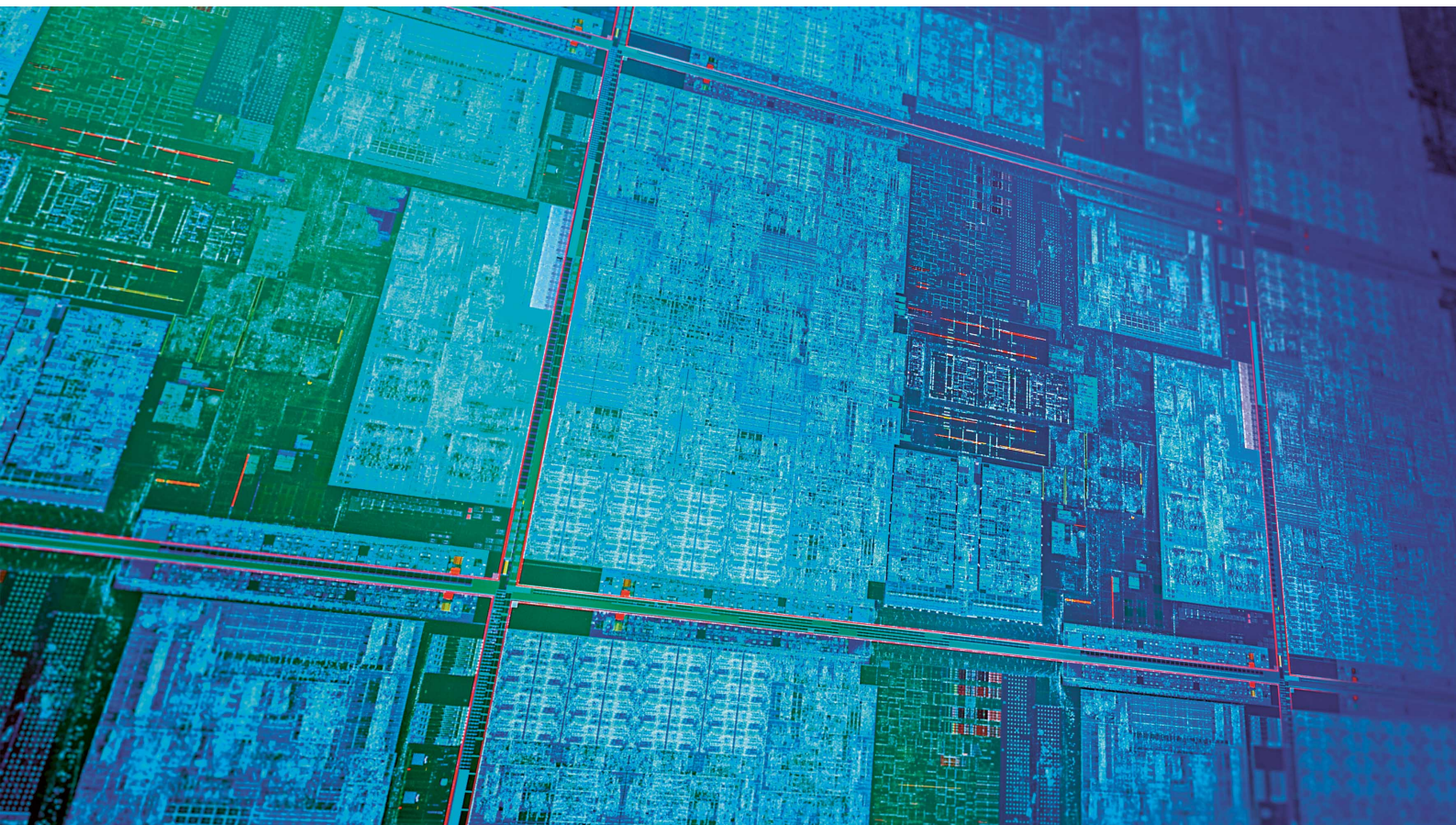




Mehr als Moore

29/03/2021 Seit Jahrzehnten schrumpfen die Transistoren in integrierten Schaltungen, wodurch immer leistungsfähigere Chips entstanden sind. Porsche Engineering berichtet, wie Forscher und Industrie daran arbeiten, um Chips auch in Zukunft immer leistungsfähiger zu machen und wie davon langfristig die Automobilindustrie profitiert.

