

Ortsaufgelöste Leckstrom-Messungen an der Silizium-Elektrolyt-Grenzfläche

Dissertation

zur Erlangung des akademischen Grades
Doktor der Naturwissenschaften
(Dr. rer. nat.)
der Technischen Fakultät
der Christian-Albrechts-Universität zu Kiel

vorgelegt von
Jens Korallus
Kiel 2001

1. Gutachter: Prof. Dr. H. Föll
2. Gutachter: Prof. Dr. F. Faupel
Tag der mündlichen Prüfung: 25.04.2001

Inhaltsverzeichnis

Einleitung	1
1 Theoretische Grundlagen	5
1.1 Grundsätzliches zu Halbleitern	5
1.1.1 Intrinsische Halbleiter	6
1.1.2 Dotierte Halbleiter	8
1.2 Halbleiter-Metall-Kontakte	9
1.2.1 Schottky-Kontakte	10
1.2.2 Ohmsche Kontakte	11
1.2.3 Reale Metall-Halbleiter-Kontakte	12
1.3 Die Fest-Flüssig-Grenzfläche	16
1.3.1 Die Struktur der Grenzfläche	16
1.3.2 Das Energiediagramm der Fest-Flüssig-Grenzfläche	18
1.3.3 Die Durchtritts-Strom-Spannungs-Kurve	20
1.3.4 Der Silizium-Flußsäure-Kontakt	22
1.3.4.1 Der Mechanismus der Wasserstoffterminierung	23
1.3.4.2 Die Kennlinie des Silizium-Flußsäure-Kontaktes	24
1.4 Leckstrommechanismen	28
2 Elymat	31
2.1 Aufbau und Funktionsweise	31
2.2 Theorie der Elymat-Modi	35
2.2.1 FPC-Mode	36
2.2.2 BPC-Mode	38
2.2.3 BSC-Mode	40
3 Aufbau des Leckstrom-Scanners	43
3.1 Funktionsprinzip	43
3.2 Komponenten	45

3.2.1	Der xyz-Manipulator	45
3.2.2	Die Schutzvorrichtung für den Piezo-Translator	46
3.2.3	Der Potentiostat	48
3.2.4	Der 'Peak'-Detektor	50
3.2.5	Die Elektroden	53
3.2.5.1	Die Platindraht-Elektrode	54
3.2.5.2	200 μm Zylinder-Elektrode	56
3.2.5.3	Die 50 μm Zylinder-Elektrode	59
3.2.5.4	Die Scheiben-Elektrode	60
3.2.6	Die elektrochemische Zelle	62
3.2.7	Die 'Dark-Box'	62
3.3	Gesamtaufbau des Leckstrom-Scanners	63
3.3.1	Spannungsfolger	64
3.3.2	Zusatzmodul für die Höhenkalibrierung	64
3.4	Kompensation der Probenschieflage	65
3.5	Der Elektrolyt	67
3.6	Testen der Elektroden	69
4	Messungen an monokristallinem Silizium	71
4.1	Vorbereitungen	71
4.1.1	Auswahl und Präparation des Probenmaterials	71
4.1.2	Auswahl der Testdefekte	72
4.1.3	Wahl der Zellenspannung	73
4.1.4	Referenz-Elektrode	75
4.2	Charakterisierung der Elektroden	76
4.2.1	Einfluß der Flußsäurekonzentration	78
4.2.2	Einfluß der Elektrodenspannung	79
4.2.3	Einfluß des Elektrodenabstandes	80
4.2.4	Diskussion	83
4.3	Flächen-'Scans'	87
4.3.1	Auswahl der Teststruktur	87
4.3.2	Flächen-'Scans' mit den einzelnen Elektroden	88
4.3.3	Messungen an Material mit höherem spez. Widerstand	92
4.3.4	Einfluß von Gasbläschen auf die Messungen	95
4.4	Maßnahmen zur Vermeidung des Einflusses von Gasbläschen auf die Messungen	98
4.4.1	Pulsen der Zellenspannung	98
4.4.2	ElektrodenSpülung	98

4.4.3	Nicht-wässriger Elektrolyt	99
4.4.3.1	Auswahl des Redoxpaares	100
4.4.3.2	Herstellung des Elektrolyten	101
4.4.3.3	Messungen	101
4.4.3.4	Fazit	103
4.4.4	'Peak'-Detektor	103
4.5	Vergleich mit dem integralen Leckstrom	106
5	Messungen an multikristallinem Silizium	109
5.1	Probenpräparation	109
5.1.1	Zersägen der Wafer	109
5.1.2	Reinigung der Wafer	110
5.1.3	Beseitigung des Säge-'Damages'	111
5.2	Eigenschaften von Korngrenzen	114
5.3	Messungen	117
5.3.1	Vergleich mit Diffusionslängenverteilung	129
5.3.2	Fazit	132
	Zusammenfassung und Ausblick	133
A	Meßprogramm	137
B	'Peak'-Detektor	141
B.1	Die Wandlerplatine	142
B.1.1	Initialisierung	142
B.1.2	'Sample'-Vorgang	144
B.1.3	DA-Wandlung	145
B.2	Die Filterplatine	145
B.2.1	Hochpaß	145
B.2.2	Tiefpaß	145
B.2.3	Spannungsfolger	147
B.3	Die Steuerplatine	147
C	Bad-Heizung	149
C.1	Sägezahn-generator	149
C.2	Pulsbreitenmodulation	152
D	Temperatur-Regler	153
D.1	Aufbau des Reglers	153

D.1.1	PID-Regler	153
D.1.2	Strombegrenzung	158
D.1.3	Modul für die Gatespannungen	158
D.2	Regelverhalten	160
Literaturverzeichnis		163

Einleitung

Bis zu Beginn der zweiten Hälfte des zwanzigsten Jahrhunderts war Silizium ein Stoff von nur geringer Bedeutung. Das änderte sich als 1947 in den Bell Laboratories der Transistor erfunden und 1948 erstmals von *John Bardeen* und *Walter Brattain* gebaut wurde. Dies war die Geburtsstunde der Halbleiter-Technologie, die mit ihrer stürmischen Entwicklung bis in ferne Zukunft wohl einmalig bleiben wird. 1960 wurde der erste integrierte Schaltkreis gebaut. 1971 stellte Intel mit dem 4004 den ersten Mikroprozessor vor, der aus mehr als 2000 Transistoren bestand. Es folgten die bedeutungsvollen Prozessoren der 80x86-Serie mit bis zu 30.000 Transistoren, denen sich die Prozessoren der Serie 80286 bis 80486 anschlossen. Prozessoren der heutigen Generation bestehen aus mehr als 40 Millionen Transistoren mit Taktfrequenzen, die die 1 GHz-Schwelle bereits überschritten haben.

Aber nicht nur leistungsfähige Mikroprozessoren erlaubt die Erfindung des Bipolar-Transistors. Eine Weiterentwicklung dieses Bipolar-Transistors war der MOS-Transistor (**M**etal **O**xide **S**emiconductor). MOS-Transistoren ermöglichten die Entwicklung von DRAMs (**D**ynamic **R**andom **A**ccess **M**emory), die in Computern als digitale Massenspeicher Verwendung finden. Die ersten DRAMs hatten eine Kapazität von einigen kByte. Die im Verlauf der Entwicklung immer besser beherrschten Prozeßschritte sowie die stetige Miniaturisierung der Strukturen führten zu heutigen DRAMs mit einer Kapazität von 128 MByte.

Voraussetzung für diese Entwicklungen war, daß es möglich wurde, Siliziumeinkristalle herzustellen, die in ihrer Reinheit und Perfektion mit nichts vergleichbar sind. Einer der wichtigsten Materialparameter ist die Minoritätslebensdauer τ_n . Die erwähnte Reinheit und Perfektion des Siliziums erlauben große Minoritätslebensdauern und damit das Zustandekommen des Transistor-Effektes. Aber geringste Verunreinigungen oder Gitterdefekte im Silizium können eine drastische Verringerung der Minoritätslebensdauer, damit das Verschwinden des Transistor-Effektes und letztlich das Versagen eines gesamten Schaltkreises zur Folge haben. Dieser Umstand wird um so bedeutsamer, je größer die Fläche ist, die ein Schaltkreis auf einem Siliziumwafer in Anspruch nimmt. Obwohl sich Weiterentwicklung und Prozessierung immer schwieriger gestalten und deshalb die Chiphersteller zunehmend an den Rand der Rentabilität drängt, behält das Mooresche Gesetz bis zum heutigen Tag seine Gültigkeit. Das Mooresche Gesetz besagt, daß sich seit 1971 die Dichte der prozessierten Transistoren auf dem Wafer jedes Jahr um den Faktor 1.4 erhöht.

Ein Aspekt, der die Schwierigkeiten bei der Weiterentwicklung verdeutlicht, ist, daß

es zunehmend schwieriger wird, die zu prozessierenden Strukturen zu verkleinern, nicht nur, weil sich durch die verringerte Chipfläche die Ausbeute erhöht, sondern auch um parasitäre Kapazitäten und Bahnwiderstände zu verringern, was höhere Taktraten erlaubt. Die um 1960 kleinsten hergestellten Strukturen betrugen $25\text{ }\mu\text{m}$ und verringerten sich mit 11 % pro Jahr [TAU86]. Heutige Prozessoren werden in $0.18\text{ }\mu\text{m}$ -Technologie gefertigt [WAL95c].

Zusätzlich werden die Anforderungen, die an das Silizium bezüglich der Reinheit gestellt werden, immer größer. Silizium, das für die ULSI-Technologie (**U**ltra **L**arge **S**caled **I**ntegration) verwendet werden soll, darf nicht mehr als 10^{10} cm^{-3} an Fremd-Atomen enthalten [FIS94, GUP97, HEN95]. Besonders Schwermetalle wie Eisen, Kupfer, Nickel und Chrom gelten als ausgesprochen schädlich: Eisen gilt als der Lebensdauerkiller schlechthin. Eisen hat nicht nur eine Verschlechterung der pn-Übergänge durch Erhöhung des Leckstromes oder das Verschwinden des Transistor-Effektes bzw. die Verringerung des Stromverstärkungsfaktors in Bipolar-Transistoren zur Folge, sondern reduziert auch durch strahlungslose Rekombination die Lumineszenz von optoelektronischen Bauelementen, die vorwiegend aus III-V- und II-VI-Halbleitern gefertigt werden. Des weiteren kann sich die Lebensdauer des Bauteiles, bedingt durch den REDM¹-Mechanismus, stark verringern [HOP94]. Kupfer und Nickel sind nur bei höheren Temperaturen in größeren Mengen im Silizium löslich. Bei Raumtemperatur hingegen beträgt die Löslichkeit nur wenige Atome/ cm^3 und sie reagieren mit Silizium zu metallischen Phasen, die ausdiffundieren und sich an der Siliziumoberfläche anlagern [IST97]. Wird an diesen Orten beispielsweise ein MOS-Transistor prozessiert, dessen Gateoxid sich auf einer dieser metallischen Einschlüsse befindet, so ist an diesen Stellen die Wahrscheinlichkeit eines elektrischen Durchschlages während des Betriebes der Schaltung stark erhöht, wodurch der MOS-Transistor und damit die Schaltung zerstört werden würden [HEL94b, IST97, LIN98a, WAL95a, WAL95d]. Selbst wenn das Gateoxid durch Schwermetall-Kontaminationen während des Betriebes der Schaltung nicht zerstört wird, können Schwermetalle den Leckstrom von pn-Übergängen in DRAMs vergrößern, wodurch kürzere Refresh-Zyklen notwendig werden, um Datenverluste zu vermeiden [KEE98, LIN98a, SZE85].

Eine besondere Bedeutung kommt der Prozeßkontrolle zu. Da die Lebensdauer so sehr empfindlich auf geringste Verunreinigungen reagiert, eignet sie sich besonders gut, den Prozeß der Chipherstellung zu überwachen [SCH98]. Für diesen Zweck wurden mehrere Verfahren entwickelt: spektroskopische oder 'direkte' Verfahren wie TXRF, SIMS, VPD-AAS und ICPMS² und 'indirekte' Verfahren wie SPV, Elymat, μ -PCD, SCA und SCP³ [BER93, DAN98, NOR96, TAR95]. Die 'indirekten' Verfahren eignen sich deshalb zur Prozeßkontrolle, weil sie schnell sind, keine Präparation des Wafers erfordern, zerstörungsfrei arbeiten und vor allem kostengünstiger sind.

¹ REDM = **R**ecombination **E**nhanced **D**islocation **M**otion

² TXRF = **T**otal **X**-**R**ay **F**luorescence, SIMS = **S**econdary **I**on **M**ass **S**pectroscopy, VPD-AAS = **V**apour **P**hase **D**ecomposition-**A**tomie **A**bsorption **S**pectroscopy, ICPMS = **I**nductive **C**oupled **P**lasma **M**ass **S**pectrometer

³ SPV = **S**urface **P**hoto **V**oltage, Elymat = **E**lectrolytical **M**etal **T**racer, μ -PCD = **M**icrowave **P**hoto **C**onductivity **D**ecay, SCA = **S**urface **C**harge **A**nalyser und SCP = **S**urface **C**harge **P**rofiler

Nicht nur die Chipherstellung wird von Schwermetallen negativ beeinflusst, auch bei der Herstellung von Solarzellen erweisen sich die genannten Schwermetalle als schädlich. Die Folge einer geringen Diffusionslänge ist eine schlechte Quantenausbeute des langwelligen Teils des Sonnenspektrums, was den Wirkungsgrad der Solarzelle verringert [GLU95]. Kupfer- und Nickel-Ausscheidungen können lokal einen Kurzschluß der Raumladungszone des pn-Überganges zur Folge haben. Diese sogenannten 'Shunts' haben ebenfalls eine Verschlechterung des Wirkungsgrades der Solarzelle zur Folge.

Aus Kostengründen wird ein Großteil der Solarzellen aus multikristallinem Material gefertigt, die geringere Wirkungsgrade haben als Solarzellen aus monokristallinem Material. Die Ursache für den geringeren Wirkungsgrad findet sich an den Korngrenzen. An diesen kann zum einen, bedingt durch die erhöhte Defektdichte sowie durch Getter-Effekte, die Diffusionslänge reduziert sein, und zum anderen können Versetzungen Teile der Raumladungszone kurzschließen, wobei die letztgenannten 'Shunts' aber nur vereinzelt auftreten [CAR01, GLU96].

Elymat-Messungen an multikristallinem Silizium zeigten, daß die genannten Defekte den Leckstrom erheblich erhöhen können [LIP97]. Weder der Elymat noch die anderen der 'indirekten' Verfahren liefern Informationen über die örtliche Verteilung der Leckstromquellen.

Gegenstand dieser Arbeit ist die Vorstellung eines ganz neuen Verfahrens, das es ermöglichen wird den Leckstrom von mono- oder multikristallinen Siliziumwafern orts aufgelöst zu messen. Dieses Verfahren ist besonders für die Solarzellenforschung interessant, da es den Forschern, neben orts aufgelösten Diffusionslängen-Messungen, beispielsweise aus Elymat-Messungen, zusätzliche Informationen über das Material liefert, aus dem Solarzellen gefertigt werden sollen. Die prinzipielle Arbeitsweise des Leckstrom-Scanners ist das lokale Messen des am Ort (x_n, y_m) über eine Silizium-Elektrolyt-Grenzfläche fließenden Stromes. Aus diesem Grund wird für diese Arbeit folgende Gliederung gewählt:

In Kapitel 1 werden die Grundlagen vermittelt, die zum Verständnis der Vorgänge an der Halbleiter-Metall- sowie der Fest-Flüssig-Grenzfläche erforderlich sind. Um die weiter oben erwähnten 'zusätzlichen Informationen', wie sie aus Leckstrom-Messungen erhalten werden, zu bewerten, bietet sich ein Vergleich mit den zugehörigen Diffusionslängen an. Orts aufgelöste Diffusionslängen-Messungen sind mit dem Elymaten möglich. Eine Beschreibung des Funktionsprinzips sowie die Theorie der verschiedenen Elymat-Meßmodi finden sich im Kapitel 2.

In Kapitel 3 werden der Aufbau des Leckstrom-Scanners und die Herstellung der einzelnen Elektroden im Detail beschrieben. Das Gelingen der Elektrodenherstellung war von fundamentaler Bedeutung für den erfolgreichen Einsatz des Leckstrom-Scanners auf multikristallinem Material, so daß auf eine ausführliche Beschreibung der Elektrodenherstellung nicht verzichtet werden darf.

In Kapitel 4 werden Messungen an monokristallinem Material vorgestellt, die sich in zwei Klassen einteilen lassen: 'Line-Scans' über Kratzer zwecks Optimierung von diversen Meßparametern und Flächen-'Scans', um die Funktionalität des Leckstrom-Scanners zu demonstrieren. Schließlich folgen in Kapitel 5 Messungen an multikristal-

linem Material. Das Erreichen von örtlicher Auflösung in der Leckstromverteilung auf Material mit natürlichen Oberflächendefekten stellte das Hauptziel der Arbeit dar, die mit einer Zusammenfassung und einem Ausblick schließt.

Kapitel 1

Theoretische Grundlagen

Dieses Kapitel vermittelt die Grundlagen, die zum Verständnis der Vorgänge an Fest-Flüssig-Grenzflächen erforderlich sind. Diese Grenzfläche ist für die durchgeführten Experimente von zentraler Bedeutung, so daß auf eine ausführliche Behandlung nicht verzichtet werden kann.

1.1 Grundsätzliches zu Halbleitern

Der wesentliche Unterschied zwischen Metallen, Halbleitern und Isolatoren liegt in der Bandstruktur. Bei Metallen liegen entweder zum Teil gefüllte Bänder vor, oder das volle Valenz- und Leitungsband überlappen sich. Hiervon scharf abgrenzen lassen sich die Halbleiter. Bei $T = 0$ K ist das Valenzband komplett gefüllt, das Leitungsband leer, wobei beide Bänder energetisch durch einen Bandabstand (Bandgap) $E_G < 2.6$ eV¹ voneinander getrennt sind [FAS87]. Gilt $E_G > 2.6$ eV, so spricht man von einem Isolator. Letztlich verhalten sich Halbleiter und Isolatoren gleich, allerdings auf unterschiedlichen Temperaturskalen: Ein Halbleiter ist bei $T = 0$ K ein perfekter Isolator, bei Zimmertemperatur hingegen ein Leiter, allerdings weit weniger gut als ein Metall. Isolatoren sind auch noch bei Zimmertemperatur sehr gute Isolatoren, und erst bei hohen Temperaturen, ggf. dicht am Schmelzpunkt, kann man vom einem, wenn auch sehr schlechten, Leiter sprechen [MCK93].

Halbleiter lassen sich in zwei Klassen einteilen: direkte und indirekte Halbleiter. Bei den direkten wie den meisten III-V- und II-VI-Halbleitern liegt das Maximum der oberen Valenzbandkante beim gleichen $\vec{k}$ -Wert wie das Minimum der unteren Leitungsbandkante. Bei den indirekten wie beispielsweise den Element-Halbleitern Silizium und Germanium liegen sie bei verschiedenen $\vec{k}$ -Werten, siehe Abb. 1.1.

Das Auftreten von Bändern, die durch eine Energielücke voneinander getrennt sind, hat zur Folge, daß die spezifische Leitfähigkeit σ sehr stark von der Temperatur abhängt. Mit zunehmender Temperatur wird das Leitungsband durch Generation von Elektronen-Loch-Paaren immer mehr mit Elektronen bevölkert. Löcher im Valenzband,

¹ Eine klare Trennung zwischen Halbleitern und Isolatoren ist nicht möglich. So finden sich in der Literatur verschiedene Angaben für die Grenzenergie: $2.50 \text{ eV} < E_G < 3.00 \text{ eV}$ [WEI95a].

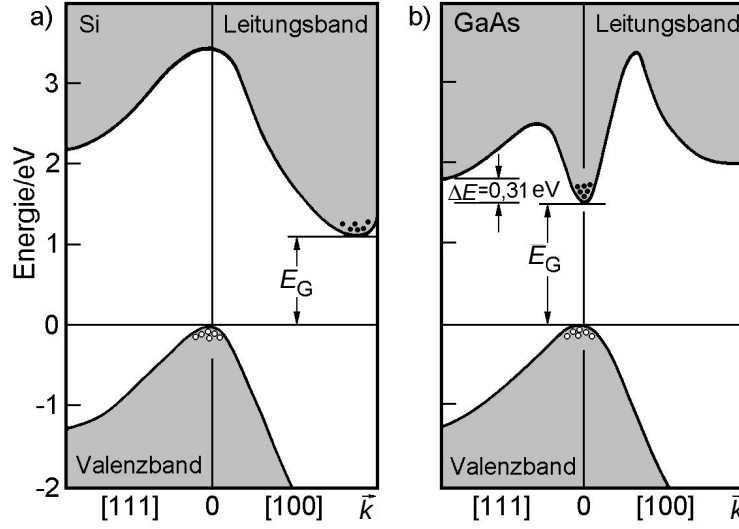


Abb. 1.1: Bandstruktur von a) Silizium und b) Galliumarsenid [SZE85]

die auch als Defektelektronen bezeichnet werden, tragen ebenfalls zur Leitfähigkeit bei:

$$\sigma = \sigma_n + \sigma_p = q(n\mu_n + p\mu_p). \quad (1.1)$$

$\mu_n = \frac{q\tau_n}{m_n^*}$ ist die Beweglichkeit der Elektronen im Leitungsband, wobei τ_n eine Relaxationszeit bedeutet, deren Größe von den Wechselwirkungen der Leitungselektronen mit atomaren Störstellen und akustischen Phononen abhängt. m_n^* ist die effektive Masse der Elektronen im Leitungsband. Die Beweglichkeit sagt aus, wie groß die mittlere Geschwindigkeit $\vec{v}_n$ der Elektronen ist, die sich längs eines elektrischen Feldes $\vec{E}$ bewegen: $\vec{v}_n = \mu_n \vec{E}$. Entsprechende Bedeutungen haben die Größen μ_p , τ_p und m_p^* für die Löcher im Valenzband.

1.1.1 Intrinsische Halbleiter

Als intrinsisch wird ein undotierter Halbleiter bezeichnet². Jedes Elektron, das ins Leitungsband gelangt ist, muß im Valenzband ein Loch hinterlassen haben.

Die Anzahl der Elektronen im Leitungsband im Energieintervall dE beträgt: $dn(E) = f_n(E)D_C(E)dE$. Die Fermi-Verteilung³ $f_n(E) = [1 + \exp(\frac{E-E_F}{kT})]^{-1}$ gibt die Wahrscheinlichkeit an, ob ein Zustand der Energie E mit einem Elektron besetzt ist. $D_C(E) = \frac{8\sqrt{2}\pi}{h^3} \cdot m_n^{*3/2} \cdot \sqrt{E - E_C}$ ist die Zustandsdichte im Leitungsband⁴. Um die

² Strenggenommen gibt es keine intrinsischen Halbleiter, da aus thermodynamischen Gründen das Herstellen eines Halbleiters ohne Fremdatome unmöglich ist. Heutige Einkristalle, besonders die der Element-Halbleiter Si und Ge, sind aber so rein, daß deren elektrische Eigenschaften denen der idealen intrinsischen Halbleiter sehr nahe kommen [MCK93].

³ Da Elektronen Fermionen sind, gehorchen sie der Fermi-Dirac-Statistik.

⁴ Die angegebene Zustandsdichte ist die des freien Elektronengases. Die reale und materialabhängige

Ladungsträgerdichte n zu berechnen, muß über die Beiträge aller Zustände im Leitungsband $dn(E)$ summiert werden. Da die Fermi-Verteilung $f_n(E)$ für $E \rightarrow \infty$ rasch gegen null konvergiert, darf die obere Integrationsgrenze nach ∞ verschoben werden⁵:

$$n = \int_{E_C}^{\infty} f_n(E) D_C(E) dE = N_C \exp\left(-\frac{E_C - E_F}{kT}\right), N_C = 2 \left(\frac{2\pi m_n^* kT}{h^2}\right)^{3/2}. \quad (1.2)$$

Entsprechend gilt für die Konzentration der Löcher im Valenzband, wobei $f_p(E) = 1 - f_n(E)$ die Wahrscheinlichkeit angibt, ob ein Zustand der Energie E unbesetzt ist:

$$p = \int_{-\infty}^{E_V} f_p(E) D_V(E) dE = N_V \exp\left(-\frac{E_F - E_V}{kT}\right), N_V = 2 \left(\frac{2\pi m_p^* kT}{h^2}\right)^{3/2}. \quad (1.3)$$

Dabei bedeuten N_C und N_V die effektiven Zustandsdichten im Leitungs- bzw. Valenzband. Als unbekannte Größe enthalten die Ladungsträgerkonzentrationen noch die Fermi-Energie E_F . Die Bildung des Produktes der Gleichungen 1.2 und 1.3 ergibt:

$$n \cdot p = N_C N_V \exp\left(-\frac{E_C - E_V}{kT}\right) = N_C N_V \exp\left(-\frac{E_G}{kT}\right) = n_i^2, \quad (1.4)$$

$n_i = 2 \left(\frac{2\pi \sqrt{m_n^* m_p^*} kT}{h^2}\right)^{3/2} \exp\left(-\frac{E_G}{2kT}\right)$ ist die intrinsische Ladungsträgerkonzentration. Das Produkt $n \cdot p$ hängt *nicht* von der Fermi-Energie E_F , sondern nur von der Temperatur T und dem Bandabstand E_G ab⁶. Gleichsetzen der Gleichungen 1.2 und 1.3 ergibt schließlich für die Fermi-Energie E_F eines intrinsischen Halbleiters:

$$E_F = \frac{1}{2}(E_C + E_V) + \frac{3}{4}kT \ln\left(\frac{m_p^*}{m_n^*}\right). \quad (1.5)$$

Für die Element-Halbleiter Silizium und Germanium gilt in guter Näherung $m_n^* \approx m_p^*$ [ADA98, MCK93]. Für diese Halbleiter liegt die Fermi-Energie E_F in der Gapmitte, unabhängig von der Temperatur. Für die Komponenten-Halbleiter gilt diese Näherung meist nicht [ELS88]. Dennoch ist die Temperaturabhängigkeit der Fermi-Energie für die Komponenten-Halbleiter sehr klein. Zum einen, weil das Verhältnis der effektiven Massen nahe 1 noch logarithmiert wird und zum anderen, weil der Faktor $kT \approx \frac{1}{40}$ eV (bei Raumtemperatur) viel kleiner als der Bandabstand ist, z. B. $E_G = 1.12$ eV für Silizium.

Zustandsdichte zeigt eine äußerst komplizierte Struktur. Da der wahre funktionale Zusammenhang zwischen Zustandsdichte und Energie im wesentlichen eine $\sqrt{E - E_C}$ -Abhängigkeit zeigt, stellt das Verwenden der Zustandsdichte des freien Elektronengases i. a. eine gute Näherung dar.

⁵Um die Ladungsträgerkonzentration n zu berechnen, darf in Gleichung 1.2 die Fermi-Verteilung $f_n(E)$ durch die Boltzmann-Verteilung $f_n^B(E) = \exp\left(-\frac{E - E_F}{kT}\right)$ genähert werden, weil für $E \geq E_C$, auch bei höheren Temperaturen, $\frac{E - E_F}{kT} \gg 1$ gilt, was die Boltzmann-Näherung erlaubt und eine analytische Auswertung des Integrals erst möglich macht.

⁶Gleichung 1.4 entspricht dem Massenwirkungsgesetz, das auch in der Chemie Anwendung findet. Dort beschreibt es z. B. die Dissoziation von Wasser in H^+ - und OH^- -Ionen. Hier steht es für die Dissoziation von kovalenten Bindungen in Elektron-Loch-Paare.

1.1.2 Dotierte Halbleiter

Dotieren von Silizium und Germanium erfolgt durch Zugabe geringer Mengen drei- oder fünfwertiger Elemente, die während des Kristallwachstums substitutionell ins Gitter eingebaut werden. Technisch relevante Konzentrationen liegen im Bereich 10^{14} - $10^{18} \frac{1}{\text{cm}^3}$.

Im Falle der Dotierung mit einem dreiwertigen Element kann eines der vier umgebenden Si-Atome nicht alle Valenzen absättigen. Das Besetzen dieser Bindungslücke mit einem Elektron erfordert eine Energie von nur einigen meV. Ein Elektron, das diese Störstelle ionisiert, hinterläßt im Valenzband ein Loch. Da diese Art von Störstellen Elektronen aufnehmen, bezeichnet man sie als *Akzeptoren* und den Halbleiter als p-leitend.

Entsprechend ergibt sich bei Dotierung mit einem fünfwertigen Element, daß die Störstelle sehr leicht ihr Elektron an das Leitungsband abgeben kann, wofür ebenfalls nur einige meV erforderlich sind. Störstellen, die ihr Elektron abgeben, werden als *Donatoren* und der Halbleiter als n-leitend bezeichnet.

Um die Ladungsträgerkonzentration von dotierten Halbleitern zu berechnen, verwendet man als Ansatz die Neutralitätsbedingung:

$$\underbrace{p + N_D^+}_{\text{pos. Ladungen}} = \underbrace{n + N_A^-}_{\text{neg. Ladungen}}.$$

N_A^- ist die Konzentration der ionisierten Akzeptoren und N_D^+ die Konzentration der ionisierten Donatoren:

$$N_A^- = \frac{N_A}{1 + \exp\left(\frac{E_A - E_F}{kT}\right)} \quad \text{und} \quad N_D^+ = \frac{N_D}{1 + \exp\left(\frac{E_F - E_D}{kT}\right)}. \quad (1.6)$$

Setzt man diese beiden sowie die Gleichungen 1.2 und 1.3 in die Neutralitätsbedingung ein, so ergibt sich eine transzendente Bestimmungsgleichung für die Fermi-Energie⁷:

$$\begin{aligned} & N_V \exp\left(-\frac{E_F - E_V}{kT}\right) + \frac{N_D}{1 + \exp\left(\frac{E_F - E_D}{kT}\right)} \\ &= N_C \exp\left(-\frac{E_C - E_F}{kT}\right) + \frac{N_A}{1 + \exp\left(\frac{E_A - E_F}{kT}\right)}. \end{aligned} \quad (1.7)$$

Für den Fall, daß der Halbleiter nur eine Sorte Dotierstoff enthält, kann Gleichung 1.7 für genügend tiefe Temperaturen, d. h. $\frac{E_F - E_A}{kT} \gg 1$ bzw. $\frac{E_D - E_F}{kT} \gg 1$, nach der Fermi-Energie aufgelöst werden. Beispielsweise für einen n-Halbleiter:

$$E_F = E_D + kT \ln \left[\frac{1}{2} \left(\sqrt{1 + \frac{4N_D}{N_C} \exp\left(\frac{E_C - E_D}{kT}\right)} - 1 \right) \right]. \quad (1.8)$$

⁷Das Verwenden der Boltzmann-Näherung ist hier nicht erlaubt, da für $E \geq E_C$ die Bedingung $\frac{E - E_F}{kT} \gg 1$ nicht immer erfüllt ist, vgl. Fußnote auf Seite 7.

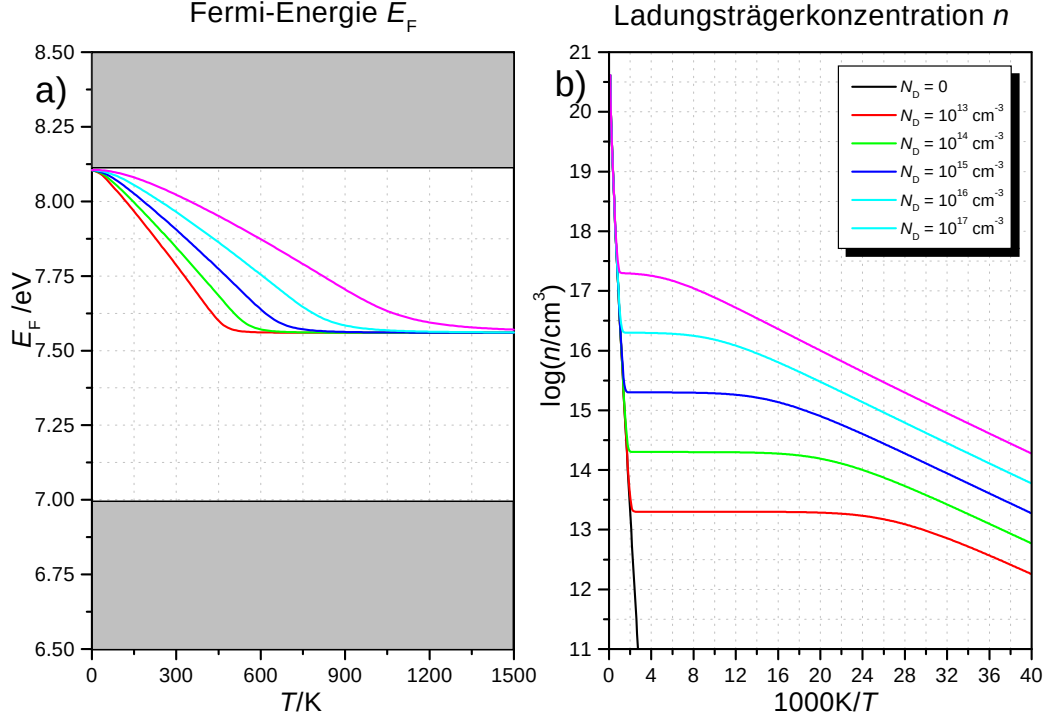


Abb. 1.2: a) Temperaturabhängigkeit der Fermi-Energie E_F für verschiedene Donorkonzentrationen N_D berechnet für $E_C - E_D = 0.03 \text{ eV}$, $E_G = 1.12 \text{ eV}$, $m_n^* = 0.3m_0$ (m_0 ist die Elektronenmasse) und $N_A = 0$. Das Niveau der unteren Valenzbandkante wurde auf 7.0 eV gelegt. b) Zugehörige Ladungsträgerkonzentrationen n .

Abb. 1.2 a) zeigt die Abhängigkeit der Fermi-Energie E_F und b) der Ladungsträgerkonzentration n von der Temperatur T für verschiedene Dotierkonzentrationen N_D . Die Ladungsträgerkonzentrationen n wurden berechnet, indem die Fermi-Energien aus a) in Gleichung 1.2 eingesetzt wurden. Je höher die Temperatur, um so mehr liegt die Fermi-Energie in der Gapmitte und um so mehr ähneln die Eigenschaften des dotierten einem intrinsischen Halbleiter. Für den Bereich in dem $|\frac{dE_F}{dT}|$ groß ist, ist die Ladungsträgerkonzentration über einen weiten Temperaturbereich annähernd konstant. Ein Vergleich mit der Ladungsträgerkonzentration für den intrinsischen Fall (schwarze Kurve in Abb. 1.2 b) zeigt, daß bereits geringe Dotierkonzentrationen besonders bei tiefen Temperaturen einen extrem starken Einfluß auf die Ladungsträgerkonzentration haben.

1.2 Halbleiter-Metall-Kontakte

Der sperrende Halbleiter-Metall-Kontakt war das erste Halbleiterbauteil überhaupt und ermöglichte um 1904 bereits zahlreiche Anwendungen. Erst 1938 schlug *Walter Schottky* vor, daß die sperrenden Eigenschaften dieser Kontakte ihre Ursache in einer Potentialbarriere haben, die durch feste Raumladungen innerhalb des Halbleiters hervorgerufen werden.

1.2.1 Schottky-Kontakte

Abb. 1.3 zeigt das Banddiagramm eines idealen Schottky-Kontaktes auf einem p-Halbleiter. Die Auslösearbeit des Metalls sowie die des Halbleiters beziehen sich auf das überall einheitliche Vakuumniveau. Da beide Auslösearbeiten in der Regel nicht identisch sind, befinden sich die Fermi-Energien von Metall und Halbleiter auf unterschiedlichen Niveaus. Werden nun Metall und Halbleiter in engen Kontakt zueinander gebracht, so findet solange ein Ladungstransfer statt, in diesem Fall eine Injektion von Elektronen aus dem Metall in den Halbleiter, bis sich die Fermi-Niveaus von Metall und Halbleiter einander angeglichen haben. Im Bereich der Oberfläche des Halbleiters wird also ein Teil der Löcher im Valenzband von Elektronen aus dem Metall neutralisiert. Da die bei Raumtemperatur stets ionisierten Akzeptoren ortsfest sind, kommt es zur Ausbildung einer negativen Raumladung im Bereich der Oberfläche des Halbleiters, die durch eine gleich große positive Oberflächenladung im Metall kompensiert wird (siehe Abb. 1.3 b).

Die sich einstellende Barrierenhöhe $q\phi_B$ ergibt sich aus dem Bandabstand E_G abzüglich der Differenz zwischen der Auslösearbeit des Metalls ϕ_M und der Affinität des Halbleiters χ :

$$q\phi_B = E_G - q(\phi_M - \chi). \quad (1.9)$$

Um die Tiefe der *Raumladungszone* W zu berechnen, verwendet man die Poisson-Gleichung, die eine örtliche Ladungsverteilung $\rho(x)$ mit dem aus ihr resultierenden

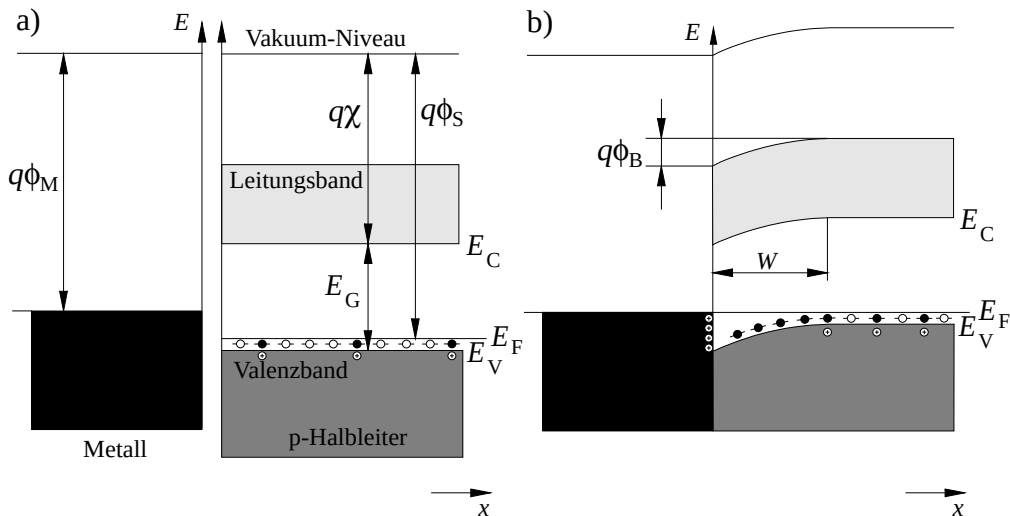


Abb. 1.3: Sperrender Halbleiter-Metall-Kontakt, a) vor und b) nach Einstellung des thermischen Gleichgewichtes.

Verlauf des elektrischen Potentials $\Phi(x)$ verknüpft⁸:

$$\frac{d^2\Phi(x)}{dx^2} = -\frac{\rho(x)}{\epsilon\epsilon_0} = \frac{qN_A}{\epsilon\epsilon_0} \quad \text{für } 0 \leq x \leq W. \quad (1.10)$$

Unter Berücksichtigung der Randbedingung $\Phi(W) = 0$ und $\left.\frac{d\Phi(x)}{dx}\right|_{x=W} = 0$ ergibt sich nach zweimaliger Integration:

$$\Phi(x) = \frac{1}{2} \cdot \frac{qN_A}{\epsilon\epsilon_0} \cdot (x - W)^2. \quad (1.11)$$

Aus der Bedingung, daß $\Phi(0) - \Phi(W) = V_{bi} - V$ sein muß, wobei V_{bi} das 'Built-in'-Potential und V die über dem Kontakt anliegende Spannung bedeuten, erhält man schließlich für die Weite der Raumladungszone W :

$$W = \sqrt{\frac{2(V_{bi} - V)\epsilon\epsilon_0}{qN_A}}. \quad (1.12)$$

Die letzte Gleichung zeigt, daß der Schottky-Kontakt sperrende Eigenschaften hat: Gilt $V > 0$, so wird die Raumladungszone um so kleiner, je größer V ist. Dieser Fall entspricht der Durchlaßrichtung. Gilt hingegen $V < 0$, wird die Tiefe der Raumladungszone mit zunehmender Spannung immer größer. In diesem Fall wird der Schottky-Kontakt in Sperrichtung betrieben.

Für nicht zu hohe Dotierungen N_A findet der Ladungstransport über die Potentialbarriere hauptsächlich durch thermische Emission statt. Dann gilt für die J - V -Kennlinie eines Schottky-Kontaktes:

$$J = J_s \left[\exp\left(\frac{qV}{kT}\right) - 1 \right], \quad \text{mit } J_s = A^*T^2 \exp\left(-\frac{q\phi_B}{kT}\right). \quad (1.13)$$

A^* ist die *effektive Richardson-Konstante* und beträgt $110 \frac{\text{A}}{\text{K}^2\text{cm}^2}$ für n-Silizium und $32 \frac{\text{A}}{\text{K}^2\text{cm}^2}$ für p-Silizium [RID74, SZE85].

1.2.2 Ohmsche Kontakte

Für den Fall, daß die Auslösearbeit des Metalls ϕ_M größer ist als die des Halbleiters ϕ_S , erhält man einen ohmschen Kontakt. Dieser ist dadurch gekennzeichnet, daß sich im Bereich der Halbleiteroberfläche nicht eine Verarmungsschicht wie beim, Schottky-Kontakt, sondern eine *Anreicherungs-schicht* bildet, da es in diesem Fall zu einer Löcherinjektion in das Valenzband des p-Halbleiters kommt. Abb. 1.4 zeigt die Bandstruktur des Kontaktes a) vor und b) nach Einstellung des thermischen Gleichgewichtes.

Ohmschen Kontakten kommt, besonders bei Leistungsbau-elementen, eine besondere Bedeutung zu. Eine äußere Spannung soll nicht über ohmschen Kontaktierungen,

⁸Daß die Ladungsverteilung $\rho(x)$ für $x \leq W$ keine Ortsabhängigkeit besitzt, ist strengenommen nicht richtig. Kurz vor der neutralen Zone werden die ionisierten Akzeptoren innerhalb eines schmalen Bereiches teilweise von Löchern kompensiert. Dadurch erfolgt das Anwachsen von $\rho(x)$ stetig und nicht sprunghaft.

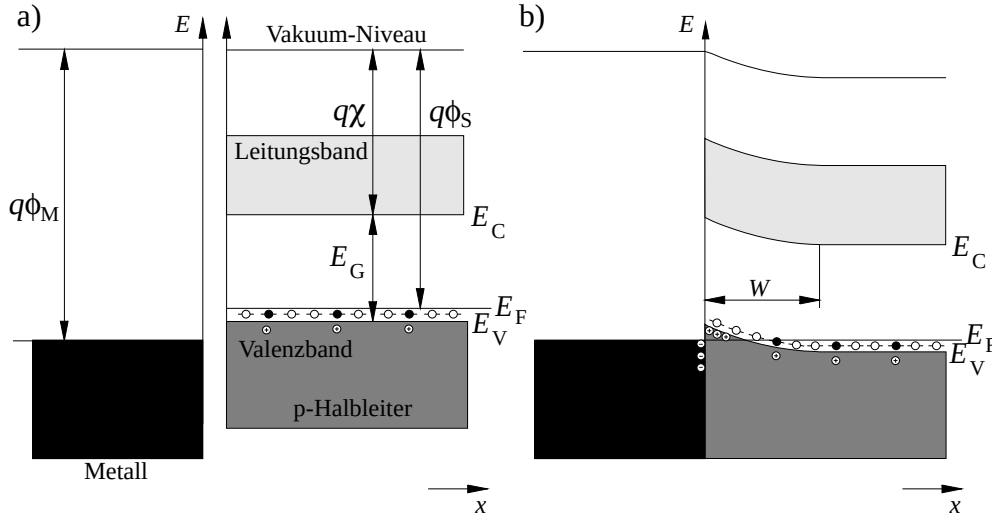


Abb. 1.4: Ohmscher Halbleiter-Metall-Kontakt, a) vor und b) nach Einstellung des thermischen Gleichgewichtes. W bedeutet die Weite der Anreicherungszone.

sondern über den aktiven Bereichen des Bauteiles abfallen. Daher sollte der *spezifische Kontaktwiderstand*

$$\rho_c = \left(\frac{\partial J}{\partial V} \right)_{V=0}^{-1} \quad (1.14)$$

möglichst klein sein. Auch gute Solarzellen erfordern ohmsche Kontakte mit kleinen spezifischen Kontaktwiderständen, denn schlechte Rückseitenkontakte erhöhen den Serienwiderstand R_s der Solarzelle. Die Folge ist ein vergleichsweise schlechter Wirkungsgrad η , da solche Widerstände den Füllfaktor verkleinern.

Abb. 1.5 zeigt weitere Möglichkeiten zur Erzeugung ohmscher Kontakte. Wird der Halbleiter bis zur Entartung dotiert, so kann die Potentialbarriere aufgrund der sehr dünnen Raumladungszone leicht von Ladungsträgern durchtunnelt werden. Diese Methode findet bei den III-V-Halbleitern mit großem Bandabstand vielfach Anwendung. Für diese Halbleiter existiert kein Metall mit genügend großer Auslösearbeit ϕ_M , so daß die Bedingung $\phi_M > \phi_S$ erfüllt wäre. Eine andere Möglichkeit besteht in der Erzeugung einer Heterostruktur aus einem Halbleiter mit großem Bandabstand und einem mit kleinem Bandabstand. Auf dem oberflächennahen Halbleiter mit kleinem Bandabstand können dann sehr leicht ohmsche Kontakte hergestellt werden. Schließlich kann z. B. durch Ionen-Implantation das oberflächennahe Gebiet des Halbleiters zerstört werden. In diesem amorphen Gebiet kann Ladungstransport über lokale 'Trap'-Zustände stattfinden [SEB82].

1.2.3 Reale Metall-Halbleiter-Kontakte

Experimentelle Untersuchungen an III-V-Halbleitern zeigten, daß die Barrierenhöhe ϕ_B *unabhängig* von der verwendeten Metallisierung ist [PIO83, RID74]. Bereits 1947 schlug

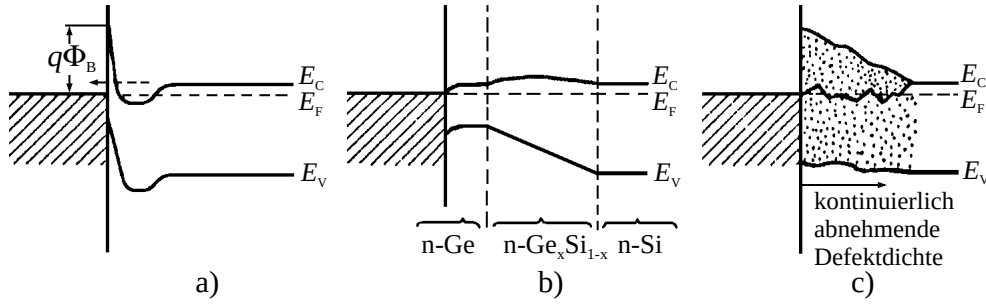


Abb. 1.5: Realisierungsmöglichkeiten ohmscher Kontakte auf n-Silizium: a) Im oberflächennahen Bereich des Halbleiters befindet sich eine n^{++} -Schicht. Die Raumladungszone ist sehr dünn, weshalb Elektronen die Potentialbarriere leicht durchtunneln können. b) Durch Erzeugung einer Ge-Si-Heterostruktur verringert sich zur Oberfläche hin der Bandabstand von $E_G = 1.12$ eV (Silizium) nach $E_G = 0.66$ eV (Germanium). c) Wird nahe der Oberfläche eine amorphe Schicht hergestellt, ist Ladungstransport über lokale Trap-Zustände möglich. Eingezeichnet ist ein möglicher Weg, den Elektronen nehmen können [KOR95, SZE81, WUE89].

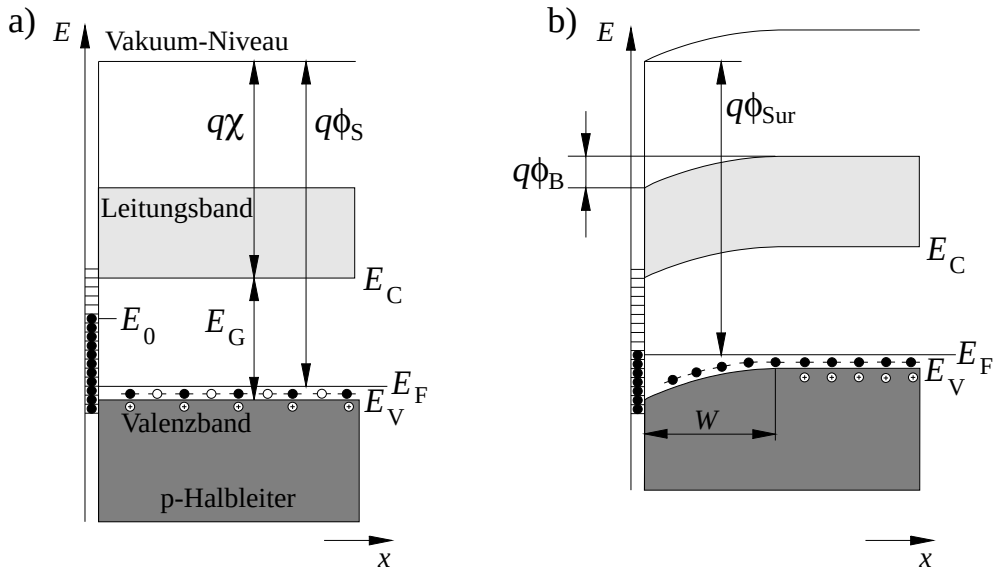


Abb. 1.6: Reale Halbleiteroberfläche a) vor und b) nach Einstellung des thermischen Gleichgewichtes.

Bardeen vor, daß nicht abgesättigte Valenzen (d. h. Änderungen in der Periodizität des Oberflächenpotentials), veränderte chemische Eigenschaften sowie Adsorbate an der Halbleiteroberfläche die *Oberflächenzustandsdichte* N_{SS} erheblich erhöhen können [BAR47].

Den Einfluß der Oberflächenzustände auf die Halbleiter-Vakuum-Grenzfläche zeigt Abb. 1.6. Wird beispielsweise ein Siliziumwafer gespalten, so befinden sich die beiden erzeugten Oberflächen unmittelbar nach der Spaltung nicht im thermischen Gleichge-

wicht. Die Oberflächenzustände, die gleichmäßig im 'Gap' verteilt sein mögen, seien bis zum Niveau E_0 mit Elektronen besetzt. Zwecks Einstellung des thermischen Gleichgewichtes besetzen Elektronen aus den Oberflächenzuständen Akzeptorniveaus⁹. Hat sich das thermische Gleichgewicht eingestellt, fällt das höchste besetzte Niveau der Oberflächenzustände mit dem Fermi-Niveau E_F zusammen. An der Halbleiteroberfläche hat sich eine Verarmungsschicht der Weite W gebildet.

Die Neutralitätsbedingung fordert, daß die Zahl der ionisierten Akzeptoren innerhalb der Verarmungszone gleich der Zahl der Überschußladungsträger in den Oberflächenzuständen ist:

$$N_{SS}(E_0 - \phi_B) = N_A W, \quad \text{oder} \quad (1.15)$$

$$E_0 - \phi_B = \frac{N_A W}{N_{SS}}. \quad (1.16)$$

Aus Gleichung 1.16 folgt, daß im Fall von $N_{SS} = \infty$ gilt: $\phi_B = E_0$, d. h. die Barrierrhöhe ϕ_B wird nur durch die Eigenschaften der Oberfläche bestimmt. Eine offene Bindung, die man auch als '*dangling bond*' bezeichnet, stellt einen Oberflächenzustand dar. Unter der Annahme, daß auf einer Si-Oberfläche jedes Si-Atom mit einem Zustand zu N_{SS} beträgt, ergibt sich $N_{SS} \approx (5 \cdot 10^{22})^{2/3} \frac{1}{\text{cm}^2 \text{ eV}} = 1.2 \cdot 10^{15} \frac{1}{\text{cm}^2 \text{ eV}}$. Für $N_A = 10^{15} \frac{1}{\text{cm}^3}$ beispielsweise wird $E_0 - \phi_B = 10^{-4} \text{ eV}$ und damit $\phi_B \approx E_0$. Tatsächlich ist der Wert für N_{SS} sehr viel kleiner durch H-Terminierung (Passivierung) und durch Rekonstruktion der '*dangling bonds*' untereinander. In beiden Fällen werden die Oberflächen-Niveaus in die Bänder hinein verschoben, was eine Verkleinerung von N_{SS} um mehrere Größenordnungen zur Folge hat.

Wird auf eine solche Oberfläche eine Metallisierung aufgebracht, entsteht ein realer Halbleiter-Metall-Kontakt. Das zugehörige Banddiagramm zeigt Abb. 1.7. Gilt $\phi_M > \frac{E_G}{q} - \chi$ (Bedingung für ohmschen Kontakt), so treten Elektronen aus den Oberflächenzuständen ins Metall über. Ist N_{SS} klein genug, so reichen die Elektronen in den Oberflächenzuständen nicht aus, um das thermische Gleichgewicht einzustellen. Es müssen weitere Elektronen aus dem Halbleiter in das Metall übertreten. Dies geschieht durch Injektion von Löchern in das Valenzband des Halbleiters, vgl. Abschnitt 1.2.2; dadurch wird der Kontakt durch Entstehung einer Anreicherungsschicht ohmsch.

Ist hingegen N_{SS} hinreichend groß, ist die Zahl an Elektronen, die sich auf Oberflächenzuständen befinden, ausreichend, um das thermische Gleichgewicht einzustellen. Hat sich dieses eingestellt, fällt das Potential an der Grenzfläche zwischen Halbleiter und Metall im Bereich atomarer Abstände um $\Delta\phi = \phi_M - \phi_{Sur}$. Die Anzahl der ins Metall übergetretenen Elektronen beträgt $N = \frac{Q}{q} = \frac{(\phi_M - \phi_{Sur})\epsilon\epsilon_0}{qa}$. Hierbei bedeutet a der Abstand der beiden geladenen Oberflächen. Gilt z. B. $\Delta\phi = 1 \text{ eV}$ und $a = 5 \cdot 10^{-10} \text{ m}$, so wird $N \approx 10^{13} \frac{1}{\text{cm}^3}$. Beträgt der Wert der Oberflächenzustandsdichte $N_{SS} \approx 10^{14} \frac{1}{\text{cm}^2 \text{ eV}}$, bedeutet dies, daß sich die Besetzungsgrenze der Oberflächenzustände lediglich um $\Delta E = \frac{1}{10} \text{ eV}$ erniedrigt, was einer Verringerung der Potentialbarriere um denselben Wert zur Folge hat. Obwohl ein Metall mit $\phi_M > \frac{E_G}{q} - \chi$ gewählt

⁹Damit Elektronen Akzeptor-Niveaus besetzen können, darf die Temperatur nicht so hoch sein, daß bereits alle Akzeptor-Niveaus mit Elektronen aus dem Valenzband besetzt sind, vgl. Abb. 1.2 b).

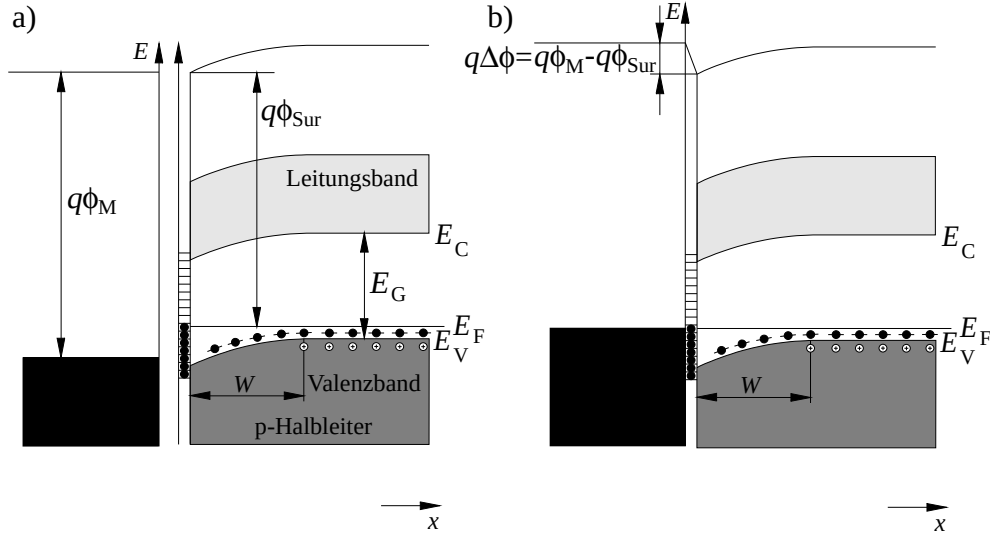


Abb. 1.7: Realer Halbleiter-Metall-Kontakt a) vor und b) nach Einstellung des thermischen Gleichgewichtes.

wurde, zeigt der Kontakt sperrende Eigenschaften¹⁰.

Neben den Oberflächenzuständen beeinflusst auch der *Schottky-Effekt* die Höhe der Potentialbarriere. Unter dem Schottky-Effekt versteht man eine Erniedrigung der Potential-Barriere bei Präsenz eines elektrischen Feldes über den Kontakt. Dabei induzieren oberflächennahe Ladungen im Halbleiter Spiegelladungen im Metall, die eine Erniedrigung der Potentialbarriere zur Folge haben. Somit wird die obere Valenzbandkante nach oben und die untere Leitungsbandkante nach unten verbogen. Das Minimum der Potentialbarriere verlagert sich ins Innere des Halbleiters. Durch diesen Bildkrafteinfluß erniedrigt sich die Höhe der effektiven Barriere um

$$\Delta\phi_{BK} = 4 \sqrt{\frac{q^2(\phi_B - \phi_0)N_A}{8\pi^2\epsilon\epsilon_d^2\epsilon_0^3}}, \quad (1.17)$$

wobei ϵ_d die dynamische Dielektrizitätskonstante des Halbleiters und ϕ_B die Höhe der Potentialbarriere bedeuten [PIO83, RID74, WUE89]. Nach Sze hat im allgemeinen der Halbleiter während des Transfers von Ladungsträgern über die Potentialbarriere genügend Zeit, um auf die Ladungsverschiebungen mit Polarisierung zu reagieren. Daher gilt gewöhnlich $\epsilon_d \approx \epsilon$ [SZE81].

¹⁰ Der oberflächennahe Bereich des Halbleiters verhält sich wie ein Metall, wenn die Oberflächenzustandsdichte N_{SS} hinreichend groß ist. Da die Besetzungsgrenze E_0 dieser Oberflächenzustände im allgemeinen im Gap liegt, erzeugt diese 'Metallisierung' einen Kontakt mit sperrenden Eigenschaften. Wird ein Metall auf diesem Halbleiter aufgebracht, erzeugt man mit diesem Metall strenggenommen nicht einen Halbleiter-Metall-Kontakt, sondern einen Metall-Metall-Kontakt mit dem Kontaktpotential $\Delta\phi = \phi_M - \phi_{Sur}$ ohne die Eigenschaften des Halbleiters mit seiner Oberflächen'metallisierung' wesentlich zu verändern.

1.3 Die Fest-Flüssig-Grenzfläche

Nicht nur für die Thematik dieser Arbeit ist die Fest-Flüssig-Grenzfläche von zentraler Bedeutung, sie ermöglicht auch technisch zahlreiche Anwendungen. Als wichtige Beispiele seien die zahlreichen naßchemischen Prozesse bei der Chipherstellung und der Blei-Akkumulator genannt, der unverzichtbarer Bestandteil eines jeden Kraftfahrzeuges ist.

1.3.1 Die Struktur der Grenzfläche

Die Strukturen der Grenzflächen zwischen Halbleiter und Elektrolyt sowie zwischen Metall und Elektrolyt unterscheiden sich elektrolytseitig nur geringfügig. Deshalb wird in diesem Abschnitt allein die Metall-Elektrolyt-Grenzfläche betrachtet.

Taucht eine Metallelektrode in einen Elektrolyten, in dem ein Salz des Metalls Me gelöst ist, so läuft im Bereich der Oberfläche folgende Reaktion ab:



Für den Fall, daß die Hinreaktion überwiegt, wird sich ein kleiner Teil der Kationen Me^{z+} auf der Metalloberfläche abscheiden und diese positiv aufladen (innere Helmholtz-Schicht). Diese positive Oberflächenladung des Metalls hat eine oberflächennahe Anziehung der Anionen zur Folge. Ist der Durchmesser der Anionen inklusive Solvathülle a , so beträgt der minimale Abstand der Anionen von der Oberfläche ungefähr $\frac{a}{2}$ (äußere Helmholtz-Schicht). Allerdings versucht die Wärmebewegung der Moleküle die äußere Helmholtz-Schicht aufzuweichen [CHA13, GOU10]. Dieser Einfluß hat zur Folge, daß sie nicht auf die unmittelbare Umgebung der Metalloberfläche begrenzt ist, sondern sich, je nach Molarität des Elektrolyten, mehr oder weniger weit in den Elektrolyten hinein erstreckt. Dieser diffuse Teil der äußeren Helmholtz-Schicht wird als Gouy-Chapman-Schicht bezeichnet.

Das Auftreten räumlich getrennter Ladungen zieht ein elektrisches Feld $\vec{E}(\vec{r})$ und wegen $\vec{E}(\vec{r}) = -\vec{\nabla}\phi(\vec{r})$ ein räumlich nicht konstantes elektrisches Potential $\phi(\vec{r})$ nach sich. Der Raum zwischen der Metalloberfläche und den Ladungsschwerpunkten der Kationen ist ladungsfrei. Für diesen Bereich hat die eindimensionale Poisson-Gleichung eine einfache Gestalt:

$$\frac{d^2\phi(x)}{dx^2} = -\frac{4\pi\rho(x)}{\epsilon\epsilon_0} = 0. \quad (1.19)$$

In diesem Bereich steigt das Potential linear von ϕ_{E} auf ϕ_{IH} . Ebenfalls ladungsfrei ist der Bereich zwischen den Ladungsschwerpunkten der Kationen und denen der Anionen, so daß in diesem Bereich das Potential linear von ϕ_{IH} auf ϕ_{AH} fällt, siehe Abb. 1.8.

Für die Ionendichte $n_i(x)$ der Sorte i ist innerhalb der diffusen Schicht die Boltzmann-Verteilung anzusetzen:

$$n_i(x) = n_i^0(x) \exp\left(-\frac{z_i q(\phi(x + x_{\text{D}}) - \phi_{\text{L}})}{kT}\right) \quad (1.20)$$

aus Gleichung 1.22:

$$\rho(\xi) = -\frac{q^2 N_A}{kT} \sum_i z_i^2 c_i^0 (\phi(\xi) - \phi_L) = -\frac{2q^2 N_A}{kT} I (\phi(\xi) - \phi_L). \quad (1.23)$$

$I = \frac{1}{2} \sum_i z_i^2 c_i^0$ ist die sogenannte *Ionenstärke* des Elektrolyten. Verwendet man zusätzlich die Abkürzung¹²

$$\frac{1}{\kappa} = \sqrt{\frac{8\pi q^2 N_A}{\epsilon \epsilon_0 kT}} I, \quad (1.24)$$

so ergibt sich schließlich für die Poisson-Gleichung:

$$\frac{d^2 \phi(\xi)}{d\xi^2} = \left(\frac{1}{\kappa}\right)^2 (\phi(\xi) - \phi_L), \quad \xi \geq 0. \quad (1.25)$$

Als Lösung dieser Gleichung erhält man unter Berücksichtigung der Randbedingungen $\phi(x_D) = \phi_{AH}$ und $\phi(x) \rightarrow \phi_L$ für $x \rightarrow \infty$:

$$\phi(x) = (\phi_{AH} - \phi_L) \exp\left(-\frac{x - x_D}{\kappa}\right) + \phi_L. \quad (1.26)$$

Gleichung 1.26 zeigt, daß mit zunehmender Konzentration des Elektrolyten die Ausdehnung der diffusen Schicht immer kleiner wird.

Abb. 1.8 zeigt den Potentialverlauf im Elektrolyten in Abhängigkeit vom Abstand. Der Potentialabfall im Elektrolyten setzt sich zusammen aus dem über der starren und dem über der diffusen Doppenschicht: $\Delta\phi = \Delta\phi_{\text{starr}} + \Delta\phi_{\text{diffus}}$. $\Delta\phi$ wird als *Galvani-Spannung* bezeichnet [WED87].

1.3.2 Das Energiediagramm der Fest-Flüssig-Grenzfläche

Die Ionen eines in einem polaren Lösungsmittel gelösten Redoxpaares (z. B. $\text{Fe}^{2+}/\text{Fe}^{3+}$) sind von einer geordneten Hülle aus Wassermolekülen umgeben, die als Solvathülle bezeichnet wird. Auch hier versucht die Wärmebewegung Unordnung in die H_2O -Moleküle der Solvathülle zu bringen. Je nach Anordnung der H_2O -Moleküle wird das Coulombfeld eines Ions mehr oder weniger stark abgeschirmt, so daß die potentiellen Energien der oxidierenden Spezies E_{ox} und der reduzierenden Spezies E_{red} in der Lösung zeitlichen Schwankungen unterworfen sind.

Die Energie-Niveaus sind gaußförmig um den wahrscheinlichsten Energiewert E_t verteilt¹³[CAM98, GER69]:

$$W(E) = \frac{1}{\sqrt{4\pi\lambda kT}} \exp\left(-\frac{(E_t - E)^2}{4\lambda kT}\right), \quad (1.27)$$

¹²Mit dieser Definition kann κ als die Dicke der diffusen Schicht interpretiert werden, d. h. am Ort $x = x_D + \kappa$ ist das Potential auf $\frac{1}{q} \cdot \phi_{AH}$ abgefallen, siehe auch Gleichung 1.26.

¹³ Der Vorfaktor stellt sicher, daß die integrale Wahrscheinlichkeit gleich eins ist.

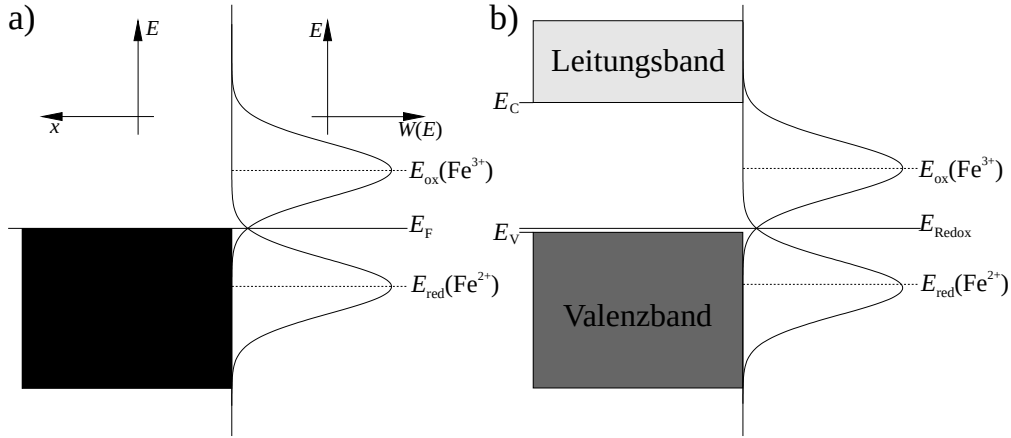


Abb. 1.9: Bandstruktur der Fest-Flüssig-Grenzfläche für ein polares Lösungsmittel: a) Metall-Elektrolyt- und b) p-Halbleiter-Elektrolyt-Kontakt jeweils im thermischen Gleichgewicht. $W(E)$ bedeutet die Wahrscheinlichkeit den Zustand E_{ox} bzw. E_{red} bei der Energie E zu finden, hier am Beispiel des Redoxpaares $\text{Fe}^{2+}/\text{Fe}^{3+}$ [CAM98, GER61, LOO81, MOR77].

wobei E_t für E_{ox} oder E_{red} steht. In dieser Gleichung ist λ die *Reorganisationsenergie*. Sie gibt die Energie an, die aufgewendet werden muß bzw. frei wird, wenn sich die äußere Solvathülle nach einem Elektronentransfer neu organisiert. Sie liegt typischerweise bei $\lambda = 1 \pm 0.5$ eV [MOR80, SCH96a].

In einem Elektrolyten liegen mindestens zwei verschiedene Ionensorten vor. Jede dieser Ionensorten besitzt im Elektrolyten ein eigenes Energie-Niveau. Positiv geladene Kationen haben, da sie ein Elektron aufnehmen können, oxidierende Eigenschaften (Elektronen-Akzeptor). Entsprechend haben Anionen reduzierende Eigenschaften, sie können ein Elektron abgeben (Elektronen-Donator). Die energetischen Lagen dieser Niveaus werden als E_{ox} bzw. E_{red} bezeichnet, wobei der Abstand der Niveaus $E_{\text{ox}} - E_{\text{red}} = 2\lambda$ beträgt (Franck-Condon-Shift¹⁴).

Für die Gleichgewichtsenergie E_{Redox} eines Elektrolyten nahe der Elektrodenoberfläche gilt die Nernstsche Gleichung

$$E_{\text{Redox}} = \frac{1}{2} (E_{\text{ox}} + E_{\text{red}}) + kT \ln \left(\frac{c_{\text{red}}}{c_{\text{ox}}} \right), \quad (1.28)$$

die im thermischen Gleichgewicht mit der Fermi-Energie E_F der Elektrode zusammenfällt¹⁵. E_{Redox} ist schwach konzentrationsabhängig, und nur wenn $c_{\text{ox}} = c_{\text{red}}$ ist, liegt E_{Redox} genau in der Mitte zwischen E_{ox} und E_{red} [CHA81, GER69, MOR77, MOR80].

¹⁴ Die Bezeichnung 'Franck-Condon-Shift' rührt daher, daß bei der Berechnung des Abstandes $E_{\text{ox}} - E_{\text{red}}$ das Franck-Condon-Prinzip Anwendung fand. Dieses nimmt an, daß der Elektronentransfer so schnell ist, daß die Ionen während dieses Vorganges quasi 'eingefroren' sind, was die Berechnungen vereinfacht.

¹⁵ Strenggenommen müssen in Gleichung 1.28 die Konzentrationen c_i durch die Aktivitäten $a_i = f_i c_i$ ersetzt werden. Für nicht zu hohe Elektrolytkonzentrationen sind die Aktivitätskoeffizienten $f_i \approx 1$ [HAM98].

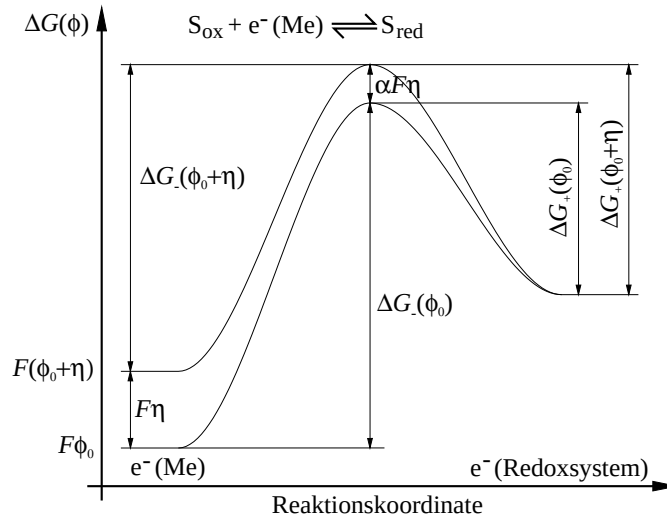


Abb. 1.10: Innere Energie eines Elektrons beim Übergang von der Elektrode zur Redoxkomponente in Lösung (und umgekehrt). Die Änderung des Potentials von ϕ_0 nach $\phi_0 + \eta$ entspricht einer Potentialverschiebung in negativer Richtung [HAM98, PAR51].

Abb. 1.9 zeigt das Banddiagramm a) eines Metall-Elektrolyt- und b) eines p-Halbleiter-Elektrolyt-Kontaktes im thermischen Gleichgewicht.

1.3.3 Die Durchtritts-Strom-Spannungs-Kurve

Wird an die Elektrode eine Spannung angelegt, die auch als *Überspannung* bezeichnet wird, so wird das Gleichgewicht an der Grenzfläche gestört. Die anodische und die kathodische Teilreaktion heben sich im zeitlichen Mittel nicht mehr auf. Eine Teilreaktion wird beschleunigt, die andere gehemmt: Es fließt ein makroskopischer Strom über den Kontakt.

In Abb. 1.10 ist die innere Energie in Abhängigkeit der Reaktionskoordinate dargestellt. Der qualitative Verlauf erklärt sich über die Theorie des aktivierten Komplexes. Hierbei wird angenommen, daß ein Teilchen, das am Ladungstransfer beteiligt ist, im Verlauf der chemischen Reaktion zwischenzeitlich einen Zustand höherer Energie einnimmt, der dem des aktivierten Komplexes entspricht. In diesem Zustand findet die chemische Reaktion statt (hier der Ladungstransfer), und das Teilchen nimmt dann wieder einen Zustand niedrigerer Energie ein.

Wird das Potential der Elektrode verändert, so ändert sich die Freie Aktivierungsenthalpie für die in anodischer Richtung ablaufende Teilreaktion:

$$\Delta G_+(\phi_0 + \eta) = \Delta G_+(\phi_0) - \alpha F \eta. \quad (1.29)$$

α ist der sogenannte *Durchtrittsfaktor* ($0 < \alpha < 1$)¹⁶. Dieser gibt Auskunft darüber,

¹⁶ Werden beim Ladungstransfer nicht ein, sondern n Elektronen übertragen, so ändert sich die Freie Aktivierungsenthalpie um $\alpha n F \eta$.

wie symmetrisch die Durchtritts-Strom-Spannungs-Kurve ist, siehe Abb. 1.11. Entsprechend ergibt sich aus Abb. 1.10 für die Änderung der Freien Aktivierungsenthalpie der in kathodischer Richtung ablaufenden Teilreaktion:

$$\Delta G_-(\phi_0 + \eta) = \Delta G_-(\phi_0) + (1 - \alpha)F\eta. \quad (1.30)$$

Für den Zusammenhang zwischen der Freien Aktivierungsenthalpie $\Delta G_+(\eta)$ und der anodischen Teilstromdichte $j_+(\eta)$ gilt [HAM98]:

$$\begin{aligned} j_+(\eta) &= Fc_{\text{red}}k_0^+ \exp\left(-\frac{\Delta G_+(\phi_0) - \alpha F\eta}{RT}\right) \\ &= j_0 \exp\left(\frac{\alpha F\eta}{RT}\right). \end{aligned} \quad (1.31)$$

Hierbei ist $j_+(0) = j_0 = Fc_{\text{red}}k_0^+ \exp\left(-\frac{\Delta G_+(\phi_0)}{RT}\right)$ die sogenannte *Austauschstromdichte*. Entsprechend gilt für die kathodische Teilstromdichte $j_-(\eta)$:

$$j_-(\eta) = -j_0 \exp\left(\frac{(1 - \alpha)F\eta}{RT}\right). \quad (1.32)$$

Die über die Grenzfläche fließende Gesamtstromdichte j ist die Summe beider Teilstromdichten:

$$j = j_+ + j_- = j_0 \left[\exp\left(\frac{\alpha F\eta}{RT}\right) - \exp\left(\frac{(1 - \alpha)F\eta}{RT}\right) \right]. \quad (1.33)$$

Dies ist die Gleichung für die Durchtritts-Strom-Spannungs-Kurve der Fest-Flüssig-Grenzfläche, die man auch als Butler-Volmer-Gleichung bezeichnet [GER99, ENG92, HAM98, SCH96a].

