

Halbleitertechnologie von A bis Z

Optoelektronik & Photonik

Grundlagen der Optoelektronik	3
<i>Warum Optoelektronik?</i>	3
<i>Direkte vs. indirekte Bandlücke</i>	3
<i>Photonenenergie und Bandlücke</i>	3
<i>Materialsystem-Übersicht</i>	4
<i>Quanteneffizienz als Bewertungsgröße</i>	5
Leuchtdioden (LED)	5
<i>Aufbau und Funktionsprinzip</i>	5
<i>Materialsysteme nach Farbe</i>	6
<i>Quantum Wells zur Effizienzsteigerung</i>	6
<i>Droop-Effekt und Auskopplungseffizienz</i>	8
<i>Von der Emission zur nutzbaren Lichtleistung</i>	9
Laserdioden	9
<i>Vom LED zum Laser: Besetzungsinversion und stimulierte Emission</i>	9
<i>Der optische Resonator und Modenbildung</i>	10
<i>Kantenemitter: Fabry-Pérot- und DFB-Laser</i>	10
<i>VCSEL: Aufbau und Vorteile</i>	11
<i>Materialsysteme und Anwendungsfelder</i>	11
Photodioden & Photodetektoren	12
<i>Funktionsprinzip: pn- und pin-Photodiode</i>	12
<i>Kenngößen: Quantenausbeute und Responsivität</i>	13
<i>Avalanche-Photodiode (APD)</i>	14
<i>Materialwahl je Wellenlängenbereich</i>	14
<i>Anwendungen: Glasfaserkommunikation und Bildsensorik</i>	15
Integration & Anwendungen	15
<i>Die optische Übertragungskette</i>	15
<i>Displaytechnik: LED-Backlight, μLED und OLED im Vergleich</i>	16
<i>Sensorik: LiDAR und Time-of-Flight</i>	16
<i>Sensorik: Pulsoximetrie und Näherungssensoren</i>	17
<i>Photonische integrierte Schaltkreise (PIC) - Ausblick</i>	18

Grundlagen der Optoelektronik

Warum Optoelektronik?

Optoelektronische Bauelemente nutzen die Wechselwirkung zwischen Licht und Halbleitern in beide Richtungen: Bei der Emission wandeln LEDs und Laserdioden elektrische Energie in Photonen um, bei der Detektion wandeln Photodioden Licht zurück in elektrische Signale. Elektronen aus dem Leitungsband rekombinieren mit Löchern im Valenzband, die freiwerdende Energie wird als Photon abgestrahlt. Bei der Detektion läuft der Prozess umgekehrt – ein absorbiertes Photon hebt ein Elektron ins Leitungsband und erzeugt so ein Elektron-Loch-Paar, das als Photostrom messbar wird.

Diese drei Bauelementklassen – LED, Laserdiode, Photodiode – bilden zusammen die Grundlage praktisch jeder optischen Datenübertragung, vieler Sensorik-Anwendungen und der modernen Displaytechnik. Anders als bei den bisher behandelten Bauelementen (MOSFET, Bipolartransistor) steht hier nicht das Schalten von Strömen im Vordergrund, sondern die kontrollierte Wandlung zwischen elektrischer und optischer Energie.

Direkte vs. indirekte Bandlücke

Ob ein Halbleiter effizient Licht emittieren kann, entscheidet sich an der Bandstruktur im k -Raum, also der Auftragung der Energie über dem Kristallimpuls der Elektronen. Bei einem direkten Halbleiter wie GaAs liegen das Minimum des Leitungsbands und das Maximum des Valenzbands beim gleichen k -Wert. Ein Elektron kann beim Übergang ins Valenzband rekombinieren, ohne seinen Impuls ändern zu müssen – die Impulserhaltung übernimmt allein das emittierte Photon, dessen Impuls im Vergleich zum Elektronenimpuls vernachlässigbar klein ist.

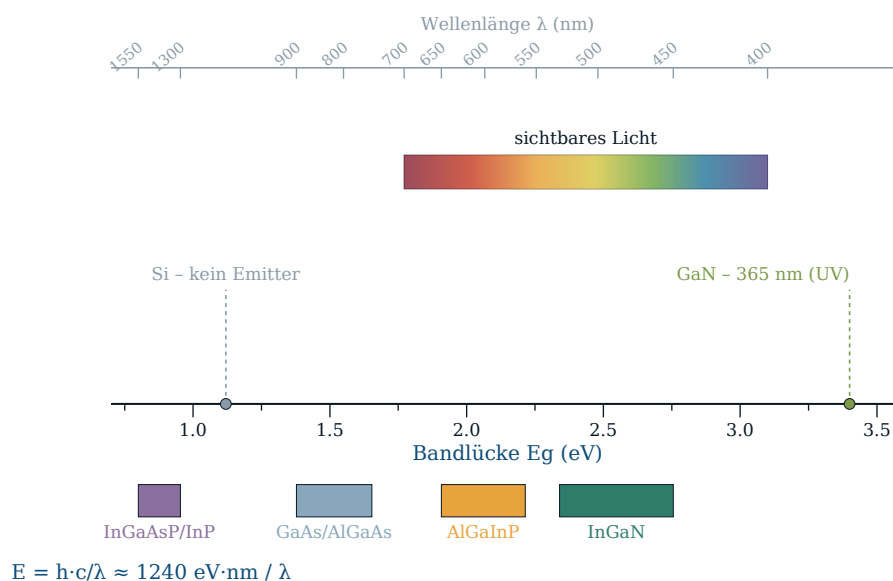
Bei einem indirekten Halbleiter wie Silicium liegen Leitungsband-Minimum und Valenzband-Maximum bei unterschiedlichen k -Werten. Die Rekombination erfordert zusätzlich ein Phonon, das die Impulsdifferenz aufnimmt oder liefert. Da ein Dreiteilchenprozess (Elektron, Loch, Phonon) erheblich unwahrscheinlicher ist als ein Zweiteilchenprozess, liegt die strahlende Rekombinationsrate in Silicium um mehrere Größenordnungen niedriger als in GaAs – die interne Quanteneffizienz liegt bei Si typischerweise im Bereich von 10^{-6} , bei GaAs nahe 1. Das ist der physikalische Kern, warum Si-LEDs praktisch nicht existieren und III-V-Verbindungshalbleiter für Lichtemission alternativlos sind.

