

Data Center

Das RZ-Netz der Zukunft

Teil 2: 25/50/100 GbE-Switching und die 100G-Transformation

von Dr. Franz-Joachim Kauffels

Nach den einführenden Überlegungen im Teil 1 kommen wir heute zunächst zur Technik. Hier ist 100 GbE und die sich daraus „natürlich“ ergebenden Datenraten 25 und 50 GbE für die nächsten Jahre das Maß aller Dinge in den meisten Rechenzentren, ob privat betrieben oder Hyperscaler. Die Technik ist schon deutlich anders als bei den 10/40 GbE-Vorgängern und man sollte sie verstehen, bevor man sie bewertet. Kern sind die 100G Switch-ASICs, die in den aktuellen Modellen führender Hersteller verbaut werden.

Eine andere Frage, die immer wieder auftaucht, ist, ob man sich nicht einfach Bare



Metal Switches mit diesen ASICs kaufen und sich somit von den etablierten Herstellern verabschieden kann. Am Beispiel von LinkedIn zeigen wir, dass das geht, aber damit auch ein Aufwand verbunden ist, den „normale“ Betreiber wohl eher weniger leisten wollen. Extrem spannend ist aber der generelle Netzaufbau des LinkedIn-RZs. Strukturell kann er durchaus auch auf deutlich kleinere Umgebungen übertragen werden und zeigt nach Ansicht des Autors, wie moderne RZs in Zukunft aussehen können.

weiter auf Seite 8

IT-Endgeräte

Internet of Things – die vierte industrielle Revolution Teil 4

von Dipl.-Inform. Petra Borowka-Gatzweiler

Dieser Teil der Serie "Internet of Things" beschreibt die aktuell abzusehende IoT Roadmap und aktuelle Trends im Internet of Things Markt.

5. Internet of Things: Roadmap und Trends

Während in der ersten Dekade dieses Jahrhunderts schnellere Logistik und Kos-

teneinsparungen in den Bereichen Überwachung, Sicherheit, Gesundheitswesen, Transport und Lebensmittelsicherheit Treiber für den IoT-Markt waren, wendet sich die zweite Dekade 2010 bis 2020 der globalen Lokalisierung von Menschen und Dingen über Geosignale, der remote Steuerung vernetzter Objekte über effiziente Elektronik sowie Software-Agenten in gro-

ßen zusammengeschlossenen Sensor-Netzen zu, wie die Übersicht in Abbildung 5.1 zeigt. Ein einzelnes solches Bot-Netz kann 50.000 und mehr Objekte umfassen.

weiter auf Seite 19

Geleit

Aktuelle Diskussionen im ComConsult Research Team: Netzwerk-Sicherheit und VXlan

auf Seite 2

Standpunkt

Ein Zertifikat für WLAN Clients?

auf Seite 18

Aktueller Kongress

ComConsult Netzwerk-Forum 2017

ab Seite 4

Aktuelle Sonderveranstaltung

Implementierung von IPv6 – Erkenntnisse und Erfahrungen

auf Seite 17

Geleit

Aktuelle Diskussionen im ComConsult Research Team: Netzwerk-Sicherheit und VXlan

Unsere in diesem Jahr recht intensiven Vorbereitungsarbeiten für das Netzwerk-Forum neigen sich dem Ende zu. Wie immer werden wir die Ergebnisse exklusiv auf der Veranstaltung vorstellen und mit den Teilnehmern diskutieren.

Aus den laufenden Diskussionen möchte ich nur zwei herausgreifen, die sich beide mit extrem wichtigen Themen befassen:

- Netzwerk-Sicherheit
- die Umsetzung von VXlan

Die aktuellen Diskussionen um Wahl-Fälschung in den USA und wie diese eventuell umgesetzt wurde - sofern sie denn stattgefunden hat - machen deutlich, dass wir weiterhin eine Verschiebung der Angriffe in Richtung auf das Endgerät bekommen. Hier ist der Einfallspunkt, über den die Angreifer agieren. Gleichzeitig machen Mikroservice-Architekturen Angriffe auf den Server immer unwahrscheinlicher. Der Angreifer würde ein enormes Wissen über die Architektur benötigen, die gleichzeitig eine dynamische ist, die sich permanent wandelt. Anders formuliert: selbst wenn ein Angreifer den Hypervisor übernehmen würde, ist es zweifelhaft, ob er mit der Information, die er dann hat, die Architektur der Applikation soweit versteht, dass sie ihm von Nutzen ist. Das ist einer der Gründe, warum wir eine deutliche Zunahme der Angriffe auf Endgeräte verzeichnen. Das Endgerät sieht Daten in Architektur-neutraler Form. Man muss die Architektur weder kennen noch verstehen und hat den Zugriff auf genau das, was spannend ist: die Daten des Unternehmens.

