



Deutsches  
Patent- und Markenamt

# Erfinderaktivitäten 2014 / 2015

Schwerpunkt: Sensorik



# Inhalt

|                                                                                                                                 |    |
|---------------------------------------------------------------------------------------------------------------------------------|----|
| <b>Vorwort</b>                                                                                                                  | 3  |
| <br><b>Treppauf und treppab – Sensortechnik in treppengängigen Fahrzeugen</b><br><i>Dr. Jan-Friedrich Süßmuth</i>               | 4  |
| <br><b>Magnetische Positionssensoren in der Anwendung</b><br><i>Dr. Christopher Root</i>                                        | 11 |
| <br><b>Halbleiter- und Festkörperlbauelement-Sensoren</b><br><i>Dr. Ralf Henninger</i>                                          | 19 |
| <br><b>Abgassensoren für die Automobilindustrie</b><br><i>Dr. Bettina Kuhn</i>                                                  | 34 |
| <br><b>Kameras in Kraftfahrzeugen</b><br><i>Dipl.-Ing. Steffen Merunka</i>                                                      | 39 |
| <br><b>Der Weg zum autonomen Fahren – Bildverarbeitung zur Fahrerunterstützung</b><br><i>Dipl.-Phys. Hartmut Wilhelms</i>       | 48 |
| <br><b>Datenbrillen</b><br><i>Dipl.-Ing. Martin Bässler</i>                                                                     | 57 |
| <br><b>Halbleiter-Bildsensorik am Beispiel von Flachdetektoren für die Röntgendiagnostik</b><br><i>Dipl.-Ing. Martin Albert</i> | 68 |
| <br><b>Talbot-Lau-Interferometer zur Phasenkontrast-Bildgebung mit Röntgenstrahlen</b><br><i>Dr. Florian Siebel</i>             | 84 |
| <br><b>Glossar</b>                                                                                                              | 92 |

# Vorwort

Sehr geehrte Leserinnen und Leser,

Sensorik – das ist viel mehr als ein Touchscreen, der konzipiert ist, um Wischbewegungen des Nutzers zu erkennen. Schon an dieser Erfindung wird deutlich, wie Sensoren in kürzester Zeit unsere Welt verändert haben: Fingerspitzengefühl wird nun der Technik anvertraut, komplexe Ablaufmechanismen werden automatisiert.

Das tägliche Leben wird inzwischen von einer Vielzahl weiterer Neuerungen im Bereich der Sensorik beeinflusst. Aber wie kann die Sensortechnik Positionen und Hindernisse erkennen oder ein Fitnessarmband meinen Blutsauerstoffgehalt bestimmen? Welche Entwicklungen gibt es auf dem Gebiet der Röntgendiagnostik? Sind Datenbrillen unsere Wegweiser der Zukunft? Ist es richtig, dass Automobile bereits vor hundert Jahren mit Kameras ausgestattet wurden, und wie ist der Entwicklungsstand heute? Die Antworten auf diese Fragen finden Sie in den nächsten Kapiteln. Nicht zuletzt gibt es ein ganz aktuelles Thema: Abgassensoren für die Automobilindustrie.

Die Artikel wurden von den mit den jeweiligen Fachgebieten befassten Patentprüferinnen und Patentprüfern geschrieben. Sie geben Ihnen einen Einblick in die Fülle von Erfindungen auf dem spannenden Gebiet der Sensorik. Am Schluss dieser Ausgabe bieten wir Ihnen erstmals ein Glossar an, in dem Sie Fachbegriffe der jeweiligen Artikel nachschlagen können.

Viel Spaß beim Lesen wünscht Ihnen

Mark Haslinger  
(Redaktion)

# Treppauf und treppab – Sensortechnik in treppengängigen Fahrzeugen

*Dr. Jan-Friedrich Süßmuth, Patentabteilung 1.21*

Handgeführte Fahrzeuge zum Transport von Lasten aller Art kommen in großem Umfang im gewerblichen wie privaten Bereich zum Einsatz. Um mit diesen Fahrzeugen auch Stufen oder Treppen ohne große Anstrengung bewältigen zu können, wurden bereits zahlreiche Treppensteigvorrichtungen mit und ohne eigenen Antrieb erdacht. Dieser Artikel zeigt an ausgewählten Beispielen, welchen Beitrag die Sensortechnik leisten kann, um dem sicheren und mühelosen Lastentransport über Treppen näherzukommen.

---

## 1 Einleitung

Spätestens beim nächsten Umzug stellt sich wieder die alte Frage: Wie kommt die Waschmaschine in den vierten Stock, wenn es keinen Aufzug gibt? Diese oder auch andere alltägliche Aufgabenstellungen, bei denen Lasten über Treppen transportiert werden müssen, werden erfahrungsgemäß immer noch in den meisten Fällen durch eine handelsübliche Sackkarre und menschliche Muskelkraft bewältigt. Die Gefahr, sich dabei Schäden am Bewegungsapparat oder am Transportgut einzuhandeln, ist jedermann bekannt. Daher haben sich bereits zahlreiche Anmelder damit beschäftigt, handgeführte Transportkarren mit Hilfseinrichtungen zu versehen, die das Befahren von Treppen auch mit schweren Lasten leicht und sicher ermöglichen sollen. Der Begriff „Transportkarre“ schließt im Folgenden verschiedene Arten handgeführter Fahrzeuge für den Lastentransport ein, beispielsweise auch Sackkarren, Zustellerwagen und Einkaufstrolleys. Ihnen ist gemeinsam, dass sie ein zu Fuß gehender Bediener führt – wodurch sie sich von Fahrzeugen unterscheiden, die den Bediener selbst transportieren.

Die Anstrengungen der Erfinder und Erfinderinnen richten sich auf diesem Gebiet im Wesentlichen auf die Bereitstellung folgender Funktionalitäten einer Transportkarre:

- normale Straßenfahrt in der Ebene

- Heben/Senken der Transportkarre mit der Last bei Fahrt über Stufen
- sicheres Halten der Transportkarre in jeder Phase der Bewegung
- Gewährleisten eines sicheren Übergangs zwischen Treppen- und Straßenfahrt

Diesen Anforderungen können handgeführte treppengängige Fahrzeuge speziell durch den Einsatz von Sensoren gerecht werden.

Die aufgeführten Beispiele aus der IPC-Klasse B62B 5/02 sollen aufzeigen, welche Bedeutung der zunehmend kleiner und preiswerter werdenden Sensortechnik für die Entlastung und Sicherheit des Menschen in diesem Bereich zukommt.

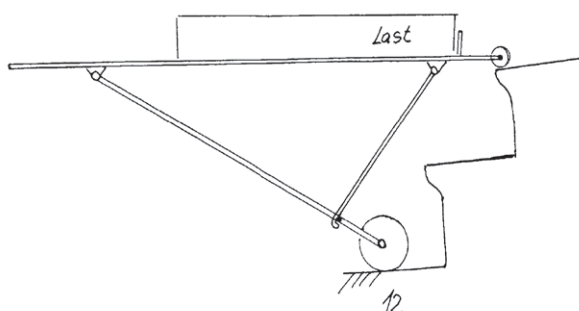
Zur Begriffsbestimmung sei angemerkt, dass im Folgenden unter Laufrädern diejenigen Räder einer Transportkarre zu verstehen sind, die bei Straßenfahrt in der Ebene allein lasttragend sind.

## 2 Mechanische Sensoren und Endlagenschalter

Eine Transportkarre kann nur wirtschaftlich erfolgreich sein, wenn sie zuverlässig und dabei dennoch kostengünstig ist. Bei jeder Anstrengung, den Nutzwert oder die Sicherheit von Transportkarren weiterzuentwickeln, dürfen diese Aspekte daher nicht unberücksichtigt bleiben.

Hieraus ergibt sich folgerichtig der Wunsch nach möglichst einfacher Technik, was selbstverständlich auch für die Sensorik gilt, wie am Beispiel einer altbewährten Form von Treppenkarre deutlich wird:

Eine klassische, rein manuell bewegte Treppenkarre besteht aus einem Gestell mit Laufrädern für die Fortbewegung in der Ebene und zusätzlichen Hilfsrollen für die Bewegung auf Treppen (siehe Figur 1). Die im Hinblick auf die zu erwartenden Stufen ausgelegte Geometrie des Gestells ermöglicht es, beim aufsteigenden Befahren von Treppen zunächst die Hilfsrollen auf einer der nächsthöheren Trittstufen abzusetzen, die Laufräder durch manuelles Anheben der Transportkarre zu entlasten und Letztere auf den Hilfsrollen solange vorwärts zu schieben, bis die Laufräder wieder auf eine Stufe gesetzt werden können, woraufhin sich das Spiel wiederholt. Diese Technik verlangt vom Bediener einiges an Kraft und insbesondere beim Fahren treppab Gefühl dafür, wann die Hilfsrollen auf die abwärts folgende Treppenstufe aufgesetzt werden können.



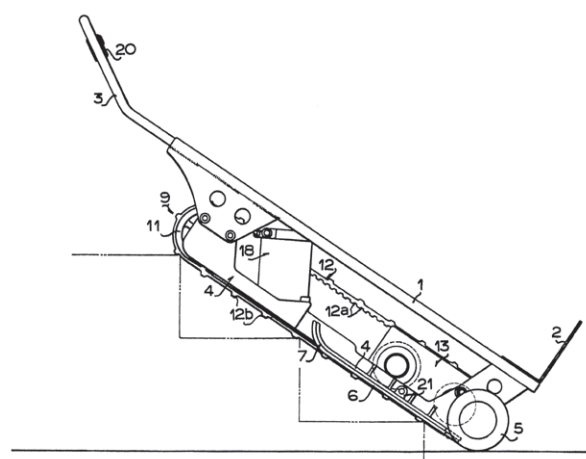
Figur 1: Manuelle Treppenkarre (aus DE 20 2012 010 141 U1)

Um dies insbesondere für den Fall zu verbessern, in dem der Benutzer in Folge sperriger Lasten keinen Blick auf die Hilfsrollen hat, schlägt der Autor der DE 20 2012 010 141 U1 einen Treppensensor einfachster, aber robuster Art vor: Eine bewegliche Zunge (in Figur 1 nicht erkennbar) ist so gestaltet, dass sie auf die nächstuntere Stufe aufschlägt, wenn die Hilfsrollen eine obere Stufe verlassen haben, wobei in dieser Phase die Laufräder das Gewicht tragen. Dem Bediener wird so akustisch angezeigt, dass die Hilfsrollen abgesenkt werden können.

Der erste Schritt, um dem Nutzer einer Transportkarre die Hebearbeit beim Befahren von Stufen abzunehmen, ist das Vorsehen elektrisch angetriebener Treppensteig-

vorrichtungen. Diese lassen sich in zwei Hauptgruppen aufteilen: kontinuierlich arbeitende Systeme (wie angetriebene Radsterne und Raupenfahrwerke) oder intermittierend arbeitende Vorrichtungen (wie Steigbeine beziehungsweise Steighebel in verschiedenen Ausgestaltungen).

Zwei Beispiele mögen zunächst die grundlegenden Konstruktionsprinzipien erläutern:



Figur 2: Transportkarre mit Raupenfahrwerk (aus DE 25 52 179 A1)

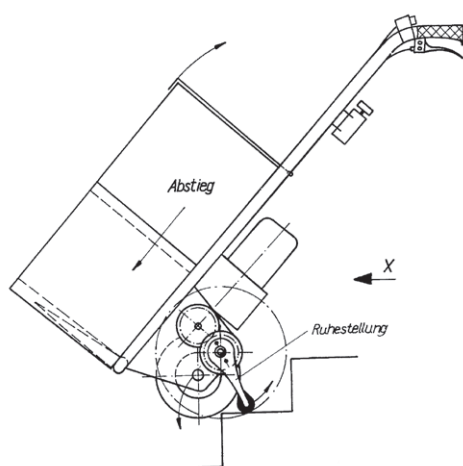
Während die Transportkarre aus Figur 2 gemäß DE 25 52 179 A1 bei einer Bewegung in der Ebene auf den nicht angetriebenen Laufrädern rollt, wird sie bei der Treppenfahrt von dem zentral angeordneten Raupenband, das auf mindestens zwei Stufenkanten aufliegt, und seitlichen Stützkufen getragen. Vorteilhaft an derartigen Raupenband-Treppensteigvorrichtungen erscheint zunächst, dass Sensorik und aufwändige Steuerungstechnik nicht zwingend erforderlich sind. Darüber hinaus funktioniert das Raupenband auf verschiedenen Treppengeometrien, sofern es nur lange genug ausgeführt ist und auf mindestens zwei Stufenkanten aufliegt. Der Benutzer muss zudem nicht mehr tun, als dem Gefährt eine Richtung zu geben und den Raupenbandantrieb erforderlichenfalls einzuschalten. Diesen Vorzügen stehen allerdings auch Nachteile gegenüber: hoher elektromechanischer Bauaufwand, eine hohe Belastung der oft empfindlichen Stufenkanten sowie unvorhersehbares Verhalten beim Übergang zur Horizontalfahrt am Ende einer Treppe durch plötzliches Kippen.





Durch die vorgestellten automatisch wirkenden Sicherungen gegen Abwärtsrollen ist schon ein wesentlicher Schritt hin zu mehr Sicherheit getan. Bei den bisher betrachteten Lösungen ist der Benutzer jedoch immer noch gefordert, die Lastkarre auf den Trittstufen auszurichten, den Schwerpunkt stets über den lasttragenden Elementen zu balancieren sowie zusätzlich den Treppensteigantrieb koordiniert mit der richtigen Drehrichtung ein- und auszuschalten.

Um dem Benutzer etwas von dieser Koordinationsübung zu entlasten sieht die Patentschrift DD 100 381 A3 (siehe Figur 5) eine Steuerung für den Motorantrieb vor. Damit muss der Bediener nur noch ein Startsignal geben, woraufhin der Steighebel für genau eine Umdrehung angeschaltet bleibt. Damit erfolgt das Heben und Absetzen auf der nächsten Stufe schon fast automatisch. Vergleichbar, aber bereits mit Endschaltern und einstellbaren Anschlägen zur Begrenzung der Hubbewegung, arbeitet die Transportkarre mit elektrischem Hubantrieb gemäß DD 117 961 A3, bestehend aus Steigbeinen mit Spindelantrieb. Die Einstellbarkeit der Anschläge bedeutet tatsächlich eine deutliche Verbesserung, da die Transportkarre nun für verschiedene Stufenhöhen optimal eingestellt werden kann.



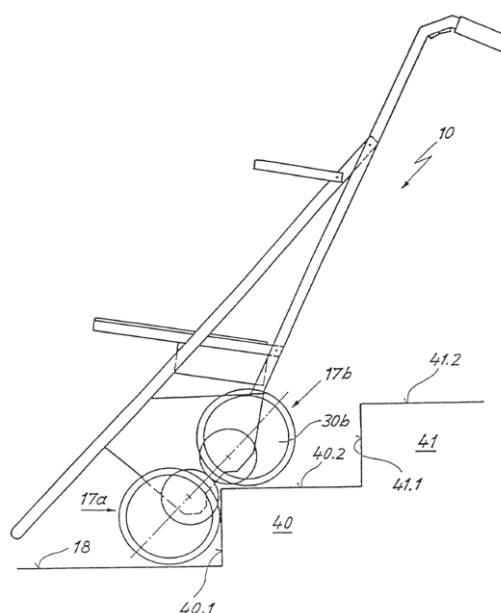
Figur 5: Transportkarre mit Kurvenscheibensteuerung (aus DD 100 381 A3)

Die bisherigen Beispiele zeigen, dass schon mit einfachen Endlagenschaltern, die erst den Anfang dessen darstellen, was heute unter Sensorik verstanden wird, der Bedienkomfort und die Sicherheit von treppengängigen Transportkarren deutlich verbessert werden kann.

### 3 Berührungslose Sensoren und Umfelderkennung

Für den elektromechanischen, lasttragenden Teil der Treppensteigvorrichtungen greifen die Anmelder nach wie vor zumeist auf bewährte Grundprinzipien zurück, da die Problematik der zu bewegenden Lasten unverändert geblieben ist. Im Bereich der Leistungselektronik und ihrer Steuerung nutzen sie jedoch die Möglichkeiten, welche sich mit dem wachsenden Potential an kostengünstigen Sensoren bieten.

Um etwa ein sanftes Aufsetzen der hier als Steigelemente dienenden Laufräderpaare durch ein geregeltes Abbremsen des Antriebs der Treppensteigvorrichtung nach DE 199 12 932 C1 sicherzustellen, schlagen deren Autoren unter anderem Infrarot- und Ultraschallsensoren vor. Diese ermitteln die Entfernung der Steigelemente von der Stufenoberfläche und stellen diese der Motorsteuerung zur Verfügung. Wahlweise kann ferner ein programmgesteuerter Ablauf für bekannte Stufenhöhen oder ein sensorgesteuerter Steigbetrieb für beliebige Stufen gewählt werden. Anzumerken ist, dass bei diesem Patent auch der schonende Transport von in ihrer Mobilität eingeschränkten Menschen im Blickfeld der Erfinder lag, siehe hierzu die Ausgestaltung als treppengängiger Rollstuhl in Figur 6.

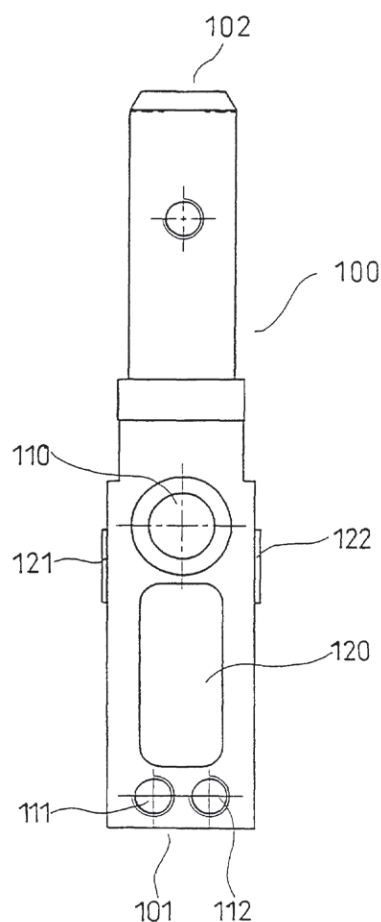


Figur 6: Treppengängiger Rollstuhl mit sensorgesteuerten Steigelementen, hier Laufräderpaar (aus DE 199 12 932 C1)

Sanftes Aufsetzen, hier bei einer Steigbeinanordnung, ist auch Gegenstand der DE 10 2012 205 204 A1, die gleichfalls als Annäherungssensoren ausgebildete Infrarotsensoren vorschlägt.

Die bisher dargestellten Transportkarren erfordern jedoch immer noch einen kräftigen, sachkundigen Bediener. Dieser muss zwar nicht mehr die Hubarbeit, so doch einiges an Haltearbeit leisten, um ein Kippen der Transportkarre treppauf oder treppab zu verhindern – dies umso mehr, wenn es sich um den Transport von Menschen handelt.

Hier setzen die Erfinder der DE 100 18 516 C1 an und sehen an der Griffstange 20 einen Haltekraftsensor 100 vor (siehe Figur 7), dort bei Ausnehmung 120. Dieser besteht aus einem elastisch deformierbaren Abschnitt der Griffstange 20 beziehungsweise deren Befestigung am Gestell. Auf diesen relativ biegeweichen Abschnitt sind Dehnmessstreifen (DMS) 121 und 122 zur



Figur 7: Haltekraftsensor in der Griffstange (aus DE 100 18 516 C1)

Ermittlung der Verformung beziehungsweise der korrelierten Kräfte appliziert. Alternativ werden zur Verformungsmessung induktive Systeme oder optoelektronische Sensoren vorgeschlagen. Um eine besonders bei DMS mögliche Drift der Nullpunktslage des Kraftsensors auszugleichen, kann ein automatischer Nullabgleich während der Aufladephasen der Spannungsquelle der Treppensteigvorrichtung vorgesehen sein. Die Haltekraft, gegebenenfalls ergänzt durch Messwerte redundanter Neigungs- und/oder Beschleunigungssensoren, wird als Indikator für bevorstehendes Kippen ausgewertet und im Gefahrenfall eine Stützvorrichtung ausgelöst. Diese redundante Sensorik ist trotz des erhöhten Aufwands hier auch vor dem Hintergrund sinnvoll, dass jegliches Fehlen von Haltekraft, wenn nämlich der Bediener den Griff völlig loslässt, sonst nicht als Gefahrenfall erkannt werden könnte.

In Fortentwicklung der auf Ermittlung des reinen Haltekraftbetrags ausgelegten Schaltungen ziehen die Autoren der DE 101 61 220 A1 auch die Wirkungsrichtungen der Kräfte am Griff als Kenngrößen für die Lage des Schwerpunkts von Last und Transportkarre in Bezug auf die Steigelemente heran. Die zugeordnete Steuerung gibt die Fortsetzung des Steigvorgangs nur frei, wenn hierfür eine sichere Lage vorliegt. Anhand eines eingelernten vorangegangenen Haltekraftverlaufs erfolgt eine Optimierung der Regelung der Geschwindigkeit des Antriebs, um einen schonenden und dennoch zügigen Transport zu ermöglichen. Über den wie vorstehend dargestellt gestalteten Haltekraftsensor hinaus können ein Geschwindigkeitssensor, ein Stufenkantensensor, ein Positionssensor zum Detektieren der momentanen Position der Steigelemente und/oder ein Neigungswinkelsensor vorgesehen sein. Neben der Berücksichtigung dieser Daten für die Regelung des Antriebs kann ein Warnsignal an die Bedienperson ausgegeben werden, wenn eine unsichere Lage droht.

Mit Blick auf die demographische Entwicklung hin zu einer älteren Gesellschaft stellen die Autoren der DE 10 2007 043 487 A1 eine Transportkarre vor, die beispielsweise als Einkaufstrolley (siehe links in Figur 8) im halbautonomen Betrieb zu nutzen ist. Unter halbautonom ist hierbei zu verstehen, dass das Fahrzeug



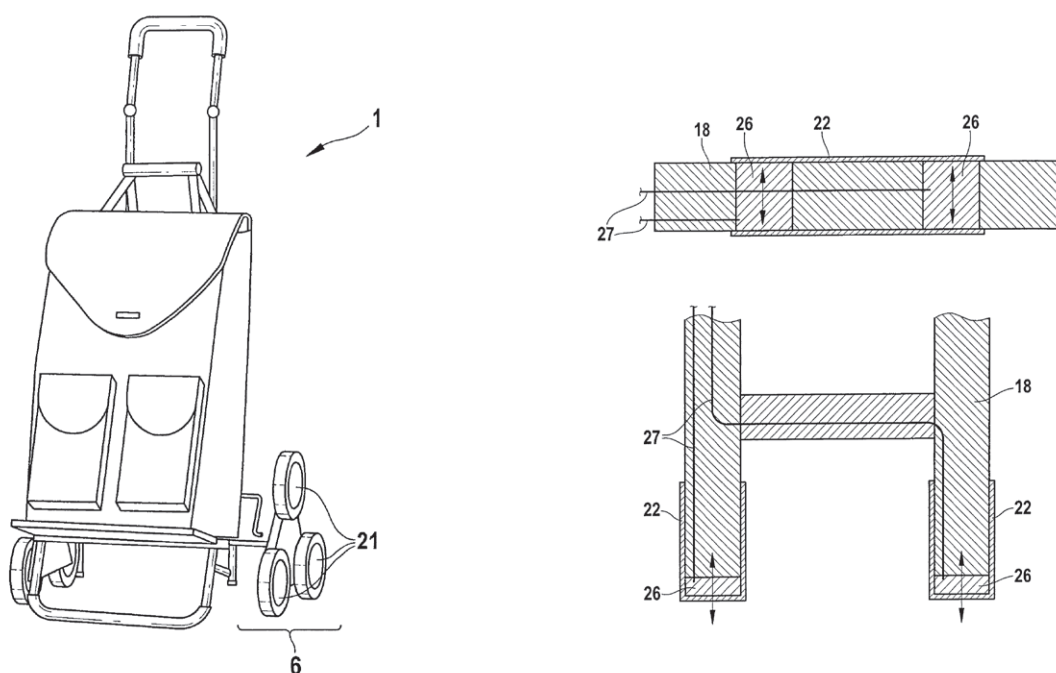
dem Benutzer unabhängig von Bodenbeschaffenheit oder Steigung folgt, sofern dieser durch Handkraft am Griff Fahrzeuggeschwindigkeit und -fahrtrichtung vorgibt. Ein programmierbarer Microcontroller wertet diese Vorgaben aus, um den elektrischen Antrieb nicht nur für das Fahren, sondern auch für das Einleiten von Lenkbewegungen zu steuern. Die von einem Radstern getragenen Laufräder werden über ein Zentralrad synchron angetrieben und ermöglichen somit auch das Befahren von Treppen. Die Sensorik zur Generierung der benötigten Sollwerte ist bei diesem Vorschlag direkt in den für Ein- oder Zweihandbedienung ausgelegten Handgriff integriert (rechts in Figur 8), wobei die Handkraft auf die Griffteile durch DMS, Spannungsteilerschaltungen oder Piezoelemente erfassbar ist.

Messung und Auswertung der Kräfte sowie ihrer Wirkungsrichtung am Handgriff sind auch beim Einkaufstrolley in Figur 9 gemäß DE 92 17 351 U1 vorgesehen. Der Kraftverlauf am Handgriff signalisiert der Steuerung die Notwendigkeit, zum Überwinden von Stufen vom elektrischen Radantrieb auf ein zusätzlich vorhandenes Raupenfahrwerk umzuschalten. Durch das von der Stufengeometrie unabhängig funktionierende Raupenfahrwerk können neben üblichen

Treppen auch Einstiegsstufen öffentlicher Verkehrsmittel erstiegen werden. Allerdings neigt eine solche Transportkarre mit Raupenfahrwerk am oberen Ende einer Treppe zum plötzlichen Kippen auf die erreichte Ebene, wie bereits eingangs erwähnt. Um dies zu vermeiden, sind im Gebrauchsmuster DE 92 17 351 U1 Stützräder (in Figur 9 nicht erkennbar) vorgesehen, die automatisch ausfahren, wenn die Steuerung aus den Kraftdaten der Griffsensorik das Bevorstehen einer solchen Kippsituation ermittelt.

Nahezu das gesamte Spektrum der heute für Fahrzeuge verfügbaren Sensorik nutzt schließlich das mit DE 10 2013 006 692 A1 angemeldete universelle autonome Fahrgestell (siehe Figur 10).

Dieses befördert beispielsweise Einkäufe auch über unebene Untergründe wie Treppen und verstaut sich nach dem Einkauf selbst im Heckbereich eines PKW, sobald es vom Benutzer durch optisch oder akustisch übertragene Befehle dazu angewiesen wird oder selbstständig einen bestimmten PKW erreicht. Darüber hinaus lässt sich das dreibeinige autonome Fahrgestell auch für Arbeiten in gefährlichen Umgebungen, wie beispielsweise bei der Bombenentschärfung, ausrüsten. Voraussetzung für derart anspruchsvolle Aufgaben

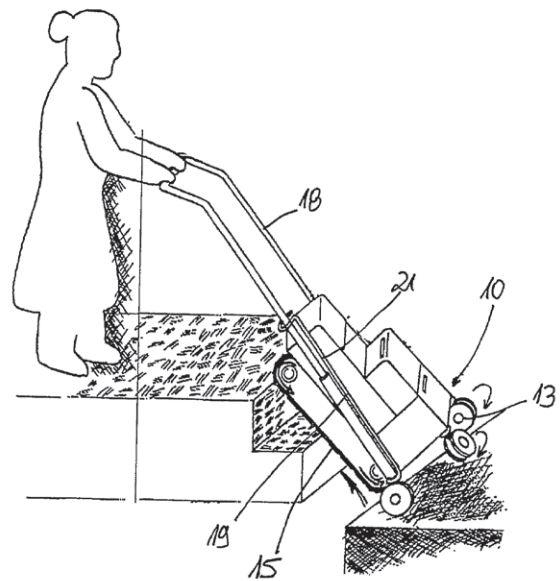


Figur 8: Halbautonomer Einkaufstrolley mit Haltekraftsensorik im Ein- und Zweihandgriff (aus DE 10 2007 043 487 A1)

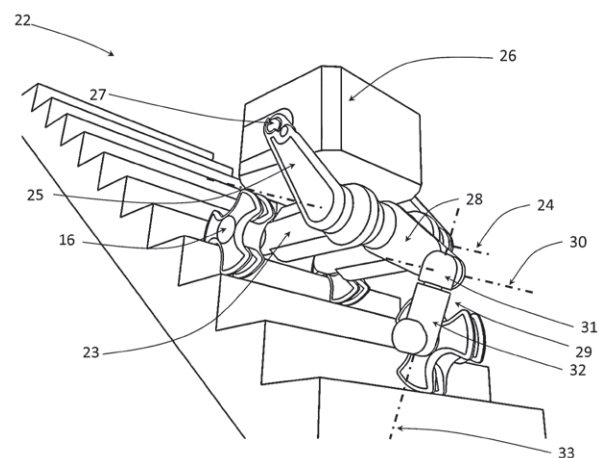
sind Mittel zur Erfassung des Umfeldes wie optische, akustische, elektromagnetische, taktile und kinetische Sensoren sowie Kombinationen davon. Speziell gestaltete Räder an dem dreibeinigen Transportgerät umfassen gegeneinander verdrehbare Segmente, die sowohl das Befahren von Treppen und sonstigen Hindernissen als auch ebenen Untergründen ermöglichen. Damit entfällt auch der Umstand, eine die Nutzlast reduzierende Treppensteigeinrichtung zusätzlich zu Laufrädern mitführen zu müssen. Ferner ist für die Anwendung in vorbestimmten Umgebungen, beispielsweise einem regelmäßig zu befahrenden Treppenhaus, eine Nahfeldkommunikation dieses Transportgeräts mit ortsfesten Sensoren vorgesehen.

#### 4 Ausblick

Die im Lichte der jeweils implementierten Sensortechnik betrachteten treppengängigen Transportvorrichtungen lassen unter Berücksichtigung ihrer zeitlichen Einordnung gegenwärtig zwei Perspektiven erkennen: Auf der einen Seite die einfach, robust und kostengünstig konstruierten Transportkarren, die zumindest für das Bewältigen von Stufen mit Elektroantrieben ausgestattet sind. Auf der anderen Seite hochkomplexe Geräte, die eher den Namen Roboter als Transportkarre verdienen und die der Öffentlichkeit bislang allenfalls in der Anwendung bei Militär oder Katastrophenschutz bekannt sind. Auf Grund der demographischen Entwicklung ist jedoch künftig von einem steigenden Bedarf an elektrifizierten treppengängigen Transportvorrichtungen besonders im privaten Bereich auszugehen. Vor dem Hintergrund der freien Verfügbarkeit kamerabestückter Drohnen zu Discounterpreisen, die vor wenigen Jahren kaum vorstellbar schien, erscheint die Hoffnung gar nicht so abwegig, dass wir beim nächsten Umzug gemächlich dabei zusehen können, wie ein sensorgesteuertes Transportsystem unsere Waschmaschine die Treppe hochträgt.



Figur 9: Einkaufstrolley mit Haltekraftsensorik  
(aus DE 92 17 351 U1)



Figur 10: Universales autonomes Fahrgestell  
(aus DE 10 2013 006 692 A1)

# Magnetische Positionssensoren in der Anwendung

Dr. Christopher Root, Patentabteilung 1.56

Präzise und robuste Positionsmesstechnik ist eine wichtige Voraussetzung für die zunehmende Automatisierung in allen Lebensbereichen. Dazu werden in der Mess- und Sensortechnik auch magnetische Eigenschaften und Effekte ausgenutzt. Der nachfolgende Artikel gibt einen Überblick über unterschiedliche Einsatzgebiete bekannter Magnetsensorik und stellt jeweils kurz dar, wie auf Grundlage einer magnetischen Wechselwirkung ein Positionssignal erzeugt wird.

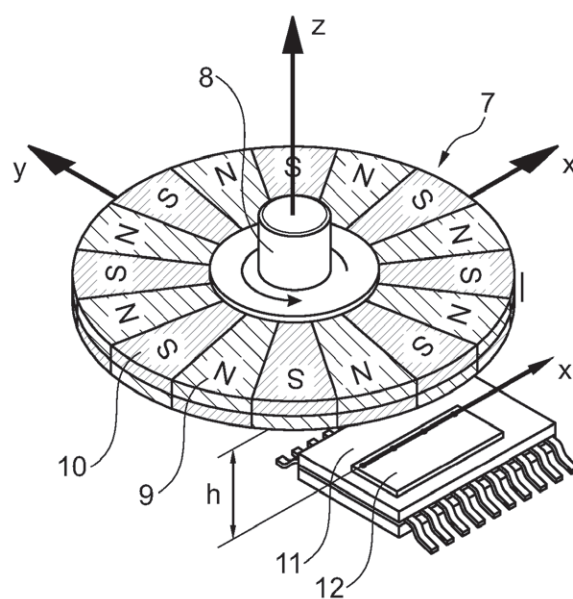
## 1 Einführung

In verschiedensten technischen Gebieten ist es von grundlegendem Interesse, die relative Lage von gegeneinander beweglichen Bauteilen zu kennen. Üblicherweise werden dabei räumliche Koordinaten, relative lineare Verschiebungen oder Rotationen bestimmt. Die erhaltenen Werte dienen einer laufenden Betriebsüberwachung und Protokollierung oder werden zur weiteren Verarbeitung einer Steuerung und Regelung zugeführt. Die Anwendungsgebiete reichen dabei von der Produktionstechnik über die Fahrzeugtechnik bis zu Gegenständen des täglichen Lebens. Die Messtechnik stellt für die Positionsbestimmung eine Vielzahl unterschiedlicher Lösungsansätze bereit, die abhängig von den Details des spezifischen Einzelfalls eingesetzt werden können. Neben den hier ausführlicher dargestellten Anwendungen, die auf magnetischen Effekten und Eigenschaften basieren, sind insbesondere auch elektrische, akustische und optische Methoden in der Praxis weit verbreitet [1].

Die Beispiele des Artikels sind im Hinblick auf die verwendete Sensorik bewusst so ausgewählt, dass sich einerseits darin die vielfältigen Einsatzmöglichkeiten der magnetischen Positionssensoren widerspiegeln. Andererseits soll aber auch eine breite Palette bekannter und weniger bekannter magnetischer Effekte, die in der technischen Entwicklung der letzten Jahre für die Positionsbestimmung Verwendung fanden, vorgestellt werden.

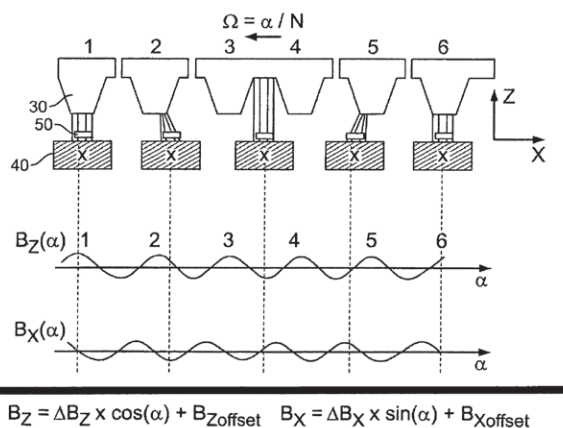
## 2 Magnetische Sensoren und Sensorsysteme

Damit eine Lage, eine Lageveränderung oder eine Bewegung eines Körpers mit magnetischen Methoden erfasst werden kann, muss die Position oder Positionsveränderung mit einem Magnetfeld oder einer Änderung eines magnetischen Feldes verknüpft werden. Dies kann auf unterschiedliche Weise erfolgen. Zum Beispiel ist eine räumliche Änderung der relativen Lage eines magnetfelderzeugenden **Geberlements** und eines Magnetfeld-Sensors in Form einer Abstandsänderung oder Rotation möglich. Die DE 10 2012 213 717 A1 zeigt in Figur 1 einen Inkrementalwegsensor zur Linear- und Winkelmessung.



Figur 1: Inkrementalwegsensor (aus DE 10 2012 213 717 A1)

Dabei wird auf einem scheibenförmigen Maßkörper 7 eine wechselnde Abfolge von Nord- und Südpolen aufgebracht. Bei einer Bewegung des Maßkörpers, dargestellt ist die Rotation der Welle 8, werden die Pole 9 und 10 am Sensor 11 vorbeigeführt. Durch die Zählung der wechselnden Polfolge kann eine Aussage über die zurückgelegte Wegstrecke oder die Anzahl von Umdrehungen gemacht werden. In einer anderen exemplarischen Herangehensweise, für die die WO 2012/010507 A1 als Beispiel dient, ist der Abstand zwischen einem Magneten und einem Sensorelement fixiert (siehe Figur 2). Ein Element 40, das auf der Rückseite des Sensors 50 angeordnet ist, erzeugt dazu ein Vorspannungsmagnetfeld. Ziele 30 aus ferromagnetischem Material bilden ein Profil aus Zähnen und Nuten. Durch die Bewegung der Ziele wird die Richtung und Dichte der Magnetfeldlinien am Ort des Sensors 50 geändert.



Figur 2: Magnet und Sensor mit einem fixen Abstand (aus WO 2012/010507 A1)

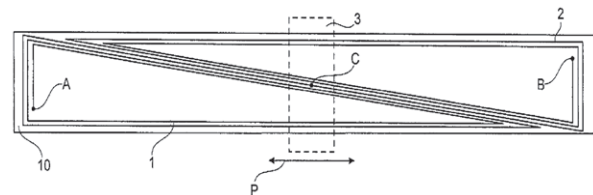
Um eine Positionsmessung zu ermöglichen, werden hier also gezielt weitere Bauteile eingesetzt, die aufgrund ihrer Form oder ihrer magnetischen Eigenschaften das messbare Magnetfeld beeinflussen. Beide Prinzipien werden im Folgenden näher vorgestellt.

## 2.1 Spulen

Spulen bieten eine bekannte und weit verbreitete Möglichkeit, Veränderungen eines magnetischen Feldes zu erfassen und damit eine Positionsveränderung zu messen. Mittels Induktion führt jede Veränderung

eines Magnetfeldes, das die Spule durchsetzt, im Zusammenwirken mit der Induktivität der Spule zu einem messbaren Signal.

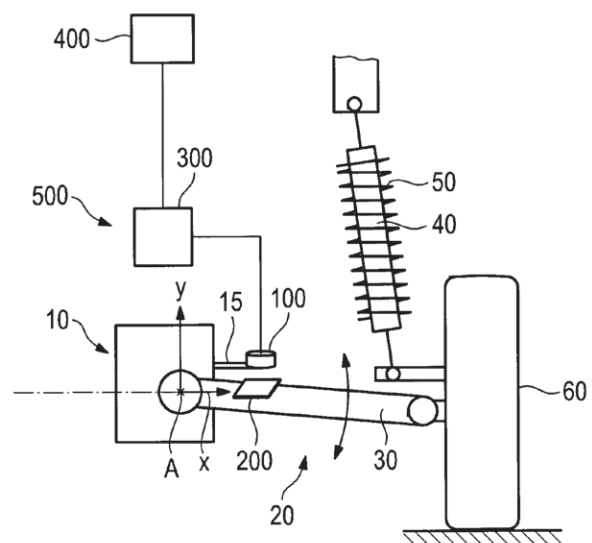
Die DE 10 2011 077 118 A1 betrifft eine geometrische Anordnung zweier Spulen 1 und 2 mit jeweils einem ausgeprägt keilförmigen Querschnitt (vergleiche Figur 3).



Figur 3: Schematisierte Draufsicht einer Weggeberanordnung (aus DE 10 2011 077 118 A1)

Der dicht oberhalb der Spulen angeordnete Schieber 3 lässt sich entlang der Richtung P verstellen und verändert dabei die Induktivität der Spulen 1 und 2 in einander entgegengesetzter Weise. Die Stellung des Schiebers wird aus dem Signalverlauf der gemessenen Spannungen ermittelt, wenn die Anordnung mit einer Wechselspannung beaufschlagt wird.

Figur 4 der Patentschrift DE 10 2012 015 035 B3 zeigt, wie der Abstand zweier zueinander beweglicher Teile 10 und 30 eines Fahrzeugs bestimmt wird.

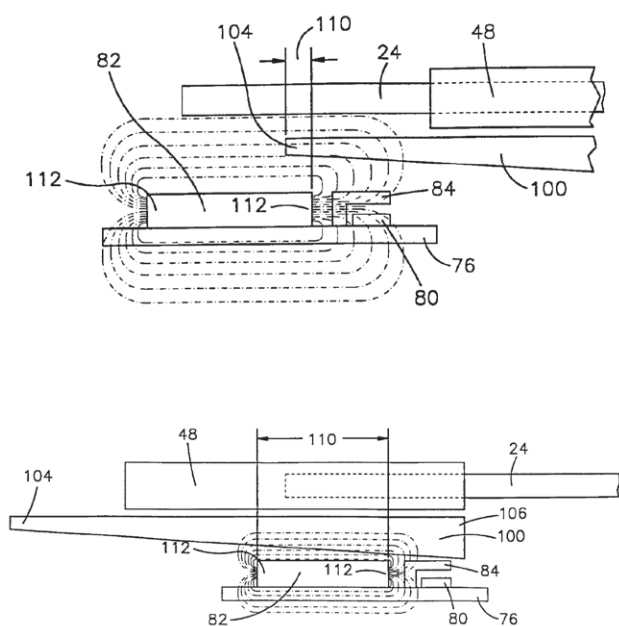


Figur 4: Abstandsbestimmung zweier zueinander beweglicher Teile (aus DE 10 2012 015 035 B3)

Ein Fahrzeugteil 10 enthält eine induktive Messeinheit 500 mit einer Spule 100, mit der ein elektromagnetisches Feld erzeugt wird. Das andere Fahrzeugteil 30 hingegen ist elektrisch leitfähig oder mit einem elektrisch leitfähigen Körper 200 ausgestattet, in dem auf die Erregung des elektromagnetischen Feldes hin Wirbelströme erzeugt werden. Dadurch verändert sich die Dämpfung eines elektrischen Schwingkreises, in den die Spule 100 eingebunden ist. Im dargestellten Ausführungsbeispiel handelt es sich bei den Fahrzeugteilen um das Chassis 10 und den Querlenker 30 des Fahrwerks 20 eines Fahrzeugs.

## 2.2 Hall-Sensoren

Das Messprinzip von Hall-Sensoren beruht auf der Anwendung des Hall-Effekts. Sie können Betrag und Richtung einer Komponente eines Magnetfeldes, die ihre Messfläche senkrecht durchsetzt, bestimmen und stellen in der gegenwärtigen Technik eine verbreitete Möglichkeit zur Magnetfeldmessung dar. Durch entsprechende Ausrichtung mehrerer Sensorflächen ist auch eine räumliche – im Sinne einer richtungsaufgelösten – Messung möglich.

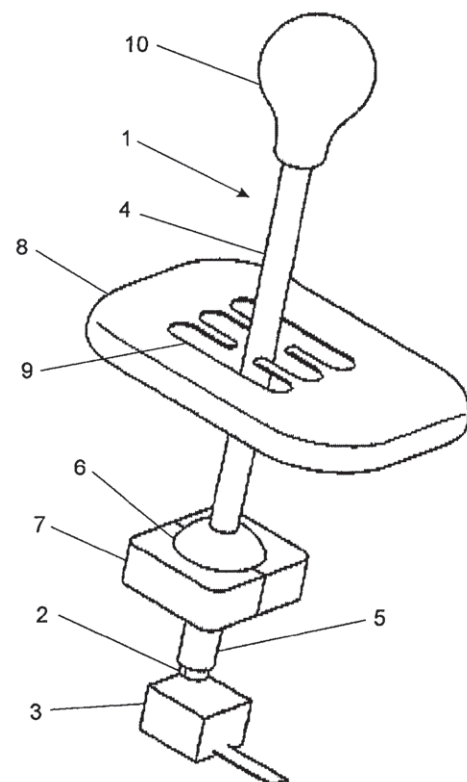


Figur 5: Schematische Ansicht zweier Stellungen einer Vorrichtung, mit der die Position eines Fahrzeugsitzes bestimmt wird (aus DE 103 57 750 A1)

In Figur 5 aus der DE 103 57 750 A1 wird eine Vorrichtung gezeigt, mit der die Position eines Fahrzeugsitzes bestimmt wird.

Dazu wird das Feld eines Magneten 82 von einem Hall-Sensor 80 detektiert. Die beiden befinden sich dabei zusammen auf einem Träger 76. Die Position des Fahrzeugsitzes wird durch ein mit dem Sitz verbundenes ferromagnetisches Bauteil 100 gewonnen. Durch die keilförmige Form des Bauteils 100 variiert bei einer Verschiebung des Sitzes die Größe des ferromagnetischen Materials in die Nähe des Magneten 82. So wird das vom Sensor 80 jeweils gemessene Magnetfeld gezielt beeinflusst. Die durch die Vorrichtung gewonnenen Daten werden für den Insassenschutz verwendet. Zum Beispiel wird das Aufblasen eines Airbags abhängig von der gemessenen Sitzposition gesteuert.

Die DE 10 2007 026 303 B4 zeigt in Figur 6 eine Vorrichtung zur Erfassung und Auswertung der Gangwahl in einem Kraftfahrzeug.

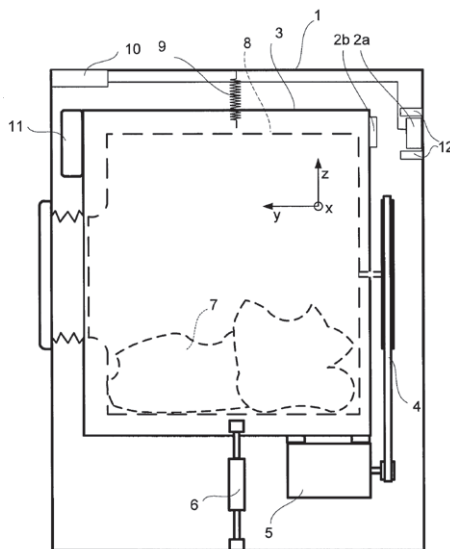


Figur 6: Schematische Darstellung einer Vorrichtung zur Erfassung der Gangwahl bei einem Kraftfahrzeug (aus DE 10 2007 026 303 B4)



Der Bedienabschnitt 4 eines Wählhebels 1 kann innerhalb der Hebelführung 8 verschiedene Positionen 9 einnehmen. Am unteren Ende des Schaltabschnitts 5 ist ein Magnet 2 befestigt. Unterhalb des Magneten 2 ist ein Sensorelement 3 angeordnet, das die räumlichen Komponenten des Magnetfeldes detektiert. Damit kann über die relative Position des Magneten 2 zum Sensorelement 3 auf die Position des Wählhebels geschlossen werden.

Die Gebrauchsmusterschrift DE 20 2007 002 626 U1 beschreibt ein Wäschebehandlungsgerät mit Detektor, wie eine Waschmaschine (siehe Figur 7) oder einen Wäschetrockner.



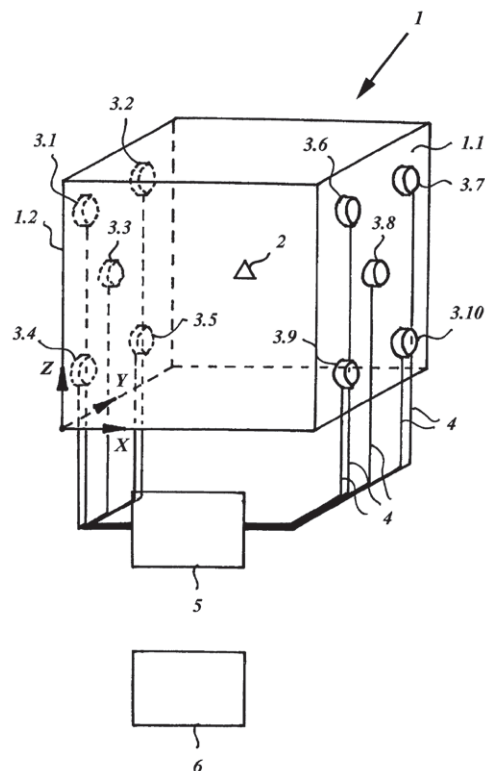
Figur 7: Waschmaschine mit Detektorvorrichtung  
(aus DE 20 2007 002 626 U1)

In einem Gehäuse 1 ist mittels einer Feder 9 ein Behälter 3 schwingend aufgehängt. Die Vorrichtungen 2a und 2b bestimmen die relative Lage des Behälters innerhalb des Gehäuses. Dafür wird mit einem Magneten 2b ein Magnetfeld erzeugt, dessen räumliche Verteilung durch eine mehrere Hall-Elemente enthaltende Detektoreinrichtung erfasst wird. Somit können durch Beladung und Unwucht entstehende Schwingungen, die zu einem Anschlagen des Behälters 3 an der Gehäusewand 1 führen können, rechtzeitig erkannt und Gegenmaßnahmen eingeleitet werden. Darüber hinaus ist es möglich, die durch das Gewicht der Wäsche verursachte Längenänderung der Feder 9 auszumessen und so die in das Gerät eingebrachte Wäsche zu wiegen.

### 2.3 Magneto-resistive Sensoren

Bei magneto-resistiven Sensoren kommt es infolge der Einwirkung eines magnetischen Feldes zu einer Veränderung ihres elektrischen Widerstands. Magneto-resistive Effekte treten in ferromagnetischen Materialien und in ferromagnetisch-metallischen Dünnschichtsystemen auf. Man unterscheidet bei den heute im Allgemeinen eingesetzten magneto-resistiven Sensoren insbesondere einen **anisotropen** magneto-resistiven Effekt (AMR – anisotropic magnetoresistance), einen großen magneto-resistiven Effekt (GMR – giant magnetoresistance) sowie einen magnetischen Tunnelwiderstand (TMR – tunnel magnetoresistance). Die einzelnen Effekte besitzen dabei unterschiedliche physikalische Ursprünge. Ebenso variieren sie in der Größe der beobachtbaren Widerstandsänderungen. Eine Übersicht mit näheren Einzelheiten hierzu findet man in der Literatur [2], [3] und [4].