Photonenenergie und Bandlücke

Die Wellenlänge des emittierten oder absorbierten Lichts hängt direkt von der Bandlücke ab: $E = h \cdot c / \lambda$. Für die Praxis hat sich eine Faustformel etabliert: $\lambda \text{ [nm]} \approx 1240 / E_g \text{ [eV]}$. Eine größere Bandlücke bedeutet kürzerwelliges, energiereicheres Licht. GaAs mit $E_g \approx 1,42 \text{ eV}$ emittiert bei rund 870 nm im nahen Infrarot, GaN mit $E_g \approx 3,4 \text{ eV}$ liegt bei etwa 365 nm im UV-Bereich.

Sichtbares Licht erfordert Bandlücken zwischen etwa 1,8 eV (rot) und 3,1 eV (violett) – ein Bereich, den kein binäres III-V-Material vollständig abdeckt. Erst durch Legierung, etwa InGaInP oder AlGaInP, lässt sich die Bandlücke kontinuierlich zwischen den Werten der binären Randverbindungen einstellen (Vegard-Regel als Näherung) und so gezielt auf die gewünschte Emissionswellenlänge trimmen, ohne das grundlegende Materialsystem zu wechseln.

Bandlücke und Emissionswellenlänge
Materialsysteme optoelektronischer Bauelemente im Überblick



Materialsystem-Übersicht

Jeder Wellenlängenbereich hat sein etabliertes Materialsystem, das sich an der jeweiligen Anwendung orientiert:

- InGaIn – blau bis grün (450–530 nm), Basis moderner weißer LEDs in Kombination mit gelbem Phosphor
- AlGaInP – rot bis gelb (560–650 nm), ergänzt InGaIn zur vollständigen RGB-Abdeckung

- GaAs / AlGaAs – nahes Infrarot (750–900 nm), Standardmaterial für einfache Laserdioden und IR-LEDs (Fernbedienungen, Näherungssensoren)
- InGaAsP / InP – Telekommunikationswellenlängen 1,3 μm und 1,55 μm , gewählt wegen der Dämpfungs- und Dispersionsminima von Quarzglasfasern in diesem Bereich
- Silicium / Germanium – kein Emitter, aber wichtige Photodetektor-Materialien im sichtbaren (Si) bzw. nahinfraroten Bereich (Ge), da Absorption anders als Emission keinen strahlenden Übergang benötigt und daher auch in indirekten Halbleitern effizient funktioniert

Diese Materialzuordnung zieht sich als roter Faden durch die folgenden Kapitel: Welches Material für LED, Laser oder Photodiode infrage kommt, ist im Kern immer eine Frage der Zielwellenlänge und der damit verknüpften Bandlücke.

Quanteneffizienz als Bewertungsgröße

Um Emitter und Detektoren später vergleichbar zu machen, lohnt sich an dieser Stelle eine Begriffsklärung. Die interne Quanteneffizienz beschreibt, welcher Anteil der injizierten Elektron-Loch-Paare tatsächlich strahlend rekombiniert, im Gegensatz zu nichtstrahlenden Verlustkanälen über Kristalldefekte oder Oberflächenzustände. Die externe Quanteneffizienz berücksichtigt zusätzlich, wie viele der intern erzeugten Photonen den Halbleiter tatsächlich verlassen, statt durch Totalreflexion oder Reabsorption im Bauelement verloren zu gehen.

Diese Unterscheidung wird im LED-Kapitel bei der Diskussion von Auskopplungseffizienz und Droop-Effekt wieder aufgegriffen und ist auch bei Photodioden zentral, dort als Responsivität ausgedrückt.

Leuchtdioden (LED)

Aufbau und Funktionsprinzip

Eine LED ist im Kern ein in Durchlassrichtung betriebener pn-Übergang aus einem direkten Halbleiter. Legt man eine Vorwärtsspannung an, senkt sich die Potentialbarriere der Raumladungszone, und Elektronen aus der n-Seite sowie Löcher aus der p-Seite werden über die Verarmungszone hinweg injiziert – man spricht von Minoritätsträgerinjektion. Im aktiven Bereich treffen beide Ladungsträgertypen aufeinander und rekombinieren strahlend, wie im Grundlagenkapitel beschrieben; die Photonenenergie entspricht dabei näherungsweise der Bandlücke des Materials, leicht reduziert durch die thermische Energieverteilung der Ladungsträger (daher ist das

Emissionsspektrum kein scharfer Peak, sondern hat eine spektrale Breite von typischerweise 20–30 nm bei Raumtemperatur). Anders als beim Si-pn-Übergang, der primär als Diode oder Solarzelle genutzt wird, ist bei der LED die Lichtemission der eigentliche Zweck, nicht Nebeneffekt.

Die Durchlassspannung einer LED korreliert direkt mit der Bandlücke: Faustregel $V_f \approx E_g/e$, leicht erhöht durch Serienwiderstände und Kontaktverluste. Eine rote AlGaInP-LED ($E_g \approx 2$ eV) hat daher eine typische Flussspannung um 1,8–2,2 V, eine blaue InGaN-LED ($E_g \approx 2,7$ eV) liegt bei 2,8–3,4 V – deutlich höher als bei einer klassischen Si-Diode ($\approx 0,7$ V), was sich direkt in der benötigten Betriebsspannung von LED-Schaltungen niederschlägt. Der Wirkungsgrad hängt entscheidend davon ab, wie viele Ladungsträger tatsächlich strahlend statt nichtstrahlend rekombinieren – gesteuert durch Dotierqualität, Defektdichte und Grenzflächenzustände im Kristall, die als Rekombinationszentren wirken und Ladungsträger ohne Photonemission "verschlucken".

