

HZD

Hessische Zentrale für Datenverarbeitung

HESSEN



TRENDBERICHT 2018



VORWORT

Hessen stellt in seiner IT erneut die Weichen in Richtung Zukunft: Die digitale Verwaltung wird aufgrund des Gesetzes zur Verbesserung des Onlinezugangs zu Verwaltungsleistungen (kurz: Onlinezugangsgesetz, OZG) ihre Leistungen weiter systematisch ausbauen und ihre Angebote bündeln. Auch wenn der Bund hier eine weitgehende Regelungskompetenz hat, erfolgt die Umsetzung wesentlich in den Portalen und Fachverfahren der Länder. Dafür sind leistungs- und zukunfts-fähige Infrastrukturen und Anwendungen auf allen Verwaltungsebenen erforderlich. Die hierzu benötigten und einsetzbaren Technologien verändern sich sehr dynamisch. Unter dem Stichwort „Digitalisierung“ entfaltet sich dabei ein breites Spektrum an Möglichkeiten und Herausforderungen, aber auch Bedrohungen. Diese Entwicklungen zu kennen, innovative Verwaltungsprodukte zu entwickeln und sicher zu betreiben ist daher für die HZD kein Selbstzweck.



// Wir wollen einen Beitrag zu einem modernen Staat leisten.

Wir wollen damit – gemeinsam mit Bund, Ländern und Kommunen – einen Beitrag leisten zu einem modernen Staat, im Interesse der Unternehmen, Organisationen und Menschen in Hessen.

Joachim Kaiser
Joachim Kaiser, Direktor der HZD

EINLEITUNG

In manchen Jahren dominieren einige große Themen die IT-Presse und den HZD-Trendbericht. Doch wie neben der Großwetterlage auch das lokale Klima eine Rolle spielt, sind es auch hier manchmal die Detailthemen, die unsere Aufmerksamkeit verdienen, denn sie sind zugleich Quelle und Ergebnis der großen Strömungen. So finden wir im Kapitel *Architektur, Lösungen & Systeme* ganz verschiedene Antworten auf die Frage, wie sich Anwendungen und Rechenleistung modularisieren und verteilen lassen. Zum Bereich *Entwicklung, Technik & Netze* gehören neben zwei weiteren Facetten der Software-defined-Techniken Spezialthemen wie tastbare Displays und neue Entwicklungen bei Minirechnern. Web-Apps, die auch offline funktionieren, Bots und neue Ansätze der grafischen Programmierung zeigen wie vielfältig der Themenbereich *Daten, Information & Software* ist. Und schließlich betrachten wir im Kapitel *Gesellschaft, Sicherheit & Recht* aus ganz unterschiedlichen Richtungen, wie wir mit unseren Daten umgehen.

Mit unseren Beschreibungen und Bewertungen der Trends wollen wir unseren Leserinnen und Lesern erneut einen interessanten Blick auf die IT und deren Einsatzmöglichkeiten in der Verwaltung eröffnen. Wir würden uns freuen, wenn Sie über Ihre Erkenntnisse und daraus gewachsene Ideen mit uns ins Gespräch kommen.

So breit wie das Themenspektrum der IT-Trends ist auch das Wissen der Kolleginnen und Kollegen, die an dieser Ausgabe beteiligt waren: Harms Becker, Gundula Bucsa, Janina Einsele, Petra Fechner, Matthias Guckler, Jason Horn, Markus Kantowski, Dr. Alberto Kohl, Adam Miosga, Peter Müller, Dr. Klaus-Dieter Niebling, Melanie Pörschke, Markus Schramm sowie Tülün Syha und schließlich haben bei Gestaltung und Produktion aus dem Team der Öffentlichkeitsarbeit Birgit Lehr und Manuel Milani mitgewirkt.

Dr. Markus Beckmann

ZU DIESEM TEXT

Mit dem Trendbericht ermöglichen wir unseren Leserinnen und Lesern einen Ausblick auf aktuelle Trends in der Informationstechnologie. Dabei wollen wir uns jedoch nicht auf eine rein fachliche Information über die technischen Hintergründe und die weitere Entwicklung beschränken. Als IT-Gesamtdienstleister für die Hessische Landesverwaltung steht für uns die strategische Bedeutung der erfassten Trends für die Verwaltung im Mittelpunkt. Daher haben wir jedes einzelne Thema im Hinblick auf seine Auswirkungen auf die Verwaltung bewertet. Der Fokus liegt dabei auf der Hessischen Landesverwaltung. Neben einem kurzen Bewertungstext werden jeweils drei Kennzahlen angegeben, die die Einordnung der Themen in IT-strategische Überlegungen erlauben.

Verwaltungsrelevanz

Die Verwaltungsrelevanz gibt an, in welchem Maß sich ein Trend auf die Verwaltung auswirken kann. Dies kann auf zweierlei Art erfolgen: Zum einen können Trends zu technischen Änderungen in der IT-Landschaft führen bzw. solche Änderungen ermöglichen. Diese Trends sind daher in dem Maß verwaltungsrelevant, wie sie sich auf einige oder alle IT-Arbeitsplätze im Land auswirken – entweder direkt am Arbeitsplatz oder durch die Gesamtinfrastruktur.

Zum anderen können IT-Trends dazu führen, dass sich Verwaltungsabläufe ändern oder ganz neue Abläufe etabliert werden (können). In diesen Fällen haben die IT-Trends also Auswirkungen auf die Kernprozesse der Verwaltung.

Die Verwaltungsrelevanz wird auf einer fünfteiligen Skala angegeben, in der die Auswirkung des Trends auf die Verwaltung bewertet sind:

■■■■■	starke
■■■■□	deutliche
■■■□□	mittlere
■■□□□	geringe
□□□□□	keine

Umsetzungsgeschwindigkeit

Die Umsetzungsgeschwindigkeit gibt an, wie schnell ein Trend in der Verwaltung umgesetzt werden kann. Sie kann als ein Maß für die Komplexität der entsprechenden Trendergebnisse gesehen werden: Je komplizierter ein Resultat oder Produkt ist, desto länger dauert es, dieses in Verwaltung und Unternehmen nutzbar zu machen. Die fünfteilige Skala gibt die Einführungsgeschwindigkeit mit folgenden Werten an:

■■■■■	sofort
■■■■□	1-2 Jahre
■■■□□	2-4 Jahre
■■□□□	mindestens 4 Jahre
□□□□□	noch nicht absehbar

Marktreife/Produktverfügbarkeit

Der Wert für Marktreife bzw. Produktverfügbarkeit gibt an, wie lange es dauern wird, bis Produkte, die auf der im Trend beschriebenen Entwicklung basieren, am Markt verfügbar sind. Die fünfteilige Skala gibt die Marktreife bzw. Produktverfügbarkeit mit folgenden Werten an:

■■■■■	sofort
■■■■□	1-2 Jahre
■■■□□	2-4 Jahre
■■□□□	mindestens 4 Jahre
□□□□□	noch nicht absehbar

INHALT

01 ARCHITEKTUREN // LÖSUNGEN // SYSTEME	8
Standard für Anwendungscontainer	10
Drum prüfe... kopfloses CMS	11
Nicht ohne meinen Server? Serverless Computing	14
Virtuelle Zäune (Geofencing)	18
Edge, Fog, Multi-Cloud	20
02 ENTWICKLUNG // TECHNIK // NETZE	22
Haptische Displays	24
Software belebt die Netze – SDN	26
– SD – jetzt auch im WAN	26
– ... und „wireless“ – SDWN	29
Minirechner – klein aber oho!	30
03 DATEN // INFORMATIONEN // SOFTWARE	32
„Wir sind die RoBOTer“	34
Von der Web-App zur App-App	37
„no“ oder „low“ – Programmieren (fast) ohne Code	40
04 GESELLSCHAFT // SICHERHEIT // RECHT	42
Unhörbare Datenübertragung per Ultraschall	44
Ein Quäntchen Sicherheit?	46
Once-only: „Einmal und nie wieder!“	49
 Trendmap	 52
Impressum	54



01

ARCHITEKTUREN // LÖSUNGEN // SYSTEME





STANDARD FÜR ANWENDUNGSCONTAINER

Wer seinen Kofferraum für den Urlaub packt, muss evtl. feststellen, dass das letzte Gepäckstück nicht mehr so einfach unterzubringen ist, und beginnt ein lästiges Puzzlespiel. Beim Transport von Waren kann man sich solche Spiele nicht erlauben, denn Lagerkapazitäten und Liegezeiten sind Kostenfaktoren. Um diese Kosten einzusparen, wurden Frachtcontainer entwickelt – zunächst für den Transport auf Schiffen. Da verschifft Container aber auch weiter transportiert werden müssen und die Bahn- und Straßeninfrastruktur in verschiedenen Teilen der Welt unterschiedlich ausfällt, wurden für Frachtcontainer mit der ISO-Norm 668 Standards definiert.

Auch für Software wurden Containertechniken entwickelt (vgl. HZD-Trendbericht 2016). Diese erlauben es, Anwendungen in leicht austauschbaren Komponenten auf einer Systemplattform zu betreiben, ohne dabei alle Systemressourcen virtualisieren zu müssen. Das ermöglicht es z. B., mehrere Instanzen einer Anwendung schnell zur Verfügung zu stellen. Allerdings gibt es verschiedenen Techniken, mit denen sich solche Container bauen lassen – z. B. Docker oder rkt (gesprochen „rock-it“). Darüber hinaus wurden Techniken entwickelt, die die Bereitstellung, Skalierung und Verwaltung von Containern bzw. Containeranwendungen automatisieren. Doch aufgrund der verschiedenen Ansätze drohte der Markt der Anwendungscontainer zu zersplittern. Um dem entgegenzuwirken und die Interoperabilität zwischen unterschiedlichen Systemen zu verbessern, haben sich zahlreiche Firmen und Organisationen unter dem Dach der Linux Foundation zusammengetan und die „Open Container Initiative“ (OCI) gegründet. Deren Ziel ist es, einen Standard für Anwendungscontainer zu etablieren.

Dabei werden fünf Prinzipien für Standardcontainer zugrunde gelegt:

1. Standardoperationen

Die Standardcontainer können über Standardoperationen generiert, gestartet und gestoppt sowie kopiert werden. Die sogenannten Snapshots, die den aktuellen Zustand eines Con-

tainers festhalten, können über Standardwerkzeuge des unterliegenden Filesystems erzeugt werden. Und sowohl Up- als auch Download für Container verwenden die Standardwerkzeuge des Netzwerks.

2. Unabhängigkeit vom Containerinhalt

Die Arbeitsweise der Standardoperationen ist unabhängig vom Inhalt der Container – z. B. einer Datenbank, einem Webserver oder einer Fachanwendung.

3. Unabhängigkeit von der Infrastruktur

Die Standardcontainer können auf allen Infrastrukturen, die den OCI-Standard unterstützen, be- und verarbeitet werden – z. B. auf Einzelplatzrechnern gebaut, in einen Cloud-Speicher hochgeladen und dann über eine Staging-Umgebung mit diversen Servern in einer öffentlichen Cloud mit mehreren Produktivinstanzen eingesetzt werden.

4. Ausrichtung auf Automation

Standardcontainer sind aufgrund der ersten drei Prinzipien besonders für die automatisierte Bereitstellung von Anwendungen geeignet.

5. Unterstützung industrieller Produktion

Durch die Standardisierung sollen Anwendungscontainer bedarfsgerecht – aber nicht bedarfsspezifisch – nutzbar werden. D. h., dass die Anwendungen in großen oder kleinen Stückzahlen schnell und verlässlich für „jeden“ Zweck bereitgestellt und betrieben werden können.

Mitte des Jahres 2017 wurden die ersten zwei Elemente des Standards freigegeben, die das Format von Containern (Image Format Specification) und deren Einbindung in die Laufzeitumgebung (Runtime Specification) spezifizieren. Weitere Spezifikationen, die noch bearbeitet werden, legen z. B. fest, wie Container als eine Ansammlung von Dateien im Filesystem organisiert werden (Filesystem Bundle), oder wie deren Integrität durch kryptografische Funk-

tionen abgesichert werden kann (Hashing for Content Integrity).

Zudem arbeitet die Initiative an einer Prüf-umgebung, mit der die Einhaltung der OCI-Spezifikationen nachgewiesen werden kann (Compliance Test Suite). Das soll durch ein Zertifizierungsprogramm unterstützt werden.

Bewertung

Die Anzahl der Faktoren, die die Entwicklung, Bereitstellung und den Betrieb von Verwaltungsanwendungen bestimmen, ist groß. Neben den fachlichen und organisatorischen Vorgaben, die komplexe Updates und termingerechte Bereitstellung erfordern können, sind dies auch tech-

nische Rahmenbedingungen wie akute Bedrohung der Sicherheit, Veränderungen an Basis-komponenten (Datenbank, Webserver etc.) oder Anpassungen der Hardware-Infrastruktur, etwa die Art und Anzahl der Server. Wenn dann die Prozesse für Entwicklung, Bereitstellung und Betrieb der Anwendungen spezifisch an derartige Veränderungen angepasst werden müssen, erfordert dies Zeit und Energie und verursacht dadurch zusätzliche Kosten. Techniken, die die Unabhängigkeit der Prozesse von spezifischen Inhalten und Infrastrukturen fördern, und eine Standardisierung, die die Automation dieser Techniken unterstützt, können somit einen wertvollen Beitrag zu einer modernen Anwendungslandschaft auch bei Verwaltungsverfahren leisten.

ANWENDUNGS-CONTAINER

Verwaltungsrelevanz:

■■■□

Umsetzungs-geschwindigkeit:

■■■□

Marktreife/Produkt-verfügbarkeit:

■■■□

DRUM PRÜFE... KOPFLOSES CMS

Hin und wieder findet man sie noch: Webseiten, die den Charme der frühen Jahre des Internets versprühen und mit „handgeklöppeltem“ HTML in 16 knalligen Farben gestaltet sind. Dabei muss man aber kein Spezialist für alle möglichen Web-Sprachen sein, um ansprechende Seiten gestalten zu können. Die meiste Arbeit nehmen uns Content-Management-Systeme (CMS) ab. Darin lassen sich Webseiten gestalten und Inhalte – getrennt von der Gestaltungsvorlage – erfassen. Das CMS fügt dann beides beim Veröffentlichen zusammen. Solche CMS gibt es von Leichtgewichten für den Privatgebrauch zum Schreiben eines Blogs bis hin zu komplexen Systemen für die Gestaltung umfangreicher Portalangebote, die z. B. Shop-Anwendungen integrieren.

Hin und wieder kommt es vor, dass ein neues CMS benötigt wird, – sei es etwa, weil das bestehende System nicht mehr weiterentwickelt wird, auf einer neuen technischen Plattform betrieben werden soll oder einfach den gestalterischen Anforderungen nicht mehr genügt. Dann müssen nicht nur Seiten und Masken neu

gestaltet werden, sondern auch die Inhalte von einem System in das andere umziehen. Und das ist i. d. R. kein Vergnügen.

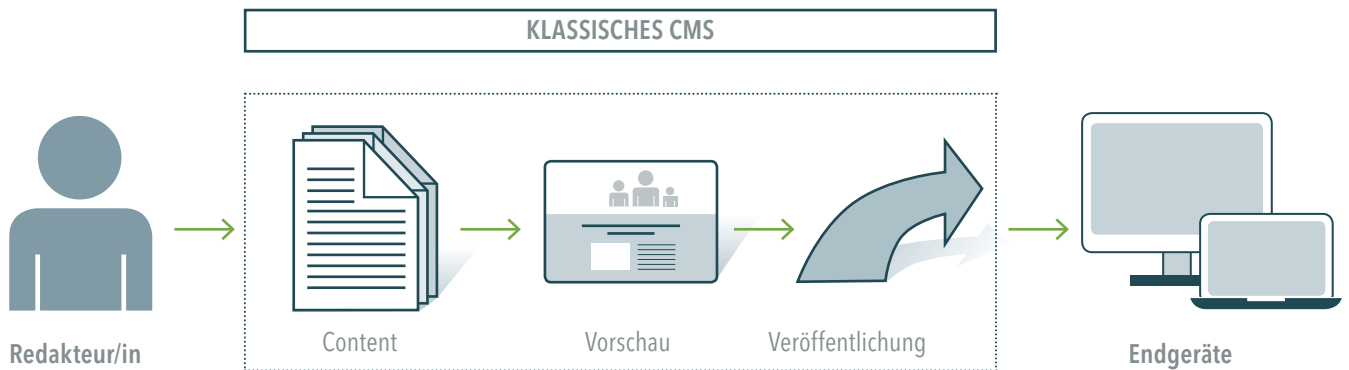
Eine andere Herausforderung bei der Verwendung eines CMS ist es, wenn für die Präsentation der Inhalte verschiedene Plattformen be-

// Einfache Webseiten, Fachanwendungen und mobile Apps können ganz unterschiedliche Anforderungen an die Aufbereitung von Content haben.

dient werden müssen. Das betrifft nicht allein die Frage des benutzten Endgeräts – Desktop, Tablet oder Smartphone – mit den verschiedenen Bildschirm- und Bedienformaten. Hierfür lassen sich mit Ansätzen für Responsive-Design oder adaptive Webseiten zumeist zufriedenstellende Ergebnisse erzielen. Es betrifft auch die Verwendung der Inhalte. Wenn diese neben „einfachen“ Webseiten auch von Fachanwendungen oder mobilen Apps genutzt werden



01



Beim klassischen CMS finden alle Schritte des Redaktionsprozesses in einem System statt.

sollen, können ganz unterschiedliche Anforderungen an die Präsentation oder gar die Bearbeitung der Inhalte entstehen. Hier stoßen klassische CMS schnell an Grenzen.

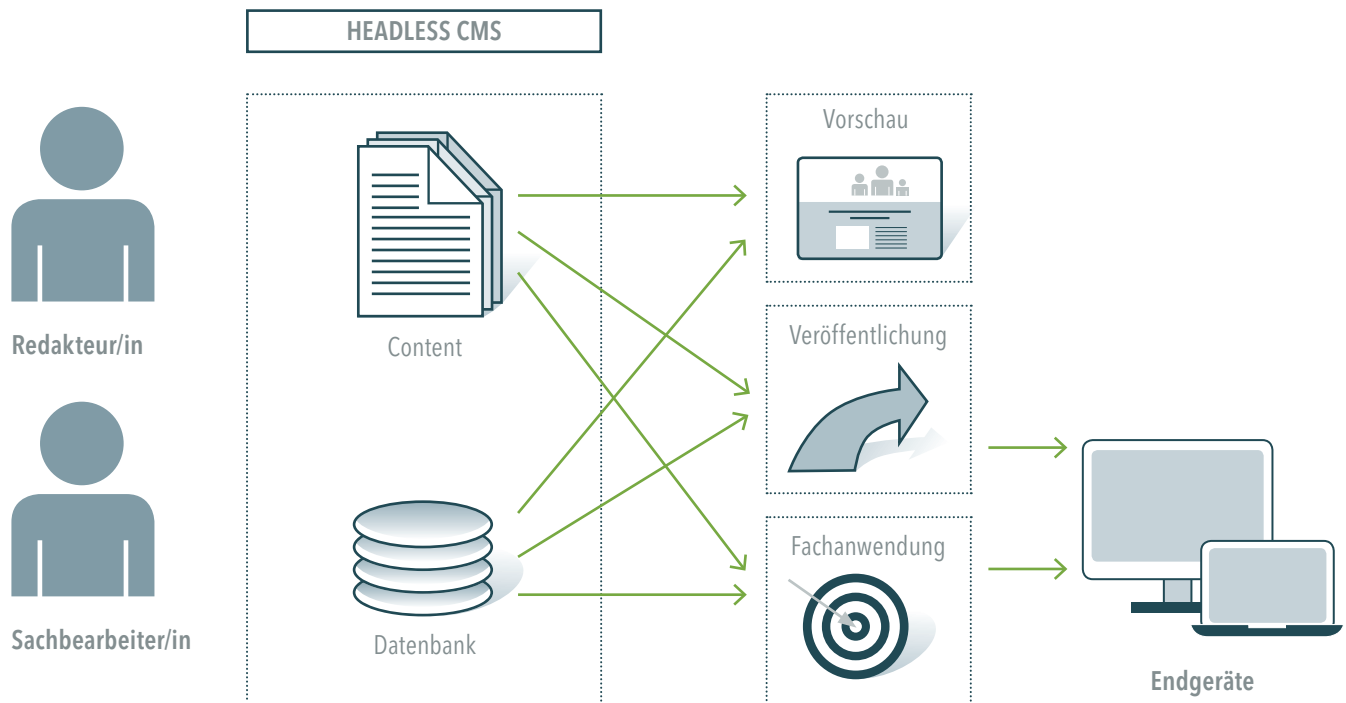
Diesen beiden Herausforderungen kann man entgegentreten, wenn man die Verwaltung der Inhalte, also das eigentliche Content-Management, von der Definition der Präsentation und von der Veröffentlichung abkoppelt – und zwar nicht nur funktional innerhalb einer Anwendung, sondern auch technisch im Rahmen der Anwendungsarchitektur.

Man muss sich dabei vor Augen führen, dass die Arbeit mit einem CMS vier wesentliche Bereiche hat:

- die Gestaltung der Darstellung,
- die Erfassung von Inhalten,
- deren Aufbewahrung und Verwaltung (z. B. Veröffentlichungszeitraum) und
- das Publizieren der Inhalte.