Die EP 2 572 167 B1 betrifft eine Vorrichtung mit der die räumlichen Koordinaten eines **Sensorknotens** 2 bestimmt werden (vergleiche Figur 8).

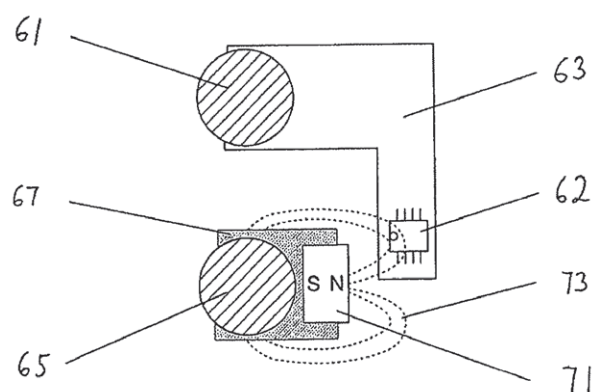


Figur 8: Behälter mit Spulen zur Bestimmung der Position eines Sensorknotens (aus EP 2 572 167 B1)



Im beschriebenen Ausführungsbeispiel ist der Sensorknoten aus drei AMR-Sensoren aufgebaut, die die entsprechenden räumlichen Anteile eines magnetischen Flusses erfassen. Dieser wird von Spulen 3.1 bis 3.10 erzeugt, die auf den Außenwänden eines Behälters 1 angebracht sind. Nacheinander werden diese Spulen kurzzeitig jeweils in zwei Richtungen bestromt. Durch Differenzbildung der Messungen mit unterschiedlicher Bestromung kann die absolute Magnetfeldstärke ermittelt und somit auf den Abstand des Sensorknotens zur jeweils felderzeugenden Spule geschlossen werden. Die EP 2 572 167 B1 schlägt vor, ein solches System beispielsweise für Gärtanks einzusetzen, deren Dimensionsmaße mehrere Meter betragen. Mit diesem System können dort eingebrachte Sonden für Temperatur, Druck oder pH-Wert zusätzlich ortsaufgelöst überwacht werden.

Die DE 10 2010 014 668 A1 behandelt einen Stellantrieb eines Ventils, bei dem die Stellung des Ventilglieds erfasst wird. Zu diesem Zweck wird an einer Feder, die das Ventilglied in eine Sicherheitsstellung drängt, eine Vorrichtung aus einem Magnetfeld-Sensor und einem **Geberelement** angebracht. Durch die homogen-elastische Verformung der Feder bei der Bewegung des Ventilglieds lässt sich auf deren gesamte Längenänderung und weiter auf die Stellung des Ventilglieds schließen. Im dargestellten Ausführungsbeispiel in Figur 9 ist ein AMR-Sensor 62 mittels eines Kopplungsarmes 63 an einer Windung 61 einer Spiralfeder angebracht. Eine weitere Windung 65 der Spiralfeder trägt einen Permanentmagneten 71.

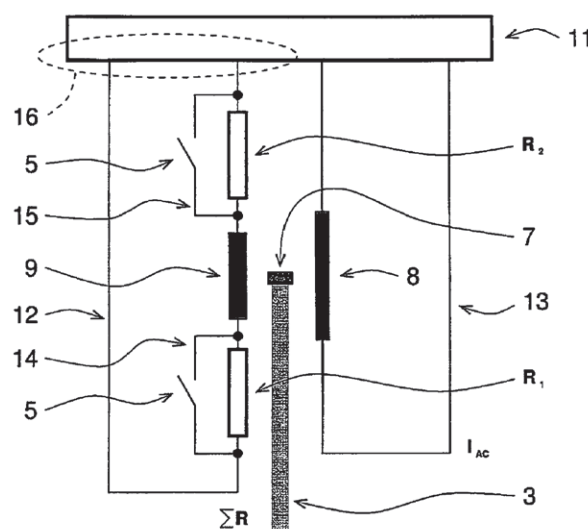


Figur 9: Darstellung einer Spiralfeder mit einem AMR-Sensor und einem Permanentmagnet (aus DE 10 2010 014 668 A1)

## 2.4 Reed-Schalter

Ein Reed-Schalter ist aus zwei ferromagnetischen Schaltungen aufgebaut, die sich in einem evakuierten oder mit Schutzgas gefüllten Glasröhrchen befinden. Die beiden Schaltungen bilden Kontakte, die beim Einwirken eines entsprechenden Magnetfeldes, je nach Konstruktion des Sensors, entweder geöffnet oder geschlossen werden [5].

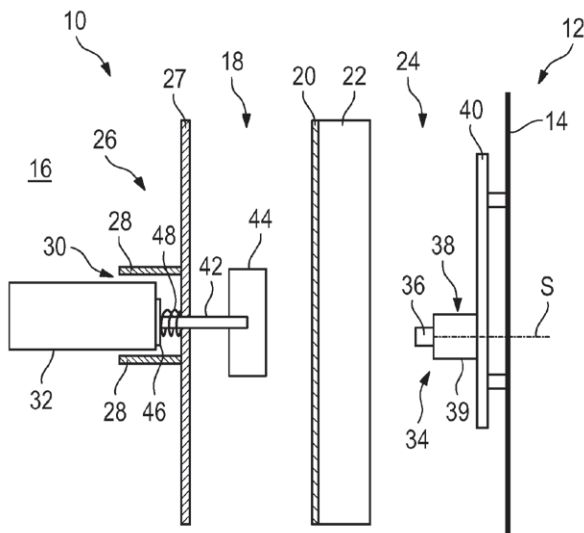
In der DE 10 2014 050 765 B9 wird ein Stellungsmesssystem beschrieben (vergleiche Figur 10), bei dem die Position eines Stabs 3, an dem ein Magnet 7 befestigt ist, entlang dem geradlinigen Weg eines Führrsystems erfasst wird.



Figur 10: Schematische Teilansicht eines Stellungsmesssystems, bei dem die minimale und die maximale Position eines Stabs jeweils mit Hilfe eines Reed-Schalters erfasst werden (aus DE 10 2010 050 765 B9)

Das System ist vorgesehen für den Einsatz in einer kerntechnischen Anlage. Damit die maximale und die minimale Auslenkung des Stabs erkannt werden, ist an den entsprechenden Positionen des Führsystems jeweils ein Reed-Schalter 5 angebracht. Je nach Position des Stabs 3 wird durch die magnetische Wechselwirkung ein Schalter 5 geschlossen und der jeweils parallel dazu geschaltete Widerstand  $R_1$  oder  $R_2$  überbrückt. Aus den sich daraus ergebenden Signalen ist die Position des Stabs ableitbar.

Das Gebrauchsmuster DE 20 2014 104 977 U1 zeigt in Figur 11, wie bei einem Gargerät eingeschobene Zubehöerteile durch einen Reed-Sensor erkannt werden.



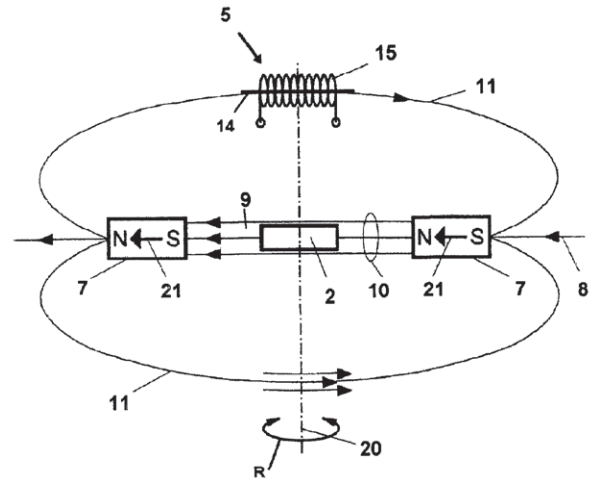
Figur 11: Detailansicht der schematischen Darstellung eines Gargeräts zur Erfassung eines Zubehöerteils (aus DE 20 2014 104 977 U1)

In den Garraum 16 können in Schienen 28 Zubehöerteile 32 eingeschoben werden. Mittels einer federnd gelagerten Verstelleinrichtung 42, 46 und 48 wird die Stellung eines Magneten 44 verändert. Hinter einer isolierten Wand 20 und 22 ist ein Sensor 36 angebracht, mit dem die entsprechende Position des Magneten 44 erfasst werden kann. Dabei kann ein Reed-, aber auch ein Hall-Sensor verwendet werden.

## 2.5 Wiegand-Sensoren

Wiegand-Sensoren, auch als Impulsdraht-Sensoren bekannt, bestehen aus einem Wiegand-Draht, der einen magnetfeldempfindlichen Impulsgeber bildet und einer in der Nähe angebrachten Spule als Impulsnehmer. Ein Wiegand-Draht ist aus einem **weichmagnetischen** Kern und einem **hartmagnetischen** Mantel aufgebaut. Unter Einwirkung eines äußeren magnetischen Feldes kommt es im Wiegand-Draht zu einer sprunghaften Ummagnetisierung zunächst des Kerns und dann gegebenenfalls des Mantels. Dies ist jeweils als Spannungsimpuls in der Spule nachweisbar ([6], [7] und [8]).

Die DE 10 2010 022 154 A1 beschreibt die Verwendung eines Wiegand-Sensors als Zählereinheit für einen Umdrehungszähler (vergleiche Figur 12).



Figur 12: Schematisierte Ansicht eines Drehgebers (aus DE 10 2010 022 154 A1)

Bei der Bewegung um die Achse 20 bilden die beiden Permanentmagnete 7 eine Geberinheit, deren homogenes Magnetfeld 9 von einer Sensoreinheit 2, hier einem Hall-Sensor, detektiert wird. Für die Umdrehungszählung ist im Streufeld 11 der Magnete 7 ein Wiegand-Sensor 5 angeordnet, dessen Impulse bei jeder vollen Umdrehung ausgewertet werden.

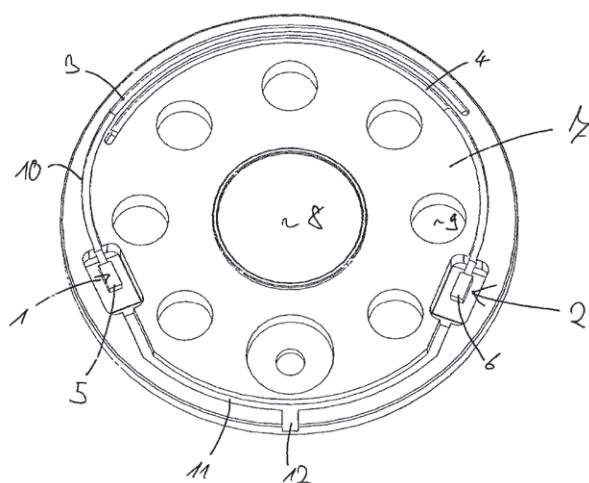
## 2.6 Magnetostruktive Sensoren

Magnetostruktive Effekte treten in ferromagnetischen Materialien auf. Sie verknüpfen ein auf einen entsprechenden Körper wirkendes Magnetfeld mit messbaren mechanischen Veränderungen, sowie auch umgekehrt eine mechanische Einwirkung auf einen Körper mit einer Änderung seiner magnetischen Eigenschaft. Man unterscheidet dementsprechend auf der einen Seite den Joule- und den Wiedemann-Effekt sowie auf der anderen Seite den Villari- und den Matteucci-Effekt [9] und [10].

In den beiden folgenden Beispielen wird eine Positionsmessung beschrieben, bei der ein ferromagnetischer Wellenleiter mit einem Stromimpuls beaufschlagt wird. An der Position eines am Wellenleiter angeord-

neten Magneten entsteht eine mechanische Welle, die sich in Richtung der Enden des Wellenleiters mit bekannter Geschwindigkeit ausbreitet. Aus der ermittelten Laufzeit der Welle kann auf diese Weise auf die Position des Magneten geschlossen werden.

Die Gebrauchsmusterschrift DE 20 2012 008 717 U1 offenbart eine Sensoranordnung zur Messung des Winkels zwischen einer Sattelzugmaschine und einem Auflieger (siehe Figur 13).



Figur 13: Magnetostriktiver Wellenleiter zweier Sensoren (aus DE 20 2012 008 717 U1)

Die redundant ausgeführte Sensoranordnung 1 und 2 ist jeweils aus zwei Teilen aufgebaut. Ein erstes Teil besteht aus einem bogenförmigen magnetostriktiven Wellenleiter 3 und 4 als Messstrecke, der sich um einen Königszapfen oder eine Königszapfenaufnahme der Sattelzugmaschine oder des Aufliegers erstreckt. Das zweite Teil ist ein Magnet, der sich bei einer Drehbewegung zwischen Sattelzugmaschine und Auflieger entlang der Messstrecke des ersten Teils bewegt. Am Kopfende jedes Wellenleiters ist jeweils eine Signalelektronik angebracht, die zum einen den oben beschriebenen Stromimpuls generiert und zum anderen die entstehende mechanische Welle, einen Torsionsimpuls, auswertet.

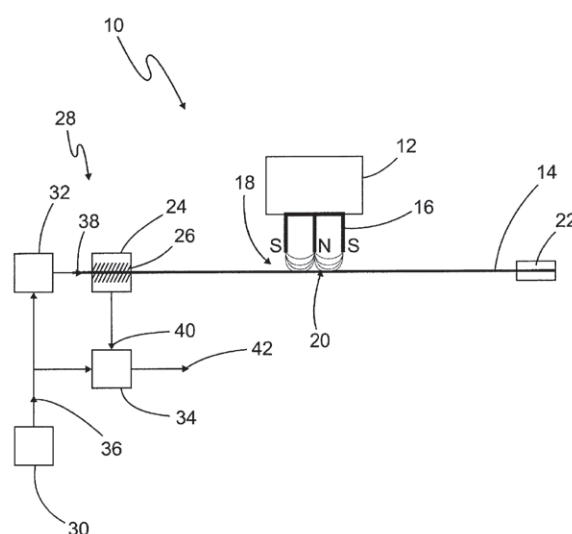
In Figur 14 der DE 10 2010 008 495 A1 wird eine Positions-Messvorrichtung gezeigt, die die Position eines Objekts 12 bezüglich eines Wellenleiters 14 ermittelt.

Dabei ist dem Objekt 12 ein Magnet 16 derart zugeordnet, dass eine Komponente 20 seines Magnetfeldes im Wellenleiter 14 wirken kann. An einem Ende des Wellenleiters ist eine Dämpfungsvorrichtung 22 angebracht. Am anderen Ende befindet sich eine Vorrichtung 24 zur Erfassung der mechanischen Welle und zur Bereitstellung eines geeigneten Stromimpulses. Für die Erfassung ist vorgesehen, den Wellenleiter 14 mit einem Eigenmagnetfeld senkrecht zu seiner axialen Richtung zu beaufschlagen. Dadurch kann in der Spule 26, wiederum durch magnetostruktive Wechselwirkung, die einlaufende Welle detektiert werden.

### 3 Zusammenfassung und Ausblick

Die Positionsbestimmung unter Einsatz magnetischer Messsensorik hat inzwischen in vielen Bereichen der Technik Einzug gehalten und ist zu einem unverzichtbaren Bestandteil geworden. Die eingesetzten Sensoren sind in großen Stückzahlen herstellbar und teilweise auch für den Einsatz unter extremen Bedingungen geeignet.

Da sich die technischen Möglichkeiten, immer mehr Daten zu erfassen und zu verarbeiten, rasant entwickeln, ist die Grundlage für eine vielfache Verwendung und weitere Verbreitung von Messsensorik vorhanden. Auch die Positionsbestimmung wird dabei weiterhin eine wesentliche Rolle spielen.



Figur 14: Positions-Messvorrichtung (aus DE 10 2010 008 495 A1).

In den dargestellten Beispielen sind die jeweils eingesetzten Sensoren teilweise in ihrer Aufgabenstellung austauschbar. Jedoch besitzt im Nebeneinander der unterschiedlichen Sensorik jeder der Sensoren unterschiedliche Vor- und Nachteile, die im Einzelfall gegeneinander abzuwägen sind. Es ist also davon auszugehen, dass auch in Zukunft verschiedene Ansätze genutzt und weiterentwickelt werden, um unter Einsatz magnetfeldempfindlicher Sensoren Positionsdaten auszumessen.

- [10] HERING, E.; SCHÖNFELDER, G.: Sensoren in Wissenschaft und Technik. 1. Auflage, Wiesbaden: Vieweg+Teubner Verlag, 2012, S. 21–24. – ISBN 978-3-8348-0169-2

### Nicht-Patentliteratur

- [1] HERING, E.; SCHÖNFELDER, G.: Sensoren in Wissenschaft und Technik. 1. Auflage, Wiesbaden: Vieweg+Teubner Verlag, 2012, S. 127–312. – ISBN 978-3-8348-0169-2
- [2] LOREIT, U., ACHENBACH, J.: Magnetoresistive Sensoren machen mobil. Automotive, 2004, Nr. 7/8, S. 24–27
- [3] GRÜNBERG, P.: Kopplung macht den Widerstand. Physik Journal, Vol. 6, 2007, Nr. 8/9, S. 33–39
- [4] WASSERMANN E.; HILLEBRANDS B.: Riesen- effekt in dünnen Schichten. Physik Journal, Vol. 6, 2007, Nr. 12, S. 23–25
- [5] Reed Technologie; URL: [http://www.meder.com/fileadmin/meder/images/marketing/Marketing\\_Resources/SME\\_Datenbuch\\_-\\_Reed\\_Technologie\\_DE.pdf](http://www.meder.com/fileadmin/meder/images/marketing/Marketing_Resources/SME_Datenbuch_-_Reed_Technologie_DE.pdf) [abgerufen am 10.04.2015]
- [6] TRÄNKLER, H.-R.; OBERMEIER, E.: Sensortechnik – Handbuch für die Praxis. 1. Auflage, Springer-Verlag, 1998, S. 504–505. – ISBN 978-3-662-09866-0
- [7] HERING, E.; SCHÖNFELDER, G.: Sensoren in Wissenschaft und Technik. 1. Auflage, Wiesbaden: Vieweg+Teubner Verlag, 2012, S. 237–239. – ISBN 978-3-8348-0169-2
- [8] Wikipedia: Wiegand-Sensor; URL: <https://de.wikipedia.org/wiki/Wiegand-Sensor> [abgerufen am 17.4.2015]
- [9] Magnetostriktion – Physikalische Grundlagen URL: [http://www.mtsensor.de/fileadmin/medien/downloads/mts\\_messprinzip.pdf](http://www.mtsensor.de/fileadmin/medien/downloads/mts_messprinzip.pdf) [abgerufen am 13.04.2015]

# Halbleiter- und Festkörperbauelement-Sensoren

*Dr. Ralf Henninger, Patentabteilung 1.33*

Autos warnen uns, wenn wir zu dicht auffahren und bleiben dank Antiblockiersystem (ABS) trotz Vollbremsung lenkbar. Fitnessarmbänder messen unseren Puls. Durch bildgebende Verfahren aus der Medizintechnik können Ärzte in unseren Körper blicken und Organen bei der Arbeit zuschauen. Unser Mobiltelefon schießt mit einer winzigen Kamera auch bei schlechten Lichtverhältnissen gute Fotos. Waschmaschinen wuchten sich beim Schleudern selber aus und mit einer Festplatte so klein wie ein Kartenspiel können wir in einer Million Büchern lesen. All das ermöglichen uns heute kleine Bauelemente (Chips) mit Halbleiter- und Festkörperbauelement-Sensoren, die zum Teil nur so klein wie ein Stecknadelkopf sind. Wie der sensitive Teil aktueller Sensoren aufgebaut ist und welche physikalischen Effekte ausgenutzt werden, soll im Folgenden vorgestellt werden.

---

## 1 Einleitung

Das Thema Sensoren betrifft wegen der vielfältigen Anwendungsbereiche ein interdisziplinäres Gebiet der Technik. Um dies zu illustrieren, werden anhand von Anwendungsbeispielen, gegliedert nach den auftretenden physikalischen Effekten, die Funktionsweise und der innere Aufbau der Sensor-Bauelemente dargestellt.

Die internationale Patentklassifikation für Patentdokumente (IPC) gliedert die Unterklasse der Halbleiter- und Festkörperbauelemente (H01L) zum Teil ebenfalls nach physikalischen Effekten. Daher wird in diesem Artikel zu Beginn jeden Abschnitts die zugehörige Hauptgruppe der IPC in Klammern angegeben. In der IPC kann man im Internet mittels [DEPATISnet](#) bequem nach solchen Klassifikationssymbolen und Stichwörtern recherchieren [1]. Zu den Anwendungsbeispielen werden somit die relevanten IPC-Stellen angegeben.

## 2 Feldeffekt- und Bipolartransistoren (H01L 29/00)

Feldeffekt- und Bipolartransistoren bestehen aus Halbleitern. Silizium ist hier nach wie vor der gebräuchlichste Halbleiter. Jeder kennt Silizium in seiner oxidierten Form, nämlich als hellen Sand am Strand. Siliziumoxid ist ein elektrischer Isolator und leitet keinen Strom.

Im Unterschied dazu hat einkristallines reines Silizium, also nicht oxidiertes Silizium, eine silbergraue Farbe und wird in der Halbleiterindustrie in der Form von Scheiben (Wafern) verarbeitet.

Silizium ist isolierend, weil alle Elektronen unbeweglich in niedrigen Energieniveaus (Valenzband) liegen. Zwischen dem Valenzband und den energetisch höher liegenden unbesetzten Zuständen (Leitungsband), in denen sich die Elektronen frei bewegen könnten, liegt eine Energielücke, die sogenannte Bandlücke. Wird jedoch in das Kristallgitter des reinen Siliziums in geringsten Mengen ein sogenannter Dotierstoff eingebaut, zum Beispiel im Verhältnis 1:10 000 000, so wird Silizium leitend. Verwendet man als Dotierstoff die Elemente Phosphor oder Arsen, die ein Valenzelektron mehr als Silizium besitzen, so leitet Silizium fast wie ein Metall die zusätzlichen Elektronen im Leitungsband gut. Man spricht hierbei von einer n-Dotierung. Verwendet man dagegen Bor als Dotierstoff, das ein Elektron weniger als Silizium aufweist, wird Silizium für Löcher leitend und man spricht von einer p-Dotierung. Als Löcher werden dabei die Stellen bezeichnet, an denen Elektronen im p-Halbleiter fehlen. Der Stromtransport im p-Halbleiter erfolgt dann vereinfacht gesprochen dadurch, dass Elektronen an die Stellen der Löcher springen und so neue Löcher hinterlassen. Auf diese Weise wandern Löcher als freie Ladungsträger im Valenzband durch den p-Halbleiter.

Ein p- und ein n-Halbleiter, die einander berühren, bilden eine pn-Diode, die den Strom nur in eine Richtung leiten kann. In Sperrrichtung bildet sich zwischen p-Halbleiter und n-Halbleiter eine hochohmige Raumladungszone, in der alle Löcher von freien Elektroden besetzt und somit keine frei beweglichen Ladungsträger mehr vorhanden sind. Eine Struktur aus zwei hintereinander geschalteten Dioden, also eine Abfolge pnp oder npn der Dotiergebiete, sperrt den Strom somit in beiden Richtungen, da unabhängig von der Polung der angelegten Spannung immer eine Teildiode sperrt.

Der elektrische Feldeffekt beruht nun darauf, dass durch eine elektrische Spannung an einer Steuerelektrode (Gateelektrode oder kurz Gate) die Verteilung der Elektronen oder Löcher im Halbleiter verändert und damit die Leitfähigkeit des Halbleiters gesteuert wird.

Der bekannteste Typ eines Feldeffekttransistors ist der MOSFET, eine Abkürzung der englischsprachigen Bezeichnung „Metal Oxide Semiconductor Field Effekt Transistor“ (Metall Oxid Halbleiter Feldeffekttransistor). Dieser Name beschreibt zugleich seinen Aufbau. Die Gateelektrode liegt durch eine Oxidschicht (Gateoxid) isoliert auf einer Halbleiterschicht und besteht aus Metall oder üblicherweise aus sehr gut leitendem polykristallinem Silizium (Dotierstoff im Verhältnis 1:1 000). Polykristallines Silizium wird kurz auch Polysilizium genannt. Die Halbleiterschicht weist die oben beschriebene pnp- oder npn-Struktur auf. Die jeweils äußeren Dotiergebiete werden Source und Drain, das Gebiet dazwischen Kanalgebiet genannt.

Beim MOSFET werden durch die Spannung an der isolierten Gateelektrode die in dem Kanalgebiet in Unterzahl vorhandenen Minoritätsladungsträger an der Grenzfläche zum Gateoxid in einem leitfähigen Inversionskanal angesammelt, der das Source- und das Draingebiet leitend miteinander verbindet. Man sagt, durch die Gatespannung wird der MOSFET aufgesteuert. Der MOSFET ist ein unipolares Bauelement, da in ihm je nach Dotierung des Kanalgebiets entweder nur Elektronen oder nur Löcher zum Stromfluss beitragen.

Auch Bipolartransistoren weisen eine pnp- oder npn-Struktur auf. Zur Unterscheidung vom MOSFET werden

die jeweils äußeren Gebiete Emitter und Kollektor sowie die mittlere Schicht Basis genannt. Im Unterschied zum MOSFET ist die Basis sehr dünn und es gibt keine Gateelektrode. Hier funktioniert die sehr dünne Basis-Schicht als Steuerelektrode.

Der Stromfluss zwischen Emitter und Kollektor kann nämlich durch das Potential an der Basis gesteuert werden. Wird das Potential an der Basis so eingestellt, dass die Diode aus Emitter und Basis in Durchlassrichtung gepolt ist, so fließt ein kleiner Steuerstrom, wobei Elektronen und Löcher in unterschiedliche Richtungen fließen. Dadurch verschwindet die Raumladungszone zwischen Basis und Emitter und der Emitter überflutet die Basis auf breiter Front mit einer Sorte von Ladungsträgern, die wiederum zu einem sehr großen Strom zwischen Emitter und Kollektor führen, solange der kleine Steuerstrom fließt. Da sowohl Elektronen als auch Löcher fließen, spricht man hier von einem bipolaren Bauelement.

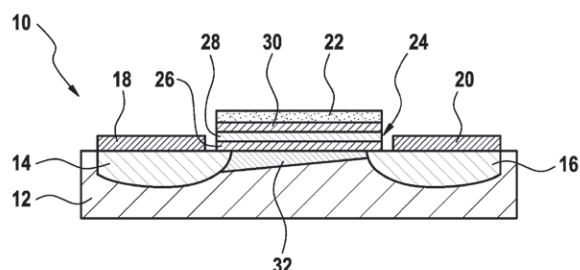
## 2.1 Ionensensitive Feldeffekttransistoren mit isoliertem Gate (G01N 27/414)

Ionensensitive Feldeffekttransistoren sind spezielle MOSFETs, sie werden auch CHEMFETs genannt. Der Aufbau solcher CHEMFETs und ähnlicher Sensoren (resistiver Sensoren) sowie deren Anwendung werden beispielsweise in folgender Literatur beschrieben ([2], [3] und [4]).

In Figur 1 aus der DE 10 2012 211 460 A1 ist der Querschnitt eines CHEMFETs dargestellt. Die Gateelektrode 22 dieses Feldeffekttransistors ist porös, damit die zu detektierenden Ionen durch sie hindurch zum Gateoxid 24 wandern können. Diese Ionen verändern durch ihre Ladung dann die elektrischen Felder im Kanalgebiet 32 und somit messbar den Kanalstrom zwischen Source 14 und Drain 16. Wird der Transistor den Ionen nicht mehr ausgesetzt, diffundieren die restlichen Ionen wieder aus der Gateelektrode. Um einer **Degradation** des CHEMFETs entgegenzuwirken, werden feste Ladungen in einer mittleren Schicht 28 des Gateoxids eingebaut, sodass die Ionen nicht bis in den Halbleiterkörper wandern können.



Dieser CHEMFET wird zum Nachweis von Gasen im Abgasstrom einer Brennkraftmaschine eingesetzt. Daher wird als Halbleitermaterial Siliziumkarbid (SiC) verwendet, das besonders temperaturstabil ist und wegen seiner großen Bandlücke selbst bei hohen Temperaturen nur geringe Leckströme zeigt.



Figur 1: Querschnitt eines sensitiven CHEMFETs für Gase (aus DE 10 2012 211 460 A1)

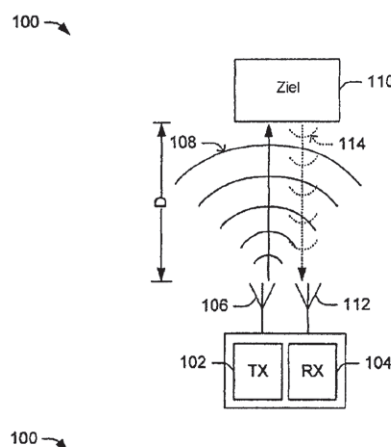
## 2.2 Abstandsradar (G01S 13/00) mit integrierten SiGeBipolartransistoren (H01L 29/737)

Zunehmend werden im Automobil Sensoren eingesetzt, die den Abstand zu den anderen Verkehrsteilnehmern und Hindernissen messen. Eine Elektronik passt auf der Grundlage dieser Messsignale die Geschwindigkeit des Fahrzeugs automatisch an und hilft so dabei, eine Kollision zu verhindern. Solche Messsysteme funktionieren zum Beispiel als Dauerstrichradar, auch FMCW-Radar (englisch: „frequency modulated continuous wave radar“) genannt. So ein Abstandsradar für Automobile ist beispielsweise aus der DE 10 2009 000 816 A1 bekannt und ist schematisch in Figur 2 dargestellt.

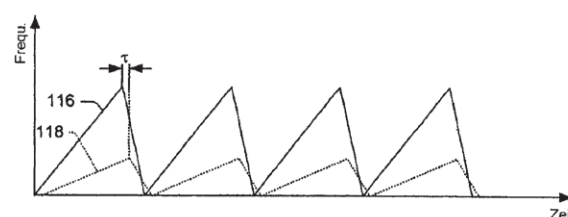
Das von einem Radarsender (englisch: „transceiver“, 102) emittierte Signal 106 wird von einem Objekt 110 als abgeschwächtes Signal 114 reflektiert und von einem Empfänger (englisch: „receiver“, 104) empfangen.

Bei dem gesendeten Signal 116 handelt es sich um eine periodische Frequenzrampe, deren Frequenz in Form eines Sägezahns in Abhängigkeit von der Zeit linear bis zu einer maximalen Frequenz zunimmt und danach wieder abfällt. Das reflektierte Signal 118 weist ebenso eine Sägezahnform auf. Die Zeitdifferenz ( $\tau$ ) zwischen den Maxima von gesendetem und empfangenem Signal

entspricht der Laufzeit des Radarsignals, das sich mit Lichtgeschwindigkeit ausbreitet, für den Hin- und Rückweg des Signals, also für die doppelte Entfernung  $D$  zum Objekt. Darüber hinaus kann aus der Verschiebung der Frequenzen durch den Doppler-Effekt zusätzlich die Relativgeschwindigkeit zwischen Sender und Objekt bestimmt werden. Im vorliegenden Fall ist das empfangene Signal zu tieferen Frequenzen verschoben.



100



Figur 2: Messprinzip eines Dauerstrich frequenzmodulierten Radarsystems (FMCW) im Automobil (aus DE 10 2009 000 816 A1)

Für das Abstandsradar im Automobil haben sich 77 GHz etabliert. Bei diesen hohen Frequenzen müssen Halbleitermaterialien verwendet werden, die eine entsprechend hohe Ladungsträgerbeweglichkeit aufweisen. Ein großer Kostenvorteil bei der Herstellung von Sensoren für ein Abstandsradar wurde durch die Entwicklung einer Fertigungstechnologie möglich, die herkömmliche preisgünstige Siliziumsubstrate für CMOS-Logik (Logik realisiert mit komplementär dotierten Feldeffekttransistoren) mit Silizium-Germanium (SiGe), das in Form von dünnen Schichten nur lokal aufgebracht wird, kombiniert. So kann mit nur einem Chip sowohl die Erzeugung beziehungsweise das Empfangen der Radarsignale als auch die Logik zur Auswertung der Signale bereitgestellt werden.

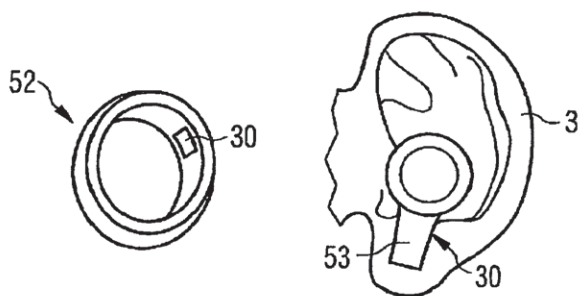


### 3 Innerer Fotoeffekt (H01L 31/00)

Der innere Fotoeffekt tritt in Halbleitern auf. Sobald der Halbleiter mit Licht einer Energie bestrahlt wird, die größer als seine Bandlücke ist, werden Elektronen vom Valenzband in das Leitungsband angeregt. Das so entstandene Leitungselektron und das Loch im Valenzband werden im elektrischen Feld eines pn-Übergangs oder einer äußeren Spannung getrennt. Dadurch entsteht ein messbarer Fotostrom.

#### 3.1 Pulsoximeter (A61B 5/1455) mit Fotodioden (H01L 31/105)

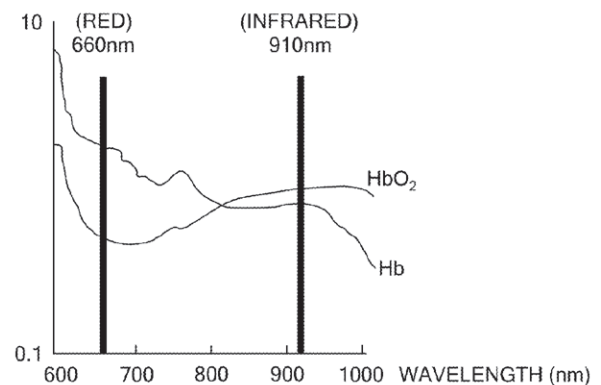
Aktuell werden Fitnessarmbänder von einer Vielzahl von Herstellern angeboten, mit denen der Puls und die Sauerstoffsättigung im Blut ununterbrochen gemessen werden können. Die dazu nötigen Messgeräte werden Pulsoximeter [5] genannt und werden heute in einer solch kleinen Bauform hergestellt, dass sie in Armbanduhren, Hörgeräte oder sogar Ringe eingebaut werden können (vergleiche Figur 5). Auch hierzu hat die Miniaturisierung in der Mikroelektronik wesentlich beigetragen [6].



Figur 5: Fingerring (links) und ein Hörgerät (rechts) mit darin untergebrachten Pulsoximetermodulen 30 (aus DE 10 2012 017 919 A1)

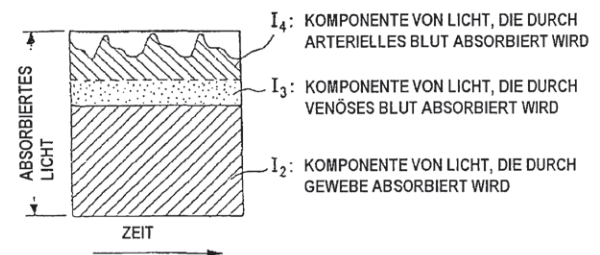
Grundlage ist der Effekt, dass sich das Licht-Absorptionsvermögen des Hämoglobins im Blut deutlich ändert, wenn es in der Lunge Sauerstoff aufgenommen hat. Das sauerstoffreiche Blut wird dann vom Herzen als arterielles Blut in den ganzen Körper gepumpt. In Figur 6 aus der US 2015/0057511 A1 ist gut zu erkennen, dass zum Beispiel rotes Licht (650 nm) von Hämoglobin, das Sauerstoff aufgenommen hat ( $\text{HbO}_2$ ), deutlich

schwächer und infrarotes Licht (900 nm) deutlich stärker absorbiert wird als von Hämoglobin ohne gebundenem Sauerstoff (Hb). Das Verhältnis des gemessenen Absorptionsvermögens des Blutes bei diesen beiden Wellenlängen ändert sich daher deutlich, wenn sich der Anteil des Hämoglobins, das Sauerstoff aufgenommen hat, ändert und ist daher ein verlässliches Maß für die Sauerstoffsättigung im Blut.



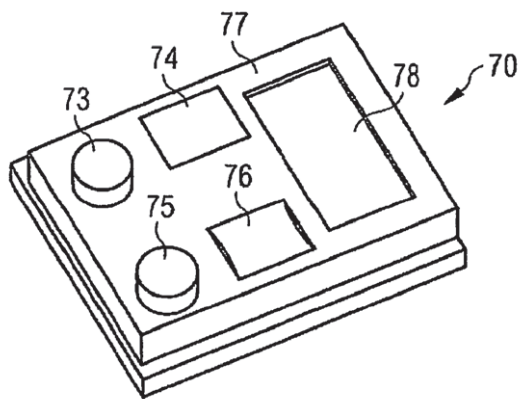
Figur 6: Absorptionsvermögen von Hämoglobin mit ( $\text{HbO}_2$ ) und ohne Sauerstoff (Hb) in Abhängigkeit der Wellenlänge (aus US 2015/0057511 A1 beziehungsweise WO 2013/148753 A1)

Aus Figur 7 wird deutlich, dass mit diesem Messverfahren auch die Pulsfrequenz ermittelt werden kann. In der schematischen Darstellung ist das Absorptionsvermögen von durchblutetem Gewebe in Abhängigkeit von der Zeit aufgetragen. Lediglich das Signal des arteriellen Blutes  $I_4$  variiert mit der Periode des Pulsschlages, denn mit jedem Pulsschlag weiten sich die Arterien bis in die feinsten Kapillaren der Haut und es erhöht sich mit dem relativen Anteil des arteriellen Blutes das Absorptionsvermögen.



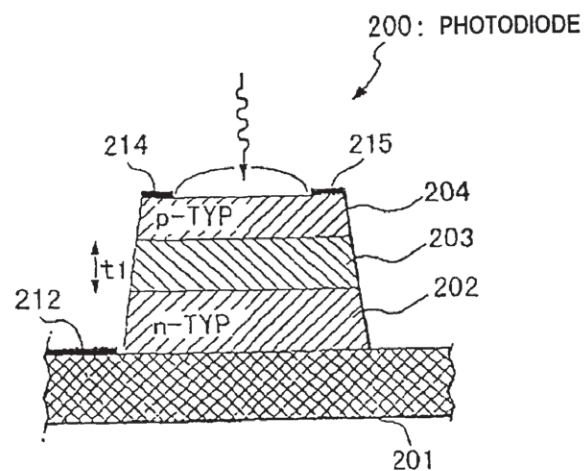
Figur 7: Zeitliche Änderung des Absorptionsvermögens von menschlicher Haut in Abhängigkeit von der Zeit für die Dauer von drei Pulsschlägen (aus DE 698 34 616 T2)

Ein hoch integrierter Chip mit der Funktion eines Puls-oximeters ist in Figur 8 aus der DE 10 2012 017 919 A1 dargestellt. In dem Chip sind zur Messung der Absorption des roten und des infraroten Lichts getrennte Leuchtdioden 73 und 75 und getrennte Fotodioden 74 und 76 und ein Bluetooth-Sender 78 zum Übermitteln des Signals an ein Anzeigegerät gezeigt. In dem Chip sind außerdem Auswerteschaltungen integriert (nicht gezeigt). In dieser Konfiguration wird das Licht von den Leuchtdioden emittiert, dringt in die Haut ein, wird durch Absorption geschwächt und diffus zu den Fotodioden zurückgestreut.



Figur 8: Integrierter Halbleiterchip eines Pulsoximeters (aus DE 10 2012 017 919 A1)

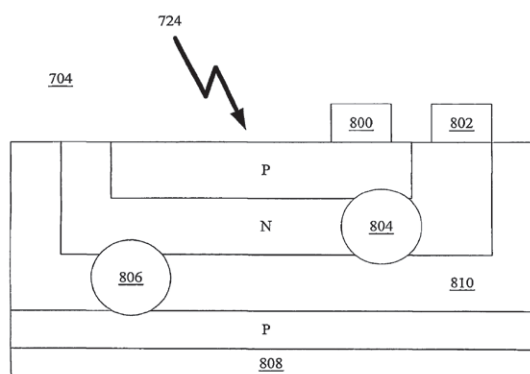
Im einfachsten Falle ist die Fotodiode eine PIN-Diode, deren Querschnitt und Beschaltung in Figur 9 dargestellt ist. Die Fotodiode wird in Sperrrichtung betrieben, das heißt der Minuspol der Spannungsquelle (kurzer dicker Strich der Spannungsquelle E) wird an die Anode beziehungsweise den p-Typ-Halbleiter (ausgefülltes Dreieck) und der Pluspol der Spannungsquelle (langer Strich der Spannungsquelle E) an die Kathode beziehungsweise den n-Typ-Halbleiter angeschlossen, sodass der nicht dotierte (intrinsische) Bereich 203 zwischen Anode und Kathode 204 und 202 vollständig an freien Ladungsträgern verarmt und ganz von einer Raumladungszone eingenommen wird. Sobald Licht auf die Fotodiode fällt, werden Elektronen-Lochpaare in der Raumladungszone erzeugt und im elektrischen Feld getrennt, sodass ein Fotostrom fließt. Die Fotodiode wird niederohmig und die Spannung ( $V_{out}$ ) über der Fotodiode sinkt dadurch stark ab. Die Veränderung der Spannung kann als Messsignal genutzt werden.



Figur 9: Querschnitt einer Fotodiode und schematische Beschaltung (aus DE 698 34 616 T2)

Nach der DE 10 2012 017 919 A1 kann ein Pulsoximeter, um Fläche zu sparen, auch nur einen Sensor aufweisen. Die Leuchtdioden werden dazu gepulst betrieben, so dass die Fotodiode zeitlich nacheinander die Signale der Leuchtdioden empfängt. Bei bekannter spektraler Empfindlichkeit der Fotodiode können die Signale in einem integrierten Schaltkreis entsprechend korrigiert werden. In der WO 2013/148753 A1 werden die Leuchtdioden zum Beispiel bei 100 Hz betrieben, damit der zeitliche Verlauf des Herzschlags noch gut aufgelöst werden kann. Im gepulsten Betrieb ist zudem das Umgebungslicht als Störgröße leicht vom Messsignal zu trennen.

Aus der europäischen Patentschrift DE 601 29 332 T2 ist eine Fotodiode für ein Pulsoximeter bekannt, die zeitgleich zwei unterschiedliche Wellenlängenbereiche messen kann (vergleiche Figur 10). Dazu werden durch unterschiedlich tiefe Wannengebiete in einem gemeinsamen Siliziumsubstrat zwei Fotodioden 804 und 806 übereinander gebildet. Diese haben eine gemeinsame Kathode 802. Die obere Fotodiode wirkt dabei als Absorptionsfilter für die untere. Die obere Fotodiode misst die kleinen Wellenlängen (hohe Energie), die untere die großen Wellenlängen.



Figur 10: Querschnitt eines Halbleiterschips mit zwei übereinander angeordneten Fotodioden 804 und 806 (aus DE 601 29 332 T2)

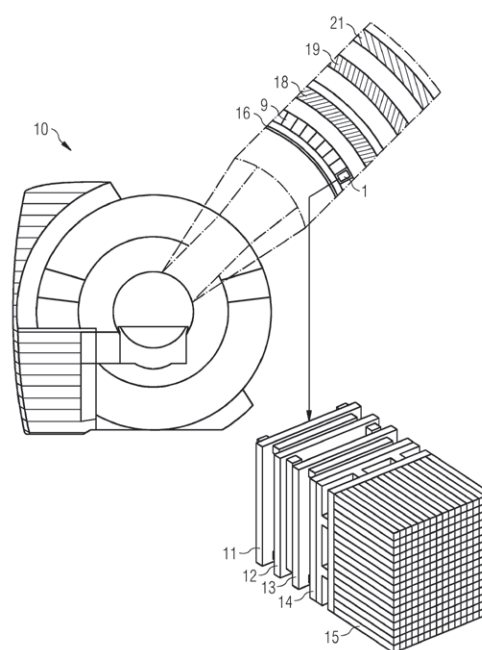
### 3.2 Nuklearmedizin, PET-MRT (G01T 1/202, G01R 33/34) mit Avalanche-Fotodioden (H01L 31/115)

In der Nuklearmedizin ist mit dem Einsatz von Avalanche-Fotodioden eine kompakte Kombination von nuklearmedizinischer Diagnostik mit Magnetresonanztomografen (MRT) möglich geworden. Denn Avalanche-Fotodioden sind, anders als herkömmliche Elektronenvervielfacherröhren, in den starken Magnetfeldern des MRT verwendbar. Kommerziell erhältlich sind bereits kombinierte PET-MRT-Anlagen. PET steht hierbei für Positronen-Emissions-Tomografie und ist ein bildgebendes Verfahren der Nuklearmedizin. Dabei wird dem Patienten in geringsten Dosen ein Radionuklid verabreicht, das in Substanzen, die im Körper am Stoffwechsel teilnehmen, chemisch gebunden wird. Diese Substanzen werden Tracer genannt. Die Tracer können so spezifisch hergestellt werden, dass die radioaktiven Atome sich beispielsweise nach dem Spritzen in die Blutbahn ausschließlich in der Schilddrüse anreichern, weil nirgendwo sonst im Körper dieser Tracer verarbeitet wird. Auf diese Weise kann spezifisch die Funktion von Organen, zum Beispiel der Schilddrüse oder des schlagenden Herzens, untersucht werden.

Als radioaktive Stoffe finden  $\beta^+$ -Strahler Anwendung, die unter Aussendung eines Positrons zerfallen. Die Positronen zerstrahlen auf kurzer Distanz in Verbindung mit einem Elektron des umgebenden Gewebes zu zwei Gamma-Quanten, die näherungsweise in einem  $180^\circ$  Winkel zueinander den Körper verlassen und in einem Detektorring detektiert werden können.

In der DE 10 2012 212 574 A1 ist ein PET-MRT gezeigt. Die Figur 11 zeigt den Blick vom Fußende der Patientenliege in die Röhre des MRT. Innerhalb der Gradientenspule 18 und der Spule 19 für das primäre Magnetfeld des MRT liegt für den Nachweis der Gamma-Quanten der Detektorring 9 des PET, der aus vielen Detektormodulen 1 aufgebaut ist.

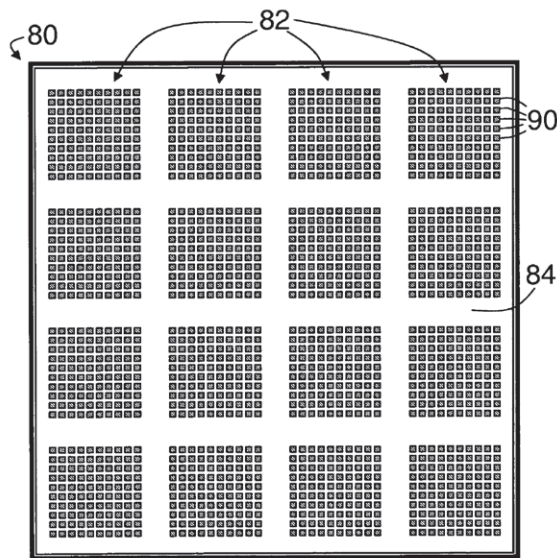
Die Abfolge der Bestandteile des Detektormoduls sind, von der Patientenliege aus gesehen, eine Matrix aus LSO-Kristallen 15 (Lutetiumoxyorthosilicat), in denen die einzelnen Gamma-Quanten in Lichtblitze umgewandelt werden, 3x3 Avalanche-Fotodioden 14, eine Hochvolt-Leiterplatte 13, eine Vorverstärker-Leiterplatte 12 und eine Treiber-Leiterplatte 11. Ein Gamma-Quant erzeugt in der Matrix aus LSO-Kristallen (Szintillator) einen Lichtblitz, der von mehreren der Avalanche-Fotodioden detektiert wird. Die Leiterplatten weisen zudem Schaltungen auf, die durch Vergleich der Signalhöhen der benachbarten Avalanche-Fotodioden feststellen, wo der Schwerpunkt des Lichtblitzes lag. Auf diese Weise ist eine hohe Ortsauflösung des Detektormoduls zu erzielen, die deutlich unter der Ausdehnung einer der 3x3 Avalanche-Fotodioden liegt (Prinzip der **Anger-Kamera**).



Figur 11: PET-MRT mit zwei Ausschnittsvergrößerungen  
Rechts oben: Abfolge von den Magnetspulen des MRT und dem PET-Detektorring / Rechts unten: PET-Detektormodul mit 3x3 Avalanche-Fotodioden (aus DE 10 2012 212 574 A1)

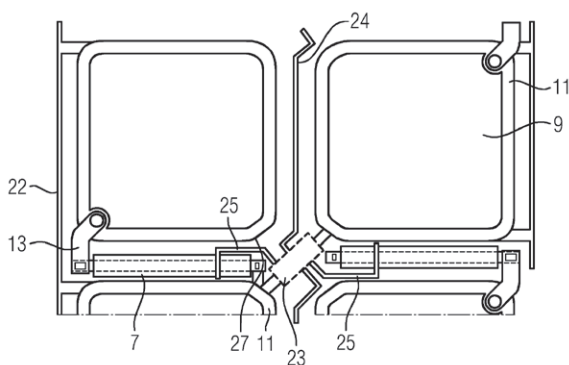


In der europäischen Patentschrift EP 1 875 273 B1 (siehe Figur 12) werden in einem Detektormodul 4x4 Avalanche-Fotodioden 82 in einem Siliziumchip 80 integriert. Dabei ist jede der Avalanche-Fotodioden aus einer Matrix von lichtempfindlichen Pixeln 90 aufgebaut.