Die Durchtritts-Strom-Spannungs-Kurve von Metallen ist in der Regel symmetrisch, d. h. für sie gilt $\alpha \approx \frac{1}{2}$. Bei Halbleiterelektroden setzt sich der über die Grenzfläche fließende Strom aus vier Teilströmen zusammen, jeweils die anodischen und kathodischen Teilströme des Valenz- und die des Leitungsbandes. Für den anodischen Teilstrom ins Leitungsband gilt $j_+^C = nFc_{\text{red}}k_0E_C D_C(E_C)$. Da die Anzahl freier Zustände nahe der unteren Leitungsbandkante sehr hoch ist, ändert sich die Besetzung des Leitungsbandes nur wenig mit der Überspannung¹⁷ η . Daher ist $j_+^C = j_0^C = \text{const.}$ Andererseits hängt der kathodische Strom aus dem Leitungsband sehr stark von der Überspannung η ab, da die Konzentration der Elektronen im Leitungsband an der Halbleiteroberfläche n_e^C , die nötig ist, um die oxidierte Spezies im Elektrolyten zu reduzieren, exponentiell mit der Überspannung steigt [GER61, SUZ83]: $j_-^C = nFc_{\text{ox}}k_0N_C \exp\left(\frac{q\eta}{kT}\right) = j_0^C \exp\left(\frac{q\eta}{kT}\right)$. Damit ist der über das Leitungsband fließende Gesamtstrom:

$$j^C = j_+^C - j_-^C = j_0^C \left[1 - \exp\left(-\frac{q\eta}{kT}\right) \right] = j_0^C \left[1 - \exp\left(-\frac{F\eta}{RT}\right) \right]. \quad (1.34)$$

¹⁷ Das gilt nur dann, wenn die Fermi-Energie E_F im Gap liegt, der Halbleiter also nicht entartet ist.

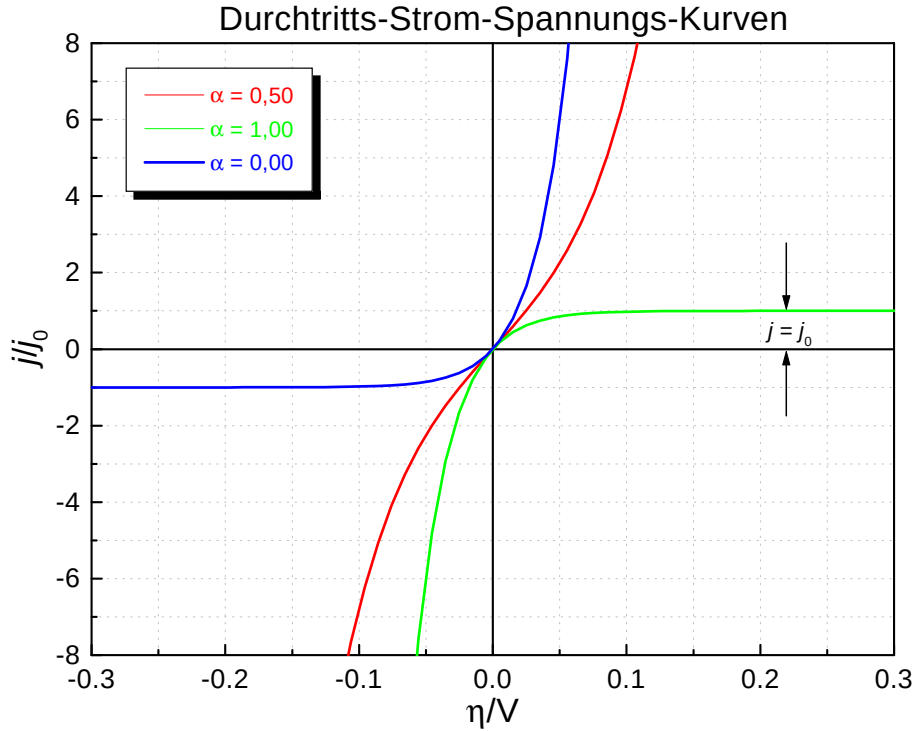


Abb. 1.11: Durchtritts-Strom-Spannungs-Kurven der Fest-Flüssig-Grenzfläche für verschiedene Durchtrittsfaktoren α .

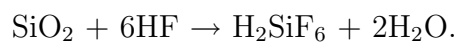
Dies ist die Butler-Volmer-Gleichung mit einem anodischen Durchtrittsfaktor von 0 und einem kathodischen Durchtrittsfaktor von 1. Ganz analog ergibt sich für den über das Valenzband fließenden Strom

$$j^V = j_0^V \left[\exp\left(\frac{F\eta}{RT}\right) - 1 \right] \quad (1.35)$$

wieder die Butler-Volmer-Gleichung, allerdings mit einem anodischen Durchtrittsfaktor von 1 und einem kathodischen Durchtrittsfaktor von 0. Zum Gesamtstrom über die Grenzfläche tragen beide Bänder bei. Welches Band den größeren Beitrag liefert, hängt von der Lage der Fermi-Energie E_F ab. Das Band, zu dem der Fermi-Energie den kleineren Abstand hat, wird den Hauptanteil zum Gesamtstrom liefern, im Falle eines p-Halbleiters also das Valenzband [BER99, MAN90, SCH96a].

1.3.4 Der Silizium-Flußsäure-Kontakt

Nach längerer Lagerung an Luft ist Silizium mit einer mehrere Nanometer dicken Schicht aus Siliziumoxid bedeckt. Dieses Oxid kann chemisch nur mit wässriger Flußsäure entfernt werden [SOM85, TAR95, TUC75, TUR58]:



Die Moleküle bzw. Ionen, die in wässriger HF vorliegen, sind HF, H^+ , F^- , HF^- und $(\text{HF})_2$. Da Flußsäure schwach dissoziierend ist, ist die Konzentration an HF-Molekülen sehr viel größer als die Konzentrationen der Ionen [BER99, SER94, VER94]:

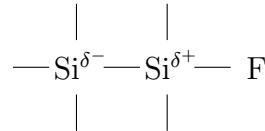


1.3.4.1 Der Mechanismus der Wasserstoffterminierung

Im Verlauf des Ätzvorganges werden die nicht abgesättigten Valenzen ('dangling bonds') des Siliziums mit Wasserstoff abgesättigt. Auf diese Weise entsteht eine für längere Zeit chemisch sehr stabile Oberfläche, die sonst, da Silizium ein sehr reaktives Element ist, sofort wieder oxidieren oder auf andere Art und Weise reagieren würde.

Lange Zeit wurde angenommen, daß die freien Valenzen der Siliziumoberfläche mit Fluorionen abgesättigt werden. Grund für diese Annahme ist die äußerst feste Bindung zwischen Silizium und Fluor: $E_{\text{SiF}} = 6.0 \text{ eV}$, wohingegen die Bindung zwischen Silizium und Wasserstoff weitaus schwächer ist: $E_{\text{SiH}} = 3.5 \text{ eV}$ [BER99, HIG90, YAB86]. Daß die Siliziumoberfläche mit Wasserstoff terminiert ist, hat kinetische und weniger thermodynamische Gründe:

Nach Entfernung des Oxides mit wässriger HF sind die 'dangling bonds' des Siliziums mit Fluor terminiert [JUD71, HOL66]. Die stark polare Si-F-Bindung polarisiert die rückseitige Si-Si-Bindung:



Diese Polarisierung ermöglicht die Reaktion mit einem HF-Molekül, wobei die $\text{Si}^{\delta-}\text{---Si}^{\delta+}$ -Bindung aufgebrochen wird. Dabei reagiert das Fluor mit dem Silizium-Atom, das den positiven Ladungsschwerpunkt trägt. Die dabei auftretende Oxidation des Fluor-Atoms zieht die Generation eines oberflächennahen Elektrons im Silizium nach sich¹⁸. Dieses Elektron wird dem Silizium bei der darauf folgenden Wasserstoffreduktion wieder entzogen. Dabei bildet sich nicht $\frac{1}{2}\text{H}_2\uparrow$, sondern der Wasserstoff reagiert mit dem Silizium-Atom und sättigt die offene Bindung ab, vgl. Abb. 1.12 b). Die nun vorhandenen zwei Fluor-Atome haben eine noch stärkere Polarisierung der rückseitigen Bindungen zur Folge. Nach Reaktionen mit zwei weiteren HF-Molekülen hat sich SiF_4 gebildet, das Si-Atom ist nun vollständig von der Oberfläche gelöst. Das SiF_4 -Molekül reagiert mit zwei HF-Molekülen zu H_2SiF_6 , der als wasserlöslicher Komplex in Lösung geht. Zurück bleibt eine Siliziumoberfläche, deren Valenzen elektrolytseitig mit Wasserstoff abgesättigt sind und eine weitere Reaktion mit HF-Molekülen verhindern, siehe Abb. 1.12 d) [CHA89].

¹⁸ Da es freie Elektronen in Elektrolyten nicht gibt, findet unmittelbar nach der Oxidation des Fluor-Atoms ein Elektronentransfer ins Silizium statt. Eine andere Möglichkeit wäre eine Reaktion mit Wasser: $\text{H}_2\text{O} + \text{e}^- \rightarrow \frac{1}{2}\text{H}_2\uparrow + \text{OH}^-$. Diese Reaktion ist deshalb weitaus unwahrscheinlicher, weil die Transferreaktion (quantenmechanischer Tunneleffekt) sehr viel schneller ist [BER99, HAM98].

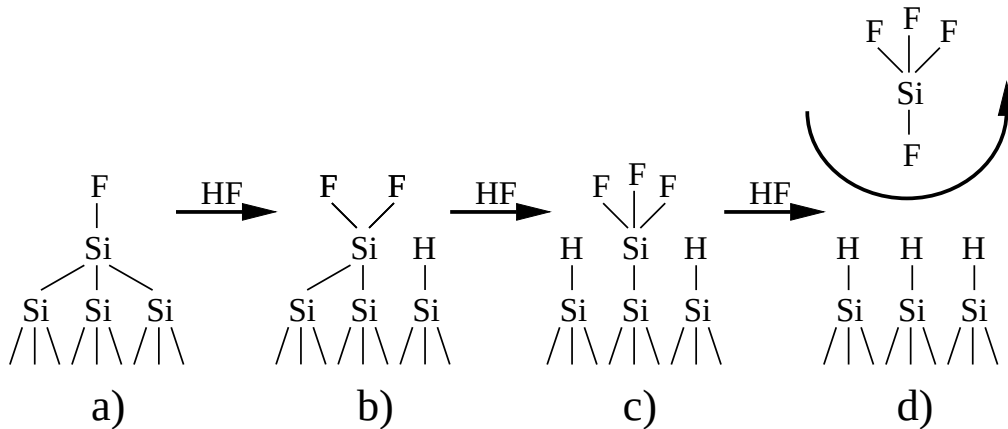


Abb. 1.12: Schematische Darstellung des Mechanismus der Wasserstoffterminierung [TRU90].

1.3.4.2 Die Kennlinie des Silizium-Flußsäure-Kontaktes

Abb. 1.13 zeigt vereinfacht den Aufbau, mit dem die Kennlinie des Silizium-Flußsäure-Kontaktes gemessen werden kann. Ein Stück Silizium taucht in einen Elektrolyten, bestehend aus verdünnter HF (1-5 Vol.%). Als Gegen- oder 'Counter'-Elektrode wird ein chemisch inertes Material verwendet, meist Platin. Weil das Potential des Elektrolyten nicht zugänglich ist (vgl. Abb. 1.8 auf Seite 17), verwendet man im allgemeinen eine Bezugs- oder Referenz-Elektrode, gegen die alle Potentiale gemessen werden, und statt der Spannungsquelle einen Potentiostaten. Neben einem definierten Bezugspunkt, der weitgehend temperaturunabhängig ist, dient eine Referenz-Elektrode auch dazu, ohmsche Verluste im Silizium und im Elektrolyten zu kompensieren. Das Potential der Referenz-Elektrode dient dem Potentiostaten, die Zellenspannung so zu regeln, daß der Einfluß der genannten ohmschen Widerstände eliminiert wird. Näheres in Abschnitt 3.2.3.

Da der Potentialabfall zwischen Referenz-Elektrode und Elektrolyt für jede Referenz-Elektrode verschieden ist, wird sich bei Verwendung einer anderen Referenz-Elektrode, die IV-Kennlinie auf der Spannungsachse etwas verschieben, ihre eigentliche Gestalt wird aber exakt die Gleiche bleiben.

Bei nicht zu schnellen Spannungsänderungen ($\frac{\Delta U}{\Delta t} < 100 \frac{\text{mV}}{\text{s}}$) erhält man Kennlinien, wie sie in Abb. 1.14 dargestellt sind. Sie zeigen, daß die Kennlinie des Silizium-Flußsäure-Kontaktes sowohl für n- als auch für p-Silizium recht kompliziert ist. Für nicht zu große Spannungen ($0 < U < 0.5 \text{ V}$) ist sie der Kennlinie des Schottky-Kontaktes sehr ähnlich. Dennoch gibt es vier wesentliche Unterschiede:

1. Beim Silizium-Flußsäure-Kontakt ist, wie auch bei jeder anderen elektrolytischen Zelle, Stromfluß zwangsweise mit einer elektrochemischen Reak-

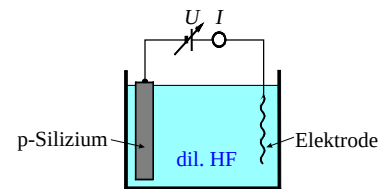


Abb. 1.13: Aufbau der Zelle zum Messen der IV-Kennlinie des Silizium-Flußsäure-Kontaktes

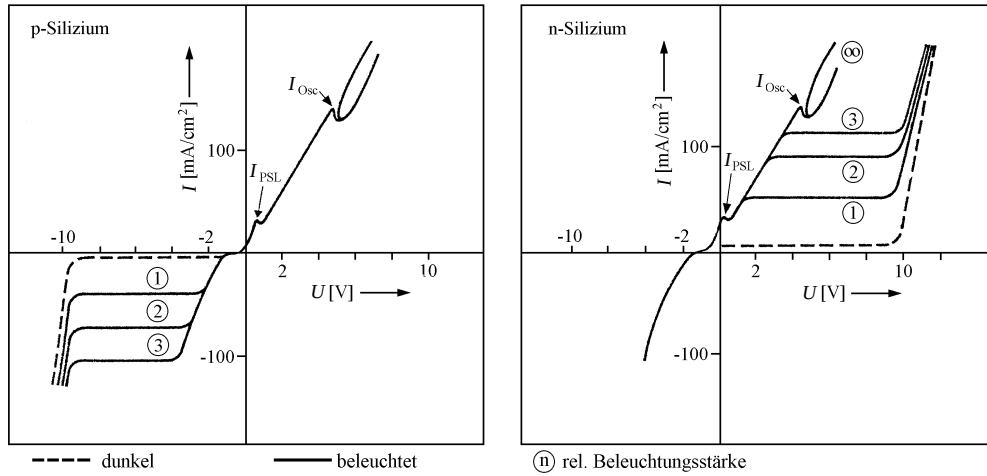


Abb. 1.14: Kennlinien des Silizium-Flußsäure-Kontaktes [FOE91].

tion verbunden: Für kathodische Ströme die Reduktion von Wasserstoff $2\text{H}_3\text{O}^+ + 2\text{e}^- \rightarrow 2\text{H}_2\text{O} + \text{H}_2 \uparrow$ und für anodische Ströme die Oxidation von Silizium $\text{Si} + 6\text{H}_2\text{O} + 4\text{h}^+ \rightarrow \text{SiO}_2 + 4\text{H}_3\text{O}^+$. In einem weiteren Schritt reagiert das Siliziumoxid mit Flußsäure: $\text{SiO}_2 + 6\text{HF} \rightarrow \text{H}_2\text{SiF}_6 + 2\text{H}_2\text{O}$. Fließen anodische Ströme, d. h. werden Löcher in den Elektrolyten injiziert, geht Silizium in Lösung [CHA79, NAK87, SOM85].

2. Die Leerlaufspannung U_{oc} (oc = 'open-circuit') ist zeitlich nicht stabil. Sie hängt empfindlich von der Beschaffenheit der Siliziumoberfläche und dem Elektrolyten (Zusammensetzung, Reinheit, etc.) ab. Kurz nach Eintauchen einer frischen Si-Elektrode ist deren Oberfläche noch nicht passiviert. Aufgrund der sehr großen Oberflächenzustandsdichte N_{SS} verhält sich die Si-Elektrode daher wie ein Metall: Nahezu das gesamte Potential fällt im Elektrolyten ab. Nach abgeschlossener H-Terminierung, wobei sich die N_{SS} um mehrere Größenordnungen verkleinert, kann sich im Silizium eine Raumladungszone ausbilden: Nun fällt fast das gesamte Potential im Silizium ab und nur noch ein kleiner Teil im Elektrolyten. Der Prozeß der H-Terminierung führt zu einer allmählichen positiven Aufladung der Si-Elektrode und letztlich zu einer Erhöhung von U_{oc} [BER99].
3. Für n-Silizium ist der fließende Photostrom *nicht* proportional zur Beleuchtungsintensität und ist bei kleinen Intensitäten größer als für p-Silizium. Ursache hierfür ist, daß für jedes lichtgenerierte Loch, das die Silizium-Elektrolyt-Grenzfläche erreicht, bis zu drei Elektronen ins Silizium injiziert werden. Dadurch ergibt sich ein bis zu viermal höherer Photostrom. Bei höheren Intensitäten bzw. bei höheren Photoströmen nimmt das Injektionsverhältnis von 4 auf 1 ab [FOE91].
4. Bei hohen kathodischen Spannungen (p-Silizium) bzw. hohen anodischen Span-

nungen (n-Silizium) kommt es, je nach Grad der Dotierung, zu einem Durchbruch. Der Silizium-Flußsäure-Kontakt übersteht dies schadlos, ein Schottky-Kontakt würde dabei zerstört werden.

Über den p-Silizium-Flußsäure-Kontakt fließt ohne Beleuchtung und bei defektfreier Oberfläche nur ein sehr kleiner kathodischer Strom $I < 0.2 \frac{\mu\text{A}}{\text{cm}^2}$. Bei Vorhandensein von Defekten, wie Korngrenzen, Versetzungen, metallischen Ausscheidungen an der Oberfläche, kann dieser Strom erheblich größer sein. In diesem Bereich der Kennlinie arbeitet der Elymat, ein Verfahren, mit dem die Minoritätslebensdauer von p-Silizium orts aufgelöst gemessen werden kann [LEH88, LIP97]. Ebenfalls in diesem Bereich der Kennlinie arbeitet der hier vorgestellte Leckstrom-Scanner.

Der anodische Bereich läßt sich in drei Bereiche unterteilen:

PSL¹⁹-Bereich: Im Bereich $0 < I \leq I_{\text{PSL}}$ bildet sich poröses Silizium. An Orten mit unvollständiger H-Terminierung, die zufällig über die Oberfläche verteilt sind, läuft die Auslösungsreaktion bevorzugt ab, und es entstehen Grübchen (Nukleationsprozeß). Diese Grübchen wachsen in die Tiefe und bilden unterhalb des Wendepunktes makroporöses Silizium, wohingegen oberhalb des Wendepunktes nanoporöses Silizium entsteht [CHR00]. Im Falle von nanoporösem Silizium bilden die einkristallinen Bereiche zwischen den Poren sogenannte 'quantum walls' mit einem Bandabstand von $E_G \approx 1.5 \text{ eV}$. Aufgrund des vergrößerten Bandabstandes mangelt es zwischen den Poren an Löchern die nötig sind, damit die Auflösungsreaktion ablaufen kann. Daher werden die Porenzwischenräume nicht aufgelöst [FOE91].

Elektropolierbereich: Für $I_{\text{PSL}} < I \leq I_{\text{Osc}}$ liegt der Elektropolierbereich vor. Die Auslösungsreaktion erfolgt nun so schnell, daß die Diffusion von HF-Molekülen zur Oberfläche geschwindigkeitsbestimmend wird. Über einer Erhöhung ist die HF-Konzentration höher, so daß hier die Auflösungsreaktion schneller erfolgt als anderswo und die Erhöhung deshalb mit der Zeit verschwindet. Des weiteren erfolgt verstärkt die Bildung von anodischem Oxid. Da anodisches Oxid amorph ist, erfolgt dessen Auflösung isotrop. Dieser diffusionslimitierte Mechanismus der Siliziumauflösung ist um so effektiver, je kleiner die Abmessungen der Unebenheiten sind. So verschwinden Unebenheiten im Nanometerbereich sehr viel schneller als im Mikrometerbereich [CHR00].

Oszillationsbereich: Gilt $I > I_{\text{Osc}}$, so treten am Kontakt Oszillationserscheinungen auf: Diese Oszillationen können gedämpft (kleine HF-Konzentrationen) oder ungedämpft (hohe HF-Konzentrationen) auftreten. Ihre Frequenz und Form hängt im wesentlichen von der HF-Konzentration und der Spannung ab.

Voraussetzung für das Auftreten von makroskopisch beobachtbaren Oszillationen ist eine geschlossene Oxidschicht auf dem Silizium, deren Dicke mit gleicher Frequenz aber phasenverschoben mitoszilliert. Jeglicher Stromfluß muß also über Nukleation durch diese Schicht passieren. Nach einem aktuellen Modell, läßt sich dieses Phänomen folgendermaßen verstehen:

Die Oxidschicht wird in einer Phase elektrochemisch gebildet und in einer anderen rein chemisch aufgelöst. Da konstant eine Spannung an der Elektrode anliegt, sich aber

¹⁹ PSL = **P**orous **S**ilicon **L**ayer

	p-Silizium		n-Silizium	
	dunkel	beleuchtet	dunkel	beleuchtet
A N O D I S C H	Beleuchtung hat keinen Einfluß Durchlaßrichtung Siliziumauflösung Besonderheiten: • PSL-Bildung • Defektätzung • Elektropolieren • Strom/Spannungs-Oszillationen		Sperrichtung Siliziumauflösung Besonderheiten: Anätzen von Defekten • Photostrom ist nicht proportional zur Intensität • I und U sind über die Intensität unabhängig einstellbar • Anätzen von Defekten in verschiedenen Modes • 'schwarzes Silizium'	
	Sperrichtung H_2 -Entwicklung, Si ist inert Besonderheiten: • sehr kleiner Leckstrom • H_2-Bläschen an Defekten • Photostrom ist proportional zur Intensität		Beleuchtung hat keinen Einfluß Durchlaßrichtung H_2 -Entwicklung, Si ist inert Besonderheiten: - keine -	
K A T H O D I S C H				

Tab. 1.1: Charakteristische Bereiche der IV-Kennlinien von p - und n -Silizium in wässriger Flußsäure [FOE91].

die Dicke der Oxidschicht ändert, ändert sich ebenfalls die Feldstärke in dem Oxid. Ist das Oxid dünn genug überschreitet die Feldstärke einen kritischen Wert, was zu einem (lokalen) Durchbruch des Oxides und der Bildung eines stromführenden Kanals führt. Mit diesem Strom ist eine elektrochemische Reaktion zur Bildung von anodischem Oxid verbunden, was die Schichtdicke lokal wieder erhöht, bis die dabei ständig abnehmende Feldstärke so klein ist, daß weiterer Stromfluß nicht mehr möglich ist und der Kanal versiegt. Einen solchen lokalen Prozeß kann man als einzelnen Oszillator verstehen. Die makroskopischen Oszillationen bedürfen der zusätzlichen Synchronisation dieser Mikrooszillatoren.

Die Größe der Oszillatoren (laterale Abmessungen sowie Dicke) ist von der Oszillationsspannung abhängig. Ab einer bestimmten Größe überlappen die Oszillatoren und interagieren miteinander. Es ist nun günstiger für sie, synchronisiert zu oszillieren, was eben zu den makroskopischen Oszillationen führt. Räumliche Interaktion resultiert also in zeitlicher Korrelation.

Zusammen mit einem desynchronisierenden Mechanismus (die Spannung um einen stromführenden Kanal ist reduziert) läßt sich das Modell mit Monte-Carlo-Simulationen nachbilden. Die simulierten Oszillationen zeigen sehr gute Übereinstimmung mit den gemessenen. Ein weiteres interessantes Ergebnis der Simulationen ist, daß die makroskopischen Oszillationen ein Perkulationsphänomen darstellen [HAS01].

In Tab. 1.1 sind die charakteristischen Bereiche der IV-Kennlinie von p- und n-Silizium aufgelistet.

1.4 Leckstrommechanismen

Grundsätzlich sind oberflächennahe null- oder mehrdimensionale Gitterdefekte, wie interstitielle oder substitutionelle Fehlerstellen, Versetzungen, Korngrenzen und/oder Ausscheidungen, Ursache eines erhöhten Sperrstromes [WAL95a]. Dennoch kann zwischen verschiedenen Mechanismen unterschieden werden:

Wird ein pn- oder ein Schottky-Kontakt in Sperrichtung betrieben, so fließt im allgemeinen nur ein sehr kleiner Sperr- oder Leckstrom. Beim pn-Übergang hängt nach *Shockley* der Sperrstrom neben den Dotierungen auch von den Diffusionslängen ab [SHO50]:

$$J_s = \frac{qD_p p_{n0}}{L_p} + \frac{qD_n n_{p0}}{L_n}. \quad (1.36)$$

p_{n0} und n_{p0} sind die Gleichgewichts-Minoritätsträgerdichten im n- bzw. im p-Gebiet, die von der Dotierung abhängig sind. Gleichung 1.36 gilt in der Form nur für Germanium. In Silizium ist der Sperrstrom größer und zudem abhängig von der Sperrspannung. Ursache hierfür ist eine Sperrstromerhöhung durch Elektron-Loch-Paar-Generation in der Raumladungszone, deren Weite spannungsabhängig ist.

Bei Schottky-Kontakten existieren je nach Grad der Dotierung drei verschiedene Leitungsmechanismen [ELS88, OLA96]:

1. die thermische Emission von Minoritäten über die Potentialbarriere,

2. die feldunterstützte Diffusion von Minoritäten über die Potentialbarriere und
3. Tunneln von Minoritäten durch die Potentialbarriere.

Für den ersten Fall beträgt der Sperrstrom (vgl. Gleichung 1.13):

$$J_s = A^* T^2 \exp\left(-\frac{q\phi_B}{kT}\right). \quad (1.37)$$

Findet Ladungstransport über feldunterstützte Diffusion statt, gilt für den fließenden Sperrstrom:

$$J_s = \frac{q^2 D_p N_V}{kT} \sqrt{\frac{2q(V - V_{bi})N_A}{\epsilon_s}} \exp\left(-\frac{q\phi_B}{kT}\right). \quad (1.38)$$

Dieser Sperrstrom weist eine Spannungsabhängigkeit auf, während derjenige, für den der Mechanismus der thermischen Emission dominiert, von der Spannung unabhängig ist [SCH38].

Tunneleffekte treten erst bei hohen Dotierkonzentrationen oberhalb von $N_A > 10^{18} \frac{1}{\text{cm}^3}$ und bei genügend tiefen Temperaturen auf. Wird der Tunneleffekt zum dominierenden Leitungsmechanismus, so verliert der Schottky-Kontakt seine sperrenden Eigenschaften und wird ohmsch, siehe Abschnitt 1.2.2.

An Orten mit Defekten kann der fließende Sperrstrom erheblich erhöht sein. Entweder findet Ladungstransport über oberflächennahe 'Trap'-Zustände N_t oder über Oberflächenzustände N_{ss} bei sonst intakter Raumladungszone statt [OBE95]. Mit steigender Defektdichte nimmt die Weite der Raumladungszone immer weiter ab, so daß, wie bereits oben erwähnt, durch Tunneleffekte dieser Bereich ohmsch wird. Diese lokalen ohmschen Kurzschlüsse der Raumladungszone werden als '*Shunts*' bezeichnet. Die Konsequenz für den Füllfaktor der IV-Kennlinie einer Solarzelle und damit für deren Wirkungsgrad, wenn ein 'Shunt' den pn-Kontakt durchdringt, zeigt Abb. 1.15.

Ein weiteres Beispiel findet sich bei DRAMs: Das Speichern eines Bits erfolgt durch Aufladung eines Kondensators über einen Transistor. Ein erhöhter Leckstrom zwischen den Elektroden des Kondensators macht kürzere 'Refresh'-Zyklen erforderlich, um Datenverluste zu vermeiden. Die Folge können erhebliche Performance-Einbußen oder gar die Unbrauchbarkeit des DRAMs sein [HOR94, STE94, SZE85].

Ein Leckstrommechanismus ist die Injektion von Minoritäten. An der Peripherie von Kontakten aktiver Bauelemente können sehr hohe elektrische Feldstärken auftreten, die eine Injektion von Minoritäten zur Folge haben können. Die Konsequenzen wären ein erhöhter Leckstrom und eine Verschlechterung der Bauteileigenschaften [OLA96].

Schließlich können an der Halbleiteroberfläche adsorbierte Ionen auf Grund des Schottky-Effektes bewirken, daß die Höhe der Potentialbarriere verringert und damit der Sperrstrom erhöht wird [SZE81], siehe Abschnitt 1.2.3.

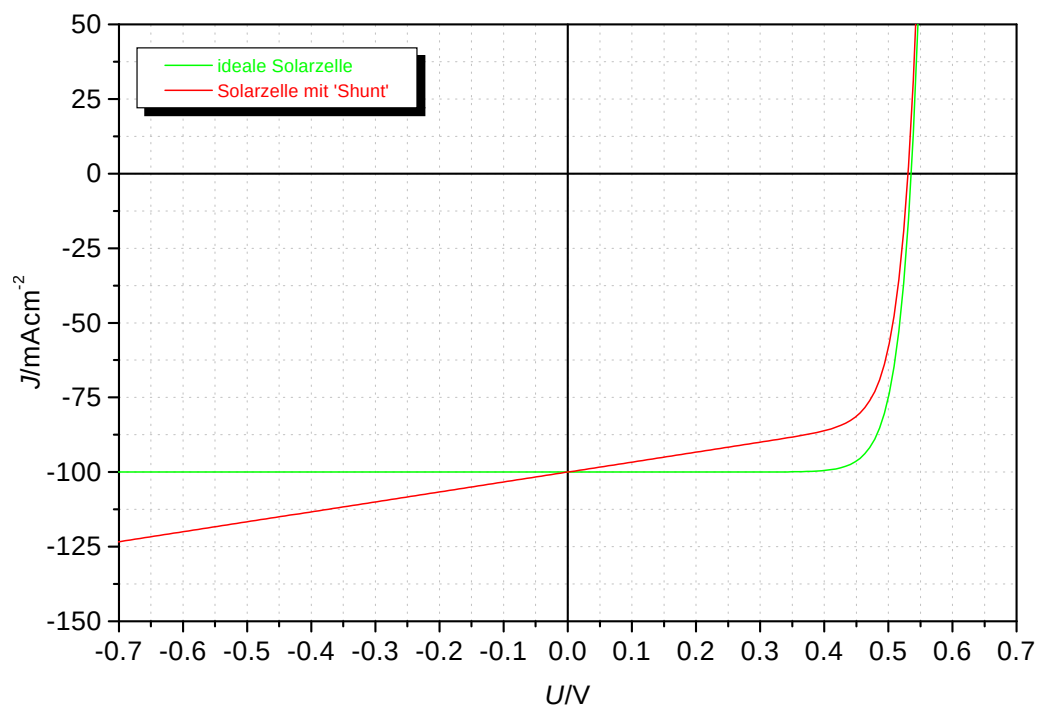


Abb. 1.15: Kennlinie einer beleuchteten Solarzelle ohne und mit lokal kurzgeschlossener Raumladungszone ('Shunt').

Kapitel 2

Elymat

Mit zunehmender Chipfläche und Integrationsdichte werden immer größere Anforderungen bezüglich Perfektion und Reinheit an das Silizium gestellt. Um akzeptable Ausbeuten bei der Chipherstellung erwarten zu können, ist es notwendig, im Verlauf der Prozessierung beispielsweise die Eisenkonzentration unter $10^{10} \frac{1}{\text{cm}^3}$ zu halten [ABA95, KEE98, SCH98]. Für diesen Zweck wurden mehrere Analyse-Tools entwickelt, die in der Kontaminationskontrolle Anwendung finden. Die drei wesentlichen sind: SPV [HAG98, HEN95, LIU96, TAR95], μ -PCD [POL98, SCH96b, TAR95] und der Elymat [EIC97, FAL94, KUS97, LEH88, OBE98, OST95, WAL95b, WIT98]. Diese Verfahren erlauben orts aufgelöste Messungen der Minoritätslebensdauer τ_n , wobei sich der Elymat gegenüber den beiden anderen Verfahren durch seine höhere Meßgeschwindigkeit auszeichnet [TAR95]. Mit Hilfe von Elymat-Messungen konnte nachgewiesen werden, daß 'Waferholder' und der Prozeß der Ionenimplantation starke Kontaminationsquellen sein können bzw. sind [BER91, POL95a, POL95b].

2.1 Aufbau und Funktionsweise

Abb. 2.1 zeigt schematisch den Aufbau des Elymaten: Eine mono- oder multikristalline p-Siliziumwafer¹ wird in die elektrolytische Doppelzelle eingespannt. Nach dem Füllen der Zellen mit verdünnter HF (ca. 1-2 Vol.%) und dem Beschalten der vorderen und/oder der hinteren Zelle mit einer kathodischen Spannung (je nach verwendetem Meßmode), wird die Wafervorderseite lokal mit einem Laserstrahl beleuchtet. Die daraus resultierenden Beleuchtungsströme werden entweder vorn-FPC-Mode (**F**rontside **P**hoto **C**urrent), hinten-BPC-Mode (**B**ackside **P**hoto **C**urrent) oder beidseitig-BSC-Mode (**B**oth **S**ide **P**hoto **C**urrent) gemessen². Da beim Messen des Beleuchtungsstroms der Dunkelstrom mitgemessen wird, d. h. der Strom der bei ausgeschaltetem Laser über

¹ Es können mit den Elymaten nur p-Siliziumwafer gemessen werden. Im Falle von n-Silizium kommt es bei der Transferreaktion an der in Sperrichtung betriebenen Si-HF-Grenzfläche zur Injektion von Elektronen ins Silizium. Damit verbunden ist die Auflösung von Silizium, vgl. Abschnitt 1.3.4.

² Für eine FPC-Messung wird die hintere Zelle, für eine BPC-Messung die vordere Zelle 'open-loop' geschaltet.

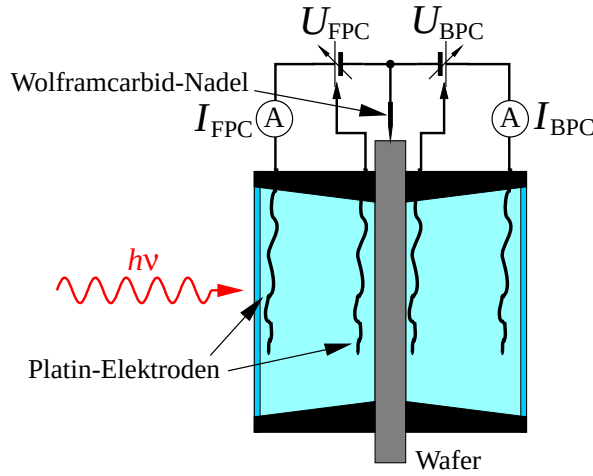


Abb. 2.1: Schematischer Aufbau des Elymaten.

die Grenzfläche fließt, muß dieser ebenfalls gemessen und anschließend vom Beleuchtungsstrom subtrahiert werden. Prinzipiell muß der Dunkelstrom nur einmal gemessen werden, da aber der Dunkelstrom zeitlich stark fluktuieren kann, besonders bei multikristallinem Silizium, wird an jeder Stelle (x, y) neben dem Beleuchtungsstrom auch der Dunkelstrom erfaßt.

Die Kontaktierung des Wafers erfolgt über Wolframcarbid-Nadeln. Wolframcarbid wird deshalb verwendet, weil dieses Material sehr hart ist und daher beim Aufsetzen der Nadeln die Oxidschicht des Siliziums durchstoßen wird, was einen guten ohmschen Kontakt ergibt. Andererseits besteht bei Verwendung von Wolframcarbid-Nadeln nicht die Gefahr einer Kontamination wie beispielsweise mit Eisen.

Die in Abb. 2.1 dicht vor dem Wafer platzierten Platindrähte sind Referenz-Elektroden. Sie dienen dazu, ohmsche Verluste im Meßkreis zu kompensieren, siehe Abb. 2.2 a). Allerdings kann mit dieser Meßanordnung nur der Widerstand des Kontaktes zwischen Silizium und Wolframcarbid R_C und der Elektrolytwiderstand R_E kompensiert werden. Hingegen nicht kompensiert werden kann der Elektrolytwiderstand zwischen Silizium und Referenz-Elektrode R'_E und der Bulkwiderstand des Silizium R_B . Die Leitfähigkeit des Elektrolyten ist im allgemeinen so groß, daß der Widerstand R'_E vernachlässigt werden kann. Nach Messen einer IV-Kennlinie wird der Einfluß des Bulkwiderstandes dadurch sichtbar, daß der Sättigungsbereich der IV-Kennlinie eine negative Steigung hat. Durch 'Raten' des Bulkwiderstandes R_B und Einrechnen in die IV-Kennlinie kann erreicht werden, daß die Steigung der IV-Kennlinie im Sättigungsbereich annähernd verschwindet. Auf diese Weise kann der Einfluß des Bulkwiderstandes nachträglich eliminiert werden. Der Einsatz von Referenz-Elektroden erlaubt also, die über der Silizium-Elektrolyt-Grenzfläche (in Abb. 2.2 a) symbolisiert durch einen schwarzen Kasten) abfallende Spannung zeitlich konstant zu halten und damit unabhängig vom Zellenstrom zu machen. Abb. 2.2 b) zeigt den Einfluß, den die Kompensation der Widerstände R_C und R_E auf den Verlauf der IV-Kennlinie hat (ohne

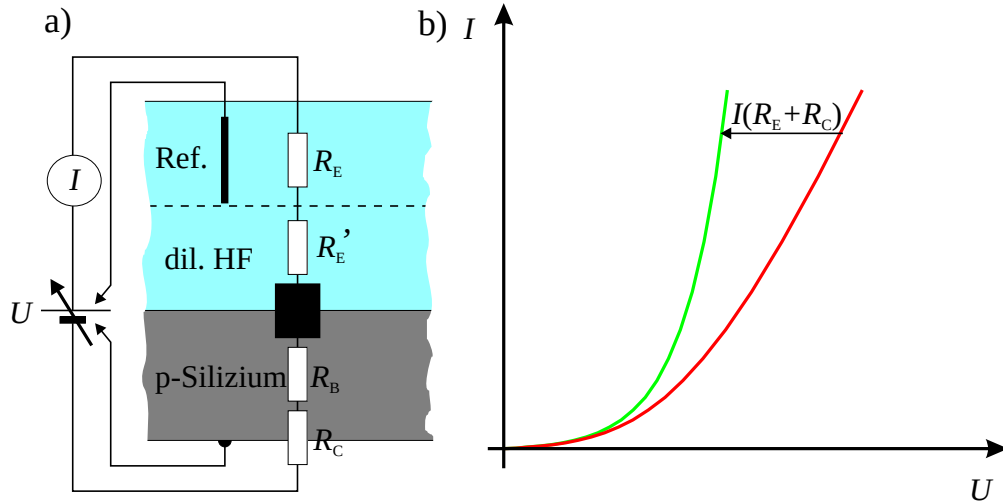


Abb. 2.2: a) Vereinfachtes Ersatzschaltbild der Elymat-Zelle. Der schwarze Kasten steht für alle Eigenschaften, die die Silizium-Elektrolyt-Grenzfläche betreffen. b) Einfluß des Elektrolytwiderstandes R_E und des Kontaktwiderstandes R_C auf die Gestalt der IV-Kennlinie der Elymat-Zelle (schematisch).

Beleuchtung).

Damit die durch das Laserlicht generierten Minoritäten über den Silizium-Elektrolyt-Kontakt ausgekoppelt werden können, muß eine Sperrspannung ausreichender Höhe an die Zelle angelegt werden. Um diese Spannung zu bestimmen, wird an mehreren Stellen des Wafers eine IV-Kennlinie des Kontaktes gemessen, siehe Abb. 2.3. Wird eine Zellenspannung gewählt, die im steigenden Bereich der Kennlinie liegt, werden nicht alle Minoritäten ausgekoppelt. In diesem Bereich wird der Elymat im sogenannten RPC-Mode (**R**estricted **P**hoto **C**urrent) betrieben. Zwar ist die Berechnung der Minoritätslebensdauer τ_n bzw. der Diffusionslänge L auch in diesem Teil der Kennlinie möglich, jedoch ist die zugehörige Theorie recht kompliziert, da nicht nur Bulk- und Oberflächeneigenschaften des Siliziums Einfluß auf die Höhe des fließenden Stromes haben, sondern auch die Elektrochemie des Elektrolyten. Die sich ergebenden nichtlinearen Differentialgleichungen können nur numerisch gelöst werden und erfordern spezielle numerische Methoden [LIP97].

Als Arbeitspunkt wird eine Zellenspannung gewählt, die im Sättigungsbereich der IV-Kennlinie liegt. Hier hängt der Strom nur noch von der Diffusionslänge L sowie der Oberflächenrekombinationsgeschwindigkeit S_b (FPC-Mode, b=back) bzw. S_f (BPC-Mode, f=front) ab. Allerdings sollte die Zellenspannung auch nicht zu hoch gewählt werden, weil, bedingt durch den mit der Spannung steigenden Dunkelstrom, sich das Signal-zu-Rausch-Verhältnis besonders bei multikristallinem Silizium, mit steigender Spannung immer weiter verschlechtert. Als Arbeitspunkt wählt man eine Spannung ca. 0.5 V rechts vom Beginn des Plateaus.

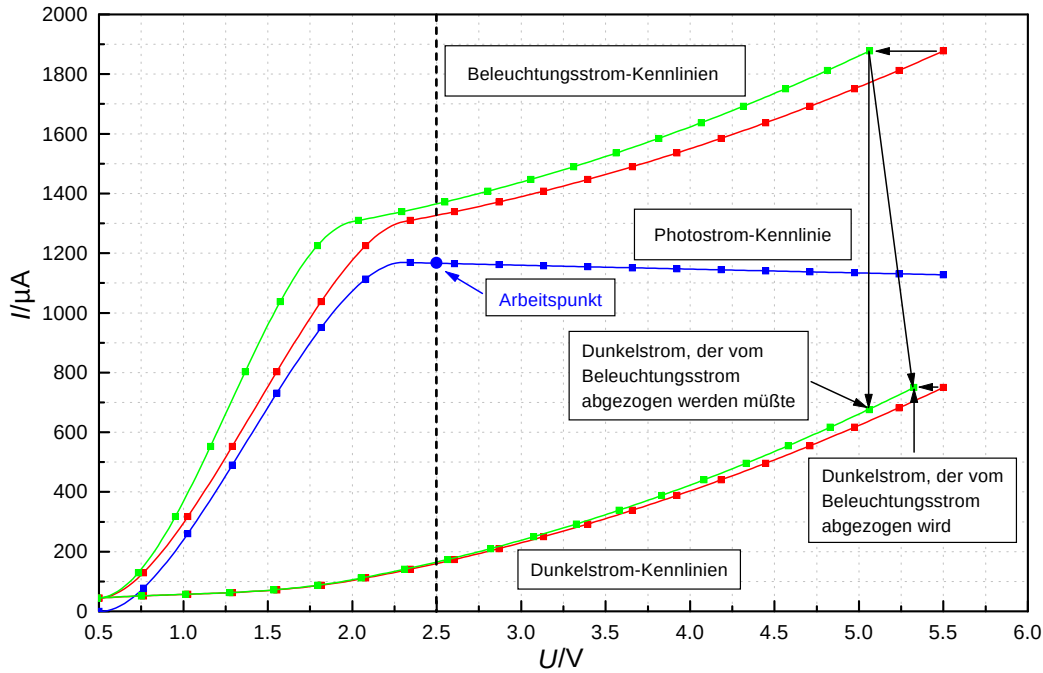


Abb. 2.3: IV-Kennlinien der Elymat-Zelle. Erläuterungen siehe Text.

Zufällig könnte die Messung der BPC-Kennlinie an einer schlechten Stelle erfolgt sein. Wird als Arbeitspunkt der Beginn des Plateaus gewählt, reicht während der Messung die Zellenspannung nicht aus, um an den Orten mit größerer hoher Diffusionslänge alle Minoritäten auszukoppeln (Messung im Bereich des RPC). Aufgrund dessen, daß die gemessenen Ströme zu klein sind, ergibt sich eine zu kleine Diffusionslänge. Aus diesem Grund wählt man nicht den Beginn des Plateaus als Arbeitspunkt, vgl. Abb. 2.3.

Die Referenz-Elektroden haben beim Messen von monokristallinem Silizium nur eine untergeordnete Bedeutung, da die Dunkelströme dieser Wafer im allgemeinen sehr klein sind. Zwingend notwendig sind die Referenz-Elektroden für multikristalline Wafer, weniger weil die Dunkelströme viel größer sind, sondern weil die Dunkelströme dieser Wafer in der Regel stark spannungsabhängig sind. Würde auf Referenz-Elektroden verzichtet werden, so erfolgt das Messen von Beleuchtungs- und Dunkelstrom bei unterschiedlichen Spannungen, da je nach Stärke des fließenden Stromes der Spannungsabfall über dem Widerstand des Kontaktes Silizium-Wolframcarbid, der in der Größenordnung $10^3 \Omega$ liegt, unterschiedlich ist. Bei konstanter Zellenspannung ändert sich damit auch in Abhängigkeit des Stromes der Spannungsabfall über dem Silizium-Elektrolyt-Kontakt, vgl. hierzu auch Abb. 2.2. Bei der Berechnung des Photostromes wird also vom Beleuchtungsstrom der falsche Dunkelstrom, nämlich ein zu großer, abgezogen. Abb. 2.3 verdeutlicht diesen Zusammenhang. Die Länge des waagerechten Pfeils entspricht der Spannung, die über dem Kontaktwiderstand abfällt. Bei Verwendung von

Referenz-Elektroden wird dieser Fehler vermieden. Die Spannung an den Referenz-Elektroden dient als 'Feedback' für die Spannungsquellen U_{FPC} und U_{BPC} , die die Zellenspannung so regeln, daß die über dem Silizium-Elektrolyt-Kontakt abfallende Spannung für die Beleuchtungs- und die Dunkelmessung stets dieselbe ist.

2.2 Theorie der Elymat-Modi

Das Modell zur Berechnung des Photostromes von *Lehmann* und *Föll* enthält die folgenden Annahmen [LEH88, WAL95b]:

1. Die Diffusion der Minoritäten durch den Halbleiter erfolgt eindimensional.
2. Die Minoritätslebensdauer τ_n am Ort (x, y) ist unabhängig von der Minoritätsträgerkonzentration $n_p(z)$ (keine 'Injection-Level'-Abhängigkeit).
3. Wechselwirkungen zwischen injizierten Minoritäten und Majoritäten werden vernachlässigt (keine ambipolare Diffusion).
4. Die Zellenspannung fällt nur über der Raumladungszone ab. Das Innere des Halbleiters als auch des Elektrolyts sind feldfrei.

Mit z als Richtung senkrecht zur Oberfläche des p-Siliziumwafers gilt damit für die Diffusionsgleichung im stationären Fall:

$$D_n \frac{\partial^2 n_p(z)}{\partial z^2} + G(z) - U(z) = \frac{dn_p(z)}{dt} = 0. \quad (2.1)$$

Hierin bezeichnen D_n die Diffusionskonstante der Minoritäten, $G(z)$ die Generations- und $U(z)$ die Rekombinationsrate. Für die Generationsrate $G(z)$ gilt nach *Walz*:

$$G(z) = (1 - R)\Phi_0 \alpha \exp(-\alpha z), \quad (2.2)$$

wobei α die Absorptionskonstante, Φ_0 den einfallenden Photonenfluß und R die Reflektivität der beleuchteten Waferoberfläche bedeutet [PLE90, WAL95b].

Zwischen der Absorptionskonstante α und der Wellenlänge λ des Lasers besteht der folgende funktionale Zusammenhang [NAR85]:

$$\alpha [\mu\text{m}^{-1}] = \left(\frac{847.32}{\lambda [\text{nm}]} - 0.7642 \right)^2. \quad (2.3)$$

Ist die Energie der Photonen kleiner als die Bandlücke des Halbleiters, was im Fall von Silizium für Wellenlängen größer als 1100 nm der Fall ist, verliert die Gleichung ihre Gültigkeit.

Für die Rekombinationsrate $U(z)$ gilt nach *Schockley-Read-Hall*:

$$U(z) = \frac{\sigma_p \sigma_n v_{\text{th}} [n_p(z) p_p(z) - n_i^2] N_T}{\sigma_n \left[n_p(z) + n_i \exp\left(\frac{E_T - E_i}{kT}\right) \right] + \sigma_p \left[p_p(z) + n_i \exp\left(-\frac{E_T - E_i}{kT}\right) \right]}, \quad (2.4)$$

hierin bezeichnen σ_p und σ_n die optischen Einfangquerschnitte für Löcher bzw. Elektronen, $v_{th} = \sqrt{\frac{3kT}{m^*}}$ die thermische Geschwindigkeit der Ladungsträger, N_T die Konzentration der Traps, E_T das Niveau der Traps, E_i das Fermi-Niveau des intrinsischen Halbleiters, $n_p(z) = n_{p0} + \Delta n_p(z)$ und $p_p(z) = p_{p0} + \Delta p_p(z)$ die Elektronen- bzw. die Löcher- und n_i die intrinsische Ladungsträgerkonzentration [ADA98, GER99, SHO52].

Mit $U(z)$ nach Gleichung 2.4 ist die Diffusionsgleichung analytisch nicht lösbar. Für nicht zu hohe Beleuchtungsintensitäten ist die Konzentration der Überschußladungsträger viel kleiner als die Gleichgewichtskonzentration: $\Delta n_p(z) \ll n_{p0}$ ('low-injection'). Für diesen Fall kann die Rekombination durch folgenden Ausdruck beschrieben werden [IST97, MUR95, SZE85]:

$$U(z) = \frac{n_p(z) - n_{p0}}{\tau_n}, \quad \text{mit} \quad \tau_n = \frac{1}{\sigma_n v_{th} N_T}. \quad (2.5)$$

Mit den Ausdrücken für $G(z)$ und $U(z)$ wird die Diffusionsgleichung zu:

$$D_n \frac{\partial^2 n_p(z)}{\partial z^2} + (1 - R) \Phi_0 \alpha \exp(-\alpha z) - \frac{n_p(z) - n_{p0}}{\tau_n} = 0, \quad (2.6)$$

die durch folgende Gleichung für $n_p(z)$ gelöst wird [OBE95, WAL95b]:

$$n_p(z) - n_{p0} = A \exp\left(-\frac{z}{L}\right) + B \exp\left(+\frac{z}{L}\right) + \frac{(1 - R) \Phi_0 \alpha \tau_n}{1 - (\alpha L)^2} \exp(-\alpha z), \quad (2.7)$$

mit $L = \sqrt{D_n \tau_n}$. Die Konstanten A und B sind bestimmt durch die Wahl der Randbedingungen.

2.2.1 FPC-Mode

Der FPC-Strom J_{FPC} setzt sich aus zwei Anteilen zusammen: $J_{FPC} = J_{FPC}^R + J_{FPC}^B$. Die Minoritäten, die jenseits der Raumladungszone generiert werden, tragen den Strom J_{FPC}^B ($B = \text{Bulk}$), diejenigen die in der Raumladungszone generiert werden den Strom J_{FPC}^R ($R = \text{Raumladungszone}$).

Wird an die vordere Zelle die Sperrspannung U_{FPC} angelegt und die hintere Zelle 'open-loop' geschaltet, sind die Randbedingungen für $n_p(z)$ im Gebiet $w \leq z \leq d$ ($w = \text{Weite der Raumladungszone}$):

$$\begin{aligned} -D_n \left. \frac{\partial n_p(z)}{\partial z} \right|_{z=d} &= S_b (n_p(d) - n_{p0}) & : \quad \text{Rekombination an der Rückseite.} \\ n_p(w) - n_{p0} &= 0 & : \quad \text{Keine Akkumulation von Minoritäten} \\ & & \quad \text{an der hinteren Grenze der Raumla-} \\ & & \quad \text{dungszone.} \end{aligned}$$

Mit diesen Randbedingungen ergibt sich für J_{FPC}^B :

$$J_{\text{FPC}}^{\text{B}} = +qD_n \left. \frac{\partial n_p(z)}{\partial z} \right|_{z=w} = \frac{(1-R)q\Phi_0(\alpha L)^2}{(\alpha L)^2 - 1}.$$

$$\frac{\exp(-\alpha d) \left[\frac{S_b}{D_n \alpha} - 1 \right] + \exp(-\alpha w) \left\{ \left[\frac{S_b}{D_n} L - \frac{1}{\alpha L} \right] \sinh\left(\frac{d-w}{L}\right) + \left[1 - \frac{S_b}{D_n \alpha} \right] \cosh\left(\frac{d-w}{L}\right) \right\}}{\cosh\left(\frac{d-w}{L}\right) + \frac{S_b}{D_n} L \sinh\left(\frac{d-w}{L}\right)}.$$
(2.8)

Definiert man

$$F(d) = \frac{1}{\alpha} \exp(-\alpha d) + \exp(-\alpha w) \left[L \sinh\left(\frac{d-w}{L}\right) - \frac{1}{\alpha} \cosh\left(\frac{d-w}{L}\right) \right] \quad (2.9)$$

läßt sich die Gleichung für $J_{\text{FPC}}^{\text{B}}$ in einer sehr kompakten Form angeben:

$$J_{\text{FPC}}^{\text{B}} = \frac{(1-R)q\Phi_0(\alpha L)^2}{(\alpha L)^2 - 1} \cdot \frac{\frac{S_b}{D_n} F(d) + \frac{\partial}{\partial d} F(d)}{\cosh\left(\frac{d-w}{L}\right) + \frac{S_b}{D_n} L \sinh\left(\frac{d-w}{L}\right)}. \quad (2.10)$$

Die Bestimmung der Diffusionslänge aus Gleichung 2.8 ist nur möglich mit Hilfe numerischen Methoden. Es ist daher sinnvoll, die Definitionslücke an der Stelle $L = \frac{1}{\alpha}$ zu beseitigen. Das Ersetzen des Zählers im rechten Bruch durch die Taylor-Näherung bis zum linearen Glied ergibt:

$$J_{\text{FPC}} = \frac{(1-R)q\Phi_0(\alpha L)^2}{\alpha L + 1} \cdot \frac{\exp(-\alpha w) \left(1 + \frac{S_b}{D_n \alpha} \right) \sinh[\alpha(d-w)] + \alpha(d-w) \left(1 - \frac{S_b}{D_n \alpha} \right) \exp(-\alpha d)}{\cosh\left(\frac{d-w}{L}\right) + \frac{S_b}{D_n} L \sinh\left(\frac{d-w}{L}\right)}. \quad (2.11)$$

Der Anteil aus der Raumladungszone $J_{\text{FPC}}^{\text{R}}$ ergibt sich einfach durch Integration der Generationsrate $G(z)$ über das Gebiet der Raumladungszone, da in dieser, aufgrund des Fehlens von Majoritäten, die Rekombinationsrate $U(z) = 0$ ist:

$$J_{\text{FPC}}^{\text{R}} = q \int_0^w G(z) dz = q \int_0^w (1-R)\Phi_0 \alpha \exp(-\alpha z) dz = (1-R)q\Phi_0 [1 - \exp(-\alpha w)]. \quad (2.12)$$

Da die Ausdehnung der Raumladungszone für übliche Dotierungen im Bereich $w = 1 \mu\text{m}$ liegt, ist sie für gewöhnlich gegenüber der Waferdicke und der Eindringtiefe des Lasers vernachlässigbar: $w \ll d$ und $\alpha w \ll 1$. In diesem Fall vereinfacht sich Gleichung 2.8 zu:

$$J_{\text{FPC}} = \frac{(1-R)q\Phi_0(\alpha L)^2}{(\alpha L)^2 - 1} \cdot \frac{\exp(-\alpha d) \left[\frac{S_b}{D_n \alpha} - 1 \right] + \left[\frac{S_b}{D_n} L - \frac{1}{\alpha L} \right] \sinh\left(\frac{d}{L}\right) + \left[1 - \frac{S_b}{D_n \alpha} \right] \cosh\left(\frac{d}{L}\right)}{\cosh\left(\frac{d}{L}\right) + \frac{S_b}{D_n} L \sinh\left(\frac{d}{L}\right)}. \quad (2.13)$$

Ist die rückseitige Oberflächenrekombinationsgeschwindigkeit S_b klein, d. h. gilt $\frac{S_b}{D_n}L \ll 1$ und ist die Eindringtiefe des Lasers α^{-1} nicht zu groß, so daß $\frac{S_b}{D_n\alpha} \ll 1$, läßt sich die Gleichung für J_{FPC} weiter vereinfachen:

$$J_{\text{FPC}} = \frac{(1-R)q\Phi_0(\alpha L)^2}{(\alpha L)^2 - 1} \cdot \frac{\cosh\left(\frac{d}{L}\right) - \frac{1}{\alpha L} \sinh\left(\frac{d}{L}\right) - \exp(-\alpha d)}{\cosh\left(\frac{d}{L}\right)}. \quad (2.14)$$

Wird ein Laser verwendet, dessen Wellenlänge im sichtbaren Bereich oder im nahen Infrarot liegt, gilt üblicherweise $\alpha d \gg 1$:

$$J_{\text{FPC}} = \frac{(1-R)q\Phi_0(\alpha L)^2}{(\alpha L)^2 - 1} \cdot \frac{\cosh\left(\frac{d}{L}\right) - \frac{1}{\alpha L} \sinh\left(\frac{d}{L}\right)}{\cosh\left(\frac{d}{L}\right)}. \quad (2.15)$$

Ist $\alpha L \gg 1$, läßt sich die letzte Gleichung weiter vereinfachen:

$$J_{\text{FPC}} = (1-R)q\Phi_0. \quad (2.16)$$

2.2.2 BPC-Mode

Beim BPC-Mode wird die Sperrspannung an die hintere Zelle gelegt und die vordere 'open-loop' geschaltet. Damit lauten die Randbedingungen für $n_p(z)$:

$$\begin{aligned} +D_n \left. \frac{\partial n_p(z)}{\partial z} \right|_{z=0} &= S_f (n_p(0) - n_{p0}) && : \text{ Rekombination an der Vorderseite.} \\ n_p(d-w) - n_{p0} &= 0 && : \text{ Keine Akkumulation von Minoritäten} \\ &&& \text{ an der vorderen Grenze der Raumladungszone.} \end{aligned}$$

Mit diesen Randbedingungen für $n_p(z)$ im Gebiet $0 \leq z \leq d-w$ ergibt sich für J_{BPC} ³:

$$\begin{aligned} J_{\text{BPC}} &= -qD_n \left. \frac{\partial n_p(z)}{\partial z} \right|_{z=d-w} = \frac{(1-R)q\Phi_0(\alpha L)^2}{(\alpha L)^2 - 1} \cdot \\ &\frac{1 + \frac{S_f}{D_n\alpha} - \exp[-\alpha(d-w)] \left\{ \left[\frac{1}{\alpha L} + \frac{S_f}{D_n}L \right] \sinh\left(\frac{d-w}{L}\right) + \left[1 + \frac{S_f}{D_n\alpha} \right] \cosh\left(\frac{d-w}{L}\right) \right\}}{\cosh\left(\frac{d-w}{L}\right) + \frac{S_f}{D_n}L \sinh\left(\frac{d-w}{L}\right)}. \quad (2.17) \end{aligned}$$

Ersetzen des Zählers des unteren Bruches durch die Taylor-Näherung bis zum linearen Glied liefert hier:

$$J_{\text{BPC}} = \frac{(1-R)q\Phi_0(\alpha L)^2}{\alpha L + 1} \cdot \frac{\exp[-\alpha(d-w)] \left(1 - \frac{S_f}{D_n\alpha} \right) \sinh[\alpha(d-w)] + \alpha d \left(1 + \frac{S_f}{D_n\alpha} \right)}{\cosh\left(\frac{d-w}{L}\right) + \frac{S_f}{D_n}L \sinh\left(\frac{d-w}{L}\right)}. \quad (2.18)$$

³ Strenggenommen muß zu J_{BPC} noch der Anteil aus der rückseitigen Raumladungszone addiert werden. Da aber für Laser mit nicht zu großen Eindringtiefen die Intensität an der Rückseite bereits sehr klein sein wird, kann der Anteil aus der Raumladungszone in sehr guter Näherung vernachlässigt werden: $J_{\text{BPC}}^B = q \int_{d-w}^d G(z) dz = (1-R)q\Phi_0 \exp(-\alpha d) [\exp(\alpha w) - 1] \approx 0$, wenn $\alpha d \gg 1$ oder $\alpha w \ll 1$.

Gewöhnlich gilt $w \ll d$, und Gleichung 2.17 vereinfacht sich zu:

$$J_{\text{BPC}} = \frac{(1-R)q\Phi_0(\alpha L)^2}{(\alpha L)^2 - 1} \cdot \frac{1 + \frac{S_f}{D_n \alpha} - \exp(-\alpha d) \left\{ \left[\frac{1}{\alpha L} + \frac{S_f}{D_n} L \right] \sinh\left(\frac{d}{L}\right) + \left[1 + \frac{S_f}{D_n \alpha} \right] \cosh\left(\frac{d}{L}\right) \right\}}{\cosh\left(\frac{d}{L}\right) + \frac{S_f}{D_n} L \sinh\left(\frac{d}{L}\right)}. \quad (2.19)$$

Wird ein Laser mit kleiner Eindringtiefe verwendet, d. h. gilt $\alpha d \gg 1$, läßt sich die letzte Gleichung für J_{BPC} erheblich vereinfachen:

$$J_{\text{BPC}} = \frac{(1-R)q\Phi_0(\alpha L)^2}{(\alpha L)^2 - 1} \cdot \frac{1 + \frac{S_f}{D_n \alpha}}{\cosh\left(\frac{d}{L}\right) + \frac{S_f}{D_n} L \sinh\left(\frac{d}{L}\right)}. \quad (2.20)$$

Ist $\alpha L \gg 1$ wird Gleichung 2.20 zu:

$$J_{\text{BPC}} = (1-R)q\Phi_0 \cdot \frac{1 + \frac{S_f}{D_n \alpha}}{\cosh\left(\frac{d}{L}\right) + \frac{S_f}{D_n} L \sinh\left(\frac{d}{L}\right)}. \quad (2.21)$$

Ist aufgrund guter Passivierung die vordere Oberflächenrekombinationsgeschwindigkeit klein, d. h. sind die Bedingungen $\frac{S_f}{D_n \alpha} \ll 1$ und $\frac{S_f}{D_n} L \ll 1$ erfüllt, ergibt sich die Gleichung für J_{BPC} in ihrer einfachsten Form:

$$J_{\text{BPC}} = (1-R)q\Phi_0 \cdot \frac{1}{\cosh\left(\frac{d}{L}\right)}. \quad (2.22)$$

Ist die Wellenlänge des Lasers im sichtbaren Bereich oder im nahen Infrarot, so ergibt eine FPC-Messung den Wert für $(1-R)q\Phi_0$, vgl. Gleichung 2.16. Diese FPC-Messung liefert also einen Normierungsfaktor für die BPC-Ströme:

$$\frac{J_{\text{BPC}}}{J_{\text{FPC}}} = \frac{1}{\cosh\left(\frac{d}{L}\right)}. \quad (2.23)$$

Der große Nutzen dieser FPC-Messung rührt daher, daß beim Normieren der BPC-Ströme die unsicheren Faktoren $(1-R)$ und Φ_0 eliminiert werden, die nur stark fehlerbehaftet bestimmt werden können.

Werden Laser mit größeren Eindringtiefen verwendet, so ist die Normierung der BPC-Ströme nicht ohne weiteres durchführbar, da durch Rekombinationsverluste im Silizium der gemessene FPC-Strom zu klein ist ($\alpha d \gg 1$ ist nicht erfüllt). Um den wahren Wert für $(1-R)q\Phi_0$ zu bestimmen, müßte Gleichung 2.14 herangezogen werden, wofür allerdings die Diffusionslänge bereits bekannt sein müßte. Eine Möglichkeit diese Schwierigkeit zu beseitigen, besteht darin, mit Hilfe einer Messung im BSC-Mode die Diffusionslänge L zu bestimmen und anschließend mit Gleichung 2.14 den Wert für $(1-R)q\Phi_0$ zu berechnen.

2.2.3 BSC-Mode

Beim BSC-Mode werden beide Zellen beschaltet. Werden die Zellenspannungen so hoch gewählt, daß im Plateau der IV-Kennlinie gemessen wird, sind die Minoritätsträgerkonzentrationen an Vorder- und Rückseite gleich null, da *alle* Minoritäten, die in das Gebiet der Raumladungszone diffundieren, ausgekoppelt werden. Die Gleichungen für den BSC-Mode ergeben sich also dadurch, daß man in Gleichung 2.8 S_b und in Gleichung 2.17 S_f gegen unendlich gehen läßt:

$$\lim_{S_b \rightarrow \infty} J_{\text{FPC}} = J_{\text{FBSC}} = \frac{(1-R)q\Phi_0(\alpha L)^2}{(\alpha L)^2 - 1} \cdot \frac{\frac{1}{\alpha} \exp(-\alpha d) + \exp(-\alpha w) \left\{ L \sinh\left(\frac{d-w}{L}\right) - \frac{1}{\alpha} \cosh\left(\frac{d-w}{L}\right) \right\}}{L \sinh\left(\frac{d-w}{L}\right)} \quad (2.24)$$

und

$$\lim_{S_f \rightarrow \infty} J_{\text{BPC}} = J_{\text{BBSC}} = \frac{(1-R)q\Phi_0(\alpha L)^2}{(\alpha L)^2 - 1} \cdot \frac{\frac{1}{\alpha} - \exp[-\alpha(d-w)] \left\{ L \sinh\left(\frac{d-w}{L}\right) + \frac{1}{\alpha} \cosh\left(\frac{d-w}{L}\right) \right\}}{L \sinh\left(\frac{d-w}{L}\right)}. \quad (2.25)$$

Quotientenbildung liefert:

$$\frac{J_{\text{BBSC}}}{J_{\text{FBSC}}} = \frac{1 - \exp[-\alpha(d-w)] \left\{ \alpha L \sinh\left(\frac{d-w}{L}\right) + \cosh\left(\frac{d-w}{L}\right) \right\}}{\exp(-\alpha d) + \exp(-\alpha w) \left\{ \alpha L \sinh\left(\frac{d-w}{L}\right) - \cosh\left(\frac{d-w}{L}\right) \right\}}. \quad (2.26)$$

Unter den Voraussetzungen $d \gg w$ und $\alpha w \ll 1$ ergibt sich schließlich eine Gleichung, die nur noch die Diffusionslänge L als Unbekannte enthält:

$$\frac{J_{\text{BBSC}}}{J_{\text{FBSC}}} = \frac{1 - \exp(-\alpha d) \left\{ \alpha L \sinh\left(\frac{d}{L}\right) + \cosh\left(\frac{d}{L}\right) \right\}}{\exp(-\alpha d) + \alpha L \sinh\left(\frac{d}{L}\right) - \cosh\left(\frac{d}{L}\right)}. \quad (2.27)$$

Beim BSC-Mode darf die Wellenlänge des Lasers nicht zu klein sein, da sonst, selbst für größere Diffusionslängen, die BBSC-Ströme sehr klein sind. Zwar ergeben größere Wellenlängen gut meßbare BBSC-Ströme, allerdings ist dann die Intensität an der Stelle $z = d$ noch so groß, daß die Reflexion des Lichtes an der Waferrückseite nicht vernachlässigt werden kann. Der reflektierte Teil des Lichtes generiert seinerseits Minoritäten, die eine Vergrößerung von sowohl J_{BBSC} als auch J_{FBSC} zur Folge haben.

Prinzipiell bereitet die Berechnung der Diffusionslänge für diesen Fall keine besonderen Schwierigkeiten, es sind lediglich die einzelnen Beiträge zum Gesamtstrom aus

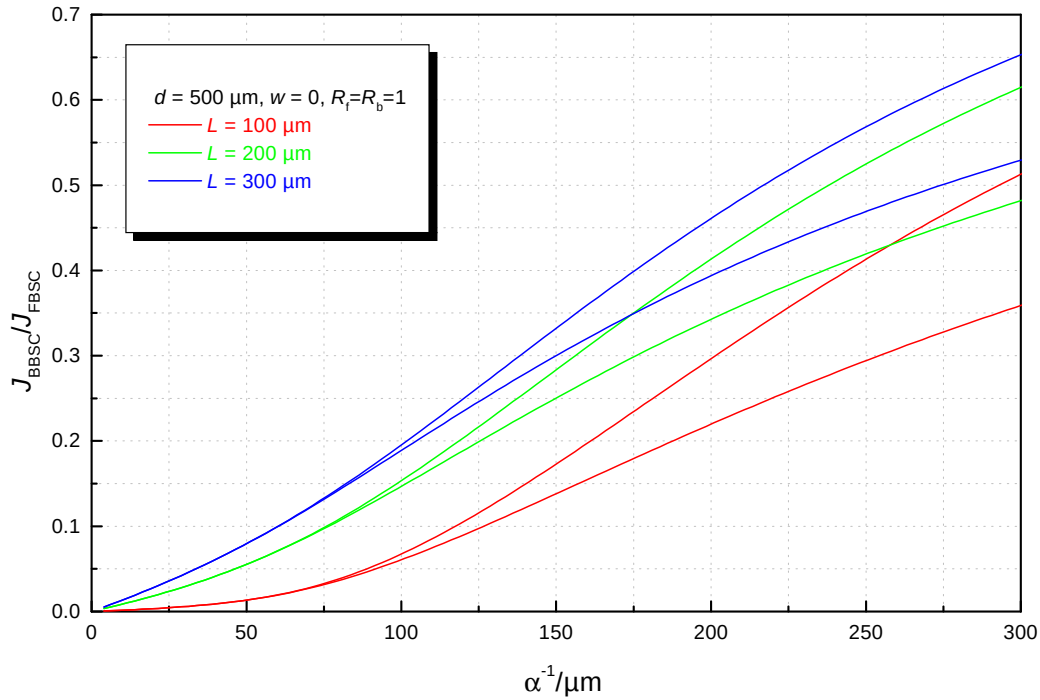


Abb. 2.4: BSC-Quotient in Abhängigkeit der Eindringtiefe α^{-1} des Lasers für verschiedene Diffusionslängen. Die untere der jeweiligen Kurven wurde mit Gleichung 2.27 berechnet. Bei den oberen Kurven, mit Gleichung 2.28 berechnet, wurden die Reflexionen an Vorder- und Rückseite des Wafers berücksichtigt.

hin- und rücklaufendem Licht aufzuaddieren⁴:

$$\frac{J_{\text{BBSC}}}{J_{\text{FBSC}}} = \frac{\sum_{\nu=0}^{\infty} \exp[-2\nu\alpha d] R_f \cdot F_{\text{BBSC}} + \exp[-(2\nu+1)\alpha d] R_b \cdot F_{\text{FBSC}}}{\sum_{\nu=0}^{\infty} \exp[-2\nu\alpha d] R_f \cdot F_{\text{FBSC}} + \exp[-(2\nu+1)\alpha d] R_b \cdot F_{\text{BBSC}}}. \quad (2.28)$$

R_f und R_b sind die Reflektivitäten der Vorder- bzw. der Rückseite und

$$\begin{aligned} F_{\text{FBSC}} &= \exp(-\alpha d) + \exp(-\alpha w) \left\{ \alpha L \sinh\left(\frac{d-w}{L}\right) - \cosh\left(\frac{d-w}{L}\right) \right\} \\ F_{\text{BBSC}} &= 1 - \exp[-\alpha(d-w)] \left\{ \alpha L \sinh\left(\frac{d-w}{L}\right) + \cosh\left(\frac{d-w}{L}\right) \right\}. \end{aligned}$$

Als unsicher gelten auch hier wieder die Reflektionskoeffizienten. Abb. 2.4 zeigt den Einfluß der Reflexionen des Lichtes an den Waferoberflächen, berechnet mit $R_f = R_b = 1$ für verschiedene Diffusionslängen. Abb. 2.4 zeigt weiter, daß für BSC-Messungen ein Laser mit einer Eindringtiefe um $\alpha^{-1} = 90 \mu\text{m}$ verwendet werden sollte.

⁴ Ist die Weite der Raumladungszone gegenüber der Waferdicke nicht vernachlässigbar, so müssen auch die Anteile beider Raumladungszonen bei der Berechnung der Ströme Berücksichtigung finden, siehe Fußnote auf Seite 38.

Kapitel 3

Aufbau des Leckstrom-Scanners

Seit Erfindung des 'Scanning Tunneling Microscopes' (STM) durch *Binnig* und *Rohrer* wurden zahlreiche weitere 'Scanning' Mikroskope entwickelt [SHI98]. Das Prinzip der meisten dieser Weiterentwicklungen besteht darin, einen Sensor (im allgemeinen eine Elektrode) über eine zu untersuchende Oberfläche zu führen und lokal eine elektrische Größe zu messen, wobei es sich in den meisten Fällen um einen Strom handelt.

Eine dieser Weiterentwicklungen ist das 'Scanning Electrochemical Microscope' (SECM) von *Bard et al.*: Hier wird eine Mikroelektrode mit einem Durchmesser zwischen $0.1\text{ }\mu\text{m}$ und $20\text{ }\mu\text{m}$ in geringem Abstand über die Oberfläche der zu untersuchenden Probe 'gescannt', die von einem Elektrolyten bedeckt wird. Zwischen Probe und Mikroelektrode wird ein festes oder variables Potential angelegt und der von den chemischen und physikalischen Eigenschaften der Probenoberfläche abhängige Elektrodenstrom gemessen [BAR92, FAN98, LEE91a, WEI95b].

Das SECM erlaubt die Untersuchungen an verschiedensten Proben mit Ortsauflösungen im Nanometerbereich, wie Metalle, Halbleiter, Isolatoren und auch biologischen wie beispielsweise Redoxreaktionen an Enzymen. Zudem lieferte das SECM Informationen über die Kinetik des Elektrontransferprozesses an der Grenzfläche [BOR94, HOR93a].

Aufbau und Funktionsweise des Leckstrom-Scanners sind einem SECM sehr ähnlich, allerdings arbeitet der Leckstrom-Scanner auf einer anderen Größenskala: Während mit dem SECM Probenflächen von ca. $1 \times 1\text{ }\mu\text{m}^2$ bis $100 \times 100\text{ }\mu\text{m}^2$ ausgemessen werden, soll der Leckstrom-Scanner bei der Oberflächen-Charakterisierung von multikristallinen $10 \times 10\text{ cm}^2$ Siliziumwafern Anwendung finden, aus denen Solarzellen gefertigt werden sollen.