Während die Gestaltung der Darstellung eine in dem Sinne einmalige oder zumindest seltene Aufgabe ist, als Webseiten nicht täglich neu- oder umgestaltet werden, gehört das Erfassen und Publizieren von Inhalten zum Tagesgeschäft vieler Web-Redaktionen. Das legt es nahe, den entsprechenden Redaktionsprozess in einer Anwendung durchgängig umzusetzen. Löst man die Darstellung und damit das Publizieren aber aus dem CMS heraus, fehlt bei der Arbeit mit einem derart „kopflösen CMS“ (engl. „headless CMS“) genau dieser Prozess. Schon die Vorschau auf einen gerade erstellten Content wird ggf. schwierig.

Es kommt also darauf an, für welche Einsatzzwecke man ein CMS verwendet, wenn man beurteilen möchte, ob ein headless CMS das Mittel der Wahl ist – und wie dieses dann gebaut ist. Um die Inhaltsverwaltung und die Präsentation trennen zu können, wird das CMS mit einer Programmierschnittstelle (API) ausgestat-



Beim headless CMS ist die Verwaltung des Contents von dessen Weiterverarbeitung getrennt.

tet, die Routinen zum Abruf definierter Daten bereitstellt. Das kann z. B. über eine REST-API unter Verwendung von XML-Strukturen oder JSON-Objekten erfolgen. Ein Vorteil dieses Ansatzes besteht darin, dass die Entwickler der Webseiten oder Anwendungen sehr flexibel in der Verwendung der Programmiersprachen sind. Auch das erleichtert ggf. die Bereitstellung von Inhalten auf verschiedenen Ausgabeplattformen – insbesondere dann, wenn eine Plattform spezifische Entwicklungswerkzeuge voraussetzt. Zudem können die Redaktions- und die Darstellungsanwendung unabhängig voneinander skaliert werden, z. B. wenn mehrere, verteilte Redaktionsteams Inhalte zu einem Portal beisteuern. Der Umfang und der Betriebsaufwand eines headless CMS mögen geringer sein, als bei Komplettsystemen. Ob aber die Bereitstellung des Gesamtsystems von der Content-Erfassung bis zur Präsentation mit einem headless CMS einfacher wird, hängt davon ab, wie „pflegeleicht“ die Anwen-

dungen für das Frontend sind. Dies muss neben den Migrationskosten erwogen werden, wenn der Wechsel auf ein headless CMS in Erwägung gezogen wird.

Der zunehmende Wunsch nach mehr Flexibilität bei der Anbindung von Frontends zur Darstellung von Inhalten hat offenbar CMS-Entwickler dazu bewogen, ihre bestehenden Systeme um APIs zu erweitern. Damit können neben der integrierten Publikations- und Darstellungskomponente bei Bedarf weitere Ausgabeplattformen angebunden werden.

Ein solches „entkoppeltes CMS“ (engl. “decoupled CMS”) bietet zumindest schon einmal mehr Flexibilität als geschlossene Systeme. Im Kern bleibt hier aber die Abhängigkeit von einer Gesamtanwendung und von einem Hersteller erhalten, sofern man nicht das integrierte Frontend ignoriert.



KOPFLOSES CMS

Verwaltungsrelevanz:

■■■□

Umsetzungs- geschwindigkeit:

■■■□

Marktreife/Produkt- verfügbarkeit:

■■■□

Es wäre wünschenswert, dass sich entlang der ganzen Prozesskette beim Content-Management – von der Erfassung bis zur Präsentation – Standards entwickeln, die den flexiblen Einsatz bzw. Austausch einzelner Komponenten nicht nur theoretisch, sondern tatsächlich unterstützen. Andernfalls scheint sich der Einsatz von decoupled oder headless CMS in erster Linie dort zu lohnen, wo Contents in verschiedenen Anwendungen verarbeitet werden, in denen sowieso in nennenswertem Umfang Entwicklungs- oder Anpassungsarbeiten anfallen.

Bewertung

Öffentliche Verwaltungen weisen zumeist eine hohe fachliche und organisatorische Komplexität auf. Dem gegenüber steht der Wunsch von Verwaltungskunden, alle benötigten Dienstleistungen „aus einer Hand“ zu beziehen. Nicht umsonst sind Verwaltungs- bzw. Bürgerportale weiterhin ein wichtiges Thema des E-Government. Dazu kommt, dass zu publizierende Inhalte von Ver-

waltungen eine hohe Bandbreite an Dynamik aufweisen. Gesetzliche Regelungen und Verwaltungsvorschriften, die über viele Jahre Bestand haben oder sporadisch fortgeschrieben werden, gehören ebenso dazu wie tagesaktuelle Meldungen. Dies allein erfordert komplexe, aber dennoch flexible CMS-Strukturen. Dazu können und sollen Verwaltungsdienstleistungen neue Wege und Möglichkeiten der Bereitstellung nutzen.

Es liegt also nahe, für die Pflege und Publikation von Inhalten eine Content-Infrastruktur zu verwenden, die die verschiedenen resultierenden Anforderungen erfüllen kann. Eine Entkopplung bzw. Trennung der großen Komponenten kann die dafür notwendige Flexibilität unterstützen. Im Fokus müssen aber die praktischen Tätigkeiten in den Redaktionen und damit letztlich die Kundenservices stehen.

NICHT OHNE MEINEN SERVER? SERVERLESS COMPUTING

Beim Betrieb von IT-Anwendungen wird in vielen Architekturmodellen der Abstand zwischen der Software und der Hardware immer größer. Wenn nicht gerade bewusst – z. B. aus Performance-Gründen – maschinennah programmiert wird, ist es uninteressant, wie viele und welche Rechner später die Arbeit zu bewältigen haben. So könnte man denken, denn beim Bau einer Anwendung, die etwas mehr tut, als einfache Berechnungen auszuführen oder überschaubare Daten von einigen wenigen Nutzern zu verarbeiten, muss sich der Entwickler doch oft Gedanken über die Plattform machen, auf der die Anwendung betrieben wird. Da stehen Fragen im Raum wie:

- Welches Datenvolumen fällt in welchen Zeiten an? Wie viel ist dauerhaft, wie viel temporär zu speichern?
- Wann arbeiten wie viele Nutzer mit dem System? Sind die räumlich und zeitlich gleichmäßig verteilt oder gibt es Zeiten hoher Last?
- Wann wird hohe Rechenleistung benötigt und was passiert in der Zwischenzeit?

Die Antworten auf diese und ähnliche Fragen sowie die Ermittlung der benötigten Rechenleistung liefern den Bedarf an entsprechenden Systemen, einschließlich einer Reserve. Während die Systeme früher 1:1 in Hardware abgebildet wurden, hat sich der Einsatz virtueller Maschinen in sehr vielen Fällen bewährt und ist damit zum

Standard geworden. Neben der einfacheren Bereitstellung der „Maschine“ spielt dabei auch die bessere Ausnutzung von Hardware eine Rolle. Insbesondere die zeitlich begrenzte Erweiterung des Systems um mehr Rechenleistung – z. B. in Zeiten mit Spitzenlast – ist kostengünstiger, als wenn man „die maximale Hardware“ ständig vorhalten muss, selbst wenn ein Teil davon immer wieder ungenutzt bleibt.

Dieser Ansatz lässt sich auf Softwarekomponenten und Anwendungen übertragen. Auch diese können zeitweise und in der Menge gestaffelt mehr oder weniger dynamisch dazu gebucht werden. Die resultierenden Modelle sind heute je nach Schwerpunkt unter den Namen Infrastructure -, Platform - oder Software as a Service (IaaS, PaaS resp. SaaS) mit vielen weiteren Spielarten Bestandteil der einschlägigen Cloud-Angebote. Mit Hilfe von Anwendungscontainern lassen sich individuelle Konfigurationen bzw. Komponenten schnell auf diesen Plattformen bereitstellen. Doch auch in einem solchen Modell muss der Entwickler bzw. der Systemarchitekt planen, was wann benötigt wird – und wie sich die dann anfallenden Kosten optimieren lassen.

Unter dem Namen „cloud native“ (etwa: „in der Cloud beheimatet“) entstehen derzeit Modelle und Plattformen, die noch stärker auf Virtualisierung und dynamische Nutzung fokussiert sind. Hierbei werden die Anwendungen als Ansammlungen einzelner Dienste bzw. Funktionen (Stichwort: „microservices“) entwickelt. Dadurch sollen die Anwendungen flexibler werden und besser zu pflegen sein. Jede derartige Softwarekomponente wird in einen Container verpackt, der sie leicht reproduzierbar macht und den Zugriff auf Ressourcen strukturiert. Die Container werden dann hochdynamisch orchestriert, d. h. für den jeweiligen Anwendungszweck zusammengestellt. Auf diese Weise soll die Ressourcennutzung – und damit auch die Kostenentwicklung – weiter optimiert werden.

Auf der Seite der Anwendungsentwicklung korrespondiert dieses Betriebsmodell mit dem Ansatz des „serverless computing“ (kurz: „server-

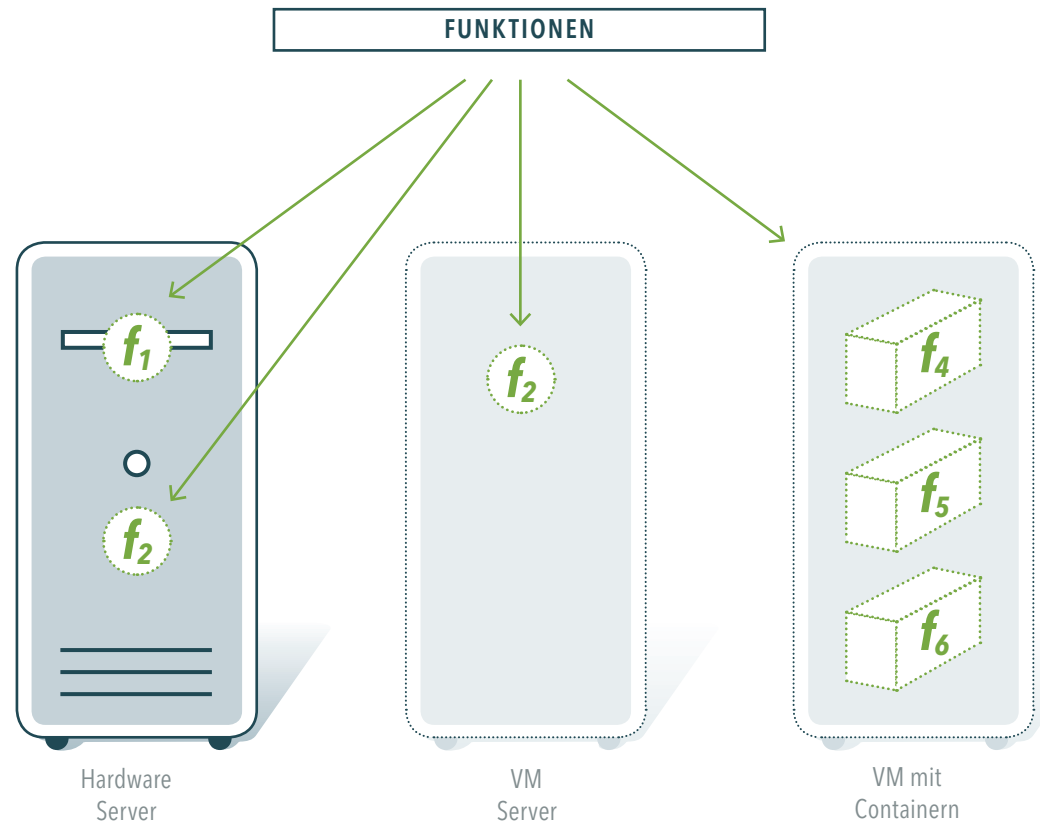
less“). Mit der Bezeichnung soll nicht suggeriert werden, dass keine Server mehr benötigt werden. „Irgendwo“ werden auch in diesem Modell Maschinen arbeiten müssen, auf denen die benötigten Anwendungscontainer laufen können, die Daten verwaltet und die Maßnahmen der Orchestrierung umgesetzt werden. Allerdings soll sich der Anwendungsentwickler keinerlei Gedanken darüber machen müssen. Stattdessen gestaltet sich das Szenario wie folgt: Wenn eine bestimmte Funktion in einer Anwendung benötigt wird, wird der entsprechende Funktionscontainer gestartet, arbeitet die Funktion ab und wird danach wieder beendet. Eine solche Funktion kann z. B. darin bestehen, dass beim Eintreffen neuer Daten „in der Anwendung“ ein spezifischer Bearbeitungsschritt erfolgt. So

// Der Anwendungsentwickler soll sich keine Gedanken um Hardware, virtuelle Maschinen oder Container machen müssen.

könnten beim Eintreffen einer Nachricht mit Anlagen die Metadaten der Anhänge (z. B. Titel, Autor, Erstelldatum, Dateiname und -größe) ausgelesen und mit der jeweiligen Anlage in ein Dokumentverzeichnis geschrieben werden. Eine solche Funktion wird eventuell nur sporadisch benötigt und die Programmroutine muss dementsprechend nur zeitweise Ressourcen belegen. Sollten jedoch vorübergehend sehr viele Nachrichten eintreffen, werden auch viele Container mit der Metadatenfunktion gestartet – und anschließend wieder beendet.

Dieses einfache Beispiel lässt ahnen, dass nicht alle Funktionen gleich gut für eine für eine Serverless-Architektur geeignet sind. Besonders gut geeignet sind Funktionen, die

- nicht unverzüglich benötigt werden, denn die Startzeit für den zugehörigen Container darf im Vergleich zur Problemstellung nicht lange dauern oder gar die „Lebenszeit“ der zu verarbeitenden Daten übersteigen.



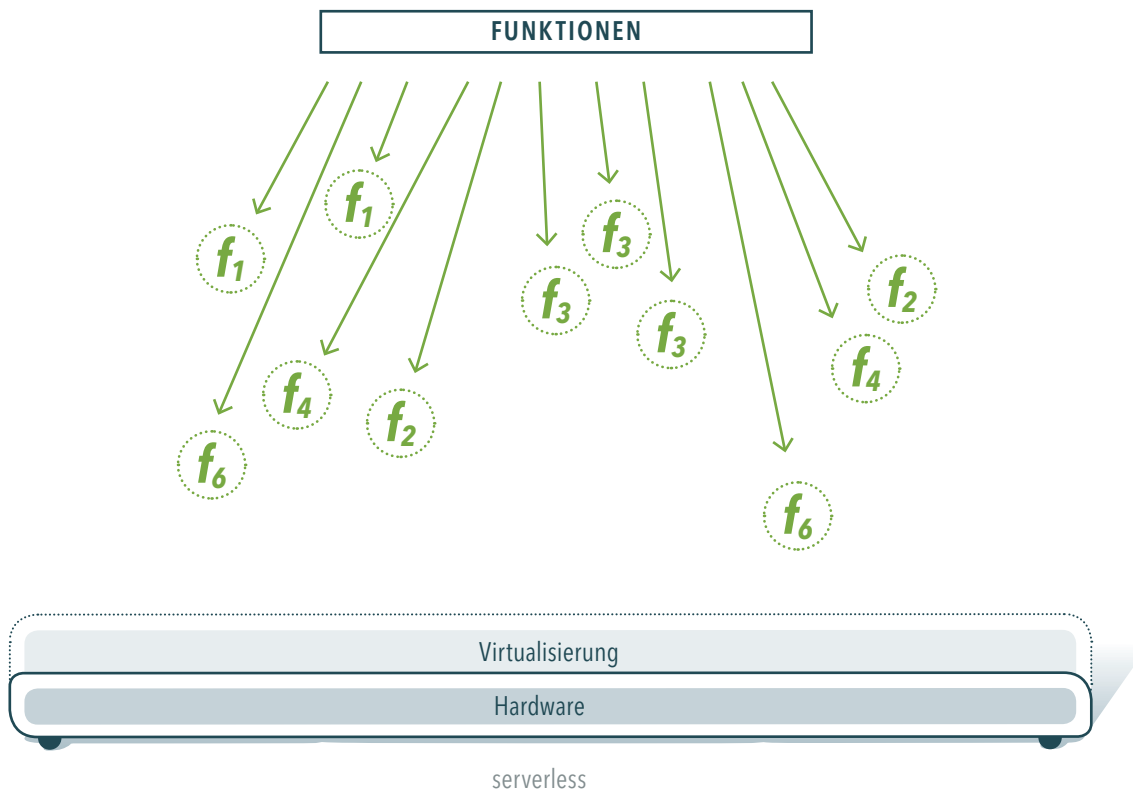
Anwendungsfunktionen werden einem bestimmten Gerät, einer virtuellen Maschine oder einem Anwendungscontainer zugeordnet.

- sich leicht parallel anwenden lassen, denn dann kann für jeden Anwendungsfall ein Container gestartet werden - bis die zur Verfügung stehenden Ressourcen ausgeschöpft sind.
- unregelmäßig und mit wechselnder Häufigkeit benötigt werden.
- asynchron arbeiten, sodass sie jeweils bei Bedarf und nicht in einem bestimmten Arbeitstakt eingesetzt werden können.
- keinen „inneren Zustand“ benötigen („stateless“), denn dieser würde mit dem Beenden des Containers verlorengehen.
- häufig an veränderte Anforderungen anzupassen sind, denn es ist ggf. einfacher, den Container auszutauschen, als die komplette Anwendung mit der geänderten Funktion neu zu generieren.

Beispiele für solche Funktionen sind:

- Synchronisierung von Datenbeständen
- Chat Bots
- ereignisgesteuertes Backup
- Initialisierung eines neuen Datensatzes
- automatische Bildbearbeitung (Größenanpassung und „Thumbnail“-Erzeugung)
- Verarbeitung von Sensordaten
- Back-Ends für mobile Anwendungen (über API-Gateways)
- Verarbeitung von SQL-Abfragen in Datenströmen

Der cloud-native-Ansatz und Serverless-Konzepte werden durch die Cloud Native Computing Foundation (CNCF), eine industrielle Vereinigung, vorangetrieben. Deren selbstgewählte Aufgabe besteht darin, eine neue Sichtweise auf eine IT zu etablieren, die daraufhin optimiert ist,



In einer Serverless-Architektur werden Anwendungsfunktionen ohne Bezug zu konkreter Hardware oder Virtualisierungstechnik implementiert und so oft gestartet wie sie benötigt werden.

in modernen verteilten System-Umgebungen mit zehntausenden selbst heilenden Rechnerknoten zu skalieren. In einer derart idealisierten IT-Landschaft werden die beiden wesentlichen Rollen bei der Bereitstellung von Anwendungen unterstützt: die Entwickler, die sich ganz auf die Realisierung der Software konzentrieren können, ohne über die Infrastruktur nachdenken zu müssen, und die Plattformbetreiber, die eben diese benötigte Infrastruktur bereitstellen, ohne die Anwendungen im Blick haben zu müssen. Aber auch wenn Anwendungsbetrieb sowie Produktivität und Kosten der eingesetzten Komponenten dabei verbessert werden, bringt diese Sichtweise einige Herausforderungen mit sich.

Das Problem der Performance wurde bereits angesprochen: Da nicht benötigte Ressourcen möglichst umgehend wieder freigegeben wer-

den sollen, müssen ggf. viele Container häufig neugestartet werden. Dadurch können sich deren Anlaufzeiten zu spürbaren Verzögerungen im Anwendungsverlauf summieren.

Für „normale“ Anwendungen scheinen die Ressourcen der Infrastruktur unbegrenzt: Es steht immer genau das zur Verfügung, was benötigt wird. Aber auch das hat natürlich seine Grenzen. Hochleistungsrechner, die ihre Spitzenleistung gleichzeitig für eine Anwendung benötigen, lassen sich mit diesem Ansatz kaum ersetzen, denn auch Serverless-Plattformen müssen mit den vorhandenen Ressourcen haushalten.

Auch die Ausrichtung auf isolierte Funktionen birgt Herausforderungen: Die einzelnen Funktionen, die als Service bereitgestellt werden, lassen sich isoliert gut testen. Deren Zusammenspiel



SERVERLESS COMPUTING

Verwaltungsrelevanz:

■■■□

Umsetzungs- geschwindigkeit:

■□□□

Marktreife/Produkt- verfügbarkeit:

■■■□

zu testen, kann aber angesichts der verteilten und asynchronen Arbeitsweise der Gesamtanwendung schwierig werden. Dazu kommt, dass die Überwachung (Monitoring) und die Fehleranalyse mittels Debugging von zustandslosen Funktionen, die an nicht vorhersagbaren Orten ausgeführt werden, ebenfalls eine Herausforderung ist.

Und schließlich stellt sich die Frage nach der Sicherheit von Serverless-Architekturen. Dabei dient die Trennung von Infrastruktur und Anwendung der Schärfung der Zuständigkeit: Da der Entwickler keinen Einfluss auf die Infrastruktur hat, ist klar, dass deren Absicherung gegen Bedrohungen z. B. auf Ebene des Betriebssystems oder bei DoS-Attacken ausschließlich dem Anbieter obliegt. Umgekehrt kann ihm die Verantwortung für die Absicherung der Anwendungskomponenten nicht zugesprochen werden. Diese bleibt beim Entwickler der Anwendung und umfasst Themen wie die Sicherheit von Softwarebibliotheken oder Berechtigungen für die Nutzung feingranularer Funktionen. Des Weiteren wird darauf hingewiesen, dass die Vielzahl der Schnittstellen der Microservices die Angriffsfläche eher vergrößert.

Serverless wird als eine logische Weiterentwicklung des Cloud-Computing betrachtet. Dabei ist auch dieser Ansatz kein Allheilmittel für jegliche IT-Anwendung.