Die Geburtsstunde der modernen Elektronik schlug am 16. Dezember 1947 in den Bell Labs in Murray Hill (New Jersey). Damals gelang es dem Physiker Walter Brattain erstmals, mit einem improvisierten Halbleiter-Bauelement eine elektrische Spannung zu verstärken. Der Transistor war geboren. Mit ihm stand zum ersten Mal eine Alternative zu den klobigen, unzuverlässigen und energiehungrigen Vakuumröhren zur Verfügung. Brattains Laboraufbau aus einem Germanium-Plättchen, einem Plastik-Dreieck, einer Goldfolie und einer Büroklammer hatte zwar äußerlich noch nichts mit modernen Chips zu tun – aber er lautete dennoch die Ara von Personal Computern, Smartphones und selbstfahrenden Autos ein.

Das neue elektronische Bauelement eignete sich als Verstärker und als Schalter – und es ließ sich mit seinesgleichen sowie anderen Komponenten wie Widerständen und Kondensatoren als integrierte Schaltung (integrated circuit, IC) auf einem einzigen Halbleiterplättchen unterbringen. In den folgenden

Jahrzehnten gelang es den Halbleiterunternehmen, die Bauelemente immer mehr zu verkleinern und eine immer größere Zahl von ihnen auf der gleichen Fläche unterzubringen. Gordon Moore sagte bereits 1965 voraus, dass die Zahl der Transistoren pro Flächeneinheit exponentiell ansteigen werde.

Einfache Skalierung kommt an ihr Ende

Moore's Vorhersage hat sich über Jahrzehnte im Wesentlichen als korrekt erwiesen. Jetzt stößt sie aber endgültig an ihre Grenzen, weil die schrittweise Verkleinerung der bewährten MOSFETs (Metal-Oxide-Semiconductor Field-Effect Transistor) als Schalter auf den Chips nicht mehr funktioniert: „Vor etwa 15 Jahren hat man erkannt, dass die einfache Skalierung an ihr Ende gekommen ist“, berichtet Dr. Heike Riel, IBM Fellow am IBM- Forschungszentrum im schweizerischen Rueschlikon. „Zuerst haben die Hersteller darum bei gleichbleibender MOSFET-Geometrie zum Beispiel Siliziumdioxid als Isoliermaterial in den Transistoren durch sogenannte High-K-Materialien ersetzt. So ließen sich Chips mit 45 Nm großen Strukturen herstellen.“

Aber auch dieser Trick konnte das Mooresche Gesetz nur einige Jahren am Leben halten. Darum setzten die Chiphersteller ab der ersten Hälfte der 2010er-Jahre auf eine neue Transistor-Architektur für noch kleinere Komponenten: den FinFET. Bei ihm ist der leitende Kanal zwischen Source- und Drain-Anschluss wie eine Flosse (englisch „fin“) geformt und wird von der Steuerelektrode (Gate) an mehreren Seiten umschlossen. „Dadurch lässt sich der Stromfluss im Transistor wesentlich besser kontrollieren“, so Riel. „FinFETs kamen ab 22 Nm Strukturgröße zum Einsatz und sind heute Standard in integrierten Schaltungen.“