Figur 12: Avalanche-Fotodioden-Chip eines PET-Detektormoduls. (aus EP 1 875 273 B1)

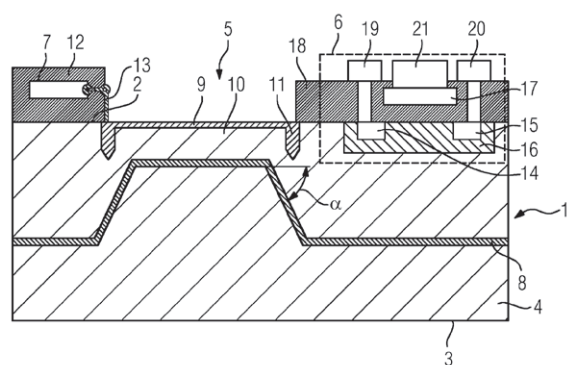
Aus der DE 10 2012 215 637 A1 ist der Aufbau der Pixel von Avalanche-Fotodioden bekannt und wird in Figur 13 gezeigt. Avalanche-Fotodioden werden wie Fotodioden in Sperrrichtung betrieben. Im Unterschied zu Fotodioden werden die Avalanche-Fotodioden allerdings nahe der Durchbruchspannung oder knapp darüber betrieben (sogenannter Geiger-Modus). Durch Licht erzeugte Elektronen-Lochpaare werden im starken elektrischen Feld getrennt und durch Stoßionisation vervielfacht (Avalanche-Effekt).



Die Sperrspannung wird zwischen dem Substratkontakt 8 und der Spannungszuführung 22 angelegt. Sie wird dann über einen Löschwiderstand 7 an die p+ dotierte Anode 9 geführt. Der Löschwiderstand und das zugehörige Pixel sind also in Serie geschaltet (ähnlich wie der Vorwiderstand vor einer Fotodiode, vergleiche Figur 9). Sobald das Pixel der Avalanche-Fotodiode durch die p+ dotierte Anode beleuchtet wird, wird dieses schlagartig leitfähig und über dem Löschwiderstand fällt fast die gesamte Sperrspannung ab. Der Mittelabgriff 25 auf dem Löschwiderstand schnell dann auf ein Zwischenpotential. Da dieser über eine Metallbahn 25 mit dem Gate eines zugeordneten MOSFETs verbunden ist, steuert dieser MOSFET auf und legt ein Triggersignal (Startsignal) an die Triggerleitung 24 der Avalanche-Fotodiode. Dieses Triggersignal startet in einer externen Auswerteschaltung eine Integration der Signale aller Pixel der einen von 4x4 Avalanche-Fotodioden des PET-Detektormoduls. Durch das Triggersignal wird zudem ermittelt, wann genau das Gamma-Quant im Modul eintraf, und mit dem Auswerte-Integral wird festgestellt, welcher Anteil des Lichtblitzes insgesamt auf die eine Avalanche-Fotodiode im Vergleich zu den anderen Avalanche-Fotodioden des PET-Detektormoduls gefallen ist.

Da der Löschwiderstand und die Pixeldiode in Serie geschaltet sind, fällt über dem beleuchteten Pixel der Avalanche-Fotodiode keine Spannung mehr ab und der Avalanche löscht sich selber. Nach einer Totzeit ist das Pixel daher wieder empfindlich.

Im rechten Teil der Figur 13 wird gezeigt, dass das vergrabene n+ dotierte Rückkontaktgebiet 8 unter



Figur 13: Pixel einer Avalanche-Fotodiode (aus DE 10 2012 215 637 A1)

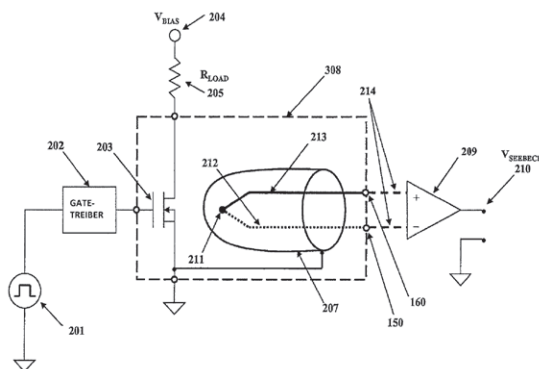


dem MOSFET 6 tiefer verläuft. Dadurch werden dort hohe Spannungen vermieden. Zudem ist das vorderseitige Anodengebiet 9 an den Rändern 11 nach unten gezogen, damit der Avalanche-Strom des Pixels nicht zu dem MOSFET gelangt.

Dadurch, dass das erste ansprechende Pixel einer der 4x4 Avalanche-Fotodioden das Triggersignal des gesamten Detektormoduls liefert, ist eine hohe zeitliche Auflösung zu erzielen, sodass die „**Time of flight (TOF)-Methode**“ angewendet werden kann, um die Ortsauflösung des PET deutlich zu verbessern. Ohne diese Technik könnte lediglich festgestellt werden, dass die Quelle der Strahlung auf der Verbindungsline der im Detektorring gegenüberliegenden und in Koinzidenz ansprechenden Detektormodule lag [7].

#### 4 Seebeck-Effekt (H01L 35/00)

Wenn die Kontaktstelle zwischen zwei Leitern aus unterschiedlichen Materialien eine andere Temperatur hat als die beiden offenen Enden der Leiter, so ist zwischen den offenen Enden eine Spannung messbar, die sogenannte Seebeck-Spannung. Die unterschiedlichen Materialien können unterschiedliche Metalle oder n- und p-leitende Halbleiter sein. An der Kontaktstelle kann auch ein weiteres Material eine ohmsche Verbindung schaffen. Die Höhe der Seebeck-Spannung hängt dabei von dem Temperaturunterschied sowie Unterschieden in den Materialeigenschaften der Leiter ab (Elektronenbeweglichkeit und Elektronenzustandsdichte). Diese Anordnung wird auch Thermoelement genannt.



Figur 14: Schaltungsschema eines in einem Leistungstransistor integrierten Thermoelements (aus DE 10 2010 028 275 A1)

#### 4.1 Leistungstransistor mit integriertem Seebeck-Effekt-Bauelement

In den Schaltungen für die Stromversorgung der starken Elektromotoren von Elektrofahrzeugen werden Halbleiter-Leistungstransistoren eingesetzt, die bei hohen Spannungen und hohen Strömen schnell schalten. Um diese Leistungstransistoren im Betrieb vor Überhitzung zu schützen, zeigt die DE 10 2010 028 275 A1, wie Seebeck-Effekt-Bauelemente in diese Leistungstransistoren integriert werden können.

In Figur 14 ist eine solche Anordnung dargestellt. Die Kontaktstelle 211 zwischen zwei Leitern 212 und 213 aus unterschiedlichen Materialien liegt dicht an einem Leistungstransistor 203, der sich im Betrieb schlagartig erwärmen kann. Zwischen den beiden kalten Enden der elektrischen Leiter 150 und 160 ist mit einem Verstärker 209 die Seebeck-Spannung messbar. Eine Abschirmung 207 gegen elektrische Störfelder, die beispielsweise von den **Kommutatoren** der Elektromotoren verursacht werden, ist erforderlich, da der Betrag der Seebeck-Spannung lediglich im Bereich von einigen Millivolt liegt.

Figur 15 zeigt eine Serienschaltung mehrerer Thermoelemente, in der sich die Seebeck-Spannungen der einzelnen Thermoelemente 1802 und 1803 zu einem größeren Ausgangssignal ( $V_{\text{Seebeck}}$ ) addieren. Die Leiter zwischen der heißen Seite A und der kalten Seite C können wegen der stromlosen Spannungsmessung eine große Länge aufweisen und werden aus unterschiedlich dotiertem Polysilizium gebildet. Besonders kostengünstig wird die Integration des Seebeck-Bauelements in den Leistungstransistoren (nicht dargestellt), weil dieses im Wesentlichen mit denselben Prozessschritten wie der Leistungstransistor hergestellt und im selben Siliziumchip integriert werden kann. Im unteren Teil der Figur 15 ist im Querschnitt die elektrische Abschirmung 701 gezeigt, die aus den im Prozess zur Verfügung stehenden Metallisierungen 701 und einem Dotiergebiet 704 im Substrat gebildet wird.

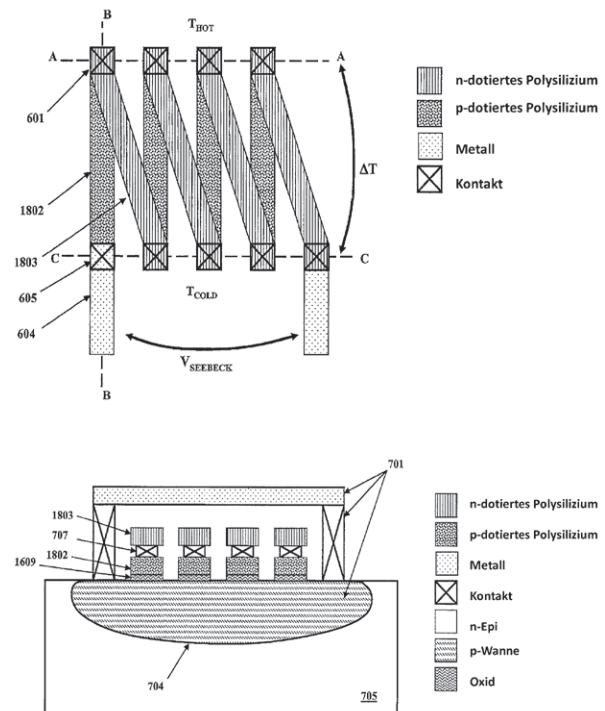
In Figur 16 ist eine weitere Möglichkeit gezeigt, wie ein Seebeck-Bauelement in den Chip eines Graben-Leistungsfeldeffekttransistors integriert werden kann. Die isolierte Gateelektrode 2406, durch welche ein leitender Inversionskanal im Kanalgebiet 2409 beeinflusst werden kann, besteht aus Polysilizium. Im unmittelbar benachbarten Graben 2401 ist das Seebeck-Bauelement aus zwei verschiedenen dotierten Polysiliziumschichten 2403 und 2404 gebildet. Das Silizium-Substrat sorgt in diesem Fall für die elektrische Abschirmung um das Seebeck-Bauelement. Die beiden Schichten des Seebeck-Bauelements haben nur am hinteren Ende des Grabens elektrischen Kontakt 2501. Wird beim Auftreten einer Temperaturänderung eine vorbestimmte Temperaturdifferenz zwischen dieser Kontaktstelle und den außerhalb des Leistungstransistors liegenden, in der Figur 16 nicht gezeigten, kalten Enden überschritten, kann der Transistor zum Schutz vor Überlastung abgeschaltet werden.

## 5 Piezoelektrischer Effekt (Piezoeffekt) (H01L 41/00)

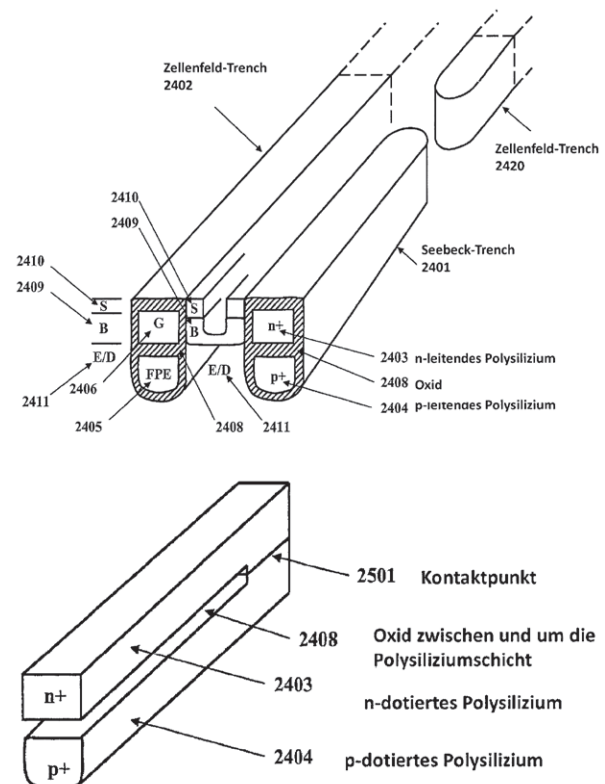
Der piezoelektrische Effekt tritt an Kristallen aus unterschiedlichen Atomsorten auf, deren Elementarzelle nicht punktsymmetrisch ist. Wird der insgesamt neutrale Kristall von gegenüberliegenden Seiten zusammengedrückt, treten gegensätzliche Oberflächenladungen an diesen gegenüberliegenden Flächen auf. Dieser Elektronendichteunterschied kann als Spannung gemessen werden, man spricht von Piezoelektrizität.

### 5.1 Ultraschallsensor mit piezoelektrischer Dünnschicht

Piezoelektrische Dünnschichten aus Zinkoxid (ZnO) können mit sehr kostensparenden Abscheideverfahren, wie zum Beispiel Sputtern, hergestellt werden. Die ZnO-Schicht wächst dabei zwar polykristallin auf, jedoch mit einer Textur. Die Kristallachsen der Kristallite dieser polykristallinen Schicht weisen damit eine Vorzugsrichtung auf und sind deshalb piezoelektrisch, vergleiche hierzu auch die Patentschrift DE 36 03 337 C2.



Figur 15: In Serie verschaltete Thermolemente aus n- und p-dotiertem Polysilizium auf einem Substrat aus Silizium (aus DE 10 2010 028 275 A1)

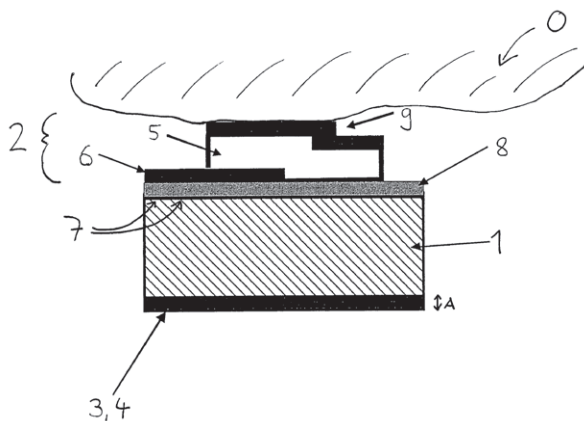


Figur 16: Thermolement aus n- und p-dotiertem Polysilizium, das in ein Zellenfeld eines Grabentransistors integriert ist (aus DE 10 2010 028 275 A1)

Piezoelektrische Keramiken müssen im Unterschied dazu nach einem **Sinterschritt** durch das Anlegen einer hohen Spannung erst polarisiert werden, vergleiche DE 10 2012 111 972 A1.

Aus der europäischen Schrift EP 2 348 503 B1 ist ein Ultraschallsensor aus solch einer piezoelektrischen Dünnschicht aus ZnO oder Aluminiumnitrid (AlN) auf einem Siliziumsubstrat bekannt. In Figur 17 ist auf dem Siliziumsubstrat 1 zwischen zwei Metall-Elektroden 6 und 9 die piezoelektrische Dünnschicht 5 zu sehen. Zwischen der unteren Elektrode und dem Substrat ist zudem eine elektrische Isolierschicht 8 angeordnet. Die piezoelektrische Schicht kann sowohl als Ultraschallsender als auch Ultraschallempfänger fungieren. Die aus dem Körper O zurückgeworfenen Schallreflexe verformen die Dünnschicht und können als piezoelektrische Spannungen gemessen werden. Allerdings werden die Schallwellen auch an der Rückseite des Substrats reflektiert und führen so zu Störsignalen.

Daher wird auf der Rückseite 3 des Siliziumsubstrats eine schallschluckende Schicht 4 aus sogenanntem „black silicon“ durch Ätzung erzeugt. Die Struktur des „black silicon“ besteht aus dicht an dicht stehenden Nadeln aus Silizium, welche stehenbleiben, wenn das Silizium durch ein reaktives Ionenätzen angeätzt wird.



Figur 17: Ultraschallsensor mit piezoelektrischer Dünnschicht (aus EP 2 348 503 B1)

## 6 Hall-Effekt und GMR-Effekt (H01L 43/00)

Der Hall-Effekt tritt bei stromdurchflossenen flächigen Leitern auf, die von einem Magnetfeld durchsetzt werden. Die längs des Stromflusses bewegten Ladungsträger werden in dem Magnetfeld aufgrund der Lorentzkraft abgelenkt, was zu einem Ladungsträgerdichteunterschied, der Hall-Spannung, senkrecht zum Stromfluss führt. Die technische Stromrichtung (von plus nach minus), die Richtung des Magnetfeldes und die Lorentzkraft stehen dabei in einem rechtshändigen Koordinatensystem senkrecht aufeinander.

Ein weiterer magnetischer Effekt ist der Riesenmagnetowiderstand-Effekt (GMR-Effekt). Für die Entdeckung des GMR-Effekts wurde 2007 der Nobelpreis an Albert Fert (Paris) und Peter Grünberg (Jülich) verliehen [8]. Die Entdeckung war, dass der elektrische Widerstand einer Schichtfolge aus wenigstens einer ersten ferromagnetischen, einer nichtmagnetischen Koppelschicht und einer zweiten ferromagnetischen Schicht sich deutlich erhöht, wenn die beiden ferromagnetischen Schichten in entgegengesetzter Richtung magnetisiert sind.

Die Widerstandsänderung tritt dabei auf, wenn der Strom parallel durch die Schichtfolge fließt, wie in der Patentschrift DE 600 23 835 T2, der europäischen Patentschrift EP 1 830 350 B1 oder der amerikanischen Patentschrift US 7 212 384 B1 beschrieben, oder wenn der Strom senkrecht durch die Schichtenfolge fließt, eine aktuelle Entwicklung wie in der amerikanischen Patentschrift US 8 675 309 B2 beschrieben.

### 6.1 3D-Magnetfeldsensor mit Hall-Effekt-Bauelement (H01L 43/06)

Um die Verteilung magnetischer Felder zu messen, werden Hall-Sensoren eingesetzt. Eine Anwendung ist zum Beispiel die genaue Positionsbestimmung von Maschinenteilen, an denen kleine Permanentmagnete befestigt werden.

Vergleiche die Patentschrift DE 10 2008 039 569 B4 zur Positionsbestimmung einer Waschmaschinen-trommel, die beim Schleudern ausgewuchtet werden kann [9]. In Antiblockiersystemen (ABS) im Auto werden als aktive Raddrehzahlfühler unter anderem auch Hall-Sensoren eingesetzt.

Bei solchen Anwendungen ist es besonders zweckmäßig einen Hall-Sensor bereitzustellen, der die Komponenten des magnetischen Feldes in allen drei Raumrichtungen (3D) am selben Punkt im Sensor misst, also einen echten 3D-Sensor.

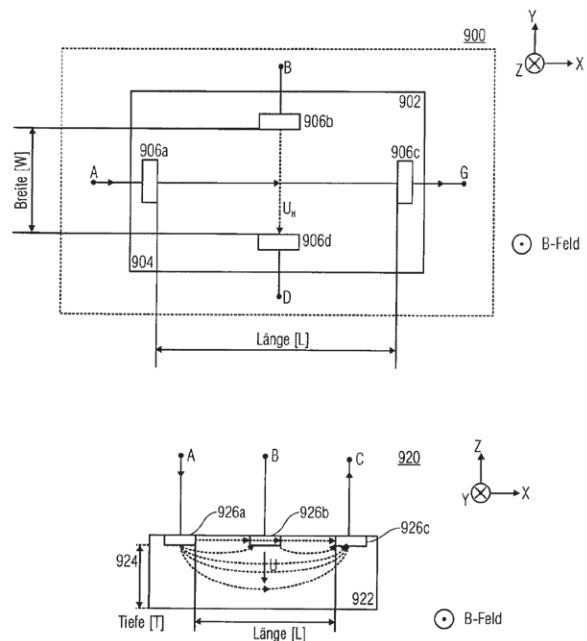
Aus der Patentschrift DE 10 2006 037 226 B4 ist bekannt, zum Bereitstellen eines 3D-Hall-Sensors horizontale und vertikale Hall-Effekt-Bauelemente in demselben Siliziumchip zu integrieren.

In Figur 18 sind auf der Oberfläche des Siliziumsubstrats links und rechts die Elektroden 906a und 906c des horizontalen Hall-Sensors gezeigt, zwischen denen der Messstrom fließt. Die zwischen diesen Elektroden fließenden Ladungsträger werden mittels der z-Komponente des Magnetfeldes (B-Feld) durch die Lorentzkraft senkrecht zum Stromfluss in der Ebene der Substratoberfläche abgelenkt. Dies führt zu einer messbaren Hall-Spannung ( $U_H$ ) zwischen den Messelektroden 906b und 906d, die ebenfalls auf der Substratoberfläche angeordnet sind.

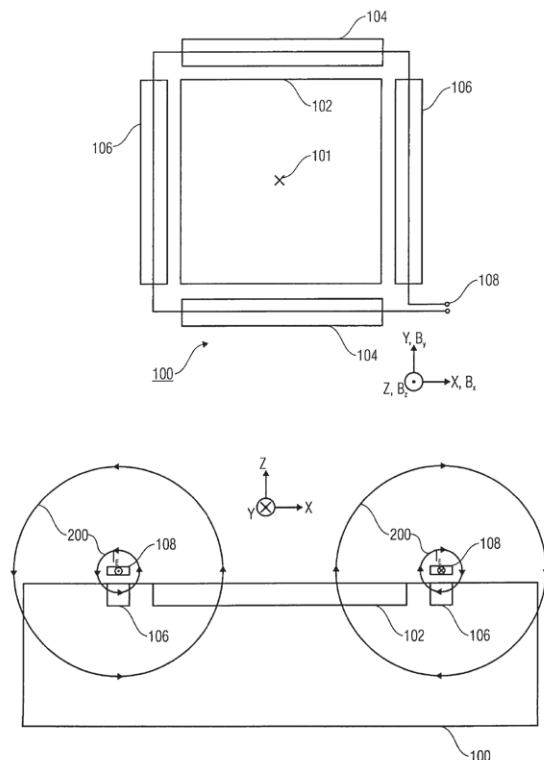
Im unteren Teil der Figur 18 werden die zwischen den Elektroden des vertikalen Hall-Sensors fließenden Ladungsträger mittels der y-Komponente des Magnetfeldes durch die Lorentzkraft je nach Stromrichtung in die Tiefe des Substrats oder an die Oberfläche abgelenkt. Dies führt zu einer messbaren Hall-Spannung (U) zwischen der Messelektrode 926b und einer nicht dargestellten Substrat-Elektrode.

Im oberen Teil der Figur 19 ist die Draufsicht auf das Siliziumsubstrat gezeigt, in welchem zwei Paare vertikaler Hall-Elemente 104 und 106 und ein horizontales Hall-Element 102 symmetrisch zu einem gemeinsamen Messzentrum 101 angeordnet sind. Das eine Paar von vertikalen Hall-Elementen 104 misst dabei die y-Komponente und das andere Paar 106 die x-Komponente des Magnetfeldes.

Das horizontale Hall-Element in der Mitte erfasst die senkrecht auf der Substratoberfläche (Zeichenebene) stehende z-Komponente des Magnetfeldes.



Figur 18: Horizontaler (oben) und vertikaler (unten) Hall-Sensor (aus DE 10 2006 037 226 B4)



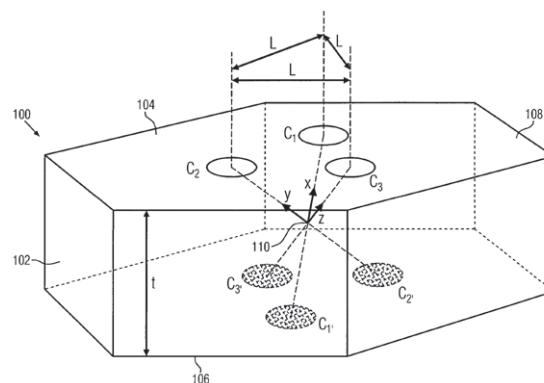
Figur 19: Integrierter 3D-Hall-Sensor (aus DE 10 2006 037 226 B4)

Eine Kalibrierleitung 108 verläuft oberhalb der vertikalen Hall-Elemente in einer geschlossenen Schleife. Durch diese kann für kurze Zeit ein Strom geschickt werden, um ein definiertes Magnetfeld im Siliziumsubstrat zu erzeugen. Unmittelbar unter der Leitung wird eine Magnetfeldkomponente in x- und y-Richtung parallel zur Substratoberfläche und am Ort des Messzentrums 101 eine Magnetfeldkomponente in z-Richtung senkrecht zur Substratoberfläche erzeugt. So kann von Zeit zu Zeit ein Selbsttest des 3D-Hall-Sensors durchgeführt werden.

Systematische Fehler, die zum Beispiel dadurch entstehen, dass bei den vertikalen Hall-Elementen die Messelektroden wegen Fertigungstoleranzen nicht in der Mitte der stromführenden Elektroden sitzen, können bei der Auswertung korrigiert werden. Dazu wird die Stromrichtung in den Paaren der vertikalen Hall-Elemente umgekehrt und der Betrag der Hall-Spannungen verglichen. Solche Verfahren beim Beschalten und Auslesen von Hall-Sensoren sind unter dem Namen „Spinning-Current-Verfahren“ bekannt.

Allerdings bleibt das Problem bestehen, dass die vertikalen Hall-Elemente in der Regel eine geringere Empfindlichkeit als die horizontalen aufweisen. Dies muss durch die Auswerteschaltung bei der Bestimmung des Messfeldes berücksichtigt werden. Ein 3D-Hall-Sensor mit isotroper Empfindlichkeit ist dagegen aus der DE 10 2013 209 514 A1 bekannt.

In Figur 20 sind auf der Vorder- und Rückseite eines Siliziumsubstrats jeweils drei Kontakte in Form von zwei gleichseitigen Dreiecken angeordnet. Die Dreiecke sind in Draufsicht um  $60^\circ$  gegeneinander verdreht. Die Seitenlänge  $L$  der Dreiecke wird auf die Dicke  $t$  des Substrats so abgestimmt, dass die Verbindungslinien (gestrichelte Linien) von schräg gegenüberliegenden Kontakten der Vorder- und Rückseite genau senkrecht aufeinander stehen, sodass sie ein rechtwinkeliges Koordinatensystem aufspannen. So können durch paarweises Nutzen der schräg gegenüberliegenden Kontakte der Vorder- und Rückseite zum Einprägen des Messstroms und zum Messen der Hall-Spannung nacheinander die Magnetfeldkomponenten entlang der drei Achsen dieses Koordinatensystems gemessen werden.



Figur 20: 3D-Hall-Sensor mit Kontakten auf gegenüberliegenden Seiten des Chips (aus DE 10 2013 209 514 A1)

## 6.2 Festplatten-Lesekopf (G11B 5/39) mit GMR-/STO-Sensor (H01L 43/08)

Die Festplatten, wie wir sie heute kennen, wurden erst möglich durch die Entdeckung des GMR-Effekts. Ein Überblick zu der Entwicklung der Festplattentechnik ist in [10] vorgestellt. Eine Weiterentwicklung eines GMR-Lesekopfes, bei dem der Strom senkrecht durch die Schichten des Sensors fließt, wird in der amerikanischen Patentschrift US 8 675 309 B2 vorgeschlagen. Dieser Sensor wird STO-Sensor genannt, wobei STO die Abkürzung für „spin-torque oscillator“ ist.

Die Figur 21 aus der amerikanischen Patentschrift US 8 675 309 B2 zeigt einen solchen STO-Sensor 200 in schematischer Darstellung, aus der sein Schichtaufbau deutlich wird. Der Sensor schwebt auf einem Luftpolster über einer sich drehenden Festplatte 252, die sich unter dem Sensor (in der Zeichenebene von rechts nach links) hinwegbewegt. Auf diese Weise werden die senkrecht zur Festplatten-Oberfläche magnetisierten (englisch: „perpendicular magnetic recording“, PMR) gespeicherten Informationseinheiten (Bits) an dem Schichtstapel des Sensors vorbeigeführt.

Wie ein GMR-Sensor besteht auch der STO-Sensor aus der Schichtfolge einer ersten ferromagnetischen Schicht 210, einer nicht magnetischen leitfähigen Koppelschicht 230, zum Beispiel aus Kupfer, und einer zweiten ferromagnetischen Schicht 220. Die zweite ferromagnetische Schicht besteht aus einem magnetisch härteren Material als die erste ferromagnetische Schicht und weist eine fest eingeprägte Magnetisierungs-



richtung auf, die in Schichtebene und zugleich senkrecht zur Festplattenoberfläche verläuft. Dagegen ist die Magnetisierungsrichtung der ersten ferromagnetischen Schicht in der Schichtebene „frei“ und wird leicht durch ein äußeres Magnetfeld in ihrer Richtung umgekehrt.

Wollte man diesen Sensor als GMR-Sensor betreiben, dürfte nur ein geringer Strom senkrecht durch den Schichtstapel fließen. Die erste ferromagnetische Schicht würde dann, immer wenn eine Grenzfläche zweier unterschiedlich magnetisierter Domänen (also zwischen Bits der Werte 0 und 1 oder 1 und 0) auf der Festplatte am Sensor vorbeibewegt werden würde, seine Magnetisierungsrichtung ändern. Zwischen zwei unterschiedlich magnetisierten Domänen sind die magnetischen Felder nämlich am größten. Diese sind der Anstoß dafür, dass die Magnetisierungsrichtung der ersten Schicht sich umkehrt und der Widerstand durch den Schichtstapel sich messbar ändert, was für den Lesevorgang der Festplatte ausgenutzt wird.

Wenn allerdings der elektrische Messstrom durch den Schichtstapel über einen kritischen Wert erhöht wird, verursacht der Messstrom selbst ein magnetisches Feld, das die Magnetisierungsrichtung der ersten Schicht

mit einer festen Frequenz oszillieren lässt. Diese Frequenz ändert sich allerdings unter dem Einfluss der magnetischen Felder zwischen den Domänen der Informationseinheiten auf der Festplatte. Bei einem STO-Sensor wird diese Änderung der Oszillationsfrequenz der ersten ferromagnetischen Schicht zum Auslesen der Informationseinheiten genutzt, anders als beim GMR-Sensor, bei dem die Widerstandsänderung gemessen wird.

## 7 Ausblick

Nach dem empirischen Gesetz von Gordon Moore (englisch: „Moore’s Law“) verdoppelt sich seit 1965 nach wie vor die Packungsdichte in kommerziellen integrierten Schaltungen und Speichermedien alle 18 Monate. Die Entwicklung neuer Halbleiterfertigungstechnologien macht dies möglich. Entsprechend werden Sensoren, die auf diesen Technologien basieren, immer kleiner und zugleich intelligenter. Der Einsatz von Halbleiter- und Festkörper-Sensoren wird sich daher weiter rasant ausweiten und sich auch in neuen Anwendungen in der Automobiltechnik, der Unterhaltungselektronik, in der Medizin, im Haushalt und ganz neuen Bereichen etablieren. Gerade in der Halbleiter- und Festkörper-technologie wurden immer wieder vermeintlich unüberwindliche physikalische Grenzen schließlich doch überwunden, weil Innovationen wieder völlig neue Innovationen nach sich gezogen haben, an die vorher nicht zu denken war.

## Nicht-Patentliteratur

- [1] Deutsches Patent- und Markenamt: Internationale Patentklassifikation. URL: <https://depatisnet.dpma.de/ipc/> [abgerufen am 8.4.2015]
- [2] WEIGLEIN, Georg: Feldeffekttransistoren als Sensoren. In: Physik in unserer Zeit, Nr. 3, 1990, S. 113–116. – ISSN 0031-9252
- [3] BÖTTGER, Peter: Ionen- oder chemischsensitive Transistoren. In: Erfinderaktivitäten 1992, Jahresbericht 1992. München: Deutsches Patent- und Markenamt, 1993, S. 33–34. – ISSN 2193-8180

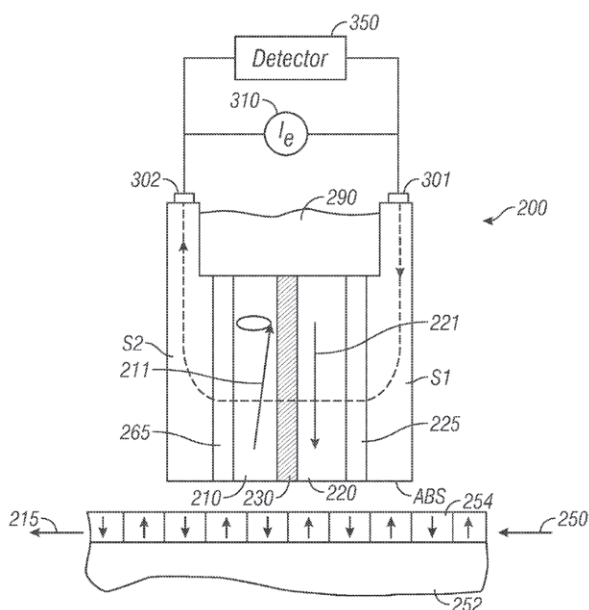


Fig. 21: Querschnitt eines STO-Lesekopfes einer Festplatte, einer Weiterentwicklung eines GMR-Festplattenlesekopfes (aus US 8 675 309 B2)

- [4] BÖTTGER, Peter: Elektronische Nasen und Zungen. In: Erfinderaktivitäten 2000. München: Deutsches Patent- und Markenamt, 2001, S. 16–18. – ISSN 2193-8180
- [5] VOGEL, Michael: Der Fingerhut. In: Physik-Journal, Nr. 01, 2014, S. 40–41. – ISSN 1617-9439
- [6] OSRAM: Light is wearable, The new SFH 7050 for Biomonitoring applications. September 2014. URL: <http://www.osram-os.com/media/resource/HIRES/541656/246267/light-is-wearable---flysheet-biomon-sensor-sfh-7050-gb.pdf> [abgerufen am 8.4.2015]
- [7] PIETRZYK, Uwe; KHODAVERDI, Maryam: Der Blick ins Gehirn, Bildgebende Verfahren in der Gehirnforschung. In: Physik in unserer Zeit, Vol. 37, 5/2006, S. 235–240. – ISSN 0031-9252
- [8] BRUNO, Patrick: Die Entdeckung des Riesen-Magnetowiderstandes, Physik-Nobelpreis 2007. In: Physik in unserer Zeit, Vol. 38, 6/2007, S. 272–273. – ISSN 0031-9252
- [9] VOGEL, Michael: Damit es rund läuft. In: Physik-Journal, Nr. 03, 2015, S. 50–51. – ISSN 1617-9439
- [10] BORN, Axel: 50 Jahre Festplattentechnologie: Highlight and end of the line?. In: Erfinderaktivitäten 2007/2008. München: Deutsches Patent- und Markenamt, 2008, S. 3–8. – ISSN 2193-8180

# Abgassensoren für die Automobilindustrie

Dr. Bettina Kuhn, Patentabteilung 1.52

Die Aktivität der Anmelderschaft auf dem Gebiet der Abgastechnik von Verbrennungsmotoren ist ungebrochen [1]. Gegenstand aktueller Entwicklungen und damit auch beliebte Patentanmeldegebiete sind: Lambdasonden (Sprung- und Breitbandsonden), Rußsensoren, Stickoxid-Sensoren, Ammoniak-Sensoren, Differenzdrucksensoren, Kohlenwasserstoff- sowie Temperatursensoren. Dieser Artikel gibt einen Überblick über etablierte Technologien und aussichtsreiche Weiterentwicklungen zu Gas- und Partikelsensoren für den Abgasstrang. Mit diesen gewinnt man Messgrößen zur Regelung der Verbrennung um Rohemissionen des Motors möglichst gering zu halten und zur On-Board-Diagnose (OBD) von Abgasnachbehandlungssystemen.

---

## 1 Einleitung

Alle Funktionen im Verbrennungsprozess von modernen Motoren und deren Abgassystemen benötigen von der Gemischbildung bis zur Überwachung von Emissionsgrenzwerten robuste und günstig herstellbare Sensoren. Die Regelung der Rohemissionen übernimmt ein Steuergerät, welches unter anderem Informationen von der Lambdasonde erhält, wie den Restsauerstoffgehalt im Abgas. Lambdasonden sind seit den 1970er Jahren in den ersten Fahrzeugen mit geregelter Katalysator auf dem Markt, werden jedoch stetig weiterentwickelt. Bei Dieselfahrzeugen neuerer Bauart werden neben der Lambdasonde aufgrund der hohen Anforderungen des Gesetzgebers an die Überwachung von Partikelfiltern und SCR-Systemen (Selective Catalytic Reduction – Abgasnachbehandlung zur Reduzierung von Stickoxiden) häufig ein Partikel- sowie ein Stickoxid-Sensor ( $\text{NO}_x$ -Sensor) eingesetzt.

Ein Differenzdrucksensor misst den Druckabfall über dem Partikelfilter. Aus seinem Sensorsignal kann der Beladungsgrad des Filters berechnet werden.

Ammoniaksensoren werden häufig in Verbindung mit einem SCR-System in Fahrzeuge eingebaut. Kohlenwasserstoffsensoren (HC-Sensoren) finden in erster Linie Anwendung in Abgasnachbehandlungssystemen, deren Partikelfilter aktiv regeneriert wird. Hierbei sprüht eine Düse sehr feine Dieseltropfen in den Abgasstrom, bei dem im Zusammenspiel mit einem

Oxidations-Katalysator die Abgastemperatur über die Rußzündtemperatur von  $550^\circ\text{C}$  gehoben wird. Ein eventueller Schlupf von unverbranntem Dieseldieselkraftstoff wird dann mit Hilfe eines HC-Sensors detektiert.

Temperatursensoren dienen der Korrektur von Sensordaten, wenn diese eine Temperaturabhängigkeit zeigen, oder als eigenständige Überwachungsmodule der Abgastemperatur.

Dieser Aufsatz beschränkt sich auf Lambdasonden, Partikelsensoren und  $\text{NO}_x$ -Sensoren als Sensorsysteme, die klassischerweise auf einem keramischen Sensorelement basieren.

## 2 Grundlegende Sensorkonzepte

### 2.1 Die Lambdasonde auf Zirkondioxidbasis

Die Luftzahl  $\lambda$  setzt die tatsächlich für eine Verbrennung zur Verfügung stehende Luftmasse  $m_{\text{real}}$  ins Verhältnis zur mindestens notwendigen stöchiometrischen Luftmasse  $m_{\text{stöchiometrisch}}$ , die für eine vollständige Verbrennung benötigt wird. Ein Lambdawert von 1 bedeutet, dass theoretisch alle Kraftstoffmoleküle vollständig mit dem vorhandenen Sauerstoff reagieren und dieser vollständig aufgebraucht werden kann. Ein Lambdawert kleiner 1 bedeutet Luftmangel und in der Konsequenz spricht man von einem fetten Gemisch, da hierbei

Kohlenwasserstoffe im Überfluss vorhanden sind. Bei einem Lambdawert größer 1 herrscht Luftüberschuss, was man als mageres Gemisch bezeichnet.

Das Prinzip der Sauerstoffpartialdruckmessung mit Hilfe eines Sensors auf Festelektrolytbasis geht erstmals aus der DDR-Patentschrift DD 37800 A1 aus dem Jahre 1965 hervor.

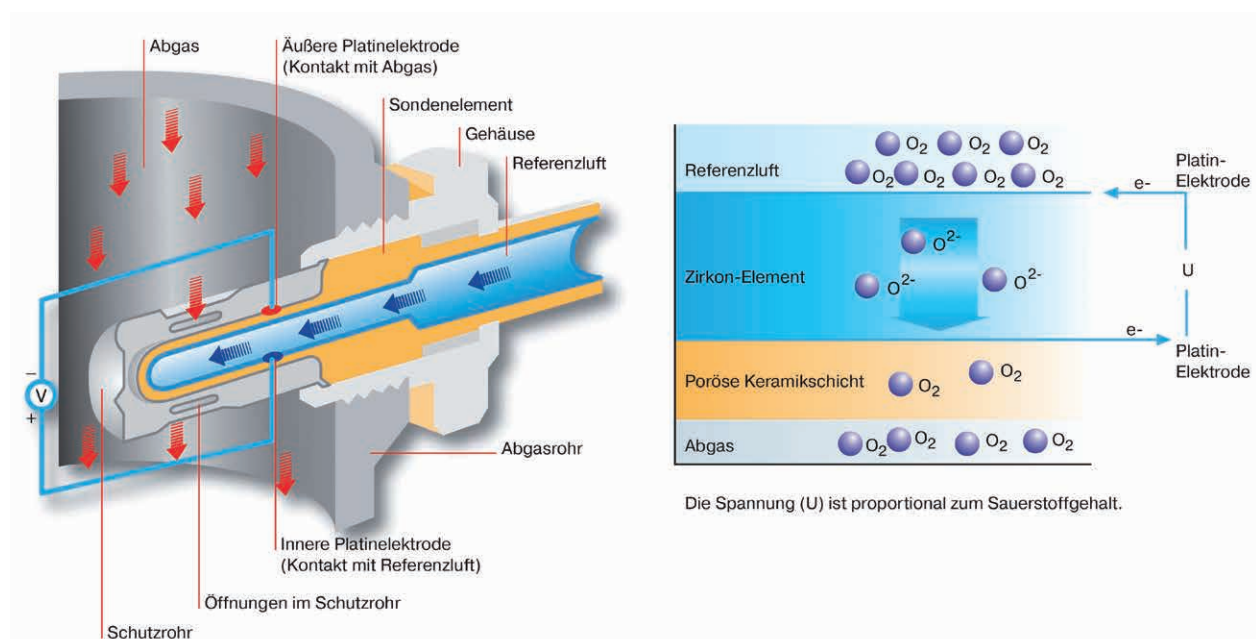
Die Lambdasonde wurde 1976 zur Dosierung der Einspritzmenge bei Ottomotoren eingeführt [2]. Basierend auf dem gängigen Design der Zündkerze sah die erste Lambdasonde aus wie ein konischer Finger und wurde auch so genannt („Fingersonde“). Diese Sonde hatte keine eigene Heizung und wurde durch das Abgas auf eine Temperatur von circa 600 °C gebracht. Lambdasonden der ersten Generation zeigen einen Sprung im Signal bei Lambda gleich 1 an. Man bezeichnet sie daher auch als Sprungsonden. Die Kennlinie ähnelt dem griechischen Buchstaben Lambda, was die Namensgebung begründet [3]. Sie regeln einen Lambdawert nahe 1 ein, was dem vollständigen Verbrennen des Kraftstoffes im stöchiometrischen Gleichgewicht entspricht. Durch den Sprung im Signal ist bei Luftzahlen nahe Lambda gleich 1 eine sehr genaue Messung möglich, bei deutlich kleineren und größeren Luftzahlen ist die Messgenauigkeit sehr viel geringer.

Das Funktionsprinzip der Lambdasonde ist in Figur 1 am Beispiel einer Fingersonde dargestellt. In einem Schutzrohr befindet sich ein keramisches Sensorelement. Der Einbau der Sonde erfolgt hier senkrecht zum Abgasstrom.

Eine Spannungsdifferenz besteht zwischen der Außenelektrode und der Referenzelektrode. Katalytisch aktives Platin an den Elektroden ermöglicht den Einbau von Sauerstoff an der Dreiphasengrenze Platin/Gas-Phase/Zirkondioxid. Der Transport von Sauerstoff-Ionen mit zweifach negativer Ladung erfolgt über den Festelektrolyten aus Zirkondioxid.

Befindet sich im Referenzluftkanal mehr Sauerstoff, wie in Figur 1 dargestellt, wird Sauerstoff an der Referenzelektrode in den Festelektrolyten (hier Zirkondioxid) eingebaut und von dort zur gegenüberliegenden Elektrode, die mit dem Abgas in Kontakt steht, transportiert. Die Spannung zwischen der äußeren Platinelektrode, die mit dem Abgas in Kontakt steht, und der inneren Referenzelektrode ist proportional zum Sauerstoffgehalt im Abgas.

Zur Verkürzung der Kaltstartphase wurde 1982 die Heizung bei Fingersonden eingeführt. 1997 wurden beheizte planare Sensorelemente auf den Markt gebracht.



Figur 1: Funktionsprinzip der Lambdasonde (aus [4])

In der Zwischenzeit wurden neben Lambdasonden vom Sprungsonden-Typ für die Regelung bei Lambda ungleich 1 auch Breitbandsonden entwickelt.

Entwicklungsziele bei Lambdasonden sind aktuell das Ansprechverhalten der Sensoren sowie deren Robustheit, Temperatur- und Langzeitstabilität zu optimieren. Ein Beispiel für die Weiterentwicklung einer Lambdasonde geht aus der Patentanmeldung DE 10 2009 055 041 A1 hervor. Hier wird eine spezielle Strategie beschrieben, die das Zeitintervall der Sonde vom Kaltstart bis zur Betriebsbereitschaft verkürzt.

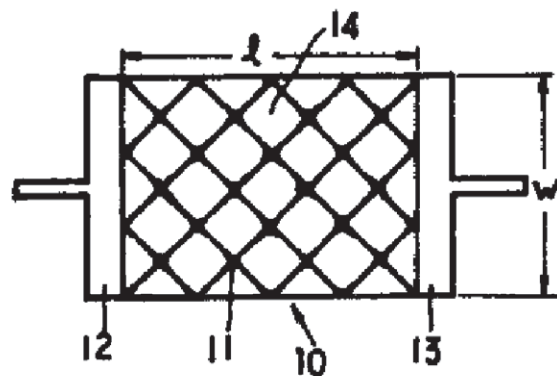
In der Phase des Aufheizens der Sonde kann bei der Verbrennung Wasserdampf entstehen, der auf kalten Oberflächen des Abgasstrangs kondensiert. Wenn nun ein Wassertropfen auf die heiße Sensoroberfläche trifft, können lokale Temperaturunterschiede zu hohen thermischen Spannungen führen, die möglicherweise zu einer Zerstörung der Sonde führen. Dieses Phänomen wird als sogenannter Thermoschock bezeichnet. Um diesen Thermoschock zu verhindern, werden nach der DE 10 2009 055 041 A1 verschiedenste Betriebsstrategien für Sonden entwickelt, wobei im Wesentlichen der Hochheizvorgang der Sonde auf ihre Betriebstemperatur (zwischen circa 650 °C und 850 °C) optimiert wird.

## 2.2 Der Partikelsensor

Partikelsensoren werden zur Überwachung des Partikelfilters (On-Board-Diagnose) meist hinter einem Partikelfilter in den Abgasstrang eingebaut. Die Funktion des Sensorelements basiert auf einer Widerstandsmessung. Bei Partikelsensoren haben sich ebenfalls Sensoren mit planaren Sensorelementen durchgesetzt, häufig auch Keramiken mit sehr geringer elektrischer Leitfähigkeit, die hohe Temperaturstabilität aufweisen. Auf ein keramisches Substrat wird auf die Unterseite ein Heizelement aufgebracht. In der Mitte des Sensors befindet sich ein Temperatursensor und auf der Oberseite eine sogenannte Interdigitalelektrode, bestehend aus zwei ineinandergreifenden Kammelektroden. Beim Anlegen einer Spannung zwischen den Elektroden bildet sich ein elektrisches Feld, welches Rußpartikel anzieht. Diese bilden dann in dendritischen Strukturen zwischen den beiden Kammelektroden eine

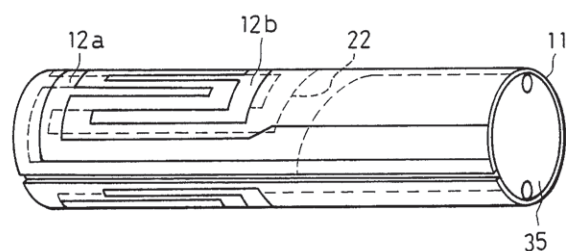
Schicht. Die Leitfähigkeit dieser Rußpfade wird über eine Elektrometerschaltung gemessen und ausgewertet.

Die japanische Patentanmeldung JP 57147043 A aus dem Jahr 1982 zeigt erstmals einen Sensor, bei dem die Anlagerung des Rußes zwischen zwei Elektroden 12 und 13 erfolgt (siehe Figur 2). Die Gitterstruktur 11 besteht aus einer Keramik, an der Ruß haften bleibt. Mit zunehmender Anlagerung von Rußpartikeln nimmt auch hier der elektrische Widerstand ab und dient als Maß für die Rußkonzentration im Abgas.



Figur 2: Netzartige Struktur zwischen zwei Elektroden eines Partikelsensors (aus JP 57147043 A)

Die erste Erwähnung eines Partikelsensors, dessen Sensorelement zylinderförmig aufgebaut ist, findet sich in der Patentschrift US 4 656 832 A aus dem Jahr 1987. Figur 3 zeigt ein Ausführungsbeispiel des Sensors als Zylinder. Auch hier hat man auf eine bereits bekannte Sensorform zurückgegriffen und die homogene Umströmung der Zylinderform genutzt. Die Elektrodenstruktur umfasst zwei kammartig ineinandergreifende Elektroden 12a und 12b. Der Widerstand zwischen diesen beiden Elektroden wird während der Rußanlagerung gemessen.