Bewertung

Anwendungen, die ein schwankendes Nutzungs- und Datenaufkommen haben, dürfte es auch in der Verwaltung zur Genüge geben. Anlass- und fallbezogene Bearbeitung bei flexibler Anpassung an sich ändernde Regelungen könnten Indikatoren dafür sein, dass ein Verfahren von einer Serverless-Architektur profitieren könnte. Das setzt aber voraus, dass Anwendungen microservice-orientiert gebaut werden und eine entsprechende Systemplattform mit den geeigneten Orchestrierungsmechanismen zur Verfügung steht. Dieser Architekturansatz vergrößert jedoch die Gesamtkomplexität durch die Integration weiterer Zwischenschichten. Ob man ihn verfolgt, ist also eine strategische Frage, die in einer heterogenen Anwendungslandschaft unter Einbeziehung vieler Akteure und im Gesamtbild beantwortet werden muss.

VIRTUELLE ZÄUNE (GEOFENCING)

Unabhängig von ihrer Gestaltung haben Zäune eine wichtige Aufgabe: Sie trennen jeweils zwei Bereiche voneinander. Das gilt für den heimischen Garten oder ein einzelnes Blumenbeet genauso wie für große, zu sichernde Anlagen oder gar Staaten. In einer globalisierten Welt scheinen sie einerseits anachronistisch, gewinnen aber andererseits als politisches Symbol in verschiedenen Teilen der Welt wieder an Bedeutung. Während sich Waren und Personen evtl. von Zäunen aufhalten oder zumindest leiten lassen, scheint dies für Daten nicht zu gelten: Die „Grenzübergänge“ für Daten und Informationen orientieren sich in der technischen Infrastruktur i. d. R. nicht an geografischen Gegebenheiten. Das führt dann

zu Problemen, wenn diese Daten in verschiedenen geografischen Regionen rechtlich unterschiedlich behandelt werden. Beispiele dafür sind der europäische Datenschutz oder Webinhalte und Dienste, die in einigen Ländern verboten sind, und Kryptowährungen werfen bei internationalen Transaktionen interessante Fragen nach Regulierung auf.

Doch auch in alltäglichen Situationen kann es hilfreich sein, wenn Anwendungen in der Lage sind zu erkennen, dass – im geografischen Sinne – eine Grenze überschritten wurde oder werden soll, selbst wenn diese nur virtuell ist. So kann z. B. der Staubsaugerroboter vor einem Sturz in den Keller gerettet oder der autonome Rasen-

mäher daran gehindert werden, die Blumen in Nachbars Garten gleich mit zu mähen.

Das Einrichten virtueller Zäune wird Geofencing („fence“ engl. für „Zaun“) genannt. Dabei werden i. d. R. Gebiete definiert, die von einer geschlossenen Linie umgeben sind. Im einfachsten Fall ist dies ein Kreis oder Quadrat. Es sind aber auch kompliziertere Linienzüge möglich, die z. B. die Umrisse eines Werksgeländes beschreiben. Dadurch wird der innere vom äußeren Bereich abgegrenzt und es kann festgelegt werden, wie sich eine Anwendung im Innen- bzw. Außenbereich oder beim Überschreiten der Grenze verhält. Im Zusammenhang mit Flugdrohnen hat das Thema Geofencing viel Aufmerksamkeit erfahren. Hier werden technische Methoden verwendet, um Flugverbotszonen zu definieren und durchzusetzen. Die Steuerungssoftware verhindert, dass die Drohne in ein solches Gebiet fliegt, das entweder dauerhaft (z. B. Flughafen) oder zeitweise (z. B. Krisengebiet oder Großveranstaltung) gesperrt ist. Geofencing kann aber auch dazu ver-

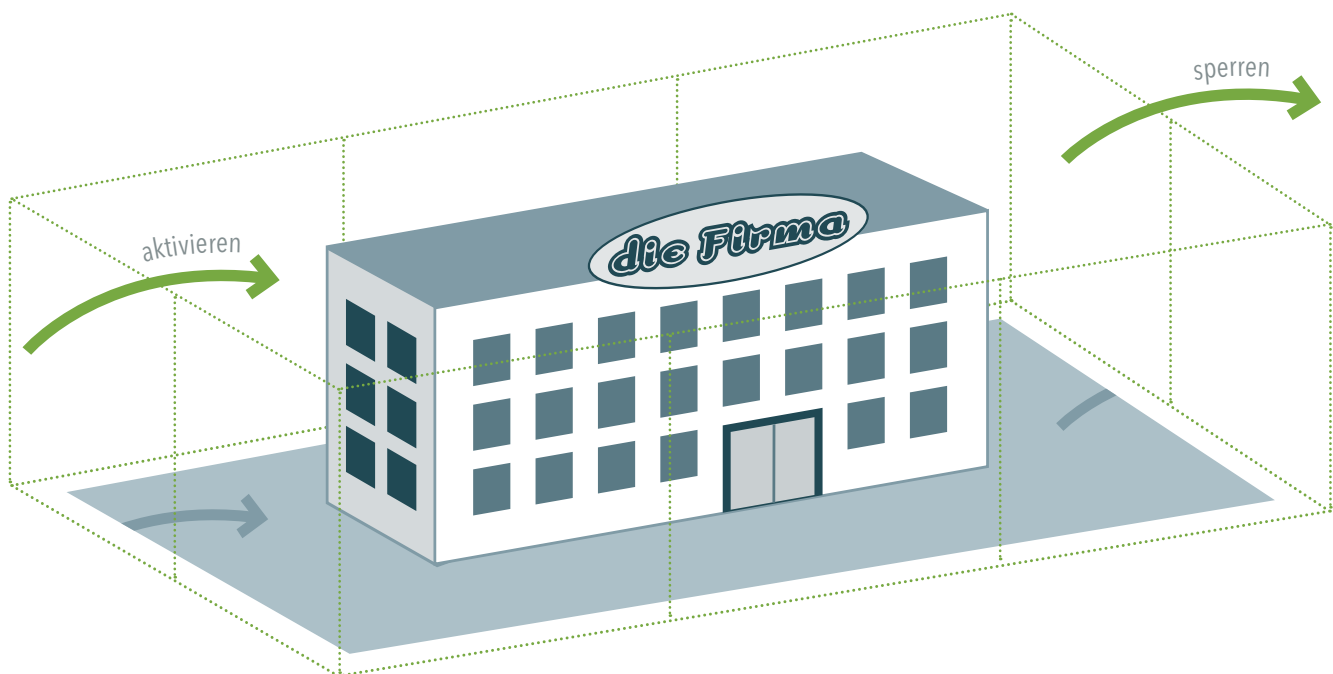
wendet werden, einen Alarm auszulösen, wenn ein Gegenstand oder eine Person ein definiertes Gebiet verlässt. So kann über die Steuerung der elektronischen Wegfahrsperre der Nutzungsbereich von Mietwagen auf das Inland oder bei einem Handy auf eine bestimmte Umgebung begrenzt werden. Auch für Personen, die Orientierungsschwierigkeiten haben oder unter elektronischer Aufenthaltsüberwachung stehen, kann Geofencing eingesetzt werden.

Um Geofencing verwenden zu können, werden folgende Komponenten benötigt:

- Informationen zur Definition des relevanten Gebiets bzw. der Grenze,
- ein System, das es ermöglicht, eine Position in der Fläche oder im Raum zu bestimmen, und
- ein Element, dessen Position in der Fläche bzw. im Raum ausgewertet wird.

Zur Positionsbestimmung können Satellitennavigationsysteme wie GPS verwendet werden.

Funktionen können beim Betreten oder Verlassen des „eingezäunten“ Gebiets aktiviert oder gesperrt werden.





GEOFENCING

Verwaltungsrelevanz:

■□□□

Umsetzungs- geschwindigkeit:

■□□□

Marktreife/Produkt- verfügbarkeit:

■□□□

Aber auch Mobilfunknetze oder bekannte WLANs sowie das im Energieverbrauch optimierte LPWAN bieten Orientierung. Bei der Verwendung ist darauf zu achten, ob die Techniken auch in Gebäuden genutzt werden können. Ggf. können dort zusätzlich akustische (s. Artikel „Ultraschall-Tracking“), optische (s. HZD-Trendbericht 2012) oder Funk-Signale (zu „Beacons“ s. HZD-Trendbericht 2015) genutzt werden. Die Anwendung, die auf die Position eines Objektes innerhalb oder außerhalb des Zaunes reagiert, kann entweder unmittelbar an dem Objekt (z. B. bei Auto oder Handy) oder als Zentralanwendung (z. B. bei Alarmierung) implementiert sein. Im letzteren Fall muss neben dem System zur Positionsbestimmung im relevanten Gebiet auf jeden Fall auch ein Netz zur Übermittlung der Positionsdaten vorhanden sein. Für den Einsatz von Geofencing bei kritischen Anwendungen ist zudem zu beachten, dass Signale zur Positionsbestimmung gestört werden können, was zu Desorientierung führt. Auch das Fälschen von Positionssignalen ist

möglich: So gab es zeitweise durch sog. GPS-Spoofing in größerem Umfang ein Problem bei einem Handy-Spiel, das ortsbezogen im Freien gespielt wurde. Mit solchen Techniken ließe sich beim Geofencing auch eine Position innerhalb oder außerhalb des spezifizierten Gebietes simulieren.

Bewertung

Es gibt in der öffentlichen Verwaltung viele Leistungen, die einen ortsbezogenen Radius besitzen. Die meisten Leistungen beziehen sich allerdings auf die offizielle Meldeadresse und nicht auf den Standort eines IT-Gerätes (z.B. Elterngeld). Interessant könnte das Geofencing jedoch im Zusammenhang mit kontextbezogener Sicherheit sein, die neben anderen Informationen auch den Ort auswertet. So könnten Anwendungen oder Funktionen z. B. nach einem zeitlichen und räumlichen Profil verfügbar gemacht oder gesperrt werden.

EDGE, FOG, MULTI-CLOUD

Als die Cloud-Technologie die Schlagzeilen der IT-Presse eroberte, verbanden viele damit die Hoffnung auf eine einfache, sichere und kostengünstige Infrastruktur. Für einzelne Anwendungen funktioniert das in der Regel ganz gut. Doch auch hier zeigte sich einmal mehr, dass es für komplexe Aufgaben nur ganz selten einfache Lösungen gibt, denn die Cloud-Technologie bringt eigene Herausforderungen mit sich – auch jenseits der vermeintlich übertriebenen Sorge um den Datenschutz. Das liegt zum einen an der Frage, was von der Cloud als Dienstleistung erwartet wird: Infrastruktur, Systemplattformen oder (Anwendungs-)Software stellen als IaaS resp. PaaS und SaaS nur eine grobe Kategorisierung der Angebote dar. Je heterogener und spezifischer die zu betreibenden Anwendungen sind, desto mehr liegt der Schwerpunkt der Clouddienste auf der automatisierten Bereitstellung von Infrastruktur.

Zum anderen spielen auch Art und Verwendung der Endgeräte, die mit der Cloud verbunden werden sollen, eine Rolle. Diese können von Terminal-artig benutzten PCs über Mobilgeräte mit Apps bis zu Elementen des Internet-of-Things aus einem breiten technischen Spektrum kommen. Und auch ihre Verwendung kann ganz unterschiedlich ausfallen. Dabei spielen Fragen eine Rolle, wie:

- Welche Daten bzw. Datenmengen fallen vor Ort an?
- Welche Daten werden vor Ort benötigt?
- Wie lange werden Daten vor Ort benötigt?
- Welche Reaktions- und Antwortzeiten sind erforderlich?
- Welche geografische Abdeckung hat die Anwendung?

Diese Themen beschäftigen die Mobilfunkanbieter schon seit geraumer Zeit, denn sie müssen den Transport von Daten zwischen den Endgeräten und den zentralen Komponenten bewerkstelligen – unabhängig davon, ob deren Ziel „die Cloud“ oder ein spezifisches Rechenzentrum ist.

Mit dem Internet der Dinge und insbesondere mit dem taktilen Internet (s. HZD-Trendbericht 2017) entwickelten sich Architekturen, die wieder mehr Rechenleistung zwecks Vorverarbeitung der Daten und Entlastung von Netzen näher am Endgerät vorsehen. Beim sog. Edge-Computing („edge“ engl. für „Rand“) wird Anwendungsintelligenz an den Rand der Gesamtinfrastruktur verlagert – und zwar an die Orte, an denen Daten anfallen. Das ist im Extremfall das Endgerät selbst. Beispiele dafür können vernetzte Fahrzeuge, per Fernservice gewartete Maschinen oder verteilte Videosysteme sein, die Bilder vor Ort analysieren anstatt das gesamte Videomaterial über das Netzwerk in die Cloud zu senden. Solche gerätebezogenen Teile der Anwendungslogik können sehr unterschiedliche Anforderungen an die dafür benötigte Hardware stellen. Diese reicht daher vom Minirechner (s. Artikel Minirechner – klein aber oho!) bis hin zu leistungsstarken Geräten (z. B. Bildverarbeitung). Verschiebt man dagegen einen größeren Teil der Anwendung – einschließlich Logik, Speicherung, Netzwerkfunktionen und Systemsteuerung – an den Rand der Infrastruktur, so spricht man von Fog-Computing. Die Komponenten, die hier dezentral platziert werden, werden Cloudlet bezeichnet und können die Rolle eines Minirechenzentrums einnehmen.

Ein weiteres Anzeichen dafür, dass hinter dem Cloud-Computing nicht die eine, alles umfassende Wolke steht, ist im Begriff der Multicloud zu sehen. Dieser konkretisiert die Verwendung mehrerer, einander entsprechender Clouddienste von verschiedenen Anbietern. Dadurch soll zum einen die Verfügbarkeit der Daten und Anwendungen erhöht werden und zugleich die Abhängigkeit von einem einzelnen Anbieter der Clouddienste reduziert werden. Zusammen mit

der Kopplung privater, öffentlicher und spezifisch bereitgestellter Clouddienste in der Hybrid-Cloud ergeben die oben beschriebenen Ansätze ein breites Spektrum möglicher Systemarchitekturen. Die Frage, ob die Verteilung und Diversifikation von Clouddiensten durch eine größere Zahl an Angriffsvektoren zusätzliche Risiken birgt, wird untersucht. Wenn mehrere Elemente aus den beschriebenen Modellen verwendet werden, steigt die Komplexität und der Koordinierungsaufwand für das Gesamtsystem. Bei umfangreichen Systemlandschaften verbietet sich dann die „manuelle“ und individuelle Konfiguration der Komponenten. Hierzu wird untersucht, wie software-definierte Techniken – insbesondere in Bezug auf Netzwerke (SDN) – den Betrieb von heterogenen Edge-, Fog- und Cloud-Strukturen unterstützen können.

Bewertung

Die beschriebenen Techniken signalisieren einen Trend weg von der IT an einem Ort und aus einer Hand hin zu diverseren Angeboten. Dabei sind die Überlegungen zur Verteilung der Aufgaben und der Leistung nicht neu. Sie sind vielmehr die Weiterentwicklung oder Fortführung von bereits erprobten Mustern mit Satelliten und Zentralsystemen – jetzt unter Nutzung von Cloud- und Datenbanktechnologien. Mit der Zunahme der Digitalisierung von Verwaltungsleistungen muss auch hier bedacht werden, was in welcher Form zentral, was dezentral passieren kann und muss. Solche Überlegungen können bei bestehenden Portalen der Verwaltung oder jenen, die im Rahmen der Umsetzung des OZG entstehen, eine Rolle spielen. Selbst wenn die dem Endanwender bereitgestellten Leistungen als ein Gesamtsystem erscheinen, können sich dahinter viele verschiedene und verteilte Architekturelemente verbergen. Dazu können auch die vorgestellten Techniken gehören – wenn auch für die öffentliche Verwaltung nicht unbedingt aus „der“ Public Cloud. Die strategische Entwicklung von Verwaltungs-IT sollte jedoch all diese Ansätze in den Blick nehmen, ihre Entwicklung verfolgen und sie auf ihre Anwendbarkeit prüfen.

EDGE, FOG,
MULTI-CLOUD

Verwaltungsrelevanz:

■■■□

Umsetzungs-
geschwindigkeit:

■□□□

Marktreife/Produkt-
verfügbarkeit:

■■■□



02

ENTWICKLUNG // TECHNIK // NETZE



HAPTISCHE DISPLAYS

In vielen Arbeitssituationen ist ein hohes Maß an Konzentration erforderlich. Das ist oft damit verbunden, ein bestimmtes Objekt im Blick zu behalten, während es bearbeitet wird – z. B. bei einer Ultraschall-Untersuchung des Arztes, aber auch bei Schreibarbeiten am Computer. Im Idealfall bedient man die Technik, ohne hinzusehen. Auch im Straßenverkehr sollte die Aufmerksamkeit der Straße und den anderen Verkehrsteilnehmern gelten. Nicht umsonst hat der Gesetzgeber das Verwenden einer ganzen Reihe von Geräten während des Fahrens in § 23 Abs. 1a StVO verboten. Dazu zählen nicht nur Mobiltelefone, sondern auch die Systeme für Navigation oder „Infotainment“. Allerdings sieht die Straßenverkehrsordnung Ausnahmen vor, nämlich dann, wenn „zur Bedienung und Nutzung des Gerätes nur eine kurze, den Straßen-, Verkehrs-, Sicht- und Wetterverhältnissen angepasste Blickzuwendung zum Gerät bei gleichzeitig entsprechender Blickabwendung vom Verkehrsgeschehen erfolgt oder erforderlich ist.“ (§ 23 Abs. 1a – 2.b) StVO). Da ist es naheliegend, dass

Bedienung per Touchscreen findet sich in immer mehr Geräten. Doch gerade die virtuellen Tastaturen, die auf hochglanzpolierten Glasoberflächen dargeboten werden, machen die Eingabe ohne ständigen Blickkontakt schwierig. Abgesehen von der geringen Größe der einzelnen Tasten, Buchstaben oder Regler fehlt vor allem jedwede Rückmeldung, ob etwas – und erst recht was – eingegeben wurde. Auf mechanischen Tastaturen, bei denen die einzelnen Tasten spürbar gegeneinander abgegrenzt sind, bieten viele Layouts noch eine zusätzliche Orientierungshilfe, indem die Buchstaben „F“ und „J“ besonders ausgeformt sind. Die Haut ist gerade in den Fingern mit zahlreichen Druck- und Tastsensoren ausgestattet, die Druckstärke (Merkel-Zellen), veränderlichen Druck (Meissner-Körper) oder Vibrationen (Vater-Pacini-Körper) registrieren. Damit ist es möglich, minimale Erhebungen (einige μm) und sehr kurze Druckimpulse (einige ms) wahrzunehmen.

Das machen sich die Hersteller haptischer Displays zunutze. Diese ertastbaren Anzeigen integrieren eine optische Darstellung von Bedienelementen – etwa durch einen farbigen Bereich – mit verformbaren Elementen, die z. B. die Erhebung einer Taste und sogar den sog. Druckpunkt bei der Bedienung spüren lassen. Dazu können verschiedene Techniken eingesetzt werden:

Elektroaktive Polymere (EAP) sind chemische Stoffe, die ihre Gestalt ändern, wenn an sie eine elektrische Spannung angelegt wird. Damit lassen sich spürbare Erhebungen bis zu 3 mm dynamisch erzeugen, die zusätzlich in Vibration versetzt werden können. Die so steuerbaren Einheiten können sehr schmal sein (ca. 100 μm), sodass sich auch kleinste Details darstellen lassen. Neben Tasten können damit z. B. auch Schieberegler virtualisiert werden. Solche Erhebungen lassen sich auch erzeugen, indem winzige Kanäle in halbelastischen Kunststoffen mit Flüssigkeit gefüllt werden. Diese Mikrofluide können sehr fein gesteuert werden.

die Hersteller solcher Systeme versuchen, neben der Spracheingabe andere Möglichkeiten zur „blinden“ Bedienung zu schaffen. Zwar lassen sich einige Steuerelemente am Lenkrad unterbringen. Aber zum einen ist dort der Platz sehr begrenzt, und zum anderen haben diese Bedienelemente einen i. d. R. kleinen und spezifischen Funktionsumfang (z. B. Radio laut/leise). Komplexere Anwendungen benötigen ihre eigenen Interaktionsmöglichkeiten, die – je nach Hersteller – unterschiedlich ausfallen.

Mit den Smartphones und Tabletcomputern haben variable Eingabeschnittstellen schon vor Jahren Einzug in die Alltags-IT gehalten. Und die

Neben der Form der Bedienelemente lassen sich auch verschiedene Oberflächenstrukturen – z. B. rau oder glatt – virtualisieren. Das Verhalten

// Haptische Displays können Anwendungen und virtuelle Welten realistischer erscheinen lassen, indem sie Bedienelemente und Oberflächenstrukturen ertastbar machen und dynamisch an die Situation anpassen.

wie das Klicken am Druckpunkt einer Taste kann etwa durch Vibrationen simuliert werden. Dadurch kann auch bei Geräten mit starrer Oberfläche ein Feedback in Abhängigkeit vom Druck gegeben werden.

Noch einen Schritt weiter gehen haptische Bedienelemente, die frei im Raum positioniert werden. Dazu muss zunächst genau erfasst werden, wo sich die Hand des Benutzers befindet. Dann können mittels Ultraschall die Tastzellen in der Hand stimuliert werden. Da einige Zellen nicht nur die Druckintensität erfassen, sondern auch sehr fein auf die Veränderung reagieren, lässt sich so eine virtuelle Form ertasten, indem man darüberstreicht. Die Stimulation geht nicht so weit, dass der Eindruck entsteht, man könne ein Bedienelement – z. B. einen Drehknopf – wirklich anfassen. Jedoch lassen sich mit der Technik Interaktionen zur Steuerung realisieren. Während die weiter oben beschriebenen Displays eher die Benutzeroberflächen von mobilen Anwendungen anreichern, könnte die Ultraschall-Anwendung virtuelle Welten ein wenig realistischer machen. Gegenstände, die mit einer VR-Brille betrachtet werden, könnten dann auch ertastet werden. Ob sich damit auch komplexe Szenen

realisieren lassen, wie Science-Fiction-Fans sie von den „Holodecks“ der Fernsehserie Star Trek kennen, scheint aber doch eher fraglich.