Materialsysteme nach Farbe

Die Emissionsfarbe folgt direkt aus der Bandlücke des Materials, weshalb jede LED-Farbe ihr eigenes etabliertes System hat. Rote und gelbe LEDs basieren auf AlGaInP-Heterostrukturen auf GaAs-Substrat, gitterangepasst durch die Wahl des Aluminium-Gallium-Verhältnisses. Infrarot- und einfache rote LEDs nutzen historisch auch GaAsP oder AlGaAs. Blaue und grüne LEDs nutzen InGaN auf Saphir- oder SiC-Substrat – trotz erheblicher Gitterfehlانpassung zum Substrat (bis zu 16 % bei GaN-auf-Saphir), was zu hohen Versetzungsdichten führt, die die Materialqualität lange limitierten.

Die technologische Reife von InGaN in den 1990er-Jahren (Shuji Nakamura, Nobelpreis Physik 2014) war die entscheidende Voraussetzung für effiziente weiße LEDs, da zuvor nur Rot und Gelb in ausreichender Effizienz verfügbar waren – das sogenannte "Green Gap" (mangelnde Effizienz grüner LEDs zwischen den gut beherrschten Systemen InGaN/blau und AlGaInP/gelb) ist bis heute ein aktives Forschungsthema. Weißes Licht entsteht heute überwiegend nicht durch additive RGB-Mischung, sondern durch eine einzelne blaue InGaN-LED (Peak ≈ 450 – 460 nm), deren Licht teilweise von einem gelben Phosphor (meist Ce:YAG, Cer-dotiertes Yttrium-Aluminium-Granat) durch Photolumineszenz in längerwelliges Licht (Peak ≈ 560 nm) umgewandelt wird. Die Kombination aus durchscheinendem Blau und angeregtem Gelb ergibt für das menschliche Auge einen weißen Farbeindruck; über Phosphor-Mischungen und -Dicke lässt sich die Farbtemperatur (warmweiß bis kaltweiß) sowie der Farbwiedergabeindex (CRI) gezielt einstellen.

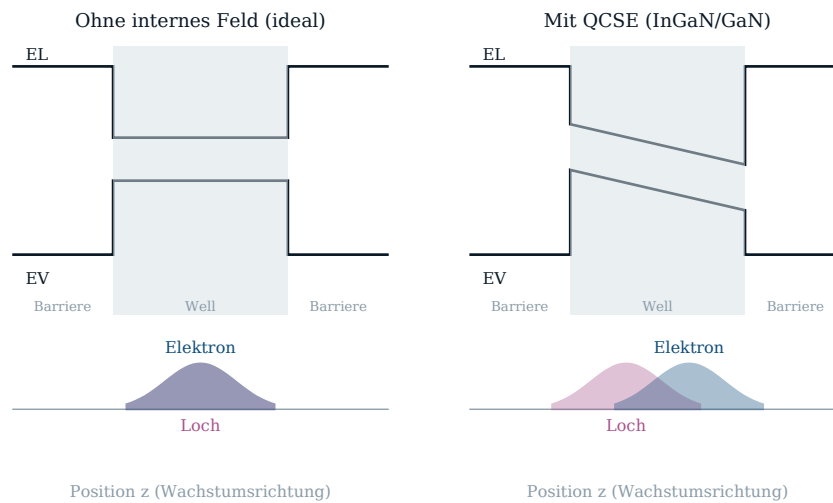
Quantum Wells zur Effizienzsteigerung

Moderne LEDs nutzen keinen einfachen homogenen pn-Übergang, sondern eine aktive Zone aus mehreren Quantum Wells (MQW, Multiple Quantum Wells): dünne Schichten (typisch 2–5 nm) eines Materials mit kleinerer Bandlücke, eingebettet zwischen Barriereschichten mit größerer Bandlücke – bei InGaN/GaN-LEDs etwa dünne InGaN-Wells zwischen GaN-Barrieren. Ladungsträger werden in diesen Potentialtöpfen quantenmechanisch räumlich eingesperrt (Quantum Confinement), wodurch sich die Zustandsdichte von einer kontinuierlichen 3D- zu einer stufenförmigen 2D-Verteilung ändert. Das erhöht die effektive Ladungsträgerdichte im aktiven Volumen und damit die strahlende Rekombinationsrate deutlich gegenüber einem dicken, unstrukturierten Übergang.

Über die Well-Dicke und Materialzusammensetzung lässt sich die effektive Bandlücke und damit die Emissionswellenlänge fein einstellen – bei InGaN/GaN-Systemen typischerweise über den Indium-Gehalt (höherer In-Anteil senkt die Bandlücke, verschiebt die Emission Richtung Grün) in Kombination mit der Well-Breite. Eine materialspezifische Komplikation bei InGaN/GaN ist der Quantum-Confined Stark-Effekt (QCSE): Da GaN entlang der c-Achse piezoelektrisch und spontan polarisiert ist, entstehen an den Well-Grenzflächen starke interne elektrische Felder (mehrere MV/cm), die Elektronen und Löcher räumlich innerhalb des Wells trennen. Das verringert die Überlappung ihrer Wellenfunktionen und damit die strahlende Rekombinationswahrscheinlichkeit – ein Effekt, der bei hohem Indium-Gehalt (grüne LEDs) besonders ausgeprägt ist und mitverantwortlich für das erwähnte Green Gap ist. Diese Heterostruktur-Bauweise ist im Prinzip dieselbe Technik, die auch bei HEMTs zum Einsatz kommt, nur mit dem Ziel der Ladungsträgereinfangung zur Rekombination statt hoher lateraler Beweglichkeit.

Quantum Well ohne und mit QCSE

Räumliche Trennung von Elektron und Loch durch das interne Feld in InGaN/GaN



Starke Überlappung → hohe strahlende Rekombinationsrate;
getrennte Wellenfunktionen → reduzierte Effizienz.