Der Gate-all-around-FET

Aber auch ihr Nachfolger steht schon bereit. Ab Strukturgrößen von fünf Nanometern soll der GAA-FET (Gate-all-around-FET) die Arbeit in den Chips übernehmen. „Bei ihm besteht der leitende Kanal zwischen Source und Drain aus mehreren parallelen Silizium-Nanodrähten, die jeweils komplett von der Gate-Elektrode umschlossen sind“, erklärt Riel. „Das ist die optimale Geometrie für die Steuerung des Stromflusses. Außerdem spart man Fläche auf den Chips, weil mehrere dieser Nanodrahtstrukturen, die den Kanal des Transistors bilden, übereinander liegen.“ Von Entwicklungen wie dem GAAFET werden künftig auch Automotive-Anwendungen profitieren – denn sowohl die neuen, leistungsstarken High Performance Computing Platforms (HCPs) als Nachfolger der vielen dezentralen Steuergeräte als auch die Spezialprozessoren für das autonome Fahren sind auf Chips mit hoher Rechenleistung angewiesen. Das Mooresche Gesetz wird aber auch der GAAFET auf Dauer nicht retten können: Jenseits von drei Nm Strukturgröße wird es eng. In drei bis vier Jahren konnte diese Grenze erreicht sein.

„Eine geniale Phase des Verbesserns kommt damit an ihr Ende“, sagt Prof. Dr. Thomas Schimmel, Direktor am Institut für Nanotechnologie des Karlsruher Instituts für Technologie (KIT). „Bisher hat man immer eine technische Lösung gefunden, die eine weitere Verkleinerung des klassischen Transistors ermöglichte – aber jetzt hat man in den Chips atomare Dimensionen erreicht. Durch den

quantenmechanischen Tunneleffekt können Elektronen Isolierungen durchqueren, was die Bauelemente unbrauchbar machen würde. Denn im Gegensatz zu den Vorstellungen der klassischen Physik können die Elektronen selbst dann Barrieren überwinden, wenn sie eigentlich nicht genügend Energie dafür besitzen.“ Aber auch das gezielte Einbringen von Fremdatomen in das hochreine Silizium während des Produktionsprozesses – „Dotieren“ genannt – funktioniert bei immer kleineren Strukturen nicht mehr zuverlässig.

Kein heißer Nachfolge-Kandidat in Sicht

Gesucht wird darum ein Nachfolger des Transistors, mit dem sich die Performance elektronischer Schaltungen in Zukunft weiter steigern lässt. IBM-Forscherin Riel zählt eine ganze Reihe von MOSFET-Alternativen auf, darunter den Carbon Nanotube Field-Effect Transistor (CNFET) und den Tunnel-FET (TFET): Im CNFET fließt der Strom durch winzige Röhren aus Kohlenstoff. Forscher vom MIT haben dieses Jahr gezeigt, dass sich die schnellen und energieeffizienten Schalter in konventionellen Chip-Fabriken herstellen lassen. TFETs sind ähnlich aufgebaut wie konventionelle Transistoren, nutzen beim Schalten aber den quantenmechanischen Tunneleffekt zu ihrem Vorteil. Sie sind energiesparend und schnell. Ob CNFET, TFET oder ein anderer Ansatz das Rennen machen wird, ist aber völlig offen. „Es gibt derzeit viel Forschung, aber noch keinen heißen Kandidaten für die Nachfolge des optimierten Silizium-MOSFET“, so Riel.

KIT-Forscher Schimmel setzt dafür langfristig auf den Einzelatom-Transistor: Bei ihm verschiebt eine Steuerelektrode ein Atom, das die winzige Lucke zwischen zwei Anschlüssen schließen und so den Stromfluss ermöglichen kann. „Im Prinzip funktioniert das wie bei einem Relais mit zwei stabilen Zuständen“, so Schimmel, der 2004 mit seinem Team den ersten Einzelatom-Transistor entwickelte. „Dadurch ist der Einzelatom-Transistor nicht nur ein Schalter, sondern auch ein nichtfluchtiger Speicher. Er konnte also zugleich auch herkömmliche RAM-Chips als Arbeitsspeicher in Computern ersetzen. Da er auch ohne Strom seinen Zustand beibehält, musste man Computer in Zukunft nicht mehr neu starten, sondern konnte nach einer Pause sofort weiterarbeiten.“

