOFFONOFFONOFFONOFFONOFFON
 OFFONOFFONOFFONOFFONOFFONOFFONOFF
 ONOFFONOFFONOFFONOFFONOFFONOFFON
 OFFONOFFONOFFONOFFONOFFONOFFONOFF
 ONOFFONOFFONOFFONOFFONOFFONOFFON
 OFFONOFFONOFFONOFFONOFFONOFFONOFF
 ONOFFONOFFONOFFONOFFONOFFONOFFON
 OFFONOFFONOFFONOFFONOFFONOFFONOFF
 ONOFFONOFFONOFFONOFFONOFFONOFFON

WE are here for you!

Nehmen Sie teil an unseren kostenlosen

Webinaren: www.we-online.de/webinare

Kurzhubtaster

Die Kurzhubtaster von Würth Elektronik zeichnen sich durch Zuverlässigkeit und eine hohe Nutzungsdauer aus. Alle metallischen Elemente werden auf die Korrosionsresistenz mit einem 48-stündigen Salz-Nebel-Test geprüft. Polyimid-Folie oder Silikoneinsätze schützen den Kurzhubtaster auch bei extremen Einsatzbedingungen. Ab Lager verfügbar. Kostenlose Muster erhältlich.

Weitere Informationen unter: www.we-online.de/switch

- IP67
- Hohe Lebensdauer
- Hohe Beständigkeit
- Korrosionsbeständige Metallkomponenten
- Test der Schaltzyklen bei Volllast
- Geringe Toleranzen durch automatisierte Produktion



SMT Vertical



SMT Side Push



THT Vertical



THT Right Angled



IP67



Illuminated Solution