Figur 3: Partikelsensor (aus US 4 656 832 A)



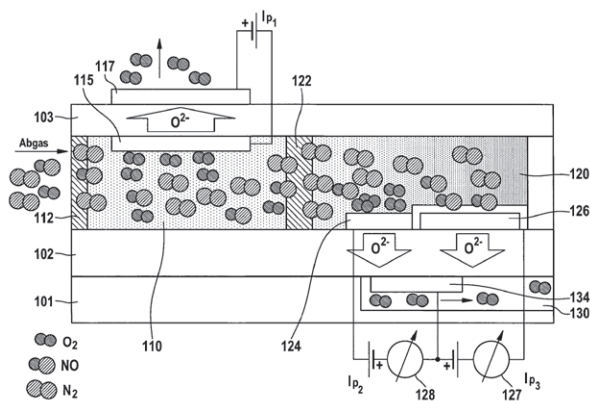
### 2.3 Der Stickoxid-Sensor ( $\text{NO}_x$ -Sensor)

Der  $\text{NO}_x$ -Sensor dient der bedarfsgerechten Regelung der eindosierten Harnstoffmenge in SCR-Systemen zur  $\text{NO}_x$ -Reduzierung sowie zur Überwachung der SCR-Komponenten beziehungsweise eines  $\text{NO}_x$ -Speicherkatalysators (On-Board-Diagnose).

Ein Vorläufer des  $\text{NO}_x$ -Sensors auf Festelektrolyt-Basis ist der Patentschrift US 3 622 487 A aus dem Jahr 1971 zu entnehmen, wobei hier als Elektrolyt eine wässrige Lösung eingesetzt wurde.

Der Aufbau eines  $\text{NO}_x$ -Sensors, wie er heute in der Serie zu finden ist, zeigt Figur 4. Ein keramisches Sensorelement mit zwei hintereinanderliegenden Kammern bestimmt den  $\text{NO}_x$ -Gehalt im Abgas. Beide Kammern sind durch eine Diffusions- oder Strömungsbarriere 122 voneinander getrennt. Das Abgas strömt durch die Diffusionsschichten von links nach rechts. In drei Pumpvorgängen reagieren Sauerstoff und Stickoxide an den Elektroden.

Das Abgas strömt durch eine erste Diffusionsbarriere in die erste Kammer des Sensors. Dort wird Sauerstoff „abgepumpt“. Die Gegenelektrode 117 befindet sich in Verbindung zum Abgas. Der Pumpstrom  $I_{p1}$  ist proportional zum Sauerstoffgehalt. Gas strömt aus der ersten Kammer über eine poröse Diffusionsschicht 122 in die zweite Kammer und an einer ersten Elektrode 124 wird erneut Sauerstoff abgepumpt ( $I_{p2}$ ), wobei als Gegenelektrode eine Luftreferenzelektrode 134 dient, die mit einem Luftreferenzraum 130 in Verbindung steht.



Figur 4:  $\text{NO}_x$ -Sensor mit zwei Kammern (aus DE 10 2009 046 232 A1)

Unterhalb einer porösen Diffusionsschicht ist eine zweite Elektrode, die  $\text{NO}_x$ -Elektrode 126 angeordnet. An dieser Elektrode reagieren Stickoxide, beispielsweise Stickstoffmonoxid (NO) zu Sauerstoff ( $\text{O}_2$ ) und Stickstoff ( $\text{N}_2$ ). Wenn der Sauerstoffgehalt in der zweiten Kammer ausreichend gering ist, kann anhand des Pumpstroms  $I_{p3}$  auf die NO-Konzentration im Abgas geschlossen werden.

### 3 Ausblick

Komplexe Abgasnachbehandlungssysteme helfen, die Emissionen von Verbrennungsmotoren im Automobilbereich zu reduzieren. Zu ihrer Überwachung und Steuerung werden eine Vielzahl etablierter und neuer Sensoren eingesetzt. Angetrieben durch die Abgasgesetzgebung wird neben der Steuerung abgasrelevanter Funktionen auch die Überwachung der Komponenten immer wichtiger.

Neben neuen Sensortechnologien sind vor allem Maßnahmen zur Langlebigkeit, zur Temperatur- und Thermoschock-Stabilität und zur verbesserten Eigen diagnose Gegenstand aktueller Patentanmeldungen. Durch den starken wirtschaftlichen Druck und die hohen Anforderungen an die Elektronikarchitektur moderner Fahrzeuge gibt es auch Bestrebungen, Sensorsignale in verschiedenster Art und Weise auszuwerten und mehrfach zu verwenden. So geht beispielsweise aus der Patentschrift US 6 925 373 B2 ein virtueller Rußsensor hervor, der sich der Auswertung diverser Sensordaten des Verbrennungsprozesses bedient. Die Komplexität einer modernen Abgasanlage mit diversen Abgasnachbehandlungssystemen und deren Überwachung hat große Bedeutung sowohl bei Fahrzeugherstellern als auch bei deren Lieferanten.

## Nicht-Patentliteratur

- [1] DEUTSCHES PATENT- UND MARKENAMT  
(Hrsg.): Jahresbericht 2014. München, 2014. S. 13
- [2] Robert Bosch GmbH (Hrsg.): Produktionsjubiläum:  
500 Millionen Lambdasonden bei Bosch gefertigt,  
Serienfertigung seit 1976. In: Presse-Information,  
Mai 2008. URL: [http://www.bosch-presse.de/  
presseforum/details.htm?txtID=3544&tk\\_id=108](http://www.bosch-presse.de/presseforum/details.htm?txtID=3544&tk_id=108)  
[abgerufen am 15.05.2015]
- [3] RIEGEL, J., NEUMANN, H., WIEDENMANN  
H.-M.: Exhaust gas sensors for automotive emis-  
sion control. In: Solid State Ionics 152-153 (2002),  
S. 783–800
- [4] NGK NTK (Hrsg.): Das Geheimnis der Lambda-  
sonde. In: Broschüre für Kunden, Ratingen 2011,  
S. 6 – mit freundlicher Genehmigung der NGK  
Spark Plug Europe GmbH, Ratingen

# Kameras in Kraftfahrzeugen

*Dipl.-Ing. Steffen Merunka, Patentabteilung 1.56*

Das Automobil und die Kamera wurden vor weit mehr als 100 Jahren erfunden. Beide miteinander zu kombinieren, hat Erfinder schon früh beschäftigt. Konstruktiv ist es an sich ja keine große Sache, eine Kamera in ein Kraftfahrzeug zu montieren. Nur zu welchem Zweck? Wo im Kraftfahrzeug? Wann und wie wird die Kamera ausgelöst? Wie werden die Bilddaten erfasst, gespeichert und verarbeitet? Wie wird die Kamera mit Energie versorgt? Und welchen Nutzen bringt all das für die fahrende Person, die weiteren Fahrzeuginsassen und andere Verkehrsteilnehmer? Der nachfolgende Beitrag wirft einen Blick auf die Historie und stellt einige aktuelle Entwicklungen vor.

---

## 1 Einleitung

Unter Kameras werden im Folgenden Aufnahmegeräte sowohl für bewegte Bilder als auch für statische Bilder verstanden, also sowohl Filmkameras als auch Fotoapparate. Einschränkungen hinsichtlich der Aufnahme-medien oder Aufzeichnungsverfahren werden keine gemacht. Die technischen Möglichkeiten reichen von der Kamera im Wortsinne, als Camera Obscura beziehungsweise Lochkamera, über die Aufzeichnung auf fotografischen Film, die elektronische Aufzeichnung auf Magnetband bis hin zur Digitalisierung und Erfassung mittels elektronischer Bildsensoren.

Der Begriff Kraftfahrzeug wird im Folgenden im Zusammenhang mit mehrspurigen Straßenfahrzeugen verwendet. Autonome Kraftfahrzeuge werden nicht gesondert betrachtet, denn mit diesen befasst sich ein anderer Beitrag der vorliegenden Erfinderaktivitäten.

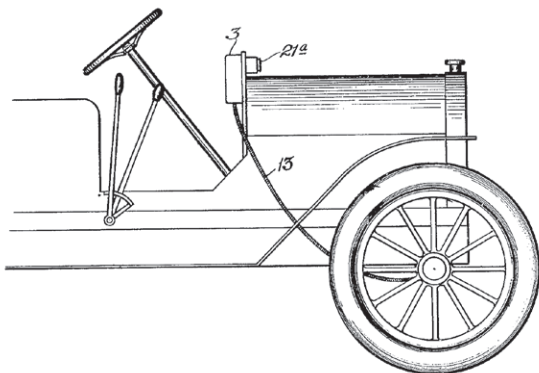
Heutige Kraftfahrzeuge weisen eine Vielzahl von Fahrerassistenzsystemen auf. Diese Systeme sind dafür ausgelegt, die Menschen bei der Fahraufgabe zu unterstützen oder in bestimmten Fahrsituationen sie zu ersetzen. Die fahrende Person wird so von Routinetätigkeiten entlastet, bei unangenehmen oder unkomfortablen Fahraufgaben unterstützt oder abgelöst. Und in kritischen oder gefährlichen Verkehrssituationen kann besser, schneller und entschlossener gehandelt werden. Derartige Fahrerassistenzsysteme sind auf die Erfassung des Fahrzeugumfelds, der Verkehrssituation und

der Fahrzeugparameter angewiesen. Dafür werden unterschiedliche Sensoren und Datenerfassungssysteme eingesetzt. Die Sensordaten werden zu Umgebungsmodellen fusioniert, welche das Fahrzeugumfeld und die Fahrsituation abbilden. Kameras spielen hierbei als Sensoren eine wichtige Rolle. Die folgenden Abschnitte geben einen Überblick über den Einsatz von Kameras in Kraftfahrzeugen und gehen dabei auch auf die historische Entwicklung ein. Es wird gezeigt, dass für den Einsatz von Kameras in Kraftfahrzeugen, die in heutigen Fahrzeugen als serienreife Komponenten verbaut werden oder sich für noch nicht marktreife Fahrerassistenzsysteme noch in der Entwicklung befinden, zum Teil bereits vor 50 oder sogar 100 Jahren in der Patentliteratur Spuren nachweisbar sind.

## 2 Historische Entwicklungen in der ersten Hälfte des 20. Jahrhunderts

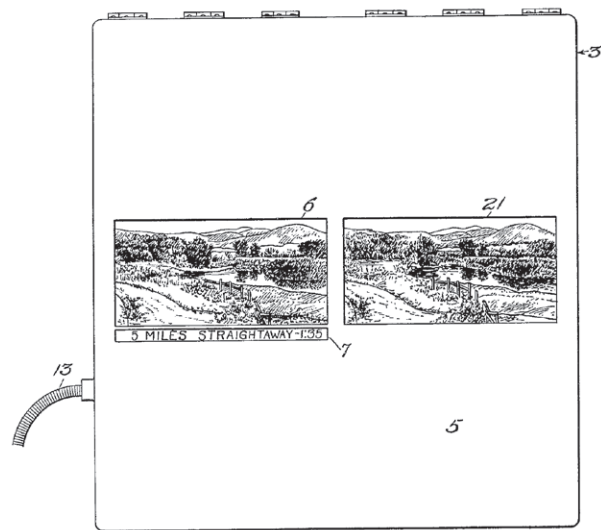
Die Kombination von Kameras und Kraftfahrzeugen hat schon früh den Ideenreichtum von Erfindern angeregt. So zeigt die Patentschrift US 990 739 aus dem Jahr 1911 ein Kraftfahrzeug mit einem Film-basierten Navigationssystem, hier Location Indicating Device. Das Gerät dient der Unterstützung beim Befahren unbekannter Strecken und stellt während der Fahrt eine synchron mit der Fortbewegung des Fahrzeugs fortschreitende Kartenansicht auf einem Bildschirm dar. Für die Erzeugung der Karte befährt ein Pilotfahrzeug mit einem streckenkundigen Fahrer die Strecke vorab

und macht in regelmäßigem Abstand Bildaufnahmen der Strecke und des Umfeldes. Diese Bildaufnahmen werden nachgelagert außerhalb des Fahrzeugs bearbeitet und mit Fahrthinweisen, beispielsweise betreffend Entfernung, Fahrtzeit und Richtung, versehen. Von den bearbeiteten Bildaufnahmen werden Reproduktionen erstellt. Diese können dann von Anderen beim Befahren einer unbekannten Strecke ausgewählt und im Navigationssystem verwendet werden. Hierzu ist in den dafür vorgesehenen Kraftfahrzeugen ein Anzeigegerät eingebaut (siehe Figur 1), das während der Fahrt die Bildaufnahmen der Reproduktion 6 darstellt und zugleich daneben über eine Lochkamera eine aktuelle Streckensicht 21 einblendet (siehe Figur 2). Aus dem Vergleich zwischen den beiden Darstellungen kann die fahrende Person den Streckenfortschritt erkennen und Navigationshinweise im Hinblick auf die aktuelle Fahrsituation entnehmen.



Figur 1: Kraftfahrzeug mit einem filmbasierten Navigationssystem (aus US 990 739)

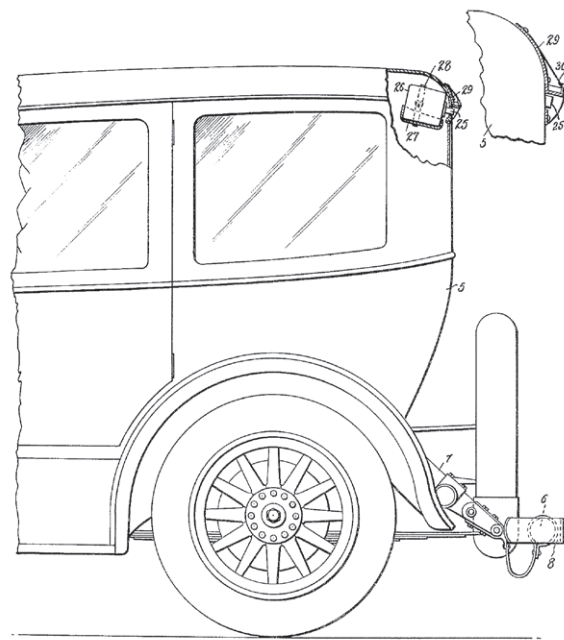
Auch Kameras zur Unfalldokumentation sind bereits in der Patentliteratur der 1920er Jahre vorhanden. So beschreibt die Patentschrift US 1 701 800 aus dem Jahr 1929 eine Kamera, die am Fahrzeugheck beziehungsweise wie in Figur 3 am Dachrahmen erhöht montiert ist und der Aufnahme eines Bildes im Falle eines Heckaufpralls dient. Die Kamera wird elektrisch über mehrere im Stoßfänger integrierte elektromechanische Schalter sowie einem Elektromagneten für die Betätigung des Verschlusses ausgelöst. Die Schalter sind so angeordnet, dass sowohl ein direkter Heckaufprall als auch ein solcher mit seitlichem Versatz zur Auslösung der Kamera führt. Ferner ist ein Blitzlichteinsatz vorgesehen, der über einen elektrischen



Figur 2: Darstellung von Navigationssicht und Streckensicht (aus US 990 739)

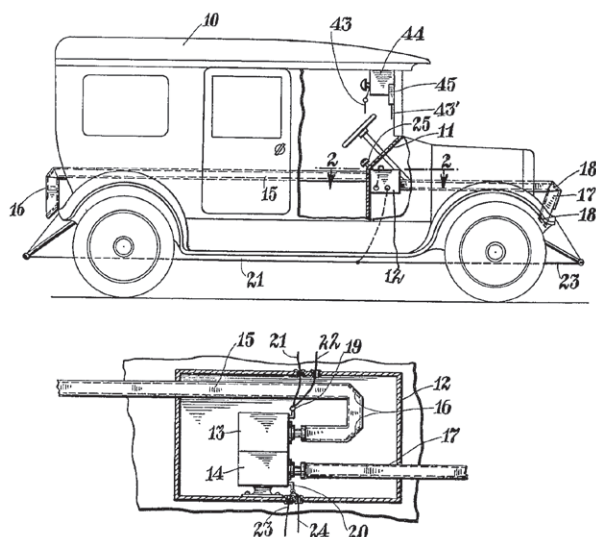
Zünder aktiviert wird. Die Vorrichtung kann über zusätzliche Schalter deaktiviert sowie zwischen Tag- und Nachtbetrieb (Blitzlichteinsatz) umgeschaltet werden.

Die Patentschrift US 1 733 783 (siehe Figur 4), ebenfalls aus dem Jahr 1929, zeigt ein Fahrzeug, welches mit Filmkameras zur Unfalldokumentation und zum Diebstahlschutz ausgestattet ist. Zwei im Armaturenbrett angeordnete Filmkameras 13 und 14 sind mit zur Front



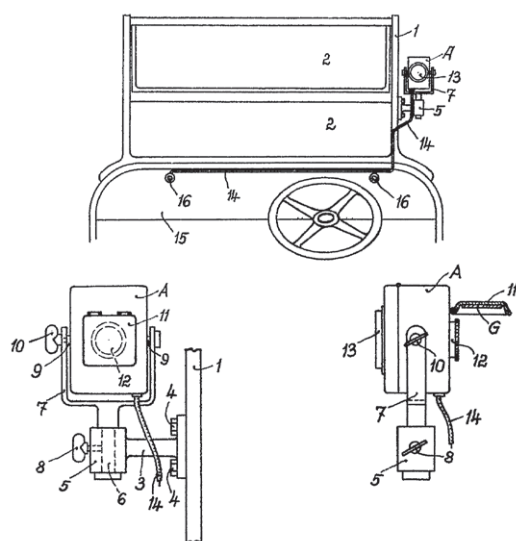
Figur 3: Fotokamera mit elektromechanischer Auslösung bei einem Heckaufprall (aus US 1 701 800)

und zum Heck verlaufenden Periskopanordnungen 15 und 17 verbunden und filmen, ausgelöst über eine Anordnung von Hebeln und Seilzügen, im Falle eines Crashes das unmittelbare Fahrzeugumfeld. Zusätzlich ist eine Innenraumkamera 44 vorgesehen, die bei nicht autorisierter Verwendung des Fahrzeugs elektrisch ausgelöst das Fahrzeuginnere filmt.



Figur 4: Fahrzeug mit Front-, Heck- und Innenkamera (aus US 1 733 783)

Die österreichische Patentschrift AT 131 738 B aus dem Jahr 1933 zeigt in Figur 5 ebenfalls eine Kamera zur Unfalldokumentation. Neben einer manuellen Auslösung der Kamera durch einen Fahrzeuginsassen



Figur 5: Fahrzeug mit Unfallüberwachungskamera und Überwachung der Brems- und Lenkaktivität (AT 131 738 B)

über Taster 16 ist hier eine automatische Auslösung in Abhängigkeit von Brems- oder Lenkbetätigung vorgesehen. Realisiert wird dies beispielsweise durch eine Überwachung der Bremshebelbetätigung mittels Fliehkraftpendel und durch einen bei starker Bremsverzögerung auslösenden Kipphebel. Es ist bemerkenswert, dass hier mit recht einfachen Mitteln bereits eine fahrsituationsabhängige Auslösung der Kamera dargestellt wurde.

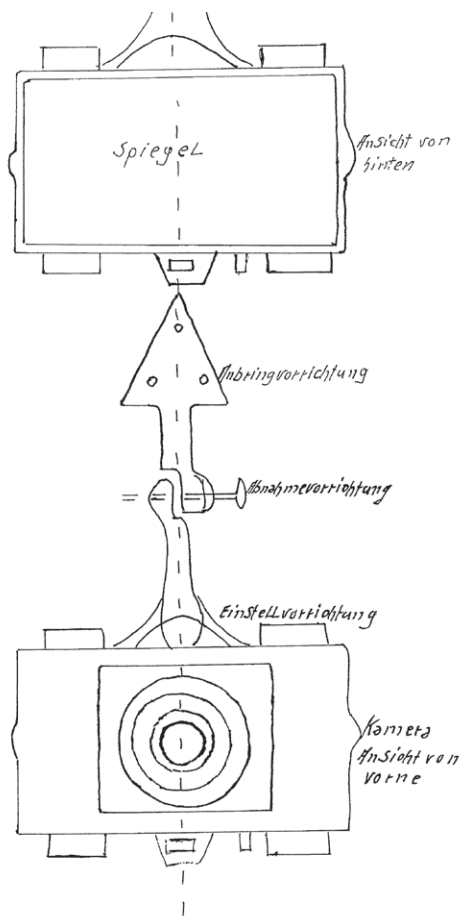
### 3 Historische Entwicklungen in der zweiten Hälfte des 20. Jahrhunderts

Im Zuge der Massenmotorisierung, die in Deutschland ab den 1950er Jahren begann, stieg der Anteil des motorisierten Individualverkehrs und das Automobil wandelte sich vom Luxusgegenstand zum Konsumgut. Vor diesem Hintergrund überrascht nicht, dass die damit einhergehenden nachteiligen Effekte wie höhere Verkehrsdichte und steigende Unfallzahlen die Menschheit angeregt haben, technische Lösungen zu entwickeln, die das Autofahren sicherer und dessen Folgen besser beherrschbar machen.

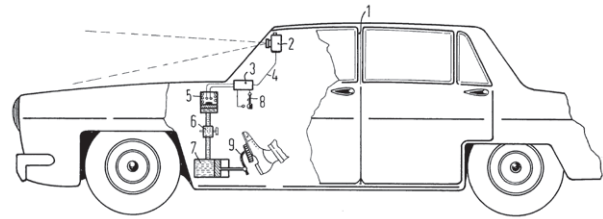
Aus dieser Zeit existieren bereits Veröffentlichungen, die als Vorläufer der heute bekannten sogenannten Dash-Cams gelten können, also Kameras, die im Fahrzeug während der Fahrt das Fahrzeugumfeld annähernd aus Fahrerperspektive aufzeichnen. So zeigt die Gebrauchsmusterschrift DE 1 725 190 U aus dem Jahr 1956 einen Autorückspiegel mit daran angeordneter Kamera (siehe Figur 6). Mit der Kamera kann während der Fahrt das Verkehrsgeschehen aufgezeichnet werden. Bei Unfällen, Sachschäden oder anderen unklaren Verkehrslagen können die Bilder als Beweismittel für Polizei, Gericht und Versicherung verwendet werden. Ferner ist die gesamte Einheit aus Rückspiegel und Kamera abnehmbar ausgebildet, um Bilder außerhalb des Fahrzeugs aufzunehmen. Datenschutzrechtliche Aspekte, die heutzutage zumindest in Deutschland die Debatte um die Verwendung von Dash-Cams zu beherrschen scheinen, werden nicht erwähnt. Bemerkenswert ist, dass in der Gebrauchsmusterschrift ausdrücklich die Verwendung als Verkehrserziehungsmittel, und zwar für die anderen Verkehrsteilnehmer, beansprucht wird.



Für die voll- oder halbautomatische Auslösung von im Fahrzeug verbauten Kameras bietet die Patentliteratur ebenfalls vielfältige Lösungen. Die Gebrauchsmusterschrift DE 1 904 397 U aus dem Jahr 1964 zeigt in Figur 7 die Auslösung durch Bremsbetätigung mittels Bremsdruckschalter, sobald der Fahrer wegen einer kritischen Verkehrssituation überdurchschnittlich stark bremst. Ferner ist ein Pendelschalter vorgesehen, der bei starker Fahrzeugverzögerung ohne eigene Bremsbetätigung, also bei äußerer Einwirkung auf das Fahrzeug, auslöst. Die Patentschrift DE 972 776 B aus dem Jahr 1959 (siehe Figur 8) zeigt die elektromechanische Auslösung der Kamera durch außen an der Karosserie angebrachte elektromechanische Kontaktschalter, beispielsweise an den Stoßfängern und am Kotflügel. Diese Anordnung hat den Vorteil, dass sie unabhängig von der Betätigung durch Fahrzeuginsassen arbeitet und auch bei Abwesenheit derselben aktiv bleiben kann, beispielsweise um abgestellte Wagen überwachen zu können.

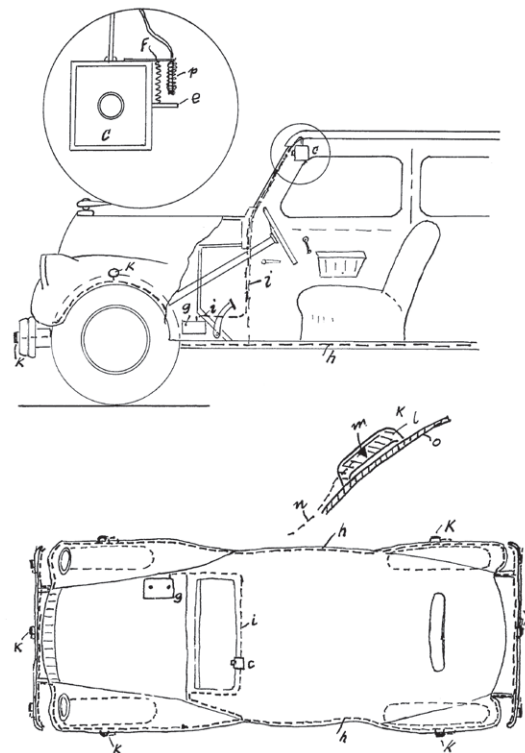


Figur 6: Rückspiegel mit daran angeordneter abnehmbarer Kamera (aus DE 1 725 190 U)

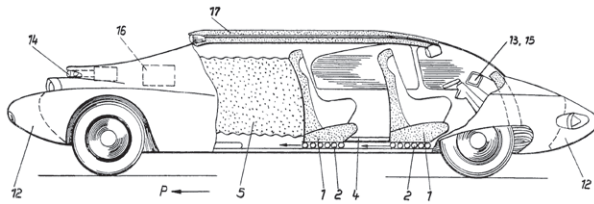


Figur 7: Kameraauslösung durch Bremsdruckschalter oder Pendelschalter (aus DE 1 904 397 U)

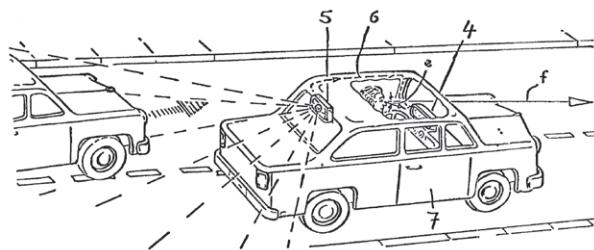
Auch für den Schutz der Fahrzeuginsassen bei Unfällen wurden Ideen entwickelt, die zur Lösung spezieller Probleme Kameras verwenden. Eine überraschende Möglichkeit zeigt Manfred von Ardenne in der DD 27 588 A1 aus dem Jahr 1964 (siehe Figur 9). Dort ist vorgesehen, zum Schutz der Insassen die Sitze entgegengesetzt zu der durch den Pfeil P angedeuteten Fahrtrichtung anzuordnen und innerhalb der Fahrgastzelle frontseitig einen Kunststoffpuffer 5 einzubringen, der beim Aufprall deformiert wird. Da bei einer derartigen Anordnung allerdings der Fahrer oder die Fahrerin keine Sicht auf die vor dem Fahrzeug liegende Fahrbahn mehr hat, soll eine Fernsehanlage 13 installiert werden, die mittels Fernsehkamera das Fahrzeugumfeld aufnimmt und



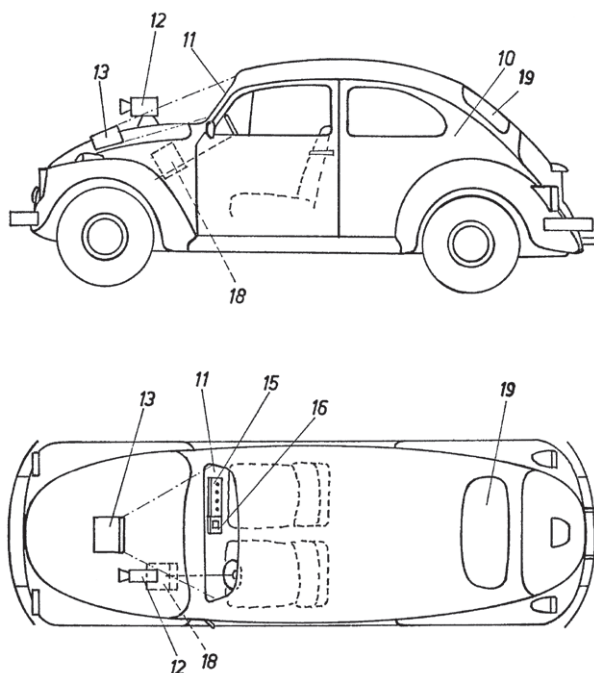
Figur 8: Elektromechanische Kameraauslösung an den Fahrzeugaußenseiten (aus DE 972 776 B)



Figur 9: Unfallsicheres Landfahrzeug mit entgegengesetzt der Fahrtrichtung angeordneten Sitzen und Panorama-Fernsehanlage zur Fahrersicht auf die Straße (aus DD 27 588 A1)



Figur 10: Kamera-Monitoranordnung zur Vermeidung des toten Winkels (aus DE 7 146 943 U)



Figur 11: Wiedergabevorrichtung im Fahrerblickfeld mit Möglichkeit zur Einblendung von Verkehrssituationen (aus DE 2 224 908 A)

eine Panoramaansicht des Straßenbilds darstellt. Für die Aufnahme einer 180° Panoramabilds sind drei Kameras mit einem horizontalen Winkelbereich von 70° vorgesehen. Sie sollen außerdem für Infrarotaufnahmen empfindlich sein, um auch bei schlechten Sichtverhältnissen, beispielsweise bei Nebel oder Dunkelheit, eine Fahrersicht zu ermöglichen.

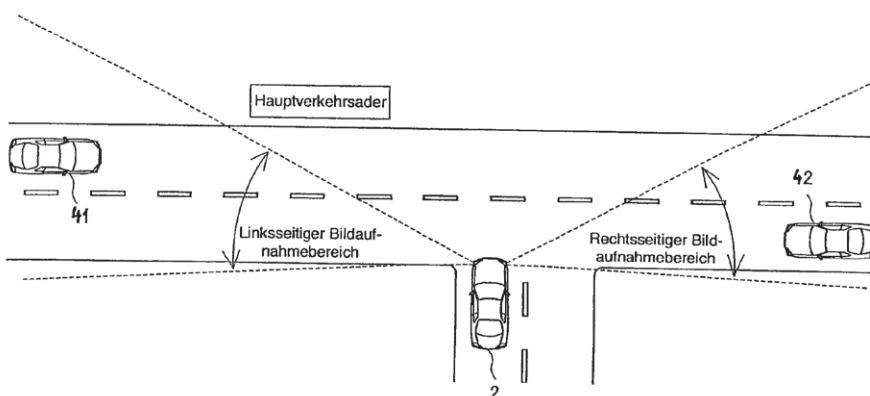
Herkömmliche Rückspiegelanordnungen weisen konstruktiv zwangsläufig einen toten Winkel auf, der die Beobachtung des Fahrzeugumfelds nach hinten und zur Seite einschränkt. Auch für dieses Problem finden sich im Stand der Technik Lösungen, die Kameras verwenden. So zeigt die Gebrauchsmusterschrift DE 7 146 943 U aus dem Jahr 1972 auf, wie eine Kamera innen im Fahrzeug an der Heckscheibe angebracht werden kann, um das rückwärtige Verkehrsgeschehen aufzunehmen. Am Armaturenbrett, im Blickfeld der fahrenden Person, ist ein Monitor gemäß Figur 10 verbaut, so dass der Fahrer oder die Fahrerin, ohne wesentlich vom Blick auf die Fahrbahn abgelenkt zu werden, das rückwärtige Verkehrsgeschehen überwachen kann. Die Kamera kann zusätzlich mit einer motorischen Schwenkvorrichtung sowie mit einer Heizung für das Abtauen der Heckscheibe ausgerüstet sein.

Auch Anordnungen ähnlich heutigen Head-up Displays, die Informationen auf der Frontscheibe im Blickfeld der fahrenden Person einblenden, sind bereits vor mehr als 40 Jahren im Stand der Technik belegt. Die DE 2 224 908 A aus dem Jahr 1973 (siehe Figur 11) zeigt eine Anordnung, die im Fahrbetrieb das Verkehrsgeschehen vor dem Fahrzeug mit einer Kamera aufnimmt und dann mittels einer Wiedergabevorrichtung auf die Frontscheibe abbildet. Das entsprechend ausgerüstete Fahrzeug ist für Fahrprüfungen und Tauglichkeitsuntersuchungen vorgesehen. Hierfür werden zusätzlich bestimmte Verkehrssituationen über ein Bordgerät auf der Wiedergabevorrichtung eingeblendet und die Fahrerreaktion beobachtet. Ferner können Umwelteinflüsse sowie Störungen am Kraftfahrzeug simuliert werden, wozu auch Stellglieder vorgesehen sind, die auf Lenkung, Antrieb oder Bremse einwirken.

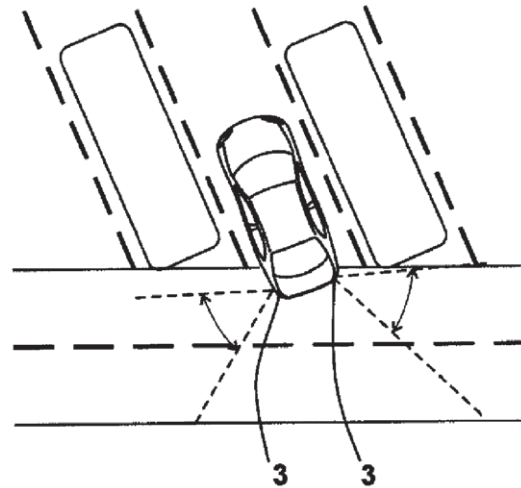
#### 4 Aktuelle Entwicklungen

Heutige Serienkraftfahrzeuge weisen bereits eine Vielzahl von Fahrerassistenzsystemen auf, die bei der Fahraufgabe unterstützen oder Teile der Fahraufgabe abnehmen können. Verbesserte und neue Fahrerassistenzsysteme erreichen in immer kürzeren Abständen die Serienreife. Viele dieser Fahrerassistenzsysteme sind auf Bildinformationen angewiesen. Für die Erfassung der Bildinformationen werden zumeist im Fahrzeug verbaute Kameras eingesetzt, wobei als Bildsensoren neben CCD-Sensoren zunehmend CMOS-basierende Sensoren verwendet werden. Vor allem die CMOS-basierenden Sensoren weisen erhebliche Vorteile wie geringe Baugröße, Integration von Auswertelogik und geringen Stromverbrauch auf und erlauben es den Konstrukteuren, kostengünstig eine Vielzahl von Kameras im Fahrzeug an Einbauorten unterzubringen, die noch vor wenigen Jahren weder konstruktiv noch wirtschaftlich realisierbar gewesen wären. Ferner erlaubt die Integration von Auswertelogik und Signalverarbeitung auf dem Bildsensor eine Vorverarbeitung, die eine Integration in die verschiedenen Fahrerassistenzsysteme ermöglicht.

Risikante Fahrsituationen, bei denen das Fahrzeugumfeld nicht oder nicht vollständig einsehbar und die Verkehrssituation somit schwer einzuschätzen ist, sind wohl jedem Fahrer bekannt, beispielsweise beim Ein- oder Ausparken, bei Kolonnenfahrten oder bei Überholabsicht. In derartigen Situationen können Kameras beziehungsweise damit ausgestattete Fahrerassistenzsysteme den Fahrer stark entlasten.



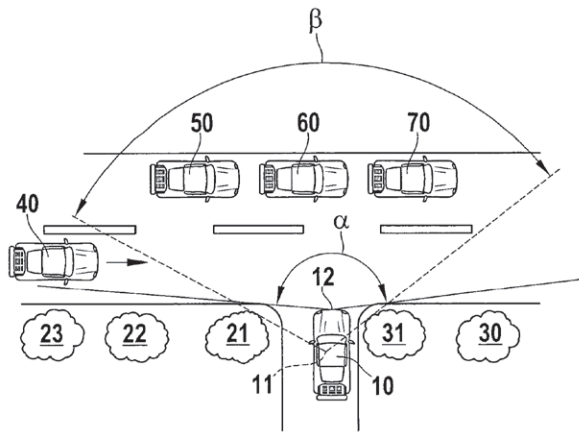
Figur 13: Frontsichtkameras für unübersichtliche Verkehrssituationen bei Einmündungen (aus DE 10 2005 013 920 B4)



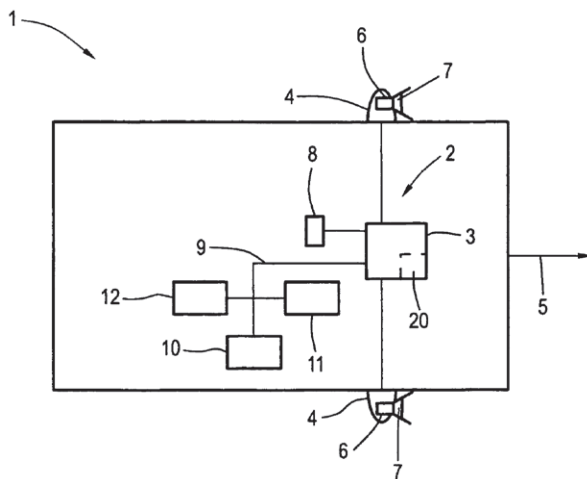
Figur 12: Kameraeinsatz bei unübersichtlicher Verkehrssituation beim rückwärts Ausparken (aus AT 509 121 A1)

Die österreichische Veröffentlichung AT 509 121 A1 zeigt in Figur 12 ein in den Scheinwerfern oder in den Heckleuchten eines Kraftfahrzeugs verbautes Infrarotkamerasystem. Dessen Kameras sind beidseitig angeordnet, seitwärts ausgerichtet und weisen einen recht großen Erfassungswinkel auf. Somit wird es ermöglicht, in einer unübersichtlichen Verkehrssituation, wie etwa beim Ausparken rückwärts aus einer quer zur Fahrbahn angeordneten Parkposition, bereits beim minimalen Ausfahren einen vollständigen Überblick über die Verkehrssituation zu erlangen. Aus der Patentschrift DE 10 2005 013 920 B4 ist ein Fahrzeug mit zwei Frontsichtkameras bekannt, welches erlaubt, bereits beim minimalen Einfahren in eine Kreuzung die Verkehrssituation erfassen zu können (siehe Figur 13). Und die DE 10 2007 055 182 A1 zeigt in Figur 14 ein Fahrzeug, welches sowohl mit Scheinwerferkameras

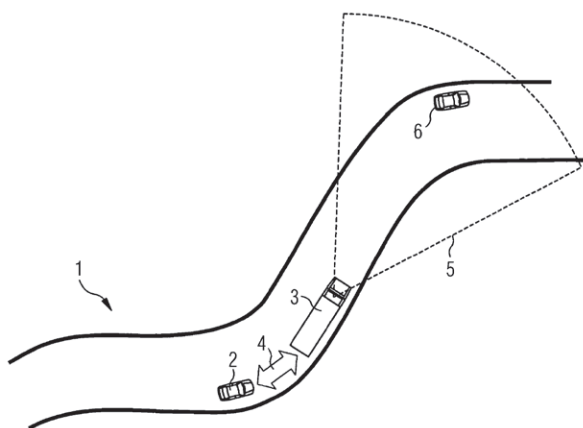
als auch mit einer mittig angeordneten Frontkamera mit Prismenspiegel ausgerüstet ist. Dadurch wird in unübersichtlichen Verkehrssituationen, beispielsweise bei von Bäumen verdeckter Sicht auf den Querverkehr, der Blickwinkel der fahrenden Person zu den Seiten hin erheblich vergrößert.



Figur 14: Frontkamera mittig am Stoßfänger  
(aus DE 10 2007 055 182 A1)



Figur 15: Kameras in Außenspiegeln verbaut für Überholassistent  
(aus EP 2 479 077 B1)

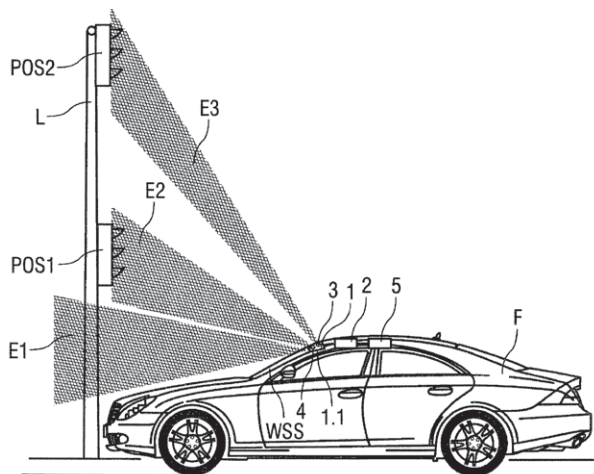


Figur 16: Übertragung von Bilddaten zwischen Fahrzeugen  
(aus DE 10 2006 055 344 A1)

Ein für Überholvorgänge vorgesehenes Fahrerassistenzsystem beschreibt die europäische Patentschrift EP 2 479 077 B1 (vergleiche Figur 15). Dort sind in den Außenspiegeln Kameras verbaut, die in Fahrtrichtung orientiert sind. Somit kann auch bei vorausfahrenden Fahrzeugen oder bei einer Kolonnenfahrt die Verkehrssituation vor dem Fahrzeug erfasst werden, ohne dass das Fahrzeug in riskanter Weise an den Fahrspurrand oder darüber hinaus bewegt werden muss. Die Bilddaten werden verwendet, um unter anderem Anzahl, Länge, Geschwindigkeit und Lückengröße der vorausfahrenden Fahrzeuge oder der Fahrzeugkolonne zu ermitteln sowie um möglichen Gegenverkehr oder Verkehrszeichen zu erfassen. Ferner werden vom Navigationssystem Streckendaten zugestellt, um den Streckenverlauf oder eine Annäherung an Kreuzungen oder Einmündungen zu ermitteln. Die gesammelten Daten werden sodann ausgewertet, um für die fahrende Person eine Überholempfehlung oder wichtiger, die Empfehlung nicht zu überholen, auszugeben.

Noch einen Schritt weiter geht die Lösung in Figur 16 aus der DE 10 2006 055 344 A1. Hierbei wird die Kamera eines vorausfahrenden Fahrzeugs verwendet, um im Nachfolgefahrzeug eine Bildansicht der eigentlich verdeckten Strecke anzuzeigen. Das vorausfahrende Fahrzeug überträgt dazu die mittels Frontkamera aufgenommenen Bilddaten über ein drahtloses ad-hoc-Netzwerk an das Nachfolgefahrzeug, welches die Bilddaten anzeigt.

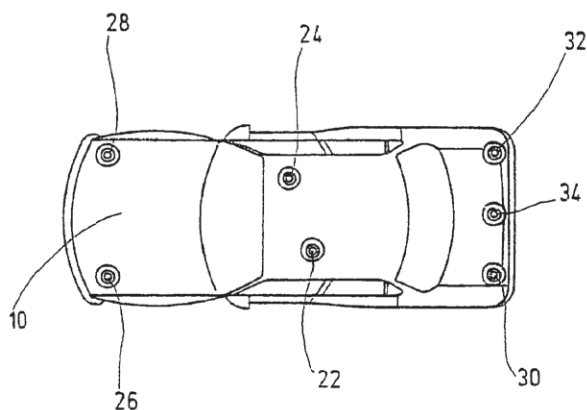
Für Verkehrssituationen vor Lichtsignalanlagen, umgangssprachlich Ampeln, ist die in der DE 10 2009 040 252 A1 (siehe Figur 17) dargestellte Kameraanordnung geeignet. Hierbei werden von einem Fahrerassistenzsystem bei einem vor einer Ampel stehenden Fahrzeug durch zwei Kameras typische Positionen im Erfassungsbereich oberhalb vor der Fahrzeugfront überwacht, um die Lichtsignalgeber zu erfassen, ohne dass die fahrende Person eine unbequeme Körperhaltung einnehmen muss. Zunächst mag diese Lösung als reine Komfortfunktion erscheinen. Bei genauerer Betrachtung erschließt sich jedoch, dass die ohne das Fahrerassistenzsystem einzunehmende unbequeme Körperhaltung zur Beobachtung der Lichtsignalgeber auch die Aufmerksamkeit beeinträchtigt und im Falle



Figur 17: Kameraerfassung von Lichtsignalen  
(aus DE 10 2009 040 252 A1)

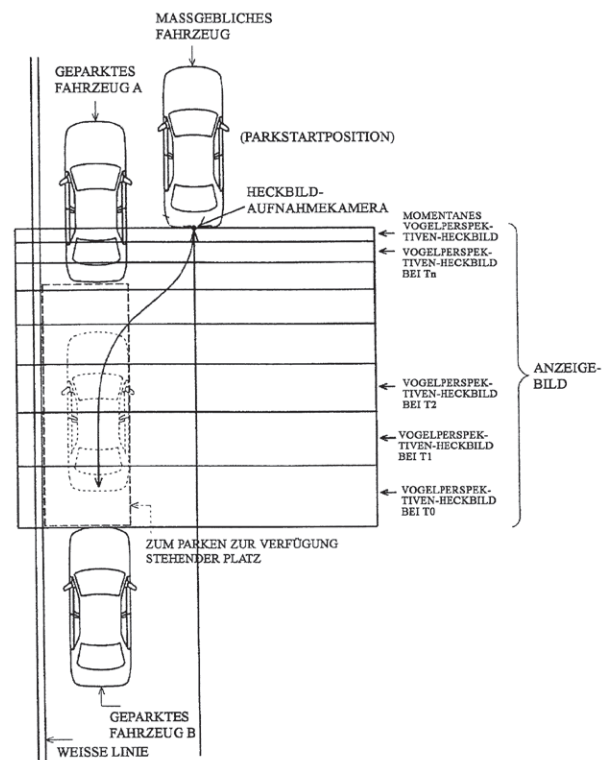
eines Crashes eine ungünstige Position mit erhöhter Verletzungsgefahr bedingen würde. Insofern ermöglicht dieses Fahrerassistenzsystem neben einer Komforterrhöhung auch eine Erhöhung der Verkehrs- und Insassensicherheit.

Die im Fahrzeug verbauten Kameras unterliegen zwangsläufig Einschränkungen hinsichtlich Einbauort, Ausrichtung und Erfassungsbereich. Diesen Einschränkungen muss sich der Konstrukteur stellen. So zeigt die Patentschrift EP 1 339 561 B1 in Figur 18 die Anordnung von mehreren Kameras auf dem Fahrzeug, wobei die Kameras auf dem Fahrzeugdach eine Fisheye- oder Spiegeloptik aufweisen, welche ein Gesichtsfeld von 360° erlauben. Damit sich die Kameras in den rele-



Figur 18: Anordnung mehrerer Umfeldkameras am Kraftfahrzeug  
(aus EP 1 339 561 B1)

vanten Winkelbereichen seitlich vom Fahrzeug nicht gegenseitig abschatten, sind sie in Längsrichtung versetzt angeordnet. Ergänzt wird die Anordnung durch weitere Kameras an Front und Heck des Fahrzeugs. Durch die versetzte Anordnung der Kameras und Überdeckung der jeweiligen Erfassungsbereiche wird eine vollständige Stereovermessung des Fahrzeugumfelds ermöglicht. Eine andere Möglichkeit zeigt die Patentschrift DE 10 2005 051 777 B4 auf (siehe Figur 19). Hierbei ist am Fahrzeug lediglich eine einzelne Heckkamera vorgesehen, die einen eingeschränkten Erfassungsbereich, vor allem in seitlicher Richtung, aufweist. Die Heckkamera erzeugt bei Vorbeifahrt an einer Parklücke unter Verwendung von Fahrzeugdynamikdaten aus mehreren Einzelaufnahmen ein zusammengesetztes Bild des Fahrzeugumfelds hinter und seitlich des Fahrzeugs und zeigt dieses als synthetische Vogelperspektivdarstellung an, was ein sicheres und komfortables Einparken in Rückwärtsfahrt ermöglicht.



Figur 19: Vogelperspektivdarstellung zusammengesetzt aus Heckkameraaufnahmen  
(aus DE 10 2005 051 777 B4)



## 5 Fazit

Bereits in der Patentliteratur aus der ersten Hälfte des 20. Jahrhunderts kann nachgewiesen werden, dass Erfinder technische Lösungen gesucht haben, um Kameras im Kraftfahrzeug zu verwenden. Es wurden Kameraanordnungen konstruiert, die beispielsweise der Unfalldokumentation, der Navigationsunterstützung oder dem Diebstahlschutz dienen sollten. Trotz der erheblichen technischen Einschränkungen der damals verfügbaren Kamerasysteme wurden neben einer manuellen Auslösung der Kameras durch eine Person auch bereits Konstruktionen mit fahrerunabhängiger Auslösung in Abhängigkeit von der Fahrsituation oder Fahrdynamik entwickelt. Mit der einsetzenden Massenmotorisierung in der zweiten Hälfte des 20. Jahrhunderts wurden Lösungen zur Umfeldbeobachtung und Verkehrsüberwachung entwickelt, beispielsweise ähnlich heutigen Dash-Cams. Für die Auslösung der Kameras wurden halb- oder vollautomatische Systeme konstruiert, die die Kameras in Abhängigkeit von Bremsbetätigung, Fahrzeugverzögerung oder Aufpralldetektion auslösen. Weitere Anordnungen betrafen den Schutz der Fahrzeuginsassen, die Umfeldüberwachung mit Kompensation des toten Winkels sowie Anzeigevorrichtungen im Kraftfahrzeug.