Mit der Konzentration auf das Endgerät als den Haupt-Risiko-Faktor müssen wir demzufolge mindestens zwei Fragen klären:

- wie können wir den Zugang eines Angreifers zum Endgerät besser kontrollieren? Immerhin haben wir eine zunehmende Zahl von mobilen Endgeräten, die sich statisch und vor allem außerhalb des Unternehmens deutlich schwerer absichern lassen. Speziell zu diesem Punkt haben wir einen Vortrag und eine Diskussion zum Forum aufgesetzt.
- wie können wir ein Endgerät, das übernommen wurde, in seinem Schadens-



Potenzial reduzieren? Hier kommen unsere schon seit langen diskutierten Zonen-Architekturen wieder ins Spiel. Wir müssen jemanden, der eingedrungen ist, in seinen Möglichkeiten einschränken bzw. durch seine Aktivitäten identifizieren.

Beide Fragen müssen wir, um ein Zukunfts-sicheres Fundament zu haben, mit der Herausforderung der Nutzung von Cloud-Applikationen kombinieren. Selbst wenn mobile Endgeräte lokal im Unternehmen sind, verlassen sie das Unternehmen, um auf eine Cloud-Anwendung in der Public-Cloud zuzugreifen. Wie können wir das besser kontrollieren? Und was machen wir, wenn das mobile Endgerät gerade im Hotel ist und von dort auf eine Cloud-Applikation zugreift? Dann ist keine einzige der im Unternehmen aufgebauten Zonen ein Teil der Kommunikation (je nach Art der Gestaltung).

Das sind Szenarien, die wir mit Ihnen diskutieren wollen.

Das Thema Netzwerk-Sicherheit ist momentan so aktuell, dass wir ihm den Vertiefungstag auf dem Forum gegeben haben. Hier gehen wir in die Details und diskutieren, welche aktuellen Möglichkeiten es im Moment gibt.

Das andere Thema, das ich an dieser Stelle ansprechen möchte, ist VXlan.

Der Einsatz von VXlan ist nach Jahren der Diskussion momentan ein sehr aktuelles

Thema auch in großen Projekten. Nachdem die Hardware der meisten Switches das Verfahren jetzt mit der aktuellen Chip-generation auch auf Commodity-Switches unterstützt, erlaubt das stark verbesserte Preis-Leistungsverhältnis den wirtschaftlichen Einsatz. Und die nutzbaren Varianten sind keineswegs auf reine VMware-Umgebungen beschränkt.

Warum ist das Thema so wichtig?

Im RZ haben wir mehrere führende Trends, die sich nicht voneinander trennen lassen:

- die weiterhin wachsende Zahl von Kernen in den CPUs gibt mehr Potenzial für noch mehr virtuelle Maschinen pro CPU bzw. pro Server (unbestätigtes Gerücht: 28 Kerne im neuen E5, das erzeugt auf einem 4-Sockel-System mehr als 100 Kerne)
- Mikroservice-Architekturen, also die Zerlegung großer monolithischer Software-Pakete in viele kleine mobile Service-Einheiten, erhöhen die Anzahl von VMs und verkleinern sie gleichzeitig. Im Ergebnis heißt das: ein neues Ökosystem von vielen miteinander kommunizierenden VMs und noch mehr VMs pro CPU. Wir erhalten damit extrem große Netzwerke innerhalb der Hypervisor
- Container-Technologien erhöhen die Portabilität weiterhin und erlauben eine Plattform-neutrale Gestaltung. Sie werden die Basis für viele Umsetzungen von Virtualität der Zukunft sein
- die Bedeutung der Private Cloud wird weiter zunehmen: die schnelle Bereitstellung neuer Kapazitäten und Lösungen wird nur auf dieser Basis umsetzbar sein
- die Public Cloud wird ebenfalls für alle Unternehmen unvermeidbar sein. Sie hat deutliche Vorteile im Bereich von Anwendungen für externe Nutzer. Ein Teil dieser Vorteile im Sinne der nutzbaren Dienste wird einfach lokal nicht existieren

Was bedeutet das für Netzwerke? Ganz einfach: hier entsteht ein massives virtuelles Ökosystem mit einer eigenen virtu-

Das RZ-Netz der Zukunft

Teil 2: 25/50/100 GbE-Switching und die 100G-Transformation

Fortsetzung von Seite 1



Dr. Franz-Joachim Kauffels ist Technologie- und Industrie-Analyst und Autor. Seit über 30 Jahren unabhängiger, kritischer und oft unbequemer Bestandteil der Netzwerkszene. Verfasser von über 20 Büchern in über 70 Ausgaben sowie über 2000 Artikeln, Videos und Reports.