3.1 Funktionsprinzip

Der Aufbau des Leckstrom-Scanners entspricht im wesentlichen einem vergrößerten SECM, siehe Abb. 3.1: Auf einen p-Siliziumwafer wird eine Meßzelle aufgesetzt und

diese anschließend mit Elektrolyt gefüllt¹. Ein unterseitig angebrachter O-Ring verhindert das Austreten von Elektrolyt, der aus verdünnter HF besteht. Näheres zum Elektrolyten in Abschnitt 3.5. Ein xyz-Manipulator erlaubt das 'Abscannen' der Probenoberfläche bei frei wählbarem Elektrodenabstand. Hierfür wird mit Hilfe eines Potentiostaten (nicht in Abb. 3.1 eingezeichnet) ein kathodisches Potential an die Elektrode gelegt und an verschiedenen Orten (x,y) der lokale Strom über den Si-HF-Kontakt gemessen. Werden die Ströme gegen x und y unter Verwendung von Falschfarben aufgetragen, ergeben sich Informationen über die örtliche Verteilung des Leckstromes.

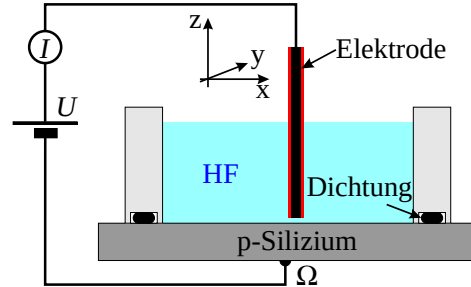


Abb. 3.1: Schematischer Aufbau des Leckstrom-Scanners.

Da eine Strommessung nur den Leckstrom über den Si-HF-Kontakt enthalten soll, muß die Messung im Dunkeln erfolgen, um Photostroman-teile im gemessenen Strom zu vermeiden.

Je nach verwendeter Elektrode zeigt der Elektroden-Strom eine mehr oder weniger starke Abstandsabhängigkeit, siehe Abb. 3.2. Da es kaum möglich ist, die Probe so zu fixieren, daß sie exakt parallel zur xy-Ebene der Elektrode liegt, muß eine Manipulation der Elektrode auch in z-Richtung möglich sein, um die Konstanz des Elektrodenabstandes während der gesamten Messung sicherzustellen.

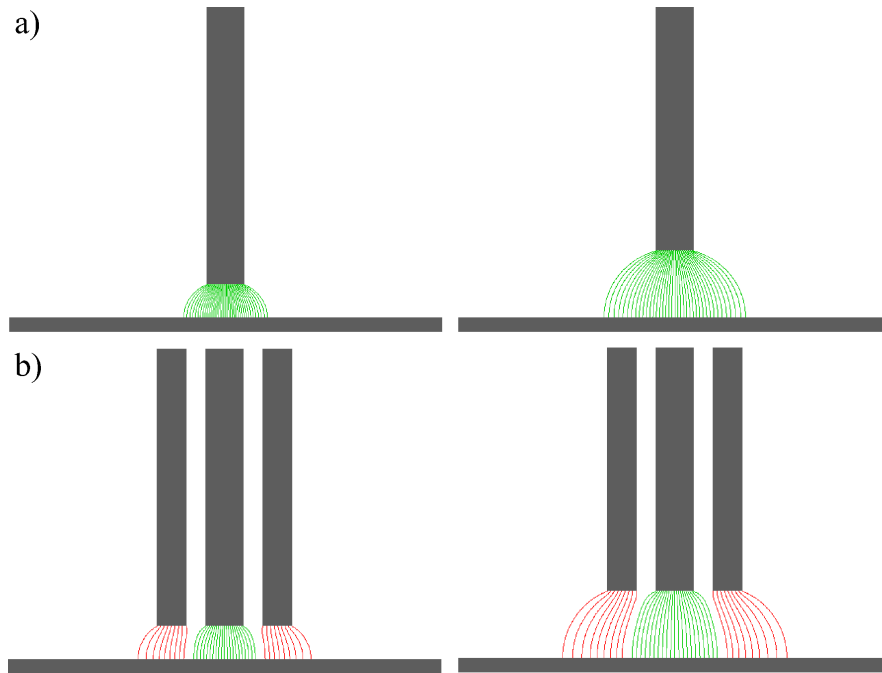


Abb. 3.2: Schematischer Feldverlauf unter a) der Platindraht-Elektrode und b) der 200 µm Zylinder-Elektrode (vgl. hierzu Abb. 3.13).

¹ Siehe Fußnote Seite 31.

3.2 Komponenten

Die Komponenten des Leckstrom-Scanners, die im folgenden beschrieben werden, sind

- der xyz-Manipulator
- die Schutzvorrichtung für den Piezo-Translator
- der Potentiostat
- der 'Peak'-Detektor
- die Elektroden

3.2.1 Der xyz-Manipulator

Alle Komponenten des xyz-Manipulators wurden von der Firma *Physik Instrumente GmbH & Co.* bezogen. Zum Einstellen der Elektrodenposition in x- und y-Richtung wurden zwei Linearversteller des Typs M-525.21 mit 200 mm Fahrweg verwendet. Die Anfahrgenauigkeit der Linearversteller beträgt 1 μm . Linearversteller kommen deshalb zum Einsatz, um die Möglichkeit besonders schneller Messungen nutzen zu können, indem man die Elektrode mit konstanter Geschwindigkeit in x- oder y-Richtung fahren läßt und in konstanten Zeitintervallen den Zellenstrom mißt. Diese Meßweise würde sich mit Schrittmotoren nicht realisieren lassen.

Als Manipulator für die z-Richtung wurde ebenfalls ein Linearversteller verwendet. Zum Einsatz kam der M-125.10 mit einem Fahrweg von 25 mm. Die Anfahrgenauigkeit dieses Linearverstellers beträgt ebenfalls 1 μm . Da für eine Korrektur des Elektrodenabstandes der Linearversteller für die z-Richtung zum einen zu langsam und zum anderen zu ungenau ist, das gilt im besonderen für das Messen mit konstanter Geschwindigkeit, wurde der Linearversteller durch einen Piezo-Translator ergänzt. Verwendet wurde der P-841.60, ein Translator mit einem Gesamtfahrweg von 90 μm und integriertem Sensor zur Positionsbestimmung. Das Ansteuern der Linearversteller über den IEEE-488-Bus erfolgte nicht direkt, sondern über den Motor-Controller C-804.50. Gleiches gilt für den Piezo-Translator, der über den Piezo-Controller P-862.50 angesteuert wurde.

Die Anfahr-Routinen der xyz-Positioniereinheit müssen sicherstellen, daß Linearversteller und Piezo-Translator so zusammenarbeiten, daß weder Piezo-Translator noch Elektrode beschädigt oder gar zerstört werden können. Dazu könnte es kommen, wenn die Elektrode beispielsweise gegen den Rand der Meßzelle läuft. Gleichfalls muß sichergestellt sein, daß bei Bewegungen in z-Richtung die Elektrode die Probenoberfläche nie berührt: Ist beispielsweise im Verlauf der Messung der Elektrodenabstand zu groß geworden, agiert im allgemeinen der Piezo-Translator und stellt den Elektrodenabstand wieder auf den richtigen Wert ein. Wenn sich aber nun der Piezo-Translator bereits im Zustand maximaler Expansion befindet, dieser also nicht in der Lage ist, den Elektrodenabstand zu korrigieren, muß mit dem Linearversteller gefahren werden, vgl. rechte Seite in Abb. 3.3. Um ein häufiges Fahren mit dem Linearversteller aus den oben genannten Gründen zu vermeiden ist es sinnvoll, den Elektrodenabstand nicht mit dem Linearversteller allein zu korrigieren, sondern mit dem Linearversteller so weit nach unten zu fahren, daß, wenn sich der Piezo-Translator in der Mittelposition befindet,

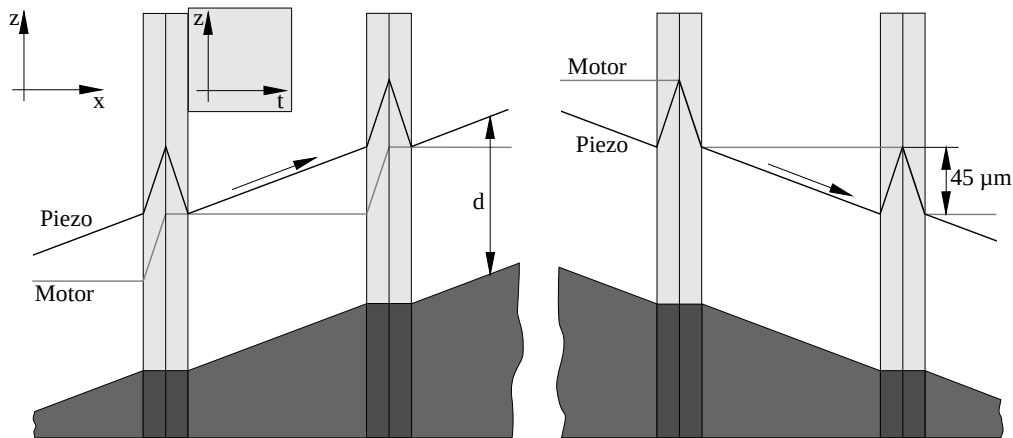


Abb. 3.3: Schematische Arbeitsweise des z-Manipulators. Die hellgrau unterlegten Bereiche repräsentieren Orte auf der Probe an denen sich der z-Manipulator in Abhängigkeit der Zeit neu konfiguriert. Hierbei kennzeichnet die graue Kurve den Weg des Linearverstellers (bez. als Motor) und die schwarze Kurve den des Piezo-Translators. Die Pfeile parallel zur schwarzen Kurve zeigen die 'Scan'-Richtung.

der Abstand der Elektrode den gewünschten Wert hat. Um ein Aufschlagen der Elektrode auf die Probe zu vermeiden, muß *erst* der Piezo-Translator in die Mittelposition gefahren werden, *bevor* mit dem Linearversteller gefahren wird.

Entsprechendes gilt, wenn der Elektrodenabstand zu klein geworden ist und sich der Piezo-Translator im Zustand minimaler Expansion befindet, wie im linken Teil von Abb. 3.3 gezeigt. In diesem Fall muß erst mit dem Linearversteller und anschließend mit dem Piezo-Translator gefahren werden. Die Abbildung zeigt weiter, daß der Abstand der Elektrode, der der schwarzen Linie entspricht, beim Fortschreiten in x-Richtung nie d unterschreitet.

3.2.2 Die Schutzvorrichtung für den Piezo-Translator

Piezo-Translator und Elektroden sind mit Abstand die empfindlichsten Komponenten des Leckstrom-Scanners. Selbst geringe auf den Piezokopf wirkende Torsionskräfte oder Kräfte mit Komponenten senkrecht zur Ausdehnungsrichtung müssen unbedingt vermieden werden, anderenfalls besteht die Gefahr der Zerstörung des Piezo-Translators. Da Piezo-Translatoren vergleichsweise teuer sind, ist das Verwenden einer Schutzvorrichtung für den Piezo-Translator auf jeden Fall sinnvoll.

Funktionsprinzip: Um die teilweise aus Metall bestehende Verbindung zwischen Piezo-Translator und Elektrode ist eine Metallhülse angebracht, die vom metallischen Teil des Verbindungsstückes elektrisch isoliert ist. Diese zwei Metallteile bilden einen Schalter wie in Abb. 3.4 gezeigt.

Im Fall der Meßbereitschaft ist das Flip-Flop gesetzt, und das Relais befindet sich im angezogenen Zustand, so daß die elektrischen Zuleitungen des Motors mit dem Controller verbunden sind.

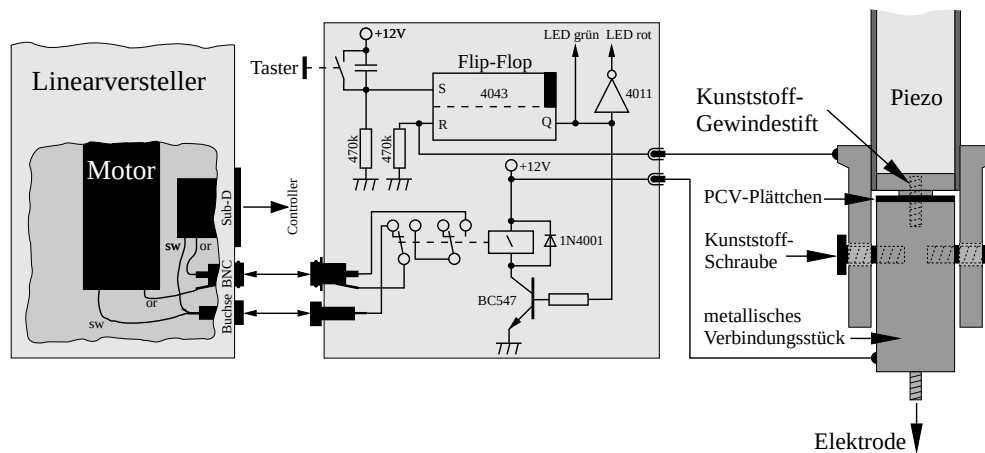


Abb. 3.4: Aufbau der Piezo-Schutzvorrichtung. Stößt die Elektrode in x- oder y-Richtung auf ein Hindernis, wird der Flip-Flop zurückgesetzt, das Relais fällt ab und trennt dadurch die elektrischen Zuleitungen zum Motor.

Stößt die Elektrode in x- oder in y-Richtung auf ein Hindernis, so berühren sich beide Metallteile, wodurch der Schalter geschlossen und das Flip-Flop zurückgesetzt wird. Dadurch fällt das Relais ab und trennt die Zuleitungen zum Motor vom Controller. Auf diese Weise kommt der Meßkopf zum Stillstand, und eine Zerstörung von Piezo-Translator und Elektrode ist damit verhindert. Im abgefallenen Zustand sind die beiden Zuleitungen zum Motor kurzgeschlossen. Dieser Kurzschluß wirkt als EMK²-Bremsen und verhindert sehr effektiv ein Nachlaufen des Motors. Die 12 V, die am inneren Teil des Schalters anliegen, haben auf die gemessenen Ströme keinen Einfluß, da sich zwischen diesem Metallteil und der eigentlichen Elektrode noch ein Distanzstück aus PVC befindet, siehe Abb. 3.17.

Sicherlich ist das Verwenden von Relais als Unterbrecher bedenklich, weil Relais zum Schalten Zeit benötigen. Da aber die Verbindung zwischen Piezo-Translator und Elektrode bewußt nicht starr konstruiert wurde (schon deshalb, damit ein Schließen des Schalters überhaupt möglich ist) und zusätzlich die Maximalgeschwindigkeit des x- und y-Linearverstellers auf einen geringen Wert eingestellt wurde (ca. $2.0 \frac{\text{cm}}{\text{s}}$), fängt die flexible Verbindung ein geringes Nachlaufen der Linearversteller ab. Für die Verwendung von Relais spricht, daß sie auf einfache Weise eine galvanische Trennung zwischen der Piezo-Schutzvorrichtung und der Elektronik der Linearversteller erlauben, so daß Wechselwirkungen zwischen beiden Komponenten ausgeschlossen sind.

Die beiden Kunststoff-Schrauben sollen verhindern, daß Drehmomente auf den Piezokopf wirken können. Da aus Gründen der Isolation die Schrauben aus Kunststoff sein müssen, der relativ weich ist, ist die Schutzwirkung der Schrauben nicht sehr groß. Aufgrund der Elektroden-Geometrie, siehe Abschnitt 3.2.5, kann im Fehlerfall das Einwirken von Drehmomenten aber ausgeschlossen werden, so daß diese Schwachstelle den Piezo-Translator kaum gefährdet.

² EMK = **E**lektromotorische **K**raft

Wesentlich schwerer wiegt die Tatsache, daß die Vorrichtung nicht verhindern kann, daß die Elektrode zu weit nach unten gefahren wird, wenn beispielsweise der z-Wert der Anfahrposition zu niedrig ist. Da aber ein Piezo-Translator sehr große Kräfte kompensieren kann, die parallel zur Expansionsrichtung wirken, sollte dieser keinen Schaden nehmen. Dies gilt aber nicht für die Elektrode, die beschädigt oder gar zerstört werden würde. Ein Drücken gegen die Verbindung zwischen Elektrode und Piezo-Translator, um die Motoren auszuschalten, würde nicht von Nutzen sein, da dabei lediglich der x- und der y-Linearversteller abgeschaltet werden würden, nicht aber der Linearversteller für die z-Richtung. Die einzigen Maßnahmen, die hier greifen, sind das rechtzeitige Drücken des am Motor-Controller rückseitig angebrachten Reset-Knopfes oder das Ziehen des Netzsteckers.

3.2.3 Der Potentiostat

In Abb. 2.3 wurde gezeigt, daß beim Elymaten der Widerstand des Kontaktes zwischen Silizium und Wolframcarbid-Nadel Ursache dafür ist, daß beim Messen der Dunkel- und Beleuchtungsströme bei jeweils gleicher Zellenspannung an verschiedenen Orten der Kennlinie des Silizium-Flußsäure-Kontaktes gemessen wird, weil der Spannungsabfall über der Grenzfläche vom Strom abhängt. Durch Verwendung von Referenz-Elektroden kann der Arbeitspunkt auf der Kennlinie stabilisiert werden, so daß dessen Lage auf der Kennlinie unabhängig von der Stärke des fließenden Zellenstromes ist.

Dieser Sachverhalt gilt genauso für den Leckstrom-Scanner, nur werden hier Variationen im Strom nicht durch Änderung der Lichteinstrahlung hervorgerufen, sondern durch die laterale Dynamik des Leckstromes. An einem Ort mit kurzgeschlossenener Raumladungszone ('Shunt') sind verhältnismäßig hohe Leckströme zu erwarten. Je größer der fließende Strom, um so höher ist der Spannungsabfall über dem Widerstand des Rückseiten-Kontaktes und um so mehr verringert sich der Potentialabfall über der Silizium-Flußsäure-Grenzfläche, woraus auch hier eine Verschiebung des Arbeitspunktes resultieren würde.

Diese Verschiebung des Arbeitspunktes kann, wie beim Elymaten, durch Verwendung einer Referenz-Elektrode vermieden werden. Das Potential der Referenz-Elektrode dient dem Potentiostaten zum Stabilisieren des Arbeitspunktes, so daß dieser unabhängig vom fließenden Zellenstrom ist.

Der Potentiostat gehört zu den wichtigsten Meß- bzw. Regelinstrumenten der Elektrochemie, so daß auf eine kurze Beschreibung nicht verzichtet werden sollte.

Abb. 3.5 zeigt den Aufbau eines Potentiostaten. Hier entspricht WE ('Working'-Elektrode) dem Silizium und CE ('Counter'-Elektrode) der nicht ortsfesten Meßelektrode. Vorgegeben über eine externe Spannungsquelle wird das gewünschte Potential U_{in} , welches über der Grenzfläche abfallen soll. Im eingeregelter Zustand verschwindet die Potentialdifferenz der Eingänge³: $U_+ - U_- = 0$. Deswegen liegt der invertierende Eingang des Verstärkers CA ('Control-Amplifier') auf Massepotential (auch als virtu-

³Im Falle eines übersteuerten Operationsverstärkers verschwindet die Potentialdifferenz seiner Eingänge nicht.

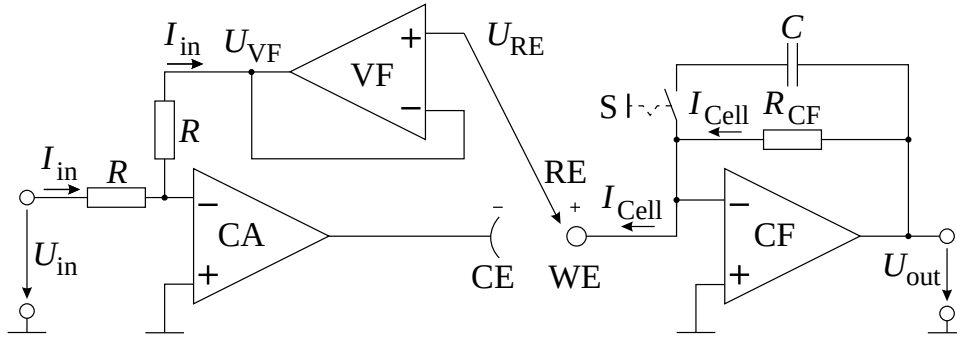
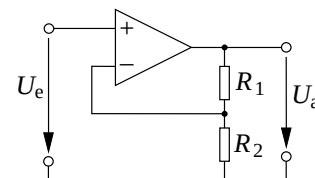


Abb. 3.5: Schaltungstechnische Realisierung eines Potentiostaten. Der Verstärker CA ('Control-Amplifier') stellt den eigentlichen Potentiostaten dar. Aufgabe des Spannungsfolgers VF ('Voltage-Follower') ist es sicherzustellen, daß das Messen des Potentials der Referenz-Elektrode stromlos erfolgt. Der Verstärker CF ('Current-Follower') arbeitet als Strom-Spannungs-Konverter. Die Spannungen U_{RE} und U_{VF} sind bezogen auf das Massepotential [HAS99, POP00].

elle Masse bezeichnet). Somit fließt über den Widerstand R Strom $I_{in} = \frac{U_{in}}{R}$. Da so gut wie kein Strom in den Operationsverstärker hineinfließt, muß durch den Widerstand R , der mit dem Ausgang des Spannungsfolgers⁴ VF ('Voltage-Follower') verbunden ist, der Strom $-I_{in}$ fließen. Damit dieser Strom fließen kann, muß am Ausgang des Spannungsfolgers und damit auch, wie in der Fußnote beschrieben, an dessen Eingang das Potential $-U_{in}$ anliegen: $U_{RE} = U_{VF} = -U_{in}$. Als Spannungsfolger sollte möglichst ein Operationsverstärker mit FET-Eingang Verwendung finden. Da solche Verstärker Eingangswiderstände in der Größenordnung $10^{14} \Omega$ haben, wird das Messen der Spannung U_{RE} stromlos erfolgen. Dies ist sehr wichtig, da bereits kleine, über die Referenz-Elektrode fließende, Ströme das Gleichgewicht an der Grenzfläche zwischen Referenz-Elektrode und Elektrolyt empfindlich stören können, wodurch die sich der Potential der Referenz-Elektrode merklich ändert.

Ist beispielsweise das Potential der Referenz-Elektrode zu groß, d. h. gilt $U_{RE} = U_{VF} > -U_{in}$, so gilt für die Potentialdifferenz der Eingänge des Verstärkers CA: $U_+ - U_- < 0$. Da die Differenz negativ ist, wird sich die Ausgangsspannung des Operationsverstärkers und damit das Potential an der 'Counter'-Elektrode verringern, was eine Vergrößerung der Potentialdifferenz zwischen 'Working'- und 'Counter'-Elektrode zur Folge hat: $U_{CE} = A_D(U_+ - U_-)$, wobei A_D die Differenzverstärkung des Operationsverstärkers bedeutet. Unter der Voraussetzung, daß die Zelle mit einem elek-

⁴ Die rechtsstehende Abbildung zeigt die Grundsaltung eines nicht-invertierenden Verstärkers. Über das Widerstandsnetzwerk R_1 und R_2 wird der Anteil $k = \frac{R_2}{R_1 + R_2}$ der Ausgangsspannung U_a auf den invertierenden Eingang rückgekoppelt. Für den Verstärkungsfaktor A des Verstärkers gilt damit: $A = \frac{U_a}{U_e} = \frac{1}{k} = 1 + \frac{R_1}{R_2}$. Wählt man $R_1 = 0$ und läßt R_2 fort ($R_2 \rightarrow \infty$) wird $A = 1$. Da für diesen Fall die Ausgangsspannung der Eingangsspannung folgt ($U_a = U_e$), bezeichnet man diese Schaltung auch als *Spannungsfolger*, die sehr oft als *Impedanzwandler* Anwendung findet, weil zwischen der Impedanz des Einganges und der des Ausganges viele Größenordnungen liegen [TIE93].



trisch leitenden Medium gefüllt ist, hat die vergrößerte Potentialdifferenz zwischen 'Working'- und 'Counter'-Elektrode im allgemeinen eine Zunahme des Zellenstromes I_{Cell} zur Folge. Ein größerer Zellenstrom führt zu einer Verkleinerung des Potentials an der Referenz-Elektrode U_{RE} und damit auch von U_{VF} . Im eingeregelter Zustand gilt daher: $U_{\text{RE}} = -U_{\text{in}}$, und zwar *unabhängig* vom fließenden Zellenstrom I_{Cell} !

Der Stromfolger CF ('Current-Follower') hat zwei Funktionen: Zum einen sorgt er dafür, daß die 'Working'-Elektrode auf definiertem Potential liegt. Dieses Potential ist gleich dem Massepotential, da der nicht-invertierende Eingang mit Masse verbunden ist. Zum anderen arbeitet dieser Verstärker als Strom-Spannungs-Konverter: Da kein Strom in den Operationsverstärker hineinfließt, muß der Strom quasi über den Widerstand R_{CF} 'ausweichen'. Damit der Strom I_{Cell} durch den Widerstand R_{CF} fließt, muß der Verstärker die Ausgangsspannung $U_{\text{out}} = I_{\text{Cell}} R_{\text{CF}}$ bereitstellen. Liegt die Counter-Elektrode auf negativem Potential, so gilt für die Ausgangsspannung des Stromfolgers: $U_{\text{out}} > 0$.

Durch Zuschalten des Kondensators C über den Schalter S wird die Bandbreite des Verstärkers reduziert. Dies führt bei kleinen Strömen zu einer Verkleinerung des Meßrauschens und verhindert zudem, daß der Potentiostat ins Schwingen geraten kann. Für das Messen hochfrequenter Ströme, wie bei der Impedanz-Spektroskopie, dürfte der Kondensator nicht zugeschaltet werden, weil bei höheren Frequenzen, aufgrund der Tiefpaß-Eigenschaften des Stromfolgers, das Meßsignal selbst gedämpft werden würde [HAS99]. Da mit dem Leckstrom-Scanner ausnahmslos nur Gleichströme gemessen werden, ist das Zuschalten des Kondensators auf jeden Fall zu befürworten.

3.2.4 Der 'Peak'-Detektor

Enthält die Zelle einen wässrigen Elektrolyten, wie er für Messungen verwendet wird, ist Stromfluß zwangsweise mit der Entstehung von Gasen verbunden: An der Siliziumelektrode die Entwicklung von Wasserstoff und an der Meßelektrode die Entwicklung von Sauerstoff. Die permanente Bildung von Gasbläschen hat zur Folge, daß das Meßsignal stark verrauscht ist. Näheres zum Elektrolyten in Abschnitt 3.5.

Aus diesem Grund wird an den Ausgang des Potentiostaten ein 'Peak'-Detektor angeschlossen, der während eines vom Experimentator frei wählbaren Zeitintervalls Δt ('Sample'-Zeit) nach dem Maximalwert des Stromes 'sucht'⁵. Die Arbeitsweise des 'Peak'-Detektors verdeutlicht Abb. 3.6.

Aus zwei Gründen wurde ein digital arbeitender 'Peak'-Detektor gewählt:

1. Da der aktuelle Maximalwert des Stromes in einem Speicher-Register gesichert wird, ergeben sich theoretisch unendlich lange Haltezeiten. Bei Verwendung einer analog arbeitenden 'Sample-Hold'-Schaltung würde der Spannungswert exponentiell mit der Zeit abfallen.

⁵ Der Potentiostat konvertiert den Strom in eine Spannung deren Wert gemessen wird. Daher läuft die Bestimmung des maximalen Stromes auf das Abtasten einer Spannung hinaus.

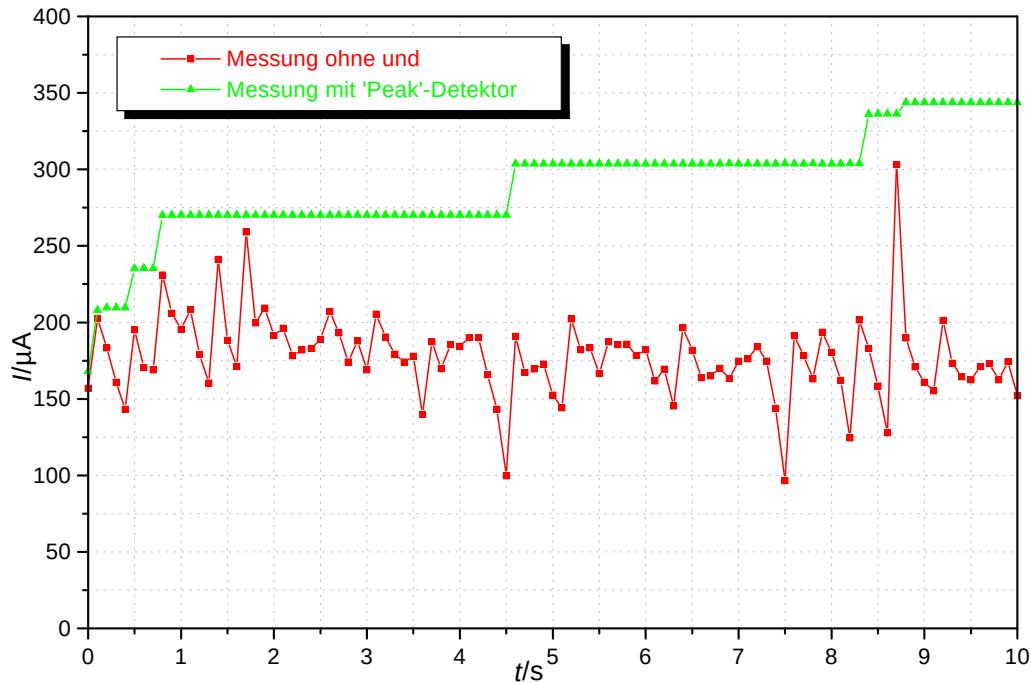


Abb. 3.6: Arbeitsweise des 'Peak'-Detektors. Die $200\text{ }\mu\text{m}$ Zylinder-Elektrode wurde in Nähe eines Kratzers im Abstand von $d = 250\text{ }\mu\text{m}$ positioniert. Wird ein Spannungswert digitalisiert, der größer ist als der gespeicherte, wird das Speicher-Register mit dem neuen Spannungswert überschrieben. Das Meßgerät, mit dem die Ausgangsspannung des Potentiostaten gemessen wird (untere Messung), mittelt pro Meßpunkt über 16.67 ms . Während dieses Zeitraumes startet der 'Peak'-Detektor fast 200 'Sample'-Vorgänge. Aus diesem Grund sind die mit dem Meßgerät gemessenen Spannungen stets kleiner die mit dem Peak-Detektor gemessenen.

2. Das Vorliegen des Meßwertes in digitaler Form erlaubt die Einsparung des Digital-Multimeters, da die Meßwerterfassung auch über die parallele Schnittstelle des Computers erfolgen kann.

Diese beiden Vorzüge eines digital arbeitenden 'Peak'-Detektors erfordern allerdings einen etwas erhöhten schaltungstechnischen Aufwand, siehe Anhang B.

Während eines vom Experimentator frei wählbaren Zeitintervalls Δt wird die Ausgangsspannung des Potentiostaten mit einer Frequenz von 11.3 kHz abgetastet. Ist der aktuelle Wert größer als der momentan gespeicherte, wird dieser Spannungswert in die Speicher-Register geschrieben, ansonsten wird der Wert ignoriert. Wenn die Zeit Δt verstrichen ist, wird der 'Sample'-Vorgang gestoppt und der Inhalt des Registers als Meßwert ausgelesen.

Zentraler Baustein des 'Peak'-Detektors, dessen Blockschaltbild in Abb. 3.7 dargestellt ist, ist der AD574A ein 12 Bit Analog-Digital-Wandler, von der Firma Analog Devices. Für einen Spannungsbereich von $\pm 10\text{ V}$ beträgt die Auflösung $\pm 4.88\text{ mV}$, was einem relativen Fehler bei der Strommessung von unter $0.05\text{ }\%$ entspricht. Der Oszil-

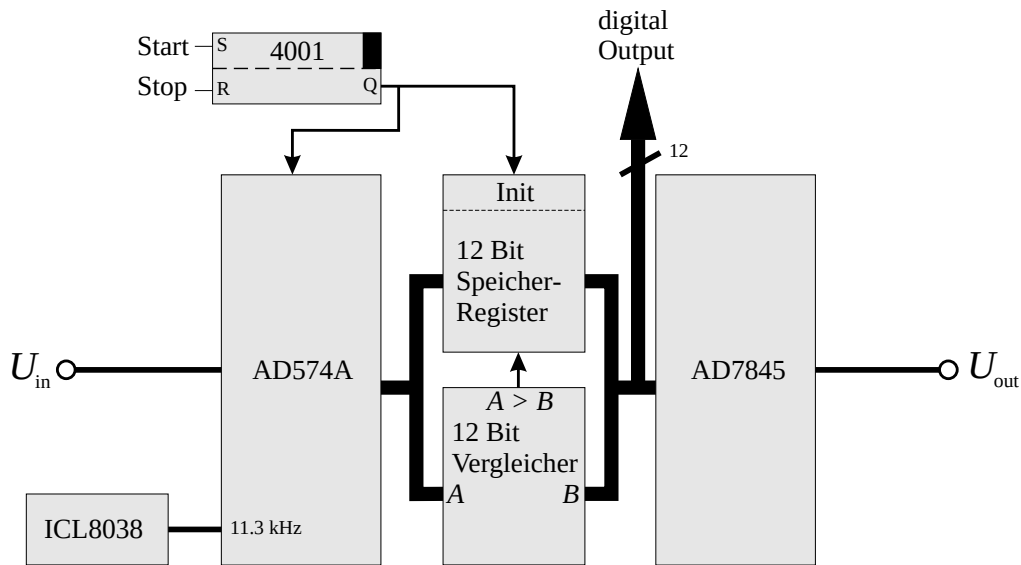


Abb. 3.7: Blockschaltbild des 'Peak'-Detektors. Nach Start des 'Sample'-Vorganges wird das Register mit Nullen gefüllt (Init). Der vom Potentiostaten gelieferte Spannungswert U_{in} wird vom AD745A mit 12 Bit Genauigkeit digitalisiert. Ist dieser Wert A größer als der momentan im Register gespeicherte Wert B, wird der Inhalt des Registers überschrieben. Beim Abtasten nach einem Minimum wird das Register während der Initialisierung mit Einsen gefüllt. Ein Speichern des aktuellen Wertes in das Register erfolgt in diesem Fall für $A < B$.

lator *ICL8038* liefert für den 'Sample'-Vorgang ein Rechtecksignal mit einer Frequenz von 11.3 kHz, der durch das Setzen des Flip-Flops (4043) gestartet wird. Gleichzeitig wird mit Setzen des Flip-Flops das 12 Bit Speicher-Register initialisiert. Dabei werden alle Bits des Registers entweder mit Nullen oder mit Einsen gefüllt, je nachdem ob das Signal nach einem Maximum (Nullen) oder einem Minimum (Einsen) abgetastet werden soll. Um eine einfache Integration des 'Peak'-Detektors in den bestehenden Aufbau des Leckstrom-Scanners zu ermöglichen, wurde dem Speicher-Register der *AD7845* ein 12 Bit Digital-Analog-Wandler, ebenfalls von der Firma *Analog Devices*, nachgeschaltet.

Mit dem Rücksetzen des Flip-Flops wird der 'Sample'-Vorgang beendet, wobei der Zustand des Registers erhalten bleibt. Sein Inhalt geht erst dann verloren, wenn ein erneuter 'Sample'-Vorgang gestartet wird (Füllung mit Nullen oder Einsen).

Im Meßprogramm wurde eine automatische Umschaltung des Meßbereiches realisiert: Ist die Ausgangsspannung des Potentiostaten kleiner als 4.9 V, sampelt der 'Peak'-Detektor im Spannungsbereich ± 5 V, in dem die Auflösung ± 2.44 mV beträgt. Übersteigt die Spannung die Schwelle von 4.9 V, wird in den ± 10 V Bereich gewechselt und die Messung wiederholt. In diesem Bereich verbleibt der 'Peak'-Detektor solange, bis die Schwelle wieder unterschritten wird.

3.2.5 Die Elektroden

Von der Qualität der hergestellten Elektroden wird abhängen, wie gut, d. h. mit welcher Ortsauflösung, Defektstrukturen auf den Siliziumwafern vermessen werden können. Deshalb muß die Herstellung der Elektroden mit größter Sorgfalt erfolgen.

Die Herstellung der Mikroelektroden für ein SECM ist relativ einfach: Man führt einen dünnen Draht aus Edelmetall, z. B. aus Platin, in ein feines Quarzröhrchen ein. Über einer heißen Flamme wird das Glasröhrchen so lange erwärmt, bis es beginnt weich zu werden. Dann wird das Röhrchen langgezogen, wobei es gleichzeitig gedreht wird. Schließlich ist der Edelmetalldraht fest von einer dünnen isolierenden Glasschicht umgeben. Um der Mikroelektrode eine definierte Geometrie zu geben, muß sie abschließend noch poliert werden, was aufgrund ihrer großen Empfindlichkeit mit besonderer Vorsicht erfolgen muß [LEE91a]. Anstelle eines Edelmetalldrahtes wird oftmals auch eine Kohlenstoff-Faser verwendet, vor allem dann, wenn sehr feine Elektroden benötigt werden. Auf diese Weise lassen sich Elektroden herstellen, bei denen der Durchmesser der elektrochemisch aktiven Region $1\text{ }\mu\text{m}$ beträgt [BIX86, HOR93b, KWA89b, MAC92]. Solche Elektroden können jedoch nicht beim Leckstrom-Scanner zum Einsatz kommen, da hier ein flußsäurehaltiger Elektrolyt verwendet werden muß, der die Elektrode binnen kurzer Zeit zerstören würde.

Im Verlauf einer früheren Diplomarbeit am Lehrstuhl wurde eine Elektrode aus einem Iridium-Stab hergestellt, einem sehr harten Edelmetall. Nachdem der Stab spitz zugeschliffen worden war, wurde er mit einer Kunststoff-Ummantelung versehen und der vordere Teil vollständig in Paraffin eingebettet. Durch leichtes Reiben über Papier wurde schließlich der unmittelbare Teil der Spitze selbst freigelegt.

Zwar konnte mit dieser Elektrode, wie in Abb. 3.8 a) gezeigt, die Korn- bzw. De-

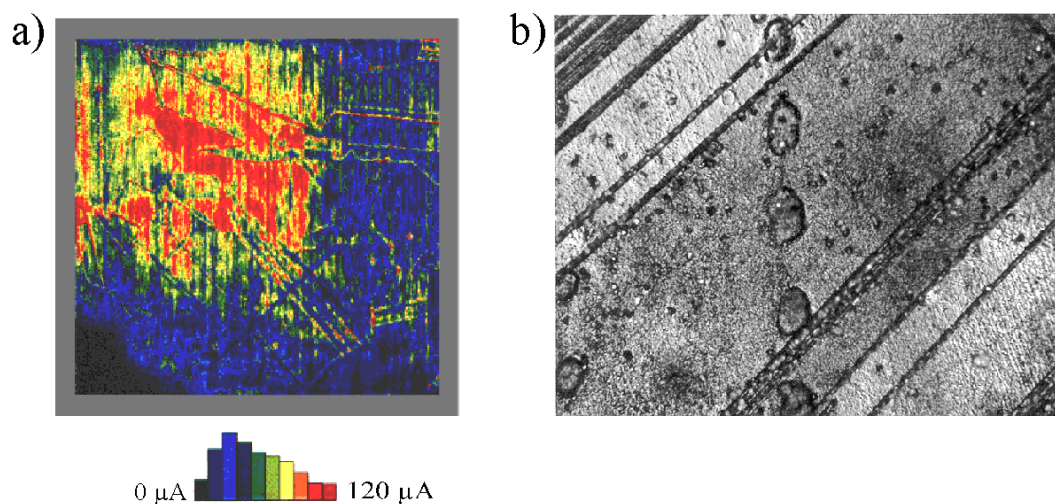


Abb. 3.8: a) 'Scan' eines multikristallinen Siliziumwafers mit einer spitzen Iridium-Elektrode. b) Rasterelektronenmikroskopische Aufnahme der Probenoberfläche nach der Messung. Deutlich zu erkennen ist das 'Scan'-Raster, entnommen aus [LIE95].

fektstruktur von multikristallinem Silizium aufgelöst werden, die Ströme, die über die Elektrode flossen, waren aber viel zu groß! Dies ergab ein Vergleich zwischen dem Strom, der sich aus der Aufsummierung aller lokalen Ströme unter Berücksichtigung der Ortsauflösung ergibt, und dem integralen Leckstrom, gemessen mit einer Platin-Schleife. Nach der Messung zeigte eine Betrachtung der Probe unter dem Mikroskop, daß sie an den Meßpunkten stark korrodiert war, siehe Abb. 3.8 b). Dies ließ vermuten, daß es durch die sehr spitze Form der Elektrode und den daraus resultierenden sehr hohen elektrischen Feldstärken zu einem Felddurchbruch durch die Silizium-Elektrolyt-Grenzfläche gekommen ist. Ein Felddurchbruch würde nicht nur die zu hohen Ströme erklären, sondern auch die lokale Degradation der Probenoberfläche [LIE95].

Der Aufbau der Elektroden muß also derartig sein, daß feldinduzierte Ströme vermieden werden. Deshalb wird auf den Einsatz spitzer Elektroden vollständig verzichtet. Stattdessen werden nur solche Elektroden verwendet, deren elektrochemisch aktive Fläche zuvor sorgfältig poliert wurde. Zwecks Vermeidung von feldinduzierten Strömen ist nicht nur die Geometrie der Elektrode relevant, sondern auch das über der Silizium-Elektrolyt-Grenzfläche abfallende kathodische Potential, das darum nicht zu hoch gewählt wird.

Im Verlauf der Arbeit wurden vier verschiedene Elektroden gebaut, deren Herstellung im folgenden beschrieben wird. Alle Materialien, die zum Bau der Elektroden Verwendung fanden, wurden, wenn nichts anderes angegeben, von der Firma *Goodfellow GmbH* bezogen.

3.2.5.1 Die Platindraht-Elektrode

Die Platindraht-Elektrode besteht lediglich aus einem isolierten 200 μm dicken Platindraht. Um ein Brechen des Pt-Drahtes während der Präparation zu vermeiden, wird der Draht über der Flamme eines Gasbrenners zur Weißglut gebracht, wodurch der Pt-Draht sehr weich wird. Als Beschichtung wird Polyurethan gewählt, ein besonders widerstandsfähiger Isolierlack. Vor dem Aufbringen des Isolierlackes wird der Pt-Draht sorgfältig mit Aceton gereinigt. Es werden insgesamt zwei Polyurethanschichten aufgebracht.

Von Vorteil ist, daß die Polyurethanschicht recht weich ist, so daß sie beim Verbiegen des Pt-Drahtes nicht beschädigt wird. Dem steht allerdings der Nachteil gegenüber, daß die Schicht glasklar ist und daher unter dem Stereo-Mikroskop nicht geprüft werden kann, ob der Draht vollständig mit Polyurethan bedeckt ist. Lecks entlang des Drahtes sind unbedingt zu vermeiden! Deshalb wird der lackisolierte Pt-Draht fast auf ganzer Länge von einem Schrumpfschlauch umgeben, lediglich die Enden des Drahtes bleiben frei.

Auch wenn die gemessenen Ströme nicht von hochfrequenten Einstreuungen von außerhalb der Zelle gestört werden sollten, da diese vom Potentiostaten herausgefiltert werden, vgl. Abb. 3.5, wird die Elektrode dennoch bis kurz vor Ende des Pt-Drahtes mit einer Abschirmung versehen. Hierzu dient Aluminiumfolie, die um die Elektrode gewickelt wird. Um die Folie elektrolytseitig vor Korrosion zu schützen, wird sie mit einem weiteren Schrumpfschlauch umgeben. Am oberen Ende wird die Folie mit ei-

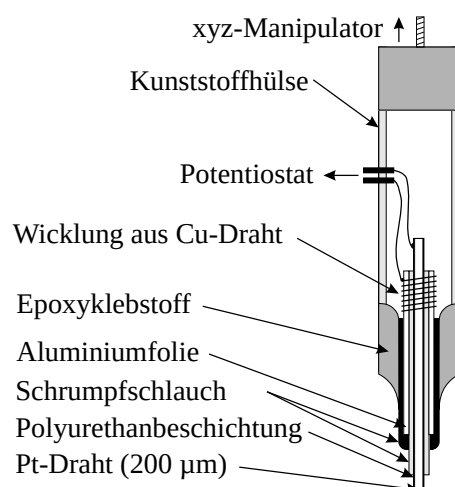


Abb. 3.9: Aufbau der Platindraht-Elektrode. Die Seiten des Drahtes werden zwecks Isolierung mit Polyurethan beschichtet.

nem dünnen Kupferdraht umwickelt, um den elektrischen Kontakt zur Abschirmung herzustellen.

Vom oberen Ende des Platindrahtes wird die Polyurethanschicht wieder entfernt. Hierfür wird kein Messer verwendet, da dann die Gefahr eines späteren Bruches sehr groß ist, sondern ein auf eine sehr hohe Temperatur (ca. 450 °C) erhitzter Lötkolben, mit dem die Lackschicht am oberen Ende aufgeweicht wird. Danach wird ebenfalls ein Kupferdraht zur Kontaktierung an den Pt-Draht gelötet. Damit die Elektrode später am xyz-Manipulator befestigt werden kann, wird die Elektrode in eine Kunststoffhülse eingesetzt und die Öffnung mit Epoxylebstoff vergossen.

Da auch das Ende des Platindrahtes mit Polyurethan bedeckt ist, besteht der letzte Schritt, der auch gleichzeitig der kritischste ist, darin, die unmittelbare Spitze der Elektrode abzutrennen. Hierfür wird eine scharfe Schere verwendet. Es sollte versucht werden, den Draht möglichst senkrecht durchzuschneiden, um die Elektrodenfläche nicht unnötig zu vergrößern. Den Aufbau der Elektrode im Detail zeigt Abb. 3.9.

Bei später hergestellten Elektroden wurde die Polyurethan- durch eine rote Acryllackschicht ersetzt. Hierfür gibt es zwei Gründe: Erstens kann die Güte der Polyurethanschicht nicht, wie oben beschrieben, optisch geprüft werden. Zweitens, und dieser Grund wiegt schwerer, löst sich die Polyurethanbeschichtung mit der Zeit aufgrund des Ethanolanteils im Elektrolyten ab, siehe Abschnitt 3.5. Von Nachteil ist allerdings, daß die Acryllackschicht vergleichsweise hart ist und daher leicht bricht, wenn der Pt-Draht gebogen wird. Dadurch würde die Elektrode unbrauchbar werden, da sie über die Risse in der Isolierung nahezu den gesamten Leckstrom des Wafers ziehen würde.

3.2.5.2 200 μm Zylinder-Elektrode

Wie die Ausführungen in Abschnitt 3.1 und die Experimente in Abschnitt 4.2 zeigen, weist das Auflösungsvermögen der Platindraht-Elektrode im Vergleich zu den anderen Elektroden eine verhältnismäßig starke Abstandsabhängigkeit auf. Aus der Tatsache, daß zwischen Siliziumwafer und Elektrode Ströme fließen, resultiert ein elektrisches Feld. Gelingt es, das Feld zu homogenisieren, wie es Abb. 3.2 schematisch zeigt, sollte sich die Abstandsabhängigkeit der Ortsauflösung erheblich verringern. Danach muß für eine Verringerung der Abstandsabhängigkeit die Elektrode aus zwei gegeneinander isolierten Metallflächen bestehen: der Meßelektrode im Zentrum und einer sie umgebenden Ring- oder Zylinder-Elektrode.

Für die Innenelektrode wird wieder ein 200 μm dicker Platindraht gewählt. Als Zylinder-Elektrode kommt ein Platin-Röhrchen mit einem Innendurchmesser von 400 μm und einem Außendurchmesser von 700 μm zum Einsatz. Anfänglich wurde die Innenelektrode mit Photolack isoliert. Photolack wurde deshalb verwendet, weil er sehr dunkel ist und daher optisch leicht geprüft werden kann, ob die Elektrode vollständig mit Lack bedeckt ist. Des weiteren ist der Lack sehr dünnflüssig, so daß nach Einsprühen des Pt-Drahtes, der zuvor mit Aceton gereinigt wurde, und mehrmaligem Schütteln, die am Pt-Draht herunterlaufenden Lacktropfen eine dünne Schicht zurücklassen. Außerdem ist der Photolack sehr hart, so daß beim Einführen des Pt-Drahtes in das Pt-Röhrchen die Wahrscheinlichkeit gering ist, die Lackschicht dabei zu beschädigen. Von Nachteil ist, daß der Pt-Draht nach dem Beschichten nicht mehr gebogen werden darf, da sonst die Lackschicht sofort absplittern würde. Vor dem Einführen der Innenelektrode wurde das Pt-Röhrchen unter Verwendung einer Einspritzespritze mit Epoxyklebstoff gefüllt. Um die Viskosität des Klebstoffes zu reduzieren, wird der Klebstoff vor dem Einspritzen mit einem Heißluftföhn erwärmt. Diese Phase der Elektrodenherstellung stellt einen kritischen Moment dar. Ein Erwärmen des Klebstoffes verringert nicht nur dessen Viskosität, sondern auch erheblich die Verarbeitungszeit, so daß nach dem Einspritzen des Klebstoffes das Einführen des Pt-Drahtes möglichst schnell erfolgen muß!



Abb. 3.10: Grundbestandteile der 200 μm Zylinder-Elektrode: links Pt-Röhrchen, Mitte Polyimid-Röhrchen und rechts 200 μm Pt-Draht. Der Abstand zweier Striche beträgt 1 mm.

Leider wurden die auf diese Weise hergestellten Elektroden nach kurzer Zeit unbrauchbar, weil das im Elektrolyt enthaltene Ethanol die Isolierung mit zunehmender Verwendungsdauer mehr und mehr auflöste. Schließlich war das Herauslösen des Lackes so weit fortgeschritten, daß sich Innen- und Außenelektrode berührten. Durch diesen Kurzschluß wurde die Elektrode unbrauchbar.

Im weiteren Verlauf der Arbeit wurde eine weitaus bessere Methode gefunden, Innen- und Außenelektrode gegeneinander zu isolieren: Über die Innenelektrode wird ein Polyimid-Röhrchen mit einem Innendurchmesser von 250 μm und einem Außendurchmesser von 300 μm gestülpt, wie

im links stehenden Bild gezeigt. Wie die Abmessungen des Polyimid-Röhrchens zeigen, findet das Ganze dann noch im Pt-Röhrchen Platz. Als Isolationsmaterial wird Polyimid verwendet, weil es chemisch sehr beständig ist, vor allem in organischen Lösungsmitteln und starken Säuren.

Das Pt-Röhrchen wird mit einer Länge von 20 cm geliefert. Für die Elektrodenherstellung muß ein ca. 15 mm langes Stück davon abgetrennt werden, wofür eine Kristallsäge verwendet wird. Auf diese Weise wird verhindert, daß das Röhrchen bei der Durchtrennung gequetscht wird. Hingegen kann nicht verhindert werden, daß die Öffnung des Röhrchens nach dem Sägevorgang verschlossen ist. Mit dem Pt-Draht selbst gelingt es aber, das Röhrchen wieder zu öffnen.

Vor dem Einführen des mit Polyimid isolierten Pt-Drahtes wird ein 0.5 mm lackisolierter Kupferdraht ans Ende des Pt-Röhrchens gelötet. Um die Lackschicht am Ende des Kupferdrahtes zu entfernen wird so verfahren wie bei der weiter oben beschriebenen Herstellung der Platindraht-Elektrode. Beim Anlöten des Kupferdrahtes ist besondere Vorsicht geboten, damit kein Lot in das Innere des Pt-Röhrchens läuft!

Um das Polyimid-Röhrchen im Pt-Röhrchen zu fixieren, wird es mit einem lösungsmittel- und säurebeständigen Lack bestrichen. Dann wird die Innenelektrode in das Röhrchen eingesetzt und durch hin- und herbewegen dafür gesorgt, daß sich der Lack gleichmäßig verteilt. Auf diese Weise wird nicht nur eine Fixierung des Polyimid-Röhrchens erreicht, sondern auch, daß sich die Innenelektrode ungefähr in der Mitte des Pt-Röhrchens befindet.

Als Elektrodenkörper werden Pipettenspitzen der Firma *Sarstedt* verwendet⁶. Dies hat folgende Gründe: Pipettenspitzen weisen für den Einsatz am Leckstrom-Scanner eine nahezu ideale Geometrie auf. Oben sind sie weit, was günstig für die Montage am xyz-Manipulator ist, und nach unten hin laufen sie spitz zu. Daß die Elektroden unten spitz sind, ist für die Charakterisierung der Elektroden (Abschnitt 4.2) von Vorteil, da so in einfacher Weise eine Justierung der Elektrode über der Teststruktur (Kratzer) möglich ist. Zudem werden die kupfernen Anschlußdrähte im Inneren der Elektrode vor dem chemisch sehr aggressiven Elektrolyt geschützt, und letztlich besteht aufgrund der Rotationssymmetrie nicht die Gefahr, daß im Fehlerfall Drehmomente auf den Piezokopf einwirken können, wie im Abschnitt 3.2.2 näher beschrieben.

Bevor die Pipettenspitze nahezu vollständig mit Epoxyklebstoff gefüllt wird, wird der vorderste Teil der Spitze abgeschnitten, da sonst das Pt-Röhrchen nicht durch die Öffnung paßt. Um sicherzustellen, daß kein Elektrolyt in das Innere der Elektrode eindringen kann, wird der Klebstoff mit einem Stück Draht so weit in die Spitze gedrückt, daß der Klebstoff aus der Öffnung austritt. Vor dem Einsetzen der Elektrode werden die Anschlußdrähte mit Schlaufen versehen, die neben

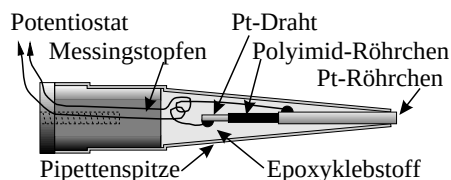


Abb. 3.11: In Pipettenspitze montierte Zylinder-Elektrode. Erläuterungen siehe Text.

⁶Die Pipettenspitzen wurden dankenswerterweise vom Institut für Allgemeine Mikrobiologie der Universität Kiel zur Verfügung gestellt.

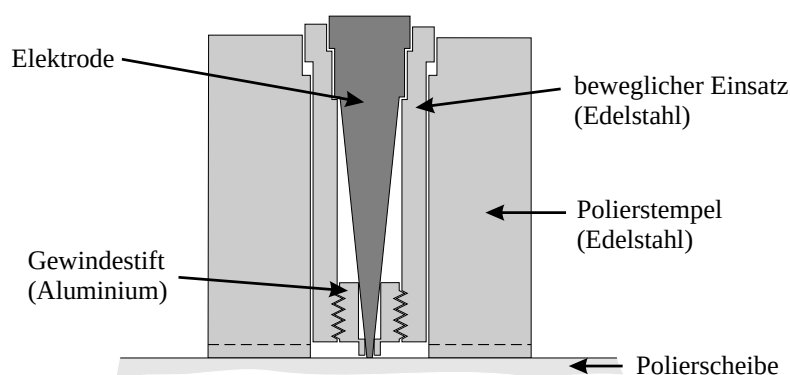


Abb. 3.12: Polierstempel zum Nachbearbeiten der Elektrodenoberfläche.

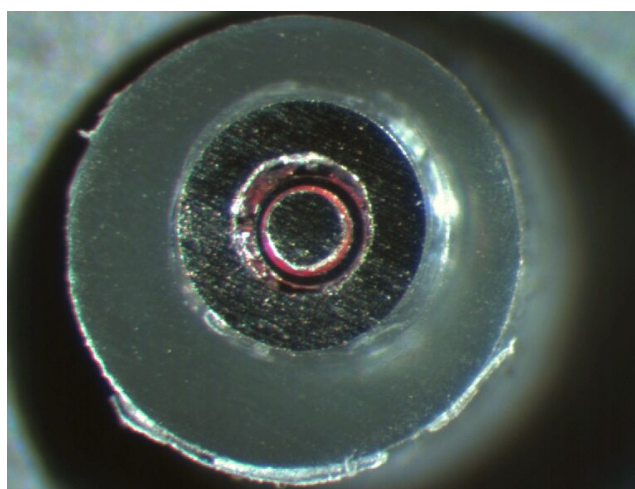


Abb. 3.13: Ansicht der fertig bearbeiteten 200 μm Zylinder-Elektrode. Gut zu erkennen ist das beide Elektrodenanteile isolierende Polyimid-Röhrchen.

dem Klebstoff selbst als zusätzliche Zugentlastung fungieren. Passend zur Pipettenspitze wird ein Messingstopfen angefertigt. Dieser Stopfen enthält ein Innengewinde über das die Montage an den xyz-Manipulator erfolgt. Zusätzlich ist der Stopfen teilweise abgeflacht, so daß die Anschlußdrähte, wie in Abb. 3.11 gezeigt, herausgeführt werden können, ohne gequetscht zu werden.

Nach dem Aushärten des Klebstoffes wird die Elektrode poliert. Hierfür wurde eigens ein Polierstempel angefertigt, da ein Polieren mit freier Hand kaum möglich ist. Der Polierstempel in Abb. 3.12 garantiert, daß die Elektrode während des Polierens stets senkrecht zur Polierscheibe orientiert ist. Die beiden gestrichelten Linien in Abb. 3.12 stehen für mehrere in die Unterseite des Stempels eingebrachte Nuten, die dafür sorgen, daß genügend Polierflüssigkeit zur Elektrode gelangt. Ein Abschneiden des vorderen Teils der Pipettenspitze ist für die 50 μm Zylinder-Elektrode, deren Herstellung im nächsten Abschnitt beschrieben wird, nicht erforderlich. Darum wurde der untere Teil der Pipettenführung als Gewindestift realisiert. Je nach zu polierender

Elektrode hat die untere Öffnung des Gewindestiftes einen anderen Durchmesser. Auf diese Weise wird verhindert, daß sich die Elektrode während des Polierens parallel zur Oberfläche der Polierscheibe bewegen kann.

Da Platin ein sehr weiches Metall ist, gestaltet sich der Poliervorgang als ausgesprochen schwierig. Hier müssen höchste Anforderungen an Sauberkeit und Sorgfalt gestellt werden. Polierplatten, mit denen zuvor beispielsweise Silizium poliert wurde, können nicht verwendet werden. Stattdessen werden mit neuen Filzen versehene Polierplatten verwendet. Begonnen wird mit nassem SiC-Schleifpapier 2500er Körnung. Dann folgt jeweils ein Polieren unter Verwendung von 10 μm , 6 μm , 3 μm sowie 1 μm Diamant-Suspension, wobei das Eigengewicht des Einsatzes als Anpreßdruck dient. Zwischen zwei Poliervorgängen werden alle Bestandteile des Polierstempels einschließlich der Elektrode sehr gründlich unter deionisiertem Wasser abgespült.

Trotz aller Bemühungen gelang es nicht, die Elektrodenoberfläche frei von Kratzern zu bekommen. Stets durchzogen mehrere Kratzer die Elektrodenoberfläche. Vermutet wird, daß sich im Verlauf des Poliervorganges kleine Epoxydharz-Partikel lösen und die Oberfläche zerkratzen. Da der getriebene Aufwand in keinem Verhältnis zum Nutzen stand, wurde schließlich anders verfahren: Statt Diamant-Suspensionen zu verwenden, wird die Elektrodenoberfläche lediglich mit SiC-Schleifpapier 4000er Körnung unter Verwendung von viel Wasser nachbearbeitet.

In einem letzten Schritt werden einige Zehntelmillimeter des untersten Teils der Pipettenspitze mit einem Skalpell entfernt. Der Grund für diesen Schritt wird in Abschnitt 3.3.2 erläutert werden. Ein Bild der fertigen 200 μm Zylinder-Elektrode zeigt Abb. 3.13.

3.2.5.3 Die 50 μm Zylinder-Elektrode

Ohne Zweifel verbessert sich die Ortsauflösung, wenn sich die Abmessungen der elektrochemisch aktiven Fläche der Elektrode verkleinern. Aus diesem Grund wurde eine zweite Zylinder-Elektrode gebaut, deren Innenelektrode aus einem 50 μm dicken Pt-Draht besteht. Damit die Außenelektrode in gewünschter Weise den Verlauf des elektrischen Feldes und damit den Einzugsbereich der Innenelektrode beeinflusst, müssen auch die Abmessungen der Außenelektrode verkleinert werden. Daher wird als Außenelektrode ein Pt-Röhrchen mit einem Innendurchmesser von 100 μm und einem Außendurchmesser von 400 μm verwendet. Selbst wenn Polyimid-Röhrchen mit geeigneten Abmessungen, wie sie für die Herstellung der 50 μm -Elektrode benötigt würden, erwerbbar wären, könnten sie dennoch nicht verwendet werden, weil es kaum gelingen würde alle drei Teile, wie in Abb. 3.10 gezeigt, ineinander zu setzen. Dafür ist der Pt-Draht zu weich.

Glücklicherweise konnte ein 50 μm Pt-Draht beschafft werden (*Goodfellow*), der bereits mit 5 μm Teflon beschichtet ist! Von dem Pt-Röhrchen, das mit einer Länge von 10 cm geliefert wird, wird mit der Kristallsäge ein ca. 15 mm langes Stück abgesägt. Auch hier ist nach dem Sägen die Öffnung des Röhrchens verschlossen. Um sie wieder zu öffnen, wird ein 80 μm dicker Wolframdraht verwendet, dessen unmittelbares Ende mit einer scharfen Schere abgeschnitten wird. Nachdem das Pt-Röhrchen mit Klebe-



Abb. 3.14: Ansicht der fertig bearbeiteten 50 μm Zylinder-Elektrode.

band auf einer massiven Unterlage fixiert worden ist, wird mit Hilfe einer Pinzette und eines Stereo-Mikroskops der Wolframdraht in das Pt-Röhrchen geführt und dann vorsichtig Stück für Stück durch das Pt-Röhrchen geschoben. Stößt der Wolframdraht auf Widerstand, wird der Draht mit der Pinzette dicht vor der Öffnung des Pt-Röhrchens fest angefaßt und mit einem sanften Ruck die Verstopfung durchstoßen. Anschließend wird das Pt-Röhrchen leicht über Papier gerieben, um den dabei entstandenen Grat zu entfernen. Unter dem Stereo-Mikroskop wird geprüft, ob die Öffnung frei und ohne Grat ist.

Nach dem Anlöten eines Kupferdrahtes an das Pt-Röhrchen wird der teflonisierte Pt-Draht in genau der gleichen Art und Weise wie der Wolframdraht durch das Pt-Röhrchen geführt. Zur Anbringung des kupfernen Anschlußdrahtes dient auch hier ein auf eine hohe Temperatur erhitzter Lötkolben, um die Teflonbeschichtung am Ende des Pt-Drahtes zu entfernen.

Das weitere Vorgehen entspricht dem bei der Herstellung der 200 μm Zylinder-Elektrode. Das Abschneiden des vordersten Teils der Pipettenspitze ist hier, wie bereits erwähnt, nicht erforderlich. Beim abschließenden Polieren der Elektrodenfläche muß, bedingt durch die um den Faktor drei kleinere Fläche, mit besonderer Vorsicht vorgegangen werden. Ein Bild der fertigen 50 μm Zylinder-Elektrode zeigt Abb. 3.14.

3.2.5.4 Die Scheiben-Elektrode

Damit die Bedingung $D \gg d$ auch für größere Elektrodenabstände d gilt, muß der Elektrodendurchmesser D vergrößert werden. Eine Möglichkeit dies zu erreichen besteht darin, die 200 μm Zylinder-Elektrode zusätzlich mit einer Platinscheibe zu versehen, die elektrisch mit dem Pt-Röhrchen verbunden ist. Für diesen Zweck wurde eine Pt-Folie (25 mm $\times$ 25 mm $\times$ 0.5 mm) beschafft, und daraus Pt-Scheiben mit einem Durchmesser von 4 mm und einer mittigen Bohrung von 0.7 mm angefertigt.

Bevor ein ca. 15 mm langes Stück vom Pt-Röhrchen abgesägt wird, wird die

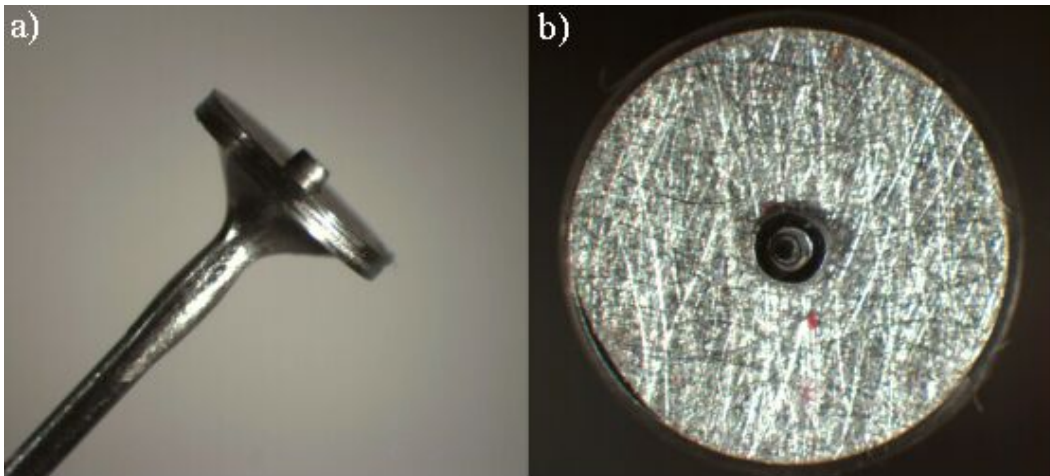


Abb. 3.15: a) Mit Epoxyklebstoff am Pt-Röhrchen befestigte Pt-Scheibe. Die leicht verdickte Stelle am Röhrchen zeigt den aufgetragenen Silberleitlack zwecks Kontaktierung beider Teile. b) Ansicht der fertig bearbeiteten Scheiben-Elektrode.

Pt-Scheibe daran befestigt, wobei sichergestellt werden muß, daß beide Teile elektrisch miteinander in Kontakt sind: Das Pt-Röhrchen wird auf einer massiven Unterlage so mit Klebeband fixiert, daß dessen Ende über den Rand der Unterlage hinausragt. Das Ende des Pt-Röhrchens wird rundherum mit Epoxyklebstoff versehen. Wichtig ist hierbei, möglichst wenig Klebstoff aufzubringen, so daß beim anschließenden Aufsetzen der Pt-Scheibe nicht die ganze Fläche mit Klebstoff benetzt wird. Die erste Verklebung soll gewährleisten, daß die Stelle zwischen Röhrchen und Scheibe abgedichtet wird. Beim Aufsetzen der Pt-Scheibe ist zum einen darauf zu achten, daß sie möglichst senkrecht zum Pt-Röhrchen steht und zum anderen darauf, daß die Scheibe nicht bündig mit dem Ende des Röhrchens abschließt, sondern daß das Röhrchen einige Zehntelmillimeter hervorsteht! Die auf dem Pt-Röhrchen befestigte Pt-Scheibe zeigt Abb. 3.15 a).

Nach dem Aushärten des Klebstoffes werden beide Teile kontaktiert. Hierfür wird Silberleitlack (*Conrad Electronic GmbH*) verwendet. Damit die Scheiben-Elektrode später nicht in ihrer Stabilität beeinträchtigt wird, wird der Kontakt nur an einer Stelle hergestellt. Nach dem Trocknen des Silberleitlackes wird mit einem Ohmmeter geprüft, ob beide Teile leitend miteinander verbunden sind. Anschließend wird der Bereich zwischen Scheibe und Röhrchen reichlich mit Epoxyklebstoff vergossen.

Die weitere Vorgehensweise ist mit der Herstellung der 200 μm Zylinder-Elektrode identisch. Beim Anlöten des Kupferdrahtes an das Pt-Röhrchen muß allerdings beachtet werden, daß der Epoxyklebstoff dabei möglichst wenig thermisch belastet wird. Um die thermische Belastung der Verklebung gering zu halten, wird die Platinscheibe während des Lötens in ein Wasserbad getaucht. Das abschließende Polieren der Elektrode muß ohne den Polierstempel durchgeführt werden. In Abb. 3.15 b) ist die fertige Scheiben-Elektrode abgebildet.

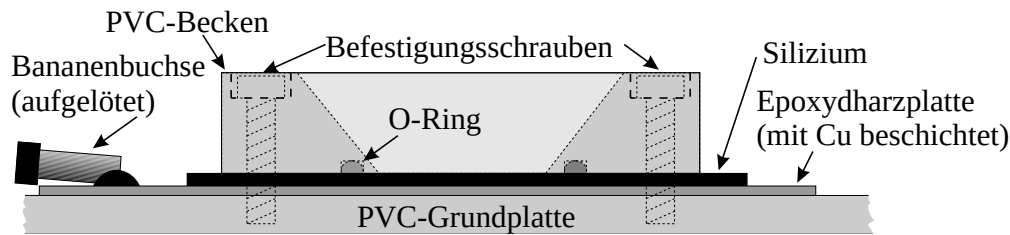


Abb. 3.16: Schematischer Aufbau der Meßzelle. Die Bananenbuchse wurde beußt flach auf die Epoxydharzplatte gelötet, damit sie und die eingesteckte Zuleitung zum Potentiostaten kein Hindernis für die Elektrode darstellen.

3.2.6 Die elektrochemische Zelle

Als Meßzelle dient ein PVC-Becken, das Flächen-’Scans’ von maximal $20\text{ mm} \times 20\text{ mm}$ erlaubt, siehe Abb. 3.16. Der ohmsche Rückseiten-Kontakt besteht aus einem Indium-Gallium-Eutektikum (Ga mit 21.3 Gew.% In [MAS92]), welches bei Raumtemperatur flüssig ist. Mit einem Wattestäbchen wird auf das Rückseite der Proben etwas des InGa-Eutektikums aufgetragen, das gut auf Silizium haftet⁷. In ersten Versuchen wurde als leitende Unterlage Aluminiumfolie verwendet. Diese zeigte aber bereits nach kurzer Meßdauer starke Degradationserscheinungen (sie zerbröselte regelrecht), so daß stattdessen eine mit Kupfer beschichtete Epoxydharzplatte verwendet wird.

Nicht ganz unproblematisch ist das Dichten der Zelle. Längere Zeit wurden hierfür weiche Schaumstoffmatten der Firma Späh verwendet. Diese erwiesen sich aber für größere Meßdauern als ungeeignet, weil sie sich mit der Zeit mit Elektrolyt vollsaugen, was einen Kurzschluß zwischen Vorder- und Rückseite des Wafers zur Folge hat.

Eine Neukonstruktion der Meßzelle erlaubt den Einsatz von weichen O-Ringen, die Dichtigkeit auch bei größeren Meßdauern gewährleisten. Hier zeigen sich aber Schwierigkeiten bei Messungen an multikristallinem Material: Multikristalline Wafer müssen zur Beseitigung des Säge-’Damages’ zuvor überätzt werden⁸. Das hat zur Folge, daß die überätzten Waferunebenheiten aufweisen, weil je nach Orientierung der Körner die Ätzrate verschieden ist. Um Dichtigkeit zu erzielen, müssen hohe Anzugsmomente bei der Montage der Meßzelle angewandt werden. Dabei zerbrechen die Wafer gelegentlich. Die Präparation multikristalliner Wafer wird in Abschnitt 5.1 beschrieben.

3.2.7 Die ’Dark-Box’

Da nur der über den Silizium-Elektrolyt-Kontakt fließende Leckstrom gemessen werden soll, gilt es, Photostromanteile in den gemessenen Strömen zu vermeiden. Aus diesem

⁷ ’Haften’ bedeutet, daß es an der Grenzfläche zwischen InGa und Silizium zur Interdiffusion kommt und sich dadurch in Oberflächennähe des Siliziums eine p^{++} bildet. Derartige Kontakte zeigen ohmsches Verhalten, vgl. Abschnitt 1.2.2 [KOR95].

⁸ Der Grund für den Säge-’Damage’ rührt daher, daß multikristallines Silizium im Blockgußverfahren hergestellt wird. Aus dem gegossenen Siliziumblock werden Säulen gesägt, die wiederum in einzelne Wafer zersägt werden.

Grund wurde der Leckstrom-Scanner in einen Abzug gestellt. Die bei Abzügen übliche transparente Frontscheibe wurde durch eine PVC-Platte ersetzt. Bohrungen, die für die elektrischen Zuleitungen eingebracht werden mußten, wurden anschließend mit Schaumstoff verschlossen.

Neben der Lichtabschirmung erfüllt der Abzug eine weitere wichtige Aufgabe: Das Absaugsystem verhindert, daß die empfindlichen Komponenten, besonders Linearverstärker und Piezo-Translator sowie der Experimentator selbst, den stark korrosiven und gesundheitsschädlichen Flußsäuredämpfen ausgesetzt sind.

3.3 Gesamtaufbau des Leckstrom-Scanners

Zusätzlich zu den im letzten Abschnitt beschriebenen Komponenten des Leckstrom-Scanners werden neben einem Computer zur Meßwerterfassung folgende Geräte benötigt:

- eine Strom-Spannungs-Quelle (Keithley 236)
- ein Programmier (HP59501A)
- eine Spannungsquelle (Grundig PN300)
- ein Digital-Multimeter (Keithley 2000)

sowie ein kleines Zusatzmodul, dessen Bedeutung weiter unten kurz beschrieben wird.

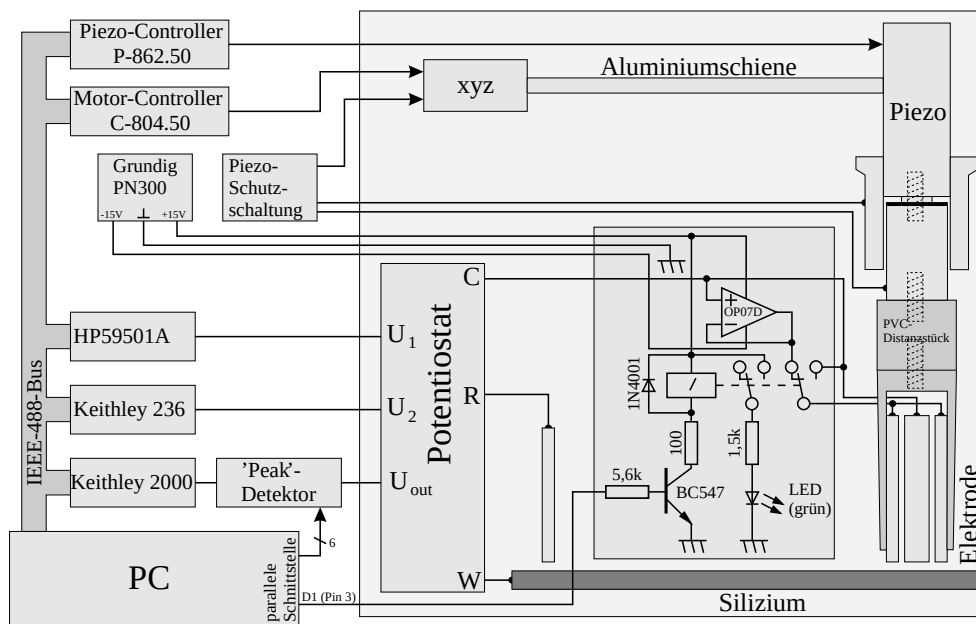


Abb. 3.17: Aufbau des Leckstrom-Scanners. Die Komponenten im hellgrau unterlegten Bereich befinden sich in der 'Dark-Box'. Das Zusatzmodul für die Höhenkalibrierung (grau unterlegte Schaltung) ist am Meßkopf befestigt.