Weiterer Vorteil: Der Einzelatom-Transistor benötigt deutlich weniger Spannung als der MOSFET und wurde dadurch im Vergleich pro Schaltvorgang nur etwa ein Zehntausendstel der Energie verbrauchen. Das Hitze-Problem heutiger Chips wäre damit gelöst, und Taktfrequenzen von bis zu 100 Gigahertz kamen in Reichweite. Einen ersten IC mit zwei seiner neuartigen Transistoren hat Schimmel schon gebaut, und für eine spätere Serienfertigung konnte man einen Mix aus bewährten Prozessen der Halbleiterindustrie und galvanischen Verfahren nutzen. „Das ist wie beim Verzinken einer Karosserie – nur eben im atomaren Maßstab“, sagt der Karlsruher Wissenschaftler.

Neue Chip- und Computer-Architekturen

Alternativ zur immer weiteren Verkleinerung der Bauelemente bieten sich auch neue Ansätze bei der Chip-Architektur an, zum Beispiel der Weg in die dritte Dimension: Um mehr Leistung in die

Schaltkreise zu packen, ließen sich mehrere Elektronik-Schichten aufeinander stapeln – schon heute praktiziert bei Flash-Speichern. Dabei konnten die Hersteller auf eine Lage mit herkömmlichen Silizium-Transistoren künftig auch eine Ebene aus sogenannten Verbindungshalbleitern wie Indiumgalliumarsenid (InGaAs) aufbringen.

Sie eignen sich für Spezialaufgaben wie besonders schnelle Verstärkung, zur Emission oder Detektion von Licht und auch als mögliche Quantenbauelemente. Viele Experten setzen auf die Integration solcher zusätzlicher Funktionen in die Chips, um das Ende des Mooreschen Gesetzes zu kompensieren. Ihre Devise lautet: Statt „More Moore“ (weitere Miniaturisierung) lieber „More than Moore“ (die Vereinigung von digitalen und nicht-digitalen Funktionen auf demselben Chip).

Das In-Memory-Computing

Eine deutlich höhere Rechenleistung und mehr Energieeffizienz verspricht das In-Memory-Computing, das die räumliche Trennung von Recheneinheit und Speicher in gängigen Computern aufheben soll. Denn damit wurde der zeitraubende und energieintensive Transport der Bytes zwischen Mikroprozessor und RAM entfallen. So lassen sich beispielsweise die Vektor-Matrix-Berechnungen in einem neuronalen Netz mithilfe einer Crossbar-Architektur durchführen – analog statt digital. Bei diesem Ansatz kreuzen sich zwei Bündel von horizontalen und vertikalen Leitungen, die jeweils als Ein- und Ausgänge des neuronalen Netzes fungieren. An ihren Kreuzungspunkten sind die Leitungen über nichtfluchtige Speicherelemente miteinander verbunden, die die Gewichtungsfaktoren (das „Wissen“) des neuronalen Netzes repräsentieren. Die Eingangswerte des neuronalen Netzes werden als analoge Spannungswerte an die horizontalen Leitungen angelegt. Die Ergebnisse der Berechnungen stehen fast augenblicklich an den vertikalen Leitungen bereit, ebenfalls in analoger Form – und ohne jeden Datentransport.

Als nichtfluchtige Speicher an den Kreuzungspunkten der Leitungen kommen unter anderem Memristoren infrage, neuartige elektronische Bauelemente, deren Widerstand durch eine von außen angelegte Spannung dauerhaft verändert werden kann und die sich mit bestehenden Herstellungsprozessen der Halbleiterindustrie kombinieren lassen. Mit Memristoren kann man die Berechnungen in neuronalen Netzen je nach Anwendung zehn- bis hundertmal schneller durchführen, zudem lässt sich die Energieeffizienz um den Faktor zehn bis tausend verbessern. Von solchen Leistungs- und Effizienzsteigerungen konnten auch autonome Fahrzeuge profitieren, in denen neuronale Netze eine Hauptrolle spielen. Das Beispiel zeigt: Auch wenn das Mooresche Gesetz bald an Grenzen stößt – die kontinuierliche Leistungssteigerung der Elektronik ist noch lange nicht am Ende.