3. 25/50/100 GbE-Switching

Im RZ ist die Gemengelage bei den 40 GbE-Anschlüssen unübersichtlich, optimal und problemangepasst ist hier gar nichts, man kann auch kurz konstatieren, dass 40 GbE genau so tot ist, wie wir es immer prognostiziert haben. Die aktuelle Generation höchst leistungsfähiger 25/50/100 GbE Switch-ASICs bietet offensichtlich die Möglichkeit zum Aufbau zukunftssicherer, skalierbarer, leistungsfähiger und gut zu verwaltender RZ-Netze jeder Größenordnung.

Dennoch gibt es offensichtlich vielfach Unklarheiten in der Diskussion. Die Anzahl der Hersteller der eigentlichen Switch-Chips für 100 G (und damit in Folge auch für die flexiblen 10/25/50/100 GbE-Lösungen) ist sehr übersichtlich. Damit stellen sich viele folgende Frage: „Wenn ohnehin in allen Switches das Gleiche drin ist und die Hersteller sich nur noch über die Software differenzieren, kann ich dann nicht einfach den billigsten Switch nehmen?“

Die Antwort: NEIN. Das zugrundeliegende Missverständnis der Leute, die so fragen ist, dass 25/50/100G-Switch Chips intern genau so konstruiert sind wie ihre 10/40-Vorgänger. Das genau ist aber nicht so. 25/50/100 G-Switches haben eine Switching Matrix im Kern und darum herum eine Vielzahl möglicher Schnittstellenmodule, Speichererweiterungen und Prozessoren. Switch-Hersteller wie Cisco, Arista und andere können sich sehr wohl auch in der Hardware Differenzierungsmerkmale schaffen. Schließlich gibt es durch das Hersteller-Duo Mellanox und Broadcom noch zwei grundsätzliche verschiedene Bauarten der eigentlichen Switch-Matrix.

Der Sinn dieses Unterkapitels ist es, diese Dinge genau, aber dennoch hoffentlich auch für jemanden, der kein Chip-Nerd ist, nachvollziehbar, darzustellen.

Auf fast schon wundersame Weise hat es Ethernet geschafft, sich mit bezogen auf den Zeitraum eigentlich recht wenigen gravierenden Änderungen (wie z.B. der Entwicklung vom Shared Medium zum Switching) eher in den Stand einer Commodity, eines grundsätzlichen Versorgungssystems, hochzuarbeiten. Systeme zur Wasser- und Stromversorgung haben trotz aller zwischenzeitlichen technischen Neuerungen den erheblichen Charme, dem Nutzer gegenüber immer in der gleichen Form zu erscheinen, lediglich mit Unterschieden in der Leistung.

Und das ist auch das Geheimnis von Ethernet. Nicht die rohen Datenraten oder die teilweise verwegenen Zusatzprotokolle wie FCoE, sondern schlicht und ergreifend das Paketformat und die grundsätzliche Verarbeitung sind der Kern des Erfolges.

In den vergangenen Jahrzehnten hatten wir uns daran gewöhnt, dass es ab und an einen Leistungssprung um den Faktor 10 zu grob den dreifachen Kosten der aktuell bestehenden Leistungsstufe gab, also von 10 MbE auf Fast Ethernet, von Fast auf Gigabit Ethernet und von Gigabit auf 10 Gigabit Ethernet. Durch den Standard zu 40 und 100 GbE gab es in diesem System eine Unterbrechung. Anlass war, dass die Provider vor einigen Jahren gerne eine 40G-Technologie haben wollten und 100G ohnehin technisch noch außer Reichweite war. An privaten RZ-Netzen ging diese Diskussion ebenso vorbei wie die Weiterentwicklung der optischen Übertragungstechnik am Standard selbst.

Bevor wir aber auf die einzelnen Alternativen eingehen, möchte ich einige Bemerkungen zur Struktur moderner Switch-ASICs machen, die für das Verständnis der Entwicklungen wesentlich sind.

3.1 Das Multi-Lane Konzept als Grundlage

Für die Konstruktion von Switch-ASICs, aber auch für die Implementierung vieler begleitender Funktionen wie Sicherheitsprüfungen oder Zähler oder Flow-Funktionen ist das Multi-Lane Konzept eine lebenswichtige Grundlage. Wir kennen seit Jahrzehnten Moore's Law, das besagt, dass sich die Anzahl der auf einer Chipfläche verfügbaren Transistorfunktionen alle 18 bis 24 Monate verdoppelt. Darauf können wir uns auch in den nächsten Jahren verlassen. Allerdings steht nirgendwo, dass die immer kleiner werdenden Transistoren immer schneller werden. Es gibt unterschiedliche VLSI-Herstellungsprozesse. Der preisgünstige CMOS-Prozess, der seit vielen Jahren die Basis für alle Entwicklungen ist, führt zu Schaltungen, die sich mit höchstens rund 3 GHz takten lassen. Natürlich gibt es auch erheblich schnellere Techniken, die aber alle deutlich geringere Integrationsgrade und höhere Kosten nach sich ziehen. Deren Einsatz wird auf das absolut notwendige Minimum beschränkt.