Für aktuelle Kraftfahrzeuge mit ihrer Vielzahl an bereits serienreifen oder in der Entwicklung zur Marktreife befindlichen Fahrerassistenzsystemen stehen elektronische Bildsensoren zur Verfügung, die die vormaligen Einschränkungen bei der Verwendung von Kameras im Kraftfahrzeug überwunden haben. Insbesondere CMOS-basierende Bildsensoren erlauben eine kostengünstige Verwendung und wegen ihrer geringen Abmessungen konstruktiv vorteilhaft die Integration von Kameras in Fahrerassistenzsystemen. Fahrerassistenzsysteme mit Bildsensoren sehen mehr, weiter und besser als der Fahrer oder die Fahrerin. Sie ermüden nicht, werden nicht unaufmerksam und ihnen entgeht nichts. Sie können Bereiche einsehen, die der fahrenden Person verborgen sind. Sie vergessen nicht, und in Kombination mit Telematiksystemen können sie sogar versuchen, einen Blick in die Zukunft oder zumindest bis hinter die nächste Kurve oder die nächste Fahrbahnkuppe zu werfen. Das alles leisten sie gleichzeitig, schneller und präziser als der Mensch.



# Der Weg zum autonomen Fahren – Bildverarbeitung zur Fahrerunterstützung

Dipl.-Phys. Hartmut Wilhelms, Patentabteilung 1.53

In jüngster Zeit hat das Thema „Autonomes Fahren“ viel Aufmerksamkeit erfahren, denn es könnte einen Umbruch im individuellen und gesellschaftlichen Umgang mit dem Automobil nach sich ziehen. Zur Fahrerunterstützung und Erfassung der „Umgebung“ spielen die Bildverarbeitung und Sensoren in Kraftfahrzeugen eine große Rolle. Der folgende Artikel gibt einen Überblick über die Bildverarbeitung, Einblick in technische Entwicklungen und einen kurzen Ausblick auf die zunehmende Verwendung von fahrspurbezogenen Kartendaten.

## 1 Einleitung

„Autonomes Fahren“ bildet einen Forschungsschwerpunkt nahezu aller großen Automobilhersteller, die ihre herkömmlich angetriebenen Fahrzeuge mit entsprechenden Sensoren ausstatten. Hinzu kommen neue Anbieter für Elektromobilität oder Internetkonzerne, welche basierend auf dem autonomen Fahren komplett neue, meist elektrisch angetriebene Fahrzeuge konstruieren [1].

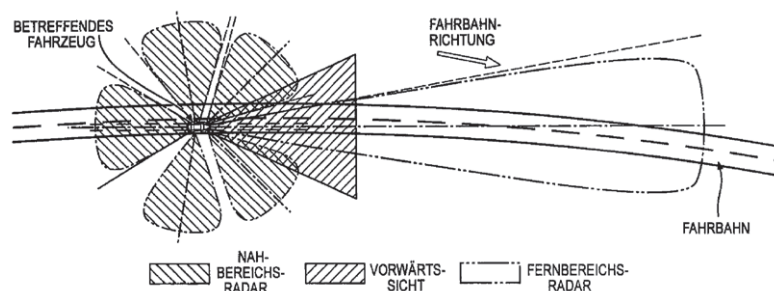
Weil dem „Autonomen Fahren“ aus rechtlichen Erwägungen auf öffentlichen Straßen noch Grenzen gesetzt sind, werden vorrangig Systeme zur Fahrerunterstützung eingesetzt, unter anderem bei:

- Geschwindigkeits- oder Abstandsregelsystemen
- Spurhalte- oder Spurwechselassistenten
- Hinderniserkennungs- und Antikollisionssystemen
- Verkehrszeichenerfassung

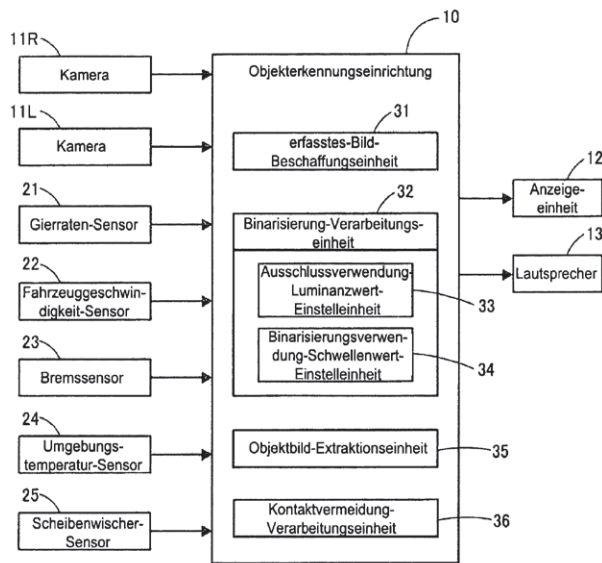
Diese Systeme nutzen gleichzeitig eine Vielzahl von Sensoren, um Informationen über die Umgebung des Fahrzeugs zu erhalten, um den Fahrer zu warnen oder unmittelbar auf das Fahrzeug einzuwirken. Die DE 10 2011 009 665 A1 zeigt

den Erfassungsbereich verschiedener Sensoren um ein Kraftfahrzeug (siehe Figur 1). Der Nah- und der Fernbereich werden dabei durch Radarsensoren und die Vorwärtssicht mit einer Kamera erfasst. Die Daten der Sensoren werden einzeln oder gemeinsam verarbeitet, um den Fahrer akustisch zu warnen oder ihm die notwendigen Informationen optisch darzustellen.

Ein solches System ist in Figur 2 dargestellt. Dabei werden die Daten der den Bereich vor dem Fahrzeug erfassenden Kameras 11R und 11L durch eine Datenverarbeitungsanlage 10 verarbeitet, um eine Kollision beziehungsweise einen Kontakt zu vermeiden 36. Dazu müssen beispielsweise vorher die Geschwindigkeit 22 und die Bilder der Sensoren erfasst 31 und vorverarbeitet 32 sowie das Objekt in den Bildern erkannt 35 werden. Diese exemplarischen Schritte werden im Folgenden etwas detaillierter gezeigt.



Figur 1: Der Erfassungsbereich verschiedener Sensoren um ein Kraftfahrzeug (aus DE 10 2011 009 665 A1)



Figur 2: Sensoreinsatz im Kraftfahrzeug zur Fahrerunterstützung (aus DE 11 2013 003 276 T5)

## 2 Bildverarbeitung

Damit die fahrende Person rechtzeitig vor einem Fußgänger gewarnt oder das Fahrzeug frühzeitig abgebremst werden kann, müssen die Bilder der verwendeten Sensoren in Sekundenbruchteilen verarbeitet werden. Angenommen, die Kamera hätte eine Auflösung von 1,4 Megapixeln (1360 x 1024). Sollte das Fahrzeug mit einer Geschwindigkeit von 180 Stundenkilometern fahren, so würde es in jeder Sekunde 50 Meter zurücklegen. Sofern alle 5 Meter eine Bildinformation benötigt würde, müssten mindestens 10 Bilder pro Sekunde erfasst und verarbeitet werden. Wenn diese Bilder mit einer Farbtiefe von 12 Bit bei 3 Farbkanälen pro Pixel ausgelesen würden, betrüge die Datenrate ungefähr 60 MByte pro Sekunde. Daher wird neben einer leistungsfähigen Hardware auch eine spezielle Bilddatenverarbeitung benötigt, um mit diesen Datenraten umgehen zu können.

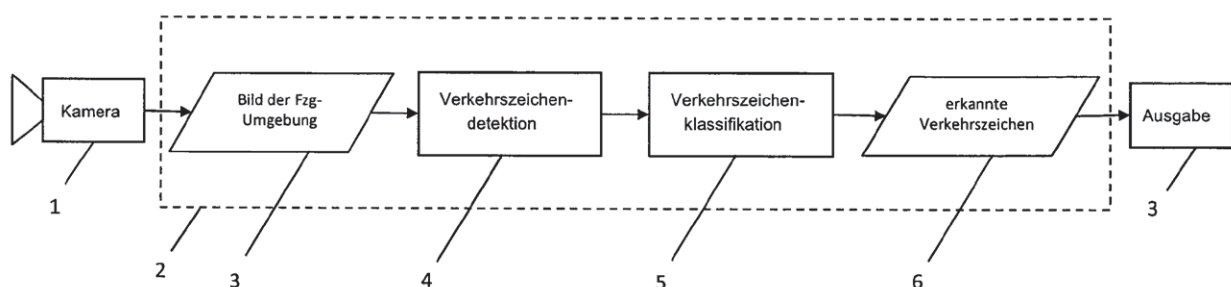
In Figur 3 ist eine Bildverarbeitungskette 2 zur Verkehrszeichenerkennung dargestellt. Dabei werden durch die Kamera 1 laufend Bilder der Umgebung des Fahrzeugs 3 aufgenommen. In jedem dieser Bilder werden nun die Bereiche detektiert 4, welche ein Verkehrszeichen enthalten können. Anschließend wird das Verkehrszeichen klassifiziert 5 und das erkannte Verkehrszeichen als Ergebnis ausgegeben.

Obwohl sich diese Bildverarbeitung je nach Anwendungsgebiet im Detail stark unterscheiden kann, lässt sie sich doch generell in mindestens zwei große Abschnitte aufteilen: „Lokalisieren des Bildbereichs“ (siehe 2.1) und „Klassifizieren des vom Bildbereich umfassten Objekts“ (siehe 2.2). Als weiterführende Lektüre sei dabei auf [2] verwiesen.

### 2.1 Lokalisierung des Objekts

Zunächst werden in jedem Bild die sogenannten „interessierenden Bereiche“ ROI („Regions of Interest“) ermittelt, in denen sich das gesuchte Objekt überhaupt befinden kann (in der Figur 3 wird dies im Rahmen der „Verkehrszeichendetektion“ 4 durchgeführt). Durch die grobe Bestimmung der ROIs kann die spätere Klassifikation erheblich beschleunigt werden, weil dann nur eine geringere Menge an Bilddaten durch den aufwändigeren Klassifikator verarbeitet werden muss.

In der Figur 4 ist das Bild einer Kreuzung dargestellt, in dem die „interessierenden Bereiche“ 500 zusammen mit einigen Kandidaten 510 dargestellt sind, welche anschließend wieder verworfen wurden. Oft werden mehrere „interessierende Bereiche“ in einem Bild identifiziert, welche sich auch überlappen können.



Figur 3: Bildverarbeitungskette zur Verkehrszeichenerkennung (aus DE 10 2009 048 066 A1)

Die Bestimmung dieser Bereiche ist jeweils vom Anwendungsgebiet abhängig und erfolgt hauptsächlich durch einfache und schnelle Algorithmen (Filter), beispielsweise durch Verstärkung bestimmter Farben (wie Rot oder Weiß bei Verkehrszeichen beziehungsweise Straßenmarkierungen), und anschließender Schwellwertbildung.

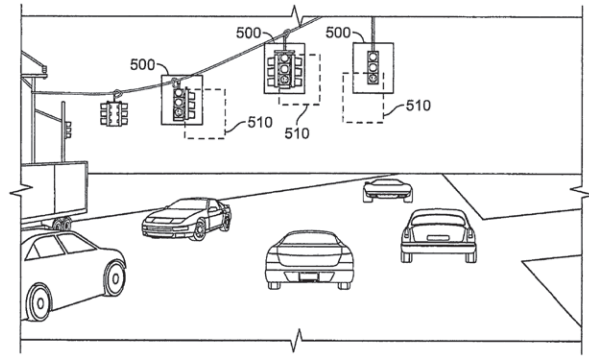
Die „interessierenden Bereiche“ können dabei, je nach gesuchtem Objekt, unterschiedliche Formen haben (rechteckig, kreisförmig, dreieckig, trapezförmig). Deshalb werden die Kandidatenbereiche bei Verkehrszeichen, Fahrzeugen oder Fahrbahnen, zusätzlich noch mit den erwarteten geometrischen Formen verglichen, um die ROIs zu finden. Hier kann auch aus dem Vorgängerbild die Kenntnis über den ungefähren Ort des gesuchten Objekts ausgenutzt werden.

## 2.2 Klassifizierung des Objekts

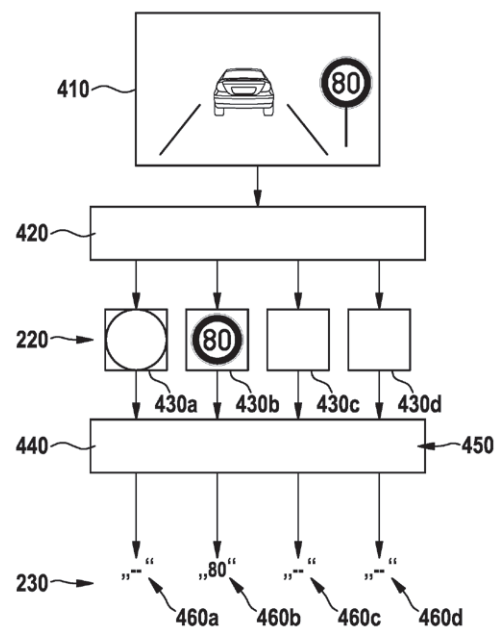
Nachdem diese „interessierenden Bereiche“ identifiziert wurden, muss das darin enthaltene Objekt jeweils einer Klasse zugeordnet werden.

In der Druckschrift DE 10 2012 206 037 A1 ist dieser Vorgang schematisch für eine Geschwindigkeitsbegrenzung dargestellt (siehe Figur 5). Dabei wird zuerst das Verkehrszeichen in dem Bild lokalisiert 420 und anschließend werden die Merkmale des Verkehrszeichens mit den Merkmalen verschiedener bekannter Geschwindigkeitsbegrenzungen 220 verglichen, um die Klasse mit der größten Übereinstimmung zu ermitteln 440, welche dann dem Verkehrszeichen zur weiteren Verwendung zugeordnet wird. Vorliegend wurde dem Verkehrszeichen eine Geschwindigkeitsbegrenzung von 80 km/h zugeordnet.

Was hier so einfach dargestellt wurde, ist natürlich viel aufwändiger. Zunächst müssen in einem ersten Schritt die Bildbereiche statistisch analysiert werden, was über die Histogramme der Gradienten möglich ist, wobei ein Gradient senkrecht zu den vorherrschenden Linien steht und dessen Länge ein Maß für die Größe der Änderung (Steilheit) angibt. Die dabei ermittelten Zahlenwerte werden als „Merkmale“ (Features) bezeichnet.



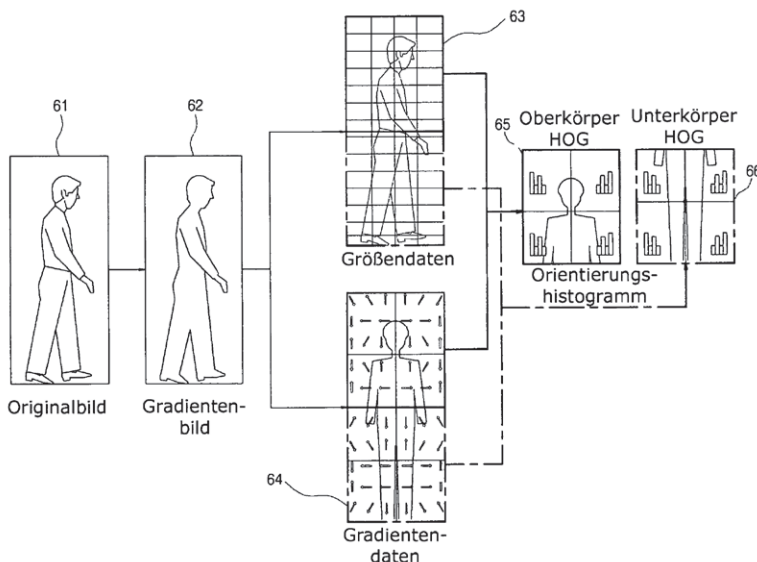
Figur 4: „Interessierende Bereiche“ für eine Bildauswertung (aus DE 10 2012 207 620 B4)



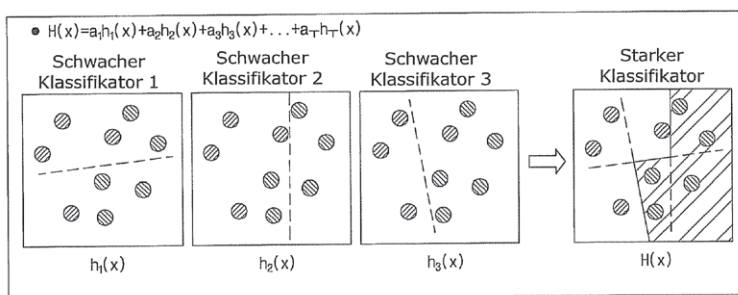
Figur 5: Beispiel für einen Verkehrszeichen-Klassifikator (aus DE 10 2012 206 037 A1)

Dieser Schritt ist in Figur 6 beispielhaft für die Erkennung eines Fußgängers dargestellt, die Technik ist jedoch in vielen weiteren Bereichen anwendbar. Zunächst wird (über eine zweidimensionale partielle Ableitung) im vorher ermittelten ROI das farbige Originalbild in ein Gradientenbild umgewandelt und anschließend in Teilbereiche unterteilt, für welche jeweils die Histogramme der Gradientenorientierungen (HOG) bestimmt werden. Die Orientierung wird dabei in 45°-Schritten bestimmt und von links nach rechts aufgetragen, wobei der jeweilige Wert (Höhe im Histogramm) von der Stärke des Gradienten in der jeweiligen Richtung abhängt.

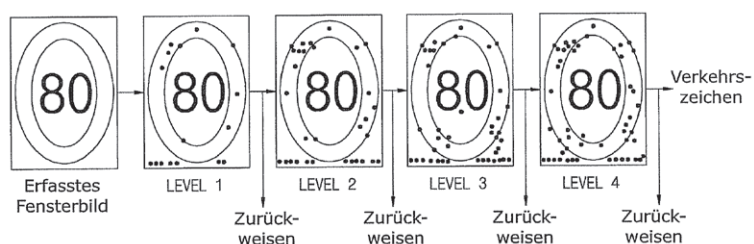
Dadurch kann jedes Teilbild mathematisch über seine vorherrschenden Kantenrichtungen beschrieben werden, wobei oft viele weitere Merkmale wie Histogramme der Farbverteilungen bestimmt und verwendet werden [3]. Die Auswahl und Anzahl der Merkmale richtet sich dabei nach den zu erkennenden Objekten.



Figur 6: Eine exemplarische Fußgängererkennung über Histogramme aus orientierten Gradienten (aus DE 10 2014 207 650 A1)



Figur 7: Kombination von schwachen Klassifikatoren zu einem starken Klassifikator (aus DE 10 2014 214 448 A1)



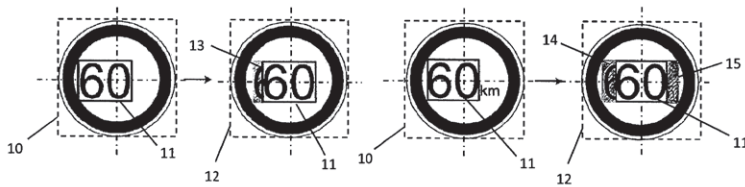
Figur 8: Klassifizierungskaskade am Beispiel eines Verkehrszeichens (aus DE 10 2014 214 448 A1)

Eine gute Übersicht zur Wahl geeigneter Merkmale und/oder Klassifikatoren ist aus der Literatur ersichtlich [4].

Die DE 10 2004 214 448 A1 beschreibt, wie die einzelnen Merkmale oder mathematische Kombinationen von mehreren Merkmalen verwendet werden, um

den Bildinhalt einer vorbekannten Klasse (zum Beispiel den unterschiedlichen Geschwindigkeitsbegrenzungen in Figur 5) zuzuordnen, also das abgebildete Objekt zu klassifizieren. In Figur 7 ist dieser Prozess der Klassifizierung über Merkmale stark vereinfacht dargestellt, wobei die gestrichelte Linie jeweils die akzeptierten von den verworfenen Merkmalen trennt. Die Kunst besteht nun darin, die Trennlinien (beziehungsweise Trennkurven) geeignet zu ermitteln und/oder eine geeignete Kombination der Merkmale zu bestimmen, wozu oft eine „Stützvektormaschine“ (*Support Vector Machine SVM*) eingesetzt wird. Es können auch mehrere Klassifikatoren kombiniert (siehe Figur 7) oder als Kaskade hintereinander (siehe Figur 8) ausgeführt werden, um das gewünschte Ergebnis zu bekommen. Dabei werden zuerst schwache Klassifikatoren angewendet (zum Beispiel ob das Verkehrsschild „genügend“ Rot aufweist), anschließend werden die Klassifikatoren immer rechenintensiver und genauer. Das Ziel bei einer solchen Kaskade besteht erneut darin, anfangs möglichst viele Kandidaten zu verwerfen und die aufwändigsten Operationen nur bei wenigen Kandidaten auszuführen.

Es gibt aber auch viele andere Klassifikatoren, welche aus den Merkmalen einen Vektor mit 100 bis 1000 Elementen bilden, und anschließend im 100- beziehungsweise 1000-dimensionalen Raum versuchen, über das Skalarprodukt den nächstkommen Klassenvektor zu bestimmen. Dazu müssen vorher die Klassenvektoren durch ein Training gelernt werden.



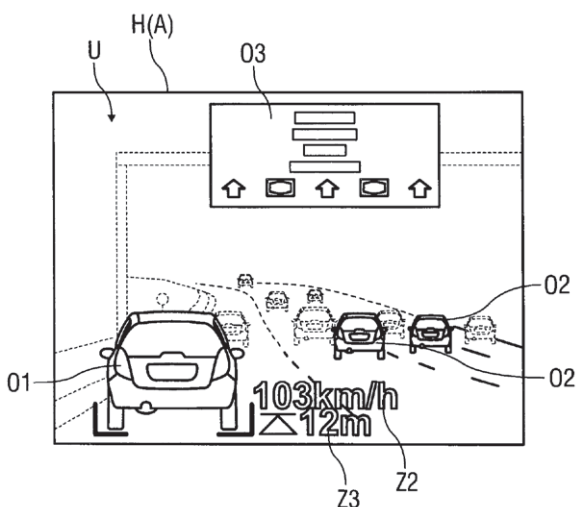
Figur 9: International unterschiedliche Versionen eines Verkehrszeichens zur Beschränkung der zulässigen Höchstgeschwindigkeit – zusammen mit dem ROI (aus DE 10 2009 048 066 A1)

Nähere Informationen sind in [5], [6] und [7] zu finden, die sich mit der Fußgängererkennung befassen. Details der Verkehrszeichenerkennung werden in [8] bis [13] und die Fahrspurerkennung wird von [14] bis [18] behandelt. Ein neuerer Ansatz zur Klassifikation mit Context Forests ist in [19] beschrieben.

### 3 Beispiele zur Erkennung und Klassifizierung

#### 3.1 Training des Klassifikators

Die DE 10 2012 206 037 A1 beschreibt das Training eines Klassifikators über eine Rückmeldung (richtig/falsch) des Beifahrers beim Vorbeifahren an einem Verkehrszeichen. In einer weiteren Ausgestaltung kann diese Trainingsarbeit auch auf viele Nutzer von Smartphones verlagert werden, so wie dies bei anderen Open-Source-Projekten oder der Open-Street-Map bereits durchgeführt wird.



Figur 10: Anzeige von Informationen auf der Windschutzscheibe (aus DE 10 2013 016 242 A1)

#### 3.2 Verkehrszeichenerkennung

Die Verkehrszeichenerkennung ist eine notwendige Voraussetzung für das autonome Fahren, damit das Fahrzeug die Gebote und Verbote selbstständig beachten kann. Bei einer Unterstützung können somit Hinweise gegeben werden, beispielsweise bei einer Überschreitung der Höchstgeschwindigkeit.

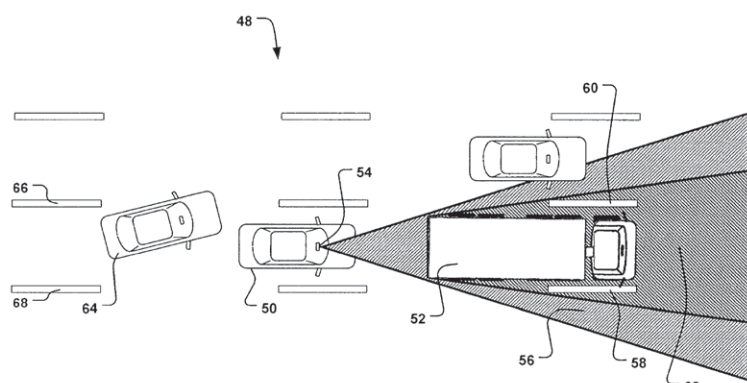
Die Anmeldung DE 10 2009 048 066 A1 beschreibt ein Verfahren, bei dem die Ziffern einer Geschwindigkeitsbegrenzung auch dann erkannt werden können, wenn diese trotz einer Vereinheitlichung gemäß dem Wiener Übereinkommen über Straßenverkehrszeichen teilweise unterschiedlich dargestellt werden (zum Beispiel mit/ohne „km“-Angabe oder in einer nicht zentrierten Darstellung des Ziffernblocks). Dabei wird der interessante Bereich (ROI) innerhalb des Verkehrszeichens dynamisch verschoben und skaliert (siehe Figur 9), um den Trainingsaufwand des Klassifikators reduzieren zu können.

Die DE 10 2011 081 456 A1 schlägt die Unterdrückung versehentlich erfasster Geschwindigkeitsbeschränkungen vor, beispielsweise wenn diese für eine andere Fahrspur gültig ist oder sich auf einem anderen Fahrzeug befindet.

Die DE 10 2014 000 264 A1 befasst sich mit der Transformation ausländischer Verkehrszeichen, wobei das durch eine Kamera erfasste Verkehrszeichen in einer internationalen Verkehrsschild-Datei erkannt und als ein bekanntes analoges deutsches Verkehrszeichen auf der Windschutzscheibe überlagert oder auf einem Display angezeigt wird. Dabei können auch Texte oder Zahlenangaben übersetzt sowie Einheiten umgerechnet werden (zum Beispiel 50ft oder 30 inch).

Die Darstellung der über die Sensoren erkannten Informationen kann nach der DE 10 2013 016 242 A1 ebenfalls auf der Windschutzscheibe oder in einer Datenbrille erfolgen (siehe Figur 10). Hier werden nicht nur die Geschwindigkeit und der Abstand eines vorausfahrenden Fahrzeugs angezeigt, sondern bedarfs-



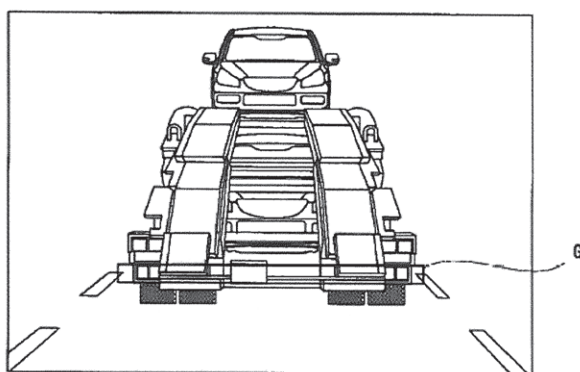


Figur 11: Verdeckte Fahrspurmarkierungen durch vorausfahrende Fahrzeuge  
(aus DE 10 2012 104 786 A1)

weise auch die Beschränkungen der einzelnen parallelen Fahrspuren. Es kann auch die Geschwindigkeit und Winkelbeschleunigung des eigenen Fahrzeugs in Verbindung mit den aus Kartendaten entnommenen Kurvenradien zu einer Schleuder-Warnung genutzt und angezeigt werden.

### 3.3 Fahrspurerkennung

Die Fahrspurerkennung ist hilfreich, um die fahrende Person vor dem (ungewollten) Verlassen der Spur zu warnen oder das Fahrzeug automatisch in dieser Spur zu halten. Zudem ist die Fahrspurerkennung eine nützliche Information für die Erkennung des Fahrbahnzustands oder von Hindernissen (zum Beispiel Fußgänger) in der Fahrbahn, denn dadurch kann der „interessierende Bereich“ (ROI) für die weitere Bildverarbeitung auf die Fahrbahn oder dessen Rand eingeschränkt werden.



Figur 12: Ein auf der Straße abgestellter Autotransporter  
(aus DE 10 2014 115 017 A1)

Üblicherweise wird eine Fahrbahn in einem Bereich unmittelbar vor dem Fahrzeug über ihre parallel in einem meist konstanten Abstand voneinander verlaufenden Markierungen erkannt, die unterbrochen oder durchlaufend sein können. Deshalb wird in erster Näherung ein trapezförmiger Bereich der Straße in Front des Fahrzeugs auf zusammenhängende Bereiche mit hellen Pixeln untersucht, die idealerweise auf geraden Linien verlaufen. Im Falle von Kurven wird ähnlich einem Kurvenlineal über Splines

versucht, diese Kurven mathematisch nachzubilden (siehe auch Figur 14). Zusätzlich werden Informationen von genauen Kartendaten des Navigationssystems verwendet, um den Verlauf der Fahrbahn vorhersagen zu können (DE 10 2010 018 333 A1).

Diese Informationen zu Anzahl und Breite der Fahrbahnen oder deren Krümmungsdaten und/oder Kurvenradien werden auch in der DE 10 2012 104 786 A1 zur Vorhersage der Fahrspur verwendet, wobei die Vorhersage noch zusätzlich durch eine Erkennung und Verfolgung (Tracking) des vorausfahrenden Fahrzeugs unterstützt wird. Die verschiedenen Vorhersagen sollen anschließend durch einen „Fusion“-Prozessor gewichtet werden, um auch im Fall der Verdeckung der Fahrspur durch ein vorausfahrendes Fahrzeug, wie in Figur 11 gezeigt, die Fahrspur sicher erkennen zu können (oder wenn das Führungsfahrzeug zum Überholen ausschert). Zusätzlich wird auch vorgeschlagen, über eine Fahrzeug-zu-Fahrzeug (V2V) Kommunikation die Informationen anderer Fahrzeuge zu nutzen.

Die DE 10 2014 115 017 A1 beschäftigt sich ebenfalls mit der Fahrspurerkennung und beschreibt dabei auch die Erkennung von außergewöhnlichen Situationen, beispielsweise das Ende der nutzbaren Fahrbahn hinter einem abgestellten Autotransporter, so wie dies in der Figur 12 dargestellt ist.





- [10] LIANG, Ming [et al.]: Traffic Sign Detection by ROI Extraction and Histogram Features-based Recognition. In: The 2013 International Joint Conference on Neural Networks (IJCNN), IEEE, 2013, S. 1–8
- [11] BARÓ, Xavier [et al.]: Traffic Sign Recognition Using Evolutionary Adaboost Detection and Forest-ECOC classification. IEEE Transactions on Intelligent Transportation Systems, 2009, 10. Jg., Nr. 1, S. 113–126
- [12] MØGELMOSE, Andreas [et al.]: Vision-Based Traffic Sign Detection and Analysis for Intelligent Driver Assistance Systems: Perspectives and Survey. In: IEEE Transactions on Intelligent Transportation Systems, 2012, 13. Jg., Nr. 4, S. 1484–1497
- [13] GUDIGAR, Anjan; CHOKKADI, Shreesh; RAGHAVENDRA, U. A review on automatic detection and recognition of traffic sign. In: Multimedia Tools and Applications, Springer, 2014, S. 1–32
- [14] THONGPAN, Narathip; KETCHAM, Mahasak: The State of the art in Development a Lane Detection for Embedded System Design. International Conference on Advanced Computational Technology and Creative Media (ICACTCM 2014)
- [15] DENG, Jiayong; HAN, Youngjoon: A real-time system of lane detection and tracking based on optimized RANSAC B-spline fitting. In: Proceedings of the 2013 Research in Adaptive and Convergent Systems. ACM, 2013, S. 157–164
- [16] MCCALL, Joel C.; TRIVEDI, Mohan M. Video-Based Lane Estimation and Tracking for Driver Assistance: Survey, System, and Evaluation. In: IEEE Transactions on Intelligent Transportation Systems, 2006, 7. Jg., Nr. 1, S. 20–37
- [17] WANG, Yifei; DAHNOUN, Naim; ACHIM, Alin: A novel system for robust lane detection and tracking. In: Signal Processing, Elsevier 2012, 92. Jg., Nr. 2, S. 319–334. ISSN 0165-1684
- [18] YIM, Young Uk; OH, Se-Young. Three-Feature Based Automatic Lane Detection Algorithm (TFAL-DA) for Autonomous Driving. In: IEEE Transactions on Intelligent Transportation Systems, 2003, 4. Jg., Nr. 4, S. 219–225
- [19] MODOLO, Davide; VEZHNEVETS, Alexander; FERRARI, Vittorio: Context Forest for efficient object detection with large mixture models. arXiv preprint arXiv:1503.00787, 2015

# Datenbrillen

*Dipl.-Ing. Martin Bässler, Patentabteilung 1.53*

Brillen werden als Sehhilfen oder Schutzbrillen verwendet. Sie eignen sich aber auch als Träger zusätzlicher Funktionen. Zu den Brillen mit zusätzlichen Funktionen gehören auch die Datenbrillen. Die Anforderungen an Datenbrillen haben sich über die letzten vier Jahrzehnte hinweg nicht wesentlich geändert. Die Datenbrille soll leichtgewichtig, handlich und stabil sein. In diesem Artikel werden Hörbrillen, Augmented-Reality-Brillen und Virtual-Reality-Brillen vorgestellt. Anhand von verschiedenen Patentanmeldungen wird gezeigt, wie die der Konstruktion zugrunde liegenden Prinzipien aussehen, und welche Anwendungsgebiete die jeweiligen Datenbrillen haben.

## 1 Hörbrillen

### 1.1 Die Optimierung von Leitungen

Bei einer Hörbrille wird ein Hörgerät in eine normal funktionierende Brille integriert. Ein Beispiel zeigt die Patentschrift DE 23 60 342 C2 (1973). Im linken Brillenbügel 2, siehe Figur 1, werden Schallsignale mit einem Mikrofon 18 aufgenommen, verstärkt und an die Induktionsspule 20 abgegeben. Von dort wird ein Induktionsfeld aufgebaut, aus welchem die Spule 9 im rechten Bügel die Signale aufnimmt und über den Verstärker 7 und einen Lautstärkeregel 8 dem Lautsprecher 6 zuführt. Der Lautsprecher ist mit einer Schallableitung 13 verbunden, die elastisch in den Gehörgang des Ohres einsetzbar ist.

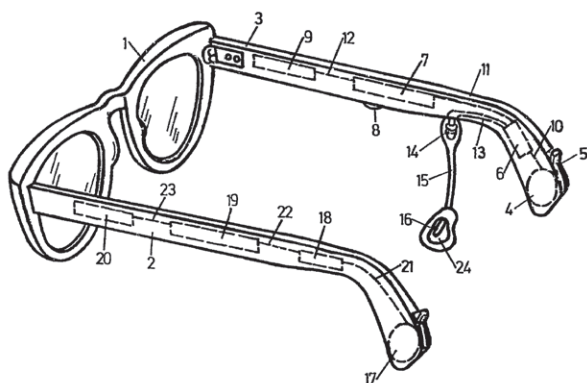
Die grundsätzlichen Herausforderungen, die eine Integration von Geräten in eine Brille betreffen, sind bereits

in dieser Patentschrift angedeutet. Die Leitungen 10, 11, 12, 21, 22 und 23 sollen in unauffälliger Weise über den Brillenrahmen geführt, die Anzahl der Leitungen minimiert, der Verlauf der Leitungen so geschaffen werden, dass sie im Gebrauch nicht gebogen werden und störanfällige Kontakte vermieden werden. Die Lösung für diese Herausforderungen ist in dem beschriebenen Beispiel die drahtlose Verbindung zwischen dem Mikrofon auf der einen Seite und dem Lautsprecher auf der anderen Seite des Kopfes.

### 1.2 Die Austauschbarkeit von Komponenten

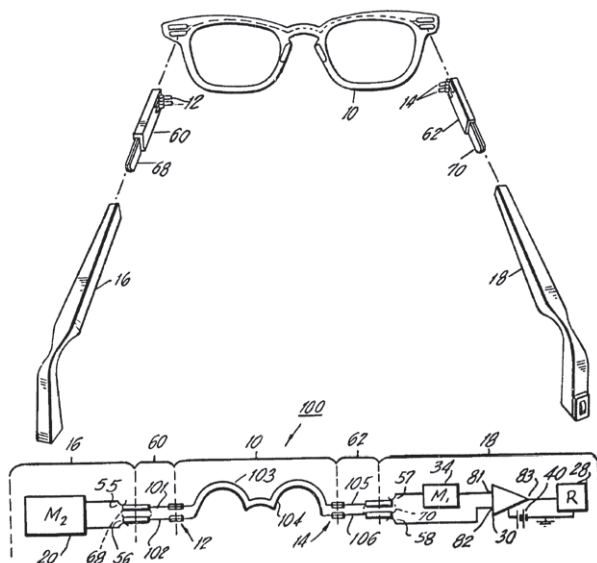
Eine weitere Aufgabe bei Hörbrillen ist, Komponenten möglichst flexibel zu gestalten, um auf veränderte Anforderungen der Nutzer ohne aufwändige Anpassungen reagieren zu können. Auch diese generelle Anforderung an Datenbrillen ist in der Schrift DE 23 04 134 B (1973) bereits formuliert.

Ist es erwünscht, die Hörhilfe sowohl für eine Seite als auch für beide Seiten zu nutzen, wird in Figur 2 folgende Lösung vorgeschlagen: In jedem Bügel ist jeweils ein Mikrofon vorhanden. Das Mikrofon M1 im rechten Brillenbügel 18 und das Mikrofon M2 im linken Brillenbügel 16 sind in Reihe geschaltet. Falls jedoch nur für eine Seite eine Gehörverstärkung gewünscht wird, so kann der andere Brillenbügel durch einen herkömmlichen Metallabschnitt ohne elektrische Bauteile flexibel ersetzt werden.



Figur 1: Hörgerät integriert in einer Brille (aus DE 23 60 342 C2)





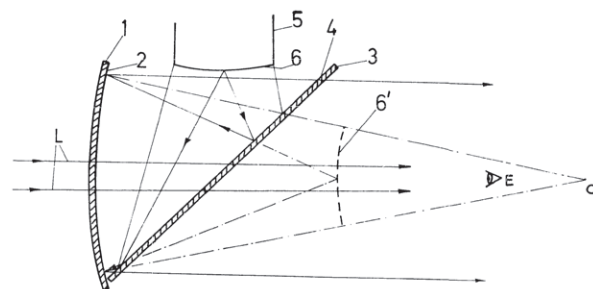
Figur 2: Hörgerät integriert in einer Brille mit auswechselbaren Brillenbügeln (aus DE 23 04 134 B)

## 2 Bildapparat-Einblendungen für Pilotinnen und Piloten

Bei vielen Datenbrillen werden externe Informationen in das Gesichtsfeld des Nutzers eingespielt. Die grundsätzliche Vorgehensweise wird anhand der Patentschrift DE 20 09 798 C3 (1970) erläutert. Es handelt sich dabei um das Einspielen von Informationen in das Blickfeld eines Piloten oder einer Pilotin während des Fluges, wobei die fliegende Person in die Lage versetzt wird, diese Informationen wahrzunehmen, ohne den Blick von der Szenerie außerhalb des Flugzeuges abzuwenden.

Die externen Informationen werden dazu auf einem Bildschirm einer Kathodenstrahlröhre präsentiert, die außerhalb des Blickfeldes der fliegenden Person liegt. Der Pilot oder die Pilotin muss jedoch die weit entfernte äußere Szenerie und den nur wenige Zentimeter entfernten Bildschirm gleichzeitig scharf sehen. Dazu muss aber die Brennweite des Auges auf unterschiedliche Entfernungen angepasst (akkommodiert) werden. Da das Auge zum Erkennen der äußeren Szenerie fern-akkommodiert sein muss, muss der Bildschirm so eingespielt werden, dass er ebenfalls vom fern-akkommodierten Auge scharf gesehen werden kann. Dazu müssen die Lichtstrahlen, die von einem Punkt des Bildschirms ausgehen, parallel zueinander am Auge ankommen. Das wird folgendermaßen gelöst:

Die Figur 3 zeigt einen ersten teildurchlässigen Reflektor 1 mit einer kugelförmig konkaven Reflexionsfläche 2. Der Krümmungsmittelpunkt C der Reflexionsfläche fällt annähernd mit dem Zentrum des Kopfs des Beobachters, und zwar in Höhe der Augen E, zusammen. Er sieht das von außen kommende Licht der äußeren Szenerie L durch diesen und einen zweiten teildurchlässigen Reflektor 3. Der zweite Reflektor liegt zwischen den Augen und dem Teilreflektor 1. Er hat eine ebene Oberfläche und ist geneigt. Oberhalb der Teilreflexionsfläche 3 ist eine Kathodenstrahlröhre 5 angeordnet, die einen kugelförmig konvexen Bildschirm 6 aufweist. Der Krümmungsradius des Bildschirms ist halb so groß wie derjenige der Reflexionsfläche 2. Der Bildschirm 5 ist so angeordnet, dass ein virtuelles Bild 6' infolge der Reflexion durch die ebene Reflexionsfläche 4 in der Brennebene der kugelförmigen Reflexionsfläche 2 erzeugt wird. Bei Bildern, die in der Brennebene einer kugelförmigen Reflexionsfläche 2 erzeugt werden, verlaufen die von der kugelförmigen Reflexionsfläche reflektierten Strahlen jedes Bildpunkts parallel [1]. Somit wird infolge dieser Stellung des virtuellen Bildes 6' das Licht, das von jedem Punkt des Bildschirms 6 ausgeht und von der Fläche des zweiten Reflektors reflektiert wird, dann von der kugelförmigen Fläche 2 als paralleler Strahl reflektiert. Der Betrachter kann somit ein Bild erkennen, weil das Auge aus den parallelen Strahlen jedes Bildpunkts ein reales Bild auf der Netzhaut erzeugen kann. Das Auge muss dazu auf Unendlich akkommodiert sein [1].



Figur 3: Schematische Darstellung eines Über-Kopf-Bildapparates (aus DE 20 09 798 C3)

## 3 Videobrillen

Die Gebrauchsmusterschrift DE 92 17 643 U1 (1992) beschreibt die Einspiegelung von zwei- oder dreidimensionalen Videobildern.

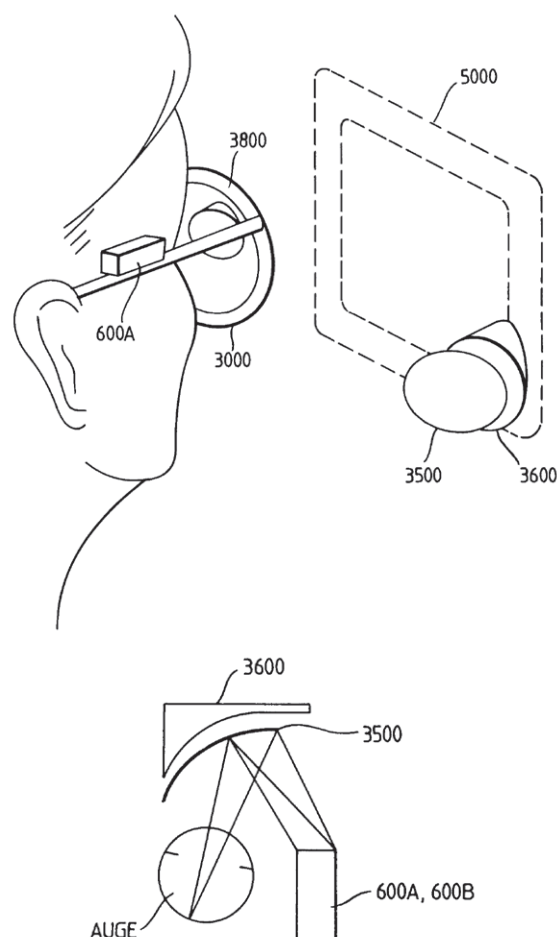
### 3.1 Die Verwendung neuer Technologien von Anzeigeeinrichtungen

Bei der Beschreibung des Standes der Technik wird gezeigt, dass die Auflösung der Videobilder, die mit einer Kamera aufgenommen und auf analogen Kathodenstrahlröhren angezeigt werden, durch die Geschwindigkeit der Erzeugung des Videobildes in den Kathodenstrahlröhren begrenzt wird. Im Gegensatz dazu besitzen neuere Anzeigeeinrichtungen wie Flachbildschirme (flat panel displays), Leuchtdioden (LEDs), *Elektrolumineszenzanzeigen*, Flüssigkristalldisplays, Plasmaanzeigen und kleine Kathodenstrahlröhren-Videodisplays (small video CRT tubes) eine hohe Taktrate und räumliche Auflösung. Viele dieser Anzeigeeinrichtungen sind für digitale Bilder ausgelegt. Bei der Digitalisierung von Bildern wird zusätzlich die Abtastrate und somit die Auflösung der Videobilder erhöht und an die hohe Auflösung der neueren Anzeigeeinrichtungen angepasst. Zusätzlich lassen sich bei Verwendung von Komparatoren und digitalen Filtern auch die Helligkeits- und Kontrastwerte im Ausgangsvideobild verändern. Die weiteren Vorteile von beispielsweise Flachbildanzeigen sind ihr geringes Gewicht und ihre dünne kompakte Form, wodurch sie leichter in eine Brille eingebaut werden können.

### 3.2 Die Erzeugung virtueller Bilder

Die geschilderten neuen Technologien ermöglichen eine vereinfachte und kompaktere Ausführung der Brillen. Die Figur 4 zeigt für beide Augen je einen Hohlspiegel 3500, der auf einer Befestigung 3600 drehbar gelagert ist. Die Befestigungen 3600 sind jeweils im oberen Teil 3800 der Brille auf die hintere Oberfläche der Brillengläser geklebt. Außerhalb des Blickfeldes des Betrachters ist für jedes Auge je eine Anzeigeeinrichtung 600A und 600B zwischen dem Hohlspiegel und dem Brennpunkt des Hohlspiegels angeordnet. Die Augen erzeugen auf der Netzhaut aus den von den jeweiligen Hohlspiegeln reflektierten Strahlen jedes Bildpunkts der jeweiligen Anzeigeeinrichtung 600A und 600B ein virtuelles (hinter dem Hohlspiegel), aufrechtes und vergrößertes Bild [1]. Das virtuelle Bild 5000 scheint für den Betrachter ungefähr 35 bis

45 Zentimeter vor dem Gesicht des Betrachters im Raum zu schweben. Es liegt also nicht im Unendlichen wie in Figur 3. Wenn den beiden Augen unterschiedliche Bilder präsentiert werden, kann das Gesamtbild auch dreidimensional erscheinen. Die Lage des Hohlspiegels 3500 ermöglicht dem Betrachter, durch den Rest der Brillengläser zu schauen.



Figur 4: Bilderzeugungssystem integriert in einer Brille (aus DE 92 17 643 U1)

### 3.3 Die Erzeugung virtueller Bilder bei verspiegelten Brillen (Virtual-Reality-Brille)

Während in Figur 4 der Betrachter noch die reale Umgebung anschauen kann, kommt nun die gesamte wahrgenommene Information aus externen Quellen (virtuelle Welt). Die gesamte Brillenfläche ist verspiegelt und wird als Reflektor genutzt. Es entsteht auf jeder Netzhaut ein virtuelles (hinter dem Reflektor), aufrechtes und vergrößertes Bild der Anzeigeeinrichtung 600A und 600B.

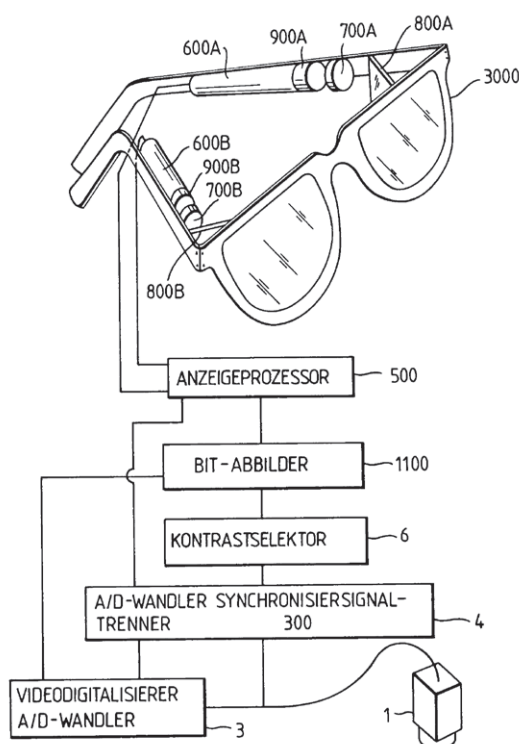
Die Figur 5 zeigt zwei hochauflösende Anzeigeeinrichtungen 600A und 600B, die Videobilder anzeigen, die von einer außerhalb der Brille angeordneten Videokamera 1 aufgenommen werden. Sie werden dann digitalisiert und mit einer höheren Abtastrate versehen 3 und 4, gegebenenfalls in Helligkeits- und Kontrastwerten verändert 6 und von einem Anzeigeprozessor 500 berechnet. In diesem Beispiel werden für die Anzeigeeinrichtungen Kathodenstrahlröhren-Video-displays verwendet. Da Kathodenstrahlröhren-Video-displays mit analogen Videobildern arbeiten, werden die Videobilder nach dem Anzeigeprozessor 500 noch analog gewandelt. Die Bilder auf der Anzeige 600A und 600B werden über eine jeweilige Projektoroptik 700A und 700B und Spiegel 800A und 800B auf den Reflektor 3000 projiziert. Das Auge erzeugt aus den vom Reflektor reflektierten Strahlen ein virtuelles, aufrechtes und vergrößertes Bild.