Bewertung

Haptische Displays befinden sich noch im Entwicklungsstadium und werden prototypisch realisiert. Die einzelnen Teiltechniken sind zum Teil schon in Consumer-Geräten zu erleben – z. B. das Feedback mittels Vibration. Auch wenn es noch geraume Zeit zu dauern scheint, bis haptische Displays mit größerem Funktionsumfang alltagstauglich werden, sind diese Entwicklungen ein weiteres Anzeichen dafür, dass sich auch komplexe Anwendungen der IT immer stärker vom klassischen Aufbau mit Computer, Bildschirm und Tastatur lösen, der auch den typischen Verwaltungsarbeitsplatz dominiert. Für spezifische Anwendungen von maßgeschneiderten haptischen Displays dürfte der Bedarf in der öffentlichen Verwaltung bisher gering sein. Hier werden sie wohl eher als künftige Standardkomponenten Einzug halten, sofern sich die Techniken etablieren.

HAPTISCHE DISPLAYS

Verwaltungsrelevanz:

■ □ □ □

Umsetzungs-

geschwindigkeit:

■ ■ ■ ■

Marktreife/Produkt-

verfügbarkeit:

■ ■ □ □

SOFTWARE BELEBT DIE NETZE

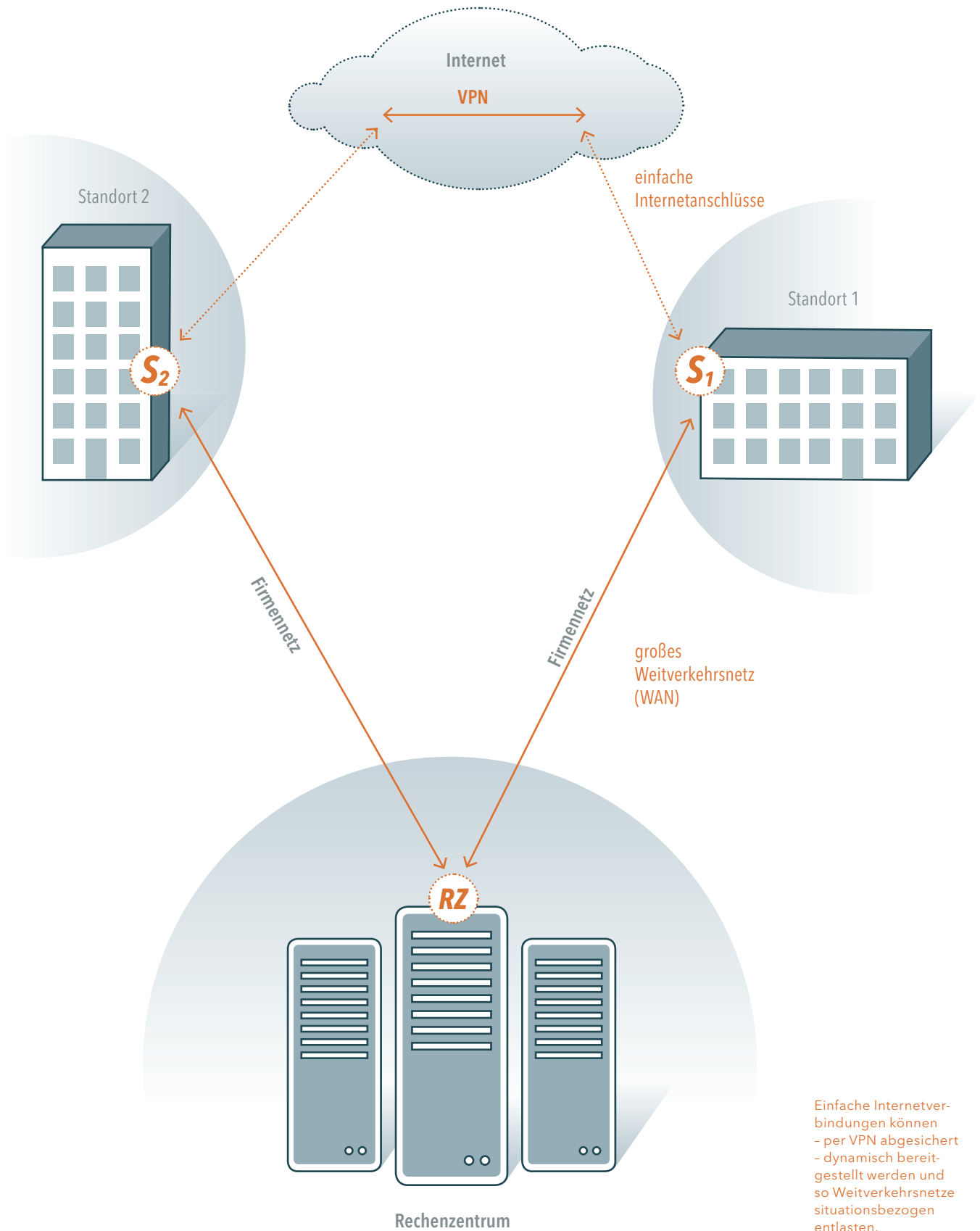
Datennetze verbinden heute viele Arbeitsplätze und Privathaushalte in aller Welt miteinander. Dank Funktechniken können wir sogar im Garten oder an abgelegenen Orten der Welt an den Datenströmen teilhaben. Damit das funktioniert, braucht es viele Leitungen und Funkkanäle – und eine Steuerung, die die Daten durch diese Verbindungen an ihren Bestimmungsort dirigiert. Vielerorts sind die Funktionen dieser Steuerung in separat eingerichteten und individuell konfigurierten Komponenten – z. B. einem Router, einem Load-Balancer oder einer Firewall – untergebracht. Hat man ein umfangreiches Netzwerk mit unterschiedlichen Anforderungen in verschiedenen Zonen, kann diese individuelle Konfiguration aufwändig werden – insbesondere, wenn sie Eingriffe vor Ort am Gerät erfordert.

Der Trend geht daher zu einer Netzwerksteuerung per Software – engl. „software defined networking“ (SDN, s. HZD-Trendberichte 2014 und 2016). Das Prinzip besteht darin, die Datenübertragungsebene (data plane) von der Steuerungsebene (control plane) zu trennen und die Steuerung von den Netzwerkkomponenten in eine zentrale Anwendung zu verlagern. Nachdem dies vorwiegend für lokale Netze (LAN) erforscht und entwickelt wurde, sollen nun auch Weitverkehrsnetze und mobile Netzwerke von dem Ansatz profitieren. Die beiden nachfolgenden Artikel stellen die Entwicklungen in diesen Bereichen dar.

SD – JETZT AUCH IM WAN

In Weitverkehrsnetze (kurz „WAN“ für engl. „wide area network“) verbinden räumlich weit auseinanderliegende Bereiche eines Netzwerks. Das können zum Beispiel die Niederlassungen einer Firma, Standorte einer Behörde oder redundant ausgelegte Rechenzentren sein. An den einzelnen Standorten sind in der Regel weitere, lokale Netzwerke (LAN) angeschlossen. Das WAN bündelt den Datenverkehr dazwischen und muss entsprechend der vorhandenen Anwendungen dimensioniert werden. Wenn neben einfachen, sporadischen Daten auch Sprache, Videokonferenzen, zeitkritische oder massenhafte Daten z. B. aus Sensornetzen übertragen werden sollen, wird eine hohe Bandbreite benötigt, um Störungen während des Transports zu vermeiden. Neben dem Totalausfall eines Netzdienstes können solche Störungen

etwa durch Überlast (engl. network congestion), Schwankungen in der Laufzeit (jitter) oder Paketverluste (losses) entstehen. Um trotzdem die nötige Servicequalität sicherzustellen, werden WANs i. d. R. spezifisch gestaltet und ggf. redundant ausgelegt. Dadurch muss die maximal benötigte Bandbreite mehrfach bereitgestellt werden, was mit entsprechenden Kosten verbunden ist. Viele Netzwerke leiten („routen“) ihren Datenverkehr über eine zentrale Stelle, selbst wenn bei der Kommunikation – z. B. zwischen zwei Standorten – ein kürzerer Weg möglich wäre. Das führt auf einigen Strecken zu mehr Datenverkehr. Eine Lösung für dieses Problem besteht darin, anstelle des zentralen Ansatzes (engl. „hub and spoke“ für „Nabe und Speichen“) einen stärker vernetzten zu wählen, bei dem alle Standorte miteinander verbunden sind



// Funktionen für Konfiguration und die Steuerung des WAN werden in einer Anwendung zentralisiert.

- oder zumindest diejenigen, die oft größere Datenmengen miteinander austauschen („fully –“ resp. „partially meshed“). Die Einrichtung und Pflege eines solchen Netzes kann jedoch aufwändig sein. In den lokalen Netzen werden die für die Bewältigung solcher Steuerungsaufgaben benötigten Funktionen daher zunehmend virtualisiert und ihre Steuerung wird durch Software zentralisiert, wodurch Software-definierte Netze (SDN) entstehen. Ähnliche Ansätze etablieren sich derzeit auch für WANs unter dem Namen software-defined WAN (SD-WAN). Dabei weisen SD-WANs insbesondere folgende Merkmale auf: Sie

- unterstützen hybride Netze mit verschiedenen Verbindungstechniken (z. B. MPLS oder LTE),
- ermöglichen das dynamische Routing auf Basis von anwendungsorientierten Regeln (policies),
- unterstützen die Verwendung von VPNs und Netzwerkdiensten.

Um dies nicht nur grundsätzlich möglich zu machen, sondern auch effizient und effektiv handhaben zu können, unterstützen sie die Einrichtung und Verwaltung des WAN-Verkehrs durch eine einfach zu bedienende Anwendung sowie die Automation der Konfigurationsprozesse. Im Idealfall soll das Anschließen eines neuen Standorts nur noch wenige „Handgriffe“ erfordern, nämlich dort eine Netzwerkkomponente aufzustellen und zu verkabeln, die sich dann automatisch ihre Konfiguration aus der zentralen Managementkomponente lädt und installiert. Auch die Einrichtung einzelner Netzdienste soll sich von durchschnittlich vier Monaten auf wenige Tage oder Stunden reduzieren.

Neben dem Komfort beim Netzwerkmanagement wird das Sparpotenzial von SD-WANs angesprochen. Dies lässt sich dann besonders gut erschließen, wenn der Datenverkehr – zumindest

in wesentlichen Teilen – anstelle von spezifischen Netzen über Internetverbindungen geleitet werden kann. Dabei können an einem Standort mehrere dieser vergleichsweise günstigen Anschlüsse von verschiedenen Providern und unterschiedlicher Qualität verwendet werden, um die benötigte Bandbreite und Verfügbarkeit zu erhalten. Ob dieser Weg gangbar und sinnvoll ist, hängt nicht nur von der Kritikalität der Daten bzw. der Anwendungen ab, sondern auch von der Dynamik der benötigten Verbindungen. Wenn spezifische Standorte mit garantierter Leistung, Kontrolle und Zuverlässigkeit versorgt werden müssen, sind die privaten Verbindungen ggf. die bessere Wahl. Wenn sich aber die Anforderungen der Standorte phasenweise ändern oder Teile der Anwendungen in dynamischen Cloud-Umgebungen betrieben werden, überwiegt evtl. der Aspekt der Flexibilität.

Bewertung

SD-WAN-Techniken und ihre Einsatzszenarien sind noch in der Entwicklung begriffen – auch wenn es schon Lösungen zu kaufen gibt. Ob und in welchen Situationen flexible Netze unter Einbeziehung öffentlicher Internetverbindungen für Anwendungen der Verwaltung einsetzbar sind, hängt vom Schutzbedarf der zu transportierenden Daten ab. Die zurückhaltende Nutzung von Diensten aus Public-Clouds seitens der Verwaltung lässt diesen Anwendungsfall von SD-WAN-Techniken eher als Ausnahme oder als Ergänzung zu bestehenden Verwaltungsnetzen erscheinen. Versteht man SD-WAN als Teil einer übergeordneten SDx-Strategie, sollte das Thema weiterverfolgt werden, denn von den wesentlichen Möglichkeiten der SD-WANs – einer komfortablen Steuerung und optimierten Nutzung der bestehenden Weitverkehrsnetze und einer höheren Agilität bei deren Einrichtung und Anpassung – würden Verwaltungen und ihre Kunden profitieren.

SD – JETZT AUCH
IM WAN

Verwaltungsrelevanz:

■■■□

Umsetzungs-
geschwindigkeit:

■■■□

Marktreife/Produkt-
verfügbarkeit:

■■■□

... UND „WIRELESS“ – SDWN

Mobile Netze besitzen gegenüber drahtgebundenen Netzen einige zusätzliche Herausforderungen für den Betrieb: z. B. eine oft sehr heterogene Landschaft mobiler Endgeräte, die Überlagerung von Signalen oder sich schnell ändernde Anforderungen an Übertragungsgeschwindigkeit und -menge. Da mobile Verbindungen für den letzten Netzabschnitt (last hop) immer häufiger genutzt werden, sind die Ausrüster von Funktechnik und die Anbieter mobiler Netze bereits gefordert, ihre Produkte möglichst flexibel an diese Dynamik anzupassen. Dabei spielt die Bereitstellung virtualisierter Netzdienste (engl. network function virtualization, NFV) im mobilen Bereich eine zunehmend wichtige Rolle. Da liegt es fast auf der Hand, diese Aufgaben von zentralen Steuerungen aus wahrzunehmen. Entsprechende Anwendungen sollen es dem Administrator ermöglichen, einen einheitlichen Überblick über den Zustand des gesamten Netzes bzw. von Netzabschnitten zu gewinnen, z. B. bezüglich Verfügbarkeit, Übertragungsqualität und -geschwindigkeiten oder Nutzerstandorten. Sowohl bei der Analyse des Netzes als auch bei der Konfiguration soll die zentrale Steuerungsanwendung die Komplexität und Heterogenität des Netzes kapseln und zwar sowohl in technischer Hinsicht als auch nutzungsbezogen. Dadurch kann das mobile Netz kurzfristig aufgrund von Nutzerrollen, Anwendungen oder des aktuellen Zustandes angepasst werden. So können z. B. massive Datenströme einer Anwendung oder „Power User“ anders bedient werden, als wenn nur vereinzelt Daten übertragen werden müssen. Oder Datenströme werden bei Engpässen auf andere Übertragungskanäle verlagert. Durch die zentrale Steuerung können der Zustand und die Möglichkeiten des Gesamtnetzes in die Entscheidungen einbezogen werden und sind nicht anhand „statischer“ technischer Gegebenheiten auf den Netzwerkkomponenten zu konfigurieren. Dort verbleiben ggf. nur noch zeitkritische Funktionen, die so nah wie möglich an den Zugängen zum mobilen Netz bearbeitet werden sollen.

Neben den Weiterentwicklungen des SDN-Standardprotokolls OpenFlow an die Bedürf-

nisse mobiler Netze werden auch Architekturen für mobile softwaredefinierte Netze entwickelt. Die gilt sowohl für den Mobilfunk – insbesondere im Hinblick auf 5G – als auch für Wi-Fi (WLAN). So werden in einem Modell z. B. virtuelle Access-Points definiert, die nutzerspezifisch bereitgestellt werden können.

// Drahtlose Verbindungen im letzten Netzabschnitt vor dem Endgerät müssen sich flexibel an verschiedene und sich schnell ändernde Anforderungen anpassen lassen.

Mit den Entwicklungen von Software-definierten Netzwerken in den Bereichen LAN und WAN sowie bei Funknetzen scheint eine einheitliche und anwendungsnahe Steuerung komplexer, heterogener Netzinfrastrukturen denkbar. Wie gut dies insgesamt gelingt, hängt davon ab, wie sich die verschiedenen Netztechniken weiterentwickeln und ob die verschiedenen SDN-Techniken zusammenwachsen. Wünschenswert wäre es.

Bewertung

Neben den verwaltungseigenen Netzen, die häufig wenig dynamisch sind, spielen für effizientes digitales Verwaltungshandeln alle Netze von und bis zu Kunden eine wichtige Rolle. Die Weiterentwicklung von SDN-Techniken für die „digitale Gesellschaft“ ist aus Verwaltungssicht also im doppelten Sinne erstrebenswert, nämlich sowohl im Hinblick auf die Steuerung eigener Netze, als auch bezüglich der Anbindung von Kunden. Insbesondere die werden zunehmend mobil. Für den öffentlichen Bereich stehen dabei neben Eigenschaften wie einer leichten Bedienung oder einfachen Bereitstellung neuer Dienste auf allen Kanälen vor allem die Steuerung der Sicherheit und Zuverlässigkeit der Netze im Fokus.

... UND
„WIRELESS“ – SDWN

Verwaltungsrelevanz:

■ ■ ■ □

Umsetzungs-
geschwindigkeit:

■ ■ ■ □

Marktreife/Produkt-
verfügbarkeit:

■ ■ ■ □

MINIRECHNER – KLEIN ABER OHO!

Wer angesichts der Überschrift denkt, dass der folgende Beitrag lediglich „Spielzeugcomputer“ beschreibt, sollte sich gut 50 Jahre zurück an Bord einer Apollo-Rakete versetzen. Hier, zwischen Himmel und Erde, versah ein Navigationscomputer seinen Dienst – der Apollo-Guidance-Computer (AGC). Dieser Computer, der eigenständig Rechenprozesse priorisieren konnte, verfügte über 4 kB RAM, 36 kB Dauerspeicher und eine CPU, die mit rd. 1 MHz getaktet war. 20 Jahre später holte der erste PC Rechenleistung an den gewöhnlichen Büroarbeitsplatz. An Bord waren 64 kB RAM, eine CPU mit 4,77 MHz Taktung und sagenhafte 360 kB Speicher auf Disketten – je nach Modell in ein oder zwei Laufwerken.

Heutzutage tragen wir das Hunderttausend- oder Millionenfache der AGC-Rechenleistung in Form von Smartphones mit uns in der Hosentasche herum, um damit Selfies zu machen und Katzenfotos anzusehen. Das macht deutlich, dass es bei der Bewertung von IT nicht allein auf die Menge der Technik ankommt, sondern auch deren Verwendung.

// Minicomputer reichen vom Universalrechner bis zum Spezialgerät. Ihre Nähe zur Hardware macht sie für viele Anwender interessant.

Die räumliche Größe von Rechnern spielt aber eine Rolle, wenn es gilt, neue Anwendungsgebiete von IT zu erschließen. Mit Maßen von ungefähr $60 \times 30 \times 15 \text{ cm}^3$ und einem Gewicht von gut 30 kg war der AGC nicht gerade handlich und die klassischen Arbeitsplatzrechner wurden nicht umsonst „Desktop-PC“ genannt. Seit einigen Jahren werden immer wieder kleine Rechner auf den Markt gebracht, die – da zu meist ohne Gehäuse – an die romantisch verklärten Garagenwerkstätten aus den Pionierzeiten der PCs erinnern. Ähnlich wie damals werden sie mit mehr oder weniger viel Zubehör ausgeliefert. Wer einen solchen Rechner in Betrieb nehmen will, muss sich durchaus mit der Frage beschäftigen, wie eigentlich die Soft-

ware auf einen Rechner gelangt und wie man ggf. vorhandene Hardwarekomponenten zum Laufen bekommt. Manche der Kleinrechner können durchaus als Heim-PC genutzt werden und verfügen über ansehnliche Leistungsmerkmale (s. Kasten).

// Bei der HZD-Hausmesse 2017 wurden für ein IoT-Szenario mehrere Raspberry Pi-Rechner mit 64-Bit-Architektur, 1 GHz-CPU, 1 GB RAM und 256 GB Dauerspeicher als Sensornetzwerk, WLAN-Hotspot und Auswertungsrechner eingesetzt. Die „RasPi“-Rechner wurden von der wohlthätigen britischen Raspberry Pi Foundation entwickelt, die sich die Förderung der Ausbildung in der Informatik (engl. „computers, computer science and related subjects“) zum Ziel gesetzt hat.

Es ist aber gerade ihre Nähe zur Hardware, die die Minicomputer für viele Anwender interessant macht. Dank Steckverbindungen z. B. mit GPIO-Anschlüssen (GPIO = general purpose input/output) ist es relativ einfach möglich, Sensoren auszulesen oder andere Geräte zu steuern.

Die Zahl der Geräte und Gerätevarianten bei den Kleinrechnern ist sehr groß und der Markt entsprechend schwer zu überschauen. Daher wollen wir im Folgenden drei Gruppen von Kleinrechnern betrachten.

Da sind zum einen die „Universalrechner“, die sich wie ein Heim-PC nutzen lassen. Diese können i. d. R. mit verschiedenen Betriebssystemen ausgestattet werden, wobei Linux-Varianten mit oder ohne grafischer Benutzeroberfläche am häufigsten sind. Das System wird zu meist auf einer SD-Karte installiert, deren Größe dann auch den verfügbaren Dauerspeicher bestimmt. Alternativ können die Systeme auch so konfiguriert werden, dass das Gerät direkt als Medienserver verwendet werden kann. Je nach Gerätetyp finden sich auf diesen Minirechnern Ethernet-Anschluss, WLAN oder Bluetooth, USB-Anschlüsse und Videoausgänge wie

HDMI, DVI-D und andere. Aus manchen Minirechnern lassen sich Cluster aufbauen, auf denen dann Anwendungscontainer oder Frameworks für verteilte Anwendungen wie Hadoop laufen können. Insbesondere der leichte Zugang zur Hardware macht diese Rechner auch für die sog. Maker-Szene interessant, in der viele originelle Anwendungen – vom smarten Spiegel mit Info-Display bis zur Heimautomation mit IoT-Technik – gebaut werden.