Bei Prozessoren führt das einfach zu mehr Cores über die Zeit und ohne ein Betriebskonzept wie die Virtualisierung hätten alle Prozessor-Hersteller große Probleme, die vielen Cores noch sinnvoll zu nutzen. Bei Speicher ist es einfach nett, immer mehr Kapazität auf immer kleinere Flächen zu bringen, wobei jetzt aktuell ja auch die dritte Dimension erklimmt wird, was zu weiteren Optimierungen führt.

Das RZ-Netz der Zukunft - Teil 2: 25/50/100 GbE-Switching und die 100G-Transformation

Aber seit dem Schritt von 1 GbE auf 10 GbE führen die beschriebenen Fakten zur Notwendigkeit, den Datenstrom zu Beginn eines Schaltkreises auseinander zu nehmen und am Ende wieder zusammen zu setzen, denn für die meisten Aufgaben einer Schaltung, die den Datenstrom manipulieren soll, muss anschaulich gesprochen die Schaltung schneller sein als der Datenstrom. Bei 10 GbE hat das dazu geführt, dass vier Lanes zu je 2,5 Gbit/s vorgesehen wurden, in die der Datenstrom zur Bearbeitung zerlegt wurde. Damit konnte man Schaltkreise verwenden, die im Bereich zwischen 2,6 ... 3,0 GHz getaktet wurden, wie es der preisgünstige CMOS-Prozess verlangt.

Im Standard IEEE 802.3 ba für 40 und 100 GbE wurde dieses Konzept verallgemeinert. Nach außen sehen wir z.B. bei 40(100) GBASE-SR die Möglichkeit, 40 GbE auf vier und 100 GbE auf 10 Fasern zu implementieren, nach innen gibt es aber weitere Differenzierungen mit dem PCS (Physical Coding Sublayer) Distribution Konzept, welches einen Strom aus 64/66b-Worten durch Einfügung von Markern in praktisch beliebig viele Lanes zerlegen kann.

Hat man jetzt in einem Schaltkreis durch die konsequente Anwendung eines Bibliothekskonzeptes mit Modulen hinreichend viele parallele Schaltungen ist es überhaupt kein Problem, z.B. einen 40 GbE Datenstrom in 16 Lanes á 2,5 Gbps zu zerlegen und diese Lanes dann in den entsprechenden Modulen weiter zu verarbeiten.

Die Einrichtungen zur Zerlegung und zum Wiederaufbau von Datenströmen heißen SerDes (Serialisierer/Dserialisierer). Bei der aktuellen Generation von Switch-ASICs kommt dann noch praktischerweise hinzu, dass diese mit Speicheranpassung arbeiten. Für die Implementierung eines Switching-Vorgangs muss man nur Speicherbereiche definieren, in die die Lanes die Daten parallel ablegen können bzw. aus ihnen wiederaufnehmen. Der eigentliche Switching-Vorgang bewegt die Daten ja gar nicht, sondern ordnet nur den durch die Menge der zu einem Input-Port gehörenden Speicherzellen einen passenden Output-Port zu, das ist aber keine Manipulation mit Daten, sondern mit Adressen der Speicherzellen.

Es ist vielleicht dem einen oder anderen Leser aufgefallen, dass die Generation der 10/40 G-Ethernet-Switches im Großen und Ganzen die gleiche L2-Switch-Latenz hat wie ihre 10G-Vorgänger.

Solange nur geschaltet wird, ist es gleichgültig, ob 4 Lanes (für 10G) oder 16 Lanes (für 40G) parallel bearbeitet werden. Unterschiede gibt es natürlich bei der Bearbeitung von L3 oder noch höheren Funktionen, das hängt dann davon ab, wie viel Rechenleistung in Form paralleler Prozessoren im Switch Chip steckt.

Die aktuelle Switch-Generation auf der Basis von 40 G Switch-ASICs wie Broadcom Trident II ® oder Mellanox Switch-X® ist deshalb auch großzügig hinsichtlich der Konfiguration. Ein Switch mit z.B. 48 40 GbE-Ports kann im Extrem auch so konfiguriert werden, dass er 192 10 GbE Ports hat. Und es kann sinnvolle Mischungen geben, wie z.B. 16 40 GbE Ports und 128 10 GbE Ports. Natürlich beherrschen die Switches dann Rate Conversion.