Droop-Effekt und Auskopplungseffizienz

Zwei Effekte begrenzen die reale Effizienz einer LED unabhängig von der internen Materialqualität. Der Droop-Effekt bezeichnet den Effizienzabfall bei hohen Stromdichten, wie er sich im sogenannten ABC-Modell beschreiben lässt: Die Rekombinationsrate setzt sich aus einem defektbedingten Shockley-Read-Hall-Term ($A \cdot n$, linear), dem gewünschten strahlenden Term ($B \cdot n^2$, quadratisch) und einem nichtstrahlenden Auger-Term ($C \cdot n^3$, kubisch) zusammen. Bei niedrigen Ladungsträgerdichten n dominiert der strahlende B-Term, doch mit steigender Stromdichte wächst der Auger-Term (Drei-Teilchen-Prozess, bei dem die Rekombinationsenergie statt an ein Photon an ein drittes Ladungsträgerteilchen übertragen wird) überproportional schneller – die interne Quanteneffizienz erreicht daher ein Maximum bei moderater Stromdichte und fällt danach wieder ab, ausgerechnet im für hohe Helligkeit gewünschten Betriebsbereich. Die genauen physikalischen Ursachen des Auger-Anteils bei InGaN sind bis heute Gegenstand aktiver Forschung.

Der zweite limitierende Faktor ist die Auskopplungseffizienz. Da der Brechungsindex von GaN ($n \approx 2,4$) deutlich über dem von Luft ($n \approx 1$) liegt, definiert das Snelliussche Brechungsgesetz einen engen Grenzwinkel der Totalreflexion – nur Licht, das nahezu senkrecht auf die Grenzfläche trifft, verlässt den Chip direkt; der überwiegende Teil wird zurückgeworfen, mehrfach im Chip reflektiert und dabei zunehmend absorbiert. Maßnahmen wie aufgeraute Oberflächen (Photonic Roughening), geformte Chip-Geometrien (z.

B. abgeschrägte Seitenflächen), photonische Kristallstrukturen an der Auskoppelfläche oder Flip-Chip-Montage mit reflektierendem Rückkontakt brechen diese Totalreflexion gezielt auf und steigern die extern nutzbare Lichtausbeute erheblich – von oft unter 50 % bei unbehandelten Chips auf über 80 % bei optimierten Designs.

Von der Emission zur nutzbaren Lichtleistung

Für die praktische Bewertung einer LED ist die Kette aus mehreren Wirkungsgraden entscheidend, die multiplikativ zur sogenannten Wall-Plug-Efficiency (elektrische Leistung → sichtbares Licht) zusammenwirken: interne Quanteneffizienz (strahlende vs. nichtstrahlende Rekombination im Kristall), Auskopplungseffizienz (Chip → Luft, siehe oben), sowie bei weißen LEDs zusätzlich die Phosphor-Konversionseffizienz. Letztere unterliegt dem Stokes-Verlust: Da ein blaues Photon (hohe Energie) in ein gelbes Photon (niedrigere Energie) umgewandelt wird, geht die Energiedifferenz zwangsläufig als Wärme im Phosphor verloren – ein prinzipbedingter, nicht vermeidbarer Effizienzverlust bei phosphorkonvertierten weißen LEDs.

Hinzu kommt der thermische Droop: Steigende Chiptemperatur im Betrieb reduziert zusätzlich die interne Quanteneffizienz, weshalb effizientes Wärmemanagement (Submount-Material, thermische Vias, Gehäusedesign) für die reale Lichtausbeute ebenso wichtig ist wie die Halbleiterphysik selbst. In Summe erreichen moderne kommerzielle Weißlicht-LEDs Lichtausbeuten von über 150–200 lm/W, verglichen mit rund 15 lm/W bei klassischen Glühlampen.

Laserdioden

Vom LED zum Laser: Besetzungsinversion und stimulierte Emission

Der fundamentale Unterschied zwischen LED und Laserdiode liegt nicht im Grundaufbau – beide basieren auf einem pn-Übergang aus direktem Halbleiter –, sondern im Emissionsmechanismus. Eine LED erzeugt Licht durch spontane Emission: Ein angeregtes Elektron rekombiniert zu einem zufälligen Zeitpunkt mit einem Loch, das entstehende Photon hat eine zufällige Phase und Richtung. Ein Laser hingegen nutzt stimulierte Emission: Ein bereits vorhandenes Photon trifft auf ein angeregtes Elektron-Loch-Paar und löst dessen Rekombination aus – das neu entstehende Photon ist dabei in Phase, Richtung und Wellenlänge exakt identisch mit dem auslösenden Photon. Dieser Verstärkungsmechanismus (Light Amplification by Stimulated Emission of Radiation) erzeugt kohärentes

Licht, im Gegensatz zum inkohärenten LED-Licht.

Damit stimulierte Emission gegenüber Absorption überwiegt, muss im aktiven Medium Besetzungsinversion herrschen: Es müssen mehr Elektronen im energetisch höheren Leitungsband-Zustand vorhanden sein als im niedrigeren Valenzband-Zustand – ein Zustand, der im thermischen Gleichgewicht nie auftritt und daher aktiv durch hohe Injektionsstromdichte erzwungen werden muss. Bei Halbleiterlasern übernimmt der stark in Durchlassrichtung gepumpte pn-Übergang diese Aufgabe: Oberhalb eines bestimmten Schwellstroms kippt das Verhältnis, optische Verstärkung setzt ein, und aus spontaner Emission wird zunehmend stimulierte Emission.

Der optische Resonator und Modenbildung

Besetzungsinversion allein erzeugt noch keinen Laser, sondern nur ein optisch verstärkendes Medium. Für den eigentlichen Lasereffekt braucht es zusätzlich Rückkopplung: einen optischen Resonator, meist realisiert durch zwei parallele, teilreflektierende Spiegelflächen beidseits der aktiven Zone. Bei einfachen Kantenemittern entstehen diese Spiegel durch die natürlichen Kristallspaltflächen des Halbleiters selbst (Bruchflächen entlang bestimmter Kristallebenen), die durch den Brechungsindexsprung zu Luft bereits eine Reflektivität von 30–35 % aufweisen – ausreichend, um Licht mehrfach durch das Verstärkungsmedium zu zirkulieren.