Den Komplettaufbau des Leckstrom-Scanners zeigt Abb. 3.17. Die Linearversteller für die x- und y-Richtung wurden so zusammengebaut, daß die Durchbiegung des oberen Linearverstellers möglichst gering ist. Auf diese Weise werden unnötig lange Fahrwege des z-Translators während einer Messung vermieden. Damit der maximale 'Scan'-Bereich von $10\text{ cm} \times 10\text{ cm}$ zugänglich ist, ist die Elektrode samt Piezo-Translator an einer Aluminiumschiene befestigt. Durch entsprechende Dimensionierung der Aluminiumschiene wurde deren Eigenfrequenz so eingestellt, daß es während der Messung möglichst zu keinen Resonanzphänomenen am Meßkopf kommen kann.

3.3.1 Spannungsfolger

Bei Messung mit einer der Zylinder- oder der Scheiben-Elektrode soll nur der Strom gemessen werden, der über die Innenelektrode fließt. Da das Potential, welches hierbei an die Innenelektrode angelegt wird, auch an der Außenelektrode anliegen muß, gilt es zu verhindern, daß der über die Außenelektrode fließende Strom in den Potentiostaten hineinfließt und somit mitgemessen wird.

Wird die Außenelektrode nicht direkt mit der 'Counter'-Elektrode des Potentiostaten verbunden, sondern über einen Spannungsfolger⁹, kann das oben Verlangte auf einfachste Weise erreicht werden.

3.3.2 Zusatzmodul für die Höhenkalibrierung

Um den Ursprung der z-Koordinate zu bestimmen, wird die Elektrode so lange nach unten gefahren, bis sie auf der Probe aufsitzt. Dieses Ereignis wird als Stromsprung registriert. Die momentane Position des z-Translators bei Auftreten des Stromsprunges wird mit dem Ursprung der z-Koordinate gleichgesetzt. In Abschnitt 3.4 wird diese Höhenkalibrierung näher beschrieben.

Da es nun kaum möglich war, den Leckstrom-Scanner und speziell die Elektroden so zu konstruieren, daß die Elektroden an jedem Ort und zu jeder Zeit exakt parallel zur Oberfläche der Probe orientiert sind, ist das Zusammenschalten von Innen- und Außenelektrode während des oben beschriebenen Vorgehens zwingend erforderlich. Sonst verhindert der Spannungsfolger, daß das Ereignis des Stromsprunges und damit das Aufsitzen der Elektrode auf dem Wafer registriert wird. Eine mögliche Orientierung der Elektrode zur Probe in übertriebener Darstellung zeigt Abb. 3.18.

Unterbleibt die Zusammenschaltung, erfolgt der Stromsprung entweder später, weil die Elektrode zufällig in die richtige Position gebogen wurde oder er erfolgt gar nicht. Im ersten Fall wäre entweder der Elektrodenabstand

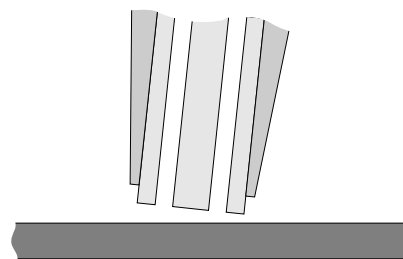


Abb. 3.18: Wahrscheinliche Orientierung der Elektrode über der Probe (stark übertriebene Darstellung).

⁹ Siehe Fußnote Seite 49.

zu klein oder die Elektrode würde während des Scannens über die Probe geschleift, wobei sie eventuell beschädigt werden würde. Im letzteren Fall würde das eine Zerstörung der Elektrode zur Folge haben!

Die Zuleitungen der Elektrode sind mit dem Zusatzmodul verbunden. Da das Modul selbst am Meßkopf befestigt wurde, fungiert es gleichzeitig als Zugentlastung für die Elektrode. Dadurch wird gewährleistet, daß der Potentiostat, der ortsfest ist, während der Messung nicht an den Zuleitungen ziehen und auf diese Weise sowohl die Position der Elektrode als auch ihren Abstand verändern kann.

3.4 Kompensation der Probenschiefelage

Die Anzugsmomente der Zellschrauben und die Langzeitdynamik des Scanners legen fest, wie die Probe zur z-Richtung fehlorientiert ist. Darum muß vor Beginn einer jeden Messung die Fehlorientierung bestimmt werden. Durch Nachführen des Piezo-Translators während einer Messung kann die Fehlorientierung kompensiert und so der Elektrodenabstand konstant gehalten werden. Die Konstanz des Elektrodenabstandes ist besonders wichtig für Messungen mit der Platindraht-Elektrode.

Um die Position $z = 0$ zu finden, wird wie folgt verfahren: Nach Initialisierung des Scanners und dem Anfahren der Startposition bietet das Meßprogramm die

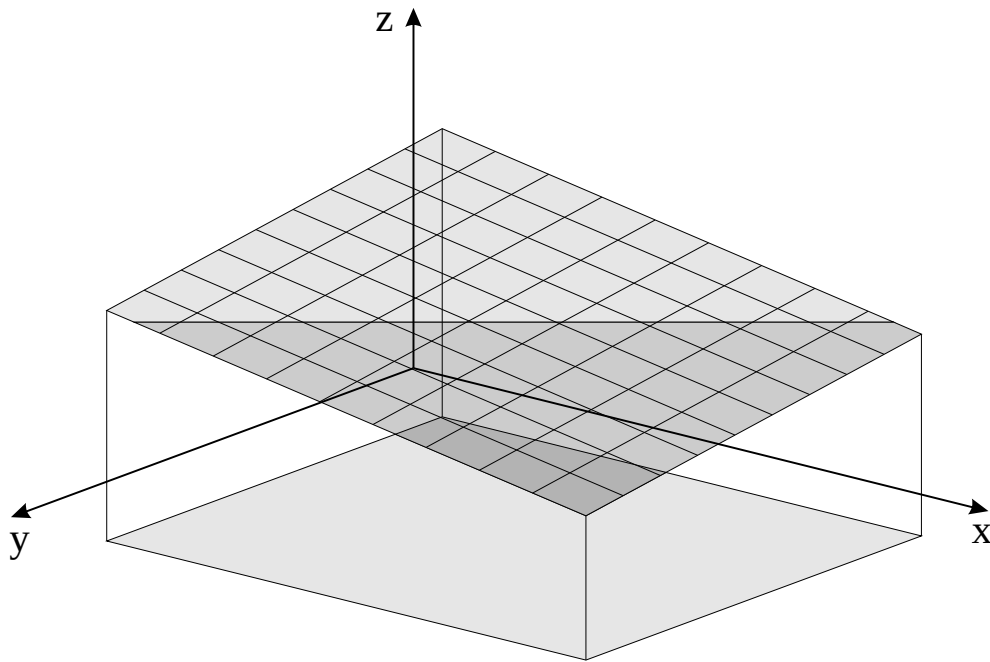


Abb. 3.19: Orientierung der Probe (Fläche mit eingezeichnetem 'Scan'-Raster) zur Fahrebene des x- und y-Linearverstellers (Fläche ohne 'Scan'-Raster). Die Punkte auf der Trennlinie zwischen der hell- und dunkelgrauen Fläche zeigen die Orte, an denen sich der Piezo-Translator in Mittellage befindet.

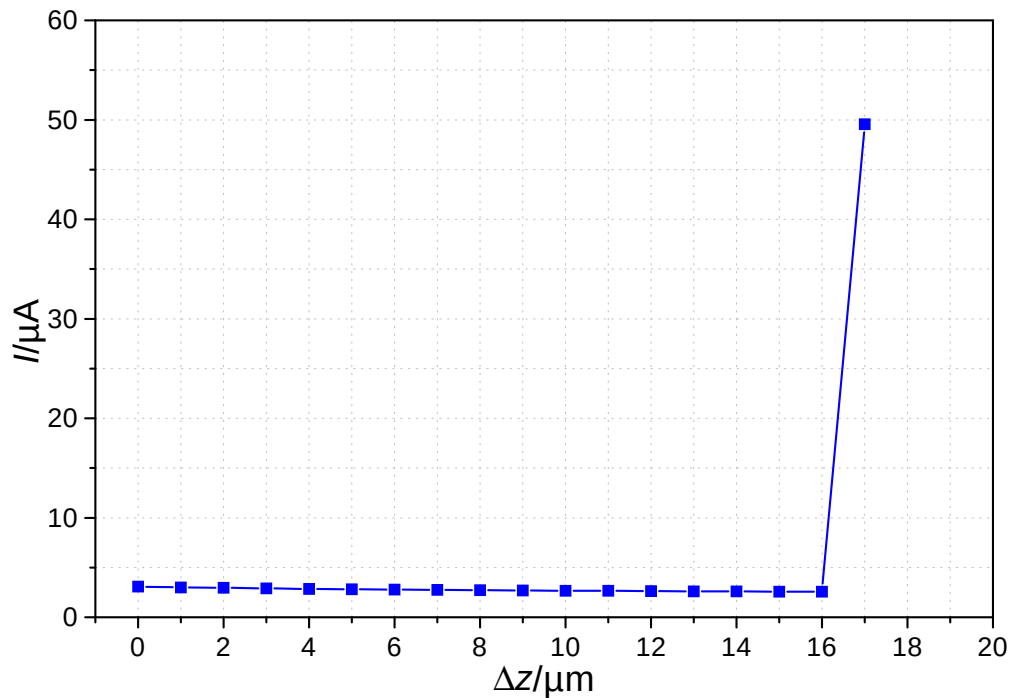


Abb. 3.20: Bestimmung von $z = 0$. Die Elektrode wird in $1 \mu\text{m}$ -Schritten nach unten gefahren. Der Stromsprung signalisiert, daß die Elektrode auf der Probe aufsitzt.

Möglichkeit, die Elektrode neu zu positionieren. Hierbei können für den x- und y-Linearversteller Schrittweiten zwischen 1 mm und $1 \mu\text{m}$ gewählt werden. Weil bei kleinen Schritten in z-Richtung mit dem Piezo-Translator gefahren wird, beträgt die kleinste Schrittweite für den z-Manipulator $0.1 \mu\text{m}$. Nachdem der Linearversteller für die z-Richtung ausgewählt wurde, muß der Abstand zwischen Probenoberfläche und Elektrode abgeschätzt und die richtige Schrittweite eingestellt werden. Dann wird die Elektrode Stück für Stück nach unten gefahren und ggf. die Schrittweite weiter verkleinert. Ist die Schrittweite auf $100 \mu\text{m}$ eingestellt, wird nach jedem Fahren mit einem Finger ganz leicht gegen die Elektrode gedrückt. Läßt sie sich bewegen, wird sie einen Schritt weiter nach unten bewegt, anderenfalls sitzt sie auf der Probe auf und wird wieder einen Schritt nach oben gefahren¹⁰. Nun wird die Schrittweite um den Faktor 10 verkleinert und die obige Vorgehensweise wiederholt. Sitzt die Elektrode erneut auf, wird sie ein oder zwei Schritte weiter nach oben gefahren und das Manipulieren der Elektrode von Hand beendet. Die Position, die die Elektrode jetzt inne hat, ist die neue Startposition, deren Werte in einer Datei gespeichert werden.

Nun wird die erste Ecke der gewählten 'Scan'-Fläche angefahren und an die Zelle eine Spannung von 2.0 V gelegt. Daraufhin beginnt das Programm, die Elektrode selbständig in $1 \mu\text{m}$ Schritten nach unten zu fahren, wobei nach jedem Schritt der Zel-

¹⁰ Das Aufsitzen der Elektrode auf der Probe ist völlig unkritisch für Elektrode, Probe und Piezo-Translator, da die Elastizität im System das Zuweitfahren von maximal $100 \mu\text{m}$ problemlos ausgleichen kann.

lenstrom gemessen wird. In dem Moment, in dem sich Elektrode und Probe berühren, wird der Silizium-Elektrolyt-Kontakt kurzgeschlossen, woraus eine sprunghafte Zunahme des Stromes resultiert, vgl. Abb. 3.20. Der Ort, an dem der Stromsprung erfolgt, markiert den Ursprung der z -Koordinate und damit $z = 0$. Bevor die nächste Ecke angefahren wird, wird die Elektrode in z -Richtung zur Startposition zurückgefahren.

Da die vier Eckpunkte in der Regel nicht auf einer Ebene liegen, wird an sie eine Ebene gefittet, auf der sich die Elektrode während der Messung bewegt. Wurde der Wafer durch Festziehen der Zellschrauben nicht übermäßig verspannt, wird auf diese Weise die Konstanz des Elektrodenabstandes an jedem Ort der Meßfläche gewährleistet, siehe Abb. 3.19. Näheres zum Meßprogramm und zur Kompensation der Proben-schiefelage findet sich in Anhang A.

3.5 Der Elektrolyt

Soll ein wässriger Elektrolyt für Leckstrom-Messungen zum Einsatz kommen, so muß dieser zwangsweise flußsäurehaltig sein: Nach längerer Lagerung an Luft ist die Siliziumoberfläche von einer mehrere Nanometer dicken Oxidschicht bedeckt, die elektrisch sehr gut isolierend ist. Chemisch kann Siliziumoxid nur mit HF aufgelöst werden. Die durch Dissoziation von Wasser stets vorhandenen OH^- -Ionen oxidieren mit der Zeit die Siliziumoberfläche, so daß es nicht möglich ist, nach dem Entfernen des Oxides für die Messungen einen wässrigen, aber flußsäurefreien Elektrolyten zu verwenden.

Damit überhaupt Strom durch die Zelle fließen kann, bedarf es wenigstens einer elektrochemisch aktiven Spezies. Bei wässrigen Elektrolyten sind es die Redoxpaare H^+/H_2 und $\text{O}_2^{2-}/\text{O}_2$ [BAG82, CAM98, PAR51, SEA91]. In sauren Medien werden bei der Elektronentransferreaktion vom Silizium in den Elektrolyten Hydroniumionen reduziert wie im rechts stehenden Bild gezeigt: $4\text{e}^- + 4\text{H}_3\text{O}^+ \rightarrow 4\text{H}_2\text{O} + 2\text{H}_2\uparrow$ und an der Platin-Elektrode H_2O oxidiert: $6\text{H}_2\text{O} \rightarrow 4\text{H}_3\text{O}^+ + \text{O}_2\uparrow + 4\text{e}^-$. Stromfluß hat also zwangsweise Wasserstoffentwicklung an der Silizium- und Sauerstoffentwicklung an der Platin-Elektrode zur Folge.

Besonders an Orten mit erhöhtem Leckstrom können sich somit auf der Meßelektrode schnell größere O_2 -Bläschen bilden. Je mehr O_2 -Bläschen an der Meßelektrode haften, um so mehr verringert sich ihre elektrochemisch aktive Fläche und damit auch der Zellenstrom, falls potentiostatisch gemessen wird. Löst sich das Bläschen von der Meßelektrode ab, wird wieder der maximale Zellenstrom fließen. Das Bilden und Ablösen der O_2 -Bläschen hat also zur Folge, daß das Meßsignal stark verrauscht ist.

Um dem störenden Einfluß der Bläschenbildung entgegenzuwirken, wird der Elektrolyt mit Zusätzen versehen, die die Oberflächenspannung des Elektrolyten reduzie-

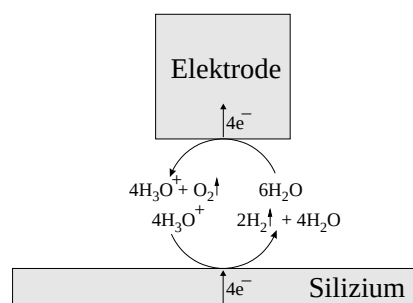


Abb. 3.21: Reaktionen der Redoxpaare H^+/H_2 und $\text{O}_2^{2-}/\text{O}_2$: Am Silizium die Reduktion von Wasserstoff und an der Elektrode die Oxidation von Wasser.

ren. Beispielsweise haben sich für Elymat-Messungen die Zugabe von Ethanol sowie eines Tensids gut bewährt [LIP97]. Folgende Zusammensetzung des Elektrolyten wurde gewählt:

1.0 Vol.% HF	:	CH ₃ CH ₂ OH	:	Tensid
1	:	1	:	1-2 Tropfen auf 100 ml.

In Abschnitt 4.2.1 wird auf die Zusammensetzung des Elektrolyten noch einmal zurückgekommen.

Bei späteren Messungen wurde auf den Ethanolanteil im Elektrolyten verzichtet. Hierfür gab es zwei Gründe: Erstens greift Ethanol, wenn auch langsam, den Epoxylebstoff an, wodurch die erste der hergestellten Scheiben-Elektroden zerstört wurde. Zweitens kann in die Elektronentransferreaktion an der Platin-Elektrolyt-Grenzfläche auch Ethanol involviert sein, in deren Verlauf es zur Schädigung der Siliziumelektrode kommt.

Da sich Ethanol wie eine schwache Säure verhält, existieren im Elektrolyten CH₃CH₂OH⁻-Ionen. Diese Ionen können ihre negative Ladung an die Platin-Elektrode abgeben, wobei sie selbst zu Ethanol-Radikalen werden. Weil Platin ein sehr reaktionsträges Element ist, werden die Radikale nicht mit dem Platin reagieren, sondern mit der Siliziumoberfläche, die auf diese Weise lokal degradiert [MOR77]. An Orten mit Kratzern werden die Ströme und damit auch die Konzentration an Ethanol-Radikalen erhöht sein, die im Bereich des Kratzers mit der Siliziumoberfläche reagieren. Dadurch wird der Kratzer scheinbar breiter, weil in näherer Umgebung die Oberflächenzustandsdichte N_{SS} stark erhöht wird. Wird wiederholt gemessen, äußert sich dies in deutlich verbreiterten Signalen über den Kratzern. In den Abschnitten 4.3 und 5.3 werden hierfür Beispiele gezeigt.

Aber auch mit ethanolfreien Elektrolyten zeigen die Proben nach mehrmaligem Messen eine geschädigte Oberfläche als Folge des Stromflusses. Allerdings sind die Schädigungen der Oberfläche weitaus geringer, als dies für einen ethanolhaltigen Elektrolyten der Fall ist. Der Grund für die stromflußbedingte Korrosion ist folgender:

Minoritätsladungsträger, die die Siliziumoberfläche erreichen, reduzieren, wie in Abb. 3.21 dargestellt, Hydroniumionen. Der reduzierte Wasserstoff wird dabei an der Siliziumoberfläche adsorbiert. Bevor sich H₂ bilden kann, müssen die auf der Oberfläche adsorbierten H-Atome auf ein zweites warten. Genausogut können sie mit H-Atomen reagieren, die Valenzen von Si-Atomen absättigen. Diese Reaktion ist auf Grund der festen Si-H-Bindung unwahrscheinlicher als die Reaktion zweier adsorbierter H-Atome (Volmer-Tafel-Reaktion), findet aber dennoch statt. Bei Reaktion des adsorbierten H-Atoms mit einem H-Atom einer Si-H-Bindung entsteht ein Si-Radikal, das entweder mit einem F⁻- oder mit einem H₃O⁺-Ion reagieren wird. Bei Reaktion mit einem F⁻-Ion würde das Si-Atom, an das Fluor gebunden ist, nach weiteren in Abb. 1.12, gezeigten Reaktionen in Lösung gehen und damit die Oberflächenrauigkeit des Siliziums und letztlich den Grad der Korrosion erhöhen. Dies zeigt, daß neben den Redoxpaaren H⁺/H₂ und O₂²⁻/O₂ auch das Redoxpaar F⁻/F an der Elektronentransferreaktion teilnimmt [BER99, SCH96a].

3.6 Testen der Elektroden

Für eine erste Überprüfung der Elektroden diente eine mit Fotolack beschichtete Kupferplatte wie sie zur Herstellung gedruckter Schaltungen verwendet wird.

Mit einem Skalpell wurde der Fotolack auf einer Länge von ca. 20 mm durchtrennt. Hierfür wurde nur wenig Kraft aufgewendet, um Gratbildung zu beiden Seiten des Kratzers zu vermeiden. Damit sich der z-Manipulator initialisieren kann, wurde mit einem in Aceton getränkten Wattestäbchen der Fotolack zu beiden Seiten des Kratzers entfernt, und zwar derart, daß sich der Kratzer genau in der Mitte zwischen den beiden lackfreien Kupferflächen befindet. Von einem der Enden der Kupferplatte wurde der Fotolack ebenfalls entfernt und zwecks Kontaktierung eine Bananenbuchse aufgelötet entsprechend Abb. 3.16.

Als Elektrolyt wurde 0.5 M KCl verwendet. Das Ergebnis eines 'Line-Scans' über die Kupferplatte mit der 200 μm Zylinder-Elektrode im Abstand von 10 μm zeigt Abb. 3.22.

Die erstaunlich steilen Flanken zeigen an, daß mit der 200 μm Zylinder-Elektrode eine gute bis sehr gute Trennung von Gebieten mit hohem und Gebieten mit geringem Leckstrom möglich sein sollte. Daß der mittlere 'Peak' mit ca. 700 μm viel breiter ist als die erwartete Auflösung von 200-300 μm erklärt sich damit, daß sich der Fotolack in unmittelbarer Nähe zu den lackfreien Stellen mit der Zeit ablöste, vermutlich als Folge des Stromflusses.

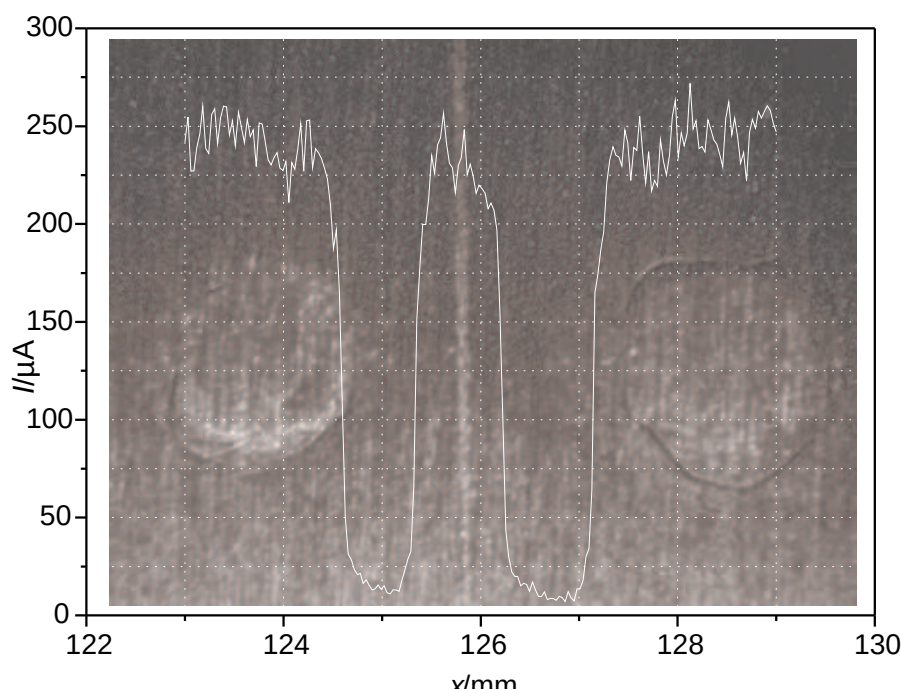


Abb. 3.22: Testen der Elektroden an einer mit Fotolack beschichteten Kupferplatte. An den etwas heller erscheinenden Bereichen wurde der Fotolack entfernt. Messung mit 200 μm Zylinder-Elektrode im Abstand von 10 μm .

Kapitel 4

Messungen an monokristallinem Silizium

Die Anwendung des Leckstrom-Scanners auf monokristalline Siliziumwafer ist weniger interessant, weil der Leckstrom solcher Wafer in der Regel sehr klein und zudem über den Wafer konstant ist. Zum Charakterisieren der Elektroden sind sie allerdings sehr gut geeignet, weil der Leckstrom aus den gezielt aufgebrachten Defektstrukturen im Verhältnis zu dem aus den defektfreien Bereichen hohe und damit gut auswertbare Signal-zu-Untergrund-Verhältnisse mit guter Reproduzierbarkeit liefern sollte.

4.1 Vorbereitungen

Bevor Messungen durchgeführt werden können, müssen einige Vorkehrungen getroffen werden. Wie bereits einleitend erwähnt, ist der Leckstrom monokristalliner Wafer im allgemeinen sehr klein, so daß Testdefekte geschaffen werden müssen, um eine örtliche Dynamik im Leckstrom zu erhalten.

Aus Messung des integralen Leckstromes ergeben sich Informationen darüber, wie die Zellenspannung für 'Line'- und die in Abschnitt 4.3 beschriebenen Flächen-'Scans' zu wählen ist.

4.1.1 Auswahl und Präparation des Probenmaterials

Für alle Messungen wird Bor-dotiertes FZ Silizium (100) mit einen spezifischen Widerstand zwischen $\rho = 2 \cdot 10^{-2} \Omega\text{cm}$ und $\rho = 1 \cdot 10^2 \Omega\text{cm}$ verwendet, das dankenswerterweise von der *Wacker Siltronic AG* zur Verfügung gestellt worden ist.

Da die Wafer in einer mit Kunststoff-Folie versiegelten Waferbox geliefert wurden, ist ein Reinigen vor der Messung nicht erforderlich.

Zum Spalten des Wafers werden ausschließlich ein Diamantgriffel und zwei mit Teflon beschichtete Pinzetten verwendet. Das Spalten des Wafers, das Aufbringen des InGa-Rückseiten-Kontaktes sowie der Einbau in die Zelle erfolgen nur mit Latex-

Handschuhen. Das Verwenden von Handschuhen ist nicht nur wichtig, um die Sauberkeit der Probe zu wahren, sondern auch deshalb, weil das InGa-Eutektikum und vor allem aber die Flußsäure giftig sind.

Nach Einbau des Wafers in die Meßzelle wird das Oxid mit 10 Vol.% HF entfernt. Anschließend wird die HF mit einer Einwegspritze abgesaugt und die Zelle ausgiebig mit deionisiertem Wasser ($\sigma < 0.5 \frac{\mu\text{S}}{\text{cm}}$) gespült. Danach wird die Zelle mit Elektrolyt gefüllt.

4.1.2 Auswahl der Testdefekte

Für eine künstliche, lokale Erhöhung des Leckstromes existieren die folgenden Möglichkeiten:

- Zerstörung der Kristallstruktur
- Erzeugung von Nickel- oder Kupfer-Präzipitaten an der Siliziumoberfläche
- Kratzen

Die Zerstörung der Kristallstruktur wird beispielsweise durch Bombardement mit Schwermetallionen oder Sputtern erreicht. Nicht nur das Präparieren des Wafers würde einen recht hohen apparativen und zeitlichen Aufwand bedeuten, sondern auch die Herstellung der Masken zum Definieren der Defektstrukturen, so daß von dieser Möglichkeit abgesehen wird.

Die Löslichkeit von Nickel und Kupfer in Silizium ist bei hohen Temperaturen relativ hoch, bei Raumtemperatur allerdings sehr klein und beträgt nur wenige Atome pro Kubikzentimeter [IST97]. Mit sinkender Temperatur diffundieren Nickel und Kupfer an die Oberfläche und bilden dort metallische Phasen: NiSi_2 bzw. Cu_3Si . Diese metallischen Ausscheidungen an der Waferoberfläche können eine beträchtliche Erhöhung des Leckstromes zur Folge haben, da sie einen Kurzschluß der Raumladungszone verursachen [BER91, HEL94a, SCH95, WIT98]. Auch das Herstellen dieser Defekte ist apparativ und zeitlich sehr aufwendig. Das Tempern muß in mehreren Schritten erfolgen und beansprucht insgesamt mehr als 12 Stunden [FAL98]. Aufgrund der hohen Temperaturen von bis zu 1000 °C, die hierfür erforderlich sind, ist das Fehlen von Reinraumbedingungen bedenklich.

Kratzen mit einem Diamantgriffel ist die mit Abstand einfachste Methode, eine lokale Leckstrom-Erhöhung zu erreichen. Von Nachteil ist allerdings, daß mit einem Diamantgriffel ein definiertes Kratzen nur sehr schwer möglich ist. So hängt die Struktur des erzeugten Defektes empfindlich davon ab, wieviel Kraft beim Kratzen aufgewandt wird und wie der Griffel zum Wafer orientiert ist. Zudem ist nicht auszuschließen, daß neben dem Kratzer selbst in näherer Umgebung wegen der lokal sehr hohen mechanischen Spannungen auch sogenannte 'Micro-Cracks' entstehen, die ebenfalls zur Erhöhung des Leckstromes beitragen.

Trotz der genannten Nachteile wurden Kratzer als Defektstrukturen gewählt. Da die Charakterisierung aller Elektroden an derselben Struktur erfolgte, war die Schwierigkeit, daß ein Kratzer ohne besondere Vorkehrungen nicht definiert aufgebracht werden kann, ohne jede Bedeutung. Für keine der vier Elektroden war die Ortsauflösung

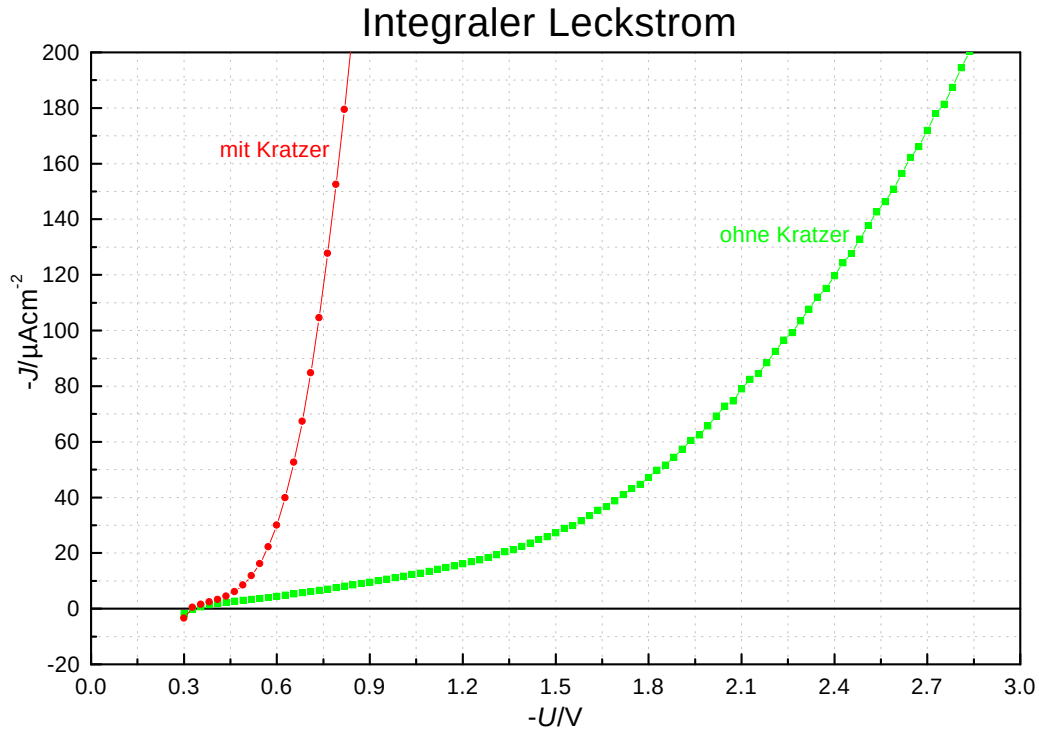


Abb. 4.1: Wirkung eines Kratzers auf den integralen Leckstrom gemessen mit einer Platin-Schleife und Referenz-Elektrode.

bei kleinen Abständen kleiner als erwartet, so daß 'Micro-Cracks', falls sie überhaupt vorhanden waren, sich nicht nachteilig auf die Messungen auswirkten.

4.1.3 Wahl der Zellenspannung

Die erhebliche Wirkung eines Kratzers auf den integralen Leckstrom, der mit einer Pt-Schleife gemessen wird, zeigt die obenstehende Abbildung. Das Signal-zu-Untergrund-Verhältnis ist um so größer, je höher die Zellenspannung ist. Dennoch sollte sie nicht zu hoch gewählt werden, weil die bei höheren Spannungen verstärkte Wasserstoff- und Sauerstoffentwicklung ein stark verrauschtes Meßsignal zur Folge hat. Ebenfalls verstärkt wird bei höheren Spannungen die stromflußbedingte Korrosion des Wafers, die in Abschnitt 3.5 näher beschrieben ist. Abb. 4.1 zeigt, daß bei einer Zellenspannung von $U_{\text{Mess}} \approx -0.8$ V bereits gute Signal-zu-Untergrund-Verhältnisse zu erwarten sind.

Nach dem Anfahren der nächsten Position wird nicht sofort gemessen, sondern erst einige 100 ms gewartet, da der Meßkopf durch das Anfahren in leichte Schwingungen versetzt wird. Um die Entwicklung von Gasbläschen und die Korrosion durch Stromfluß während dieser Phase gering zu halten, wurde bei ersten Messungen nach der Strommessung die Zellenspannung auf ein kleineres kathodisches Potential U_{Bias} gelegt: Die IV-Kennlinie in Abb. 4.1 zeigt, daß bei Spannungen oberhalb von $U = -0.3$ V anodi-

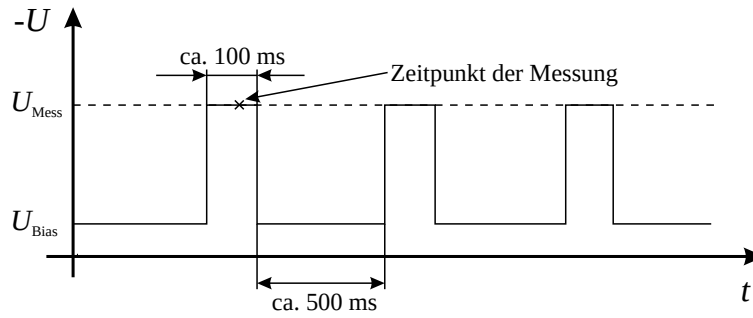


Abb. 4.2: Um die Bildung von H_2 - und O_2 -Bläschen sowie die stromflußbedingte Korrosion des Wafers zu verringern, wird zwischen zwei Messungen die Zellenspannung verringert.

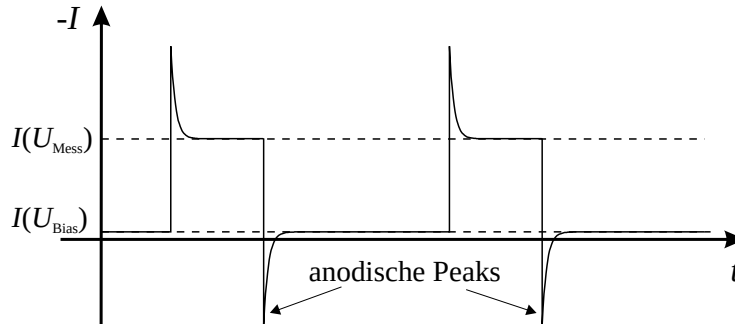


Abb. 4.3: Strom-Zeit-Verlauf für eine gepulste Zellenspannung. Die Strom-*'Peaks'* werden durch das Umladen von Raumladungszone und Helmholtz-Kapazitäten hervorgerufen.

sche Ströme fließen, die PSL-Bildung zur Folge haben (vgl. auch Abb. 1.14). Deshalb wurden nicht 0 V an die Zelle gelegt, sondern eine etwas höhere kathodische Spannung.

Bei der Wahl von U_{Bias} gilt es zu beachten, daß die Wasserstoffterminierung erst nach Stunden abgeschlossen ist. Solange der Prozeß der Passivierung fortschreitet, werden sich H_3O^+ -Ionen an der Oberfläche des Wafers entladen und diese zunehmend positiv aufladen, woraus eine Verschiebung des *'open-circuit'*-Potentials U_{oc} zu höheren kathodischen Spannungen resultiert [BER99]. Aus diesem Grund wird für die Messungen $U_{\text{Bias}} \approx -0.5$ V gewählt.

Die Zeitabhängigkeit der Zellenspannung ist in Abb. 4.2 dargestellt. Wird die Spannung U_{Mess} an die Zelle gelegt, fließt kurzzeitig ein größerer Strom, bedingt durch das Umladen von Raumladungszone und Helmholtz-Kapazitäten. Aus diesem Grund erfolgte die Messung erst 50-100 ms später.

Die Tatsache, daß beim Ändern der Zellenspannung Umladungsströme fließen, war letztlich ausschlaggebend dafür, das Pulsen der Zellenspannung zu unterlassen. Eine Untersuchung des Strom-Zeit-Verlaufes mit Hilfe eines Oszilloskops zeigte, daß das Setzen der Zellenspannung auf U_{Bias} einen anodischen Strom-*'Peak'* nach sich zieht, der die Waferoberfläche korrodieren läßt, siehe Abb. 4.3.

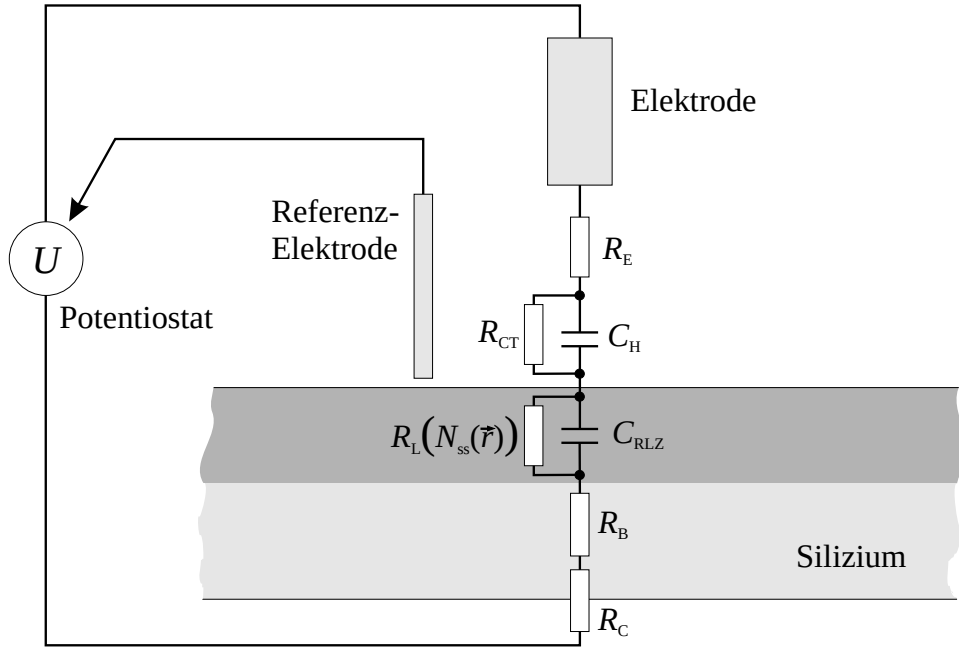


Abb. 4.4: Vereinfachtes Ersatzschaltbild der Silizium-Elektrolyt-Grenzfläche [ROE95, SEA91, WOL86]. Erläuterungen siehe Text.

4.1.4 Referenz-Elektrode

Ein vereinfachtes¹ Ersatzschaltbild der Silizium-Elektrolyt-Grenzfläche zeigt Abb. 4.4. Hierin bedeuten:

- R_C : den Kontaktwiderstand zwischen Silizium und InGa
- R_B : den Widerstand des Bulk-Materials
- C_{RLZ} : die Kapazität der Raumladungszone
- R_L : den Widerstand eines ohmschen Pfades durch die Raumladungszone ('Shunt'), der sich am Ort $\vec{r}$ befindet
- $N_{ss}(\vec{r})$: die Oberflächenzustandsdichte
- C_H : die Kapazität der Helmholtz-Schicht
- R_{CT} : einen Widerstand, der den elektrochemischen Prozeß des Ladungstransfers an der Grenzfläche beschreibt (CT = Charge-Transfer)
- R_E : den Elektrolytwiderstand

Da die über einem Kratzer gemessenen Ströme sehr viel größer sein werden als fernab, ist eine Referenz-Elektrode erforderlich, um die Spannungsabfälle über den Widerständen R_C , R_B und R_L zu kompensieren und damit die Spannung über der Raumladungszone konstant zu halten. Damit die Spannung über der Raumladungszone am

¹Das Ersatzschaltbild berücksichtigt nicht, daß Minoritätsladungsträger für längere Zeiten in oberflächennahen 'Traps' gefangen sein können. Die zugehörigen Relaxationszeiten τ liegen im Millisekunden-Bereich und gewinnen daher erst bei Impedanz-Messungen oberhalb von 1 kHz an Bedeutung [POP00, WOL86].

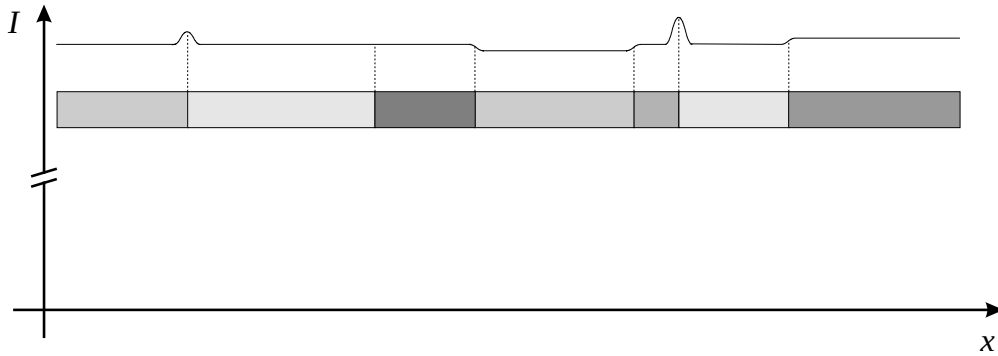


Abb. 4.5: Erwartete Ortsabhängigkeit des Stromes bei einem 'Line-Scan' über einem multikristallinen Wafer. Die einzelnen Grautöne symbolisieren unterschiedliche Kornorientierungen.

Ort der Elektrode konstant ist, müßte die Referenz-Elektrode in die Meßelektrode integriert werden, was ohne weiteres nicht realisierbar ist. Da bei Messungen an Kratzern nur relative Ströme interessieren, ist die Konstanz der Spannung über der Raumladungszone nicht von Wichtigkeit, so daß auf die Verwendung einer Referenz-Elektrode verzichtet werden kann.

Ist die Leitfähigkeit des Elektrolyten hinreichend klein, so werden ohmsche Verluste im Elektrolyten nur einen geringen Einfluß auf die Halbwertsbreite des Strom-'Peaks' über dem Kratzer haben. Deshalb sollte sich der Verzicht der Referenz-Elektrode bei genügend kleiner Elektrolyt-Leitfähigkeit nur unwesentlich auf die Abstandsabhängigkeit des Auflösungsvermögens auswirken, siehe Abschnitt 4.2.3.

Für Messungen an multikristallinem Material kann gänzlich auf die Referenz-Elektrode verzichtet werden. Nicht nur, weil die Ströme um mehrere Größenordnungen niedriger sind, sondern auch weil die Dynamik in den Strömen sehr viel kleiner ist als bei Wafern mit Kratzern, siehe Abb. 4.5.

4.2 Charakterisierung der Elektroden

Abb. 4.6 zeigt eine der ersten Messungen bei der die Lokalisierung eines Kratzers erfolgreich war. Mit Hilfe von 'Line-Scans' über einem Doppelkratzer wird der Einfluß von

- der Flußsäurekonzentration c_{HF}
- der Elektrodenspannung und U_{E}
- dem Elektrodenabstand d

auf das Auflösungsvermögen A und das Signal-zu-Untergrund-Verhältnis S für die vier Elektroden untersucht.

Als Probenmaterial werden Wafer mit einem spezifischen Widerstand von $\rho = 0.02 \, \Omega\text{cm}$ verwendet. Kratzer auf diesem Material ändern ihre Eigenschaften, als

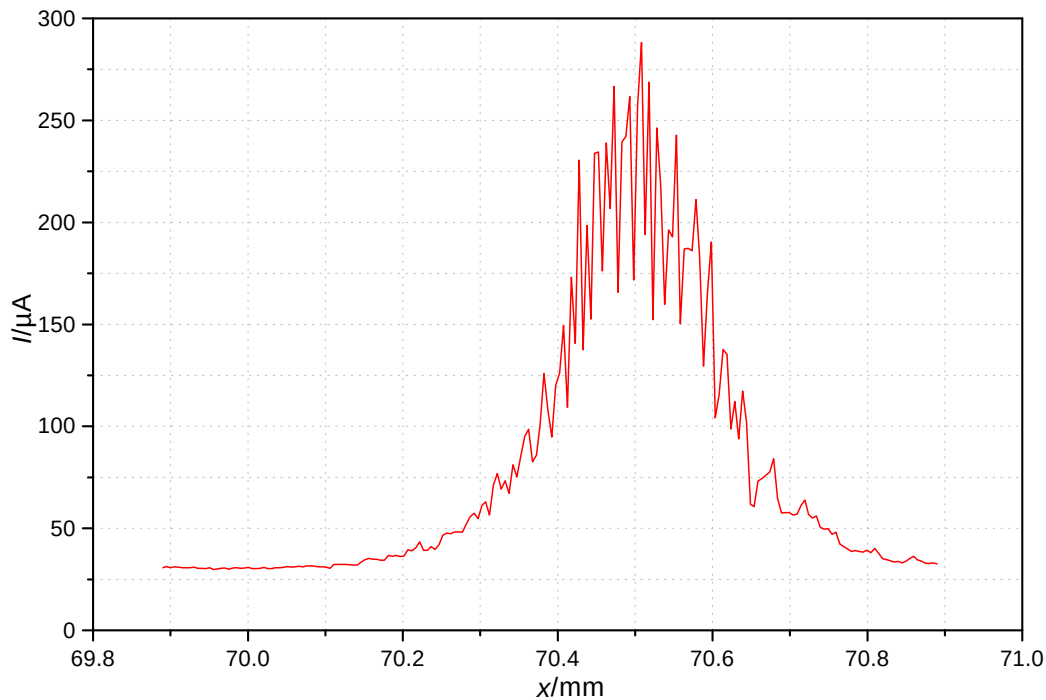
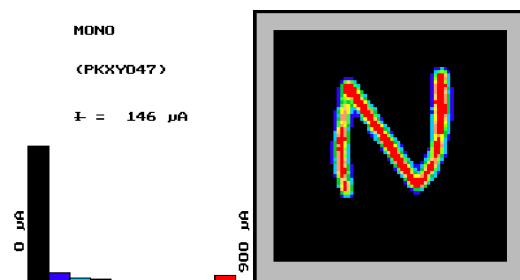


Abb. 4.6: 'Line-Scan' über Kratzer gemessen mit der Platindraht-Elektrode.

Leckstromquelle zu wirken mit der Zeit weniger, als Kratzer auf Wafern mit höherem spezifischen Widerstand. Offenbar resultiert der anfänglich hohe Leckstrom der letztgenannten Wafer allein aus unabgesättigten Bindungen, die durch Kratzen in großer Zahl erzeugt werden. Diese offenen Bindungen werden mit der Zeit mit Wasserstoff abgesättigt, wodurch der Kratzer an Wirkung verliert. Ein Beispiel für die allmähliche Passivierung eines Kratzers auf Material mit 6-10 Ωcm findet sich in Abb. 4.23. Aufgrund der hohen Dotierung des verwendeten Materials scheinen hier Tunneleffekte im Bereich der zerstörten Kristallstruktur zu dominieren, so daß der Kratzer seine Wirkung auch über längere Zeit beibehält².

² Beobachtet wird, daß Kratzer auf diesem Material die Fähigkeit als Leckstrom-Quelle zu wirken mit der Zeit *gar nicht* verlieren. Messungen an Proben mit Kratzern, die viele Monate alt sind, zeigten die Defektstruktur immer noch sehr deutlich!

So zeigt das rechte Bild eine Messung mit der 200 μm Zylinder-Elektrode an einem 'N', das 27 Monate zuvor gekratzt worden war. Die älteren Messungen an diesem 'N' zeigt Abb. 4.25 auf Seite 97.



4.2.1 Einfluß der Flußsäurekonzentration

Unter Verwendung der 200 μm Zylinder-Elektrode wird der Einfluß der Flußsäurekonzentration c_{HF} auf die Höhe der Strom-’Peaks’ I_{max} , den Untergrund I_{off} , das Auslösungsvermögen A sowie das Signal-zu-Untergrund-Verhältnis S untersucht. Für alle ’Line-Scans’ beträgt der Elektrodenabstand 10 μm . Die Ergebnisse sind in Abb. 4.7 dargestellt.

Relevant für die Wahl der HF-Konzentration im Elektrolyten sind die Parameter A und S : Zwar steigen die ’Peak’-Höhen ungefähr proportional mit dem Logarithmus der HF-Konzentration, da aber der Untergrund I_{off} mit zunehmender Konzentration stark anwächst, führt das insgesamt zu einer Abnahme des Signal-zu-Untergrund-Verhältnisses S für $c_{\text{HF}} > 2.3 \text{ Vol.}\%$. Bezüglich des Auslösungsvermögens A findet sich bei $c_{\text{HF}} = 0.4 \text{ Vol.}\%$ ein relativ flaches Minimum.

Für Messungen an monokristallinem Material ist eine gute Ortsauflösung wichtiger als ein hohes Signal-zu-Untergrund-Verhältnis. Dennoch wurde nicht $c_{\text{HF}} = 0.4 \text{ Vol.}\%$ als HF-Konzentration gewählt, sondern $c_{\text{HF}} = 1.0 \text{ Vol.}\%$. Der Grund hierfür ist, daß sich mit steigender HF-Konzentration das Signal-zu-Untergrund-Verhältnis bedeutend stärker erhöht, als sich das Auflösungsvermögen verringert. Im Gegensatz dazu sind für Messungen an multikristallinem Material hohe Signal-zu-Untergrund-Verhältnisse

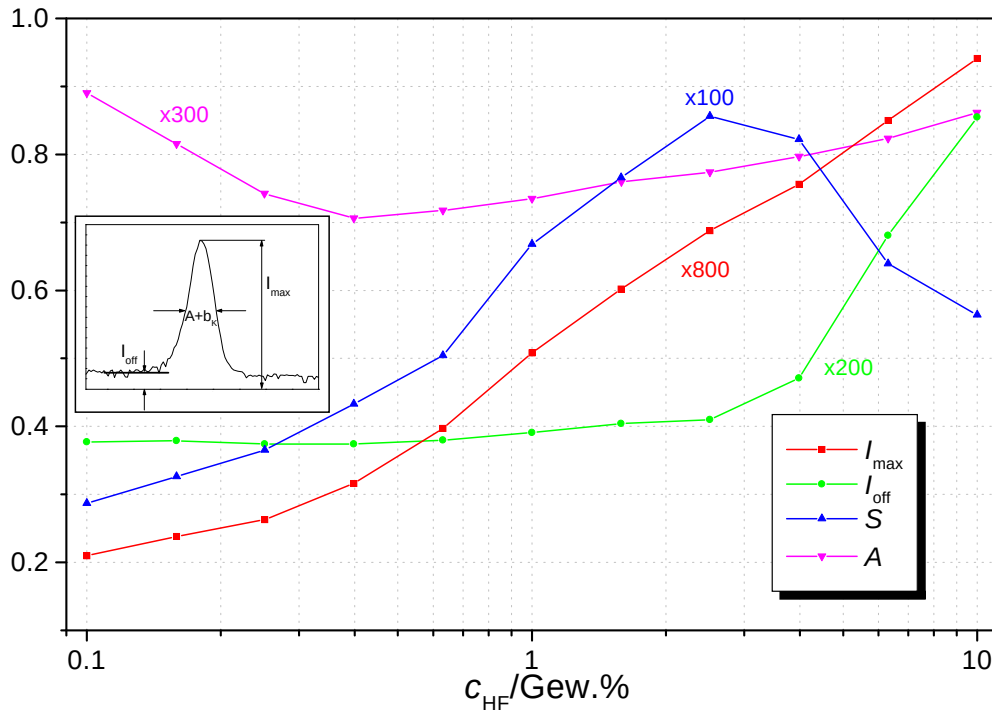


Abb. 4.7: ’Peak’-Höhen I_{max} , Untergrund I_{off} , Auflösungsvermögen A und Signal-zu-Untergrund-Verhältnis S in Abhängigkeit der Flußsäurekonzentration c_{HF} gemessen mit der 200 μm Zylinder-Elektrode. Der Elektrodenabstand beträgt 10 μm . Im Inset bedeutet b_K die Breite des Kratzers.

wichtiger. Für solche Messungen sollte eine HF-Konzentration von $c_{\text{HF}} = 2.3 \text{ Vol.}\%$ verwendet werden, auch wenn das auf Kosten einer geringfügig schlechteren Ortsauflösung geht.

4.2.2 Einfluß der Elektrodenspannung

Um zu untersuchen, welchen Einfluß die Elektrodenspannung U_E auf das Signal-zu-Untergrund-Verhältnis S hat, wird für alle 'Line-Scans' ein Elektrodenabstand von $d = 10 \text{ }\mu\text{m}$ gewählt. Der Grund hierfür ist, daß kleine Elektrodenabstände die besten Signal-zu-Untergrund-Verhältnisse ergeben, wie in Abschnitt 4.2.3 gezeigt wird.

Die Ergebnisse sind in Abb. 4.8 dargestellt. Für alle Elektroden ist etwa eine Mindestspannung von $U_E \approx 3 \text{ V}$ erforderlich, um den aus dem Kratzer kommenden Strom vom Untergrund zu trennen. Oberhalb dieser Spannung steigt das Signal-zu-Untergrund-Verhältnis S für alle Elektroden linear mit der Zellenspannung U_E . Wie erwartet, sind die Signal-zu-Untergrund-Verhältnisse von $200 \text{ }\mu\text{m}$ Zylinder- und Scheiben-Elektrode bei kleinen Zellenspannungen fast identisch, wobei das der Scheiben-Elektrode etwas größer ist. Erst bei größeren Elektrodenabständen werden die Unterschiede deutlicher.

Für die Untersuchungen sind vom Aufbau her folgende Grenzen gesetzt: Mit dem Potentiostaten können Ströme von höchstens $1400 \text{ }\mu\text{A}$ gemessen werden. Der Program-

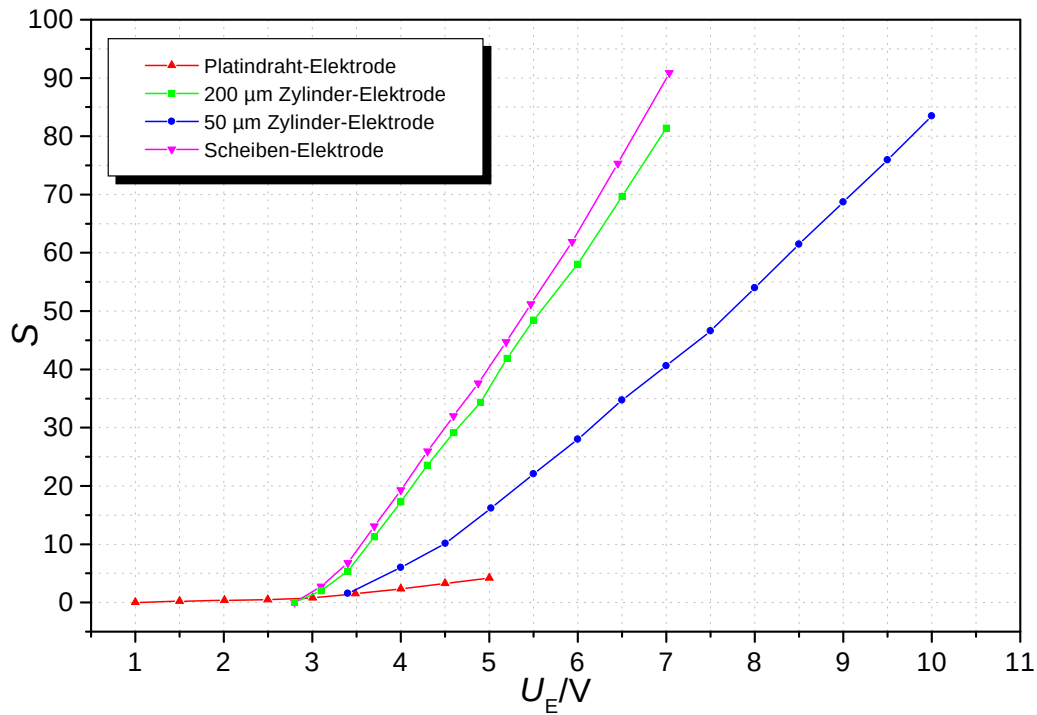


Abb. 4.8: Abhängigkeit des Signal-zu-Untergrund-Verhältnisses von der Elektrodenspannung. Der Elektrodenabstand beträgt $10 \text{ }\mu\text{m}$.

mer HP59501A (vgl. Abb. 3.17) kann maximal eine Spannung von 9.99 V liefern, so daß der Maximalwert für die Elektrodenspannung ebenfalls $U_E = 9.99$ V beträgt. Das ist der Grund dafür, daß das Signal-zu-Untergrund-Verhältnis der Platindraht-Elektrode nur bis zu einer Spannung von $U_E = 5.0$ V bestimmt werden kann, weil oberhalb dieser Spannung die Ströme den Grenzwert von $I = 1400 \mu\text{A}$ überschreiten, wenn sich die Elektrode über einer der beiden Kratzer befindet. Gleiches gilt für die $200 \mu\text{m}$ Zylinder- und die Scheiben-Elektrode. Für diese Elektroden wird die Grenze oberhalb von $U_E = 7.0$ V überschritten. Entsprechend ist das Limit der Elektrodenspannung der Grund dafür, daß das Signal-zu-Untergrund-Verhältnis der $50 \mu\text{m}$ Zylinder-Elektrode nur bis zu einer Spannung von $U_E = 9.99$ V bestimmt werden kann.

Diese meßtechnischen Grenzen des Leckstrom-Scanners haben nur Bedeutung, falls mit der Platindraht-Elektrode gemessen wird. Für diese Elektrode steigen, aufgrund der fehlenden Fokussierung des elektrischen Feldes die Ströme mit zunehmender Spannung U_E stark, so daß das Limit des Stromes, wie oben beschrieben, bereits für $U_E \approx 5.0$ V erreicht wird. Dieser geringe Wert für die maximal mögliche Zellenspannung U_E sowie die fehlende Feldfokussierung haben zur Folge, daß mit der Platindraht-Elektrode, im Vergleich zu den anderen Elektroden, nur sehr kleine Signal-zu-Untergrund-Verhältnis erreicht werden können, siehe Abb. 4.8.

4.2.3 Einfluß des Elektrodenabstandes

Zwecks Bestimmung der Abstandsabhängigkeit des Auflösungsvermögens wird für alle Elektroden der Elektrodenabstand beginnend mit $d = 10 \mu\text{m}$ erst in kleineren ($10 \mu\text{m}$) und dann in größeren Schritten ($50 \mu\text{m}$) erhöht.

Ist A die Halbwertsbreite des Signals über dem Kratzer, so können, wie die rechts stehende Abbildung verdeutlicht, zwei parallel verlaufende Kratzer mit Leckstrom-Scanner dann aufgelöst werden, wenn ihr Abstand mindestens A beträgt³.

Deshalb wird das Auflösungsvermögen A der jeweiligen Elektrode bei gegebenem Elektrodenabstand d mit der Halbwertsbreite des Strom-‘Peaks’ über dem Kratzer abzüglich der Breite des Kratzers b_K gleichgesetzt, vgl. Inset in Abb. 4.7. Das Ergebnis dieser Untersuchung zeigt Abb. 4.10. Die Breite des Kratzers b_K wird unterm Lichtmikroskop ausgemessen.

Wie erwartet, zeigt das Auflösungsvermögen der Platindraht-Elektrode eine sehr starke Abstandsabhängigkeit, wohingegen die der anderen Elektroden weitaus schwächer sind. Was allerdings überrascht, ist die Tatsache, daß

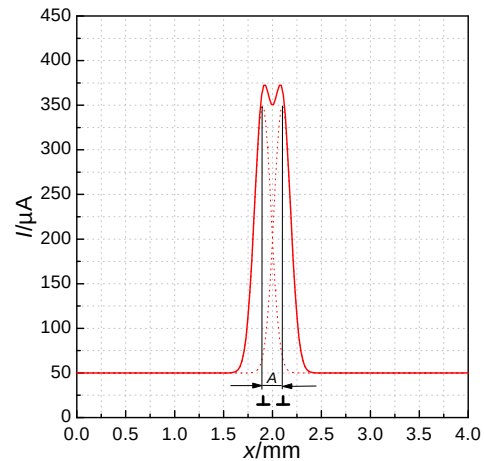


Abb. 4.9: Zur Definition des Auflösungsvermögens A . Erläuterungen siehe Text.

³ Für die Betrachtung werden die Kratzer als unendlich dünn angenommen: $b_K = 0$.

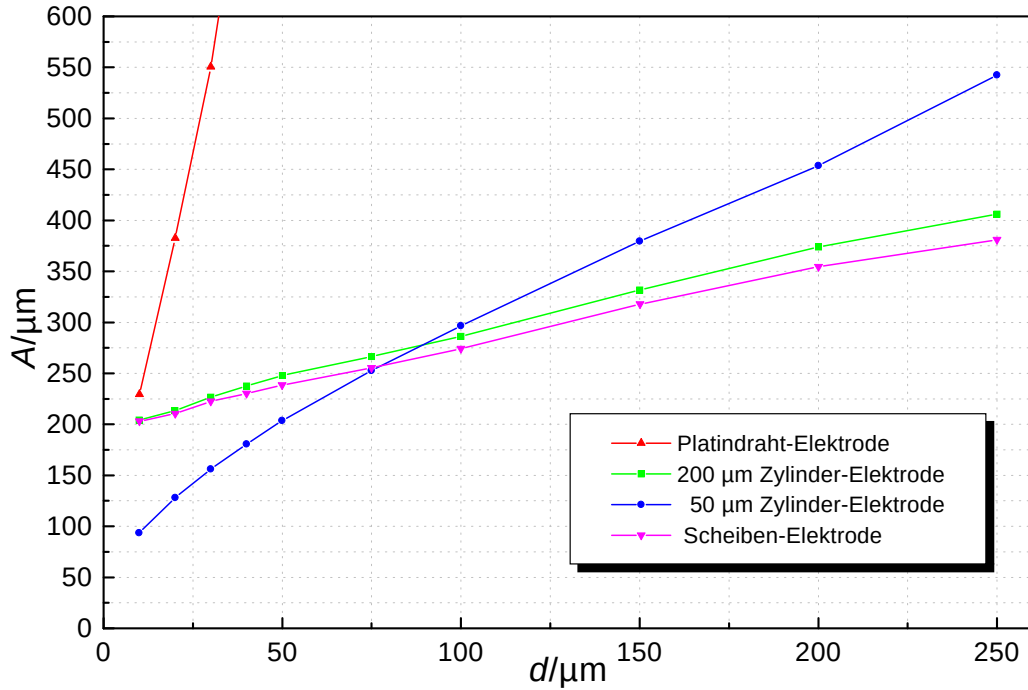


Abb. 4.10: Abhängigkeit des Auflösungsvermögens vom Elektrodenabstand.

die Abstandsabhängigkeit der Scheiben-Elektrode nur geringfügig schwächer ist als die der 200 μm Zylinder-Elektrode. Wie in Abschnitt 3.1 erläutert, wird für die Scheiben-Elektrode erwartet, daß ihr Auflösungsvermögen für nicht zu große Elektrodenabstände nahezu konstant ist. Da die 50 μm Zylinder-Elektrode bezüglich ihrer Abmessungen der Platindraht-Elektrode am ähnlichsten ist, weist ihr Auflösungsvermögen eine stärkere Abstandsabhängigkeit auf als das der 200 μm Zylinder- und der Scheiben-Elektrode.

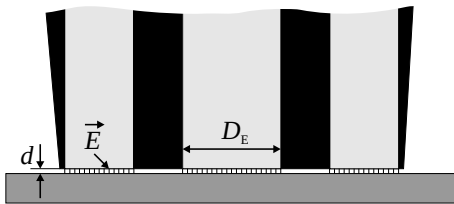


Abb. 4.11: Für $d \rightarrow 0$ ist das Auflösungsvermögen A identisch mit dem Durchmesser der Innenelektrode D_E .

Je dichter sich die Elektrode über der Probe befindet, um so homogener wird das elektrische Feld $\vec{E}$ zwischen Probe und Elektrode sein, wie dies die links stehende Abbildung am Beispiel der 200 μm Zylinder-Elektrode zeigt. Darum muß bei Extrapolation auf $d = 0$ das Auflösungsvermögen A mit dem Durchmesser der Innenelektrode D_E identisch sein:

$$\lim_{d \rightarrow 0} A(d) = D_E.$$

Mit Ausnahme der Platindraht-Elektrode wird dieser Grenzwert von allen Elektroden erreicht.

Ungenauigkeiten im Elektrodenabstand machen sich bei der Platindraht-Elektrode wesentlich stärker bemerkbar als bei den anderen, in diesem Fall etwa 6 μm . Wie in

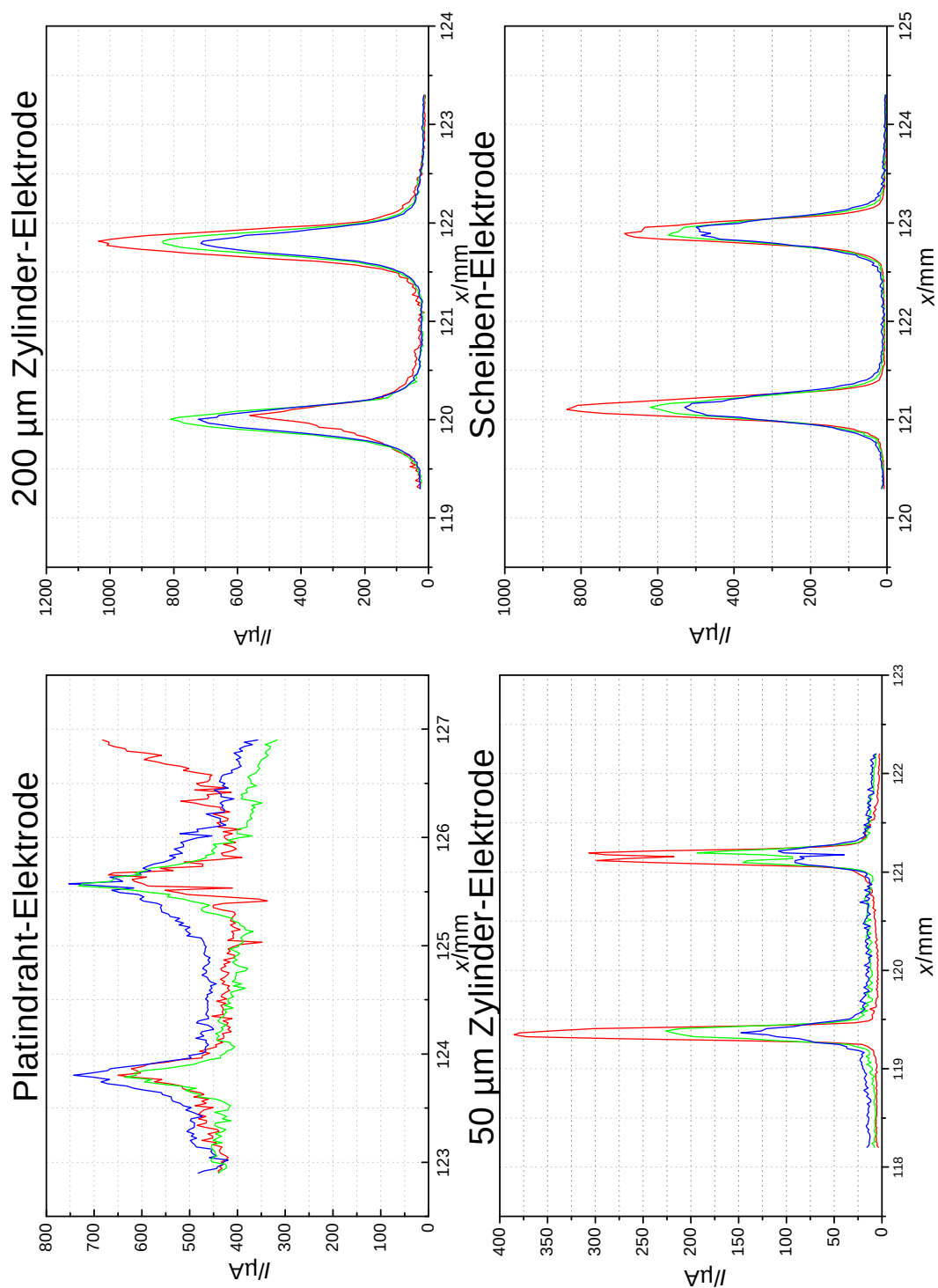


Abb. 4.12: 'Line-Scans' über Doppelkratzer für verschiedene Elektrodenabstände: 10 μm (rot), 20 μm (grün) und 30 μm (blau).

Abschnitt 3.4 näher beschrieben, beträgt die Genauigkeit des Elektrodenabstandes d am Start- und Endpunkt des 'Line-Scans' $\Delta d = \pm 0.5 \mu\text{m}$. Wird beispielsweise beim Montieren der Zelle der Wafer leicht verspannt, ist es durchaus wahrscheinlich, daß der tatsächliche Elektrodenabstand am Ort der Kratzer $6 \mu\text{m}$ kleiner ist als der im Meßprogramm eingestellte.

Eine andere Erklärung ist, daß sich zwischen der Elektrodenspitze und dem Wafer Lackpartikel befunden haben könnten, die während der Initialisierung des z-Translators (siehe Abschnitt 3.4) den Kurzschluß zwischen Elektrode und Wafer verzögerten, was die Einstellung eines zu kleinen Elektrodenabstandes zur Folge haben würde⁴. Das würde auch erklären, daß der Grenzwert zu klein ist.

Abb. 4.12 zeigt eine Zusammenstellung der 'Line-Scans' aller Elektroden für verschiedene Elektrodenabstände. Bei einem Vergleich der Elektroden untereinander zeigt die Platindraht-Elektrode die gravierendsten Unterschiede: Neben einem sehr verrauschten Signal und vergleichsweise großen 'Peak'-Breiten ist auch der Untergrund um ca. zwei Größenordnungen höher als bei den anderen Elektroden. Der Stromanstieg ab der Position $x = 126.6 \text{ mm}$ weist darauf hin, daß die Elektrode die Siliziumoberfläche berührte. Das spricht für die weiter oben gegebene Erklärung, daß bei Extrapolation auf $d = 0 \mu\text{m}$ in Abb. 4.10 für die Platindraht-Elektrode ein zu kleiner Grenzwert erhalten wird.

Die Form der 'Line-Scans' der $200 \mu\text{m}$ Zylinder-Elektrode sind zu denen der Scheiben-Elektrode sehr ähnlich. Dennoch bestehen zwei deutliche Unterschiede: Zum einen sind die Ströme, die mit der Scheiben-Elektrode gemessen werden, kleiner und zum anderen ist mit dieser Elektrode eine noch schärfere Trennung zwischen dem Kratzer und seiner Umgebung möglich. Daß beim 'Line-Scan' mit der $200 \mu\text{m}$ Zylinder-Elektrode im Abstand von $10 \mu\text{m}$ die Ströme über dem Kratzer, welcher sich am Ort $x = 120.0 \text{ mm}$ befindet, kleiner sind als für die anderen Abstände, kann in einfacher Weise erklärt werden: Wie in Abschnitt 3.5 erläutert, können sich H_2 - und O_2 -Bläschen nachteilig auf die Messungen auswirken. So wird des öfteren beobachtet, daß nach einer Messung in einigen Bereichen des Kratzers Bläschen sitzen, erkennbar als feine Bläschenkette. Bleiben sie über längere Zeit darin haften, resultiert daraus bei einer erneuten Messung zwangsläufig ein reduziertes Signal.

4.2.4 Diskussion

Mit größer werdendem Elektrodenabstand war zu erwarten, daß sich der Strom über die Innenelektrode erhöht, da die Fläche, aus der die Elektrode ihren Strom zieht, mit dem Elektrodenabstand steigt. Beobachtet wird aber, daß sich der Strom mit zunehmendem Abstand verkleinert! Dieses Verhalten kann wie folgt erklärt werden:

Abb. 4.13, zeigt wie sich die Potentiale im Stromkreis (vgl. Abb. 4.4) verteilen, wenn die Elektrode vom Kratzer lateral weit entfernt ist und wenn sich die Elektro-

⁴ Da von der $200 \mu\text{m}$ und der $50 \mu\text{m}$ Zylinder-Elektrode die untersten Zehntelmillimeter der Pipettenspitze abgeschnitten wurden, ist das Verpassen von $z = 0$ sowohl für diese, als auch für die Scheiben-Elektrode unmöglich, siehe Abschnitt 3.3.2.

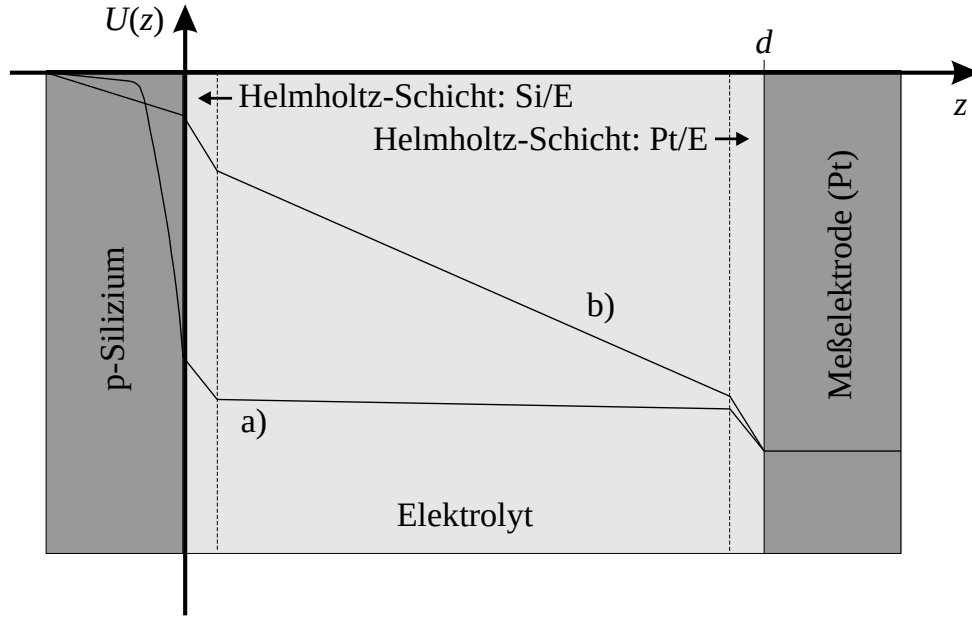


Abb. 4.13: Potentialverteilung im Stromkreis: a) Wenn die Elektrode lateral weit vom Kratzer entfernt ist, und b) wenn sich die Elektrode über dem Kratzer befindet. Nicht berücksichtigt bleibt der Spannungsabfall über dem Kontaktwiderstand R_C . Der spezifische Widerstand von Platin wird null gesetzt. Weitere Erläuterungen siehe Text.

de direkt über einem Kratzer befindet. Im ersten Fall ist die Raumladungszone im Silizium intakt, so daß über ihr der größte Teil der Zellenspannung abfällt. Weitere Anteile der Zellenspannung werden über den beiden Helmholtz-Schichten sowie über dem Elektrolytwiderstand R_E abfallen. Bedingung dafür, daß Strom durch die Zelle fließt, ist, daß die Summe beider Potentialabfälle über den Helmholtz-Schichten größer als die Zersetzungsspannung des Wassers ist:

$$U_{\text{Si/E}} + U_{\text{Pt/E}} > U_Z = 1.23 \text{ V}$$

Bekannt ist, daß die Zersetzung von Wasser oberhalb von 1.23 V thermodynamisch möglich wird. Soll jedoch Strom durch die Zelle fließen, bedarf es nach der Butler-Volmer-Gleichung (Gleichung 1.33) einer etwas höheren Spannung. Die Differenzen $\eta_{\text{Si/E}}$ und $\eta_{\text{Pt/E}}$, die auch als *Überspannungen* bezeichnet werden, dienen dazu, die Kinetik an der Silizium-Elektrolyt- bzw. der Platin-Elektrolyt-Grenzfläche zu treiben [POP00]. Ist die Bedingung erfüllt, wird die Redoxreaktion, wie in Abb. 3.21 gezeigt, ablaufen und Strom durch die Zelle fließen, begleitet von Wasserstoffentwicklung am Silizium und Sauerstoffentwicklung an der Meßelektrode.