3.2 100 G Switching

Wie schon erwähnt, legen auch die Provider keinen großen Wert mehr auf 40 G. Für den Ausbau der Netze, vor allem vor dem Hintergrund von 5G, müssen stärkere Mittel her. Also gibt es schon seit einiger Zeit von führenden Herstellern passende 100 G-Switches für Provider, wie z.B. Cisco CRX-1, Juniper T-, MX-, EX-Reihen oder Arista 7500 E, letzterer ist durchaus auch für die Anwendung in Rechenzentren gedacht. In 2016 sind dann

sehr sinnvolle Alternativen für 100G-Switches für das RZ hinzugetreten.

Sieht man aber genauer hin, ist das Design der Switches völlig anders als bei den monolithischen 40 G Switch ASICs. Aktuelle 100 G Switches arbeiten mit einem Kern für das Switching und Ingress bzw. Egress Modulen, die den Kern umgeben und all die Funktionen bereitstellen, die für die Manipulation der Datenpakete benutzt werden. Das hat den enormen Vorteil, dass man den Kern auf maximale L2-Leistung optimieren kann und bei der Gestaltung der Randmodule eine sehr hohe Flexibilität walten lassen kann.

Abbildung 3.1 zeigt das Basisdesign, in diesem Fall mit einem Clos-Netz als Kern. Das Clos-Netz hat ja den Vorteil, dass man ein einmal bestehendes Switching-Basis-Element rekursiv wiederverwenden kann, wobei die Stufenzahl nur logarithmisch steigt.

Das Design geht zurück auf Arbeiten der Firma Dune, die von Broadcom gekauft wurde. Die wesentliche Leistungssteigerung eines Mehrstufen-Mehrfach-Verbindungsnetzwerkes nach Clos geschieht durch dynamisches Routing, welches es ermöglicht, einen Datenstrom auch in parallelen Wegen durch den Switching-Kern zu führen, siehe Abbildung 3.2. Das har-

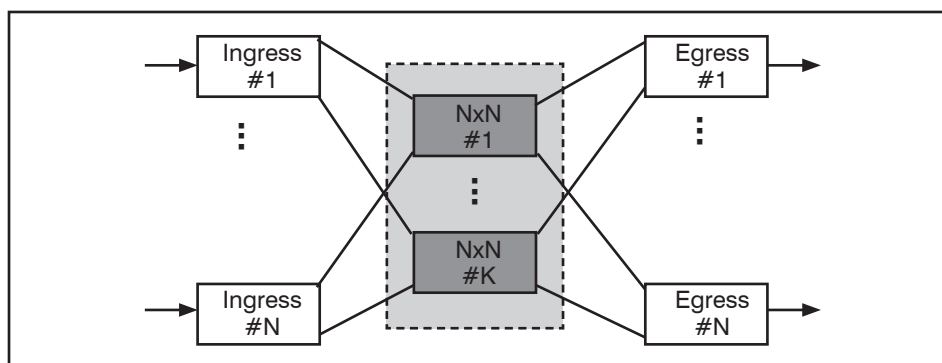


Abbildung 3.1: 100G-Switch Basis-Struktur

Quelle: Broadcom

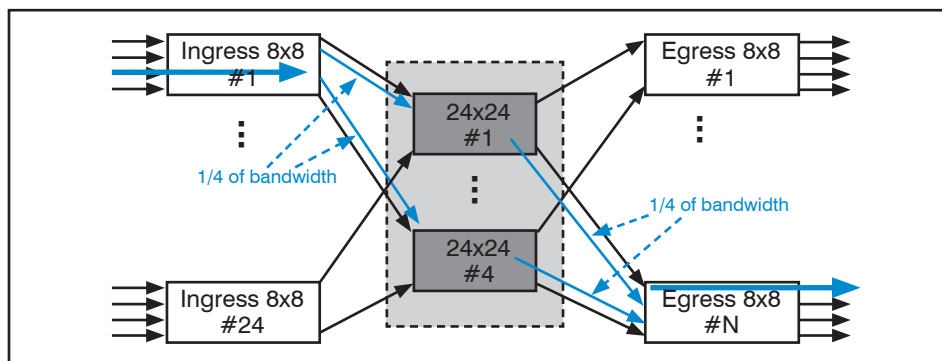


Abbildung 3.2: DUNE Clos Netzwerk mit dynamischem Routing

Quelle: Broadcom

Endgeräte

Internet of Things – die vierte industrielle Revolution Teil 4

Fortsetzung von Seite 1



Dipl.-Inform. Petra Borowka-Gatzweiler leitet das Planungsbüro UBN und gehört zu den führenden deutschen Beraterinnen für Kommunikationstechnik. Sie verfügt über langjährige erfolgreiche Praxiserfahrung bei der Planung und Realisierung von Netzwerk-Lösungen und ist seit vielen Jahren Referentin der ComConsult Akademie. Ihre Kenntnisse, internationale Veröffentlichungen, Arbeiten und Praxisorientierung sowie herstellerunabhängige Position sind international anerkannt.