Nur Licht, dessen Wellenlänge ein ganzzahliges Vielfaches der Resonatorlänge ergibt (stehende Welle zwischen den Spiegeln), wird konstruktiv verstärkt – alle anderen Wellenlängen interferieren destruktiv und werden unterdrückt. Dies führt zur charakteristischen longitudinalen Modenstruktur eines Lasers: statt eines breiten, kontinuierlichen Spektrums wie bei der LED zeigt sich ein Kamm scharfer, eng benachbarter Emissionslinien, deren Abstand durch die Resonatorlänge bestimmt wird. Übersteigt die optische Verstärkung innerhalb des Resonators die Summe aller Verluste (Spiegeltransmission, interne Absorption, Streuung), setzt Laseroszillation ein – der Schwellstrom I_{th} markiert genau diesen Übergangspunkt und ist eine der wichtigsten Kenngrößen jeder Laserdiode.

Kantenemitter: Fabry-Pérot- und DFB-Laser

Die einfachste Bauform ist der Fabry-Pérot-Laser: zwei parallele Spaltflächen als Resonatorspiegel, Licht wird seitlich aus der Chipkante ausgekoppelt (daher "Kantenemitter" oder Edge-Emitting Laser). Da der Resonator meist mehrere hundert Mikrometer lang ist, unterstützt er gleichzeitig mehrere longitudinale Moden, deren Wellenlängenabstand invers proportional zur Resonatorlänge ist –

ein Fabry-Pérot-Laser emittiert daher typischerweise multimodig, mit mehreren Spektrallinien im Abstand von etwa 0,1–0,3 nm.

Für Anwendungen, die eine einzelne, extrem stabile Wellenlänge benötigen – allen voran die Glasfaserkommunikation, wo chromatische Dispersion sonst das Signal verschmiert – kommt der DFB-Laser (Distributed Feedback) zum Einsatz. Statt der einfachen Endspiegel wird hier ein periodisches Beugungsgitter direkt in die Schichtstruktur nahe der aktiven Zone integriert (Gitterperiode im Bereich weniger hundert Nanometer, hergestellt meist per Elektronenstrahlolithographie). Dieses Gitter reflektiert wellenlängenselektiv über die gesamte Resonatorlänge verteilt statt nur an den Endflächen und erzwingt so longitudinal-monomodigen Betrieb mit einer Seitenmodenunterdrückung von typischerweise über 30–40 dB. DFB-Laser auf InGaAsP/InP-Basis bei 1,55 μm sind der Industriestandard für Langstrecken-Glasfaserübertragung.

VCSEL: Aufbau und Vorteile

Eine grundlegend andere Geometrie verfolgt der VCSEL (Vertical-Cavity Surface-Emitting Laser): Statt Licht seitlich aus der Kante auszukoppeln, emittiert er senkrecht zur Waferoberfläche. Der Resonator liegt vertikal, gebildet aus zwei DBR-Spiegeln (Distributed Bragg Reflector) – Stapeln aus alternierenden Schichten unterschiedlichen Brechungsindex, jeweils eine Viertelwellenlänge dick, die durch konstruktive Interferenz Reflektivitäten von über 99 % erreichen können, da einzelne Grenzflächenreflexionen bei Halbleitern zu gering für einen brauchbaren Resonator wären. Da die aktive Zone dabei nur wenige Mikrometer lang ist statt mehrerer hundert Mikrometer bei Kantenemittern, sind extrem hohe Spiegelreflektivitäten nötig, um trotz der kurzen Verstärkungsstrecke die Laserschwelle zu erreichen.

Die vertikale Bauweise bringt entscheidende praktische Vorteile: VCSEL lassen sich wie LEDs auf Waferebene testen, bevor sie vereinzelt werden (bei Kantenemittern ist das erst nach dem Spalten möglich), sie emittieren ein rundes, wenig divergentes Strahlprofil (statt des elliptischen Strahls von Kantenemittern, was die Faserkopplung vereinfacht), und sie lassen sich dicht gepackt zu zweidimensionalen Arrays anordnen. Diese Eigenschaften machen VCSEL zur dominanten Technologie für kurze Datacom-Strecken (Rechenzentren, 850-nm-GaAs-VCSEL über Multimode-Glasfaser) und – in großer Stückzahl – für 3D-Sensorik wie Gesichtserkennung und LiDAR-Beleuchtung in Smartphones, wo Arrays aus tausenden VCSEL gleichzeitig strukturiertes Licht projizieren.

Materialsysteme und Anwendungsfelder

Die Materialwahl folgt denselben Grundsätzen wie bei der LED, ergänzt um resonatorspezifische Anforderungen. GaAs/AlGaAs-Systeme decken 780–980 nm ab (CD/DVD-Laser historisch, heute vor allem Pump Laser für Faserlaser und VCSEL für Datacom/Sensorik). InGaAsP/InP dominiert bei 1,3 μm und 1,55 μm für Telekommunikation, da diese Wellenlängen mit den Dämpfungs- und Dispersionsminima von Quarzglasfasern zusammenfallen. GaN-basierte Laserdioden decken den blauvioletten Bereich um 405 nm ab und ermöglichen die hochauflösende Blu-ray-Speichertechnik, da kürzere Wellenlängen eine kleinere Beugungsbegrenzung und damit höhere Speicherdichte erlauben.

Über die klassische Datenübertragung hinaus haben sich Laserdioden in drei weiteren Feldern etabliert: In der Materialbearbeitung ersetzen leistungsstarke Diodenlaser-Barren (arrayartig angeordnete Kantenemitter) zunehmend Gaslaser beim Schweißen, Schneiden und als Pumpquelle für Festkörperlaser, dank hoher elektrischer Effizienz (Wandlungswirkungsgrad oft über 50 %). In der Sensorik nutzt LiDAR (Light Detection and Ranging) kurze Laserpulse zur Laufzeitmessung für Umgebungserfassung in Automotive- und Robotik-Anwendungen – hier verschieben sich VCSEL zunehmend in den Vordergrund gegenüber klassischen Kantenemittern. Und in der Displaytechnik dienen rote, grüne und blaue Laserdioden als Lichtquelle für Laserprojektoren, wo ihre hohe spektrale Reinheit einen deutlich größeren Farbraum ermöglicht als LED- oder Lampenprojektoren.