Bewegt sich die Meßelektrode über einen Kratzer, werden sich die Spannungsverhältnisse im Stromkreis, wie Abb. 4.13 b) zeigt, völlig ändern. Wie anfangs erwähnt, werden durch Kratzen lokal Kristalldefekte induziert, wodurch an den betroffenen Stellen die Oberflächenzustandsdichte N_{SS} um viele Größenordnungen erhöht wird. Weil es am Ort eines Kratzers praktisch keine Raumladungszone mehr gibt, wird sich das Si-

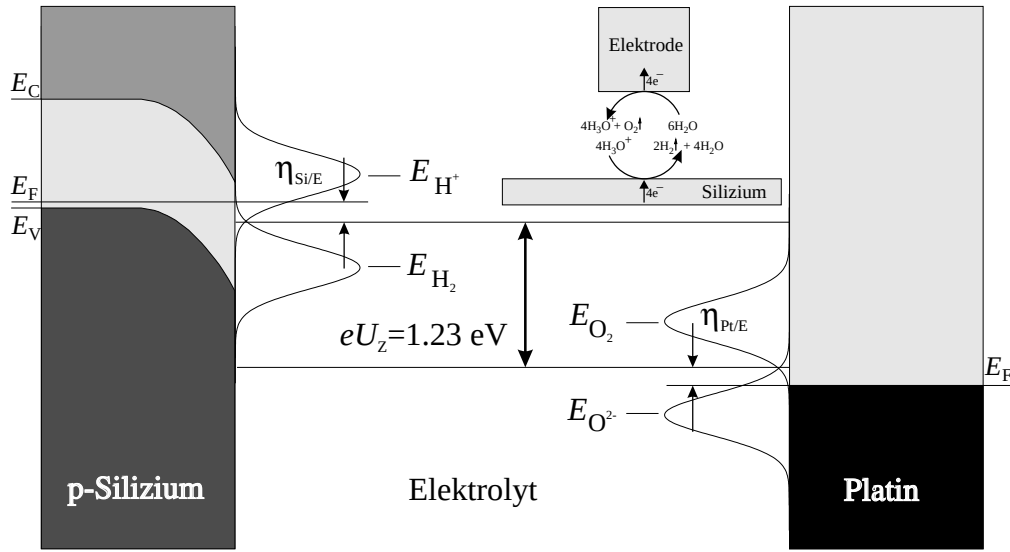


Abb. 4.14: Energiediagramm des Systems p-Silizium/Elektrolyt/Messelektrode.

litzium dort wie ein Metall verhalten. Am Ort eines Kratzers wird im Silizium deshalb ein sehr viel kleinerer Teil der Zellenspannung abfallen als an einem Ort ohne Kratzer. Der Anteil der Spannung, der im Silizium abfällt, wird im wesentlichen bestimmt durch die Spannungsabfälle über dem Kontaktwiderstand R_C und dem Bulkwiderstand R_B , siehe Abb. 4.4. Da nun mehr Spannung im Elektrolyten abfallen muß, werden sich die Spannungsabfälle über den beiden Helmholtz-Schichten und damit die Überspannungen $\eta_{Si/E}$ und $\eta_{Pt/E}$ entsprechend vergrößern. Eine Erhöhung der Überspannungen hat nach der Butler-Volmer-Gleichung eine Vergrößerung des Zellenstromes zur Folge, wodurch auch der Spannungsabfall über dem Elektrolytwiderstand R_E steigt.

Da Flußsäure schwach dissozzierend ist (vgl. Abschnitt 1.3.4), ist die Leitfähigkeit σ_E von 1 Vol.% HF relativ gering⁵. Darum wird über dem Elektrolytwiderstand R_E ein Großteil der Zellenspannung abfallen, wenn sich die Elektrode über einem Kratzer befindet, siehe Abb. 4.13. Wegen

$$R_E(d) = \frac{4}{\pi \sigma_E} \cdot \frac{d}{A(d)^2}$$

vergrößert sich der Elektrolytwiderstand proportional zu $\frac{d}{A(d)^2}$ [GER95, WAL79]. Strenggenommen hat das durchströmte Volumen unter der Elektrode die Form eines Konus. Mit dem Durchmesser der Innenelektrode D_E gilt für alle Elektroden stets: $d \ll D_E$. Mit Ausnahme der Platindraht-Elektrode (keine Feldfokussierung!) kann aus diesem Grund das durchströmte Volumen in guter Näherung durch einen Zylinder mit dem Durchmesser $A(d)$ ersetzt werden.

⁵ Verwendet wird 48 Vol.% HF, die mit deionisiertem Wasser auf 1 Vol.% HF verdünnt wird. Die Leitfähigkeit des Wassers ist ebenfalls sehr gering und beträgt lediglich $\sigma < 0.5 \frac{\mu S}{cm}$.

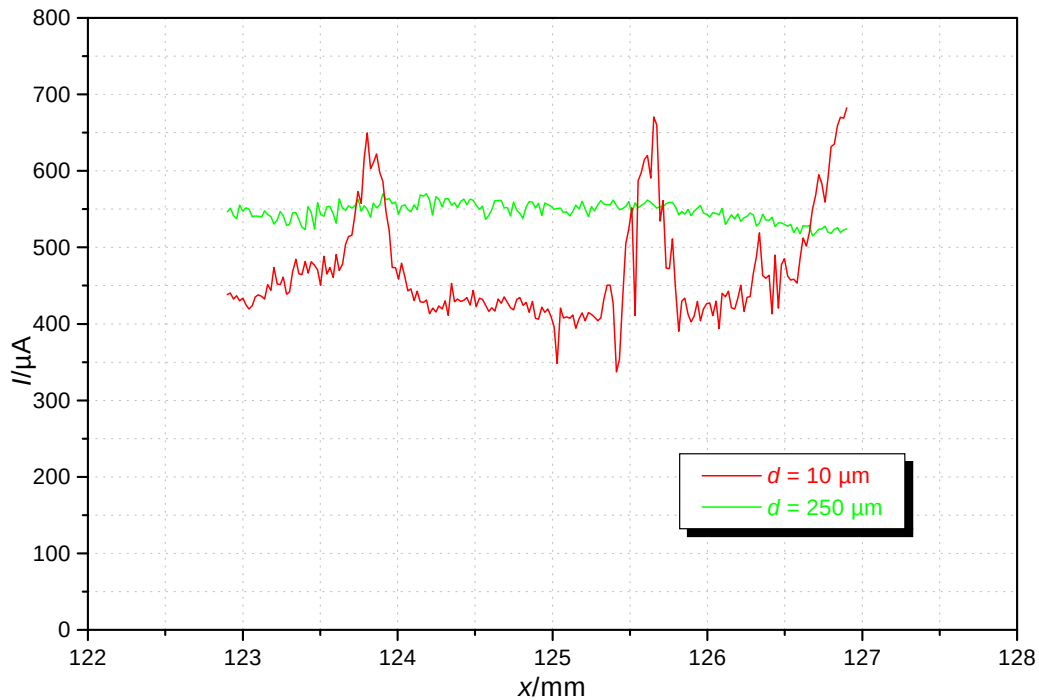


Abb. 4.15: Der Einfluß des Elektrodenabstandes auf den Rauschpegel bei Messung mit der Platindraht-Elektrode.

Da sich das Verhältnis zwischen dem Elektrodenabstand d und $A(d)^2$ für die $50 \mu\text{m}$ Zylinder-Elektrode mit zunehmendem Abstand schneller vergrößert als für die $200 \mu\text{m}$ Zylinder- oder die Scheiben-Elektrode, verringern sich die Ströme am Ort des Kratzers für die $50 \mu\text{m}$ Zylinder-Elektrode, wie Abb. 4.12 zeigt, in stärkerem Maße als für die anderen beiden Elektroden.

Daß die 'Line-Scans' der Platindraht-Elektrode, wie Abb. 4.12 zeigt, im Vergleich zum Signal über dem Kratzer so stark verrauscht sind, hat folgende Ursache: Nur der integrale Leckstrom aus einer hinreichend großen Fläche ist im zeitlichen Mittel konstant. Wird der aus einem kleinen Bereich fließende Strom gemessen, werden zeitliche Fluktuationen sichtbar, die dem weißen Rauschen ähnlich sind [CAR01]. So zeigt Abb. 4.15, daß sich bei Vergrößerung des Elektrodenabstandes der Rauschpegel etwas verkleinert. Daß sich der Rauschpegel nur etwas verkleinert, kann wie folgt erklärt werden:

Die Bildung von H_2 - und O_2 -Bläschen trägt ebenfalls dazu bei, daß das Meßsignal verrauscht ist. Je höher der Strom, um so höher wird auch die Rate der Bläschenbildung sein. Dies gilt aufgrund der festen Elektrodenfläche in besonderem Maße für die Bildung von O_2 -Bläschen an der Meßelektrode. Trotz vergleichbarer Ströme über den Kratzern, bezüglich Platindraht- und Scheiben-Elektrode, ist der Rauschpegel bei Messung mit der Scheiben-Elektrode sehr viel kleiner (praktisch null), als wenn mit der Platindraht-Elektrode gemessen wird. Die Ursache hierfür ist, daß die Messung mit der Scheiben-Elektrode bei höherer Spannung durchgeführt wird. Über einem Kratzer wird

aufgrund der dort erhöhten Stromdichte auch die Rate der Bläschenbildung sehr hoch sein. Die daraus resultierenden hohen Frequenzen des Rauschsignals werden komplett vom Potentiostaten herausgefiltert, siehe Abschnitt 3.2.3.

4.3 Flächen-’Scans’

4.3.1 Auswahl der Teststruktur

Als Teststruktur wird ein gekratztes ’N’ gewählt, siehe Abb. 4.16 a). Die Liniensegmente, die sich unter einem spitzen Winkel schneiden, eignen sich sehr gut, um zu überprüfen, wie gut die einzelnen Elektroden in der Lage sind, die Segmente in Abhängigkeit des Abstandes örtlich zu trennen.

Ein weiterer Vorteil ist, daß das ’N’ in einem Zug gekratzt werden kann. So wäre beispielsweise das Kratzen eines ’A’ mit Schwierigkeiten verbunden, weil das Kratzen des waagerechten Querstriches kaum so zu bewerkstelligen ist, daß die Endpunkte auf den beiden schrägen Linien zu liegen kommen.

Abb. 4.16 b) zeigt einen Flächen-’Scan’ von jenem ’N’, dessen lichtmikroskopische Vergrößerung in a) abgebildet ist. Die Messung erfolgte mit der 200 μm Zylinder-Elektrode im Abstand von $d = 10 \mu\text{m}$.

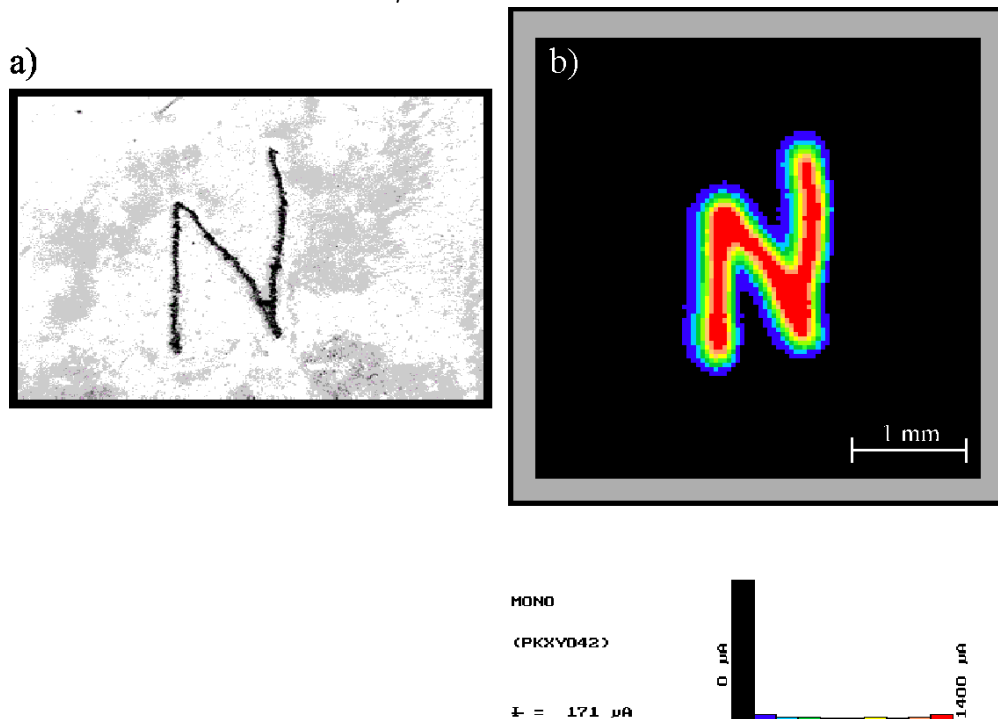


Abb. 4.16: a) Foto des eingekratzten ’N’. b) Zugehöriger Flächen-’Scan’, gemessen mit der 200 μm Zylinder-Elektrode. Der Elektrodenabstand beträgt 10 μm .

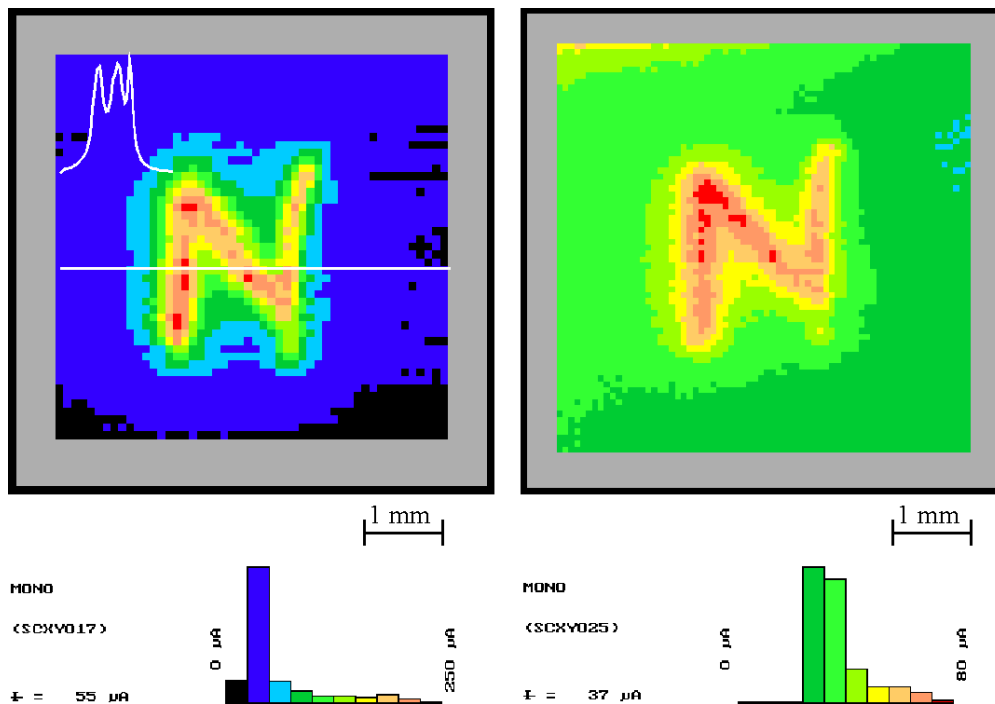


Abb. 4.17: Flächen-’Scan’ mit Platindraht-Elektrode. Im rechten Bild war der Elektrodenabstand über den Wafer nicht konstant. Oben links im linken Bild eingezeichnet ist ein ’Line-Scan’ entlang der weißen Linie.

4.3.2 Flächen-’Scans’ mit den einzelnen Elektroden

So weit im folgenden nichts anderes angegeben wird, haben alle Wafer, die als Proben für Flächen-’Scans’ verwendet wurden, einen spezifischen Widerstand von $0.02 \Omega\text{cm}$.

Mit Abstand am schwierigsten ist es, das ’N’ mit der Platindraht-Elektrode zu messen. Neben meßtechnischen Problemen können auch die in Abschnitt 3.2.5.1 diskutierten Schwierigkeiten bei Herstellung der Platindraht-Elektrode als Ursachen genannt werden:

1. Die sehr starke Abstandsabhängigkeit des Auflösungsvermögens der Platindraht-Elektrode, siehe Abb. 4.10, erfordert, daß die Elektrodennachführung sehr präzise erfolgen muß.
2. Es gestaltet sich ziemlich schwierig, die Elektroden-Herstellung so zu bewerkstelligen, daß am Ende die fertige Elektrode keine ’Leaks’ in der Isolierung des Pt-Drahtes aufweist.

Wird Polyurethan als Isolationsmaterial verwendet, besteht eine zusätzliche Schwierigkeit darin, daß sich die Polyurethanschicht in ethanolhaltigen Elektrolyten mit der Zeit ablöst.

3. Allzu leicht kann es passieren, daß sich die Spitze der Elektrode während der Prozedur zur Bestimmung von $z = 0$, siehe Abschnitt 3.4, beim Aufsetzen auf den Wafer geringfügig verbiegt, etwa weil Lackpartikel den Kurzschluß zwischen

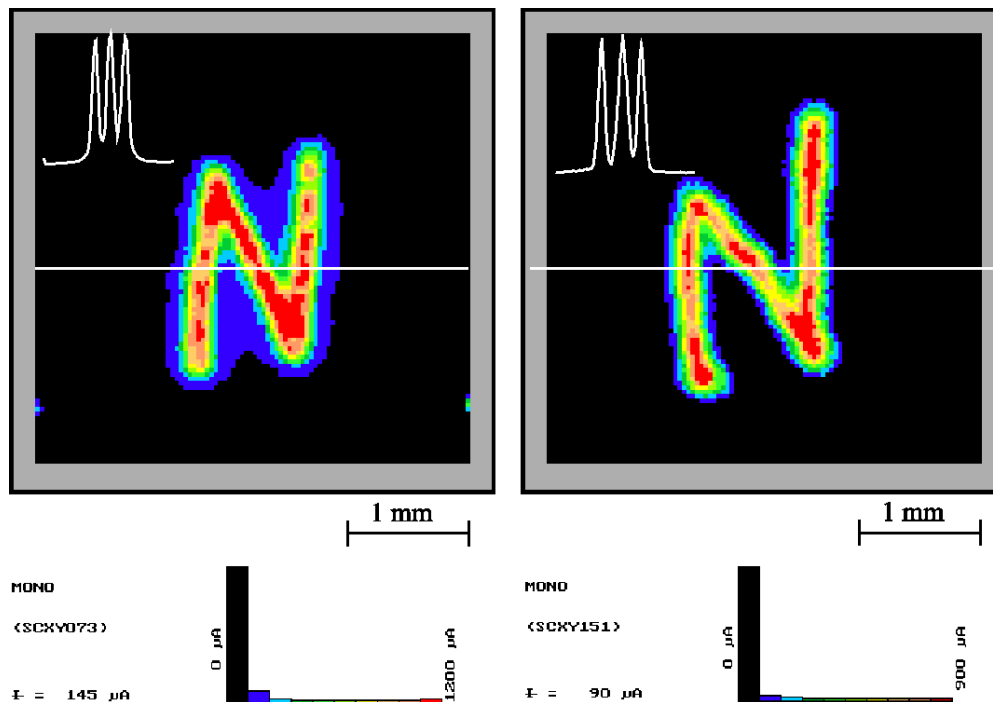


Abb. 4.18: Flächen-‘Scans’ mit der 200 μm Zylinder-Elektrode. Für beide Messungen beträgt der Elektrodenabstand 10 μm .

Elektrode und Wafer verhindern. Die Konsequenz ist, daß der Flächen-‘Scan’ eine Drift im Strom zeigt, weil infolgedessen der Abstand über dem Wafer nicht konstant ist. Vergleiche hierzu die rechte Messung in Abb. 4.17.

4. Die Elektrodenfläche ist im Vergleich zu den anderen Elektroden weit weniger definiert, da sie nach Fertigstellung nicht weiter bearbeitet werden kann. Wie in Abschnitt 3.2.5.1 beschrieben, ist daher im letzten Schritt nach Abschneiden der Elektrodenspitze mit einer scharfen Schere der Erhalt einer guten Elektrode reine Glückssache.

Nach Ersetzen der Polyurethanbeschichtung durch eine farbige Lackschicht ist die Isolierung zwar lösungsmittelfest, das Abschneiden der Elektrodenspitze gestaltet sich aber noch schwieriger, weil der Lack aufgrund seiner Härte während des Schneidens sehr leicht splintern kann, wodurch die Elektrode unbrauchbar werden würde.

Darauf hingewiesen werden muß, daß das ‘Scannen’ eines ‘N’ mit der Platindraht-Elektrode, so wie es das linke Bild in Abb. 4.17 zeigt, mit der Güte und einem Signal-zu-Untergrund-Verhältnis von 5(!) nur ein einziges Mal gelang! Leider wurde die Elektrode, wie oben beschrieben, nach einigen Messungen unbrauchbar. Es ist fast unmöglich unter Laborbedingungen die Platindraht-Elektrode reproduzierbar herzustellen, so daß mehrere Versuche, eine solche Elektrode mit vergleichbarer Qualität herzustellen, erfolglos blieben.

Besonders gute Messungen mit der ersten 200 μm Zylinder-Elektrode zeigt

Abb. 4.18. Ein Vergleich mit dem 'Line-Scan' der Platindraht-Elektrode als Inset in Abb. 4.17 zeigt, daß die 200 μm Zylinder-Elektrode eine sehr viel schärfere Auflösung der Teststruktur erlaubt. Grund hierfür ist das fokussierte elektrische Feld (vgl. Abb. 3.2) unter der 200 μm Zylinder-Elektrode, aus dem sehr steile Stromflanken resultieren, wenn sich die Elektrode über einen Kratzer hinweg bewegt.

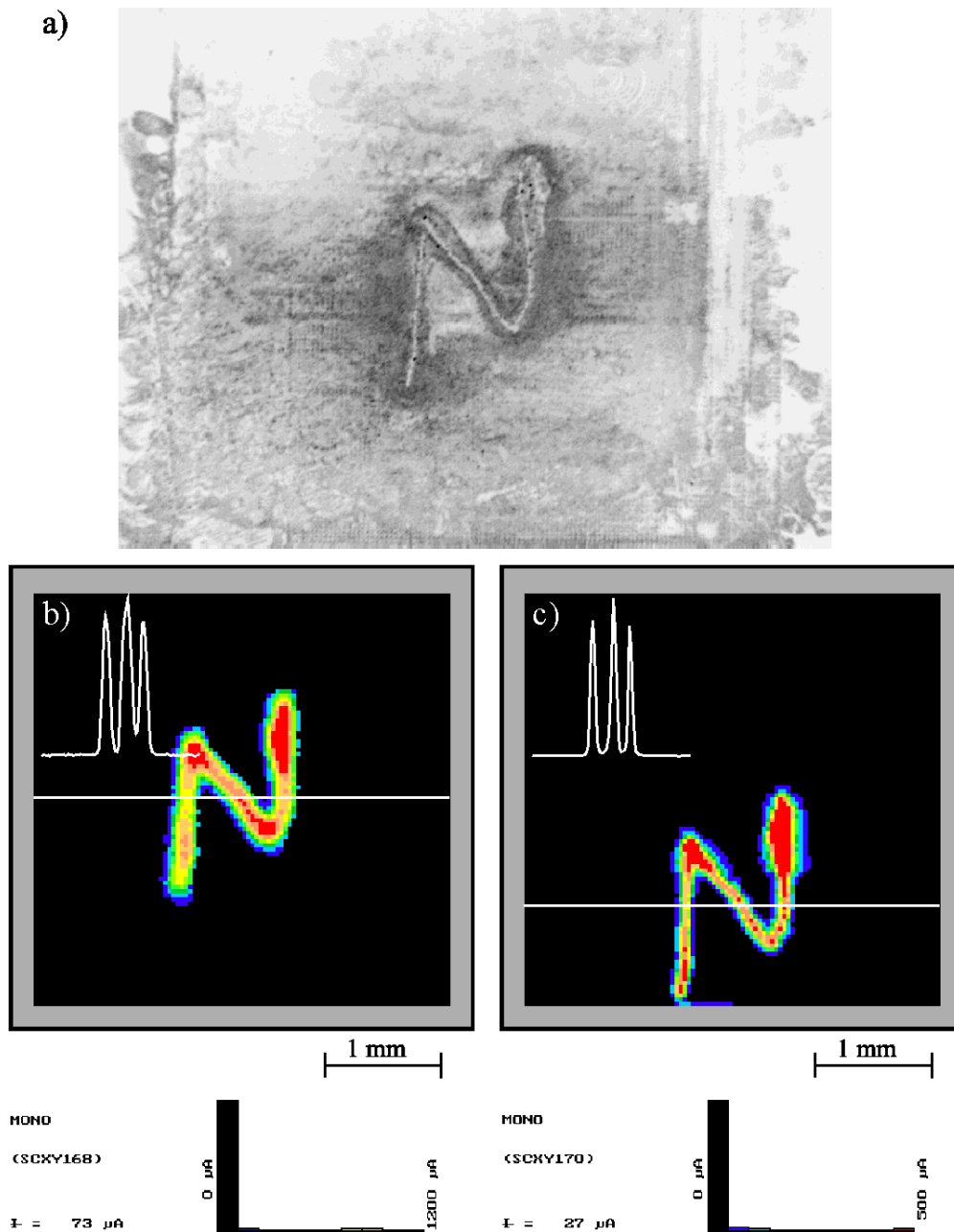


Abb. 4.19: Vergleich beider Zylinder-Elektroden. a) Das Foto zeigt deutlich den Doppelkratzer im oberen rechten Teil des 'N'. Ebenfalls deutlich zu erkennen ist das 'Scan'-Raster. b) Messung mit der 200 μm und c) Messung mit der 50 μm Zylinder-Elektrode. Für beide Messungen beträgt der Elektrodenabstand 10 μm .

Ein direkter Vergleich der 200 μm mit der 50 μm -Zylinder-Elektrode ist in Abb. 4.19 dargestellt. Für beide Messungen beträgt der Elektrodenabstand $d = 10 \mu\text{m}$. Wie erwartet, wird das ‘N’ von der 50 μm Zylinder-Elektrode bedeutend schärfer aufgelöst als mit der 200 μm Zylinder-Elektrode. Die Verschiebung der Bilder rührt daher, daß die Elektroden nicht exakt senkrecht zur Waferoberfläche orientiert waren. Beim Kratzen geriet eine der parallelen Linien zu kurz, so daß durch Nachkratzen versucht wurde, sie zu verlängern. Statt einer Verlängerung wurde ein zweiter Kratzer aufgebracht, der zum ersten parallel verläuft, siehe Abb. 4.19 a). Aus diesem Grund sind die Messungen an diesem Bereich einander sehr ähnlich.

Die lichtmikroskopische Aufnahme des ‘N’ zeigt deutlich das ‘Scan’-Raster! In Abschnitt 3.5 wird diskutiert, daß an der elektrochemischen Reaktion zwischen Meßelektrode und Wafer auch Ethanol, der für diese Messungen dem Elektrolyten beigemischt wurde, beteiligt ist und zu lokaler Korrosion der Waferoberfläche führt. Während mit ethanolhaltigen Elektrolyten nach einer Messung das ‘Scan’-Raster oftmals sichtbar war, wurde es grundsätzlich *nie* sichtbar, wenn mit einem ethanolfreien Elektrolyten gemessen worden war, selbst dann nicht, wenn mehrfach gemessen wurde.

Das Ergebnis des ersten Flächen-‘Scans’ mit der Scheiben-Elektrode zeigt das linke Bild in Abb. 4.20. Etwa nach der Hälfte der Messung stieg der Strom sprunghaft an und verblieb dort bis zum Ende, wobei die Ortsauflösung komplett verloren ging. Die Ursache hierfür ist, daß das Ethanol im Elektrolyten den Epoxylebstoff mit der Zeit auflöste. Zwecks Kontaktierung von Pt-Röhrchen und Pt-Scheibe wurde bei Herstellung der ersten Scheiben-Elektrode nicht Silberleitlack verwendet, wie in Abschnitt 3.2.5.4 beschrieben. Stattdessen wurde der Kontakt über drei Kupferdrähte hergestellt, die zwischen beide Teile gelötet wurden und die zudem auch noch für eine sehr gute mechanische Stabilität sorgten. Nach dem Zusammenlöten der einzelnen Teile wurde abschließend das kupferne Dreieck mit Epoxylebstoff vergossen. Bedingt durch die Auflösung des Klebstoffes durch Ethanol wurden während der Messung die Kupferdrähte teilweise wieder freigelegt, woraufhin sich auf der Siliziumoberfläche elektrochemisch Kupfer abschied. Aus dieser Kupferschicht auf dem Silizium resultierte letztlich der hohe Strom und der Verlust der Ortsauflösung.

Bei Herstellung der zweiten Scheiben-Elektrode erfolgte die Kontaktierung beider Elektrodenteile nicht über ein kupfernes Dreieck, sondern so, wie in Abschnitt 3.2.5.4 beschrieben. Die Messung mit der neuen Scheiben-Elektrode zeigt das rechte Bild in Abb. 4.20. Wie bereits die ‘Line-Scans’ der Scheiben-Elektrode zeigten (siehe Abb. 4.12), erlaubt diese Elektrode ein besonders scharfes Auflösen der Teststruktur! Ein derart steiler Anstieg des Stromes, wenn sich die Elektrode dem Kratzer nähert, wird nur mit dieser Elektrode erreicht. Eine Betrachtung der Teststruktur unter dem Lichtmikroskop zeigte, daß im oberen Teil des ‘N’ nicht durchgängig gekratzt worden war. Deshalb sind an diesem Ort die Ströme kleiner. Im Kratzer festsitzende H_2 -Bläschen sind also nicht Ursache des verringerten Stromes.

Häufig wird beobachtet, daß die Reproduzierbarkeit der Messungen sehr schlecht ist, wenn ein ethanolhaltiger Elektrolyt verwendet wird. Von Messung zu Messung wird das Signal über den Kratzern immer breiter, bis es kaum noch zu erkennen ist. Ursache hierfür ist die in Abschnitt 3.5 beschriebene Degradation der Siliziumoberfläche, verur-

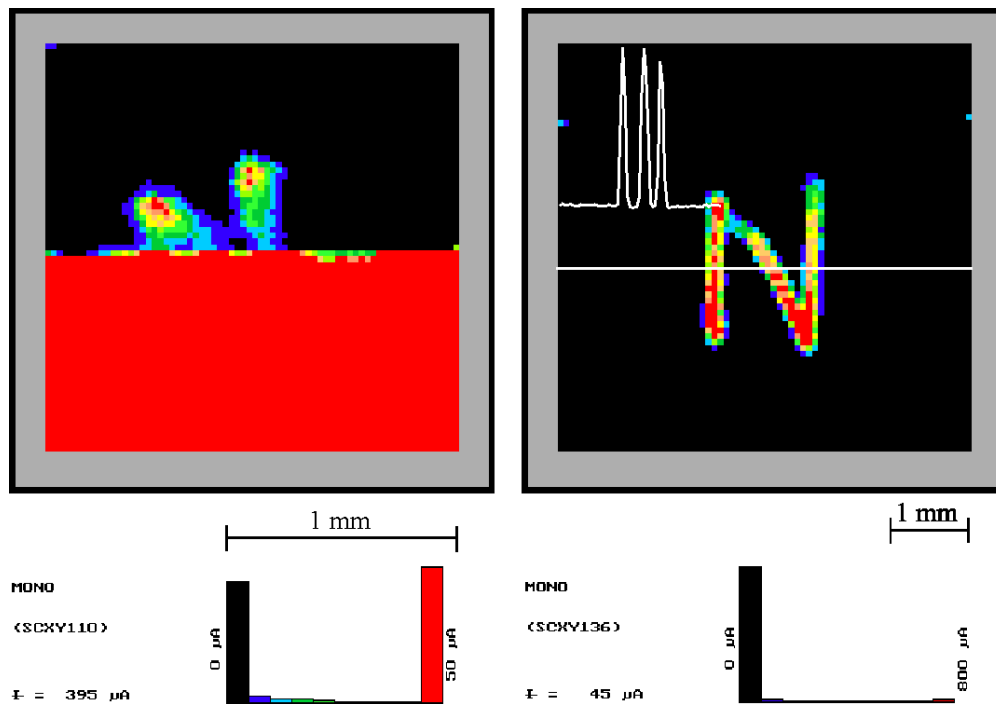


Abb. 4.20: Flächen-‘Scan’ mit Scheiben-Elektrode. Linkes Bild: Während der Messung wurde die Elektrode zerstört, weil das Ethanol im Elektrolyten den Epoxyklebstoff auflöste. Rechtes Bild: Messung in ethanolfreiem Elektrolyt mit neuer Scheiben-Elektrode.

sacht durch Ethanol-Radikale, die als Folge des Stromflusses gebildet werden. Abb. 4.21 zeigt zwei Beispiele anhand derer die Degradation der Waferoberfläche während der Messung besonders deutlich sichtbar wird.

Wird dem Elektrolyten kein Ethanol zugesetzt, kommt es grundsätzlich *nie* zur Degradation der Waferoberfläche, was sich in sehr guter Reproduzierbarkeit der Meßergebnisse äußert. Zwei Beispiele, die das zum Ausdruck bringen, zeigt Abb. 4.22.

4.3.3 Messungen an Material mit höherem spez. Widerstand

An dieser Stelle soll anhand von Messungen an Material mit höherem spezifischen Widerstand demonstriert werden, warum für die Charakterisierung der Elektroden in Abschnitt 4.2 hochdotiertes Silizium ($0.02 \Omega\text{cm}$) verwendet worden ist. Das Silizium, der in diesem Abschnitt gezeigten Messungen, hat einen spezifischen Widerstand von $6\text{-}10 \Omega\text{cm}$.

Zwar reicht nach Gleichung 1.12 die Raumladungszone von niedrigdotiertem Material an defektfreien Stellen tief ins Silizium, so daß sich das Signal über einem Kratzer sehr deutlich vom Untergrund abheben sollte, allerdings ändern sich die Eigenschaf-

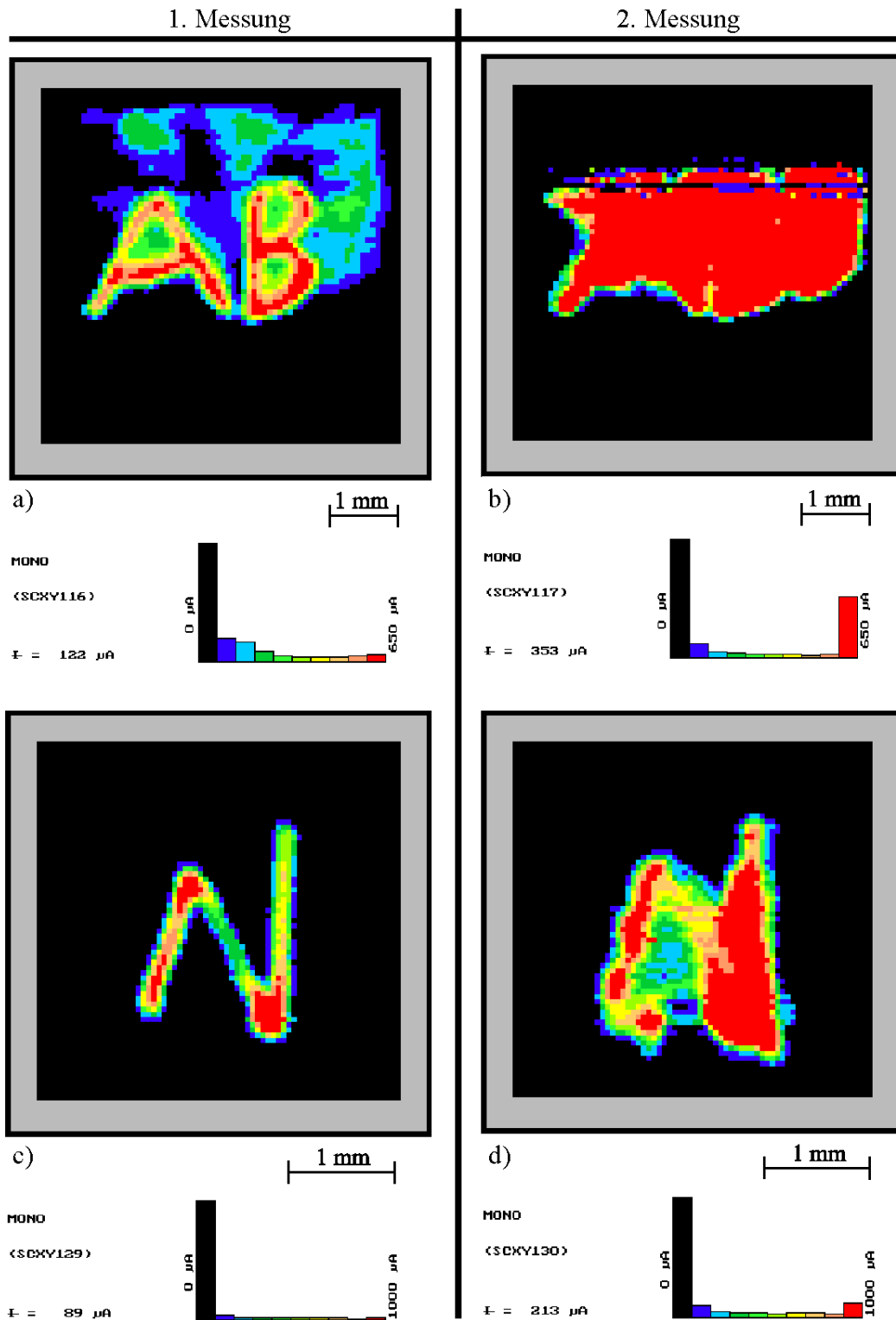


Abb. 4.21: Zwei Beispiele für schlechte Reproduzierbarkeit. 1. Beispiel: a) und b), 2. Beispiel: c) und d). Sehr wahrscheinlich verursachen Ethanol-Radikale eine lokale Degradation der Waferoberfläche, deren Konzentration an Orten hoher Ströme besonders groß ist.

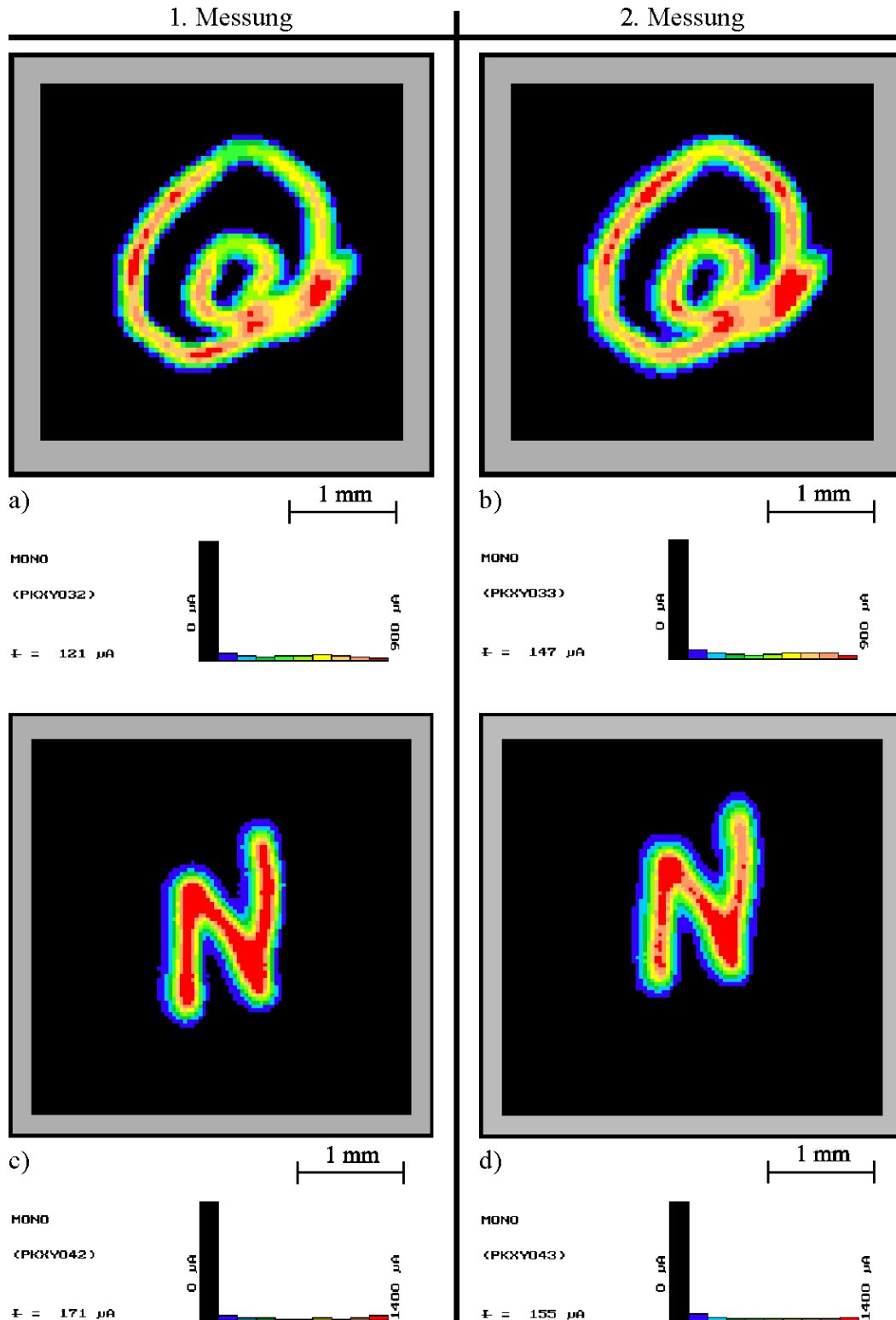


Abb. 4.22: Jeweils zwei Messungen an derselben Teststruktur. 1. Beispiel: a) und b), 2. Beispiel: c) und d). Grund für die sehr gute Reproduzierbarkeit ist das Verwenden eines ethanolfreien Elektrolyten.

ten von Kratzern auf diesem Material mit der Zeit relativ stark. Die Eigenschaft, als Leckstromquelle zu wirken, besitzt der Kratzer nur kurz nachdem er aufgebracht worden ist. Da aber die zahlreichen offenen Bindungen, die durch das Kratzen erzeugt werden, im Verlauf der Zeit mit Wasserstoff abgesättigt werden, verliert der Kratzer zusehends diese Eigenschaft. Den Vorgang der allmählichen Passivierung eines Kratzers mit der Zeit verdeutlicht Abb. 4.23. Unmittelbar nach Aufbringen des Kratzers sind die Ströme an dessen Ort deutlich erhöht. Bereits nach 30 min konnte der Kratzer nicht mehr lokalisiert werden!

Daß mit dem Leckstrom-Scanner auch Teststrukturen auf hochohmigem Material, trotz allmählicher Passivierung, lokalisiert werden können, zeigt Abb. 4.24.

4.3.4 Einfluß von Gasbläschen auf die Messungen

Gestützt durch Ergebnisse aus zahlreichen Messungen und Hinweisen in der Literatur kann als sicher gelten, daß das Ethanol im Elektrolyten bei Stromfluß zu erheblichen Korrosionserscheinungen an der Siliziumelektrode führt, siehe Abschnitt 3.5. Da der Leckstrom-Scanner möglichst zerstörungsfrei arbeiten soll, wurde im fortgeschrittenen Stadium dieser Arbeit für Messungen auf den Ethanolanteil im Elektrolyten verzichtet.

Ohne Ethanol ist die Oberflächenspannung des Elektrolyten wesentlich höher. Die

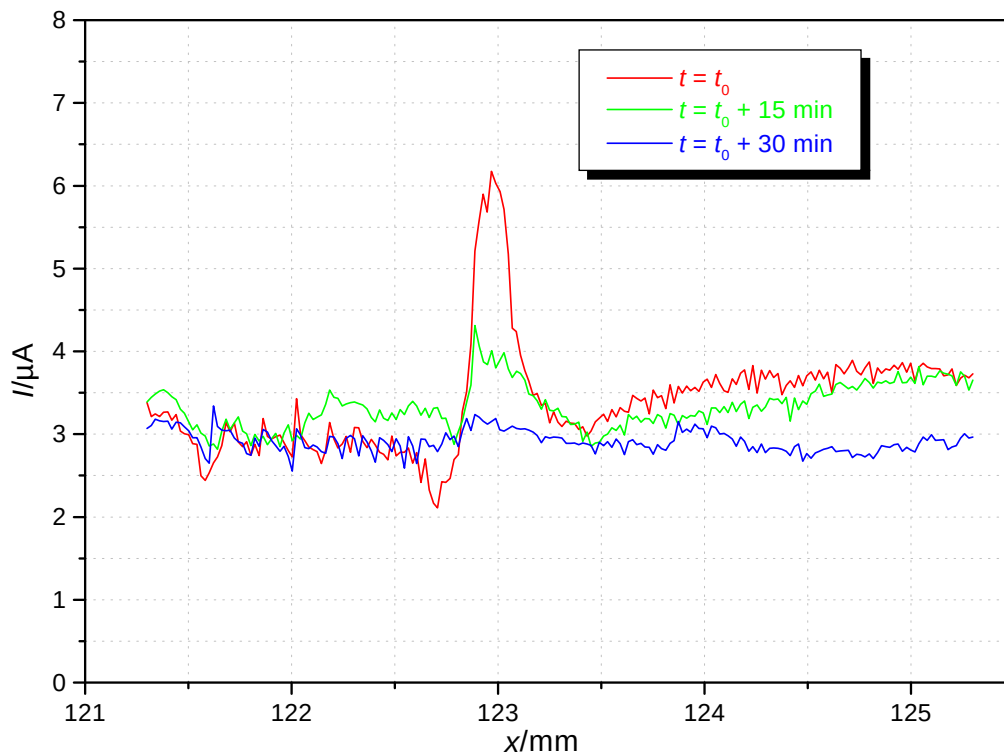


Abb. 4.23: ‘Line-Scans’ über einem Kratzer auf Material mit 6-10 Ωcm . Sofort nach Aufbringen des Kratzers (rot), 15 min später (grün) und 30 min (blau) später. Messung mit der 200 μm Zylinder-Elektrode. Der Elektrodenabstand beträgt 10 μm .

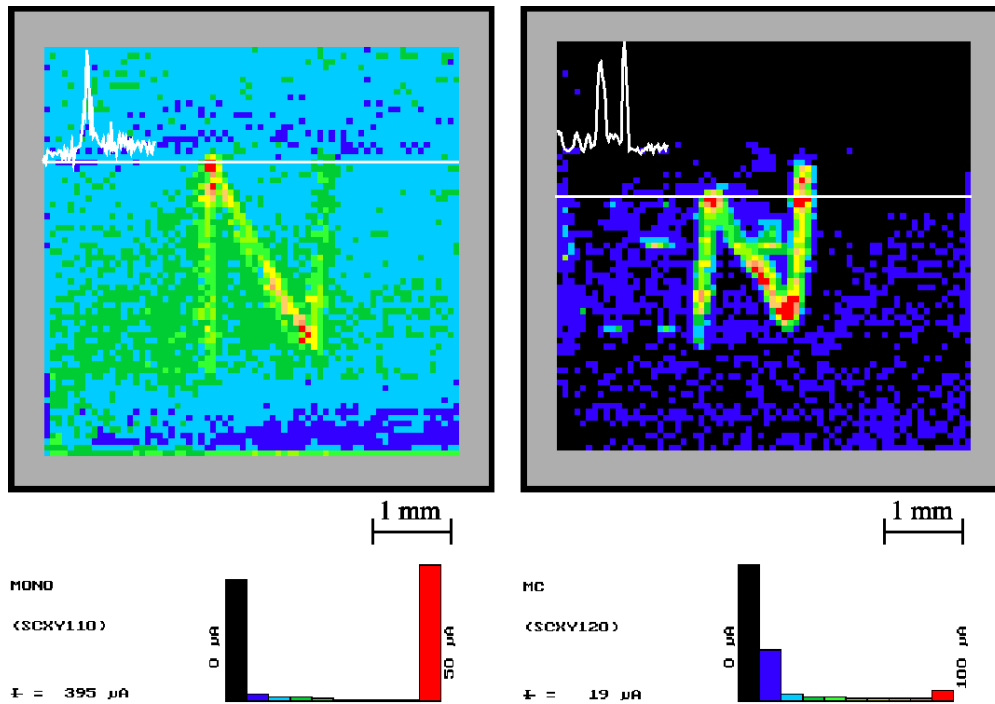


Abb. 4.24: Zwei Messungen an Material mit $6\text{--}10\ \Omega\text{cm}$, gemessen mit der $200\ \mu\text{m}$ Zylinder-Elektrode. Der Elektrodenabstand beträgt $10\ \mu\text{m}$.

Konsequenz ist, daß Gasbläschen, die sich während einer Messung zwischen Wafer und Meßelektrode bilden, länger haften bleiben. Während der Einfluß von H_2 -Bläschen, die sich an der Waferoberfläche bilden, auf die Messungen im allgemeinen unwesentlich ist, können sich O_2 -Bläschen, die an der Meßelektrode gebildet werden, sehr störend auf die Messung auswirken: Ist das Bläschen klein, bleibt es dort haften und wird im Verlauf der Messung 'mitgeschleift', wobei es sich zunehmend vergrößert. Mit größer werdendem O_2 -Bläschen wird der Zellenstrom stetig abnehmen, weil sich die elektrochemisch aktive Fläche der Elektrode immer weiter verkleinert. Irgendwann wird sich das Bläschen ablösen, woraufhin kurzzeitig wieder der maximale Strom fließen wird, vgl. hierzu Abb. 3.6.

Wie bereits weiter oben diskutiert, ist der Einfluß der permanenten Bläschenbildung auf die Messung dann gering, wenn das Signal-zu-Untergrund-Verhältnis groß ist, was bei Messungen an Kratzern im allgemeinen der Fall ist. Für Messungen an multikristallinem Material sollte der Einfluß aber erheblich sein, weil für dieses Material eine sehr viel kleinere Dynamik in den Strömen erwartet wird, siehe Abb. 4.5. Das ständige Wechselspiel zwischen Bildung und Ablösung von Sauerstoffbläschen hat zur Folge, daß das eigentliche Meßsignal komplett im Rauschen verschwindet. Beispiele, die den nachteiligen Einfluß der Bläschenbildung illustrieren, zeigt Abb. 4.25.

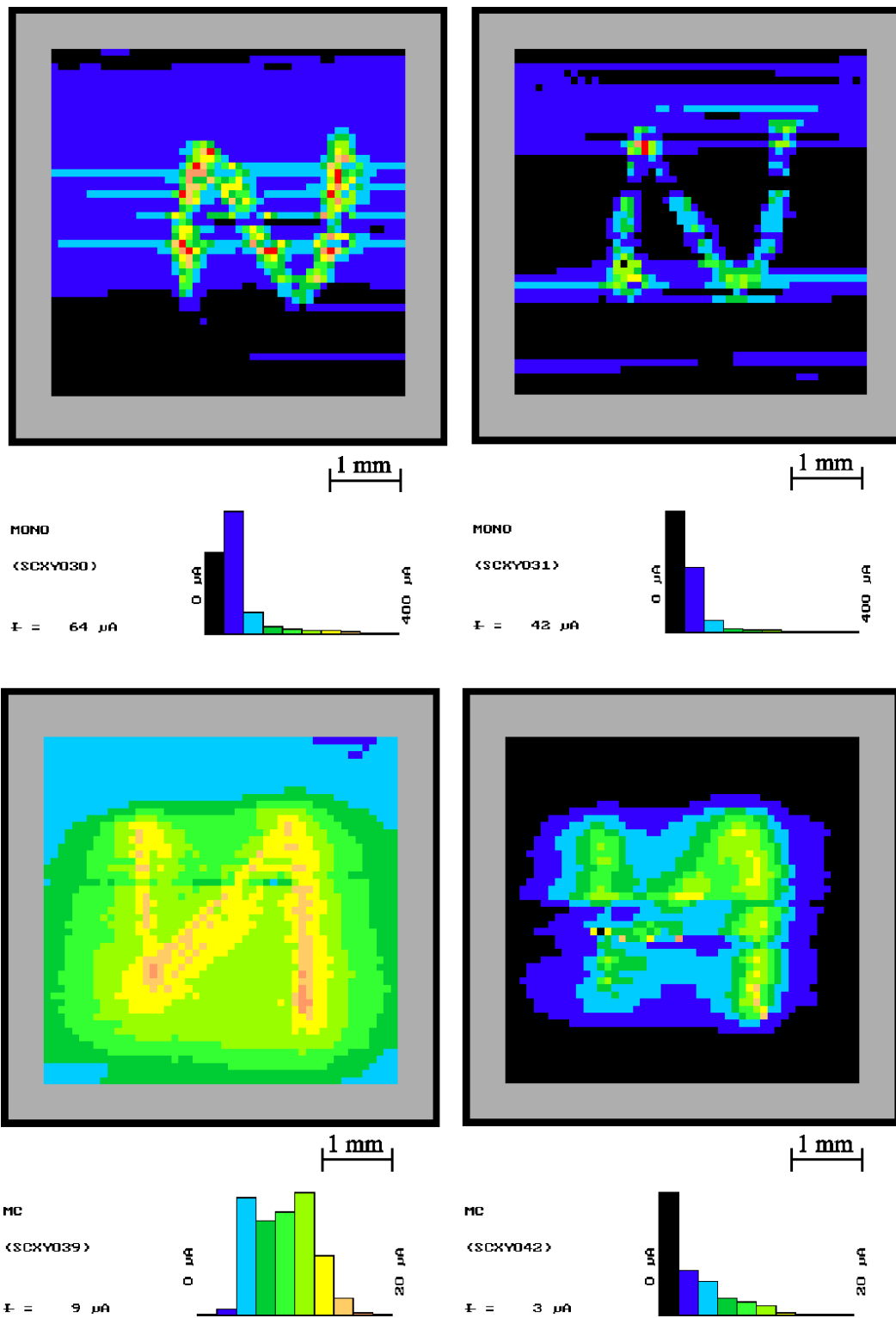


Abb. 4.25: Beispiele für Messungen bei denen sich die Bildung von O_2 -Bläschen unter der Meßelektrode deutlich auf die Ströme auswirkt. In jeder Messung ist die 'Scan'-Richtung gut zu erkennen. Alle Messungen wurden mit der Platindraht-Elektrode durchgeführt. Der Elektrodenabstand beträgt $10 \mu m$.

4.4 Maßnahmen zur Vermeidung des Einflusses von Gasbläschen auf die Messungen

Da sich H_2 - und besonders O_2 -Bläschen nachteilig auf die Messungen auswirken können, wie Abb. 4.25 anhand einiger Beispiele zeigt, wurde mit verschiedenen Ansätzen versucht, den Einfluß der Gasbläschen auf die Messungen entweder erheblich zu reduzieren oder gar gänzlich zu beseitigen. Die Ansätze im Einzelnen waren:

- Pulsen der Zellenspannung.
- Einsatz einer Elektroden-spülung.
- Ersetzen des wässrigen Elektrolyten durch einen nicht-wässrigen.
- Einsatz eines 'Peak'-Detektors für die Strom-Messungen.

4.4.1 Pulsen der Zellenspannung

Bereits in Abschnitt 4.1.3 wurde diskutiert, daß das Pulsen der Zellenspannung zwar den Grad der Bläschenbildung verringert, gleichzeitig aber ein verstärktes Korrodieren der Waferoberfläche (Bildung von PSL) zur Folge hat, weil beim Umladen von Raumladungszonen und Helmholtz-Kapazitäten kurzzeitig anodische Ströme durch die Zelle fließen, siehe Abb. 4.3. Aus diesem Grund wurde das Pulsen der Zellenspannung kurz nach Einführung wieder aufgegeben.

4.4.2 Elektroden-spülung

Durch Einstellung eines Strömungsprofils unmittelbar am Ort der Elektroden-spitze wurde versucht, den Einfluß von H_2 - und O_2 -Bläschen auf die Meßergebnisse zu verringern. Die hierfür erforderliche Modifizierung der Meßzelle zeigt Abb. 4.26.

Für eine genaue Justierung des Kunststoff-Röhrchens am Ort der Elektrode wird es durch ein entsprechend zurechtgebogenes Kupferrohr geführt. Das Kupferrohr wird auf ein Kupferblech gelötet, das selbst am Meßkopf befestigt ist. Zum Umwälzen des Elektrolyten wird eine Perestaltik-Pumpe verwendet, die speziell für den Einsatz in aggressiven Medien konstruiert ist.

Trotz zahlreicher Versuche mit verschiedenen Orientierungen des Kunststoff-Röhrchens zur Elektrode sowie variierten Pumpleistungen konnte nicht die geringste Verbesserung der Meßergebnisse erzielt werden. Ursache hierfür ist vermutlich, daß der eigentliche Bereich zwischen Elektrode und Silizium, selbst für größere Elektroden-abstände, nicht oder nur sehr schlecht durchströmt wird. Denkbar ist auch, daß die O_2 -Bläschen an der Meßelektrode sehr fest haften, so daß sie sich kaum fortspülen lassen.

Verbesserungen sind zu erwarten, wenn die Spülung mit in die Elektrode integriert wird, da dann sichergestellt ist, daß der Bereich zwischen Meßelektrode und Silizium von Elektrolyt durchströmt wird. Das zu realisieren wäre aber aufgrund der Feinheit der Elektroden, was in besonderem Maße für die 50 μm Zylinder-Elektrode gilt, mit

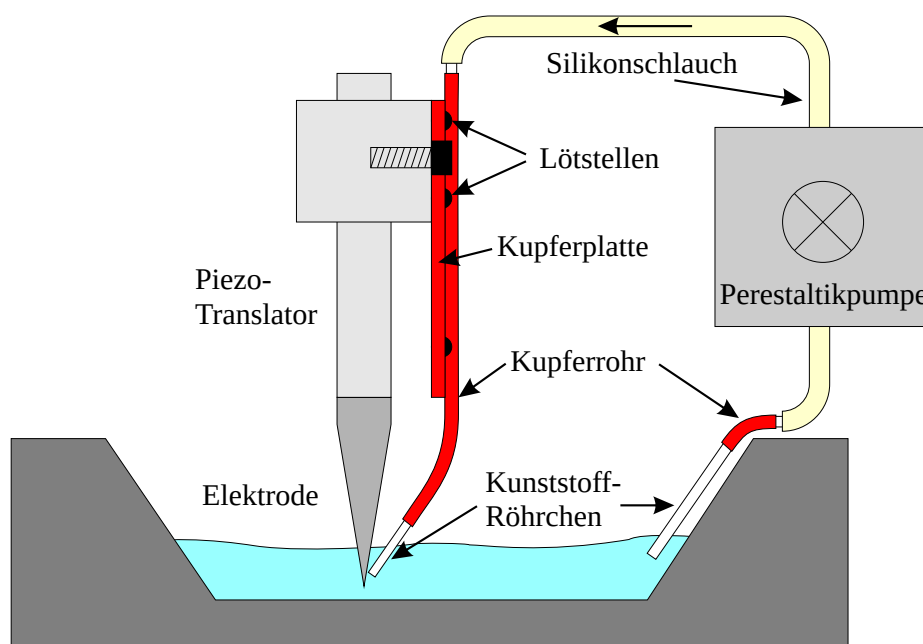


Abb. 4.26: Erweiterung der Meßzelle um eine Elektroden-spülung. Das Kupferrohr ermöglicht eine sehr präzise Justierung des Kunststoff-Röhrchens vor der Elektroden-spitze.

erheblichen technischen Schwierigkeiten verbunden. Fraglich ist, ob überhaupt ausreichend hohe Strömungsgeschwindigkeiten unter der Elektrode erreicht werden können, weil die sehr feinen Röhrchen, die für eine Realisierung hierfür erforderlich sind, sehr große Strömungswiderstände haben würden. Außerdem besteht die Gefahr, daß die Meßelektrode von der Oberfläche des Siliziums fortgedrückt wird, was sich nachteilig auf die Meßergebnisse auswirken würde. Im ungünstigsten Fall entstehen Schwingungen senkrecht zur Siliziumoberfläche, die zum einen die Messung aufgrund der zahlreichen Kurzschlüsse zwischen Meßelektrode und Wafer unbrauchbar machen würden, und zum anderen eine Beschädigung der Elektrode zur Folge haben könnten.

4.4.3 Nicht-wässriger Elektrolyt

Findet Ladungstransport im Elektrolyten über Redoxpaare statt, deren reduzierte oder oxidierte Form nicht im Elektrolyten verbleibt, können Gase entstehen, wie dies bei den Redoxpaaren H^+/H_2 und $\text{O}_2^{2-}/\text{O}_2$ der Fall ist. Deshalb wurde versucht, diese Redoxpaare durch ein anderes zu ersetzen. Hierfür müssen Redoxpaar und Lösungsmittel folgenden Anforderungen genügen:

1. Sowohl die oxidierte als auch die reduzierte Form des Redoxpaares müssen im Elektrolyten verbleiben, d. h. sie dürfen weder Gase bilden noch dürfen sie sich auf einer der Elektroden abscheiden.

2. Die Produkte der elektrochemischen Reaktion müssen im Lösungsmittel löslich sein und dürfen nicht mit ihm oder den Elektroden chemisch reagieren.
3. Das Standardpotential des Redoxpaares muß im kathodischen Bereich der p-Silizium-Elektrode liegen.
4. Das Lösungsmittel muß wasserfrei sein, darf nicht mit den Elektroden chemisch reagieren und darf sich in Umgebung vom Standardpotential des Redoxpaares nicht zersetzen.

Redoxpaar		Std. Pot. vs. SCE	Literatur- angaben
Cobaltocene/Cobalticinium	$\text{Co}(\text{Cp})_2/\text{Co}(\text{Cp})_2^+$	-0.94 V	[CHA80], [CHA81]
Chromocene/Chromicinium	$\text{Cr}(\text{Cp})_2/\text{Cr}(\text{Cp})_2^+$	-0.67 V	[CHA81]
Nickelocene/Nickelicinium	$\text{Ni}(\text{Cp})_2/\text{Ni}(\text{Cp})_2^+$	-0.09 V	[CHA81]
Benzoquinone/Monoanion	BQ/BQ^-	-0.51 V	[CHA81], [CHA83a], [LAS76], [NAG82]
Anthraquinone/Monoanion	AQ/AQ^-	-0.94 V	[CHA81], [CHA83a], [LAS76], [NAG82], [ODE90], [TUR80]
Nitrobenzene/Monoanion	NB/NB^-	-1.15 V	[CHA81], [LI95]

Tab. 4.1: Auswahl geeigneter Redoxpaare für Untersuchungen an der p-Silizium-Elektrolyt-Grenzfläche. $(\text{Cp})_2$ ist eine Abkürzung für Bis(cyclopentadienyl).

Tab. 4.1 zeigt eine Zusammenstellung von einigen Redoxpaaren, die in Acetonitril (CH_3CN) gelöst, den oben stehenden Anforderungen genügen und die für verschiedenartige Untersuchungen an der p-Silizium-Elektrolyt-Grenzfläche vielfach Anwendung finden, siehe Literaturangaben.

Reines Acetonitril selbst hat eine sehr kleine Leitfähigkeit. Um zu erreichen, daß der Großteil der im Elektrolyten abfallenden Spannung über den beiden Helmholtz-Schichten abfällt, wird im allgemeinen dem Elektrolyten noch ein Leitsalz zugesetzt. Für die Experimente wird Tetrabutylammonium-perchlorat (TBAP) als Leitsalz verwendet.

4.4.3.1 Auswahl des Redoxpaares

Als Kriterium für die Wahl des Redoxpaares könnte dessen Standardpotential herangezogen werden. Je weiter im kathodischen Bereich gemessen wird, desto besser wird eine Trennung zwischen Signal und Untergrund möglich sein, siehe Abb. 4.1. Dennoch wird nicht das Redoxpaar NB/NB^- gewählt, sondern die Redoxpaare $\text{Co}(\text{Cp})_2/\text{Co}(\text{Cp})_2^+$ und $\text{Cr}(\text{Cp})_2/\text{Cr}(\text{Cp})_2^+$, trotz der geringeren kathodischen Standardpotentiale. Grund für die Bevorzugung ist folgender:

Da vor jeder Messung ein Entfernen des Oxides mit HF zwingend notwendig ist, kann es kaum gelingen, die Zelle frei von Wasser zu halten. Hinzu kommt, daß das Herstellen von wasserfreiem Acetonitril schwer zu realisieren ist. Getrocknetes Acetonitril,

wie es für die Experimente verwendet wird, hat einen Restwassergehalt⁶ von $50 \frac{\mu\text{L Wasser}}{\text{L}}$. Dieser Restwassergehalt hat eine relativ zügige Oxidation der Waferoberfläche zur Folge. Mit einer Diffusionskonstanten von Wasser in Acetonitril von $D \approx 10^{-6} \frac{\text{cm}^2}{\text{s}}$ hat sich bereits nach $\tau = 36 \text{ s}$ auf dem Silizium eine $d = \sqrt{D\tau} = 30 \text{ \AA}$ dicke Schicht aus SiO_2 gebildet, was etwa 10 Monolagen entspricht [CHA80, CHA81]!

Cobaltocene und Chromocene sind sehr leicht oxidierbar, so daß sie im Elektrolyten enthaltenes Wasser oder gelösten Sauerstoff 'wegfangen'. Die Reaktionsprodukte sind in Acetonitril nicht lösliche Oxide oder Hydroxide, die sich nicht auf Silizium oder Platin abscheiden. Auf diese Weise wird die allmähliche Oxidation der Siliziumoberfläche zwar nicht verhindert, jedoch sehr stark reduziert.

Einen gegenteiligen Effekt besitzt das Redoxpaar NB/NB^- . Dieses Redoxpaar reagiert ebenfalls mit dem Wasser im Elektrolyten, allerdings gehören zu den Reaktionsprodukten auch OH^- -Ionen, die die Oxidation der Siliziumoberfläche beschleunigen⁷ [CHA81]!

4.4.3.2 Herstellung des Elektrolyten

Folgende Zusammensetzung des Elektrolyten wird verwendet: 0.1 M TBAP und 0.01 M Cobaltocene bzw. Chromocene in Acetonitril [CHA80, CHA81]. Das Cobaltocene wurde von der Firma *Alta*, das Chromocene von der Firma *Boulder Scientific Company* in Colorado bezogen. Zuvor wird das Leitsalz im Trockenschrank in einer PE-Weithals-Flasche bei $T = 70 \text{ }^\circ\text{C}$ und einem Restdruck von $p < 10^{-2} \text{ Torr}$ 12 Stunden lang getrocknet. Das Abwiegen des Redoxpaares und das Einfüllen in die Weithals-Flasche hat möglichst schnell zu erfolgen. Abschließend wird das Acetonitril zugegeben und danach die Flasche schnellstmöglich verschlossen.

4.4.3.3 Messungen

Die oben erwähnte PE-Weithals-Flasche dient gleichzeitig als Testzelle, deren Deckel hierfür mit fünf Durchführungen versehen wurde. Drei Durchführungen für die Elektroden und zwei weitere für Gasleitungen, um das Innere der Zelle in eine Schutzgas-Atmosphäre zu bringen.

Für erste Messungen wurde auf die Verwendung von p-Silizium als 'Working'-Elektrode verzichtet und stattdessen ein ca. $15 \text{ mm} \times 10 \text{ mm}$ großes Platinblech gewählt. Als Referenz-Elektrode dient ein $200 \text{ }\mu\text{m}$ dicker Pt-Draht, der mit einem Polyimid-Röhrchen isoliert wird, vgl. Abb. 3.10. Als 'Counter'-Elektrode wird ebenfalls ein $200 \text{ }\mu\text{m}$ dicker Pt-Draht verwendet. Alle Elektroden werden jeweils für 10 min mit Aceton und Methanol im Ultraschallbad gereinigt, dessen Wasserbad hierfür auf eine Temperatur von $T = 50 \text{ }^\circ\text{C}$ erhitzt wird [LU90]. Anschließend wird der isolierte Pt-Draht mit etwas Epoxylebstoff auf dem Pt-Blech fixiert. Während des Trocknens

⁶ Es ist möglich, den Restwassergehalt im Elektrolyten auf unter $15 \frac{\mu\text{L Wasser}}{\text{L}}$ zu reduzieren (Karl-Fischer-Methode) [CHA81, CHA89]. Hierfür fehlten aber technisch die Voraussetzungen.

⁷ Reagiert das Redoxpaar AQ/AQ^- mit Wasser, ist die Bildung von OH^- -Ionen weniger intensiv, als dies für die Paare BQ/BQ^- und NB/NB^- der Fall ist [CHA81].

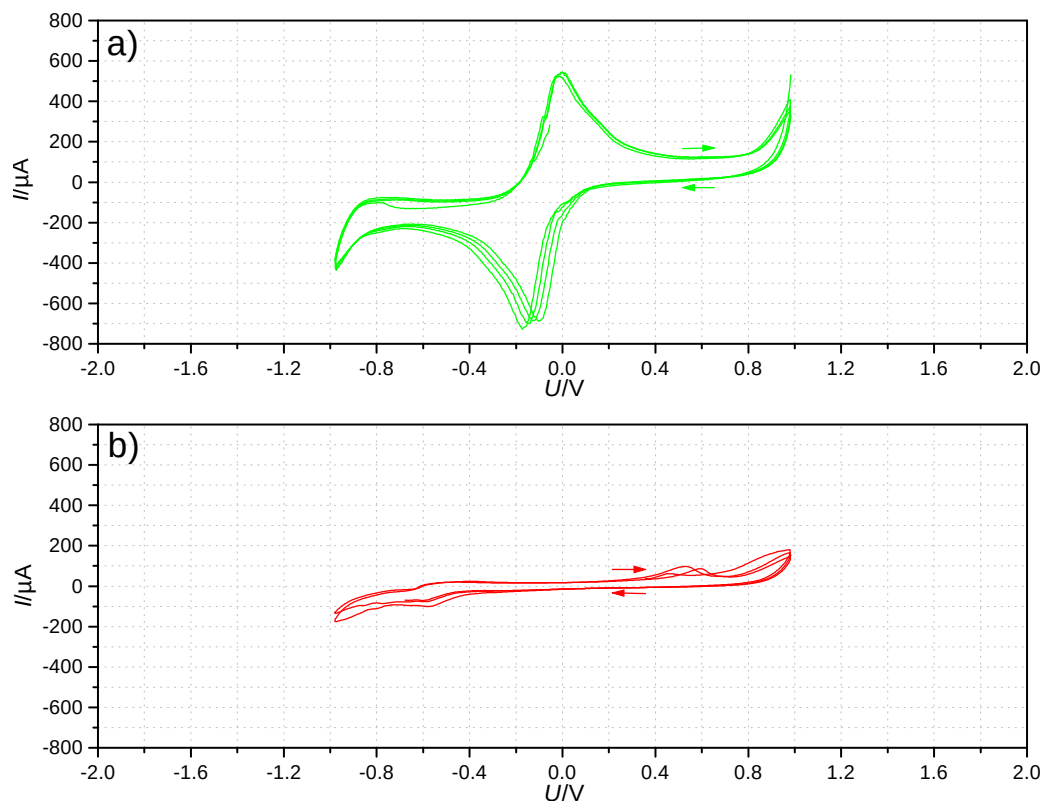


Abb. 4.27: *Cyclovoltammogramm des Systems Platin/Elektrolyt/Platin. Als Elektrolyt wurde Acetonitril mit 0.1 M TBAP und 0.01 M Cobaltocene verwendet. a) Messung sofort nach Herstellung des Elektrolyten und b) 12 Stunden später. Für beide Messungen beträgt die Vorschubgeschwindigkeit $20 \frac{\text{mV}}{\text{s}}$.*

des TBAP befindet sich der mit Elektroden und Schläuchen versehene Deckel der Weithals-Flasche ebenfalls im Trockenschrank.

Sofort nach Herstellung des Elektrolyten wird das Innere der Zelle in eine Stickstoff-Atmosphäre gebracht. Das Ende des N_2 -Auslasses befindet sich in einem Gefäß das mit Acetonitril gefüllt ist. Das Aufsteigen von N_2 -Bläschen zeigt an, daß die Zelle permanent von Stickstoff durchströmt wird.

Abb. 4.27 zeigt ein Cyclovoltammogramm mit Cobaltocene als Redoxpaar, das sofort nach Beendigung der Vorbereitungen und eines, das ca. 12 Stunden später aufgenommen wurde. Die erste Messung zeigt, daß die Redoxreaktion sehr gut abläuft, wohingegen bei der zweiten Messung die Redoxreaktion fast vollständig zum Erliegen gekommen ist. Messungen mit Chromocene zeigten das gleiche Resultat.

Es ist nicht davon auszugehen, daß während der Messung Wasser in die Zelle eingedrungen ist. Sehr viel wahrscheinlicher ist, daß im Verlauf der Vorbereitungen, trotz größter Bemühungen, 'zu feucht' gearbeitet worden ist, so daß mit der Zeit die gesamte Menge an Chromocene bzw. Cobaltocene dafür verbraucht wurde, daß Wasser im Elektrolyten zu beseitigen. Die Vermutung wird dadurch untermauert, daß 12 Stunden später der anfänglich klare Elektrolyt sehr trübe geworden war, was auf nicht lösliche Produkte aus der Reaktion zwischen Chromocene bzw. Cobaltocene und Wasser oder Sauerstoff hindeutet.

4.4.3.4 Fazit

Da Acetonitril stark hygroscopisch ist, ist das Öffnen des Gefäßes sowie das Abwiegen von Chromocene und Cobaltocene aufgrund der sehr starken Oxidierbarkeit an Luft bedenklich. Für eine wasser- und sauerstofffreie Herstellung des Elektrolyten scheint eine schutzgashaltige 'Glove-Box' unverzichtbar zu sein. Nur in einer solchen Box sollte das Öffnen sämtlicher Chemikalien-Behälter erfolgen [BYK82, CHA80, CHA81, CHA89, FLA98, LAS76, NAG82, ODE90].

Weiter ist zu bedenken, daß die Präparation der Zelle mit p-Silizium als 'Working'-Elektrode bedeutend aufwendiger ist, da dann noch dafür gesorgt werden muß, daß der InGa-Rückseiten-Kontakt nicht mit Elektrolyt benetzt wird. Soll letztlich für Leckstrom-Messungen ein nicht-wässriger Elektrolyt zum Einsatz kommen, scheint es aufgrund der beweglichen 'Counter'-Elektrode notwendig zu sein, den *gesamten* Leckstrom-Scanner unter Schutzgas zu setzen.

Auch wenn die ersten Ergebnisse vielversprechend sind, wurde das Vorhaben, einen nicht-wässrigen Elektrolyten zu verwenden, aufgegeben. Die Notwendigkeit, absolut wasser- und sauerstofffrei zu arbeiten, um so akzeptable Zeiträume für die Verwendbarkeit des Elektrolyten zu erzielen, würde den Gesamtaufbau des Leckstrom-Scanners erheblich verkomplizieren und Kosten intensivieren. Nicht zu verachten wären die Betriebskosten des Leckstrom-Scanners, da Cobaltocene, Chromocene und TBAP teure Chemikalien sind. Zudem würde eine 'Glove-Box' die Bewegungsfreiheit des Experimentators stark einschränken.