Im Gegensatz zum Hohlspiegel mit einer einheitlichen kugeligen Wölbung besitzt der Reflektor asymmetrische Krümmungen, die zu außeraxialen Verzerrungen (off axis) führen, die korrigiert werden müssen. Es gibt drei verschiedene Möglichkeiten der Korrektur der Ver-

zerrungen im Reflektor. In einer ersten Ausführung wird, wie in der Figur 5 beschrieben, ein kohärentes optisches Faserbündel 900A und 900B zwischen den Anzeigevorrichtungen und dem Reflektor positioniert. Das kohärente optische Faserbündel erzeugt durch eine asymmetrische Vergrößerung die notwendige Verzerrung, die von der Verzerrung des Reflektors aufgehoben wird. In einer zweiten Ausführung kann das kohärente optische Faserbündel 900A und 900B entfernt werden. In diesem Fall stellt der **Bit-Abbilder (remapper)** 1100 das Bild so ein, dass eine Verzerrung erzeugt wird, die von der Verzerrung des Reflektors aufgehoben wird. In einer dritten Ausführung wird ein Spektrometer in der **Fastie-Elbert-Konfiguration** verwendet, um eine zum Reflektor gegenläufig außer-axiale Konfiguration (off axis) des Bildes zu erzeugen. Einzelheiten können der Literatur entnommen werden ([2] und [3]).

#### 4 Autarke Datenbrillen

In der Patentschrift DE 44 36 528 C2 (1994) sind zwei Ausführungsformen einer Brille gezeigt, um ein auf einer Anzeigeeinrichtung 2 dargestelltes Bild zu betrachten (siehe Figur 6). Bei der ersten Ausführungsform ist die Anzeigeeinrichtung in der Brille so angeordnet, dass eine direkte Sicht auf die Anzeigeeinrichtung im unteren Bereich des Gesichtsfeldes möglich ist. Die Anzeigeeinrichtung ist hohlspiegelartig gewölbt. Die zweite Ausführungsform benützt die auch in der Figur 4 gezeigte Projektion des von der Anzeigeeinrichtung abgegebenen Bildes 2 über einen an einer Halterung 1 befestigten Hohlspiegel 4 auf das Auge. Wie bereits in Kapitel 3.2 beschrieben, nimmt hierbei der Benutzer ein virtuelles Bild wahr, das vergrößert ist. Im Gegensatz zur Figur 4 ist die Anzeigeeinrichtung 2 in die verstärkte Querstrebe 6 der Halterung 1 integriert und nicht am Brillenbügel befestigt. Ferner ist die Bildfläche senkrecht nach unten ausgerichtet. Die zweite Ausführungsform wird gegenüber der ersten Ausführungsform bevorzugt. Denn durch die Verwendung des Hohlspiegels 4 wird eine erleichterte Anbringung der Anzeigeeinrichtung 2 geschaffen, da das schwere Display 2 nach oben auf die Halterung 1 verlegt werden kann. Weiterhin wird erreicht, dass



Figur 5: Verspiegelte Brille – Virtual-Reality-Brille  
(aus DE 92 17 643 U1)

der Benutzer das virtuelle Bild durch die Nutzung des Hohlspiegels 4 vergrößert wahrnimmt. Bei beiden Ausführungsformen erhält der Benutzer sowohl eine Ansicht der realen Umgebung 5 als auch eine Ansicht des Bildes auf der Anzeigeeinrichtung 2. So kann er sich durch Heben oder Senken des Blickes entweder auf die Bilder in der Anzeigeeinrichtung oder auf die Sicht in die reale Umgebung konzentrieren. Brillengläser 5 können in die Halterung 1 integriert werden. Durch die hohlspiegelartige Wölbung der Anzeigeeinrichtung 2 in der ersten Ausführungsform oder durch die hohlspiegelartig gewölbte Fläche des Spiegels 4 in der zweiten Ausführungsform entsteht beim Benutzer der Eindruck eines Bildes, das eine gewisse Distanz vom Auge entfernt ist. Je nach Wölbung der Fläche kann dieser Abstand zwischen 20 Zentimetern und fünf Metern betragen. In der Regel wird die Wölbung so gestaltet sein, dass das Bild in einer Entfernung von ein bis zwei Metern erscheint.

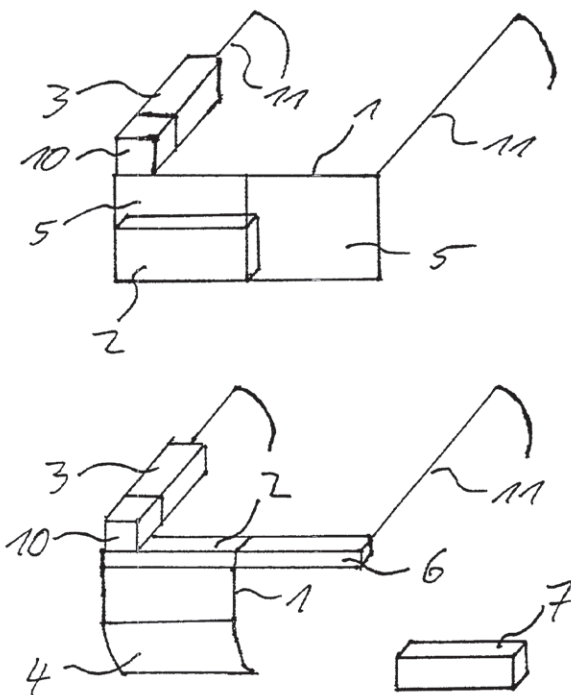
Weiterhin ist eine Miniaturkamera 3 mit Abmessungen von wenigen Zentimetern mit einem austauschbaren Objektiv 10 am Brillenbügel 11 möglichst zentral über dem Auge befestigt. Die damit aufgenommenen Bilder werden über eine Steuervorrichtung digitalisiert, die Helligkeits- und Kontrastwerte gegebenenfalls geän-

dert und an die Anzeigevorrichtung 2 weitergegeben. Während in den Figuren 4 und 5 die Kamera und die Steuervorrichtung außerhalb der Brille angeordnet sind, kann die Steuerung in der Figur 6 baulich nun auch in die auf der Brille angeordnete Videokamera 3, die verstärkte Querstrebe 6 oder in die Anzeigevorrichtung 2 integriert sein. Durch die Befestigung der Videokamera an der Halterung 1 folgt die Videokamera in ihrer Blickrichtung den Kopfbewegungen des Benutzers. Dadurch erhält der Benutzer neben einer normalen Ansicht eines anvisierten Objekts in der realen Umgebung mit Hilfe der vergrößernden Optik des Objektivs 10 zusätzlich eine vergrößerte Ansicht des anvisierten Objekts. Die in die Brille integrierte Steuerung empfängt ebenfalls die Daten einer zweiten Videokamera 7 und kann auch diese Daten auf der Anzeigeeinrichtung 2 ausgeben. Die zweite Kamera, die mit der Anzeigeeinrichtung über Funk- oder Fernsteuerung in Datenverbindung steht, kann zum Beispiel in einem Endoskop für Bauchoperationen oder in einem Gerät zur interoralen Untersuchung angeordnet sein.

## 5 Datenbrillen mit integrierter Kommunikationsschnittstelle

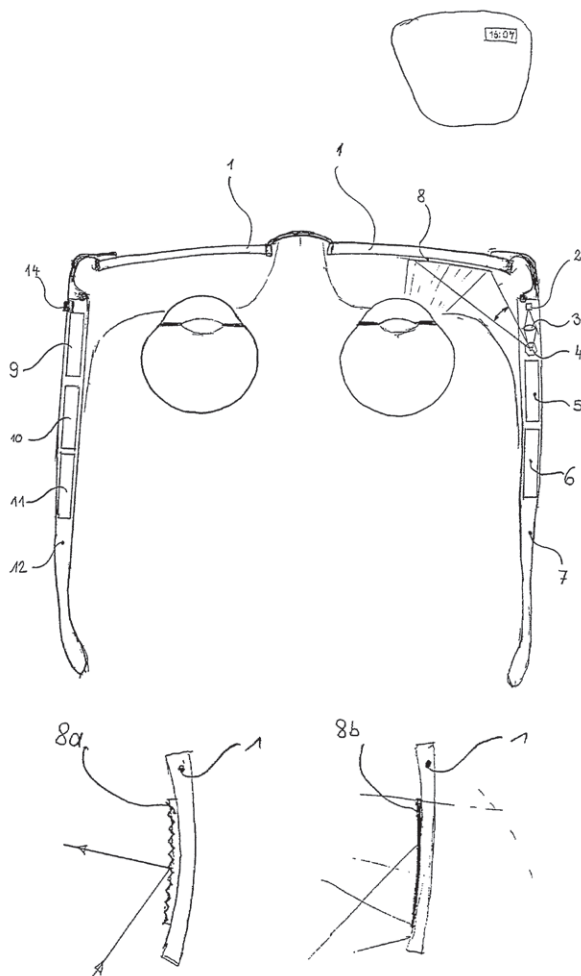
### 5.1 Anzeigeeinheit für Handheld-Computer, Eingabemöglichkeit über Augensensor

In Figur 7 ist aus der DE 196 25 435 A1 (1996) eine Datenbrille gezeigt, die neben der Bildwiedergabeeinheit auch viele weitere Funktionen integriert. Die Funktionen betreffen neben der Uhrzeit, Datum und Personenruf 6 auch die Kommunikation 9 mit einem Handheld-Computer, einem intelligenten Mobilfunkgerät, das neben der Telefonkommunikation auch die Funktionen einer Mailbox, einer Datei- und Terminverwaltung und einer Online-Kommunikation übernimmt, oder einem intelligenten, uhrenartigen Terminal am Handgelenk. Der Zweck dieser Kommunikation ist die Anzeige in einer komfortableren, vergrößerten Form (siehe Figur 7 am Beispiel der Anzeige der Uhrzeit 16:04 Uhr). Die Übertragungsstrecke kann kabelgebunden oder kabellos sein. Zur Stromversorgung werden Batterien, Akkus 11 oder Solarzellen verwendet. Auch weitere Eingabemöglichkeiten für die Datenbrille sind



Figur 6: Frontansicht autarker Datenbrillen (aus DE 44 36 528 C2)

beschrieben. So wird zum Beispiel die Augenstellung des Benutzers ermittelt. Mit ihr lässt sich eine Eingabe von Daten dadurch steuern, dass der Blick auf einen im Display dargestellten Buchstaben, auf ein Zeichen, Namen oder grafisches Icon interpretiert wird. Die Eingabe kann auch durch zweimalige Lidbewegungen, Drücken eines Knopfes 14 am Brillenbügel 12 oder am Handgelenk-Computer validiert werden. Das ausgesendete Licht einer Leuchtdiodenzeile 2 wird über eine Sammellinse oder -optik 3 und über eine Spiegelanordnung (prismatischer Drehspegel, Schwingspiegel, Kippspiegel-Array) 4 so abgelenkt, dass das Bild auf die augenseitige Oberfläche des Brillenglases 1 projiziert wird. Die Breite des Bildes auf dem Brillenglas wird durch den einstellbaren Schwenkwinkel der Spiegelanordnung 4 bestimmt. Bei der Reflexion des Lichts ins Auge wird für diese Breite des Bildes eine flächenhafte Sammeloptik 8 auf der augenseitigen Oberfläche der Brillengläser verwendet.



Figur 7: Gesamtaufbau einer Datenbrille (aus DE 196 25 435 A1)

Für die Ausbildung der flächenhaften Sammeloptik sind mehrere Alternativen denkbar. Die erste Alternative ist die Nutzung eines Hohlspiegels, der auch aus der Figur 4 bekannt ist, wobei hier der Hohlspiegel in das Brillenglas eingeschliffen wird. Die zweite Alternative ist das Aufkleben oder Einpressen einer Stufenlinse (**Fresnel-Linse**) mit hinterlegter Spiegel folie 8a. Die dritte Alternative ist das Aufkleben oder Aufdampfen einer Optikvorrichtung, die auch bei Hologrammen Verwendung findet, mit oder auch ohne hinterlegter Spiegelfolie 8b.

## 5.2 Zusätzliche Integration von Audioschnittstellen

Die DE 198 44 301 A1 (1998) beschreibt eine Datenbrille, die in einer kontaktbehafteten oder kontaktlosen Verbindung mit einer Sende- und Empfangseinheit, beispielsweise einem UMTS-Mobilfunkgerät, steht. Dabei sind in der Datenbrille zusätzlich zu der Bildwiedergabeeinheit ein Mikrofon und ein Lautsprecher für den individuellen Tonempfang des Nutzers integriert, die mit der Sende- und Empfangseinheit des UMTS-Mobilfunkgeräts verbunden sind.

## 5.3 Virtual Reality-Geräte

In der DE 198 44 301 A1 werden die sogenannten Head Mounted Devices (HMD) vorgestellt. Head Mounted Devices stellen die Anzeigeeinheiten für spezielle Simulationszwecke oder Virtual Reality dar und ermöglichen eine Bilddarstellung in einer für den Betrachter angenehmen Größe. Das Prinzip der Bildgebungseinheit ist hierbei in der Figur 5 gezeigt. Die Augen erzeugen auf der Netzhaut aus den von den verspiegelten Brillenflächen reflektierten Strahlen jedes Bildpunkts der Anzeigeeinrichtung ein virtuelles, aufrechtes und vergrößertes zwei- oder dreidimensionales Bild. Head Mounted Devices haben ein hohes Gewicht, so dass die Anzeigeeinheit nur auf einem den Kopf umschließenden Gestell mit Gurten, auf einem Helm oder an einem größeren Kopfhörer befestigt werden kann. Um das Gewicht des Head Mounted Devices nicht unnötig zu erhöhen, wird die Aufbereitung der anzuzeigenden Videobilder, wie auch in der Figur 5



gezeigt, in den meisten Modellen außerhalb des Head Mounted Devices in einer Datenverarbeitungsanlage durchgeführt. Auch kann die Bildgebungseinheit vom Umgebungslicht abgeschirmt werden.

## 5.4 Augmented Reality-Geräte

Augmented-Reality nennt man die Erweiterung (englisch: Augmentation) der physischen Welt durch optisch überlagerte Informationen aus dem Computer (überlagerte Information). Insofern sind die in den Kapiteln 3.2, 4 und 5.1 vorgestellten Datenbrillen Augmented-Reality-Brillen. In der in Kapitel 5.2 erwähnten Druckschrift sind sowohl Viewfinder, eine Weiterentwicklung der Head Mounted Devices [4] als auch zwei weitere Ausführungsformen von Datenbrillen beschrieben.

## 6 Optische Komponenten

### 6.1 Anforderungen an die Bildgebungseinheit

In der in Kapitel 5.2 erwähnten Druckschrift sind auch Maßnahmen beschrieben, den Displays optische Komponenten vorzuschalten. Die optischen Komponenten dienen zum einen einer effizienten Lichtausbeute und zum anderen der Vergrößerung des Bildes. So wird zwischen einem **Elektrolumineszenz-Display** und einer Linse eine Glasplatte angeordnet, die aus parallel dicht benachbart angeordneten Glasfasern besteht. Bedingt durch einen angepassten Brechungsindex und durch die Parallelisierung des austretenden Lichts erhöht sich die Lichtausbeute der elektrolumineszierenden Schicht. So ist eine um den Faktor zwei- bis drei höhere Leuchtdichte am Auge im Vergleich zu einem herkömmlichen Aufbau feststellbar. Außerdem führt diese Anordnung zu einer besseren Abschirmung gegen Fremd- und Streulicht und zu einem robusten, vollverklebten Aufbau, der nur gering anfällig ist gegen Bruch und Dejustage. Darüber hinaus sind **Elektrolumineszenz-Displays** in einem größeren Temperaturbereich einsetzbar, robust in ihrer Handhabbarkeit und flach ausgebildet. Jedoch ist es mit dieser Anordnung nicht gelungen, eine leistungsarme und billige Anordnung zu schaffen, da **Elektrolumineszenz-Displays** selbst-

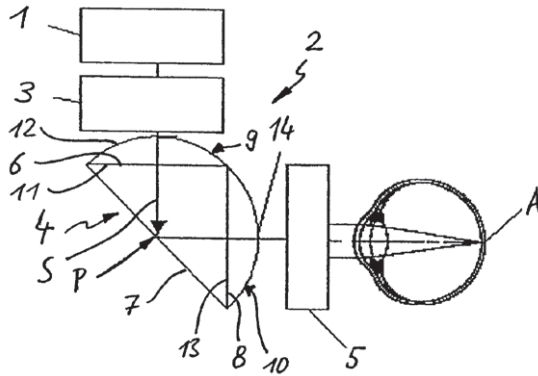
emissiv und teuer in der Herstellung sind. Flüssigkristall-Displays wären hierzu eine Alternative, denn sie benötigen weniger Leistung zum Betreiben und sind billiger herzustellen. Allerdings erfordern sie einen tieferen optischen Aufbau, so dass hiermit die oben genannten Vorteile nicht ohne Weiteres gewährleistet werden können.

### 6.2 Nutzung von Prismen

Der Vorteil einer im Unendlichen sichtbaren Überlagerung von Informationen mit der Wirklichkeit ist in Kapitel 2 bereits beschrieben: Durch Fern-Akkommodation der Augen kann der Nutzer gleichzeitig die reale Szenerie und die überlagerten Informationen sehen, ohne dass er gezwungen ist, seine Augen zu wenden oder deren Brennweiteinstellung zu verändern. Die Nutzung von Prismen stellt eine alternative Form zur Verwendung von Hohlspiegeln dar.

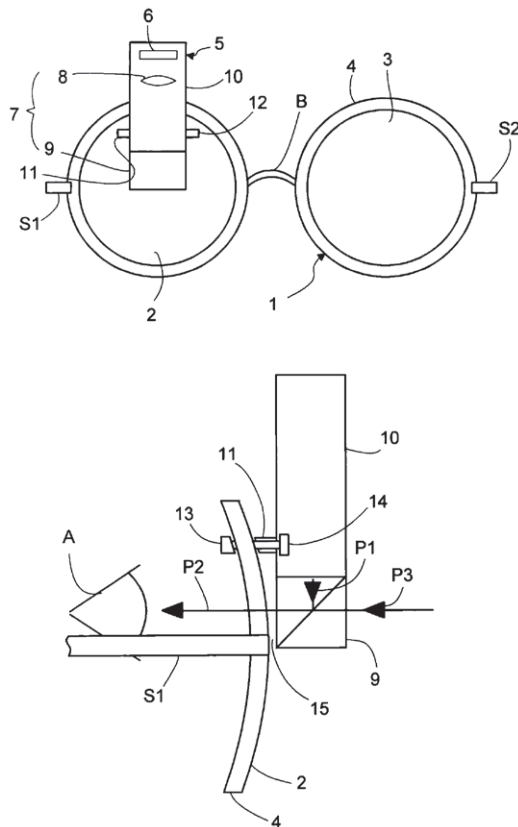
Wie man sich den Aufbau eines Umlenkprismas vorstellen muss, ist in der DE 10 2004 020 818 A1 (2004) beschrieben. Die Figur 8 zeigt eine Anzeigevorrichtung mit einem Bildelement 1 und einer ersten Optikgruppe 3, einem Umlenkelement 4 und einer zweiten Optikgruppe 5. Das Bildelement kann hierbei ein selbstleuchtendes Bildelement, wie zum Beispiel eine OLED (organische LED), oder ein nicht selbstleuchtendes Bildelement sein, das reflektiv oder transmissiv ausgebildet sein kann (zum Beispiel LCD-Modul, LCoS-Modul, Kippspiegelmatrix). Das Umlenkelement 4 umfasst ein 90°-Umlenkprisma P mit einer ersten 6, zweiten 7 und dritten 8 Prismenseite sowie eine erste 9 und zweite Plankonvexlinse 10, die jeweils eine Planfläche 11 und 13 sowie eine konvexe Fläche 12 und 14 aufweisen. Durch das Umlenkelement werden die von der Anzeigevorrichtung kommenden Strahlen S um 90° abgelenkt. Durch die Kombination von ebenen und gewölbten Reflexionsflächen (siehe dazu Kapitel 2) im Umlenkelement 4 sowie den Bearbeitungen in den beiden Optikgruppen 3 und 5 erreichen das Auge die von einem Bildpunkt des Bildelements 1 ausgehenden Strahlen in paralleler Form. Das Auge A erzeugt aus diesen parallelen Strahlen ein Bild auf der Netzhaut. Es muss dazu auf Unendlich akkommodiert sein.

Diese Anordnung ist leicht und kostengünstig herstellbar und somit für die Integration in Brillengestelle attraktiv.



Figur 8: Schematische Darstellung einer Anzeigevorrichtung (aus DE 10 2004 020 818 A1)

Die Patentschrift DE 102 31 427 B4 (2002) zeigt in Figur 9 eine Augmented-Reality-Datenbrille, die eine Überlagerung von Informationen mit der Wirklichkeit bietet. Dabei werden in einem Einkoppelement 9 die Strahlen



Figur 9: Vorderansicht und vergrößerte Seitenansicht einer Augmented-Reality-Datenbrille (aus DE 102 31 427 B4)

P1 aus dem Bilderzeugungselement 6 in einem Gehäuse 10 mit den Strahlen aus der Umgebung P3 überlagert. Das Gehäuse ist an der dem Auge abgewandten Seite eines der Brillengläser 2 über mehrere Durchgangsböhrungen geschraubt 11 und 13. Dabei ist zwischen dem Einkoppelement 9 und dem Brillenglas ein Luftspalt 15 vorhanden. Dies ist besonders vorteilhaft für den Fall, dass die Brille 1 eine Brille zur Korrektur einer Fehlsichtigkeit ist, da aufgrund des Luftspalts das Einkoppelement die Korrekturwirkung des tragenden Brillenglases nicht beeinflusst. In der Patentschrift ist auch die Möglichkeit beschrieben, dasselbe Gehäuse 10 auch an das zweite Brillenglas 3 anzuschrauben, womit unterschiedliche Bilder erzeugt und den Augen zugeführt werden, wodurch eine dreidimensionale Darstellung ermöglicht wird.

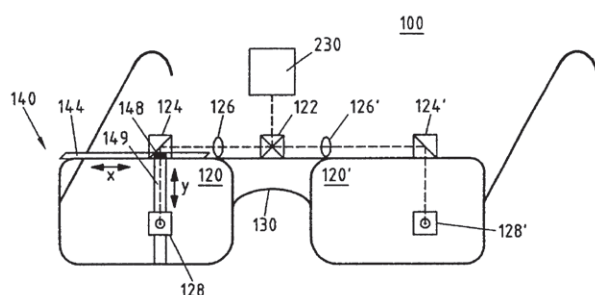
### 6.3 Variable Befestigung von Umlenkeinrichtungen

In Figur 4 sind die optischen Endelemente auf die Brillengläser drehbar geklebt. Bei der Figur 7 sind die optischen Endelemente in die Brillengläser fest eingeklebt. Mit Hilfe von drehbaren Spiegeln muss der Strahlengang auf die festgeklebten oder eingeklebten optischen Endelemente einjustiert werden. Die optischen Endelemente in den Figuren 4, 6 und 7 liegen jedoch nicht in der Hauptblickrichtung des Benutzers, so dass der Benutzer hinschauen muss, um die eingeblendeten Informationen zu sehen. In den Figuren 3 und 9 muss der Benutzer nicht extra hinschauen, da teildurchlässige Reflektoren dafür sorgen, dass die optischen Endelemente in der Hauptblickrichtung liegen, so dass für den Benutzer die angezeigten Bilder virtuell (hinter dem optischen Endelement) mit der Wirklichkeit überlagert erscheinen.

Es gibt verschiedene weitere Ansätze, die optischen Endelemente in die Hauptblickrichtung des Benutzers für eine Augmented-Reality-Datenbrille zu bringen. Das führt auch zu Anpassungen an den individuellen Augenabstand.

Die Druckschrift DE 101 32 872 A1 (2001) zeigt in der Figur 10 ein Durchsichtssystem 100 mit einer Bildquelle 230, wobei der Strahlengang über Umlenkeinrichtungen

122 und 124 und einem Linsensystem 126 an ein optisches Endelement 128 umgelenkt wird. An dem Gestell 130 ist eine Verstelleinrichtung 140 zum Verändern der Position je eines optischen Endelements 128 parallel zum Gesichtsfeld des Benutzers befestigt. Dies geschieht auf Schlitten 148 mit einem Führungselement 149 in einer Führungsschiene 144.



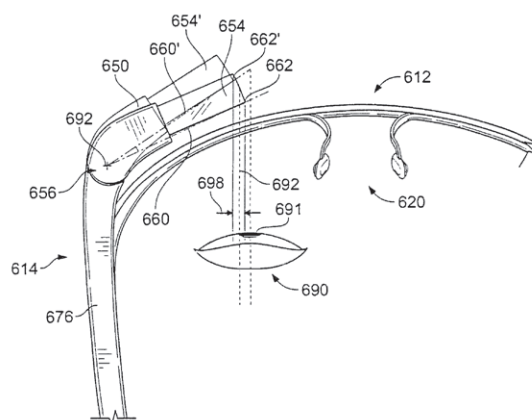
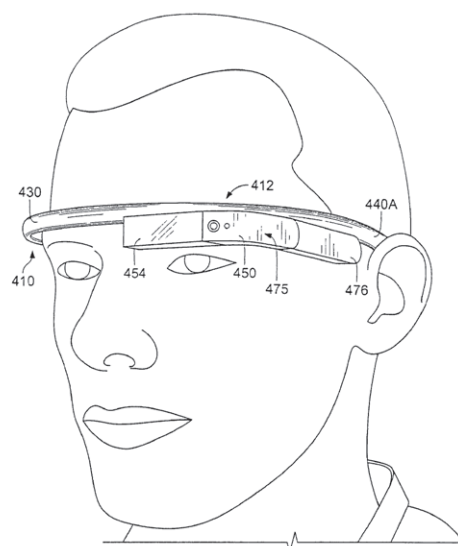
Figur 10: Durchsichtssystem (aus DE 101 32 872 A1)

Das Verschieben der optischen Endelemente entlang einer vorgegebenen Bahnkurve ist in der DE 10 2006 001 505 A1 (2006) offenbart.

Eine sehr flexible Anordnungsmöglichkeit von Umlenkeinrichtungen ist in der Gebrauchsmusterschrift DE 20 2012 003 317 U1 (2012) beschrieben. Die Figur 11 zeigt eine Vorrichtungsanordnung 410 mit einem Band 430, der den Kopf umspannt. Die Anzeigeeinrichtung 412 besteht aus einer Anzeige unter Nutzung eines teildurchlässigen Prismensystems 454 und 654, einem Ellenbogenabschnitt 450 und 650, einem Anzeigeabschnitt 475 und einem Armabschnitt 476. Sie ist mit dem Band auf einer, hier der linken Seite des Kopfes 440A befestigt. Durch diese Anordnung ist es möglich, die Anzeige 454 und 654 direkt in der Sichtlinie oder auch außerhalb der Sichtlinie anzuordnen.

In der Ausführungsform der Figur 11 liegt die Anzeige auf Augenbrauenhöhe des Benutzers, die man somit auch als „Head Up Display“ (HUD) bezeichnen kann. So muss der Benutzer nach oben schauen, um die eingeblendeten Informationen zu sehen [5]. Ferner kann es bei manchen Formen der Anzeige 654 und bei bestimmten Anwendungen erwünscht sein, eine seitliche Position der Anzeige einzustellen. Dabei kann die Anzeige 654' so angeordnet werden, dass sie sich zumindest teilweise über eine Außenkante (durch eine

Linie angegeben) der Pupille 691 des Trägers hinaus erstreckt. Das Gelenk 656 kann es dem Benutzer ermöglichen, den Ellenbogenabschnitt 650 zu drehen, so dass sich die Anzeige 654, während sie vom Auge 690 fort nach außen bewegt wird, auch um eine Strecke 698 entlang einer seitlichen Richtungskomponente bewegt. Veränderungen können auch innerhalb der Anzeigeeinrichtung 412 vorgenommen werden. So kann man durch die Drehung des Ellenbogenabschnitts 650 um die Achse 692 bewirken, dass die Fläche 660 näher zum Auge 690 hin oder weiter davon weg liegt. Dies kann es dem Benutzer ermöglichen, die Anzeige 654 für eine bequeme Betrachtung eines darauf präsentierten Bildes einzustellen und die Anzeige in einem solchen Abstand zu positionieren, dass sie beispielsweise die Augenbrauen oder Augenwimpern des Benutzers nicht berührt.



Figur 11: Vorrichtungsanordnung mit einem Kopf umspannenden Band (aus DE 20 2012 003 317 U1)

## 7 Ein- und Ausgabeeinrichtungen in der Datenbrille

Die Datenbrille aus Figur 11 besitzt im Ellenbogenabschnitt eine Kamera, ein Mikrofon und einen Kamera-Auslöser. Es können auch weitere Sensoren wie ein Augensensor, Beschleunigungssensoren, **Gyrosensor** und Kompass integriert sein. Im Armabschnitt ist ein Mainboard inklusive Prozessor, RAM und Flash-Speicher, Micro-USB und Status-LED, Bluetooth und WLAN, Einschalter und ein Touchpad enthalten. Ein Akku und ein Knochenschall-Lautsprecher befinden sich am Bandende am Ohr. Die Ein- und Ausgabeeinrichtungen für diese Datenbrille sind somit diejenigen, die auch Smartphones oder Smartwatches besitzen. Somit können unterschiedliche Funktionen der Datenbrille intuitiv bedient werden, wie man es von Smartphones oder Smartwatches her kennt. So kann man den Informationsaustausch mit dem über Funk verbundenen Internet oder einem Smartphone beispielsweise mit Wischbewegungen am Brillenbügel oder mit Sprachbefehlen steuern. Und die Datenbrille hält die Hände frei, die Interaktion wird unmittelbarer [5]. Beispielsweise kann man beim Fahrradfahren durch ein Nicken die Datenbrille einschalten, durch ein Zwinkern Fotos und Videos aufnehmen, sich mit Wischbewegungen am Brillenbügel und mit Sprachbefehlen navigieren lassen, über ein mit der Datenbrille verbundenes Smartphone Mitteilungen versenden oder telefonieren, ohne dass man die Hand vom Lenker nehmen muss ([6] und [7]).

Wie bereits im Text zu den Figuren 4, 6 und 7 beschrieben, kann die Augmented-Reality-Datenbrille als erweiterte Anzeige von internen und externen Kameras und als Headset dienen. Sie kann auch Kommunikationsgeräte steuern. Dabei übertragen die Kommunikationsgeräte die anzuzeigenden Daten, die durch Sensoren gemessen und gegebenenfalls für die Übertragung in den Kommunikationsgeräten zwischengespeichert wurden. Durch die Kamera kann man Markierungen scannen, die Navigationshinweise oder Anleitungen für die Anzeige liefern. Hinweise können mit Klängausgaben oder Vibrationen am Brillengestell versehen sein. Handgesten kann man durch die Kamera beispielsweise auch vor der Datenbrille erkennen. Die Datenbrille wird hauptsächlich in den Bereichen eingesetzt, wo man die Hand für andere Arbeiten

braucht. Das macht den Einsatz in verschiedensten Bereichen wie der Logistikbranche und in Krankenhäusern, aber auch bei Installateuren, Handwerkern [8], Klavierschülern [9], Spargelstechern (DE 10 2013 010 491 A1) (2013), Golfspielern (US 2015/0 057 095 A1) (2013) sowie Autofahrern [10] interessant.

Die Virtual-Reality-Datenbrille aus der Figur 5 kann ebenso mit Sensoren ausgestattet werden. So sind die meisten Virtual-Reality-Datenbrillen mit Beschleunigungssensoren, **Gyrosensor** und Kompass ausgestattet. So kann man mit Hilfe der Sensoren die Simulationsprogramme mit dem Kopf steuern. Sparen kann man sich diese Sensoren durch die Nutzung von Virtual-Reality-Gestellen, wenn man das Smartphone mit der Anzeige und den entsprechenden Sensoren einer Virtual-Reality-Brille, soweit vorhanden, in einer Virtual-Reality-Halterung nutzt und aufsetzt. Die Integration von Augensensoren in Virtual-Reality-Brillen ist auch möglich, um mit den Augen Simulationsprogramme steuern zu können. Es gibt erweiterte Möglichkeiten, mit dem Simulationsprogramm in Beziehung zu treten, genau dann, wenn sich der Benutzer genauso durch den virtuellen Raum wie durch den realen Raum bewegt. Die Bewegungen des Benutzers in der realen Umgebung müssen erfasst und mit dem Simulationsprogramm abgestimmt und übernommen werden. Dazu müssen Sensoren verwendet werden, die die reale Körperhaltung, Kopf- und Handbewegungen messen. Dies kann über optische Sensoren in der Brille und am Ende von zwei Hand-Controllern geschehen, die rotes gepulstes Laser-Licht aus gegenüberliegenden Ecken eines leeren Raumes erfassen, um die genaue Position der Sensoren im realen Raum zu bestimmen [11]. Umgekehrt kann die Brille mit Sendern ausgestattet sein [12]. Aber auch die Auswertung von außerhalb der Brille aufgestellten Tiefenkamerabildern oder die Auswertung der Änderungen von Magnetfeldern kann zu der Ermittlung der Positionen führen. Die Eingabe für ein Zeichenprogramm kann beispielsweise aber auch über die Eingabe auf einem Touchpad erfolgen.

## 8 Ausblick

Die derzeitige Entwicklung von Datenbrillen stellt eine Mischform aus Augmented-Reality und Virtual-Reality dar. Es wird auch von einer Holografiebrille gesprochen, die jedoch mit echter Holografie nichts gemein hat. Ihre geschwungenen Frontgläser sind teildurchlässig, so dass man gleichzeitig die reale Umgebung und die eingeblendeten virtuellen Objekte oder Informationen sehen kann. Die eingeblendeten virtuellen Objekte oder Informationen müssen dabei in Übereinstimmung mit den Bildern der realen Umgebung sein. Diverse Sensoren in der Brille ermitteln hierzu die Position des Benutzers, seiner Hände und die Blickrichtung seiner Augen in der realen Umgebung. Die Brillen können auch mit Projektoren für Virtual Displays ausgestattet sein [13]. Durch das Scannen von Markierungen in Form von QR-Codes in der realen Umgebung können reale Umgebungen bei Benutzung dieser Brillen zu virtuellen Umgebungen werden [11]. Somit kann der Benutzer durch Markierungen, durch Eingabegesten oder durch die Eingabe über ein Touchpad die Abfolge der eingeblendeten oder projizierten virtuellen Objekte oder Informationen steuern. Navigationshinweise oder Anleitungen können in die reale Umgebung in der Betrachtungsrichtung eingespielt werden. Dadurch dass die eingeblendeten oder projizierten virtuellen Objekte oder Informationen dreidimensional in der Betrachtungsrichtung erscheinen, wird auch ein größerer Sehkomfort durch die nicht ermüdende Ansicht der virtuellen Objekte oder Informationen gewährleistet, mit denen man direkt interagieren kann.

### Nicht-Patentliteratur

- [1] Wikipedia: Hohlspiegel, URL: <http://de.wikipedia.org/wiki/Hohlspiegel> [recherchiert am 23. März 2015]
- [2] Fastie-Ebert configuration, The Spectroscopy Net, 19. November 2006. URL: [http://www.the-spectroscopynet.eu/?Spectrometers:Monochromator:Fastie-Ebert\\_configuration](http://www.the-spectroscopynet.eu/?Spectrometers:Monochromator:Fastie-Ebert_configuration) [recherchiert am 23. März 2015]
- [3] SINGH, Shyam: Diffraction gratings: aberrations and applications. In: Optics and Laser Technology, Vol. 31, April 1999, S. 195–218
- [4] ZIEGLER, Peter-Michael: Blickpunkte, In: Magazin für computer technik (c't) 2013, Heft 13 vom 3. Juni 2013, S. 70–71
- [5] PORTECK, Stefan; SOKOLOV, Daniel AJ; ZOTA, Volker: Glass durchschaut, In: Magazin für computer technik (c't) 2013, Heft 13 vom 3. Juni 2013, S. 62–68
- [6] JANSSEN, Jan-Keno: Sind wir Cyborg?, In: Magazin für computer technik (c't) 2013, Heft 12 vom 21. Mai 2013, S. 74–76
- [7] SCHMUNDT, Hilmar, KRUG, Matthias und SPERBER, Sandra: Fotobrille Google Glass: Star-Fotografen testen die Zukunft, URL: <http://www.spiegel.de/netzwelt/gadgets/google-glass-star-fotografen-elliott-erwitt-und-bruce-gilden-testen-a-1015979.html> [recherchiert am 23. März 2015]
- [8] RÜDLIN, Nicole: Die Datenbrille ist nicht tot, sondern lebendiger denn je, URL: <http://t3n.de/news/datenbrille-tot-google-glass-interview-595510> [recherchiert am 24. März 2015]
- [9] WAGNER, Paul: Erfindung aus Kiel: Klavierspiel durch die Datenbrille, URL: <http://www.kn-online.de/Schleswig-Holstein/Schule-Studium/Hochschule/Erfindung-aus-Kiel-Klavierspiel-durch-die-Datenbrille> [recherchiert am 24. März 2015]
- [10] Heise Online: BMW: Röntgenbrille als Parkhilfe, URL: <http://www.heise.de/newsticker/meldung/BMW-Roentgenbrille-als-Parkhilfe-2545646.html> [recherchiert am 23. März 2015]
- [11] AUSTINAT, Roland, GEISELMANN, Hartmut: Parallelgesellschaft, In: Magazin für computer technik (c't) 2015, Heft 8 vom 21. März 2015, S. 20–24
- [12] JANSSEN, Jan-Keno: Brett vorm Kopf, In: Magazin für computer technik (c't) 2015, Heft 7 vom 7. März 2015, S. 96–97
- [13] JANSSEN, Jan-Keno, KULHLMANN, Ulrike: Blicken statt klicken, In: Magazin für computer technik (c't) 2015, Heft 5 vom 7. Februar 2015, S. 58–59



# Halbleiter-Bildsensorik am Beispiel von Flachdetektoren für die Röntgendiagnostik

Martin Albert, Patentabteilung 1.31

In den letzten Jahrzehnten hat die Bild- und Videotechnik eine immer weitere Verbreitung in den verschiedensten Lebensbereichen gefunden. Mit entscheidend für diese Entwicklung sind die Fortschritte in der Halbleitertechnologie und Mikroelektronik, wodurch insbesondere eine Reduktion der Herstellungs- und Betriebskosten, aber auch eine Verringerung der räumlichen Abmessungen und eine Vereinfachung der Benutzung von Geräten ermöglicht wurde. Wesentlichen Anteil dabei haben die Bildsensoren, die nunmehr zum überwiegenden Teil als höchstintegrierte mikroelektronische Schaltkreise ausgebildet sind. Der vorliegende Beitrag befasst sich mit Halbleiterbildaufnehmern, die als Flächensensoren eine Bildsensormatrix (Array) mit einer zweidimensionalen Anordnung von Sensorelementen umfassen.

---

## 1 Wellenlängenbereiche der elektromagnetischen Strahlung

Die hier betrachteten Bildsensorarrays sind dafür vorgesehen, Bilder elektromagnetischer Strahlung aufzunehmen, wobei verschiedene Sensoren verwendet werden, die für unterschiedliche Wellenlängenbereiche geeignet und für entsprechende Strahlung empfindlich sind. Der optische Bereich zwischen 100 nm und  $10^6$  nm Wellenlänge entspricht Frequenzen von 3000 bis 0,3 THz. Für den weiteren Bereich des elektromagnetischen Spektrums ist im Hinblick auf Bildsensoren insbesondere der Röntgenbereich (X-Strahlung, X-Ray) beispielsweise in der Medizintechnik von Bedeutung.

## 2 Bildgebende Röntgendetektoren

### 2.1 Historische Entwicklung der Röntgentechnik

Nachdem ihm die Universität Würzburg die Habilitation wegen seines fehlenden Abiturs noch verweigert hatte, wurde Wilhelm Conrad Röntgen dort 1888 zum Professor für Physik ernannt und fünf Jahre später zum Rektor der Universität gewählt. Am 8. November 1895 entdeckte er die von ihm selbst als „X-Strahlen“ und im Englischen weiterhin als „x-rays“ bezeichnete elektro-

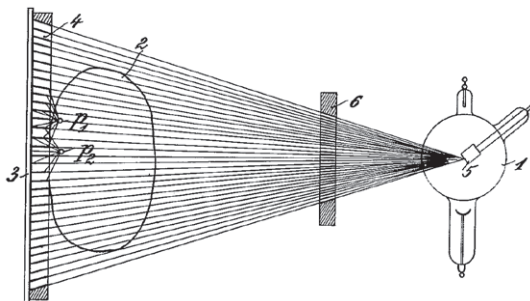
magnetische Strahlung, die im Deutschen aber nach ihm benannt wurde. Schon am 22. Dezember 1895 fertigte er eine Aufnahme der Hand seiner Frau an. Somit stellte sich schnell die Nutzbarkeit der Röntgenstrahlen für die medizinische Diagnostik heraus. Im Jahre 1901 erhielt er für seine Entdeckung den ersten Nobelpreis für Physik überhaupt. Die Versuchsanordnung mit der technischen Strahlenquelle und den zugehörigen Nachweisapparaturen, die in der Erstveröffentlichung „Über eine neue Art von Strahlen“ in „Aus den Sitzungsberichten der Würzburger Physik.-medic. Gesellschaft 1895.“ beschrieben ist, ließ Röntgen nicht patentieren: „Doch Röntgen war es wichtiger, dass die neuen Strahlen überall schnell zum Wohle der Menschen eingesetzt werden konnten, statt sie zu seinem Vorteil zu vermarkten.“ [1]. Röntgengeräte fanden innerhalb weniger Jahre weite Verbreitung in der medizinischen Diagnostik.

### 2.2 Streustrahlenraster

In der Röntgendiagnostik werden keine weit entfernten Objekte aufgenommen, die „fremdes“ Licht oft nur diffus in alle Richtungen reflektieren, sondern das aufzunehmende Objekt wird aus geringer Entfernung von einer kontrollierbaren Strahlenquelle durchstrahlt

und der Strahlungsempfänger ist – angepasst an die Größe des Objekts – sehr großflächig. Eine Fokussierung der Röntgenstrahlung zum Empfänger hin und eine verkleinernde Abbildung sind deshalb bei der Bildaufnahme nicht nötig.

Im aufzunehmenden Objekt, beispielsweise einem menschlichen Körper, entsteht aufgrund des Compton-Effekts Streustrahlung, die bei der Aufnahme mehrere nebeneinander liegende Bildsensorelemente durchquert. Um Artefakte aufgrund von schräg auf den Sensor auftreffenden hochenergetischen Teilchen der ionisierenden Strahlung zu reduzieren, werden typischerweise sogenannte „Streustrahlen-Raster“ als Kollimatoren unmittelbar vor dem Röntgendetektor eingesetzt, um die Strahlung zu „parallelisieren“. Das von Gustav Bucky eingeführte und nach ihm als „Bucky-Blende“ benannte und bis heute verwendete Streustrahlen-Raster 4 (siehe Figur 1) führt zur Kontrasterhöhung und zu weniger Unschärfe in den Röntgenaufnahmen. Dieses Raster kann aus rechteckigen oder wabenförmigen, parallel zur Nutzstrahlung ausgerichteten Lamellen aus Bleifolie bestehen, die die Streustrahlung absorbieren.



Figur 1: Streustrahlen-Raster für eine Röntgenvorrichtung (aus US 1,164,987 A)

Der Lamellenabstand kann beispielsweise 0,1 mm und die Dicke der Lamellen 0,05 mm betragen. Da dieses „grid“ einen Teil der Strahlung absorbiert, muss die Dosis entsprechend erhöht werden, um die Absorption zu kompensieren. Zur Vermeidung einer Abbildung der Lamellen im Röntgenbild führte Hollis E. Potter 1917 eine Bewegungseinrichtung ein, die das Raster während der Röntgenbelichtung in Schwingungen versetzt. In der DE 202 20 461 U1 ist ein digitaler Röntgendetektor mit einem zufallsverteilten Streustrahlenraster beschrieben.

## 2.3 Digitale Röntgenbild-Systeme

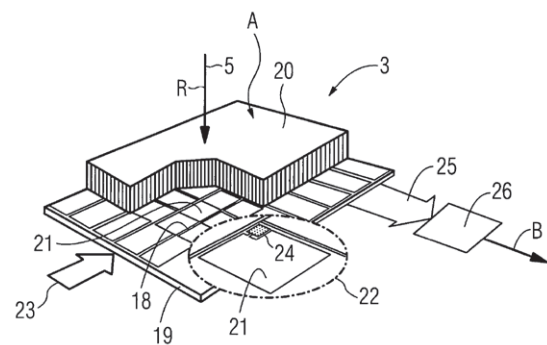
### 2.3.1 Speicherfolien-Systeme

Seit etwa 1980 haben digitale Systeme in Form von wiederverwendbaren Speicherfolien-Systemen in der Röntgendiagnostik Einzug gehalten, bei denen Schichten aus photostimulierbarem Phosphor nach einer Röntgenbelichtung ein latentes Bild über längere Zeit speichern können. Anschließend wird das Bild in einem Lesegerät mittels eines Lasers ausgelesen und nach erfolgter Digitalisierung einem Computer zugeführt, so dass eine digitale Bildverarbeitung möglich ist [2]. Ein weiterer Vorteil ist die in weiten Bereichen lineare Antwort des Systems auf die Röntgen-Intensität, aber die Bildqualität erreicht nur annähernd diejenige konventioneller Film-Folien-Systeme. Der nächste Schritt war die Entwicklung großflächiger digitaler Detektoren, die unmittelbar digitale Bilddaten liefern.

### 2.3.2 Flachdetektoren

Der Flachdetektor FD, gelegentlich auch als „Flachbild-detektor“ und im Englischen als „flat-panel detector“ FPD oder „flat panel imager“ bezeichnet, erobert seit 1999 das gesamte Spektrum der Röntgenaufnahme-systeme [3].

Bei den Halbleiter-Bildaufnehmer-Arrays, insbesondere den Flachdetektoren, wird unterschieden zwischen direkten Flachdetektorsystemen, bei denen Röntgenstrahlung unmittelbar elektrische Ladungen generiert, und indirekten Flachdetektorsystemen (siehe Figur 2),



Figur 2: Indirekter digitaler Röntgendetektor mit Szintillatorschicht (aus DE 103 43 496 A1)

bei denen die Röntgenstrahlung in einem Szintillator, der auch als Konverter bezeichnet wird, zunächst sichtbares Licht erzeugt, das erst in einem zweiten Schritt in lichtempfindlichen Photodetektoren in elektrische Ladungen umgewandelt wird.

Direkt wandelnde Flachdetektoren FD wurden ab circa 1995 entwickelt; die Einführung begann ab etwa 2000. Sie werden auf Basis amorphen Siliziums (a-Si) aufgebaut, wobei die hochenergetischen Teilchen der Strahlung beispielsweise in einer Schicht aus amorphem Selen (a-Se) unmittelbar elektrische Ladungen generieren. Bei indirekt wandelnden Flachdetektoren wird die Strahlung (im Wellenlängenbereich  $< 1 \text{ nm}$ ) in einem Szintillator, zum Beispiel einem Cäsium-Iodid-(CsI-) Kristall, in energetisch niedrigere elektromagnetische Strahlung größerer Wellenlänge, insbesondere in sichtbares Licht, umgewandelt.

Im Jahr 2004 erreichten „dynamische“ Flachdetektoren eine Bildwiederholfrequenz von 30 Bildern in der Sekunde [4], so dass sie auch Bildsequenzen beziehungsweise Videos „in Echtzeit“ zur Darstellung und Verarbeitung von Bewegtbildern aufnehmen können. Flachdetektoren ermöglichen durch unmittelbare Gewinnung und Weiterleitung der Bilddaten mit sofortiger Verfügbarkeit der digitalen Röntgenbilder einen verbesserten „Workflow“ ohne Transport und Auslesen von Röntgenkassetten und ohne Filmentwicklung. Aufgrund der hohen Sensitivität ist eine Reduktion der Röntgendosis von bis zu 75 % bei bestimmten Aufnahmen möglich [5]. Ein weiterer Vorteil ist die große Bildsignaldynamik und damit eine hohe Bildqualität. Aufgrund der universellen Einsetzbarkeit der Flachdetektoren besteht erstmals die Chance mit einer einzigen Technologie alle Applikationen in der Röntgendiagnostik abzudecken [3].

## 2.4 Großflächige Detektoren

Einteilige monolithisch (das heißt als Einkristall) integrierte Bildsensoren in der für medizinische Röntgenbildgebung erforderlichen Größe von  $20 \times 20 \text{ cm}$ , also im Bereich der Fläche von Halbleiter-Wafern, und mit einer ausreichend niedrigen Fehleranzahl herzustellen, ist kommerziell problematisch. Da bei Detektoren für

die medizinische Diagnostik nur eine geringe Zahl von Pixeldefekten erlaubt ist, würde eine weitere Vergrößerung der Chipfläche zu einer stark abnehmenden Ausbeute von Bildsensor-Arrays mit tolerierbaren Defekten führen.

Auch Versuche, eine Vielzahl von CCD-Bildsensorchips in einer größeren Matrix zusammenzuschalten und mit einer Konverter-Schicht (Szintillator) zu kombinieren, wodurch sich grundsätzlich der Anteil der aufgrund von nicht tolerierbaren Defekten zu verwerfenden Chips im Vergleich zu einem entsprechend großen einteiligen Bildsensor reduzieren lässt, scheiterten aufgrund von Problemen an den Stoßstellen zwischen den einzelnen Halbleiterchips ([6] und [7]).

In der Medizintechnik arbeiteten mehrere Gerätehersteller zusammen an der Entwicklung eines zunächst  $20 \times 20 \text{ cm}$  großen Bildaufnehmers mit  $1024 \times 1024$  Bildpunkten, der in Dünnschichttechnologie (thin film) auf einem Glasträger aufgebaut ist.