Photodioden & Photodetektoren

Funktionsprinzip: pn- und pin-Photodiode

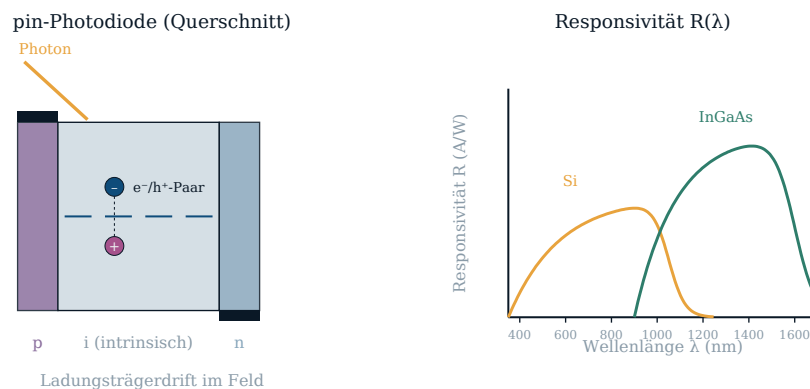
Eine Photodiode kehrt das LED-Prinzip um: Statt Ladungsträger zu injizieren und Licht zu erzeugen, wird einfallendes Licht absorbiert und in einen elektrischen Strom umgewandelt. Trifft ein Photon mit ausreichender Energie ($E \geq E_g$) auf den Halbleiter, wird ein Elektron aus dem Valenzband ins Leitungsband gehoben und erzeugt so ein Elektron-Loch-Paar. In der Raumladungszone eines in Sperrrichtung betriebenen pn-Übergangs trennt das dort vorhandene elektrische Feld dieses Paar sofort, bevor es rekombinieren kann – Elektron und Loch driften zu den jeweiligen Kontakten und erzeugen einen messbaren Photostrom, proportional zur einfallenden Lichtleistung.

In der einfachen pn-Diode ist die Raumladungszone jedoch sehr dünn (typisch $< 1 \mu\text{m}$), sodass der Großteil des Lichts erst tiefer im feldfreien Bulk-Material

absorbiert wird, wo erzeugte Ladungsträgerpaare durch langsame Diffusion statt schnelle Drift zu den Kontakten gelangen müssen – das begrenzt sowohl Empfindlichkeit als auch Schaltgeschwindigkeit. Die pin-Photodiode löst dieses Problem, indem zwischen p- und n-Schicht eine dicke, intrinsische (undotierte) i-Schicht eingefügt wird. Das elektrische Feld liegt dadurch über eine deutlich größere Distanz (oft 10–50 µm) an, praktisch das gesamte einfallende Licht wird im feldstarken Bereich absorbiert, und die Ladungsträger driften statt zu diffundieren – das erhöht die Quantenausbeute erheblich und ermöglicht Bandbreiten im Gigahertz-Bereich.

pin-Photodiode und spektrale Responsivität

Aufbau der pin-Photodiode und materialabhängige Empfindlichkeit über der Wellenlänge



In der i-Schicht liegt das elektrische Feld über die gesamte Dicke an –
Ladungsträger driften statt zu diffundieren, das erhöht Ausbeute und Bandbreite.

Kenngößen: Quantenausbeute und Responsivität

Zwei Größen beschreiben, wie effizient eine Photodiode Licht in Strom umsetzt. Die Quantenausbeute η gibt an, welcher Anteil der einfallenden Photonen tatsächlich ein detektierbares Elektron-Loch-Paar erzeugt (Werte nahe 1 bei gut designten Bauelementen). Die praktisch relevantere Größe ist die Responsivität $R = I_{ph}/P_{opt}$ in A/W, die den erzeugten Photostrom direkt zur eingestrahnten optischen Leistung ins Verhältnis setzt: $R = \eta \cdot e / (h \cdot f) = \eta \cdot \lambda / 1240 \text{ nm}$ (mit R in A/W und λ in nm). Eine Silicium-Photodiode erreicht bei 850 nm typischerweise $R \approx 0,5\text{--}0,6 \text{ A/W}$, eine InGaAs-Photodiode bei 1550 nm etwa $R \approx 0,9\text{--}1,0 \text{ A/W}$.

Die Responsivität ist wellenlängenabhängig und fällt oberhalb einer materialspezifischen Grenzwellenlänge (bestimmt durch die Bandlücke, $\lambda_g \approx 1240 \text{ nm}/E_g$) steil ab, da Photonen unterhalb der Bandlückenenergie nicht mehr absorbiert werden können. Zu kurzen Wellenlängen hin sinkt die Responsivität ebenfalls, da hochenergetische Photonen bereits sehr

oberflächennah absorbiert werden, wo Oberflächenrekombination einen Teil der erzeugten Ladungsträger vor der Trennung wieder vernichtet. Jedes Material hat daher ein charakteristisches Responsivitätsfenster, das die nutzbare Detektionsbandbreite definiert.

Avalanche-Photodiode (APD)

Für Anwendungen mit sehr schwachen Lichtsignalen – etwa Langstrecken-Glasfaserempfänger oder Einzelphotonendetektion – reicht die Empfindlichkeit einer einfachen pin-Photodiode oft nicht aus, da das erzeugte Signal im elektronischen Rauschen der nachfolgenden Verstärkerstufe untergeht. Die Avalanche-Photodiode (APD) löst dieses Problem durch interne Verstärkung: Sie wird mit einer deutlich höheren Sperrspannung betrieben (typisch 20–200 V, materialabhängig), wodurch in einem hochfeldigen Multiplikationsbereich der Photodiode ein sehr starkes elektrisches Feld entsteht.

Photogenerierte Ladungsträger werden in diesem Feld so stark beschleunigt, dass sie beim Stoß mit Gitteratomen genug Energie besitzen, um durch Stoßionisation weitere Elektron-Loch-Paare zu erzeugen – diese wiederum werden beschleunigt und erzeugen ihrerseits neue Paare. Der resultierende Lawineneffekt (daher "Avalanche") liefert einen internen Multiplikationsfaktor M von typischerweise 10 bis über 100, bevor das Signal überhaupt den externen Verstärker erreicht. Der Multiplikationsfaktor lässt sich über die angelegte Sperrspannung fein einstellen, muss aber sorgfältig gegen zusätzliches Multiplikationsrauschen (Excess Noise, durch die statistische Natur des Lawinenprozesses) abgewogen werden – ein zu hoher Verstärkungsfaktor verschlechtert das Signal-Rausch-Verhältnis wieder, statt es zu verbessern.