Hatten wir in den 90er Jahren noch etwa eine Million vernetzter "Dinge" (siehe die Übersicht in Abbildung 5.2), so sind für 2017 schon 20 bis 22 Milliarden vernetzte Objekte prognostiziert (Quelle: World Economic Forum 2015, IHS Markit 2017). Bis 2020 soll diese Anzahl auf 40 bis 50 Milliarden ansteigen, das ist innerhalb von 3 Jahren mehr als eine Verdoppelung. Digitalisierung in der industriellen Fertigung, Roboter, Smart Cities, Smart Homing und Smart Devices am menschlichen Körper im so genannten BAN (Uhren, Armbänder, Multimedia-Funktionskleidung) sind riesige Märkte.

Die größte Gruppe bilden Consumer Geräte mit 8 Milliarden und einem hohen Wachstum von geschätzten 16 Prozent CAGR zwischen 2015 und 2025 (Quelle: IHS Markit 2017). Sie werden gefolgt von 6 Milliarden Kommunikations-Geräten mit immerhin noch 8,5 Prozent CAGR. An dritter Stelle stehen industrielle Objekte mit 3,6 Milliarden; sie weisen jedoch mit einem CAGR von 27,8 Prozent das höchste Wachstum auf. Während vernetzte Computer zwar immer noch die Position 4 innerhalb der vernetzten Objektkategorien einnehmen, sind sie mit -2 Prozent CAGR erkennbar auf dem abstei-

genden Ast. Medizintechnik und Automotive folgen mit 319 Millionen und 202 Millionen vernetzten Geräten in sehr großem Abstand, zeigen aber mit 17,8 Prozent und 22 Prozent CAGR ebenfalls ein sehr hohes Wachstum. Den kleinsten Marktanteil haben Militär und Luftfahrt mit 5 Millionen vernetzten Geräten und 12,9 Prozent CAGR. (siehe hierzu Abbildung 5.3)

Im September 2016 hat das Weiße Haus in USA 80 Millionen USD für eine Smart Cities Initiative bereitgestellt. Im Oktober 2016 hat das DOT weitere 65 Millionen USD für Mobilitäts- und Transport-Projek-