An dieser Stelle kommt dann eine zweite Gruppe von Minirechnern, die Development-Boards, ins Spiel. Bei vielen hardwarenahen Anwendungen werden Anschlüsse für Bildschirm, Maus und Tastatur oder Netzwerk gar nicht benötigt, sondern stören im Gegenteil durch Energieverbrauch oder Zeitverzögerungen. Dem entsprechend spartanisch sind diese Entwicklungsrechner ausgestattet. Da aber auch sie über Steckverbindungen verfügen, die Zugriff auf das System erlauben, gibt es zahlreiche Hardware-Module, die per Kabel angeschlossen oder direkt nach dem Rucksackprinzip mit ihrer Platine aufgesetzt werden können. In Folge der kompakten Bauweise lässt sich mit diesen Kleinstcomputern Rechenleistung in andere Geräte einbauen. Während bei den universellen Computern der ersten Gruppe die Programmierung auch auf dem Gerät erfolgen kann, wird bei dieser Gerätegruppe die Software i. d. R. extern bereitgestellt und dann in einen nicht-flüchtigen Speicher übertragen. Aus dem wird sie dann beim Einschalten des Minirechners gestartet, und dieser dann autonom ohne Benutzerschnittstelle betrieben.

Wer mit PC und Smartphone aufwächst, muss erst einmal ein Gespür dafür bekommen, dass Computer programmiert werden können und müssen. Die Geräte sind von sich aus „dumm“ und nicht für jede denkbare Aufgabe gibt es eine fertige App. Um das technische Verständnis von Computern über die anwendungsbezogene Medienkompetenz hinaus gerade bei Kindern und Jugendlichen zu fördern, wurden verschiedene Lerncomputer entwickelt. Diese sind ähnlich kompakt wie die Development-Boards und verfügen über einfache LED-Anzei-

// Minicomputer sind gut als Lerngeräte geeignet, wenn es darum geht, Abläufe im Rechner und bei der Verbindung mit der Umwelt kennen zu lernen.

gen und einige Sensoren. Darüber hinaus können weitere Sensoren oder andere Komponenten wie Displays oder Motoren über GPIO-Steckverbindung oder ähnliche Kopplungen angeschlossen werden. Das Kennenlernen der Technik und ihrer Nutzungsmöglichkeiten wird von Programmierumgebungen unterstützt, die auf die Hardware abgestimmt sind. Diese erlauben zum Teil die grafische Programmierung, bei der Anweisungen, Laufscheifen und andere Programmelemente – z. B. für die Sensoren – als Bausteine zusammengefügt werden. Auf diese Weise lassen sich schnell und einfach erste Schritte machen (z. B. blinkende LED), aber auch komplexe Anwendungen (z. B. Wetterstation) bauen.

Bewertung

Computer haben sich heute zu Alltagshelfern entwickelt, die mal mehr, mal weniger sichtbar überall ihre Dienste tun. In den meisten Fällen reicht es aus, wenn sie da sind und funktionieren. Im Büroalltag einer Verwaltung sind sie ein normales Arbeitsmittel und verrichten in manchen Verfahren noch an Stellen ihren Dienst, an denen sie nicht einmal mehr wahrgenommen werden. Insofern sind Minicomputer als unsichtbare Geräte auch jetzt schon Teile der Verwaltungs-IT. Das „Selbermachen“ steht dabei nicht im Fokus. Wenn es aber darum geht, innovative Ideen für den IT-Einsatz in der Verwaltung – jenseits von elektronischen Akten – zu entwickeln, können mithilfe der Minicomputer spannende Einblicke in die faszinierenden Möglichkeiten von Technik gewonnen werden. Und darüber kann man dann auch mit Kindern und Enkeln fachsimpeln.

MINIRECHNER

Verwaltungsrelevanz:

■ □ □ □

Umsetzungs- geschwindigkeit:

■ ■ ■ ■

Marktreife/Produkt- verfügbarkeit:

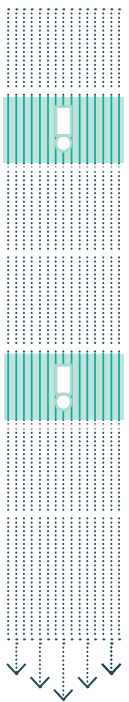
■ ■ ■ ■



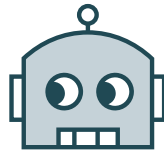
03

DATEN // INFORMATIONEN // SOFTWARE

Datenstrom



BOT



(a) Bot „beobachtet“ Datenstrom.
Bei bestimmten Ereignissen werden
Daten aufgenommen.

Datenstrom

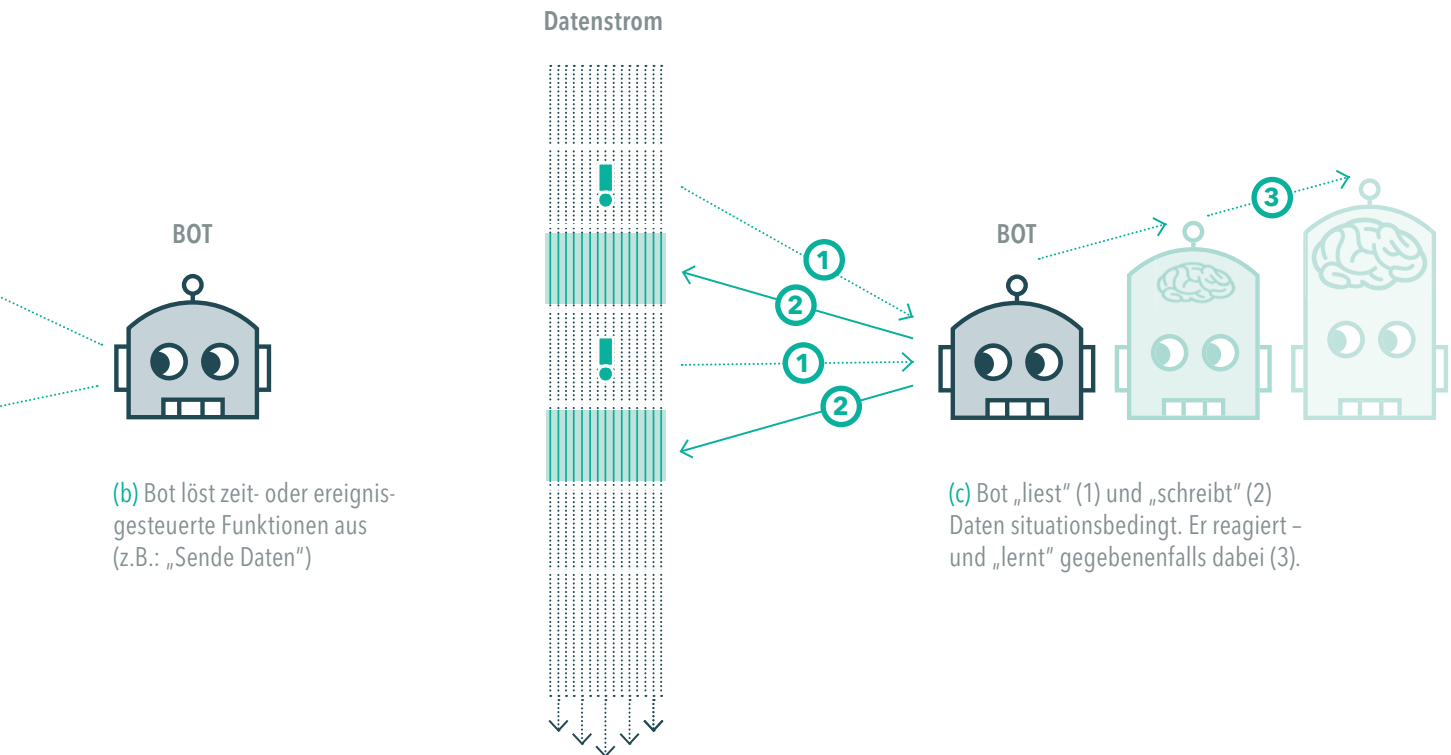


„WIR SIND DIE ROBOTER“

// Ich bin dein Sklave, ich bin dein Arbeiter,
wir sind auf alles programmiert und was
du willst, wird ausgeführt“

So sang vor gut 40 Jahren die Elektropop-Band Kraftwerk in ihrem Lied „Roboter“. Weitere zwölf Jahre davor hatte der Informatiker Joseph Weizenbaum sein Programm ELIZA entwickelt, das heute als der erste Chat-Bot (etwa „Dialog-Roboter“) gilt. Beide Werke weisen auf wesentliche Aspekte hin, die Roboter ausmachen: Roboter sollen für uns arbeiten und sie sollen sich „irgendwie klug“ verhalten. Schon lange vor C3PO aus den Star-Wars-Filmen wurden Roboter als Maschinenmenschen gestaltet – mehr oder weniger kastenförmig, oft mit kantigen Bewegungen und einer monoton knarrenden Blechstimme. Inzwischen verfügen zwei und vierbeinige Roboter über eine bemerkenswerte Beweglichkeit. Sie können menschen- bzw. tierähnlich laufen, Lasten tragen, über Hindernisse springen oder selbständig wieder aufstehen, wenn sie umgestoßen werden. Unabhängig von

der Motorik hat sich bei Robotern eine weitere Fähigkeit entwickelt, die sich jenseits von Literatur und Film erstmal in dem oben erwähnten ELIZA-Programm zeigte, nämlich die Dialogfähigkeit. Das Bedürfnis, mit einem Roboter auch über die Programmierung und das Ausführen des Gewollten hinaus per Sprache interagieren zu können, scheint groß. Das Bindeglied zwischen der mechanischen und der dialogfähigen Variante des Roboters ist Software. Auch ohne die Integration in Menschmaschinen hat diese Kategorie von Programmen in den letzten Jahren unter dem Namen „Bots“ großen Aufschwung erlebt. Dabei gehören viele Bot-Techniken schon seit geraumer Zeit zum Standard-Repertoire der Informationsverarbeitung. Da sind z. B. Webcrawler, die eigenständig Informationen von Webseiten sammeln und in Linksammlungen oder Suchmaschinen zur Verfügung stellen. Schon diese einfache Anwendung zeigt, dass Software-Roboter zwei Seiten haben können: die Sammelei kann gutartig sein und der Verbreitung von Wissen in einem spezifischen Fachgebiet dienen. Aber auch die Vorbereitung von Spam-Verteilung oder kriminellere Aktivi-



Bots können z. B. Informationen aus Datenströmen sammeln (a), Informationen verbreiten (b) oder in „intelligenten“ Dialogen auf Situationen reagieren (c).

täten können unterstützt werden, wenn gezielt personenbezogene Daten gesammelt werden. Dabei agieren Bots nicht nur jeweils allein, sondern vernetzen sich auch untereinander. Zumeist werden über solche Bot-Netze Schadprogramme verbreitet, über die Massen-Spam versandt oder verteilte Denial-of-Service-Attacken durchgeführt werden. Unabhängig von der Intention, mit der die Bots entwickelt werden, dienen sie dazu, Routinearbeiten schnell und einfach zu erledigen und sind somit letztlich Automationswerkzeuge. Während viele derartige Bots im Stillen agieren, sind in den letzten Jahren zunehmend Bots entstanden, die Kommunikation automatisieren sollen. Insbesondere in den sozialen Netzwerken werden solche Bots verwendet, um Themen zu platzieren und Nachrichten zu verbreiten. Solche „social bots“ können z. B. Nachrichten aus einer externen Datenquelle so aufbereiten, dass die im Netzwerk verbreitet werden – etwa als Kurznachricht der definierten maximalen Länge und mit den passenden Schlagworten (sog. Hashtags). Andere social bots finden Nachrichten, die bestimmte Schlagworte oder Begriffe enthalten, und verteilen

diese weiter. Dadurch kann Werbung verbreitet oder eine politische Position verstärkt werden. Da entsprechende Botschaften glaubwürdiger wirken, wenn sie von einer Person stammen und nicht als Bot-Nachricht erkennbar sind, wird von den Bot-Betreibern versucht, ihren Kommunikationsagenten ein menschliches Aussehen zu geben – mit Hilfe von Bildern (sog. Avatare) und eingestreuten anderen Nachrichten, oder indem sie die Meldungen anderer Nutzer abonnieren (ihnen „folgen“). So sollen der Wahlkampf bei der US-amerikanischen Präsidentenwahl und die Abstimmung über den Brexit in Großbritannien in den sozialen Netzwerken wesentlich durch Bots beeinflusst worden sein.

Die menschliche Ausstrahlung einer Software-Maschine wird dann erhöht, wenn sie nicht nur einseitig Nachrichten verbreitet, sondern auch – anscheinend intelligent – auf Botschaften oder Fragen eines tatsächlich menschlichen Partners reagiert. Um feststellen zu können, ob eine Maschine intelligent ist, entwickelte der Logiker und Mathematiker Alan Turing bereits in den 1950er Jahren einen Test – zumindest als Gedan-

kenexperiment. Danach sollte eine Maschine als intelligent gelten, wenn sie Dialoge mit Menschen so gestalten könne, dass ihr Gesprächspartner nicht erkennen kann, ob er mit einem Menschen oder eben einer Maschine kommuniziert. Die Aussagekraft dieses Tests im Hinblick auf die Frage, was Intelligenz ist, ist zwar umstritten. Jedoch gilt der Test auch ein halbes Jahrhundert später noch als so interessant, dass er in den 2000er Jahren in verschiedenen Wettbewerben um intelligente Dialogsysteme angewendet wurde. 2014 wurde dann erstmals verlautbart, ein Chat-Bot namens „Eugene Goostman“ habe den Turing-Test bestanden. Diese Feststellung gilt jedoch als fraglich, da der Chat-Bot einen Sprecher imitiert, der nur über eingeschränkte Kenntnisse in der verwendeten Sprache verfügt.

Viele Chat-Bots wurden in der Vergangenheit anhand von schlagwort- und phrasenorientierten, starren Regelwerken realisiert. Seit die künstliche Intelligenz (KI) dank leistungsfähiger Rechner und nahezu unbegrenzter Rechenleistung in Cloud-Systemen eine Renaissance erlebt, werden auch Dialogsysteme damit ausgestattet und sollen der maschinellen Kommunikation mit menschlichen Partnern eine neue Qualität geben. Dabei sollen die Systeme selbständig lernen, Gesprächsinhalte und -situationen zu erkennen und passend darauf zu reagieren. Sofern der Dialog nicht nur per Tastatur geführt wird, sondern durch Spracheingabe und ggf. mit Videounterstützung, sollen künftig neben der rein syntaktischen Prüfung des Textes auch Stimmlage oder Mimik und Gestik eines menschlichen Gesprächspartners in die Analyse einbezogen werden.

Chat-Bots sollen vor allem in der Kundenkommunikation eingesetzt werden und sind schon jetzt auf manchen Internetportalen zu finden. Auch hier geht es um die Erledigung von Routineaufgaben, nämlich das Anliegen eines potenziellen Kunden derart vor zu qualifizieren, dass nur noch spezifische Fragen durch einen menschlichen Kundenberater – und zwar einen fachlich zuständigen – zu beantworten sind. Dabei geht die Phantasie mancher Protagonisten

soweit, dass in naher Zukunft Kundendialoge ganz ohne menschliche Beteiligung auskommen: Der Kunden-Bot, der die Bedürfnisse, Anforderungen und Lebensumstände seines Besitzers kennt, kommuniziert dann mit den künstlichen Intelligenzen von Dienstleistern, um das tägliche Leben zu regeln. Das mag uns heute noch unrealistisch vorkommen. Aber die Geschichte zeigt, dass manches Gedankenexperiment doch irgendwann Realität wird – wenn auch zunächst in kleinen Schritten.

Bewertung

Die öffentliche Verwaltung hat im Rahmen der Sachbearbeitung viele Routineaufgaben. Deren Automation wird in Fachverfahren ggf. schon unterstützt. Sicherlich ließen sich weitere – insbesondere verfahrensübergreifende – Arbeitsschritte durch Bots unterstützen und gerade in der Kommunikation mit Verwaltungskunden könnten deren Dienste zum beiderseitigen Vorteil eingesetzt werden: Kunden werden im Dialog spezifischer durch Antragsprozesse begleitet, als es anhand von Formularen – „... dann weiter mit Nr. 18a ...“ – möglich ist. Wenn ein solcher Bot einen echten Dialog unterstützt, können evtl. auf dieser Ebene auch schon Fragen des Kunden bearbeitet werden. Das ist auf jeden Fall besser, als ein falsch ausgefülltes Formular bearbeiten zu müssen. Eine bessere Datenqualität käme auf jeden Fall auch der Sachbearbeitung zugute. Auch Rückfragen an den Verwaltungskunden könnten im automatisierten Dialog bearbeitet werden. Dies kann dann auch asynchron erfolgen – ohne dass Sachbearbeiter und Verwaltungskunde gleichzeitig an dem Vorgang arbeiten, wie es etwa bei einer telefonischen Rückfrage erforderlich ist.

Für den Erfolg von Bots dürfte wesentlich sein, wie gut sie in der Lage sind, sich auf verschiedene Situationen einzustellen. Ausschlaggebend wird aber letztlich die Frage sein, an welcher Stelle wohl doch der Dialog mit einem menschlichen Experten einsetzen muss.

WIR SIND DIE
ROBOTER

Verwaltungsrelevanz:

■■■□

Umsetzungs-
geschwindigkeit:

■■■□

Marktreife/Produkt-
verfügbarkeit:

■■■□

VON DER WEB-APP ZUR APP-APP

Damit Computer auch jenseits der Systemtechnik Arbeit verrichten können, werden sie mit Anwendungssoftware bestückt. Auf Englisch heißt diese „application (software)“, und spätestens seit Smartphones und Tablets den Verbrauchermarkt für Computerhardware bestimmen, ist die Abkürzung „app“ für diese Art der Software weit verbreitet. Solche Anwendungen sind – sofern nicht ihr Quellcode verfügbar ist – in der Regel an technische Plattformen gebunden und können nicht ohne Weiteres auf Geräten aus anderen „Ökosystemen“ genutzt werden. Dementsprechend haben sich unter dem Namen App-Store Verteilportale etabliert, aus denen die „native“ Anwendungssoftware bezogen werden kann oder muss – je nach Plattform.

Einen ganz anderen Weg der Leistungsbereitstellung gehen Anwendungen, die über ein Webangebot genutzt werden können. Hierbei ist es in der Regel das Ziel, möglichst universelle Angebote ohne Bezug zu einem spezifischen System oder einer Gerätekategorie zu machen. So werden heute schon viele Webanwendungen so gestaltet, dass sie auf verschiedenen Endgeräten zwar evtl. unterschiedlich aussehen, aber gut nutzbar sind („Responsive Design“ zur Anpassung an Bildschirmgröße oder -ausrichtung). Mit Hilfe moderner Entwicklungswerkzeuge lässt sich heutzutage die Systemtechnik der Endgeräte – z. B. die eingebaute Kamera – in „Web-Apps“ besser verwenden, als noch vor einigen Jahren. Trotzdem haben die Web-Apps gegenüber nativen Apps einen wesentlichen Nachteil: Sie benötigen eine permanente Netzverbindung, um zu funktionieren.

Wer für möglichst viele Nutzer eine optimale Anwendung zur Verfügung stellen möchte, muss also mehrfache Arbeit investieren, um verschiedene Endgeräte-Klassen, Ökosysteme und „Netzsituationen“ zu bedienen.

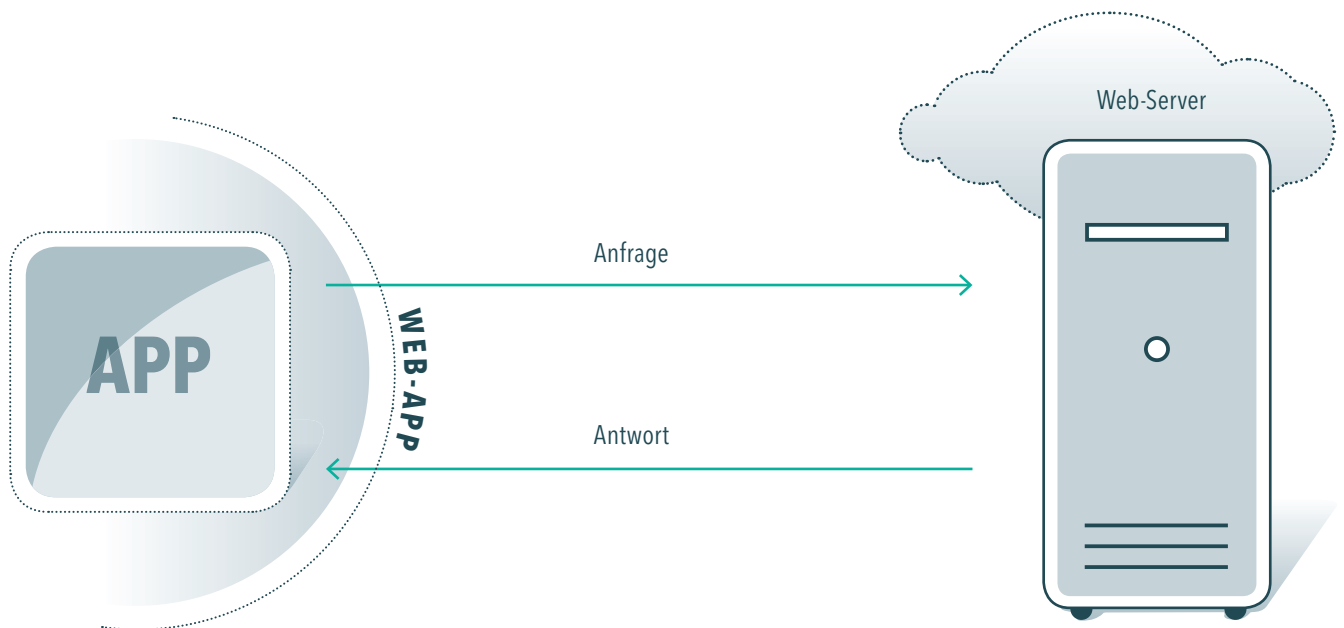
Abhilfe soll eine neue Generation von Web-Apps schaffen, die gleich mehrere Herausforderungen auf einmal angeht. Progressive Web Apps (PWA) sollen stets aktuell, anpassungsfähig und verlinkbar sein. Sie sollen sich auch ohne Netzverbindung und sicher nutzen lassen

und sich im Gebrauch wie native Apps anfühlen, d. h. insbesondere den Zugriff auf die Ressourcen des Endgerätes zulassen.