MOSFET, FinFET und GAAFET

Der MOSFET (links) ist seit Jahrzehnten das Arbeitspferd als Schalter für die Digitaltechnik und hat durch seine permanente Verkleinerung das Mooresche Gesetz am Leben erhalten. Die Spannung zwischen Gate- und Source-Elektrode bestimmt bei ihm den Strom, der durch den Kanal von Source

nach Drain fließt. Beim FinFET (Mitte) hat der Kanal die Form einer Flosse (englisch „fin“), sodass das Gate ihn an drei Seiten umschließen kann. Das verbessert die Steuerung des Stromflusses im Vergleich zum MOSFET, bei dem das Gate nur von oben auf den Kanal einwirken kann. Beim GAAFET (rechts) umschließt das Gate den Kanal aus Silizium-Nanodrahten vollständig. Das ist die optimale Geometrie, um den Stromfluss zu kontrollieren.

Künftige Alternativen zum herkömmlichen Transistor-Design

Beim TFET (links) sind Source und Drain im Gegensatz zum MOSFET unterschiedlich dotiert. Er nutzt den quantenmechanischen Tunneleffekt: Bei ihm bestimmt die Spannung zwischen Gate und Source, ob Ladungsträger die energetische Barriere zwischen Source und Drain „durchtunneln“ können und ein Stromfluss möglich ist. Beim CNFET (mitte) besteht der Kanal zwischen Source und Drain aus Kohlenstoff-Nanoröhrchen. Auch hier bestimmt die Gate-Source-Spannung den Stromfluss. Beim Einzelatom-Transistor (rechts) verschiebt die Spannung zwischen Source und Gate ein einzelnes Atom, das den Stromkreis zwischen Source und Drain entweder schließt oder öffnet (grüne/rote Position).

Das Mooresche Gesetz

Vor 55 Jahren machte Gordon Moore eine bemerkenswerte Voraussage: Die Zahl der Transistoren pro Chip wird sich in Zukunft jedes Jahr verdoppeln, so der damalige Forschungsdirektor des US-Halbleiterherstellers Fairchild Semiconductor und spätere Intel-Mitgründer 1965 in der Zeitschrift „Electronics“. Bereits 1975 wurde es darum möglich sein, rund 65.000 von ihnen auf einem winzigen Siliziumplättchen unterzubringen. Das „Mooresche Gesetz“ wurde immer wieder leicht angepasst, hat sich im Grundsatz aber als korrekt erwiesen und avancierte zur Richtschnur der Halbleiterhersteller: Bis heute gelingt es ihnen immer wieder, in kurzen Abständen die Zahl der Transistoren pro Flächeneinheit zu verdoppeln. In Zukunft wird das aber nicht mehr möglich sein. Eine höhere Leistungsfähigkeit der Chips muss dann auf andere Weise erreicht werden.

Zusammengefasst

Die weitere Verkleinerung des MOSFET-Transistors und seiner Varianten FinFET und GAAFET dürfte in den kommenden Jahren an Grenzen stoßen. Um Chips auch in Zukunft immer leistungsfähiger zu machen, arbeiten Forschung und Industrie sowohl an neuen Transistor-Designs wie Tunnel-FETs als auch an neuen Architekturen wie In-Memory-Computing.

Info

Text: Christian Buck

Text erstmalig erschienen im Porsche Engineering Magazin, Nr. 1/2021

MEDIA ENQUIRIES



Frederic Damköhler

Senior Manager Corporate Communications Porsche Engineering
+49 (0) 711 / 911 16361
frederic.damkoehler@porsche.de

Linksammlung

Link zu diesem Artikel

<https://newsroom.porsche.com/de/2021/innovation/porsche-engineering-mooresches-gesetz-mikroelektronik-chips-23344.html>

Externe Links

<https://www.porscheengineering.com/peg/de/>