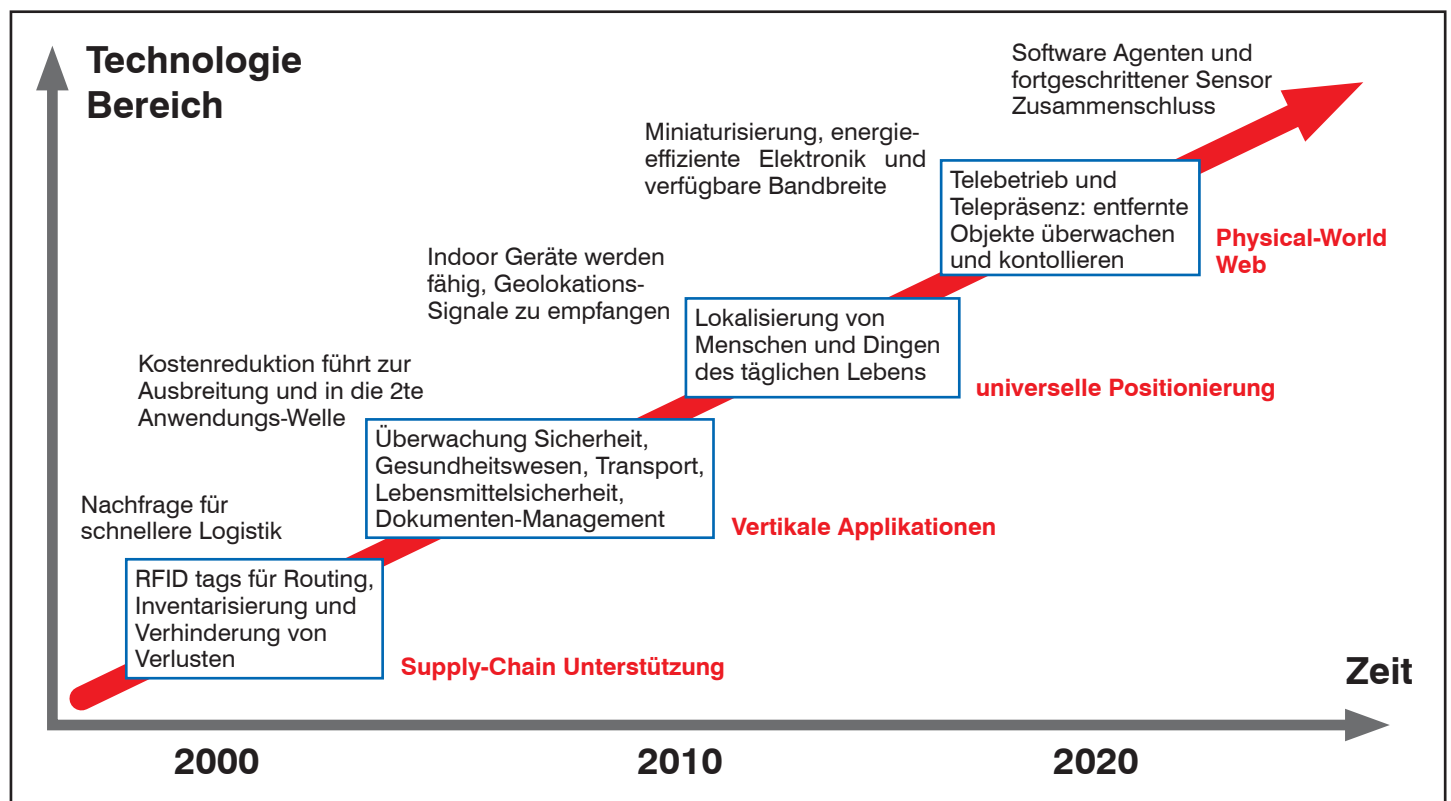


Abbildung 5.1: IoT Roadmap der Technologien und Applikationen

Quelle: SRI Consulting Business Intelligence

Internet of Things – die vierte industrielle Revolution - Teil 4

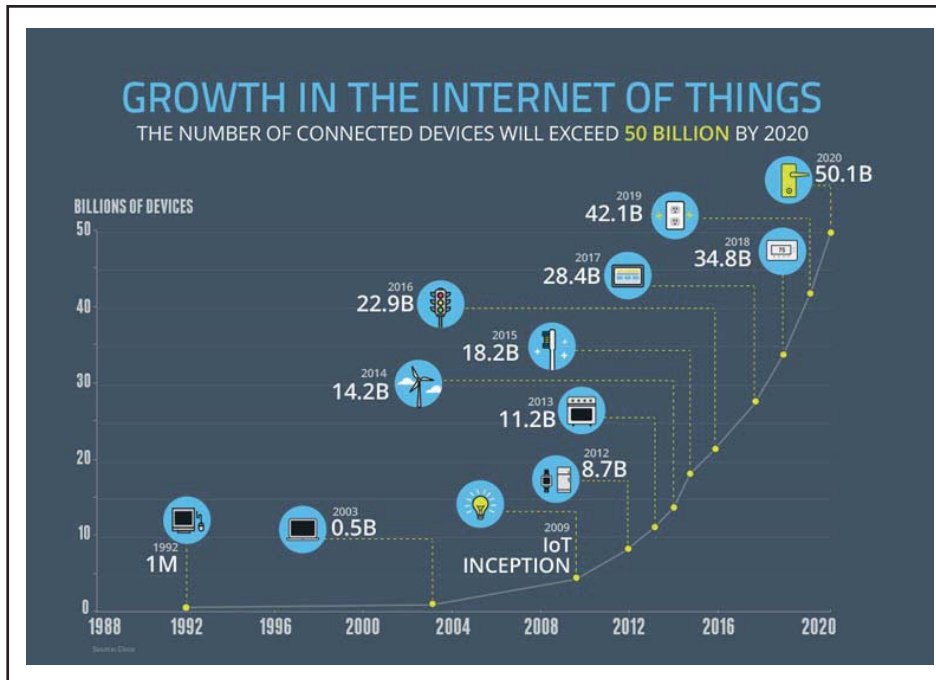


Abbildung 5.2: IoT Roadmap vernetzter Objekte

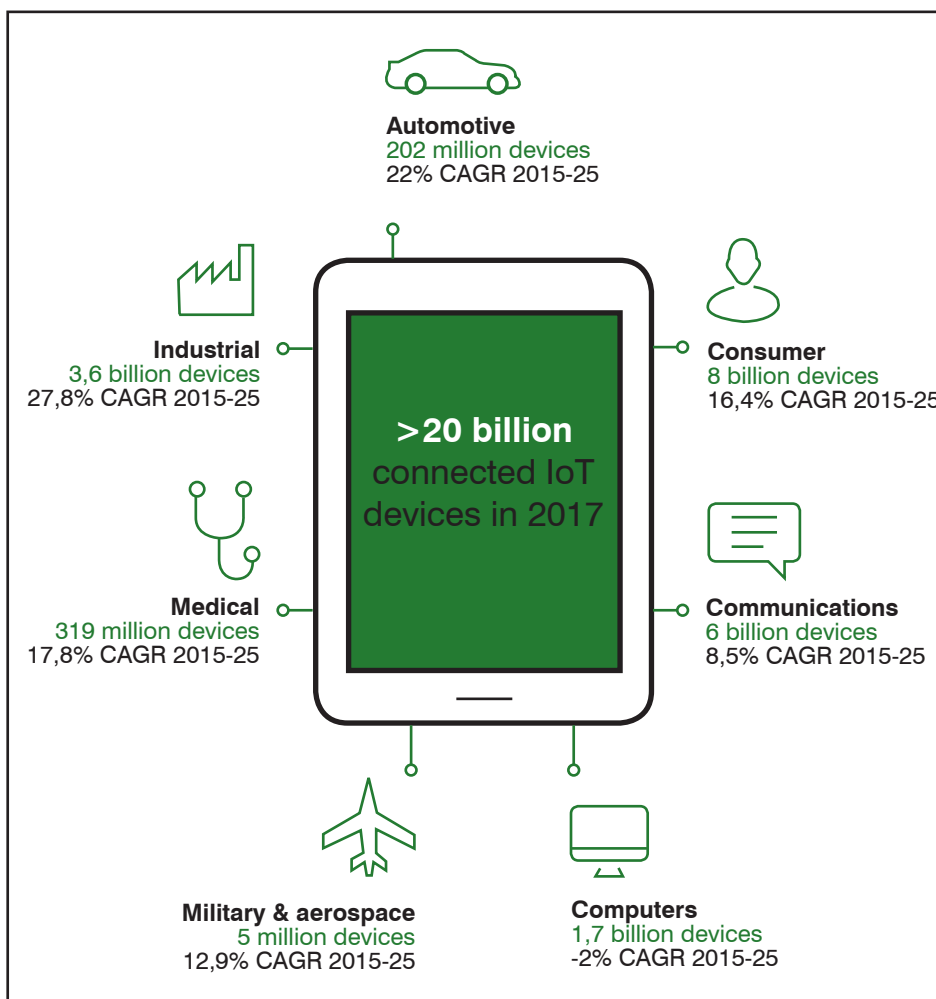


Abbildung 5.3: Verteilung vernetzter Objekte auf verschiedene Märkte

te in Aussicht gestellt (nun ja – hier bleibt für beide Bereiche abzuwarten, ob der neue Präsident Trump dies weiterhin für gut befindet ...). In Europa stellt das Horizon 2020 Programm knapp 70 Millionen EUR für ein Städtebeleuchtungs-Projekt zur Verfügung, für Nachhaltigkeit in Städten mit naturnahen Lösungen werden 44 Millionen EUR ausgelobt.

Funktechnologien und der Internet of Things Markt

Funkverbindungen und mobile Geräte sind Kernfaktoren des IoT Marktes. Wi-Fi ist hier mit IEEE 802.11ad (auch als Wi-Gig oder MGWS zu finden) im High End Bereich der Funktechnologien einzuordnen. Smart Home Plattformen erreichen durch die zunehmende Adaption von Wi-Fi als Standard-Konnektivität Interoperabilität und werden bis 2020 die Steuerung von 319 Millionen vernetzten Geräten übernehmen. Wi-Fi-fähige Home Anwendungen werden von 37 Millionen in 2017 bis 2020 auf 221 Millionen anwachsen – das entspricht einem traumhaften CAGR von 82 Prozent. Eine Übersicht über den Smart Home Bereich zeigt Abbildung 5.4.