Der Anspruch ist hoch und lässt sich so ganz ohne Internetverbindung auch nicht realisieren. Kern der PWA-Technik ist nämlich ein umfangreicher Caching-Ansatz, bei dem Anwendungskomponenten beim ersten Start im Gerät gespeichert werden. Vergleichbar ist diese Zwischenspeicherung mit der von konventionellen Webseiten im Browser Cache, um Dauer und Menge der Datenübertragung beim Wiederaufrufen einer Webseite im Browser zu sparen. Ob die Progressive Web App einen Teil der Anwendung dann aus dem Cache nimmt oder neu lädt, steuert ein Software-Skript, der sog. Service-Worker. Dieser fängt die Netzwerkanfragen der Anwendung auf und leitet sie entweder an die gespeicherte Anwendung oder den Webserver weiter. Der Service-Worker wird in JavaScript implementiert. Aus Gründen der Sicherheit läuft er in einem eigenen globalen Kontext und kann keine Teile des Dokumenten-Objekt-Modells (DOM) verändern, das die Struktur der Webseite beschreibt. Außerdem funktioniert der Service-Worker nur dann, wenn die Webseite über eine verschlüsselte HTTPS-Verbindung geladen wurde.

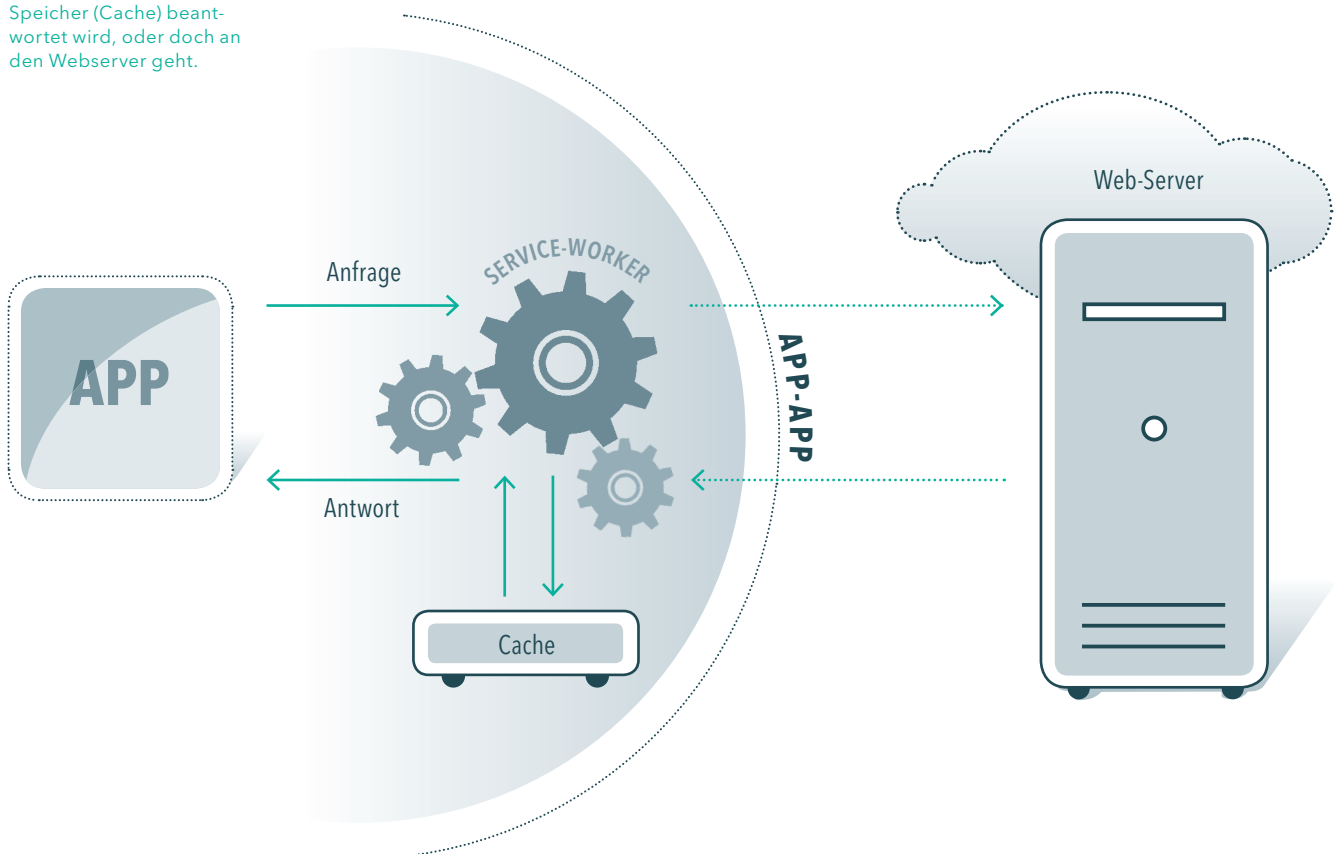
Ob es sich bei einer Web-Anwendung um eine herkömmliche Seite oder eine Progressive Web App handelt, kann der Browser am Web-App-Manifest erkennen. Ein solches Manifest ist eine im JSON-Format angelegte Datei, in der Eigenschaften der Webseite beschrieben werden. Dazu gehören Elemente wie ein Name und ein Kurzname der Anwendung, eine Beschreibung, die Sprache, Farben und Darstellungsoptionen oder auch Referenzen auf Icons. Darin kann auch der Service-Worker mit Quelle und Anwendungszusammenhang spezifiziert werden.

Das Progressive an PWAs besteht darin, dass die neuen Möglichkeiten nur dann verwendet werden, wenn sie von der Einsatzumgebung auch unterstützt werden. Ist ein Browser z. B. nicht in der Lage einen Service-Worker zu be-



In der klassischen Anwendung werden Anfragen der App vom Webserver beantwortet.

Bei der Progressive Web App entscheidet der Service-Worker, ob die Anfrage aus dem internen Speicher (Cache) beantwortet wird, oder doch an den Webserver geht.



nutzen, verhält sich die Anwendung wie eine herkömmliche Web-App und ist dann z. B. auf eine Netzverbindung angewiesen.

Neben der Offline-Fähigkeit sollen PWAs noch eine weitere Eigenschaft haben, die sie wie native, mobile Apps erscheinen lassen: Sie sollen installierbar sein. Das bedeutet, dass sich die PWA mit einem eigenen Icon auf dem Endgerät neben den anderen Anwendungen ablegen lässt. Dadurch kann sie direkt aus der Systemumgebung gestartet und ggf. ohne den Anwendungsrahmen des Browsers verwendet werden.

Progressive Web Apps werden durch die nachfolgend genannten Eigenschaften charakterisiert. Dabei scheinen die ersten sieben obligatorisch zu sein, die weiteren Eigenschaften sind nur bei einigen Definitionen vorhanden: PWAs sind ...

- Connectivity Independent: Sie lassen sich ohne permanente Netzverbindung verwenden.
- Discoverable: Sie lassen sich anhand ihrer Beschreibung im Manifest als PWA identifizieren – und können so z. B. in App-Katalogen gesammelt werden.
- Installable: PWAs lassen sich auf Ebene der Apps auf dem Endgerät bereitstellen und müssen nicht aus dem Browser gestartet werden.
- Linkable: Auf die PWA selbst und auf ihre Inhalte kann durch URLs referenziert werden.
- Re-engageable: Die PWA kann Push-Nachrichten senden, auch wenn sie nicht läuft.
- Responsive: Sie passt sich an die Darstellungsmöglichkeiten des jeweiligen Gerätes an.
- Safe: Die PWA lässt sich ausschließlich über HTTPS-Verbindungen nutzen.
- App-like: Die PWA lässt sich wie eine native App verwenden – einschließlich Audio, Video, Kamera, Touch-Input, Gerätesensoren usw.
- Fresh: Im Online-Betrieb kann die jeweils neuste Version der PWA bezogen werden.

- Progressive: PWAs passen sich an die Fähigkeiten des Browsers an und lassen sich insbesondere mit älteren Browsern als normale Web-App verwenden.

Mit der Bereitstellung von PWAs im Web und deren App-artiger Erscheinung könnte sich ein neuer Vertriebsweg für mobile Anwendungen etablieren, der die eingangs erwähnten App-Stores umgeht. Die Anbieter der Ökosysteme erwirtschaften jedoch auf diesem Weg einen Teil ihrer Erlöse, indem sie an den Verkaufsgebühren beteiligt werden. Trotzdem scheinen alle großen Plattformen das Konzept der Progressive Web Apps zu unterstützen. Die Einbindung der Techniken in spezifische Browser ist jedoch noch unterschiedlich weit fortgeschritten.

Bewertung

Das Thema Progressive Web App verzeichnet seit einiger Zeit zunehmendes Interesse. Die Möglichkeiten, Anwendungen mit einer guten Nutzungserfahrung auf einfachem Weg zu verbreiten – und das weitgehend losgelöst von technischen oder organisatorischen Rahmenbedingungen der Plattformbetreiber – scheinen verlockend. Zudem kann so jeweils die aktuellste Version der Software genutzt werden – auch ohne lange Updates, die sonst seitens der nutzenden Organisation verteilt werden müssen. Öffentliche Verwaltungen könnten von der Technik profitieren, wenn sie einzelne Anwendungen und Dienste auch außerhalb umfassender Portale bereitstellen wollen. Gerade für den mobilen Einsatz, bei dem nicht die umfängliche Anwendung von Fachverfahren im Vordergrund steht, könnten Progressive Web Apps eine Lücke in den Angeboten der Verwaltungen zu schließen helfen. Es lohnt sich also auf jeden Fall, die Entwicklung der Technik weiter zu verfolgen und ihre Möglichkeiten systematisch zu untersuchen.

VON DER WEB-APP
ZUR APP-APP

Verwaltungsrelevanz:

■■■■□

Umsetzungs-
geschwindigkeit:

■■■■■

Marktreife/Produkt-
verfügbarkeit:

■■■■□

„NO“ ODER „LOW“ – PROGRAMMIEREN (FAST) OHNE CODE

Programmierer sind schwer zu finden und Programmierer für mobile Anwendungen sind noch schwerer zu finden. Doch zum Glück gibt es jetzt neue Werkzeuge, mit denen jeder – auch ohne Programmierkenntnisse – Anwendungen und mobile Apps bauen kann. So lassen sich grob einige Marketingsprüche und Fachartikel zusammenfassen, die sich mit dem Thema „low-code“ oder „no-code“ befassen. Unter diesen Schlagworten werden Werkzeuge und Plattformen angepriesen, die den normalen IT-Anwender in die Lage versetzen sollen, mit Hilfe einer grafischen Oberfläche Software zu bauen. Dazu muss der „citizen developer“ nur per drag-and-drop die gewünschten Funktionen kombinieren und mit ein paar Klicks steht die App im Web und in den Softwareportalen für mobile, native Anwendungen bereit, ohne dass dafür auch nur eine Zeile Code programmiert werden muss („no-code“) – oder höchstens ganz wenig („low-code“). Doch eigentlich dürfte es den Hype um „no-code“ und „low-code“ gar nicht geben, denn das gleiche Problem – sieht man von der mobilen Komponente ab – wurde schon ab den 1980er Jahren mit CASE-Tools

waltung ein Beispiel: Auch wenn im E-Government noch viel mehr denkbar – und machbar – wäre, sind hier zahlreiche Verfahren oder Verfahrensschritte digitalisiert worden. Etwas näher am Puls der Zeit sind komplexe Infrastrukturen, die sich – durch Software gesteuert – akut und ggf. vorausschauend an verschiedene Situationen anpassen können. „Intelligente“ Netze waren vor 30 Jahren auch noch mehr Fiction als Realität.

Zum anderen hat Software an vielen Stellen Einzug in den Alltag gehalten. Die viel zitierte Digitalisierung findet nicht nur über komplexe Systeme und vollständige Fachverfahren statt. Im privaten wie auch geschäftlichen Umfeld sind es Massen an kleinen mobilen Helfern, die das Ausmaß der Veränderung spürbar machen. So scheinen neben Telefonbüchern, Stadtplänen und Terminkalendern auch Einkaufszettel, Fahrkarten und Zeitungen zunehmend von der papiergebundenen Bildfläche zu verschwinden. Dazu kommen zahlreiche Informationsangebote, die es uns ermöglichen, Dinge zu finden, die sonst fast nur zufällig zu entdecken oder mit Aufwand zu beschaffen waren, wie z. B. die Öffnungszeiten eines speziellen Ladens, die nicht online zu finden sind.

// Die Digitalisierung wird von komplexen Systemen und vielen kleinen mobilen Anwendungen getragen. Die dafür benötigte Software muss mit angemessenem Aufwand plan- und realisierbar sein.

(CASE: computer aided software engineering), den Programmiersprachen der vierten Generation (4GL) und der schnellen Anwendungsentwicklung (engl. „rapid application development“) gelöst. Sind no- und low-code also wieder nur neue Marketingmaschinen? Oder steckt doch mehr dahinter?

Zumindest ist klar, dass seit den 1980er Jahren der Bedarf an Software deutlich größer geworden ist. Zum einen traut man sich und der Technik heute zu, Probleme mit Software zu lösen, die damals nur mit viel Phantasie denkbar waren. Vielleicht ist hierfür gerade die öffentliche Ver-

All diese mal kleinen, mal größeren digitalen Helfer und Anwendungen wollen programmiert werden. Und da macht es sich durchaus bemerkbar, dass die Zahl der qualifizierten Entwickler, die eine solche Anwendung von Grund auf bauen können, nicht beliebig groß ist. Zudem wäre ggf. der Weg von einer App-Idee bis zu einem solchen Entwicklung oft so aufwändig, dass es unattraktiv würde, sie umsetzen zu lassen: Problem spezifizieren, Entwickler suchen, Angebot einholen, sich über Rechte an der Idee und der App einigen ... Vertrag, Projektplan, Abnahme ...

An dieser Stelle kommen dann die no- und low-code-Plattformen ins Spiel. Die bieten die Entwicklung von Anwendungen – insbesondere im mobilen Bereich – mit wenigen Klicks nicht nur für den professionellen Programmierer an. Einmal gebaut, kann die Anwendung für verschiedene technische Systeme bereitgestellt werden.

Es gibt eine ganze Reihe von Anbietern solcher Plattformen, die einen in Dauer und Umfang begrenzten Zugang zu ihren Entwicklungsumgebungen ermöglichen. Dank leistungsfähiger Cloud- und Netzwerk-Strukturen lassen sich solche PaaS-Angebote nach nur wenigen Klicks im Browser ausprobieren. Der Weg von einer Idee zu einem Software-Experiment ist also nicht mehr unbedingt mit der Beschaffung und Installation einer entsprechenden Entwicklungsumgebung verbunden, wie es in den 1980er-Jahren noch die Regel war.

In der Tat bieten solche Entwicklungsplattformen zahlreiche vorgefertigten Bausteine, die in der Regel als grafische Elemente zusammengefügt werden und so den Rahmen einer Anwendung ergeben. Insbesondere der Aufbau von Benutzeroberflächen sowie die Anbindung von Datenquellen und externen Anwendungen über REST-APIs werden dabei unterstützt. Wer die Beispiele der oft vorhandenen Online-Tutorien durchlaufen hat und seine eigene Idee umsetzen will, steht dann aber schnell vor Fragen wie: Welches Beispiel kann ich als Vorlage verwenden? Welche Module muss ich ggf. austauschen? (Wo) Gibt es für meine Aufgabe eine passende Komponente? Wie muss ich die anpassen, um meine Aufgabe zu lösen? Hier wird also oft die Kenntnis von Programmiersprachen und Software-Bibliotheken durch die Notwendigkeit ersetzt, passende Bausteine und Muster in den angebotenen Werkzeugkästen zu finden und zu konfigurieren. Und da setzt dann auch die Kritik einiger Marktbeobachter an: Wer eine no- oder low-code-Plattform für die Anwendungsentwicklung einsetzt, kann schnell in die Abhängigkeit vom Anbieter geraten („vendor lock in“). Diese wird besonders spürbar, wenn die auf der Plattform entworfene Anwendung außerhalb davon weiterentwickelt und z. B. um fachliche Logik in konventionell programmiertem Code ergänzt wird. Dann ist die Möglichkeit zum Re-Import in die grafische Entwicklungsumgebung nicht immer garantiert.

Ist den no-code und low-code-Plattformen also ein ähnliches Schicksal bestimmt, wie den 4GL-Tools und ähnlichen Ansätzen der Vergangen-

heit – nämlich nach einiger Zeit sang und klaglos wieder zu verschwinden? Das vorherzusagen dürfte so schwierig sein wie die Frage nach der Gesamtentwicklung der IT in den nächsten Jahren zu beantworten. Man kann aber trotzdem sagen, dass no- und low-code-Plattformen einen guten Ansatz bieten, um für aktuelle Anwendungsfälle und Techniken Lösungen zu entwerfen und ggf. auch marktfähig zu machen. Dabei

// Bei der Entwicklung komplexer Anwendungen können no- und low-code-Plattformen ein Baustein sein, der ggf. nur phasenweise zum Einsatz kommt.

gilt es genau zu prüfen, welches Betriebsmodell man nutzt – Cloud oder doch „on premise“, an welchen Stellen im Entwicklungsprozess sie ansetzen – Prototyp- vs. operative Entwicklung, oder welche Möglichkeiten, in den „klassischen“ Quellcode einzugreifen, gewünscht oder benötigt werden.

Bewertung

Der oben skizzierte Weg von einer Idee zur klassischen Softwarelösung ist gerade im öffentlichen Bereich stark ausgeprägt und hat evtl. noch die ein oder andere zusätzliche Schleife. Hier könnten Werkzeuge zur Ideenentwicklung und Erprobung von Lösungswegen sicher dazu beitragen, Softwarelösungen schneller zu konzipieren und umzusetzen. Nun ist aber nicht jedermann zum Entwickler geboren – egal ob unter Zuhilfenahme grafischer oder textbasierter Plattformen. Die passenden Werkzeuge in den Händen der richtigen Mitarbeiter könnten aber sehr wohl helfen, Lösungen für die digitale Verwaltung auf den Weg zu bringen. In diesem Sinne sind no-/low-code-Plattformen als Innovationswerkzeuge interessant.

NO CODE

Verwaltungsrelevanz:

■ ■ □ □

Umsetzungs-

geschwindigkeit:

■ ■ ■ ■

Marktreife/Produkt-

verfügbarkeit:

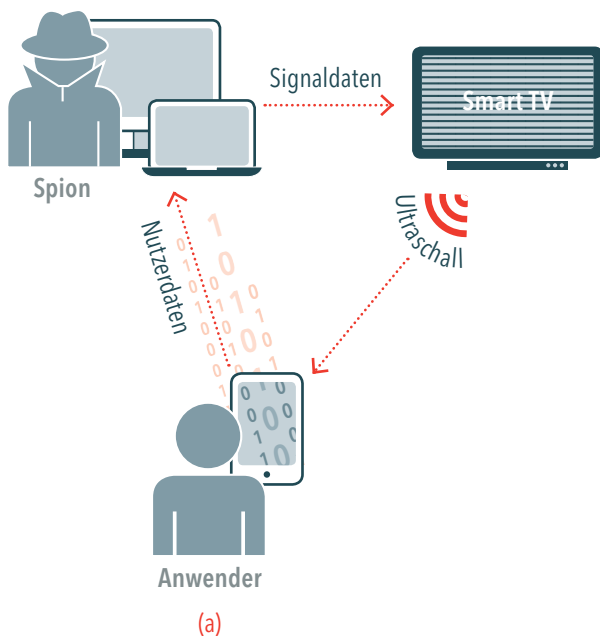
■ ■ ■ ■



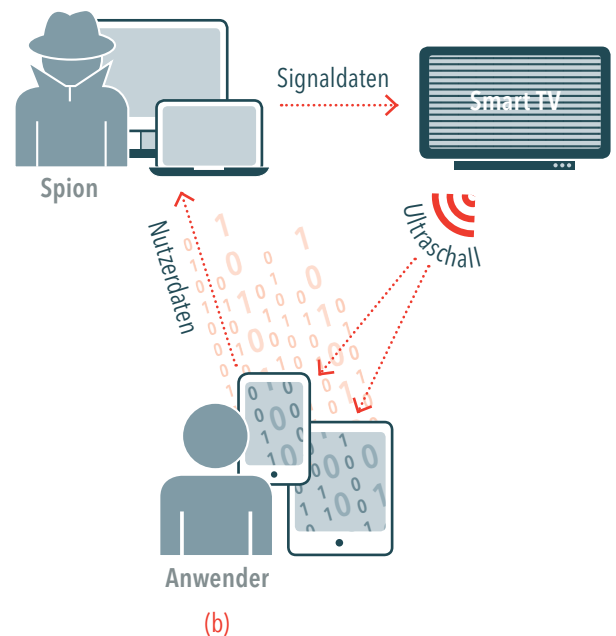


04

GESELLSCHAFT // SICHERHEIT // RECHT



Mit Ultraschall-Tracking lässt sich ermitteln, welche Medien genutzt werden. (a).