4.4.4 'Peak'-Detektor

Das Einfügen eines 'Peak'-Detektors zwischen Potentiostat und Digital-Multimeter (Keithley 2000), siehe Abb. 3.17, war die *einzig*e Maßnahme, die auf eine Verbesserung der Meßergebnisse führte. Näheres zum 'Peak'-Detektor in Abschnitt 3.2.4 und in Anhang B.

Erste Messungen zeigten, daß das Signal-zu-Rausch-Verhältnis um so schlechter wird, je weiter die Elektrode von der Oberfläche des Siliziums entfernt ist. Nur mit Hilfe des 'Peak'-Detektors war es möglich, für größere Elektrodenabstände auswertbare Ergebnisse, wie sie in Abschnitt 4.2 vorgestellt wurden, zu erzielen.

Abb. 4.28 zeigt beispielhaft den Einfluß, den der 'Peak'-Detektor auf das Resultat der Messung hat. Gemessen wurde mit der 200 μm Zylinder-Elektrode, die einen Abstand von $d = 250 \mu\text{m}$ von der Oberfläche des Wafers hat. Die 'Sample'-Zeit Δt beträgt 5 s.

Während die Messung ohne 'Peak'-Detektor lediglich zeigt, daß der Wafer im gemessenen Bereich defektbehaftet ist, ist es bei Messung mit 'Peak'-Detektor möglich, die Defektstruktur aufzulösen und das sogar ausgesprochen deutlich! Daß der Strom, der ohne 'Peak'-Detektor gemessen wird, an einigen Orten größer ist zeigt, daß der wahre Leckstrom noch etwas höher ist, als der, der mit dem 'Peak'-Detektor gemessen wird. Diese Deutung wird untermauert durch Abb. 3.6. Die Abbildung zeigt, daß

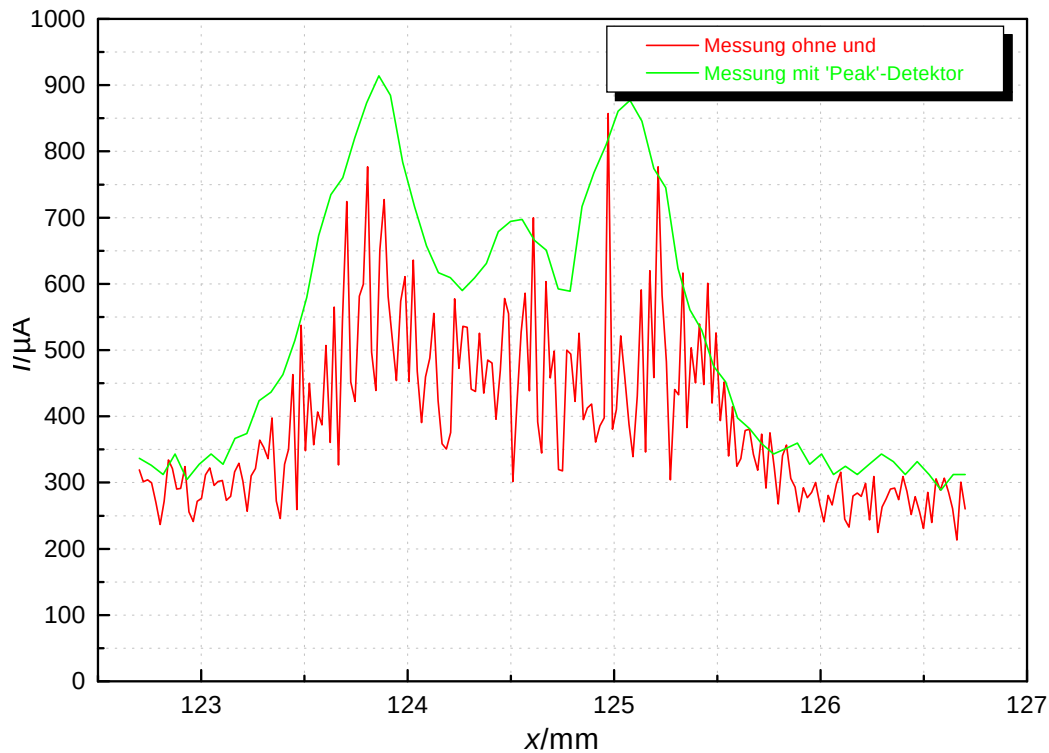


Abb. 4.28: Einfluß des 'Peak'-Detektors auf die Messung an einem Doppelkratzer. Die Messung erfolgte mit der 200 μm Zylinder-Elektrode im Abstand von 250 μm . Die 'Sample'-Zeit beträgt $\Delta t = 5$ s.

sich der Spannungswert am Ausgang des 'Peak'-Detektors selbst 8 s nach Beginn des 'Sample'-Vorganges noch erheblich vergrößern kann!

Grund für das Auftreten des kleinen 'Peaks' zwischen den beiden großen ist folgender: Beim Versuch, den zweiten Kratzer aufzubringen, berührte der Diamantgriffel aus Unachtsamkeit vorzeitig, d. h. an einer anderen Stelle als der gewünschten, die Oberfläche des Wafers, die dabei Schaden nahm.

Nur das Verwenden der relativ langen 'Sample'-Zeit von $\Delta t = 5$ s erlaubte es, die Defektstruktur in Abb. 4.28 mit der gezeigten Deutlichkeit aufzulösen. Um akzeptable Meßzeiten zu erreichen, sollten für Flächen-'Scans', speziell für die an multikristallinem Material, kleinere 'Sample'-Zeiten verwendet werden, auch wenn sich das in einem etwas stärker verrauschten Leckstrombild niederschlägt.

Abb. 4.29 zeigt eine Meßserie, bei der der Elektrodenabstand von Messung zu Messung vergrößert wurde. Alle Messungen wurden mit dem 'Peak'-Detektor durchgeführt. Die 'Sample'-Zeit beträgt $\Delta t = 1$ s. Auch hier zeigt sich wieder die sehr klare Abhebung des Signals vom Untergrund, wenn der Elektrodenabstand klein ist. Mit größer werdendem Abstand sinkt das Signal-zu-Untergrund-Verhältnis, da der Untergrund zunimmt, das Signal über dem Kratzer sich hingegen verringert. Diese Ergebnisse der Flächen-'Scans' für größere Elektrodenabstände, wie sie in Abb. 4.29 gezeigt sind, wären ohne 'Peak'-Detektor undenkbar!

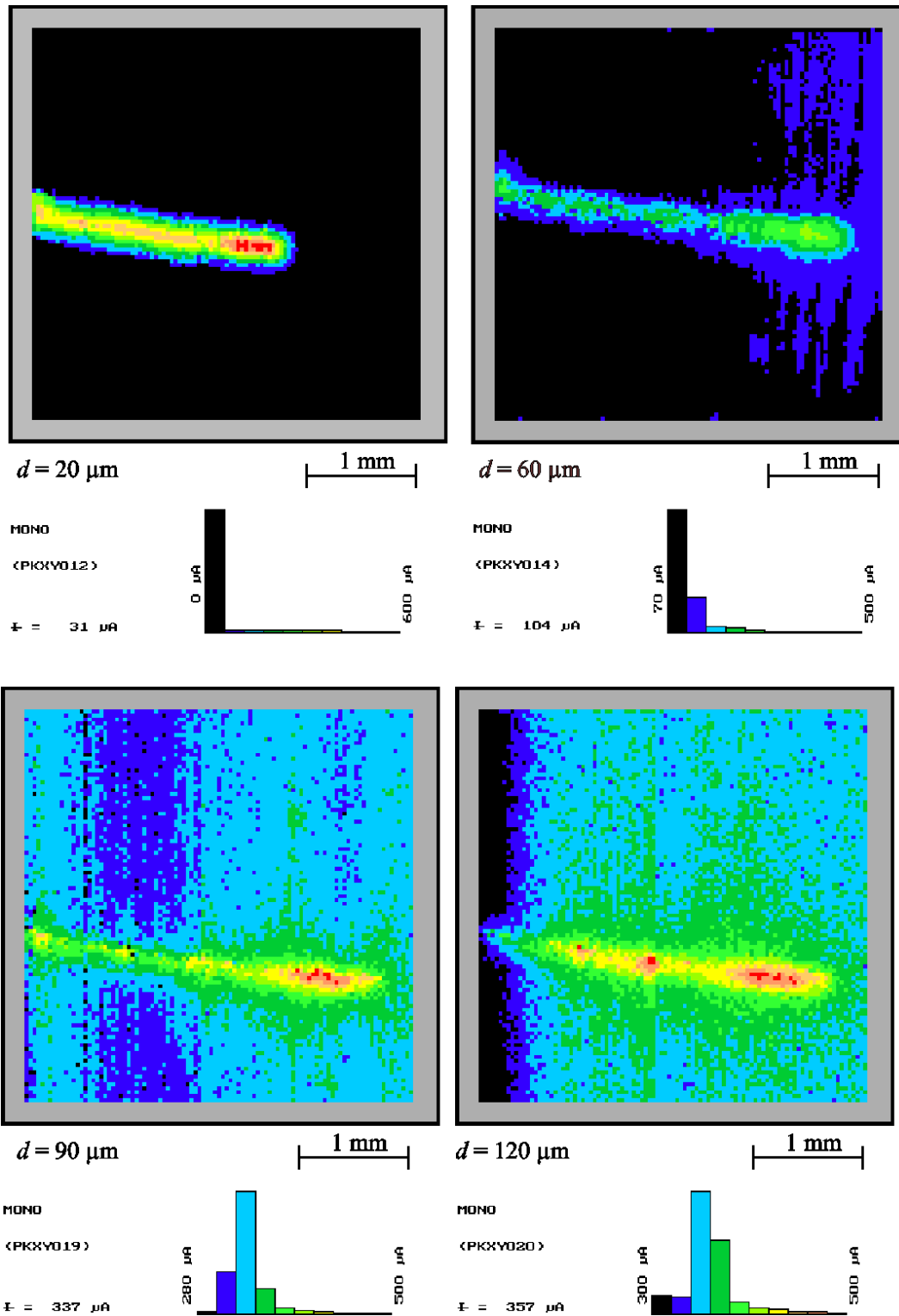


Abb. 4.29: Flächen-Scans über Kratzer mit verschiedenen Elektrodenabständen, gemessen mit der $200 \mu m$ Zylinder-Elektrode. Alle Messungen wurden mit 'Peak'-Detektor durchgeführt. Die 'Sample'-Zeit beträgt $\Delta t = 1 s$.

4.5 Vergleich mit dem integralen Leckstrom

Sind die gemessenen Ströme nicht feldinduziert, was bei der in Abb. 3.8 gezeigten Messung der Fall war, dann müßte die Summe aller lokalen Ströme der in Abb. 4.30 a) gezeigten Messung, multipliziert mit einem Geometriefaktor $0 < g < 1$, der dem Verhältnis zwischen der Fläche des gelben Quadrates und der des violetten Kreises in Abb. 4.30 b) entspricht, den integralen Leckstrom I_Σ ergeben:

$$g \cdot \sum_{n,m=1}^N I(x_n, y_m) = I_\Sigma. \quad (4.1)$$

N bedeutet die Anzahl der Meßpunkte pro 'Line-Scan'. Grund für die Multiplikation der Summe mit g ist, daß sich auf der Meßfläche die Gebiete, aus denen die Elektrode den Strom zieht, überlappen, entsprechend der violetten Kreisfläche in Abb. 4.30 b):

$$I(x_n, y_m) \approx I_0(x_n, y_m) + \sum_{\substack{|\nu - n|, |\mu - m| \leq N \\ \nu \neq n, \mu \neq m}} I(x_\nu, y_\mu). \quad (4.2)$$

In dieser Gleichung bedeutet $I_0(x_n, y_m)$ den Strom, der aus der gelben Fläche in Abb. 4.30 b) fließt. $I_0(x_\nu, y_\mu)$ sind die Ströme, die aus all denjenigen Quadraten fließen, die das gelbe umgeben.

Der Geometriefaktor könnte mit Hilfe von Abb. 4.10 bestimmt werden. Um aber Einflüsse aus einer eventuellen Ungenauigkeit des Elektrodenabstandes auszuschließen, wird das Auflösungsvermögen direkt aus der Messung bestimmt, siehe 'Line-Scan' als Inset in Abb. 4.30 a). Die Schrittweite Δx ist bestimmt durch die 'Scan'-Länge L und N :

$$\Delta x = \frac{L}{N-1} = \frac{4 \text{ mm}}{70-1} = 57.97 \mu\text{m}.$$

Mit $A = 229 \mu\text{m}$ ergibt sich nach der obenstehenden Definition für den Geometriefaktor g :

$$g = \frac{4}{\pi} \cdot \left(\frac{\Delta x}{A} \right)^2 = 0.082. \quad (4.3)$$

Unmittelbar nach der Messung wird der integrale Leckstrom I_Σ mit Hilfe des Pt-Blechtes aus Abschnitt 4.4.3.3 gemessen, das für die Messung genau über dem Kratzer positioniert wird. Um einen Kurzschluß zwischen Pt-Blech und Silizium zu vermeiden, werden an der Unterseite des Pt-Blechtes schmale Glimmer-Streifen, deren Dicke $50 \mu\text{m}$ beträgt, mit Sekundenkleber befestigt, siehe Abb. 4.30 c). Der sich aus dieser Messung ergebende Geometriefaktor beträgt:

$$g = \frac{I_\Sigma}{\sum_{n,m=1}^N I(x_n, y_m)} = \frac{21.52 \text{ mA}}{320.22 \text{ mA}} = 0.067. \quad (4.4)$$

Ein Vergleich beider Geometriefaktoren zeigt, daß sie sich lediglich um den Faktor 1.22 unterscheiden. Hierbei ist der Geometriefaktor, der sich aus der Messung ergibt,

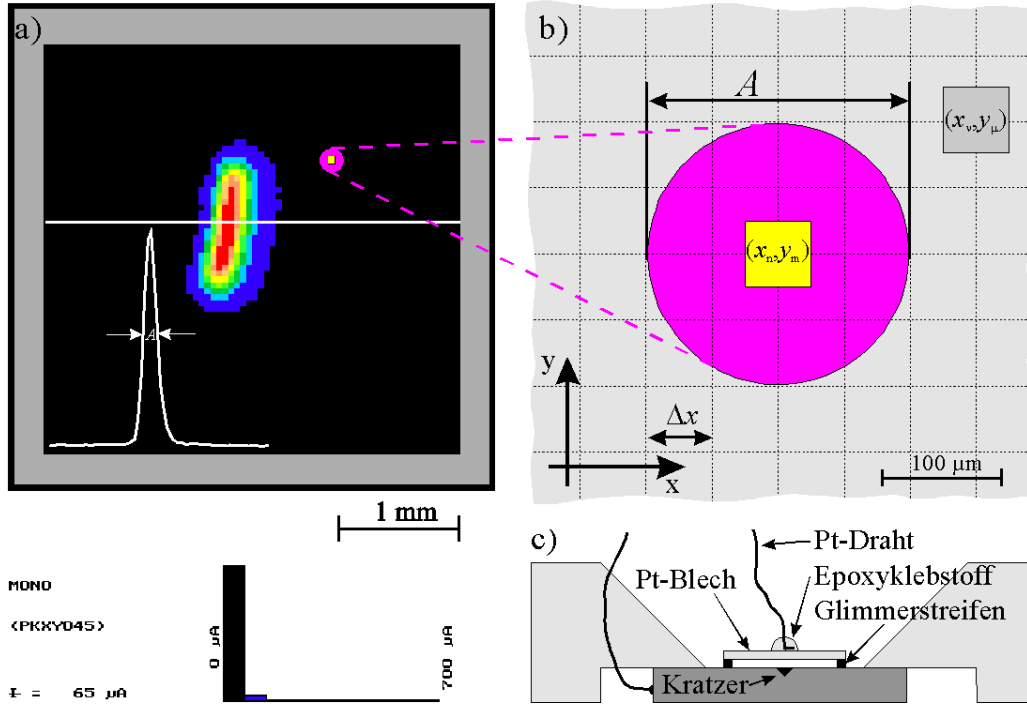


Abb. 4.30: a) Messung an Kratzer gemessen mit der 200 µm Zylinder-Elektrode im Abstand von 10 µm. b) Meßfläche mit 'Scan'-Raster. Die violette Fläche kennzeichnet den Bereich aus dem die Elektrode lokal den Strom zieht. c) Meßanordnung zur Bestimmung des integralen Leckstromes I_Σ . Pt-Draht und Pt-Blech wurden nicht zusammengelötet, sondern nur mit Epoxyklebstoff fixiert.

kleiner als der, der durch die Geometrie des Systems bestimmt wird. Dies überrascht insofern, als daß die Fläche des Pt-Blechtes größer ist als die gemessene Fläche auf dem Silizium. Grund für die Abweichung ist, daß der Wert für den integralen Leckstrom als zu klein bestimmt wurde:

1. Der Abstand des Pt-Blechtes ist bei Messung des integralen Leckstromes mit 50 µm, zuzüglich der Schichtdicke des Sekundenklebers, die mehrere 10 µm betragen kann, bedeutend weiter von der Oberfläche entfernt, als die Meßelektrode. Ohmsche Verluste im Elektrolyten haben also zur Folge, daß der Wert für I_Σ zu klein ist. Vergleiche hierzu auch Abb. 4.12.
2. Die gebildeten H_2 - und O_2 -Bläschen können aufgrund der sehr großen Fläche des Pt-Blechtes nicht so leicht entweichen, als dies für Bläschen der Fall ist, die unter der Meßelektrode gebildet werden. Dadurch verkleinert sich die effektive Meßfläche und damit ebenfalls I_Σ .

Trotz geringer Abweichung beider Geometriefaktoren können feldinduzierte Komponenten im gemessenen Strom ausgeschlossen werden. Damit sind die an den Orten (x_n, y_m) gemessenen Ströme in der Tat die lokalen Leckströme $I(x_n, y_m)$ über die Silizium-Elektrolyt-Grenzfläche!

Kapitel 5

Messungen an multikristallinem Silizium

Der Leckstrom-Scanner als leistungsfähige Einheit hat sein Funktionieren anhand zahlreicher Messungen in Kapitel 4 unter Beweis gestellt. Die eigentliche Anwendung des Leckstrom-Scanners gilt dem multikristallinen Material mit seiner Vielzahl natürlicher Defekte. Der Wunsch, multikristalline Wafer auf ihre laterale Leckstromverteilung zu untersuchen, war die eigentliche Motivation für den Aufbau des Leckstrom-Scanners.

5.1 Probenpräparation

Die Präparation multikristalliner Wafer erfordert einen höheren Aufwand. Die Schritte, die im einzelnen durchgeführt werden müssen, sind

- das Zersägen der Wafer,
- die Reinigung der Wafer und
- die Beseitigung des Säge-’Damages’,

wobei die Reihenfolge der Schritte davon abhängt, ob vor der Leckstrom-Messung mit dem Elymaten eine Diffusionslängen-Messung durchgeführt werden soll. Da mit Elymaten nur $10\text{ cm} \times 10\text{ cm}$ Wafer gemessen werden konnten, hat in diesem Fall erst die Beseitigung des Säge-’Damages’ und dann das Zersägen sowie die Reinigung für die anschließende Leckstrom-Messung zu erfolgen. Näheres zum Elymaten in Kapitel 2.

5.1.1 Zersägen der Wafer

Die Meßzelle des Leckstrom-Scanners kann keine $10\text{ cm} \times 10\text{ cm}$ Wafer aufnehmen, sondern nur Wafer mit maximal $5\text{ cm} \times 5\text{ cm}$, so daß aus den $10\text{ cm} \times 10\text{ cm}$ Wafern kleinere Stücke herauspräpariert werden müssen. Im Gegensatz zu monokristallinen Wafern können multikristalline Wafer nicht durch Spalten in einzelne Probenstücke zerteilt werden. Da ein multikristalliner Wafer beim Versuch ihn zu spalten, unkontrolliert splintern würde, müssen hierfür andere Techniken zum Einsatz kommen.

Für die Zerteilung des Wafers wird eine Kristallsäge verwendet. Um die Gefahr des Berstens während des Sägens zu minimieren, wird der Wafer auf einer ca. 10 cm × 10 cm großen Aluminiumplatte mit Kerzenwachs fixiert, wobei zu beachten ist, daß der Wafer an jedem Ort mit Wachs benetzt wird. Um dies sicherzustellen, ist wie folgt zu verfahren: Die Aluminiumplatte wird auf einer Heizplatte so weit erwärmt, daß aufgetragenes Wachs gut verläuft. Nach gleichmäßiger Verteilung des Wachses wird der Wafer auf die Aluminiumplatte gebracht und durch leichtes Drücken dafür gesorgt, daß überschüssiges Wachs zwischen Aluminiumplatte und Wafer beseitigt wird. Nach dem Sägevorgang wird der Wafer durch erneutes Erwärmen von der Aluminiumplatte gelöst.

5.1.2 Reinigung der Wafer

Als erstes müssen Wachsreste vom Wafer entfernt werden, was in zwei Schritten erfolgt: Im ersten Schritt wird der Wafer in deionisiertem Wasser ausgekocht, wodurch der größte Teil des Wachses entfernt werden kann. Anschließend wird der Wafer in einem zweiten Schritt mit Chloroform bei ca. 50 °C für 10 min im Ultraschallbad gereinigt. Chloroform als Reinigungsmittel hat sich als besonders effektiv erwiesen, was die Beseitigung hartnäckiger Verschmutzungen wie Wachsreste betrifft [KOR95, MOR94]. Auch zum Beseitigen von kleinen Schmutzpartikeln, die elektrostatisch sehr fest an der Waferoberfläche haften können und daher schwer zu beseitigen sind, leistet das Ultraschallbad gute Dienste. Das Erwärmen der Reinigungssubstanz (im allgemeinen ein organisches Lösungsmittel) auf annähernd Siedetemperatur fördert den Reinigungsprozeß zusätzlich [LU90]. Nach Beendigung der Ultraschallreinigung wird der Wafer mit *demselben* Lösungsmittel, also mit Chloroform (in einer Spritzflasche), gründlich abgespült. Weil sich Chloroform nicht mit Wasser mischt, muß vor dem Abspülen unter deionisiertem Wasser noch ein Spülen mit Aceton erfolgen [KOR95].

Durch Reinigen in Chloroform, das sich, wie oben beschrieben, sehr gut zum Beseitigen von Wachsresten eignet, wird der Wafer aber andererseits mit organischem 'Schmutz' kontaminiert. Das Entfernen von Chloroform sowie von verbliebenen ionischen Stoffen oder durch intensives Spülen mit deionisiertem Wasser zusätzlich aufgebrachten ionischen Partikeln, erfordert eine sogenannte *basische RCA-Reinigung* [KER70].

Bezüglich ionischer Kontaminierungen beseitigt eine basische RCA-Reinigung im wesentlichen Schwermetallionen, aber nicht die Ionen der Alkalimetalle. Diese Ionen können unter Einfluß elektrischer Felder oder bei erhöhten Temperaturen auf der Siliziumoberfläche wandern und metallische Inseln bilden, die lokal Inversionsschichten oder erhöhten Leckstrom zur Folge haben können. Sehr effektiv beseitigen lassen sich Ionen der Alkalimetalle mit einer sogenannten *sauren RCA-Reinigung* [KER70].

In beiden Fällen ist H₂O₂ Bestandteil der Reinigungslösungen, das die anhaftenden organischen Moleküle und Ionen oxidiert. Da die

Reaktionsprodukte meist nicht wasserlöslich sind, ist ein sogenannter *Komplexbildner* notwendig, damit die oxidierten Stoffe als wasserlösliche Komplexe in Lösung gehen

können und ein erneutes Abscheiden auf der Oberfläche unterbunden wird. Die Rolle des Komplexbildners kommt bei der basischen RCA den NH_4OH - und bei der sauren RCA den HCl -Molekülen zu [BOE79, GHA82, HOW85, MAL94].

Tab. 5.1 zeigt die Zusammensetzungen der Reinigungslösungen für die basische und die saure RCA-Reinigung, ergänzt durch Angaben zu Kontaminierungen, die mit der jeweiligen RCA-Reinigung beseitigt werden können.

	basische RCA	saure RCA
beseitigt	organischen Schmutz und Schwermetallionen wie Cu, Ag, Ni, Co und Au	Ionen der Alkalimetalle Li, Na und K
Zusammensetzung	$\text{H}_2\text{O}:\text{H}_2\text{O}_2:\text{NH}_4\text{OH}$	$\text{H}_2\text{O}:\text{H}_2\text{O}_2:\text{HCl}$
Anteile	5:1:1 - 7:2:1	6:1:1 - 8:2:1
verwendet	5:2:1	6:2:1

Tab. 5.1: *Basische und saure RCA-Reinigung. Die Reinigung erfolgt bei 75-85 °C für 10-20 min. Nach jedem Schritt erfolgt das Spülen mit deionisiertem Wasser [KER70, LIN98a, ROS95, YAN96].*

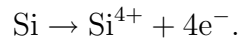
Um die Reinigung der Wafer und das Beseitigen des Säge-’Damages’ in einfacher Weise durchführen zu können, wurde hierfür ein spezielles Bad konstruiert, das gleichermaßen zum Reinigen und Abätzen der Wafer verwendet werden kann. Dieses besteht aus zwei PVC-Gefäßen in die je eine Platin-Heizung integriert ist, siehe Abb. 5.3. Auf diese Weise kann auf die Verwendung kontaminationsträchtiger externer Heizelemente verzichtet werden, die zudem thermisch schlecht an das Bad gekoppelt sind. Die Beize zum Beseitigen des Säge-’Damages’ ist äußerst aggressiv, so daß für ihre Erwärmung gänzlich auf handelsübliche Heizelemente verzichtet werden mußte. Eine Beschreibung des Aufbaus der Platin-Heizung findet sich in Anhang C. Eine Möglichkeit, die Beize zu erwärmen besteht darin, das Ätzbad mit einer leistungsstarken Lampe zu beleuchten [LIP97]. Aufgrund der sehr schlechten thermischen Ankopplung kann die Aufwärmphase, je nach Umgebungstemperatur und Füllmenge, mehrere Stunden betragen!

5.1.3 Beseitigung des Säge-’Damages’

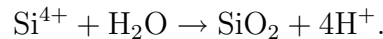
Im Gegensatz zu monokristallinem Silizium, das durch Ziehverfahren aus der Schmelze gewonnen wird, wird multikristallines Silizium, das nur für die Herstellung von Solarzellen(-Modulen) technische Relevanz hat, im Blockgußverfahren hergestellt. Bei diesem Verfahren wird das flüssige Silizium in einen erhitzten Behälter gegossen, dessen Wände anschließend kontrolliert auf Raumtemperatur abgekühlt werden. Zwar haben aus diesem Material prozessierte Solarzellen einen kleineren Wirkungsgrad ($\approx 15.5\%$) als solche aus monokristallinem Material ($\approx 17.5\%$), dafür sind die Herstellungskosten aber deutlich geringer [NOW95]! Der gegossene Siliziumblock wird in Säulen und diese

wiedermum werden in die einzelnen Wafer zersägt. Da die gewonnenen Wafer quadratisch sind, erlauben sie, integriert in Solarzellen-Module, eine sehr viel bessere Ausnutzung der bestrahlten Fläche, als dies mit monokristallinem Material möglich wäre, dessen Wafer rund sind. Der durch den Sägeprozeß erzeugte sogenannte *Säge-’Damage’* des Wafers muß vor der Messung entfernt werden, da sonst die natürliche Defektstruktur des Wafers mittels einer Leckstrom-Messung nicht aufgelöst werden kann.

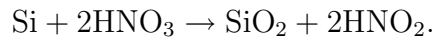
Die Beseitigung des Säge-’Damages’ erfolgt chemisch mit einer Politurbeize, wofür eine Mischung aus Fluß-, Salpeter- und Essigsäure verwendet wird. Um Siliziumatome aus dem Kristallverbund zu lösen, bedarf es einer oxidierenden Spezies, die die Oberflächenatome ionisiert:



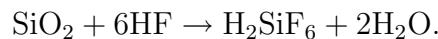
Diese oxidierende Spezies ist Salpetersäure, die die vier Elektroden des Siliziums aufnimmt¹. Daraufhin reagiert das vierfach positiv geladene Siliziumatom mit Wasser, wobei Siliziumoxid entsteht:



In einer Reaktion zusammengefaßt, kann die Oxidation des Siliziums geschrieben werden als:



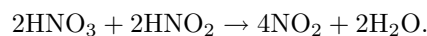
Wird das Silizium in konzentrierte oder verdünnte Salpetersäure getaucht, wird die Auflösungsreaktion kurz nach Beginn wieder zum Erliegen kommen. Der Grund hierfür ist, daß Siliziumoxid in Wasser nahezu unlöslich ist. Die Schicht aus Siliziumoxid würde also eine weitere Reaktion mit der Salpetersäure verhindern. Deshalb wird eine Mischung aus Salpeter- und Flußsäure verwendet. Flußsäure ist als einzige Chemikalie in der Lage, Siliziumoxid aufzulösen, vgl. Abschnitt 1.3.4:



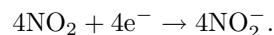
Für die Einstellung der Ätzrate sollte zum Verdünnen des Salpetersäure-Flußsäure-Gemisches nicht Wasser verwendet werden, weil bei zu großem Wasseranteil die Salpetersäure ihre oxidierende Wirkung verliert. Deshalb wird Essigsäure verwendet, die selbst nicht in die chemische Reaktion involviert ist [TUC75].

Die Abhängigkeit der Ätzrate von der Zusammensetzung bei $T = 25^\circ\text{C}$ zeigt Abb. 5.1. Wie zu erwarten, wird die Ätzrate um so höher, je kleiner der Essigsäureanteil und je näher das Verhältnis zwischen Salpeter- und Flußsäure bei 25:75 liegt².

¹ Die Reduktion von HNO_3 erfolgt nicht direkt, sondern über einen Zwischenschritt: Zunächst reagiert Salpetersäure mit salpetriger Säure:



Die vier gebildeten NO_2 -Moleküle nehmen schließlich die Elektroden des Siliziums auf [TUC75]:



² Zwei HNO_3 -Moleküle sind erforderlich, um ein SiO_2 -Molekül zu schaffen. Um dieses SiO_2 -Molekül

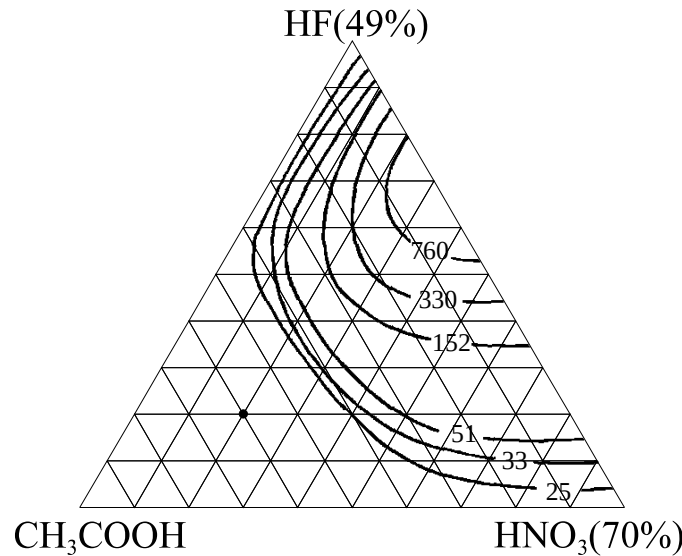


Abb. 5.1: Linien konstanter Ätzraten von Silizium im System $\text{HNO}_3\text{:HF:CH}_3\text{COOH}$ bei $T = 25\text{ }^\circ\text{C}$. Die Zahlen auf den Kurven bezeichnen die Ätzrate in $\frac{\mu\text{m}}{\text{min}}$ [BOG67, TUC75]. Der Punkt markiert die Zusammensetzung der Beize, die zum Beseitigen des Säge-'Damages' verwendet wird.

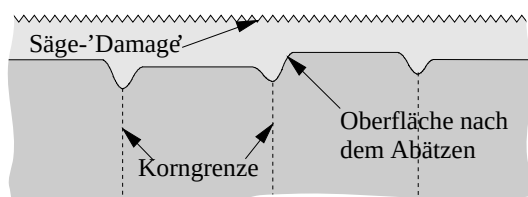


Abb. 5.2: Das Ätzen in der Politurbeize führt zu Unebenheiten der Waferoberfläche. In Umgebung von Korngrenzen entstehen ausgeprägte Vertiefungen [BOG67].

Vertiefungen entstehen, siehe Abb. 5.2. Wird der Wafer während der Handhabung geringfügig verbogen, konzentrieren sich unter diesen Vertiefungen mechanische Spannungen, so daß die Gefahr eines Bruches erhöht ist!

Zur Beseitigung des Säge-'Damages' ist das Abtragen von 10-20 μm ausreichend [LIP97]. Folgende Zusammensetzung für die Politurbeize, die auch als CP6 bezeichnet

Ist der Essigsäureanteil nicht zu hoch, erfolgt das Ätzen im System $\text{HNO}_3\text{:HF:CH}_3\text{COOH}$ reaktionslimitiert³. Da bei Reaktionslimitierung die lokalen Ätzraten von den Kornorientierungen abhängig sind, sollte vom Wafer möglichst wenig Material abgetragen werden, um ausgeprägte Unebenheiten über dem Wafer zu vermeiden. Außerdem bewirkt Reaktionslimitierung, daß an defektbehafteten Bereichen, wie Korngrenzen, deutliche Ver-

in Lösung zu bringen, müssen sechs HF-Moleküle aufgewendet werden. Daher ist das optimale stöchiometrische Verhältnis zwischen HNO_3 und HF 1:3.

³ Um Diffusionslimitierung zu erhalten, bedarf es eines sehr hohen Essigsäureanteils, weil die an die Reaktion gekoppelte Gasentwicklung eine Durchmischung der Beize bewirkt, die das Diffusionsprofil zerstört.

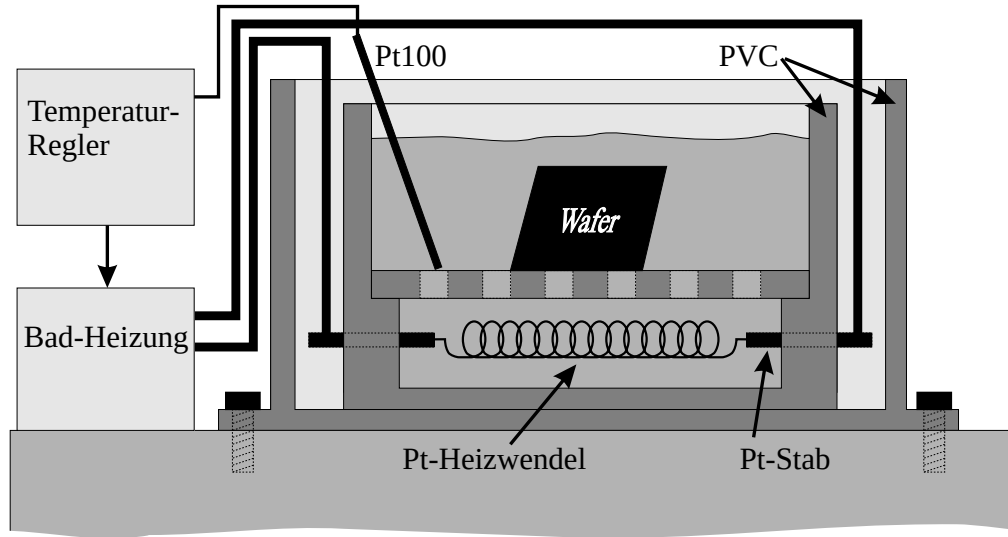
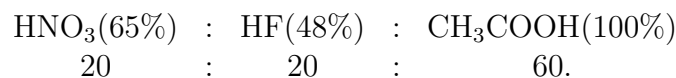


Abb. 5.3: Temperaturgeregeltes Bad zum Beseitigen des Säge-‘Damages’ und für die RCA-Reinigung multikristalliner Wafer. Das Thermoelement (Pt100) ist von einer eng anliegenden Teflonschicht umgeben. Die mit Löchern versehene PVC-Platte verhindert direkten Kontakt zwischen Heizwendel und Wafer. Das äußere PVC-Gefäß soll bei Undichtigkeiten die austretende Flüssigkeit auffangen und ist zur Erhöhung der Sicherheit fest mit einer massiven Unterlage verschraubt.

wird [BOG67], ist zu verwenden:



Nach Abb. 5.1 beträgt die Ätzrate für diese Zusammensetzung ca. $5 \frac{\mu\text{m}}{\text{min}}$.

Die Wafer werden in derselben Vorrichtung, in der auch die Reinigung erfolgt bei $40 - 50^\circ\text{C}$ geätzt, siehe Abb. 5.3. Kurz nachdem der Wafer in das Ätzbad getaucht worden ist, kommt es zu heftiger Gasentwicklung. Während dieser Phase wird der Säge-‘Damage’ mit hoher Rate weggeätzt. Das Nachlassen der Gasentwicklung nach maximal zwei Minuten zeigt an, daß der Säge-‘Damage’ entfernt worden ist. Der Wafer sollte dann noch für ca. 10-15 s in der Beize belassen und der Ätzvorgang erst danach beendet werden. Anschließend wird der Wafer unter deionisiertem Wasser abgespült. Da die frisch überätzten Oberflächen sehr sauber sind, kann auf eine anschließende Reinigung verzichtet werden.

5.2 Eigenschaften von Korngrenzen

Wie in Abschnitt 5.1.3 näher beschrieben, wird multikristallines Silizium im Blockgußverfahren hergestellt. Der auf diese Weise hergestellte Siliziumblock besteht aus zahlreichen monokristallinen Bereichen gleicher Struktur (*Kristallite*), die voneinander

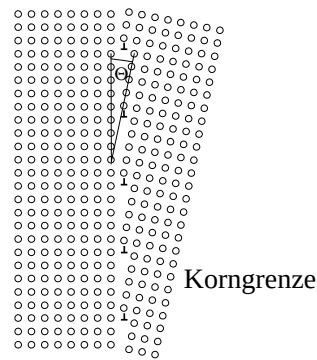


Abb. 5.4: Eine Kleinwinkel-Korngrenze entsteht durch Aneinanderreihung von Versetzungen $\perp$ mit einer Fehlorientierung der benachbarten Kristallbereiche um den Winkel Θ .

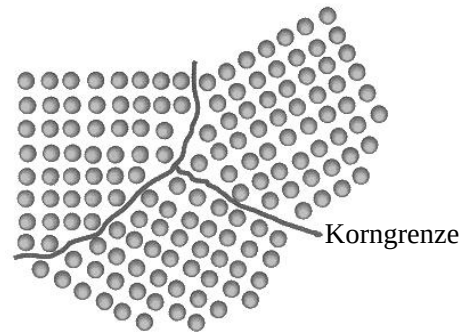


Abb. 5.5: Großwinkel-Korngrenzen. Die in der Nähe der Korngrenzen gelegenen Atome sind in ihrer Gleichgewichtslage gestört (aus [ASK96]).

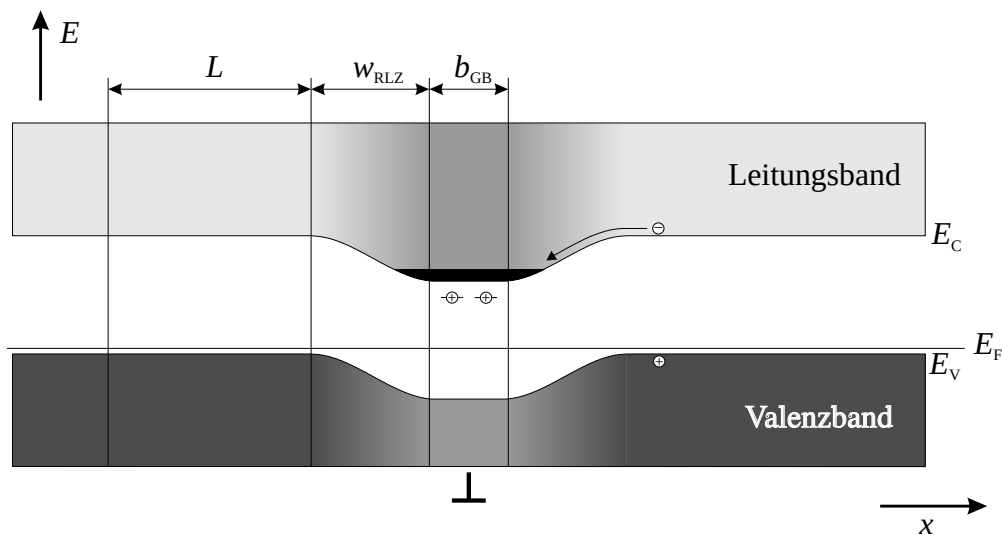


Abb. 5.6: Banddiagramm einer zwei Kristallite trennenden Korngrenze in p-Silizium. Die Korngrenze sammelt die in ihrer näheren Umgebung befindlichen Elektronen ein, wodurch sich am Ort der Korngrenze ein 'Elektronensee' ausbildet (symbolisiert durch die schwarzen Fläche an der unteren Leitungsbandkante am Ort der Korngrenze). Weitere Erläuterungen siehe Text.

durch Grenzflächen getrennt sind. Diese Grenzflächen werden als *Korngrenzen* bezeichnet und bestehen aus schmalen Übergangszonen mit gestörter atomarer Anordnung, worin die fehlgelagerten Atome in ihrer Umgebung Druck- oder Zugspannungen bewirken [ASK96].

Unterscheiden sich die Orientierungen von zwei benachbarten Kristalliten um weniger als $\Theta = 15^\circ$, so spricht man von einer *Kleinwinkel-Korngrenze* [BAR73]. Sie bestehen modellhaft aus aneinander gereihten Versetzungen, die eine schwache Fehlори-

entierung der angrenzenden Gitterbereiche verursachen, siehe Abb. 5.4. Das Modell der aneinandergereihten Versetzungen verliert seine Gültigkeit, wenn sich die Orientierungen zweier benachbarter Kristallite um mehr als $\Theta = 15^\circ$ unterscheiden, weil die Versetzungen, die die Korngrenze bilden, dann so dicht beisammen liegen, daß sie kaum mehr unterscheidbar sind. Diese Korngrenzen werden als *Großwinkel-Korngrenzen* bezeichnet, wofür Abb. 5.5 Beispiele zeigt [ASK96, BAR73].

In p-Silizium können die Versetzungen und damit die Korngrenzen Donator-Eigenschaften haben. Bei genügend hohen Temperaturen geben die Versetzungen der Korngrenze Elektronen an das Leitungsband ab, die dort die p-Leitung des Siliziums teilweise kompensieren. Dadurch werden sich in Umgebung der Korngrenze die Bandkanten nach unten verschieben, so daß dort der Abstand des Fermi-Niveaus E_F von der oberen Valenzbandkante E_A etwas größer ist, vgl. Abb. 5.6. Die Raumladungszonen der Weite w_{RLZ} , die zu beiden Seiten der Korngrenze verlaufen, haben eine sammelnde Wirkung auf Minoritätsladungsträger. Für alle Minoritätsladungsträger, deren Abstand von der äußeren Grenze der Raumladungszone weniger als die Diffusionslänge L beträgt, ist die Wahrscheinlichkeit von null verschieden, vom 'absaugenden' elektrischen Feld der Korngrenze erfaßt zu werden, siehe Abb. 5.6. Für Elektron-Loch-Paare, die innerhalb der Weite b_{GB} thermisch generiert werden, ist die Wahrscheinlichkeit relativ hoch, daß das Loch vom elektrischen Feld der Raumladungszone aus der näheren Umgebung der Korngrenze vertrieben wird. Deshalb ist die Lebensdauer der Minoritätsladungsträger im 'Elektronensee' vergleichsweise hoch, weil sie aufgrund der sehr geringen Löcherdichte im Valenzband nicht ohne weiteres rekombinieren können. Durchstößt eine derartige Korngrenze die Oberfläche des Kristalls, so können die Ladungsträger des 'Elektronensees', begrenzt durch die Geschwindigkeit des Ladungstransferprozesses an der Silizium-Elektrolyt-Grenzfläche, als Leckstrom den Kristall verlassen.

Jede Versetzungslinie, die die Oberfläche eines Polykristalls durchstößt, bedeutet einen Oberflächendefekt. Die Oberflächenzustandsdichte am Ort einer Großwinkel-Korngrenze ist bedeutend höher, als am Ort einer Kleinwinkel-Korngrenze. Eine lokal erhöhte Oberflächenzustandsdichte bedeutet aber nicht zwangsweise auch einen erhöhten Leckstrom. Die Höhe des fließenden Leckstromes wird unter anderem davon abhängen, ob die Versetzungen elektrisch geladen sind. Tragen die Defekte einer Korngrenze an der Oberfläche eine negative Ladung, so verschwindet bei hinreichend hoher Defektdichte die Weite der Raumladungszone, weil dann die positive Flächenladung im Elektrolyten allein durch die negative Oberflächenladung kompensiert werden kann und hierfür ionisierte Akzeptoren nicht mehr von Nöten sind. Auch in diesem Fall wird am Ort einer Korngrenze ein erhöhter Leckstrom über die Grenzfläche des Kontaktes fließen.

Da die Oberflächenzustandsdichte eines jeden Kristallits von seiner Orientierung zur Waferoberfläche abhängt, wird auch der Leckstrom aus einem Kristallit von dessen Orientierung abhängig sein. Wenn beispielsweise die Leckströme zweier benachbarter Körner sehr verschieden sind, so wird dies auch für die Orientierungen beider Kristallite gelten. In diesem Fall werden beide Körner durch eine Großwinkel-Korngrenze voneinander getrennt sein.

Umgekehrt folgt aus der Gleichheit der Leckströme zweier benachbarter Kristallite nicht, daß sie von einer Kleinwinkel-Korngrenze getrennt werden. Denn aus zwei völlig unterschiedlichen Kristallitorientierungen kann eine nahezu identische Oberflächenzustandsdichte resultieren, womit auch die Leckströme annähernd gleich wären.

5.3 Messungen

Die laterale Dynamik im Leckstrom von multikristallinen Wafern wird im wesentlichen hervorgerufen durch deren natürliche Defektstruktur (Korngrenzen, 'Shunts' etc.). Allerdings können andere Einflüsse den Leckstrom erheblich vergrößern, so daß das Auflösen der Leckstromstruktur auch ohne 'Peak'-Detektor möglich ist. Dies ist beispielsweise für Messungen an definiert aufgetragenen Kratzern der Fall, siehe Kapitel 4. Ein nicht vollständig beseitigter Säge-'Damage', Kurzschlüsse zwischen Wafer und Elektrode während der Messung, ein zerbrochener Wafer und letztlich auch Kratzer zählen zu den Einflüssen, die die Kornstruktur eines multikristallinen Wafers überdecken können und deshalb auch ohne 'Peak'-Detektor nachweisbar sind. Auf den folgenden Seiten werden anhand von Messungen, die vor dem Bau des 'Peak'-Detektors

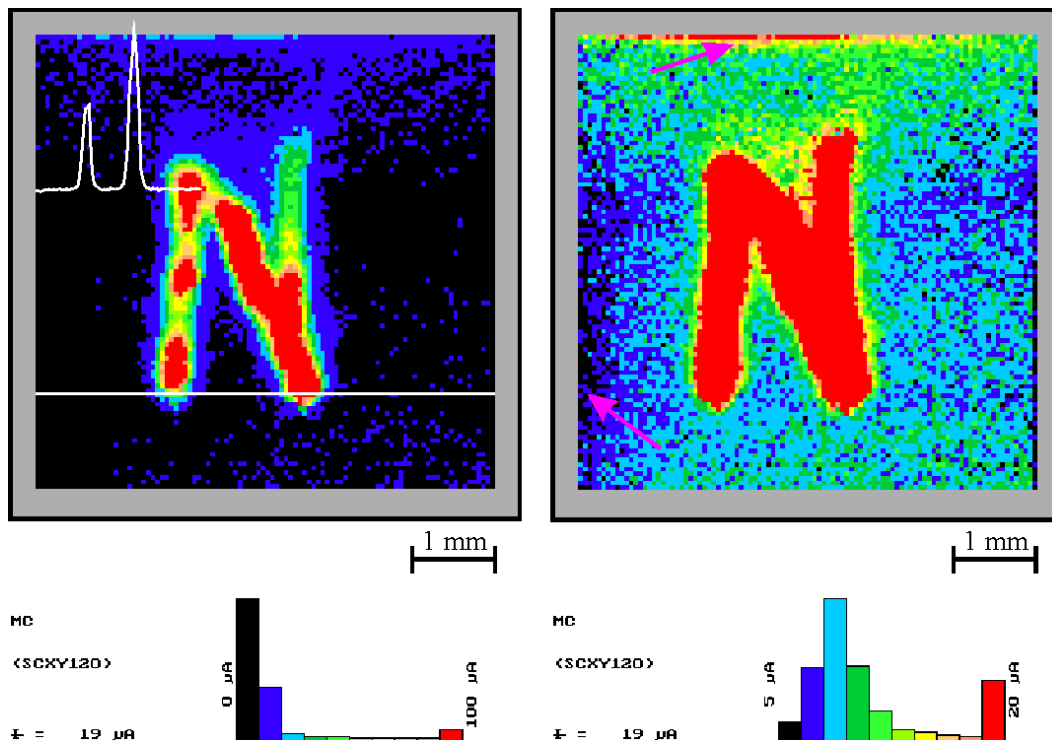


Abb. 5.7: Messung an multikristallinem Silizium mit eingekratztem 'N'. Im rechten Bild wurde umskaliert, um die Struktur des Untergrundes besser sichtbar zu machen. Die violetten Pfeile markieren stärkere Abweichungen vom Mittelwert des Untergrundes. Messung ohne 'Peak'-Detektor mit der 200 μm Zylinder-Elektrode. Der Elektrodenabstand beträgt 10 μm .

durchgeführt worden sind, Beispiele gezeigt.

Messungen an monokristallinem Silizium mit eingekratzten Teststrukturen liefern, besonders an Material mit geringem spezifischen Widerstand, im allgemeinen sehr gute Ergebnisse. Das verwendete multikristalline Silizium hatte einen spezifischen Widerstand von ca. $5 \Omega\text{cm}$, so daß das Messen einer eingekratzten Teststruktur auch auf diesem Material möglich sein sollte. Daher wurde in einer der ersten Messungen versucht, ein in dieses Material gekratztes 'N' zu messen. Das Resultat der Messung zeigt Abb. 5.7.

Das linke Bild zeigt, daß sich das 'N' überraschend deutlich vom Untergrund abhebt. Ursache der an einigen Stellen des 'N' verringerten Ströme war nicht die Gasbläschenentwicklung zwischen Elektrode und Wafer, obwohl ohne 'Peak'-Detektor (siehe Abschnitt 3.2.4) gemessen wurde, sondern unzureichendes Kratzen. Dafür spricht auch, daß die Ströme am Ort des 'N' kaum verrauscht sind.

Das rechte Bild in Abb. 5.7 zeigt dieselbe Messung, allerdings wurde sie umskaliert, um die Struktur des Untergrundes besser aufzulösen. Daß die Messung (entfernt vom 'N') nur Rauschen und kaum Struktur zeigt liegt daran, daß ohne 'Peak'-Detektor gemessen wurde. Die Messung wurde mit der $200 \mu\text{m}$ Zylinder-Elektrode durchgeführt, so daß das Gebiet, aus dem die Elektrode den Strom zieht nur relativ schwach vom Abstand abhängt, siehe Abb. 4.10. Des weiteren sind aufgrund der kleinen Ströme auch die Spannungsabfälle über dem Elektrolytwiderstand vergleichsweise klein. Solange Elektrode und Wafer einander nicht berühren, sollten sich Änderungen des Elektrodenabstandes, beispielsweise hervorgerufen durch Waferunebenheiten oder einen in der Zelle leicht verspannten Wafer, in einer Messung nicht bemerkbar machen. Damit können die geringen Leckstromvariationen einzelnen Körnern des Wafers zugeordnet werden.

Daß Dickenunterschiede über dem Wafer nicht immer vernachlässigbar sind und zu Problemen bei der Messung führen können, zeigt Abb. 5.8. Nach dem Zersägen des Wafers und der RCA-Reinigung wurde in diesem Beispiel versäumt, den Temperatur-Regler auf 45°C zu stellen, so daß das Überätzen des Wafers bei 75°C erfolgte! Dies führte neben bedeutend größeren Unebenheiten über dem Wafer auch dazu, daß der Wafer relativ dünn wurde.

Die vergleichsweise starken Unebenheiten des Wafers hatten zur Folge, daß die vier Punkte mit $z = 0$ nicht auf einer Ebene lagen. In diesem Fall war das Korn oben rechts in Abb. 5.8 'zu hoch', bedingt durch eine im Vergleich zu den anderen drei Eckpunkten kleinere Ätzrate. Die Konsequenz war, daß die Elektrode während der Messung an der gegenüberliegenden Ecke beim Nachführen auf den Wafer gefahren (siehe Abb. 5.8 unten links) und dadurch ein Kurzschluß zwischen Elektrode und Wafer verursacht wurde.

Neben dem eben geschilderten Problem, das aus Unebenheiten des Wafers resultieren kann, stellt auch der Einbau des Wafers in die Zelle einen kritischen Moment im Verlauf der Meßvorbereitungen dar. Leicht kann der Wafer beim Festziehen der Zellschrauben zerbrechen. Abb. 5.9 zeigt ein Beispiel, wo dies beim Festziehen der oberen rechten Zellschraube geschah. Daß der Wafer in der Tat zerbrochen war, zeigte sich

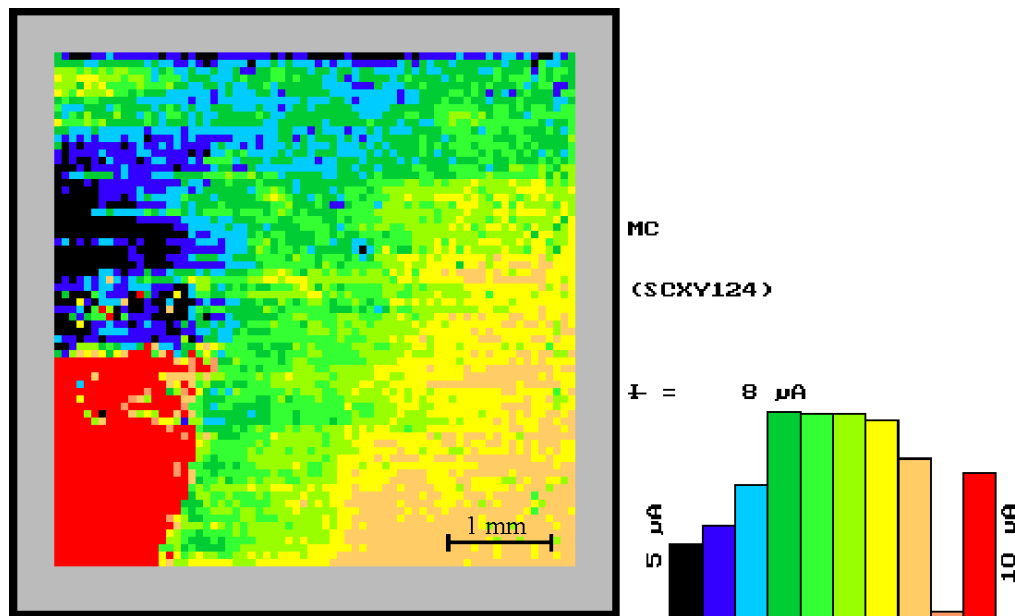


Abb. 5.8: Berührung von Meßelektrode und Wafer als Folge eines zu großen Materialabtrages während des Ätzens in der Politurbeize. Messung ohne 'Peak'-Detektor mit der 200 μm Zylinder-Elektrode. Der Elektrodenabstand beträgt 10 μm . Weitere Erläuterungen siehe Text.

nach den Messungen, als der Wafer aus der Zelle ausgebaut wurde.

Da es sich hierbei um einen 'ordnungsgemäß' präparierten Wafer handelte, ist es unwahrscheinlich, daß während dieser Messung dasselbe passierte wie in Abb. 5.8, d. h. daß die Elektrode dort auf dem Wafer aufsaß, wo die Ströme sehr viel höher sind als an anderen Stellen der Messung. Als während einer Wiederholungsmessung die Ströme an derselben Stelle wie zuvor ebenfalls stark anstiegen, wurde die Messung vorübergehend gestoppt und mit dem Finger ganz leicht gegen die Elektrode gedrückt. Hätte die Elektrode mechanischen Kontakt zum Wafer, würde das im Finger als leicht kratzendes Gefühl zu spüren sein⁴! Da dem nicht so war, konnte ein mechanischer Kontakt zwischen Elektrode und Wafer ausgeschlossen werden.

Die Stromdrift im linken Bild von Abb. 5.9 zeigt, daß die Oberflächenpassivierung noch nicht abgeschlossen war und daher mit der Messung zu früh begonnen wurde. Deshalb nahmen die Ströme vom Beginn der Messung (Meßpunkt oben links) bis zum Ende der Messung (Meßpunkt unten rechts) etwas ab. Die Zeit, die vom Beginn der ersten Messung bis zum Beginn der zweiten Messung verstrich, war lang genug, um die Oberflächenpassivierung abzuschließen. Aus diesem Grund zeigt die zweite Messung keine Drift in den Strömen, die zudem kleiner sind.

Um zu prüfen, ob der Leckstrom-Scanner überhaupt in der Lage ist, die natürliche

⁴ Daß das Aufsitzen der Elektrode als leicht kratzendes Gefühl im Finger zu spüren ist, wenn sie ganz leicht bewegt wird, wird gerade dafür ausgenutzt, beispielsweise nach einem Elektroden-Wechsel, die Elektrode so zu justieren, daß während der Initialisierung des z-Translators nur kurze Wege gefahren werden müssen. Näheres in Abschnitt 3.4.

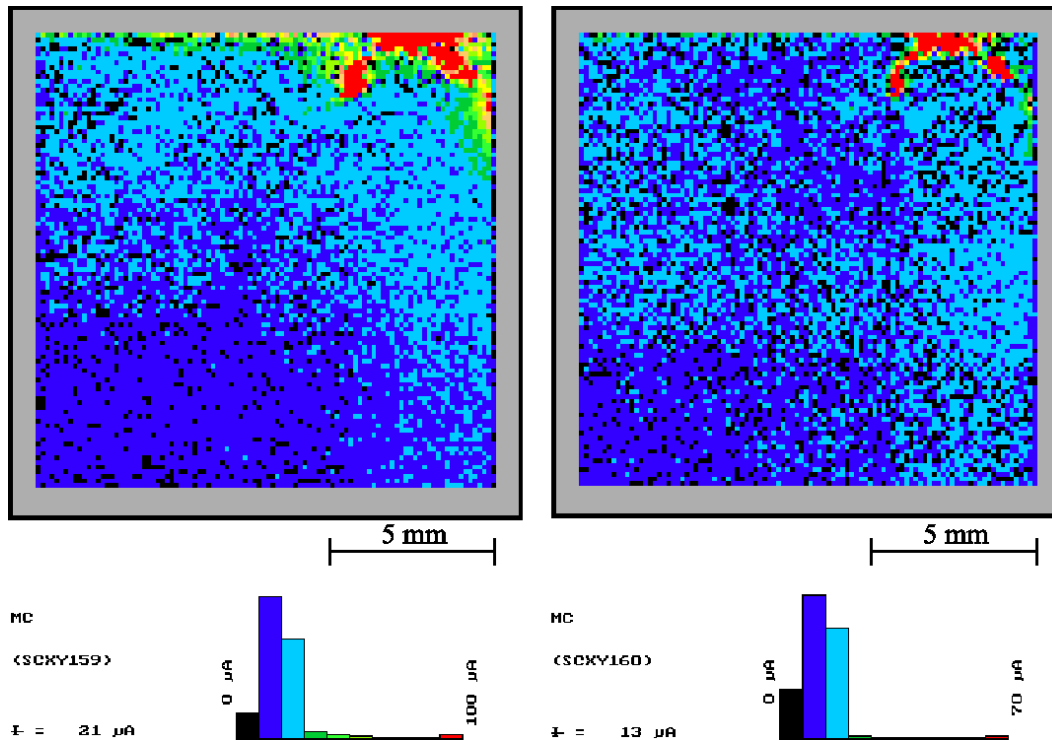


Abb. 5.9: Messung an einem multikristallinen Wafer mit der 200 μm Zylinder-Elektrode und ohne 'Peak'-Detektor. Der Elektrodenabstand beträgt 10 μm . Die lokal erhöhten Ströme konnten, wie das rechte Bild zeigt, reproduziert werden. Nähere Erläuterungen zu den Stromerhöhungen siehe Text.

Leckstromstruktur eines multikristallinen Wafers aufzulösen, wurden zahlreiche Versuche durchgeführt, die Struktur mit 'Line-Scans' zu erfassen, statt mit zeitaufwendigen Flächen-'Scans'. Hierfür wurde bezüglich der Waferpräparation nicht das volle Programm abgearbeitet, wie in Abschnitt 5.1 beschrieben, sondern Waferbruchstücke verwendet, die lediglich noch einer RCA-Reinigung unterworfen werden mußten⁵.

Eine Meßserie, in der die Leckstromverteilung des Wafers deutlich und mit sehr guter Reproduzierbarkeit aufgelöst werden konnte, zeigt Abb. 5.10. Daß fast alle Versuche die Leckstromstruktur aufzulösen nicht erfolgreich waren, hat im wesentlichen zwei Ursachen:

- Wie bereits einleitend erwähnt, wurden all diese Messungen ohne 'Peak'-Detektor durchgeführt, der erst später gebaut und in den Meßaufbau integriert worden war und deshalb
- bestand bezüglich des Probenmaterials die eigentliche Schwierigkeit darin, Bruchstücke zu finden, die 'schlecht genug' waren, d. h. eine lateral hohe Dynamik im Leckstrom hatten, weil zu diesem Zeitpunkt die Meßapparatur noch

⁵ Die Bruchstücke stammten von Wafern, deren Säge-'Damage' bereits für frühere Messungen abgeätzt worden war. Da sie z. T. für lange Zeit in einer 'Abfallbox' gelagert wurden, war eine RCA-Reinigung unumgänglich.

nicht sensitiv genug war, um die Leckstromstruktur 'normaler' Wafer auszulösen.

Mit Sicherheit kann ausgeschlossen werden, daß bei allen Messungen, die zuvor durchgeführt wurden, der Elektrodenabstand größer war als der in der Bildunterschrift von Abb. 5.10 angegebene und deshalb die laterale Leckstromstruktur nicht aufgelöst werden konnte. Bei den zahlreichen Messungen an monokristallinem Material hatte sich die 200 μm Zylinder-Elektrode als sehr zuverlässig und langlebig erwiesen, so daß eventuelle 'zeitliche Änderungen' der Elektroden-Eigenschaften ebenfalls nicht in Betracht zu ziehen sind⁶.

Abb. 5.10 liefert eine Bestätigung der erwarteten Ortsabhängigkeit des Leckstromes, die schematisch in Abb. 4.5 gezeigt wurde. Die schwarzen Linien ordnen Leckstrom-Änderungen im 'Line-Scan' Korngrenzen auf dem Foto zu, wobei in der Messung lediglich die unterschiedlichen Leckströme der Körner, nicht aber die Korngrenzen selbst, sichtbar werden.

Nach den Ausführungen in Abschnitt 5.2 steht zu erwarten, daß in Umgebung einer Korngrenze der Leckstrom erhöht ist, da jede Korngrenze an der Kristalloberfläche einer Aneinanderreihung von Kristalldefekten entspricht. Daß die Korngrenzen, in Abb. 5.10 durch die gestrichelten Linien markiert, in der Messung nicht sichtbar sind, kann zwei Ursachen haben:

1. Wie in Abschnitt 5.2 diskutiert, besteht die Möglichkeit, daß die unmittelbaren Umgebungen der betrachteten Korngrenzen eine geringe Defektdichte aufweisen und deshalb der Leckstrom dort lokal, entweder gar nicht oder nur ganz geringfügig erhöht ist.
2. Das weiter oben bereits erwähnte Fehlen des 'Peak'-Detektors. Es ist nämlich durchaus möglich, daß die Leckströme an den betrachteten Stellen zwar erhöht sind, sich aber aufgrund des starken Meßrauschens nicht vom Leckstrom der angrenzenden Körner abheben.

Der Leckstrom-Scanner erlaubt, wie Abb. 5.10 beweist, die Leckstromstruktur multikristalliner Wafer mit 'Line-Scans' aufzulösen. Deshalb wurde versucht, Gleiches mit Flächen-'Scans' zu erreichen. Trotz zahlreicher Versuche war die Ausbeute an brauchbaren Messungen äußerst bescheiden. Ein Grund für die geringe Zahl an verwertbaren Messungen ist die bereits weiter oben erwähnte, im Vergleich zu monokristallinen Wafern, etwas schwierigere Handhabung der Wafer. Multikristalline Wafer sind ungefähr nur halb so dick wie monokristalline, von denen zudem noch etwas Material abgeätzt werden muß. Daher ist die Gefahr, daß die multikristallinen Proben beim Einbau in die Zelle brechen höher, als dies für Proben aus monokristallinem Material der Fall ist.

Ein weiteres Problem ist der z-Linearverstärker. Muß während der Messung mit dem z-Linearverstärker gefahren werden, weil sich der Piezo-Translator im Zustand

⁶ Die zeitliche Konstanz der Elektroden-Eigenschaften gilt nicht für die Platindraht-Elektrode. In der Regel verliert diese Elektrode mit der Zeit die Fähigkeit, Ströme aus einem kleinen Bereich der Oberfläche des Wafers zu ziehen, weil die Isolierung des Platindrahtes sehr empfindlich ist und leicht beschädigt werden kann.

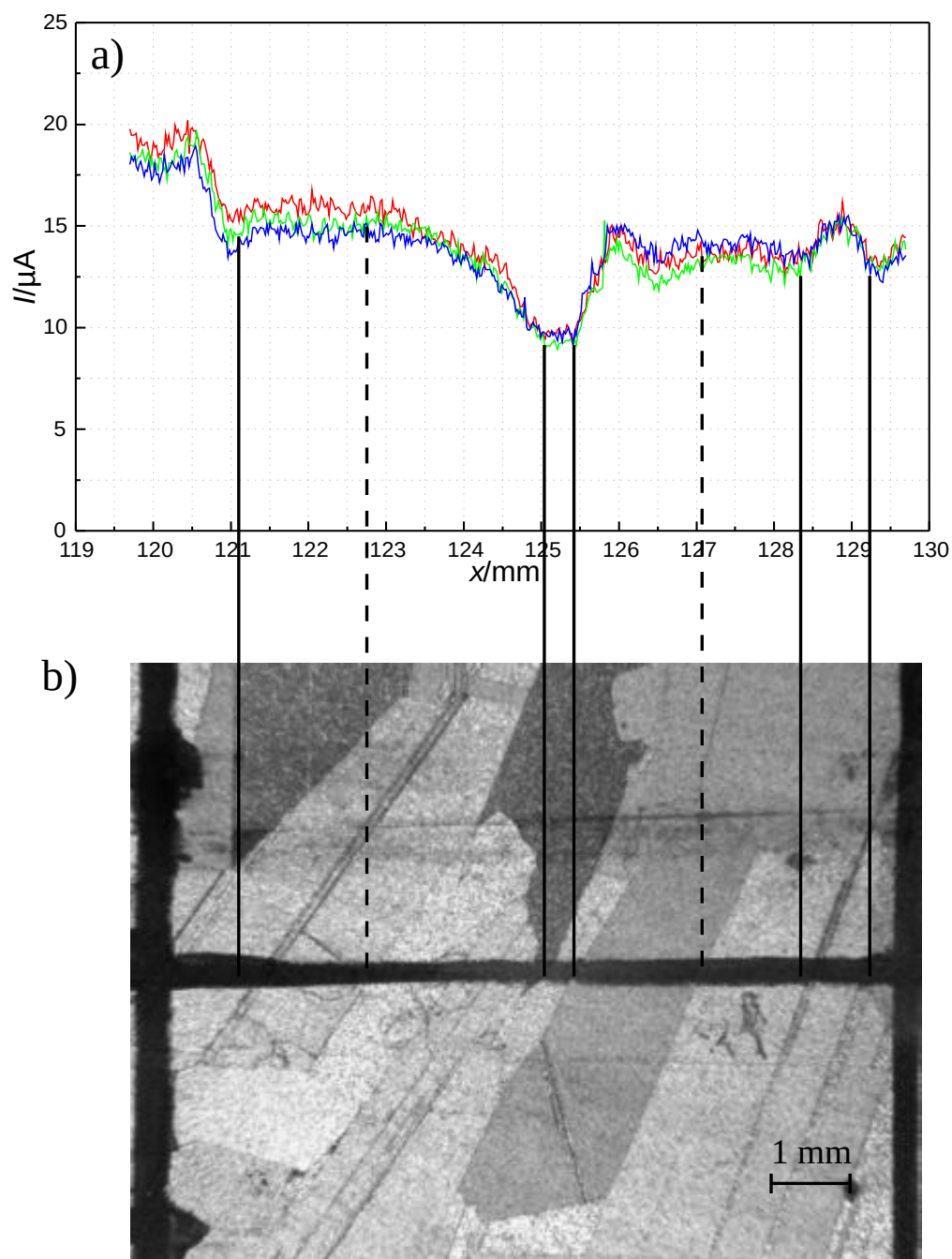


Abb. 5.10: a) 'Line-Scan' über einen multikristallinen Wafer mit der $200\ \mu\text{m}$ Zylinder-Elektrode und ohne(!) 'Peak'-Detektor. Der Elektrodenabstand beträgt $10\ \mu\text{m}$. b) zugehörige lichtmikroskopische Aufnahme des Wafers mit eingezeichnetem 'Scan'-Bereich.

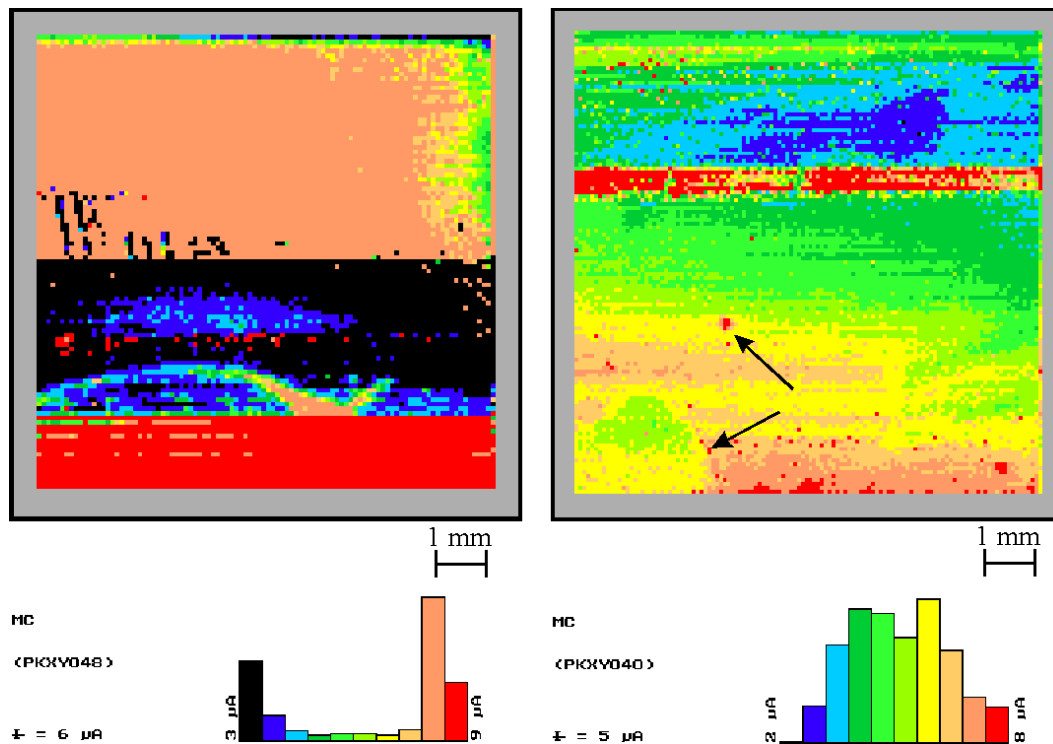


Abb. 5.11: Zwei Beispiele für fehlerhaftes Nachführen der Elektrode, bedingt durch leichtes Verkanten des Fahrschlittens beim Fahren des z-Linearverstellers. Messung mit der 200 μm Zylinder-Elektrode im Abstand von 10 μm und 'Peak'-Detektor.

minimaler oder maximaler Expansion befindet, kann das beträchtliche Änderungen des Elektrodenabstandes zur Folge haben. Der Grund hierfür ist nicht Spiel in der Mechanik des z-Linearverstellers, dessen Anfahrngenauigkeit 1 μm beträgt, sondern die relativ große Entfernung zwischen Meßkopf und Linearversteller: Der Meßkopf ist, wie in Abschnitt 3.3 näher beschrieben, an einer 30 cm langen Aluminiumschiene befestigt, siehe Abb. 3.17. Wird mit dem Linearversteller gefahren, so besteht die Gefahr, daß der Fahrschlitten beim Fahren geringfügig verkantet, weil durch einseitige Belastung Drehmomente auf ihn einwirken. Das kann Änderungen im Elektrodenabstand von bis zu 20 μm zur Folge haben⁷!

Diese Schwierigkeit zeigte sich nicht bei den Messungen an monokristallinem Material. Hier wurden maximal Flächen von 6 mm $\times$ 6 mm gemessen, so daß die Konstanz des Elektrodenabstandes mit dem Piezo-Translator allein sichergestellt werden konnte. Auf multikristallinem Material hingegen wurden zahlreiche Versuche unternommen, größere Flächen zu messen (max. 15 mm $\times$ 15 mm).

⁷ Die Problematik des Verkantens tritt klar in Erscheinung, wenn der Leckstrom-Scanner nach einer Messung für eine weitere mit unveränderten Parametern neu initialisiert wurde. Muß für die erste Messung die Elektrode um die Strecke Δz nach unten gefahren werden, um einen mechanischen Kontakt zwischen Elektrode und Wafer herbeizuführen, kann für die zweite Messung ein Weg von $\Delta z \pm 20 \mu\text{m}$ von Nöten sein!

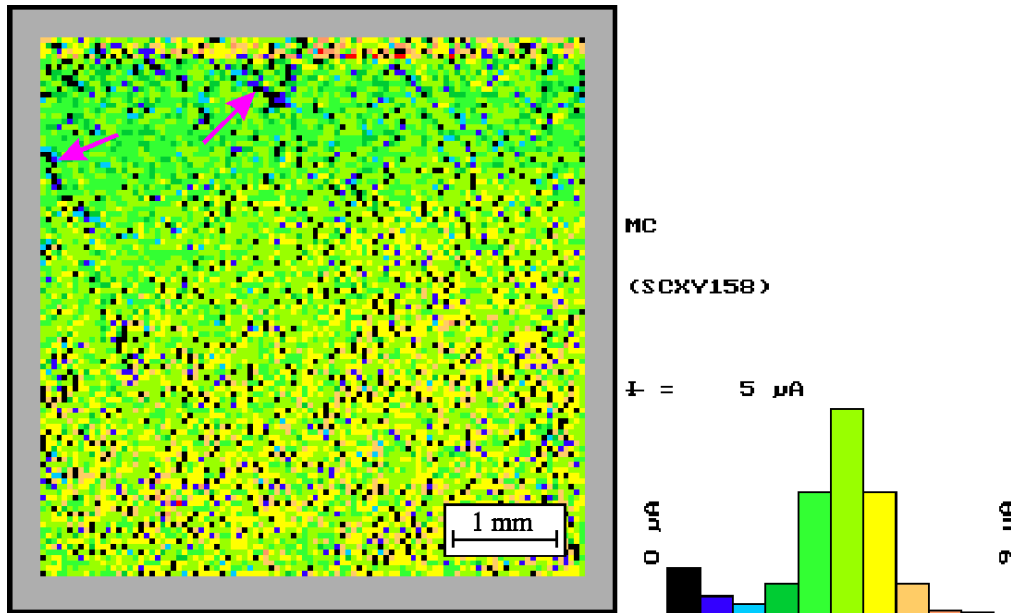


Abb. 5.12: Messung auf multikristallinem Silizium mit der $200 \mu m$ Zylinder-Elektrode im Abstand von $10 \mu m$ und ohne 'Peak'-Detektor. Die Kornstruktur des Wafers ist nur vereinzelt und schwach sichtbar, markiert durch die Pfeile.