Materialwahl je Wellenlängenbereich

Wie bei LED und Laserdiode bestimmt die Bandlücke, welches Material für welchen Spektralbereich geeignet ist – allerdings mit einer wichtigen Umkehrung der Logik: Für die Absorption ist keine direkte Bandlücke zwingend erforderlich, da anders als bei der strahlenden Rekombination kein Photon emittiert werden muss. Silicium (indirekte Bandlücke, $E_g \approx 1,12$ eV) ist deshalb trotz seiner Ungeeignetheit als Lichtquelle ein exzellenter, kostengünstiger Photodetektor im sichtbaren und nahen Infrarot bis etwa 1000–1100 nm – der Wellenlängenbereich, den auch CMOS-Bildsensoren in Digitalkameras und Smartphones nutzen.

Für Telekommunikationswellenlängen (1,3 μm und 1,55 μm), oberhalb der Grenzwellenlänge von Silicium, kommt $\text{In}_{0,53}\text{Ga}_{0,47}\text{As}$ (gitterangepasst auf InP-Substrat) zum Einsatz – dasselbe Materialsystem, das auch bei den zugehörigen

DFB-Laserdioden verwendet wird, was Sender und Empfänger einer Glasfaserstrecke materialtechnisch eng verwandt macht. Germanium deckt einen ähnlichen Infrarotbereich ab und gewinnt durch seine CMOS-Kompatibilität an Bedeutung für monolithisch integrierte Silicium-Photonik-Schaltkreise. Für UV- und tiefblaue Detektion kommen breitbandige Materialien wie GaN oder SiC zum Einsatz, deren große Bandlücke sichtbares und infrarotes Licht praktisch transparent durchlässt und damit eine intrinsische Blindheit gegenüber Störlicht bietet – vorteilhaft etwa bei UV-Flammendetektoren.

Anwendungen: Glasfaserkommunikation und Bildsensorik

Die wirtschaftlich bedeutendste Anwendung von Photodioden ist der Empfänger in der Glasfaserkommunikation: Am Ende jeder Glasfaserstrecke wandelt eine InGaAs-pin- oder -APD-Photodiode das optische Signal zurück in ein elektrisches, das dann elektronisch weiterverarbeitet wird – die Empfängerempfindlichkeit (minimale detektierbare optische Leistung bei gegebener Bitfehlerrate) ist dabei neben der Senderleistung die zentrale Kenngröße, die die maximale Reichweite eines Glasfasersystems ohne Zwischenverstärkung bestimmt.

Die zweite große Anwendungsdomäne ist die Bildsensorik: CMOS-Bildsensoren in Kameras nutzen Millionen einzelner Silicium-Photodioden (Pixel), jede mit vorgeschalteten Farbfiltern (Bayer-Muster aus Rot-, Grün- und Blaufiltern) für die Farbauflösung, kombiniert mit integrierter Ausleseelektronik direkt am Pixel. Photodioden dienen darüber hinaus als einfache Lichtsensoren in unzähligen Alltagsanwendungen – automatische Bildschirmhelligkeit, Belichtungsmessung, Rauchmelder – sowie als Grundlage der Photovoltaik, wo dieselbe pn-Übergangsphysik statt für Signaldetektion für die großflächige Energiegewinnung genutzt wird, allerdings ohne Sperrspannung und mit auf Leistungsausbeute statt Geschwindigkeit optimierter Geometrie.

Integration & Anwendungen

Die optische Übertragungskette

Die vorangegangenen drei Kapitel behandelten LED, Laserdiode und Photodiode als einzelne Bauelemente – in der Praxis funktionieren sie jedoch fast nie isoliert, sondern als Glieder einer Übertragungskette: Sender (LED oder Laserdiode) → Übertragungsmedium (Glasfaser oder freier Raum) → Empfänger (Photodiode). Diese Kette macht deutlich, warum die Materialwahl in den

vorherigen Kapiteln so eng gekoppelt war: Ein InGaAsP-DFB-Laser bei 1,55 μm ergibt nur dann ein funktionierendes System, wenn am anderen Ende eine InGaAs-Photodiode mit passender Responsivität in genau diesem Wellenlängenfenster sitzt – Sender und Empfänger sind kein zufälliges Materialpaar, sondern eine aufeinander abgestimmte Systemkomponente.

Bei der Systemauslegung addieren sich die Wirkungsgrade und Verluste jeder Stufe: die Wall-Plug-Efficiency des Senders, die Einkoppleffizienz in die Faser (typischerweise der verlustreichste Schritt bei Kantenemittern, ein Grund für die zunehmende VCSEL-Nutzung mit ihrem runden, gut koppelbaren Strahlprofil), die Faserdämpfung über die Streckenlänge sowie schließlich die Empfängerempfindlichkeit der Photodiode. Die Leistungsbilanz (Link-Budget) aus Sendeleistung minus aller Verluste muss am Ende oberhalb der minimal detektierbaren Leistung der Photodiode liegen – reicht die Marge nicht aus, kommen Zwischenverstärker (optische Faserverstärker) oder empfindlichere APD-Empfänger zum Einsatz.

Displaytechnik: LED-Backlight, μLED und OLED im Vergleich

In der Displaytechnik kommen LEDs in mehreren, technisch grundverschiedenen Rollen zum Einsatz. Bei klassischen LCD-Bildschirmen dient eine Reihe weißer InGaN-LEDs (siehe Kapitel 2) lediglich als Hintergrundbeleuchtung – die eigentliche Bildinformation entsteht durch nachgeschaltete Flüssigkristallzellen, die das LED-Licht pixelweise durchlassen oder blockieren. Bei OLED-Displays hingegen ist jedes Pixel selbst ein Lichtemitter: organische Halbleiterschichten emittieren direkt Licht, wodurch sich Kontrast (echtes Schwarz durch vollständiges Abschalten einzelner Pixel) und Blickwinkelstabilität deutlich verbessern – allerdings zulasten von maximaler Helligkeit und Langzeitstabilität gegenüber anorganischen Halbleitern.