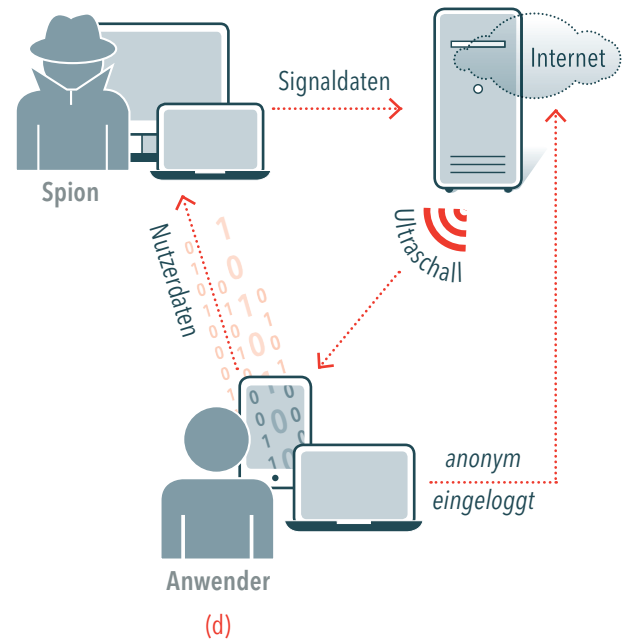
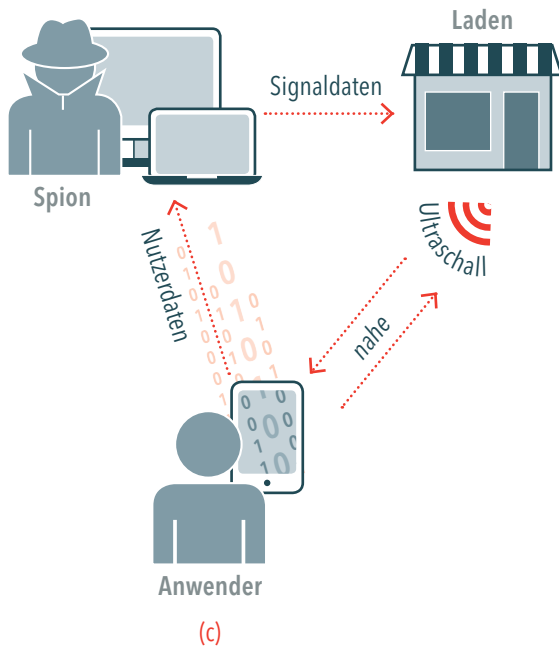
Cross-Device-Tracking erlaubt das Zusammenführen verschiedener Nutzungsprofile (b).

UNHÖRBARE DATENÜBERTRAGUNG PER ULTRASCHALL

„Ultraschall“ bezeichnet den Bereich von Tönen, deren Frequenz oberhalb des menschlichen Hörvermögens liegt. Als Richtwert dafür wird oft 20.000 Hz angegeben. Die tatsächliche Grenze der wahrnehmbaren Frequenzen verändert sich aber im Laufe des Lebens. Technische Anwendungen gibt es seit langem z. B. in der Medizin (Sonografie) oder in Ultraschallreinigungsgeräten. Zudem sind viele herkömmliche Lautsprecher und Mikrofone in der Lage, Ultraschalltöne zu erzeugen bzw. aufzunehmen. Diese Tatsache lässt sich dazu verwenden, Nachrichten zwischen verschiedenen Geräten akustisch, aber für den Menschen nicht hörbar zu übertragen. Und das kann man mit aber auch ohne Wissen der Anwender tun. Dass

dies nicht nur ein theoretisches Konzept ist, haben Forscher aufgedeckt, die in zahlreichen mobilen Apps entsprechende Komponenten zum Empfangen und Auswerten von Ultraschallsignalen fanden. Die über 230 identifizierten Anwendungen sind vermutlich auf mehreren Millionen Mobilgeräten installiert. Aber wozu soll das gut sein?

Die Nachverfolgung (Tracking) von Ultraschallsignalen über Mobilgeräte hat verschiedene Anwendungen: Relativ einfach hierbei ist die Medienauswertung. Dabei werden in Medien – z. B. Werbung, Fernsehsendungen, Webseiten – Ultraschallsignale versteckt und ausgestrahlt. Wenn ein mobiles Gerät die Signale „hört“ und diese



Cross-Device-Tracking erlaubt das Zusammenführen verschiedener Nutzungsprofile (c).

Anonyme Anwender können durch Ultraschall-Tracking enttarnt werden (d).

Tatsache weitermeldet, lässt sich daraus auf den Medienkonsum des Gerätebesitzers schließen.

Wird das gleiche Signal von mehreren Geräten eines Anwenders empfangen, kann der „Ultraschallspion“ erkennen, welche Geräte eine Person verwendet. Dieses sog. Cross-Device-Tracking erlaubt es dann, vollständigere Profile von Nutzern zu erstellen, weil sich etwa dienstliche und private Geräte oder mobile und stationäre Geräte einander zuordnen lassen. Dadurch können dann z. B. verschiedene Analyseergebnisse des Surfverhaltens zusammengeführt werden.

Beim Location-Tracking werden Ultraschallsignale nicht an beliebigen Orten wie beim Fernseh-

hen ausgesendet, sondern in einem räumlich eng begrenzten Gebiet – etwa im Umfeld eines Ladens. Verrät ein Mobilgerät, dass es die Signale empfängt, befindet es sich offenbar in der Nähe. Diese Methode wird verwendet, um Anwendern ortsspezifische Informationen – z. B. Werbung für die Produkte des Ladens – zu senden.

Und schließlich kann die Technik zur Enttarnung anonymer Anwender genutzt werden. Bei der De-Anonymisierung muss der bis dato unbekannte Anwender dazu gebracht werden, einen präparierten Dienst zu verwenden, der dann ein für ihn spezifisch konstruiertes Ultraschallsignal aussendet. Wird dieses Signal von einem Gerät



DATENÜBERTRAGUNG PER ULTRASCHALL

Verwaltungsrelevanz:

■ ■ □ □

Umsetzungs- geschwindigkeit:

■ ■ ■ ■

Marktreife/Produkt- verfügbarkeit:

■ ■ ■ ■

empfangen, dass sich einer bekannten Person zuordnen lässt, kann der Anwender enttarnt werden. Wenn die anonyme Anwendung z. B. in einem Internetcafé über einen PC erfolgt, kann das neben dem Gerät liegende, „ungeschützte“ Smartphone mit einer entsprechenden App zum Verräter werden.

Während Anwendungen für das Location-Tracking oft in den Apps der entsprechenden Ladenketten implementiert sind, können weitere Anwendungen der Ultraschallverfolgung in beliebigen Anwendungen stecken, denn die Technik lässt sich dank entsprechender Softwaremodule bei der Entwicklung mobiler Apps leicht einbinden. Dies geschieht i. d. R. für den Anwender nicht erkennbar und die Apps bieten dann auch keine Option, den „Dienst“ abzuschalten. Da bleibt für den Anwender nur die Möglichkeit, den Zugriff von Apps auf das Mikrofon sehr restriktiv zu handhaben. In Zeiten der vermehrten Verbreitung von Sprachassistenten scheint dies jedoch ein schwieriges Unterfangen.

Bewertung

Selbst wenn Anwender von mobilen Geräten der Verwendung von Apps und der Freigabe des Mi-

krofons zustimmen, dürfte die Ermittlung und „Analyse“ ihres Verhaltens durch Ultraschall-Tracking in den meisten Fällen eher fragwürdig sein, denn diese „Funktionen“ sind wohl in den seltensten Fällen Teil des Anwendungszwecks. Im Umgang mit mobilen Geräten ist also auch in einem weiteren Punkt Vorsicht geboten. Dies gilt umso mehr, wenn dienstlich genutzte Anwendungen betroffen sind.

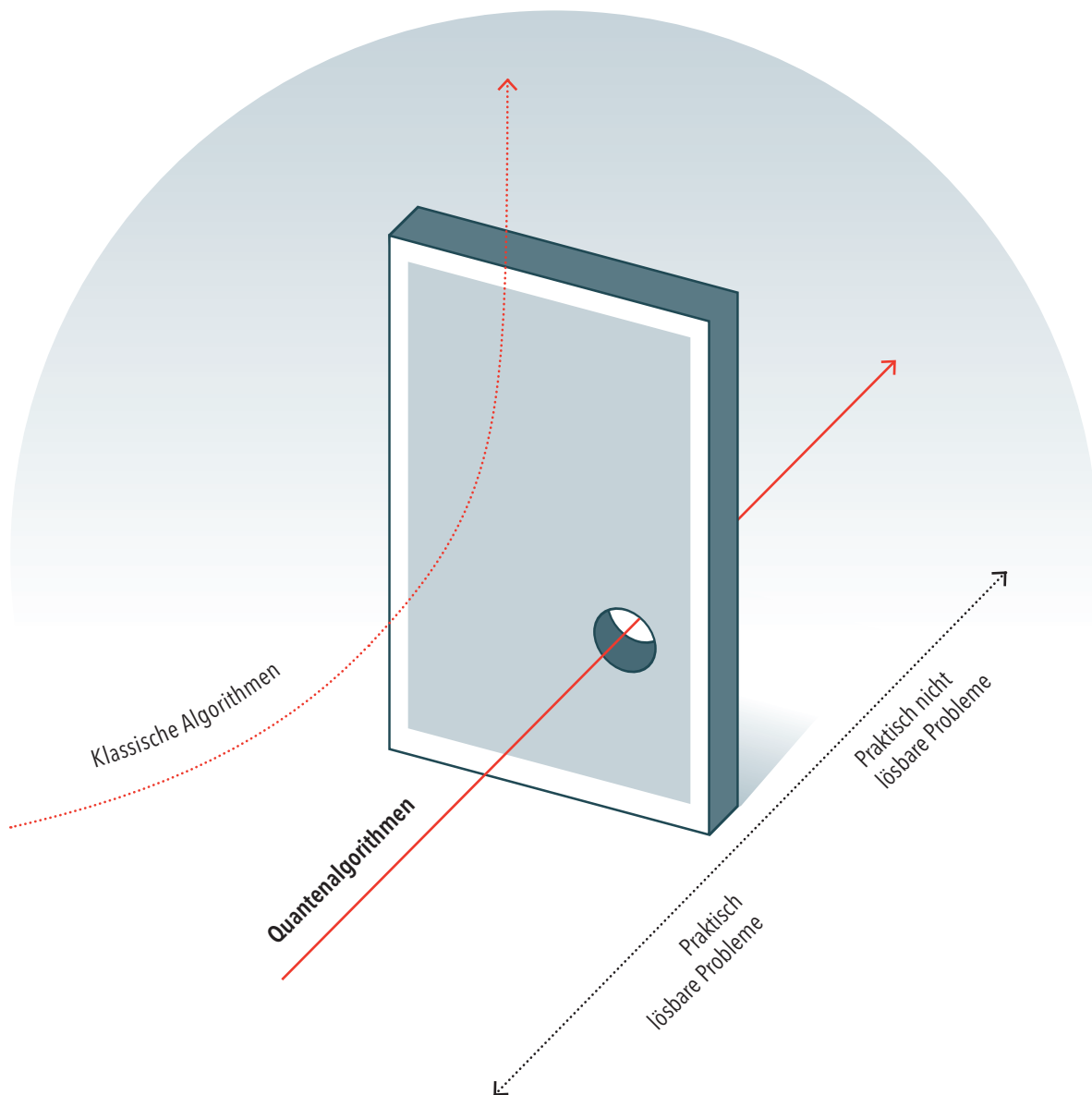
Die reine Übermittlung von Informationen per Ultraschall könnte aber auch anderweitig genutzt werden, z. B. um Daten in Umgebungen zu übertragen, in denen Datenfunk (z. B. per Bluetooth, WLAN oder Mobilnetz) nicht möglich oder unerwünscht ist. Auch der Einsatz von Ultraschall-Beacons wäre denkbar, um Funktionen von Anwendungen nur in einer definierten unmittelbaren Umgebung zu ermöglichen. Ob dies wegen der relativ kurzen Reichweite sinnvoll und hinreichend unempfindlich gegen akustische Störungen ist, bliebe zu prüfen.

Der Ultraschallbereich stellt einen akustischen Übertragungskanal dar. Es wäre wünschenswert, diesen ein- und ausschalten zu können und und zwar unabhängig von der Verwendung von Lautsprecher und Mikrofon im hörbaren Bereich.

EIN QUÄNTCHEN SICHERHEIT?

In der ersten Hälfte des letzten Jahrhunderts experimentierten Wissenschaftler an verschiedenen Orten mit allerlei merkwürdig anmutenden Bauteilen, um automatische Rechenmaschinen zu entwickeln. Die zum Teil monströsen Aufbauten aus mechanischen Schaltteilen (Zuse Z1), elektromagnetischen Relais (Zuse Z3), Elektronenröhren (ENIAC, Colossus, ABC=Atanasoff-Berry-Computer), Lochstreifen zur Programmierung (Zuse Z3, Mark I) oder Laufzeitspeichern (UNIVAC I) mochten den Zeitgenossen genauso merkwürdig vorkommen, wie sie uns erscheinen, wenn wir von Computern sprechen.

Anfang dieses Jahrhunderts experimentieren wieder Wissenschaftler an verschiedenen Orten mit allerlei merkwürdig anmutenden Bauteilen, um automatische Rechenmaschinen zu entwickeln. Auch diese Aufbauten wirken mit ihrer Größe und ihren Bauteilen für uns befremdlich, die wir doch fast alle bereits leistungsfähige Computer in der Hosentasche herumtragen. Heute sind es allerdings keine Relais, Röhren oder Siliziumchips, die da verbaut werden, sondern z. B. Ionenfallen oder schleifenförmige Supraleiter, die in speziellen Gehäusen nahe an den absoluten Nullpunkt (0 K = -273,15 °C) heruntergekühlt werden. Hier werden Quantencomputer gebaut, die anders als „normale“



Computer nicht nur mit einer binären Logik – also den Zuständen „0“ und „1“ – arbeiten, sondern auch einen überlagerten Zustand kennen – „gleichzeitig 0 und 1“. Solche Computer wären in der Lage, manche informationstechnischen oder mathematischen Aufgaben sehr effizient zu lösen, die mit heutigen Computern nur sehr aufwändig zu erledigen sind. Während es bei der Suche in Datenbanken ärgerlich ist, wenn diese lange dauert, machen es sich andere Aufgabenfelder zunutze, dass Lösungen nicht in vernünftiger Zeit gefunden werden können, selbst wenn man „alle verfügbaren“ Computer zu Hilfe nähme. So beruhen viele Verschlüsselungstechniken auf solchen mathematischen Kopfnüssen:

Das kryptografische RSA-Verfahren etwa wird – oft in Kombination mit weiteren Verfahren – an zahlreichen Stellen angewendet, z. B. X.509-Zertifikate, IPsec, TLS oder SSH, S/MIME oder in Verbindung mit dem RFID-Chip des deutschen Reisepasses. Die mathematische Herausforderung bei der Entschlüsselung entsprechender Nachrichten besteht darin, sehr große Zahlen in ihre Faktoren zu zerlegen. Ohne den geheimen Teil des entsprechenden Schlüssels ist dieses Problem mit heutigen Computern nicht in realistischer Zeit zu lösen.

Quantencomputer jedoch könnten diese Aufgabe effizient erledigen. Der dafür entwickelte

Quantenalgorithmen können Probleme lösen, bei denen klassische Algorithmen an der Barriere der praktischen Lösbarkeit scheitern, weil der Aufwand zu hoch wird.



QUÄNTCHEN SICHERHEIT

Verwaltungsrelevanz:

■ ■ ■ □

Umsetzungs- geschwindigkeit:

□ □ □ □

Marktreife/Produkt- verfügbarkeit:

■ □ □ □

Shore-Algorithmus verwendet allerdings Werkzeuge der Quanteninformatik (Quanten-Fouriertransformation) und lässt sich daher auf heutigen Computern nicht implementieren und auch nicht effizient simulieren. Dass die ersten experimentellen Quantencomputer noch keine Gefahr für das RSA-Verfahren darstellen, liegt daran, dass für ernsthafte Anwendungen des Shore-Algorithmus relativ viele „Quantenzellen“ – sog. Qubits in Analogie zu den Speicher-Bits – benötigt werden. Im Jahr 2001 gelang es erstmals, mit einem Quantenrechner mit sieben Qubits die Zahl 15 in ihre Faktoren 3 und 5 zu zerlegen. Seither wurden weitere Quantenrechner mit einigen Qubits realisiert. So können Wissenschaftler heute bereits online mit einem Quantencomputer mit fünf bzw. 20 Qubits experimentieren; Rechner mit 50 Qubits und mehr sind in Arbeit.

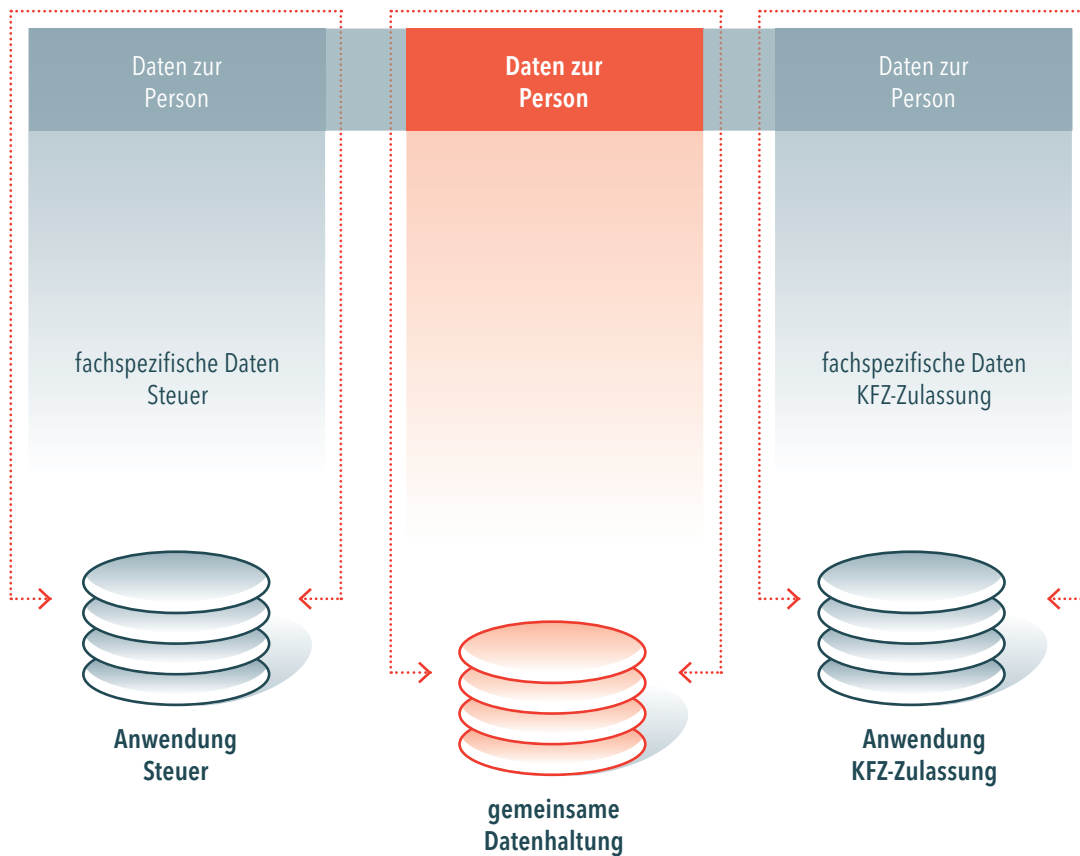
Dass die Rechner nur so wenige Qubits nutzen können, liegt nicht daran, dass das einzelne Qubit so kompliziert ist. Um Quantenalgorithmen rechnen zu können, müssen sie aber miteinander verschränkt sein. D. h., dass mehrere Qubits gemeinsam einen Zustand definieren, der aber nicht durch die einzelnen Qubits „zusammengeschaltet“ werden kann. In diesem kohärenten Zustand arbeitet das System quantentechnisch. Allerdings sind die Zeiten, in denen dieser Zustand anhält, noch sehr kurz. Die Verlängerung von 50µs auf 90µs gilt schon als ein Erfolg.

Aufgrund der schwierigen Bedingungen, unter denen quantentechnisch gerechnet werden kann, werden Zweifel geäußert, dass kommerzielle Rechnersysteme mit 1.000 oder gar 2.000 Qubits tatsächlich durchgängig als Quantencomputer arbeiten. Trotzdem liefern diese Systeme offenbar sehr hohe Rechenleistungen – wenn auch unter ungewöhnlichen Bedingungen: Der supraleitende Prozessor wird auf 0,015 K gekühlt und in einem extremen Hochvakuum betrieben. Durch eine starke Abschirmung werden störende Einflüsse durch das Magnetfeld um Faktor 50.000 reduziert. Und dabei soll sich die Leistungsaufnahme mit 25 kW selbst unter Last im Vergleich mit anderen Supercomputern sparsam ausnehmen, die gerne in den Megawatt-Bereich vorstoßen.