## 2.5 Dünnschicht- oder Dünnschichttechnologie und TFT-Technologie

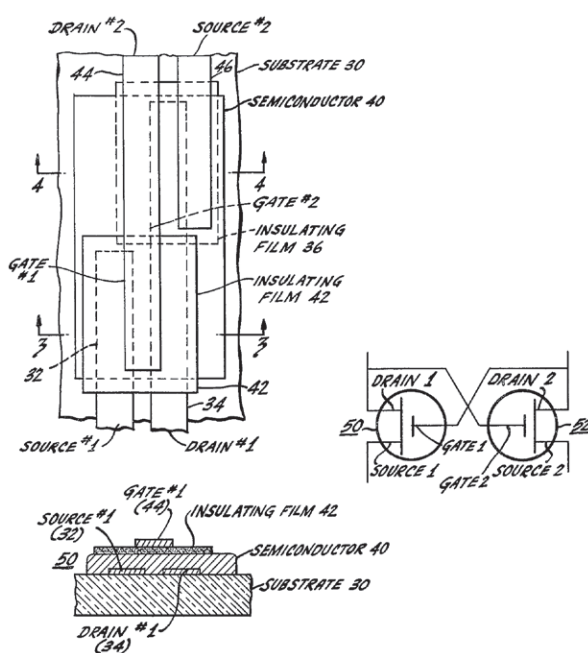
### 2.5.1 Dünnschicht- oder Dünnschichttechnologie

Die Dünnschicht- oder Dünnschichttechnologie ist seit den 40er Jahren des vergangenen Jahrhunderts bekannt. Erste systematische Arbeiten wurden im Rahmen der Entwicklung von elektronischen Annäherungszündern für Geschosse betrieben, wobei 1943 erste lineare Dünnschicht-Netzwerke entstanden [8]. Der Begriff wird verwendet bei Schichtdicken bis circa  $1 \mu\text{m}$ , nach anderer Angabe bei Schichtdicken bis zu mehreren Mikrometern. Die Produkte dieser Technologie dienten ursprünglich als Träger für elektronische Bauteile; sie erlaubte die Herstellung „gedruckter“ Schaltungen mit Leiterbahnen, wobei durch die Materialauswahl, Leiterbahnbreite und Leiterbahnführung beziehungsweise durch das Layout in bestimmten Grenzen auch passive Bauelemente (Widerstände, Kondensatoren und Induktivitäten) in der Dünnschichttechnologie selbst realisierbar waren. Mehrlagige keramische Substrate ermöglichen in Verbindung mit Durchkontaktierungen komplexe MLC-Anwendungen (Multi-Layer Ceramic) (US 4,751,349 A), die beispielsweise als Träger für un-

gehäuste monolithische Halbleiterchips zum Einsatz in Mainframe-Computern (Multichip Modules MCM) entwickelt und gefertigt wurden, wobei die die „aktiven“ elektronischen Bauelemente beinhaltenden Chips „kopf-über“ mittels der sogenannten „Flip-Chip“-Löttechnik montiert und kontaktiert wurden (US 5,914,533 A).

### 2.5.2 Thin-Film Transistor TFT

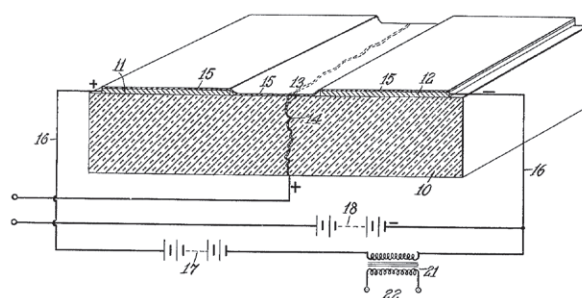
Um 1961 gelang Paul K. Weimer die Herstellung der ersten Dünnschichttransistoren (Thin-Film Transistor TFT) unter Verwendung mikrokristalliner Schichten aus Cadmiumsulfid CdS [9]. Dabei handelt es sich nach Figur 3 um spezielle Feldeffekttransistoren (FET) mit isoliertem Gate (IGFET, MISFET).



Figur 3: TFT-Transistor (aus US 3,191,061 A)

Das Konzept für dieses Bauelement wurde bereits 1925 von Julius Edgar Lilienfeld erfunden (siehe Figur 4) [10], konnte damals aber wegen mangelnder technologischer Beherrschung des Halbleitermaterials Kupfersulfid CuS noch nicht realisiert werden. Lilienfeld erfand auch den Sperrschicht-FET und den Bipolar-Transistor (US 1,900,018 A, US 1,877,140 A) [11], wofür Shockley, Brattain und Bardeen 1956 den Nobelpreis für Physik erhielten. Oskar Heil meldete 1934 ebenfalls einen Feldeffekttransistor zum Patent an (GB 439,457 A),

wobei kein Nachweis über den Bau eines funktionierenden Bauelements existiert.



Figur 4: Konzept des ersten Transistors (aus US 1,745,175 A)

### 2.5.3 Technologischer Fortschritt durch Einsatz hydrogenisierten amorphes Siliziums

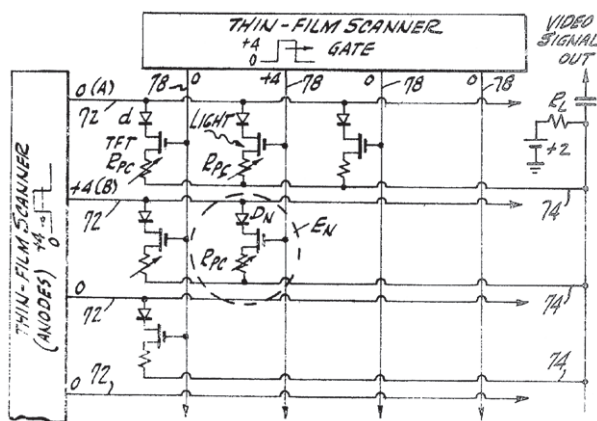
Mit den von Paul K. Weimer gebauten TFTs konnte noch keine Massenproduktion auf großformatigen Substraten realisiert werden. Ein technologischer Durchbruch gelang 1979 mit dem ersten aus hydrogenisiertem amorphem Silizium (a-Si:H) hergestellten funktionierenden TFT (IGFET) mit einem Siliziumnitrid-Gate, weil der Herstellungsprozess einfach ist und weil das Bauteil bei Raumtemperatur unter den Bedingungen der Erdatmosphäre stabil ist [10] und [12]. Bereits 1975 beschrieben Neudeck und Malhotra einen „amorphous silicon thin film transistor“ und erwähnten dabei auch die Möglichkeit der Optimierung durch dünnere Oxide oder durch das Hydrogenisieren des amorphen Siliziums während der Herstellung [13]. Auch hatte Alain Deneuville schon 1977 einen bipolaren TFT aus a-Si:H zum Patent angemeldet (US 4,127,861 A).

Heute werden in dieser auf amorphen Halbleitern basierenden Technologie nicht nur großformatige Active Matrix-Flat Panel-Detektoren aufgebaut, sondern beispielsweise auch Active Matrix-Flat Panel-Displays beziehungsweise TFT-LCDs, aber auch amorphe Dünnschicht-Solarzellen. Neben der Möglichkeit, Detektoren mit sehr großer Fläche herzustellen, bietet die TFT-Technologie für Röntgendetektoren gegenüber monolithisch integrierten Detektormatrizen den wichtigen Vorteil, dass die Alterung der Dünnschicht-Transistoren infolge der Röntgenbestrahlung weniger ausgeprägt ist als die von Transistoren aus kristallinem Silizium. Nachteil dagegen ist eine geringere Fertigungs-Ausbeute als bei Standard-Halbleitertechnologien wie CMOS oder

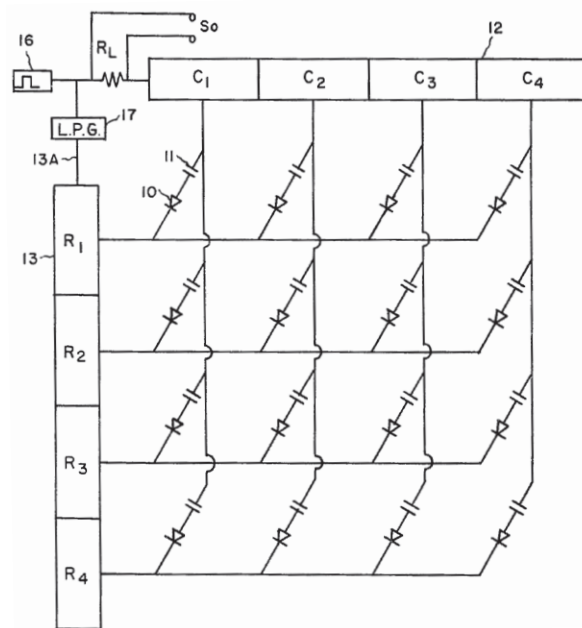
auch CCD, insbesondere dann, wenn Dünnschicht-Fotodioden und Dünnschicht-Transistoren unterschiedliche Materialien und damit auch zusätzliche photolithographische Belichtungsmasken und Prozessschritte für die Herstellung benötigen.

#### 2.5.4 Bildsensor-Arrays in TFT-Technologie

Parallel zu dieser technologischen Weiterentwicklung meldete Paul K. Weimer schon ab 1965 mehrere Patente an, mit denen in TFT-Technologie aufgebaute Bildsensor-Arrays geschützt werden sollten (siehe Figur 5).



Figur 5: Bildsensor-Array in TFT-Technologie (aus US 3,488,508 A)





ringert sich die gespeicherte Ladung nur sehr langsam; bei Lichteinfall hängt die Geschwindigkeit des Abfalls der gespeicherten Ladung (aufgrund des induzierten Photostroms, der zur Lichtintensität proportional ist) linear von der Intensität des auftreffenden Lichts ab. Die Verringerung der Ladung ist proportional zum zeitlichen Integral der Belichtung über die verstrichene Belichtungszeit und wird folglich als Integrationszeit bezeichnet. Weckler schlug vor, Arrays monolithisch zu integrieren und zum Auslesen die einzelnen Photodetektoren (in einer geeigneten Reihenfolge) periodisch abzutasten. Als Vorteile nannte er neben einer höheren Empfindlichkeit und einem verbesserten Signal-Rausch-Verhältnis die lineare Abhängigkeit des gemessenen Signals von der Lichtintensität über mehrere Größenordnungen hinweg, die elektronische Steuerbarkeit der Empfindlichkeit (durch Ansteuerung der Belichtungszeit) und die einfache Integration in Bildsensor-Arrays. Diese Erkenntnisse sind für die weit überwiegende Anzahl heutiger Bildsensoren – insbesondere auch für die meisten monolithisch integrierten CCD- und CMOS-APS-Sensoren – von außerordentlicher Bedeutung ([14] und [15]).

### 2.5.6 Zukünftige Weiterentwicklungen der TFT-Technologie

Bisherige Weiterentwicklungen betrafen die Vergrößerung der Fläche des Glassubstrats, die evolutionäre Verbesserung der Fertigungsanlagen, die Vereinfachung der Herstellungsprozesse und ein besseres Verständnis der Prozesse für große Flächen. Weiterhin wird die Verbesserung der bisher im Hinblick auf effiziente Produktionsprozesse für großflächige Schaltungen noch weniger weit fortgeschrittenen TFT-Technologie mit polykristallinem Silizium (poly-Si) statt des bisher verwendeten amorphen Siliziums (a-Si) angestrebt, wobei thermische und laserbasierte Methoden zur Umwandlung durch Kristallisation untersucht werden. Poly-Si-TFTs sind vorteilhaft gegenüber a-Si-TFTs bezüglich der höheren Ladungsträgerbeweglichkeit und höheren Zuverlässigkeit. Ziel ist es, analog zu den monolithisch integrierten „system-on-chip“-Bauelementen komplette „system-on-glass“-Einheiten in Dünnschichttechnik zu integrieren und somit bisher

getrennt gefertigte und angeordnete periphere Schaltkreise – neben den obligatorischen I/O-Schaltkreisen auch CPUs und Speicher – auf das Substrat zu verlagern. Zudem wird nach alternativen Materialien gesucht, die die hohe Ladungsträgerbeweglichkeit des poly-Si mit weniger aufwändigen Fertigungsverfahren und somit niedrigeren Herstellungskosten vereinen. Eine mögliche Entwicklung ist dabei der Einsatz organischer TFTs (OTFT) beziehungsweise organischer Feldeffekttransistoren (OFET) auf Basis von organischen halbleitenden Polymer-Materialien, die ohne Vakuumtechnologie bei Raumtemperatur zum Beispiel mittels sehr preiswerter (Ink Jet-) Druckverfahren hergestellt werden können. Auch wenn minimale Fertigungskosten in der Massenproduktion im Hinblick auf einen Einsatz in der medizinischen Diagnostik nicht von überragender Bedeutung sind, so könnten sich doch Einsatzfelder für mechanisch flexible, biegbare Substrate ergeben (vergleiche DE 103 33 821 A1, für medizinische CT- und insbesondere auch für industrielle Anwendungen).

## 2.6 Flachdetektoren in TFT-Technologie für die medizinische Diagnostik

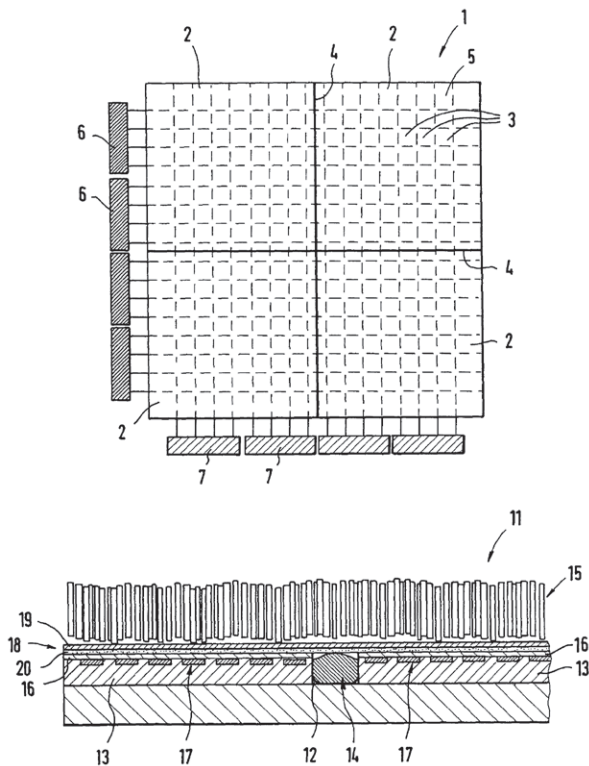
### 2.6.1 Entwicklung bis zur Einführungsphase

Besondere Herausforderungen bei der Entwicklung der Flachdetektoren (Flat Panel Detector FPD) stellten die möglichst fehlerfreie Herstellung von etwa einer Million Pixelschaltungen sowie die Problematik der schnellen Alterung der Transistoren, die der Röntgenstrahlung mit hoher Dosis ausgesetzt sind, dar [6]. Die klinische Erprobung begann etwa im Jahr 2000. Ein in verschiedenen Produkten eingesetzter Detektor erreichte im Jahr 2002 ein Format von 43 x 43 cm bei einer 3 x 3k Matrix (das entspricht circa 9 Millionen Pixeln) und einer Fläche der quadratischen Pixel von 143 x 143  $\mu\text{m}$  [5]. Als Konverter-Schicht kommen dabei Szintillatoren aus einem CsI-Film zum Einsatz.

### 2.6.2 Szintillator bei indirekten Flachdetektoren

Mit Natrium dotiertes Cäsium-Iodid CsI:Na zeichnet sich dadurch aus, dass es sich durch Kristallwachstum

auf der Sensoroberfläche so herstellen lässt, dass sich eine Säulenstruktur aus Kristalliten (Nadeln mit einem Durchmesser von 5 bis 10 µm, siehe Figur 8) ausbildet, die wie Lichtleitfasern das Szintillationslicht auf den Sensor leiten.



Figur 8: Aus mehreren Panels zusammengesetzter Festkörper-Bilddetektor mit einer gemeinsamen Szintillatorschicht (aus DE 199 14 701 B4)

Aufgrund der daraus resultierenden geringen Streuung können relativ dicke Schichten von typischerweise 500 µm Dicke hergestellt werden, in denen bis zu 70 % der Röntgenquanten nachgewiesen werden können (Quanteneffizienz DQE), ohne dass dadurch die Schärfe der aufgenommenen Bilder zu sehr leidet. Dabei ist die Zahl der Photonen des im Szintillationskristall von einem einzigen absorbierten Röntgen- oder Gamma-Quant verursachten sogenannten „elektromagnetischen Schauers“ und des letztendlich erzeugten Lichtblitzes proportional zur Energie, die das Quant im Szintillator abgibt [6].

Es treten aber – abhängig vom Material – nicht alle einfallenden hochenergetischen Quanten mit dem Konvertermaterial in Wechselwirkung, so dass nur ein Teil der Quanten im Szintillator absorbiert wird. Eine

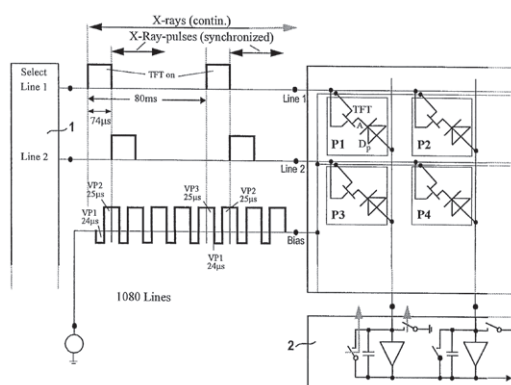
Reflexionsschicht für sichtbares Licht, die für Röntgenstrahlung transparent sein muss, auf der dem Sensor gegenüberliegenden Seite der Konverter-Schicht sorgt dafür, dass auch die vom Sensor weg gerichtet ausgesandten Szintillations-Photonen zurück auf den Sensor reflektiert werden und so detektiert werden können. Abgedeckt wird der Detektor durch eine auf den Reflektor aufgetragene Schutzschicht beispielsweise aus Graphit. Alternative Materialien für Szintillatoren sind Sinterkeramiken aus Gadoliniumoxysulfid  $Gd_2O_2S$  oder  $GdO_2$  [4].

### 2.6.3 Aktive TFT-Matrix (AM) mit Photodioden für indirekte Flachdetektoren

Die von der auf einem indirekten Flachdetektor aufgetragenen Konverterschicht emittierten Lichtblitze beziehungsweise Photonen treffen auf den eigentlichen Sensor, der auf die Farbe dieses Lichts, beispielsweise Grün, bestmöglich abgestimmt sein sollte, das heißt seine maximale Empfindlichkeit bei der entsprechenden Wellenlänge aufweisen sollte. Die Dünnschicht- oder auch Dünnschicht-Technologie, die für die Herstellung dieses Sensors eingesetzt wird, ist grundsätzlich seit langem bekannt [16], erlangte aber mit den LCDs und später den TFT-Displays eine enorme Weiterentwicklung und Verbreitung, wobei auch sehr großflächige „Substrate“ mit elektronischen Schaltungen versehen werden können. Durch die Verwendung von amorphem Silizium (a-Si), das Halbleitereigenschaften aufweist, wurde die Herstellung elektronischer Bauelemente wie Dioden (thin film diode TFD), insbesondere auch Photodioden, und „aktiver“ Transistoren (thin film transistor TFT) ermöglicht. Speziell für den Einsatz in Röntgendetektoren eignen sich die aus amorphem Silizium hergestellten Bauelemente in den „aktiven Matrizen“ besonders gut, weil sie eine hohe Unempfindlichkeit gegenüber Röntgenstrahlung aufweisen. Bei den indirekt arbeitenden Flachdetektoren werden Photonen aus der Konversionsschicht in aus amorphem Silizium bestehenden Photodioden des Sensors in einem zweiten Umwandlungsschritt in elektrische Ladungen gewandelt. Der Füllfaktor der typischerweise 150 µm großen Pixel beträgt circa 60 bis 70 %, das heißt die photosensitive Fläche der Photodiode nimmt diesen

Anteil der gesamten Pixelfläche ein. Ein absorbiertes Röntgenquant aus einer Röntgenröhre mit einer Beschleunigungsspannung von beispielsweise 53 kV und mit einer Energie von somit 53 keV kann in der Photodiode eine Ladung von etwa 800 bis 1000 Elektronen erzeugen. Jede in Dünnschicht-Technologie hergestellte Pixelzelle enthält neben der Photodiode (PD/TFD) einen Feldeffekttransistor (FET/TFT) als Schalttransistor. Über „Bias-Leitungen“ werden die Anoden der Photodioden mit einer geeigneten Vorspannung versorgt. Über jede Zeile der gesamten Matrix führt eine Gateleitung, über die die TFTs der Pixel dieser Zeile angesteuert werden können. Die TFTs aller Pixel jeder Spalte sind über jeweils eine Datenleitung miteinander verbunden, über die die in der Photodiode eines Pixels während der Röntgenbelichtung erzeugte und auf ihrer Sperrschichtkapazität akkumulierte Ladung ausgelesen wird, wenn der Schalttransistor mittels eines geeigneten Gate-Potentials leitend geschaltet wird. Durch sequentiell zeilenweise Ansteuerung lässt sich eine Zeile nach der anderen auslesen, und zwar jeweils alle Pixel einer Zeile parallel über die Datenleitungen (Spaltenleitungen). Die aus der Matrix ausgelesenen Ladungen werden in speziellen Schaltkreisen rauscharm verstärkt und vor der Weiterverarbeitung einer Analog-Digital-Wandlung (ADW/ADU/ADC) unterzogen. Die Bildinformation erzeugenden Schichten der Detektorelemente haben eine Dicke von etwa 0,5 bis 2  $\mu\text{m}$ ; das Glassubstrat als Trägerschicht ist etwa 1 bis 2 mm dick [4].

Ein Beispiel aus der Patentliteratur zeigt Figur 9 aus der DE 196 31 385 A1. Diese Druckschrift beschreibt

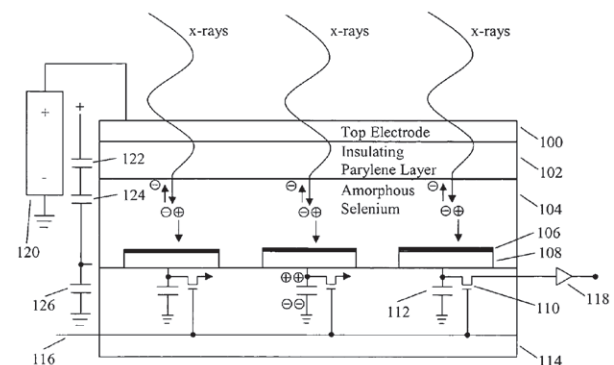


Figur 9: Schematischer Aufbau und Betriebsverfahren für einen TFT-Festkörper-Bildwandler für Röntgenstrahlung (aus DE 196 31 385 A1)

einen TFT-Festkörper-Bildwandler mit einer Photodioden-Matrix insbesondere auch zur Erfassung von Röntgenstrahlung und Verfahren zu dessen Betrieb, wobei auf die Notwendigkeit hingewiesen wird, jeweils angepasst an verschiedene Betriebsarten mittels geeigneter Dunkel- und Gainbilder eine Korrektur der aufgenommenen Bilder vorzunehmen.

## 2.6.4 Direkte Flachdetektoren

Direkte Flachdetektoren benötigen keinen Szintillator, weil ein „Direktkonverter“ absorbierte Röntgenquanten unmittelbar in freie Elektronen wandelt. Das am weitesten verbreitete Material hierfür ist wie in Figur 10 amorphes Selen (a-Se). Es wird direkt auf eine aktive Matrix aus amorphem Silizium (a-Si) aufgebracht, bei der jede Pixelzelle neben dem Transistor (statt einer Photodiode wie bei indirekten Flachdetektoren) eine Elektrode umfasst. Abhängig vom Einsatzbereich und damit abhängig vom verwendeten Röntgenspektrum sind Schichtdicken des Sells von bis zu 1 mm erforderlich, wobei ein absorbiertes Röntgenquant einige 100 bis über 1000 Elektronen generiert. Diese Ladungsträger müssen mittels eines elektrischen Feldes auf die Elektrode des zugehörigen Pixels transportiert werden. Dazu wird eine Hochspannung von etwa 10 kV an die 1 mm dicke Sellschicht angelegt, damit nicht zu viele Ladungen unterwegs verloren gehen. Das Auslesen der auf den Elektroden akkumulierten Elektronen und die Weiterverarbeitung der Signale können in gleicher Weise geschehen wie bei den indirekten Flachdetektoren.



Figur 10: Direkt wandelnder Röntgen-Flachdetektor (aus US 7,304,308 B2)



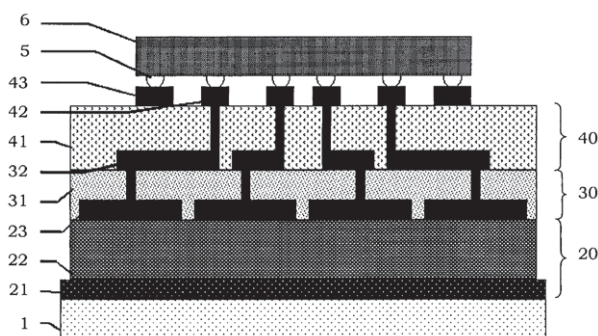


Defektpixel betrachtet werden und erhöhen die Zahl der zu korrigierenden Bildpunkte beträchtlich.

### 2.6.5.2 Mehrlagiger Aufbau und Modularisierung

Die DE 10 2005 045 901 A1 beschreibt einen direkten oder indirekten Röntgen-Flachbilddetektor mit einer auf einem Glassubstrat angeordneten, aus „Pixel-Auslese-einheiten“ zusammengesetzten, aktiven Auslesematrix. Unter dem Substrat liegen mechanisch sehr empfindliche Elektronikplatinen. Der Flachdetektor ist für den mobilen Einsatz konzipiert und beinhaltet eine drahtlose Datenübertragungs-Schnittstelle, so dass er völlig frei platziert werden kann. Um eine Beschädigung der Platinen beispielsweise durch Verbiegen oder andere mechanische Beanspruchungen zu vermeiden, sind die Elektronikplatinen modular aufgebaut. Auch die DE 10 2007 046 258 A1 offenbart einen portablen Flachdetektor in einem Schutzgehäuse, wobei eine drahtlose Daten- und Bildübertragung vorgesehen ist. Solche in einem schützenden Gehäuse eingebauten FPDs können die gleiche Größe und Dicke haben wie konventionelle Filmkassetten; dann können sie diese in entsprechenden Schubladen, die umgangssprachlich im Englischen als „Bucky Tray“ oder kurz als „Bucky“ bezeichnet werden, ersetzen.

Es ist vorteilhaft, die beispielsweise aus monolithischen ICs bestehende periphere Elektronik in einer weiteren Ebene unter oder über dem Dünnschichtschaltungs-Substrat anzuordnen und so eine „dreidimensionale Hybridschaltung“ aufzubauen. Dadurch entsteht ein sehr kompakter mehrlagiger Aufbau mit kurzen Leitungswegen. Dies ist insbesondere bei „dynamisch“ betriebenen, das



Figur 12: Schichtförmiger Aufbau einer Sensoranordnung (aus DE 101 42 531 A1)

heißt Bildsequenzen aufnehmenden Flachdetektoren wichtig, weil dort die zum Auslesen der Bilddaten erforderliche Bandbreite eine Begrenzung darstellt. Diese Platzierung der peripheren Elektronikschaltkreise vereinfacht auch das nahtlose Zusammenfügen von mehreren Panels zu einem größeren Flachdetektor. Figur 12 aus der DE 101 42 531 A1 zeigt einen Flachdetektor mit einer Trennung zwischen dem Sensor-Panel und den ICs der Auswerteelektronik in verschiedenen Ebenen.

Die DE 10 2007 022 197 A1 beschreibt ein Detektorelement für Röntgenstrahlung, das aus einer Ebene mit einem Detektionselement zur Umwandlung der Röntgenstrahlung in elektrische Signale, einer Ebene mit einer Mehrzahl von elektronischen „Bauelementen“ und mit einer in der Mitte liegenden Ebene mit einem „Umkontaktierungs-Substrat“ aufgebaut ist, wobei über Letzteres die Kontaktierung und die elektrisch leitende Verbindung der Detektor-Schicht und der Bauelemente-Schicht erfolgt. Dabei ist eine „zumindest in deren Funktion gegen eine Einwirkung der Röntgenstrahlung resistente“ Ausbildung der elektronischen Bauelemente zur Vorverarbeitung der Signale des Detektionselements, bei denen es sich um ASICs handeln kann, vorgesehen, wobei „keinerlei Abschirmung der Bauelemente gegen die Röntgenstrahlung erforderlich“ sein soll. Der Vorteil beim Einsatz solcher Detektoren beispielsweise im Bereich der medizinischen Röntgendiagnostik ist offenkundig.

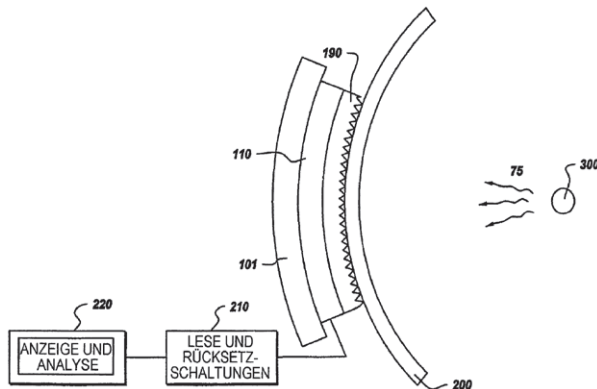
### 2.6.6. Organische flexible Sensoren

Es ist bekannt, Dünnschichtschaltungen mit mechanisch flexiblen, organischen Substraten und Halbleitermaterialien aufzubauen.

Die DE 102 44 177 A1 offenbart einen Bilddetektor für eine Röntgeneinrichtung mit rückseitig kontaktierten organischen Bildsensoren. Als Material für die Photosensoren wird ein organisches Photodioden-Material verwendet.

Figur 13 aus der DE 103 33 821 A1 zeigt eine mechanisch flexible Bildgebungseinrichtung auf einem biegbaren Substrat 101 aus einem organischen Polymer für medizinische CT- und insbesondere auch für industrielle Anwendungen.





Figur 13: Flexible digitale Bildgebungseinrichtung  
(aus DE 103 33 821 A1)

### 3 Korrekturverfahren für Bildsensor-Arrays

#### 3.1 Sensorbedingte Fehler

Halbleiterschaltkreise können Defekte aufgrund von Fertigungsfehlern enthalten. Bei Festkörper-Bildsensoren führen solche Schäden zu unerwünschten Bildfehlern. Je nach Anwendungsbereich werden an Bildsensor-Arrays hohe Qualitäts-Anforderungen gestellt. Bei Systemen zur radiologischen Bildgebung sind nur geringe Fehler zulässig; Fehler müssen vermieden oder zumindest korrigiert werden. Die danach verbliebenen Fehler entscheiden über die in der Fertigung erzielten Ausbeuten und damit letztendlich auch in hohem Maße über den wirtschaftlichen Erfolg.

Fertigungsbedingte oder sich im Betrieb einstellende oder verschlechternde Mängel beziehungsweise Fehler können sich in kontinuierlichen Änderungen von Parametern über die Fläche des Sensor-Arrays oder in un stetigen, nur statistisch zu beschreibenden Variationen der Parameter von Pixel zu Pixel niederschlagen. Bei bestimmten Fehlern ändert sich das Verhalten eines betroffenen Sensorelement-Systems, das angibt, wie das Sensorausgangssignal auf eintreffende Strahlung unter verschiedenen Randbedingungen antwortet, nur graduell und ist dann durch eine für jedes Pixel individuelle Korrektur von Parametern grundsätzlich korrigierbar. Sind diese Fehler allerdings zu groß, so führt eine rechnerische Korrektur nicht mehr zu verwertbaren Ergebnissen, weil die berechneten Ergebnisse zu unsicher sind.

Fertigungsbedingte oder im Laufe des Betriebs erworbene Fehler können sich aber auch als Defekte äußern, die bei Bildsensoren zu schadhafte Pixeln führen. Diese Fehler sind rechnerisch nicht kompensierbar, weil keine (Sensor-)Information vorliegt, die die Grundlage für Korrekturberechnungen bilden könnte. Bei Bildsensoren sind auf eintreffendes Licht oder Strahlung nicht reagierende Pixel als nicht korrigierbar defekte Sensorelemente anzusehen. Bei Bildsensor-Arrays ist auch kein Ersatz durch redundante Hardwarekomponenten möglich.

Das Zusammenfügen mehrerer „Panels“ zu einem großen Flachdetektor, auch als „Butting“ bezeichnet, führt darüber hinaus trotz großer Sorgfalt zu zahlreichen zusätzlichen defekten Pixeln beziehungsweise fehlerhaften Pixelwerten an den in Zeilen- und Spaltenrichtung verlaufenden Nahtstellen.

Sowohl korrigierbare Fehler als auch nicht korrigierbare Defekte in Bildsensor-Arrays können bereits seit dessen Herstellung bestehen oder aber erst nachträglich während der Betriebsphase des Sensors entstehen und sich insbesondere auch zeitlich verändern. Dabei können bei den verwendeten Materialien immanente, mit unterschiedlichsten Zeitkonstanten ablaufende Veränderungen und Alterungsprozesse stattfinden. Die Alterungsprozesse können aber auch durch äußere Einflüsse hervorgerufen oder beschleunigt werden. Beispielsweise verursacht eine Strahlenbelastung sowohl bei monolithisch integrierten Sensoren als auch, wenn auch weniger stark, bei TFT-Detektoren Schäden und führt dadurch zu einer schleichenden „künstlichen Alterung“. Dass Sensoren hochenergetischer Strahlung, beispielsweise Gamma-, Röntgen- oder sonstiger Teilchenstrahlung, ausgesetzt werden, ist bei medizinischen Röntgendetektoren unvermeidlich. Treffer hochenergetischer Strahlungsquanten können aber nicht nur mehr oder weniger langsam verlaufende Material- und damit Parameteränderungen bewirken, sondern auch sofort und unmittelbar eintretende, vorübergehende oder bleibende Schäden verursachen und Defekte hervorrufen.

Echte Hardware-Defekte in der Pixelmatrix können in anderen benachbarten oder schaltungstechnisch verknüpften Pixeln, die selbst grundsätzlich korrekt funktionieren, zu Folgefehlern führen. Defekte, die

zum „Überlaufen“ akkumulierter Ladungsträger eines Pixels zu Nachbarpixeln führen, können dort falsche Pixelwerte verursachen. Bekannt von Röntgen-Flachdetektoren sind Defekte zum Beispiel in Pixeln oder in Spaltenleitungen, die das Auslesen von Pixelinformation von in Bezug auf die Richtung des Ladungstransports beim Auslesen vor der Fehlerstelle liegenden Bildsensorelementen verhindern; dies führt zu typischen Fehlermustern.

Weiterhin können auch Übersprech-Effekte in grundsätzlich der Spezifikation entsprechend funktionierenden Detektor-Matrizen aufgrund unzureichender örtlicher oder zeitlicher Entkopplung sowie ungünstiger Betriebsbedingungen zu einer **Degradation** der Nutzsignale führen. So können Röntgen-Flachdetektoren verschiedene das aufgenommene Bild verfälschende Effekte aufweisen, die beispielsweise beim Auslesen der Dünnfilm-TFT-Detektoren durch Übersprecheffekte auf zu Nachbarpixeln gehörende Ladungspakete in der Spaltenleitung entstehen und zu langgezogenen Artefakten führen, die auch als „Kometenartefakte“ bezeichnet werden. Weiterhin können Speichereffekte im Halbleitermaterial von Flachdetektoren dazu führen, dass Reste vorhergehender Aufnahmen in neu aufgenommene Röntgenbilder mit eingehen; diese Artefakte werden als „Geistbildartefakte“ bezeichnet.

### 3.2 Pixeldefekte und Defektkorrektur

Bei Bildsensor-Arrays mit mehreren Millionen Pixelzellen sind in der Regel einzelne Bildpunkte, zusammenhängende Gruppen von Pixeln oder aber komplette Zeilen oder Spalten oder Teile davon defekt. Dies betrifft alle Arten von Bildsensoren, ist aber beispielsweise bei Flachdetektoren für die medizinische Diagnostik von besonderer Relevanz. Bei Röntgen-Detektoren werden aufwändige Verfahren als Software auf leistungsfähigen Rechnern implementiert, um die auftretenden Pixeldefekte zu korrigieren. Würde man nämlich alle Arrays mit Pixeldefekten verwerfen, so wäre eine (aus Herstellersicht) kommerzielle Nutzung aufgrund der extrem hohen Ausschussquote nicht möglich. Deshalb werden nur Arrays mit nicht zufriedenstellend korrigierbaren Defekten ausgesondert; die verbliebenen Defekte werden soweit möglich rechnerisch korrigiert.

Je nach Stabilität der Materialien von Sensor-Arrays können Pixeldefekte nicht nur von Anfang an fertigungsbedingt vorhanden sein, sondern sich auch erst später während der Betriebsphase des Sensors einstellen. Dieses „Driften“, das besonders bei bestimmten Infrarot-Sensoren auftritt, erfordert die regelmäßige Durchführung von Kalibriermaßnahmen. Bei Röntgen-detektoren spielen hingegen anfänglich vorhandene Fertigungsfehler die entscheidende Rolle, so dass defekte Pixel vorab „in der Fabrik“ in einem „Kalibrier-vorgang“ identifiziert und in eine sogenannte Defektpixel-Karte („defective pixel map“) eingetragen werden können. Im Betrieb ist es dann möglich, Werte defekter Pixel in Echtzeit durch Werte einzelner benachbarter Pixel oder durch interpolierte Werte mehrerer einander gegenüber liegender Nachbarpixel zu ersetzen („bad pixel replacement“) oder auch komplexere Korrekturalgorithmen anzuwenden. Doch auch bei Röntgen-detektoren kann die Bildqualität mit der Zeit nachlassen, wobei korrigierbare Fehler durch wiederholte Kalibrierungen während Pausen in der Betriebsphase des Detektors oder auch „im Hintergrund“ während des Betriebs beseitigt werden können. Die Zunahme der Zahl nicht korrigierbarer Fehler kann irgendwann dazu führen, dass einzelne Panels oder der gesamte Detektor ausgetauscht werden müssen oder der Detektor instand gesetzt werden muss ([17] und [18]).

Um defekte Pixel zu identifizieren, ist es im einfachsten Fall möglich, die von ihnen gelieferten Intensitätswerte mit denen benachbarter Pixel zu vergleichen. Treten größere Abweichungen bei einzelnen Bildpunkten auf, als es das Bildsignal (aufgrund der darin enthaltenen Ortsfrequenzen) erwarten lässt, so handelt es sich um „Ausreißer“ („outliers“, „Salz und Pfeffer-Rauschen“). Es ist offensichtlich, dass diese Wertabweichungen grundsätzlich auch andere Ursachen, insbesondere äußere Störungen, haben können als Pixelfehler. Ein Ausreißer kann dann als Fehler identifiziert und herausgefiltert werden. Ein Ersetzen des Pixelwerts durch die Werte benachbarter fehlerfreier Pixel oder durch Interpolationswerte aus Werten der Pixelumgebung stellt dann nichts anderes als eine Art Glättung oder Tiefpassfilterung aufgrund eines Plausibilitätstests dar, wobei der Korrekturwert nur eine „kosmetische“ Näherungslösung ist, die dem Betrachter den Eindruck einer höheren Bildqualität vermittelt. Insbesondere

bei einer maschinellen Weiterverarbeitung und einer automatischen Bildauswertung sollte erwogen werden, einen höheren Aufwand zu treiben, um defekte Pixel sicher zu identifizieren.

Bei bildgebenden Systemen der Röntgendiagnostik sind neben den Defekten des Sensorsystems gegebenenfalls noch Abschattungen der von der Röntgenröhre ausgehenden Strahlung durch ein – insbesondere feststehendes – Streustrahlenraster (Bucky-Blende, „grid“) zu berücksichtigen.

Nachdem „ausgefallene“, „tote“ Pixel kein Signal liefern, das Rückschlüsse auf die eintreffende Strahlung erlaubt, ist bei ihnen auch keine Korrektur des Signals möglich. Bildinformation, die vom Sensor nicht aufgenommen wird, ist verloren und kann somit auch grundsätzlich nicht wiedergewonnen werden. Liefern Pixel dagegen ein schwaches Signal und kann man die Abschwächung ermitteln, so ist eine echte Korrektur möglich. Maßnahmen hierzu sind unter anderem die FPN-Korrektur und die Ungleichförmigkeits-Korrektur (NUC).

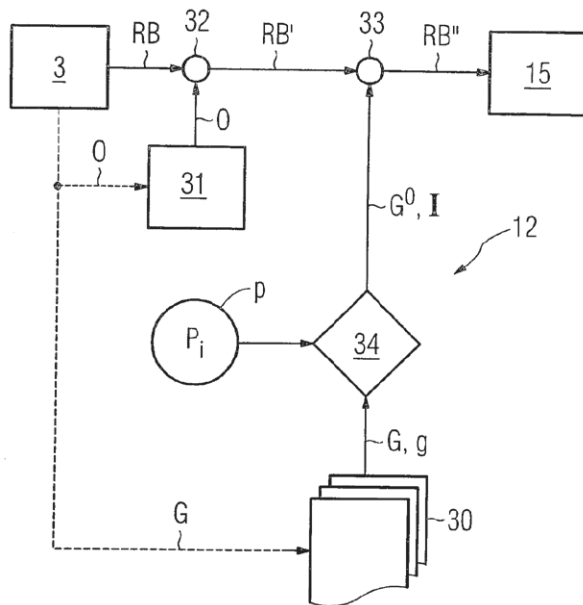
### 3.3 Festmusterrauschen, Sensorkalibrierung und Inhomogenitätskorrektur

Die Korrektur von Inhomogenitäten und Festmusterrauschen wird im Englischen als „non-uniformity correction“ NUC bezeichnet. Die Korrektur des Festmusterrauschens („fixed pattern noise“ FPN) ist grundsätzlich ein Querschnitts-Thema der Bildsensorik, wenn auch mit besonderer Relevanz je nach Art des Detektors. Die einzelnen photosensitiven Pixelzellen einer Bildsensor-Matrix sind aufgrund von Fertigungstoleranzen nicht völlig identisch. So können sich etwa über die Fläche des Sensors hinweg bestimmte Materialeigenschaften kontinuierlich ändern und zu Ungleichförmigkeiten führen. Die Sensoren können auch durch die Herstellungsverfahren bedingte spezielle Muster wie zum Beispiel Streifen aufweisen. Durch Unterschiede beispielsweise in den Spalten-Ausleseschaltungen und Spaltenverstärkern treten zudem typischerweise vertikale Streifen im Bild auf, weil sich die Abweichungen auf alle Pixel der betreffenden Spalte auswirken. Weiterhin gibt es unstetige Änderungen von Parametern der einzelnen Pixel von Pixelzelle zu Pixelzelle. Die nicht vom aufzunehmenden Objekt abhängigen

und somit eine Störung darstellenden Strukturen in den Bildern werden als Festmusterrauschen (FPN) bezeichnet. Dieses dem Sensor immanente stochastische Rauschen ist somit nicht durch zeitliche Schwankungen, die bereits in den aufgenommenen physikalischen Größen, insbesondere der Intensität der zu erfassenden Strahlung, enthalten sein können, sondern durch örtliche Schwankungen gekennzeichnet. Beide Arten des Rauschens mit dem Signal-Rausch-Verhältnis SNR beeinträchtigen die Bildqualität und den nutzbaren Dynamikbereich der Intensität in gleicher Weise.

Theoretisch ist es denkbar, verschiedene Parameter, welche die Eigenschaften eines Bildsensors-Systems bestimmen, messtechnisch zu erfassen und diese zur Kompensation von Sensorfehlern zu verwenden. Beispielsweise hängt die Photosensitivität („photo responsivity“) der einzelnen Bildsensorelemente von einer Vielzahl von Parametern ab, von denen einige vorab identifiziert und einer Kompensation zugeführt werden können; es ist aber in der Praxis unmöglich, dies für alle Parameter durchzuführen. Eine solche analytische Methode kann nämlich bei Berücksichtigung einer steigenden Zahl von Parametern mathematisch sehr aufwändig werden und trotzdem unvollständig bleiben. Dann verbleibt auch nach einer solchen Kompensation ein Bild mit einem überlagerten, von unbekannten Umgebungsparametern abhängigen Muster als Festmusterrauschen FPN. Somit beschränkt man sich bei der analytischen Vorgehensweise zumeist auf wenige einflussreiche Parameter, wie die Temperatur, die Belichtungszeit und in der Röntgendiagnostik die Beschleunigungsspannung der Röntgenröhre, und verwendet diese Parameter in einem üblichen NUC-Verfahren. Die DE 103 43 496 A1 offenbart im Hinblick auf digitale Röntgendetektoren, dass typische Parameter von Röntgendiagnostik-Systemen wie beispielsweise das Röntgenspektrum oder die Generatorspannung verwendet werden, um ein passendes Gain-Bild aus einer Menge bereitstehender Gain-Bilder auszuwählen. Der in Figur 14 vorgestellte Algorithmus beinhaltet eine Kalibrierungsphase vor dem laufenden Betrieb oder im Hintergrund des laufenden Betriebs einer Röntgenvorrichtung und eine Korrekturphase während des laufenden Betriebs. Das Röntgenbild RB wird sowohl mit einem Offset-Bild O (Aufnahme ohne Röntgen-

strahlung) als auch mit einem Gain-Bild  $G^0$  verknüpft. Das Gain-Bild wird in Abwesenheit des Messobjekts aufgenommen und spiegelt die variierende Detektoreffizienz der verschiedenen Sensorflächen wieder.



Figur 14: Algorithmus mit Kalibrierungsphase und Korrekturphase (aus DE 103 43 496 A1)

Ein additiver, örtlich von Pixelzelle zu Pixelzelle schwankender „Offset“, insbesondere auch ein Offset zwischen Pixeln bei „Dunkelheit“, ist die wichtigste Komponente des FPN. Diese Komponente wird auch als DSNU (dark signal non-uniformity) bezeichnet. Insbesondere aufgrund des sogenannten Dunkelstroms („dark current“) liefern die einzelnen Pixelzellen meist auch ohne Belichtung ein temperaturabhängiges und davon abgesehen ein idealerweise zeitlich konstantes Signal. Eine naheliegende Methode der Korrektur der additiven Fehler ist die rechnerische Offset-Korrektur durch Subtraktion eines Dunkelbilds („dark frame“). Eine Korrektur des Dunkelstroms beziehungsweise des „Dunkelrauschens“ wird oft dadurch erreicht, dass bestimmte Pixelbereiche („dark pixel“, „dark photo-sensors“, „dummy photosensors“) als Dunkelreferenz-zonen abgeschirmt und so vor Einfall von Strahlung geschützt werden („shielding“), so dass der Dunkelstrom unmittelbar den gelieferten Messwert der abgedunkelten Pixelzellen bestimmt und dieser direkt in die Korrektur der Nutzpixel-Messwerte einfließen kann. Die DE 199 15 851 A1 beschreibt ein Korrekturverfahren

für eine Pixelmatrix, insbesondere einen aus mehreren „Kacheln“ mittels „Butting“ zusammengesetzten Festkörperbildsensor für Röntgenstrahlung, mit einer die aktive Pixelfläche rundherum umgebenden Dunkelreferenzzone DRZ, wobei anhand von Signalen der Pixel der Dunkelreferenzzone Offset-Korrekturwerte ermittelt werden (siehe Figur 11).

Bei vielen Pixeldesigns, besonders auch solchen, die unmittelbar auf einem Photostrom basieren, ist das Messsignal – abgesehen von einem Offset – proportional zur Lichtintensität. Somit ist die Empfindlichkeit eines einzelnen Pixels konstant und die Kennlinie des Pixels ist demnach linear. Aufgrund von beispielsweise Fertigungstoleranzen in den Pixelzellen selbst variiert auch die Empfindlichkeit örtlich von Pixelzelle zu Pixelzelle. Im Zusammenhang mit der NUC wird im Englischen oft auch von „Gain“ – in Ergänzung zum „Offset“ – gesprochen. Als Ursache sind beispielsweise Variationen der photosensitiven Fläche der Photodioden von Pixel zu Pixel und speziell bei indirekten Röntgen-Flachdetektoren zum Beispiel die Ungleichmäßigkeiten des in Form von Kristallen aufgewachsenen Szintillators und der Klebstoffschichten zu berücksichtigen.

Bei bildgebenden Systemen der Röntgendiagnostik kommen zu den Ungleichmäßigkeiten des Sensorsystems noch winkelabhängige Abschwächungen und Ungleichförmigkeiten der von der Röntgenröhre ausgehenden Strahlung aufgrund des sogenannten Heel-Effekts hinzu.

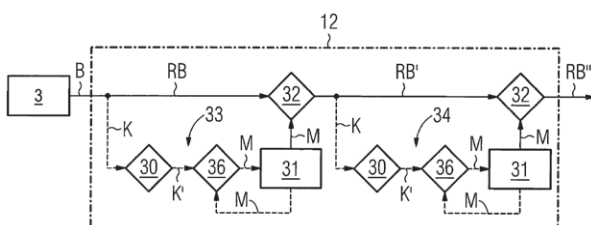
Die Komponente des FPN, die die Variation der Ansprechempfindlichkeit (responsivity variation) zwischen Pixeln unter Bestrahlung betrifft, wird auch als PRNU (photo response non-uniformity, auch pixel response non-uniformity) bezeichnet.