Aber nicht nur Wi-Fi, sondern auch WAN-Technologien im Low End Bereich wie LPWAN werden sich zu starken Treibern für den Internet of Things Markt entwickeln: verspricht das Low Power WAN doch niedrige Kosten, niedrigen Strombedarf und hohe Reichweite für Netze, die Millionen von Geräten zusammenschalten können, die vorher praktisch unnetzbar waren.

LPWAN entwickelt sich somit zu einer starken Konkurrenz für die Short Range Verfahren (Wi-Fi, Bluetooth, ZigBee und ähnliche), da das Verfahren zum einen kostengünstiger und einfacher ist, zum anderen eine höhere Reichweite hat. Smart Metering, Smart Building, intelligente Agrikultur (Smart Farming) und Umgebungs-Sensoren nutzen Low Power WAN Verfahren. (siehe Abbildung 5.5)

Auch im 5G Funkmarkt gibt es einige heiß gehandelte IoT Anwendungs-Szenarien: Intelligente Agrikultur nutzt seit einigen Jahren in erheblich steigendem Maß vernetzte Sensortechnologien mit LPWAN Vernetzung. Das reicht von einfacher Wassertanküberwachung bis hin zu speziellen Sensoren, die den Feuchtigkeitsgrad und die chemische Zusammensetzung des Bodens überwachen können. Darüber hinaus ist Landwirtschaft einer der größten prognostizierten Märkte für Drohnen-Einsatz.

Höhere (5G) Datenraten werden Video-