Auch wenn Quantencomputer noch mehr wie Science-Fiction und weniger wie praktikable Technik erscheinen, ist es sinnvoll, sich schon jetzt Gedanken über die Konsequenzen zu machen, wenn sie einmal alltagstauglich werden. Im Fokus der Post-Quanten-Kryptografie steht die Frage, wie kryptografische Verfahren aussehen können, wenn tatsächlich leistungsfähige Quantencomputer zur Verfügung stehen. Weitere Fragen in diesem Zusammenhang sind z. B.: Welche Auswirkungen hat es, wenn bestehende Sicherungssysteme nicht mehr wirksam sind (z. B. bei der Langzeitarchivierung)? Oder wie können Systeme so gebaut werden, dass ein massenhafter Austausch von Sicherheitskomponenten unter Wahrung der Integrität möglich wird? Die Ereignisse rund um Spectre und Meltdown führen vor Augen, welche Auswirkungen das flächendeckende Auftreten von Sicherheitslücken haben kann. Noch ist der Quantencomputer ein komplexes Ungetüm, das sich zumeist in Forschungslaboren aufhält. Aber vielleicht ist ja der Weg zum Quantenrechner in der Hosentasche doch kürzer, als wir denken.

Bewertung

Welche Sicherheitsmechanismen für die IT in den Verwaltungsverfahren eingesetzt werden, sollte eigentlich egal sein, solange sie tatsächlich sicher sind. Solange! Auch wenn Quantencomputer noch in weiter Ferne sind, muss überlegt werden, wie Post-Quanten-Kryptografie in den Verwaltungen aussehen bzw. genutzt werden kann. Auch hier wird deutlich, dass nicht in einzelnen, isolierten Techniken gedacht werden darf. Vielmehr muss die Frage von sicheren Verwaltungsverfahren und -daten sowohl systemtechnisch als auch zeitlich umfassend bedacht werden. Es geht also nicht nur darum quantenkryptografische Verfahren zu entwickeln, sondern auch zu überlegen, wie diese sicher für bestehende und künftige Datenbestände und Verfahren nutzbar gemacht werden können.



Daten, die in verschiedenen Verfahren benötigt werden, werden nur einmal gespeichert.

ONCE-ONLY: „EINMAL UND NIE WIEDER!“

Ach wäre das schön, wenn man seine Daten nicht immer und immer wieder in irgendwelche Formulare oder Bildschirmmasken eintragen müsste. Obwohl dieses Phänomen auch im privatwirtschaftlichen Bereich anzutreffen ist, wird das ständige Wiederholen der gleichen Angaben vornehmlich mit der öffentlichen Verwaltung in Verbindung gebracht. Dabei wird diese von ihren Kunden als eine Einheit angesehen und es ist schwer vermittelbar, dass eine Angabe gegenüber dem Amt A nicht auch beim Amt B vorliegt. Ggf. muss sich der Kunde sogar von Amt A seine Angaben bescheinigen lassen, um sie bei Amt B verwenden zu können. Dass dahinter verschiedene Aufgaben und Zuständigkeiten stecken und die Daten jeweils für einen bestimmten Zweck erhoben werden, ist da nur wenig tröstlich.

Um in dieser Sache Erleichterungen für Bürgerinnen und Bürger und vor allem für Unternehmen, die ebenfalls Daten mehrfach übermitteln

müssen, zu schaffen, hat die EU-Kommission 2016 das Thema „einmalige Erfassung“ (engl. „once-only“) unter dem Stichwort Elektronische Behördendienste in ihre Strategie für einen digitalen Binnenmarkt für Europa aufgenommen. Im EU-eGovernment-Aktionsplan 2016-2020 zur Beschleunigung der Digitalisierung der öffentlichen Verwaltung findet es sich als zweiter von sieben Grundsätzen wieder (s. Kasten).

// „Grundsatz der einmaligen Erfassung: Öffentliche Verwaltungen sollten sicherstellen, dass die Menschen und Unternehmen ihnen dieselben Informationen nur einmal übermitteln. Soweit zulässig, sollten sie diese Daten – unter vollständiger Beachtung der Datenschutzvorschriften – intern mehrmals verwenden, um eine unnötige zusätzliche Belastung der Bürgerinnen und Bürger und der Unternehmen zu vermeiden.“
[EU-eGovernment-Aktionsplan 2016-2020]



// Die Verwendbarkeit vorhandener elektronischer Daten steigert die Attraktivität papierloser Vorgänge.

Die Überlegungen enden aber nicht mit der Vermeidung von mehrfach erfassten Daten. Das Thema ist eng verknüpft mit den Bemühungen um ein „One-Stop-Government“, das alle nötigen Verwaltungsleistungen in einem Portal bündelt. Von einer solchen zentralen Anlaufstelle könnten dann auch selbstständig Vorgänge initiiert werden. In diesem Zusammenhang wird gerne auf das österreichische Verfahren der Familienbeihilfe verwiesen. Hier wird nach Meldung einer Geburt automatisch geprüft, ob ein Rechtsanspruch auf die Gewährung der Familienbeihilfe besteht, und ggf. wird die Zahlung initiiert. Dabei müssen das Personenstandswesen und die Finanzverwaltung Daten austauschen und einen gemeinsamen Prozess implementiert haben. Dieses Beispiel macht deutlich, dass für eine echte Entlastung von Personen und Organisationen durch Automation nicht nur die immer wiederkehrenden Daten vorrätig sein müssen, sondern auch die fallbezogenen Informationen (hier: zur Begründung des Anspruchs). Selbst wenn das nicht in allen Fällen realisierbar ist, wäre oft schon allein die Möglichkeit, teilweise vorausgefüllte Formulare verwenden zu können, eine Erleichterung.

Ein weiteres Thema, das die Bereitstellung und Mehrfachnutzung von Daten betrifft, ist die Frage nach der Identität der handelnden Person bzw. Organisation. Deren Feststellung wird noch (zu) oft mit dem persönlichen Einreichen von Formularen verbunden. So ist es nicht verwunderlich, dass bei den Initiativen der EU zu once-only die eIDAS-Verordnung im Hinblick auf elektronische Identitäten eine wichtige Rolle spielt, da so der Gang zum Amt entfallen könnte.

Wenn verschiedene Verwaltungen derart über Zuständigkeitsgrenzen hinweg zusammenarbeiten sollen, erfordert das eine Anpassung von Systemen und Abläufen – und ggf. von Regelungen. Dies gilt innerhalb einer Kommune und erst Recht auf Ebene der EU. Dafür bietet das once-

only-Prinzip aber auch für die Verwaltungen einige Chancen, die solche Anstrengungen rechtfertigen. Dazu gehören:

- Prozessoptimierung – neben der rein zeitlichen auch in der strukturellen Dimension,
- bessere Datenqualität durch verlässliche, „amtliche“ Quellen und damit
- höhere Effizienz der Vorgangsbearbeitung, indem Rückfragen und Klärungsbedarf auf fachlich-inhaltliche Aspekte reduziert werden.

Zudem steigert die Verwendbarkeit von vorhandenen elektronischen Daten die Attraktivität papierloser Vorgänge, was ebenfalls einer schnelleren Bearbeitung zugute kommt.

Diesen Potenzialen stehen aber auch einige Herausforderungen gegenüber. Technisch wäre der Aufbau von zentralen Registern keine große Herausforderung. Da bestehende elektronische Verfahren für ihre Daten aber bereits eine Lösung haben, müssen für die gemeinsame Nutzung die Datenbestände entweder zusammengeführt oder Daten fallweise ausgetauscht werden. Beides setzt voraus, dass sich die Beteiligten auf die Bedeutung der Daten und deren Format einigen. Die Entwicklungen um die Standardisierung von XML in der öffentlichen Verwaltung (XÖV) sind Beispiele dafür auf nationaler Ebene. Sowohl die technische als auch die fachliche Abstimmung müssen organisatorisch gesteuert werden, denn zunächst bringen sie nur Aufwände mit sich, die oft zusätzlich zum Tagesgeschäft zu leisten sind – und das in verschiedenen Zuständigkeitsbereichen. Ggf. kommt eine weitere Abstimmungsebene hinzu, wenn die rechtlichen Rahmenbedingungen für die zu beteiligenden Verfahren unterschiedlich sind. Dies spielt insbesondere auf der internationalen Ebene eine große Rolle.

Eine besondere Herausforderung im rechtlichen Zusammenhang stellt der Datenschutz dar. Nun soll das once-only-Prinzip nicht ausschließlich und nicht einmal vorwiegend für personenbezogene Daten gelten. Trotzdem steht dieses Thema oft im Fokus der Diskussion und Bericht-

erstattung über once-only. Das sog. Volkszählungsurteil des Bundesverfassungsgerichts von 1983 zur informationellen Selbstbestimmung und das Urteil des Europäischen Gerichtshofs im Zusammenhang mit der Vorratsdatenspeicherung von 2016 setzen einen engen Rahmen für die Verwendung personenbezogener Daten. Zentrale Register, die Daten über den akuten Zweck hinaus für mögliche weitere Anwendungsfälle enthalten, sind daher ohne Einwilligung der betroffenen Person nicht zulässig. Auch eine Verordnung der EU-Kommission, durch die der Austausch von Behördendaten über ein Single Digital Gateway ermöglicht wird, ist aus Sicht des Europäischen Datenschutzbeauftragten noch nachzubessern.

Neben verschiedenen Aktivitäten auf jeweils nationaler Ebene – hier werden neben Österreich insbesondere Estland und die Niederlande als Beispiele genannt – hat die EU Fördermittel aus dem Programm Horizon 2020 zur Verfügung gestellt, um Anwendungen für das once-only-Prinzip im internationalen Rahmen zu unterstützen. Hier steht das Projekt TOOP (The Once-Only-Principle Project) im Fokus. Darin wollen 51 Partner aus 21 Ländern eine generische föderierte Architektur entwickeln, über die dann mehr als 60 Informationssysteme Unternehmensdaten miteinander austauschen können. Die drei pilotierten Anwendungsfälle dafür sind:

- Grenzübergreifende E-Services für die Wirtschaft (z. B. elektronische Vergabe öffentlicher Aufträge; hier wirkt auch die Metropolregion Rhein-Neckar mit)
- Aktualisierung von Daten bei länderübergreifend vernetzten Firmen (z. B. im Handels- oder Genossenschaftsregister)
- Online-Zertifikate für Schiffe und Mannschaften (z. B. für die Hafenbehörden)

Ebenfalls Horizon-gefördert untersucht das Projekt Stakeholder Community of the Once-Only Principle For Citizens (SCOOP4C) Möglichkeiten, wie das once-only-Prinzip bei Bürgerdiensten angewendet werden kann. Zu den Aufgaben gehören die Entwicklung einer Roadmap zur Umsetzung von once-only für Bürger und die

Verbreitung des Themas durch den Austausch von themenrelevanten Personen.

Die beiden Initiativen sind mit 30 resp. 24 Monaten relativ lang angelegt und zeigen anhand ihrer Ziele und Aufgaben, dass es dabei um Grundlagenarbeit geht. Die umfassende konkrete Anwendung einer generischen Architektur und die Umsetzung von Maßnahmen entlang einer Roadmap werden danach viele bestehende Verfahren und neue Projekte beeinflussen.

Bewertung

Die Vermeidung von Mehrfacharbeit und von Medienbrüchen stehen schon seit vielen Jahren in den Auftragsbüchern der E-Government-Initiativen. Elektronisch vorhandene und wiederverwendbare Daten könnten einen wesentlichen Beitrag dazu leisten, in diesen Bemühungen weiter voranzukommen. Wenn die Umsetzung dann auch noch flächendeckend gelingt – sowohl national (und damit auch regional) als auch EU-weit – könnte dies insgesamt einen enormen Fortschritt für das E-Government ergeben, denn wo Daten kompatibel sind, können auch elektronische Verfahren leichter miteinander verknüpft werden. Die fach- oder gar vorgangsspezifischen Daten und Aufgaben werden dann immer noch fallweise und lokal anfallen. Wenn ihre Verarbeitung aber künftig entsprechend der Grundsätze aus dem EU-Aktionsplan „standardmäßig digital“ und „einmalige Erfassung“ in der (Weiter-)Entwicklung von Fachverfahren von Anfang an gestaltet wird, könnten innovative Verwaltungsdienstleistungen leichter entstehen. Durch koordinierte E-Government-Angebote, wie sie mit der Umsetzung des Onlinezugangsgesetzes (OZG) geschaffen werden, erhält die „einfache“ Datenhaltung über den reinen Komfortgedanken hinaus eine strategische Bedeutung. Daher bleibt zu wünschen, dass sich die entsprechende Praxis mit nicht personenbezogenen Daten erfolgreich und zeitnah etablieren kann, und dass für die Frage der Zweckbindung von personenbezogenen Daten im Hinblick auf once-only innerhalb der öffentlichen Verwaltung eine gute Lösung gefunden wird.

ONCE ONLY

Verwaltungsrelevanz:

■■■■

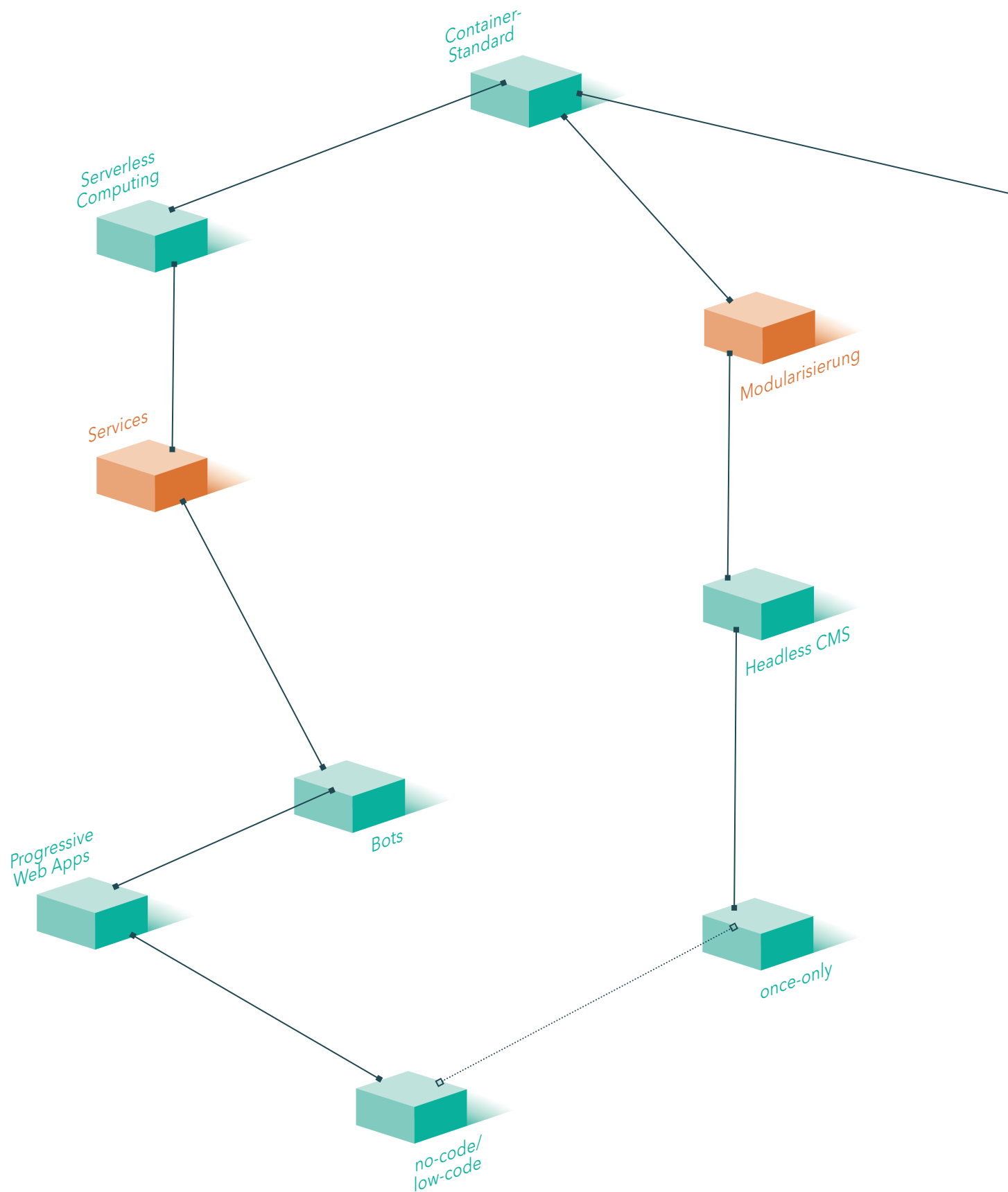
Umsetzungs-
geschwindigkeit:

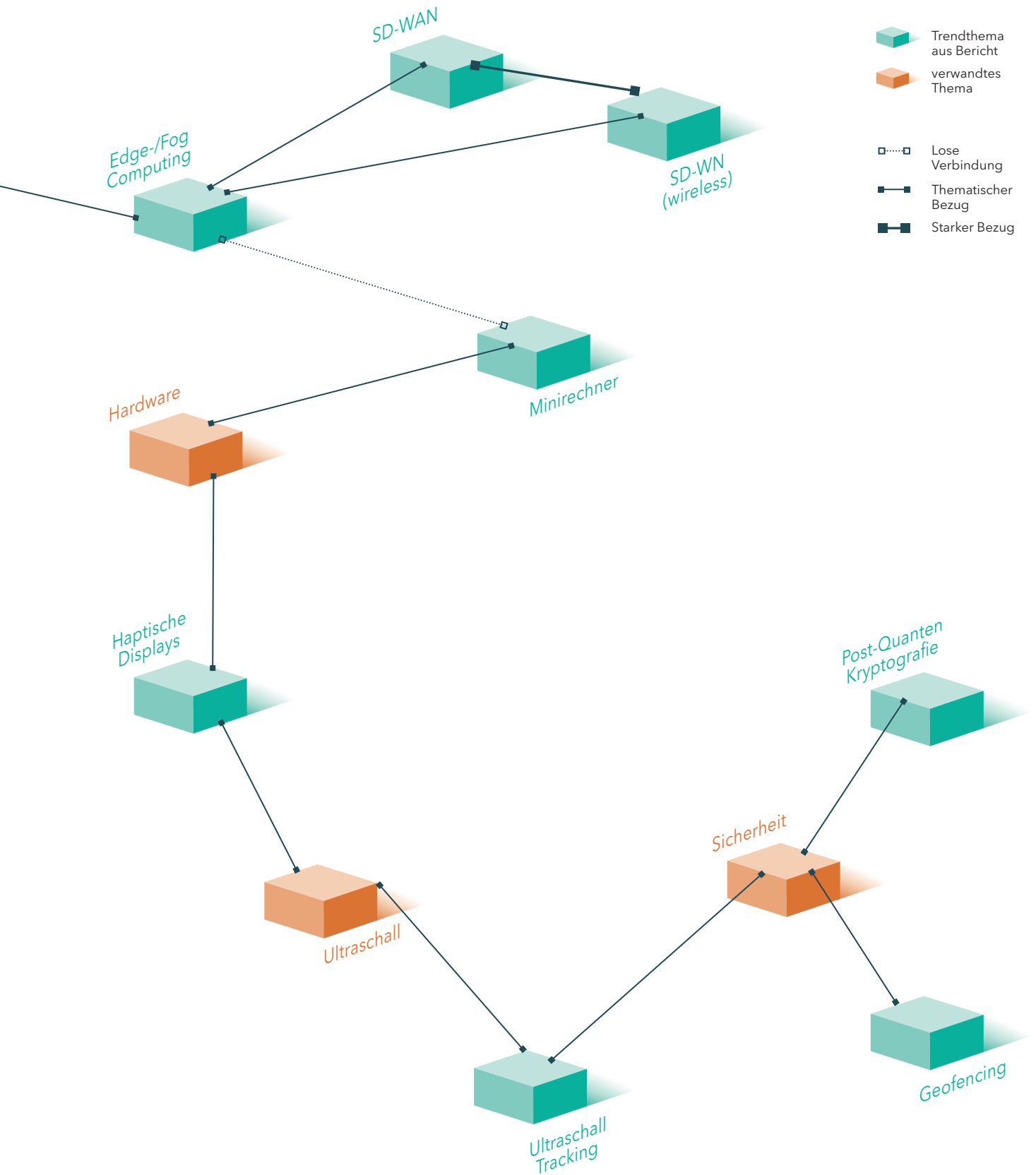
■□□□

Marktreife/Produkt-
verfügbarkeit:

■□□□

TRENDMAP





IMPRESSUM

Herausgeber

Hessische Zentrale für Datenverarbeitung
Mainzer Straße 29
65185 Wiesbaden
Telefon: 0611 340-0
E-Mail: info@hzd.hessen.de
www.hzd.hessen.de

Verantwortlich

Dr. Markus Beckmann
Telefon: 0611 340-1280
E-Mail: markus.beckmann@hzd.hessen.de

Gestaltung

Agentur 42, Konzept & Design

Kapitelillustrationen

Titel, S. 8/9 ff: © SFIO CRACHO/© spainter_vfx
- fotolia.com; S. 22/23 ff: © mooshny/©
pinkeyes - fotolia.com; S. 32/33 ff: © ivanmol-
lov/© mooshny- fotolia.com; S. 42/43 ff:
© rcfotostock/© spainter-vfx- fotolia.com

Druck

Druckerei Chmielorz GmbH;
www.druckerei-chmielorz.de

Erscheinungstermin

Juni 2018

Vervielfältigung und Verbreitung, auch
auszugsweise, mit Quellenangabe gestattet.





Mainzer Straße 29 | 65185 Wiesbaden
Telefon: 0611 340-0 | E-Mail: info@hzd.hessen.de
www.hzd.hessen.de