Stellt die DSNU insbesondere aufgrund des Dunkelstroms den Hauptanteil des FPN, so kann die Offset-Korrektur ausreichend sein. Eine Verbesserung der Ergebnisse ergibt sich in vielen Fällen, wenn neben der Offset-Korrektur noch eine Gain-Korrektur erfolgt. Bei linearen Pixelkennlinien, bei denen sich das FPN ausschließlich aus DSNU und PRNU zusammensetzt, ist so eine vollständige NUC möglich.

Vor der fortlaufenden Korrektur im Betrieb ist die „Kalibrierung“ des Bildsensors erforderlich, bei der Korrekturwerte für jedes Pixel bestimmt und in einem beispielsweise als Tabelle („table“) oder Karte („map“) bezeichneten Speicher abgelegt werden.

Soll nur eine Offset-Korrektur durchgeführt werden, so kann beispielsweise durch eine Abdunkelung der Pixel die Belichtung sämtlicher Pixel unterbunden werden, was eine (spezielle) homogene Belichtung darstellt. Dabei können die Offset-Korrekturwerte als Dunkelreferenzbild ermittelt und zur späteren Subtraktion in die „Offset Map“ geschrieben werden. Man spricht in diesem Fall auch von einer „Einpunkt-Kalibrierung“. Soll eine Offset- und eine Gain-Korrektur durchgeführt werden, so müssen dem Sensor zwei homogene Szenen unterschiedlicher Intensität dargeboten werden. Dabei handelt es sich um die fachübliche „Zweipunkt-Kalibrierung“ [19].

Ein Beispiel für komplexere Kalibrieralgorithmen zeigt Figur 15 aus der DE 103 43 787 B4. Diese Druckschrift stellt ein spezielles Verfahren zur Gain-Kalibrierung 34 und/oder Offset-Kalibrierung 33 digitaler Röntgen-Flachdetektoren dar, wobei die Kalibrierbilder K mit einem Glättungsfilter 30 bearbeitet werden.



Figur 15: Verfahren zur Kalibrierung eines digitalen Röntgendetektors mit Offset- und Gain-Korrektur (aus DE 103 43 787 B4)

#### 4 Fazit

Am Beispiel der hier vorgestellten Röntgen-Flachdetektoren zeigt sich, wie die technologische Entwicklung Nutzen für den Menschen bringen kann, beispielsweise in Form minimierter Strahlenbelastung bei medizinischen Untersuchungen, und gleichzeitig volkswirtschaftlich vorteilhaft sein kann. Eine wesentliche Triebfeder hierbei ist der Fortschritt in der Halbleitertechnologie, der sich sowohl bei mikroelektronisch

aufgebauten Detektoren selbst zeigt als auch in der Bereitstellung hochintegrierter Prozessoren, die es beispielsweise ermöglichen, komplexe Korrekturalgorithmen für die Bildsensorik in Echtzeit auszuführen.

#### Nicht-Patentliteratur

- [1] LOSSAU, Norbert: Röntgen verzichtete auf ein Patent. In: DIE WELT, 03.12.2001
- [2] KRESTEL, Erich (Hrsg.): Bildgebende Systeme für die medizinische Diagnostik. 2. Auflage, Berlin; München: Siemens AG, 1988. – ISBN 3-8009-1505-7
- [3] SPAHN, M.; HEER, V.; FREYTAG, R.: Flachbild-detektoren in der Röntgendiagnostik. In: Der Radiologe, Bd. 43, Nr. 5 (2003), S. 340–350
- [4] HERTRICH, Peter H.: Röntgenaufnahmetechnik. Erlangen: Publicis Corporate Publishing, 2004. – ISBN 3-89578-209-2
- [5] STROTZER, Michael; VÖLK, Markus; FEUER-BACH, Stefan: Flachdetektoren in der digitalen Radiographie. In: Deutsches Ärzteblatt, Jg. 99, Heft 38, 20. Sept. 2002, S. A 2484–A 2488
- [6] DÖSSEL, Olaf: Bildgebende Verfahren in der Medizin – Von der Technik zur medizinischen Anwendung. Berlin; Heidelberg: Springer, 2000. – ISBN 978-3-662-06047-6
- [7] JORDEN, Paul R.; MORRIS, David G.; POOL, Peter J.: Technology of Large Focal Planes of CCDs. In: Proc. SPIE 5167: Focal Planes for Space Telescopes, January 2004, S. 72–82
- [8] LEWICKI, Andreas: Einführung in die Mikroelektronik. München; Wien: R. Oldenbourg, 1966.
- [9] WEIMER, Paul K.: The TFT – A New Thin-Film Transistor. In: Proceedings of the IRE, Vol. 50, No. 6, June 1962, S. 1462–1469
- [10] KUO, Yue: Thin Film Transistor Technology – Past, Present, and Future. In: The Electrochemical Society – Interface, Spring 2013, S. 55–61
- [11] CULLIS, Roger: Patents, Inventions and the Dynamics of Innovation: A Multidisciplinary Study. Cheltenham, UK; Northampton, MA, USA: Edward Elgar Publ., 2007. – ISBN 978-1-84542-958-4



- [12] LE COMBER, P. G.; SPEAR, W. E.; GHAITH, A.: Amorphous Silicon Field-Effect Device and Possible Application. *Electronic Letters*, Vol. 15, No. 6, 15th March 1979, S. 179–181
- [13] NEUDECK, Gerold W.; MALHOTRA, Arun K.: An Amorphous Silicon Thin Film Transistor: Theory and Experiment. In: *Solid State Electronics*, 1976, Vol. 19, S. 721–729
- [14] WECKLER, Gene P.: A Silicon Photodevice to Operate in a Photon Flux Integrated Mode. In: *International Electron Devices Meeting, IEDM 1965*, S. 38–39
- [15] WECKLER, Gene P.: Operation of p-n Junction Photodetectors in a Photon Flux Integrating Mode. In: *IEEE Journal of Solid-State Circuits*, Vol. SC-2, No. 3, September 1967, S. 65–73
- [16] LÜDER, Ernst: *Bau hybrider Mikroschaltungen. Einführung in die Dünn- und Dickschichttechnologie*. Berlin; Heidelberg; New York: Springer, 1977. – ISBN 978-3-540-08289-7
- [17] LÓPEZ-ALONSO, José Manuel; ALDA, Javier: Bad Pixel Identification by Means of Principal Components Analysis. In: *Optical Engineering* 41(9), September 2002, S. 2152–2157
- [18] CAREY, S. [et al.]: Stability of the Infrared Array Camera for the Spitzer Space Telescope. In: *Proc. SPIE Vol. 7010*, 2008, 70102V
- [19] HOLST, Gerald C.: *CCD Arrays, Cameras and Displays*. 2. Auflage. SPIE Press, 1998. – ISBN 978-0819428530

# Talbot-Lau-Interferometer zur Phasenkontrast-Bildgebung mit Röntgenstrahlen

Dr. Florian Siebel, Patentabteilung 1.52

Die Bildgebung mit Röntgenstrahlen ist eine Standardmethode in der Medizin und der zerstörungsfreien Prüfung. Dabei macht man sich zunutze, dass Röntgenstrahlen beim Durchtritt durch Materie teilweise absorbiert werden. Darüber hinaus beeinflusst Materie aber auch die Phaseneigenschaften der Röntgenstrahlen. Unter Verwendung von Röntgensensoren und von speziellen Gitterstrukturen gelingt es in Talbot-Lau-Interferometern, die Phaseneigenschaften der Röntgenstrahlen zu messen und zusätzlich zu den klassischen Röntgenbildern komplementäre Phasenkontrast-Bilder der inneren Struktur der untersuchten Materie zu erzeugen.

---

## 1 Einleitung

Die Entdeckung hochenergetischer Strahlen im Jahre 1895 durch Wilhelm Conrad Röntgen kann als Meilenstein für die Medizin und die zerstörungsfreie Prüfung angesehen werden, die Röntgen 1901 den ersten Nobelpreis für Physik einbrachte. Mit Röntgenstrahlen ist es möglich, für gewöhnliches Licht undurchsichtige Objekte zu durchstrahlen und bei räumlich variablem Absorptionsvermögen der Objekte Bilder der inneren Struktur dieser Objekte zu erzeugen. Bei Röntgenstrahlen handelt es sich um elektromagnetische Wellen, die im Vergleich zu gewöhnlichem Licht, das heißt elektromagnetischen Wellen im sichtbaren Wellenlängenbereich, weitaus kurzwelliger und damit energiereicher sind. Wie bei elektromagnetischen Wellen allgemein lässt sich die Wechselwirkung von Röntgenstrahlen mit einem Objekt durch einen komplexen Brechungsindex  $n = 1 - \delta + i\beta$  des Objekts beschreiben. Der Imaginärteil  $\beta$  beschreibt das Absorptionsvermögen des Objekts und hängt stark von der Ordnungszahl des durchstrahlten Materials ab. Das Dekrement  $\delta$  des Realteils des Brechungsindex beeinflusst die Phase der elektromagnetischen Wellen und ist im Wellenlängenbereich von Röntgenstrahlen in der Regel betragsmäßig sehr viel größer als der Imaginärteil  $\beta$ . Durch Auswertung dieser Phaseninformation lassen sich prinzipiell auch für Materialien mit ähnlichen Ordnungszahlen kontrastreiche Bilder erzeugen.

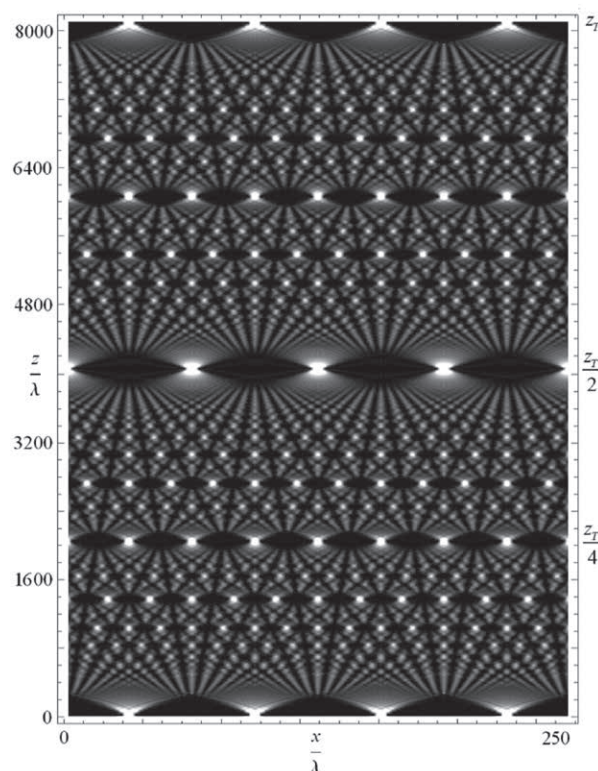
Das Prinzip der Phasenkontrast-Bildgebung geht auf Frits Zernike zurück, der für seine Entwicklungen eines Phasenkontrast-Mikroskops für sichtbares Licht ab etwa 1930 im Jahre 1953 den Nobelpreis für Physik erhielt [1]. Wegen fehlender röntgenoptischer Komponenten dauerte es jedoch noch einige Zeit, bis die Phasenkontrast-Bildgebung vom sichtbaren Wellenlängenbereich auf Röntgenstrahlen übertragen werden konnte. Als Pionierleistung auf diesem Gebiet gelten die Entwicklungen von Ulrich Bonse und Michael Hart [2], die auch in der Patentschrift US 3 446 961 A beschrieben sind. Beim Bonse-Hart-Interferometer wird ein Röntgenstrahl in der Art eines Mach-Zehnder-Interferometers mit Hilfe eines Kristalls in zwei Röntgenstrahlen aufgeteilt, die unterschiedliche optische Wege durchlaufen, bevor sie auf einem Analysator-Kristall interferieren.

Inzwischen existieren mehrere unterschiedliche Methoden zur Phasenkontrast-Bildgebung mit Röntgenstrahlen. Dieser Artikel beschränkt sich auf das Verfahren der gitterbasierten Phasenkontrast-Bildgebung und speziell auf Talbot-Lau-Interferometer. Bei dieser Art von Interferometern werden, im Gegensatz zum Bonse-Hart-Interferometer, mit Hilfe von Gittern die Röntgenstrahlen auf leicht versetzten Pfaden eines gemeinsamen optischen Weges geführt („shearing interferometer“), um Phasendifferenzen zwischen diesen versetzten Pfaden messen zu können. Im Folgenden

werden die zu Grunde liegenden optischen Effekte näher beschrieben.

### 1.1 Talbot-Effekt

Der Talbot-Effekt beruht auf der Nahfeld-Beugung (Fresnel-Beugung) an einem Beugungsgitter. Er wurde 1836 von Henry Fox Talbot für Sonnenlicht entdeckt [3] [4]. Bestrahlt man ein Gitter mit kohärenten Strahlen, so ergeben sich hinter dem Gitter in bestimmten Abständen (ganzzahligen Vielfachen des so genannten Talbot-Abstandes  $z_T$ ) Abbildungen, deren Helligkeitsverteilung der Struktur des Gitters entspricht (Selbstbilder). Für eine planare Strahlung und unter der Annahme, dass die Wellenlänge  $\lambda$  der Strahlung klein gegenüber der Gitterkonstanten  $d$  des Gitters ist, ergibt sich der Talbot-Abstand  $z_T$  zu  $z_T = 2d^2/\lambda$ . Darüber hinaus entstehen auch in Bruchteilen des Talbot-Abstandes regelmäßige Muster.



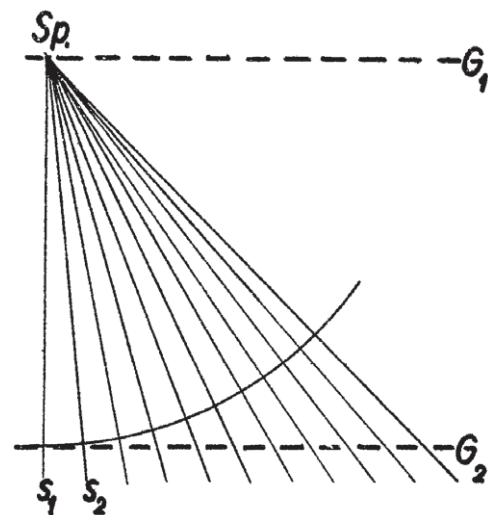
Figur 1: Der optische Talbot-Effekt für monochromatisches Licht (aus Wikipedia [5])

Figur 1 zeigt den Talbot-Effekt für monochromatische Strahlung im sichtbaren Wellenlängenbereich. Das

am unteren Rand der Figur angeordnete, aber nicht explizit gezeigte Gitter wird von unten mit kohärenter Strahlung bestrahlt. In Abständen  $z = z_T$ ,  $z = z_T/2$ ,  $z = z_T/4$ , ... hinter dem Gitter ergeben sich dabei regelmäßige Interferenzmuster.

### 1.2 Talbot-Lau-Effekt

Wie der Talbot-Effekt ist der Talbot-Lau-Effekt ein Beugungseffekt. Er geht auf den Physiker Ernst Lau zurück und wurde noch vor der wissenschaftlichen Veröffentlichung [6] in der Patentschrift DE 747 545 A beschrieben [7].



Figur 2: Parallelanordnung zweier Gitter G1 und G2 zur Erzeugung des Talbot-Lau-Effekts (aus DE 747 545 A)

Ordnet man gemäß Figur 2 zwei Liniengitter  $G_1$  und  $G_2$  so hintereinander an, dass ihre Linienspalte parallel zueinander ausgerichtet sind, und bestrahlt diese Gitteranordnung so mit Licht, dass die Lichtstrahlen zuerst durch die Linienspalte des ersten Gitters  $G_1$  fallen, dann entstehen analog zum Talbot-Effekt hinter dem zweiten Gitter  $G_2$  in geeigneten Abständen (Vielfachen des Talbot-Abstandes) Selbstbilder des zweiten Gitters  $G_2$  beziehungsweise Interferenzmuster. Dabei interferieren Strahlen unterschiedlicher Strahlrichtung (zum Beispiel  $S_1$  und  $S_2$ ), die durch einen gemeinsamen Spalt (zum Beispiel  $S_p$ ) des ersten Gitters  $G_1$  getreten sind, wegen des Gangunterschieds hinter dem zweiten Gitter  $G_2$ . Im Unterschied zum Talbot-Effekt sind die Anforderungen an die Kohärenz der Lichtquelle sehr viel

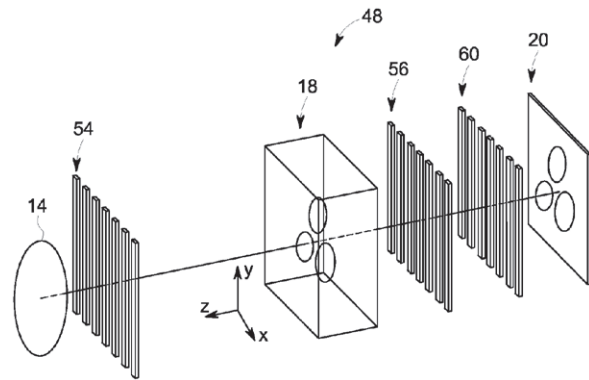
geringer, da die einzelnen Linienspalte des ersten Gitters  $G_1$  eine räumliche Kohärenz der von den Linienspalten austretenden Lichtstrahlen sicherstellen. In der Patentschrift DE 747 545 A wird für Wellenlängen im sichtbaren Bereich bereits vorgeschlagen, durch eine Anordnung eines Stoffes zwischen den Gittern  $G_1$  und  $G_2$  mit Hilfe der Interferenzmuster den Brechungsindex  $n$  des Stoffes zu vermessen und somit Phaseninformationen der Strahlen auszuwerten.

## 2 Talbot-Lau-Interferometer für Röntgenstrahlen

Nach der Einführung in die grundlegenden Effekte, die historisch für sichtbares Licht untersucht wurden, wird im Folgenden der Standard-Aufbau eines Talbot-Lau-Interferometers für Röntgenstrahlen beschrieben. Wegen der typischen Wellenlängen von Röntgenstrahlen benötigt man für ein solches Interferometer entsprechend feinere Gitterstrukturen. Um beispielsweise für Röntgenstrahlen mit einer Wellenlänge von  $\lambda = 10^{-10}$  m ein Selbstbild eines Beugungsgitters im Talbot-Abstand  $z_T = 0,1$  m zu erhalten, muss das Beugungsgitter eine Gitterperiode  $d$  von etwa  $2,2 \mu\text{m}$  besitzen.

### 2.1 Standard-Aufbau eines Talbot-Lau-Interferometers für Röntgenstrahlen

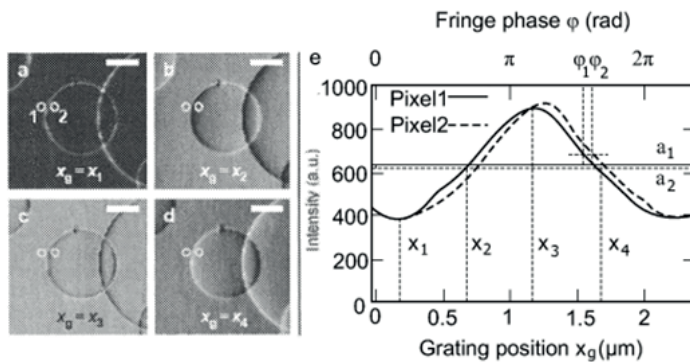
Wie aus Figur 3 (entnommen der US 2014/0169524 A1) ersichtlich ist, entspricht der Standard-Aufbau eines Talbot-Lau-Interferometers für Röntgenstrahlen weitgehend dem Aufbau von Ernst Lau. Er umfasst zur Untersuchung eines Objekts 18 eine inkohärente Röntgenquelle 14, ein Absorptionsgitter 54, auch Quellengitter genannt, nahe der Röntgenquelle zur Erzeugung von räumlich kohärenten Einzelstrahlen sowie ein Beugungsgitter 56. Mit Hilfe der Kombination aus inkohärenter Röntgenquelle 14 und Quellengitter 54 gemäß Ernst Lau wird das zu untersuchende Objekt 18 mit räumlich kohärenten Röntgenstrahlen durchstrahlt, so dass sich hinter dem Beugungsgitter 56 in Vielfachen des Talbot-Abstandes gemäß dem Talbot-Effekt Beugungsstrukturen (Interferenzmuster) ergeben, die Phaseninformationen und damit Informationen über den Realteil des Brechungsindex des Objekts 18 enthalten.



Figur 3: Standard-Aufbau eines Talbot-Lau-Interferometers für Röntgenstrahlen (aus US 2014/0169524 A1)

Da das typische Auflösungsvermögen eines Röntgensensors beziehungsweise Röntgendetektors 20 in der Regel nicht ausreicht, die entstehenden Beugungsstrukturen in der Größenordnung von  $1 \mu\text{m}$  aufzulösen, wird zusätzlich ein weiteres Absorptionsgitter 60 mit einer Gitterperiode in der Größe der Beugungsstrukturen eingefügt. Denn durch das zusätzliche Absorptionsgitter 60, auch Analysorgitter genannt, werden Moiré-Muster hinter diesem Gitter erzeugt, die wegen der den Beugungsstrukturen ähnlichen Periodizität sehr viel größer sind und mit einem gewöhnlichen Röntgendetektor 20 aufgenommen werden können. Zur Bestimmung der Phase wird beim so genannten Phase-Stepping-Verfahren das Absorptionsgitter 60 mehrfach um Bruchteile einer Gitterperiode transversal zu den Gitterlinien versetzt (in x-Richtung in Figur 3) und mit dem Röntgendetektor 20 jeweils ein Bild aufgenommen. Durch das leichte Versetzen des Absorptionsgitters 60 relativ zu den Beugungsstrukturen im Phase-Stepping-Verfahren verändern sich die entstehenden Moiré-Muster, wodurch man Informationen über die Phase erhalten kann. Dies wird anhand von Figur 4 aus der WO 2011/070489 A1 verdeutlicht.

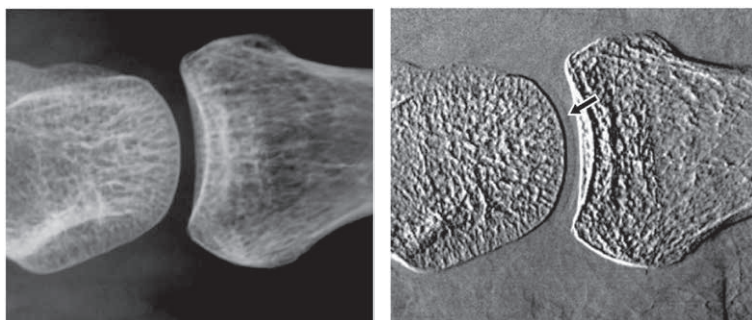
Auf der linken Seite der Figur 4 sind vier mit dem Röntgendetektor 20 gewonnene Abbildungen für vier verschiedene Positionen  $x_g = x_1, \dots, x_4$  des Absorptionsgitters 60 gezeigt. Trägt man die für die Detektorpixel 1 und 2 (Figur 4, Bild a links) gewonnenen Intensitätswerte als Funktion der Position des Absorptionsgitters 60 auf, so ergeben sich periodische Verläufe, aus deren



Figur 4: Phase-Stepping-Verfahren für ein Talbot-Lau-Interferometer  
Links: Bilder a, ..., d eines Objekts bestehend aus Polystyren-Kugeln für vier verschiedene Positionen des Absorptionsgitters 60  
Rechts: Gemessene Intensität des Röntgendetektors als Funktion der Gitterposition des Absorptionsgitters 60 für zwei unterschiedliche Detektorpixel (aus WO 2011/070489 A1)

Versatz sich die relativen Phasen  $\phi_1$  und  $\phi_2$  für die Detektorpixel 1 und 2 bestimmen lassen. Gleichzeitig lassen sich durch Mittelwertbildung des Intensitätsverlaufs über eine Periode  $2\pi$  die konventionellen, auf der Absorption beruhenden Intensitätswerte  $a_1$  und  $a_2$  bestimmen. Für eine genauere Beschreibung des Phase-Stepping-Verfahrens wird auf [8] verwiesen. Da die gemessenen relativen Phasen  $\phi_1$  und  $\phi_2$  proportional zum Gradienten des Phasenprofils der Wellenfronten sind, spricht man allgemein auch von einem differentiellen Phasenkontrast-Verfahren.

Figur 5 aus der JP 2014-117485 A zeigt ein Anwendungsbeispiel. Während auf der linken Seite der Figur ein konventionelles Absorptionskontrast-Bild der Verbindung zweier Handknochen dargestellt ist, zeigt die rechte Seite ein differentielles Phasenkontrast-Bild desselben



Figur 5: Absorptionskontrast-Bild (links) und differentielles Phasenkontrast-Bild (rechts) der Verbindung zweier Handknochen / Nur im differentiellen Phasenkontrast-Bild ist der Knorpel zwischen den Knochen zu erkennen (siehe Pfeil) (aus JP 2014-117485 A)

Ausschnitts. Im differentiellen Phasenkontrast-Bild sind zusätzliche Details zu erkennen.

Darüber hinaus lassen sich neben den zueinander komplementären Informationen über den Absorptionskontrast und den Phasenkontrast mit dem Standard-Aufbau eines Talbot-Lau-Interferometers zusätzlich Informationen über Streustrahlung, so genannte Dunkelfeldbilder, gewinnen [9].

## 2.2 Frühe Entwicklungen

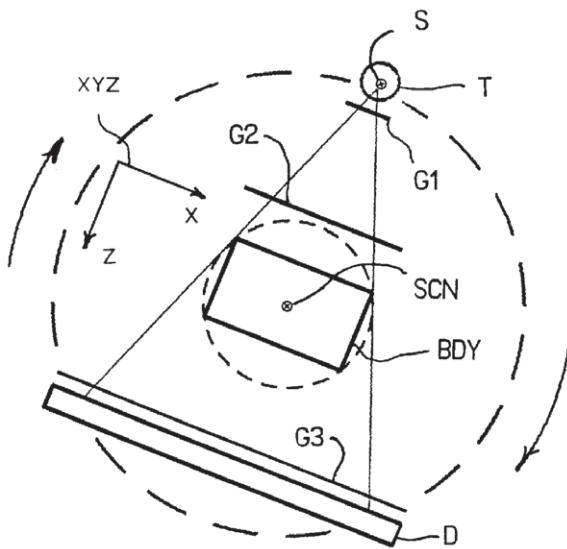
In den frühen Entwicklungen gitterbasierter Interferometer wurden in der Regel spezielle Röntgenquellen verwendet, die ein ausreichendes Maß an Kohärenz sicherstellen, so dass auf ein Quellengitter nach Ernst Lau verzichtet werden kann (zum Beispiel [10], [11], [12], EP 1 623 671 A1). Solche Interferometer ohne Quellengitter werden Talbot-Interferometer genannt. Aufbauend auf diesen Entwicklungen wurden 2006 Untersuchungen mit einem Talbot-Lau-Interferometer für Röntgenstrahlen beschrieben [13]. Diese werden in der wissenschaftlichen Literatur allgemein als Durchbruch auf dem Gebiet angesehen, da bei einem solchen Talbot-Lau-Interferometer nun gewöhnliche Röntgenröhren verwendet werden können. Weniger bekannt ist hingegen (vergleiche aber zum Beispiel [14]), dass Talbot-Lau-Interferometer für Röntgenstrahlen, wie unten näher ausgeführt, schon 1998 patentiert wurden

(US 5 812 629 A) und bereits 1992 vorgeschlagen wurden [15]. Darüber hinaus sind Talbot-Lau-Interferometer für Röntgenstrahlen in den Offenlegungsschriften EP 1 447 046 A1 und EP 1 731 099 A1 beschrieben.

Die Patenschrift US 5 812 629 A des Physikers John F. Clauser [16] aus dem Jahr 1998 beschreibt sehr detailliert die wesentlichen Aspekte eines Talbot-Lau-Interferometers für Röntgenstrahlen.



Neben radiographischen Phasenkontrast-Abbildungen lassen sich auch dreidimensionale Phaseninformationen eines Objekts gewinnen. Figur 6 zeigt ein Computertomografiegerät bestehend aus einer Röntgenröhre T als Strahlenquelle S und Gittern G1, G2 und G3 sowie einem Detektor D, wobei das Computertomografiegerät um eine Achse SCN gedreht wird. Damit lassen sich dreidimensionale Bilder der inneren Struktur des Körpers BDY gewinnen. Wie Figur 6 weiter zeigt, kann der zu untersuchende Körper BDY im Gegensatz zur Ausführung nach Figur 3 auch zwischen dem Beugungsgitter G2 und dem Analysatorgitter G3 angeordnet sein.



Figur 6: Computertomografiegerät mit einem Talbot-Lau-Interferometer (aus US 5 812 629 A)

## 2.3 Jenseits des Standard-Aufbaus

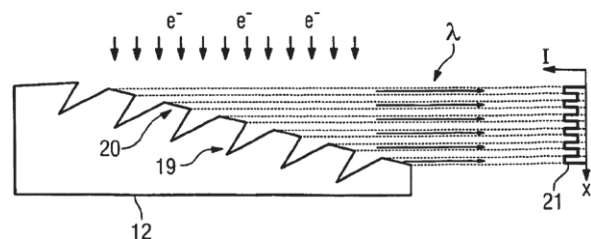
Im Folgenden wird anhand von Patentdokumenten auf weitere Entwicklungen im Umfeld von Talbot-Lau-Interferometern eingegangen. Wegen der Fülle der Patentanmeldungen in den letzten Jahren auf diesem Gebiet können hier nur ausgewählte Aspekte exemplarisch dargestellt werden.

### 2.3.1 Strukturierte Röntgenquellen

Das Quellengitter eines Talbot-Lau-Interferometers hat die Verwendung von herkömmlichen inkohärenten Röntgenröhren mit einem relativ großen Brennfleck

erst möglich gemacht. Problematisch ist dabei jedoch, dass ein solches Quellengitter bereits einen großen Teil der Strahlung einer Röntgenröhre absorbiert, so dass für reale Anwendungen die verfügbare Intensität der Röntgenstrahlen nicht immer ausreicht [14]. Daher wurde auch die Verwendung spezieller Röntgenquellen vorgeschlagen.

Beispielsweise kann gemäß der EP 1 803 398 B1 die Gitterstruktur des Quellengitters durch eine gitterförmige Anordnung des Elektronenstrahlbrennflecks auf der Anode der Röntgenröhre ersetzt werden. In dem in Figur 7 gezeigten Ausführungsbeispiel wird dies durch Einkerbungen 19 auf der Oberfläche der Anode erreicht. Die auf die Anode treffenden Elektronen  $e^-$  erzeugen Röntgenstrahlung  $\lambda$ , wobei die Plateaus 20 zwischen den Einkerbungen 19 zu einer erhöhten Intensität der entstehenden Röntgenstrahlung  $\lambda$  führen, wie das Intensitätsprofil 21 zeigt. Darüber hinaus wurden auch verteilte Anordnungen einzeln ansteuerbarer Röntgenquellen entwickelt.



Figur 7: Anode einer Röntgenröhre mit Einkerbungen 19 zwischen Plateaus 20, was zu dem räumlich variablen Intensitätsverlauf 21 von Röntgenstrahlen  $\lambda$  führt (aus EP 1 803 398 B1)

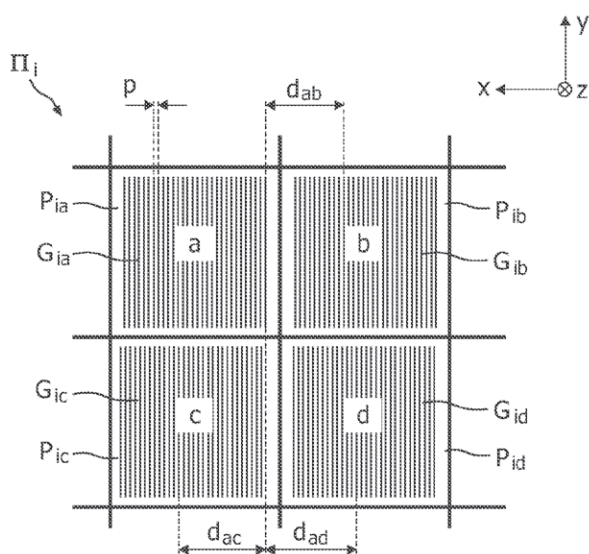
### 2.3.2 Zweidimensionale Gitterstrukturen

Während im Standardaufbau eines Talbot-Lau-Interferometers gemäß Figur 3 eindimensionale Gitter (Linien-gitter) Verwendung finden, lässt sich darüber hinaus die Theorie auch auf zweidimensionale Gitter anwenden, wie in der Patentschrift US 5 812 629 A beschrieben wird. Eine Vielzahl von Patentdokumenten beschäftigt sich mit der Herstellung und Geometrie der Gitter.

### 2.3.3 Analysatorgitter und Detektor

Wie in der US 5 812 629 A dargelegt wird, lässt sich die räumlich periodische Struktur des Analysatorgitters im Zusammenwirken mit einem beliebigen Detektor auch durch einen geeigneten räumlich periodischen Detektor ersetzen.

Das im Standardaufbau eines Talbot-Lau-Interferometers nach Figur 3 verwendete Phase-Stepping-Verfahren zur Abtastung der Phaseninformationen stellt in der Praxis sehr hohe Anforderungen an die Genauigkeit der Mechanik. Denn das Analysatorgitter muss mehrfach um Bruchteile der Gitterperiode in der Größenordnung von  $1\ \mu\text{m}$  versetzt werden. Zudem erlauben medizinische Anwendungen oft keine verlängerten Aufnahmezeiten, so dass das Phase-Stepping-Verfahren als ein wesentliches Hindernis für die Anwendung von Talbot-Lau-Interferometern in der Medizin angesehen wurde. In der Patentschrift US 8 576 983 B2 wird daher vorgeschlagen, das zeitliche Abtasten im Phase-Stepping-Verfahren durch ein räumliches Abtasten mit einer geeigneten Kombination von Analysatorgitter und Detektor zu ersetzen. Figur 8 zeigt ein aus vier Detektorpixeln  $P_{ia}$  bis  $P_{id}$  bestehendes Makropixel  $\Pi_i$ , wobei vor jedem der Detektorpixel  $P_{ia}$  bis  $P_{id}$  ein eigenes Analysatorgitter  $G_{ia}$  bis  $G_{id}$  ange-



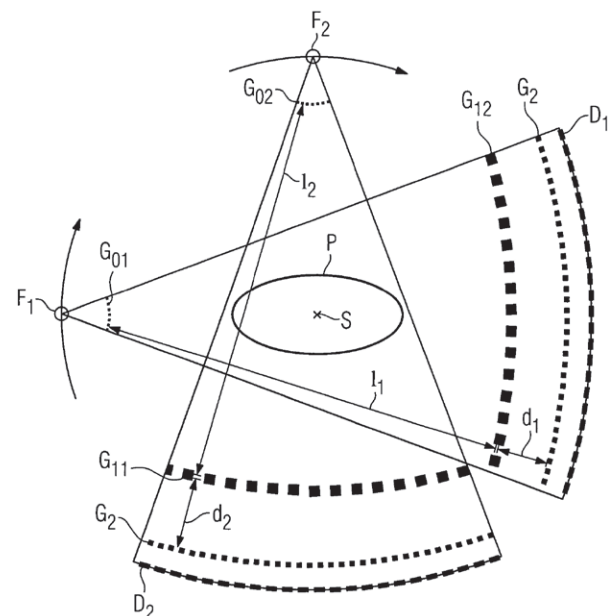
Figur 8: Makropixel  $\Pi_i$  aus einer Kombination von Detektorpixeln  $P_{ia}$  bis  $P_{id}$  und Analysatorgittern  $G_{ia}$  bis  $G_{id}$ . Diese Analysatorgitter besitzen alle dieselbe Periode  $p$ , unterscheiden sich aber in ihrer relativen Phase (aus US 8 576 983 B2)

ordnet ist. Diese Analysatorgitter  $G_{ia}$  bis  $G_{id}$  besitzen alle dieselbe Periode  $p$ , unterscheiden sich aber in ihrer relativen Phase. Dabei sind die Abstände  $d_{ab}$ ,  $d_{ac}$  und  $d_{ad}$  zwischen den Analysatorgittern  $G_{ia}$  bis  $G_{id}$  so gewählt, dass sie, modulo der Periode  $p$ , gleichmäßig über die Periode verteilt sind. Somit lassen sich mit einer einzigen Aufnahme die relevanten Phaseninformationen gewinnen. Das Phase-Stepping-Verfahren kann in diesem Fall entfallen.

### 2.3.4 Dual-Energy-Verfahren für die Phasenkontrast-Bildgebung

Wie die einfache Formel für den Talbot-Abstand zeigt, hängt die Lage der hinter dem Beugungsgitter entstehenden Interferenzmuster von der Wellenlänge der Röntgenstrahlung, das heißt von der Energie der Röntgenstrahlung ab. Für die Bestrahlung eines Objekts mit Röntgenstrahlung unterschiedlicher Energien muss daher im Allgemeinen die Lage der Gitter speziell auf diese Energien angepasst werden.

Figur 9 aus der DE 10 2006 015 356 A1 zeigt ein Computertomografiegerät mit zwei um  $90^\circ$  zueinander versetzten Talbot-Lau-Interferometern. Die zwei Talbot-Lau-Interferometer bestehen aus Röntgenquellen



Figur 9: Computertomografiegerät mit zwei um  $90^\circ$  zueinander versetzten Talbot-Lau-Interferometern mit unterschiedlichen Gittersätzen (aus DE 10 2006 015 356 A1)

mit Fokuspunkten  $F_1$  und  $F_2$ , Quellengittern  $G_{01}$  und  $G_{02}$ , Beugungsgittern  $G_{11}$  und  $G_{12}$ , Analysatorgittern  $G_2$  sowie Detektoren  $D_1$  und  $D_2$ . Dabei sind die Abstände  $l_1$  und  $l_2$  zwischen den Quellengittern und den Beugungsgittern sowie die Abstände  $d_1$  und  $d_2$  zwischen den Beugungsgittern und den Analysatorgittern an die jeweilige Energie der Röntgenquellen angepasst. Mit einer solchen Anordnung lassen sich energiespezifische Informationen über die Phase beziehungsweise den Brechungsindex gewinnen, analog zum Dual-Energy-Verfahren bei konventionellen, die Absorption messenden Röntgengeräten. Zum Erhalt computertomografischer Projektionen werden die beiden Talbot-Lau-Interferometer mit gleicher Geschwindigkeit um das Untersuchungsobjekt, hier ein Patient P, rotiert.

### 3 Fazit und Ausblick

Mit der Hilfe von Talbot-Lau-Interferometern ist es möglich, neben den konventionellen Absorptionskontrast-Bildern auch Phasenkontrast-Bilder der inneren Struktur eines Objekts bei Durchstrahlung mit Röntgenstrahlen zu gewinnen. Wie weiter aufgezeigt wurde, haben die wesentlichen Entwicklungen von Talbot-Lau-Interferometern in der Patentliteratur ihren Niederschlag gefunden. Nicht zuletzt aus der großen Anzahl von Patentanmeldungen ist ersichtlich, dass es sich um ein sehr aktives Gebiet der Forschung und Entwicklung handelt.

Trotz aller Bemühungen ist es bisher nicht gelungen, die Phasenkontrast-Bildgebung mit Röntgenstrahlen als eine Standardalternative zur gewöhnlichen Absorptionskontrast-Bildgebung für den medizinischen Alltag zu etablieren [17] [18]. Vor allem für die Untersuchung von Weichgewebe wird die Phasenkontrast-Bildgebung jedoch als sehr vielversprechend angesehen. Gleichzeitig dürfte die Methode in der Zukunft auch verstärkt in der zerstörungsfreien Prüfung Verwendung finden (zum Beispiel [19]). Die Etablierung von Anwendungsmöglichkeiten von Talbot-Lau-Interferometern zur Phasenkontrast-Bildgebung mit Röntgenstrahlen befindet sich derzeit in einer entscheidenden Phase.

### Nicht-Patentliteratur

- [1] Erfindungen von Nobelpreisträgern – Frits Zernike. Online-Präsentation des Deutschen Patent- und Markenamts (2013). URL: <http://www.dpma.de/service/galerie/nobel/nobelpreisphysik/zernike/index.html> [abgerufen am 20.03.2015]
- [2] BONSE, U; HART, M.: An X-Ray Interferometer. In: Applied Physics Letters, Vol. 6, No. 8, 1965, S. 155–156
- [3] TALBOT, H. F.: LXXXVI. Facts relating to Optical Science. No. IV. In: Philosophical Magazine Series 3, Vol. 9, No. 56, 1836. 401–407
- [4] H. F. Talbot ist darüber hinaus für seine Leistungen auf dem Gebiet der Fotografie bekannt, siehe zum Beispiel das britische Patent 8842 von 1841.
- [5] [http://commons.wikimedia.org/wiki/File:Optical\\_Talbot\\_Carpet.png](http://commons.wikimedia.org/wiki/File:Optical_Talbot_Carpet.png) [frei unter der Creative Commons Lizenz, abgerufen am 20.03.2015]
- [6] LAU, E.: Beugungserscheinungen an Doppellrastern. In: Annalen der Physik, Bd. 2, Folge 6, 1948, S. 417–423
- [7] Nebenbei sei angemerkt, dass mehrere Patentdokumente auf den Erfinder Ernst Lau zurückgehen, darunter die Patentschrift DE 1 422 125 C zu Gleitsichtbrillen (zusammen mit G. Jaeckel und R. Rieker).
- [8] WEITKAMP, T. [et al.]: X-ray phase imaging with a grating interferometer. In: Optics Express, Vol. 13, No. 16, 2005, S. 6296–6304
- [9] PFEIFFER, F. [et al.]: Hard-X-ray dark-field imaging using a grating interferometer. In: Nature Materials, Vol. 7, No. 2, 2008, S. 134–137
- [10] RESS, D. [et al.]: Measurement of Laser-Plasma Electron Density with a Soft X-ray Laser Deflectometer. In: Science, Vol. 265, No. 5171, 1994. 514–517
- [11] DAVID, C. [et al.]: Differential x-ray phase contrast imaging using a shearing interferometer. In: Applied Physics Letters, Vol. 81, No. 17, 2002, S. 3287–3289
- [12] MOMOSE, A. [et al.]: Demonstration of X-Ray Talbot Interferometry. In: Japanese Journal of Applied Physics, Vol. 42, No. 7B, 2003, S. L866–L868

- [13] PFEIFFER, F. [et al.]: Phase retrieval and differential phase-contrast imaging with low-brilliance X-ray sources. In: *Nature Physics*, Vol. 2, No. 4, 2006, S. 258–261
- [14] OLIVO, A.: Recent Patents in X-Ray Phase Contrast Imaging. In: *Recent Patents on Biomedical Engineering*, Vol. 3, No. 2, 2010, S. 95–106
- [15] CLAUSER, J.F.; REINSCH, M.W.: New Theoretical and Experimental Results in Fresnel Optics with Applications to Matter-Wave and X-Ray Interferometry. In: *Applied Physics B*, Vol. 54, No. 5, 1992, S. 380–395
- [16] John F. Clauser ist hauptsächlich für seine experimentellen Arbeiten zu den Grundlagen der Quantenmechanik bekannt (Bellsche Ungleichung), für die er 2010 den Wolf Prize in Physik erhalten hat.
- [17] PFEIFFER, F.: Grating-based X-ray phase contrast for biomedical imaging applications. In: *Zeitschrift für Medizinische Physik*, Vol. 23, No. 3, 2013, S. 176–185
- [18] OLIVO, A., ROBINSON, I.: ‘Taking X-ray phase contrast imaging into mainstream applications’ and its satellite workshop ‘Real and reciprocal space X-ray imaging’. In: *Philosophical Transactions of the Royal Society A*, Vol. 372, No. 2010, 2014. 20130349
- [19] KOTTLER, C. [et al.]: X-ray Grating-based Phase Contrast CT for Non-Destructive Testing and Evaluation. In: *Proceedings of the Conference on Industrial Computed Tomography*, Wels, 2012, S. 129–134. – ISBN 978-3844012811

# Glossar

## **Anger-Kamera**

Seite 25

Nach dem Erfinder Hal Anger benannter ortssensitiver Gammadetektor aus einem flächigen Szintillator-Kristall, in dem ein einfallendes Gamma-Quant einen Lichtblitz auslöst, der von dahinter im Raster angeordneten lichtempfindlichen Verstärkerröhren (Photomultipliern) detektiert wird. Die Signalstärke der einzelnen Photomultiplier wird durch eine Verstärkerschaltung in zwei Signale überführt, die die x- und y-Koordinate des Lichtblitzes angeben.

## **anisotrop**

Seite 14

Eine Eigenschaft wird als anisotrop bezeichnet, wenn sie von der Richtung oder Orientierung im Raum abhängig ist.

## **Bit-Abbilder (remapper)**

Seite 60

Rasterbilder (Bitmaps) stellen Bilder mit einem rechteckigen Gitter aus Bildelementen (Pixeln) dar. Jedem Bildelement (Pixel) ist eine bestimmte Position und ein Farbwert zugewiesen. Bei einem Bit-Abbilder (remapper) werden Bildelemente (Pixel) eines Rasterbilds umgestaltet.

## **Degradation**

Seite 20, 79

Schleichende Veränderung von Eigenschaften

## **Elektrolumineszenz-Display**

Seite 63

Das Licht wird durch Anlegen einer elektrischen Spannung bzw. eines elektrischen Feldes in einem Festkörper erzeugt. Lumineszenzstrahler zeigen je nach Bauart völlig verschiedene Spektren, angefangen beim Linienspektrum über Bandenspektren bis zu kontinuierlichen Spektren. Zu Elektrolumineszenzstrahlern gehören anorganische Leuchtdioden (LEDs) und organische Leuchtdioden (OLEDs) sowie Elektrolumineszenz-Folien.

## **Fastie-Elbert-Konfiguration**

Seite 60

Veränderung der Strahlengänge durch Beugung an einem Gitter eines Spektrometers, das in der Brennebene eines Hohlspiegels angeordnet ist [2].

## **Fresnel-Linse**

Seite 62

Eine Verringerung des Volumens gegenüber einer großen Linse geschieht bei der Fresnel-Linse durch eine Aufteilung des Volumens in ringförmige Bereiche. In jedem dieser Bereiche wird die Dicke verringert, sodass die Linse eine Reihe ringförmiger Stufen erhält. Da Licht nur beim Passieren der Linsen-Oberflächen gebrochen wird, ist der Brechungswinkel nicht von der Dicke, sondern nur von dem Winkel zwischen den beiden Oberflächen abhängig.

## **Geberelement**

Seite 11, 15

In diesem Zusammenhang wird damit ein Magnet bezeichnet, dessen Feld für die Positionsbestimmung ausgemessen und analysiert wird.

## **Gyrosensor**

Seite 66

Ein Gyrosensor ist ein Beschleunigungs- oder Lagesensor, der auf kleinste Beschleunigungen, Drehbewegungen oder Lageänderungen reagiert. Das Prinzip des Gyrosensors basiert auf der Massenträgheit und wird u.a. in Fliehkraftreglern eingesetzt.

## **hartmagnetisch**

Seite 16

Hartmagnetische Werkstoffe bilden die Basis für aus dem Alltag bekannte Permanentmagnete. Sie lassen sich durch äußere magnetische Felder kaum ent- oder ummagnetisieren.

## **Kommutator**

Seite 27

Die Spulen auf dem drehenden Anker eines Elektromotors werden durch den Kommutator mit Strom versorgt. Zum Kommutator gehören feststehende leitende Bürsten und die bewegten Kontakte der Spulen. Der Kommutator dient dazu, die Stromrichtung durch die Spulenwicklungen im Laufe jeder Drehung des Ankers periodisch umzukehren, so der Elektromotor durch wechselnde Magnetfelder angetrieben wird.

## **Sensorknoten**

Seite 14

Ein Sensorknoten besteht aus einem Prozessor, der mit Speicher und Sensoren in der Regel auf einem Chip untergebracht ist. Das System wird mit Batterie betrieben und ist mit einem Modul zur Kommunikation, beispielsweise via Funk, ausgestattet.

## **Sinterschritt**

Seite 29

Sintern bezeichnet eine Verbindungstechnik. Dabei wird unter Druck und Temperatur ein Rohling aus Pulver in einen festen Körper umgewandelt, ohne dass der Rohling aufgeschmolzen wird.

## **„Time of flight (TOF)-Methode“ beim PET**

Seite 27

Bei einem TOF-PET wird mit einer hohen zeitlichen Auflösung der zeitliche Abstand erfasst, in dem zwei Detektormodule, die sich im Detektorring gegenüberliegen, ansprechen. Bei einer zeitlichen Auflösung von zum Beispiel 550 Picosekunden, kann der Ort, von dem die beiden Gamma-Quanten von dem zerfallenen  $\beta^+$ -Radionuklid ausgesendet wurden, auf  $\pm 8$  cm genau bestimmt werden (in 550 Picosekunden legt das Gamma-Quant mit Lichtgeschwindigkeit etwa 16 cm zurück). Diese zusätzliche Information führt im Ergebnis zu kontrastreicherem, schärferen Bildern.

## **weichmagnetisch**

Seite 16

Weichmagnetische Materialien lassen sich vergleichsweise leicht durch ein auf sie wirkendes Magnetfeld magnetisieren. Anschaulich gesprochen verstärken sie so das Magnetfeld. Weil sie sich auch relativ leicht ummagnetisieren lassen (kleine Hysterese), werden sie typischerweise in Transformatoren und Elektromotoren eingesetzt.



# Impressum

**Herausgeber**

Deutsches Patent- und Markenamt  
Zweibrückenstraße 12  
80331 München

Telefon +49 89 2195-0  
[www.dpma.de](http://www.dpma.de)

**Stand**

Oktober 2015

**Bildnachweis**

fotolia.com: Andy Dean Photography