Zwei Beispiele für Messungen mit zeitweiligen Kurzschlüssen zwischen Elektrode und Wafer als Folge des oben beschriebenen Problems mit der Mechanik des z-Linearverstellers zeigt Abb. 5.11. Dennoch gelang es in beiden mit 'Peak'-Detektor durchgeführten Messungen die Kornstruktur teilweise deutlich aufzulösen. Bezüglich der Messung im rechten Bild sind zwei Dinge beachtenswert:

1. Mit dieser Messung gelang es, einen 'Shunt' zu lokalisieren (oberer Pfeil), der den Wirkungsgrad einer Solarzelle erheblich verringern kann, vgl. Abb. 1.15 in Abschnitt 1.4. Die Fähigkeit, 'Shunts' in multikristallinem Material zu lokalisieren, stellt eine Besonderheit des Leckstrom-Scanners dar. So ist beispielsweise mit dem Elymaten das Lokalisieren von 'Shunts' nicht möglich!
2. An der Korngrenze des unten links sichtbaren Kornes ist der Leckstrom deutlich erhöht (unterer Pfeil)!

Hier zeigt sich die Bedeutung des 'Peak'-Detektors, die gemessene Leckstromstruktur des Wafers von Rauschen weitgehend frei zu halten, abermals. Beim Vergleich der Messung in Abb. 5.12, die vor Integration des 'Peak'-Detektors in den Meßaufbau durchgeführt worden war, mit der in Abb. 5.11, wird dies besonders deutlich. Das vergleichsweise starke Meßrauschen in Abb. 5.12 läßt die Kornstruktur fast vollständig darin verschwinden. Lediglich im oberen Teil der Messung ist sie, markiert durch die Pfeile, ansatzweise zu erkennen.

In Abb. 4.21 wurde der zerstörerische Einfluß, den Ethanol bei Stromfluß auf monokristallines Silizium hat, demonstriert.

Die Meßserie in Abb. 5.13 zeigt, daß auf multikristallinem Silizium von Messung zu Messung ebenfalls ein ausgeprägter Degradationsmechanismus ablaufen kann, falls

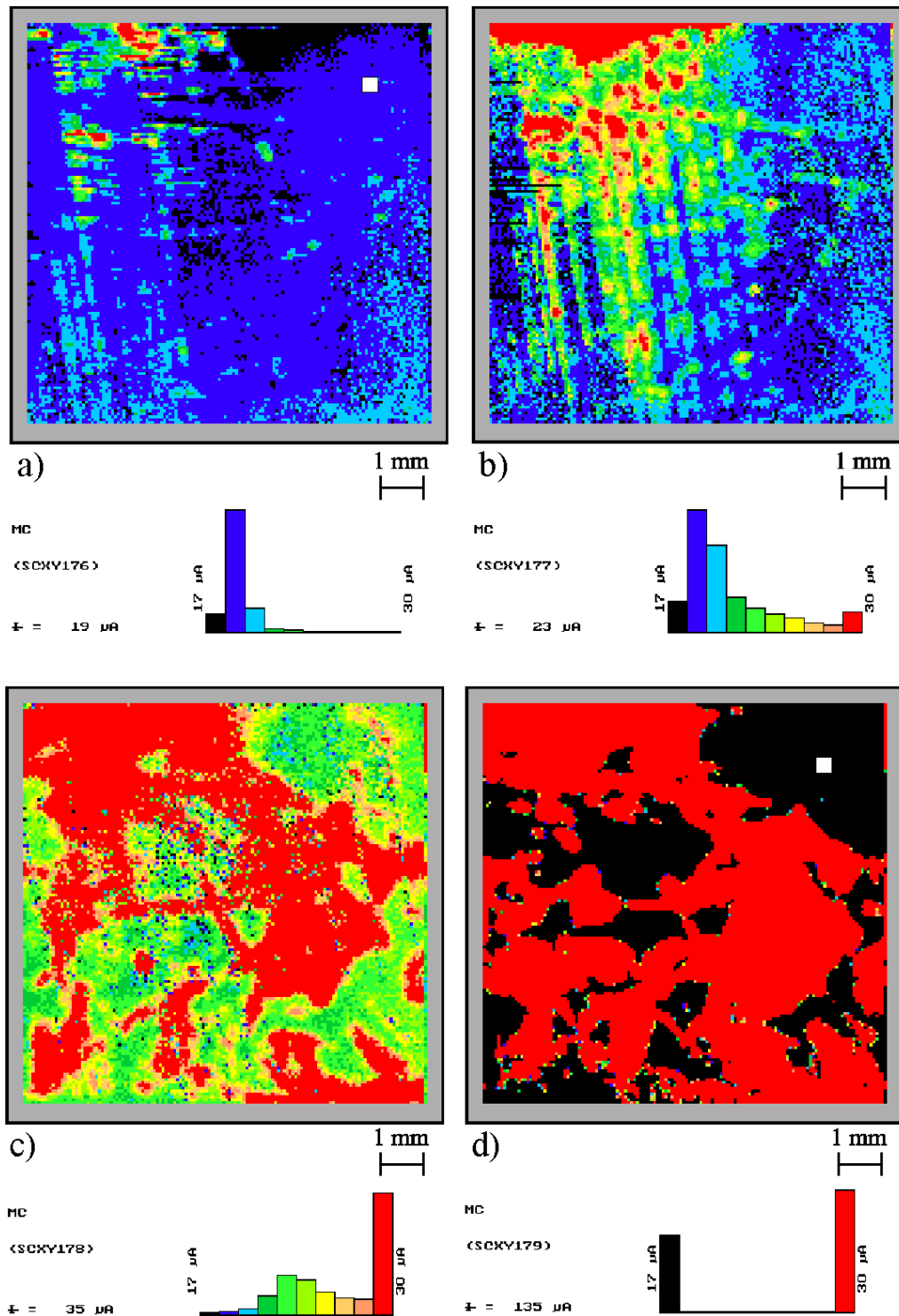


Abb. 5.13: Ein weiteres Beispiel für eine Meßserie, die den zerstörerischen Einfluß von Ethanol auf die Siliziumelektrode demonstriert. Messung mit der $200 \mu m$ Zylinder-Elektrode im Abstand von $10 \mu m$. a) 1. Messung, b) 2. Messung, c) 3. Messung und d) 4. Messung.

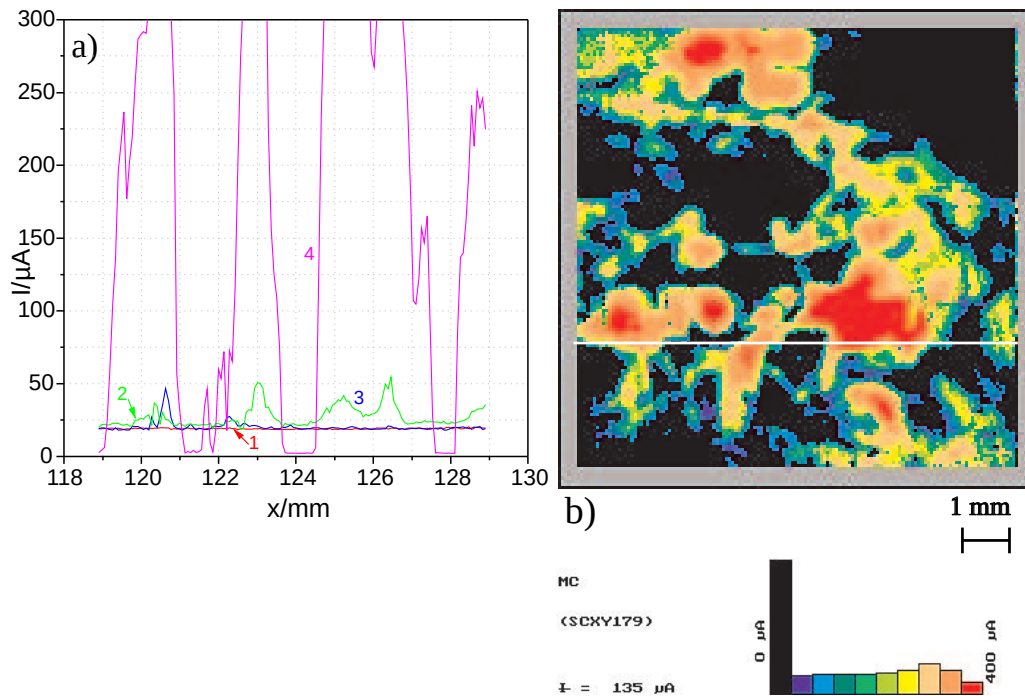


Abb. 5.14: a) 'Line-Scans' der Messungen aus Abb. 5.13 an der in b) markierten Stelle (4. Messung, umskaliert). Beachtenswert ist, daß bei der vierten Messung an vielen Stellen die Ströme kleiner sind als bei den anderen Messungen!

ein ethanolhaltiger Elektrolyt verwendet wird⁸:

Der Mechanismus läuft bevorzugt an Orten mit erhöhtem Strom ab, weil dort das an den Stromfluß gekoppelte Angebot an Ethanol-Radikalen ebenfalls erhöht ist. Aus diesem Grund tritt die Struktur in der zweiten Messung noch deutlicher hervor. In Bild c) ist die Struktur wieder fast vollständig verschwunden, da die Degradation zunehmend heftiger vonstatten geht und deshalb auch solche Bereiche geschädigt werden, deren Leckstrom anfänglich relativ klein war.

Bei der dritten Messung in Abb. 5.13 c) beginnen sich Leckstrom-'Cluster' zu bilden: Während der Strom aus dem einen 'Cluster' mit zunehmender Meßzeit immer weiter ansteigt, nimmt der Strom aus dem anderen deutlich ab! So beträgt der Strom, der in Bild a) aus der mit dem Quadrat markierten Stelle floß, 19.3 μA , während aus derselben Stelle in Bild d) nur noch 2.2 μA flossen, siehe auch Abb. 5.14. Ein Vergleich der Bilder c) und d) zeigt, daß die 'Cluster' mit reduziertem Leckstrom von Messung zu Messung immer kleiner werden und wahrscheinlich ganz verschwinden würden, wenn

⁸ Es sei nochmals darauf hingewiesen, daß im Verlauf der Arbeit nicht bewiesen wurde, daß Ethanol Ursache der Waferdegradationen ist. Die Behauptung, daß Ethanol die Ursache ist, rechtfertigt sich allein aus der Beobachtung, daß die Degradation des Wafers bei mehrfacher Messung *nie* beobachtet wurde, wenn der Elektrolyt ethanolfrei war (siehe Abb. 4.22). Zudem finden sich in der Literatur Hinweise, die Ethanol der zerstörerischen Wirkung an Halbleiterelektroden über den Mechanismus der Radikalbildung beschuldigen [MOR77]!

die Messung noch einige Male wiederholt worden wäre.

Sehr wahrscheinlich sind die erkennbaren Strukturen des Wafers gar nicht dessen Kornstruktur, sondern Restspuren des Säge-’Damages’, der nicht vollständig abgeätzt worden war⁹. Dafür sprechen zum einen die von oben nach unten verlaufenden Streifen, die in b) besonders deutlich sichtbar sind, und zum anderen die vergleichsweise hohen Ströme.

Wie bereits weiter oben erwähnt, wurden fast alle Messungen mit der 200 μm Zylinder-Elektrode durchgeführt, weil deren Herstellung vergleichsweise einfach und die Elektrode selbst relativ robust ist. Diese Eigenschaften besitzt die 50 μm Zylinder-Elektrode nicht. Da sie sehr leicht beschädigt werden kann, wurden mit ihr weit weniger Messungen durchgeführt als der mit 200 μm Zylinder-Elektrode. Die *einzigsten* Messungen, die mit der 50 μm Zylinder-Elektrode an multikristallinem Silizium gelangen, zeigen die Abb. 5.15 b) und c). Unter dem Lichtmikroskop wurde zuvor eine Stelle auf dem Wafer ausgewählt und fotografiert, die sich durch eine hohe Dichte an Korngrenzen auszeichnet und daher für eine Leckstrom-Messung als interessant befunden wurde, siehe Abb. 5.15 a).

Wie ein Vergleich mit dem Foto in Abb. 5.15 a) zeigt, konnte an vielen Stellen der Messungen in b) und c), die Kornstruktur deutlich aufgelöst werden. Die gemessene Struktur ist in der Tat die Leckstromverteilung des Wafers und nicht nur ein Abbild des Höhenprofils: Wie in Abschnitt 5.1.3 beschrieben ätzt die Politurbeize bevorzugt defektbehaftete Bereiche an, so daß das Ätzen an Korngrenzen mit etwas höherer Rate erfolgt, siehe Abb. 5.2. Wenn Bild 5.15 b) nur das Höhenprofil des Wafers widerspiegeln würde, dann müßten die Ströme aus den Korngrenzen kleiner sein als die aus den Körnern selbst, weil über einer Korngrenze der Elektrodenabstand etwas erhöht und dadurch der Spannungsabfall über den Elektrolytwiderstand größer ist (vgl. hierzu Abb. 4.12).

Des weiteren gilt zu beachten, daß der Strom über dem Kratzer in Abb. 4.12 deshalb stark vom Abstand abhängt, weil die dort fließenden Ströme sehr hoch sind und darum die über dem Elektrolytwiderstand abfallende Spannung erheblich ist. Die Ströme in Abb. 5.15 sind um zwei bis drei Größenordnungen niedriger, so daß der Spannungsabfall über dem Elektrolytwiderstand für diese Messung bei gleichem Elektrodenabstand sehr viel kleiner ist. Damit muß die Abstandsabhängigkeit des Stromes insgesamt weitaus geringer ausfallen, d. h. die gemessenen Ströme werden von Unebenheiten des Wafers kaum beeinflußt.

Zwei aneinandergrenzende Kristallite mit fast derselben Orientierung werden durch eine Kleinwinkel-Korngrenze voneinander getrennt, deren Defektdichte viel kleiner ist als die einer Großwinkel-Korngrenze (vgl. Abschnitt 5.2). Haben zwei Kristallite annähernd dieselbe Orientierung, werden sie auffallendes Licht nahezu in derselben Weise reflektieren, so daß die sie trennende Korngrenze bei Betrachtung des Wafers kaum oder gar nicht sichtbar ist.

Überträgt man diesen Sachverhalt auf die in Abb. 5.15 b) mit dem Pfeil markierte

⁹ Die Probe war ein Bruchstück aus der ’Abfallbox’, vgl. Fußnote auf Seite 120. Auf welche Art und Weise der Säge-’Damage’ dieses Bruchstückes beseitigt worden war, ist nicht bekannt.

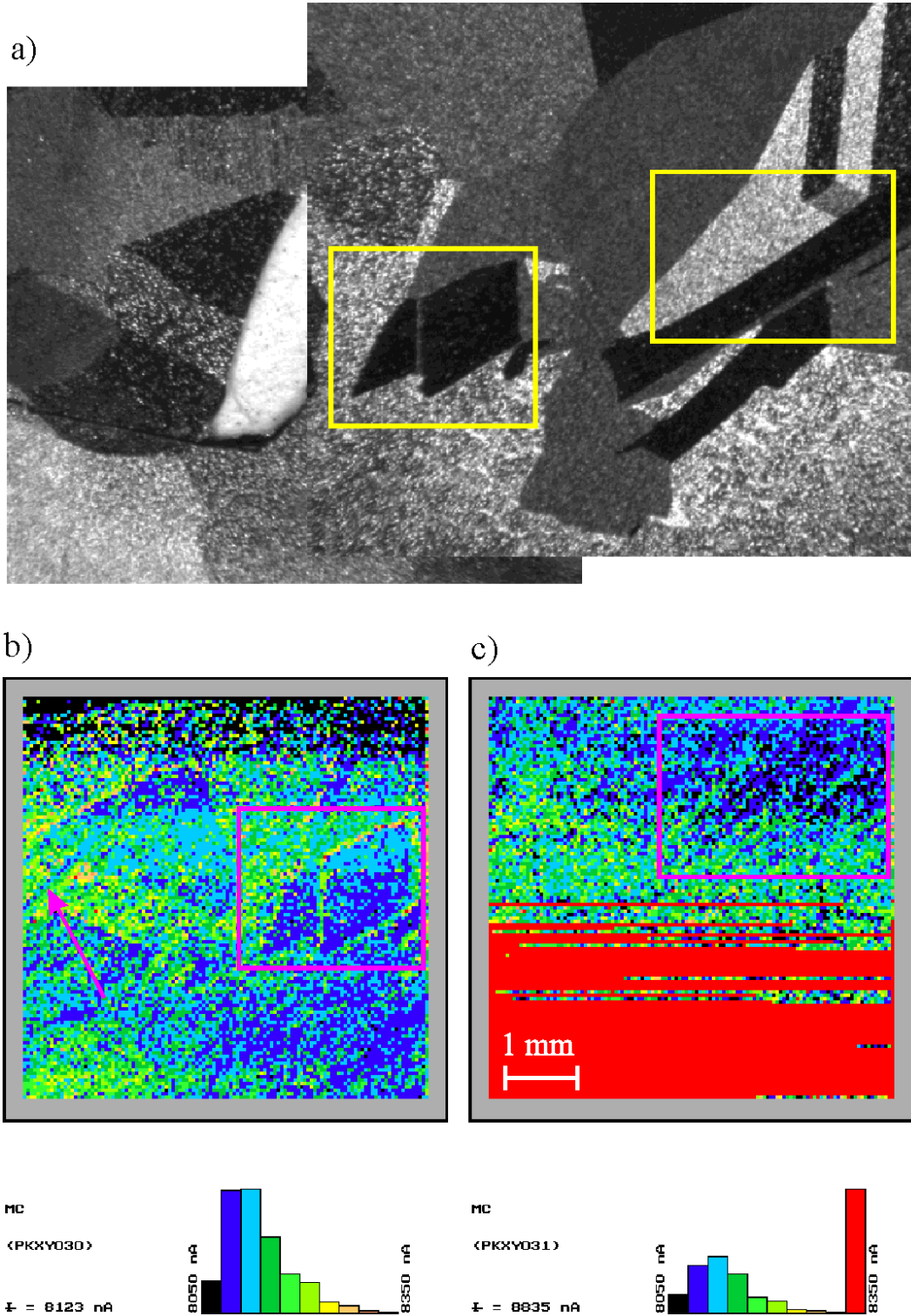


Abb. 5.15: Aufgelöste Kornstruktur eines multikristallinen Wafers, gemessen mit der $50 \mu\text{m}$ Zylinder-Elektrode im Abstand von $10 \mu\text{m}$ und 'Peak'-Detektor. a) Lichtmikroskopische Aufnahme des Wafers. Die 'Sample'-Zeiten betragen in b) 4 s und in c) 2 s. Der Pfeil weist auf eine Korngrenze, die auf dem Foto nicht sichtbar ist.

Korngrenze, die zwei 'schlechte' Körner voneinander trennt, so könnte es sich hierbei wahrscheinlich in der Tat um eine Kleinwinkel-Korngrenze handeln. Wie das Foto zeigt, reflektieren beide Kristallite das Licht in der gleichen Weise, so daß die beide Kristallite trennende Korngrenze auf dem Foto nicht sichtbar ist. Für eine Kleinwinkel-Korngrenze spricht weiter, daß der aus ihr fließende Leckstrom kleiner ist als der aus den anderen Korngrenzen. Da offenbar beide Kristallite fast dieselbe Orientierung haben, müßten auch, wie in Abschnitt 5.2 diskutiert wird, die Oberflächenzustandsdichten und damit die Leckströme annähernd dieselben sein. Abb. 5.15 b) zeigt, daß das tatsächlich der Fall ist.

Ebenso interessant wegen der hohen Korngrenzendichte ist der Bereich mit den zwei langen 'schwarzen' Korngrenzen in Abb. 5.15, so daß an dieser Stelle eine weitere Messung durchgeführt worden war. Um zu prüfen, inwieweit die 'Sample'-Zeit das Ergebnis beeinflußt, wurde sie für die zweite Messung halbiert. Das Ergebnis der Messung in Abb. 5.15 c) zeigt, daß die Kornstruktur fast vollständig im Rauschen verschwindet. Lediglich die Strukturen im vom gelben Rechteck umrandeten Bereich konnten Körnern auf dem Foto zugeordnet werden.

Leider ist die untere Hälfte der Messung unbrauchbar, weil dort die Elektrode die Waferoberfläche berührte. Da die Messung über Nacht durchgeführt wurde, wird vermutet, daß ein Absinken der Temperatur im Labor während der Nacht Ursache dafür war, daß sich der Meßaufbau thermisch *geringfügig* verzog, was letztlich das Aufsitzen der Elektrode auf dem Wafer zur Folge hatte. Da für eine konstante Temperatur in Umgebung des Leckstrom-Scanners nicht gesorgt wird, kann es bei allen zuvor durchgeführten Messungen ebenfalls zu thermisch bedingten Änderungen des Elektrodenabstandes gekommen sein. Allerdings blieben die Temperaturänderungen stets so klein, daß der Kurzschluß zwischen Elektrode und Wafer ausblieb.

Um zu überprüfen, ob die beiden unmittelbar aneinandergrenzenden Gebiete mit erhöhtem Leckstrom (siehe Pfeil) sich tatsächlich Strukturen auf dem Wafer zuordnen lassen, wurde die Probe nach der zweiten Messung für das Anfertigen eines zweiten Fotos ausgebaut, wobei sie allerdings zerbrach.

5.3.1 Vergleich mit Diffusionslängenverteilung

Für einen Vergleich zwischen der Leckstrom- mit der Diffusionslängenverteilung eines multikristallinen Wafers wurde mit dem Elymaten eine BPC-Messung durchgeführt. Für Elymat-Messungen wird ein roter Laser mit einer Wellenlänge von $\lambda = 670 \mu\text{m}$ verwendet, dessen Eindringtiefe sich mit Gleichung 2.3 zu $\alpha^{-1} = 4.0 \mu\text{m}$ berechnet.

Mit Hilfe einer Mikrometerschraube wird die Waferdicke d an den vier Ecken bestimmt. Da im allgemeinen das Abätzen des Säge-'Damages' zur Folge hat, daß die Dicke über dem Wafer leicht variiert, wird im Programm für die Auswertung von Elymat-Messungen mit der jeweils aktuellen Waferdicke gerechnet, die sich aus den Dicken an den vier Eckpunkten und der Länge des Wafers bestimmt. Unter der Annahme, daß $S_f = 0$ ist, wird mit Gleichung 2.23 die Diffusionslänge L berechnet, wobei

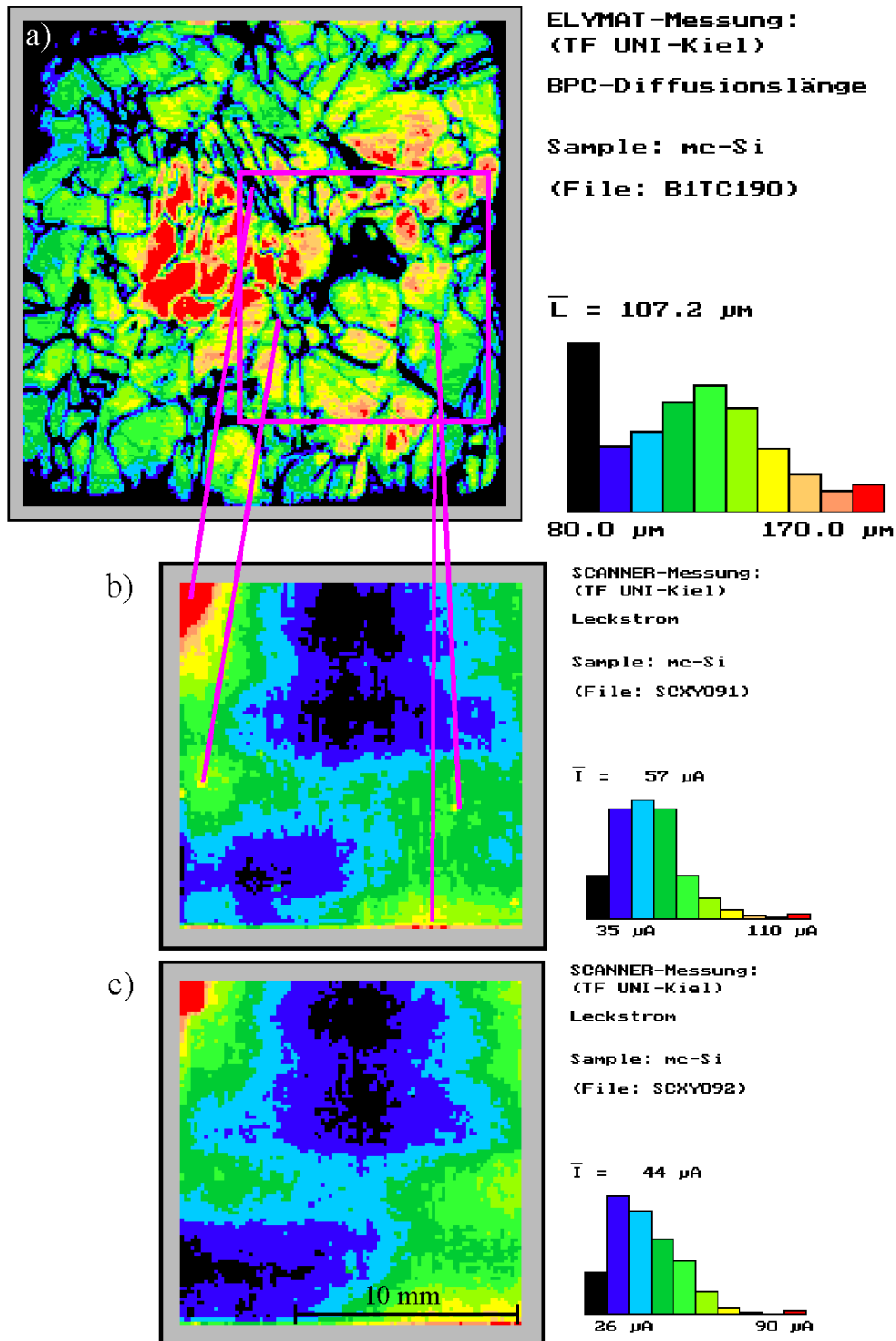


Abb. 5.16: a) Mit dem Elymaten gemessene Diffusionslängenverteilung. b) Zugehörige Leckstromverteilung des in a) markierten Bereiches, gemessen mit der 200 μm Zylinder-Elektrode im Abstand von 10 μm . c) Wiederholte Messung zur Überprüfung der Reproduzierbarkeit.

sich der Wert für J_{FPC} mit Hilfe einer FPC-Messung bestimmen läßt ¹⁰.

Das Ergebnis der Elymat-Messung zeigt Abb. 5.16 a). Interessant für eine Leckstrom-Messung sind solche Bereiche, die sich durch eine große Dynamik in der Diffusionslänge auszeichnen. Ein Bereich für den die Bedingung erfüllt wird, ist in Abb. 5.16 a) durch das Quadrat markiert. Innerhalb dieses Bereiches wurden nacheinander zwei Leckstrom-Messungen durchgeführt.

Ein Vergleich der Diffusionslängen- mit der Leckstrom-Messung zeigt:

1. Die Strukturen der Diffusionslängen- und der Leckstrom-Messung sind vollkommen unterschiedlich.
2. Es gibt mehrere Bereiche an denen der Leckstrom groß und die Diffusionslänge klein ist.

Eine mögliche Erklärung dafür, daß an einem Ort mit erhöhtem Leckstrom auch stets die Diffusionslänge reduziert ist, könnte in den elektrischen Eigenschaften der Korngrenzen zu finden sein. Daß die durch Lichteinstrahlung generierten Minoritäten (Überschußladungsträger) an Korngrenzen rekombinieren führt dazu, daß in ihrer näheren Umgebung die Diffusionslänge verringert ist, siehe Abb. 5.16 a). Die Rekombination von Minoritäten an Korngrenzen sagt aber noch nichts über die Höhe des dort fließenden Leckstromes aus. Wie in Abschnitt 5.2 näher erläutert, wird der Leckstrom davon abhängig sein, ob die Versetzungen der Korngrenze elektrisch geladen sind oder ob sie einen ohmschen Kanal durch die Waferoberfläche darstellen, durch den Minoritätsladungsträger abfließen können, siehe Abb. 5.6. Trifft einer der Mechanismen zu, so würde an den betreffenden Korngrenzen der Leckstrom erhöht sein. Abb. 5.16 zeigt, daß dies an vier Orten der Fall sein könnte.

Die Messung in Abb. 5.15 zeigt, daß der Leckstrom-Scanner erlaubt, die Kornstruktur eines multikristallinen Wafers aufzulösen. Da aber das Leckstrombild in Abb. 5.16 eine relativ hohe Dynamik in den Strömen aufweist, wird die Kornstruktur des Wafers sozusagen 'überstrahlt'. Selbst wenn die Messung keine so hohe Dynamik in den Strömen zeigen würde, hätte die Kornstruktur sehr wahrscheinlich dennoch nicht aufgelöst werden können, weil die Messung zu einem Zeitpunkt durchgeführt worden war, an dem noch ohne 'Peak'-Detektor gemessen wurde.

¹⁰ Strenggenommen gilt $S_f = 0$ nie! Für sehr gut passivierte und defektfreie Oberflächen können Oberflächenrekombinationsgeschwindigkeiten von unter $0.25 \frac{\text{cm}}{\text{s}}$ erreicht werden [YAB86], wobei für H-terminierte Oberflächen Werte im Bereich $100 - 200 \frac{\text{cm}}{\text{s}}$ realistischer sind. Ist $\frac{L}{d}$ und damit die Diffusionslänge L nicht zu groß, hat bei einer BPC-Messung die vordere Oberflächenrekombinationsgeschwindigkeit S_f nur einen geringen Einfluß auf die Diffusionslänge L [LIP97].

5.3.2 Fazit

Abb. 5.15 zeigt, daß sich die Ströme aus den Korngrenzen nur geringfügig von denen aus den Körnern unterscheiden. Daher ist ein Auflösen der Kornstruktur nur dann möglich, wenn

1. der Leckstrom der einzelnen Körner nicht zu unterschiedlich ist,
2. der Säge-’Damage’ komplett beseitigt wurde und
3. das Meßrauschen hinreichend klein ist.

Die Erfüllung der dritten Bedingung erfordert entweder den Einsatz des ’Peak’-Detektors, dessen ’Sample’-Zeit, wie Abb. 5.15 nahelegt, mindestens bei 2 s liegen sollte oder das Verwenden eines nicht-wässrigen Elektrolyten, siehe Abschnitt 4.4.3.

Weil multikristalline Wafer etwa nur halb so dick wie monokristalline Wafer sind ist der Einbau der Probe in die Zelle nicht ganz unproblematisch. Um die Zelle zu dichten, müssen relativ hohe Anzugsmomente für die Zellenschrauben verwendet werden. Es kann durchaus vorkommen, daß die Probe dabei zerbricht.

Da das Verwenden langer ’Sample’-Zeiten von mindestens 2 s offenbar unumgänglich ist (siehe Abb. 5.15), scheint es fraglich, ob der Leckstrom-Scanner als ’Monitoring-Tool’, für große Messungen, d. h. Messungen mit beispielsweise 500×500 Meßpunkten, überhaupt eingesetzt werden sollte. Bei Meßzeiten von knapp $500^2 \cdot 2 \text{ s} = 140 \text{ h}$ würde der Leckstrom-Scanner als ’Monitoring-Tool’ wohl seinen Zweck verfehlen. Der Elymat beispielsweise, der vorwiegend als ’Monitoring-Tool’ in der Kontaminationskontrolle Anwendung findet, arbeitet um Größenordnungen schneller!

Zusammenfassung und Ausblick

Zielsetzung dieser Arbeit war der Aufbau eines Meßplatzes, mit dem der Leckstrom über die Silizium-Elektrolyt-Grenzfläche orts aufgelöst gemessen werden kann (bezeichnet als *Leckstrom-Scanner*).

Trotz der simplen Funktionsweise, die einem 'Scanning Electrochemical Microscope' (SECM) stark ähnelt, galt es, zum Teil schwerwiegende Probleme zu lösen, um zur Information des lokal fließenden Leckstromes zu gelangen: Der Erhalt von Ortsauflösung erforderte die Konstruktion von Elektroden, deren Meßfläche auf ein kleines Gebiet begrenzt ist. Die Platindraht-Elektrode, die lediglich aus einem isolierten 200 μm dicken Platindraht besteht, ist bezüglich ihres Aufbaus die einfachste Elektrode, mit der Ortsauflösung erreicht worden ist. Da ihr Auflösungsvermögen extrem stark vom Elektrodenabstand abhängt, sind für Messungen sehr kleine Abstände von Nöten. Allerdings führen geringste Schäden in der Isolierung (beispielsweise feine Risse) zum vollständigen Verlust der Ortsauflösung! Zudem hängt die Eigenschaft der Platindraht-Elektrode, Variationen im Leckstrom lateral aufzulösen, sehr empfindlich von der Beschaffenheit der Elektrodenfläche ab, die nur schlecht definiert werden kann.

Die im Zusammenhang mit der Platindraht-Elektrode genannten Probleme motivierten die Entwicklung von insgesamt drei *abgeschirmten* Elektroden: Die 200 μm Zylinder-, die 50 μm Zylinder- und die Scheiben-Elektrode. Allen drei Elektroden ist gemeinsam, daß die innere Elektrode von einer zweiten Elektrode in Form eines Zylinders umgeben wird. Während einer Messung befinden sich beide Teilelektroden auf demselben Potential, wobei aber nur der durch die innere Elektrode fließende Strom gemessen wird! Da die äußere Elektrode bewirkt, daß das elektrische Feld der inneren Elektrode auf die Probe fokussiert wird, konnte die Abstandsabhängigkeit des Auflösungsvermögens erheblich reduziert werden! Solange beide Teilelektroden keinen direkten elektrischen Kontakt zueinander haben, beeinflußt eine schadhafte Isolierung der inneren Elektrode das Auflösungsvermögen in keiner Weise, so daß auch das Problem der Isolierung gelöst werden konnte.

Anfängliche Probleme mit der Isolierung beider Teilelektroden, die sich mit der Zeit mehr und mehr ablöste und schließlich in einem Kurzschluß mündete, konnten durch Einsatz von Polyimid auf sehr erfolgreiche Weise gelöst werden. Mit Polyimid als Isolationsmaterial wurde eine 200 μm Zylinder-Elektrode gebaut, die sich als sehr robust und langlebig erwies und mit der ein Großteil der vorgestellten Messungen durchgeführt wurde.

Zwecks Verbesserung der Ortsauflösung wurde die 50 μm Zylinder-Elektrode kon-

struiert. Aufgrund der Tatsache, daß Platin ein sehr weiches Metall ist, gestaltete sich das Nachbearbeiten der Elektrodenfläche als besonders schwierig. Zudem konnten geringste Deformierungen der Elektroden spitze leicht zu einem Kurzschluß zwischen beiden Teilelektroden führen, der nur schwer wieder beseitigt werden kann.

Die Scheiben-Elektrode ist vom Aufbau her mit der 200 μm Zylinder-Elektrode identisch, wobei die Elektroden spitze um eine Platinscheibe ergänzt ist. Sie sollte eine Homogenisierung des elektrischen Feldes zwischen Elektrode und Wafer bewirken und damit das Auflösungsvermögen über einen weiten Bereich unabhängig vom Elektrodenabstand machen. Allerdings zeigte sich, daß die Scheiben-Elektrode lediglich eine etwas schwächere Abstandsabhängigkeit besitzt als die 200 μm Zylinder-Elektrode.

Ein schwerwiegendes Problem für die Qualität der Messungen war die an den Stromfluß gekoppelte und damit unvermeidbare Entwicklung von H_2 - und O_2 -Bläschen zwischen Elektrode und Wafer, die zur Folge hatten, daß die Ströme stark verrauscht waren. Infolgedessen war es in der Regel unmöglich, die Leckstromstruktur multikristalliner Wafer aufzulösen, die komplett im Rauschen verschwand. Durch Einsatz eines 'Peak'-Detektors, konnte der Problematik der Gasbläschenentwicklung sehr erfolgreich begegnet werden. Um die Messung weitgehend frei von Rauschen zu halten, waren allerdings relativ lange 'Sample'-Zeiten um 2 s notwendig. Auf diese Weise gelang es erstmalig, die Kornstruktur eines multikristallinen Wafers in Form eines Leckstrombildes zu messen!

Eine weitere Möglichkeit, die Problematik der Gasbläschenentwicklung zu eliminieren, besteht in der Verwendung eines nicht-wässrigen Elektrolyten, dem ein Leitsalz und ein Redoxpaar zugesetzt werden. Erste Versuche waren sehr erfolgreich, allerdings erfordert der Umgang mit diesem Elektrolyten eine absolut trockene und sauerstofffreie Umgebung ('Glove-Box' mit Schutzgas-Atmosphäre)! Der erhöhte apparative Aufwand würde sich dadurch rechtfertigen, daß dann auf den 'Peak'-Detektor verzichtet werden kann, was die Meßgeschwindigkeit rapide erhöhen würde. Dieser Aspekt gewinnt an Bedeutung, wenn beispielsweise das komplette Leckstrombild eines 10 cm $\times$ 10 cm Wafers erstellt werden soll. Um den Leckstrom an jedem Ort zu erfassen, wäre unter Verwendung der 200 μm Zylinder-Elektrode eine Messung mit 500 $\times$ 500 Punkten von Nöten. Mit einem wässrigen Elektrolyten und einer 'Sample'-Zeit von 2 s ergäbe das eine Meßzeit von ca. 140 h! Diese lange Meßzeit würde den Einsatz des Leckstrom-Scanners als 'Monitoring-Tool' verbieten. Im Vergleich hierzu arbeitet der Elymat, der als 'Monitoring-Tool' während der Prozessierung von Siliziumwafern Anwendung findet, um Größenordnungen schneller!

Im Zusammenhang mit großen Messungen gilt dem Meßprogramm Beachtung: Dieses ist lediglich ein einfaches DOS-Programm, daß seine Existenz allein dadurch rechtfertigt, während der Entwicklungsphase des Leckstrom-Scanners als Testprogramm gedient zu haben. Deshalb ist die Entwicklung eines beispielsweise unter Windows laufenden Programms angebracht, um den Maßstäben der heutigen Zeit gerecht zu werden. Zudem lassen sich die oben erwähnten Messungen mit 500 $\times$ 500 Meßpunkten unter Windows weitaus bequemer handhaben als unter DOS.

Der Einbau multikristalliner Wafer in die Meßzelle war nicht ganz unproblematisch.

Erfolgte das Festziehen der Zellschrauben zu fest oder ungleichmäßig, zerbrach der Wafer; waren sie zu locker, war die Zelle undicht. Die Beseitigung dieses Problems erfordert eine komplette Neukonstruktion. So sollte der Zellenaufsatz nicht frei beweglich sein, sondern über Führungsschienen bewegt werden, so daß beim Schließen der Zelle der Aufsatz exakt planparallel zum Wafer geführt wird. Des weiteren sollten Zellenaufsatz sowie die Waferauflage eine starre Einheit bilden, damit im geschlossenen Zustand mechanische Spannungen im Wafer möglichst klein gehalten werden.

Das Problem mit der Mechanik des z-Linearverstellers (leichtes Verkanten des Schlittens beim Fahren) verbietet zur Zeit Messungen mit einer 'Scan'-Fläche von mehr als $6\text{ mm} \times 6\text{ mm}$, da nur bis zu dieser Größe die Konstanz des Elektrodenabstandes mit dem Piezo-Translator allein sichergestellt werden kann. Würde die neue Zelle die Möglichkeit bieten, sich innerhalb eines kleinen Winkelbereiches frei orientieren zu lassen, könnten die Waferauflage sowie die xy-Fahrebene der Linearversteller aufeinander einjustiert werden. Dadurch sollte der derzeitig verwendete Piezo-Translator mit einem Gesamtfahrweg von $90\text{ }\mu\text{m}$ auch für Messungen mit einer 'Scan'-Fläche von bis $10\text{ cm} \times 10\text{ cm}$ für die Elektrodennachführung ausreichend sein.

Beim derzeitigen Aufbau des Leckstrom-Scanners läuft der Experimentator Gefahr, die Elektrode zu beschädigen oder gar zu zerstören, wenn er im Verlauf der Elektrodenjustierung diese zu weit nach unten fährt. Deshalb sollte der Aufbau durch eine Vorrichtung ergänzt werden, die dieses Zuweitfahren zuverlässig unterbindet.

Anhang A

Meßprogramm

Die Durchführung einer Messung umfaßt im wesentlichen die folgenden Schritte:

- das Setzen der Meßparameter
- die Initialisierung der Meßgeräte
- die Positionierung der Elektrode
- die eigentliche Messung

Die Arbeitsweise des Meßprogramms soll anhand des in Abb. [A.1](#) gezeigten Flußdiagramms erläutert werden.

Unmittelbar nach dem Start des Programms werden die aktuellen Meßparameter aus einer Datei gelesen und am Bildschirm angezeigt. Nach einer eventuellen Änderung der Meßparameter, wie beispielsweise Elektrodenabstand, Elektrodenspannung, Zahl der Meßpunkte pro 'Line-Scan' usw., werden die Parameter in der Datei gespeichert. Wird die Startposition für die z-Richtung geändert, muß unbedingt darauf geachtet werden, daß hierfür nicht versehentlich ein zu großer Wert gewählt wird, weil sonst beim Anfahren der Startposition die Elektrode in die Probe gefahren und dabei beschädigt oder gar zerstört werden würde! Ist die daraufhin folgende Initialisierung der Meßgeräte erfolgreich, wird die Elektrode zur Startposition gefahren, ansonsten wird das Programm mit einer Fehlermeldung beendet.

Anschließend bietet das Programm die Möglichkeit, die Position der Elektrode zu manipulieren, beispielsweise um sie für eine Messung an einem 'N' genau zu justieren. Auch hier gilt zu beachten, daß beim Herunterfahren der Elektrode nicht versehentlich eine zu große Schrittweite gewählt wird, wenn sich die Elektrode bereits dicht über der Probe befindet! Wird die Justierung der Elektrode beendet, sollte die z-Position derart sein, daß bei der im nächsten Schritt erfolgenden Bestimmung der Ebenenparameter nur kurze Wege in z-Richtung gefahren werden müssen. Kurze Wege sind deshalb wichtig, weil dann das Fahren mit dem Piezo-Translator allein ausreichend ist, um den Kurzschluß zwischen Elektrode und Probe herbeizuführen. Muß zusätzlich mit dem z-Linearversteller gefahren werden, wird die Bestimmung der Ebenenparameter und damit die spätere Elektrodennachführung ungenau. Grund hierfür ist, daß der

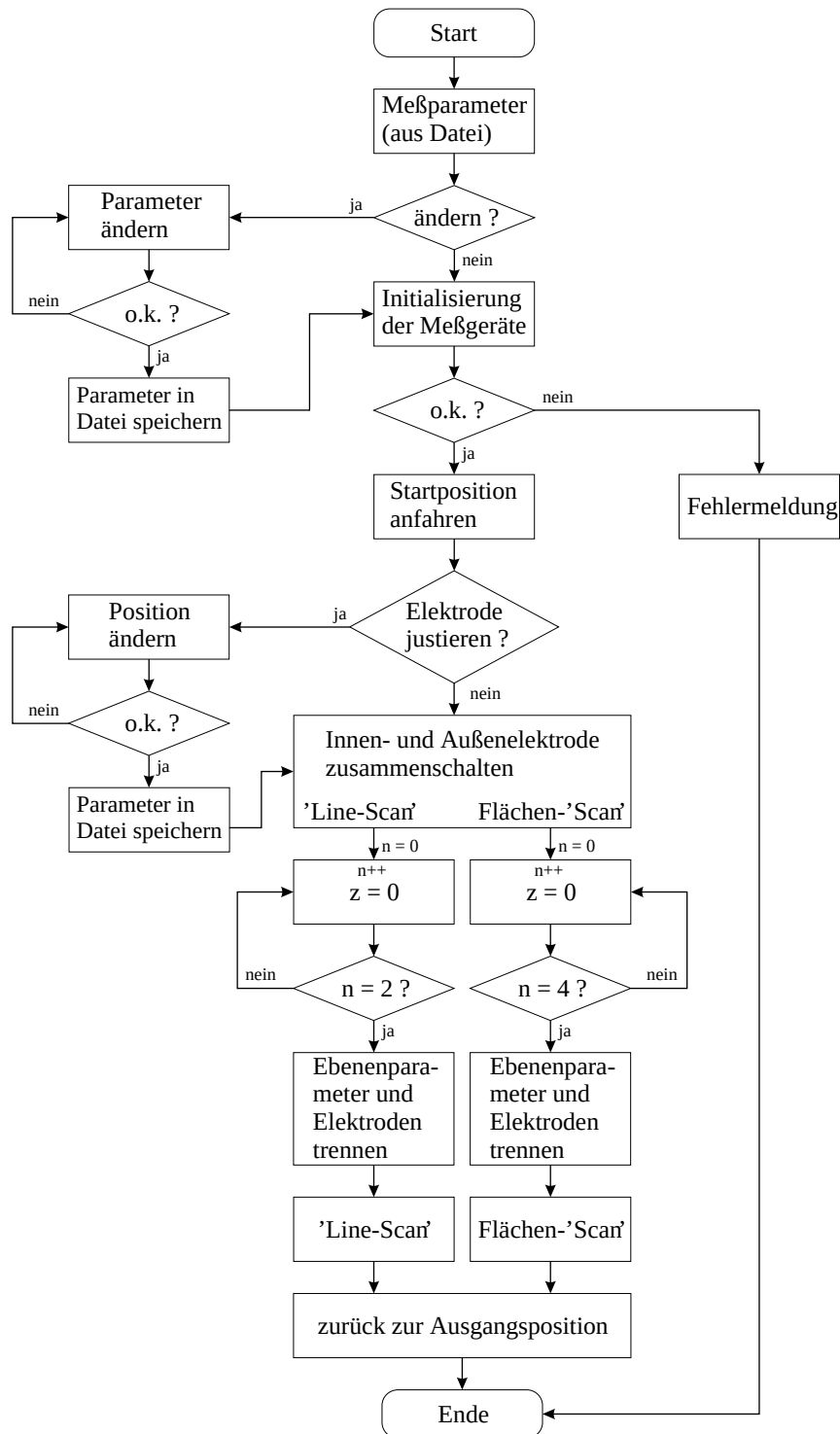


Abb. A.1: Flußdiagramm des Meßprogramms.

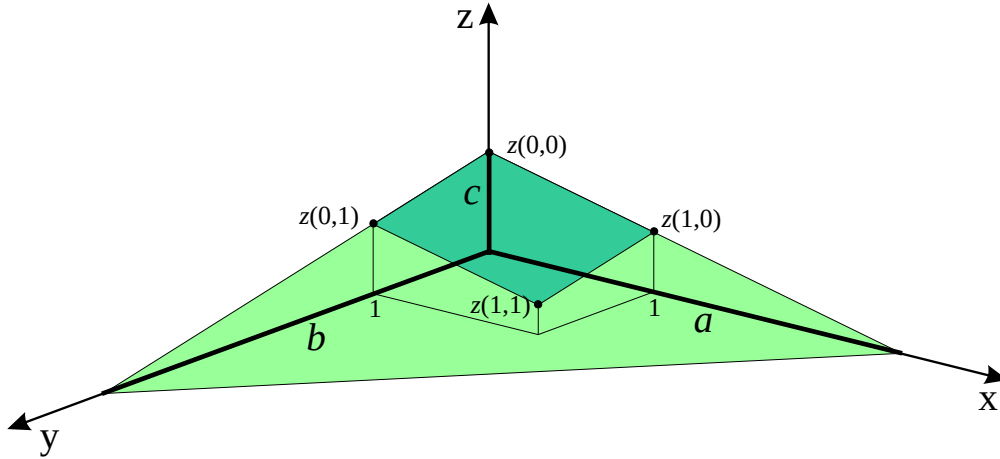


Abb. A.2: Zur Bestimmung der Ebenenparameter für die Elektroden-Nachführung. Erläuterungen siehe Text.

Schlitten des z-Linearverstellers beim Fahren leicht verkantet, was Ungenauigkeiten im Elektrodenabstand von bis zu $20\ \mu\text{m}$ zur Folge haben kann¹!

Je nachdem, ob ein 'Line'- oder ein Flächen-'Scan' durchgeführt werden soll, werden als nächstes entweder die Höhen an Start- und Endpunkt ('Line-Scan') oder die Höhen der vier Eckpunkte (Flächen-'Scan') bestimmt. Soll die Messung mit einer der Zylinder- oder der Scheiben-Elektrode erfolgen, müssen die Innen- und die Außen-elektrode hierfür zusammengeschaltet werden. Das Zusammenschalten erfolgt durch ein Relais, das über Bit 1 des Daten-Registers der parallelen Schnittstelle angesteuert wird, siehe Abb. 3.17. Nachdem die Elektrodenspannung auf 2 V gesetzt wurde, wird die Elektrode in $1\ \mu\text{m}$ -Schritten nach unten gefahren. Kommt es zum Kurzschluß zwischen Elektrode und Probe, wird dies als Stromsprung registriert (siehe Abb. 3.20) und die Abwärtsbewegung gestoppt. Die aktuelle z-Position wird in der Variablen z_ν gespeichert. Bevor die nächste Eckposition angefahren wird, wird die Elektrode in z-Richtung wieder zur Startposition zurückgefahren.

Im Falle eines Flächen-'Scans' liegen die vier Eckpunkte mit $z = 0$ im allgemeinen nicht auf einer Ebene. Deshalb wird an die vier Eckpunkte eine Ebene gefittet. Die Achsenabschnittsform der Ebene lautet

$$\frac{x}{a} + \frac{y}{b} + \frac{z}{c} = 1, \quad (\text{A.1})$$

wobei a , b und c die Strecken sind, die von der Ebene auf den Koordinatenachsen abgeschnitten werden, siehe Abb. A.2. Aus den vier zuvor bestimmten z-Werten $z_\nu = z(x_\nu, y_\nu)$ gilt es, die Ebenenparameter a , b und c zu bestimmen, wobei die Punkte (x_ν, y_ν) auf den Eckpunkten eines Quadrates mit der Kantenlänge 1 liegen, siehe Abb. A.2. Folgende Funktion gilt es unter Variation der Parameter a , b und c zu mi-

¹ Siehe Fußnote auf Seite 123.

nimieren:

$$F(a, b, c) = \sum_{\nu=1}^4 \left(z_{\nu} - c \left[1 - \frac{x_{\nu}}{a} - \frac{y_{\nu}}{b} \right] \right)^2 = \min! \quad (\text{A.2})$$

Hierfür wird ein Standardverfahren herangezogen, das das Minimieren in mehreren Dimensionen erlaubt [PRE89].

Mit den so bestimmten Ebenenparametern kann nun an jedem Ort der Meßfläche die Abweichung Δz vom Mittelniveau z_M , das sich aus dem arithmetischen Mittel der vier Eckpunkte z_{ν} ergibt, berechnet und die Elektrode entsprechend nachgeführt werden. Da die Ebenenparameter für ein Quadrat mit der Kantenlänge 1 bestimmt wurden, müssen für die Berechnung von Δz die x- und y-Koordinaten der aktuellen Elektrodenposition normiert werden. Mit der 'Scan'-Länge L und den Startpositionen der x- und y-Richtung x_0 bzw. y_0 ergibt sich für die normierten Koordinaten:

$$\xi_x = \frac{x - x_0}{L} \quad \text{und} \quad \xi_y = \frac{y - y_0}{L}. \quad (\text{A.3})$$

Damit berechnet sich Δz wie folgt:

$$\Delta z = c \left[1 - \frac{\xi_x}{a} - \frac{\xi_y}{b} \right] - z_M. \quad (\text{A.4})$$

Abschließend wird der z-Manipulator derart konfiguriert, daß der Piezo-Translator in Mittelposition ist, wenn sich die Elektrode an einem Ort befindet, für den $\Delta z = 0$ gilt, vgl. Abb. 3.19.

Mit Bestimmung der Ebenenparameter ist der Vorgang der Initialisierung abgeschlossen. Nachdem Innen- und Außenelektrode wieder elektrisch voneinander getrennt wurden, wird die eigentliche Messung gestartet. Nach der Messung wird die Elektrode zurück zur Ausgangsposition gefahren und das Programm beendet.

Anhang B

'Peak'-Detektor

Bei Leckstrom-Messungen mit wässrigen Elektrolyten kann sich die stromflußbedingte Entwicklung von Wasserstoff- und Sauerstoffbläschen sehr nachteilig auswirken, was in stark verrauschten Messungen zum Ausdruck kommt. Aus diesem Grund wird zwischen dem Potentiostaten und dem Digital-Multimeter (Keithley 2000) ein 'Peak'-Detektor zwischengeschaltet, der während eines vom Experimentator frei wählbaren Zeitintervalls Δt ('Sample'-Zeit) nach dem Maximalwert des Stromes 'sucht'¹, siehe Abb. 3.17. Während der 'Peak'-Detektor bei Messungen an Kratzern aufgrund der großen Signal-zu-Untergrund-Verhältnisse eine untergeordnete Rolle spielt, ist er für Messungen an multikristallinem Material zwingend erforderlich!

Realisiert wurde ein digital arbeitender 'Peak'-Detektor zum einen, weil dieser unendlich lange Haltezeiten² erlaubt und zum anderen, weil dadurch das Digital-Multimeter entbehrlich wird. Das Digital-Multimeter wird deshalb entbehrlich, weil im 'Peak'-Detektor der Meßwert bereits in digitaler Form mit 12 Bit Genauigkeit vorliegt und zwecks Meßwerterfassung beispielsweise über die parallele Schnittstelle des Computers eingelesen werden kann.

Während in Abschnitt 3.2.4 lediglich die Funktionsweise des 'Peak'-Detektors anhand des Blockschaltbildes beschrieben wird, soll im folgenden die schaltungstechnische Realisierung angegeben und beschrieben werden.

Der 'Peak'-Detektor ist aufgebaut aus drei Modulen:

- der Wandlerplatine mit dem 12 Bit Speicher-Register
- der Platine mit den Filtern
- der Steuerplatine zum Betrieb an der parallelen Schnittstelle.

¹ Der Potentiostat konvertiert den Strom in eine Spannung, deren Wert gemessen wird. Daher läuft die Bestimmung des maximalen Stromes auf das Abtasten einer Spannung hinaus.

² Bei analog arbeitenden sogenannten 'Sample-Hold'-Gliedern fällt die Ausgangsspannung exponentiell mit der Zeit.

B.1 Die Wandlerplatine

Ein digitalisierter Spannungswert wird nur dann im 12 Bit Speicher-Register gespeichert, wenn der aktuelle vom 12 Bit AD-Wandler *AD574A* digitalisierte Spannungswert größer oder kleiner ist als der momentan gespeicherte, je nachdem, ob das Signal nach einem Maximum oder Minimum abgetastet wird. Damit dies in der Weise funktioniert, muß der 'Peak'-Detektor zuvor initialisiert werden: Soll der Maximalwert der Spannung registriert werden, so muß das Speicher-Register vor dem Start des 'Sample'-Vorganges mit dem Minimalwert der Spannung, d. h. mit Nullen, gefüllt werden. Entsprechend muß das Speicher-Register mit Einsen gefüllt werden, wenn der Minimalwert der Spannung registriert werden soll. Die Schaltung der Wandlerplatine zeigt Abb. [B.1](#).

B.1.1 Initialisierung

Der 'Peak'-Detektor befinde sich in dem Zustand, den er nach Beendigung des 'Sample'-Vorganges innehat³: Die Flip-Flops FF1, FF2 und FF3 sind alle im rückgesetzten Zustand ($\overline{Q}$ auf H und Q auf L). Solange sich das Flip-Flop FF1 im rückgesetzten Zustand befindet, ist der Reset-Eingang des Dezimalzählers *4017* auf H und damit der Zähler gelöscht, d. h. der Ausgang 0 (Pin 3) ist auf H und alle anderen Ausgänge sind auf L.

Durch Setzen des Flip-Flops FF1 wird der Vorgang der Initialisierung gestartet, der Reset-Eingang des *4017* (Pin 15) geht auf L und gibt den Zähler frei. Der *4017* wird durch einen als Oszillator beschalteten *NE555* mit 1 kHz getaktet. Die 1 kHz liegen an Pin 3 des *NE555* an. Damit dauert die Phase der Initialisierung insgesamt 4 ms. Mit jedem Takt schreitet der Zähler einen Schritt vorwärts:

Ausgang 1 auf H: Da der *AD574A* in der Phase noch keine Spannungen wandelt, ist der Statusausgang (Pin 28) des *AD574A* permanent auf L. Damit geht der Ausgang des NOR-Gatters *7402* auf L. Des weiteren wird das Flip-Flop FF3 gesetzt, wodurch $\overline{Q}$ auf L geht. Solange $\overline{Q}$ auf L ist, wird auch der Ausgang des AND-Gatters *7408* und damit auch die beiden Eingänge für die Datenfreigabe des 12 Bit Speicher-Registers *74HC173* IE1⁴ und IE2 auf L sein⁵. $\overline{Q}$ auf L bewirkt weiter, daß sich die Ausgänge der Bus-Treiber *74LS126* im hochohmigen Zustand befinden. Auf diese Weise sind die Werte der Eingangsdaten $D0 \dots D11$ bestimmt durch die Pullup-Widerstände ($12 \times 47 \text{ k}\Omega$), deren gemeinsamer Anschluß entweder auf Masse (für Füllung mit Nullen) oder auf +5 V (für Füllung mit Einsen) liegt. Der Ausgang $\overline{Q}$ des *4001* kann nicht 12 TTL-Lasten treiben. Darum wurde eine Treiberstufe bestehend aus drei $4.7 \text{ k}\Omega$ Widerständen und zwei npn-Transistoren *BC547* zwischengeschaltet.

Ausgang 2 auf H: Geht der Ausgang 2 des *4017* (Pin 4) auf H, geht der Ausgang 1

³ Im folgenden bedeutet 'H' logisch High und 'L' logisch Low.

⁴ IE = Input Enable

⁵ Das 12 Bit Speicher-Register ist aufgebaut aus drei einzelnen 4 Bit Speicher-Registern *74HC173* bei denen die Steuereingänge IE1, IE2 und Clock jeweils zusammengeschaltet sind. Deshalb wird im Text auf die Angabe von Pin-Nummern verzichtet. Beim 4 Bit Register *74HC173* liegt IE1 an Pin 10, IE2 an Pin 9 und Clock an Pin 7.

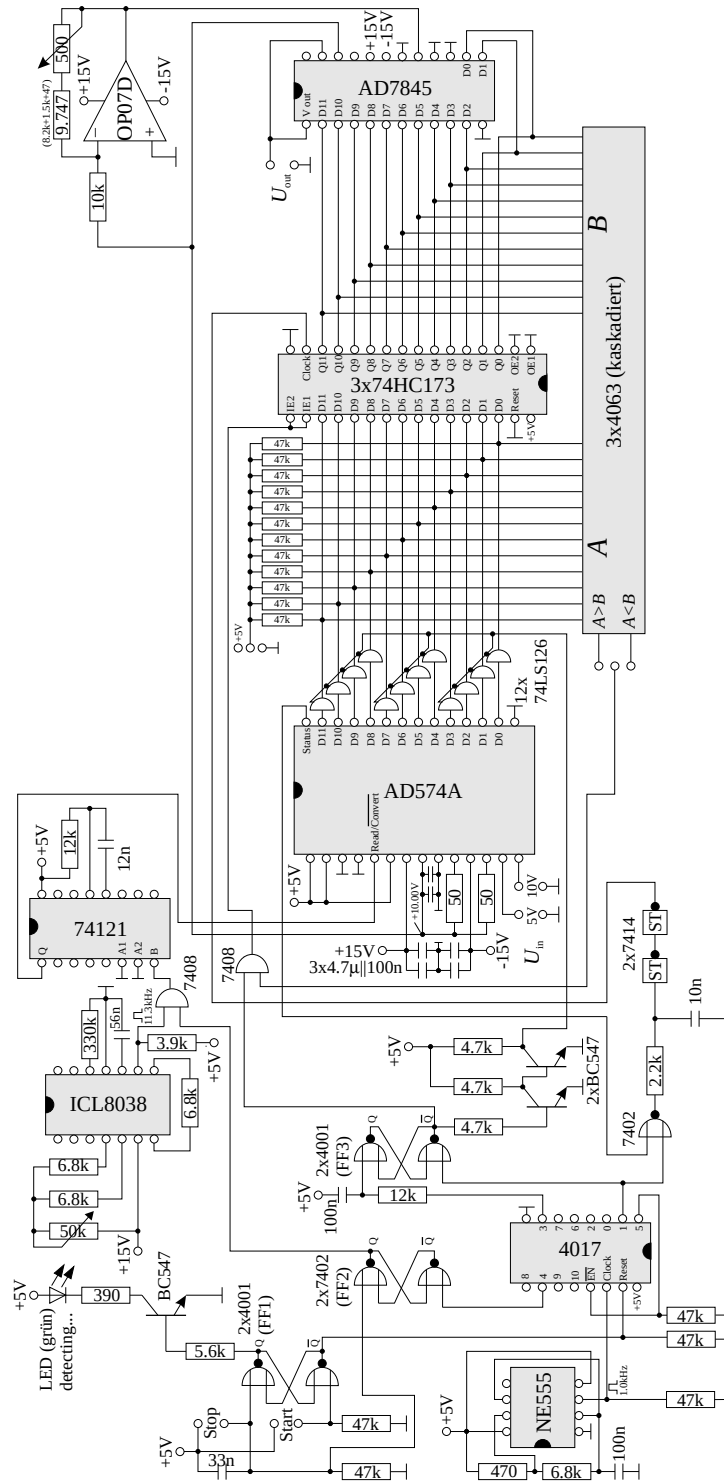


Abb. B.1: Hauptplatine des 'Peak'-Detektors mit Wandlern und Speicher-Register.

(Pin 2) auf L. Da Pin 28 des *AD574A* während der Initialisierung permanent L ist, resultiert daraus am Ausgang des NOR-Gatters 7402 und damit am Clock-Eingang des 74HC173 eine steigende Flanke, die dazu führt, daß die Daten $D0 \dots D11$ in das Register geschrieben werden (IE1 und IE2 sind L, weil das Flip-Flop FF3 gesetzt ist).

Ausgang 3 auf H: Das Flip-Flop FF3 wird zurückgesetzt. Dadurch verhalten sich die Bus-Treiber 74LS126 wieder transparent bezüglich der vom *AD574A* kommenden Daten. Da nun $\overline{Q}$ auf H ist, wird der Ausgang des 7408 bestimmt durch das Resultat des 12 Bit Größenvergleichers: Soll beispielsweise der Maximalwert der Spannung bestimmt werden, so ist der Ausgang $A > B$ des Größenvergleichers mit dem Eingang des 7408 verbunden. Ist die Bedingung $A > B$ für den zuletzt digitalisierten Spannungswert erfüllt, so geht der Ausgang $A > B$ auf L. Damit liegen die Eingänge IE1 und IE2 des 74HC173 ebenfalls auf L, und bei der folgenden steigenden Flanke am Clock-Eingang werden die Daten ins Register geschrieben. Ist die Bedingung $A > B$ nicht erfüllt, werden die Daten bei der steigenden Flanke am Clock-Eingang ignoriert, weil IE1 und IE2 in diesem Fall H sind.

Ausgang 4 auf H: In diesem Moment wird der 'Sample'-Vorgang gestartet. Das Flip-Flop FF2 wird gesetzt und weil Q nun auf H ist, erscheint das Taktsignal am Ausgang des AND-Gatters (7408). Als Taktgenerator dient der *ICL8038*, der ein Rechtecksignal mit einer Frequenz von 11.3 kHz liefert. Die Frequenz des Oszillators kann über einen weiten Bereich mit dem 50 k Ω Trimpotentiometer eingestellt werden. Zwar wird der *ICL8038* mit +15 V betrieben, die Ausgangsspannung an Pin 9 variiert aber nur zwischen 0 und +5 V, was ein direktes Verbinden mit dem Eingang des AND-Gatters erlaubt.

Ausgang 5 auf H: Da der Ausgang 5 (Pin 1) mit dem $\overline{EN}$ -Eingang⁶ (Pin 13) des 4017 verbunden ist, sperrt sich der Zähler selbst. Damit ist die Initialisierung des 'Peak'-Detektors beendet.

B.1.2 'Sample'-Vorgang

Zum Starten einer Digitalisierung muß der Read/ $\overline{\text{Convert}}$ -Eingang des *AD574A* für mindestens 200 ns auf L sein. Erst nachdem Read/ $\overline{\text{Convert}}$ wieder H ist, beginnt die Digitalisierung. Um die Zeit, in der Read/ $\overline{\text{Convert}}$ L ist, präzise einstellen zu können, wurde dem Taktgenerator *ICL8038* das Mono-Flop 74121 nachgeschaltet. Widerstand und Kapazität wurden so dimensioniert, daß der Ausgang Q, der mit dem Read/ $\overline{\text{Convert}}$ des *AD574A* verbunden ist, bei jedem Takt für 250 ns L ist.

600 ns nachdem an den Ausgängen $D0 \dots D11$ des *AD574A* gültige Daten anliegen, geht der Statusausgang auf L. Da der 12 Bit Größenvergleicher für den Vergleich der aktuellen Daten mit denen im Register gespeicherten etwas Zeit benötigt, erscheint die steigende Flanke am Clock-Eingang nicht sofort, nachdem der Statusausgang L geworden ist, sondern einige μs später! Nach Beendigung der Initialisierung ist der Ausgang 5 des 4017 auf H und alle anderen Ausgänge auf L (s.o.). Deshalb verhält sich das NOR-Gatter 7402 wie ein Inverter, wobei das dem 7402 nachgeschaltete RC-Glied

⁶ EN = Enable

die Zeitverzögerung bewirkt. Der Schmitt-Trigger *74LS14* soll sicherstellen, daß der Pegelwechsel mit hoher Flankensteilheit erfolgt, damit undefinierte Zwischenzustände am Clock-Eingang des *74HC173* verhindert werden. Der zweite Schmitt-Trigger hat lediglich die Funktion eines Inverters.

B.1.3 DA-Wandlung

Damit der 'Peak'-Detektor in einfacher Weise in den bestehenden Meßaufbau integriert werden konnte, wurde das Ergebnis des 'Sample'-Vorganges in ein analoges Signal gewandelt. Hierfür wurde der 12 Bit DA-Wandler *AD7845* verwendet.

Über den invertierenden Verstärker *OP07D* wird dem DA-Wandler die Referenzspannung vom *AD574A* (Pin 8) zugeführt. Sind alle Eingänge des *AD7845* *D0...D11* auf H, läßt sich mit dem 500 Ω Trimpmpotentiometer die Ausgangsspannung des *AD7845* (Pin 1) sehr präzise auf 10.00 V einstellen.

B.2 Die Filterplatine

Strenggenommen hätte auf die Filter verzichtet werden können. Da sie aber in jedem Fall eine Verbesserung bei geringem schaltungstechnischen Mehraufwand darstellen, wurden sie dennoch verwendet. Die Schaltung der Filterplatine zeigt [B.2](#).

B.2.1 Hochpaß

Aufgabe des Hochpasses ist, die Eingangsspannung U_{in} vom Gleichspannungsmittelwert zu befreien. Realisiert wurde ein aktives Filter 4. Ordnung mit Einfachmitkopplung und Butterworth-Charakteristik. Die Dimensionierung der Widerstände und Kondensatoren wurde derart gewählt, daß die 3 dB-Grenzfrequenz des Filters 100 Hz beträgt [[TIE93](#)].

Für den Einsatz am Leckstrom-Scanner hat der Hochpaß gar keine Bedeutung, weil dort nur die DC-Komponente des Stromes interessiert. Verwendet wird der Hochpaß beispielsweise dann, wenn die Spitzenwerte des Rauschpegels einer Spannungsquelle bestimmt werden sollen. Der 'Peak'-Detektor ist also auch abseits des Leckstrom-Scanners universell einsetzbar.

B.2.2 Tiefpaß

Unsicher ist das niedrigwertigste Bit des *AD7845*, dessen Wert im allgemeinen zwischen L und H hin und her wechselt. Das hat zur Folge, daß der Ausgangsspannung des *AD7845* (Pin 1) ein Rauschsignal im kHz-Bereich und einem 'Peak-to-Peak'-Wert von einigen Millivolt überlagert ist.

Aufgabe des Tiefpasses, der aus dem Hochpaß durch Vertauschung von Widerständen und Kondensatoren hervorgeht, ist die Ausgangsspannung des *AD7845* von



Abb. B.2: Filterplatine.

diesem Rauschsignal zu befreien. Die 3 dB-Grenzfrequenz des Filters beträgt ebenfalls 100 Hz [TIE93].

B.2.3 Spannungsfolger

Der Spannungsfolger⁷, in den die Spannung U_{in} eingespeist wird, wäre auch verzichtbar, da an ihn der Strom-Monitor des Potentiostaten angeschlossen wird, der einen kleinen Ausgangswiderstand hat. Auch hier wurde wieder das Ziel der universellen Einsetzbarkeit des 'Peak'-Detektors verfolgt. Soll beispielsweise die Amplitude eines LC- oder Quarz-Oszillators bestimmt werden, so ist der Spannungsfolger wichtig, da solche Oszillatoren möglichst wenig belastet werden sollten.

Wird der ± 5 V Meßbereich verwendet, so muß die Ausgangsspannung des AD7845 durch 2 dividiert werden, wofür ein Spannungsteiler Anwendung findet. Mit dem 500 Ω Trimpotentiometer läßt sich der Faktor der Division auf genau 2 einstellen. Hier ist das Nachschalten eines Spannungsfolgers sehr wichtig, weil die Spannung am Abgriff des Spannungsteilers ihren Wert um so mehr ändert, je stärker die Belastung ist.

B.3 Die Steuerplatine

Der 'Peak'-Detektor kann entweder über die Schaltelemente der Frontplatte oder über die parallele Schnittstelle des Computers gesteuert werden. Der Wechsel zwischen beiden Bedienungsmöglichkeiten erfolgt über die Steuerplatine in Abb. B.3.

Bei nicht vorhandener Verbindung zwischen 'Peak'-Detektor und Computer, bewirkt der 39 k Ω Pullupwiderstand am $\overline{\text{SI}}$ -Eingang⁸ (Bit 3 des Kontroll-Registers, Pin 17 des 25 pol. SUB-D-Steckers) ein L am Ausgang des Inverters 4069, was zur Folge hat, daß sich alle Relais in der oberen sowie der unteren Reihe in Abb. B.3 im abgefallenen Zustand befinden und die gelbe Remote-LED nicht leuchtet. In diesem Fall ist die Bedienung des 'Peak'-Detektors über die Schaltelemente der Frontplatte möglich.

Ein H am $\overline{\text{SI}}$ -Eingang hingegen deaktiviert die Schaltelemente der Frontplatte, die gelbe Remote-LED leuchtet und die Steuerung erfolgt über die Bits $D3 \dots D7$ des Daten-Registers (Pin 5 bis Pin 9 des SUB-D-Steckers).

⁷ Siehe Fußnote auf Seite 49.

⁸ SI = **S**elect-**I**nterface

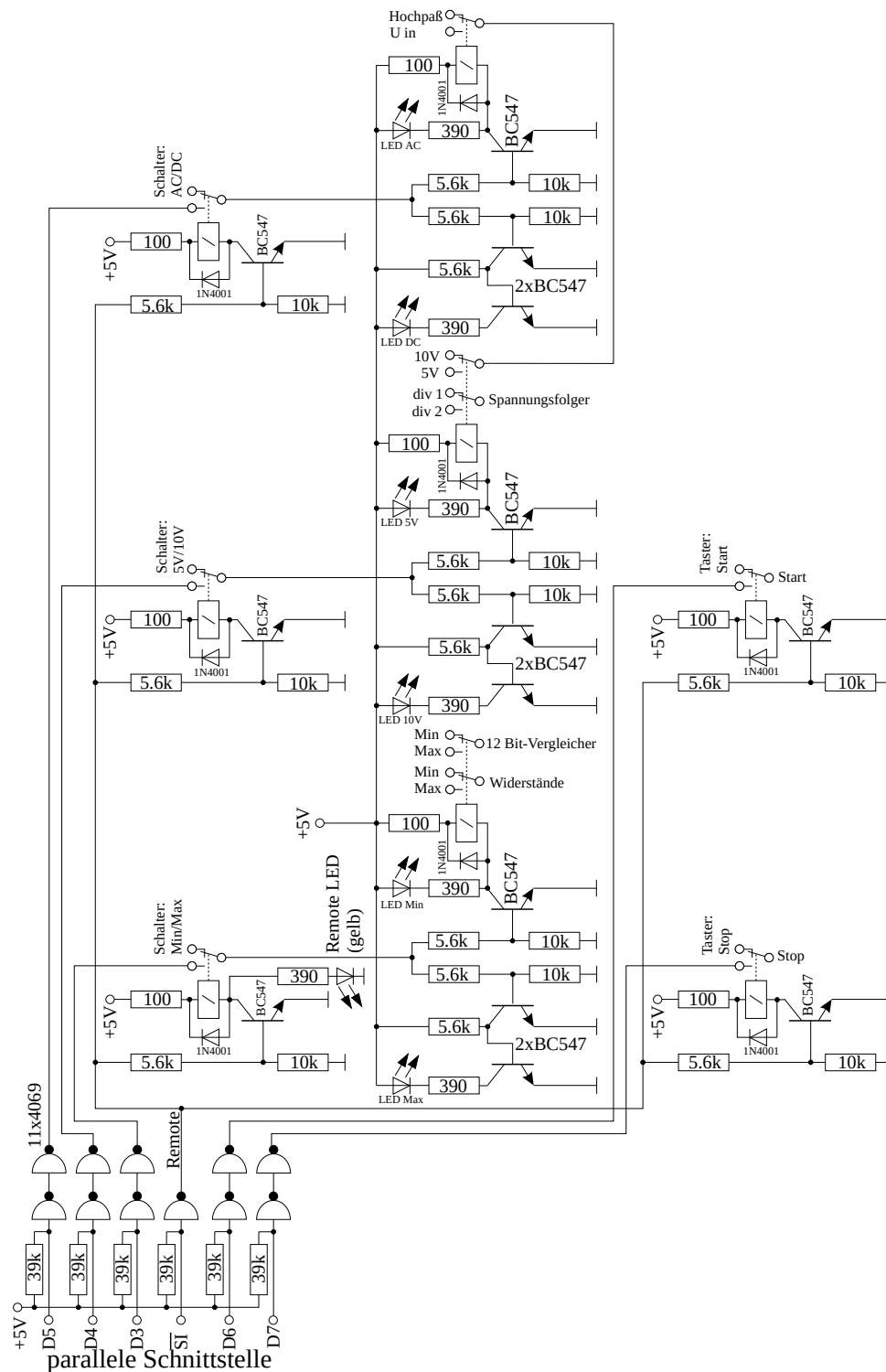


Abb. B.3: Steuerplatine zum Wechseln zwischen manueller Bedienung und Steuerung über die parallele Schnittstelle.