μLED -Displays verfolgen einen dritten Ansatz: Statt organischer Emmitter werden winzige, anorganische InGaN/AlGaInP-LED-Chips (Kantenlänge oft unter 50 μm) direkt als individuell ansteuerbare Pixel verwendet – dieselbe Halbleitertechnologie aus Kapitel 2, nur massiv miniaturisiert und in extrem hoher Stückzahl auf ein Trägersubstrat übertragen ("Mass Transfer", einer der größten Fertigungsengpässe dieser Technologie). μLED s vereinen dadurch die Kontrastvorteile selbstemittierender Pixel mit der Helligkeit, Lebensdauer und Materialstabilität anorganischer III-V-Halbleiter, sind aber wegen der Fertigungskomplexität aktuell noch auf Premium-Anwendungen und große Displaydiagonalen beschränkt.

Sensorik: LiDAR und Time-of-Flight

LiDAR-Systeme (Light Detection and Ranging) nutzen die in Kapitel 3 und 4 behandelten Bauelemente in einer präzisen Zeitmessungsanwendung: Eine Laserdiode – meist ein VCSEL oder Kantenemitter im Nahinfrarot (905 nm oder augensicherere 1550 nm) – sendet einen kurzen Lichtpuls aus, eine schnelle Photodiode (oft eine APD wegen der geringen Rücksignalstärke) registriert die Reflexion vom Zielobjekt. Aus der Laufzeit Δt zwischen Aussenden und Empfangen ergibt sich die Distanz direkt über $d = c \cdot \Delta t / 2$ – bei Lichtgeschwindigkeit entspricht das einer Zeitauflösung im Pikosekundenbereich für Zentimeter-Distanzgenauigkeit, was hohe Anforderungen an die Schaltgeschwindigkeit sowohl der Laserdiode als auch der Photodiode-Ausleseelektronik stellt.

Automotive-LiDAR-Systeme rastern dabei ein ganzes Sichtfeld ab, entweder mechanisch (rotierender Spiegel) oder zunehmend solid-state mittels VCSEL-Arrays, die elektronisch statt mechanisch verschiedene Winkel adressieren – ein direkter Anwendungsfall der in Kapitel 3 beschriebenen VCSEL-Array-Fähigkeit. Bei den 1550-nm-Systemen kommt der Vorteil der bereits für Telekommunikation etablierten InGaAsP/InP-Technologie zum Tragen: Diese Wellenlänge lässt sich mit deutlich höherer Leistung betreiben als 905 nm, ohne die Augensicherheitsgrenzwerte zu überschreiten, da sie von der Hornhaut absorbiert statt auf die Netzhaut fokussiert wird.

Sensorik: Pulsoximetrie und Näherungssensoren

Ein medizintechnisches Anwendungsbeispiel für die Materialkombination aus Kapitel 2 und 4 ist die Pulsoximetrie: Ein Fingerclip-Sensor durchleuchtet das Gewebe mit zwei LEDs unterschiedlicher Wellenlänge (rot, ≈ 660 nm, und infrarot, ≈ 940 nm) und misst mit einer dahinterliegenden Photodiode die durchgelassene Lichtintensität. Sauerstoffgesättigtes und -entsättigtes Hämoglobin absorbieren diese beiden Wellenlängen unterschiedlich stark; aus dem Verhältnis der pulsierenden Absorptionssignale (verursacht durch das arterielle Blutvolumen bei jedem Herzschlag) berechnet das Gerät die Sauerstoffsättigung, ganz ohne Blutentnahme.

Näherungssensoren in Smartphones funktionieren nach einem ähnlichen, aber einfacheren Prinzip: Eine IR-LED (meist GaAs-basiert, siehe Kapitel 2) sendet unsichtbares Licht aus, eine benachbarte Silicium-Photodiode misst die Reflexion – ist ein Objekt (etwa das Ohr beim Telefonieren) nah genug, steigt die reflektierte Intensität deutlich an und schaltet den Bildschirm ab. Dieselbe Grundarchitektur, erweitert um strukturiertes Licht aus VCSEL-Arrays (siehe Kapitel 3), bildet auch die Basis für 3D-Gesichtserkennungssysteme, die statt

einer einfachen Ja/Nein-Näherungsmessung ein vollständiges Tiefenbild der Szene rekonstruieren.

Photonische integrierte Schaltkreise (PIC) – Ausblick

Alle bisher behandelten Bauelemente – Laser, Photodiode, teils auch Modulatoren – existieren heute überwiegend als diskrete, einzeln gehäuste Komponenten, die über Fasern oder freien Raum miteinander verbunden werden. Photonische integrierte Schaltkreise (PICs) verfolgen denselben Miniaturisierungs- und Integrationsgedanken, der die Mikroelektronik seit Jahrzehnten prägt: mehrere optische Funktionen – Laser, Wellenleiter, Modulatoren, Photodetektoren – werden auf einem einzigen Chip kombiniert, verbunden durch strukturierte Wellenleiter statt diskreter Fasern.

Die Silicium-Photonik ist dabei der aktuell wichtigste Trend: Da Silicium selbst nicht emittiert (Kapitel 1), aber als Wellenleiter- und Detektormaterial (in Kombination mit Germanium, siehe Kapitel 4) hervorragend geeignet ist und zudem auf etablierten CMOS-Fertigungslinien hergestellt werden kann, lassen sich passive und detektierende photonische Funktionen kostengünstig integrieren. Die Lichtquelle selbst bleibt dabei meist ein separat gefertigter III-V-Laserchip (InP-basiert), der nachträglich auf den Silicium-Chip gebondet oder über Hybridintegration angekoppelt wird – ein aktives Forschungsfeld, das die in diesem Bereich behandelten III-V-Materialsysteme direkt mit der Silicium-Fertigungswelt der übrigen Kapitel auf der Webseite verknüpft. Anwendungstreiber sind vor allem Rechenzentren, wo PICs die elektrische Verkabelung zwischen Servern zunehmend durch optische Verbindungen mit höherer Bandbreite und geringerem Energieverbrauch pro Bit ersetzen.