Anhang C

Bad-Heizung

Während die Temperatur bei der RCA-Reinigung weniger kritisch ist, sollte die Temperatur der Politurbeize zum Abätzen des Säge-’Damages’ einigermaßen genau eingehalten werden, siehe Abschnitt 5.1.3. Ist die Temperatur zu hoch, besteht die Gefahr, daß von dem Wafer zu viel Material abgetragen wird. Bedingt durch die je nach Kornorientierung unterschiedlichen Ätzraten, resultieren daraus übermäßig starke Waferunebenheiten, die sich nachteilig auf die Messung auswirken können. Zudem werden die Wafer dadurch unnötig dünn, was deren Handhabung erschwert. Deshalb wurde für die Beseitigung des Säge-’Damages’ und für die RCA-Reinigung ein beheizbares Ätz- bzw. Reinigungsbad aufgebaut, das die Konstanz der Temperatur während des Ätzens bzw. des Reinigens sicherstellt.

Die Politurbeize ist äußerst aggressiv, so daß das Heizen des Bades mit handelsüblichen Hezelementen nicht in Frage kommt. Eine Methode, das Bad zu heizen, besteht darin, das Bad mit einer leistungsstarken Lampe zu erwärmen [LIP97]. Die Methode hat aber zwei wesentliche Nachteile: Erstens ist die Temperatur des Bades nicht genau definiert und zweitens, dieser Nachteil wiegt bedeutend schwerer, kann das Heizen des Bades viele Stunden betragen! Mit der Platin-Heizung reduziert sich die Aufheizphase auf wenige Minuten!

Zum Regeln der Temperatur wird ein Temperatur-Regler von der Firma *Eydam* verwendet. Da der Regler ausgangsseitig nur 220 V schalten kann, wird zwischen Temperatur-Regler und der Platin-Heizwendel noch eine Bad-Heizung zwischengeschaltet, vgl. Abb. 5.3. Die Schaltung der Bad-Heizung zeigt Abb. C.1.

C.1 Sägezahn-generator

Damit die Heizleistung einstellbar ist, wird die heruntertransformierte Ausgangsspannung des Temperatur-Reglers nicht direkt an das Gate des Power-MOSFET¹ BUZ71A abgeschlossen: Um beispielsweise die halbe Heizleistung zu erhalten, sollte der BUZ71A nicht halb geöffnet werden, weil dann über dem Drain-Source-Widerstand R_{DS} ein

¹ MOSFET = Metal-Oxid-Semiconductor Field-Effect Transistor

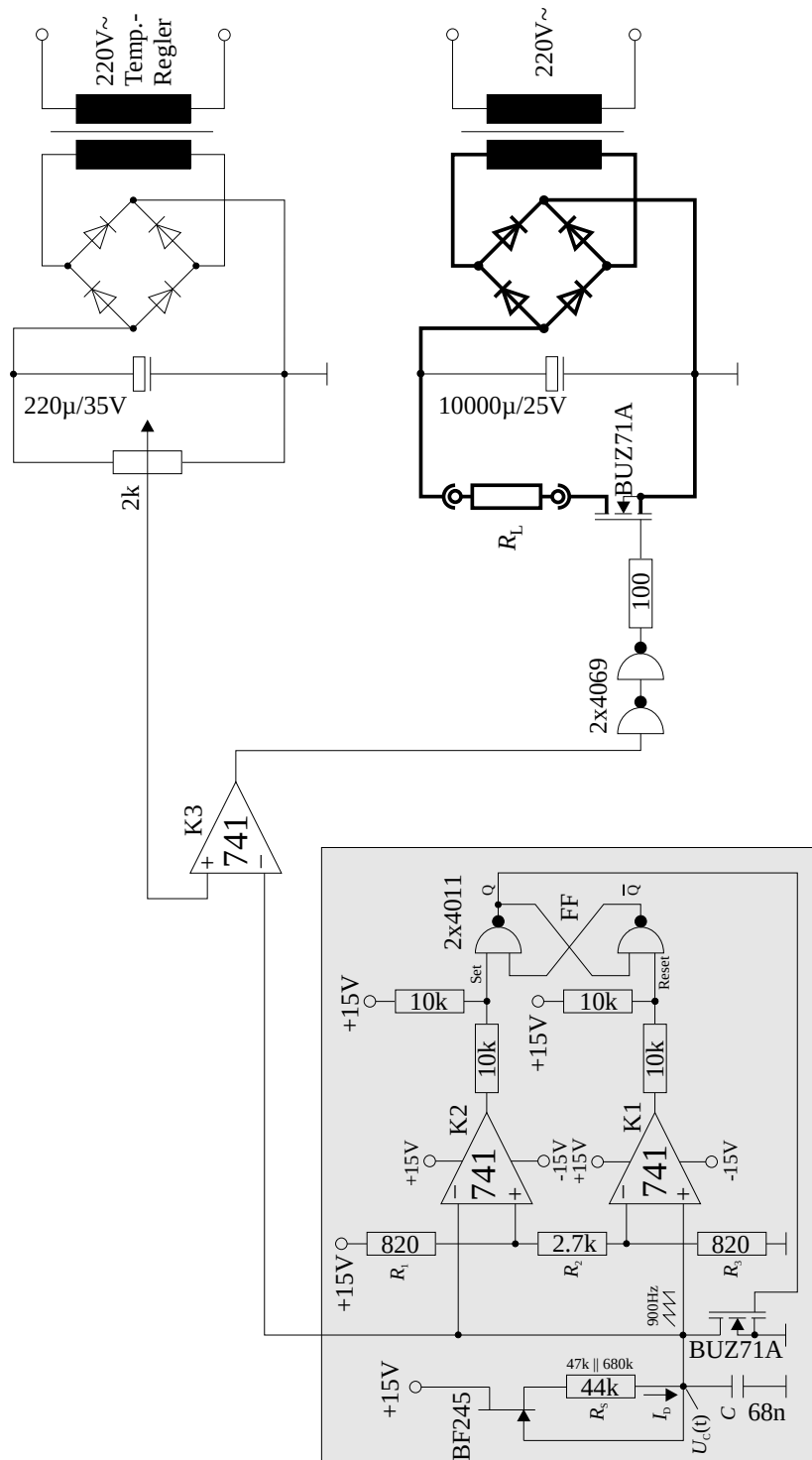


Abb. C.1: Schaltung der Bad-Heizung. Der Schaltungsteil im grau unterlegten Bereich ist ein Sägezahn-generator. Die verdickt gezeichneten Linien kennzeichnen den Lastkreis. R_L bedeutet die Platin-Heizwendel.

großer Teil der Spannung abfällt *und* große Ströme durch den Lastkreis fließen. Insgesamt bedeutet das sehr hohe Verlustleistungen am Power-MOSFET, die sich kaum abkühlen lassen.

Deshalb erfolgt die Einstellung der Heizleistung über einen *Plusbreitenmodulator*. Hierbei wird an das Gate des BUZ71A ein Rechtecksignal fester Frequenz, aber variabler Pulsbreite angelegt. Da der BUZ71A entweder voll zu- oder voll aufgesteuert ist, bleiben die Verlustleistungen klein. Im voll aufgesteuerten Zustand beträgt der Drain-Source-Widerstand des BUZ71A $R_{DS(on)} = 100 \text{ m}\Omega$.

Der Schaltungsteil im grau unterlegten Bereich in Abb. C.1 ist ein Sägezahngenerator. Für die Beschreibung seiner Funktion wird vom Zustand des vollständig entladenen Kondensators ausgegangen.

Die Kombination aus dem JFET² BF245, dem Sourcewiderstand R_S und dem Kondensator C stellt eine *Konstantstromquelle* dar. Der Konstantstrom I_D ist über R_S einstellbar. Der Sourcewiderstand bewirkt, daß das Gate durch $I_D R_S$ in Sperrrichtung vorgespannt ist, wodurch I_D reduziert und der JFET näher an den Abschnür-(Pinch-off-)Punkt gebracht wird. Des weiteren bewirkt der Sourcewiderstand eine stromsensitive Rückkopplung, was sich positiv auf die Konstanz des Stromes auswirkt [HOR89].

Ist Q die im Kondensator C gespeicherte Ladung, so gilt für die über C abfallende Spannung U_C :

$$U_C = \frac{Q}{C}. \quad (\text{C.1})$$

Mit $I_D = \frac{dQ}{dt}$ folgt daraus:

$$\frac{dU_C}{dt} = \frac{1}{C} \cdot \frac{dQ}{dt} = \frac{1}{C} \cdot I_D. \quad (\text{C.2})$$

Integration über die Zeit t ergibt:

$$U_C(t) = \frac{I_D}{C} \cdot \int dt = \frac{I_D}{C} \cdot t. \quad (\text{C.3})$$

D. h. die Kondensatorspannung steigt *linear* mit der Zeit t !

Ist U_C klein, so ist die Ausgangsspannung des Komparators³ K1 auf -15 V ($U_+ - U_- < 0$) und die des Komparators K2 auf +15 V ($U_+ - U_- > 0$). Damit ist der Reset-Eingang des Flip-Flops L (0 V) und der Set-Eingang H (+15 V). Das Flip-Flop befindet sich damit im rückgesetzten Zustand. Da der MOSFET sperrt (Ausgang Q ist auf L), lädt sich der Kondensator C mit der Zeit auf, wobei die Kondensatorspannung U_C nach Gleichung C.3 linear mit der Zeit steigt. Er wird deshalb parallel zum Kondensator der BUZ71A geschaltet, weil im gesperrten Zustand sein R_{DS} -Widerstand sehr hochohmig ist. Daher fließen über R_{DS} während der Ladephase des Kondensators keine Ladungen ab.

² JFET = **J**unction **F**ield-**E**ffect **T**ransistor

³ Der 741 ist, wie jeder andere Operationsverstärker auch, nicht in der Lage, die Ausgangsspannung über den vollen Bereich der Betriebsspannung auszusteuern, sondern einige Volt weniger. Um die Beschreibung der Funktion durch umständlichen Text nicht unnötig zu verkomplizieren, wird im folgenden davon ausgegangen, daß der Aussteuerbereich $\pm 15 \text{ V}$ betrage.

Ist U_C so weit gestiegen, daß $U_C > \frac{R_3}{R_1+R_2+R_3} \cdot 15 \text{ V}$ wird, springt die Spannung am Ausgang des Komparators K1 auf +15 V, womit der Reset-Eingang des Flip-Flops auf H geht. Da nun beide Eingänge H sind, ändert sich am Zustand der Ausgänge des Flip-Flops nichts⁴, d. h. Q bleibt auf L und U_C steigt weiterhin linear mit der Zeit.

Überschreitet die Kondensatorspannung U_C die Schwelle $\frac{R_2+R_3}{R_1+R_2+R_3} \cdot 15 \text{ V}$, springt die Spannung am Ausgang des Komparators K2 auf -15 V, da nun $U_+ - U_- < 0$ geworden ist. Weil der Set-Eingang L geht, wird das Flip-Flop gesetzt und bringt den MOSFET in den leitenden Zustand. Da der MOSFET für $U_{GS} = +15 \text{ V}$ voll aufgesteuert ist, beträgt der Drain-Source-Widerstand $R_{DS} = 100 \text{ m}\Omega$. Über den R_{DS} -Widerstand entlädt sich der Kondensator sehr rasch. Wird $U_C < \frac{R_2+R_3}{R_1+R_2+R_3} \cdot 15 \text{ V}$ sind beide Eingänge auf H und das Flip-Flop bleibt im gesetzten Zustand. Erst wenn U_C die Schwelle von $\frac{R_3}{R_1+R_2+R_3} \cdot 15 \text{ V}$ unterschreitet, geht der Reset-Eingang L und setzt das Flip-Flop zurück. Der MOSFET sperrt und der Zyklus beginnt von neuem.

C.2 Pulsbreitenmodulation

Die 220 V, die der Temperatur-Regler auf den Ausgang schaltet, werden gleichgerichtet und an das 2 k Ω Trimpmpotentiometer angelegt. Die Stellung des Abgriffes bestimmt die Spannung, die am nicht-invertierenden Eingang des Komparators K3 anliegt, und damit die Pulsbreite des am Ausgang anliegenden Rechtecksignals.

Wie in der Fußnote auf Seite 151 beschrieben, steuert der 741 den Ausgang nicht bis zur Betriebsspannung aus. Da der High-Pegel des CMOS-Inverters 4069 dicht an der Betriebsspannung liegt, werden zwei von ihnen zwischen den Ausgang des 741 und das Gate des BUZ71A geschaltet. Auf diese Weise springt die Spannung am Gate zwischen 0 V und (fast) +15 V, wobei die mittlere Heizleistung über das 2 k Ω Trimpmpotentiometer eingestellt werden kann.

⁴ Sind beide Eingänge eines aus zwei NAND-Gattern aufgebauten Flip-Flops auf L, ist das ein nicht definierter Zustand. Das Flip-Flop soll sozusagen gleichzeitig gesetzt und zurückgesetzt werden. Für ein aus zwei NOR-Gattern aufgebautes Flip-Flop, wie in Abb. B.1, ergibt sich ein nicht definierter Zustand, wenn beide Eingänge auf H sind.

Anhang D

Temperatur-Regler

In Abschnitt 2.2.3 wurde erläutert, daß für Elymat-Messungen im BSC-Mode ein Laser mit einer Eindringtiefe von $\alpha^{-1} = 90 \mu\text{m}$ verwendet werden sollte. Mit Gleichung 2.3 errechnet sich die Wellenlänge des Lasers zu $\lambda = 974 \text{ nm}$. Ein Laser mit dieser Wellenlänge konnte nicht erworben werden, deshalb wurde für BSC-Messungen eine Laserdiode mit einer Wellenlänge von $\lambda = 995 \text{ nm}$ verwendet. Die Laserdiode, mit einer Leistung von 20 mW, wurde von der Firma *TOPAG Lasertechnik GmbH* bezogen. Da die Wellenlänge mit $0.3 \frac{\text{nm}}{\text{K}}$ relativ stark von der Temperatur der aktiven Region des Lasers abhängt, muß die Laserdiode während des Betriebes temperaturgeregelt werden.

D.1 Aufbau des Reglers

Weil ein Temperatur-Regler mit stufenlosem Analogausgang nicht erwerbbar war, wurde für die Temperatur-Regelung ein PID-Regler konstruiert, der aus drei Komponenten aufgebaut ist:

- dem eigentlichen PID-Regler,
- einer Strombegrenzung und
- einem Modul, daß die Gatespannungen der Power-MOSFETs liefert.

Die drei Komponenten werden im folgenden beschrieben.

D.1.1 PID-Regler

Die Schaltung des PID-Reglers zeigt Abb. D.1. Die Kombination aus dem JFET *BF245*, dem $390 \text{ k}\Omega$ Widerstand und dem $100 \text{ k}\Omega$ Trimpotentiometer liefert einen Konstantstrom für den Thermistor R_ϑ . Die Funktion dieser Konstantstromquelle wurde bereits in Abschnitt C.1 beschrieben. Mit dem Trimpotentiometer wird der Konstantstrom auf $10 \mu\text{A}$ eingestellt. Hierfür wird der Thermistor durch einen Widerstand mit genau $100 \text{ k}\Omega$ ersetzt. Fließen $10 \mu\text{A}$ beträgt der Spannungsabfall über dem $100 \text{ k}\Omega$ Widerstand $U_{100 \text{ k}\Omega} = 10 \mu\text{A} \cdot 100 \text{ k}\Omega = 1 \text{ V}$. Der Konstantstrom von $10 \mu\text{A}$ ist ausreichend und sollte nicht wesentlich höher gewählt werden. Sonst besteht die Gefahr, daß der

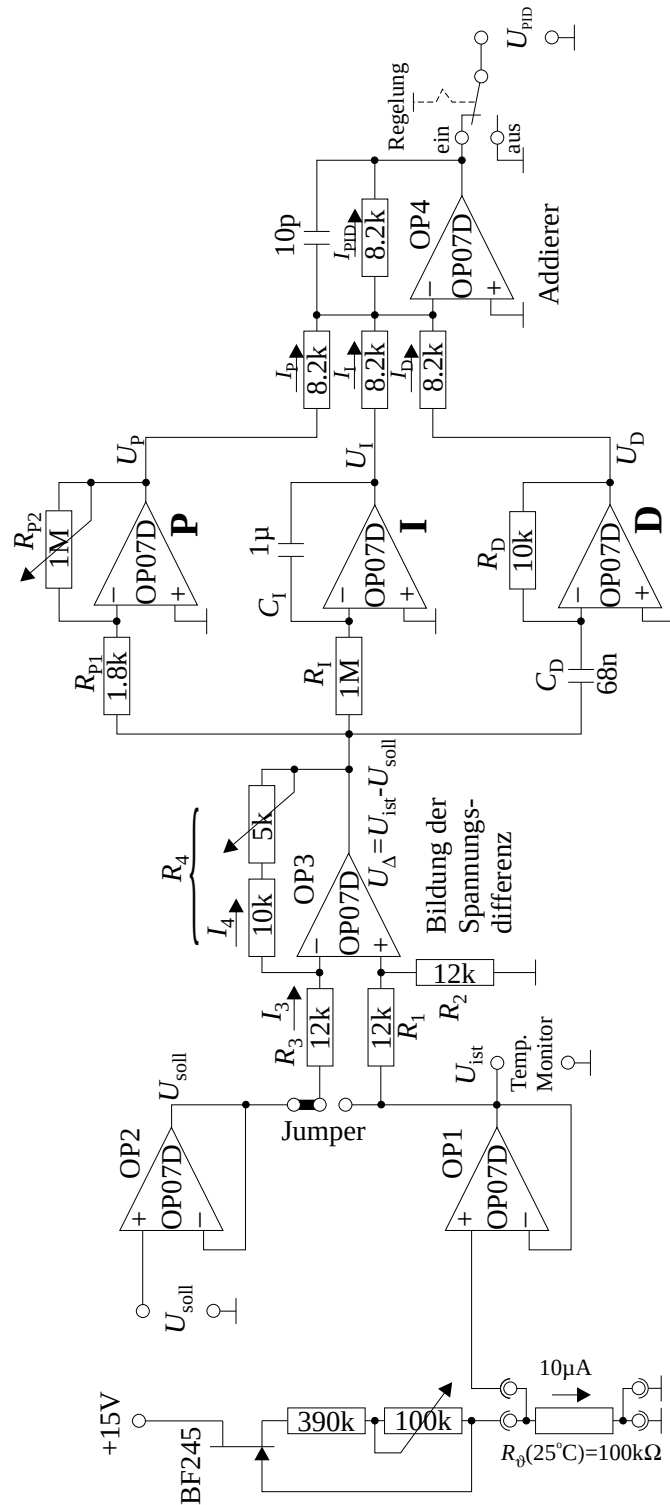


Abb. D.1: Hauptplatine des Temperatur-Reglers, die den eigentlichen PID-Regler enthält.

390 k Ω Widerstand und das Trimpotentiometer vom Strom 'geheizt' werden. Dadurch könnte sie ihre Werte und damit den Konstantstrom ändern. Letztlich würde das eine fehlerhafte Temperatur-Messung bedeuten. Um Spannungsabfälle über den Zuleitungen nicht mitzumessen, wird die Spannung direkt am Thermistor gemessen (Methode der Vier-Punkt-Messung), siehe Abb. D.1. Der Spannungsfolger¹ OP1 trägt dafür Sorge, daß die Messung stromlos erfolgt. Am *Temp. Monitor* kann die aktuelle Temperatur in Form der Spannung U_{ist} abgelesen werden.

Die Temperatur, die die Laserdiode haben soll, wird in Form der Spannung U_{soll} von außen vorgegeben, beispielsweise $U_{\text{soll}} = 1\text{V}$ für $T = 25^\circ\text{C}$. Hierfür ist eine Spannungsquelle zu verwenden, deren Spannung möglichst stabil sein sollte. Für den Fall, daß die Spannungsquelle nicht belastet werden darf, sorgt der Spannungsfolger OP2 dafür, daß die Belastung unterbleibt.

Die Spannungen U_{soll} und U_{ist} werden dem OP3 zugeführt, der die Regeldifferenz $U_{\Delta} = U_{\text{ist}} - U_{\text{soll}}$ bildet: Sind alle Widerstände gleich groß, beträgt die Spannung am nicht-invertierenden Eingang $U_+ = \frac{R_2}{R_1+R_2} \cdot U_{\text{ist}} = \frac{U_{\text{ist}}}{2}$. Der Strom I_3 durch den Widerstand R_3 berechnet sich zu: $I_3 = \frac{U_{\text{soll}} - U_+}{R_3}$. Entsprechend ergibt sich für den Strom durch den Widerstand R_4 : $I_4 = \frac{U_{\Delta} - U_+}{R_4}$, wobei U_{Δ} die Ausgangsspannung des OP3 ist. Da kein Strom in einem Operationsverstärker hineinfließt gilt $I_3 + I_4 = 0$. Mit $R_3 = R_4$ ergibt sich daraus: $2U_+ = U_{\text{soll}} + U_{\Delta}$. Mit dem obigen Wert für U_+ ergibt sich schließlich für die Ausgangsspannung des OP3: $U_{\Delta} = U_{\text{ist}} - U_{\text{soll}}$.

Um den Subtrahierer abzugleichen, wird der Jumper umgesetzt. Dadurch wird erreicht, daß beide Eingangsspannungen gleich sind. Mit dem 5 k Ω Trimpotentiometer wird die Ausgangsspannung des OP3 auf 0 V eingestellt.

Die Regeldifferenz U_{Δ} wird nun dem Proportional-, dem Integral- und dem Differenzglied zugeführt, die alle als invertierende Verstärker beschaltet sind. Im eingeregten Zustand liegt, wegen $U_+ - U_- = 0$, der invertierende Eingang auf Massepotential (virtuelle Masse). Unter Berücksichtigung der Bedingung, daß am invertierenden Eingang die Summe aller Ströme gleich Null ist, berechnen sich die Ausgangsspannungen der einzelnen Glieder wie folgt:

$$\begin{aligned} \mathbf{P - Glied} &: U_P(t) = -\frac{R_{P2}}{R_{P1}} \cdot U_{\Delta}(t), \\ \mathbf{I - Glied} &: U_I(t) = -\frac{1}{R_I C_I} \cdot \int_{t_0}^t U_{\Delta}(\tau) d\tau, \\ \mathbf{D - Glied} &: U_D(t) = -R_D C_D \cdot \frac{dU_{\Delta}(t)}{dt}. \end{aligned}$$

Die drei Ausgangsspannungen werden im OP4 aufaddiert. Da der Addierer auch als invertierender Verstärker beschaltet ist, wird die jeweilige Phasendrehung um 180° in den Gliedern zuvor, siehe Minuszeichen in den obigen Gleichungen, wieder rückgängig gemacht. Es gilt $I_P + I_I + I_D + I_{PID} = 0$. Multipliziert man die Gleichung mit $R (= 8.2 \text{ k}\Omega)$

¹ Siehe Fußnote auf Seite 49.

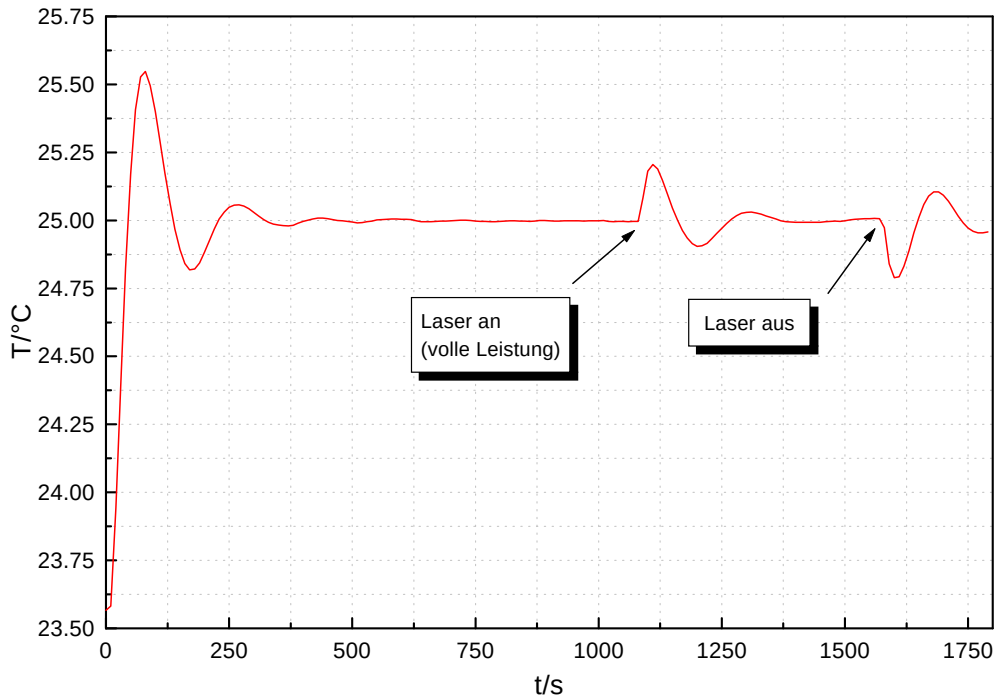


Abb. D.2: Bei zu kleiner Verstärkung des P-Gliedes, ändert sich beim An- oder Ausschalten die Temperatur des Lasers.

folgt:

$$U_{\text{PID}} = -(U_{\text{P}} + U_{\text{I}} + U_{\text{D}}).$$

Der 10 pF Kondensator koppelt den Verstärker für hohe Frequenzen gegen. Auf diese Weise wird verhindert, daß die Schaltung schwingt². Über den Schalter kann die Regelung ausgeschaltet werden, indem U_{PID} mit Masse verbunden wird.

An den Ausgang U_{PID} wird die Leistungsstufe (Stellglied bestehend aus Power-MOSFETs und Peltier-Element) angeschlossen. Da während einer Elymat-Messung der Laser ständig ein- und ausgeschaltet wird, wurde der Regler so dimensioniert, daß er sehr hart regelt. Sonst würde sich beim Ein- oder Ausschalten des Lasers dessen Temperatur merklich ändern, siehe Abb. D.2. Weil kurz nach Inbetriebnahme, wegen $|U_{\text{soll}} - U_{\text{ist}}| \gg 0$, sehr hohe Ströme durch das Peltier-Element fließen würden, muß der Peltier-Strom begrenzt werden. Die Funktion der Strombegrenzung wird im nächsten Abschnitt beschrieben.

² Eigentlich hätte der 10 pF Kondensator auch fortgelassen werden können, weil das zu regelnde System Tiefpaßeigenschaften hat. Dies verhindert, daß die Schaltung schwingen kann. Die relativ große Wärmekapazität macht das System thermisch träge, so daß es auf hochfrequente Schwingungen seitens des Verstärkers OP4 nicht reagieren kann. D. h. die Spannung am Thermistor kann keine hochfrequenten Anteile enthalten.

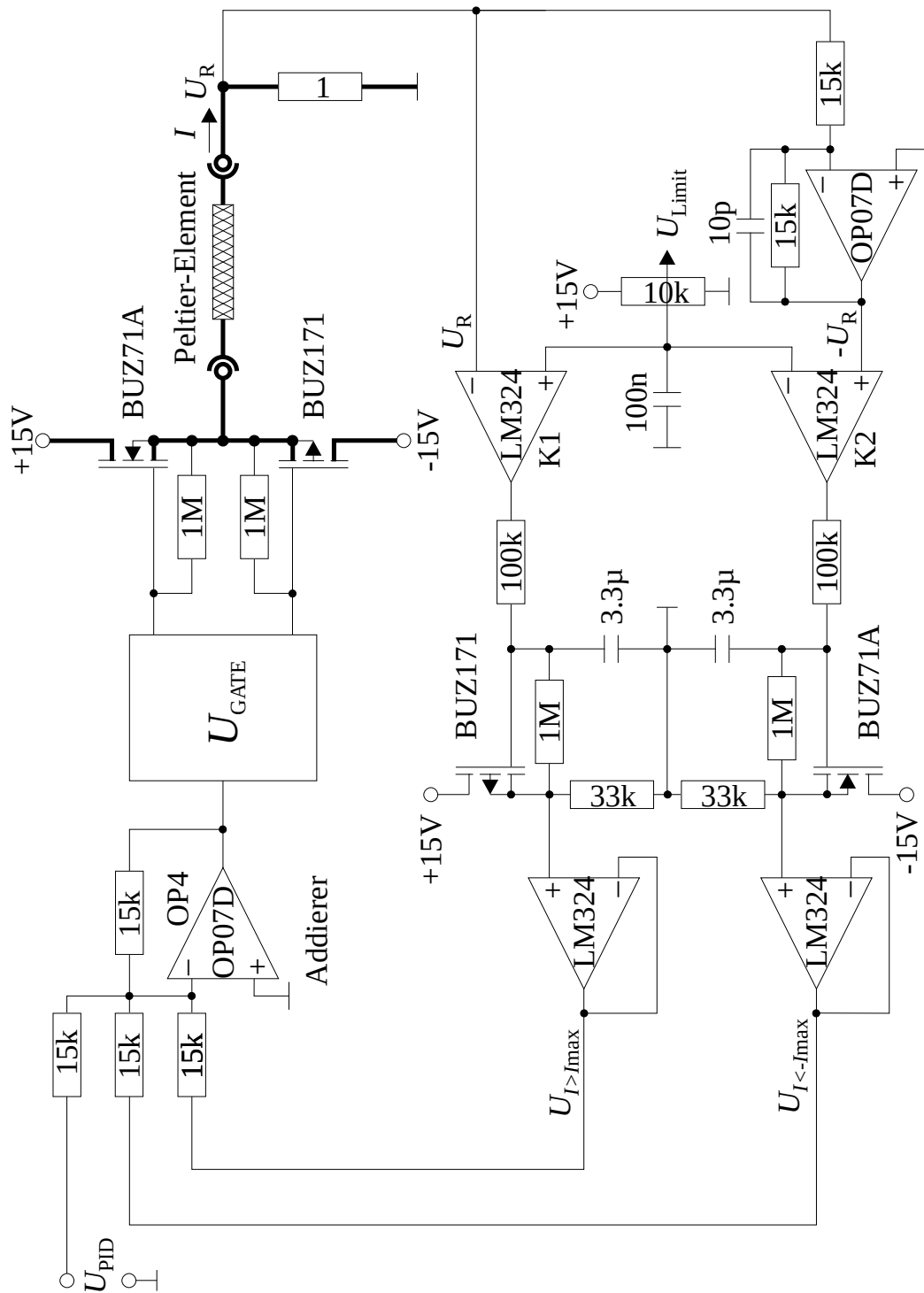


Abb. D.3: Diese Schaltung begrenzt den Strom durch das Peltier-Element.

D.1.2 Strombegrenzung

Das Stromlimit I_{Limit} kann mit dem 10 k Ω Trimpotentiometer eingestellt werden. Da der Strom-Spannungs-Konverter aus einem 1 Ω Widerstand besteht, gilt für die einzustellende Spannung U_{Limit} am Abgriff des Trimpotentiometers: $U_{\text{Limit}} = I_{\text{Limit}} \cdot 1 \Omega$. Da nur ein Trimpotentiometer vorhanden ist, ist das Stromlimit I_{Limit} sowohl für positive als auch für negative Ströme gleich.

Ist der Strom null oder klein, gilt für den Komparator K1: $U_+ - U_- = U_{\text{Limit}} - U_R > 0$ und für den Komparator K2: $U_+ - U_- = -U_R - U_{\text{Limit}} < 0$. Damit betragen die Ausgangsspannungen der Komparatoren³ +15 V bzw. -15 V. Beide MOSFETs sperren und die Ausgangsspannungen der Spannungsfolger betragen $U_{I>I_{\text{max}}} = U_{I<-I_{\text{max}}} = 0$ V. Da in diesem Fall zur Ausgangsspannung der PID-Platine U_{PID} zwei Nullen addiert werden, beträgt die Ausgangsspannung des Addierers $-U_{\text{PID}}$. Die Phasendrehung um 180° wird einfach durch Vertauschen der Anschlüsse des Peltier-Elementes wieder kompensiert.

Im folgenden gelte $U_{\text{PID}} < 0$. Dann ist die Ausgangsspannung des Addierers und damit der Strom I positiv, weil der *BUZ71A* durchsteuert, der *BUZ171* hingegen sperrt. Werden nun $|U_{\text{PID}}|$ und damit der Strom I so groß, daß der Wert von I_{Limit} überschritten wird, gilt für den Komparator K1: $U_+ - U_- = U_{\text{Limit}} - U_R < 0$ und dessen Ausgangsspannung springt auf -15 V. Das RC-Glied aus 100 k Ω Widerstand und 3.3 μF Kondensator sorgt dafür, daß der *BUZ171* weich durchschaltet. Weil damit $U_+ > 0$, gilt $U_{I>I_{\text{max}}} > 0$ und zur negativen Spannung U_{PID} wird eine positive Spannung addiert. Insgesamt verringert sich die Ausgangsspannung des Addierers und verhindert ein weiteres Anwachsen des Stromes.

Für den Fall, daß die Strombegrenzung wirkt, weil $|U_{\text{PID}}|$ zu groß ist, muß die Schaltung schwingen, da durch Verringerung der Ausgangsspannung des Addierers kurz darauf wieder $I < I_{\text{Limit}}$ gelten wird. Daraufhin wirkt die Strombegrenzung nicht mehr, so daß kurze Zeit später wieder $I > I_{\text{Limit}}$ gilt, die Strombegrenzung wirkt wieder usw. Durch Einfügen des RC-Gliedes wird erreicht, daß die Schaltung nicht zu wild schwingt. Von Vorteil ist, daß sich die Schaltung der Strombegrenzung in gar keiner Weise bemerkbar macht, wenn $|I| < I_{\text{Limit}}$ erfüllt ist. Grund hierfür sind die beiden Komparatoren K1 und K2, die Nullen zu U_{PID} addieren, wenn $|I| < I_{\text{Limit}}$ gilt.

D.1.3 Modul für die Gatespannungen

Werden die Gate-Anschlüsse der MOSFETs zusammengeschaltet, so existiert ein relativ weiter Spannungsbereich, in dem *beide* MOSFETs sperren. D. h. über diesem Spannungsbereich hat die Übertragungskennlinie des Stellgliedes die Steigung null, was zur Folge hat, daß die Temperatur-Regelung nicht richtig funktioniert: Der Regler würde stark schwingen. Um das zu vermeiden, müssen die Gate-Anschlüsse vorgespannt werden⁴. Die Schaltung, die das bewerkstelligt zeigt Abb. D.4. Mit dem Trimpoten-

³ Siehe Fußnote auf Seite 151.

⁴ Es würde auch genügen, nur ein Gate vorzuspannen. Allerdings gilt dann nicht mehr $U_{\text{PID}} = 0$, wenn der Strom I durch das Peltier-Element gleich null ist. Um die Symmetrie des Stellgliedes zu wahren und, weil der schaltungstechnische Mehraufwand gering ist, wurde beiden Gates vorgespannt.

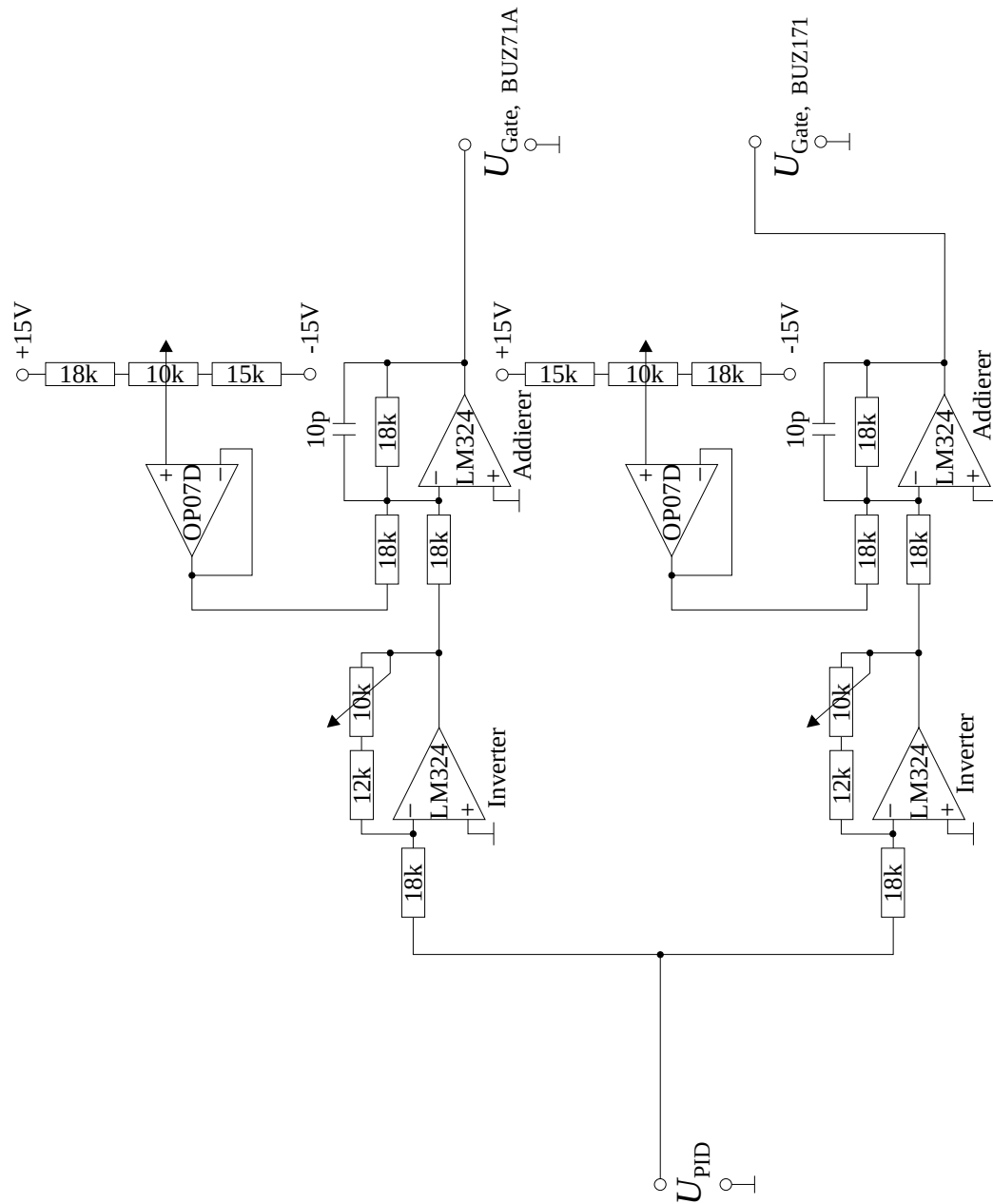


Abb. D.4: Modul zum Vorspannen der Gate-Anschlüsse der MOSFETs.

tiometer im Rückkopplungsweig der Inverter, kann die Steilheit im negativen und positiven Zweig der Übertragungskennlinie eingestellt werden. Mit den beiden anderen Trimpotentiometern können die Gates vorgespannt und damit die Zweige auf der Spannungsachse verschoben werden.

Es ist ratsam, einen kleinen Knick in der Übertragungskennlinie zu belassen, denn je näher beide Zweige zusammenrücken, um so größer wird der Querstrom durch die Leistungsstufe. Wird der Querstrom zu groß, werden die MOSFETs während des Betriebes unnötig thermisch belastet, was gar in eine Zerstörung der Leistungsstufe münden kann!

D.2 Regelverhalten

Abb. D.2 zeigt, daß das An- und Ausschalten des Lasers zu merklichen Temperaturänderungen führt, wenn die Verstärkung des P-Gliedes zu klein ist. Darum wurde der Regler, wie bereits weiter oben erwähnt, so dimensioniert, daß er sehr hart regelt. Aus diesem Grund zeigt der Temperaturverlauf des Lasers kurz nach Inbetriebnahme des Reglers ein starkes Überschwingen, siehe Abb. D.5 a). Abb. D.5 b) zeigt die gleiche Messung, allerdings umskaliert zwecks Bestimmung der Regelgenauigkeit. Hiernach bleiben die Abweichungen der Temperatur unter $\Delta T = 0.005 \text{ }^{\circ}\text{C}$, entsprechend einer relativen Abweichung von $\frac{\Delta T}{T} < 0.02 \text{ } \%$!

Während dieser Messung war der Laser aber permanent ausgeschaltet. Um zu prüfen, ob die Temperatur des Lasers auch dann stabil ist, wenn er wechselweise an- und ausgeschaltet wird, was bei Elymat-Messungen stets der Fall ist, wurde dies anhand einer zweiten Messung untersucht. Hierfür wurde die Temperatur über einen Zeitraum von 2 Stunden gemessen, wobei in immer kürzer werdenden Abständen der Laser an- und ausgeschaltet wurde. Im eingeschalteten Zustand emittierte der Laser mit maximaler Leistung. Das Resultat der Messung zeigt Abb. D.6. Trotz etwas größerer Abweichungen ist die erreichte Temperatur-Stabilität von $\frac{\Delta T}{T} < 0.10 \text{ } \%$! für Elymat-Messungen völlig ausreichend. Die Ungenauigkeiten in der Temperatur von $\Delta T = 0.025 \text{ }^{\circ}\text{C}$ führen zu einer Änderung der Laserwellenlänge von nur $\Delta \lambda < 0.01 \text{ nm}$. Diese Ungenauigkeit in der Wellenlänge ist insofern vernachlässigbar, als daß andere Unsicherheiten, wie beispielsweise die Bestimmung der Waferdicke, wesentlich schwerer wiegen!

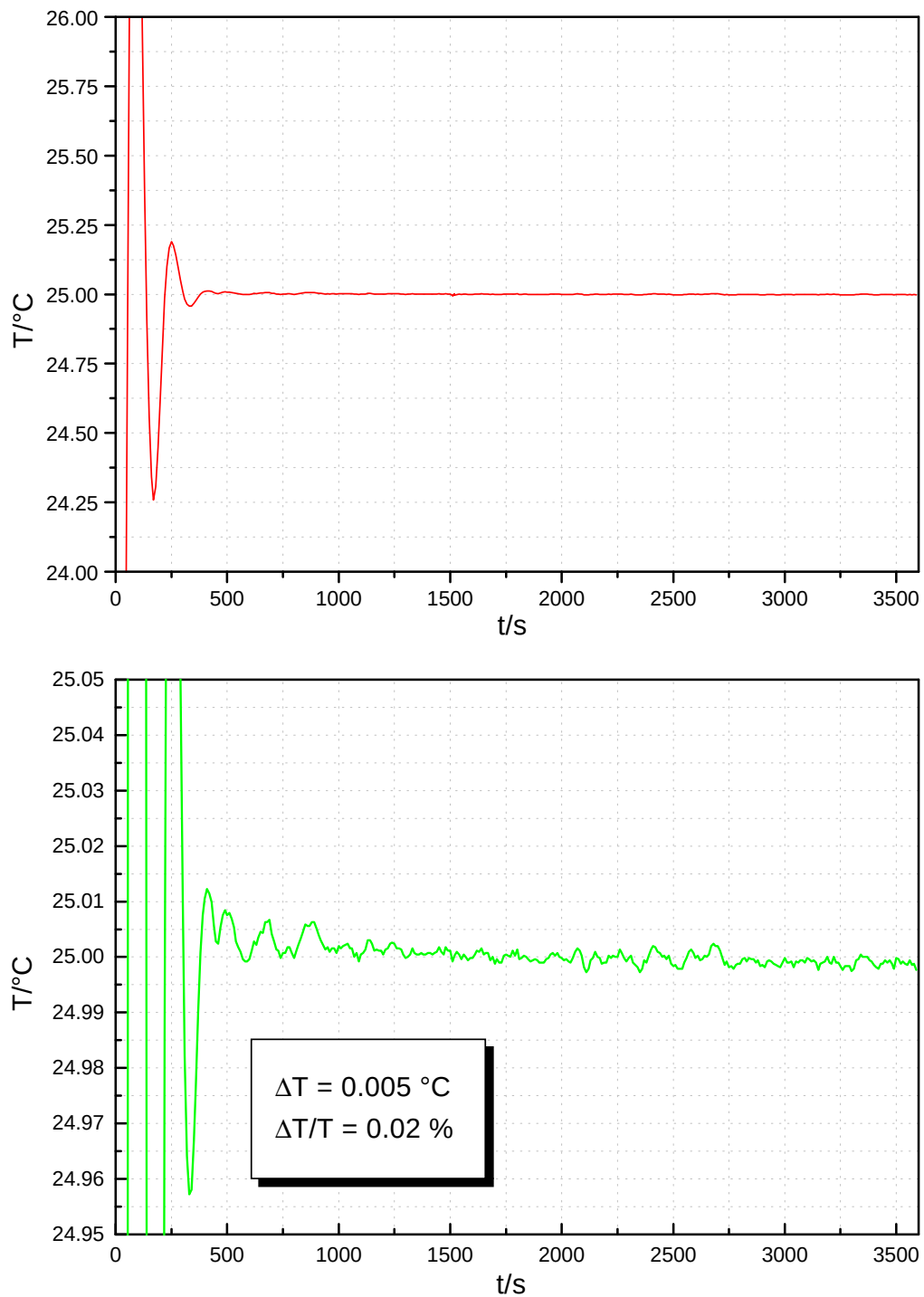


Abb. D.5: a) Regelverhalten des Temperatur-Reglers. b) Abbildung umskaliert, um die Regelschwingung sichtbar zu machen.

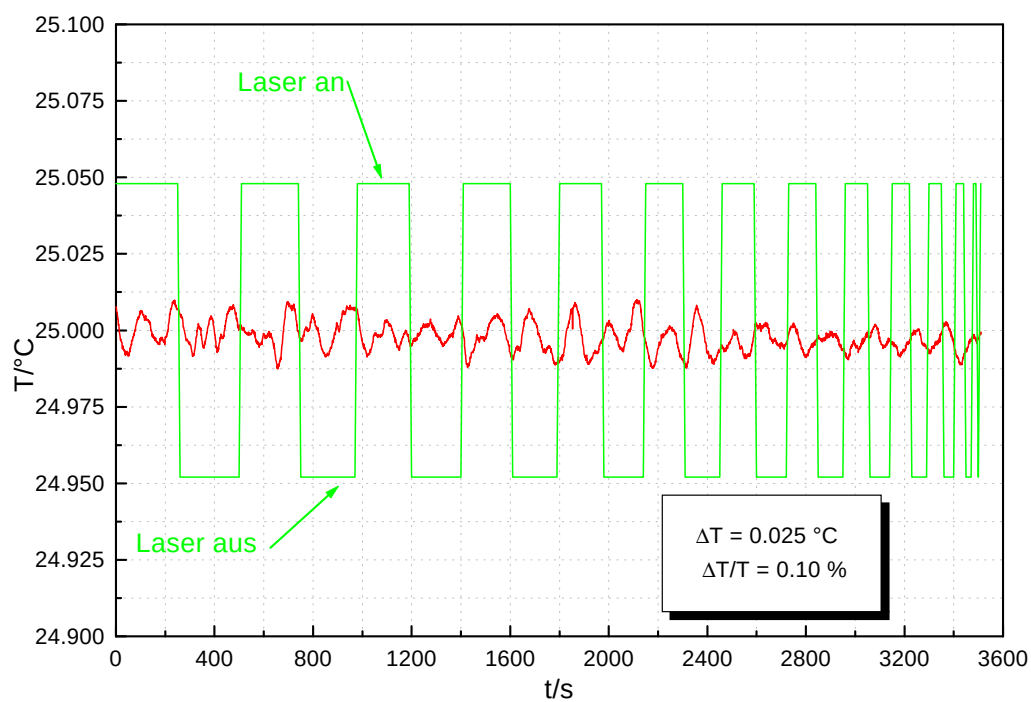


Abb. D.6: Temperatur in Abhängigkeit der Zeit, wenn der Laser währenddessen an- und ausgeschaltet wird. Im eingeschalteten Zustand emittierte der Laser mit maximaler Leistung.

Literaturverzeichnis

- [ABA95] W.W. Abadeer, A. LeBlanc, A. Kluwe, P. Schulz, R. Vollertsen, V. Penka, S. Nadahara, A. Antreasyan, D. Cote, W. Cote, M. Levy, H. Akatasu and R. Ludeke, Mater. Res. Soc. Symp. Proc. (1995), **386** (Ultraclean Semiconductor Processing Technology and Surface Chemical Cleaning and Passivation), 261-9 [31](#)
- [ADA98] B. Adamowicz and H. Hasegawa, Jpn. J. Appl. Phys. Vol. 37 (1998) pp. 1631-1637 [7](#), [36](#)
- [ASK96] D.R. Askeland: "Materialwissenschaften", Spektrum Akademischer Verlage, Heidelberg, Berlin, Oxford 1996 [115](#), [116](#)
- [BAG82] J.A. Baglio, G.S. Calabrese, E. Kamieniecki, R. Kershaw, C.P. Kubiak, A.J. Ricco, A. Wold, M.S. Wrighton and G.D. Zoski, *J. Electrochem. Soc.*, **129**(7) 1461-1472 (1982) [67](#)
- [BAR47] J. Bardeen, Phys. Rev. **71**, p. 717, 1947 [13](#)
- [BAR73] C.R. Barrett, W.D. Nix and A.S. Tetelman, "The Principles of Engineering Materials", Prentice-Hall, Inc., Englewood Cliffs, New Jersey 1973 [115](#), [116](#)
- [BAR92] A.J. Bard and F.R. Fan, Faraday Discuss., 1992, **94**, 1-22 [17](#), [43](#)
- [BER91] W. Bergholz, G. Zoth, G. Götz and A. Saliov, Diffus. Defect Data, Pt. B (1991), **19-20** (Gettering Defect Eng. Semicond. Technol.), 109-20 [31](#), [72](#)
- [BER93] W. Bergholz, D. Landsmann, P. Schauburger and B. Schoepperl, Proc.-Electrochem. Soc. (1993) **93-15** (Crystalline Defects and Contamination), 69-86 [2](#)
- [BER99] V. Bertagna, R. Erre, F. Rouelle and M. Chemla, *Journal of The Electrochemical Society*, **146**(1), 83-90, (1999) [22](#), [23](#), [25](#), [68](#), [74](#)
- [BIX86] J.W. Bixler, A.M. Bond, P.A. Lay, W. Thormann, P. van den Bosch, M. Fleischmann and B.S. Pons, *Analytical Chimica Acta*, 187 (1986) 67-77 [53](#)
- [BOE79] K. Böhm and B. Fischer, J. Appl. Phys. **50**(8), 1979 [111](#)
- [BOG67] A.F. Bogenschütz: "Ätzpraxis für Halbleiter", Carl Hanser Verlag, München 1967 [113](#), [114](#)
- [BOR94] K. Borgwarth, D.G. Ebling and J. Heinze, Phys. Chem., 98, 1317-1321, 1994, No. 10 [43](#)
- [BYK82] H. Byker, V.E. Wood and A.E. Austin, *J. Electrochem. Soc.*, **129**(9), p.1982-1987, September 1982 [103](#)
- [CAM98] S.A. Campbell and H.J. Lewerenz: "Semiconductor Micromachining", Fundamental Electrochemistry and Physics, John Wiley & Sons Ltd., Chichester, 1998 [18](#), [19](#), [67](#)
- [CAR01] J. Carstensen, persönliche Mitteilung [3](#), [86](#)
- [CHA79] J.-N. Chazalviel, Surface Science 88 (1979) 204-220 [25](#)
- [CHA13] D.L. Chapman, *Philos. Mag.* **25** (1913) 475 [16](#)
- [CHA80] J.-N. Chazalviel and T.B. Truong, *J. Electroanal. Chem.*, 114(1980) 299-303 [100](#), [101](#), [103](#)
- [CHA81] J.-N. Chazalviel and T.B. Truong, *J. Am. Chem. Soc.* **1981**, *103*, 7447-7451 [19](#), [100](#), [101](#), [103](#)

- [CHA83a] J.-N. Chazalviel, M. Stefenel and T.B. Truong, Surface Science 134 (1983) 865-885 Noth Holland Publishing Company 100
- [CHA89] J.-N. Chazalviel and F. Ozanam, *J. Electroanal. Chem.*, 269(1989) 251-266 23, 101, 103
- [CHR00] M. Christophersen, persönliche Mitteilung 26
- [DAN98] A. Danel, U. Straube, G. Kamarinos, E. Kamieniecki and F. Tardif, Proc.-Electrochem. Soc. (1998), 97-35 (Cleaning Technology in Semiconductor Device Manufacturing), 400-407q 2
- [EIC97] P. Eichinger and M. Rommel, Proc.-Electrochem. Soc. (1997), 97-22 (Crystalline Defects and Contamination: Their Impact and Control in Device Manufacturing II), 363-375 31
- [ELS88] A. Elshabini-Raid, J. He, G.L Stirk and S. Al-Mazrooh, *Solid State Electronics* Vol. 31, No. 9, pp. 1413-1420, 1988 7, 28
- [ENG92] R.C. Engstrom, B. Small and L. Kattan, *Anal. Chem.*, 1992, 64, 241-244 21
- [FAL94] T. Falter, D. Hellmann and P. Eichinger, Proc. SPIE-Int. Soc. Opt. Eng. (1994), 2337 (Optical Characterization Techniques for High-Performance Microelectronic Device Manufacturing), 104-11 31
- [FAL98] R. Falster and G. Borionetti, ASTM Spec. Tech. Publ. (1998), STP 1340 (Recombination Lifetime Measurements in Silicon), 226-249 72
- [FAN98] F.R.F. Fan, D. Clifff and A.J. Bard, *Anal. Chem.*, 1998, 70, 2941-2948 43
- [FAS87] G. Fasching: "Werkstoffe für die Elektrotechnik", Springer Verlag, Wien, New York 1987 5
- [FIS94] E.E. Fisch, R.H. Gaylord and S.A. Estes: Proc.-Electrochem. Soc. (1994), 94-7, (Proceedings of the Third International Symposium on Cleaning Technology in Semiconductor Device Manufacturing, 1993), 514-21 2
- [FLA98] J.C. Flake, M.M. Rieger, G.M. Schmid and P.A. Kohl, *J. Electrochem. Soc.*, 146(5) 1960-1965 (1999) 103
- [FOE91] H. Föll, Appl. Phys. A 53, 8-19 (1991) 25, 26, 27
- [GER61] H. Gerischer, Zeitschrift für Physikalische Chemie Neue, Folge, Bd. 27, S. 48-79 (1961) 19, 21
- [GER69] H. Gerischer, *Surface Sci.* 18, 97 (1969) 18, 19
- [GER95] H. Vogel: "Gerthsen Physik", Springer Verlag, Berlin, Heidelberg, New York 1995 85
- [GER99] P. Geranzani, M. Porrini, G. Borionetti, R. Orizio and R. Falster, *Journal of The Electrochemical Society*, 146(9) 3494-3499 (1999) 21, 36
- [GHA82] Ghandi: "Fabrication Principles Silicon and Gallium Arsenide", New York 1982 111
- [GLU95] S.W. Glunz, W. Warta, J. Appl. Phys. 77(1), 1 April 1995 3
- [GLU96] S.W. Glunz, C. Hebling and W. Warta, Diffus. Defect Data, Pt. B (1996), 51-52 (Polycrystalline Semiconductors IV), 69-74 3
- [GOU10] G. Gouy, *J. Phys.* 9 (1910) 457 16
- [GUP97] D.C. Gupta, Proc.-Electrochem. Soc. (1997), 97-12, (Diagnostic Techniques for Semiconductor Materials and Devices), 267-279 2
- [HAG98] T. Hagen, S. Grafström, J. Kowalski and R. Neumann, Appl. Phys. A: Mater. Sci. Process. (1998), A66 (Suppl., Pt. 2, Scanning Tunneling Microscopy/Spectroscopy and Related Teschniques), S973-S976 31
- [HAM98] C.H. Hamann und W. Vielstich: "Elektrochemie", Wiley-VCH, Weinheim, 1998 19, 20, 21, 23
- [HAS99] G. Hasse Diplomarbeit, Kiel 1999 49, 50
- [HAS01] G. Hasse, persönliche Mitteilung 28
- [HEL94a] D. Hellmann, M. Rother, M. Hill, W. Bergholz and M. Riedlbauer: Proc.-Electrochem. Soc. (1994), 94-9 (Contamination Control and Defect Reduction in Semiconductor Manufacturing III), 285-99 72

- [HEL94b] D. Hellmann, T. Falter, R. Berger, E. Burte, Proc. SPIE-Int. Soc. Opt. Eng. (1994), **2334** (Microelectronics Manufacturability, Yield and Reliability) , 161-70 [2](#)
- [HEN95] W.B. Henley, Proc. SPIE-Int. Soc. Opt. Eng. (1995), 2638 (Optical Characterization Techniques for High-Performance Microelectronic Device Manufacturing II), 172-82 [2](#), [31](#)
- [HIG90] G.S. Higashi, Y.J. Chabal, G.W. Trucks and K. Raghavachari, Appl. Phys. Lett. **56**(1), 12 February 1990 [23](#)
- [HOL66] P.J. Holmes and J.E. Snell, Microelec. Reliab. **5**, 337 (1966) [23](#)
- [HOP94] A.A. Hopgood, J. Appl. Pys. **76**(1) 1994 [2](#)
- [HOR89] P. Horowitz and W Hill: "Die Hohe Schule der Elektronik Teil 1 Analogtechnik", Elektor-Verlag GmbH, New Nork 1989 [151](#)
- [HOR93a] B.R. Horrocks, M.V. Mirkin, D.T. Pierce and A.J. Bard, *Anal. Chem.* **1993**, *65*, 1213-1224 [43](#)
- [HOR93b] B.R. Horrocks, D. Schmidtke, A. Heller and A.J. Bard, *Anal. Chem.* **1993**, *65*, 3605-3614 [53](#)
- [HOR94] M. Horikawa and T. Mizutani: " $n^+ - p$ Junction Leakage Current Generated by Micro-Defects in Denuded Zone of 64MbDRAM Silicon Substrates", Proceedings of the Seventh International Symposium on Silicon Materials Science and Technology, pp. 987-996, 1994 [29](#)
- [HOW85] M.J. Howes and D.V.Morgan: "Gallium Arsenide: Materials, Devices and Circuits", John Wiley and Sons Ltd., Chichester, New York, Brisbane Toronto, Singapore, 1985 [111](#)
- [IST97] A.A. Istratov, C. Flink, H. Hieslmair, T. Heiser and E.R. Weber, Appl. Phys. Lett. **71**(15), 13 October 1997 [2](#), [36](#), [72](#)
- [JUD71] J.S. Judges, *Journal of The Electrochemical Society*, **118**, p.1772, (1971) [23](#)
- [KEE98] G.K. Keefe and W.M. Hughes, ASTM Spec. Tech. Publ. (1998), STP 1340 (Recombination Lifetime Measurements in Silicon), 112-122 [2](#), [31](#)
- [KER70] W. Kern and D.A. Puotinen, RCA Review, June 1970, p. 187-206 [110](#), [111](#)
- [KOR95] J. Korallus Diplomarbeit, Clausthal-Zellerfeld 1995 [13](#), [62](#), [110](#)
- [KUS97] W. Kusian, H.-C. Ostendorf, J. Palm and A.L. Endrös, Conf. Rec. IEEE Photovoltaic Spec. Conf. (1997), 26th, 131-134 [31](#)
- [KWA89b] J. Kwad and A.J. Bard, *Anal. Chem.*, **1989**, *61*, 1794-1799 [53](#)
- [LAS76] D. Laser and A.J. Bard, The Journal of Physical Chemistry, Vol. 80, No. 5, 1976 [100](#), [103](#)
- [LEE91a] C. Lee, C.J. Miller and A.J. Bard, Analytical Chemistry, Vol. 63, No. 1, January 1, 1991, p. 78-83 [43](#), [53](#)
- [LEH88] V. Lehmann and H. Föll, *J. Electrochem. Soc.* **135**, 2381 (1988) [26](#), [31](#), [35](#)
- [LI95] Analytical Chemistry, Vol.67, No.3, February 1, 1995 [100](#)
- [LIE95] S. Liebert Diplomarbeit, Kiel 1995 [53](#), [54](#)
- [LIN98a] J. Linn, G. Rouse, W. Wereb, R. Leggett, S. Slasor, R. Lowry, R. Cameron and R. Canfield, ASTM Spec. Tech. Publ. (1998), STP 1340 (Recombination Lifetime Measurements in Silicon), 283-293 [2](#), [111](#)
- [LIP97] W. Lippik Dissertation, Kiel 1997 [3](#), [26](#), [33](#), [68](#), [111](#), [113](#), [131](#), [149](#)
- [LIU96] Q. Liu, H.E. Ruda, G.M. Chen and M. Simard-Normandin, J. Appl. Phys. **79**(10), 15 May 1996 [31](#)
- [LOO81] B.H.Loo, K.W. Frese and S.Roy Morrison, Surface Science 109 (1981) 75-81 [19](#)
- [LU90] Y.C. Lu, T.S. Kalkur and C.A. Paz de Araujo, Journal of Electronic Materials, Vol. 19, No. 1, 1990 [101](#), [110](#)
- [MAC92] S.A.S. Machado, L.H. Mazo and L.A. Avaca, *J. Braz. Chem. Soc.*, Vol. 3, N° 3, 1992 [53](#)
- [MAL94] V. Malina, V. Micheli, J. Kohout and D. Berková, Semicond. Sci. Technol. **9**, pp. 1523-1528, 1994 [111](#)

- [MAN90] D. Mandler and A.J. Bard, *Langmuir* **1990**, 6, 1489-1494 [22](#)
- [MAS92] T.B. Massalski, H. Okamoto, P.R. Subramanian and L. Kacprzak: " *Binary Alloy Phase Diagrams*", 2nd Ed. ASM International 1992 [62](#)
- [MCK93] J.P. McKelvey: " *Solid State Physics*", Krieger Publishing Company, Malabar 1993 [5](#), [6](#), [7](#)
- [MEM66] R. Memming and G. Schwandt, *Surface Science* **5** (1966) 97-110 North Holland Publishing Co., Amsterdam [17](#)
- [MOR77] S.R. Morrison: " *The Chemical Physics of Surfaces*" (Plenum, NewYork 1977) [17](#), [19](#), [68](#), [126](#)
- [MOR80] S.R. Morrison: " *Electrochemistry at Semiconductor and Oxidized Metal Electrodes*", Plenum Press, New York 1980 [17](#), [19](#)
- [MOR94] S.P. Morgan and D.V. Morgan, *Semicond. Sci. Technol.*, **9**, 1994, p.2278-2284 [110](#)
- [MUR95] Y. Murakami, H. Abe and T. Shingyouji, *Jpn. J. Appl. Phys.* Vol.34(1995) pp.1477-1482 [36](#)
- [NAG82] G. Nagasubramanian, B.L. Wheeler, F.-R.F. Fan and A.J. Bard, *J. Electrochem. Soc.*, **129**(8) 1742-1745 (1982) [100](#), [103](#)
- [NAK87] Y. Nakato, T. Ueda, Y. Egi and H. Tsubomura, *J. Electrochem. Soc.*, **134**(2) 353-358 (1987) [25](#)
- [NAR85] E.S. Nartowitz, A.M. Goodman, *J. Electrochem. Soc. Solid State Science and Technology*, Vol. 132, No. 12, pp. 2992-2997, 1985 [35](#)
- [NOR96] M. Simard-Normandin, *Proc. SPIE-Int. Soc. Opt. Eng.* (1996), 2877 (Optical Characterization Techniques for High-Performance Microelectronic Device Manufacturing), III), 186-197 [2](#)
- [NOW95] M.J. Nowlan, S.J. Hogan and G. Darkazilli, *A Cost Analysis for High Volume Crystalline Silicon Photovoltaic Module Manufacturing*, Proc. of the 13th European Photovoltaic Conference, p. 2334 Nizza 1995 [111](#)
- [OBE95] G. Obermeier, J. Hage, N. Hilger and D. Huber, *Proc. SPIE-Int. Soc. Opt. Eng.* (1995), 2638 (Optical Characterization Techniques for High-Performance Microelectronic Device Manufacturing II), 104-12 [29](#), [36](#)
- [OBE98] G. Obermeier, J. Hage and D. Huber, *ASTM Spec. Tech. Publ.* (1998), STP 1340 (Recombination Lifetime Measurements in Silicon), 305-317 [31](#)
- [OLA96] K.R. Olasupo and M.K. Hatalis, *IEEE Transactions on Electron Devices*, Vol. 43, No. 8, August 1996 [28](#), [29](#)
- [ODE90] D.M. Odell and W.J. Bowyer, *Anal. Chem.*, **1990**, 62, 1619-1623 [100](#), [103](#)
- [OST95] H.-C. Ostendorf and A.L. Endrös, *Mater. Res. Soc. Symp. Procl.* (1995), 378 (Defect and Impurity Engineered Semiconductors and Devices), 597-602 [31](#)
- [PAR51] R. Parson, *Trans. Faraday. Soc.*, **47**, (1951), 1332 [20](#), [67](#)
- [PIO83] A. Piotrowska, A. Guivarc'h and G. Pelous, *Solid-State Electronics*, Vol. 26, No. 3, pp. 231-236, 1985 [12](#), [15](#)
- [PLE90] Y.V. Pleskov: " *Solar Energy Conversion*", Springer Verlag, Berlin, Heidelberg 1990 [35](#)
- [POL95a] M.L. Polignano, C. Bresolin, F. Cazzaniga, A. Sabbadini and G. Queirolo, *Proc. SPIE-Int. Soc. Opt. Eng.* (1995), 2638 (Optical Characterization Techniques for High-Performance Microelectronics Device Manufacturing II), 14-26 [31](#)
- [POL95b] M.L. Polignano, C. Bresolin, F. Cazzaniga and G. Queirolo, *Ion Implant. Technol.-94*, Proc. Int. Conf., 10th (1995), Meeting Date 1994, 592-5. Editor(s): Coffa, Salvatore, Publisher: Elsevier, Amsterdam, Neth. [31](#)
- [POL98] M.L. Polignano, F. Cazzaniga, A. Sabbadini, F. Zanderigo and F. Priolo, *Materials Science in Semiconductor Processing* 1 (1998), 119-130 [31](#)
- [POP00] G. Popkirov, persönliche Mitteilung [49](#), [75](#), [84](#)
- [PRE89] W.H. Press, B.P. Flannery, S.A. Teukolsky and W.T. Vetterling: " *Numerical Recipes in Pascal*", Cambridge University Press, Cambridge, New York, Port Chester, Melbourne, Sydney 1989 [140](#)

- [RID74] J.L. Rideout, Solid-State Electronics, Vol. 18, pp. 541-550, 1974 [11](#), [12](#), [15](#)
- [ROE95] H. Röppischer, Yu. A. Bumai and B. Feldmann, *J. Electrochem. Soc.*, **142**(2) 650-655 (1995) [75](#)
- [ROS95] J.J. Rosato, R.M. Hall, T.B. Parry, P.G. Lindquist and T.D. Jarvis, Mater. Res. Soc. Symp. Proc. (1995), 386 (Ultraclean Semiconductor Processing Technology and Surface Chemical Cleaning and Passivation), 47-54 [111](#)
- [SCH38] W. Schottky: "Halbleitertheorie der Sperrschicht", *Naturwissenschaften*, **26**, 843 (1938) [29](#)
- [SCH95] H.-J. Schulze and G. Deboy, Proc. SPIE-Int. Soc. Opt. Eng. (1995), 2638 (Optical Characterization Techniques for High-Performance Microelectronic Device Manufacturing II), 234-45 [72](#)
- [SCH96a] W. Schmickler: "Grundlagen der Elektrochemie", Vieweg, New York, Oxford, 1996 [19](#), [21](#), [22](#), [68](#)
- [SCH96b] A. Schönecker, J.A. Eikelboom, A.R. Burgers, P. Lölgén, C. Leguijt and W.C. Sinke, *J. Appl. Phys.* **79**(1), 1 February 1996 [31](#)
- [SCH98] D.K. Schroder, "New Life in Detecting Defects", *Circuits & Devices*, November 1998 [2](#), [31](#)
- [SEA91] P.C. Searson and X.G. Zhang, *Electrochimica Acta*, Vol. 36, No 3/4, pp. 499-503, 1991 [67](#), [75](#)
- [SEB82] T. Sebestyen, Solid-State Electronics, Vol. 25, No.7, pp. 543-550. 1982 [12](#)
- [SER94] C. Serre, S. Barret and R. Hérino, *Journal of Electrochemical Chemistry*, **370** (1994) 145-149 [23](#)
- [SHI98] Guangda Shi, L.F. Garfias-Mesias and W.H. Smyrl, *J. Electrochem. Soc.* **145**(6) 2011-2016 (1998) [43](#)
- [SHO50] W. Shockley, "The Theory of pn Junctions in Semiconductors and pn Junction Transistors", *Bell Syst. Tech. J.*, **28**, 435 (1949); *Electrons and Holes in Semiconductors*, D. Van Nostrand, Princeton, N.J., 1950 [28](#)
- [SHO52] W. Shockley and W.T. Read: "Statistics of the Recombination of Holes and Electrons", *Phys. Rev.*, **87**, 835(1952) [36](#)
- [SOM85] M. Somogyi, Symposium on Electronics Technology, 1985, p. 409-414 [22](#), [25](#)
- [STE94] A.E. Stephens: "Controlling the Oxygen Precipitation in Non-Epi Wafers to Prevent Charge Leakage from DRAM Trench Capacitors", Proceedings of the Seventh International Symposium on Silicon Materials Science and Technology, pp. 908-919, 1994 [29](#)
- [SUZ83] T. Suzuki, Y. Osaka and M. Hirose, *Japanese Journal of Applied Physics*, Vol. 22, No. 5, May 1983 [21](#)
- [SZE81] S.M. Sze: "Physics of Semiconductor Devices", John Wiley & Sohns, Inc. New York 1981 [13](#), [15](#), [29](#)
- [SZE85] S.M. Sze: "Semiconductor Devices Physics and Technology", John Wiley & Sohns, Inc. New York 1985 [2](#), [6](#), [11](#), [29](#), [36](#)
- [TAR95] F. Tardif, J.P. Joly and D. Walz, Proc.-Electrochem. Soc. (1995), 95-30 (Analytical Techniques for Semiconductor Materials and Process Characterization II), 299-315 [2](#), [22](#), [31](#)
- [TAU86] R.N. Tauber and S. Wolf: *Silicon Processing for the VLSI Era*, Vol. 1, Lattice Press, Sunset Beach, California, 1986 [2](#)
- [TIE93] U. Tietze und Ch. Schenk: "Halbleiter-Schaltungs-Technik", Springer-Verlag, Berlin, Heidelberg, New York, 1993 [49](#), [145](#), [147](#)
- [TRU90] G.W. Trucks, K. Raghavachari, G.S. Higashi and Y.J. Chabal, *Phys. Rev. Lett.* **64**(4) 1990 [24](#)
- [TUC75] B Tuck, *Journal of Mathematical Science*, **10**, (1975) 321-339 [22](#), [112](#), [113](#)
- [TUR58] D.R. Turner, *Journal of The Electrochemical Society*, **105**(7), 402-408, (1958) [22](#)
- [TUR80] J.A. Turner, J. Manassen and A.J. Nozik, *Appl. Phys. Lett.* **35**(5), 1 September 1980 [100](#)
- [VER94] S. Verhaverbeke, I. Teerlinck, C. Vinckier, G. Stevens, R. Cartuyvels and M.M. Heyns, *J. Electrochem. Soc.* **141** 2852 (1987) [23](#)
- [WAL79] W. Walcher: "Praktikum der Physik", B.G. Teubner Stuttgart 1979 [85](#)

- [WAL95a] D. Walz, J.-P. Joly and G. Kamarinos, Appl. Phys. A, 62, 345-353 (1996) [2](#), [28](#)
- [WAL95b] D Walz, G. Le Carval, J.-P. Joly and G. Kamarinos, Semicond. Sci. Technol. **10** (1995) 1022-1033 [31](#), [35](#), [36](#)
- [WAL95c] D. Walz, J.P. Joly, G. Kamarinos and K. Barla, Proc.-Electrochem. Soc. (1995), 95-30 (Analytical Techniques for Semiconductor Materials and Process Charakterization II), 35-43 [2](#)
- [WAL95d] D Walz, J.-P. Joly, M. Suarez, J. Palleau and G. Kamarinos, Proc.-Electrochem. Soc. (1995), 95-30 (Analytical Techniques for Semiconductor Materials and Process Characterization II), 64-72 [2](#)
- [WED87] G. Wedler: "*Lehrbuch der Physikalischen Chemie*", VCH, Weinheim 1987 [17](#), [18](#)
- [WEI95a] Ch. Weißmantel und C. Hamann: "*Grundlagen der Festkörperphysik*", Johann Ambrosius Barth Verlag , Heidelberg, Leibzig, 1995 [5](#)
- [WEI95b] C. Wei, A.J. Bard G. Nagy and K. Toth, *Anal. Chem.*, **1995**, 67, 1346-1356 [43](#)
- [WIT98] G. Wittstock, B. Gründig, B. Strehlitz and K. Zimmer, *Electroanalysis* **1998**, 10, No. 8 [31](#), [72](#)
- [WOL86] M. Wolovelsky, J. Levy, Y. Goldstein, A. Many, S.Z. Weisz and O. Resto, Surface Science 171 (1986) 442-464 [75](#)
- [WUE89] H.J. Würfl Dissertation, Darmstadt 1989 [13](#), [15](#)
- [YAB86] E. Yablonovitch, D.L. Allara, C.C. Chang, T. Gmitterand T.B. Bright, Phys. Rev. Lett. **57**(2) 1986 [23](#), [131](#)
- [YAN96] J. Yang, W. Yang, Y. Bai, C. Shen, D. Wang, X. Chai, T. Li, C. Li and A. Pan, Thin Solids Films 284-285 (1996) 477-480 [111](#)

Danksagung

Herrn Prof. Dr. H. Föll danke ich für die Stellung des sehr interessanten Themas, deren Bearbeitung mir über die Jahre viele Mühen aber auch viel Freude bereitete und meine Kenntnisse auf verschiedenen Gebieten der Halbleiterphysik, Materialwissenschaft, Elektrochemie und Programmierung vertieften.

Herrn Dr. habil. G. Popkirov danke ich für die langjährige und äußerst fruchtbare Zusammenarbeit. Seine weitreichenden Kenntnisse in Elektrochemie speziell im Umgang mit elektrochemischer Meßtechnik mündeten in viele wertvolle Anregungen und Verbesserungsvorschläge, die geradlinig den Weg des Erfolges zeichneten.

Herrn Dr. J. Carstensen danke ich besonders für seine allzeitige Diskussionsbereitschaft. Seine Beiträge, speziell zu Problemen theoretischer Natur, erhellten das Verborgene und verhalfen mir zum Verständnis zahlreicher physikalischer Vorgänge und Hintergründe.

Allen Mitgliedern der Werkstatt der Technischen Fakultät danke ich für die stets schnell und sehr präzise durchgeführten Arbeiten, die Voraussetzung für ein Gelingen der zahlreichen Experimente waren.

Und schließlich danke ich allen anderen Mitgliedern meiner Arbeitsgruppe für die gemeinsam verbrachte Zeit, die mir bis in ferne Zukunft in angenehmer Erinnerung bleiben wird.